

Geometria e algebra lineare

Bruno Martelli

Bruno Martelli
Dipartimento di matematica
Università di Pisa

`people.dm.unipi.it/martelli`

Versione 4, settembre 2025

A Mariette e Sabatino

Il testo è rilasciato con la licenza *Creative Commons*-BY-NC-SA¹. Le figure sono tutte di pubblico dominio, eccetto le seguenti, che hanno una licenza CC-BY-SA² e sono state scaricate da Wikimedia Commons ed eventualmente modificate successivamente:

- Figura 1.2 (partizione di un insieme), creata da Wshun;
- Figura 1.5 (numeri complessi), creata da Wolfkeeper;
- Figura 1.6 (numeri complessi), creata da Kmhmkmh;
- Figura 1.10 (numeri razionali), creata da Cronholm144;
- Figura 3.9 (prodotto riga per colonna), creata da Bilou;
- Figura 4.3 (cubi distorti) creata da Irrons;
- Figure 3.5, 3.6 e 9.2 (regola della mano destra), creata da Acdx;
- Figura 13.10 (toro), creata da YassineMrabet;
- Figure 13.12 e 13.14 (quadriche), create da Sam Derbyshire;
- Figura 13.13-(sinistra, iperboloide rigato) creata da Ag2gaeh.

La copertina raffigura un'opera di Mariette Michelle Egreteau. Si ringraziano Mariette Michelle Egreteau e Silvio Martelli per i contributi dati alla veste grafica.

¹È possibile distribuire, modificare, creare opere derivate dall'originale, ma non a scopi commerciali, a condizione che venga riconosciuta la paternità dell'opera all'autore e che alla nuova opera vengano attribuite le stesse licenze dell'originale (quindi ad ogni derivato non sarà permesso l'uso commerciale).

²Come la licenza precedente, ma con anche la possibilità di uso commerciale.

Indice

Introduzione	1
Capitolo 1. Nozioni preliminari	3
1.1. Gli insiemi	4
1.2. Funzioni	12
1.3. Polinomi	21
1.4. Numeri complessi	25
1.5. Strutture algebriche	34
Esercizi	36
Complementi	38
1.I. Infiniti numerabili e non numerabili	38
1.II. Costruzione dei numeri reali	40
Capitolo 2. Spazi vettoriali	43
2.1. Lo spazio euclideo	43
2.2. Spazi vettoriali	46
2.3. Dimensione	60
Esercizi	75
Capitolo 3. Sistemi lineari	79
3.1. Algoritmi di risoluzione	79
3.2. Teorema di Rouché - Capelli	85
3.3. Determinante	93
3.4. Algebra delle matrici	104
Esercizi	111
Capitolo 4. Applicazioni lineari	115
4.1. Introduzione	115
4.2. Nucleo e immagine	123
4.3. Matrice associata	130
4.4. Endomorfismi	137
Esercizi	146
Complementi	147
4.I. Spazio duale	147
Capitolo 5. Autovettori e autovalori	151

5.1. Definizioni	151
5.2. Teorema di diagonalizzabilità	163
Esercizi	170
Capitolo 6. Forma di Jordan	173
6.1. Forma di Jordan	173
6.2. Teorema di Cayley – Hamilton	186
6.3. Polinomio minimo	190
Esercizi	195
Complementi	196
6.1. Forma di Jordan reale	196
Capitolo 7. Prodotti scalari	199
7.1. Introduzione	199
7.2. Matrice associata	210
7.3. Sottospazio ortogonale	213
7.4. Classificazione dei prodotti scalari	219
7.5. Isometrie	227
Esercizi	231
Complementi	233
7.1. Tensori	233
Capitolo 8. Prodotti scalari definiti positivi	239
8.1. Nozioni geometriche	239
8.2. Isometrie	253
Esercizi	261
Complementi	264
8.1. Angoli fra tre vettori nello spazio	264
Capitolo 9. Lo spazio euclideo	267
9.1. Prodotto vettoriale	267
9.2. Sottospazi affini	272
9.3. Affinità	291
Esercizi	307
Capitolo 10. Poligoni e poliedri	311
10.1. Poligoni	311
10.2. Poliedri	332
Esercizi	341
Complementi	342
10.1. Coordinate baricentriche	342
Capitolo 11. Teorema spettrale	347
11.1. Prodotti hermitiani	347
11.2. Endomorfismi autoaggiunti	350
11.3. Il teorema	352

Esercizi	355
Complementi	356
11.I. Dimostrazione del criterio di Cartesio	356
Capitolo 12. Geometria proiettiva	359
12.1. Lo spazio proiettivo	359
12.2. Completamento proiettivo	367
12.3. Proiettività	372
Esercizi	377
Complementi	377
12.I. I Teoremi di Desargues e di Pappo	377
12.II. Dualità	380
Capitolo 13. Quadriche	383
13.1. Introduzione	383
13.2. Coniche	389
13.3. Quadriche proiettive	409
13.4. Quadriche in $\mathbb{R}^3$	418
Esercizi	427
Complementi	429
13.I. Il Teorema di Pascal	429
Soluzioni di alcuni esercizi	431
Indice analitico	439

Introduzione

*L'algèbre n'est qu'une géométrie écrite,
la géométrie n'est qu'une algèbre figurée.*

Sophie Germain

Il presente libro contiene una introduzione agli argomenti trattati abitualmente negli insegnamenti di geometria e algebra lineare dei corsi di studio universitari di tipo scientifico.

La matematica contemporanea può essere suddivisa sommariamente in tre settori: l'*algebra* concerne i numeri, i simboli e le loro manipolazioni tramite le quattro operazioni; la *geometria* riguarda lo studio delle figure nel piano e nello spazio; l'*analisi* si basa sul calcolo infinitesimale e si occupa di quegli ambiti (successioni, funzioni, derivate, integrali) collegati al concetto di limite.

Nella storia del pensiero scientifico degli ultimi secoli, i progressi più rilevanti sono stati fatti nel momento in cui si è scoperto che fenomeni apparentemente scollegati sono in realtà descrivibili pienamente nel quadro di un unico formalismo matematico: la gravitazione di Newton descrive sia il moto di caduta di un grave che quello di rivoluzione dei pianeti; le equazioni di Maxwell fondono elettricità e magnetismo in un'unica teoria; la meccanica quantistica spiega alcuni fenomeni fisici bizzarri e fornisce un quadro solido alla chimica ed in particolare alla tabella periodica degli elementi, eccetera. Ciascun processo di fusione di due ambiti scientifici differenti ha portato ad una migliore comprensione dei fenomeni di entrambi.

In matematica si è sviluppato un processo unificante di questo tipo con la costruzione del *piano cartesiano*. Interpretando un punto del piano come una coppia (x, y) di numeri reali, abbiamo implicitamente iniziato a fondere la geometria euclidea con una parte dell'algebra chiamata *algebra lineare*. La geometria si occupa di figure quali punti, rette, piani, coniche, poligoni, poliedri. L'algebra lineare tratta invece sistemi di equazioni in più variabili di primo grado (cioè lineari), equazioni di secondo grado (riducendole quanto possibile ad uno studio di tipo lineare), ed oggetti algebrici più complessi come le *matrici* ed i *vettori*.

La geometria e l'algebra lineare sono oggi così unite che si fa fatica ormai a decidere quale argomento sia "algebra lineare" e quale sia "geometria". Il punto di contatto fra i due ambiti è la nozione di *vettore*, a

cui si può dare contemporaneamente una valenza geometrica (un punto nel piano o una freccia) ed algebrica (una sequenza di numeri). I vettori giocano un ruolo centrale in questo testo.

La suddivisione in capitoli del libro rispecchia la struttura di un insegnamento standard di geometria e algebra lineare. Iniziamo revisionando alcuni preliminari algebrici (insiemi, polinomi, funzioni, eccetera) e quindi passiamo a definire la nozione di *spazio vettoriale* che è fondamentale in tutta la trattazione. Impariamo come risolvere i sistemi di equazioni lineari e introduciamo alcune funzioni particolari dette *applicazioni lineari*. Questo ci porta naturalmente allo studio di *autovalori e autovettori*. Il Capitolo 6 sulla forma di Jordan è opzionale e può essere saltato (come del resto tutti gli argomenti presentati come complementi alla fine dei capitoli).

Passiamo quindi a studiare i *prodotti scalari* ed infine applichiamo tutti gli strumenti costruiti nei capitoli precedenti per studiare più approfonditamente la geometria euclidea con i suoi protagonisti: punti, rette, piani, poligoni, poliedri, coniche e quadriche. Un ruolo a parte è giocato dal *teorema spettrale*, un risultato profondo di algebra lineare che ha applicazioni in vari ambiti della matematica e della scienza. Nell'ultima parte del libro introduciamo una geometria non euclidea, detta *geometria proiettiva*, ottenuta aggiungendo i punti all'infinito allo spazio euclideo. Il Capitolo 10 è dedicato a poligoni e poliedri e contiene numerosi teoremi di geometria piana e solida. I Capitoli 12 e 13 riguardanti geometria proiettiva e quadriche sono entrambi opzionali e in buona parte indipendenti l'uno dall'altro.

Il libro contiene vari esercizi, alcuni posizionati lungo la trattazione ed altri alla fine dei capitoli. Teoria ed esercizi sono entrambi essenziali per una piena comprensione del testo. Nella scrittura mi sono posto due obiettivi, a volte non semplici da conciliare: descrivere in modo trasparente e rigoroso i passaggi logici che formano il corpo di ogni tipo di ragionamento astratto, con un particolare accento sulle motivazioni che hanno portato i matematici a seguire una strada invece che un'altra, e fornire una notevole quantità di esempi e di strumenti utili ad applicare proficuamente queste nozioni per affrontare problemi concreti in vari ambiti della scienza.

Bruno Martelli

Nozioni preliminari

L'aspetto che caratterizza maggiormente una teoria matematica è l'approccio *assiomatico-deduttivo*: si parte da alcuni *concetti primitivi* che non vengono definiti, si fissano degli *assiomi*, cioè dei fatti che vengono supposti veri a priori, e quindi sulla base di questi si sviluppano dei *teoremi*.

Nella geometria euclidea i concetti primitivi sono alcune nozioni geometriche basilari come quella di punto, retta e piano. Gli assiomi sono quelli formulati da Euclide nei suoi *Elementi* tra il IV e il III secolo a.C., perfezionati poi da Hilbert nel 1899.

Nella geometria analitica che seguiamo in questo libro, i concetti primitivi e gli assiomi da cui partiamo sono invece quelli della *teoria degli insiemi*. Partiamo con gli insiemi numerici $\mathbb{Z}$, $\mathbb{Q}$ e $\mathbb{R}$ formati dai numeri interi, razionali e reali. Interpretiamo quindi $\mathbb{R}$ come una retta e *definiamo* il piano e lo spazio cartesiano come gli insiemi $\mathbb{R}^2$ e $\mathbb{R}^3$ formati da coppie (x, y) e triple (x, y, z) di numeri reali. Questo approccio più astratto ci permette di definire subito lo spazio n -dimensionale $\mathbb{R}^n$ per qualsiasi n . Con la teoria degli insiemi e l'algebra possiamo definire e studiare rigorosamente oggetti pluridimensionali che la nostra intuizione non può afferrare.

In questo primo capitolo introduciamo alcune nozioni preliminari che saranno usate in tutto il libro. Supponiamo che la lettrice abbia già dimestichezza con gli insiemi e con l'algebra che viene insegnata nelle scuole superiori. Più raramente useremo alcuni teoremi di analisi.

Iniziamo richiamando la teoria degli *insiemi*, quindi passiamo alle *funzioni*, che hanno un ruolo fondamentale in tutta la matematica moderna. Fra le funzioni, spiccano per semplicità i *polinomi*. Lo studio delle radici dei polinomi ci porta quindi ad allargare l'insieme $\mathbb{R}$ dei numeri reali a quello $\mathbb{C}$ dei *numeri complessi*. Infine, tutti questi insiemi numerici hanno delle operazioni (somma e prodotto) che soddisfano alcune proprietà algebriche (commutativa, associativa, eccetera). Studieremo più in astratto gli insiemi dotati di operazioni di questo tipo e ciò ci porterà a definire delle strutture algebriche note come *gruppi*, *anelli* e *campi*.

I complementi contengono un paio di approfondimenti: una discussione sugli insiemi infiniti, in cui mostriamo che i numeri reali sono "di più" degli interi e dei razionali (pur essendo infiniti sia gli uni che gli altri) e la loro costruzione rigorosa.

1.1. Gli insiemi

La teoria degli insiemi è alla base di tutta la matematica moderna.

1.1.1. Gli insiemi numerici. Un insieme è generalmente indicato con due parentesi graffe $\{\}$, all'interno delle quali sono descritti tutti i suoi elementi. A volte si descrivono solo alcuni elementi e si usano dei puntini $\dots$ per indicare il resto. Ad esempio, questo è l'insieme dei *numeri naturali*:

$$\mathbb{N} = \{0, 1, 2, 3, \dots\}$$

L'insieme $\mathbb{N}$ contiene infiniti elementi. Se aggiungiamo i numeri negativi otteniamo l'insieme dei *numeri interi*

$$\mathbb{Z} = \{\dots, -2, -1, 0, 1, 2, \dots\}.$$

Se oltre ai numeri interi consideriamo tutti i numeri esprimibili come frazioni $\frac{a}{b}$, otteniamo l'insieme dei *numeri razionali*

$$\mathbb{Q} = \left\{ \dots, -\frac{2}{5}, \dots, \frac{1}{8}, \dots, \frac{9}{4}, \dots \right\}.$$

Sappiamo inoltre che in aritmetica esiste un insieme ancora più grande, chiamato $\mathbb{R}$, formato da tutti i *numeri reali*.

Questi insiemi numerici dovrebbero già essere familiari fin delle scuole superiori. Da un punto di vista più rigoroso, l'insieme $\mathbb{N}$ è un concetto primitivo, mentre gli insiemi $\mathbb{Z}$, $\mathbb{Q}$ e $\mathbb{R}$ sono costruiti ciascuno a partire dal precedente in un modo opportuno. La costruzione dell'insieme $\mathbb{R}$ dei numeri reali è abbastanza complessa ed è descritta nella Sezione 1.II.

1.1.2. Dimostrazione per assurdo. Possiamo subito convincerci del fatto che molti numeri reali che ci sono familiari non sono razionali. Quella che segue è la prima proposizione del libro ed è illuminante perché contiene un esempio di *dimostrazione per assurdo*.

Proposizione 1.1.1. *Il numero $\sqrt{2}$ non è razionale.*

Dimostrazione. Supponiamo per assurdo che $\sqrt{2}$ sia razionale. Allora $\sqrt{2} = \frac{a}{b}$, dove $\frac{a}{b}$ è una frazione. Elevando al quadrato e moltiplicando per b^2 entrambi i membri otteniamo

$$a^2 = 2b^2.$$

Questa uguaglianza però non è possibile: il doppio di un quadrato non è mai un quadrato (esercizio: usare la decomposizione in fattori primi). $\square$

In una dimostrazione per assurdo, si nega la tesi e si dimostra che questo porta ad un assurdo. Ne deduciamo che la tesi non può essere falsa, e quindi è vera per esclusione.

1.1.3. Sottoinsiemi. Introduciamo un po' di notazione insiemistica. Se x è un elemento dell'insieme A , scriviamo

$$x \in A$$

per dire che x appartiene ad A . Se questo non accade, scriviamo $x \notin A$. Ad esempio

$$2 \in \mathbb{Z}, \quad \sqrt{2} \notin \mathbb{Q}.$$

Un *sottoinsieme* di A è un insieme B formato da alcuni elementi di A . Se B è un sottoinsieme di A , scriviamo

$$B \subset A.$$

In questo caso tutti gli elementi di B sono anche elementi di A . Ad esempio

$$\{1, 3, 5\} \subset \{-1, 1, 3, 4, 5\}.$$

Se esiste almeno un elemento di A che non sia contenuto in B , allora diciamo che A contiene B *strettamente* e possiamo usare il simbolo

$$B \subsetneq A.$$

Ad esempio gli insiemi numerici descritti precedentemente sono ciascuno contenuto strettamente nel successivo:

$$\mathbb{N} \subsetneq \mathbb{Z} \subsetneq \mathbb{Q} \subsetneq \mathbb{R}.$$

D'altro canto, se valgono entrambi i contenimenti

$$A \subset B, \quad B \subset A$$

allora chiaramente $A = B$, cioè A e B sono esattamente lo stesso insieme di elementi. Questo fatto verrà usato spesso nelle dimostrazioni presenti in questi libro.

Il simbolo $\emptyset$ indica l'*insieme vuoto*, cioè l'insieme che non contiene nessun elemento.

1.1.4. Unione, intersezione e differenza. Dati due insiemi A e B , possiamo considerare la loro *intersezione* e la loro *unione*, indicate rispettivamente:

$$A \cap B, \quad A \cup B.$$

L'intersezione consiste di tutti gli elementi che stanno *sia* in A che in B , mentre l'unione consiste di quelli che stanno in A *oppure* in B . Ad esempio, se $A = \{1, 2, 3\}$ e $B = \{3, 5\}$, allora

$$A \cap B = \{3\}, \quad A \cup B = \{1, 2, 3, 5\}.$$

La *differenza* $A \setminus B$ è l'insieme formato da tutti gli elementi di A che non stanno in B . Se B è un sottoinsieme di A , la differenza $A \setminus B$ è anche chiamata il *complementare* di B in A ed è indicata con B^c .

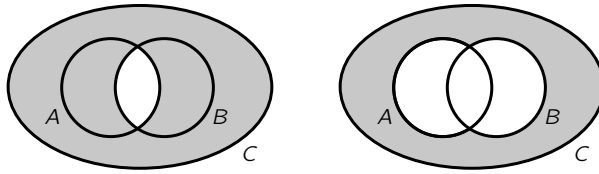


Figura 1.1. Le leggi di De Morgan.

Proposizione 1.1.2 (Leggi di De Morgan). *Siano A e B due sottoinsiemi di un insieme C . Valgono le seguenti uguaglianze:*

$$(A \cap B)^c = A^c \cup B^c$$

$$(A \cup B)^c = A^c \cap B^c$$

Qui X^c indica il complementare di X in C , cioè $X^c = C \setminus X$.

Dimostrazione. In ciascuna uguaglianza, entrambi gli insiemi indicano lo stesso sottoinsieme di C descritto (in grigio) nella Figura 1.1. $\square$

1.1.5. Notazione insiemistica. Un insieme A è spesso definito nel modo seguente:

$$A = \{\text{tutti quegli elementi } x \text{ tali che valga una certa proprietà } P\}.$$

Tutto ciò in linguaggio matematico si scrive più brevemente così:

$$A = \{x \mid P\}.$$

La barra verticale $\mid$ è sinonimo di "tale che". Ad esempio, intersezione, unione e complemento possono essere descritti in questo modo:

$$A \cap B = \{x \mid x \in A \text{ e } x \in B\},$$

$$A \cup B = \{x \mid x \in A \text{ oppure } x \in B\},$$

$$A \setminus B = \{x \mid x \in A \text{ e } x \notin B\}$$

Si possono anche usare i simboli $\wedge$ e $\vee$ come sinonimi di *e* ed *oppure*, ed i due punti $:$ al posto della barra $\mid$. Ad esempio, la prima legge di De Morgan (Proposizione 1.1.2) può essere dimostrata notando che entrambi gli insiemi $(A \cap B)^c$ e $A^c \cup B^c$ possono essere descritti nel modo seguente:

$$\{x \in C \mid x \notin A \text{ oppure } x \notin B\}.$$

Questa notazione per gli insiemi è molto flessibile. Ad esempio l'insieme P dei numeri pari può essere descritto nel modo seguente:

$$P = \{2n \mid n \in \mathbb{N}\}.$$

Otteniamo ovviamente

$$P = \{0, 2, 4, \dots\}.$$

Analogamente, l'insieme dei numeri dispari è

$$D = \{2n + 1 \mid n \in \mathbb{N}\} = \{1, 3, 5, \dots\}.$$

Infine, l'insieme dei quadrati è

$$Q = \{n^2 \mid n \in \mathbb{N}\} = \{0, 1, 4, 9, 16, \dots\}.$$

Un numero pari generico si scrive come $2n$, mentre un numero dispari generico si scrive come $2n + 1$. Usando questa notazione si possono dimostrare agevolmente dei teoremi, ad esempio questo:

Proposizione 1.1.3. *Ogni numero dispari è differenza di due quadrati.*

Dimostrazione. Un generico numero dispari si scrive come $2n + 1$, per qualche $n \in \mathbb{N}$. Questo è effettivamente la differenza di due quadrati successivi:

$$2n + 1 = (n + 1)^2 - n^2.$$

La dimostrazione è conclusa. □

1.1.6. Quantificatori. Due simboli che giocano un ruolo fondamentale in matematica sono i *quantificatori*:

$$\forall \quad \exists$$

I simboli sono una A ed una E rovesciate ed indicano le espressioni *per ogni* (*for All*) e *esiste* (*Exists*). I quantificatori sono essenziali nella formulazione dei teoremi, e più in generale di affermazioni matematiche che possono essere vere o false. Ad esempio, l'espressione

$$\forall x \in \mathbb{R} \exists y \in \mathbb{R} : 2y = x$$

dice che qualsiasi numero reale x può essere diviso per due, ed è un'affermazione vera. I due punti “:” sono sinonimo di “tale che”. D'altro canto, la stessa espressione con $\mathbb{Z}$ al posto di $\mathbb{R}$:

$$\forall x \in \mathbb{Z} \exists y \in \mathbb{Z} : 2y = x$$

è falsa perché se $x = 1$ non esiste nessun $y \in \mathbb{Z}$ tale che $2y = 1$.

Il simbolo $\exists!$ indica l'espressione *esiste ed è unico*. L'affermazione

$$\forall x \in \mathbb{R} \exists! y \in \mathbb{R} : 2y = x$$

continua ad essere vera: ogni numero reale è il doppio di un unico numero reale. Invece l'affermazione

$$\forall x \in \mathbb{R} : x > 0 \exists! y \in \mathbb{R} : y^2 = x$$

è falsa: ogni numero reale positivo x ha effettivamente una radice quadrata reale y , però questa non è unica perché le radici di x sono sempre due $\pm y$.

1.1.7. Dimostrazioni. Come facciamo a capire se una data affermazione matematica sia vera o falsa? Non c'è una regola generale, ma ci sono alcune indicazioni importanti: una affermazione del tipo

Per ogni x in un insieme A vale una certa proprietà P

può essere vera o falsa. Per dimostrare che è vera, serve una *dimostrazione* che mostri che la proprietà P è verificata da tutti gli elementi x nell'insieme A . Per dimostrare che è falsa, è invece sufficiente esibire un singolo $x \in A$ per cui P non sia soddisfatta.

Ad esempio, l'affermazione

$$\forall x \in \mathbb{N}, x^2 \geq 0$$

è vera, perché i quadrati sono sempre positivi. L'affermazione

$$\forall x \in \mathbb{N}, x^2 \geq 5$$

non è vera, perché non è soddisfatta ad esempio dal valore $x = 1$.

1.1.8. Prodotto cartesiano. Il *prodotto cartesiano* di due insiemi A e B è un nuovo insieme $A \times B$ i cui elementi sono tutte le coppie (a, b) dove a è un elemento qualsiasi di A e b è un elemento qualsiasi di B . Più brevemente:

$$A \times B = \{(a, b) \mid a \in A, b \in B\}.$$

Ad esempio, se $A = \{1, 2\}$ e $B = \{-1, 3, 4\}$, allora

$$A \times B = \{(1, -1), (1, 3), (1, 4), (2, -1), (2, 3), (2, 4)\}.$$

Più in generale, se A e B sono insiemi finiti con m e n elementi rispettivamente, allora il prodotto cartesiano $A \times B$ contiene mn elementi.

Il prodotto $A \times A$ è anche indicato con A^2 . Il caso $A = \mathbb{R}$ è particolarmente interessante perché ha una forte valenza geometrica. L'insieme $\mathbb{R}$ dei numeri reali può essere interpretato come una retta. Il prodotto cartesiano $\mathbb{R}^2 = \mathbb{R} \times \mathbb{R}$ è l'insieme

$$\mathbb{R}^2 = \{(x, y) \mid x \in \mathbb{R}, y \in \mathbb{R}\}$$

formato da tutte le coppie (x, y) di numeri reali. In altre parole, $\mathbb{R}^2$ non è nient'altro che il *piano cartesiano* già studiato alle superiori: un elemento di $\mathbb{R}^2$ è un punto identificato dalla coppia (x, y) .

Possiamo definire in modo analogo il prodotto di un numero arbitrario di insiemi. Il prodotto cartesiano di k insiemi $A_1, \dots, A_k$ è l'insieme

$$A_1 \times \dots \times A_k$$

i cui elementi sono le sequenze $(a_1, \dots, a_k)$ di k elementi in cui ciascun a_i è un elemento dell'insieme A_i , per ogni $i = 1, \dots, k$. Possiamo scrivere:

$$A_1 \times \dots \times A_k = \{(a_1, \dots, a_k) \mid a_i \in A_i, \forall i\}.$$

Se $k = 2$ ritroviamo il prodotto cartesiano di due insiemi già definito sopra. Se tutti gli insiemi coincidono, il prodotto $A \times \dots \times A$ può essere

indicato semplicemente come A^k . Come sopra, il caso geometricamente interessante è quello in cui $A = \mathbb{R}$ e quindi $\mathbb{R}^k$ è l'insieme

$$\mathbb{R}^k = \{(x_1, \dots, x_k) \mid x_i \in \mathbb{R} \forall i\}.$$

Ad esempio, se $k = 3$ otteniamo

$$\mathbb{R}^3 = \{(x, y, z) \mid x, y, z \in \mathbb{R}\}.$$

Questo insieme è l'analogo tridimensionale del piano cartesiano e può essere chiamato lo *spazio cartesiano*. Possiamo pensare ad ogni punto di $\mathbb{R}^3$ come ad un punto dello spazio con tre coordinate x, y, z .

1.1.9. Ragionamento per induzione. Uno degli strumenti più raffinati della matematica è il *ragionamento per induzione*, con il quale è possibile dimostrare in poche righe teoremi piuttosto complessi. Il ragionamento funziona nel modo seguente. Sia $P(n)$ una certa proposizione matematica che dipende da un numero naturale $n \geq 1$. Il nostro scopo è dimostrare che $P(n)$ è vera per ogni n . Per ottenere ciò, è sufficiente completare due passi:

- (1) Dimostrare la proposizione $P(1)$.
- (2) Per n generico, dare per buona $P(n-1)$ e dimostrare $P(n)$.

La proposizione $P(n-1)$ che viene data per buona e che quindi viene usata per dimostrare $P(n)$ è detta *ipotesi induttiva*. Il ragionamento è ben illustrato con un esempio.

Proposizione 1.1.4. *La somma dei primi n numeri naturali è*

$$1 + 2 + 3 + \dots + n = \frac{(n+1)n}{2}.$$

Dimostrazione. Dimostriamo l'uguaglianza per induzione su n . Il passo iniziale $n = 1$ è facile: l'uguaglianza da dimostrare è semplicemente $1 = 1$. Adesso supponiamo che l'uguaglianza sia vera per $n-1$ e la dimostriamo per n . Scriviamo:

$$1 + 2 + 3 + \dots + n = (1 + 2 + 3 + \dots + n-1) + n.$$

Usando l'ipotesi induttiva, sappiamo che

$$1 + 2 + 3 + \dots + n-1 = \frac{n(n-1)}{2}.$$

Il secondo membro dell'equazione precedente adesso diventa

$$1 + 2 + 3 + \dots + n = \frac{n(n-1)}{2} + n = \frac{n(n-1) + 2n}{2} = \frac{(n+1)n}{2}$$

ed abbiamo concluso. □



Figura 1.2. Una partizione di un insieme X è una suddivisione di X in sottoinsiemi disgiunti. Una partizione è di fatto equivalente a una relazione di equivalenza.

1.1.10. Relazione di equivalenza. Introduciamo adesso una nozione che verrà usata sporadicamente in questo libro. Sia X un insieme.

Definizione 1.1.5. Una *relazione di equivalenza* su X è una relazione $\sim$ che occorre fra alcune coppie di elementi di X e che soddisfa queste proprietà:

- (1) $x \sim x \quad \forall x \in X$ (riflessività)
- (2) se $x \sim y$ allora $y \sim x$, $\forall x, y \in X$ (simmetria)
- (3) se $x \sim y$ e $y \sim z$, allora $x \sim z$, $\forall x, y, z \in X$ (proprietà transitiva)

Diciamo che x e y sono *in relazione fra loro* se $x \sim y$.

Esempio 1.1.6. Fissiamo $n \in \mathbb{N}$. Ricordiamo che un numero intero $x \in \mathbb{Z}$ è *divisibile per n* se esiste un $k \in \mathbb{Z}$ per cui $kn = x$. Definiamo una relazione sull'insieme $\mathbb{Z}$ degli interi nel modo seguente:

$$x \sim y \iff x - y \text{ divisibile per } n.$$

Mostriamo che $\sim$ è una relazione di equivalenza verificando i tre assiomi:

- (1) $x \sim x$ perché $x - x = 0$ è divisibile per n ;
- (2) se $x \sim y$ allora $y \sim x$, infatti se $x - y$ è divisibile per n allora lo è anche il suo opposto $y - x$;
- (3) se $x \sim y$ e $y \sim z$, allora $x \sim z$, infatti se $x - y$ e $y - z$ sono divisibili per n lo è chiaramente anche la loro somma $x - y + y - z = x - z$.

Una *partizione* di un insieme X è una suddivisione di X in sottoinsiemi disgiunti, come nella Figura 1.2. Data una partizione, possiamo definire una semplice relazione di equivalenza su X nel modo seguente: $x \sim y \iff x$ e y appartengono allo stesso sottoinsieme della partizione.

Viceversa, data una relazione di equivalenza otteniamo una partizione in sottoinsiemi formati da elementi che sono in relazione fra loro. La definizione più formale è descritta sotto.

Osservazione 1.1.7. La nozione di relazione di equivalenza è molto più concreta di quanto sembri ed è spesso lo strumento usato per partizionare un insieme in sottoinsiemi. Ad esempio, la nozione di *specie biologica* è definita (semplificando molto) nel modo seguente. Sia X l'insieme di tutti gli animali sulla terra. Diciamo che due animali $x, y \in X$ sono in relazione $x \sim y$ se questi (o alcuni loro parenti stretti) possono accoppiarsi fra loro e generare dei figli fertili: due asini sono in relazione, ma un asino ed un cavallo no. Questa ovviamente non è una definizione matematica rigorosa, ma concretamente funziona bene nella maggior parte dei casi: con una buona approssimazione, la relazione $\sim$ è una relazione di equivalenza e da questa discende una partizione di X in sottoinsiemi come gatti, cani, ecc. Ciascun sottoinsieme è per definizione una *specie biologica*.

Definiamo formalmente come passare da una relazione d'equivalenza $\sim$ ad una partizione di X . In realtà useremo questa costruzione solo in un paio di punti in questo libro, quindi può essere saltata ad una prima lettura.

Per ogni $x \in X$ definiamo $U_x \subset X$ come l'insieme di tutti gli $y \in X$ tali che $x \sim y$. Un insieme del tipo $U_x \subset X$ è detto *classe di equivalenza*. Usando la proprietà transitiva si vede facilmente che tutti gli elementi di U_x sono in relazione fra loro, ma non sono mai in relazione con nessun elemento al di fuori di U_x . Se $x, y \in X$ possono accadere due casi:

- se $x \sim y$, allora $U_x = U_y$,
- se $x \not\sim y$, allora $U_x \cap U_y = \emptyset$.

Due classi di equivalenza distinte sono anche disgiunte e l'unione di tutte le classi di equivalenza è X . Quindi le classi di equivalenza formano una partizione di X .

Esempio 1.1.8. La relazione $\sim$ dell'Esempio 1.1.6 produce una partizione di $\mathbb{Z}$ in n classi di equivalenza $U_0, U_1, \dots, U_{n-1}$, dove

$$U_i = \{i + kn \mid k \in \mathbb{Z}\}.$$

L'insieme U_i è formato da tutti quei numeri interi che, se divisi per n , danno come resto i . In particolare U_0 è l'insieme dei numeri divisibili per n .

Ad esempio, per $n = 3$ otteniamo

$$U_0 = \{\dots, -3, 0, 3, 6, \dots\},$$

$$U_1 = \{\dots, -2, 1, 4, 7, \dots\},$$

$$U_2 = \{\dots, -1, 2, 5, 8, \dots\}.$$

Se $\sim$ è una relazione di equivalenza su X , l'*insieme quoziente* $X/\sim$ è l'insieme i cui elementi sono le classi di equivalenza di X . Ciascuna classe di equivalenza è adesso considerata un singolo elemento in $X/\sim$.

Esempio 1.1.9. Con la relazione $\sim$ dell'Esempio 1.1.6, l'insieme quoziente $\mathbb{Z}/\sim$ è un insieme formato da n elementi:

$$\mathbb{Z}/\sim = \{U_0, U_1, \dots, U_{n-1}\}.$$

Ciascun U_i è un sottoinsieme di $\mathbb{Z}$.

Esempio 1.1.10. Secondo l'assiomatica moderna, i numeri razionali $\mathbb{Q}$ sono costruiti a partire dagli interi $\mathbb{Z}$ nel modo seguente. Sia $\mathbb{Z}^* = \mathbb{Z} \setminus \{0\}$ l'insieme dei numeri interi non nulli. Consideriamo il prodotto cartesiano $\mathbb{Z} \times \mathbb{Z}^*$ formato dalle coppie (p, q) di interi con $q \neq 0$. Definiamo una relazione di equivalenza su $\mathbb{Z} \times \mathbb{Z}^*$ nel modo seguente:

$$(p, q) \sim (p', q') \iff pq' = p'q.$$

Si verifica facilmente che $\sim$ è una relazione di equivalenza. Definiamo infine $\mathbb{Q}$ come l'insieme quoziente:

$$\mathbb{Q} = (\mathbb{Z} \times \mathbb{Z}^*) / \sim.$$

Cosa c'entra questa definizione astratta con l'usuale insieme dei numeri razionali? Il collegamento è il seguente: interpretiamo una coppia (p, q) come una frazione $\frac{p}{q}$. La relazione di equivalenza è necessaria qui perché, come sappiamo tutti, in realtà lo stesso numero razionale può essere espresso come due frazioni differenti $\frac{p}{q}$ e $\frac{p'}{q'}$, ad esempio $\frac{2}{3}$ e $\frac{4}{6}$, e questo accade precisamente quando $pq' = p'q$, nell'esempio $2 \cdot 6 = 3 \cdot 4$.

Quindi alla domanda "che cos'è un numero razionale?" rispondiamo "è una classe di equivalenza di frazioni", cioè una classe di equivalenza di coppie (p, q) dove la relazione $\sim$ è quella descritta sopra.

1.2. Funzioni

Dopo aver richiamato la teoria degli insiemi e la notazione matematica, introduciamo alcuni fra gli oggetti più usati in matematica: le funzioni.

1.2.1. Definizione. Siano A e B due insiemi. Una *funzione* da A in B è una legge f che trasforma qualsiasi elemento x di A in un qualche elemento y di B . L'elemento y ottenuto da x tramite f è indicato come

$$y = f(x).$$

L'insieme di partenza A è detto *dominio* e l'insieme di arrivo B è detto *codominio*. Per indicare f useremo la nozione seguente:

$$f: A \longrightarrow B.$$

È importante ricordare che gli insiemi A e B sono parte integrante della funzione f , non è cioè possibile definire una funzione senza chiarire con precisione quali siano il suo dominio ed il suo codominio.

Ad esempio, possiamo prendere $A = B = \mathbb{R}$ e definire le funzioni

$$f(x) = x^2 - 1, \quad f(x) = \sin x.$$

Analogamente, possiamo definire $A = \{x \in \mathbb{R} \mid x \geq 0\}$ e $f: A \rightarrow \mathbb{R}$ come

$$f(x) = \sqrt{x}.$$

In questo caso è però fondamentale chiarire quale delle due radici di x stiamo considerando, perché $f(x)$ deve dipendere da x senza ambiguità. Generalmente si suppone che $f(x) = \sqrt{x}$ sia la radice positiva.

Osservazione 1.2.1. La nozione di funzione è molto generale e non si limita a considerare solo quelle funzioni che si possono scrivere esplicitamente usando le quattro operazioni $+, -, \times, :$ ed altre funzioni note come ad esempio quelle trigonometriche. Ad esempio, possiamo scegliere $A = B = \mathbb{N}$ e definire $f(n)$ come l' $(n+1)$ -esimo numero primo. Oppure scegliere $A = \mathbb{R}$, $B = \mathbb{N}$ e definire $f(x)$ come la 127-esima cifra di x nel suo sviluppo decimale. In entrambi i casi è impossibile (oppure molto difficile o inutile) scrivere f come una funzione algebrica esplicita. Dobbiamo rassegnarci al fatto che la maggior parte delle funzioni non sono scrivibili esplicitamente usando le usuali operazioni algebriche.

Due funzioni f e g sono considerate *uguali* se hanno lo stesso dominio e lo stesso codominio, e se $f(x) = g(x)$ per ogni elemento x del dominio. Quindi ad esempio le funzioni

$$f: \mathbb{R} \rightarrow \mathbb{R}, f(x) = x^2, \quad g: \mathbb{N} \rightarrow \mathbb{R}, g(x) = x^2$$

non sono uguali perché hanno domini diversi, mentre le funzioni

$$f: \mathbb{N} \rightarrow \mathbb{N}, f(x) = 1, \quad g: \mathbb{N} \rightarrow \mathbb{N}, g(x) = (x-1)^2 - x^2 + 2x$$

sono invece in realtà la stessa funzione e scriviamo $f = g$.

Un modo molto semplice per modificare una funzione data consiste nel restringere il dominio ad un sottoinsieme. Se $f: A \rightarrow B$ è una funzione e $A' \subset A$ è un sottoinsieme, la *restrizione* di f a A' è la funzione $f': A' \rightarrow B$ definita esattamente come f , ponendo cioè $f'(x) = f(x)$ per ogni $x \in A'$. La restrizione f' è indicata generalmente con il simbolo $f|_{A'}$.

Ad esempio, la restrizione di $f: \mathbb{R} \rightarrow \mathbb{R}, f(x) = x^2$ al sottoinsieme $\mathbb{N}$ è indicata con $f|_{\mathbb{N}}$ ed è ovviamente la funzione

$$f|_{\mathbb{N}}: \mathbb{N} \rightarrow \mathbb{R}, \quad f|_{\mathbb{N}}(x) = x^2.$$

1.2.2. Immagine e controimmagine. Il concetto di funzione è estremamente importante in matematica e nelle scienze, ed ha quindi diritto ad un vocabolario tutto suo, che viene purtroppo raramente introdotto nelle scuole e con cui il lettore deve acquistare familiarità.

Sia $f: A \rightarrow B$ una funzione. L'*immagine* di un elemento $x \in A$ è l'elemento $f(x)$ associato a x tramite f . Più in generale, se $C \subset A$ è un qualsiasi sottoinsieme, l'*immagine* di C è l'insieme

$$f(C) = \{f(x) \mid x \in C\} \subset B.$$

L'immagine $f(C)$ è un sottoinsieme del codominio B , ed è l'unione di tutte le immagini di tutti gli elementi di C . Ad esempio, se $f: \mathbb{R} \rightarrow \mathbb{R}, f(x) = x^2$ e $C = \mathbb{N}$, allora

$$f(\mathbb{N}) = \{0, 1, 4, 9, 16, \dots\}.$$

L'*immagine* della funzione f è per definizione l'immagine $f(A)$ dell'intero dominio A . Ad esempio l'immagine della funzione $f: \mathbb{R} \rightarrow \mathbb{R}, f(x) = x^2$ appena descritta è la semiretta $f(\mathbb{R}) = [0, +\infty)$ formata da tutti i numeri reali positivi o nulli.

Esiste un'altra nozione che è in un certo senso opposta a quella di immagine. Se $y \in B$ è un elemento del codominio, la sua *controimmagine* è il sottoinsieme $f^{-1}(y) \subset A$ del dominio che consiste di tutti gli elementi x la cui immagine è y , in altre parole:

$$f^{-1}(y) = \{x \in A \mid f(x) = y\}.$$

Più in generale, se $C \subset B$ è un sottoinsieme del codominio, la sua *controimmagine* è il sottoinsieme $f^{-1}(C) \subset A$ del dominio che consiste di tutti gli elementi la cui immagine è contenuta in C . In altre parole:

$$f^{-1}(C) = \{x \in A \mid f(x) \in C\}.$$

Esempio 1.2.2. Consideriamo la funzione $f: \mathbb{R} \rightarrow \mathbb{R}, f(x) = x^2$ e verifichiamo i fatti seguenti:

- la controimmagine di 4 è l'insieme con due elementi $f^{-1}(4) = \{-2, 2\}$;
- la controimmagine del segmento chiuso $C = [4, 9]$ è l'unione di due segmenti chiusi $f^{-1}(C) = [-3, -2] \cup [2, 3]$.
- la controimmagine del segmento chiuso $[-10, -5]$ è l'insieme vuoto $f^{-1}([-10, -5]) = \emptyset$.

Il lettore è invitato ad assimilare bene queste nozioni che compariranno in molti punti del libro. Raccomandiamo in particolare di non accettare acriticamente gli esempi, ma di verificarne la correttezza.

Esempio 1.2.3. Consideriamo la funzione $f: \mathbb{R} \rightarrow \mathbb{R}, f(x) = \sin x$. L'immagine di f è il segmento chiuso $[-1, 1]$. La controimmagine di $\mathbb{N}$ è l'insieme

$$f^{-1}(\mathbb{N}) = \left\{ \dots, -\frac{3\pi}{2}, -\pi, -\frac{\pi}{2}, 0, \frac{\pi}{2}, \pi, \frac{3\pi}{2}, \dots \right\} = \left\{ \frac{n\pi}{2} \mid n \in \mathbb{Z} \right\}.$$

Esercizio 1.2.4. Dimostra che vale sempre l'inclusione $C \subset f^{-1}(f(C))$. Costruisci un esempio in cui il contenimento $\subset$ è un'uguaglianza = ed un altro esempio in cui il contenimento è stretto $\subsetneq$.

1.2.3. Funzioni iniettive, suriettive e bigettive. Sia $f: A \rightarrow B$ una funzione. Definiamo alcune nozioni che saranno di fondamentale importanza nei capitoli successivi.

Definizione 1.2.5. La funzione f è

- *iniettiva* se due elementi distinti $x \neq x'$ del dominio hanno sempre immagini distinte. In altre parole se

$$\forall x, x' \in A, \quad x \neq x' \implies f(x) \neq f(x');$$

- *suriettiva* se ogni elemento del codominio è immagine di almeno un elemento del dominio, cioè

$$\forall y \in B, \quad \exists x \in A: f(x) = y.$$

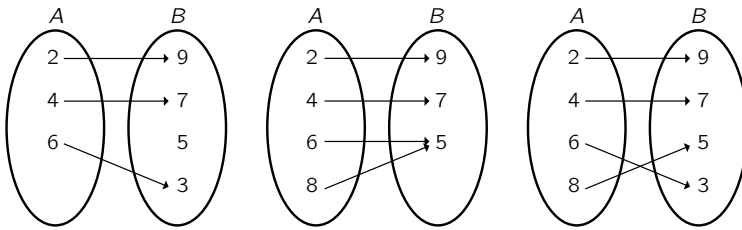


Figura 1.3. Una funzione iniettiva (sinistra), suriettiva (centro) e biettiva (destra).

In altre parole, se l'immagine di f è l'intero codominio B , cioè

$$f(A) = B.$$

Si veda la Figura 1.3. Qualche esempio:

Esempio 1.2.6. La funzione $f: \mathbb{N} \rightarrow \mathbb{N}, f(x) = x^2$ è iniettiva, perché i quadrati di due numeri naturali diversi sono sempre diversi; non è suriettiva, perché esistono dei numeri naturali (ad esempio il numero 2) che non sono quadrati e quindi $f(\mathbb{N}) \subsetneq \mathbb{N}$.

Esempio 1.2.7. La funzione $f: \mathbb{R} \rightarrow [0, +\infty), f(x) = x^2$ è suriettiva perché qualsiasi numero reale $y \geq 0$ ha una radice quadrata; non è iniettiva perché questa radice non è unica, infatti $f(-x) = f(x)$.

I due esempi precedenti ci mostrano ancora quanto sia importante la scelta del dominio e del codominio nello studio delle proprietà di una funzione.

Una funzione $f: A \rightarrow B$ che è sia iniettiva che suriettiva è detta *biettiva*; in questo caso si dice anche che f è una *bigezione* o una *corrispondenza biunivoca*. Se f è una bigezione, ogni elemento $y \in B$ ha *esattamente* una controimmagine: ne ha almeno una perché è suriettiva, e non può averne più di una perché è iniettiva. Questo è un ragionamento cruciale che il lettore deve assimilare bene prima di proseguire.

Se $f: A \rightarrow B$ è una bigezione, possiamo definire la *funzione inversa*

$$f^{-1}: B \longrightarrow A$$

nel modo seguente: per ogni $y \in B$, l'elemento $f^{-1}(y)$ è proprio quell'unico x tale che $f(x) = y$. Nella Figura 1.3 a destra, la funzione inversa f^{-1} è ottenuta da f semplicemente invertendo il verso delle frecce. La funzione inversa f^{-1} è anch'essa una bigezione.

1.2.4. Composizione di funzioni. È utile pensare ad una funzione $f: A \rightarrow B$ come ad una trasformazione che prende come *input* un qualsiasi elemento $x \in A$ e restituisce come *output* la sua immagine $f(x) \in B$. L'elemento $f(x)$ può essere a sua volta trasformato usando un'altra funzione

$g: B \rightarrow C$ e questa successione di due trasformazioni è detta *composizione di funzioni*.

Più formalmente, date due funzioni

$$f: A \longrightarrow B, \quad g: B \longrightarrow C$$

definiamo la loro *composizione* h come una nuova funzione

$$h: A \longrightarrow C$$

nel modo seguente:

$$h(x) = g(f(x)).$$

La funzione composta h è indicata come

$$h = g \circ f.$$

Si usa quindi il simbolo $\circ$ per indicare l'operazione di composizione di due funzioni. Notiamo che due funzioni f e g possono essere composte solo se il codominio di f è uguale al dominio di g . Notiamo anche che la composizione di due funzioni va letta da destra a sinistra: la funzione $h = g \circ f$ è ottenuta applicando prima f e poi g .

Esempio 1.2.8. Componendo le funzioni $f: \mathbb{R} \rightarrow \mathbb{R}, f(x) = x^3$ e $g: \mathbb{R} \rightarrow \mathbb{R}, g(y) = \sin y$ otteniamo $h = g \circ f, h(x) = \sin x^3$. Notiamo che la composizione generalmente non è commutativa: se scambiamo l'ordine delle funzioni otteniamo un'altra funzione $j = f \circ g, j(x) = (\sin x)^3$.

Proposizione 1.2.9. *Siano $f: A \rightarrow B$ e $g: B \rightarrow C$ due funzioni.*

- (1) *Se f e g sono iniettive, allora anche $g \circ f$ è iniettiva.*
- (2) *Se f e g sono suriettive, allora anche $g \circ f$ è suriettiva.*
- (3) *Se f e g sono bigettive, allora anche $g \circ f$ è bigettiva e la sua inversa è la funzione $f^{-1} \circ g^{-1}$.*

Dimostrazione. (1). Supponiamo che f e g siano iniettive e dimostriamo che anche $g \circ f$ lo è. Dobbiamo quindi mostrare la proposizione

$$x \neq x' \implies g(f(x)) \neq g(f(x')) \quad \forall x, x' \in A.$$

Siccome f è iniettiva, $x \neq x' \implies f(x) \neq f(x')$. Siccome g è iniettiva, $f(x) \neq f(x') \implies g(f(x)) \neq g(f(x'))$. La dimostrazione è conclusa.

(2). Supponiamo che f e g siano suriettive e dimostriamo che anche $g \circ f$ lo è. Dato $z \in C$ qualsiasi, dobbiamo mostrare che esiste un $x \in A$ tale che $g(f(x)) = z$. Poiché g è suriettiva, esiste un $y \in B$ tale che $g(y) = z$. Poiché f è suriettiva, esiste un $x \in A$ tale che $f(x) = y$. Quindi $g(f(x)) = z$.

(3). La composizione $g \circ f$ è bigettiva per i punti (1) e (2) già dimostrati. L'inversa è $f^{-1} \circ g^{-1}$ perché per ogni $x \in A$ otteniamo

$$f^{-1}(g^{-1}(g(f(x)))) = f^{-1}(f(x)) = x.$$

La dimostrazione è conclusa. □

Esercizio 1.2.10. Siano $f: A \rightarrow B$ e $g: B \rightarrow C$ due funzioni.

- (1) Se $g \circ f$ è iniettiva, allora f è iniettiva.
- (2) Se $g \circ f$ è suriettiva, allora g è suriettiva.

1.2.5. Permutazioni. Sia X un insieme di n elementi. Una *permutazione* è una bigezione

$$\sigma: X \longrightarrow X.$$

Proposizione 1.2.11. *Ci sono $n!$ possibili permutazioni per X .*

Dimostrazione. Costruiamo una permutazione σ definendo in ordine le immagini $\sigma(1), \sigma(2), \dots, \sigma(n)$. L'immagine $\sigma(1)$ è un qualsiasi elemento di X e quindi abbiamo n possibilità; successivamente, $\sigma(2)$ è un qualsiasi elemento di X diverso da $\sigma(1)$, e abbiamo $n-1$ possibilità; andando avanti in questo modo possiamo costruire σ in $n \cdot (n-1) \cdots 2 \cdot 1 = n!$ modi differenti. $\square$

Notiamo che l'inversa σ^{-1} di una permutazione σ è sempre una permutazione, e la composizione $\sigma \circ \tau$ di due permutazioni σ e τ è una permutazione. La permutazione *identità* è la permutazione id che fissa ciascun elemento, cioè $\text{id}(x) = x \forall x \in X$.

A meno di rinominare gli elementi di X , possiamo supporre per semplicità che X sia l'insieme formato dai numeri naturali da 1 a n :

$$X = \{1, \dots, n\}.$$

Vediamo come possiamo scrivere e studiare una permutazione σ . Un modo consiste nello scrivere la tabella

$$\begin{bmatrix} 1 & 2 & \cdots & n \\ \sigma(1) & \sigma(2) & \cdots & \sigma(n) \end{bmatrix}$$

oppure più semplicemente $[\sigma(1) \ \sigma(2) \ \cdots \ \sigma(n)]$.

Un altro metodo consiste nello scrivere σ come prodotto di *cicli*. Se $a_1, \dots, a_k$ sono elementi distinti di X , il *ciclo*

$$(a_1 \ \cdots \ a_k)$$

indica la permutazione che trasla ciclicamente gli elementi $a_1, \dots, a_k$ e lascia fissi tutti gli altri, cioè tale che:

$$\begin{aligned} \sigma(a_1) &= a_2, & \sigma(a_2) &= a_3, & \dots & \sigma(a_{k-1}) &= a_k, & \sigma(a_k) &= a_1, \\ \sigma(a) &= a, & \forall a &\notin \{a_1, \dots, a_k\}. \end{aligned}$$

Due cicli $(a_1 \dots a_k)$ e $(b_1 \dots b_h)$ sono *indipendenti* se $a_i \neq b_j$ per ogni i, j . Ogni permutazione si scrive come prodotto (cioè composizione) di cicli indipendenti. Ad esempio:

$$\begin{bmatrix} 1 & 2 & 3 & 4 & 5 & 6 \\ 4 & 5 & 3 & 6 & 2 & 1 \end{bmatrix}$$

si scrive come prodotto di cicli

$$(1 \ 4 \ 6)(2 \ 5)(3) = (1 \ 4 \ 6)(2 \ 5).$$

Il prodotto di cicli come composizione va inteso da destra a sinistra, come le funzioni. I cicli di ordine uno possono chiaramente essere rimossi. L'ordine dei cicli indipendenti non è importante: $(1\ 4\ 6)(2\ 5)$ e $(2\ 5)(1\ 4\ 6)$ sono la stessa permutazione. Inoltre il ciclo $(1\ 4\ 6)$ può essere scritto anche come $(4\ 6\ 1)$ oppure $(6\ 1\ 4)$, ma *non* come $(1\ 6\ 4)$.

La notazione come prodotto di cicli ha il pregio di funzionare molto bene con le operazioni di inversione e composizione. Per scrivere l'inversa di una permutazione basta invertire i cicli. L'inversa della permutazione $(1\ 4\ 6)(2\ 5)$ è $(5\ 2)(6\ 4\ 1)$. Per comporre due permutazioni basta affiancare i cicli da destra a sinistra: se componiamo $(2\ 4\ 1)(3\ 5\ 6)$ e $(1\ 4\ 6)(2\ 5)$ otteniamo

$$(1\ 4\ 6)(2\ 5)(2\ 4\ 1)(3\ 5\ 6) = (1\ 5)(6\ 3\ 2).$$

Infatti leggendo da destra a sinistra vediamo che $4 \rightarrow 1 \rightarrow 4$ sta fisso mentre $1 \rightarrow 2 \rightarrow 5$ e $5 \rightarrow 6 \rightarrow 1$.

1.2.6. Segno di una permutazione. Vediamo adesso che ciascuna permutazione σ ha un *segno* che può essere 1 oppure -1 . Questo fatto sarà usato solo nella definizione del determinante nella Sezione 3.3 e quindi può essere saltato ad una prima lettura.

Un ciclo di ordine 2 è chiamato una *trasposizione*. Una trasposizione è una permutazione che scambia due elementi e lascia fissi tutti gli altri. Si dimostra facilmente che ciascun ciclo è prodotto di trasposizioni, infatti

$$(a_1 \dots a_k) = (a_1\ a_k)(a_1\ a_{k-1}) \cdots (a_1\ a_2).$$

Conseguentemente, ogni permutazione può essere ottenuta come prodotto di un certo numero di trasposizioni. Facciamo un esempio concreto: prendiamo 4 assi da un mazzo di carte e li poniamo sul tavolo uno dopo l'altro, secondo la successione CQFP (cuori quadri fiori picche). Adesso permutiamo le carte in modo da ottenere la successione QFPC. Possiamo ottenere questa successione come composizione di trasposizioni, ad esempio in questo modo:

$$\text{CQFP} \rightarrow \text{QCFP} \rightarrow \text{QFCP} \rightarrow \text{QFPC}$$

Ci sono anche altri modi, ad esempio:

$$\text{CQFP} \rightarrow \text{CQPF} \rightarrow \text{FQPC} \rightarrow \text{QFPC}$$

oppure

$$\text{CQFP} \rightarrow \text{FQCP} \rightarrow \text{FPCQ} \rightarrow \text{PFCQ} \rightarrow \text{QFCP} \rightarrow \text{QFPC}$$

Abbiamo trasformato CQFP in QFPC in tre modi diversi, con un numero di trasposizioni che è 3 nei primi due casi e 5 nel terzo. Notiamo che questo numero è dispari in tutti e tre gli esempi. Questo non è un caso: la proposizione seguente ci dice in particolare che non è possibile trasformare CQFP in QFPC con un numero pari di trasposizioni.

Proposizione 1.2.12. Se una permutazione σ si scrive in due modi diversi come prodotto di n e m trasposizioni, il numero $n - m$ è sempre pari (cioè n e m sono entrambi dispari o entrambi pari).

Dimostrazione. Abbiamo per ipotesi

$$\begin{aligned}\sigma &= (a_1 \ a_2)(a_3 \ a_4) \cdots (a_{2n-1} \ a_{2n}), \\ \sigma &= (b_1 \ b_2)(b_3 \ b_4) \cdots (b_{2m-1} \ b_{2m}).\end{aligned}$$

Quindi

$$\text{id} = \sigma \circ \sigma^{-1} = (a_1 \ a_2) \cdots (a_{2n-1} \ a_{2n})(b_{2m} \ b_{2m-1}) \cdots (b_2 \ b_1).$$

Dimostriamo che l'identità id non è realizzabile come prodotto di un numero dispari $h = 2k + 1$ di trasposizioni: questo implica che $m + n$ è pari e quindi $m - n$ è pari. Procediamo per induzione su $k \geq 0$. Se $k = 0$, è ovvio che id non è realizzabile come una singola trasposizione. Supponiamo il fatto dimostrato per $k - 1$ e passiamo a $k \geq 1$. Supponiamo che

$$\text{id} = (a_1 \ a_2) \cdots (a_{2h-1} \ a_{2h})$$

con $h = 2k + 1$. Siccome la permutazione è l'identità, il termine a_1 deve comparire, oltre che nella prima trasposizione a sinistra, almeno una seconda volta (letto da sinistra a destra):

$$\text{id} = (a_1 \ a_2) \cdots (a_1 \ a_i) \cdots (a_{2h-1} \ a_{2h}).$$

Notiamo le seguenti uguaglianze, in cui lettere diverse indicano numeri diversi:

$$(c \ d)(a \ b) = (a \ b)(c \ d), \quad (b \ c)(a \ b) = (a \ c)(b \ c).$$

Possiamo usare queste "mosse" per modificare due trasposizioni successive senza cambiare il numero totale $h = 2k + 1$ di trasposizioni. Usando le mosse con $a = a_1$, possiamo spostare il secondo a_1 verso sinistra di un passo alla volta, finché non arriva in seconda posizione e otteniamo:

$$\text{id} = (a_1 \ a_2)(a_1 \ a_3) \cdots (a_{2h-1} \ a_{2h}).$$

Se $a_2 = a_3$, abbiamo $(a_1 \ a_2)(a_1 \ a_2) = \text{id}$ e possiamo cancellare le prime due trasposizioni. Troviamo una successione di $2(k - 1) + 1$ elementi e giungiamo ad un assurdo per l'ipotesi induttiva.

Se $a_2 \neq a_3$ possiamo sostituire $(a_1 \ a_2)(a_1 \ a_3)$ con $(a_1 \ a_3)(a_2 \ a_3)$. In questo modo la successione di trasposizioni contiene un a_1 in meno di prima e ripartiamo da capo. Siccome prima o poi gli a_1 finiscono, ad un certo punto ricadremo nel caso precedente $a_2 = a_3$. $\square$

Definiamo il *segno* $\text{sgn}(\sigma)$ di una permutazione σ come

$$\text{sgn}(\sigma) = (-1)^n$$

dove σ si decompone in n trasposizioni. Il segno è 1 oppure -1 ed è ben definito grazie alla Proposizione 1.2.12.

Un ciclo di ordine n ha segno $(-1)^{n-1}$. Una trasposizione ha segno -1 . L'identità ha segno 1.



Figura 1.4. Una configurazione irrisolvibile del gioco del 15.

Osservazione 1.2.13. Usando il segno di una permutazione si può dimostrare che la configurazione mostrata nella Figura 1.4 del gioco del 15 non è risolvibile (cosa ben nota a molti bambini nati nell'era pre-digitale). Si procede in questo modo. La prima cosa da notare è che per trasformare la configurazione in figura in quella giusta dobbiamo fare un numero pari di mosse: questo è dovuto al fatto che, se pensiamo ad una scacchiera 4×4 con caselle bianche e nere, la casella vuota è inizialmente in una casella nera (in basso a destra), e ad ogni passaggio salta da una casella nera ad una bianca e viceversa. Siccome alla fine deve tornare su una nera (sempre in basso a destra), deve fare un numero pari di salti.

Ogni configurazione del gioco può essere descritta come una permutazione σ dell'insieme $X = \{1, \dots, 15, 16\}$ dove 16 indica in realtà la casella vuota. La casella i è occupata dal tassello $\sigma(i)$. Scopo del gioco è ottenere la permutazione identità id , in cui la casella i è occupata dal tassello $i \ \forall i \in X$.

La configurazione σ mostrata nella Figura 1.4 è una trasposizione (14 15) e quindi ha segno negativo $\text{sgn}(\sigma) = -1$. Ad ogni mossa che facciamo componiamo σ con una trasposizione e quindi cambiamo segno alla σ . Dopo un numero pari di mosse otterremo sempre una permutazione con segno negativo, e quindi mai l'identità. Se partiamo dalla configurazione mostrata in figura, il gioco del 15 non si può risolvere.

Il segno cambia in modo controllato per inversione e composizioni:

Proposizione 1.2.14. *Per ogni $\sigma, \tau \in S_n$ abbiamo:*

$$\text{sgn}(\sigma^{-1}) = \text{sgn}(\sigma), \quad \text{sgn}(\sigma \circ \tau) = \text{sgn}(\sigma) \cdot \text{sgn}(\tau).$$

Dimostrazione. Se σ e τ sono scritte come prodotto di m e n trasposizioni, possiamo scrivere σ^{-1} e $\sigma \circ \tau$ come prodotto di m e $m+n$ trasposizioni. $\square$

1.3. Polinomi

I polinomi sono tipi particolarmente semplici di funzioni ottenute combinando numeri e variabili ed usando solo le operazioni $+$, $-$ e $\times$.

1.3.1. Definizione. Ricordiamo che un *monomio* è una espressione algebrica che ha una parte numerica (il coefficiente) ed una parte letterale; ad esempio questi sono monomi:

$$4x, \quad -2xy, \quad \sqrt{5}x^3.$$

Il *grado* di un monomio è la somma degli esponenti presenti sulle parti letterali: i tre monomi descritti sopra hanno grado 1, 2 e 3. Un monomio di grado zero è semplicemente un numero.

Un *polinomio* è una somma di monomi, ad esempio:

$$7 + 3x^2 - \sqrt{2}y^3.$$

Un polinomio è *ridotto in forma normale* se è scritto come somma di monomi con parti letterali differenti e coefficienti non nulli, oppure è il polinomio 0. Per ridurre un polinomio in forma normale è sufficiente raccogliere i monomi con la stessa parte letterale e quindi eliminare quelli con coefficiente nullo. Il *grado* di un polinomio scritto in forma normale è il massimo grado dei suoi monomi. I polinomi possono essere sommati e moltiplicati fra loro nel modo usuale.

Ci interessano particolarmente i polinomi in cui compare una sola variabile x . Un polinomio di questo tipo viene indicato con $p(x)$ oppure più semplicemente con p . Un polinomio $p(x)$ con una sola variabile può essere sempre descritto ordinando i suoi monomi da quello di grado più alto a quello di grado più basso: otteniamo quindi una scrittura del tipo

$$p(x) = a_n x^n + \cdots + a_1 x + a_0$$

dove n è il grado di $p(x)$ e $a_n \neq 0$. Ad esempio:

$$x^3 - 2x + 5, \quad 4x^2 - 7.$$

Il coefficiente a_0 è detto *termine noto* del polinomio. Un polinomio di grado zero è semplicemente un numero a_0 .

1.3.2. Divisione con resto fra polinomi. I polinomi assomigliano ai numeri interi: possono essere sommati, moltiplicati, e si possono fare le divisioni con resto.

Prendiamo due numeri interi, ad esempio 26 e 11. Se dividiamo 26 per 11 otteniamo come quoziente 2 e come resto 4. In altre parole, otteniamo:

$$26 = 2 \cdot 11 + 4.$$

Notiamo che il resto 4 è ovviamente sempre più piccolo del divisore 11. Analogamente, dati due polinomi $p(x)$ (il *dividendo*) e $d(x)$ (il *divisore*), esistono sempre (e sono unici) due polinomi $q(x)$ (il *quoziente*) e $r(x)$ (il *resto*) per cui

$$p(x) = q(x)d(x) + r(x)$$

con la proprietà che il resto $r(x)$ abbia grado strettamente minore del divisore $d(x)$. Le divisioni fra polinomi si risolvono con carta e penna esattamente con la stessa procedura usata per i numeri interi. Ad esempio, se dividiamo $p(x) = x^3 + 1$ per $d(x) = x^2 - 1$ otteniamo

$$x^3 + 1 = x(x^2 - 1) + (x + 1)$$

e la divisione ha come quoziente x e come resto $x + 1$.

Diciamo che il numero intero 7 divide 14 ma non divide 15, perché la divisione di 14 per 7 ha resto nullo, mentre la divisione di 15 per 7 invece ha un certo resto. Usiamo la stessa terminologia per i polinomi: se la divisione fra due polinomi $p(x)$ e $d(x)$ ha resto nullo, allora $p(x) = q(x)d(x)$ per qualche quoziente $q(x)$ e diciamo che $d(x)$ divide $p(x)$. Possiamo usare la barra verticale $|$ come sinonimo di "divide" e scrivere ad esempio:

$$9 \mid 18, \quad (x + 1) \mid (x^3 + 1).$$

Notiamo che effettivamente $(x^3 + 1) = (x^2 - x + 1)(x + 1)$.

1.3.3. Radici di un polinomio. Ricordiamo adesso una delle definizioni più importanti dell'algebra. Se $p(x)$ è un polinomio e a è un numero, indichiamo con $p(a)$ il numero che otteniamo sostituendo a al posto di x . Ad esempio, se $p(x) = x^2 - 3$, allora $p(-2) = 4 - 3 = 1$.

Definizione 1.3.1. Un numero a è *radice* di un polinomio $p(x)$ se

$$p(a) = 0.$$

Ad esempio, il numero -1 è radice del polinomio $p(x) = x^3 + 1$ perché $p(-1) = 0$. La determinazione delle radici di un polinomio è uno dei problemi più classici dell'algebra. A questo scopo enunciamo un criterio.

Proposizione 1.3.2. Il numero a è radice di $p(x)$ se e solo se

$$(x - a) \mid p(x).$$

Dimostrazione. Se dividiamo $p(x)$ per $(x - a)$, otteniamo

$$p(x) = q(x)(x - a) + r(x)$$

dove $q(x)$ è il quoziente e $r(x)$ il resto. Sappiamo che il grado di $r(x)$ è strettamente minore di quello di $(x - a)$, che è uno: quindi $r(x)$ ha grado zero, in altre parole è una costante che scriviamo semplicemente come r_0 . Quindi

$$p(x) = q(x)(x - a) + r_0.$$

Se sostituiamo a al posto di x , otteniamo

$$p(a) = q(a)(a - a) + r_0 = 0 + r_0 = r_0.$$

Quindi a è radice di $p(x)$ se e solo se $r_0 = 0$. D'altra parte $r_0 = 0$ se e solo se $(x - a)$ divide $p(x)$ e quindi concludiamo. $\square$

Introduciamo un'altra definizione che useremo spesso in questo libro.

Definizione 1.3.3. La *molteplicità* di una radice a di un polinomio $p(x)$ è il massimo numero k tale che $(x - a)^k$ divide $p(x)$.

Informalmente, la molteplicità misura “quante volte” a è radice di $p(x)$.

Esempio 1.3.4. Il polinomio $x^3 - 1$ ha la radice 1 con molteplicità 1 perché

$$x^3 - 1 = (x - 1)(x^2 + x + 1)$$

e $(x - 1)$ non divide $x^2 + x + 1$, semplicemente perché 1 non è radice di $x^2 + x + 1$. Analogamente il polinomio

$$x^3 - 2x^2 + x = (x - 1)^2 x$$

ha la radice 1 con molteplicità 2 e la radice 0 con molteplicità 1.

Osserviamo un fatto semplice: se moltiplichiamo un polinomio $p(x)$ per una costante k diversa da zero, otteniamo un altro polinomio $q(x) = kp(x)$ che ha le stesse radici di $p(x)$ con le stesse molteplicità. Per questo motivo, quando studiamo le radici di un polinomio di grado n del tipo

$$p(x) = a_n x^n + a_{n-1} x^{n-1} + \cdots + a_1 x + a_0$$

possiamo dividerlo per $a_n \neq 0$; in questo modo ciascun coefficiente a_i si trasformerà in un nuovo coefficiente $b_i = a_i/a_n$ e otterremo un nuovo polinomio

$$q(x) = x^n + b_{n-1} x^{n-1} + \cdots + b_1 x + b_0$$

che ha il pregio di avere il primo coefficiente pari a 1. Un polinomio di questo tipo è detto *monico*.

Proposizione 1.3.5. Sia $p(x) = q_1(x)q_2(x)$. Le radici di $p(x)$ contate con molteplicità sono l'unione di quelle di $q_1(x)$ e di $q_2(x)$.

Dimostrazione. Se a ha molteplicità m_1 in $q_1(x)$ e m_2 in $q_2(x)$, allora

$$q_1(x) = (x - a)^{m_1} d_1(x), \quad q_2(x) = (x - a)^{m_2} d_2(x)$$

con $d_1(a) \neq 0$ e $d_2(a) \neq 0$. Quindi

$$p(x) = (x - a)^{m_1+m_2} d_1(x) d_2(x)$$

con $d_1(a) d_2(a) \neq 0$ e deduciamo che a ha molteplicità $m_1 + m_2$ in $p(x)$. $\square$

Esempio 1.3.6. I polinomi

$$q_1(x) = x^2 - 2x + 1, \quad q_2(x) = x^2 - 1$$

possono essere scritti come

$$q_1(x) = (x - 1)^2, \quad q_2(x) = (x + 1)(x - 1).$$

Questi hanno rispettivamente la radice 1 con molteplicità 2 e le radici $-1, 1$ entrambe con molteplicità 1. Il prodotto

$$p(x) = q_1(x)q_2(x) = (x - 1)^3(x + 1)$$

ha la radice 1 con molteplicità 3 e la radice -1 con molteplicità 1.

Teorema 1.3.7. *Un polinomio $p(x)$ di grado $n \geq 1$ ha al più n radici, contate con molteplicità.*

Dimostrazione. Dimostriamo il teorema per induzione su n . Per quanto appena visto, a meno di dividere tutto per il primo coefficiente possiamo supporre che il polinomio $p(x)$ sia monico. Se $n = 1$, il polinomio è del tipo $p(x) = x + a_0$ ed ha chiaramente una sola radice $-a_0$. Quindi la tesi è soddisfatta.

Supponiamo la tesi vera per $n-1$ e la dimostriamo per n . Se $p(x)$ non ha radici, siamo a posto. Se ha almeno una radice a , allora per la Proposizione 1.3.2 possiamo dividere $p(x)$ per $x - a$ e ottenere un altro polinomio $q(x)$, cioè vale $p(x) = (x - a)q(x)$. Il polinomio $q(x)$ ha grado $n - 1$ e quindi per ipotesi induttiva ha al più $n - 1$ radici contate con molteplicità. Per la Proposizione 1.3.5, le radici di $p(x)$ contate con molteplicità sono esattamente quelle di $q(x)$ più a . Quindi $p(x)$ ha al più $n - 1 + 1 = n$ radici e abbiamo concluso. $\square$

Un polinomio di grado 1 è sempre del tipo $p(x) = ax + b$ ed ha quindi sempre una sola radice $x = -b/a$. Un polinomio di grado 2 è del tipo

$$p(x) = ax^2 + bx + c$$

e come sappiamo bene dalle scuole superiori le sue radici dipendono da

$$\Delta = b^2 - 4ac$$

nel modo seguente.

Proposizione 1.3.8. *Si possono presentare i casi seguenti:*

- Se $\Delta > 0$, il polinomio $p(x)$ ha due radici distinte

$$x_{\pm} = \frac{-b \pm \sqrt{\Delta}}{2a}$$

entrambe di molteplicità uno.

- Se $\Delta = 0$, il polinomio $p(x)$ ha una sola radice

$$x = -\frac{b}{2a}$$

con molteplicità due.

- Se $\Delta < 0$, il polinomio $p(x)$ non ha radici reali.

Dimostrazione. Possiamo riscrivere il polinomio $p(x)$ in questo modo

$$p(x) = a \left(x + \frac{b}{2a} \right)^2 + c - \frac{b^2}{4a} = a \left(x + \frac{b}{2a} \right)^2 - \frac{\Delta}{4a}.$$

Se $\Delta > 0$, è sufficiente sostituire $x_{\pm}$ in $p(x)$ per verificare che sono radici:

$$p(x_{\pm}) = a \left(\frac{-b \pm \sqrt{\Delta}}{2a} + \frac{b}{2a} \right)^2 - \frac{\Delta}{4a} = a \frac{\Delta}{4a^2} - \frac{\Delta}{4a} = 0.$$

Se $\Delta = 0$, otteniamo

$$p(x) = a \left(x + \frac{b}{2a} \right)^2$$

e quindi la radice $-b/(2a)$ ha chiaramente molteplicità due. Se $\Delta < 0$,

$$p(x) = a \left(x + \frac{b}{2a} \right)^2 - \frac{\Delta}{4a}$$

è sempre strettamente positivo o strettamente negativo per ogni x , a seconda che a sia positivo o negativo. Quindi non può mai essere nullo, in altre parole $p(x)$ non ha radici. $\square$

Come sapevamo dalle superiori, esistono polinomi di secondo grado senza radici reali. Come vedremo nella prossima sezione, è possibile ovviare a questo “problema” aggiungendo all’insieme $\mathbb{R}$ altri numeri, così da formare il più vasto insieme $\mathbb{C}$ dei *numeri complessi*.

1.4. Numeri complessi

I numeri complessi sono un ampliamento dell’insieme dei numeri reali $\mathbb{R}$, costruito con lo scopo di ottenere migliori proprietà algebriche.

1.4.1. Definizione. Un *numero complesso* è un oggetto algebrico che si scrive nel modo seguente:

$$a + bi$$

dove a e b sono numeri reali arbitrari e i è un nuovo simbolo chiamato *unità immaginaria*. Ad esempio, questi sono numeri complessi:

$$\sqrt{7}, \quad 2 + i, \quad 23i, \quad 4 - i, \quad -1 + \pi i.$$

I numeri complessi si sommano e si moltiplicano nel modo usuale, tenendo a mente un’unica nuova relazione:

$$i^2 = -1.$$

La somma e la moltiplicazione di due numeri complessi $a + bi$ e $c + di$ si svolge quindi nel modo seguente:

$$(a + bi) + (c + di) = a + c + (b + d)i,$$

$$(a + bi) \cdot (c + di) = ac + bci + adi + bdi^2 = ac - bd + (ad + bc)i.$$

Nel prodotto abbiamo usato che $i^2 = -1$. Ad esempio:

$$(7 + i) \cdot (4 - i) = 29 - 3i.$$

L’insieme dei numeri complessi è indicato con il simbolo $\mathbb{C}$. Abbiamo esteso la nostra sequenza di insiemi numerici:

$$\mathbb{N} \subsetneq \mathbb{Z} \subsetneq \mathbb{Q} \subsetneq \mathbb{R} \subsetneq \mathbb{C}.$$

1.4.2. Coniugio, norma e inverso. Sia $z = a + bi$ un numero complesso. I numeri a e b sono detti rispettivamente la *parte reale* e la *parte immaginaria* di z . Il numero z è reale, cioè appartiene al sottoinsieme $\mathbb{R} \subset \mathbb{C}$, se e solo se la sua parte immaginaria è nulla.

Il *coniugio* di $z = a + bi$ è il numero complesso

$$\bar{z} = a - bi$$

ottenuto da z cambiando il segno della sua parte immaginaria. Notiamo che $z = \bar{z}$ precisamente quando $b = 0$, cioè quando z è reale. Scriviamo quindi

$$z \in \mathbb{R} \iff z = \bar{z}.$$

Il *modulo* di $z = a + bi$ è il numero reale

$$|z| = \sqrt{a^2 + b^2}.$$

Il modulo $|z|$ è nullo quando $z = 0$, cioè quando $a = b = 0$, ed è strettamente positivo se $z \neq 0$. Notiamo inoltre che

$$z \cdot \bar{z} = (a + bi) \cdot (a - bi) = a^2 + b^2 = |z|^2.$$

Mostriamo adesso un fatto non banale: come nei numeri razionali e reali, ogni numero complesso $z \neq 0$ ha un *inverso* z^{-1} rispetto all'operazione di moltiplicazione, dato da

$$z^{-1} = \frac{\bar{z}}{|z|^2}.$$

Infatti se moltiplichiamo z e z^{-1} otteniamo

$$z \cdot z^{-1} = \frac{z\bar{z}}{|z|^2} = \frac{|z|^2}{|z|^2} = 1.$$

Esempio 1.4.1. L'inverso di i è $-i$, infatti $i \cdot (-i) = 1$. L'inverso di $2 + i$ è

$$(2 + i)^{-1} = \frac{2 - i}{|2 + i|^2} = \frac{2 - i}{5}.$$

Si verifica che effettivamente

$$(2 + i) \cdot \frac{2 - i}{5} = 1.$$

1.4.3. Il piano complesso. Mentre i numeri reali $\mathbb{R}$ formano una retta, i numeri complessi $\mathbb{C}$ formano un piano detto *piano complesso*. Ogni numero complesso $a + bi$ può essere identificato con il punto di coordinate (a, b) nel piano cartesiano o equivalentemente come un vettore applicato nell'origine 0 e diretto verso (a, b) , come illustrato nella Figura 1.5-(sinistra).

Gli assi delle ascisse e delle ordinate sono rispettivamente l'*asse reale* e l'*asse immaginaria* di $\mathbb{C}$. L'asse reale è precisamente il sottoinsieme $\mathbb{R} \subset \mathbb{C}$ formato dai numeri reali. L'asse immaginaria consiste di tutti i numeri complessi del tipo bi al variare di $b \in \mathbb{R}$.

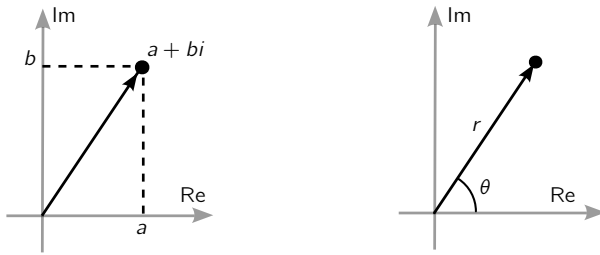


Figura 1.5. Il piano complesso (sinistra) e le coordinate polari (destra).

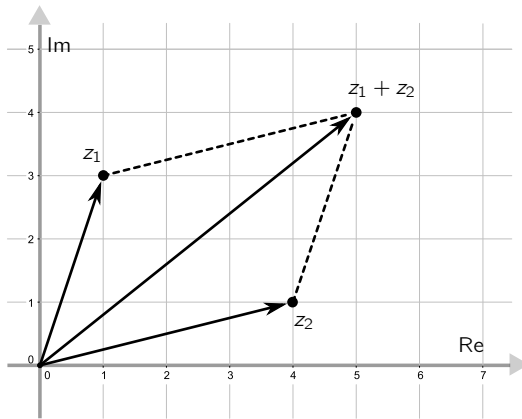


Figura 1.6. La somma $z_1 + z_2$ di due numeri complessi z_1 e z_2 può essere calcolata con la regola del parallelogramma. Qui $z_1 = 1 + 3i$, $z_2 = 4 + i$ e quindi $z_1 + z_2 = 5 + 4i$.

La somma $z_1 + z_2$ di due numeri complessi z_1 e z_2 viene calcolata interpretando z_1 e z_2 come vettori e sommandoli quindi con l'usuale regola del parallelogramma, come mostrato nella Figura 1.6.

Il prodotto $z_1 \cdot z_2$ di due numeri complessi è apparentemente più complicato, ma può essere visualizzato agevolmente usando le *coordinate polari*, che ora richiamiamo.

1.4.4. Coordinate polari. Come ricordato nella Figura 1.5-(destra), un punto (x, y) diverso dall'origine del piano cartesiano può essere identificato usando la lunghezza r del vettore corrispondente e l'angolo θ formato dal vettore con l'asse reale. Le *coordinate polari* del punto sono la coppia (r, θ) . Per passare dalle coordinate polari (r, θ) a quelle cartesiane (x, y)

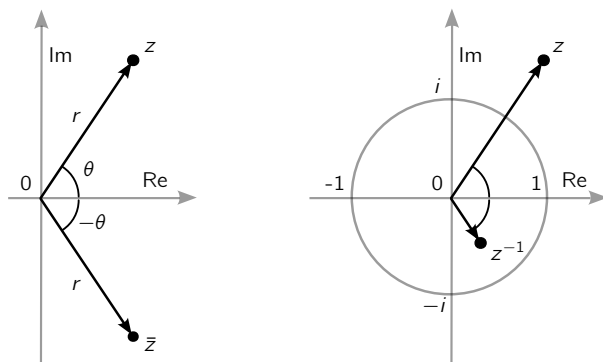


Figura 1.7. Il coniugio $\bar{z}$ di z si ottiene specchiando z lungo l'asse reale (sinistra). L'inverso z^{-1} di z ha argomento $-\theta$ opposto a quello θ di z e modulo $|z^{-1}|$ inverso rispetto a $|z|$. La circonferenza unitaria è mostrata in figura (destra).

basta usare le formule

$$x = r \cos \theta, \quad y = r \sin \theta.$$

Viceversa,

$$r = \sqrt{x^2 + y^2}, \quad \theta = \arccos \frac{x}{\sqrt{x^2 + y^2}}.$$

Tornando ai numeri complessi, un numero $z = x + yi$ può essere scritto in coordinate polari come

$$z = x + yi = r \cos \theta + (r \sin \theta)i = r(\cos \theta + i \sin \theta).$$

Notiamo che

$$|z| = \sqrt{x^2 + y^2} = r.$$

Il modulo di z è quindi la lunghezza del vettore che descrive z . Il coniugio $\bar{z} = a - ib$ è il punto ottenuto cambiando il segno della coordinata immaginaria: geometricamente questo corrisponde a riflettere il punto rispetto all'asse reale. In coordinate polari, questo corrisponde a cambiare θ in $-\theta$ lasciando fisso r . Si veda la Figura 1.7-(sinistra).

Tornando alle coordinate polari, è comodo scrivere

$$e^{i\theta} = \cos \theta + i \sin \theta.$$

In questo modo ogni numero complesso $z \neq 0$ si scrive come

$$z = r e^{i\theta}.$$

Il numero $r = |z|$ è il modulo di z mentre l'angolo θ è detto *fase* o *argomento* di z .

Il motivo profondo per cui introduciamo inaspettatamente qui la costante e di Nepero è dovuto alle rappresentazioni delle funzioni e^x , $\sin x$ e

$\cos x$ come serie di potenze. Giustificare questa scelta ci porterebbe troppo lontano; per noi è sufficiente considerare questa misteriosa esponenziale complessa $e^{i\theta}$ come un simbolo che vuol dire semplicemente $\cos \theta + i \sin \theta$. Si tratta di una simbologia azzeccata, perché $e^{i\theta}$ ha le proprietà usuali dell'esponenziale:

Proposizione 1.4.2. *Vale la relazione*

$$e^{i(\theta+\varphi)} = e^{i\theta} \cdot e^{i\varphi}.$$

Dimostrazione. Otteniamo

$$\begin{aligned} e^{i(\theta+\varphi)} &= \cos(\theta + \varphi) + i \sin(\theta + \varphi) \\ &= \cos \theta \cos \varphi - \sin \theta \sin \varphi + i(\sin \theta \cos \varphi + \cos \theta \sin \varphi) \\ &= (\cos \theta + i \sin \theta) \cdot (\cos \varphi + i \sin \varphi) \\ &= e^{i\theta} \cdot e^{i\varphi}. \end{aligned}$$

La dimostrazione è completa. □

Adesso capiamo perché le coordinate polari sono particolarmente utili quando moltiplichiamo due numeri complessi. Se

$$z_1 = r_1 e^{i\theta_1}, \quad z_2 = r_2 e^{i\theta_2}$$

allora il loro prodotto è semplicemente

$$z_1 z_2 = r_1 r_2 e^{i(\theta_1 + \theta_2)}.$$

In altre parole:

*Quando si fa il prodotto di due numeri complessi,
i moduli si moltiplicano e gli argomenti si sommano.*

Notiamo in particolare che se $z = r e^{i\theta} \neq 0$, il suo inverso è

$$z^{-1} = r^{-1} e^{-i\theta}.$$

L'inverso z^{-1} ha argomento $-\theta$ opposto a quello θ di z e ha modulo $|z^{-1}| = r^{-1}$ inverso rispetto a $|z| = r$, si veda la Figura 1.7-(destra).

I numeri complessi $e^{i\theta}$ al variare di θ sono precisamente i punti che stanno sulla circonferenza unitaria, determinati dall'angolo θ . In particolare per $\theta = \pi$ otteniamo la celebre *identità di Eulero*:

$$e^{i\pi} = -1.$$

Notiamo infine che due numeri complessi non nulli espressi in forma polare

$$r_0 e^{i\theta_0}, \quad r_1 e^{i\theta_1}$$

sono lo stesso numero complesso se e solo se valgono entrambi questi fatti:

- $r_0 = r_1$,
- $\theta_1 = \theta_0 + 2k\pi$ per qualche $k \in \mathbb{Z}$.

1.4.5. Proprietà dei numeri complessi. I numeri complessi hanno numerose proprietà. Queste si dimostrano facilmente: in presenza di un prodotto, è spesso utile usare la rappresentazione polare. Lasciamo la dimostrazione di queste proprietà per esercizio.

Esercizio 1.4.3. Valgono i fatti seguenti per ogni $z, w \in \mathbb{C}$:

$$|z + w| \leq |z| + |w|, \quad |zw| = |z||w|, \quad |z^{-1}| = \frac{1}{|z|},$$

$$|z| = |\bar{z}|, \quad \overline{z + w} = \bar{z} + \bar{w}, \quad \overline{zw} = \bar{z}\bar{w}.$$

1.4.6. Radici n -esime di un numero complesso. Sia $z_0 = r_0 e^{i\theta_0}$ un numero complesso fissato diverso da zero. Mostriamo come risolvere l'equazione

$$z^n = z_0$$

usando i numeri complessi. In altre parole, determiniamo tutte le radici n -esime di z_0 . Scriviamo la variabile z in forma polare come $z = r e^{i\theta}$. L'equazione adesso diventa

$$r^n e^{in\theta} = r_0 e^{i\theta_0}.$$

L'equazione è soddisfatta precisamente se valgono entrambi questi fatti:

- $r = \sqrt[n]{r_0}$,
- $n\theta = \theta_0 + 2k\pi$ per qualche $k \in \mathbb{Z}$.

La seconda condizione può essere riscritta richiedendo che

$$\theta = \frac{\theta_0}{n} + \frac{2k\pi}{n}$$

per qualche $k \in \mathbb{Z}$. Otteniamo quindi precisamente n soluzioni distinte: per i valori $k = 0, 1, \dots, n-1$ otteniamo gli argomenti

$$\theta = \frac{\theta_0}{n}, \frac{\theta_0}{n} + \frac{2\pi}{n}, \dots, \frac{\theta_0}{n} + \frac{2(n-1)\pi}{n}.$$

Le n soluzioni dell'equazione $z^n = z_0$ hanno tutte lo stesso modulo $\sqrt[n]{r_0}$ e argomenti che variano in una sequenza di angoli separati da un passo costante $\frac{2\pi}{n}$. Geometricamente, questo significa che le soluzioni formano i vertici di un poligono regolare centrato nell'origine con n lati e raggio $\sqrt[n]{r_0}$. Si veda la Figura 1.8.

Esempio 1.4.4. L'equazione $z^n = 1$ ha come soluzioni

$$z = e^{i\frac{2k\pi}{n}}$$

dove $k = 0, 1, \dots, n-1$. Queste n soluzioni sono i vertici di un poligono regolare di raggio 1 con n lati, avente 1 come vertice. Questi numeri complessi sono le *radici n -esime dell'unità*.

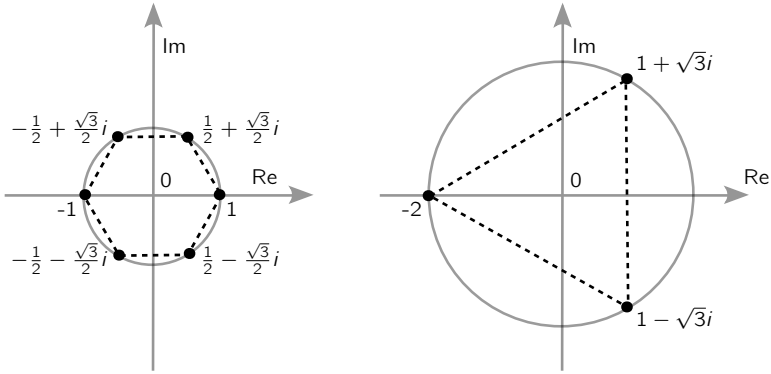


Figura 1.8. Le radici seste di 1, cioè le soluzioni di $z^6 = 1$, sono i vertici di un esagono regolare centrato nell'origine con un vertice in 1 (sinistra). Le soluzioni di $z^3 = -8$ sono i vertici di un triangolo equilatero centrato nell'origine e di raggio $\sqrt[3]{8} = 2$ (destra).

Esempio 1.4.5. Come mostrato nella Figura 1.8-(destra), le tre soluzioni dell'equazione $z^3 = -8$ hanno modulo $\sqrt[3]{8} = 2$ e argomento $\frac{\pi}{3}, \pi$ e $\frac{5\pi}{3}$. Si tratta dei numeri complessi

$$\begin{aligned} z_1 &= 2e^{i\frac{\pi}{3}} = 2\left(\cos\frac{\pi}{3} + i\sin\frac{\pi}{3}\right) = 1 + \sqrt{3}i, \\ z_2 &= 2e^{i\pi} = -2, \\ z_3 &= 2e^{i\frac{5\pi}{3}} = 2\left(\cos\frac{5\pi}{3} + i\sin\frac{5\pi}{3}\right) = 1 - \sqrt{3}i. \end{aligned}$$

Esercizio 1.4.6. Calcola le soluzioni dell'equazione $z^4 = i$.

1.4.7. Teorema fondamentale dell'algebra. Veniamo infine al vero motivo per cui abbiamo introdotto i numeri complessi.

Teorema 1.4.7 (Teorema fondamentale dell'algebra). *Ogni polinomio a coefficienti complessi ha almeno una radice complessa.*

Esistono varie dimostrazioni di questo teorema, che possono essere agevolmente trovate in rete. Quelle più accessibili usano degli strumenti analitici che si discostano da quanto trattato in questo libro e quindi le omettiamo.

Corollario 1.4.8. *Un polinomio $p(x)$ a coefficienti complessi di grado n ha esattamente n radici, contate con molteplicità.*

Dimostrazione. Dimostriamo il corollario per induzione su n . Possiamo supporre senza perdita di generalità che $p(x)$ sia monico. Se $n = 1$ allora $p(x) = x - a_0$ ha una sola radice $x = a_0$.

Dimostriamo il caso generico n dando per buono il caso $n-1$. Sappiamo per il teorema fondamentale dell'algebra che $p(x)$ ha almeno una radice a . Quindi per la Proposizione 1.3.2 possiamo scrivere $p(x) = (x-a)q(x)$ dove $q(x)$ è un altro polinomio di grado $n-1$. Per l'ipotesi induttiva $q(x)$ ha $n-1$ radici contate con molteplicità, e quindi $p(x)$ ha queste $n-1$ radici più a , quindi $p(x)$ ha esattamente n radici (sempre contate con molteplicità). $\square$

Per un polinomio di secondo grado

$$p(x) = ax^2 + bx + c$$

le due radici complesse si trovano usando la nota formula

$$x_{\pm} = \frac{-b \pm \sqrt{\Delta}}{2a}.$$

Qui $\pm\sqrt{\Delta}$ indica le due radici quadrate complesse di Δ , che coincidono solo se $\Delta = 0$.

Esempio 1.4.9. Le radici del polinomio $x^2 + 1$ sono $\pm i$. Le radici del polinomio $x^2 + (1-i)x - i$ sono

$$x_{\pm} = \frac{-1+i \pm \sqrt{2}i}{2} \implies x_{\pm} = \frac{-1+i \pm (1+i)}{2} \implies x_+ = i, x_- = -1.$$

Scriviamo il Corollario 1.4.8 in un'altra forma:

Corollario 1.4.10. *Ogni polinomio $p(x)$ a coefficienti complessi si spezza come prodotto di polinomi di primo grado:*

$$p(x) = a_n(x - z_1) \cdots (x - z_n)$$

dove a_n è il coefficiente più grande di $p(x)$ e $z_1, \dots, z_n$ sono le radici complesse di $p(x)$ contate con molteplicità.

1.4.8. Polinomi a coefficienti reali. Sappiamo che un polinomio $p(x)$ di grado n ha esattamente n soluzioni complesse contate con molteplicità. Se $p(x)$ ha coefficienti reali, possiamo dire qualcosa di più sulle sue radici complesse.

Proposizione 1.4.11. *Sia $p(x)$ un polinomio a coefficienti reali. Se z è una radice complessa di $p(x)$, allora $\bar{z}$ è anch'essa radice di $p(x)$.*

Dimostrazione. Il polinomio è

$$p(x) = a_n x^n + \cdots + a_1 x + a_0$$

e per ipotesi i coefficienti $a_n, \dots, a_0$ sono tutti reali. Se z è radice, allora

$$p(z) = a_n z^n + \cdots + a_1 z + a_0 = 0.$$

Applicando il coniugio ad entrambi i membri e l'Esercizio 1.4.3 troviamo

$$\overline{a_n z^n} + \cdots + \overline{a_1 z} + \overline{a_0} = \bar{0} = 0.$$

Siccome i coefficienti sono reali, il coniugio di a_i è sempre a_i e quindi

$$a_n \bar{z}^n + \cdots + a_1 \bar{z} + a_0 = 0.$$

In altre parole, anche $\bar{z}$ è radice di $p(x)$. □

Possiamo dedurre un teorema di spezzamento per i polinomi a coefficienti reali che preveda una scomposizione in fattori di grado 2 e 1. I fattori di grado 2 hanno radici complesse z e $\bar{z}$ coniugate fra loro, quelli di grado 1 hanno radici reali.

Corollario 1.4.12. *Ogni polinomio $p(x)$ a coefficienti reali si scompone nel modo seguente:*

$$p(x) = q_1(x) \cdots q_k(x) \cdot (x - x_1) \cdots (x - x_h)$$

dove $q_1(x), \dots, q_k(x)$ sono polinomi di grado due a coefficienti reali con $\Delta < 0$, e $x_1, \dots, x_h$ sono le radici reali di $p(x)$ contante con molteplicità.

Dimostrazione. Ragioniamo come sempre per induzione sul grado n di $p(x)$. Se $p(x)$ ha grado 1 allora $p(x) = x - x_1$ e siamo a posto.

Se $p(x)$ ha grado n , ha qualche radice complessa z . Se z è reale, scriviamo $z = x_1$, spezziamo $p(x) = (x - x_1)q(x)$ e concludiamo per induzione su $q(x)$. Se z è complesso, allora sappiamo che anche $\bar{z}$ è soluzione. Quindi $p(x) = (x - z)(x - \bar{z})q(x)$. Se $z = a + bi$ allora

$$\begin{aligned} q_1(x) &= (x - z)(x - \bar{z}) = (x - (a + bi))(x - (a - bi)) \\ &= x^2 - 2ax + a^2 + b^2 \end{aligned}$$

è un polinomio a coefficienti reali con $\Delta = -4b^2 < 0$. Scriviamo $p(x) = q_1(x)q(x)$ e concludiamo per induzione su $q(x)$. □

Abbiamo capito che le radici complesse non reali di un polinomio $p(x)$ a coefficienti reali sono presenti a coppie coniugate $z, \bar{z}$. In particolare, sono in numero pari. Ne deduciamo il fatto seguente.

Proposizione 1.4.13. *Un polinomio $p(x)$ a coefficienti reali di grado dispari ha sempre almeno una soluzione reale.*

Dimostrazione. Sappiamo che $p(x)$ ha grado n dispari e ha n soluzioni complesse contate con molteplicità. Di queste, un numero pari non sono reali. Quindi restano un numero dispari di soluzioni reali – quindi almeno una c'è. □

Osservazione 1.4.14. Esiste un'altra dimostrazione di questa proposizione che usa l'analisi. Siccome $p(x)$ ha grado dispari, i limiti $\lim_{x \rightarrow +\infty} p(x)$ e $\lim_{x \rightarrow -\infty} p(x)$ sono entrambi infiniti, ma con segni opposti. Quindi la funzione $p(x)$ assume valori sia positivi che negativi. Dal teorema di esistenza degli zeri segue che la funzione $p(x)$ assume anche il valore nullo per qualche $x_0 \in \mathbb{R}$. Quindi x_0 è radice di $p(x)$.

1.5. Strutture algebriche

Quando abbiamo introdotto i numeri complessi, abbiamo notato che questi hanno diverse proprietà algebriche simili a quelle dei numeri reali. In questa sezione esplicitiamo queste proprietà e diamo un nome agli insiemi dotati di operazioni che le soddisfano.

1.5.1. Gruppi. Un *gruppo* è un insieme G dotato di una *operazione binaria*, cioè di una funzione che associa ad ogni coppia a, b di elementi in G un nuovo elemento di G che indichiamo con $a * b$. Il simbolo $*$ indica l'operazione binaria. L'operazione deve soddisfare i seguenti tre assiomi:

- (1) $\exists e \in G : e * a = a * e = a \forall a \in G$ (esistenza dell'*elemento neutro* e),
- (2) $a * (b * c) = (a * b) * c \forall a, b, c \in G$ (*proprietà associativa*),
- (3) $\forall a \in G, \exists a' \in G : a * a' = a' * a = e$ (esistenza dell'*inverso*).

Ad esempio, l'insieme $\mathbb{Z}$ dei numeri interi con l'operazione di somma è un gruppo. Infatti:

- (1) esiste l'elemento neutro 0, per cui $0 + a = a + 0 = a \forall a \in \mathbb{Z}$,
- (2) vale la proprietà associativa $a + (b + c) = (a + b) + c \forall a, b, c \in \mathbb{Z}$,
- (3) ogni numero intero $a \in \mathbb{Z}$ ha un inverso $a' = -a$, per cui $a + (-a) = (-a) + a = 0$.

Notiamo invece che *non* sono gruppi:

- l'insieme $\mathbb{Z}$ con la moltiplicazione, perché non è soddisfatto (3),
- l'insieme $\mathbb{N}$ con la somma, sempre perché non è soddisfatto (3).

Il gruppo G è *commutativo* se vale anche la *proprietà commutativa*

$$a * b = b * a \forall a, b \in G.$$

I numeri interi $\mathbb{Z}$ formano un gruppo commutativo. Formano un gruppo commutativo anche gli insiemi numerici $\mathbb{Q}$, $\mathbb{R}$ e $\mathbb{C}$ con l'operazione di somma. Un altro esempio di gruppo è l'insieme

$$S_n$$

formato da tutte le $n!$ permutazioni dell'insieme $X = \{1, \dots, n\}$, con l'operazione $\circ$ di composizione, si veda la Sezione 1.2.5. Sappiamo infatti che:

- (1) esiste l'elemento neutro $\text{id} \in S_n$, la permutazione identità $\text{id}(i) = i$,
- (2) vale la proprietà associativa $\rho \circ (\sigma \circ \tau) = (\rho \circ \sigma) \circ \tau, \forall \rho, \sigma, \tau \in S_n$,
- (3) ogni permutazione σ ha una inversa σ^{-1} .

Il gruppo S_n delle permutazioni è detto *gruppo simmetrico*. A differenza dei gruppi numerici $\mathbb{Z}$, $\mathbb{Q}$, $\mathbb{R}$ e $\mathbb{C}$, notiamo che S_n contiene un numero finito di elementi e non è commutativo se $n \geq 3$: è facile trovare permutazioni che non commutano, ad esempio le trasposizioni $(1\ 2)$ e $(2\ 3)$ non

commutano:

$$(2 \ 3 \ 1) = (1 \ 2)(2 \ 3) \neq (2 \ 3)(1 \ 2) = (1 \ 3 \ 2).$$

In algebra si dimostrano vari teoremi sui gruppi. Il primo fatto da notare è che in un gruppo l'inverso di un qualsiasi elemento $a \in G$ è sempre unico.

Proposizione 1.5.1. Se $\exists a, a', a'' \in G$ tali che

$$a * a' = a' * a = e, \quad a * a'' = a'' * a = e,$$

allora $a' = a''$.

Dimostrazione. Troviamo

$$a' = a' * e = a' * (a * a'') = (a' * a) * a'' = e * a'' = a''.$$

La dimostrazione è conclusa. $\square$

Un altro fatto importante è che con i gruppi si può *semplificare*. Se troviamo una espressione del tipo

$$a * b = a * c$$

con $a, b, c \in G$, allora possiamo moltiplicare per l'inverso a^{-1} a sinistra in entrambi i membri e ottenere

$$a^{-1} * a * b = a^{-1} * a * c \implies b = c.$$

Notiamo che abbiamo usato l'esistenza dell'inverso e la proprietà associativa.

1.5.2. Anelli. Negli insiemi numerici $\mathbb{Z}, \mathbb{Q}, \mathbb{R}$ e $\mathbb{C}$ ci sono in realtà due operazioni $+$ e $\times$. Introduciamo adesso una struttura algebrica che prevede la coesistenza di due operazioni binarie.

Un *anello* è un insieme A dotato di due operazioni binarie $+$ e $\times$ che soddisfano questi assiomi:

- (1) A è un gruppo commutativo con l'operazione $+$,
- (2) $a \times (b \times c) = (a \times b) \times c \ \forall a, b, c \in A$ (*proprietà associativa del $\times$*),
- (3) $a \times (b + c) = (a \times b) + (a \times c); (b + c) \times a = (b \times a) + (c \times a), \ \forall a, b, c \in A$ (*proprietà distributiva*),
- (4) $\exists 1 \in A : 1 \times a = a \times 1 = a \ \forall a \in A$ (*elemento neutro per il $\times$*).

Gli interi $\mathbb{Z}$ con le operazioni di somma e prodotto formano un anello. L'elemento neutro della somma è lo 0, mentre quello del prodotto è 1. Anche gli insiemi $\mathbb{Q}, \mathbb{R}$ e $\mathbb{C}$ sono un anello con le operazioni di somma e prodotto. Un anello è *commutativo* se vale la proprietà commutativa anche per il prodotto:

$$a \times b = b \times a \ \forall a, b \in A.$$

Gli anelli $\mathbb{Z}, \mathbb{Q}, \mathbb{R}$ e $\mathbb{C}$ sono tutti commutativi. In un anello A gli elementi neutri per le operazioni $+$ e $\times$ sono generalmente indicati con 0 e 1. Dagli assiomi di anello possiamo subito dedurre qualche teorema. Ad esempio:

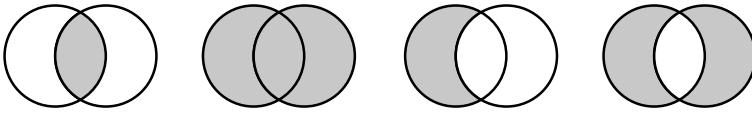


Figura 1.9. Intersezione, unione, differenza e differenza simmetrica di due insiemi.

Proposizione 1.5.2. *Vale $0 \times a = a \times 0 = 0$ per ogni $a \in A$.*

Dimostrazione. Abbiamo

$$0 \times a = (0 + 0) \times a = 0 \times a + 0 \times a.$$

Semplificando deduciamo che $0 = 0 \times a$. Dimostrazione analoga per $a \times 0$. $\square$

1.5.3. Campi. Definiamo infine un'ultima struttura algebrica, che è forse quella più importante in questo libro. Un *campo* è un anello commutativo K in cui vale anche il seguente assioma:

$$\forall a \in K, a \neq 0, \exists a' \in K : a \times a' = a' \times a = 1 \text{ (esiste inverso per il prodotto)}$$

Chiediamo quindi che ogni elemento a diverso da zero abbia anche un inverso rispetto al prodotto. Esempi di campi sono $\mathbb{Q}$, $\mathbb{R}$ e $\mathbb{C}$. Notiamo invece che $\mathbb{Z}$ *non* è un campo, perché ovviamente 2 non ha un inverso per il prodotto dentro l'insieme $\mathbb{Z}$. Riassumendo:

- $\mathbb{N}$ non è un gruppo.
- $\mathbb{Z}$ è un anello ma non è un campo.
- $\mathbb{Q}$, $\mathbb{R}$ e $\mathbb{C}$ sono campi.

In modo analogo alla Proposizione 1.5.1 dimostriamo che, in un qualsiasi anello A , se un elemento $a \in A$ ha un inverso per il prodotto, questo è unico. L'inverso per il prodotto di a è generalmente indicato con a^{-1} .

Esercizi

Esercizio 1.1. Sia $n \in \mathbb{N}$. Dimostra che $\sqrt{n}$ è razionale $\iff n$ è il quadrato di un numero naturale.

Esercizio 1.2. Dimostra che esistono infiniti numeri primi, nel modo seguente. Se per assurdo ne esistesse solo un numero finito $p_1, \dots, p_n$, allora mostra che il numero $p_1 \cdots p_n + 1$ non è divisibile per nessuno di loro.

Esercizio 1.3. La *differenza simmetrica* di due insiemi A e B è l'insieme

$$A \Delta B = (A \setminus B) \cup (B \setminus A).$$

Si veda la Figura 1.9. Mostra che Δ è commutativa e associativa:

$$A \Delta B = B \Delta A, \quad (A \Delta B) \Delta C = A \Delta (B \Delta C).$$

Esercizio 1.4. Dato un insieme X , l'*insieme delle parti* $\mathcal{P}(X)$ è l'insieme formato da tutti i sottoinsiemi di X , incluso il vuoto $\emptyset$ e X stesso. Ad esempio, se $X = \{1, 2\}$,

$$\mathcal{P}(X) = \{\emptyset, \{1\}, \{2\}, \{1, 2\}\}.$$

Mostra che se X contiene n elementi allora $\mathcal{P}(X)$ ne contiene 2^n .

Esercizio 1.5. Mostra per induzione l'uguaglianza seguente per ogni $n \geq 1$:

$$\sum_{k=1}^n k^2 = \frac{n(n+1)(2n+1)}{6}.$$

Esercizio 1.6. Mostra per induzione che qualsiasi numero del tipo $n^3 + 2n$ con $n \in \mathbb{N}$ è divisibile per 3.

Esercizio 1.7. Dimostra per induzione su n che n rette a coppie non parallele dividono il piano in $\frac{(n+1)n}{2} + 1$ regioni differenti.

Esercizio 1.8. Determina l'unica fra le seguenti che non è una relazione di equivalenza su $\mathbb{R}$:

- $x \sim y \iff x - y \in \mathbb{Z}$,
- $x \sim y \iff x^3 - 4x = y^3 - 4y$,
- $x \sim y \iff x \neq y + 1$.

Esercizio 1.9. Sia $f: A \rightarrow B$ una funzione e $S, T \subset A$ due sottoinsiemi. Dimostra:

$$f(S \cup T) = f(S) \cup f(T),$$

$$f(S \cap T) \subset f(S) \cap f(T).$$

Costruisci un esempio in cui l'inclusione $\subset$ è stretta e un altro esempio in cui è un'uguaglianza.

Esercizio 1.10. Sia $f: A \rightarrow B$ una funzione e $S, T \subset B$ due sottoinsiemi. Dimostra:

$$f^{-1}(S \cup T) = f^{-1}(S) \cup f^{-1}(T),$$

$$f^{-1}(S \cap T) = f^{-1}(S) \cap f^{-1}(T).$$

Esercizio 1.11. Sia X un insieme finito. Mostra che una funzione $f: X \rightarrow X$ è iniettiva $\iff$ è suriettiva. Mostra che questo fatto non è necessariamente vero se X è infinito.

Esercizio 1.12. Sia $p(x)$ un polinomio e $p'(x)$ la sua derivata. Sia a una radice di $p(x)$ con molteplicità $m \geq 1$. Mostra che a è radice di $p'(x)$ se e solo se $m \geq 2$, ed in questo caso a è radice di $p'(x)$ con molteplicità $m - 1$.

Esercizio 1.13. Determina tutte le radici del polinomio complesso $z^4 = -16$.

Esercizio 1.14. Determina tutti i numeri complessi z tali che $z^4 = \bar{z}^3$.

Esercizio 1.15. Sia $n \in \mathbb{N}$, $n \geq 1$ fissato. Considera l'insieme

$$C_n = \{0, 1, \dots, n-1\}$$

e definisci la seguente operazione binaria $*$ su C_n : per ogni $x, y \in C_n$, l'elemento $x * y \in C_n$ è il resto della divisione di $x + y$ per n . Mostra che C_n con questa operazione $*$ è un gruppo commutativo con elemento neutro 0. Il gruppo C_n si chiama il *gruppo ciclico* di ordine n . L'operazione $*$ è generalmente chiamata *somma* ed indicata con il simbolo $+$.

Esercizio 1.16. Sia $n \geq 2$. Sul gruppo ciclico C_n definito nell'esercizio precedente definiamo analogamente il *prodotto* $x \cdot y$ come il resto della divisione di xy per n . Mostra che C_n con le operazioni di somma e prodotto appena definite è un anello commutativo. Per quali valori di n secondo te C_n è un campo?

Esercizio 1.17. Considera l'insieme di numeri complessi

$$S = \{z \in \mathbb{C} \mid |z| = 1\}.$$

Mostra che S è un gruppo con l'operazione di prodotto. Chi è l'elemento neutro?

Esercizio 1.18. Un elemento $a \in A$ in un anello A è *invertibile* se esiste un inverso a^{-1} per il prodotto, cioè un elemento tale che $a \times a^{-1} = a^{-1} \times a = 1$. Dimostra che il prodotto ab di due elementi invertibili è anch'esso invertibile, e che $(ab)^{-1} = b^{-1}a^{-1}$.

Complementi

1.1. Infiniti numerabili e non numerabili. Nella matematica basata sulla teoria degli insiemi, gli infiniti non sono tutti uguali: ci sono infiniti più grandi di altri. Descriviamo questo fenomeno in questa sezione, concentrandoci soprattutto sugli insiemi numerici $\mathbb{N}, \mathbb{Z}, \mathbb{Q}$ e $\mathbb{R}$. Questi insiemi sono tutti infiniti, ma mentre i primi tre hanno tutti lo stesso tipo di infinito, l'ultimo è in un certo senso più grande dei precedenti.

Tutto parte come sempre dalla teoria degli insiemi e dalle funzioni.

Definizione 1.5.3. Un insieme infinito X è detto *numerabile* se esiste una bigezione fra $\mathbb{N}$ e X .

Gli insiemi numerabili sono tutti in bigezione con $\mathbb{N}$ e quindi anche in bigezione fra loro: possiamo dire che contengono "lo stesso numero di elementi", anche se questo numero è chiaramente infinito.

Concretamente, scrivere una bigezione $f: \mathbb{N} \rightarrow X$ equivale a rappresentare gli elementi di X come una successione infinita

$$X = \{f(0), f(1), f(2), \dots\}.$$

Ad esempio, il sottoinsieme $P \subset \mathbb{N}$ formato dai numeri pari è numerabile perché possiamo rappresentare gli elementi di P come una successione infinita:

$$P = \{0, 2, 4, 6, 8, \dots\}$$

La bigezione $f: \mathbb{N} \rightarrow P$ è definita semplicemente chiedendo che $f(n)$ sia l' n -esimo elemento della successione. In questo caso $f: \mathbb{N} \rightarrow P$ può essere scritta esplicitamente come $f(n) = 2n$.

Nonostante P sia contenuto strettamente in $\mathbb{N}$, esiste una bigezione fra i due insiemi! Questa è una proprietà che possono avere solo gli insiemi infiniti: se X è finito e $Y \subsetneq X$ è strettamente contenuto in X , chiaramente non può mai esistere una bigezione fra X e Y .

Un altro esempio di insieme numerabile è l'insieme $\mathbb{Z}$ dei numeri interi. Infatti possiamo metterli in successione nel modo seguente:

$$\mathbb{Z} = \{0, 1, -1, 2, -2, 3, -3, 4, \dots\}$$

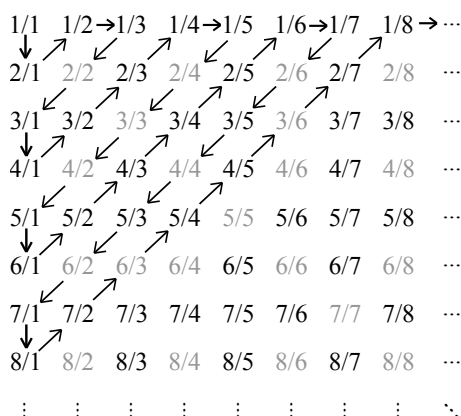


Figura 1.10. I numeri razionali $\mathbb{Q}$ sono un insieme numerabile. Per rappresentare l'insieme dei numeri razionali positivi come una successione è sufficiente seguire il percorso qui indicato, saltando i numeri (più chiari) che sono già stati incontrati precedentemente (ricordiamo che frazioni diverse possono indicare lo stesso numero). Con uno schema simile si possono inserire anche i razionali negativi o nulli e quindi ottenere una biezione fra $\mathbb{N}$ e $\mathbb{Q}$.

Un esempio ancora più sorprendente è quello dei razionali $\mathbb{Q}$. I numeri razionali possono essere messi "in fila" come indicato nella Figura 1.10. Nonostante i numeri razionali sembrano molti di più dei naturali, in realtà sono "lo stesso numero".

A questo punto può sorgere il dubbio che semplicemente tutti gli insiemi infiniti siano numerabili. Questo non è il caso:

Proposizione 1.5.4. *L'insieme $\mathbb{R}$ dei numeri reali non è numerabile.*

Dimostrazione. Supponiamo per assurdo che sia possibile scrivere $\mathbb{R}$ come successione infinita di numeri:

$$\mathbb{R} = \{x_1, x_2, \dots\}.$$

Scriviamo ciascun numero in forma decimale, ad esempio:

$$x_1 = 723,1291851\dots$$

$$x_2 = 12,8231452\dots$$

$$x_3 = 0,3798921\dots$$

$$x_4 = 110,0023140\dots$$

...

Ad essere precisi esistono alcuni numeri reali che possono essere rappresentati in due modi come forma decimale: quelli che ad un certo punto si stabilizzano con una successione di 9 o di 0. Ad esempio $0,18563\bar{9} = 0,18564$. In questo caso scegliamo la successione che si stabilizza con 0.

Indichiamo con n_i la i -esima cifra dopo la virgola di x_i . Definiamo adesso un nuovo numero reale x nel modo seguente:

$$x = 0, a_1 a_2 a_3 a_4 a_5 \dots$$

con questa regola: se $n_i \neq 1$ allora $a_i = 1$; se $n_i = 1$ allora $a_i = 2$.

Adesso notiamo che x non può essere uguale a nessun numero della successione $x_1, \dots$ e quindi giungiamo ad un assurdo. $\square$

La tecnica usata nella dimostrazione è l'*argomento diagonale di Cantor*. Abbiamo scoperto che i numeri reali $\mathbb{R}$ sono "di più" dei numeri naturali, interi o razionali. Sono sempre infiniti, ma di un ordine superiore. Non può esistere nessuna bigezione fra $\mathbb{N}$ e $\mathbb{R}$.

Quali applicazioni pratiche può avere un risultato del genere? Questo fatto ha effettivamente delle conseguenze nella vita di tutti i giorni: siccome $\mathbb{R}$ non è numerabile, non è possibile in alcun modo scrivere un programma al computer che tratti i numeri reali con precisione assoluta. In informatica i numeri reali vengono sempre approssimati, e non è possibile che sia altrimenti.

1.11. Costruzione dei numeri reali. A scuola impariamo che un numero reale è un numero che può avere infinite cifre dopo la virgola. Questa definizione è corretta (c'è solo una piccola ambiguità da tenere a mente, riguardante il fatto che numeri del tipo $5,973\bar{9}$ e $5,974$ sono in realtà lo stesso) ma non ci dice nulla su come siano definite la somma ed il prodotto di due numeri reali.

Per definire somma e prodotto è più comodo definire i numeri reali in un altro modo, che ha anche il pregio di essere più intrinseco ed indipendente dalla scelta della base 10 che abbiamo adottato solo perché la selezione naturale ci ha gentilmente fornito di 10 dita. Useremo il linguaggio dell'analisi. Supponiamo che i numeri razionali siano già definiti, come descritto precedentemente nell'Esempio 1.1.10.

Come si fa usualmente in analisi, diciamo che una successione (a_n) di numeri razionali $a_n \in \mathbb{Q}$ *tende* ad un certo numero razionale a_∞ se per ogni numero razionale $\varepsilon > 0$ esiste un N tale che $|a_n - a_\infty| < \varepsilon$ per ogni $n > N$.

Una *successione di Cauchy* di numeri razionali è una successione (a_n) di numeri razionali $a_n \in \mathbb{Q}$ con questa proprietà: per ogni numero razionale $\varepsilon > 0$ esiste un $N > 0$ per cui

$$|a_m - a_n| < \varepsilon$$

per ogni $m, n > N$. Intuitivamente, una successione di Cauchy di numeri razionali è una successione che vorrebbe tendere a qualcosa ma non è detto

che il limite esista: se il limite non c'è, lo dobbiamo creare, e lo facciamo adesso in modo astratto usando le successioni stesse per definire nuovi numeri.

Diciamo che due successioni di Cauchy (a_n) e (b_n) sono *equivalenti*, e in questo caso scriviamo $(a_n) \sim (b_n)$, se la successione delle differenze $(a_n - b_n)$ tende a zero.

Si dimostra facilmente che $\sim$ è una relazione di equivalenza fra successioni di Cauchy di numeri razionali. Siamo pronti per definire i numeri reali.

Definizione 1.5.5. L'insieme $\mathbb{R}$ dei *numeri reali* è l'insieme delle classi di equivalenza di successioni di Cauchy di numeri razionali.

Sia (a_n) una successione di Cauchy di numeri razionali. Se questa converge ad un numero razionale a_∞ , allora la successione rappresenta semplicemente il numero a_∞ . Se invece non converge a nessun numero razionale, rappresenta un nuovo numero reale, che non è contenuto in $\mathbb{Q}$: un numero *irrazionale*.

Esempio 1.5.6. La successione di numeri razionali

$$a_n = \left(1 + \frac{1}{n}\right)^n$$

è una successione di Cauchy che non converge ad un numero razionale. Quindi essa stessa definisce un nuovo numero reale e , la costante di Nepero.

Con questa interpretazione, un numero reale che abbia sviluppo decimale

$$723,1291851\dots$$

corrisponde ad una successione di Cauchy (a_n) di numeri razionali

$$a_1 = 723,1$$

$$a_2 = 723,12$$

$$a_3 = 723,129$$

$$a_4 = 723,1291$$

$\dots$

La successione che definisce questo numero reale chiaramente non è unica. Passiamo adesso a definire somma e prodotto di due numeri reali.

Definizione 1.5.7. Siano a e b due numeri reali, rappresentati da due successioni di Cauchy (a_n) e (b_n) di numeri razionali. La somma $a + b$ ed il prodotto $a \cdot b$ sono le successioni di Cauchy $(a_n + b_n)$ e $(a_n b_n)$.

Si verifica con un po' di pazienza che queste operazioni sono ben definite e che $\mathbb{R}$ con queste due operazioni è un campo. A differenza di $\mathbb{Q}$, l'insieme $\mathbb{R}$ dei numeri reali è *completo*: ogni successione di Cauchy in $\mathbb{R}$

converge. Questo vuol dire che se facciamo la stessa costruzione di prima con successioni di Cauchy in $\mathbb{R}$ non aggiungiamo più nessun altro punto: nel passare da $\mathbb{Q}$ ad $\mathbb{R}$ abbiamo già tappato tutti i buchi.

Un aspetto fondamentale degli insiemi numerici $\mathbb{Z}, \mathbb{Q}$ e $\mathbb{R}$, che li differenzia da $\mathbb{C}$, è che questi insiemi sono *ordinati*: esiste una nozione di maggiore e minore fra numeri, per cui se a, b sono due numeri distinti deve valere $a > b$ oppure $b > a$. In generale diciamo che $a > b$ se $a - b > 0$, quindi per definire questa nozione è sufficiente chiarire quali siano i numeri positivi (cioè maggiori di zero) e quali i numeri negativi (minori di zero). In $\mathbb{Z}$, quelli positivi sono $1, 2, 3, \dots$. In $\mathbb{Q}$, i numeri positivi sono le frazioni $\frac{p}{q}$ in cui p e q hanno lo stesso segno. In $\mathbb{R}$, un numero a è positivo se è rappresentabile da una successione di Cauchy (a_n) di numeri razionali tutti positivi. Su $\mathbb{C}$ invece *non* è possibile definire una nozione ragionevole di numero positivo e negativo.

Notiamo infine che il campo $\mathbb{R}$ soddisfa la *proprietà di Archimede*:

*Dati due numeri reali positivi $a < b$,
esiste sempre un $n \in \mathbb{N}$ tale che $na > b$.*

È possibile dimostrare che $\mathbb{R}$ è l'unico campo ordinato completo e archimedeo (cioè in cui sia verificata la proprietà di Archimede).

Spazi vettoriali

Nel capitolo precedente abbiamo già intravisto il piano cartesiano $\mathbb{R}^2$, lo spazio cartesiano $\mathbb{R}^3$, ed il più misterioso spazio n -dimensionale $\mathbb{R}^n$. Lo spazio $\mathbb{R}^n$ è il nostro *spazio euclideo*, l'universo dentro al quale vogliamo definire e studiare la geometria euclidea.

Il nostro scopo è quindi costruire e studiare oggetti geometrici dentro $\mathbb{R}^n$. In realtà ci accorgiamo subito che è più conveniente prendere un punto di vista ancora più astratto: invece di esaminare “solo” $\mathbb{R}^n$, ci proponiamo di definire e studiare qualsiasi spazio che abbia le stesse proprietà algebriche di $\mathbb{R}^n$. Uno spazio di questo tipo è detto *spazio vettoriale*.

Il motivo di questa ulteriore astrazione è che, inaspettatamente, esistono molti insiemi di oggetti matematici che sono spazi vettoriali oltre a $\mathbb{R}^n$: le soluzioni di un sistema lineare, gli spazi di funzioni, di matrici, eccetera. È quindi più conveniente raggruppare tutte queste entità in un'unica teoria.

2.1. Lo spazio euclideo

Nella geometria di Euclide, tutto ha inizio con l'introduzione di alcuni *concetti primitivi* quali i punti, le rette e i piani, che non vengono definiti e per i quali valgono alcuni assiomi. Nella geometria analitica adottata in questo libro, partiamo invece dalla teoria degli insiemi ed in particolare dalla retta reale $\mathbb{R}$, e costruiamo una geometria a partire da questa.

2.1.1. Definizione. Partendo da $\mathbb{R}$ e usando l'operazione di prodotto cartesiano fra insiemi, definiamo subito la nozione di spazio euclideo in dimensione arbitraria. Sia $n \geq 1$ un numero naturale.

Definizione 2.1.1. Lo *spazio euclideo* n -dimensionale è l'insieme $\mathbb{R}^n$.

Ricordiamo che $\mathbb{R}^n$ è il prodotto cartesiano $\mathbb{R} \times \cdots \times \mathbb{R}$ di n copie di $\mathbb{R}$. Un elemento dell'insieme $\mathbb{R}^n$ è una successione $(x_1, \dots, x_n)$ di n numeri reali.

L'insieme $\mathbb{R}^2$ è il piano cartesiano, in cui ogni punto è determinato da una coppia (x, y) . Analogamente $\mathbb{R}^3$ è lo spazio cartesiano, in cui ogni punto è una terna (x, y, z) di numeri reali, si veda la Figura 2.1. Quando n è arbitrario (ad esempio $n = 4$) lo spazio $\mathbb{R}^n$ è uno “spazio a n dimensioni”.

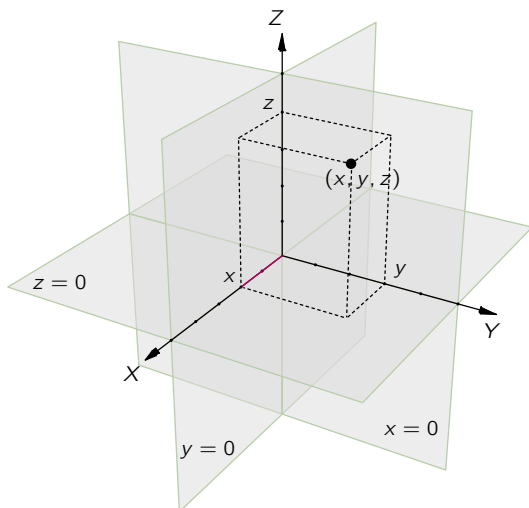


Figura 2.1. Lo spazio cartesiano $\mathbb{R}^3$. Ciascun punto è una terna (x, y, z) di numeri reali. La figura mostra i tre assi e i tre piani coordinati che contengono gli assi, descritti dalle equazioni $x = 0$, $y = 0$ e $z = 0$.

Gli spazi $\mathbb{R}^2$ e $\mathbb{R}^3$ possono essere visualizzati e studiati usando la nostra intuizione geometrica. D'altro canto, il nostro cervello non è programmato per visualizzare $\mathbb{R}^n$ per $n \geq 4$, ma possiamo comunque studiare $\mathbb{R}^n$ usando l'algebra, ed è esattamente quello che faremo in questo libro.

L'elemento $(0, \dots, 0)$ è l'*origine* dello spazio euclideo $\mathbb{R}^n$ e viene indicato semplicemente con i simboli 0 oppure O .

2.1.2. Punti o vettori? Un punto $x \in \mathbb{R}^n$ è per definizione una sequenza $(x_1, \dots, x_n)$ di numeri. Geometricamente, possiamo pensare a x come ad un punto, oppure come ad un *vettore* (cioè una freccia) che parte dall'origine 0 e arriva in x . Si veda la Figura 2.2. Entrambi i punti di vista sono ammessi e a seconda del contesto penseremo a x come ad un punto o ad un vettore. In alcuni casi, se pensiamo ad x come ad un punto possiamo usare le lettere P, Q , mentre se pensiamo ad x come vettore usiamo le lettere v, w .

Per un motivo che sarà chiaro in seguito quando introdurremo le matrici, scriveremo sempre la sequenza che identifica il punto x in verticale:

$$x = \begin{pmatrix} x_1 \\ \vdots \\ x_n \end{pmatrix}.$$

Un oggetto del genere è detto un *vettore colonna*. Diciamo che i numeri $x_1, \dots, x_n$ sono le *coordinate* del punto (o del vettore) x .

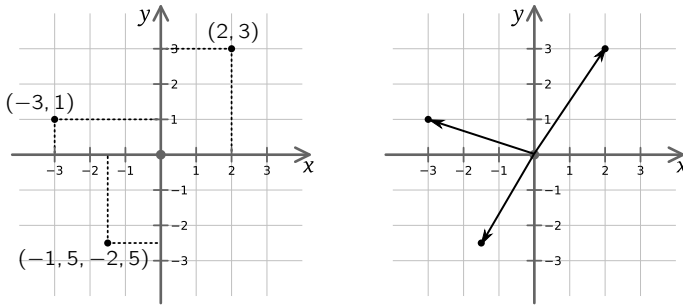


Figura 2.2. Ciascun punto $(x, y) \in \mathbb{R}^2$ può essere interpretato come un vettore applicato sull'origine.

2.1.3. Somme fra vettori. Lo spazio euclideo $\mathbb{R}^n$ non è soltanto un insieme: ha anche una struttura algebrica che introduciamo adesso. Definiamo innanzitutto una somma fra vettori. Se

$$x = \begin{pmatrix} x_1 \\ \vdots \\ x_n \end{pmatrix}, \quad y = \begin{pmatrix} y_1 \\ \vdots \\ y_n \end{pmatrix}$$

sono due vettori di $\mathbb{R}^n$, allora definiamo la loro somma semplicemente facendo la somma delle singole coordinate:

$$x + y = \begin{pmatrix} x_1 + y_1 \\ \vdots \\ x_n + y_n \end{pmatrix}.$$

Per $\mathbb{R}^2$, questa corrisponde precisamente all'usuale somma di vettori con la regola del parallelogramma, come mostrato nella Figura 2.3-(sinistra). Nella figura è descritta l'operazione

$$\begin{pmatrix} 1 \\ 2 \end{pmatrix} + \begin{pmatrix} 3 \\ 1 \end{pmatrix} = \begin{pmatrix} 4 \\ 3 \end{pmatrix}.$$

2.1.4. Prodotto per scalare. Un'altra operazione consiste nell'allungare, accorciare o ribaltare un singolo vettore. Dato un vettore

$$x = \begin{pmatrix} x_1 \\ \vdots \\ x_n \end{pmatrix}$$

ed un numero reale $\lambda \in \mathbb{R}$, possiamo definire un nuovo vettore

$$\lambda x = \begin{pmatrix} \lambda x_1 \\ \vdots \\ \lambda x_n \end{pmatrix}.$$

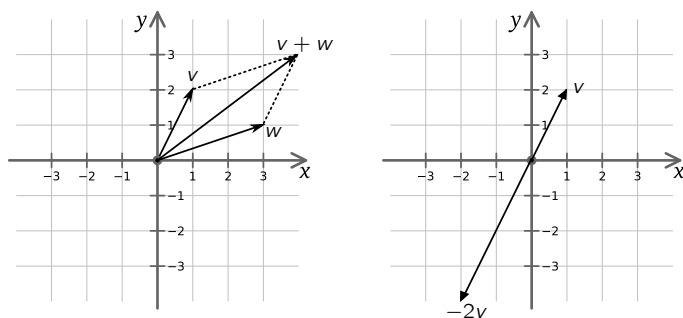


Figura 2.3. La somma $v + w$ di due vettori $v, w \in \mathbb{R}^2$ e la moltiplicazione di v per uno scalare λ (qui $\lambda = -2$).

In $\mathbb{R}^2$ possiamo visualizzare λx come il vettore ottenuto da x allungandolo o accorciandolo di un fattore $|\lambda|$ e ribaltandolo se $\lambda < 0$. Si veda la Figura 2.3. Nella figura è descritta l'operazione

$$-2 \begin{pmatrix} 1 \\ 2 \end{pmatrix} = \begin{pmatrix} -2 \\ -4 \end{pmatrix}.$$

Il numero reale λ usato in questo contesto è detto *scalare* e l'operazione che consiste nel moltiplicare un vettore x per uno scalare λ è detta semplicemente *prodotto per scalare*.

2.1.5. Proprietà algebriche. Si verifica facilmente che lo spazio euclideo $\mathbb{R}^n$ è un gruppo commutativo con l'operazione di somma appena descritta. L'elemento neutro è il vettore zero

$$0 = \begin{pmatrix} 0 \\ \vdots \\ 0 \end{pmatrix}$$

che corrisponde all'origine. L'inverso del vettore x è il vettore $-x$ che ha tutte le coordinate cambiate di segno.

La somma ed il prodotto per scalare inoltre soddisfano altre proprietà:

- (1) $\lambda(x + y) = \lambda x + \lambda y$
- (2) $(\lambda + \mu)x = \lambda x + \mu x$
- (3) $(\lambda\mu)x = \lambda(\mu x)$
- (4) $1x = x$

Tutte queste proprietà sono valide per ogni $x, y \in \mathbb{R}^n$ e per ogni $\lambda, \mu \in \mathbb{R}$. Sono tutte conseguenze delle usuali proprietà associative e distributive dei numeri reali.

2.2. Spazi vettoriali

Nella sezione precedente abbiamo iniziato a studiare lo spazio euclideo $\mathbb{R}^n$, che generalizza il piano cartesiano $\mathbb{R}^2$ in ogni dimensione. Abbiamo

definito la somma fra vettori ed il prodotto per scalare e abbiamo verificato che queste due operazioni soddisfano alcuni assiomi.

In questa sezione daremo un nome a qualsiasi struttura algebrica che soddisfi questi assiomi: la chiameremo uno *spazio vettoriale*. Il motivo di questa astrazione è che vi sono in matematica molti altri spazi che soddisfano gli stessi assiomi di $\mathbb{R}^n$, ed è quindi molto più conveniente considerarli tutti assieme.

2.2.1. Definizione. Fissiamo un campo $\mathbb{K}$. Questo è generalmente il campo $\mathbb{K} = \mathbb{R}$ dei numeri reali, ma può anche essere ad esempio il campo $\mathbb{K} = \mathbb{C}$ dei numeri complessi oppure il campo $\mathbb{K} = \mathbb{Q}$ dei numeri razionali. Gli elementi di $\mathbb{K}$ sono detti *scalari*.

Uno *spazio vettoriale* su $\mathbb{K}$ è un insieme V di elementi, detti *vettori*, dotato di due operazioni binarie:

- una operazione detta *somma* che associa a due vettori $v, w \in V$ un terzo vettore $v + w \in V$;
- una operazione detta *prodotto per scalare* che associa ad un vettore $v \in V$ e ad uno scalare $\lambda \in \mathbb{K}$ un vettore $\lambda v \in V$.

Queste due operazioni devono soddisfare gli stessi assiomi che abbiamo verificato precedentemente per lo spazio euclideo, cioè:

- (1) L'insieme V è un gruppo commutativo con la somma $+$
- (2) $\lambda(v + w) = \lambda v + \lambda w$
- (3) $(\lambda + \mu)v = \lambda v + \mu v$
- (4) $(\lambda\mu)v = \lambda(\mu v)$
- (5) $1v = v$

Le proprietà (2)-(5) devono valere per ogni $v, w \in V$ e ogni $\lambda, \mu \in \mathbb{K}$.

L'elemento neutro del gruppo V è indicato con il simbolo 0 ed è detto l'*origine* dello spazio vettoriale V . Non va confuso con lo zero 0 del campo $\mathbb{K}$. In questo libro il simbolo 0 può indicare varie cose differenti ed il significato preciso sarà sempre chiaro dal contesto.

A partire da questi assiomi possiamo dimostrare un piccolo teorema.

Proposizione 2.2.1. *Vale la relazione $0v = 0$.*

Come dicevamo, il simbolo 0 può indicare oggetti molto diversi: nell'uguaglianza $0v = 0$, il primo 0 è l'elemento neutro di $\mathbb{K}$ ed il secondo 0 è l'origine di V .

Dimostrazione. Vale

$$0v = (0 + 0)v = 0v + 0v$$

e semplificando deduciamo che $0v = 0$. □

2.2.2. Lo spazio $\mathbb{K}^n$. L'esempio principale di spazio vettoriale su $\mathbb{R}$ è ovviamente lo spazio euclideo $\mathbb{R}^n$ già incontrato precedentemente. Più in generale, è possibile definire per qualsiasi campo $\mathbb{K}$ uno spazio $\mathbb{K}^n$.

Sia $n \geq 1$ un numero naturale. Lo spazio $\mathbb{K}^n$ è l'insieme delle sequenze $(x_1, \dots, x_n)$ di numeri in $\mathbb{K}$. Gli elementi $v \in \mathbb{K}^n$ sono descritti generalmente come vettori colonna

$$v = \begin{pmatrix} x_1 \\ \vdots \\ x_n \end{pmatrix} \in \mathbb{K}^n.$$

La somma e la moltiplicazione per scalare sono definiti termine a termine esattamente come visto sopra per $\mathbb{R}^n$, cioè:

$$\begin{pmatrix} x_1 \\ \vdots \\ x_n \end{pmatrix} + \begin{pmatrix} y_1 \\ \vdots \\ y_n \end{pmatrix} = \begin{pmatrix} x_1 + y_1 \\ \vdots \\ x_n + y_n \end{pmatrix}, \quad \lambda \begin{pmatrix} x_1 \\ \vdots \\ x_n \end{pmatrix} = \begin{pmatrix} \lambda x_1 \\ \vdots \\ \lambda x_n \end{pmatrix}.$$

Facciamo alcuni esempi con dei vettori di $\mathbb{C}^2$:

$$\begin{pmatrix} 1+i \\ -2 \end{pmatrix} + \begin{pmatrix} 3i \\ 1-i \end{pmatrix} = \begin{pmatrix} 1+4i \\ -1-i \end{pmatrix}, \quad (2+i) \begin{pmatrix} 3 \\ 1-i \end{pmatrix} = \begin{pmatrix} 6+3i \\ 3-i \end{pmatrix}.$$

2.2.3. Lo spazio $\mathbb{K}[x]$ dei polinomi. Come abbiamo accennato, ci sono altri oggetti matematici molto comuni che formano degli spazi vettoriali. Ad esempio, i polinomi.

Indichiamo con $\mathbb{K}[x]$ l'insieme di tutti i polinomi con coefficienti in un certo campo $\mathbb{K}$. Due polinomi possono essere sommati, e moltiplicando un polinomio per uno scalare otteniamo un polinomio. Ad esempio:

$$(x^3 - 2x + 1) + (4x^4 + x - 3) = 4x^4 + x^3 - x - 2, \quad 3(x^3 - 2x) = 3x^3 - 6x.$$

L'insieme $\mathbb{K}[x]$ è quindi naturalmente equipaggiato con le due operazioni di uno spazio vettoriale. Inoltre tutti gli assiomi (1)-(5) sono facilmente verificati. Quindi $\mathbb{K}[x]$ è uno spazio vettoriale su $\mathbb{K}$.

2.2.4. Lo spazio $F(X, \mathbb{K})$ delle funzioni. Oltre ai polinomi, anche le funzioni formano spesso uno spazio vettoriale.

Sia X un insieme qualsiasi e $\mathbb{K}$ un campo. Consideriamo l'insieme $F(X, \mathbb{K})$ formato da tutte le funzioni $f: X \rightarrow \mathbb{K}$. Vediamo come anche $F(X, \mathbb{K})$ sia naturalmente equipaggiato con le due operazioni di uno spazio vettoriale. Date due funzioni $f, g: X \rightarrow \mathbb{K}$ possiamo definire la loro somma $f + g$ come la funzione che manda x in $f(x) + g(x)$. In altre parole:

$$(f + g)(x) = f(x) + g(x).$$

A sinistra abbiamo scritto la funzione $f + g$ fra parentesi. Notiamo che anche $f + g$ è una funzione da X in $\mathbb{K}$. Analogamente, se $f: X \rightarrow \mathbb{K}$ è una funzione e $\lambda \in \mathbb{K}$ è uno scalare, possiamo definire una nuova funzione λf che manda x in $\lambda f(x)$. In altre parole:

$$(\lambda f)(x) = \lambda f(x).$$

Anche qui abbiamo scritto la funzione λf fra parentesi. Si può verificare facilmente che l'insieme $F(X, \mathbb{K})$ soddisfa gli assiomi (1)-(5) ed è quindi uno spazio vettoriale.

L'elemento neutro di $F(X, \mathbb{K})$ è la *funzione nulla* 0, quella che fa zero su tutti gli elementi di X , cioè $0(x) = 0 \forall x \in X$.

Ad esempio, l'insieme $F(\mathbb{R}, \mathbb{R})$ di tutte le funzioni da $\mathbb{R}$ in $\mathbb{R}$ forma uno spazio vettoriale. Notiamo il livello di astrazione richiesto in questa sezione: una funzione, che può essere un oggetto abbastanza complicato, è interpretata come un punto oppure un vettore in uno spazio molto grande $F(X, \mathbb{R})$ che contiene tutte le possibili funzioni. Per quanto sembri inutilmente astratto, questo approccio globale è molto utile nello studio delle funzioni che capitano spesso in natura.

2.2.5. Le matrici. Dopo i polinomi e le funzioni, introduciamo un terzo oggetto matematico molto importante che rientra anch'esso nel quadro generale degli spazi vettoriali: le *matrici*.

Sia come sempre $\mathbb{K}$ un campo fissato. Una *matrice* con m righe e n colonne a coefficienti in $\mathbb{K}$ è una tabella rettangolare del tipo

$$A = \begin{pmatrix} a_{11} & \cdots & a_{1n} \\ \vdots & \ddots & \vdots \\ a_{m1} & \cdots & a_{mn} \end{pmatrix}$$

in cui tutti gli mn numeri $a_{11}, \dots, a_{mn}$ sono elementi del campo $\mathbb{K}$. Diciamo brevemente che A è una matrice $m \times n$. La matrice A ha effettivamente m *righe* che indicheremo con il simbolo $A_1, \dots, A_m$, del tipo

$$A_i = (a_{i1} \quad \cdots \quad a_{in})$$

e n *colonne* che indicheremo con $A^1, \dots, A^n$ del tipo

$$A^j = \begin{pmatrix} a_{1j} \\ \vdots \\ a_{mj} \end{pmatrix}.$$

Ad esempio, le seguenti sono matrici a coefficienti in $\mathbb{R}$:

$$\begin{pmatrix} 1 & \sqrt{2} \\ 0 & -5 \\ 7 & \pi \end{pmatrix}, \quad (5 \quad 0 \quad \sqrt{3}).$$

Le matrici giocheranno un ruolo fondamentale in tutto il libro. Per il momento ci limitiamo a notare che due matrici A e B della stessa taglia $m \times n$

$$A = \begin{pmatrix} a_{11} & \cdots & a_{1n} \\ \vdots & \ddots & \vdots \\ a_{m1} & \cdots & a_{mn} \end{pmatrix}, \quad B = \begin{pmatrix} b_{11} & \cdots & b_{1n} \\ \vdots & \ddots & \vdots \\ b_{m1} & \cdots & b_{mn} \end{pmatrix}$$

possono essere sommate e dare quindi luogo ad una nuova matrice $A + B$ così:

$$A + B = \begin{pmatrix} a_{11} + b_{11} & \cdots & a_{1n} + b_{1n} \\ \vdots & \ddots & \vdots \\ a_{m1} + b_{m1} & \cdots & a_{mn} + b_{mn} \end{pmatrix}.$$

Ancora una volta, si tratta semplicemente di sommare separatamente i coefficienti. Analogamente è possibile moltiplicare la matrice A per uno scalare $\lambda \in \mathbb{K}$ e ottenere una nuova matrice

$$\lambda A = \begin{pmatrix} \lambda a_{11} & \cdots & \lambda a_{1n} \\ \vdots & \ddots & \vdots \\ \lambda a_{m1} & \cdots & \lambda a_{mn} \end{pmatrix}.$$

Anche qui, abbiamo semplicemente moltiplicato per λ tutti i coefficienti della matrice. Indichiamo con $M(m, n, \mathbb{K})$ l'insieme di tutte le matrici $m \times n$ con coefficienti in $\mathbb{K}$, con le operazioni di somma e prodotto per scalare appena descritte. Quando il campo è sottinteso scriviamo $M(m, n)$.

Si verifica facilmente che l'insieme $M(m, n, \mathbb{K})$ soddisfa gli assiomi (1)-(5) ed è quindi uno spazio vettoriale.

Notiamo che una matrice $m \times 1$ è semplicemente un vettore colonna e quindi $M(m, 1, \mathbb{K}) = \mathbb{K}^m$. Un vettore colonna è una matrice formata da una sola colonna.

2.2.6. Sottospazio vettoriale. I quattro esempi di spazi vettoriali appena descritti $\mathbb{K}^n$, $\mathbb{K}[x]$, $F(X, \mathbb{K})$ e $M(m, n, \mathbb{K})$ sono in un certo senso i più importanti perché contengono quasi tutti gli spazi vettoriali che vedremo in questo libro. Definiamo adesso rigorosamente in che senso uno spazio vettoriale può contenerne un altro.

Sia V uno spazio vettoriale su un campo $\mathbb{K}$. Un *sottospazio vettoriale* di V è un sottoinsieme $W \subset V$ che soddisfa i seguenti tre assiomi:

- (1) $0 \in W$.
- (2) Se $v, v' \in W$, allora anche $v + v' \in W$.
- (3) Se $v \in W$ e $\lambda \in \mathbb{K}$, allora $\lambda v \in W$.

La prima proprietà dice che W deve contenere l'origine; la seconda può essere riassunta dicendo che W deve essere chiuso rispetto alla somma, e la terza asserendo che W deve essere chiuso rispetto al prodotto per scalare.

Il punto fondamentale qui è il seguente: un sottospazio vettoriale $W \subset V$ è esso stesso uno spazio vettoriale. Infatti, la somma ed il prodotto per scalare in W sono ben definiti grazie agli assiomi (2) e (3), e tutti gli assiomi da spazio vettoriale di V si mantengono ovviamente anche per W con la stessa origine 0 di V .

2.2.7. I sottospazi banale e totale. Ogni spazio vettoriale V ha due sottospazi molto particolari che è bene tenere a mente fin da subito:

- Il *sottospazio banale* $W = \{0\}$, formato da un punto solo, l'origine.
- Il *sottospazio totale* $W = V$, formato da tutti i vettori di V .

Tutti gli altri sottospazi W di V hanno una posizione intermedia fra quello banale e quello totale: sono sempre contenuti in V e contengono sempre $\{0\}$. In altre parole:

$$\{0\} \subset W \subset V.$$

A questo punto è relativamente facile descrivere dei sottospazi vettoriali dei quattro spazi modello $\mathbb{K}^n$, $\mathbb{K}[n]$, $F(X, \mathbb{K})$ e $M(m, n, \mathbb{K})$ descritti precedentemente e ottenere così molti nuovi spazi vettoriali. Questo è quello che faremo nelle prossime pagine.

2.2.8. Sistemi lineari omogenei. Una *equazione lineare* è una equazione con n variabili $x_1, \dots, x_n$ del tipo

$$a_1x_1 + \dots + a_nx_n = b$$

dove $a_1, \dots, a_n, b$ sono numeri in un certo campo $\mathbb{K}$. L'equazione è *omogenea* se $b = 0$. Un punto x di $\mathbb{K}^n$ è *soluzione* dell'equazione se le sue coordinate $x_1, \dots, x_n$ la soddisfano, cioè se sostituendole nell'espressione $a_1x_1 + \dots + a_nx_n$ si ottiene proprio b .

Un *sistema lineare* è un insieme di k equazioni lineari in n variabili

$$\begin{cases} a_{11}x_1 + \dots + a_{1n}x_n = b_1, \\ \vdots \\ a_{k1}x_1 + \dots + a_{kn}x_n = b_k. \end{cases}$$

Il sistema è *omogeneo* se $b_1 = \dots = b_k = 0$. Un punto $x \in \mathbb{K}^n$ è *soluzione* del sistema se è soluzione di ciascuna equazione. Le soluzioni del sistema formano un certo sottoinsieme $S \subset \mathbb{K}^n$.

Proposizione 2.2.2. *Le soluzioni $S \subset \mathbb{K}^n$ di un sistema lineare omogeneo formano un sottospazio vettoriale di $\mathbb{K}^n$.*

Dimostrazione. Consideriamo un sistema lineare omogeneo

$$\begin{cases} a_{11}x_1 + \dots + a_{1n}x_n = 0, \\ \vdots \\ a_{k1}x_1 + \dots + a_{kn}x_n = 0. \end{cases}$$

Dobbiamo verificare i 3 assiomi di sottospazio per l'insieme S .

- (1) Il vettore 0 è soluzione, infatti per ogni equazione troviamo

$$a_{i1}0 + \dots + a_{in}0 = 0.$$

Quindi $0 \in S$.

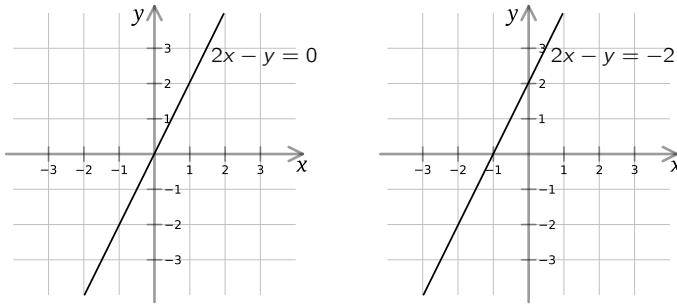


Figura 2.4. La retta $2x - y = 0$ è un sottospazio vettoriale di $\mathbb{R}^2$, mentre la retta $2x - y = -2$ no.

- (2) Se due vettori x e y sono soluzioni, allora anche $x+y$ è soluzione. Infatti per ogni equazione troviamo

$$\begin{aligned} a_{i1}(x_1 + y_1) + \cdots + a_{in}(x_n + y_n) &= \\ a_{i1}x_1 + \cdots + a_{in}x_n + a_{i1}y_1 + \cdots + a_{in}y_n &= 0 + 0 = 0. \end{aligned}$$

- (3) Se x è soluzione e $\lambda \in \mathbb{K}$ è uno scalare, allora anche λx è soluzione. Infatti per ogni equazione troviamo

$$\begin{aligned} a_{i1}(\lambda x_1) + \cdots + a_{in}(\lambda x_n) &= \\ \lambda(a_{i1}x_1 + \cdots + a_{in}x_n) &= \lambda 0 = 0. \end{aligned}$$

Quindi S è un sottospazio vettoriale di $\mathbb{K}^n$. □

Ad esempio, le soluzioni dell'equazione $2x - y = 0$ in $\mathbb{R}^2$ formano la retta mostrata nella Figura 2.4-(sinistra). Le soluzioni dell'equazione $z = 0$ in $\mathbb{R}^3$ formano il piano orizzontale mostrato nella Figura 2.1. Le soluzioni del sistema $\{x = 0, z = 0\}$ in $\mathbb{R}^3$ formano una retta, l'asse y descritto nella stessa figura. Tutti questi sono esempi di sottospazi vettoriali di $\mathbb{R}^2$ e $\mathbb{R}^3$, perché sono descritti da un sistema lineare omogeneo.

Osservazione 2.2.3. Le soluzioni S di un sistema lineare *non* omogeneo non sono un sottospazio vettoriale perché non contengono l'origine. Ad esempio l'equazione $2x - y = -2$ in $\mathbb{R}^2$ descrive una retta non passante per l'origine: questa non è un sottospazio vettoriale. Si veda la Figura 2.4.

2.2.9. Combinazioni lineari. Sia V uno spazio vettoriale qualsiasi. Siano $v_1, \dots, v_k \in V$ dei vettori arbitrari. Una *combinazione lineare* dei vettori $v_1, \dots, v_k$ è un qualsiasi vettore v che si ottiene come

$$v = \lambda_1 v_1 + \cdots + \lambda_k v_k$$

dove $\lambda_1, \dots, \lambda_k$ sono scalari arbitrari. Ad esempio, se $V = \mathbb{R}^3$ e

$$v_1 = \begin{pmatrix} 1 \\ 0 \\ 0 \end{pmatrix}, \quad v_2 = \begin{pmatrix} 0 \\ 1 \\ 0 \end{pmatrix}$$

allora una combinazione lineare arbitraria di questi due vettori è il vettore

$$\lambda_1 v_1 + \lambda_2 v_2 = \lambda_1 \begin{pmatrix} 1 \\ 0 \\ 0 \end{pmatrix} + \lambda_2 \begin{pmatrix} 0 \\ 1 \\ 0 \end{pmatrix} = \begin{pmatrix} \lambda_1 \\ 0 \\ 0 \end{pmatrix} + \begin{pmatrix} 0 \\ \lambda_2 \\ 0 \end{pmatrix} = \begin{pmatrix} \lambda_1 \\ \lambda_2 \\ 0 \end{pmatrix}.$$

Notiamo che questo vettore sta nel piano orizzontale $z = 0$. Osserviamo anche che al variare di λ_1 e λ_2 , facendo tutte le combinazioni lineari di v_1 e v_2 otteniamo precisamente tutti i punti del piano orizzontale $z = 0$. Questo ci porta alla definizione seguente.

2.2.10. Sottospazio generato. Sia V uno spazio vettoriale e siano $v_1, \dots, v_k \in V$ vettori arbitrari. Il *sottospazio generato* da $v_1, \dots, v_k$ è il sottoinsieme di V formato da tutte le combinazioni lineari dei vettori $v_1, \dots, v_k$. Questo sottospazio viene indicato con il simbolo

$$\text{Span}(v_1, \dots, v_k).$$

In inglese *span* vuol dire proprio generare (in questo contesto). In simboli:

$$\text{Span}(v_1, \dots, v_k) = \{ \lambda_1 v_1 + \dots + \lambda_k v_k \mid \lambda_1, \dots, \lambda_k \in \mathbb{K} \}.$$

Proposizione 2.2.4. *Il sottoinsieme $\text{Span}(v_1, \dots, v_k)$ è un sottospazio vettoriale di V .*

Dimostrazione. Dobbiamo mostrare che $W = \text{Span}(v_1, \dots, v_k)$ soddisfa i 3 assiomi di sottospazio.

- (1) $0 \in W$, infatti usando $\lambda_1 = \dots = \lambda_k = 0$ troviamo

$$0v_1 + \dots + 0v_k = 0 + \dots + 0 = 0 \in W.$$

- (2) Se $v, w \in W$, allora $v + w \in W$. Infatti, per ipotesi

$$v = \lambda_1 v_1 + \dots + \lambda_k v_k$$

$$w = \mu_1 v_1 + \dots + \mu_k v_k$$

e quindi raccogliendo otteniamo

$$v + w = (\lambda_1 + \mu_1)v_1 + \dots + (\lambda_k + \mu_k)v_k.$$

Anche la somma $v + w$ è combinazione lineare di $v_1, \dots, v_k$ e quindi $v + w \in W$.

- (3) Se $v \in W$ e $\lambda \in \mathbb{K}$, allora $\lambda v \in W$. Infatti se

$$v = \lambda_1 v_1 + \dots + \lambda_k v_k$$

allora

$$\lambda v = (\lambda \lambda_1)v_1 + \dots + (\lambda \lambda_k)v_k$$

è anch'esso combinazione lineare di $v_1, \dots, v_k$.

La dimostrazione è conclusa. □

Ad esempio, se v è un singolo vettore di V , otteniamo

$$W = \text{Span}(v) = \{\lambda v \mid \lambda \in \mathbb{K}\}.$$

Per esempio se $V = \mathbb{R}^2$ e $v = \begin{pmatrix} 1 \\ 2 \end{pmatrix}$ troviamo che

$$W = \text{Span}(v) = \left\{ t \begin{pmatrix} 1 \\ 2 \end{pmatrix} \mid t \in \mathbb{R} \right\} = \left\{ \begin{pmatrix} t \\ 2t \end{pmatrix} \mid t \in \mathbb{R} \right\}.$$

Notiamo che W è precisamente la retta descritta dall'equazione $y = 2x$ già mostrata nella Figura 2.4-(sinistra).

2.2.11. Forma cartesiana e forma parametrica. Un sottospazio vettoriale di $\mathbb{K}^n$ che sia descritto come luogo di zeri di un sistema di equazioni lineari omogenee è detto in *forma cartesiana*. Un sottospazio vettoriale di $\mathbb{K}^n$ descritto come sottospazio generato da alcuni vettori è invece detto in *forma parametrica*.

Vedremo che qualsiasi sottospazio vettoriale di $\mathbb{K}^n$ può essere descritto in entrambi i modi. Ad esempio, il piano $W = \{z = 0\}$ in $\mathbb{R}^3$ può essere descritto in forma cartesiana tramite l'equazione $z = 0$, oppure in forma parametrica come il sottospazio generato dai due vettori v_1 e v_2 indicati nella Sezione 2.2.9. Otteniamo in questo modo:

$$W = \text{Span}(v_1, v_2) = \left\{ t \begin{pmatrix} 1 \\ 0 \\ 0 \end{pmatrix} + u \begin{pmatrix} 0 \\ 1 \\ 0 \end{pmatrix} \mid t, u \in \mathbb{R} \right\} = \left\{ \begin{pmatrix} t \\ u \\ 0 \end{pmatrix} \mid t, u \in \mathbb{R} \right\}.$$

La forma parametrica è *esplicita* perché indica chiaramente come sono fatti i punti di W , al variare di alcuni parametri (in questo caso t e u). La forma cartesiana è *implicita* perché descrive W come luogo di soluzioni di una equazione (o di un sistema di equazioni). In molti casi la forma parametrica è preferibile proprio perché esplicita, ma può anche capitare che quella implicita sia più comoda, a seconda del contesto.

Vedremo nelle prossime pagine come passare da una descrizione in forma cartesiana ad una in forma parametrica, e viceversa.

2.2.12. Polinomi con restrizioni. Siano $\mathbb{K}$ un campo e n un numero naturale qualsiasi. Indichiamo con $\mathbb{K}_n[x]$ il sottoinsieme di $\mathbb{K}[x]$ formato da tutti quei polinomi che hanno grado minore o uguale a n . In altre parole:

$$\mathbb{K}_n[x] = \{a_n x^n + \cdots + a_1 x + a_0 \mid a_0, a_1, \dots, a_n \in \mathbb{K}\}.$$

Questa scrittura può essere interpretata in un altro modo: il sottoinsieme $\mathbb{K}_n[x]$ è il sottospazio di $\mathbb{K}[x]$ generato dai polinomi $x^n, \dots, x, 1$, cioè:

$$\mathbb{K}_n[x] = \text{Span}(x^n, \dots, x, 1).$$

In particolare, ne deduciamo che $\mathbb{K}_n[x]$ è un sottospazio vettoriale di $\mathbb{K}[x]$.

Ci sono altri modi di imporre restrizioni sui polinomi e ottenere così dei sottospazi vettoriali di $\mathbb{K}[x]$. Ad esempio:

Proposizione 2.2.5. *Sia $a \in \mathbb{K}$ un numero fissato. I polinomi $p(x)$ tali che $p(a) = 0$ formano un sottospazio vettoriale di $\mathbb{K}[x]$.*

Dimostrazione. Sia $W \subset \mathbb{K}[x]$ il sottoinsieme formato da tutti i polinomi $p(x)$ tali che $p(a) = 0$. Come al solito, verifichiamo i 3 assiomi di sottospazio:

- (1) $0 \in W$, infatti il polinomio 0 si annulla ovunque e in particolare in a .
- (2) Se $p, q \in W$, allora $p + q \in W$. Infatti se due polinomi p e q si annullano in a , lo fa anche la loro somma:

$$(p + q)(a) = p(a) + q(a) = 0 + 0 = 0.$$

- (3) Se $p \in W$ e $\lambda \in \mathbb{K}$, allora $\lambda p \in W$. Infatti se $p(a) = 0$, allora anche $\lambda p(a) = 0$.

La dimostrazione è conclusa. □

Osservazione 2.2.6. I polinomi $p(x)$ tali che $p(a) = 1$ non formano un sottospazio vettoriale di $\mathbb{K}[x]$ perché non contengono il polinomio 0.

2.2.13. Funzioni con restrizioni. Sia X un insieme e $\mathbb{K}$ un campo, e $F(X, \mathbb{K})$ lo spazio vettoriale formato da tutte le funzioni da X in $\mathbb{K}$. Come per i polinomi, è anche qui possibile costruire sottospazi di $F(X, \mathbb{K})$ imponendo delle restrizioni. Possiamo ad esempio imporre che le funzioni si annullino in un qualsiasi sottoinsieme $Y \subset X$.

Esercizio 2.2.7. Sia $Y \subset X$ un sottoinsieme qualsiasi. Il sottoinsieme

$$W = \{f \in F(X, \mathbb{K}) \mid f(x) = 0 \forall x \in Y\} \subset F(X, \mathbb{K})$$

è un sottospazio vettoriale di $F(X, \mathbb{K})$.

Traccia. Verificare i 3 assiomi. □

Possiamo richiedere che le funzioni soddisfino qualche buona proprietà, come essere continue, derivabili oppure integrabili:

Proposizione 2.2.8. *Il sottoinsieme di $F(\mathbb{R}, \mathbb{R})$ formato dalle funzioni continue (oppure derivabili, oppure integrabili) è un sottospazio vettoriale.*

Dimostrazione. Dobbiamo verificare i 3 assiomi. È sufficiente notare che (1) la funzione costantemente nulla è continua (e derivabile, integrabile), che (2) la somma di due funzioni continue (oppure derivabili, integrabili) è anch'essa continua (oppure derivabile, integrabile), e infine che (3) moltiplicando una funzione continua (derivabile, integrabile) per uno scalare otteniamo ancora una funzione continua (derivabile, integrabile). Queste sono tutte proprietà dimostrate durante il corso di analisi. □

Abbiamo trovato una catena di sottospazi:

$$\{\text{funzioni derivabili}\} \subset \{\text{funzioni continue}\} \subset \{\text{funz. integrabili}\} \subset F(\mathbb{R}, \mathbb{R}).$$

2.2.14. Matrici diagonali, triangolari, simmetriche e antisimmetriche. Ricordiamo che $M(m, n, \mathbb{K})$ è lo spazio vettoriale formato dalle matrici $m \times n$ a coefficienti in $\mathbb{K}$. Come nei paragrafi precedenti, possiamo definire dei sottospazi di $M(m, n, \mathbb{K})$ imponendo opportune restrizioni.

Data una matrice A , indichiamo con A_{ij} oppure a_{ij} i suoi coefficienti. Introduciamo delle definizioni che ci accompagneranno in tutto il libro. Una matrice $n \times n$ è detta *quadrata*.

Definizione 2.2.9. Una matrice A quadrata $n \times n$ è:

- *diagonale* se $a_{ij} = 0 \ \forall i \neq j$;
- *triangolare superiore* se $a_{ij} = 0 \ \forall i > j$;
- *triangolare inferiore* se $a_{ij} = 0 \ \forall i < j$;
- *triangolare* se è triangolare inferiore o superiore;
- *simmetrica* se $a_{ij} = a_{ji} \ \forall i, j$;
- *antisimmetrica* se $a_{ij} = -a_{ji} \ \forall i, j$.

Ad esempio, le matrici seguenti sono (in ordine) diagonale, triangolare superiore, triangolare inferiore, simmetrica e antisimmetrica:

$$\begin{pmatrix} 2 & 0 \\ 0 & -1 \end{pmatrix}, \quad \begin{pmatrix} 1 & 9 \\ 0 & \sqrt{2} \end{pmatrix}, \quad \begin{pmatrix} -1 & 0 \\ 7 & 2 \end{pmatrix}, \quad \begin{pmatrix} -1 & 2 \\ 2 & 4 \end{pmatrix}, \quad \begin{pmatrix} 0 & 1 \\ -1 & 0 \end{pmatrix}.$$

Si faccia attenzione alle definizioni: la matrice

$$\begin{pmatrix} 0 & 0 \\ 0 & 1 \end{pmatrix}$$

è al tempo stesso diagonale, triangolare superiore, triangolare inferiore e simmetrica. La matrice nulla 0 è di tutti e 5 i tipi: è diagonale, triangolare superiore, triangolare inferiore, simmetrica e antisimmetrica.

Gli elementi $a_{11}, a_{22}, \dots, a_{nn}$ in una matrice quadrata A formano la *diagonale principale* di A . Notiamo che (nel caso in cui $\mathbb{K}$ sia uno dei campi usuali $\mathbb{Q}, \mathbb{R}, \mathbb{C}$) una matrice antisimmetrica ha necessariamente elementi nulli sulla diagonale principale, perché $a_{ii} = -a_{ii}$ implica $a_{ii} = 0$.¹

Indichiamo lo spazio $M(n, n, \mathbb{K})$ delle matrici quadrate $n \times n$ con $M(n, \mathbb{K})$ oppure ancora più semplicemente con $M(n)$. Indichiamo con

$$D(n), \quad T^s(n), \quad T^i(n), \quad S(n), \quad A(n)$$

rispettivamente i sottoinsiemi di $M(n)$ formati dalle matrici diagonali, triangolari superiori, triangolari inferiori, simmetriche e antisimmetriche.

Proposizione 2.2.10. *I sottoinsiemi $D(n), T^s(n), T^i(n), S(n), A(n)$ sono tutti sottospazi vettoriali di $M(n)$.*

Dimostrazione. Mostriamo l'enunciato per $S(n)$ e lasciamo gli altri casi per esercizio. I 3 assiomi del sottospazio sono verificati, perché

¹Dimostrazione: da $a_{ii} = -a_{ii}$ segue che $2a_{ii} = 0$, quindi si divide per 2 e si ottiene $a_{ii} = 0$. Questo procedimento è corretto, tranne nel caso molto particolare in cui nel campo $\mathbb{K}$ in realtà 2, che per definizione è $1 + 1$, non è invertibile perché $1 + 1 = 0$. Questo non accade nei campi $\mathbb{Q}, \mathbb{R}, \mathbb{C}$ dove ovviamente $1 + 1 \neq 0$, ma ci sono dei campi in cui $1 + 1 = 0$.

- (1) la matrice nulla è simmetrica,
- (2) la somma $A + B$ di due matrici simmetriche A e B è anch'essa simmetrica, infatti

$$(A + B)_{ij} = A_{ij} + B_{ij} = A_{ji} + B_{ji} = (A + B)_{ji}.$$

- (3) analogamente si verifica che se A è simmetrica e $\lambda \in \mathbb{K}$ allora λA è simmetrica.

La dimostrazione è conclusa. $\square$

2.2.15. Intersezione di sottospazi. Sia V uno spazio vettoriale e $U, W \subset V$ due sottospazi. Da un punto di vista insiemistico è naturale considerare la loro intersezione $U \cap W$.

Proposizione 2.2.11. L'intersezione $U \cap W$ di due sottospazi vettoriali $U, W \subset V$ è sempre un sottospazio vettoriale.

Dimostrazione. Come sempre, verifichiamo i 3 assiomi.

- (1) $0 \in U \cap W$, perché 0 appartiene sia a U che a W .
- (2) $v, w \in U \cap W \implies v + w \in U \cap W$. Infatti $v, w \in U$ implica $v + w \in U$ e $v, w \in W$ implica $v + w \in W$, quindi $v + w \in U \cap W$.
- (3) $v \in U \cap W \implies \lambda v \in U \cap W \forall \lambda \in \mathbb{K}$ è analogo al precedente.

La dimostrazione è conclusa. $\square$

Facciamo alcuni esempi.

Esempio 2.2.12. Il piani $U = \{x = 0\}$ e $W = \{y = 0\}$ in $\mathbb{R}^3$ mostrati in Figura 2.1 si intersecano nella retta $U \cap W$, l'asse z , che può essere descritta in forma cartesiana come insieme delle soluzioni di questo sistema:

$$\begin{cases} x = 0 \\ y = 0 \end{cases}$$

oppure esplicitamente in forma parametrica come

$$U \cap W = \left\{ \begin{pmatrix} 0 \\ 0 \\ t \end{pmatrix} \mid t \in \mathbb{R} \right\}.$$

Più in generale, se U e W sono sottospazi di $\mathbb{K}^n$ descritti in forma cartesiana come soluzioni di sistemi lineari omogenei, l'intersezione $U \cap W$ è descritta sempre in forma cartesiana unendo i sistemi lineari in uno solo. Questo fatto è conseguenza di un principio insiemistico più generale:

*Affiancare equazioni in un sistema
corrisponde a fare l'intersezione delle soluzioni.*

Esempio 2.2.13. Fra i sottospazi di $M(n)$ valgono le seguenti relazioni:

$$D(n) = T^s(n) \cap T^i(n), \quad \{0\} = S(n) \cap A(n).$$

La seconda relazione dice che una matrice A che è contemporaneamente simmetrica e antisimmetrica è necessariamente nulla. Infatti $a_{ij} = a_{ji}$ e $a_{ij} = -a_{ji}$ insieme implicano $a_{ij} = 0$, per ogni i, j .

2.2.16. L'unione non funziona. Dopo aver considerato l'intersezione $U \cap W$ di due sottospazi $U, W \subset V$, è naturale adesso considerare la loro unione $U \cup W$. Ci accorgiamo però immediatamente che l'unione $U \cup W$ di due sottospazi molto spesso *non* è un sottospazio, come mostra l'esempio seguente.

Esempio 2.2.14. Se $V = \mathbb{R}^2$ e $U = \{x = 0\}$ e $W = \{y = 0\}$ sono le due rette che definiscono gli assi di $\mathbb{R}^2$, la loro unione $U \cup W$ non è certamente un sottospazio perché non è soddisfatto l'assioma (2): ad esempio, se prendiamo $u = \begin{pmatrix} 1 \\ 0 \end{pmatrix}$ e $w = \begin{pmatrix} 0 \\ 1 \end{pmatrix}$, vediamo che $u, w \in U \cup W$ ma la loro somma $u + w = \begin{pmatrix} 1 \\ 1 \end{pmatrix}$ non appartiene a $U \cup W$.

Esercizio 2.2.15. Siano $U, W \subset V$ sottospazi vettoriali. Mostra che $U \cup W$ è un sottospazio se e solo se $U \subset W$ oppure $W \subset U$.

Abbandoniamo quindi questa strada e ci accingiamo a definire un'altra operazione che sarà molto più utile in seguito.

2.2.17. Somma di sottospazi. Siano come sopra $U, W \subset V$ due sottospazi di uno spazio vettoriale V .

Definizione 2.2.16. La *somma* $U + W$ dei due sottospazi U e W è il sottoinsieme di V seguente:

$$U + W = \{u + w \mid u \in U, w \in W\} \subset V.$$

La somma $U + W$ è l'insieme di tutti i vettori di V che possono essere scritti come somma $u + w$ di un vettore $u \in U$ e di un altro $w \in W$.

Proposizione 2.2.17. *L'insieme $U + W$ contiene sia U che W .*

Dimostrazione. Prendendo $u \in U$ e $0 \in W$ otteniamo $u = u + 0 \in U + W$ e analogamente facciamo scrivendo $w \in W$ come $0 + w \in U + W$. $\square$

Proposizione 2.2.18. *L'insieme $U + W$ è sottospazio vettoriale di V .*

Dimostrazione. Come sempre, verifichiamo i 3 assiomi.

- (1) $0 \in U + W$ perché $0 \in U, 0 \in W$ e scriviamo $0 = 0 + 0$.
- (2) $v, v' \in U + W \implies v + v' \in U + W$. Infatti per definizione $v = u + w$ e $v' = u' + w'$ con $u, u' \in U$ e $w, w' \in W$. Ne segue che

$$v + v' = (u + u') + (w + w')$$

con $u + u' \in U$ e $w + w' \in W$, quindi anche $v + v' \in U + W$.

- (3) $v \in U + W, \lambda \in \mathbb{K} \implies \lambda v \in U + W$. Analogamente al precedente.

La dimostrazione è conclusa. $\square$

Prima di fornire esempi concreti, descriviamo una proposizione utile per chiarire operativamente come funzioni la somma fra sottospazi.

Proposizione 2.2.19. Siano $u_1, \dots, u_h$ e $w_1, \dots, w_k$ vettori di uno spazio vettoriale V . Se

$$U = \text{Span}(u_1, \dots, u_h), \quad W = \text{Span}(w_1, \dots, w_k)$$

allora

$$U + W = \text{Span}(u_1, \dots, u_h, w_1, \dots, w_k).$$

Dimostrazione. L'insieme $U + W$ è formato dagli $u + w$ dove $u \in U$ e $w \in W$. Per ipotesi due generici $u \in U$ e $w \in W$ si scrivono come

$$u = \lambda_1 u_1 + \dots + \lambda_h u_h, \quad w = \mu_1 w_1 + \dots + \mu_k w_k$$

e quindi

$$u + w = \lambda_1 u_1 + \dots + \lambda_h u_h + \mu_1 w_1 + \dots + \mu_k w_k.$$

Abbiamo dimostrato in questo modo che $u + w$ è sempre combinazione lineare dei vettori $u_1, \dots, u_h, w_1, \dots, w_k$. Abbiamo quindi mostrato l'inclusione

$$U + W \subset \text{Span}(u_1, \dots, u_h, w_1, \dots, w_k).$$

Mostriamo adesso l'inclusione opposta. Ogni vettore v che sia combinazione lineare dei vettori $u_1, \dots, u_h, w_1, \dots, w_k$ si scrive come

$$v = \lambda_1 u_1 + \dots + \lambda_h u_h + \mu_1 w_1 + \dots + \mu_k w_k.$$

Definiamo

$$u = \lambda_1 u_1 + \dots + \lambda_h u_h, \quad w = \mu_1 w_1 + \dots + \mu_k w_k$$

e troviamo che $v = u + w$. Quindi abbiamo anche

$$U + W \supset \text{Span}(u_1, \dots, u_h, w_1, \dots, w_k).$$

I due insiemi sono uguali. □

Diciamo che dei vettori $v_1, \dots, v_k \in W$ sono dei *generatori* di W se

$$W = \text{Span}(v_1, \dots, v_k).$$

Abbiamo visto precedentemente che affiancare equazioni in un sistema corrisponde a fare l'intersezione dei sottospazi. La proposizione che abbiamo appena dimostrato può essere riassunta in modo analogo:

Affiancare dei generatori corrisponde a fare la somma di sottospazi.

Esempio 2.2.20. Consideriamo i vettori in $\mathbb{R}^3$:

$$v_1 = \begin{pmatrix} 1 \\ 0 \\ 0 \end{pmatrix}, \quad v_2 = \begin{pmatrix} 0 \\ 1 \\ 1 \end{pmatrix}.$$

I sottospazi $U = \text{Span}(v_1)$ e $W = \text{Span}(v_2)$ sono due rette, gli assi x e y . Lo spazio somma $U + W = \text{Span}(v_1, v_2)$ è il piano orizzontale che li contiene, descritto dall'equazione $z = 0$.

Osservazione 2.2.21. Nella dimostrazione della Proposizione 2.2.19, per mostrare che due insiemi A e B sono uguali, abbiamo verificato entrambe le inclusioni $A \subset B$ e $B \subset A$. Questa è una tecnica ricorrente in matematica.

2.3. Dimensione

Sappiamo dalla scuola che un punto, una retta e un piano hanno dimensione 0, 1 e 2. In questa sezione definiamo la nozione di dimensione di uno spazio vettoriale in modo rigoroso.

2.3.1. (In) dipendenza lineare. Sia V uno spazio vettoriale su $\mathbb{K}$ e siano $v_1, \dots, v_k \in V$ alcuni vettori. Diciamo che questi vettori sono *linearmente dipendenti* se esistono dei coefficienti $\lambda_1, \dots, \lambda_k \in \mathbb{K}$, non tutti nulli, tali che

$$\lambda_1 v_1 + \dots + \lambda_k v_k = 0.$$

Se i vettori $v_1, \dots, v_k$ sono linearmente dipendenti, allora è possibile esprimere uno di loro in funzione degli altri. Infatti, per ipotesi esiste almeno un $\lambda_i \neq 0$ e dopo aver diviso tutto per λ_i e spostato gli addendi otteniamo

$$v_i = -\frac{\lambda_1}{\lambda_i} v_1 - \dots - \frac{\lambda_k}{\lambda_i} v_k$$

dove tra i vettori di destra ovviamente non compare v_i . Quindi:

Proposizione 2.3.1. *I vettori $v_1, \dots, v_k$ sono dipendenti $\iff$ uno di loro è esprimibile come combinazione lineare degli altri.*

Esempio 2.3.2. I vettori $\begin{pmatrix} 1 \\ 2 \end{pmatrix}$, $\begin{pmatrix} 1 \\ 1 \end{pmatrix}$ e $\begin{pmatrix} 2 \\ 0 \end{pmatrix}$ in $\mathbb{R}^2$ sono dipendenti perché

$$-2 \begin{pmatrix} 1 \\ 2 \end{pmatrix} + 4 \begin{pmatrix} 1 \\ 1 \end{pmatrix} - \begin{pmatrix} 2 \\ 0 \end{pmatrix} = \begin{pmatrix} 0 \\ 0 \end{pmatrix}.$$

In questo caso ciascuno dei tre vettori è esprimibile come combinazione lineare degli altri due:

$$\begin{pmatrix} 1 \\ 2 \end{pmatrix} = 2 \begin{pmatrix} 1 \\ 1 \end{pmatrix} - \frac{1}{2} \begin{pmatrix} 2 \\ 0 \end{pmatrix}, \quad \begin{pmatrix} 1 \\ 1 \end{pmatrix} = \frac{1}{4} \begin{pmatrix} 2 \\ 0 \end{pmatrix} + \frac{1}{2} \begin{pmatrix} 1 \\ 2 \end{pmatrix}, \quad \begin{pmatrix} 2 \\ 0 \end{pmatrix} = -2 \begin{pmatrix} 1 \\ 2 \end{pmatrix} + 4 \begin{pmatrix} 1 \\ 1 \end{pmatrix}.$$

I vettori $v_1, \dots, v_k$ sono *linearmente indipendenti* se non sono linearmente dipendenti. Questa importante condizione può essere espressa nel modo seguente: dei vettori $v_1, \dots, v_k$ sono linearmente indipendenti se e solo se

$$\lambda_1 v_1 + \dots + \lambda_k v_k = 0 \implies \lambda_1 = \dots = \lambda_k = 0.$$

In altre parole, l'unica combinazione lineare dei $v_1, \dots, v_k$ che può dare il vettore nullo è quella banale in cui tutti i coefficienti $\lambda_1, \dots, \lambda_k$ sono nulli.

I casi $k = 1$ e $k = 2$ sono facili da studiare. Da quanto abbiamo detto, segue facilmente che:

- un vettore v_1 è dipendente $\iff v_1 = 0$;
- due vettori v_1, v_2 sono dipendenti $\iff$ sono multipli, cioè se esiste un $k \in \mathbb{K}$ tale che $v_1 = kv_2$ oppure $v_2 = kv_1$.

I vettori $v_1 = \begin{pmatrix} 1 \\ 1 \end{pmatrix}$ e $v_2 = \begin{pmatrix} -2 \\ -2 \end{pmatrix}$ di $\mathbb{R}^2$ sono dipendenti; $w_1 = \begin{pmatrix} 1 \\ 2 \end{pmatrix}$ e $w_2 = \begin{pmatrix} 2 \\ 1 \end{pmatrix}$ sono indipendenti perché non sono multipli.

Il caso con tre (o più) vettori v_1, v_2, v_3 è più complesso. I vettori

$$v_1 = \begin{pmatrix} 1 \\ 1 \\ 0 \end{pmatrix}, \quad v_2 = \begin{pmatrix} 0 \\ 1 \\ 1 \end{pmatrix}, \quad v_3 = \begin{pmatrix} 1 \\ 0 \\ -1 \end{pmatrix}$$

di $\mathbb{R}^3$ sono dipendenti, perché $v_1 - v_2 - v_3 = 0$. Notiamo che a coppie i tre vettori sono sempre indipendenti (non sono mai multipli), ma tutti e tre non lo sono, e questo non è certamente un fatto che si vede immediatamente come nel caso $k = 2$. D'altra parte, i vettori

$$e_1 = \begin{pmatrix} 1 \\ 0 \\ 0 \end{pmatrix}, \quad e_2 = \begin{pmatrix} 0 \\ 1 \\ 0 \end{pmatrix}, \quad e_3 = \begin{pmatrix} 0 \\ 0 \\ 1 \end{pmatrix}$$

sono indipendenti: se una combinazione lineare produce il vettore nullo

$$\lambda_1 e_1 + \lambda_2 e_2 + \lambda_3 e_3 = 0$$

allora riscriviamo entrambi i membri come vettori e scopriamo che

$$\begin{pmatrix} \lambda_1 \\ \lambda_2 \\ \lambda_3 \end{pmatrix} = \lambda_1 \begin{pmatrix} 1 \\ 0 \\ 0 \end{pmatrix} + \lambda_2 \begin{pmatrix} 0 \\ 1 \\ 0 \end{pmatrix} + \lambda_3 \begin{pmatrix} 0 \\ 0 \\ 1 \end{pmatrix} = \begin{pmatrix} 0 \\ 0 \\ 0 \end{pmatrix}.$$

Da questo deduciamo che $\lambda_1 = \lambda_2 = \lambda_3 = 0$. Quindi l'unica combinazione lineare di e_1, e_2, e_3 che dà il vettore nullo è quella banale. Quindi i tre vettori e_1, e_2, e_3 sono indipendenti.

Osservazione 2.3.3. Nell'esempio fatto precedentemente, ciascuno dei tre vettori dipendenti v_1, v_2, v_3 può essere scritto come combinazione degli altri due, infatti $v_1 = v_2 + v_3$, oppure $v_2 = v_1 - v_3$, oppure $v_3 = v_1 - v_2$. Possono capitare anche dei vettori $v_1, \dots, v_k$ dipendenti in cui solo alcuni dei k vettori possono essere combinazioni di altri. La costruzione di questi esempi (in $\mathbb{R}^3$ o in qualche altro spazio vettoriale) è lasciata per esercizio.

Proposizione 2.3.4. *Se $v_1, \dots, v_k$ sono indipendenti, allora qualsiasi sottoinsieme di $\{v_1, \dots, v_k\}$ è anch'esso formato da vettori indipendenti.*

Dimostrazione. Supponiamo per assurdo che alcuni dei $v_1, \dots, v_k$ siano fra loro dipendenti. Per comodità di notazione supponiamo siano i primi h . Allora esiste una combinazione lineare non banale nulla:

$$\lambda_1 v_1 + \dots + \lambda_h v_h = 0.$$

Questa si estende ad una combinazione lineare non banale nulla di tutti:

$$\lambda_1 v_1 + \dots + \lambda_h v_h + 0v_{h+1} + \dots + 0v_k = 0.$$

E quindi anche i vettori $v_1, \dots, v_k$ sarebbero dipendenti, un assurdo. $\square$

In particolare, ne deduciamo che se $v_1, \dots, v_k$ sono indipendenti allora:

- i vettori v_i sono tutti diversi da zero;
- i vettori v_i sono a coppie non multipli.

Attenzione però: l'esempio sopra con $v_1, v_2, v_3 \in \mathbb{R}^3$ mostra che queste due condizioni sono necessarie ma non sufficienti affinché i vettori $v_1, \dots, v_k$ siano indipendenti quando $k \geq 3$.

Esercizio 2.3.5. Considera i vettori seguenti di $\mathbb{R}^3$:

$$v_1 = \begin{pmatrix} 1 \\ 1 \\ 2 \end{pmatrix}, \quad v_2 = \begin{pmatrix} -1 \\ 1 \\ -1 \end{pmatrix}, \quad v_3 = \begin{pmatrix} 1 \\ 5 \\ 4 \end{pmatrix}, \quad v_4 = \begin{pmatrix} 0 \\ 1 \\ 1 \end{pmatrix}.$$

Mostra che v_1, v_2, v_3 sono dipendenti e v_1, v_2, v_4 indipendenti.

2.3.2. Basi. Introduciamo adesso una delle definizioni più importanti del libro. Sia V uno spazio vettoriale.

Definizione 2.3.6. Una sequenza $v_1, \dots, v_n \in V$ di vettori è una *base* se sono soddisfatte entrambe queste condizioni:

- (1) i vettori $v_1, \dots, v_n$ sono indipendenti;
- (2) i vettori $v_1, \dots, v_n$ generano V .

Ricordiamo che la seconda condizione vuol dire che $V = \text{Span}(v_1, \dots, v_n)$, cioè qualsiasi vettore di V è esprimibile come combinazione lineare dei vettori $v_1, \dots, v_n$. Facciamo adesso alcuni esempi fondamentali.

Proposizione 2.3.7. *Gli elementi*

$$e_1 = \begin{pmatrix} 1 \\ 0 \\ \vdots \\ 0 \end{pmatrix}, \quad e_2 = \begin{pmatrix} 0 \\ 1 \\ \vdots \\ 0 \end{pmatrix}, \quad \dots \quad e_n = \begin{pmatrix} 0 \\ 0 \\ \vdots \\ 1 \end{pmatrix}$$

formano una base di $\mathbb{K}^n$ detta base canonica.

Dimostrazione. Mostriamo che i vettori $e_1, \dots, e_n$ sono indipendenti. Supponiamo di avere una combinazione lineare nulla

$$\lambda_1 e_1 + \dots + \lambda_n e_n = 0.$$

Questa tradotta in vettori diventa

$$\begin{pmatrix} \lambda_1 \\ \lambda_2 \\ \vdots \\ \lambda_n \end{pmatrix} = \begin{pmatrix} 0 \\ 0 \\ \vdots \\ 0 \end{pmatrix}.$$

Ne deduciamo che $\lambda_1 = \dots = \lambda_n = 0$ e quindi $e_1, \dots, e_n$ sono indipendenti.

Mostriamo che i vettori $e_1, \dots, e_n$ generano $\mathbb{K}^n$. Un generico vettore $x \in \mathbb{K}^n$ si può effettivamente scrivere come combinazione lineare di

$e_1, \dots, e_n$ nel modo seguente:

$$x = \begin{pmatrix} x_1 \\ x_2 \\ \vdots \\ x_n \end{pmatrix} = x_1 \begin{pmatrix} 1 \\ 0 \\ \vdots \\ 0 \end{pmatrix} + x_2 \begin{pmatrix} 0 \\ 1 \\ \vdots \\ 0 \end{pmatrix} + \dots + x_n \begin{pmatrix} 0 \\ 0 \\ \vdots \\ 1 \end{pmatrix} = x_1 e_1 + x_2 e_2 + \dots + x_n e_n.$$

La dimostrazione è conclusa. $\square$

Ricordiamo lo spazio $\mathbb{K}_n[x]$ dei polinomi di grado minore o uguale a n .

Proposizione 2.3.8. *Gli elementi $1, x, x^2, \dots, x^n$ formano una base di $\mathbb{K}_n[x]$, detta base canonica.*

Dimostrazione. Lo spirito della dimostrazione è lo stesso di quella precedente. Dimostriamo che i vettori $1, x, \dots, x^n$ sono indipendenti: se

$$\lambda_0 \cdot 1 + \lambda_1 x + \dots + \lambda_n x^n = 0$$

allora chiaramente $\lambda_0 = \dots = \lambda_n = 0$. D'altro canto, ciascun polinomio di grado minore o uguale a n si scrive come

$$p(x) = a_n x^n + \dots + a_1 x + a_0$$

e questa scrittura è già una combinazione lineare dei $1, x, \dots, x^n$ con coefficienti $a_0, a_1, \dots, a_n$. Quindi gli elementi $1, x, \dots, x^n$ generano $\mathbb{K}_n[x]$. $\square$

Ricordiamo anche lo spazio $M(m, n, \mathbb{K})$ delle matrici $m \times n$ a coefficienti in $\mathbb{K}$. Per ogni $1 \leq i \leq m$ e $1 \leq j \leq n$, indichiamo con e_{ij} la matrice $m \times n$ che ha tutti zeri ovunque, tranne un 1 nella casella di riga i e colonna j . Ad esempio, nel caso delle matrici quadrate 2×2 troviamo:

$$e_{11} = \begin{pmatrix} 1 & 0 \\ 0 & 0 \end{pmatrix}, \quad e_{12} = \begin{pmatrix} 0 & 1 \\ 0 & 0 \end{pmatrix}, \quad e_{21} = \begin{pmatrix} 0 & 0 \\ 1 & 0 \end{pmatrix}, \quad e_{22} = \begin{pmatrix} 0 & 0 \\ 0 & 1 \end{pmatrix}.$$

Proposizione 2.3.9. *Le matrici e_{ij} con $1 \leq i \leq m$ e $1 \leq j \leq n$ formano una base di $M(m, n, \mathbb{K})$, detta base canonica.*

Dimostrazione. La dimostrazione è del tutto analoga a quella per $\mathbb{K}^n$, solo con la complicazione notazionale che abbiamo due indici i, j invece che uno solo. Mostriamo che le matrici e_{ij} sono indipendenti. Supponiamo di avere una combinazione lineare nulla

$$\sum_{i,j} \lambda_{ij} e_{ij} = 0.$$

Qui sommiamo su $1 \leq i \leq m$ e $1 \leq j \leq n$, quindi ci sono mn addendi. Esplicitando le matrici, l'equazione diventa

$$\begin{pmatrix} \lambda_{11} & \dots & \lambda_{1n} \\ \vdots & \ddots & \vdots \\ \lambda_{m1} & \dots & \lambda_{mn} \end{pmatrix} = \begin{pmatrix} 0 & \dots & 0 \\ \vdots & \ddots & \vdots \\ 0 & \dots & 0 \end{pmatrix}.$$

Da questa ricaviamo che $\lambda_{ij} = 0$ per ogni i, j . Mostriamo che le matrici e_{ij} generano: qualsiasi matrice A si scrive come

$$A = \begin{pmatrix} a_{11} & \cdots & a_{1n} \\ \vdots & \ddots & \vdots \\ a_{m1} & \cdots & a_{mn} \end{pmatrix} = \sum_{i,j} a_{ij} e_{ij}.$$

La dimostrazione è conclusa. $\square$

Abbiamo descritto delle basi *canoniche* per gli spazi vettoriali $\mathbb{K}^n$, $\mathbb{K}_n[x]$ e $M(m, n, \mathbb{K})$. Queste però non sono certamente le uniche basi disponibili: uno spazio vettoriale ha generalmente una infinità di basi diverse e come vedremo successivamente sarà spesso importante scegliere la base giusta per risolvere i problemi che troveremo.

Esercizio 2.3.10. I vettori $\begin{pmatrix} -1 \\ 1 \end{pmatrix}$ e $\begin{pmatrix} 2 \\ 1 \end{pmatrix}$ formano una base di $\mathbb{R}^2$.

2.3.3. Coordinate di un vettore rispetto ad una base. A cosa servono le basi? Servono per dare un nome a tutti i vettori dello spazio. A questo scopo useremo la proposizione seguente.

Proposizione 2.3.11. *Sia V uno spazio vettoriale e sia $v_1, \dots, v_n$ una base di V . Ogni vettore $v \in V$ si scrive in modo unico come*

$$v = \lambda_1 v_1 + \cdots + \lambda_n v_n.$$

Dimostrazione. Sappiamo che i $v_1, \dots, v_n$ generano, quindi sicuramente v si può scrivere come combinazione lineare dei $v_1, \dots, v_n$. Supponiamo che lo si possa fare in due modi diversi:

$$v = \lambda_1 v_1 + \cdots + \lambda_n v_n = \mu_1 v_1 + \cdots + \mu_n v_n.$$

Spostando tutto a sinistra otteniamo

$$(\lambda_1 - \mu_1)v_1 + \cdots + (\lambda_n - \mu_n)v_n = 0.$$

Siccome i vettori sono indipendenti, necessariamente $\mu_i = \lambda_i$ per ogni i . $\square$

I coefficienti $\lambda_1, \dots, \lambda_n$ sono le *coordinate* di v rispetto alla base.

Esempio 2.3.12. Le coordinate di un vettore $x \in \mathbb{K}^n$ rispetto alla base canonica $e_1, \dots, e_n$ di $\mathbb{K}^n$ sono i suoi coefficienti $x_1, \dots, x_n$, visto che

$$x = x_1 e_1 + \cdots + x_n e_n.$$

Se cambiamo base questo non è più vero. Ad esempio, se prendiamo $v_1 = \begin{pmatrix} 1 \\ 1 \end{pmatrix}$, $v_2 = \begin{pmatrix} -1 \\ 1 \end{pmatrix}$ come base di $\mathbb{R}^2$, il vettore $w = \begin{pmatrix} 2 \\ 0 \end{pmatrix}$ si scrive come $w = v_1 - v_2$ e quindi le sue coordinate rispetto alla base v_1, v_2 sono 1, -1.

Esempio 2.3.13. Le coordinate di un polinomio $p(x) = a_n x^n + \cdots + a_1 x + a_0$ rispetto alla base canonica $1, x, \dots, x^n$ sono $a_0, a_1, \dots, a_n$.

Esempio 2.3.14. Le coordinate di una matrice $\begin{pmatrix} a & b \\ c & d \end{pmatrix}$ rispetto alla base canonica $e_{11}, e_{12}, e_{21}, e_{22}$ di $M(2, \mathbb{K})$ sono a, b, c, d .

Esempio 2.3.15. I vettori $\begin{pmatrix} -1 \\ 1 \end{pmatrix}, \begin{pmatrix} 2 \\ 1 \end{pmatrix}$ sono una base di $\mathbb{R}^2$ per l'Esercizio 2.3.10. Descriviamo adesso un procedimento che ci permette di capire quali siano le coordinate di un vettore generico $\begin{pmatrix} x \\ y \end{pmatrix}$ rispetto a questa base. Scriviamo

$$\begin{pmatrix} x \\ y \end{pmatrix} = t \begin{pmatrix} -1 \\ 1 \end{pmatrix} + u \begin{pmatrix} 2 \\ 1 \end{pmatrix}$$

e troviamo i parametri t, u in funzione di x, y . Troviamo un sistema lineare

$$\begin{cases} -t + 2u = x, \\ t + u = y. \end{cases}$$

Nel prossimo capitolo descriveremo degli algoritmi per risolvere qualsiasi sistema lineare. Questo è abbastanza semplice e troviamo la soluzione

$$t = \frac{-x + 2y}{3}, \quad u = \frac{x + y}{3}.$$

Quindi le coordinate di $\begin{pmatrix} x \\ y \end{pmatrix}$ rispetto alla base scelta sono $\frac{-x+2y}{3}$ e $\frac{x+y}{3}$.

2.3.4. Dimensione. Lo scopo principale di questa sezione è dimostrare questo teorema.

Teorema 2.3.16. *Due basi dello stesso spazio vettoriale V contengono lo stesso numero n di elementi.*

Questo teorema è fondamentale, perché ci permette di definire in modo rigoroso un concetto intuitivo, quello di dimensione.

Definizione 2.3.17. Se uno spazio vettoriale V ha una base $v_1, \dots, v_n$, diciamo che V ha *dimensione* n . Se V non ha una base allora diciamo che ha dimensione ∞ .

La definizione è ben posta proprio grazie al Teorema 2.3.16: lo spazio vettoriale V ha tante basi, ma hanno tutte lo stesso numero n di elementi. La dimostrazione di questo teorema si basa sul lemma seguente, che avrà anche altre conseguenze interessanti.

Lemma 2.3.18. *Sia V uno spazio vettoriale. Supponiamo di avere:*

- *dei vettori $v_1, \dots, v_n \in V$ che generano V , e*
- *dei vettori $w_1, \dots, w_n \in V$ indipendenti.*

Allora anche i vettori $w_1, \dots, w_n$ generano V .

Dimostrazione. Sappiamo che $V = \text{Span}(v_1, \dots, v_n)$ e vogliamo mostrare che $V = \text{Span}(w_1, \dots, w_n)$. Otteniamo questo posizionando ciascun w_i al posto di un v_j , iterativamente per $i = 1, \dots, n$.

Operiamo nel modo seguente. I vettori $v_1, \dots, v_n$ generano V , quindi in particolare w_1 è combinazione lineare di questi:

$$w_1 = \lambda_1 v_1 + \dots + \lambda_n v_n.$$

Poiché $w_1 \neq 0$, almeno un λ_i è non nullo: a meno di riordinare i vettori $v_1, \dots, v_n$, supponiamo che $\lambda_1 \neq 0$. Allora dividendo per λ_1 otteniamo

$$v_1 = \frac{1}{\lambda_1} w_1 - \frac{\lambda_2}{\lambda_1} v_2 - \dots - \frac{\lambda_n}{\lambda_1} v_n.$$

Quindi v_1 è combinazione lineare dei $w_1, v_2, \dots, v_n$. Questo implica facilmente che i vettori $w_1, v_2, \dots, v_n$ generano V , cioè possiamo sostituire v_1 con w_1 e

$$V = \text{Span}(w_1, v_2, \dots, v_n).$$

Iteriamo questo procedimento per n volte. Al passo numero s , supponiamo di aver già dimostrato che $V = \text{Span}(w_1, \dots, w_{s-1}, v_s, \dots, v_n)$. Allora

$$w_s = \lambda_1 w_1 + \dots + \lambda_{s-1} w_{s-1} + \lambda_s v_s + \dots + \lambda_n v_n.$$

Esiste almeno un $i \geq s$ con $\lambda_i \neq 0$. Se così non fosse, questa sarebbe una relazione di dipendenza fra i soli $w_1, \dots, w_s$, ma questo è escluso perché sono indipendenti. A meno di riordinare i vettori supponiamo che $\lambda_s \neq 0$. Possiamo dividere tutto per λ_s e spostare gli addendi in modo da ottenere:

$$v_s = \frac{1}{\lambda_s} w_s - \frac{\lambda_1}{\lambda_s} w_1 - \dots - \frac{\lambda_{s-1}}{\lambda_s} w_{s-1} - \frac{\lambda_{s+1}}{\lambda_s} v_{s+1} - \dots - \frac{\lambda_n}{\lambda_s} v_n.$$

Quindi come sopra possiamo sostituire v_s con w_s e ottenere l'uguaglianza $V = \text{Span}(w_1, \dots, w_s, v_{s+1}, \dots, v_n)$. Dopo n passi troviamo $V = \text{Span}(w_1, \dots, w_n)$ e la dimostrazione è conclusa. $\square$

Possiamo adesso dedurre facilmente il Teorema 2.3.16.

Dimostrazione. Supponiamo per assurdo che uno spazio vettoriale V contenga due basi

$$v_1, \dots, v_n, \quad w_1, \dots, w_m$$

con $n \neq m$. Sia ad esempio $n < m$.

Per ipotesi i vettori $v_1, \dots, v_n$ generano V e i vettori $w_1, \dots, w_m$ sono indipendenti. Dal lemma precedente deduciamo che $w_1, \dots, w_n$ generano V . Quindi w_{n+1} può essere espresso come combinazione lineare dei $w_1, \dots, w_n$, e allora i vettori $w_1, \dots, w_m$ non sono indipendenti: assurdo. $\square$

Indichiamo la dimensione di uno spazio vettoriale V con il simbolo $\dim V$. Possiamo calcolare la dimensione di tutti gli spazi vettoriali dei quali abbiamo esibito una base. Troviamo:

$$\dim \mathbb{K}^n = n,$$

$$\dim \mathbb{K}_n[x] = n + 1,$$

$$\dim M(m, n, \mathbb{K}) = mn.$$

Ricordiamo che $\dim V = \infty$ se V non ha una base.

Proposizione 2.3.19. *Lo spazio $\mathbb{K}[x]$ ha dimensione infinita.*

Dimostrazione. Supponiamo per assurdo che lo spazio $\mathbb{K}[x]$ abbia una base $p_1(x), \dots, p_n(x)$. Sia N il massimo dei gradi dei $p_1(x), \dots, p_n(x)$. È chiaro che il polinomio x^{N+1} non può essere ottenuto come combinazione lineare dei $p_1(x), \dots, p_n(x)$, e questo è un assurdo. $\square$

Notiamo infine che per definizione uno spazio vettoriale V ha dimensione 0 se e solo se consiste solo dell'origine $V = \{0\}$.

2.3.5. Algoritmo di completamento. Abbiamo definito la dimensione di uno spazio vettoriale V come il numero di elementi in una sua base. Adesso abbiamo bisogno di algoritmi concreti per determinare una base di V in vari contesti. Iniziamo con una proposizione utile.

Proposizione 2.3.20. *Sia V uno spazio vettoriale e $v_1, \dots, v_k \in V$ dei vettori indipendenti. Sia $v_{k+1} \in V$. Allora:*

$$v_1, \dots, v_{k+1} \text{ sono indipendenti} \iff v_{k+1} \notin \text{Span}(v_1, \dots, v_k).$$

Dimostrazione. $(\Rightarrow)$ Se per assurdo $v_{k+1} \in \text{Span}(v_1, \dots, v_k)$, allora otteniamo $v_{k+1} = \lambda_1 v_1 + \dots + \lambda_k v_k$ e quindi i vettori $v_1, \dots, v_{k+1}$ sono dipendenti.

$(\Leftarrow)$ Supponiamo che

$$\lambda_1 v_1 + \dots + \lambda_{k+1} v_{k+1} = 0.$$

Se $\lambda_{k+1} = 0$, otteniamo una combinazione lineare nulla dei primi $v_1, \dots, v_k$, quindi $\lambda_1 = \dots = \lambda_k = 0$. Se $\lambda_{k+1} \neq 0$, dividiamo tutto per λ_{k+1} ottenendo una dipendenza di v_{k+1} dai precedenti, che è assurdo. $\square$

Descriviamo adesso l'*algoritmo di completamento a base*. Sia V uno spazio vettoriale di dimensione n . Siano $v_1, \dots, v_k \in V$ dei vettori indipendenti qualsiasi, con $k < n$. Possiamo sempre completare $v_1, \dots, v_k$ ad una base di V nel modo seguente. Siccome $k < n$, i vettori $v_1, \dots, v_k$ non sono una base di V , in particolare non generano V . Allora esiste un $v_{k+1} \in V$ tale che

$$v_{k+1} \notin \text{Span}(v_1, \dots, v_k).$$

Aggiungiamo v_{k+1} alla lista, che adesso diventa $v_1, \dots, v_{k+1}$. La nuova lista è ancora formata da vettori indipendenti per la Proposizione 2.3.20. Continuiamo finché non otteniamo n vettori indipendenti $v_1, \dots, v_n$. Questi devono essere una base per il Lemma 2.3.18.

Esempio 2.3.21. Consideriamo i vettori indipendenti in $\mathbb{R}^3$

$$v_1 = \begin{pmatrix} 1 \\ 1 \\ 0 \end{pmatrix}, \quad v_2 = \begin{pmatrix} -1 \\ 0 \\ 1 \end{pmatrix}.$$

Per completare la coppia v_1, v_2 ad una base di $\mathbb{R}^3$ è sufficiente aggiungere un terzo vettore qualsiasi che non sia contenuto nel piano $\text{Span}(v_1, v_2)$. Ad esempio, $v_3 = e_1$ funziona (ricordiamo che e_1, e_2, e_3 sono i vettori della base canonica di $\mathbb{R}^3$). Quindi v_1, v_2, e_1 è una base di $\mathbb{R}^3$.

2.3.6. Algoritmo di estrazione. Descriviamo adesso l'*algoritmo di estrazione di una base*. Sia V uno spazio vettoriale e siano $v_1, \dots, v_m$ dei generatori di V . Vogliamo estrarre da questo insieme di generatori una base per V .

L'algoritmo funziona nel modo seguente. Se i generatori $v_1, \dots, v_m$ sono indipendenti, allora formano già una base per V e siamo a posto. Altrimenti, esiste almeno un v_i che può essere espresso come combinazione degli altri. Se lo rimuoviamo dalla lista, otteniamo una lista di $m-1$ vettori che continuano a generare V . Dopo un numero finito di passi otteniamo una lista di vettori indipendenti, e questi sono una base di V .

Esempio 2.3.22. Consideriamo i vettori di $\mathbb{R}^3$

$$v_1 = \begin{pmatrix} 1 \\ 1 \\ 0 \end{pmatrix}, \quad v_2 = \begin{pmatrix} 0 \\ 1 \\ 1 \end{pmatrix}, \quad v_3 = \begin{pmatrix} -1 \\ 1 \\ 1 \end{pmatrix}, \quad v_4 = \begin{pmatrix} 0 \\ 0 \\ 1 \end{pmatrix}.$$

Si verifica facilmente che questi generano $\mathbb{R}^3$, ma non sono indipendenti, ad esempio $v_1 - 2v_2 + v_3 + v_4 = 0$. Ciascuno dei quattro vettori è esprimibile come combinazione degli altri tre: rimuovendo uno qualsiasi di questi quattro vettori otteniamo una base di $\mathbb{R}^3$.

Concludiamo con una proposizione che risulta utile in molti casi concreti quando dobbiamo mostrare che un certo insieme di vettori è una base. In generale, per dimostrare che n vettori sono una base di V , dobbiamo dimostrare che generano e che sono indipendenti. La proposizione seguente dice che, se sappiamo già che $\dim V = n$, allora una qualsiasi delle due proprietà è sufficiente, perché implica l'altra.

Proposizione 2.3.23. *Siano $v_1, \dots, v_n$ vettori di uno spazio V di dimensione n . I fatti seguenti sono tutti equivalenti:*

- (1) $v_1, \dots, v_n$ generano V ;
- (2) $v_1, \dots, v_n$ sono indipendenti;
- (3) $v_1, \dots, v_n$ sono una base per V .

Dimostrazione. (3) $\Rightarrow$ (2) è ovvio. (2) $\Rightarrow$ (1) è il Lemma 2.3.18.

(1) $\Rightarrow$ (3). Se $v_1, \dots, v_n$ non fossero indipendenti, potremmo estrarre da questi una base con meno di n vettori, contraddicendo il Teorema 2.3.16. $\square$

Usando questo criterio possiamo dimostrare il fatto seguente.

Proposizione 2.3.24. *I vettori*

$$v_1 = \begin{pmatrix} a_{11} \\ 0 \\ 0 \\ \vdots \\ 0 \end{pmatrix}, \quad v_2 = \begin{pmatrix} a_{12} \\ a_{22} \\ 0 \\ \vdots \\ 0 \end{pmatrix}, \quad v_3 = \begin{pmatrix} a_{13} \\ a_{23} \\ a_{33} \\ \vdots \\ 0 \end{pmatrix}, \quad \dots, \quad v_n = \begin{pmatrix} a_{1n} \\ a_{2n} \\ a_{3n} \\ \vdots \\ a_{nn} \end{pmatrix}$$

sono una base di $\mathbb{K}^n \iff i$ numeri $a_{11}, \dots, a_{nn}$ sono tutti diversi da zero.

Dimostrazione. Mostriamo che

$$v_1, \dots, v_n \text{ sono indipendenti} \iff a_{11}, \dots, a_{nn} \neq 0.$$

Supponiamo che

$$\lambda_1 v_1 + \dots + \lambda_n v_n = 0.$$

Questa uguaglianza fra vettori è equivalente al sistema lineare omogeneo

$$\left\{ \begin{array}{l} a_{11}\lambda_1 + a_{12}\lambda_2 + a_{13}\lambda_3 + \dots + a_{1n}\lambda_n = 0, \\ a_{22}\lambda_2 + a_{23}\lambda_3 + \dots + a_{2n}\lambda_n = 0, \\ a_{33}\lambda_3 + \dots + a_{3n}\lambda_n = 0, \\ \vdots \\ a_{nn}\lambda_n = 0. \end{array} \right.$$

Mostriamo per induzione su n che

$$\exists! \text{ soluzione } \lambda_1 = \dots = \lambda_n = 0 \iff a_{11}, \dots, a_{nn} \neq 0.$$

Per $n = 1$ il sistema è semplicemente l'equazione $a_{11}\lambda_1 = 0$ ed è chiaro che se $a_{11} \neq 0$ allora $\lambda_1 = 0$ è l'unica soluzione, mentre se $a_{11} = 0$ allora qualsiasi numero $\lambda_1 = t \in \mathbb{K}$ è soluzione.

Supponiamo l'asserzione vera per $n - 1$ e la dimostriamo per n . Se $a_{nn} \neq 0$, allora l'ultima equazione $a_{nn}\lambda_n = 0$ implica che $\lambda_n = 0$. Sostituendo $\lambda_n = 0$ in tutte le equazioni precedenti ci riconduciamo al caso $n - 1$ e concludiamo per l'ipotesi induttiva. Se invece $a_{nn} = 0$, allora $\lambda_1 = 0, \dots, \lambda_{n-1} = 0, \lambda_n = t$ è soluzione per ogni $t \in \mathbb{K}$ e quindi non è vero che il sistema ha un'unica soluzione. Abbiamo concluso. $\square$

2.3.7. Sottospazi vettoriali. La proposizione seguente ci racconta un fatto molto intuitivo: uno spazio non può contenere un sottospazio di dimensione più grande di lui.

Proposizione 2.3.25. Sia V uno spazio vettoriale e $U \subset V$ un sottospazio. Vale $0 \leq \dim U \leq \dim V$. Inoltre:

- $\dim U = 0 \iff U = \{0\}$.
- $\dim U = \dim V \iff U = V$.

Dimostrazione. Sia $v_1, \dots, v_k$ base di U . Per l'algoritmo di completamento, questa può essere estesa ad una base $v_1, \dots, v_n$ di V . In particolare $n \geq k$, cioè $\dim V \geq \dim U$. Inoltre $k = 0 \iff U = \{0\}$ e $n = k \iff U = V$ $\square$

La proposizione ci dice anche che gli unici sottospazi di V di dimensione minima 0 e massima $\dim V$ sono precisamente il sottospazio banale $\{0\}$ ed il sottospazio totale V stesso.

Corollario 2.3.26. Uno spazio vettoriale V di dimensione 1 contiene solo due sottospazi: il sottospazio banale $\{0\}$ e il sottospazio totale V .

Gli spazi vettoriali di dimensione 1 e 2 vengono chiamati *rette* e *piani vettoriali*. I sottospazi di $\mathbb{R}^3$ sono, per dimensione crescente: $\{0\}$, le rette vettoriali, i piani vettoriali, e $\mathbb{R}^3$ stesso.

2.3.8. Formula di Grassmann. Concludiamo questa sezione con una formula che lega le dimensioni di due sottospazi, della loro intersezione e della loro somma.

Teorema 2.3.27 (Formula di Grassmann). *Sia V uno spazio vettoriale di dimensione finita. Siano $U, W \subset V$ sottospazi. Vale la formula*

$$\dim(U + W) + \dim(U \cap W) = \dim U + \dim W.$$

Dimostrazione. Sia $v_1, \dots, v_k$ una base di $U \cap W$. Grazie all'algoritmo di completamento, possiamo estenderla ad una base di U :

$$v_1, \dots, v_k, u_1, \dots, u_s$$

e anche ad una base di W :

$$v_1, \dots, v_k, w_1, \dots, w_t.$$

Adesso vogliamo dimostrare che

$$v_1, \dots, v_k, u_1, \dots, u_s, w_1, \dots, w_t$$

è una base di $U + W$. Se ci riusciamo, abbiamo finito perché otteniamo

$$\begin{aligned} \dim(U + W) &= k + s + t = (k + s) + (k + t) - k \\ &= \dim U + \dim W - \dim U \cap W. \end{aligned}$$

Per dimostrare che quella è una base, dobbiamo verificare che i vettori generino $U + W$ e che siano indipendenti. Per mostrare che generano, consideriamo un vettore generico $u + w \in U + W$. Poiché u si scrive come combinazione dei $v_1, \dots, v_k, u_1, \dots, u_s$ e w dei $v_1, \dots, v_k, w_1, \dots, w_t$, è chiaro che $u + w$ si scrive come combinazione dei $v_1, \dots, v_k, u_1, \dots, u_s, w_1, \dots, w_t$.

Per mostrare l'indipendenza, supponiamo di avere una combinazione lineare nulla:

$$\underbrace{\lambda_1 v_1 + \dots + \lambda_k v_k}_v + \underbrace{\mu_1 u_1 + \dots + \mu_s u_s}_u + \underbrace{\eta_1 w_1 + \dots + \eta_t w_t}_w = 0.$$

Per semplicità chiamiamo v, u, w i vettori indicati, così otteniamo:

$$v + u + w = 0.$$

Da questa deduciamo che

$$w = -v - u.$$

Notiamo che $w \in W$ e $-v - u \in U$. Siccome sono uguali, entrambi w e $-v - u$ stanno in $U \cap W$. In particolare $w \in U \cap W$. Allora w può essere scritto come combinazione lineare dei $v_1, \dots, v_k$, del tipo

$$\eta_1 w_1 + \dots + \eta_t w_t = \alpha_1 v_1 + \dots + \alpha_k v_k.$$

Spostando tutto a sinistra otteniamo

$$-\alpha_1 v_1 - \cdots - \alpha_k v_k + \eta_1 w_1 + \cdots + \eta_t w_t = 0.$$

Siccome $v_1, \dots, v_k, w_1, \dots, w_t$ è una base di W , i coefficienti $\alpha_1, \dots, \alpha_k, \eta_1, \dots, \eta_t$ sono tutti nulli. Quindi $w = 0$ e allora anche $v + u = 0$, cioè

$$\lambda_1 v_1 + \cdots + \lambda_k v_k + \mu_1 u_1 + \cdots + \mu_s u_s = 0.$$

Analogamente, siccome $v_1, \dots, v_k, u_1, \dots, u_s$ è base di U , anche i coefficienti $\lambda_1, \dots, \lambda_k, \mu_1, \dots, \mu_s$ sono nulli. La dimostrazione è conclusa. $\square$

2.3.9. Somma diretta. Sia V uno spazio vettoriale. Diciamo che due sottospazi $U, W \subset V$ sono in *somma diretta* se $U \cap W = \{0\}$. In questo caso indichiamo la loro somma con $U \oplus W$. Dalla formula di Grassmann ricaviamo

$$\dim U \oplus W = \dim U + \dim W.$$

Siamo particolarmente interessati al caso $V = U \oplus W$. Possiamo esprimere il fatto che $V = U \oplus W$ in vari modi equivalenti.

Proposizione 2.3.28. *Siano $U, W \subset V$ sottospazi. Sono equivalenti:*

- (1) $V = U \oplus W$.
- (2) $\dim V = \dim(U + W) = \dim U + \dim W$.
- (3) *Per qualsiasi scelta di basi $u_1, \dots, u_k$ di U e $w_1, \dots, w_h$ di W , l'unione $u_1, \dots, u_k, w_1, \dots, w_h$ è una base di V .*
- (4) *Ogni vettore $v \in V$ si scrive in modo unico come $v = u + w$ per qualche $u \in U$ e $w \in W$. Si veda la Figura 2.5.*

Dimostrazione. (1) $\Rightarrow$ (2) È la formula di Grassmann.

(2) $\Rightarrow$ (3) L'unione delle basi genera $U+W$, che ha dimensione proprio $k+h$. Per la Proposizione 2.3.23, $k+h$ generatori sono una base di $U+W$. Inoltre $U+W = V$ perché hanno la stessa dimensione e $U+W \subset V$.

(3) $\Rightarrow$ (4) Ogni vettore si scrive in modo unico come combinazione di $u_1, \dots, u_k, w_1, \dots, w_h$ e quindi in modo unico come $u + w$.

(4) $\Rightarrow$ (1) Otteniamo subito $V = U + W$. Se per assurdo ci fosse un $v \in U \cap W$ diverso da zero, questo si scriverebbe in modo non unico come $v = v + 0 = 0 + v$. $\square$

In una somma diretta $V = U \oplus W$, ogni vettore $v \in V$ si scrive in modo unico come somma $v = u + w$ con $u \in U$ e $w \in W$. Una rappresentazione grafica di questo fatto è mostrata nella Figura 2.5.

Notiamo ora un criterio che può essere utile per dimostrare che $V = U \oplus W$, un po' analogo alla Proposizione 2.3.23.

Proposizione 2.3.29. *Siano $U, W \subset V$ due sottospazi. Supponiamo che $\dim V = \dim U + \dim W$. I fatti seguenti sono tutti equivalenti:*

- (1) $V = U \oplus W$,
- (2) $V = U + W$,
- (3) $U \cap W = \{0\}$.

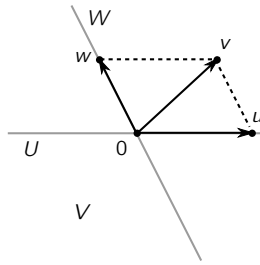


Figura 2.5. Una somma diretta $V = U \oplus W$. Ogni vettore $v \in V$ si scrive in modo unico come $v = u + w$ dove $u \in U$ e $w \in W$.

Dimostrazione. Per la formula di Grassmann abbiamo

$$\dim(U + W) + \dim(U \cap W) = \dim U + \dim W = \dim V.$$

Da questa formula deduciamo che (2) $\Leftrightarrow$ (3), infatti

$$V = U + W \Leftrightarrow \dim V = \dim(U + W) \Leftrightarrow \dim(U \cap W) = 0 \Leftrightarrow U \cap W = \{0\}.$$

La dimostrazione è conclusa. $\square$

Se sappiamo che $\dim V = \dim U + \dim W$, per dimostrare che $V = U \oplus W$ è sufficiente verificare che $V = U + W$ oppure $U \cap W = \{0\}$. Spesso la seconda condizione è più semplice da provare.

Esempio 2.3.30. Se U e W sono due rette vettoriali in $V = \mathbb{R}^2$, sappiamo che $\dim V = 2 = 1 + 1 = \dim U + \dim W$. Se sono distinte, si intersecano solo nell'origine. Ne deduciamo che se U e W sono due rette distinte qualsiasi in $\mathbb{R}^2$ vale sempre $\mathbb{R}^2 = U \oplus W$.

Analogamente, siano U un piano e W una retta vettoriali in $\mathbb{R}^3$. Se questi si intersecano solo nell'origine allora vale $\mathbb{R}^3 = U \oplus W$.

Esempio 2.3.31. In $\mathbb{K}^n$ i sottospazi

$$U = \text{Span}(e_1, \dots, e_k), \quad W = \text{Span}(e_{k+1}, \dots, e_n)$$

sono in somma diretta e $\mathbb{K}^n = U \oplus W$.

Esempio 2.3.32. Consideriamo $V = \mathbb{R}^2$ e due rette

$$U = \text{Span} \begin{pmatrix} 1 \\ 0 \end{pmatrix}, \quad W = \text{Span} \begin{pmatrix} -1 \\ 2 \end{pmatrix}.$$

Sappiamo che $\mathbb{R}^2 = U \oplus W$ e quindi ogni vettore $\begin{pmatrix} x \\ y \end{pmatrix} \in \mathbb{R}^2$ si scrive in modo unico come combinazione lineare di un vettore in U e un vettore in W . Effettivamente:

$$\begin{pmatrix} x \\ y \end{pmatrix} = \left(x + \frac{y}{2}\right) \begin{pmatrix} 1 \\ 0 \end{pmatrix} + \frac{y}{2} \begin{pmatrix} -1 \\ 2 \end{pmatrix}.$$

2.3.10. Trasposta di una matrice. Per il prossimo esempio di somma diretta abbiamo bisogno di introdurre una semplice operazione sulle matrici.

La *trasposta* di una matrice $A \in M(m, n, \mathbb{K})$ è la matrice

$${}^tA \in M(n, m, \mathbb{K})$$

definita scambiando righe e colonne, cioè:

$$({}^tA)_{ij} = A_{ji}.$$

Ad esempio:

$$A = \begin{pmatrix} 2 & 1 \\ -1 & 0 \\ 5 & 7 \end{pmatrix} \implies {}^tA = \begin{pmatrix} 2 & -1 & 5 \\ 1 & 0 & 7 \end{pmatrix}.$$

Valgono le proprietà seguenti:

$${}^t(A + B) = {}^tA + {}^tB, \quad {}^t(\lambda A) = \lambda {}^tA.$$

Se $A \in M(n)$, allora anche ${}^tA \in M(n)$. Notiamo che A è simmetrica $\iff {}^tA = A$ e inoltre A è antisimmetrica $\iff {}^tA = -A$.

2.3.11. Matrici simmetriche e antisimmetriche. Torniamo alle somme dirette di sottospazi: vogliamo descrivere un esempio con le matrici.

Esempio 2.3.33. Ricordiamo che $S(n)$ e $A(n)$ sono le matrici $n \times n$ simmetriche e antisimmetriche. Dimostriamo che

$$M(n) = S(n) \oplus A(n).$$

Dobbiamo verificare due cose:

- $M(n) = S(n) + A(n)$. Infatti ogni matrice $B \in M(n)$ si scrive come somma di una matrice simmetrica e di una antisimmetrica con il seguente trucco:

$$B = \frac{B + {}^tB}{2} + \frac{B - {}^tB}{2}.$$

Effettivamente si vede facilmente che il primo addendo è simmetrico ed il secondo antisimmetrico.

- $S(n) \cap A(n) = \{0\}$. Infatti solo la matrice nulla è contemporaneamente simmetrica e antisimmetrica.

Esercizio 2.3.34. Mostra che:

- Le matrici $e_{ij} + e_{ji}$ con $i < j$ e e_{ii} formano una base di $S(n)$,
- Le matrici $e_{ij} - e_{ji}$ con $i < j$ formano una base di $A(n)$.

Ad esempio, le matrici

$$\begin{pmatrix} 0 & 1 \\ 1 & 0 \end{pmatrix}, \quad \begin{pmatrix} 1 & 0 \\ 0 & 0 \end{pmatrix}, \quad \begin{pmatrix} 0 & 0 \\ 0 & 1 \end{pmatrix}$$

formano una base di $S(2)$, che ha dimensione 3, mentre la matrice

$$\begin{pmatrix} 0 & 1 \\ -1 & 0 \end{pmatrix}$$

forma una base di $A(2)$, che ha dimensione 1. Deduci che

$$\dim S(n) = 1 + \cdots + n = \frac{n(n+1)}{2}, \quad \dim A(n) = 1 + \cdots + n - 1 = \frac{(n-1)n}{2}.$$

Nota che effettivamente $\dim S(n) + \dim A(n) = n^2 = \dim M(n)$.

2.3.12. Somma (diretta) di più sottospazi. La nozione di somma diretta è molto utile e la ritroveremo varie volte in seguito. È possibile estenderla da 2 ad un numero arbitrario k di sottospazi.

Iniziamo estendendo l'operazione di somma. Sia V uno spazio vettoriale e $V_1, \dots, V_k \subset V$ alcuni sottospazi. Indichiamo con

$$V_1 + \cdots + V_k$$

il sottoinsieme di V formato da tutti i vettori del tipo

$$v_1 + \cdots + v_k$$

dove $v_i \in V_i$. Si verifica come nel caso $k = 2$ che $V_1 + \cdots + V_k$ è un sottospazio vettoriale di V .

Diremo che i sottospazi $V_1, \dots, V_k$ sono in *somma diretta* se

$$V_i \cap (V_1 + \cdots + V_{i-1} + V_{i+1} + \cdots + V_k) = \{0\} \quad \forall i = 1, \dots, k.$$

Si richiede che ciascun sottospazio intersechi la somma di tutti gli altri soltanto nell'origine. In questo caso scriviamo la loro somma con il simbolo

$$V_1 \oplus \cdots \oplus V_k.$$

La definizione di somma diretta sembra inutilmente tecnica, ma è quella giusta, come mostra il seguente fatto, analogo alla Proposizione 2.3.28.

Proposizione 2.3.35. *Siano $V_1, \dots, V_k \subset V$ alcuni sottospazi. I fatti seguenti sono equivalenti:*

- (1) $V = V_1 \oplus \cdots \oplus V_k$.
- (2) $\dim V = \dim(V_1 + \cdots + V_k) = \dim V_1 + \cdots + \dim V_k$.
- (3) *Per qualsiasi scelta di basi $v_1^i, \dots, v_{h_i}^i$ di ciascun V_i , la loro unione al variare di $i = 1, \dots, k$ è una base di V .*
- (4) *Ogni vettore $v \in V$ si scrive in modo unico come*

$$v = v_1 + \cdots + v_k$$

per qualche $v_i \in V_i$.

Dimostrazione. (1) $\Rightarrow$ (2) La prima uguaglianza è ovvia; per la seconda, mostriamo che

$$\dim(V_1 + \cdots + V_i) = \dim V_1 + \cdots + \dim V_i$$

per ogni $i = 1, \dots, k$ per induzione su i . Per $i = 1$ non c'è nulla da dimostrare, ed il caso generico è trattato nel modo seguente:

$$\begin{aligned}\dim(V_1 + \dots + V_i) &= \dim((V_1 + \dots + V_{i-1}) + V_i) \\ &= \dim(V_1 + \dots + V_{i-1}) + \dim V_i \\ &= \dim V_1 + \dots + \dim V_{i-1} + \dim V_i.\end{aligned}$$

Nella seconda uguaglianza abbiamo usato la formula di Grasmann, con

$$V_i \cap (V_1 + \dots + V_{i-1} + V_{i+1} + \dots + V_k) = \{0\} \implies V_i \cap (V_1 + \dots + V_{i-1}) = \{0\}.$$

Nella terza uguaglianza abbiamo usato l'ipotesi induttiva.

(2) $\Rightarrow$ (3) L'unione delle basi genera $V_1 + \dots + V_k$ ed è formata esattamente da $\dim(V_1 + \dots + V_k)$ vettori, quindi è una base di $V_1 + \dots + V_k$.

(3) $\Rightarrow$ (4) Ogni elemento $v \in V$ si scrive in modo unico come combinazione lineare dell'unione delle basi e quindi anche come combinazione lineare $v = v_1 + \dots + v_k$ con $v_i \in V_i$.

(4) $\Rightarrow$ (1) Otteniamo subito $V = V_1 + \dots + V_k$. Se per assurdo ci fosse un $v \in V_i \cap (V_1 + \dots + V_{i-1} + V_{i+1} + \dots + V_k)$ diverso da zero, questo si scriverebbe in modo non unico come somma di elementi $v_j \in V_j$.

La dimostrazione è conclusa. $\square$

Esempio 2.3.36. Consideriamo $V = \mathbb{R}^5$ e i sottospazi $V_1 = \text{Span}(e_1)$, $V_2 = \text{Span}(e_2, e_3)$ e $V_3 = \text{Span}(e_4, e_5)$. Allora $V = V_1 \oplus V_2 \oplus V_3$.

Osservazione 2.3.37. Quando $k \geq 3$, non è sufficiente che valga $V_1 \cap \dots \cap V_k = \{0\}$ affinché gli spazi siano in somma diretta. Ad esempio, un qualsiasi insieme di k rette vettoriali distinte in $\mathbb{R}^2$ soddisfa questa ipotesi, ma per $k \geq 3$ queste rette non sono mai in somma diretta.

Esercizi

Esercizio 2.1. Sia V uno spazio vettoriale su $\mathbb{K}$ e X un insieme. Mostra che l'insieme $F(X, V)$ ha una naturale struttura di spazio vettoriale su $\mathbb{K}$.

Esercizio 2.2. Siano V e W due spazi vettoriali su $\mathbb{K}$. Definisci su $V \times W$ la somma $(v, w) + (v', w') = (v + v', w + w')$ ed il prodotto per scalare $\lambda(v, w) = (\lambda v, \lambda w)$. Mostra che $V \times W$ con queste operazioni è uno spazio vettoriale.

Esercizio 2.3. Siano $v_1, \dots, v_k \in V$. Mostra che $\text{Span}(v_1, \dots, v_k)$ è l'intersezione di tutti i sottospazi di V contenenti i vettori $v_1, \dots, v_k$.

Esercizio 2.4. Siano $v_1, \dots, v_k \in V$ indipendenti. Mostra che

$$\text{Span}(v_1, \dots, v_h) \cap \text{Span}(v_{h+1}, \dots, v_k) = \{0\}.$$

Esercizio 2.5. Sia V uno spazio vettoriale e $U, W \subset V$ sottospazi. Mostra che:

- (1) $U \cap W$ è l'unione di tutti i sottospazi contenuti in U e W
- (2) $U + W$ è l'intersezione di tutti i sottospazi che contengono U e W .

Esercizio 2.6. Determina delle basi dei sottospazi seguenti di $\mathbb{R}^3$:

$$U = \{x + y + z = 0\}, \quad W = \{x - 2y - z = 0\}, \quad U \cap W, \quad U + W.$$

Esercizio 2.7. Sia $a \in \mathbb{K}$ un valore fissato. Sia $W \subset \mathbb{K}_n[x]$ il sottoinsieme formato dai polinomi che si annullano in a . Trova una base per W e calcolane la dimensione.

Esercizio 2.8. Trova una base di $W = \{x_1 + x_2 + \cdots + x_n = 0\} \subset \mathbb{R}^n$.

Esercizio 2.9. Completa i vettori seguenti a base di $\mathbb{R}^4$:

$$v_1 = \begin{pmatrix} 1 \\ 1 \\ -1 \\ -1 \end{pmatrix}, \quad v_2 = \begin{pmatrix} 0 \\ -1 \\ 1 \\ 0 \end{pmatrix}$$

Esercizio 2.10. Mostra che le matrici seguenti di $M(2, \mathbb{R})$ sono indipendenti e completale a base di $M(2, \mathbb{R})$.

$$\begin{pmatrix} 1 & -1 \\ 0 & 1 \end{pmatrix}, \quad \begin{pmatrix} 0 & 1 \\ -1 & 1 \end{pmatrix}, \quad \begin{pmatrix} -1 & 1 \\ -1 & 0 \end{pmatrix}.$$

Scrivi le coordinate della matrice $e_{11} = \begin{pmatrix} 1 & 0 \\ 0 & 0 \end{pmatrix}$ rispetto alla base che hai trovato.

Esercizio 2.11. Sia $W \subset M(2, \mathbb{R})$ il sottospazio generato dalle tre matrici descritte nell'esercizio precedente. Per quali $t \in \mathbb{R}$ la matrice $\begin{pmatrix} 0 & 1 \\ t & 0 \end{pmatrix}$ appartiene a W ?

Esercizio 2.12. Completa i seguenti polinomi a base di $\mathbb{C}_3[x]$:

$$x + i, \quad x^2 - ix.$$

Scrivi le coordinate del polinomio x^2 rispetto alla base che hai trovato.

Esercizio 2.13. Estrai da questa lista di generatori una base di $M(2, \mathbb{C})$:

$$\begin{pmatrix} 1 & i \\ 0 & 1 \end{pmatrix}, \quad \begin{pmatrix} i & -1 \\ 1 & 0 \end{pmatrix}, \quad \begin{pmatrix} 0 & -1 \\ 0 & 0 \end{pmatrix}, \quad \begin{pmatrix} 0 & 0 \\ i & 0 \end{pmatrix}, \quad \begin{pmatrix} -i & 1+i \\ 0 & 1 \end{pmatrix}.$$

Esercizio 2.14. Scrivi le coordinate dei vettori

$$\begin{pmatrix} 1 \\ 0 \end{pmatrix}, \quad \begin{pmatrix} 0 \\ 1 \end{pmatrix}, \quad \begin{pmatrix} 2 \\ 7 \end{pmatrix}$$

rispetto alla base $\begin{pmatrix} 1 \\ 2 \end{pmatrix}, \begin{pmatrix} -1 \\ 1 \end{pmatrix}$ di $\mathbb{R}^2$.

Esercizio 2.15. Dimostra che $\begin{pmatrix} a \\ b \end{pmatrix}, \begin{pmatrix} c \\ d \end{pmatrix}$ è una base di $\mathbb{R}^2 \iff ad - bc \neq 0$.

Esercizio 2.16. Mostra che i vettori $v_1 = e_1, v_2 = e_1 + e_2, \dots, v_n = e_1 + \cdots + e_n$ sono una base di $\mathbb{R}^n$.

Esercizio 2.17. Sia X un insieme finito di n elementi. Per ogni $x \in X$ definiamo la funzione $\delta_x: X \rightarrow \mathbb{K}$ in questo modo:

$$\delta_x(y) = \begin{cases} 1 & \text{se } y = x, \\ 0 & \text{se } y \neq x. \end{cases}$$

Mostra che le funzioni δ_x al variare di $x \in X$ sono una base di $F(X, \mathbb{K})$. Deduci che $\dim F(X, \mathbb{K}) = n$.

Esercizio 2.18. Siano V e W spazi vettoriali di dimensione finita. Mostra che $\dim(V \times W) = \dim V + \dim W$ (vedi l'Esercizio 2.2).

Esercizio 2.19. Sia X un insieme con n elementi e V uno spazio vettoriale di dimensione m . Mostra che $\dim F(X, V) = mn$ (vedi l'Esercizio 2.1).

Esercizio 2.20. Sia V uno spazio vettoriale di dimensione n . Mostra che per ogni $0 \leq k \leq n$ esiste un sottospazio $W \subset V$ di dimensione k .

Esercizio 2.21. Sia V uno spazio vettoriale di dimensione finita e $U \subset V$ un sottospazio. Un *complementare algebrico* per U è un sottospazio $W \subset V$ tale che $V = U \oplus W$. Mostra che esiste sempre un complementare algebrico per U .

Esercizio 2.22. Determina 4 vettori $v_1, v_2, v_3, v_4 \in \mathbb{R}^4$ con queste proprietà:

- (1) Se elimini un qualsiasi vettore, i tre rimanenti sono sempre indipendenti.
- (2) I sottospazi $U = \text{Span}(v_1, v_2)$ e $W = \text{Span}(v_3, v_4)$ non sono in somma diretta.

Esercizio 2.23. Siano U e W sottospazi di $\mathbb{R}^5$, entrambi di dimensione 3. Quali sono le possibili dimensioni di $U \cap W$? Descrivi degli esempi.

Esercizio 2.24. Considera i due sottospazi di $\mathbb{R}^3$, dipendenti da una variabile $t \in \mathbb{R}$:

$$U = \text{Span} \begin{pmatrix} 1 \\ 1 \\ 1 \end{pmatrix}, \quad W_t = \text{Span} \left(\begin{pmatrix} 1 \\ 0 \\ t \end{pmatrix}, \begin{pmatrix} 0 \\ 1 \\ 2 \end{pmatrix} \right).$$

Determina per quali valori di $t \in \mathbb{R}$ vale $U \oplus W = \mathbb{R}^3$.

Esercizio 2.25. Siano U una retta vettoriale e W un piano vettoriale in $\mathbb{R}^3$. Dimostra che i fatti seguenti sono tutti equivalenti:

$$\mathbb{R}^3 = U \oplus W \iff U \cap W = \{0\} \iff (U \not\subseteq W \text{ e } W \not\subseteq U).$$

Esercizio 2.26. Mostra che qualsiasi spazio vettoriale V di dimensione n contiene n sottospazi $V_1, \dots, V_n$ tutti di dimensione 1 tali che

$$V = V_1 \oplus \dots \oplus V_n.$$

Esercizio 2.27. Siano U, V, W sottospazi di $\mathbb{R}^4$, tutti di dimensione 3. Siano $a = \dim(U \cap V)$, $b = \dim(V \cap W)$, $c = \dim(U \cap W)$, $d = \dim(U \cap V \cap W)$.

- (1) Dimostra che $a \in \{2, 3\}$.
- (2) Dimostra che $d > 0$.
- (3) Costruisci un esempio di U, V, W in cui $d = 2$.
- (4) Costruisci un esempio di U, V, W in cui $d = 1$.
- (5) Determina tutti i possibili valori per la quadrupla (a, b, c, d) al variare di U, V, W .

Sistemi lineari

Nel capitolo precedente abbiamo usato alcuni strumenti algebrici per definire dei concetti geometrici, come quelli di spazio vettoriale, retta, piano e dimensione. In questo capitolo ci muoviamo un po' in direzione opposta: usiamo le nozioni geometriche appena introdotte per studiare un problema algebrico. Il problema algebrico in questione è la determinazione delle soluzioni di un sistema di equazioni lineari.

L'analisi dei sistemi lineari ci porterà naturalmente allo studio delle matrici. La teoria delle matrici è alla base della matematica e della scienza moderne, e porta con sé un folto numero di nozioni e strumenti, fra i quali spiccano il *rango* ed il *determinante*.

3.1. Algoritmi di risoluzione

Introduciamo in questa sezione un algoritmo che può essere usato per risolvere qualsiasi sistema lineare. L'ingrediente principale consiste in alcune mosse con le quali possiamo modificare una matrice in modo controllato, note come *mosse di Gauss*.

3.1.1. Mosse di Gauss. Un *sistema lineare* è un insieme di k equazioni lineari in n variabili

$$\begin{cases} a_{11}x_1 + \cdots + a_{1n}x_n = b_1, \\ \vdots \\ a_{k1}x_1 + \cdots + a_{kn}x_n = b_k. \end{cases}$$

I numeri a_{ij} sono i *coefficienti* e i b_i sono i *termini noti* del sistema. Sia i coefficienti, che i termini noti, che le variabili sono in un certo campo fissato $\mathbb{K}$. Possiamo raggruppare i coefficienti e i termini noti efficacemente in una matrice $k \times n$ e un vettore colonna:

$$A = \begin{pmatrix} a_{11} & \cdots & a_{1n} \\ \vdots & \ddots & \vdots \\ a_{k1} & \cdots & a_{kn} \end{pmatrix}, \quad b = \begin{pmatrix} b_1 \\ \vdots \\ b_k \end{pmatrix}.$$

Possiamo quindi unire tutto in un'unica matrice $k \times (n+1)$:

$$C = (A \mid b)$$

Il nostro scopo è determinare l'insieme $S \subset \mathbb{K}^n$ delle soluzioni del sistema. A questo scopo, notiamo che esistono alcune mosse che cambiano la matrice C senza modificare l'insieme S delle soluzioni del sistema. Queste sono le seguenti e sono note come *mosse di Gauss*:

- (I) Scambiare due righe.
- (II) Moltiplicare una riga per un numero $\lambda \neq 0$.
- (III) Aggiungere ad una riga un'altra riga moltiplicata per λ qualsiasi.

Indicando con C_i la i -esima riga di C , possiamo scrivere le mosse così:

- (I) $C_i \longleftrightarrow C_j$,
- (II) $C_i \longrightarrow \lambda C_i$, $\lambda \neq 0$,
- (III) $C_i \longrightarrow C_i + \lambda C_j$.

Alcuni esempi di composizioni di mosse di tipo (I), (II) e (III):

$$\begin{pmatrix} 1 & 2 \\ 3 & 0 \\ 4 & -2 \end{pmatrix} \xrightarrow{(I)} \begin{pmatrix} 4 & -2 \\ 3 & 0 \\ 1 & 2 \end{pmatrix} \xrightarrow{(II)} \begin{pmatrix} 2 & -1 \\ 3 & 0 \\ 1 & 2 \end{pmatrix} \xrightarrow{(III)} \begin{pmatrix} 2 & -1 \\ 5 & 4 \\ 1 & 2 \end{pmatrix}.$$

Torniamo al sistema lineare iniziale, determinato dalla matrice $C = (A | b)$.

Proposizione 3.1.1. *Le mosse di Gauss su $C = (A | b)$ non mutano l'insieme $S \subset \mathbb{K}^n$ delle soluzioni del sistema lineare.*

Dimostrazione. La mossa (I) scambia due equazioni e questo chiaramente non tocca S . La mossa (II) moltiplica una equazione lineare

$$a_{i1}x_1 + \cdots + a_{in}x_n = b_i$$

per una costante $\lambda \neq 0$, trasformandola quindi in

$$\lambda a_{i1}x_1 + \cdots + \lambda a_{in}x_n = \lambda b_i.$$

Anche qui le soluzioni sono chiaramente le stesse di prima. L'effetto della mossa (III) sulle equazioni lineari corrispondenti alle righe i e j è il seguente: cambiamo le due equazioni

$$\begin{aligned} a_{i1}x_1 + \cdots + a_{in}x_n &= b_i \\ a_{j1}x_1 + \cdots + a_{jn}x_n &= b_j \end{aligned}$$

con

$$\begin{aligned} a_{i1}x_1 + \cdots + a_{in}x_n &= b_i \\ (a_{j1} + \lambda a_{i1})x_1 + \cdots + (a_{jn} + \lambda a_{in})x_n &= b_j + \lambda b_i \end{aligned}$$

È facile convincersi che $x \in \mathbb{K}^n$ è soluzione delle prime due se e solo se è soluzione delle seconde due. $\square$

3.1.2. Algoritmo di Gauss. Introduciamo un algoritmo che può essere usato per risolvere qualsiasi sistema lineare. L'algoritmo trasforma una arbitraria matrice C in una matrice particolare detta *a scalini*.

Sia C una matrice qualsiasi. Per ogni riga C_i di C , chiamiamo *pivot* il primo elemento non nullo della riga C_i . Una *matrice a scalini* è una matrice avente la proprietà seguente: il pivot di ogni riga è sempre strettamente

più a destra del pivot della riga precedente. Ad esempio, sono matrici a scalini:

$$\begin{pmatrix} 1 & 0 & 5 \\ 0 & -1 & -1 \end{pmatrix}, \quad \begin{pmatrix} 7 & -2 & 9 \\ 0 & 0 & 1 \\ 0 & 0 & 0 \end{pmatrix}$$

mentre invece *non* sono matrici a scalini:

$$\begin{pmatrix} 1 & 2 & -1 \\ 4 & 0 & 6 \\ 0 & 7 & 0 \end{pmatrix}, \quad \begin{pmatrix} 0 & 8 \\ 0 & 1 \\ 0 & 0 \\ 0 & 0 \end{pmatrix}.$$

L'*algoritmo di Gauss* è un algoritmo che trasforma qualsiasi matrice C in una matrice a scalini usando le 3 mosse di Gauss descritte precedentemente. Funziona così:

- (1) Se $C_{11} = 0$, scambia se puoi la prima riga con un'altra riga in modo da ottenere $C_{11} \neq 0$. Se $C_{i1} = 0$ per ogni i , passa al punto (3).
- (2) Per ogni riga C_i con $i \geq 2$ e con $C_{i1} \neq 0$ sostituisci C_i con la riga

$$C_i - \frac{C_{i1}}{C_{11}} C_1.$$

In questo modo la nuova riga C_i avrà $C_{i1} = 0$.

- (3) Hai ottenuto $C_{i1} = 0$ per ogni $i \geq 2$. Continua dal punto (1) lavorando sulla sottomatrice ottenuta togliendo la prima riga e la prima colonna.

Vediamo un esempio. Partiamo dalla matrice

$$C = \begin{pmatrix} 0 & 1 & 1 & 0 \\ 1 & 1 & 2 & -3 \\ -1 & 2 & 1 & 1 \end{pmatrix}.$$

Notiamo che $C_{11} = 0$. Quindi scambiamo C_1 e C_2 in modo da ottenere $C_{11} \neq 0$. Ad ogni passaggio chiamiamo sempre la nuova matrice con la lettera C per semplicità. Otteniamo

$$C = \begin{pmatrix} 1 & 1 & 2 & -3 \\ 0 & 1 & 1 & 0 \\ -1 & 2 & 1 & 1 \end{pmatrix}.$$

Ora notiamo che $C_{31} \neq 0$. Cambiamo C_3 con $C_3 + C_1$ in modo da ottenere $C_{31} = 0$. Il risultato è

$$C = \begin{pmatrix} 1 & 1 & 2 & -3 \\ 0 & 1 & 1 & 0 \\ 0 & 3 & 3 & -2 \end{pmatrix}.$$

La prima riga e la prima colonna adesso sono a posto, quindi lavoriamo sul resto. Notiamo che $C_{32} \neq 0$. Sostituiamo C_3 con $C_3 - 3C_2$ in modo da

ottenere $C_{32} = 0$. Il risultato è

$$C = \begin{pmatrix} 1 & 1 & 2 & -3 \\ 0 & 1 & 1 & 0 \\ 0 & 0 & 0 & -2 \end{pmatrix}.$$

Adesso la matrice è a scalini e l'algoritmo termina.

3.1.3. Algoritmo di Gauss - Jordan. Per un motivo che sarà chiaro in seguito, dopo aver ridotto la matrice a scalini, è a volte utile fare delle ulteriori mosse di Gauss per fare in modo che tutti i numeri *sopra* i pivot siano nulli. Questo è sempre possibile usando mosse di Gauss di tipo (III), aggiungendo la riga che contiene il pivot (moltiplicata per una costante opportuna) alle righe precedenti.

Ad esempio, nella matrice C ottenuta precedentemente abbiamo $C_{12} \neq 0$. Per ottenere $C_{12} = 0$ possiamo sostituire C_1 con $C_1 - C_2$. Il risultato è una nuova matrice

$$C = \begin{pmatrix} 1 & 0 & 1 & -3 \\ 0 & 1 & 1 & 0 \\ 0 & 0 & 0 & -2 \end{pmatrix}.$$

Adesso notiamo che $C_{14} \neq 0$. Per ottenere $C_{14} = 0$ sostituiamo C_1 con $C_1 - \frac{3}{2}C_3$ e otteniamo infine

$$C = \begin{pmatrix} 1 & 0 & 1 & 0 \\ 0 & 1 & 1 & 0 \\ 0 & 0 & 0 & -2 \end{pmatrix}.$$

Come ultimo ritocco, vedremo che sarà anche utile ottenere che tutti i pivot abbiano valore 1. È sempre possibile ottenere questo semplicemente con mosse di Gauss di tipo (II).

Nel nostro esempio, i pivot C_{11} e C_{22} hanno già valore 1, però $C_{34} = -2$. Per ottenere $C_{34} = 1$ dividiamo la terza riga per -2. Il risultato finale è

$$C = \begin{pmatrix} 1 & 0 & 1 & 0 \\ 0 & 1 & 1 & 0 \\ 0 & 0 & 0 & 1 \end{pmatrix}.$$

L'algoritmo che abbiamo appena descritto si chiama *algoritmo di Gauss - Jordan*. Consiste in due fasi:

- (1) Trasformare la matrice a scalini tramite algoritmo di Gauss.
- (2) Ottenere solo zeri sopra i pivot con mosse (III) e tutti i pivot uguali a 1 con mosse (II).

Facciamo un altro esempio di applicazione dell'algoritmo di Gauss - Jordan. Sopra ogni freccia indichiamo la mossa di Gauss corrispondente.

$$\begin{pmatrix} 1 & -1 & 3 \\ 0 & 2 & 2 \\ 1 & 0 & 4 \end{pmatrix} \xrightarrow{C_3 \rightarrow C_3 - C_1} \begin{pmatrix} 1 & -1 & 3 \\ 0 & 2 & 2 \\ 0 & 1 & 1 \end{pmatrix} \xrightarrow{C_3 \rightarrow C_3 - \frac{1}{2}C_2} \begin{pmatrix} 1 & -1 & 3 \\ 0 & 2 & 2 \\ 0 & 0 & 0 \end{pmatrix}$$

$$\xrightarrow{C_1 \rightarrow C_1 + \frac{1}{2}C_2} \begin{pmatrix} 1 & 0 & 4 \\ 0 & 2 & 2 \\ 0 & 0 & 0 \end{pmatrix} \xrightarrow{C_2 \rightarrow \frac{1}{2}C_2} \begin{pmatrix} 1 & 0 & 4 \\ 0 & 1 & 1 \\ 0 & 0 & 0 \end{pmatrix}$$

L'algoritmo di Gauss - Jordan è garantito funzionare sempre; d'altra parte, per ottenere una matrice a scalini (eventualmente con zeri sopra i pivot e pivot uguali a 1) non è necessario seguire pedissequamente l'algoritmo: qualsiasi combinazione di mosse di Gauss è lecita purché si arrivi al risultato.

3.1.4. Risoluzione di un sistema lineare. Mostriamo adesso come sia possibile descrivere esplicitamente lo spazio $S \subset \mathbb{K}^n$ delle soluzioni di un dato sistema lineare.

Il sistema è descritto da una matrice $C = (A \mid b)$. Usando l'algoritmo di Gauss - Jordan, possiamo trasformare C in una matrice a scalini, in cui ogni pivot ha valore 1 e tutti i numeri sopra i pivot sono nulli. La matrice C sarà indicativamente di questo tipo:

$$\left(\begin{array}{cccc|ccc} 0 & 1 & ? & 0 & 0 & ? & ? \\ 0 & 0 & 0 & 1 & 0 & ? & ? \\ 0 & 0 & 0 & 0 & 1 & ? & ? \end{array} \right).$$

Adesso guardiamo le colonne che contengono i pivot, che in questo esempio sono la seconda, la quarta e la quinta. Ciascuna di queste colonne contiene un 1 al posto del pivot e 0 in tutte le altre caselle.

Ci sono due casi da considerare. Se la colonna dei termini noti b contiene un pivot, allora la matrice è indicativamente di questo tipo:

$$\left(\begin{array}{cccc|ccc} 0 & 1 & ? & 0 & 0 & ? & ? \\ 0 & 0 & 0 & 1 & 0 & ? & ? \\ 0 & 0 & 0 & 0 & 0 & 0 & 1 \end{array} \right).$$

La riga che contiene l'ultimo pivot rappresenta l'equazione $0 = 1$ che chiaramente non può avere soluzione. In questo caso l'insieme S delle soluzioni è vuoto, cioè $S = \emptyset$.

Consideriamo adesso il caso in cui l'ultima colonna non contenga pivot. La matrice C è indicativamente di questo tipo:

$$\left(\begin{array}{cccc|ccc} 0 & 1 & a_{13} & 0 & 0 & a_{16} & b_1 \\ 0 & 0 & 0 & 1 & 0 & a_{26} & b_2 \\ 0 & 0 & 0 & 0 & 1 & a_{36} & b_3 \end{array} \right).$$

Ricordiamo che ciascuna colonna corrisponde ad una variabile $x_1, \dots, x_n$, eccetto l'ultima che contiene i termini noti. Adesso assegniamo un parametro $t_1, t_2, \dots$ a ciascuna variabile che corrisponde ad una colonna che

non contiene un pivot. Nel nostro caso, scriviamo $x_1 = t_1$, $x_3 = t_2$ e $x_6 = t_3$. Conseguentemente il sistema lineare è di questa forma:

$$\begin{cases} x_2 + a_{13}t_2 + a_{16}t_3 = b_1, \\ x_4 + a_{26}t_3 = b_2, \\ x_5 + a_{36}t_3 = b_3. \end{cases}$$

Spostiamo a destra dell'uguale tutti i nuovi parametri e otteniamo:

$$\begin{cases} x_2 = b_1 - a_{13}t_2 - a_{16}t_3, \\ x_4 = b_2 - a_{26}t_3, \\ x_5 = b_3 - a_{36}t_3. \end{cases}$$

Al sistema aggiungiamo anche le equazioni corrispondenti alle assegnazioni dei parametri liberi, che nel nostro caso sono $x_1 = t_1$, $x_3 = t_2$ e $x_6 = t_3$. Il risultato è un sistema del tipo:

$$\begin{cases} x_1 = t_1, \\ x_2 = b_1 - a_{13}t_2 - a_{16}t_3, \\ x_3 = t_2, \\ x_4 = b_2 - a_{26}t_3, \\ x_5 = b_3 - a_{36}t_3, \\ x_6 = t_3. \end{cases}$$

Il sistema è risolto. I parametri $t_1, t_2, \dots$ sono liberi e possono assumere qualsiasi valore in $\mathbb{K}$, e le variabili $x_1, \dots, x_n$ dipendono da questi parametri liberi come indicato.

Notiamo che è sempre possibile usare una notazione parametrica vettoriale per esprimere le soluzioni, del tipo

$$\begin{pmatrix} x_1 \\ x_2 \\ x_3 \\ x_4 \\ x_5 \\ x_6 \end{pmatrix} = \begin{pmatrix} t_1 \\ b_1 - a_{13}t_2 - a_{16}t_3 \\ t_2 \\ b_2 - a_{26}t_3 \\ b_3 - a_{36}t_3 \\ t_3 \end{pmatrix} = \begin{pmatrix} 0 \\ b_1 \\ 0 \\ b_2 \\ b_3 \\ 0 \end{pmatrix} + t_1 \begin{pmatrix} 1 \\ 0 \\ 0 \\ 0 \\ 0 \\ 0 \end{pmatrix} + t_2 \begin{pmatrix} 0 \\ -a_{13} \\ 1 \\ 0 \\ 0 \\ 0 \end{pmatrix} + t_3 \begin{pmatrix} 0 \\ -a_{16} \\ 0 \\ -a_{26} \\ -a_{36} \\ 1 \end{pmatrix}.$$

In generale, le soluzioni sono sempre del tipo

$$x = x_0 + t_1 v_1 + \dots + t_h v_h$$

dove $x_0, v_1, \dots, v_h \in \mathbb{K}^n$ sono vettori fissati, e i v_i sono tutti della forma

$$v_i = \begin{pmatrix} ? \\ \vdots \\ ? \\ 1 \\ 0 \\ \vdots \\ 0 \end{pmatrix}$$

con la posizione della cifra 1 che scende al crescere di i . I vettori $v_1, \dots, v_h$ sono sempre indipendenti per la Proposizione 2.3.24 e inoltre $h = n - r$ dove n è il numero di variabili e r è il numero di pivot.

Esempio 3.1.2. Risolviamo il sistema seguente, già ridotto con l'algoritmo di Gauss - Jordan:

$$\begin{cases} x_1 + 3x_2 + 4x_5 = 1, \\ x_3 - 2x_4 = 3. \end{cases}$$

Le variabili corrispondenti ai pivot sono x_1 e x_3 . Le variabili libere sono x_2 , x_4 e x_5 . Quindi scriviamo

$$\begin{cases} x_1 = 1 - 3t_1 - 4t_3, \\ x_2 = t_1, \\ x_3 = 3 + 2t_2, \\ x_4 = t_2, \\ x_5 = t_3. \end{cases}$$

3.2. Teorema di Rouché - Capelli

Nella sezione precedente abbiamo visto come risolvere un sistema lineare tramite l'algoritmo di Gauss - Jordan. In questa studiamo il problema in modo più teorico. Introduciamo alcune definizioni importanti: quella di *sottospazio affine* di uno spazio vettoriale e di *rango* di una matrice. Infine dimostreremo il *Teorema di Rouché - Capelli* che caratterizza geometricamente l'insieme delle soluzioni di un sistema lineare.

3.2.1. Il sistema omogeneo associato. Consideriamo un sistema di equazioni lineari

$$(1) \quad \begin{cases} a_{11}x_1 + \dots + a_{1n}x_n = b_1, \\ \vdots \\ a_{k1}x_1 + \dots + a_{kn}x_n = b_k. \end{cases}$$

Il *sistema omogeneo associato* è quello ottenuto semplicemente mettendo a zero tutti i termini noti b_i , cioè:

$$(2) \quad \begin{cases} a_{11}x_1 + \dots + a_{1n}x_n = 0, \\ \vdots \\ a_{k1}x_1 + \dots + a_{kn}x_n = 0. \end{cases}$$

Scriviamo la matrice $A = (a_{ij})$ dei coefficienti ed il vettore $b = (b_i)$ dei termini noti. La matrice $C = (A|b)$ rappresenta il sistema lineare iniziale (1), mentre $(A|0)$, oppure più semplicemente A , rappresenta il sistema lineare omogeneo associato (2).

Sia $S \subset \mathbb{K}^n$ l'insieme delle soluzioni del sistema lineare iniziale e $S_0 \subset \mathbb{K}^n$ l'insieme delle soluzioni del sistema omogeneo associato. Ricordiamo che S_0 è un sottospazio vettoriale di $\mathbb{K}^n$ per la Proposizione 2.2.2, mentre S invece non lo è perché non contiene l'origine – a meno che l'insieme

lineare iniziale non sia già omogeneo e in questo caso $S = S_0$, si veda l'Osservazione 2.2.3.

I due insiemi S e S_0 sono strettamente correlati:

Proposizione 3.2.1. *Se $S \neq \emptyset$, allora S è ottenuto prendendo una qualsiasi soluzione $x \in S$ e aggiungendo a questa tutti i vettori di S_0 .*

Dimostrazione. Sia $x \in S$ una soluzione del sistema (1) e $x' \in S_0$ una soluzione del sistema (2). Si vede facilmente che $x + x'$ è anch'essa soluzione di (1), infatti

$$a_{i1}(x_1 + x'_1) + \cdots + a_{in}(x_n + x'_n) = b_i + 0 = b_i.$$

Se invece x'' è soluzione di (1), allora $x' = x'' - x$ è soluzione di (2) perché

$$a_{i1}(x''_1 - x_1) + \cdots + a_{in}(x''_n - x_n) = b_i - b_i = 0.$$

Quindi le soluzioni di (1) si ottengono tutte precisamente aggiungendo ad un fissato $x \in S$ le soluzioni $x' \in S_0$ di (2). $\square$

La soluzione x è a volte detta *soluzione particolare*. Le soluzioni di un sistema si ottengono aggiungendo ad una soluzione particolare tutte le soluzioni del sistema omogeneo associato.

Esempio 3.2.2. Consideriamo il sistema in $\mathbb{R}^3$

$$\begin{cases} x - y + z = 1, \\ y - z = 2. \end{cases}$$

Risolvendolo troviamo che le soluzioni S sono del tipo

$$\begin{pmatrix} 3 \\ t \\ t-2 \end{pmatrix} = \begin{pmatrix} 3 \\ 0 \\ -2 \end{pmatrix} + \begin{pmatrix} 0 \\ t \\ t \end{pmatrix}$$

al variare di $t \in \mathbb{R}$. Le soluzioni S_0 del sistema omogeneo associato

$$\begin{cases} x - y + z = 0, \\ y - z = 0 \end{cases}$$

sono precisamente i vettori del tipo

$$\begin{pmatrix} 0 \\ t \\ t \end{pmatrix}.$$

Le soluzioni in S sono tutte ottenute prendendo la soluzione particolare

$$\begin{pmatrix} 3 \\ 0 \\ -2 \end{pmatrix}$$

ed aggiungendo a questa tutte le soluzioni in S_0 del sistema omogeneo associato. Qualsiasi elemento di S può fungere da soluzione particolare:

ad esempio potremmo prendere

$$\begin{pmatrix} 3 \\ 2 \\ 0 \end{pmatrix}$$

e quindi descrivere lo stesso insieme S di soluzioni in questo modo:

$$\begin{pmatrix} 3 \\ 2 \\ 0 \end{pmatrix} + \begin{pmatrix} 0 \\ u \\ u \end{pmatrix} = \begin{pmatrix} 3 \\ 2+u \\ u \end{pmatrix}.$$

Si ottiene effettivamente lo stesso insieme S di prima, sostituendo $t = u + 2$.

3.2.2. Sottospazio affine. Cerchiamo adesso di rispondere alla seguente domanda: che tipo di luogo geometrico è l'insieme S di soluzioni di un sistema lineare? La risposta è che S è sempre un *sottospazio affine*.

Sia V uno spazio vettoriale. Un *sottospazio affine* di V è un qualsiasi sottoinsieme del tipo

$$S = \{x + v \mid v \in W\}$$

dove x è un punto fissato di V e $W \subset V$ è un sottospazio vettoriale.

Esempio 3.2.3. Abbiamo appena visto che le soluzioni S di un sistema di equazioni lineari formano l'insieme vuoto oppure un sottospazio affine di $\mathbb{K}^n$. Infatti se $S \neq \emptyset$ allora

$$S = \{x + v \mid v \in S_0\}$$

dove x è un punto qualsiasi di S e S_0 è l'insieme delle soluzioni del sistema lineare omogeneo associato, che è sempre un sottospazio vettoriale.

Per indicare un sottospazio affine S , usiamo semplicemente la scrittura

$$S = x + W = \{x + v \mid v \in W\}.$$

Lo stesso sottospazio affine può essere descritto in più modi diversi. Chiamiamo subito quando due scritture differenti descrivono lo stesso sottospazio:

Proposizione 3.2.4. *Gli spazi affini $x + W$ e $x' + W'$ coincidono se e solo se $W = W'$ e $x - x' \in W$.*

Dimostrazione. Se $W = W'$ e $x - x' \in W$, allora chiaramente ogni vettore $x + v$ con $v \in W$ può essere scritto anche come $x' + (x - x' + v)$ con $x - x' + v \in W$ e quindi $x + W \subset x' + W$. Analogamente $x + W \supset x' + W$.

D'altro canto, se $x + W = x' + W'$ allora $W = x' - x + W'$. Questo implica che $x' - x \in W$. Analogamente $x' - x \in W'$ e quindi $W = x' - x + W' = W'$. $\square$

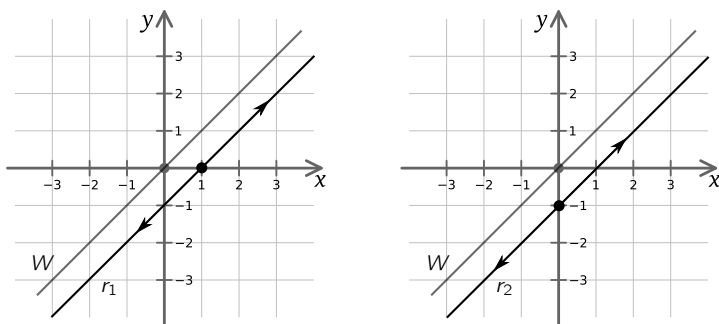


Figura 3.1. Le rette affini r_1 e r_2 sono la stessa retta.

Esempio 3.2.5. Se $W = \text{Span}\begin{pmatrix} 1 \\ 1 \end{pmatrix}$ in $\mathbb{R}^2$, le due rette affini

$$r_1 = \begin{pmatrix} 1 \\ 0 \end{pmatrix} + W = \left\{ \begin{pmatrix} 1+t \\ t \end{pmatrix} \mid t \in \mathbb{R} \right\},$$

$$r_2 = \begin{pmatrix} 0 \\ -1 \end{pmatrix} + W = \left\{ \begin{pmatrix} u \\ u-1 \end{pmatrix} \mid u \in \mathbb{R} \right\}$$

sono in realtà la stessa retta di equazione $y = x - 1$. Si veda la Figura 3.1.

Nella descrizione di uno spazio affine S come $x+W$, lo spazio vettoriale W è determinato da S ed è detto la *giacitura* di S . Il punto x invece è un qualsiasi punto di S .

Definizione 3.2.6. La *dimensione* di S è la dimensione della giacitura W .

Un sottospazio affine di dimensione 0 è un qualsiasi punto di V . Un sottospazio affine di dimensione 1 o 2 è detto *retta affine* e *piano affine*. Lo spazio $\mathbb{R}^3$ contiene ovviamente molte rette affini e molti piani affini. Questi verranno studiati più approfonditamente nel Capitolo 9.

Tornando ai nostri sistemi lineari, abbiamo appena scoperto che lo spazio delle soluzioni S , se non è vuoto, è un sottospazio affine con una certa dimensione. La sua giacitura è lo spazio S_0 delle soluzioni del sistema omogeneo associato. Vale $\dim S = \dim S_0$ per definizione.

3.2.3. Rango. Sappiamo come risolvere un sistema lineare e che l'insieme S delle soluzioni può essere vuoto o un sottospazio affine di una certa dimensione. Cerchiamo adesso delle tecniche per capire rapidamente se un sistema lineare abbia soluzioni, e in caso affermativo la loro dimensione.

Introduciamo una nozione che comparirà spesso in questo libro. Sia A una matrice $m \times n$ a coefficienti in $\mathbb{K}$. Ricordiamo che indichiamo con $A^1, \dots, A^n$ le colonne di A . Ciascun A^i è un vettore in $\mathbb{K}^m$.

Definizione 3.2.7. Il *rango* di A è la dimensione dello spazio

$$\text{Span}(A^1, \dots, A^n) \subset \mathbb{K}^m.$$

Il rango di A è la dimensione dello spazio generato dalle colonne. Viene comunemente indicato con $\text{rk}(A)$. Una formulazione equivalente è la seguente.

Proposizione 3.2.8. *Il rango di A è il massimo numero di colonne linearmente indipendenti di A .*

Dimostrazione. È un fatto generale che la dimensione di un sottospazio $W = \text{Span}(v_1, \dots, v_k)$ è il massimo numero di vettori indipendenti che possiamo trovare fra $v_1, \dots, v_k$. Segue dall'algoritmo di estrazione. $\square$

Notiamo adesso che il rango è insensibile alle mosse di Gauss.

Proposizione 3.2.9. *Se modifichiamo A per mosse di Gauss sulle righe, il suo rango non cambia.*

Dimostrazione. Se esiste una relazione di dipendenza lineare fra alcune colonne, questa si conserva con le mosse di Gauss sulle righe (esercizio). Quindi il massimo numero di colonne indipendenti non cambia. $\square$

Corollario 3.2.10. *Il rango di A è il numero di pivot in una sua qualsiasi riduzione a scalini.*

Dimostrazione. Poiché il rango non cambia con mosse di Gauss, possiamo supporre che A sia già ridotta a scalini e inoltre usare la seconda parte dell'algoritmo di Gauss-Jordan per fare in modo che le k colonne contenenti i pivot siano precisamente i primi k vettori $e_1, \dots, e_k$ della base canonica di $\mathbb{K}^m$. Tutte le altre colonne sono combinazioni lineari di questi. Lo spazio generato dalle colonne è quindi $\text{Span}(e_1, \dots, e_k)$ ed ha dimensione k pari al numero di pivot. $\square$

Corollario 3.2.11. *Se A è una matrice $m \times n$ allora $\text{rk}A \leq \min\{m, n\}$.*

Una matrice A ha *rango massimo* se $\text{rk}A = \min\{m, n\}$.

3.2.4. Teorema di Rouché - Capelli. Enunciamo e dimostriamo un teorema che fornisce un algoritmo pratico ed efficace per capire se un sistema lineare ha soluzioni, e in caso affermativo la loro dimensione.

Consideriamo come al solito un sistema lineare

$$(3) \quad \begin{cases} a_{11}x_1 + \dots + a_{1n}x_n = b_1, \\ \vdots \\ a_{k1}x_1 + \dots + a_{kn}x_n = b_k. \end{cases}$$

Scriviamo come sempre la matrice $A = (a_{ij})$ dei coefficienti, il vettore $b = (b_i)$ dei termini noti e la matrice completa $C = (A | b)$. Iniziamo con una proposizione.

Proposizione 3.2.12. *Il sistema (3) ha soluzioni se e solo se*

$$b \in \text{Span}(A^1, \dots, A^n).$$

Dimostrazione. Il sistema (3) può essere riscritto nel modo seguente:

$$x_1 A^1 + \dots + x_n A^n = b.$$

A questo punto è chiaro che esistono soluzioni se e solo se b è combinazione lineare delle colonne $A^1, \dots, A^n$, con coefficienti $x_1, \dots, x_n$. $\square$

Passiamo al teorema più importante di questa sezione.

Teorema 3.2.13 (Rouché - Capelli). *Il sistema (3) ha soluzioni se e solo se*

$$\text{rk}(A | b) = \text{rk}(A).$$

In caso affermativo, lo spazio delle soluzioni $S \subset \mathbb{K}^n$ è un sottospazio affine di dimensione $n - \text{rk}(A)$.

Dimostrazione. Sappiamo che il sistema ha soluzioni se e solo se

$$b \in \text{Span}(A^1, \dots, A^n).$$

Se questo accade, allora

$$\text{Span}(A^1, \dots, A^n, b) = \text{Span}(A^1, \dots, A^n)$$

e quindi $\text{rk}(A | b) = \text{rk}(A)$. Se invece questo non accade, allora

$$\text{Span}(A^1, \dots, A^n, b) \supsetneq \text{Span}(A^1, \dots, A^n)$$

e quindi $\text{rk}(A | b) > \text{rk}(A)$.

Se ci sono soluzioni, sappiamo che la dimensione di S è pari alla dimensione dello spazio S_0 delle soluzioni del sistema lineare omogeneo associato. Dopo aver effettuato l'algoritmo di Gauss - Jordan, abbiamo già dimostrato nella Sezione 3.1.4 che S_0 è generato da un numero di vettori indipendenti pari a n meno il numero di pivot, cioè $n - \text{rk}(A)$. Quindi $\dim S = \dim S_0 = n - \text{rk}(A)$. $\square$

Nel corollario seguente supponiamo che $\mathbb{K}$ sia un campo infinito, come ad esempio $\mathbb{Q}$, $\mathbb{R}$ e $\mathbb{C}$.

Corollario 3.2.14. *Il sistema (4) ha 0, 1 oppure ∞ soluzioni. Più precisamente, le soluzioni sono*

- 0 se $\text{rk}(A | b) > \text{rk} A$,
- 1 se $\text{rk}(A | b) = \text{rk} A = n$,
- ∞ se $\text{rk}(A | b) = \text{rk} A < n$.

Esempio 3.2.15. Consideriamo un sistema lineare dipendente da un parametro $k \in \mathbb{R}$

$$\begin{cases} x + ky = 4 - k \\ kx + 4y = 4 \end{cases}$$

Vogliamo sapere al variare di $k \in \mathbb{R}$ se ci siano soluzioni e, in caso affermativo, che dimensione abbiano. Applichiamo l'algoritmo di Gauss su $(A|b)$:

$$\left(\begin{array}{cc|c} 1 & k & 4-k \\ k & 4 & 4 \end{array} \right) \longrightarrow \left(\begin{array}{cc|c} 1 & k & 4-k \\ 0 & 4-k^2 & 4-4k+k^2 \end{array} \right)$$

Per calcolare il rango è sufficiente fermarsi qui, non è necessario proseguire con l'algoritmo di Gauss - Jordan che è utile solo se vogliamo determinare precisamente le soluzioni. Notiamo che $4-4k+k^2 = (k-2)^2$. La matrice è ridotta a scalini per tutti i possibili valori di $k \in \mathbb{R}$ e ci sono sempre due pivot, tranne per $k = 2$ in cui ce n'è uno solo. Quindi il rango di $(A|b)$ è 1 per $k = 2$ e 2 per $k \neq 2$.

La matrice dei coefficienti A con le mosse di Gauss è diventata

$$\begin{pmatrix} 1 & k \\ 0 & 4-k^2 \end{pmatrix}$$

ed ha rango 1 per $k = \pm 2$ e rango 2 per $k \neq \pm 2$. Quindi:

- Se $k = -2$, allora $\text{rk}(A|b) \neq \text{rk}(A)$ e non ci sono soluzioni.
- Se $k = 2$, allora $\text{rk}(A|b) = \text{rk}(A) = 1$, quindi le soluzioni formano un sottospazio affine di $\mathbb{R}^2$ di dimensione $2 - 1 = 1$, cioè una retta affine. Ci sono infinite soluzioni.
- Se $k \neq \pm 2$, allora $\text{rk}(A|b) = \text{rk}(A) = 2$, quindi le soluzioni formano un sottospazio affine di $\mathbb{R}^2$ di dimensione $2 - 2 = 0$, cioè un punto. C'è una soluzione sola.

3.2.5. Sistema lineare omogeneo. Consideriamo un sistema lineare omogeneo

$$\begin{cases} a_{11}x_1 + \cdots + a_{1n}x_n = 0 \\ \vdots \\ a_{k1}x_1 + \cdots + a_{kn}x_n = 0 \end{cases}$$

con matrice dei coefficienti A . Il Teorema di Rouché - Capelli per i sistemi omogenei assume una forma più semplice:

Corollario 3.2.16. *Un sistema lineare omogeneo ha sempre soluzione. Il sottospazio vettoriale $S \subset \mathbb{K}^n$ delle soluzioni ha dimensione $n - \text{rk}(A)$.*

Notiamo le due differenze principali fra i sistemi lineari omogenei e non omogenei: nei primi le soluzioni ci sono sempre e formano un sottospazio vettoriale, mentre nei secondi le soluzioni possono non esserci e se ci sono formano un sottospazio affine.

Sia $W \subset \mathbb{K}^m$ lo spazio vettoriale generato dalle colonne di A . Per definizione $\text{rk}(A) = \dim W$ e quindi per il Teorema di Rouché - Capelli

$$(4) \quad \dim S + \dim W = n.$$

Notiamo che $S \subset \mathbb{K}^n$ e $W \subset \mathbb{K}^m$ sono sottospazi vettoriali di due spazi diversi $\mathbb{K}^n$ e $\mathbb{K}^m$. È utile capire cosa succeda ad entrambi S e W se cambiamo la matrice A con mosse di Gauss sulle righe e sulle colonne:

Proposizione 3.2.17. *Valgono i fatti seguenti:*

- Con mosse di Gauss sulle righe, S è immutato ma W può variare.
- Con mosse di Gauss sulle colonne, W è immutato ma S può variare.

Dimostrazione. Il caso delle righe è già stato analizzato. Per le colonne, per definizione $W = \text{Span}(A^1, \dots, A^n)$. Se scambiamo due colonne A^i, A^j o cambiamo A^i con λA^i , $\lambda \neq 0$, chiaramente W non muta. È facile verificare che non muta neppure se cambiamo A^i con $A^i + \lambda A^j$. $\square$

Corollario 3.2.18. *Le dimensioni di S e W non mutano né per mosse di Gauss sulle righe né sulle colonne.*

Dimostrazione. Segue da (4). $\square$

Abbiamo scoperto che combinando mosse di Gauss su righe e colonne i sottospazi S e W possono mutare, ma le loro dimensioni rimangono costanti.

3.2.6. Rango per righe e per colonne. Abbiamo appena dimostrato il fatto seguente. Sia A una matrice $m \times n$.

Corollario 3.2.19. *Il rango $\text{rk}A$ di una matrice A non cambia per mosse di Gauss sulle righe o sulle colonne.*

Definiamo il *rango per righe* di A come la dimensione dello spazio generato dalle righe $\text{Span}(A_1, \dots, A_m) \subset \mathbb{K}^n$. Il *rango per colonne* è l'usuale rango. In altre parole, il rango per righe di A è il rango della trasposta tA .

Proposizione 3.2.20. *Per ogni matrice A il rango per righe è uguale al rango per colonne.*

Dimostrazione. Se riduciamo A a scalini, entrambi i ranghi non cambiano per il Corollario 3.2.19. In una matrice a scalini il rango per colonne è il numero r di pivot. Quello per righe è al massimo r perché nella matrice a scalini ci sono solo r righe non nulle. Quindi il rango per righe è sempre $\leq$ del rango per colonne. Lavorando con tA troviamo analogamente che il rango per colonne è sempre $\leq$ di quello per righe. Quindi sono uguali. $\square$

Corollario 3.2.21. *Per ogni matrice A vale $\text{rk}({}^tA) = \text{rk}A$.*

Il rango di una matrice è uguale al rango della sua trasposta.

3.2.7. Sottomatrici. Sia A una matrice $m \times n$. Una *sottomatrice* è una matrice B di taglia $k \times h$ ottenuta da A cancellando $m - k$ righe e $n - h$ colonne. Ad esempio la matrice

$$\begin{pmatrix} 1 & 2 & 3 & 4 \\ 5 & 6 & 7 & 8 \\ 9 & 10 & 11 & 12 \end{pmatrix}$$

contiene $3 \times 4 = 12$ sottomatrici di taglia 2×3 e $3 \times 6 = 18$ sottomatrici di taglia 2×2 . Tra queste ultime ci sono

$$\begin{pmatrix} 1 & 2 \\ 5 & 6 \end{pmatrix}, \quad \begin{pmatrix} 2 & 4 \\ 10 & 12 \end{pmatrix}, \quad \begin{pmatrix} 1 & 4 \\ 5 & 8 \end{pmatrix}, \quad \dots$$

Mostriamo che il rango di una sottomatrice non può essere più grande del rango della matrice iniziale.

Proposizione 3.2.22. *Se B è una sottomatrice di A , allora $\text{rk} B \leq \text{rk} A$.*

Dimostrazione. Se prendiamo r vettori colonna $v_1, \dots, v_r \in \mathbb{K}^m$ che sono linearmente dipendenti e cancelliamo $m - k$ righe di questi, otteniamo r vettori colonna $w_1, \dots, w_r \in \mathbb{K}^k$ che sono ancora linearmente dipendenti! La relazione di dipendenza che valeva prima

$$\lambda_1 v_1 + \dots + \lambda_r v_r = 0$$

continua infatti a valere ancora adesso:

$$\lambda_1 w_1 + \dots + \lambda_r w_r = 0.$$

Da questo deduciamo che se B ha r colonne indipendenti, allora le corrispondenti colonne in A sono anch'esse indipendenti (per quanto appena visto, se fossero dipendenti, lo sarebbero anche quelle di B di partenza). Poiché B ha certamente $\text{rk}(B)$ colonne indipendenti, ne deduciamo che A ha $\text{rk} B$ colonne indipendenti e quindi $\text{rk} A \geq \text{rk} B$. $\square$

3.3. Determinante

Introduciamo adesso una nozione importante che si applica solo alle matrici quadrate, quella di *determinante*.

3.3.1. Definizione. Sia A una matrice quadrata $n \times n$. Il *determinante* di A è il numero

$$\det A = \sum_{\sigma \in S_n} \text{sgn}(\sigma) a_{1\sigma(1)} \cdots a_{n\sigma(n)}.$$

Ricordiamo che S_n indica l'insieme delle $n!$ permutazioni di $\{1, \dots, n\}$, si veda la Sezione 1.2.5. Quindi questa è una sommatoria su $n!$ elementi. Il termine $\text{sgn}(\sigma) = \pm 1$ indica il segno della permutazione σ ed è 1 oppure -1 a seconda di σ . In seguito indichiamo la permutazione σ con il simbolo $[\sigma(1) \cdots \sigma(n)]$.

La definizione di $\det A$ non è facile da decifrare: cerchiamo di comprenderla meglio esaminando i casi $n = 1, 2, 3$. Se $n = 1$ la matrice A è un numero

$$A = (a_{11}),$$

l'insieme S_1 contiene solo la permutazione $\text{id} = [1]$, che ha segno positivo, e otteniamo semplicemente

$$\det A = a_{11}.$$

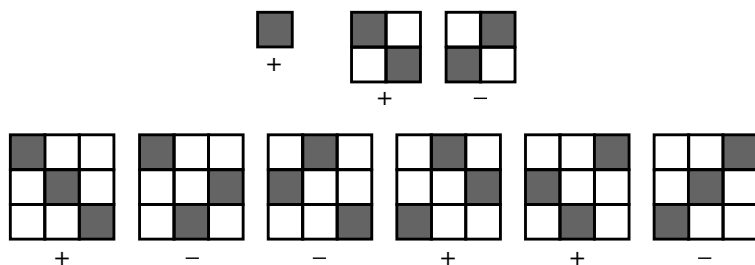


Figura 3.2. Rappresentazione grafica del determinante di una matrice $n \times n$ per $n = 1, 2$ e 3 .

Se $n = 2$ la matrice A è

$$A = \begin{pmatrix} a_{11} & a_{12} \\ a_{21} & a_{22} \end{pmatrix},$$

l'insieme S_2 contiene due permutazioni $\text{id} = [1 \ 2]$ e $[2 \ 1]$. Queste sono l'identità (che ha segno 1) e una trasposizione (che ha segno -1). Quindi

$$\det A = a_{11}a_{22} - a_{12}a_{21}.$$

Se $n = 3$ la matrice è

$$A = \begin{pmatrix} a_{11} & a_{12} & a_{13} \\ a_{21} & a_{22} & a_{23} \\ a_{31} & a_{32} & a_{33} \end{pmatrix}.$$

L'insieme S_3 contiene sei permutazioni:

$$\text{id} = [1 \ 2 \ 3], \quad [1 \ 3 \ 2], \quad [2 \ 1 \ 3], \quad [2 \ 3 \ 1], \quad [3 \ 1 \ 2], \quad [3 \ 2 \ 1].$$

Le tre trasposizioni hanno segno -1 e le altre tre permutazioni $+1$. Quindi

$$\det A = a_{11}a_{22}a_{33} - a_{11}a_{23}a_{32} - a_{12}a_{21}a_{33} + a_{12}a_{23}a_{31} + a_{13}a_{21}a_{32} - a_{13}a_{22}a_{31}.$$

La Figura 3.2 mostra una rappresentazione grafica del determinante nei casi descritti $n = 1, 2, 3$. Ciascun quadrato grande corrisponde ad una permutazione ed il suo segno è indicato sotto. Ad esempio, il determinante delle matrici seguenti

$$(3), \quad \begin{pmatrix} 1 & 2 \\ -1 & 4 \end{pmatrix} \quad \begin{pmatrix} 1 & 2 & 1 \\ 2 & 1 & 2 \\ -1 & 0 & 1 \end{pmatrix}$$

è rispettivamente:

$$3, \quad 4 + 2 = 6, \quad 1 - 0 - 4 + (-4) + 0 - (-1) = -6.$$

Osservazione 3.3.1. Ispirati dalla Figura 3.2, riformuliamo la definizione di determinante con una notazione più visiva. Consideriamo una matrice quadrata A come $n \times n$ caselle con dei numeri al loro interno. Chiamiamo

una *colorazione* di A la colorazione di n sue caselle, in modo che ogni riga e ogni colonna contengano esattamente una sola casella colorata.¹

C'è una chiara corrispondenza biunivoca fra permutazioni in S_n e colorazioni. Data una $\sigma \in S_n$, otteniamo una colorazione c dipingendo le caselle $1\sigma(1), \dots, n\sigma(n)$. Il *segno* della colorazione c è per definizione il segno della corrispondente permutazione σ . Il *peso* $p(c)$ della colorazione c è il prodotto dei numeri nelle caselle colorate, moltiplicato per il segno. Con questo linguaggio, il determinante di A è semplicemente la somma dei pesi di tutte le colorazioni: se indichiamo con $\text{Col}(A)$ l'insieme delle colorazioni, scriviamo

$$\det A = \sum_{c \in \text{Col}(A)} p(c).$$

Vediamo subito che il determinante non cambia per trasposizione.

Proposizione 3.3.2. *Vale $\det({}^t A) = \det A$.*

Dimostrazione. Seguiamo l'interpretazione visiva descritta nell'Osservazione 3.3.1. Data una colorazione $c \in \text{Col}(A)$, possiamo farne la trasposta e ottenere una colorazione ${}^t c \in \text{Col}({}^t A)$. Le due permutazioni corrispondenti sono una l'inversa dell'altra e quindi hanno lo stesso segno per la Proposizione 1.2.14. I due pesi sono gli stessi, cioè $p(c) = p({}^t c)$. Sommando sulle colorazioni si ottiene la tesi. $\square$

3.3.2. Matrici triangolari. Il calcolo del determinante di una matrice triangolare è particolarmente semplice: è il prodotto dei valori sulla diagonale principale.

Proposizione 3.3.3. *Sia $A \in M(n)$ una matrice triangolare superiore*

$$A = \begin{pmatrix} a_{11} & a_{12} & \dots & a_{1n} \\ 0 & a_{22} & \dots & a_{2n} \\ \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \dots & a_{nn} \end{pmatrix}.$$

Vale

$$\det A = a_{11} a_{22} \cdots a_{nn}.$$

Dimostrazione. Usiamo l'interpretazione visiva dell'Osservazione 3.3.1. La matrice è triangolare, quindi l'unica colorazione c che può avere peso non nullo è quella in cui le caselle della diagonale principale sono tutte colorate, che corrisponde alla permutazione identità ed ha quindi segno 1, quindi $p(c) = a_{11} \cdots a_{nn}$. In tutte le altre colorazioni c c'è almeno una casella sotto la diagonale colorata; su questa casella $a_{ij} = 0$ e quindi $p(c) = 0$. Allora $\det A = a_{11} \cdots a_{nn}$. $\square$

Lo stesso risultato vale per le matrici triangolari inferiori.

¹In altre parole, una colorazione è la configurazione di n torri su una scacchiera $n \times n$ che non possono mangiarsi l'una con l'altra.

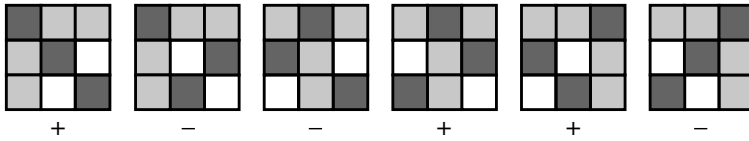


Figura 3.3. Dimostrazione dello sviluppo di Laplace.

3.3.3. Matrice identità. Introduciamo una matrice che avrà un ruolo fondamentale in tutto il libro.

Definizione 3.3.4. La *matrice identità* di taglia $n \times n$ è la matrice

$$I_n = \begin{pmatrix} 1 & 0 & \cdots & 0 \\ 0 & 1 & \cdots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \cdots & 1 \end{pmatrix}$$

i cui coefficienti sono 1 sulla diagonale principale e 0 altrove.

Per adesso ci limitiamo a notare che la matrice identità I_n è diagonale e ha determinante $\det(I_n) = 1$.

3.3.4. Sviluppo di Laplace. Ci sono vari algoritmi che permettono di calcolare il determinante di una matrice più agevolmente che con la cruda definizione. Uno di questi è noto come lo *sviluppo di Laplace* e funziona nel modo seguente.

Sia A una matrice $n \times n$ con $n \geq 2$. Indichiamo con C_{ij} la sottomatrice $(n-1) \times (n-1)$ ottenuta da A rimuovendo la i -esima riga e la j -esima colonna.

Teorema 3.3.5 (Sviluppo di Laplace). *Per ogni i fissato vale*

$$\det A = \sum_{j=1}^n (-1)^{i+j} a_{ij} \det C_{ij}.$$

Dimostrazione. Il caso $n = 3$ e $i = 1$ è nella Figura 3.3 ed è

$$\det A = a_{11} \det C_{11} - a_{12} \det C_{12} + a_{13} \det C_{13}.$$

L'idea chiave della dimostrazione è tutta nella figura, che invitiamo ad osservare attentamente. Per ogni $j = 1, \dots, n$, le colorazioni della sottomatrice C_{ij} , che sommate danno $\det C_{ij}$, contribuiscono alle colorazioni della matrice grande A con un peso che va moltiplicato per a_{ij} e per un segno $(-1)^{i+j}$. $\square$

Vediamo un esempio. Per calcolare il determinante seguente, sviluppiamo lungo la prima riga (cioè prendiamo $i = 1$) e otteniamo

+	-	+
-	+	-
+	-	+

Figura 3.4. Il segno $(-1)^{i+j}$ nello sviluppo di Laplace.

$$\begin{aligned} \det \begin{pmatrix} 1 & -1 & 0 \\ 2 & -1 & 5 \\ 1 & 1 & -1 \end{pmatrix} &= 1 \cdot \det \begin{pmatrix} -1 & 5 \\ 1 & -1 \end{pmatrix} - (-1) \cdot \det \begin{pmatrix} 2 & 5 \\ 1 & -1 \end{pmatrix} \\ &\quad + 0 \cdot \det \begin{pmatrix} 2 & -1 \\ 1 & 1 \end{pmatrix} = -4 - 7 + 0 = -11. \end{aligned}$$

È conveniente sviluppare lungo una riga che contiene degli zeri.

Notiamo inoltre che grazie al fatto che $\det({}^tA) = \det A$, è anche possibile usare lo sviluppo di Laplace sulle colonne. Ad esempio, sviluppando la matrice seguente sulla seconda colonna otteniamo:

$$\det \begin{pmatrix} 1 & 0 & 1 \\ 2 & 0 & 1 \\ \pi & 3 & \sqrt{7} \end{pmatrix} = (-1) \cdot 3 \cdot \det \begin{pmatrix} 1 & 1 \\ 2 & 1 \end{pmatrix} = 3.$$

Nello sviluppo si deve sempre fare attenzione al segno $(-1)^{i+j}$ associato alla casella ij , che varia come su di una scacchiera, si veda la Figura 3.4.

Osservazione 3.3.6. Usando lo sviluppo di Laplace sulla prima colonna, è possibile ridimostrare la formula del determinante per le matrici triangolari per induzione sulla taglia n della matrice.

3.3.5. Mosse di Gauss. Veniamo adesso al motivo principale per cui abbiamo introdotto il determinante di una matrice: la sua compatibilità con le mosse di Gauss.

Proposizione 3.3.7. Sia A una matrice $n \times n$ e sia A' ottenuta da A tramite una mossa di Gauss. Il determinante cambia nel modo seguente:

- (1) Se A' è ottenuta da A scambiando due righe, $\det(A') = -\det(A)$.
- (2) Se A' è ottenuta da A moltiplicando una riga di A per uno scalare λ , allora $\det(A') = \lambda \det(A)$.
- (3) Se A' è ottenuta da A aggiungendo ad una riga il multiplo di un'altra riga, allora $\det(A') = \det A$.

Dimostrazione. Dimostriamo la proposizione per induzione su $n \geq 2$. Se $n = 2$, tutte queste proprietà discendono facilmente dalla formula $\det \begin{pmatrix} a & b \\ c & d \end{pmatrix} = ad - bc$. Dimostriamo il teorema per $n \geq 3$ supponendo che sia vero per $n - 1$.

Siccome $n \geq 3$, esiste sempre una riga A_i che non è implicata nella mossa di Gauss. Con lo sviluppo di Laplace lungo la riga i -esima troviamo

$$\det A = \sum_{j=1}^n (-1)^{i+j} a_{ij} \det C_{ij},$$

$$\det A' = \sum_{j=1}^n (-1)^{i+j} a_{ij} \det C'_{ij}.$$

La matrice C'_{ij} è di taglia $(n-1) \times (n-1)$ ed è ottenuta da C_{ij} tramite lo stesso tipo di mossa di Gauss che trasforma A in A' . Quindi per l'ipotesi induttiva

$$\det(C'_{ij}) = -\det(C_{ij}), \quad \lambda \det(C_{ij}) \quad \text{oppure} \quad \det(C_{ij})$$

a seconda del tipo di mossa di Gauss. Raccogliendo troviamo la stessa relazione tra $\det(A')$ e $\det A$. $\square$

Poiché il determinante non cambia per trasposizione, la stessa conclusione è valida anche per mosse di Gauss sulle colonne. Mostriamo ora una formula un po' più generale. Sia A una matrice $n \times n$ e supponiamo che la i -esima colonna sia esprimibile come combinazione lineare di vettori colonna del tipo

$$A^i = \lambda_1 v_1 + \cdots + \lambda_h v_h$$

per qualche $v_j \in \mathbb{K}^n$ e $\lambda_j \in \mathbb{K}$. Sia B_j la matrice ottenuta da A sostituendo la colonna A^i con v_j , per $j = 1, \dots, h$.

Proposizione 3.3.8. *Vale $\det A = \lambda_1 \det(B_1) + \cdots + \lambda_h \det(B_h)$.*

Dimostrazione. Stessa dimostrazione della Proposizione 3.3.7. $\square$

La stessa formula vale anche con le righe al posto delle colonne.

Esempio 3.3.9. Poiché

$$\begin{pmatrix} 2 \\ 3 \\ -1 \end{pmatrix} = 2 \begin{pmatrix} 1 \\ 1 \\ 0 \end{pmatrix} + \begin{pmatrix} 0 \\ 1 \\ 0 \end{pmatrix} - \begin{pmatrix} 0 \\ 0 \\ 1 \end{pmatrix}$$

applicando la proposizione alla seconda colonna della matrice, troviamo

$$\det \begin{pmatrix} 1 & 2 & -1 \\ 8 & 3 & 1 \\ 0 & -1 & -9 \end{pmatrix} = 2 \det \begin{pmatrix} 1 & 1 & -1 \\ 8 & 1 & 1 \\ 0 & 0 & -9 \end{pmatrix} + \det \begin{pmatrix} 1 & 0 & -1 \\ 8 & 1 & 1 \\ 0 & 0 & -9 \end{pmatrix} - \det \begin{pmatrix} 1 & 0 & -1 \\ 8 & 0 & 1 \\ 0 & 1 & -9 \end{pmatrix}$$

Esempio 3.3.10. Nel calcolo del determinante di una matrice, è spesso utile combinare mosse di Gauss e sviluppo di Laplace. Ad esempio:

$$\begin{aligned} \det \begin{pmatrix} 5 & 1 & 1 & 2 & 1 \\ 1 & 4 & 1 & 1 & 1 \\ 1 & 3 & 3 & 1 & 1 \\ 1 & 1 & 1 & 6 & 1 \\ 1 & 1 & 1 & 1 & 1 \end{pmatrix} &= \det \begin{pmatrix} 4 & 0 & 0 & 1 & 0 \\ 0 & 3 & 0 & 0 & 0 \\ 0 & 2 & 2 & 0 & 0 \\ 0 & 0 & 0 & 5 & 0 \\ 1 & 1 & 1 & 1 & 1 \end{pmatrix} = \det \begin{pmatrix} 4 & 0 & 0 & 1 \\ 0 & 3 & 0 & 0 \\ 0 & 2 & 2 & 0 \\ 0 & 0 & 0 & 5 \end{pmatrix} \\ &= 5 \det \begin{pmatrix} 4 & 0 & 0 \\ 0 & 3 & 0 \\ 0 & 2 & 2 \end{pmatrix} = 5 \cdot 4 \cdot 3 \cdot 2 = 120. \end{aligned}$$

Nella prima uguaglianza abbiamo tolto l'ultima riga da tutte le altre, quindi abbiamo sviluppato lungo alcune righe o colonne e infine abbiamo usato la formula per le matrici triangolari.

3.3.6. Definizione assiomatica. Esiste una definizione alternativa del determinante che non necessita delle permutazioni ed è quindi spesso ritenuta tecnicamente meno laboriosa. Questa definizione segue un tipo di ragionamento tipicamente matematico: il determinante non è descritto esplicitamente ma dedotto in modo astratto da alcuni assiomi. L'approccio è il seguente.

Teorema 3.3.11. *Esiste un'unica funzione $\det: M(n, \mathbb{K}) \rightarrow \mathbb{K}$ che assegni ad una qualsiasi matrice quadrata $A \in M(n, \mathbb{K})$ un numero $\det A \in \mathbb{K}$ tale che:*

- (1) *Se A' è ottenuta da A tramite una mossa di Gauss, allora $\det A'$ è ottenuto da $\det A$ come descritto nella Proposizione 3.3.7,*
- (2) *$\det(I_n) = 1$.*

Dimostrazione. Mostriamo prima l'unicità. Dobbiamo verificare che $\det A$ è determinato dagli assiomi.

Sia A' matrice ottenuta da A tramite l'algoritmo di Gauss-Jordan. Per l'assioma (1), il valore di $\det A$ dipende univocamente da $\det A'$. Se $\text{rk} A' < n$, allora A' ha una riga nulla e si deduce facilmente che $\det A' = 0$. Infatti moltiplicando questa riga per 0 si ottiene sempre A' ed il determinante si annulla per l'assioma (1), quindi $\det A' = 0$. Se $\text{rk} A' = n$, allora $A' = I_n$ e per l'assioma (2) abbiamo $\det A' = 1$.

L'unicità è dimostrata. Per l'esistenza, basta usare lo sviluppo di Laplace come *definizione* di determinante e verificare per induzione su n che lo sviluppo soddisfa i due assiomi. Lasciamo questa parte per esercizio. $\square$

3.3.7. Matrici di rango massimo. A cosa serve il determinante? Innanzitutto a capire se una matrice quadrata ha rango massimo.

Proposizione 3.3.12. *Per ogni $A \in M(n, \mathbb{K})$ abbiamo*

$$\det A \neq 0 \iff \text{rk} A = n.$$

Dimostrazione. Usiamo l'algoritmo di Gauss con A . Ad ogni mossa di Gauss, il rango si mantiene e il determinante cambia per moltiplicazione per una costante non nulla. Quindi possiamo supporre che A sia ridotta a scalini. La matrice A è anche triangolare superiore e quindi

$$\operatorname{rk} A = n \iff a_{11} \cdots a_{nn} \neq 0 \iff \det A \neq 0$$

per la Proposizione 3.3.3. $\square$

3.3.8. Calcolo del rango. Più in generale, il determinante può essere usato per calcolare il rango di una matrice qualsiasi. Ricordiamo che il determinante è definito solo per le matrici quadrate, mentre il rango è definito per tutte.

Sia A una matrice $m \times n$. Un *minore* di A di ordine k è una sottomatrice quadrata $k \times k$ ottenuta cancellando $m - k$ righe e $n - k$ colonne.

Proposizione 3.3.13. *Il rango di A è il massimo ordine di un suo minore B con $\det B \neq 0$.*

Dimostrazione. Scriviamo $r = \operatorname{rk} A$. Sia B un minore con $\det B \neq 0$. Il minore B ha ordine $\operatorname{rk} B$ per la Proposizione 3.3.12. Sappiamo dalla Proposizione 3.2.22 che $\operatorname{rk} B \leq r$, quindi B ha ordine $\leq r$.

D'altro canto esiste un minore B di ordine r con $\det B \neq 0$ ottenuto nel modo seguente. Sappiamo che esistono r colonne indipendenti di A , e sia $A' \subset A$ la sottomatrice $m \times r$ formata da queste. Poiché $\operatorname{rk} A' = r$, ed il rango per righe è uguale al rango per colonne, la sottomatrice A' ha anche r righe indipendenti. Queste formano il minore B cercato: ha ordine e rango r e quindi $\det B \neq 0$. $\square$

Esempio 3.3.14. Usiamo il determinante per capire per quali $k \in \mathbb{R}$ il seguente sistema abbia soluzioni reali:

$$\begin{cases} y + kz &= -k, \\ 2x + (k-3)y + 4z &= k-1, \\ x + ky - kz &= k+1. \end{cases}$$

La matrice dei coefficienti è

$$A = \begin{pmatrix} 0 & 1 & k \\ 2 & k-3 & 4 \\ 1 & k & -k \end{pmatrix}$$

e ha determinante $(k+4)(k+1)$. Quindi per $k \neq -1, -4$ sia A che $C = (A|b)$ hanno rango massimo e il sistema ha una sola soluzione.

Per $k = -1$ la matrice C diventa

$$C = \left(\begin{array}{ccc|c} 0 & 1 & -1 & 1 \\ 2 & -4 & 4 & -2 \\ 1 & -1 & 1 & 0 \end{array} \right).$$

Troviamo $\operatorname{rk} A = \operatorname{rk} C = 2$ perché è facile trovare una relazione tra le righe di C , oppure perché tutti e 4 i minori 3×3 di C hanno determinante

nullo, e il minore $\begin{pmatrix} 0 & 1 \\ 2 & -4 \end{pmatrix}$ ha determinante diverso da zero. Quindi il sistema ha soluzione, e le soluzioni formano un sottospazio affine di dimensione $3 - 2 = 1$.

Per $k = -4$ la matrice C diventa

$$C = \left(\begin{array}{ccc|c} 0 & 1 & -4 & 4 \\ 2 & -7 & 4 & -5 \\ 1 & -4 & 4 & -3 \end{array} \right).$$

Troviamo $\text{rk}A = 2$ perché A contiene il minore $\begin{pmatrix} 0 & 1 \\ 2 & -7 \end{pmatrix}$ che ha $\det \neq 0$. D'altra parte $\text{rk}C = 3$ perché C contiene un minore 3×3 con determinante non nullo:

$$\begin{aligned} \det \begin{pmatrix} 0 & -4 & 4 \\ 2 & 4 & -5 \\ 1 & 4 & -3 \end{pmatrix} &= 4 \det \begin{pmatrix} 0 & -1 & 4 \\ 2 & 1 & -5 \\ 1 & 1 & -3 \end{pmatrix} = 4 \det \begin{pmatrix} 0 & -1 & 4 \\ 0 & -1 & 1 \\ 1 & 1 & -3 \end{pmatrix} \\ &= 4 \det \begin{pmatrix} -1 & 4 \\ -1 & 1 \end{pmatrix} = 4 \cdot 3 = 12. \end{aligned}$$

Durante il calcolo abbiamo usato le mosse di Gauss per semplificare la matrice. Quindi per $k = -4$ il sistema non ha soluzioni. Riassumendo:

- per $k = -4$ non ci sono soluzioni;
- per $k = -1$ ci sono ∞ soluzioni, che formano una retta affine in $\mathbb{R}^3$;
- per tutti gli altri valori esiste un'unica soluzione.

3.3.9. Base di $\mathbb{K}^n$. Una applicazione concreta del determinante è la seguente: per capire se n vettori $v_1, \dots, v_n \in \mathbb{K}^n$ sono una base, si affiancano e si costruisce una matrice quadrata $A = (v_1 \cdots v_n)$. Quindi:

Proposizione 3.3.15. *I vettori $v_1, \dots, v_n$ sono una base $\Leftrightarrow \det A \neq 0$.*

Dimostrazione. Per la Proposizione 2.3.23, i vettori $v_1, \dots, v_n$ sono una base $\Leftrightarrow$ generano $\mathbb{K}^n \Leftrightarrow \text{rk}(A) = n \Leftrightarrow \det(A) \neq 0$. $\square$

Corollario 3.3.16. *I vettori $\begin{pmatrix} a \\ b \end{pmatrix}, \begin{pmatrix} c \\ d \end{pmatrix} \in \mathbb{K}^2$ sono una base $\Leftrightarrow ad - bc \neq 0$.*

Osservazione 3.3.17. Con questo criterio possiamo ridimostrare agevolmente la Proposizione 2.3.24. Se affianchiamo i vettori $v_1, \dots, v_n$ otteniamo una matrice triangolare A con valori $a_{11}, \dots, a_{nn}$ sulla diagonale principale, quindi con $\det A = a_{11} \cdots a_{nn}$. I vettori formano una base $\Leftrightarrow a_{11} \cdots a_{nn} \neq 0$, cioè se e solo se i valori $a_{11}, \dots, a_{nn}$ sono tutti diversi da zero.

3.3.10. Basi positive e negative, area e volume. Siano $v_1, \dots, v_n \in \mathbb{R}^n$ dei vettori, e sia $A = (v_1 \cdots v_n)$. Abbiamo appena visto che i vettori sono indipendenti $\Leftrightarrow \det A \neq 0$. Nel caso in cui siano indipendenti, quale informazione geometrica ci dà il numero reale $\det A$?

Il numero reale $\det A$ porta con sé due informazioni. La prima è il segno: se $\det A > 0$, diciamo che la base $v_1, \dots, v_n$ è *positiva*, e se $\det A < 0$

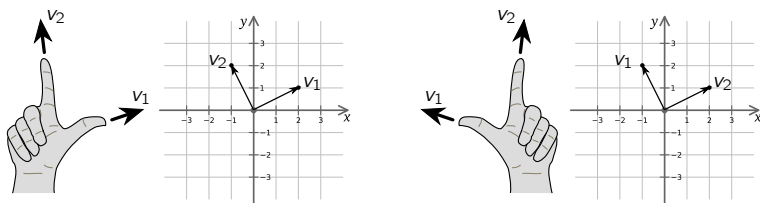


Figura 3.5. Una base positiva ed una negativa nel piano cartesiano $\mathbb{R}^2$.

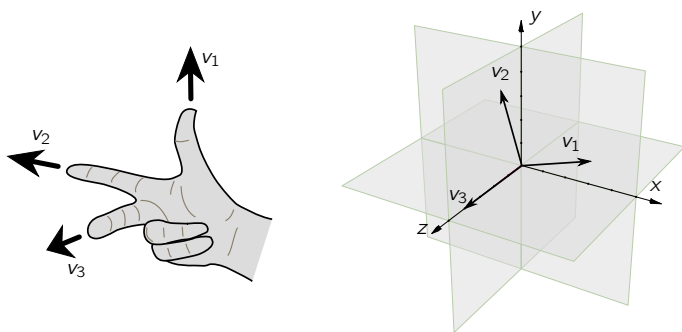


Figura 3.6. Una base positiva nello spazio $\mathbb{R}^3$.

è *negativa*. La base canonica $e_1, \dots, e_n$ ha $\det A = 1$ e quindi è positiva. Se scambiamo due elementi qualsiasi della base canonica, otteniamo una nuova base con $\det A = -1$ che è quindi negativa. In generale se scambiamo due vettori trasformiamo una base positiva in una negativa e viceversa.

Nel piano $\mathbb{R}^2$ una base v_1, v_2 è positiva se è possibile allineare i due vettori rispettivamente con il pollice e l'indice della mano destra, mentre è negativa nel caso in cui si debba usare la mano sinistra: come mostrato in Figura 3.5, la base $\begin{pmatrix} 2 \\ 1 \end{pmatrix}, \begin{pmatrix} 1 \\ 2 \end{pmatrix}$ è positiva; se li scambiamo otteniamo $\begin{pmatrix} 1 \\ 2 \end{pmatrix}, \begin{pmatrix} 2 \\ 1 \end{pmatrix}$ che è invece una base negativa. Questo è coerente con il segno di $\det A$:

$$\det \begin{pmatrix} 2 & 1 \\ 1 & 2 \end{pmatrix} = 3 > 0, \quad \det \begin{pmatrix} 1 & 2 \\ 2 & 1 \end{pmatrix} = -3 < 0.$$

Nello spazio $\mathbb{R}^3$ la situazione è simile. Una base v_1, v_2, v_3 è positiva se è possibile allineare i tre vettori rispettivamente con il pollice, l'indice ed il medio della mano destra, mentre è negativa se dobbiamo usare la mano sinistra. Nella Figura 3.6 vediamo che la base

$$v_1 = \begin{pmatrix} 3 \\ 1 \\ 0 \end{pmatrix}, \quad v_2 = \begin{pmatrix} -1 \\ 3 \\ 0 \end{pmatrix}, \quad v_3 = \begin{pmatrix} 0 \\ 0 \\ 4 \end{pmatrix}$$

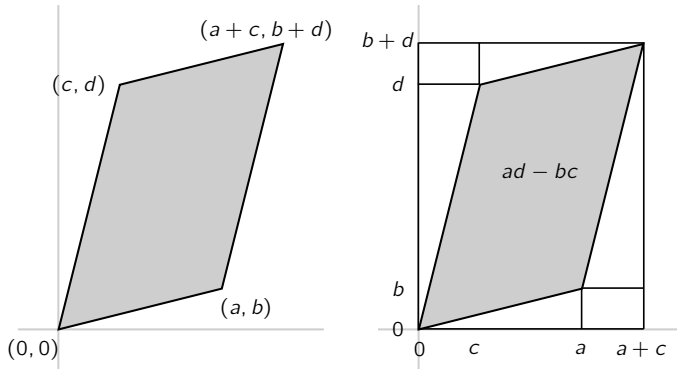


Figura 3.7. L'area del parallelogramma è il valore assoluto del determinante $|ad - bc|$.

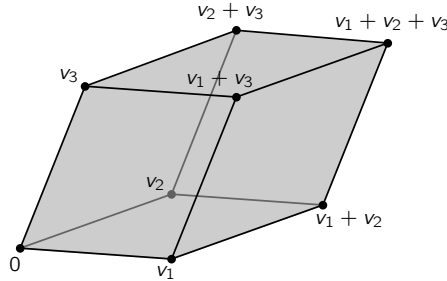


Figura 3.8. Il volume del parallelepipedo è il valore assoluto del determinante $|\det(v_1 v_2 v_3)|$.

è positiva, coerentemente con il fatto che

$$\det \begin{pmatrix} 3 & -1 & 0 \\ 1 & 3 & 0 \\ 0 & 0 & 4 \end{pmatrix} = 40 > 0.$$

La seconda informazione riguarda l'area o il volume del parallelogramma o parallelepipedo generato da $v_1, \dots, v_n$ nei casi $n = 2$ e 3 . Per $n = 2$ il valore assoluto $|\det A| = |ad - bc|$ è l'area del parallelogramma generato da v_1 e v_2 , come mostrato nella Figura 3.7 usando

$$(a+c)(b+d) - 2bc - 2\frac{ab}{2} - 2\frac{cd}{2} = ad - bc.$$

Analogamente per $n = 3$ il valore assoluto del determinante è il volume del parallelepipedo generato da v_1, v_2 e v_3 mostrato nella Figura 3.8. Dimosteremo questa formula successivamente nella Proposizione 9.1.9.

Osservazione 3.3.18. Le aree e i volumi di oggetti in $\mathbb{R}^2$ e $\mathbb{R}^3$ sono definiti con gli strumenti dell'analisi. In questo libro studieremo soltanto

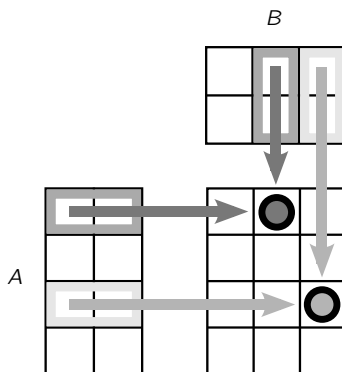


Figura 3.9. Il prodotto riga per colonna AB .

oggetti geometrici semplici come poligoni e poliedri: per studiare questi basterà sapere come calcolare l'area o il volume di oggetti di base, quali i triangoli e tetraedri, perché qualsiasi altro oggetto che vedremo si decompone in oggetti di questo tipo. Per calcolare aree e volumi di altri oggetti più complessi, quali ad esempio cerchi, ellissi e sfere, è necessario usare gli integrali.

3.4. Algebra delle matrici

I sistemi lineari possono essere scritti in una forma più compatta usando una nuova operazione algebrica che introduciamo e studiamo in questa sezione: il *prodotto* fra matrici.

3.4.1. Prodotto fra matrici. Se A è una matrice $m \times n$ e B è una matrice $n \times p$ il *prodotto* AB è una nuova matrice $m \times p$ definita nel modo seguente: l'elemento $(AB)_{ij}$ della nuova matrice AB è

$$(AB)_{ij} = \sum_{k=1}^n A_{ik} B_{kj} = A_{i1} B_{1j} + \cdots + A_{in} B_{nj}.$$

Questo tipo di prodotto fra matrici si chiama *prodotto riga per colonna* perché l'elemento $(AB)_{ij}$ si ottiene facendo un opportuno prodotto fra la riga i -esima A_i di A e la colonna j -esima B^j di B , come suggerito dalla Figura 3.9. Ad esempio:

$$\begin{pmatrix} 1 & 2 \\ -1 & 1 \\ 0 & 3 \end{pmatrix} \cdot \begin{pmatrix} -1 & 2 & 0 & 1 \\ 3 & 0 & 3 & 0 \end{pmatrix} = \begin{pmatrix} 5 & 2 & 6 & 1 \\ 4 & -2 & 3 & -1 \\ 9 & 0 & 9 & 0 \end{pmatrix}.$$

Ricordiamo che il prodotto AB si può fare solo se il numero di colonne di A è pari al numero di righe di B . In particolare, se A è una matrice $m \times n$

e x è una matrice $n \times 1$, cioè un vettore colonna $x \in \mathbb{K}^n$, allora il prodotto Ax è una matrice $m \times 1$, cioè un vettore colonna in $\mathbb{K}^m$. Ad esempio:

$$\begin{pmatrix} 1 & 2 \\ -1 & 1 \\ 0 & 3 \end{pmatrix} \cdot \begin{pmatrix} 1 \\ -1 \end{pmatrix} = \begin{pmatrix} -1 \\ -2 \\ -3 \end{pmatrix}.$$

3.4.2. Notazione compatta per i sistemi lineari. Grazie al prodotto fra matrici è possibile scrivere un sistema lineare in una forma molto più compatta. Possiamo scrivere il sistema

$$\begin{cases} a_{11}x_1 + \cdots + a_{1n}x_n = b_1, \\ \vdots \\ a_{k1}x_1 + \cdots + a_{kn}x_n = b_k, \end{cases}$$

nella forma

$$\begin{pmatrix} a_{11} & \cdots & a_{1n} \\ \vdots & \ddots & \vdots \\ a_{k1} & \cdots & a_{kn} \end{pmatrix} \begin{pmatrix} x_1 \\ \vdots \\ x_n \end{pmatrix} = \begin{pmatrix} b_1 \\ \vdots \\ b_k \end{pmatrix}$$

e quindi più brevemente

$$Ax = b.$$

Come sempre A è la matrice $k \times n$ dei coefficienti e $b \in \mathbb{K}^k$ il vettore dei termini noti.

Esempio 3.4.1. Il sistema lineare

$$\begin{cases} 2x_1 + 3x_2 = 5 \\ x_2 - 4x_3 = 1 \end{cases}$$

può essere scritto usando il prodotto fra matrici nel modo seguente:

$$\begin{pmatrix} 2 & 3 & 0 \\ 0 & 1 & -4 \end{pmatrix} \begin{pmatrix} x_1 \\ x_2 \\ x_3 \end{pmatrix} = \begin{pmatrix} 5 \\ 1 \end{pmatrix}.$$

3.4.3. Proprietà. Studiamo le proprietà algebriche del prodotto fra matrici. Ricordiamo che il prodotto AB ha senso solo se il numero di colonne di A coincide con il numero di righe di B .

Proposizione 3.4.2. *Valgono le proprietà seguenti, per ogni A, B, C matrici per cui i prodotti e le somme abbiano senso e per ogni $\lambda \in \mathbb{K}$.*

- (1) $A(B + C) = AB + AC$ e $(A + B)C = AC + BC$ (distributività)
- (2) $A(BC) = (AB)C$ (associatività)
- (3) $\lambda(AB) = (\lambda A)B = A(\lambda B)$

Dimostrazione. Per la (1), troviamo

$$(A(B+C))_{ij} = \sum_{k=1}^n A_{ik}(B+C)_{kj} = \sum_{k=1}^n A_{ik}B_{kj} + \sum_{k=1}^n A_{ik}C_{kj} = (AB)_{ij} + (BC)_{ij}$$

e l'altra uguaglianza è analoga. Per la (2), verifichiamo che

$$\begin{aligned}(A(BC))_{ij} &= \sum_{k=1}^n A_{ik}(BC)_{kj} = \sum_{k=1}^n \sum_{h=1}^m A_{ik}B_{kh}C_{hj} \\ ((AB)C)_{ij} &= \sum_{h=1}^m (AB)_{ih}C_{hj} = \sum_{h=1}^m \sum_{k=1}^n A_{ik}B_{kh}C_{hj}\end{aligned}$$

e scambiando le sommatorie si ottiene lo stesso totale. Il punto (3) è più facile ed è lasciato per esercizio. $\square$

Esercizio 3.4.3. Vale la relazione

$${}^t(AB) = {}^tB {}^tA.$$

3.4.4. L'anello delle matrici quadrate. Consideriamo lo spazio vettoriale $M(n, \mathbb{K})$ delle matrici $n \times n$ a coefficienti in $\mathbb{K}$, che indichiamo brevemente con $M(n)$. Il prodotto fra matrici è un'operazione binaria in $M(n)$, perché il prodotto di due matrici $A, B \in M(n)$ si può sempre fare ed il risultato è una nuova matrice $AB \in M(n)$. Abbiamo definito gli anelli nella Sezione 1.5.2.

Proposizione 3.4.4. L'insieme $M(n, \mathbb{K})$ delle matrici $n \times n$ è un anello con le operazioni di somma e di prodotto fra matrici. L'elemento neutro per la somma è la matrice nulla, l'elemento neutro per il prodotto è la matrice identità I_n .

Dimostrazione. Sappiamo già che $M(n)$ è un gruppo commutativo con la somma. Per la Proposizione 3.4.2, il prodotto è associativo e distributivo rispetto alla somma. Resta da mostrare che I_n è un elemento neutro per il prodotto, e cioè che

$$I_n A = A I_n = A \quad \forall A \in M(n).$$

Lo verifichiamo adesso. Indichiamo I_n con I per semplicità. Troviamo:

$$(IA)_{ij} = \sum_{k=1}^n I_{ik}A_{kj} = I_{ii}A_{ij} = A_{ij}$$

perché I_{ik} è uguale a zero se $i \neq k$ e ad uno se $i = k$. L'altra uguaglianza è dimostrata in modo analogo. $\square$

La matrice identità I_n è a volte indicata semplicemente con I . Lo spazio vettoriale $M(n)$ delle matrici quadrate è al tempo stesso uno spazio vettoriale e un anello.

3.4.5. Matrici (non) commutanti e (non) invertibili. Notiamo adesso due fatti importanti, che ci mostrano che l'anello $M(n, \mathbb{K})$ è ben lontano dall'essere un campo se $n \geq 2$. Ricordiamo che un campo è un anello in cui il prodotto è commutativo e ogni elemento diverso da zero ha un inverso rispetto al prodotto. Ebbene entrambe queste proprietà sono false in $M(n)$ se $n \geq 2$.

Se $n \geq 2$ l'anello $M(n)$ non è commutativo: è abbastanza raro che due matrici A e B commutino. Ad esempio, se $A = \begin{pmatrix} 1 & 0 \\ 0 & 0 \end{pmatrix}$ e $B = \begin{pmatrix} 0 & 1 \\ 0 & 0 \end{pmatrix}$, allora

$$AB = \begin{pmatrix} 0 & 1 \\ 0 & 0 \end{pmatrix} = B, \quad BA = \begin{pmatrix} 0 & 0 \\ 0 & 0 \end{pmatrix} = 0.$$

Ci sono inoltre molte matrici quadrate che non hanno un'inversa per il prodotto. Una matrice quadrata A per cui esista un'inversa per il prodotto è detta *invertibile*. In questo caso l'inversa è indicata con A^{-1} e per definizione vale $AA^{-1} = A^{-1}A = I$ dove I è la matrice identità, elemento neutro del prodotto.

Ad esempio, le due matrici A e B appena descritte non sono invertibili. Se per assurdo esistesse una inversa per A , che chiameremmo A^{-1} , dall'equazione $BA = 0$ moltiplicando entrambi i membri a destra per A^{-1} otterremmo $B = 0$, che è assurdo. Un ragionamento analogo vale per B .

Vedremo nelle prossime pagine come determinare le matrici invertibili. Intanto mostriamo che il prodotto di due matrici invertibili è sempre invertibile. Questo vale in qualsiasi anello, si veda l'Esercizio 1.18.

Proposizione 3.4.5. *Se $A, B \in M(n)$ sono invertibili, anche AB lo è e*

$$(AB)^{-1} = B^{-1}A^{-1}.$$

Dimostrazione. Verifichiamo che effettivamente

$$(B^{-1}A^{-1})(AB) = B^{-1}(A^{-1}A)B = B^{-1}B = I,$$

$$(AB)(B^{-1}A^{-1}) = A(BB^{-1})A^{-1} = AA^{-1} = I.$$

La dimostrazione è conclusa. □

Osservazione 3.4.6. In assenza di inversa, non si può sempre “semplificare” con il prodotto: con le matrici A e B di sopra abbiamo $AB = B$, ma non è possibile eliminare la B e ottenere $A = I_2$, infatti $A \neq I_2$. Con la somma invece si può sempre fare: se $C + D = D$ allora $C = 0$, perché aggiungiamo $-D$ ad entrambi i membri.

3.4.6. Teorema di Binet. Fra le numerose proprietà del determinante, c'è una forte compatibilità con il prodotto fra matrici: il determinante di AB è semplicemente il prodotto dei determinanti di A e B . Questo fatto molto utile è noto come *Teorema di Binet*.

Teorema 3.4.7 (Teorema di Binet). *Date $A, B \in M(n)$, vale*

$$\det(AB) = \det A \cdot \det B.$$

Dimostrazione. Sia come sempre B_i la riga i -esima di B . Abbiamo

$$\det(AB) = \det \begin{pmatrix} a_{11}B_1 + \dots + a_{1n}B_n \\ \vdots \\ a_{n1}B_1 + \dots + a_{nn}B_n \end{pmatrix}.$$

Se usiamo numerose volte la Proposizione 3.3.8 sulle righe possiamo scrivere questa quantità come somma di n^n addendi del tipo

$$a_{1\sigma(1)} \cdots a_{n\sigma(n)} \det \begin{pmatrix} B_{\sigma(1)} \\ \vdots \\ B_{\sigma(n)} \end{pmatrix}$$

al variare di $\sigma(1), \dots, \sigma(n) \in \{1, \dots, n\}$. Sappiamo che se due righe sono uguali allora il determinante si annulla: quindi ci possiamo limitare a considerare solo i casi in cui $\sigma(1), \dots, \sigma(n)$ sono tutti distinti, in altre parole σ è una permutazione. Quindi otteniamo

$$\det(AB) = \sum_{\sigma \in S_n} a_{1\sigma(1)} \cdots a_{n\sigma(n)} \det \begin{pmatrix} B_{\sigma(1)} \\ \vdots \\ B_{\sigma(n)} \end{pmatrix}.$$

In ogni addendo, la permutazione σ è prodotto di un certo numero k di trasposizioni. Quindi scambiando le righe k volte si trasforma la matrice di destra in B . Ad ogni scambio il determinante cambia di segno, e ricordando che $\text{sgn}(\sigma) = (-1)^k$ otteniamo

$$\det(AB) = \sum_{\sigma \in S_n} \text{sgn}(\sigma) a_{1\sigma(1)} \cdots a_{n\sigma(n)} \det \begin{pmatrix} B_1 \\ \vdots \\ B_n \end{pmatrix} = \det A \det B.$$

La dimostrazione è conclusa. $\square$

Corollario 3.4.8. *Per ogni $A \in M(n, \mathbb{K})$ invertibile vale $\det A \neq 0$ e*

$$\det(A^{-1}) = \frac{1}{\det A}.$$

Dimostrazione. Osserviamo che $I_n = A^{-1}A$ implica che

$$1 = \det(I_n) = \det(A^{-1}A) = \det(A^{-1}) \det A.$$

La dimostrazione è conclusa. $\square$

Osservazione 3.4.9. In generale $\det(A + B) \neq \det A + \det B$. Ad esempio $\det \begin{pmatrix} 0 & 0 \\ 0 & 0 \end{pmatrix} = 0$ e $\det \begin{pmatrix} 1 & 0 \\ 0 & 1 \end{pmatrix} + \det \begin{pmatrix} -1 & 0 \\ 0 & -1 \end{pmatrix} = 1 + 1 = 2$.

3.4.7. Calcolo dell'inversa di una matrice. Introduciamo un algoritmo che permette di calcolare l'inversa di una matrice A , quando esiste. Si basa, come molte altre cose in questo capitolo, sulle mosse di Gauss.

L'algoritmo funziona nel modo seguente. Si prende la matrice $n \times (2n)$

$$C = (A \mid I_n)$$

ottenuta affiancando A e la matrice identità I_n . Trasformiamo quindi C usando l'algoritmo di Gauss - Jordan. Se $\text{rk} A = n$, alla fine dell'algoritmo

tutti i pivot stanno nelle prime n colonne. In particolare abbiamo ottenuto una nuova matrice $n \times (2n)$

$$C' = (I | B)$$

dove B è una matrice $n \times n$.

Proposizione 3.4.10. Vale $B = A^{-1}$.

Dimostrazione. Per definizione di inversa, la colonna i -esima di A^{-1} è l'unico vettore $x_0 \in \mathbb{K}^n$ tale che $Ax_0 = e_i$. Il sistema lineare $Ax = e_i$ è descritto dalla matrice $(A | e_i)$, con unica soluzione x_0 . Se facciamo delle mosse di Gauss, l'insieme $S = \{x_0\}$ delle soluzioni non cambia.

Siccome $\text{rk} A = n$, l'algoritmo di Gauss - Jordan trasforma $(A | e_i)$ in $(I_n | B^i)$ e siccome x_0 è ancora soluzione troviamo che adesso $I_n x_0 = B^i$, cioè $B^i = x_0$. Quindi la i -esima colonna di B è uguale all' i -esima colonna di A^{-1} , per ogni i . Quindi $B = A^{-1}$. $\square$

Come sempre, non è necessario seguire pedissequamente l'algoritmo di Gauss - Jordan: qualsiasi sequenza di mosse di Gauss va bene purché trasformi A in I_n . Ad esempio, determiniamo l'inversa di $A = \begin{pmatrix} 2 & 1 \\ 1 & 1 \end{pmatrix}$ con mosse di Gauss:

$$\left(\begin{array}{cc|cc} 2 & 1 & 1 & 0 \\ 1 & 1 & 0 & 1 \end{array} \right) \rightarrow \left(\begin{array}{cc|cc} 1 & 0 & 1 & -1 \\ 1 & 1 & 0 & 1 \end{array} \right) \rightarrow \left(\begin{array}{cc|cc} 1 & 0 & 1 & -1 \\ 0 & 1 & -1 & 2 \end{array} \right)$$

Quindi $A^{-1} = \begin{pmatrix} 1 & -1 \\ -1 & 2 \end{pmatrix}$. Nel frattempo, abbiamo dimostrato questo fatto:

Proposizione 3.4.11. Una matrice $A \in M(n)$ è invertibile $\Leftrightarrow \det A \neq 0$.

Dimostrazione. Se A è invertibile, allora $\det A \neq 0$ per il Teorema di Binet. D'altro canto, se $\det A \neq 0$ allora $\text{rk} A = n$ e l'algoritmo precedente per determinare A^{-1} funziona. $\square$

Si può calcolare l'inversa di una matrice anche usando il determinante. Data $A \in M(n)$, la *matrice dei cofattori* è la matrice $\text{cof } A \in M(n)$ il cui elemento i, j è

$$\text{cof}(A)_{ij} = (-1)^{i+j} \det C_{ij}$$

dove C_{ij} è il minore di A ottenuto cancellando la riga i -esima e la colonna j -esima. Possiamo scrivere lo sviluppo di Laplace in questo modo:

$$\det A = \sum_{j=1}^n a_{ij} \text{cof}(A)_{ij}.$$

Usiamo adesso la matrice dei cofattori per determinare A^{-1} .

Proposizione 3.4.12. Se A è invertibile, allora

$$A^{-1} = \frac{1}{\det A} {}^t(\text{cof } A).$$

Dimostrazione. Dobbiamo dimostrare che

$$A \cdot {}^t(\text{cof } A) = (\det A)I_n.$$

In altre parole, il prodotto riga per colonna

$$A_i \cdot ({}^t \text{cof } A)^j = \sum_{h=1}^n A_{ih} ({}^t \text{cof } A)_{hj} = \sum_{h=1}^n A_{ih} \text{cof } (A)_{jh}$$

deve fare $\det A$ se $i = j$ e 0 se $i \neq j$. Se $i = j$, la sommatoria è semplicemente lo sviluppo di Laplace di A e quindi fa $\det A$. Se $i \neq j$, la sommatoria è lo sviluppo di Laplace della matrice A' ottenuta da A sostituendo A_i al posto di A_j , e fa $\det A' = 0$ perché A' ha due righe uguali. $\square$

Esempio 3.4.13. La matrice inversa di $A = \begin{pmatrix} a & b \\ c & d \end{pmatrix}$ con $\det A \neq 0$ è

$$A^{-1} = \frac{1}{ad - bc} \begin{pmatrix} d & -b \\ -c & a \end{pmatrix}.$$

3.4.8. Sistema lineare con A invertibile. Notiamo adesso un fatto semplice ma rilevante. Sia

$$Ax = b$$

un sistema lineare con A matrice quadrata (quindi ci sono tante variabili quante equazioni). Se A è invertibile, possiamo moltiplicare entrambi i membri a sinistra per A^{-1} e ottenere

$$x = A^{-1}b.$$

In questo modo scopriamo immediatamente che il sistema ha una sola soluzione e l'abbiamo determinata. Se invece A non è invertibile ovviamente questo metodo non funziona.

Proposizione 3.4.14 (Regola di Cramer). *Se $\det A \neq 0$, il sistema $Ax = b$ ha un'unica soluzione x , determinata nel modo seguente:*

$$x_i = \frac{\det B_i}{\det A} \quad \forall i = 1, \dots, n$$

dove B_i è la matrice ottenuta sostituendo la i -esima colonna di A con b .

Dimostrazione. Sappiamo che $x = A^{-1}b$, quindi

$$x_i = (A^{-1})_i \cdot b = \frac{{}^t((\text{cof } A)^i) \cdot b}{\det A} = \frac{\det B_i}{\det A}.$$

L'ultima uguaglianza è lo sviluppo di Laplace di B_i lungo la i -esima colonna. La dimostrazione è conclusa. $\square$

La regola di Cramer è molto compatta, ma la sua applicazione è spesso più dispendiosa in termini di conti dell'algoritmo di Gauss - Jordan.

3.4.9. Sistema con $n - 1$ equazioni e n variabili. Esaminiamo brevemente un tipo di sistemi particolarmente semplice da trattare. Sia

$$Ax = 0$$

un sistema lineare omogeneo con $x \in \mathbb{K}^n$ e $A \in M(n - 1, n)$. Sia d_i il determinante del minore di A ottenuto cancellando la i -esima colonna. Ricordiamo che $\text{rk} A = n - 1 \iff$ i valori $d_1, \dots, d_n$ non sono tutti nulli.

Proposizione 3.4.15. Se $\text{rk} A = n - 1$, lo spazio S delle soluzioni di $Ax = 0$ è la retta $S = \text{Span}(v)$ con

$$v = \begin{pmatrix} d_1 \\ -d_2 \\ d_3 \\ \vdots \\ (-1)^n d_n \end{pmatrix}.$$

Dimostrazione. Lo spazio S è una retta vettoriale per il Teorema di Rouché - Capelli. Per ogni $i = 1, \dots, n - 1$, sia $B \in M(n, n)$ la matrice ottenuta aggiungendo la riga A_i in fondo ad A . La matrice B ha due righe uguali e quindi $\det B = 0$. Sviluppando $\det B$ lungo l'ultima riga si ottiene la relazione

$$a_{i1}d_1 - a_{i2}d_2 + \dots + (-1)^n a_{in}d_n = 0.$$

Quindi v è soluzione dell' i -esima equazione per ogni i , e quindi del sistema. Poiché $v \neq 0$ troviamo $S = \text{Span}(v)$. $\square$

Esercizi

Esercizio 3.1. Determina le soluzioni dei sistemi lineari:

$$\begin{cases} 3x_1 + x_3 - x_4 = 0 \\ x_2 - 2x_3 - x_4 = -4 \\ 4x_1 + x_2 + 2x_3 - 3x_4 = 0 \\ 3x_1 - x_2 + x_4 = 1 \end{cases} \quad \begin{cases} 2x_2 + 2x_3 + 6x_4 + x_5 = 7 \\ 2x_1 + x_2 + 7x_3 + 11x_4 = 4 \\ x_1 + 3x_3 + 4x_4 + x_5 = 4 \end{cases}$$

Esercizio 3.2. Determina le soluzioni del sistema lineare seguente, al variare dei parametri $\alpha, \beta \in \mathbb{R}$.

$$\begin{cases} x + \alpha y + \beta z = 1 \\ 2x + 3\alpha y + 2\beta z = 2 \\ \beta x + \alpha(2 + \beta)y + (\alpha + \beta + \beta^2)z = 2\beta \end{cases}$$

Esercizio 3.3. Calcola il determinante delle matrici seguenti:

$$\begin{pmatrix} 1 & 2 & 3 & 4 & 5 \\ 2 & 3 & 4 & 5 & 5 \\ 3 & 4 & 5 & 5 & 5 \\ 4 & 5 & 5 & 5 & 5 \\ 5 & 5 & 5 & 5 & 5 \end{pmatrix}, \begin{pmatrix} 1 & 2 & 3 & 4 & 5 & 6 \\ 3 & 1 & 4 & 0 & 0 & 0 \\ 2 & 1 & 5 & 9 & 1 & 0 \\ 11^{12} & 9 & 0 & 0 & 0 & 0 \\ 3 & 1 & 2 & 5 & 0 & 0 \\ 2 & 0 & 0 & 0 & 0 & 0 \end{pmatrix}, \begin{pmatrix} 4 & -1 & 3 & 1 & 2 \\ 1 & 2 & 2 & 1 & 0 \\ 4 & \sqrt{3} & 4 & 1 & 1 \\ -1 & 1 & 0 & 1 & 0 \\ 1 & -1 & 0 & 1 & 2 \end{pmatrix}.$$

Esercizio 3.4. Considera i sistemi lineari dipendenti da un parametro $c \in \mathbb{R}$:

$$\begin{cases} 2x_1 - x_2 + x_3 + 6x_4 = 13 \\ x_1 + x_2 + 5x_3 + 3x_4 = c - 14 \\ x_2 + 3x_3 + x_4 = c \\ x_1 + 2x_3 + 3x_4 = 7 \end{cases} \quad \begin{cases} x + y + cz = 1 \\ x + y + c^3z = 3 \\ 2x + 2y + (1+c)z = 1 \\ x + z = 0 \end{cases}$$

$$\begin{cases} x - y + 2z = 1 \\ 2x + (c-2)y + 2z = 1 \\ -x + y + (c^2 - c - 2)z = c - 1 \end{cases} \quad \begin{cases} x + cy - 2z = c \\ -x + z = 1 - c \\ -2x - 3cy + (c+4)z = -c \end{cases}$$

Per ciascun sistema, determina per quali $c \in \mathbb{R}$ il sistema ha soluzione e la dimensione dello spazio delle soluzioni al variare di $c \in \mathbb{R}$.

Esercizio 3.5. . Calcola l'inversa, se esiste, delle seguenti matrici:

$$\begin{pmatrix} 1 & i \\ -i & -1 \end{pmatrix}, \quad \begin{pmatrix} 3 & -1 & -5 \\ 3 & -1 & 1 \\ 2 & -2 & 1 \end{pmatrix}, \quad \begin{pmatrix} 0 & 0 & 0 & 1 \\ 0 & 0 & 1 & 1 \\ 0 & 1 & 3 & 3 \\ 1 & 3 & 3 & 5 \end{pmatrix}.$$

Esercizio 3.6. Considera il sottospazio $W = \{x_1 + 2x_2 + 3x_3 + 4x_4 + 5x_5 = 0\}$ di $\mathbb{C}^5$. Determina tre sottospazi $U, V, Z \subset W$, tutti di dimensione almeno 1, tali che $W = U \oplus V \oplus Z$.

Esercizio 3.7. Sia U il sottospazio di $\mathbb{R}^4$ generato dai vettori

$$\begin{pmatrix} 1 \\ 2 \\ 3 \\ 4 \end{pmatrix}, \quad \begin{pmatrix} 4 \\ 3 \\ 2 \\ 1 \end{pmatrix}.$$

Sia V il sottospazio di $\mathbb{R}^4$ delle soluzioni del sistema

$$\begin{cases} x_1 + x_2 + x_3 + x_4 = 0 \\ x_2 + 4x_3 - x_4 = 0 \\ 3x_1 + 5x_2 + 11x_3 + x_4 = 0 \end{cases}$$

Calcola le dimensioni di $U, V, U \cap V$ e $U + V$.

Esercizio 3.8. Al variare di $k \in \mathbb{R}$, sia $U_k \subset \mathbb{R}^4$ il sottospazio formato dalle soluzioni del sistema

$$\begin{cases} x + y + kz + kt = 0 \\ 2x + (2-k)y + 3kz = 0 \\ (2-k)x + 2y + 4kt = 0 \end{cases}$$

Sia inoltre

$$W = \text{Span} \left(\begin{pmatrix} 1 \\ 1 \\ 2 \\ 3 \end{pmatrix}, \begin{pmatrix} 1 \\ 0 \\ -1 \\ -1 \end{pmatrix}, \begin{pmatrix} -2 \\ 1 \\ 1 \\ 1 \end{pmatrix}, \begin{pmatrix} -2 \\ 2 \\ 4 \\ 5 \end{pmatrix} \right).$$

- (1) Calcola la dimensione di U_k al variare di $k \in \mathbb{R}$.
- (2) Costruisci una base di $Z = W \cap U_0$.
- (3) Determina un sottospazio $Z' \subset \mathbb{R}^4$ tale che $Z \oplus Z' = \mathbb{R}^4$.

Esercizio 3.9. Sia $A \in M(n)$ una matrice invertibile. Dimostra che anche tA è invertibile e $({}^tA)^{-1} = {}^t(A^{-1})$.

Esercizio 3.10. Data $A \in M(n, \mathbb{K})$ e $\lambda \in \mathbb{K}$, mostra che $\det(\lambda A) = \lambda^n \det A$.

Esercizio 3.11. Dimostra che una matrice $A \in M(17, \mathbb{R})$ antisimmetrica ha determinante nullo.

Esercizio 3.12 (Matrice di Vandermonde). Siano $\alpha_1, \dots, \alpha_n \in \mathbb{K}$. Dimostra l'uguaglianza seguente:

$$\det \begin{pmatrix} 1 & \alpha_1 & \alpha_1^2 & \cdots & \alpha_1^{n-1} \\ 1 & \alpha_2 & \alpha_2^2 & \cdots & \alpha_2^{n-1} \\ \vdots & \vdots & \vdots & \ddots & \vdots \\ 1 & \alpha_n & \alpha_n^2 & \cdots & \alpha_n^{n-1} \end{pmatrix} = \prod_{1 \leq i < j \leq n} (\alpha_j - \alpha_i).$$

Esercizio 3.13. Sia $A \in M(n)$ una matrice a blocchi del tipo

$$A = \begin{pmatrix} B & C \\ 0 & D \end{pmatrix}$$

dove B e D sono sottomatrici quadrate di ordine $k \times k$ e $(n-k) \times (n-k)$. La sottomatrice 0 in basso a sinistra è tutta nulla e di ordine $(n-k) \times k$. Dimostra che

$$\det A = \det B \cdot \det D.$$

Suggerimento. Dimostra l'esercizio per induzione su n ed effettuando lo sviluppo di Laplace sulla prima colonna.

Esercizio 3.14. Sia A una matrice a blocchi come nell'esercizio precedente:

$$A = \begin{pmatrix} B & C \\ 0 & D \end{pmatrix}$$

Dimostra che $\text{rk} A \geq \text{rk} B + \text{rk} D$. Costruisci degli esempi in cui vale l'uguaglianza e degli esempi in cui vale la disuguaglianza stretta.

Suggerimento. Usa l'esercizio precedente e la Proposizione 3.3.13.

Applicazioni lineari

Nel Capitolo 2 abbiamo introdotto la struttura algebrica più importante del libro, quella di *spazio vettoriale*. Abbiamo descritto molti esempi di spazi vettoriali e definito alcuni concetti fondamentali come quelli di combinazione e indipendenza lineare, sottospazio, base e dimensione.

Dopo aver esaminato gli spazi vettoriali da un punto di vista statico, ci accingiamo adesso a studiarli da una prospettiva più dinamica. Introduciamo in questo capitolo le *applicazioni lineari*: queste sono funzioni fra spazi vettoriali che ne preservano la struttura. Il linguaggio delle applicazioni lineari è particolarmente efficace perché può essere utilizzato sia in algebra, con i sistemi lineari, i polinomi e le funzioni, che in geometria, nello studio delle trasformazioni geometriche del piano e dello spazio.

4.1. Introduzione

Un'applicazione lineare è una funzione fra spazi vettoriali compatibile con le operazioni di somma e prodotto per scalare.

4.1.1. Definizione. Siano V e W due spazi vettoriali sullo stesso campo $\mathbb{K}$. Una funzione

$$f: V \longrightarrow W$$

è *lineare* se valgono gli assiomi seguenti:

- (1) $f(0) = 0$.
- (2) $f(v + w) = f(v) + f(w)$ per ogni $v, w \in V$.
- (3) $f(\lambda v) = \lambda f(v)$ per ogni $v \in V$ e $\lambda \in \mathbb{K}$.

Notiamo che in (1) il primo 0 è l'origine di V ed il secondo 0 è l'origine di W . Quindi f deve mandare l'origine di V nell'origine di W .

In matematica i termini *funzione*, *applicazione* e *mappa* sono generalmente sinonimi, anche se il termine *applicazione* è usato soprattutto per indicare le applicazioni lineari.

4.1.2. Combinazioni lineari. Dagli assiomi di funzione lineare possiamo subito dedurre un fatto importante. Se $f: V \rightarrow W$ è lineare ed un vettore $v \in V$ è espresso come combinazione lineare di alcuni vettori di V :

$$v = \lambda_1 v_1 + \cdots + \lambda_k v_k$$

allora, poiché f è lineare, troviamo che la sua immagine è

$$\begin{aligned} f(v) &= f(\lambda_1 v_1 + \cdots + \lambda_k v_k) = f(\lambda_1 v_1) + \cdots + f(\lambda_k v_k) \\ &= \lambda_1 f(v_1) + \cdots + \lambda_k f(v_k). \end{aligned}$$

Notiamo quindi che f manda una qualsiasi combinazione lineare di vettori $v_1, \dots, v_k$ in una combinazione lineare (con gli stessi coefficienti) delle loro immagini $f(v_1), \dots, f(v_k)$.

4.1.3. Esempi basilari. Facciamo alcuni esempi importanti.

Esempio 4.1.1 (Funzione nulla e identità). Dati due spazi vettoriali V, W su $\mathbb{K}$, la *funzione nulla* è la funzione $f: V \rightarrow W$ che è costantemente nulla, cioè tale che $f(v) = 0$ per ogni v . La funzione nulla è lineare: i 3 assiomi di linearità per f sono soddisfatti in modo banale.

Dato uno spazio vettoriale V su $\mathbb{K}$, la *funzione identità* è la funzione $\text{id}_V: V \rightarrow V$ che manda ogni vettore in se stesso, cioè $\text{id}_V(v) = v \forall v \in V$. La funzione identità è lineare: anche qui i 3 assiomi sono ovviamente soddisfatti. A volte sottointendiamo V e indichiamo id_V con id .

Esempio 4.1.2. La funzione $f: \mathbb{R} \rightarrow \mathbb{R}$ data da $f(x) = 3x$ è lineare. Verifichiamo i tre assiomi:

- (1) $f(0) = 0$,
- (2) $f(x + x') = 3(x + x') = 3x + 3x' = f(x) + f(x') \quad \forall x, x' \in \mathbb{R}$,
- (3) $f(\lambda x) = 3\lambda x = \lambda 3x = \lambda f(x) \quad \forall x \in \mathbb{R}, \forall \lambda \in \mathbb{R}$.

Più in generale, la funzione $f(x) = kx$ è lineare per ogni $k \in \mathbb{R}$ fissato. Notiamo che per $k = 0$ otteniamo la funzione nulla $f(x) = 0$ e per $k = 1$ la funzione identità $f(x) = x$, già considerate nell'esempio precedente. Invece non sono lineari le seguenti funzioni da $\mathbb{R}$ in $\mathbb{R}$:

$$f(x) = 2x + 1, \quad f(x) = x^2.$$

Per la prima non vale l'assioma (1), per la seconda non vale l'assioma (2). Effettivamente dimostreremo che le uniche funzioni lineari da $\mathbb{R}$ in $\mathbb{R}$ sono quelle del tipo $f(x) = kx$. Gli assiomi (1), (2) e (3) sono piuttosto restrittivi.

L'esempio appena descritto può essere generalizzato in vari modi.

Esempio 4.1.3. Sia V uno spazio vettoriale e $\lambda \in \mathbb{K}$ uno scalare fissato. La funzione $f: V \rightarrow V$ data da $f(v) = \lambda v$ è lineare. Se $\lambda = 0$ otteniamo la funzione nulla e se $\lambda = 1$ otteniamo l'identità.

Esempio 4.1.4. Generalizziamo l'Esempio 4.1.2 prendendo un campo $\mathbb{K}$ qualsiasi e variando la dimensione del dominio. Siano $a_1, \dots, a_n \in \mathbb{K}$ costanti fissate. La funzione

$$f: \mathbb{K}^n \rightarrow \mathbb{K}, \quad f(x) = a_1 x_1 + \cdots + a_n x_n$$

è lineare. Infatti

- (1) $f(0) = a_1 0 + \cdots + a_n 0 = 0$,

$$(2) \quad f(x + x') = a_1(x_1 + x'_1) + \cdots + a_n(x_n + x'_n) = a_1x_1 + \cdots + a_nx_n + a_1x'_1 + \cdots + a_nx'_n = f(x) + f(x').$$

$$(3) \quad f(\lambda x) = a_1\lambda x_1 + \cdots + a_n\lambda x_n = \lambda(a_1x_1 + \cdots + a_nx_n) = \lambda f(x).$$

Ad esempio la funzione $f: \mathbb{R}^2 \rightarrow \mathbb{R}$, $f\begin{pmatrix} x \\ y \end{pmatrix} = 2x - 3y$ è lineare.

Esempio 4.1.5. Vogliamo generalizzare ulteriormente l'esempio precedente variando anche la dimensione del codominio. Prendiamo una matrice $A = (a_{ij})$ di taglia $m \times n$ e definiamo

$$L_A: \mathbb{K}^n \longrightarrow \mathbb{K}^m$$

usando il prodotto fra matrici e vettori nel modo seguente:

$$L_A(x) = Ax.$$

Nel dettaglio:

$$L_A(x) = Ax = \begin{pmatrix} a_{11} & \cdots & a_{1n} \\ \vdots & \ddots & \vdots \\ a_{m1} & \cdots & a_{mn} \end{pmatrix} \cdot \begin{pmatrix} x_1 \\ \vdots \\ x_n \end{pmatrix} = \begin{pmatrix} a_{11}x_1 + \cdots + a_{1n}x_n \\ \vdots \\ a_{m1}x_1 + \cdots + a_{mn}x_n \end{pmatrix}.$$

La lettera L sta per "left" e ci ricorda che L_A è l'applicazione che prende come input un vettore $x \in \mathbb{K}^n$ e lo moltiplica a sinistra per A , restituendo come output $Ax \in \mathbb{K}^m$. La linearità di L_A discende dalle proprietà delle matrici, elencate nella Proposizione 3.4.2. Verifichiamo infatti che:

$$(1) \quad L_A(0) = A0 = 0.$$

$$(2) \quad L_A(x + x') = A(x + x') = Ax + Ax' = L_A(x) + L_A(x').$$

$$(3) \quad L_A(\lambda x) = A(\lambda x) = \lambda Ax = \lambda L_A(x).$$

Ad esempio, se $A = \begin{pmatrix} 3 & -1 \\ 1 & 2 \end{pmatrix}$ otteniamo la funzione $L_A: \mathbb{R}^2 \rightarrow \mathbb{R}^2$ data da:

$$L_A\begin{pmatrix} x \\ y \end{pmatrix} = \begin{pmatrix} 3x - y \\ x + 2y \end{pmatrix}.$$

È spesso utile scrivere l'applicazione L_A anche in una forma alternativa che usa le colonne $A^1, \dots, A^n$ di A .

Proposizione 4.1.6. *Vale l'uguaglianza seguente:*

$$L_A(x) = x_1A^1 + \cdots + x_nA^n.$$

Dimostrazione. Si vede immediatamente che

$$\begin{aligned} L_A(x) &= \begin{pmatrix} a_{11}x_1 + \cdots + a_{1n}x_n \\ \vdots \\ a_{m1}x_1 + \cdots + a_{mn}x_n \end{pmatrix} = x_1 \begin{pmatrix} a_{11} \\ \vdots \\ a_{m1} \end{pmatrix} + \cdots + x_n \begin{pmatrix} a_{1n} \\ \vdots \\ a_{mn} \end{pmatrix} \\ &= x_1A^1 + \cdots + x_nA^n. \end{aligned}$$

La dimostrazione è conclusa. □

Ricordiamo che $e_1, \dots, e_n$ è la base canonica di $\mathbb{K}^n$. Ne deduciamo che le colonne di A sono precisamente le immagini dei vettori $e_1, \dots, e_n$.

Corollario 4.1.7. La colonna A^i è l'immagine di e_i , cioè

$$L_A(e_i) = A^i \quad \forall i = 1, \dots, n.$$

Esaminando le colonne della matrice A si vede subito dove va la base canonica. Ad esempio, se

$$A = \begin{pmatrix} 0 & 1 \\ 1 & 0 \end{pmatrix}$$

guardando le colonne capiamo subito che l'applicazione L_A scambia i due vettori e_1 ed e_2 della base canonica. Effettivamente, scrivendo esplicitamente

$$L_A \begin{pmatrix} x \\ y \end{pmatrix} = \begin{pmatrix} 0 & 1 \\ 1 & 0 \end{pmatrix} \begin{pmatrix} x \\ y \end{pmatrix} = \begin{pmatrix} y \\ x \end{pmatrix}$$

notiamo che L_A scambia le coordinate x e y e quindi i relativi assi.

Esempio 4.1.8. Due casi molto particolari sono importanti:

- Se $A = 0$ è la matrice $m \times n$ nulla, allora $L_A: \mathbb{K}^n \rightarrow \mathbb{K}^m$ è l'applicazione nulla.
- Se $A = I_n$ è la matrice identità $n \times n$, allora $L_A: \mathbb{K}^n \rightarrow \mathbb{K}^n$ è l'identità.

Infatti $L_0(x) = 0x = 0$ e $L_{I_n}(x) = I_n x = x$ per ogni $x \in \mathbb{K}^n$.

Abbiamo notato un fatto fondamentale: una matrice A di taglia $m \times n$ definisce una applicazione lineare $L_A: \mathbb{K}^n \rightarrow \mathbb{K}^m$. Le applicazioni lineari L_A compariranno ovunque in questo libro e vedremo molti esempi concreti nelle prossime pagine.

4.1.4. Esempi più elaborati. Costruiamo adesso alcune applicazioni lineari più sofisticate di varia natura. I numerosi esempi presenti in questa sezione mostrano che molte operazioni riguardanti spazi vettoriali, matrici, polinomi, sistemi lineari, funzioni, trasformazioni geometriche del piano e dello spazio, sono in realtà tutte applicazioni lineari ed è quindi conveniente considerarle con lo stesso linguaggio unificato. Il prezzo da pagare per questa unificazione è come sempre un livello di astrazione un po' più elevato rispetto a quello a cui siamo abituati; come ricompensa otterremo una teoria potente con numerose implicazioni algebriche e geometriche.

Esempio 4.1.9 (L_A sulle matrici). Sia A matrice $m \times n$. L'applicazione

$$L_A: M(n, k) \longrightarrow M(m, k)$$

definita da

$$L_A(X) = AX$$

è una applicazione lineare. Notiamo che adesso X è una matrice $n \times k$ e $L_A(X) = AX$ è una matrice $m \times k$. La linearità segue dalle proprietà algebriche delle matrici dimostrate nella Proposizione 3.4.2:

- (1) $L_A(0) = A0 = 0$.
- (2) $L_A(X + X') = A(X + X') = AX + AX' = L_A(X) + L_A(X')$.

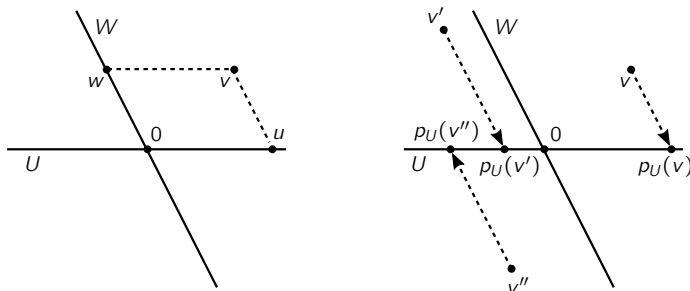


Figura 4.1. La proiezione $p_U: V \rightarrow U$ su un sottospazio determinata da una somma diretta $V = U \oplus W$.

$$(3) \quad L_A(\lambda X) = A(\lambda X) = \lambda AX = \lambda L_A(X).$$

Quando $k = 1$ questa applicazione lineare è la $L_A: \mathbb{K}^n \rightarrow \mathbb{K}^m$ già vista.

Esempio 4.1.10 (Trasposizione di matrici). La mappa

$$f: M(m, n) \longrightarrow M(n, m)$$

$$f(A) = {}^tA$$

che prende una matrice e la traspone è una applicazione lineare, perché

$${}^t0 = 0, \quad {}^t(A + B) = {}^tA + {}^tB, \quad {}^t(\lambda A) = \lambda {}^tA.$$

Esempio 4.1.11 (Proiezione). Sia $V = U \oplus W$ uno spazio vettoriale espresso come somma diretta dei sottospazi U e W . Definiamo una applicazione lineare

$$p_U: V \longrightarrow U$$

detta *proiezione* su U nel modo seguente: per la Proposizione 2.3.28, ogni vettore $v \in V$ si scrive in modo unico come $u + w$ con $u \in U$ e $w \in W$, e noi definiamo

$$p_U(v) = u.$$

Intuitivamente, il sottospazio U è il luogo su cui si proiettano tutti i punti, mentre il sottospazio complementare W indica la *direzione* lungo cui i punti vengono proiettati, come mostrato nella Figura 4.1.

La proiezione è effettivamente lineare: se $v = u + w$ e $v' = u' + w'$ allora $v + v' = u + u' + w + w'$ e quindi $p_U(v + v') = u + u' = p_U(v) + p_U(v')$. Le altre due proprietà sono dimostrate in modo analogo.

Scriviamo la proiezione in un caso concreto, quello della somma diretta $\mathbb{R}^2 = U \oplus W$ già analizzata nell'Esempio 2.3.32. Abbiamo

$$U = \text{Span} \begin{pmatrix} 1 \\ 0 \end{pmatrix}, \quad W = \text{Span} \begin{pmatrix} -1 \\ 2 \end{pmatrix}.$$

Dalla discussione dell'Esempio 2.3.32 segue che

$$p_U \begin{pmatrix} x \\ y \end{pmatrix} = \begin{pmatrix} x + \frac{y}{2} \\ 0 \end{pmatrix}.$$

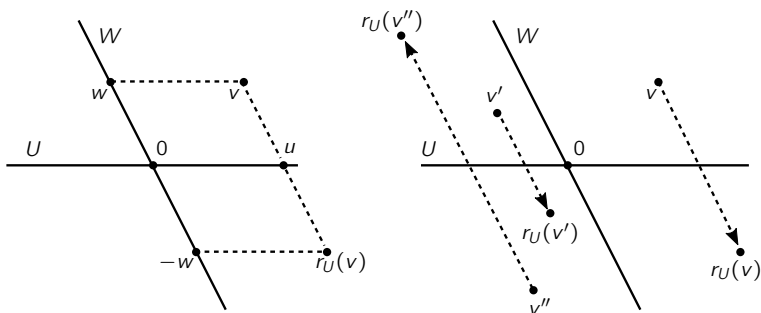


Figura 4.2. La riflessione $r_U: V \rightarrow V$ lungo un sottospazio determinata da una somma diretta $V = U \oplus W$.

Esempio 4.1.12 (Riflessione). Sia $V = U \oplus W$ come nell'esempio precedente. La *riflessione* lungo U è l'applicazione lineare

$$r_U: V \longrightarrow V$$

definita nel modo seguente: ogni vettore $v \in V$ si scrive in modo unico come $u + w$ con $u \in U$ e $w \in W$, e noi definiamo

$$r_U(v) = u - w.$$

Intuitivamente, lo spazio U è lo specchio lungo cui si riflettono i punti, mentre W è la direzione lungo cui riflettiamo. Si veda la Figura 4.2.

Scriviamo la riflessione con la somma diretta $\mathbb{R}^2 = U \oplus W$ dell'Esempio 2.3.32. Dai conti fatti precedentemente segue che

$$r_U \begin{pmatrix} x \\ y \end{pmatrix} = \left(x + \frac{y}{2}\right) \begin{pmatrix} 1 \\ 0 \end{pmatrix} - \frac{y}{2} \begin{pmatrix} -1 \\ 2 \end{pmatrix} = \begin{pmatrix} x+y \\ -y \end{pmatrix}.$$

Esempio 4.1.13 (Valore di un polinomio in un punto). Sia $x_0 \in \mathbb{K}$ un valore fissato. Consideriamo l'applicazione

$$T: \mathbb{K}[x] \longrightarrow \mathbb{K}$$

$$T(p) = p(x_0)$$

che assegna ad ogni polinomio p il suo valore $p(x_0)$ nel punto x_0 . L'applicazione T è lineare, infatti:

- (1) $T(0) = 0(x_0) = 0$,
- (2) $T(p + q) = (p + q)(x_0) = p(x_0) + q(x_0) = T(p) + T(q)$
- (3) $T(\lambda p) = \lambda p(x_0) = \lambda T(p)$.

Esempio 4.1.14 (Valori di una funzione in alcuni punti). Generalizziamo l'esempio precedente: invece che un punto solo ne prendiamo tanti, e invece che polinomi consideriamo funzioni più generali.

Sia X un insieme e $x_1, \dots, x_n \in X$ dei punti fissati. L'applicazione

$$T: F(X, \mathbb{K}) \longrightarrow \mathbb{K}^n$$

$$T(f) = \begin{pmatrix} f(x_1) \\ \vdots \\ f(x_n) \end{pmatrix}.$$

associa ad una funzione $f: X \rightarrow \mathbb{K}$ il vettore formato dai suoi valori nei punti $x_1, \dots, x_n$. Questa funzione è lineare (esercizio).

Esempio 4.1.15 (Derivata di un polinomio). La funzione

$$T: \mathbb{R}[x] \longrightarrow \mathbb{R}[x]$$

$$T(p) = p'$$

che associa ad un polinomio

$$p(x) = a_n x^n + a_{n-1} x^{n-1} + \dots + a_1 x + a_0$$

la sua derivata

$$p'(x) = n a_n x^{n-1} + (n-1) a_{n-1} x^{n-2} + \dots + a_1$$

è lineare. Si può dimostrare questo fatto direttamente oppure invocare il fatto che la derivata è sempre lineare per dei teoremi di analisi: se f' è la derivata di una funzione f valgono sempre le proprietà di linearità, che sono:

$$0' = 0, \quad (f + g)' = f' + g', \quad (\lambda f)' = \lambda f'.$$

Esempio 4.1.16 (Restrizioni). Se $f: V \rightarrow W$ è una applicazione lineare e $U \subset V$ è un sottospazio, possiamo costruire una applicazione lineare

$$f|_U: U \longrightarrow W$$

semplicemente restringendo il dominio di f da V a solo U . L'applicazione ristretta è indicata con $f|_U$ ed è chiamata la *restrizione* di f a U .

Esempio 4.1.17 (Coordinate di un vettore rispetto ad una base). Sia V uno spazio vettoriale su $\mathbb{K}$ e $\mathcal{B} = \{v_1, \dots, v_n\}$ una base¹ per V . Per ciascun $v \in V$, indichiamo con

$$[v]_{\mathcal{B}} = \begin{pmatrix} \lambda_1 \\ \vdots \\ \lambda_n \end{pmatrix} \in \mathbb{K}^n$$

il vettore colonna formato dalle coordinate $\lambda_1, \dots, \lambda_n$ di v rispetto alla base $\mathcal{B}$, definite nella Sezione 2.3.3 come gli unici scalari per cui

$$v = \lambda_1 v_1 + \dots + \lambda_n v_n.$$

¹Ad essere precisi, la notazione matematica $\{\}$ tra parentesi graffe indica un insieme *non ordinato* di vettori, mentre una base è in realtà un insieme *ordinato*. Con questo vogliamo dire che per noi sarà importante dire che v_1 è il primo vettore, v_2 il secondo, ecc. In particolare $\{v_1, v_2\}$ e $\{v_2, v_1\}$ sono due basi diverse. Generalmente gli insiemi ordinati vengono indicati con parentesi tonde $()$ invece che graffe $\{\}$, e quindi volendo essere rigorosi dovremmo scrivere $\mathcal{B} = (v_1, \dots, v_n)$. Preferiamo però usare comunque le graffe perché le tonde sono già usate abbondantemente per i vettori numerici e le matrici.

Otteniamo in questo modo una funzione

$$T: V \longrightarrow \mathbb{K}^n, \quad T(v) = [v]_B.$$

Verifichiamo che la funzione T è lineare:

- (1) $T(0) = 0$,
- (2) $T(v+v') = T(v) + T(v')$, infatti se $v = \lambda_1 v_1 + \cdots + \lambda_n v_n$ e $v' = \lambda'_1 v_1 + \cdots + \lambda'_n v_n$ allora $v + v' = (\lambda_1 + \lambda'_1)v_1 + \cdots + (\lambda_n + \lambda'_n)v_n$.
- (3) $T(\lambda v) = \lambda T(v)$, infatti se $v = \lambda_1 v_1 + \cdots + \lambda_n v_n$ allora ovviamente $\lambda v = \lambda \lambda_1 v_1 + \cdots + \lambda \lambda_n v_n$.

4.1.5. Uso delle basi. Siano V e W due spazi vettoriali su $\mathbb{K}$ arbitrari. Cerchiamo adesso un procedimento generale per scrivere tutte le applicazioni lineari da V in W . Come in altre parti di questo libro, lo strumento chiave è la scelta di una base per V .

La proposizione seguente dice che per costruire una applicazione lineare $V \rightarrow W$ è sufficiente fissare i valori su una qualsiasi base fissata per V .

Proposizione 4.1.18. Sia $v_1, \dots, v_n$ una base per V . Per ogni $w_1, \dots, w_n \in W$ esiste ed è unica una applicazione lineare

$$f: V \longrightarrow W$$

tale che $f(v_i) = w_i \forall i = 1, \dots, n$.

Dimostrazione. Mostriamo l'esistenza. Costruiamo f nel modo seguente: ogni $v \in V$ si scrive in modo unico come

$$v = \lambda_1 v_1 + \cdots + \lambda_n v_n$$

e noi definiamo $f: V \rightarrow W$ in modo molto naïf, sostituendo v_i con w_i :

$$f(v) = \lambda_1 w_1 + \cdots + \lambda_n w_n.$$

Si vede facilmente che questa funzione f è effettivamente lineare. D'altro canto, se $f: V \rightarrow W$ è una funzione lineare, otteniamo

$$f(v) = \lambda_1 f(v_1) + \cdots + \lambda_n f(v_n) = \lambda_1 w_1 + \cdots + \lambda_n w_n,$$

quindi la f definita prima è l'unica applicazione lineare che rispetti le condizioni date. La dimostrazione è conclusa. $\square$

4.1.6. Applicazioni da $\mathbb{K}^n$ a $\mathbb{K}^m$. Dimostriamo adesso un fatto a cui abbiamo già accennato: le applicazioni L_A determinate da matrici $A \in M(m, n, \mathbb{K})$ sono tutte e sole le applicazioni lineari da $\mathbb{K}^n$ in $\mathbb{K}^m$.

Proposizione 4.1.19. Per ogni $f: \mathbb{K}^n \rightarrow \mathbb{K}^m$ lineare esiste un'unica $A \in M(m, n, \mathbb{K})$ tale che $f = L_A$.

Dimostrazione. Sia $f: \mathbb{K}^n \rightarrow \mathbb{K}^m$ una funzione lineare. Ricordiamo che $e_1, \dots, e_n$ è la base canonica di $\mathbb{K}^n$. Sia A la matrice la cui colonna A^i è uguale all'immagine $f(e_i)$ di e_i . Brevemente:

$$A = (f(e_1) \cdots f(e_n)).$$

Per il Corollario 4.1.7, per ogni $i = 1, \dots, n$ otteniamo

$$L_A(e_i) = A^i = f(e_i).$$

Poiché $e_1, \dots, e_n$ è una base, per la Proposizione 4.1.18 abbiamo $f = L_A$.

Per quanto riguarda l'unicità, se $A \neq A'$ allora $A_{ij} \neq A'_{ij}$ per qualche i, j e quindi $L_A(e_j) \neq L_{A'}(e_j)$. Quindi A diverse danno L_A diverse. $\square$

4.2. Nucleo e immagine

Ad un'applicazione lineare vengono naturalmente associati due sottospazi vettoriali del dominio e del codominio, detti il *nucleo* e l'*immagine*.

4.2.1. Nucleo. Sia $f: V \rightarrow W$ un'applicazione lineare. Il *nucleo* di f è il sottoinsieme di V definito nel modo seguente:

$$\ker f = \{v \in V \mid f(v) = 0\}.$$

Il nucleo è l'insieme dei vettori che vengono mandati in zero da f . Il termine "ker" è un'abbreviazione dell'inglese *kernel*, che vuol dire nucleo.

Proposizione 4.2.1. *Il nucleo $\ker f$ è un sottospazio vettoriale di V .*

Dimostrazione. Dobbiamo verificare i 3 assiomi di sottospazio.

- (1) $0 \in \ker f$, infatti $f(0) = 0$.
- (2) $v, w \in \ker f \Rightarrow v + w \in \ker f$, infatti $f(v + w) = f(v) + f(w) = 0 + 0 = 0$.
- (3) $v \in \ker f, \lambda \in \mathbb{K} \Rightarrow \lambda v \in \ker f$, infatti $f(\lambda v) = \lambda f(v) = \lambda 0 = 0$.

La dimostrazione è conclusa. $\square$

Ricordiamo che una funzione $f: V \rightarrow W$ è iniettiva se $v \neq v' \Rightarrow f(v) \neq f(v')$. Capire se un'applicazione lineare è iniettiva è abbastanza facile: basta esaminare il nucleo. Sia $f: V \rightarrow W$ un'applicazione lineare.

Proposizione 4.2.2. *La funzione $f: V \rightarrow W$ è iniettiva $\Leftrightarrow \ker f = \{0\}$.*

Dimostrazione. ($\Rightarrow$) Sappiamo già che $f(0) = 0$. Siccome f è iniettiva, $v \neq 0 \Rightarrow f(v) \neq 0$. Quindi $\ker f = \{0\}$.

($\Leftarrow$) Siano $v, v' \in V$ diversi. Siccome $\ker f = \{0\}$, troviamo che $v - v' \neq 0 \Rightarrow f(v) - f(v') = f(v - v') \neq 0$, e quindi $f(v) \neq f(v')$. $\square$

Esempio 4.2.3. Sia A una matrice $m \times n$. Il nucleo dell'applicazione $L_A: \mathbb{K}^n \rightarrow \mathbb{K}^m$ è precisamente l'insieme S delle soluzioni del sistema lineare omogeneo $Ax = 0$. Quindi L_A è iniettiva precisamente se $S = \{0\}$.

Data una matrice A , a volte scriviamo $\ker A$ al posto di $\ker L_A$.

Esempio 4.2.4. Consideriamo una somma diretta $V = U \oplus W$. Il nucleo della proiezione p_U definita nell'Esempio 4.1.11 è precisamente W .

4.2.2. Immagine. Sia $f: V \rightarrow W$ un'applicazione lineare. Come tutte le funzioni, la f ha una immagine che indichiamo con $\text{Im } f \subset W$.

Proposizione 4.2.5. *L'immagine $\text{Im } f$ è un sottospazio vettoriale di W .*

Dimostrazione. Dimostriamo come sempre i 3 assiomi.

- (1) $0 \in \text{Im } f$ perché $f(0) = 0$.
- (2) $f(v), f(v') \in \text{Im } f \Rightarrow f(v) + f(v') \in \text{Im } f$ perché $f(v) + f(v') = f(v + v')$.
- (3) $f(v) \in \text{Im } f, \lambda \in \mathbb{K} \Rightarrow \lambda f(v) \in \text{Im } f$ perché $\lambda f(v) = f(\lambda v)$.

La dimostrazione è conclusa. $\square$

Ovviamente f è surgettiva $\iff \text{Im } f = W$. Concretamente, possiamo spesso determinare l'immagine di f usando la proposizione seguente.

Proposizione 4.2.6. *Se $v_1, \dots, v_n$ sono generatori di V , allora*

$$\text{Im } f = \text{Span}(f(v_1), \dots, f(v_n)).$$

Dimostrazione. Ogni vettore $v \in V$ si scrive come $v = \lambda_1 v_1 + \dots + \lambda_n v_n$ per qualche $\lambda_1, \dots, \lambda_n$. Quindi

$$\begin{aligned} \text{Im } f &= \{f(\lambda_1 v_1 + \dots + \lambda_n v_n) \mid \lambda_i \in \mathbb{K}\} \\ &= \{\lambda_1 f(v_1) + \dots + \lambda_n f(v_n) \mid \lambda_i \in \mathbb{K}\} \\ &= \text{Span}(f(v_1), \dots, f(v_n)). \end{aligned}$$

La dimostrazione è conclusa. $\square$

Corollario 4.2.7. *Sia A una matrice $m \times n$. L'immagine dell'applicazione $L_A: \mathbb{K}^n \rightarrow \mathbb{K}^m$ è precisamente lo spazio generato dalle colonne:*

$$\text{Im } L_A = \text{Span}(A^1, \dots, A^n).$$

Dimostrazione. Abbiamo $A^i = L_A(e_i)$ e $e_1, \dots, e_n$ generano $\mathbb{K}^n$. $\square$

Corollario 4.2.8. *Per ogni matrice A vale $\text{rk}(A) = \dim \text{Im } L_A$.*

Il rango di una matrice A è uguale alla dimensione dell'immagine di L_A . A volte scriviamo brevemente $\text{Im } A$ al posto di $\text{Im } L_A$.

4.2.3. Teorema della dimensione. Dimostriamo adesso un teorema centrale in algebra lineare.

Teorema 4.2.9 (Teorema della dimensione). *Sia $f: V \rightarrow W$ una funzione lineare. Se V ha dimensione finita n , allora*

$$\dim \ker f + \dim \text{Im } f = n.$$

Dimostrazione. Sia $v_1, \dots, v_k$ una base di $\ker f$. Possiamo sempre completarla a base $v_1, \dots, v_n$ di V . Il nostro scopo è dimostrare che $f(v_{k+1}), \dots, f(v_n)$ formano una base di $\text{Im } f$. Se ci riusciamo, abbiamo finito perché otteniamo $\dim \ker f = k$ e $\dim \text{Im } f = n - k$.

Dimostriamo che $f(v_{k+1}), \dots, f(v_n)$ generano $\text{Im } f$. Sappiamo dalla Proposizione 4.2.6 che $f(v_1), \dots, f(v_n)$ generano $\text{Im } f$. Sappiamo anche che $f(v_1) = \dots = f(v_k) = 0$, quindi possiamo rimuoverli dalla lista e ottenere che effettivamente $f(v_{k+1}), \dots, f(v_n)$ generano $\text{Im } f$.

Dimostriamo che $f(v_{k+1}), \dots, f(v_n)$ sono indipendenti. Supponiamo

$$\lambda_{k+1}f(v_{k+1}) + \dots + \lambda_n f(v_n) = 0.$$

Dobbiamo mostrare che $\lambda_{k+1} = \dots = \lambda_n = 0$. Per linearità di f , otteniamo

$$f(\lambda_{k+1}v_{k+1} + \dots + \lambda_n v_n) = 0.$$

Quindi $\lambda_{k+1}v_{k+1} + \dots + \lambda_n v_n \in \ker f$. Siccome $v_1, \dots, v_k$ è base di $\ker f$, possiamo scrivere

$$\lambda_{k+1}v_{k+1} + \dots + \lambda_n v_n = \alpha_1 v_1 + \dots + \alpha_k v_k$$

per qualche α_i . Spostando tutto a sinistra troviamo

$$-\alpha_1 v_1 - \dots - \alpha_k v_k + \lambda_{k+1}v_{k+1} + \dots + \lambda_n v_n = 0.$$

I vettori $v_1, \dots, v_n$ però sono indipendenti, quindi i coefficienti α_i, λ_j sono tutti nulli. La dimostrazione è conclusa. $\square$

Nel caso particolare di una applicazione lineare $L_A: \mathbb{K}^n \rightarrow \mathbb{K}^m$, il teorema della dimensione era già stato dimostrato con il Teorema di Rouché - Capelli per i sistemi lineari omogenei nella Sezione 3.2.5. Abbiamo quindi due dimostrazioni molto diverse dello stesso teorema: una usa le matrici e le mosse di Gauss, l'altra direttamente i vettori ed è più generale.

Esempio 4.2.10. Nella proiezione $p_U: V \rightarrow U$ definita nell'Esempio 4.1.11 a partire da una somma diretta $V = U \oplus W$, troviamo

$$\ker p_U = W, \quad \text{Im } p_U = U.$$

Quindi $V = U \oplus W = \text{Im } p_U \oplus \ker p_U$ e $\dim V = \dim \text{Im } p_U + \dim \ker p_U$.

Esempio 4.2.11. Se $A = \begin{pmatrix} 0 & 1 \\ 0 & 0 \end{pmatrix}$, per $L_A: \mathbb{K}^2 \rightarrow \mathbb{K}^2$ troviamo:

$$\text{Im } L_A = \text{Span}(A^1, A^2) = \text{Span}(e_1),$$

$$\ker L_A = \left\{ \begin{pmatrix} x \\ y \end{pmatrix} \in \mathbb{K}^2 \mid y = 0 \right\} = \text{Span}(e_1).$$

Quindi $\dim \ker L_A + \dim \text{Im } L_A = 1 + 1 = 2 = \dim \mathbb{K}^2$.

Nei prossimi esempi vediamo come usare il teorema della dimensione per determinare agevolmente il nucleo o l'immagine di una applicazione lineare.

Esempio 4.2.12. Sia $f: \mathbb{K}_2[x] \rightarrow \mathbb{K}$ la mappa $f(p) = p(2)$ dell'Esempio 4.1.13. Vogliamo determinare una base di $\ker f$.

Per ogni $\lambda \in \mathbb{K}$ prendiamo il polinomio costante $p(x) = \lambda \in \mathbb{K}$ e notiamo che $f(p) = \lambda$. Quindi f è suriettiva. Quindi $\dim \text{Im } f = 1$ e allora $\dim \ker f = \dim \mathbb{K}_2[x] - \dim \text{Im } f = 3 - 1 = 2$.

Sappiamo che $x - 2, x^2 - 2x \in \ker f$. Sono indipendenti, quindi $x - 2, x^2 - 2x$ sono una base di $\ker f$ per la Proposizione 2.3.23.

Esempio 4.2.13. Sia $T: \mathbb{R}_n[x] \rightarrow \mathbb{R}_n[x]$ la funzione derivata $T(p) = p'$ già considerata nell'Esempio 4.1.15. Vogliamo determinarne l'immagine.

Per definizione $\ker T = \{p(x) \mid p'(x) = 0\}$ è la retta formata dai polinomi costanti. Quindi $\dim \ker T = 1$ e allora $\dim \operatorname{Im} T = \dim \mathbb{R}_n[x] - \dim \ker T = n + 1 - 1 = n$. Per costruzione l'immagine è contenuta nel sottospazio $\mathbb{R}_{n-1}[x]$ formato dai polinomi di grado $\leq n - 1$, che ha anche lui dimensione n . Quindi

$$\operatorname{Im} T = \mathbb{R}_{n-1}[x].$$

Possiamo usare il teorema della dimensione per dedurre delle caratterizzazioni dell'injectività e della suriettività.

Corollario 4.2.14. *Sia $f: V \rightarrow W$ un'applicazione lineare. Vale*

$$\dim \operatorname{Im} f \leq \dim V.$$

Inoltre:

- (1) f *iniettiva* $\iff \dim \operatorname{Im} f = \dim V$.
- (2) f *suriettiva* $\iff \dim \operatorname{Im} f = \dim W$.

Dimostrazione. Ricordiamo che

$$\begin{aligned} f \text{ iniettiva} &\iff \ker f = \{0\} \iff \dim \ker f = 0, \\ f \text{ suriettiva} &\iff \operatorname{Im} f = W \iff \dim \operatorname{Im} f = \dim W. \end{aligned}$$

Per il teorema della dimensione $\dim \ker f = 0 \iff \dim \operatorname{Im} f = \dim V$. $\square$

Corollario 4.2.15. *Sia $f: V \rightarrow W$ un'applicazione lineare e $v_1, \dots, v_n$ una base per V . Valgono i fatti seguenti:*

- f *iniettiva* $\iff f(v_1), \dots, f(v_n)$ *sono indipendenti*.
- f *suriettiva* $\iff f(v_1), \dots, f(v_n)$ *generano W* .

Esempio 4.2.16. Sia $A \in M(m, n, \mathbb{K})$. La funzione $L_A: \mathbb{K}^n \rightarrow \mathbb{K}^m$ è

- *iniettiva* $\iff \operatorname{rk} A = n$,
- *suriettiva* $\iff \operatorname{rk} A = m$.

4.2.4. Composizione di applicazioni lineari. La composizione di due applicazioni lineari è sempre lineare.

Proposizione 4.2.17. *Se $f: V \rightarrow W$ e $g: W \rightarrow Z$ sono funzioni lineari, anche la composizione $g \circ f: V \rightarrow Z$ lo è.*

Dimostrazione. Valgono i 3 assiomi di linearità per $g \circ f$:

- (1) $g(f(0)) = g(0) = 0$.
- (2) $v, v' \in V \Rightarrow g(f(v + v')) = g(f(v) + f(v')) = g(f(v)) + g(f(v'))$.
- (3) $v \in V, \lambda \in \mathbb{K} \Rightarrow g(f(\lambda v)) = g(\lambda f(v)) = \lambda g(f(v))$.

La dimostrazione è conclusa. $\square$

Per le applicazioni di tipo L_A la composizione corrisponde precisamente al prodotto fra matrici:

Proposizione 4.2.18. *Siano $A \in M(k, m, \mathbb{K})$ e $B \in M(m, n, \mathbb{K})$. Consideriamo $L_A: \mathbb{K}^m \rightarrow \mathbb{K}^k$ e $L_B: \mathbb{K}^n \rightarrow \mathbb{K}^m$. Vale la relazione*

$$L_A \circ L_B = L_{AB}.$$

Dimostrazione. Segue dall'associatività del prodotto fra matrici:

$$L_A(L_B(x)) = A(Bx) = (AB)x = L_{AB}(x).$$

La dimostrazione è conclusa. $\square$

Il prodotto fra matrici codifica in modo efficace la composizione di due applicazioni lineari.

4.2.5. Isomorfismi. Dopo aver studiato le applicazioni lineari iniettive e suriettive, è naturale occuparsi adesso di quelle bigettive. Queste sono così importanti da meritare un nome specifico.

Definizione 4.2.19. Un'applicazione lineare $f: V \rightarrow W$ è un *isomorfismo* se è bigettiva.

Ricordiamo che una funzione f è bigettiva se e solo se è contemporaneamente iniettiva e suriettiva. Sappiamo che se $f: V \rightarrow W$ è bigettiva, esiste una inversa $f^{-1}: W \rightarrow V$.

Proposizione 4.2.20. *Se una funzione lineare $f: V \rightarrow W$ è bigettiva, l'inversa $f^{-1}: W \rightarrow V$ è anch'essa lineare.*

Dimostrazione. Valgono i 3 assiomi:

- (1) Chiaramente $f^{-1}(0) = 0$.
- (2) $f^{-1}(w + w') = f^{-1}(w) + f^{-1}(w')$. Infatti scriviamo $v = f^{-1}(w)$, $v' = f^{-1}(w')$ e notiamo che $f(v + v') = f(v) + f(v') = w + w'$. Quindi $f^{-1}(w + w') = v + v' = f^{-1}(w) + f^{-1}(w')$.
- (3) $f^{-1}(\lambda w) = \lambda f^{-1}(w)$. Infatti scriviamo $v = f^{-1}(w)$ e notiamo che $f(\lambda v) = \lambda f(v) = \lambda w$. Quindi $f^{-1}(\lambda w) = \lambda v = \lambda f^{-1}(w)$.

La dimostrazione è conclusa. $\square$

Ricordiamo che una matrice quadrata $A \in M(n)$ è invertibile se ha un'inversa A^{-1} per il prodotto. C'è uno stretto legame fra matrici invertibili e isomorfismi.

Proposizione 4.2.21. *Una applicazione lineare $L_A: \mathbb{K}^n \rightarrow \mathbb{K}^n$ è un isomorfismo $\iff A$ è invertibile.*

Dimostrazione. L'applicazione lineare L_A è un isomorfismo se e solo se esiste una L_B che la inverte, cioè tale $L_A \circ L_B = L_B \circ L_A$ sia l'identità. Tradotto in matrici vuol dire

$$AB = BA = I_n$$

che è esattamente la nozione di invertibilità per A . $\square$

Dato uno spazio vettoriale V ed una base $\mathcal{B}$ per V , ricordiamo la mappa $T: V \rightarrow \mathbb{K}^n$ che associa ad ogni $v \in V$ le sue coordinate $T(v) = [v]_{\mathcal{B}}$, definita nell'Esempio 4.1.17. La mappa T è un esempio importante di isomorfismo.

Proposizione 4.2.22. *La mappa $T: V \rightarrow \mathbb{K}^n$ è un isomorfismo.*

Dimostrazione. Se $\mathcal{B} = \{v_1, \dots, v_n\}$, l'inversa $T^{-1}: \mathbb{K}^n \rightarrow V$ è

$$T^{-1}(x) = x_1 v_1 + \dots + x_n v_n.$$

Si verifica che effettivamente $T^{-1} \circ T = \text{id}_V$ e $T \circ T^{-1} = \text{id}_{\mathbb{K}^n}$. $\square$

Il fatto seguente segue dal Corollario 4.2.14.

Corollario 4.2.23. *Sia $f: V \rightarrow W$ un'applicazione lineare.*

- (1) *Se f è iniettiva, allora $\dim V \leq \dim W$.*
- (2) *Se f è suriettiva, allora $\dim V \geq \dim W$.*
- (3) *Se f è un isomorfismo, allora $\dim V = \dim W$.*

La proposizione seguente è analoga alle Proposizioni 2.3.23 e 2.3.29, ed è utile concretamente per dimostrare che una data f è un isomorfismo.

Proposizione 4.2.24. *Sia $f: V \rightarrow W$ un'applicazione lineare. Se $\dim V = \dim W$, allora i fatti seguenti sono equivalenti:*

- (1) *f è iniettiva,*
- (2) *f è suriettiva,*
- (3) *f è un isomorfismo.*

Dimostrazione. Per il teorema della dimensione

$$f \text{ iniettiva} \iff \dim \ker f = 0 \iff \dim \text{Im } f = n \iff f \text{ suriettiva}$$

dove $n = \dim V = \dim W$. La dimostrazione è conclusa. $\square$

Osservazione 4.2.25. Notiamo delle forti analogie fra le Proposizioni 2.3.23, 2.3.29 e 4.2.24. In tutti e tre i casi abbiamo una definizione (base, somma diretta, isomorfismo) che prevede due condizioni (indipendenti e generatori, somma totale e intersezione banale, iniettiva e suriettiva), ed un criterio che dice che *se le dimensioni sono giuste* allora una qualsiasi delle due condizioni in realtà implica l'altra.

Esempio 4.2.26. La trasposizione $f: M(n, m, \mathbb{K}) \rightarrow M(n, m, \mathbb{K})$ data da $f(A) = {}^t A$ considerata nell'Esempio 4.1.10 è un isomorfismo. La sua inversa è ancora la trasposizione, visto che ${}^t({}^t A) = A$.

Osservazione 4.2.27. Per una matrice quadrata $A \in M(n, \mathbb{K})$, tutti i fatti seguenti sono equivalenti:

$$\begin{aligned} L_A \text{ iniettiva} &\iff L_A \text{ suriettiva} \iff L_A \text{ isomorfismo} \iff A \text{ invertibile} \\ &\iff \det A \neq 0 \iff \text{rk } A = n \iff \forall b \in \mathbb{K}^n \exists! x \in \mathbb{K}^n : Ax = b. \end{aligned}$$

Tornando alle composizioni, dimostriamo una disuguaglianza che è a volte un'uguaglianza in presenza di isomorfismi.

Proposizione 4.2.28. *Siano $f: V \rightarrow W$ e $g: W \rightarrow Z$ lineari. Allora:*

- $\dim \operatorname{Im}(g \circ f) \leq \min\{\dim \operatorname{Im} g, \dim \operatorname{Im} f\}$.
- Se f è un isomorfismo, allora $\dim \operatorname{Im}(g \circ f) = \dim \operatorname{Im} g$.
- Se g è un isomorfismo, allora $\dim \operatorname{Im}(g \circ f) = \dim \operatorname{Im} f$.

Dimostrazione. Vale $\operatorname{Im}(g \circ f) \subset \operatorname{Im} g$, quindi $\dim \operatorname{Im}(g \circ f) \leq \dim \operatorname{Im} g$. Inoltre $\operatorname{Im}(g \circ f) = \operatorname{Im}(g|_{\operatorname{Im} f})$ (si veda l'Esempio 4.1.16 per la definizione di restrizione) e quindi $\dim \operatorname{Im}(g \circ f) \leq \dim \operatorname{Im} f$ per il Corollario 4.2.14.

Se f è un isomorfismo, allora $(g \circ f) \circ f^{-1} = g$ e applicando la disuguaglianza otteniamo anche $\dim \operatorname{Im} g \leq \dim \operatorname{Im}(g \circ f)$. L'altro caso è analogo. $\square$

4.2.6. Rango del prodotto. Usiamo le proprietà delle applicazioni lineari per mostrare agevolmente alcune proprietà del rango di una matrice.

Proposizione 4.2.29. *Siano A e B due matrici che si possono moltiplicare. Valgono i fatti seguenti:*

- $\operatorname{rk}(AB) \leq \min\{\operatorname{rk} A, \operatorname{rk} B\}$.
- Se A è invertibile, allora $\operatorname{rk}(AB) = \operatorname{rk} B$.
- Se B è invertibile, allora $\operatorname{rk}(AB) = \operatorname{rk} A$.

Dimostrazione. Appliciamo la Proposizione 4.2.28 a L_A e L_B . $\square$

4.2.7. Spazi vettoriali isomorfi. Diciamo che due spazi vettoriali V e W sullo stesso campo $\mathbb{K}$ sono *isomorfi* se esiste un isomorfismo $f: V \rightarrow W$. Per capire se due spazi vettoriali sono isomorfi è sufficiente calcolare le loro dimensioni.

Proposizione 4.2.30. *Due spazi vettoriali V e W sullo stesso campo $\mathbb{K}$ sono isomorfi $\iff \dim V = \dim W$.*

Dimostrazione. $(\Rightarrow)$ È conseguenza del Corollario 4.2.23.

$(\Leftarrow)$ Per ipotesi V e W hanno due basi $v_1, \dots, v_n$ e $w_1, \dots, w_n$ entrambe di n elementi. Definiamo un'applicazione lineare $f: V \rightarrow W$ imponendo $f(v_i) = w_i$, si veda la Proposizione 4.1.18. L'applicazione lineare f è un isomorfismo per il Corollario 4.2.15. $\square$

Il fatto seguente è anche conseguenza della Proposizione 4.2.22.

Corollario 4.2.31. *Tutti gli spazi vettoriali su $\mathbb{K}$ di dimensione n sono isomorfi a $\mathbb{K}^n$.*

Ad esempio $\mathbb{K}_n[k]$ è isomorfo a $\mathbb{K}^{n+1}$ e $M(m, n, \mathbb{K})$ è isomorfo a $\mathbb{K}^{mn}$.

4.3. Matrice associata

Introduciamo in questa sezione delle tecniche che ci permettono di rappresentare vettori e applicazioni lineari fra spazi vettoriali arbitrari come vettori numerici e matrici. In questa rappresentazione giocano un ruolo fondamentale le basi.

4.3.1. Definizione. Se fissiamo una base $\mathcal{B} = \{v_1, \dots, v_n\}$ di uno spazio V , possiamo rappresentare qualsiasi vettore $v \in V$ come un vettore numerico, quello delle sue coordinate $[v]_{\mathcal{B}} \in \mathbb{K}^n$. Adesso facciamo lo stesso con le applicazioni lineari: fissando opportune basi, possiamo rappresentare qualsiasi applicazione lineare come una matrice.

Sia $f: V \rightarrow W$ un'applicazione lineare fra spazi vettoriali definiti su $\mathbb{K}$. Siano inoltre

$$\mathcal{B} = \{v_1, \dots, v_n\}, \quad \mathcal{C} = \{w_1, \dots, w_m\}$$

due basi rispettivamente di V e di W . Sappiamo che

$$f(v_1) = a_{11}w_1 + \dots + a_{m1}w_m,$$

$$\vdots$$

$$f(v_n) = a_{1n}w_1 + \dots + a_{mn}w_m.$$

per qualche insieme di coefficienti $a_{ij} \in \mathbb{K}$. Definiamo la *matrice associata* a f nelle basi $\mathcal{B}$ e $\mathcal{C}$ come la matrice $m \times n$

$$A = (a_{ij})$$

che raggruppa questi coefficienti, e la indichiamo con il simbolo

$$A = [f]_{\mathcal{C}}^{\mathcal{B}}$$

che ci ricorda che A dipende da f , da $\mathcal{B}$ e da $\mathcal{C}$. Notiamo che la base in partenza $\mathcal{B}$ sta in alto e quella in arrivo $\mathcal{C}$ sta in basso. Notiamo anche che la j -esima colonna di A contiene le coordinate di $f(v_j)$ rispetto a $\mathcal{C}$, cioè

$$A^j = \begin{pmatrix} a_{1j} \\ \vdots \\ a_{mj} \end{pmatrix} = [f(v_j)]_{\mathcal{C}}.$$

Per prima cosa, valutiamo come sempre L_A .

Proposizione 4.3.1. *La matrice associata a L_A rispetto alle basi canoniche di $\mathbb{K}^n$ e $\mathbb{K}^m$ è proprio A .*

Dimostrazione. Per costruzione $f(e_j) = a_{1j}e_1 + \dots + a_{nj}e_n$. □

4.3.2. Esempi. Ci familiarizziamo con la nozione di matrice associata con qualche esempio.

Esempio 4.3.2. Consideriamo l'applicazione lineare

$$f: \mathbb{R}_2[x] \longrightarrow \mathbb{R}^2, \quad f(p) = \begin{pmatrix} p(2) \\ p(-2) \end{pmatrix}$$

che assegna ad ogni polinomio i suoi valori in 2 e in -2 . Scriviamo la matrice associata a f nelle basi canoniche $\mathcal{C}' = \{1, x, x^2\}$ di $\mathbb{R}_2[x]$ e $\mathcal{C} = \{e_1, e_2\}$ di $\mathbb{R}^2$. Calcoliamo:

$$f(1) = \begin{pmatrix} 1 \\ 1 \end{pmatrix}, \quad f(x) = \begin{pmatrix} 2 \\ -2 \end{pmatrix}, \quad f(x^2) = \begin{pmatrix} 4 \\ 4 \end{pmatrix}.$$

Quindi la matrice associata è

$$[f]_{\mathcal{C}}^{\mathcal{C}'} = \begin{pmatrix} 1 & 2 & 4 \\ 1 & -2 & 4 \end{pmatrix}.$$

Se in arrivo prendiamo la base $\mathcal{B} = \left\{ \begin{pmatrix} 1 \\ -1 \end{pmatrix}, \begin{pmatrix} 0 \\ 1 \end{pmatrix} \right\}$ invece della canonica $\mathcal{C}$, allora dobbiamo anche calcolare le coordinate di ciascun vettore immagine rispetto a $\mathcal{B}$. Risolvendo dei semplici sistemi lineari troviamo:

$$f(1) = \begin{pmatrix} 1 \\ 1 \end{pmatrix} = \begin{pmatrix} 1 \\ -1 \end{pmatrix} + 2 \begin{pmatrix} 0 \\ 1 \end{pmatrix},$$

$$f(x) = \begin{pmatrix} 2 \\ -2 \end{pmatrix} = 2 \begin{pmatrix} 1 \\ -1 \end{pmatrix} + 0 \begin{pmatrix} 0 \\ 1 \end{pmatrix},$$

$$f(x^2) = \begin{pmatrix} 4 \\ 4 \end{pmatrix} = 4 \begin{pmatrix} 1 \\ -1 \end{pmatrix} + 8 \begin{pmatrix} 0 \\ 1 \end{pmatrix}$$

e la matrice associata diventa quindi

$$[f]_{\mathcal{B}}^{\mathcal{C}'} = \begin{pmatrix} 1 & 2 & 4 \\ 2 & 0 & 8 \end{pmatrix}.$$

Esempio 4.3.3. Consideriamo l'applicazione lineare

$$f: \mathbb{C}^2 \longrightarrow \mathbb{C}^3, \quad f \begin{pmatrix} x \\ y \end{pmatrix} = \begin{pmatrix} x - y \\ 2x \\ y \end{pmatrix}.$$

La matrice associata a f rispetto alle basi canoniche è semplicemente quella dei coefficienti:

$$\begin{pmatrix} 1 & -1 \\ 2 & 0 \\ 0 & 1 \end{pmatrix}.$$

Se invece come basi prendiamo in partenza

$$v_1 = \begin{pmatrix} 1 \\ 1 \end{pmatrix}, \quad v_2 = \begin{pmatrix} 1 \\ -1 \end{pmatrix}$$

e in arrivo

$$w_1 = \begin{pmatrix} 1 \\ 1 \\ 0 \end{pmatrix}, \quad w_2 = \begin{pmatrix} 0 \\ 1 \\ 1 \end{pmatrix}, \quad w_3 = \begin{pmatrix} 1 \\ 0 \\ 1 \end{pmatrix}$$

allora per calcolare la matrice associata dobbiamo fare dei conti. Dobbiamo determinare le coordinate di $f(v_1)$ e $f(v_2)$ rispetto alla base w_1, w_2, w_3 . Risolvendo dei sistemi lineari troviamo

$$f(v_1) = \begin{pmatrix} 0 \\ 2 \\ 1 \end{pmatrix} = \frac{1}{2} \begin{pmatrix} 1 \\ 1 \\ 0 \end{pmatrix} + \frac{3}{2} \begin{pmatrix} 0 \\ 1 \\ 1 \end{pmatrix} - \frac{1}{2} \begin{pmatrix} 1 \\ 0 \\ 1 \end{pmatrix},$$

$$f(v_2) = \begin{pmatrix} 2 \\ 2 \\ -1 \end{pmatrix} = \frac{5}{2} \begin{pmatrix} 1 \\ 1 \\ 0 \end{pmatrix} - \frac{1}{2} \begin{pmatrix} 0 \\ 1 \\ 1 \end{pmatrix} - \frac{1}{2} \begin{pmatrix} 1 \\ 0 \\ 1 \end{pmatrix}.$$

Quindi la matrice associata a f rispetto alle basi v_1, v_2 e w_1, w_2, w_3 è

$$[f]_{w_1, w_2, w_3}^{v_1, v_2} = \begin{pmatrix} \frac{1}{2} & \frac{5}{2} \\ \frac{3}{2} & -\frac{1}{2} \\ -\frac{1}{2} & -\frac{1}{2} \end{pmatrix} = \frac{1}{2} \begin{pmatrix} 1 & 5 \\ 3 & -1 \\ -1 & -1 \end{pmatrix}.$$

4.3.3. Proprietà. Abbiamo scoperto come caratterizzare qualsiasi applicazione lineare (fra spazi di dimensione finita) come una matrice. Adesso mostriamo le proprietà di questa caratterizzazione, che sono notevoli.

Intanto notiamo che dalla matrice associata possiamo facilmente calcolare l'immagine di qualsiasi vettore. Sia $f: V \rightarrow W$ un'applicazione lineare e siano $\mathcal{B} = \{v_1, \dots, v_n\}$ e $\mathcal{C} = \{w_1, \dots, w_m\}$ basi di V e W .

Proposizione 4.3.4. *Per ogni $v \in V$ troviamo*

$$[f(v)]_{\mathcal{C}} = [f]_{\mathcal{C}}^{\mathcal{B}} [v]_{\mathcal{B}}.$$

Dimostrazione. Dato $v = \lambda_1 v_1 + \dots + \lambda_n v_n$, troviamo

$$f(v) = \lambda_1 f(v_1) + \dots + \lambda_n f(v_n)$$

e quindi

$$[f(v)]_{\mathcal{C}} = \lambda_1 [f(v_1)]_{\mathcal{C}} + \dots + \lambda_n [f(v_n)]_{\mathcal{C}} = [f]_{\mathcal{C}}^{\mathcal{B}} [v]_{\mathcal{B}}.$$

La dimostrazione è completa. $\square$

Per trovare le coordinate di $f(v)$ rispetto a $\mathcal{C}$ è sufficiente moltiplicare la matrice associata a f per le coordinate di v rispetto a $\mathcal{B}$.

Osservazione 4.3.5. Scriviamo $x = [v]_{\mathcal{B}}$, $A = [f]_{\mathcal{C}}^{\mathcal{B}}$, $y = [f(v)]_{\mathcal{C}}$. allora

$$y = Ax = L_A(x).$$

Abbiamo quindi scoperto un fatto importante: dopo aver scelto due basi per V e W , qualsiasi applicazione lineare $V \rightarrow W$ può essere interpretata in coordinate come un'applicazione del tipo $L_A: \mathbb{K}^n \rightarrow \mathbb{K}^m$. È sufficiente sostituire i vettori v e $f(v)$ con le loro coordinate x e y , ed usare la matrice associata A .

Esempio 4.3.6. Nell'Esempio 4.3.2, abbiamo ottenuto

$$[f]_{\mathcal{C}}^{\mathcal{C}'} = \begin{pmatrix} 1 & 2 & 4 \\ 1 & -2 & 4 \end{pmatrix}$$

rispetto alle basi canoniche. Possiamo usare questa matrice per calcolare in coordinate l'immagine di qualsiasi polinomio. Il polinomio generico $p(x) = ax^2 + bx + c$ ha come coordinate rispetto a $\mathcal{C}'$ i suoi coefficienti (in ordine inverso). Quindi la sua immagine $f(p)$ ha coordinate

$$\begin{pmatrix} 1 & 2 & 4 \\ 1 & -2 & 4 \end{pmatrix} \begin{pmatrix} c \\ b \\ a \end{pmatrix} = \begin{pmatrix} 4a + 2b + c \\ 4a - 2b + c \end{pmatrix}$$

rispetto alla base canonica $\mathcal{C}$ di $\mathbb{R}^2$. Come verifica, possiamo calcolare la f direttamente dalla definizione e ottenere effettivamente:

$$f(ax^2 + bx + c) = \begin{pmatrix} 4a + 2b + c \\ 4a - 2b + c \end{pmatrix}.$$

Dimostriamo un fatto analogo con la composizione. Siano $f: U \rightarrow V$ e $g: V \rightarrow W$ due applicazioni lineari. Siano $\mathcal{B}, \mathcal{C}$ e $\mathcal{D}$ basi di U, V e W .

Proposizione 4.3.7. *Troviamo*

$$[g \circ f]_{\mathcal{D}}^{\mathcal{B}} = [g]_{\mathcal{D}}^{\mathcal{C}} [f]_{\mathcal{C}}^{\mathcal{B}}.$$

Dimostrazione. Sia $\mathcal{B} = \{v_1, \dots, v_n\}$. La colonna i -esima di $[g \circ f]_{\mathcal{D}}^{\mathcal{B}}$ è

$$[g(f(v_i))]_{\mathcal{D}} = [g]_{\mathcal{D}}^{\mathcal{C}} [f(v_i)]_{\mathcal{C}}$$

per la proposizione precedente. D'altra parte $[f(v_i)]_{\mathcal{C}}$ è la colonna i -esima di $[f]_{\mathcal{C}}^{\mathcal{B}}$ e quindi concludiamo. $\square$

Abbiamo scoperto che nel passaggio dalle funzioni lineari alle matrici la composizione di funzioni diventa moltiplicazione fra matrici. Ora mostriamo che l'applicazione identità (che indichiamo sempre con id_V) si trasforma sempre nella matrice identità, purché si scelga la stessa base in partenza ed in arrivo.

Proposizione 4.3.8. *Sia $\mathcal{B}$ una base di uno spazio vettoriale V di dimensione n . Troviamo*

$$[\text{id}_V]_{\mathcal{B}}^{\mathcal{B}} = I_n.$$

Dimostrazione. Sia $\mathcal{B} = \{v_1, \dots, v_n\}$. Poiché $f(v_i) = v_i$, la matrice associata è I_n per costruzione. $\square$

Adesso mostriamo che qualsiasi matrice è matrice associata di un opportuna applicazione lineare. Siano $\mathcal{B}, \mathcal{C}$ basi per V e W , con $n = \dim V$ e $m = \dim W$.

Proposizione 4.3.9. *Qualsiasi matrice $A \in M(m, n)$ è la matrice associata $A = [f]_{\mathcal{C}}^{\mathcal{B}}$ ad un'unica applicazione lineare $f: V \rightarrow W$.*

Dimostrazione. Siano $\mathcal{B} = \{v_1, \dots, v_n\}$ e $\mathcal{C} = \{w_1, \dots, w_m\}$. È sufficiente definire f imponendo che

$$f(v_i) = a_{1i}w_1 + \dots + a_{mi}w_m.$$

La f esiste ed è unica per la Proposizione 4.1.18. $\square$

Abbiamo quindi ottenuto una corrispondenza biunivoca fra applicazioni lineari e matrici (sempre dopo aver fissato delle basi). Mostriamo adesso che questa corrispondenza associa isomorfismi e matrici invertibili. Sia $f: V \rightarrow W$ una applicazione lineare, e siano $\mathcal{B}$ e $\mathcal{C}$ basi per V e W .

Corollario 4.3.10. *La funzione f è un isomorfismo $\iff$ la matrice associata $[f]_{\mathcal{C}}^{\mathcal{B}}$ è invertibile, ed in questo caso la sua inversa è $[f^{-1}]_{\mathcal{B}}^{\mathcal{C}}$.*

Dimostrazione. Se f è un isomorfismo, ha una inversa $f^{-1}: W \rightarrow V$. Dalle proposizioni precedenti deduciamo che

$$I_n = [\text{id}_V]_{\mathcal{B}}^{\mathcal{B}} = [f^{-1} \circ f]_{\mathcal{B}}^{\mathcal{B}} = [f^{-1}]_{\mathcal{B}}^{\mathcal{C}} [f]_{\mathcal{C}}^{\mathcal{B}}$$

e analogamente $I_m = [\text{id}_W]_{\mathcal{C}}^{\mathcal{C}} = [f]_{\mathcal{C}}^{\mathcal{B}} [f^{-1}]_{\mathcal{B}}^{\mathcal{C}}$. Quindi effettivamente $[f^{-1}]_{\mathcal{B}}^{\mathcal{C}}$ è la matrice inversa di $[f]_{\mathcal{C}}^{\mathcal{B}}$.

D'altro canto, se $A = [f]_{\mathcal{C}}^{\mathcal{B}}$ è invertibile, grazie alla proposizione precedente esiste una $g: W \rightarrow V$ tale che $[g]_{\mathcal{B}}^{\mathcal{C}} = A^{-1}$ e deduciamo analogamente che g è l'inversa di f , quindi f è un isomorfismo. $\square$

4.3.4. Hom. Descriviamo un altro legame forte fra applicazioni lineari e matrici. Siano V e W spazi vettoriali su $\mathbb{K}$ di dimensione n e m . Sia

$$\text{Hom}(V, W)$$

l'insieme di tutte le applicazioni lineari $V \rightarrow W$. Il nome "hom" deriva dal fatto che in matematica il termine *omomorfismo* è usato come sinonimo di applicazione lineare.

Non dovrebbe sorprendere il fatto che $\text{Hom}(V, W)$ abbia una naturale struttura di spazio vettoriale. Date $f, g \in \text{Hom}(V, W)$ e $\lambda \in \mathbb{K}$ definiamo $f + g: V \rightarrow W$ e $\lambda f: V \rightarrow W$ nel modo usuale:

$$(f + g)(v) = f(v) + g(v),$$

$$(\lambda f)(v) = \lambda f(v).$$

Lo zero in $\text{Hom}(V, W)$ è l'applicazione nulla. Siano $\mathcal{B}$ e $\mathcal{C}$ basi di V e W .

Proposizione 4.3.11. *La mappa*

$$T: \text{Hom}(V, W) \longrightarrow M(m, n, \mathbb{K}), \quad T(f) = [f]_{\mathcal{C}}^{\mathcal{B}}$$

è un isomorfismo di spazi vettoriali.

Dimostrazione. La linearità di T segue facilmente dalla definizione di matrice associata. È un isomorfismo per la Proposizione 4.3.9. $\square$

Corollario 4.3.12. *Lo spazio $\text{Hom}(V, W)$ ha dimensione mn .*

Se $V = W$, allora $\text{Hom}(V, V)$ ha anche un'altra operazione, quella di composizione fra funzioni indicata con $\circ$. Si verifica facilmente che $\text{Hom}(V, V)$ è un anello con le operazioni $+$ e $\circ$, in cui l'elemento neutro per la somma è la funzione nulla e l'identità per la composizione è la funzione identità id_V .

La mappa $T: \text{Hom}(V, V) \rightarrow M(n)$ è quindi una funzione fra anelli. Se $\mathcal{B} = \mathcal{C}$, cioè se scegliamo la stessa base in partenza ed in arrivo, per quanto visto nella sezione precedente la mappa T è compatibile anche con le operazioni di prodotto, nel senso che:

$$T(f \circ g) = T(f)T(g), \quad T(\text{id}) = I_n, \quad T(f^{-1}) = T(f)^{-1}.$$

Riassumendo, l'operatore T che associa ad ogni funzione lineare la sua matrice associata è compatibile con gran parte delle operazioni algebriche viste finora.

4.3.5. Matrice di cambiamento di base. Abbiamo visto che la matrice associata ad un'applicazione lineare dipende in maniera essenziale dalle basi scelte. Vogliamo capire meglio come controllare questo fenomeno, e a tale scopo introduciamo uno strumento noto come *matrice di cambiamento di base*.

Sia V uno spazio vettoriale e $\mathcal{B} = \{v_1, \dots, v_n\}$ e $\mathcal{C} = \{w_1, \dots, w_n\}$ due basi di V . La *matrice di cambiamento di base* da $\mathcal{B}$ a $\mathcal{C}$ è la matrice

$$A = [\text{id}_V]_{\mathcal{C}}^{\mathcal{B}}.$$

La matrice A rappresenta la funzione identità nelle basi $\mathcal{B}$ e $\mathcal{C}$. Quindi

$$v_j = a_{1j}w_1 + \dots + a_{nj}w_n \quad \forall j = 1, \dots, n.$$

La colonna j -esima di A contiene le coordinate di v_j rispetto a $\mathcal{C}$, cioè

$$A^j = [v_j]_{\mathcal{C}}.$$

La matrice di cambiamento di base A ha nelle sue colonne le coordinate di ciascun elemento di $\mathcal{B}$ rispetto a $\mathcal{C}$. Chiaramente l'inversa

$$A^{-1} = [\text{id}_V]_{\mathcal{B}}^{\mathcal{C}}$$

è la matrice di cambiamento di base da $\mathcal{C}$ a $\mathcal{B}$, e ha nelle sue colonne le coordinate di ciascun elemento di $\mathcal{C}$ rispetto a $\mathcal{B}$.

La matrice di cambiamento di base è lo strumento che funge da dizionario tra una base e l'altra. Dalla Proposizione 4.3.4 ricaviamo infatti:

Proposizione 4.3.13. *Per ogni $v \in V$ vale*

$$[v]_{\mathcal{C}} = A[v]_{\mathcal{B}}.$$

Possiamo usare la matrice di cambiamento di base A per passare agevolmente dalle coordinate dei vettori in $\mathcal{B}$ a quelle in $\mathcal{C}$. Vediamo adesso come usare le matrici di cambiamento di base per modificare le coordinate delle applicazioni lineari.

Sia $f: V \rightarrow W$ una applicazione lineare. Siano $\mathcal{B}_1, \mathcal{B}_2$ due basi di V e $\mathcal{C}_1, \mathcal{C}_2$ due basi di W . Applicando la Proposizione 4.3.7 troviamo:

$$[f]_{\mathcal{C}_2}^{\mathcal{B}_2} = [\text{id}_W]_{\mathcal{C}_2}^{\mathcal{C}_1} \cdot [f]_{\mathcal{C}_1}^{\mathcal{B}_1} \cdot [\text{id}_V]_{\mathcal{B}_1}^{\mathcal{B}_2}.$$

Per passare da $[f]_{\mathcal{C}_1}^{\mathcal{B}_1}$ a $[f]_{\mathcal{C}_2}^{\mathcal{B}_2}$ è sufficiente moltiplicare a sinistra e a destra per le matrici di cambiamento di base.

Esempio 4.3.14. Nell'Esempio 4.3.2 abbiamo analizzato un'applicazione lineare $f: \mathbb{R}_2[x] \rightarrow \mathbb{R}^2$ e calcolato la matrice associata rispetto alle basi $\mathcal{C}' = \{1, x, x^2\}$ e $\mathcal{C} = \{e_1, e_2\}$. Il risultato era il seguente:

$$[f]_{\mathcal{C}}^{\mathcal{C}'} = \begin{pmatrix} 1 & 2 & 4 \\ 1 & -2 & 4 \end{pmatrix}.$$

Successivamente abbiamo cambiato $\mathcal{C}$ con $\mathcal{B} = \left\{ \begin{pmatrix} 1 \\ -1 \end{pmatrix}, \begin{pmatrix} 0 \\ 1 \end{pmatrix} \right\}$. La matrice di cambiamento di base da $\mathcal{B}$ a $\mathcal{C}$ è molto semplice, perché le coordinate di un vettore rispetto a $\mathcal{C}$ sono il vettore stesso e quindi basta affiancare i vettori di $\mathcal{B}$; otteniamo

$$[\text{id}]_{\mathcal{C}}^{\mathcal{B}} = \begin{pmatrix} 1 & 0 \\ -1 & 1 \end{pmatrix}.$$

Adesso calcoliamo l'inversa e troviamo

$$[\text{id}]_{\mathcal{B}}^{\mathcal{C}} = \begin{pmatrix} 1 & 0 \\ 1 & 1 \end{pmatrix}.$$

Quindi la matrice associata a f nelle basi $\mathcal{C}'$ e $\mathcal{B}$ è

$$[f]_{\mathcal{B}}^{\mathcal{C}'} = [\text{id}]_{\mathcal{B}}^{\mathcal{C}} [f]_{\mathcal{C}}^{\mathcal{C}'} = \begin{pmatrix} 1 & 0 \\ 1 & 1 \end{pmatrix} \begin{pmatrix} 1 & 2 & 4 \\ 1 & -2 & 4 \end{pmatrix} = \begin{pmatrix} 1 & 2 & 4 \\ 2 & 0 & 8 \end{pmatrix}$$

che coincide effettivamente con quanto trovato nell'Esempio 4.3.2.

Esempio 4.3.15. Nell'Esempio 4.3.3 abbiamo considerato un'applicazione lineare $f: \mathbb{C}^2 \rightarrow \mathbb{C}^3$ e calcolato la matrice associata rispetto alle basi canoniche $\mathcal{C}$ e $\mathcal{C}'$ di $\mathbb{R}^2$ e $\mathbb{R}^3$:

$$[f]_{\mathcal{C}'}^{\mathcal{C}} = \begin{pmatrix} 1 & -1 \\ 2 & 0 \\ 0 & 1 \end{pmatrix}.$$

Successivamente abbiamo cambiato entrambe le basi con

$$\mathcal{B} = \left\{ \begin{pmatrix} 1 \\ 1 \end{pmatrix}, \begin{pmatrix} 1 \\ -1 \end{pmatrix} \right\}, \quad \mathcal{B}' = \left\{ \begin{pmatrix} 1 \\ 1 \\ 0 \end{pmatrix}, \begin{pmatrix} 0 \\ 1 \\ 1 \end{pmatrix}, \begin{pmatrix} 1 \\ 0 \\ 1 \end{pmatrix} \right\}.$$

Calcoliamo le matrici di cambiamento di base:

$$[\text{id}_{\mathbb{C}^2}]_{\mathcal{C}}^{\mathcal{B}} = \begin{pmatrix} 1 & 1 \\ 1 & -1 \end{pmatrix}, \quad [\text{id}_{\mathbb{C}^3}]_{\mathcal{C}'}^{\mathcal{B}'} = \begin{pmatrix} 1 & 0 & 1 \\ 1 & 1 & 0 \\ 0 & 1 & 1 \end{pmatrix}.$$

L'inversa della seconda è

$$[\text{id}_{\mathbb{C}^3}]_{\mathcal{B}'}^{\mathcal{C}'} = \frac{1}{2} \begin{pmatrix} 1 & 1 & -1 \\ -1 & 1 & 1 \\ 1 & -1 & 1 \end{pmatrix}.$$

Troviamo infine

$$\begin{aligned} [f]_{\mathcal{B}'}^{\mathcal{B}} &= [\text{id}_{\mathbb{C}^3}]_{\mathcal{B}'}^{\mathcal{C}'} \cdot [f]_{\mathcal{C}'}^{\mathcal{C}} \cdot [\text{id}_{\mathbb{C}^2}]_{\mathcal{C}}^{\mathcal{B}} \\ &= \frac{1}{2} \begin{pmatrix} 1 & 1 & -1 \\ -1 & 1 & 1 \\ 1 & -1 & 1 \end{pmatrix} \begin{pmatrix} 1 & -1 \\ 2 & 0 \\ 0 & 1 \end{pmatrix} \begin{pmatrix} 1 & 1 \\ 1 & -1 \end{pmatrix} = \frac{1}{2} \begin{pmatrix} 1 & 5 \\ 3 & -1 \\ -1 & -1 \end{pmatrix} \end{aligned}$$

coerentemente con quanto già visto nell'Esempio 4.3.3.

4.4. Endomorfismi

Trattiamo adesso più approfonditamente le applicazioni lineari $V \rightarrow V$ in cui dominio e codominio coincidono. Queste applicazioni vengono chiamate *endomorfismi*.

4.4.1. Definizione. Sia V uno spazio vettoriale.

Definizione 4.4.1. Un *endomorfismo* di V è un'applicazione lineare $f: V \rightarrow V$.

Abbiamo già visto molti esempi di endomorfismi nelle pagine precedenti. Ne ricordiamo alcuni molto brevemente. In ciascuno spazio vettoriale V , la moltiplicazione per un $\lambda \in \mathbb{K}$ fissato $f(v) = \lambda v$ è un endomorfismo. Gli endomorfismi di $\mathbb{K}^n$ sono tutti del tipo L_A dove A è una matrice quadrata $n \times n$. La derivata è un endomorfismo di $\mathbb{R}_n[x]$ e la trasposta è un endomorfismo di $M(n, \mathbb{K})$. Moltiplicando a sinistra per una matrice fissata $A \in M(n, \mathbb{K})$ si ottiene un altro endomorfismo di $M(n, \mathbb{K})$. Se $V = U \oplus W$, allora possiamo definire due endomorfismi p_U e r_U di V , la proiezione su U e la riflessione lungo U . Altri endomorfismi più geometrici di $\mathbb{R}^2$ e $\mathbb{R}^3$ verranno descritti fra poco.

4.4.2. Matrice associata. Quando rappresentiamo un endomorfismo $f: V \rightarrow V$ come matrice è ragionevole prendere la stessa base $\mathcal{B}$ per V in partenza e in arrivo. Otteniamo quindi una matrice $A = [f]_{\mathcal{B}}^{\mathcal{B}}$. Questa rappresentazione ha la buona proprietà di trasformare la composizione fra endomorfismi nel prodotto fra matrici:

$$[f \circ g]_{\mathcal{B}}^{\mathcal{B}} = [f]_{\mathcal{B}}^{\mathcal{B}} \cdot [g]_{\mathcal{B}}^{\mathcal{B}}.$$

Abbiamo anche visto che se $\mathcal{C}$ è un'altra base per V e $M = [\text{id}_V]_{\mathcal{C}}^{\mathcal{B}}$ è la matrice di cambiamento di base da $\mathcal{B}$ a $\mathcal{C}$, otteniamo

$$[f]_{\mathcal{B}}^{\mathcal{B}} = M^{-1} [f]_{\mathcal{C}}^{\mathcal{C}} M.$$

Ricapitolando, se A e B sono le matrici associate a f rispetto alle basi $\mathcal{B}$ e $\mathcal{C}$ e M è la matrice di cambiamento di base, allora

$$(5) \quad A = M^{-1}BM.$$

Ricordiamo ancora una volta che il prodotto fra matrici non è commutativo: se lo fosse, otterremmo sempre $A = M^{-1}BM = M^{-1}MB = B$. Invece le matrici A e B sono spesso ben differenti fra loro.

Esempio 4.4.2. Consideriamo la riflessione $r_U: \mathbb{R}^2 \rightarrow \mathbb{R}^2$ dell'Esempio 4.1.12. Avevamo espresso la funzione esplicitamente nel modo seguente:

$$r_U \begin{pmatrix} x \\ y \end{pmatrix} = \begin{pmatrix} x+y \\ -y \end{pmatrix}.$$

Rispetto alla base canonica $\mathcal{C} = \{e_1, e_2\}$, troviamo

$$[r_U]_{\mathcal{C}}^{\mathcal{C}} = \begin{pmatrix} 1 & 1 \\ 0 & -1 \end{pmatrix}.$$

Prendiamo adesso come base $\mathcal{B} = \left\{ \begin{pmatrix} 1 \\ 0 \end{pmatrix}, \begin{pmatrix} -1 \\ 2 \end{pmatrix} \right\}$, due vettori che generano i sottospazi U e W . La matrice di cambiamento di base da $\mathcal{B}$ a $\mathcal{C}$ è semplicemente

$$M = [\text{id}]_{\mathcal{C}}^{\mathcal{B}} = \begin{pmatrix} 1 & -1 \\ 0 & 2 \end{pmatrix}$$

e la sua inversa è

$$M^{-1} = [\text{id}]_{\mathcal{B}}^{\mathcal{C}} = \begin{pmatrix} 1 & \frac{1}{2} \\ 0 & \frac{1}{2} \end{pmatrix}.$$

Quindi la matrice associata a r_U nella base $\mathcal{B}$ è

$$[r_U]_{\mathcal{B}}^{\mathcal{B}} = M^{-1}[r_U]_{\mathcal{C}}^{\mathcal{C}}M = \begin{pmatrix} 1 & \frac{1}{2} \\ 0 & \frac{1}{2} \end{pmatrix} \begin{pmatrix} 1 & 1 \\ 0 & -1 \end{pmatrix} \begin{pmatrix} 1 & -1 \\ 0 & 2 \end{pmatrix} = \begin{pmatrix} 1 & 0 \\ 0 & -1 \end{pmatrix}.$$

Possiamo verificare che il risultato sia giusto. Le colonne di $\begin{pmatrix} 1 & 0 \\ 0 & -1 \end{pmatrix}$ sono le coordinate delle immagini dei vettori della base $\mathcal{B}$ rispetto alla stessa base $\mathcal{B}$. La prima colonna sta ad indicare che r_U manda il primo vettore di $\mathcal{B}$ in se stesso e la seconda indica che r_U manda il secondo vettore di $\mathcal{B}$ nel suo opposto. Quindi $r_U \begin{pmatrix} 1 \\ 0 \end{pmatrix} = \begin{pmatrix} 1 \\ 0 \end{pmatrix}$ e $r_U \begin{pmatrix} -1 \\ 2 \end{pmatrix} = \begin{pmatrix} 1 \\ -2 \end{pmatrix}$. Questo è effettivamente quello che fa la riflessione, che fissa i punti di U e ribalta quelli di W .

Esercizio 4.4.3. Calcola $[p_U]_{\mathcal{B}}^{\mathcal{B}}$ dove $p_U: \mathbb{R}^2 \rightarrow \mathbb{R}^2$ è la proiezione dell'Esempio 4.1.11 e $\mathcal{B}$ è la stessa base usata nell'esempio precedente.

Esempio 4.4.4. Sia V uno spazio vettoriale di dimensione n e sia $\lambda \in \mathbb{K}$. La matrice associata all'endomorfismo $f(v) = \lambda v$ è

$$[f]_{\mathcal{B}}^{\mathcal{B}} = \lambda I_n$$

rispetto a *qualsiasi* base $\mathcal{B} = \{v_1, \dots, v_n\}$, perché $f(v_i) = \lambda v_i$ per ogni i e quindi $[f(v_i)]_{\mathcal{B}} = \lambda e_i$. Questo è coerente con il fatto che

$$M^{-1}(\lambda I_n)M = \lambda M^{-1}I_n M = \lambda M^{-1}M = \lambda I_n$$

per qualsiasi matrice invertibile M . Cambiando base la matrice associata rimane sempre λI_n .

La relazione $A = M^{-1}BM$ scoperta in (5) ci porta a studiare più in generale questa operazione fra matrici quadrate.

4.4.3. Similitudine fra matrici. Lavoriamo con l'anello $M(n)$ delle matrici quadrate $n \times n$. Diciamo che due matrici $A, B \in M(n)$ sono *simili* o *coniugate* se esiste una matrice invertibile $M \in M(n)$ tale che

$$A = M^{-1}BM.$$

Se A e B sono simili scriviamo $A \sim B$. Rimandiamo alla Sezione 1.1.10 per la nozione di relazione di equivalenza.

Proposizione 4.4.5. *La similitudine è una relazione di equivalenza.*

Dimostrazione. Dobbiamo dimostrare tre proprietà:

- (1) $A \sim A$. Infatti prendendo $M = I_n$ otteniamo $A = I_n^{-1}AI_n$.
- (2) $A \sim B \Rightarrow B \sim A$. Infatti se $A = M^{-1}BM$, allora moltiplicando entrambi i membri a sinistra per M e a destra per M^{-1} troviamo $B = MAM^{-1}$, e definendo $N = M^{-1}$ otteniamo $B = N^{-1}AN$.
- (3) $A \sim B, B \sim C, \Rightarrow A \sim C$. Infatti

$$A = M^{-1}BM, B = N^{-1}CN \Rightarrow A = M^{-1}N^{-1}CNM = (NM)^{-1}C(NM).$$

La dimostrazione è completa. $\square$

L'insieme $M(n)$ delle matrici quadrate è quindi partizionato in sottoinsiemi disgiunti formati da matrici simili fra loro.

Esempio 4.4.6. Per ogni $\lambda \in \mathbb{K}$, la matrice λI_n è simile solo a se stessa: come abbiamo già notato, $M^{-1}(\lambda I_n)M = \lambda I_n$ per qualsiasi matrice invertibile M . In particolare le matrici nulla 0 e identità I_n sono simili solo a loro stesse.

Esempio 4.4.7. Le matrici $\begin{pmatrix} 1 & k \\ 0 & 1 \end{pmatrix}$ per $k \neq 0$ sono tutte simili fra loro. Per dimostrare questo fatto usiamo la proprietà transitiva e ci limitiamo a mostrare che ciascuna $\begin{pmatrix} 1 & k \\ 0 & 1 \end{pmatrix}$ con $k \neq 0$ è simile a $\begin{pmatrix} 1 & 1 \\ 0 & 1 \end{pmatrix}$. Infatti se $M = \begin{pmatrix} 1 & 0 \\ 0 & k \end{pmatrix}$ otteniamo

$$M^{-1} \begin{pmatrix} 1 & 1 \\ 0 & 1 \end{pmatrix} M = \begin{pmatrix} 1 & 0 \\ 0 & k^{-1} \end{pmatrix} \begin{pmatrix} 1 & 1 \\ 0 & 1 \end{pmatrix} \begin{pmatrix} 1 & 0 \\ 0 & k \end{pmatrix} = \begin{pmatrix} 1 & k \\ 0 & 1 \end{pmatrix}.$$

Le matrici $\begin{pmatrix} 1 & k \\ 0 & 1 \end{pmatrix}$ con $k \neq 0$ non sono però simili alla matrice $\begin{pmatrix} 1 & 0 \\ 0 & 1 \end{pmatrix}$, perché sappiamo che la matrice identità è simile solo a se stessa.

Dimostreremo nelle prossime pagine che due matrici simili hanno molto in comune. Iniziamo con il rango e il determinante.

Proposizione 4.4.8. *Se $A \sim B$ allora $\text{rk}(A) = \text{rk}(B)$ e $\det(A) = \det(B)$. In particolare A è invertibile $\iff B$ è invertibile.*

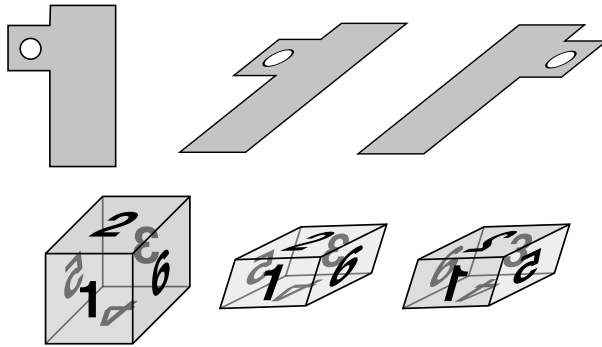


Figura 4.3. L'effetto su di un oggetto in $\mathbb{R}^2$ e in $\mathbb{R}^3$ (sinistra) di una trasformazione con determinante positivo (centro) e negativo (destra).

Dimostrazione. Usiamo la Proposizione 4.2.29 e il Teorema di Binet:

$$\det(M^{-1}BM) = \det(M^{-1}) \det B \det M = \frac{1}{\det M} \det B \det M = \det B.$$

La dimostrazione è completa. $\square$

Quindi ad esempio le matrici $\begin{pmatrix} 1 & 2 \\ 1 & 1 \end{pmatrix}$ e $\begin{pmatrix} -1 & 2 \\ 1 & 1 \end{pmatrix}$ non sono simili perché hanno determinanti diversi. Il motivo per cui siamo interessati alla relazione di similitudine è il seguente, dimostrato nelle pagine precedenti.

Proposizione 4.4.9. *Due matrici $[f]_{\mathcal{B}}^{\mathcal{B}}$ e $[f]_{\mathcal{C}}^{\mathcal{C}}$ associate allo stesso endomorfismo $f: V \rightarrow V$ ma con basi $\mathcal{B}, \mathcal{C}$ diverse sono simili.*

Come corollario abbiamo questo fatto cruciale:

Qualsiasi oggetto associato ad una matrice quadrata A che non cambi per similitudine è ben definito sugli endomorfismi.

Chiariamo questo passaggio con un esempio importante.

4.4.4. Determinante di un endomorfismo. Definiamo il *determinante* di un endomorfismo $f: V \rightarrow V$ come il determinante della matrice associata

$$\det f = \det[f]_{\mathcal{B}}^{\mathcal{B}}$$

rispetto a *qualsiasi* base $\mathcal{B}$. Per quanto abbiamo appena visto infatti il determinante è invariante per similitudine: quindi il determinante della matrice $[f]_{\mathcal{B}}^{\mathcal{B}}$ non dipende da $\mathcal{B}$ e quindi è ben definito.

Esempio 4.4.10. L'endomorfismo $f(v) = \lambda v$ in uno spazio V di dimensione n ha determinante λ^n , perché $\det(\lambda I_n) = \lambda^n$.

Sappiamo che $f: V \rightarrow V$ è un isomorfismo $\iff \det f \neq 0$. Nel caso in cui $V = \mathbb{R}^2$ oppure $\mathbb{R}^3$, il determinante di f ha una chiara interpretazione geometrica, a cui abbiamo già accennato nella Sezione 3.3.10 e che studieremo fra poco.

4.4.5. Traccia di una matrice quadrata. Introduciamo adesso un'altra quantità che non cambia per similitudine, la *traccia*.

La *traccia* di una matrice quadrata $A \in M(n)$ è il numero

$$\operatorname{tr} A = A_{11} + \cdots + A_{nn}.$$

La traccia di A è la somma dei valori sulla diagonale principale di A .

Proposizione 4.4.11. Se $A, B \in M(n)$, allora $\operatorname{tr}(AB) = \operatorname{tr}(BA)$.

Dimostrazione. Troviamo

$$\begin{aligned} \operatorname{tr}(AB) &= \sum_{i=1}^n (AB)_{ii} = \sum_{i=1}^n \sum_{j=1}^n A_{ij} B_{ji} \\ \operatorname{tr}(BA) &= \sum_{j=1}^n (BA)_{jj} = \sum_{j=1}^n \sum_{i=1}^n B_{ji} A_{ij} \end{aligned}$$

Il risultato è lo stesso. □

Corollario 4.4.12. Due matrici simili hanno la stessa traccia.

Dimostrazione. Se $A = M^{-1}BM$, troviamo

$$\operatorname{tr} A = \operatorname{tr}(M^{-1}BM) = \operatorname{tr}((M^{-1}B)M) = \operatorname{tr}(M(M^{-1}B)) = \operatorname{tr} B.$$

La dimostrazione è completa. □

Quindi due matrici con tracce diverse, come ad esempio $\begin{pmatrix} 1 & 1 \\ 0 & 1 \end{pmatrix}$ e $\begin{pmatrix} -1 & 1 \\ 0 & 1 \end{pmatrix}$, non sono simili. Notiamo però che esistono matrici con la stessa traccia e lo stesso determinante che non sono simili: ad esempio $\begin{pmatrix} 1 & 1 \\ 0 & 1 \end{pmatrix}$ e $\begin{pmatrix} 1 & 0 \\ 0 & 1 \end{pmatrix}$.

Se $f: V \rightarrow V$ è un endomorfismo e V ha dimensione finita, possiamo definire la *traccia* di f come la traccia della matrice associata $[f]_{\mathcal{B}}^{\mathcal{B}}$ rispetto a qualche base $\mathcal{B}$. Per il corollario precedente la traccia di f effettivamente non dipende dalla base $\mathcal{B}$ scelta.

Esempio 4.4.13. L'endomorfismo $f(v) = \lambda v$ in uno spazio V di dimensione n ha traccia $n\lambda$, perché $\operatorname{tr}(\lambda I_n) = n\lambda$. In particolare l'identità ha traccia n .

4.4.6. Trasformazioni lineari del piano e dello spazio. Un endomorfismo f invertibile di $\mathbb{R}^2$ o $\mathbb{R}^3$ è detto *trasformazione lineare del piano o dello spazio*. Le trasformazioni hanno un ruolo importante nella geometria moderna.

Il determinante di una trasformazione lineare f di $\mathbb{R}^2$ o $\mathbb{R}^3$ ha una chiara interpretazione geometrica. Se $\det f < 0$, la trasformazione f manda basi positive in basi negative (si veda la Sezione 3.3.10) e quindi “specchia”

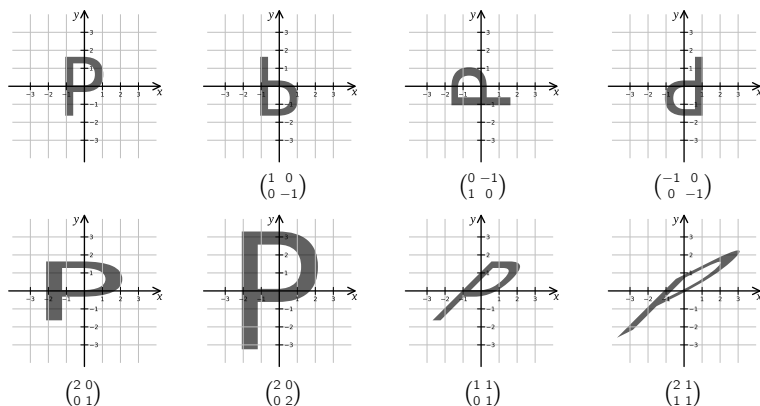


Figura 4.4. Sette trasformazioni del piano. Ciascuna trasformazione è un isomorfismo di $\mathbb{R}^2$ determinato dalla matrice 2×2 invertibile indicata in basso; viene mostrato il suo effetto sulla lettera P, originariamente nella posizione in alto a sinistra.

gli oggetti contenuti in $\mathbb{R}^2$ o $\mathbb{R}^3$, trasformando ad esempio mani destre in mani sinistre. La Figura 4.3 mostra l'effetto su di un oggetto di una trasformazione con determinante positivo e negativo.

D'altro canto (si veda sempre la Sezione 3.3.10) il valore assoluto $|\det f|$ del determinante indica come vengono trasformate le aree (in $\mathbb{R}^2$) o i volumi (in $\mathbb{R}^3$) degli oggetti. Ciascun oggetto $C \subset \mathbb{R}^2$ o $\mathbb{R}^3$ viene trasformato in un oggetto $f(C)$ la cui area o volume è esattamente $|\det f|$ volte quella originaria di C . In particolare se $\det f = \pm 1$ allora la trasformazione f non cambia l'area o il volume degli oggetti, ma può distorcerli lo stesso come vedremo fra poco.

4.4.7. Trasformazioni del piano. Alcuni esempi di trasformazioni lineari del piano sono visualizzati nella Figura 4.4. Ciascun l'isomorfismo è del tipo $L_A: \mathbb{R}^2 \rightarrow \mathbb{R}^2$ e la matrice A corrispondente è indicata sotto la figura. La lettrice è invitata ad analizzare attentamente la matrice di ciascuna trasformazione e a convincersi geometricamente che l'effetto sulla lettera P sia effettivamente quello descritto in figura. Ricordiamo che le colonne della matrice A sono le immagini dei vettori e_1 e e_2 . Alcune di queste matrici verranno studiate approfonditamente nelle prossime pagine.

Notiamo che i primi tre esempi non distorcono la figura, e sono rispettivamente una riflessione lungo l'asse x , una rotazione antioraria di $\frac{\pi}{2}$, e una riflessione rispetto all'origine, che è la stessa cosa che una rotazione di angolo π . Analizzeremo questo tipo di trasformazioni fra poco. I quattro esempi successivi invece distorcono la figura: il primo la allarga orizzontalmente di un fattore 2, il secondo è una omotetia sempre di fattore 2,

e le due trasformazioni successive sembrano distorcere la figura in modo maggiore.

I determinanti delle sette trasformazioni mostrate sono $-1, 1, 1, 2, 4, 1$ e 1 . Solo nella prima trasformazione (la riflessione) la lettera P viene specchiata. Inoltre l'area della P è preservata in tutte le trasformazioni eccetto la quarta e la quinta in cui viene moltiplicata per 2 e per 4 rispettivamente. Nelle ultime due trasformazioni c'è un forte effetto di distorsione ma l'area della P rimane la stessa!

4.4.8. Rotazioni nel piano. Introduciamo una classe importante di trasformazioni di $\mathbb{R}^2$, le *rotazioni*.

Definizione 4.4.14. Una *rotazione* di angolo θ è la trasformazione $L_A: \mathbb{R}^2 \rightarrow \mathbb{R}^2$ determinata dalla matrice $A = \text{Rot}_\theta$, con

$$\text{Rot}_\theta = \begin{pmatrix} \cos \theta & -\sin \theta \\ \sin \theta & \cos \theta \end{pmatrix}.$$

Proposizione 4.4.15. La trasformazione L_A è effettivamente una *rotazione antioraria del piano di angolo θ intorno all'origine*.

Dimostrazione. Rappresentiamo un punto $\begin{pmatrix} x \\ y \end{pmatrix}$ in coordinate polari $x = \rho \cos \varphi$ e $y = \rho \sin \varphi$. Calcoliamo la sua immagine:

$$\begin{aligned} L_A \begin{pmatrix} x \\ y \end{pmatrix} &= \begin{pmatrix} \cos \theta & -\sin \theta \\ \sin \theta & \cos \theta \end{pmatrix} \begin{pmatrix} \rho \cos \varphi \\ \rho \sin \varphi \end{pmatrix} = \begin{pmatrix} \rho(\cos \theta \cos \varphi - \sin \theta \sin \varphi) \\ \rho(\sin \theta \cos \varphi + \cos \theta \sin \varphi) \end{pmatrix} \\ &= \begin{pmatrix} \rho \cos(\theta + \varphi) \\ \rho \sin(\theta + \varphi) \end{pmatrix}. \end{aligned}$$

La trasformazione manda il punto generico con coordinate polari (ρ, φ) nel punto con coordinate polari $(\rho, \theta + \varphi)$. Tutti i punti vengono ruotati in senso antiorario dell'angolo θ . $\square$

Notiamo che la matrice Rot_θ ha determinante $\cos^2 \theta + \sin^2 \theta = 1$. Le rotazioni non specchiano gli oggetti e mantengono le aree.

4.4.9. Riflessioni ortogonali nel piano. Introduciamo un'altra classe di trasformazioni del piano, le *riflessioni ortogonali*. Fissiamo un angolo θ e consideriamo la retta vettoriale r che forma un angolo $\frac{\theta}{2}$ con l'asse x .

Definizione 4.4.16. Una *riflessione ortogonale* rispetto alla retta r è la trasformazione $L_A: \mathbb{R}^2 \rightarrow \mathbb{R}^2$ determinata dalla matrice $A = \text{Rif}_\theta$, con

$$\text{Rif}_\theta = \begin{pmatrix} \cos \theta & \sin \theta \\ \sin \theta & -\cos \theta \end{pmatrix}.$$

Proposizione 4.4.17. La trasformazione L_A è effettivamente una *riflessione ortogonale del piano rispetto a r* .

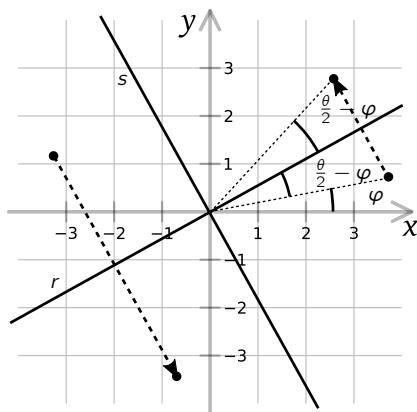


Figura 4.5. Una riflessione rispetto alla retta r che forma un angolo $\frac{\theta}{2}$ con l'asse x . Il punto (ρ, φ) viene mandato in $(\rho, \theta - \varphi)$.

Dimostrazione. Rappresentiamo un punto $\begin{pmatrix} x \\ y \end{pmatrix}$ in coordinate polari $x = \rho \cos \varphi$ e $y = \rho \sin \varphi$. Calcoliamo la sua immagine:

$$\begin{aligned} L_A \begin{pmatrix} x \\ y \end{pmatrix} &= \begin{pmatrix} \cos \theta & \sin \theta \\ \sin \theta & -\cos \theta \end{pmatrix} \begin{pmatrix} \rho \cos \varphi \\ \rho \sin \varphi \end{pmatrix} = \begin{pmatrix} \rho(\cos \theta \cos \varphi + \sin \theta \sin \varphi) \\ \rho(\sin \theta \cos \varphi - \cos \theta \sin \varphi) \end{pmatrix} \\ &= \begin{pmatrix} \rho \cos(\theta - \varphi) \\ \rho \sin(\theta - \varphi) \end{pmatrix}. \end{aligned}$$

La trasformazione manda il punto generico con coordinate polari (ρ, φ) nel punto con coordinate polari $(\rho, \theta - \varphi)$. Questa è una riflessione ortogonale rispetto a r , si veda la Figura 4.5. $\square$

Notiamo che la matrice Rif_θ ha determinante $-\cos^2 \theta - \sin^2 \theta = -1$. Effettivamente una riflessione specchia gli oggetti ma mantiene le aree.

Osservazione 4.4.18. Sia s la retta ortogonale a r , cioè facente un angolo $\frac{\theta}{2} + \frac{\pi}{2}$ con l'asse x , come nella Figura 4.5. La riflessione ortogonale descritta qui corrisponde alla riflessione determinata dalla decomposizione in somma diretta $\mathbb{R}^2 = r \oplus s$ definita nell'Esempio 4.1.12. Come già osservato nell'Esempio 4.4.2 la riflessione può essere rappresentata più agevolmente rispetto ad un'altra base $\mathcal{B} = \{v_1, v_2\}$ in cui $v_1 \in r$ e $v_2 \in s$. Poiché $f(v_1) = v_1$ e $f(v_2) = -v_2$, la matrice associata alla riflessione rispetto alla base $\mathcal{B}$ è semplicemente

$$\begin{pmatrix} 1 & 0 \\ 0 & -1 \end{pmatrix}.$$

4.4.10. Omotetie. Una *omotetia* in $\mathbb{R}^2$ oppure $\mathbb{R}^3$ di centro l'origine è semplicemente la trasformazione

$$f(x) = \lambda x$$

dove $\lambda > 0$ è il *coefficiente* dell'omotetia. Abbiamo $f = L_A$ dove $A = \lambda I_n$, con $n = 2$ oppure 3 , cioè

$$A = \begin{pmatrix} \lambda & 0 \\ 0 & \lambda \end{pmatrix} \quad \text{oppure} \quad A = \begin{pmatrix} \lambda & 0 & 0 \\ 0 & \lambda & 0 \\ 0 & 0 & \lambda \end{pmatrix}.$$

L'omotetia dilata gli oggetti di un fattore λ . Aree e volumi vengono modificati di un fattore $\det(A) = \lambda^2$ oppure λ^3 .

4.4.11. Rotazioni intorno ad un asse nello spazio. Fra le trasformazioni dello spazio, spiccano le *rotazioni intorno ad un asse*. Studieremo queste trasformazioni nel dettaglio in seguito: per adesso ci limitiamo a definire la rotazione di un angolo θ intorno all'asse z come la trasformazione $L_A: \mathbb{R}^3 \rightarrow \mathbb{R}^3$ in cui la matrice A è la seguente:

$$A = \begin{pmatrix} \cos \theta & -\sin \theta & 0 \\ \sin \theta & \cos \theta & 0 \\ 0 & 0 & 1 \end{pmatrix} = \begin{pmatrix} \text{Rot}_\theta & 0 \\ 0 & 1 \end{pmatrix}.$$

La seconda scrittura più compatta descrive la A come matrice a blocchi: Rot_θ è una sottomatrice 2×2 e i due "0" sono due sottomatrici 2×1 e 1×2 nulle.

Proposizione 4.4.19. *La trasformazione L_A è effettivamente una rotazione antioraria dello spazio di angolo θ intorno all'asse z .*

Dimostrazione. Rappresentiamo un punto di $\mathbb{R}^3$ in *coordinate cilindriche*: queste sono semplicemente le coordinate polari per x e y con la z lasciata immutata. Passiamo cioè da (x, y, z) a (ρ, φ, z) con $x = \rho \cos \varphi$ e $y = \rho \sin \varphi$. Calcoliamo l'immagine del punto generico con coordinate cilindriche (ρ, φ, z) :

$$\begin{aligned} L_A \begin{pmatrix} x \\ y \\ z \end{pmatrix} &= \begin{pmatrix} \cos \theta & -\sin \theta & 0 \\ \sin \theta & \cos \theta & 0 \\ 0 & 0 & 1 \end{pmatrix} \begin{pmatrix} \rho \cos \varphi \\ \rho \sin \varphi \\ z \end{pmatrix} \\ &= \begin{pmatrix} \rho(\cos \theta \cos \varphi - \sin \theta \sin \varphi) \\ \rho(\sin \theta \cos \varphi + \cos \theta \sin \varphi) \\ z \end{pmatrix} = \begin{pmatrix} \rho \cos(\theta + \varphi) \\ \rho \sin(\theta + \varphi) \\ z \end{pmatrix}. \end{aligned}$$

La trasformazione manda il punto con coordinate cilindriche (ρ, φ, z) nel punto con coordinate cilindriche $(\rho, \theta + \varphi, z)$. Quindi tutti i punti vengono ruotati in senso antiorario dell'angolo θ intorno all'asse z . $\square$

Vedremo nei prossimi capitoli come scrivere una rotazione di angolo θ intorno ad un asse arbitrario.

Esercizi

Esercizio 4.1. Sia $T: \mathbb{R}_2[x] \rightarrow \mathbb{R}_3[x]$ la funzione $T(p) = xp(x)$. Mostra che T è una funzione lineare e determina il nucleo e l'immagine di T .

Esercizio 4.2. Considera l'applicazione lineare $f: \mathbb{R}^3 \rightarrow \mathbb{R}^3$ data da

$$f \begin{pmatrix} x \\ y \\ z \end{pmatrix} = \begin{pmatrix} 2x - y + z \\ x + y + z \\ y - z \end{pmatrix}$$

Scegli una base $\mathcal{B}$ del piano $U = \{x + y - z = 0\}$ e scrivi la matrice associata alla restrizione $f|_U: U \rightarrow \mathbb{R}^3$ rispetto a $\mathcal{B}$ e alla base canonica di $\mathbb{R}^3$.

Esercizio 4.3. Considera l'applicazione lineare $f: \mathbb{R}^2 \rightarrow \mathbb{R}^3$,

$$f \begin{pmatrix} x \\ y \end{pmatrix} = \begin{pmatrix} 2x + y \\ x - y \\ 3x + 2y \end{pmatrix}$$

Determina il nucleo e l'immagine di f . Scrivi la matrice associata a f rispetto alle basi canoniche di $\mathbb{R}^2$ e $\mathbb{R}^3$ e rispetto alle basi seguenti:

$$\left\{ \begin{pmatrix} 1 \\ -1 \end{pmatrix}, \begin{pmatrix} 1 \\ 1 \end{pmatrix} \right\}, \quad \left\{ \begin{pmatrix} 1 \\ 1 \\ 0 \end{pmatrix}, \begin{pmatrix} 0 \\ 1 \\ 1 \end{pmatrix}, \begin{pmatrix} 1 \\ 0 \\ 1 \end{pmatrix} \right\}.$$

Esercizio 4.4. Scrivi la matrice di cambiamento di base $[\text{id}]_{\mathcal{B}}^{\mathcal{B}'}$ fra le due basi seguenti di $\mathbb{R}_2[x]$:

$$\mathcal{B} = \{1 + x + x^2, x, 1 - x^2\}, \quad \mathcal{B}' = \{1 - x, x - x^2, 1\}.$$

Esercizio 4.5. Sia $f: V \rightarrow W$ un'applicazione lineare. Mostra che:

- Se $U \subset V$ è un sottospazio, allora $f(U)$ è un sottospazio di W .
- Se $Z \subset W$ è un sottospazio, allora $f^{-1}(Z)$ è un sottospazio di V .
- Se f è un isomorfismo, allora $\dim f(U) = \dim U$ per ogni sottospazio $U \subset V$.

Esercizio 4.6. Sia $T: \mathbb{R}_n[x] \rightarrow \mathbb{R}_n[x]$ la funzione $T(p) = (x + 1)p'(x)$ dove $p'(x)$ è la derivata di $p(x)$. Scrivi T rispetto alla base canonica, determina nucleo e immagine di T , calcola la sua traccia e il suo determinante in funzione di n .

Esercizio 4.7. Sia $A = \begin{pmatrix} 1 & 2 \\ 3 & 4 \end{pmatrix}$. Sia $L_A: M(2, \mathbb{R}) \rightarrow M(2, \mathbb{R})$ l'applicazione lineare $L_A(X) = AX$.

- (1) Determina nucleo e immagine di L_A .
- (2) Scrivi la matrice associata a L_A rispetto alla base canonica di $M(2, \mathbb{R})$.
- (3) Calcola il determinante di L_A .
- (4) Mostra che per una $A \in M(2)$ generica, l'applicazione $L_A: M(2) \rightarrow M(2)$ è un isomorfismo $\iff A$ è invertibile. Mostra inoltre che $\det L_A = (\det A)^2$.

Esercizio 4.8. Sia $A \in M(2)$ fissata e $T: M(2) \rightarrow M(2)$ l'applicazione $T(X) = AX - XA$. Mostra che T è lineare. Inoltre:

- (1) Se $A = \begin{pmatrix} 1 & 1 \\ 0 & 1 \end{pmatrix}$, determina il nucleo e l'immagine di T .
- (2) Dimostra che per qualsiasi A l'applicazione T non è mai un isomorfismo.
- (3) Dimostra che per qualsiasi A l'immagine di T ha dimensione ≤ 2 .

Esercizio 4.9. Determina una matrice $A \in M(3, \mathbb{R})$ tale che l'endomorfismo L_A soddisfi le proprietà seguenti: $\text{Im } L_A = \{x - 2y + z = 0\}$ e $\ker L_A = \text{Span}(e_1 + e_2 + e_3)$.

Esercizio 4.10. Siano $f: \mathbb{C}^3 \rightarrow \mathbb{C}^2$ e $g: \mathbb{C}^2 \rightarrow \mathbb{C}^3$ due applicazioni lineari. La composizione $f \circ g$ può essere un isomorfismo? E la composizione $g \circ f$? In caso affermativo fornisci un esempio, in caso negativo dimostra che non è possibile.

Esercizio 4.11. Siano $U = \{x + y = 0\}$ e $W = \text{Span}(e_2 + e_3)$ sottospazi di $\mathbb{R}^3$. Dimostra che $\mathbb{R}^3 = U \oplus W$ e scrivi la matrice associata alle proiezioni p_U e p_W indotte dalla somma diretta rispetto alla base canonica di $\mathbb{R}^3$.

Esercizio 4.12. Sia $f: V \rightarrow W$ un'applicazione lineare fra due spazi V e W di dimensione finita. Mostra che esistono sempre due basi $\mathcal{B}$ e $\mathcal{C}$ per V e W tale che

$$[f]_{\mathcal{C}}^{\mathcal{B}} = \begin{pmatrix} I_r & 0 \\ 0 & 0 \end{pmatrix}$$

dove $r = \dim \text{Im } f$ e la descrizione è a blocchi.

Esercizio 4.13. Sia $f: V \rightarrow W$ un'applicazione lineare fra due spazi V e W di dimensione finita. Mostra che:

- f è iniettiva $\iff$ esiste una $g: W \rightarrow V$ lineare tale che $g \circ f = \text{id}_V$.
- f è suriettiva $\iff$ esiste una $g: W \rightarrow V$ lineare tale che $f \circ g = \text{id}_W$.

Esercizio 4.14. Siano $x_0, \dots, x_n \in \mathbb{K}$ punti diversi. Considera l'applicazione lineare $T: \mathbb{K}_n[x] \rightarrow \mathbb{K}^{n+1}$,

$$T(p) = \begin{pmatrix} p(x_0) \\ \vdots \\ p(x_n) \end{pmatrix}.$$

Dimostra che T è un isomorfismo. Qual è il legame con l'Esercizio 3.12?

Esercizio 4.15. Sia V uno spazio vettoriale con base $\mathcal{B} = \{v_1, \dots, v_n\}$ e $f: V \rightarrow V$ un endomorfismo. Sia $\mathcal{B}'$ ottenuta da $\mathcal{B}$ scambiando v_i e v_j . Mostra che $[f]_{\mathcal{B}'}^{\mathcal{B}'}$ è ottenuta da $[f]_{\mathcal{B}}^{\mathcal{B}}$ scambiando le righe i e j e le colonne i e j .

Esercizio 4.16. Siano V, W spazi vettoriali di dimensione m e n e siano $U \subset V$ e $Z \subset W$ sottospazi di dimensione k e h . Considera il sottoinsieme

$$S = \{f: V \rightarrow W \mid f(U) \subset Z\} \subset \text{Hom}(V, W).$$

Mostra che S è un sottospazio vettoriale di $\text{Hom}(V, W)$ di dimensione $mn + hk - kn$.

Complementi

4.1. Spazio duale. Sia V uno spazio vettoriale su $\mathbb{K}$ di dimensione n . Lo *spazio duale* di V è lo spazio

$$V^* = \text{Hom}(V, \mathbb{K})$$

formato da tutte le applicazioni lineari $V \rightarrow \mathbb{K}$. Sappiamo dalla Sezione 4.3.4 che V^* è uno spazio vettoriale della stessa dimensione n di V . Gli elementi di V^* sono a volte chiamati *covettori* e sono ampiamente utilizzati in matematica ed in fisica.

Esempio 4.4.20. Un covettore di $\mathbb{K}^n$ è un elemento di $(\mathbb{K}^n)^*$, cioè una applicazione lineare $\mathbb{K}^n \rightarrow \mathbb{K}$. Ricordiamo che le applicazioni lineari $\mathbb{K}^n \rightarrow \mathbb{K}$ sono tutte del tipo $L_A: \mathbb{K}^n \rightarrow \mathbb{K}$ per qualche matrice $A \in M(1, n)$. Una matrice A di tipo $1 \times n$ è semplicemente un vettore riga. Quindi un covettore in questo contesto è semplicemente un vettore riga.

Riassumendo, i vettori di $\mathbb{K}^n$ sono i vettori colonna, mentre i covettori di $(\mathbb{K}^n)^*$ sono i vettori riga, che codificano applicazioni lineari $\mathbb{K}^n \rightarrow \mathbb{K}$. Ad esempio, il covettore

$$(1 \quad 2 \quad -1) \in (\mathbb{R}^3)^*$$

indica l'applicazione lineare

$$f: \mathbb{R}^3 \longrightarrow \mathbb{R}, \quad f \begin{pmatrix} x \\ y \\ z \end{pmatrix} = x + 2y - z.$$

Lo spazio duale di $\mathbb{K}^n$ è pienamente descritto nell'esempio precedente. Per uno spazio vettoriale V più generale, dobbiamo lavorare come sempre con le basi. Sia $v_1, \dots, v_n$ una base di V . Questa base induce una *base duale* $v_1^*, \dots, v_n^*$ di V^* nel modo seguente: definiamo

$$v_i^*: V \longrightarrow \mathbb{K}$$

come l'unica applicazione lineare tale che $v_i^*(v_j)$ sia uguale a zero se $i \neq j$ e uno se $i = j$. Questa richiesta identifica v_i^* per la Proposizione 4.1.18.

Proposizione 4.4.21. *Gli elementi $v_1^*, \dots, v_n^*$ formano una base di V^* .*

Dimostrazione. Sappiamo già che $\dim V^* = n$, quindi ci basta dimostrare che $v_1^*, \dots, v_n^*$ generano V^* . Sia $f \in V^*$ una applicazione lineare $f: V \rightarrow \mathbb{K}$. Definiamo $\lambda_i = f(v_i)$ e consideriamo l'applicazione lineare

$$g = \lambda_1 v_1^* + \dots + \lambda_n v_n^*.$$

Notiamo che

$$g(v_i) = \lambda_1 v_1^*(v_i) + \dots + \lambda_n v_n^*(v_i) = \lambda_i v_i^*(v_i) = \lambda_i.$$

Per costruzione $f(v_i) = \lambda_i = g(v_i)$ per ogni i . Siccome g e f coincidono su una base, coincidono ovunque. Quindi $f = g$ è generata dai $v_1^*, \dots, v_n^*$. $\square$

Se fissiamo una base di V , questa induce una base di V^* .

Esercizio 4.4.22. Considera la base di $\mathbb{R}^3$ data da

$$v_1 = \begin{pmatrix} 1 \\ 1 \\ 0 \end{pmatrix}, \quad v_2 = \begin{pmatrix} 1 \\ 0 \\ 1 \end{pmatrix}, \quad v_3 = \begin{pmatrix} 0 \\ 1 \\ 1 \end{pmatrix}.$$

Determina la base duale v_1^*, v_2^*, v_3^* di $(\mathbb{R}^3)^* = M(1, 3)$.

Definizione 4.4.23. Se $U \subset V$ è un sottospazio vettoriale, l'*annullatore* $\text{Ann}(U)$ è il sottoinsieme di V^* formato da tutte quelle funzioni che si annullano completamente su U , cioè

$$\text{Ann}(U) = \{v^* \in V^* \mid v^*(u) = 0 \forall u \in U\}.$$

L'esercizio seguente è un caso particolare dell'Esercizio 4.16.

Esercizio 4.4.24. Mostra che se $U \subset V$ con $\dim U = k$ e $\dim V = n$ allora $\text{Ann}(U) \subset V^*$ è un sottospazio vettoriale di dimensione $n - k$.

Cosa accade se perseveriamo nell'astrazione e costruiamo lo spazio duale dello spazio duale? Si ottiene lo spazio di partenza. Nell'ottica dell'Esempio 4.4.20 questo fatto non dovrebbe sorprendere: passare da $\mathbb{K}^n$ a $(\mathbb{K}^n)^*$ è come sostituire i vettori colonna con i vettori riga, e se si ripete ancora questo processo ritroviamo i vettori colonna iniziali.

Definizione 4.4.25. Lo *spazio biduale* è il duale del duale di V , cioè

$$V^{**} = (V^*)^*.$$

La proposizione seguente asserisce che esiste un isomorfismo *canonico* fra V e V^{**} . Questo aggettivo sta a sottolineare il fatto che questo isomorfismo non dipende da nessuna scelta, ed è una proprietà rilevante. Sappiamo che esiste un isomorfismo tra V e V^* , perché questi due spazi hanno la stessa dimensione e quindi sono isomorfi. Non esiste però nessun isomorfismo *canonico* fra V e V^* : ci sono tanti possibili isomorfismi e nessun modo di sceglierne uno in modo naturale. Invece tra V e V^{**} l'isomorfismo naturale c'è.

Proposizione 4.4.26. *Esiste un isomorfismo canonico tra V e V^{**} .*

Dimostrazione. L'isomorfismo è il seguente: ad ogni vettore $v \in V$ associamo la funzione $V^* \rightarrow \mathbb{K}$ che manda $v^* \in V^*$ in $v^*(v) \in \mathbb{K}$.

Abbiamo costruito una funzione $T: V \rightarrow V^{**}$. Formalmente:

$$T(v): V^* \rightarrow \mathbb{K}, \quad T(v)(v^*) = v^*(v).$$

Si verifica facilmente che T è lineare. Mostriamo che T è un isomorfismo. Sia $v_1, \dots, v_n$ una base di V , $v_1^*, \dots, v_n^*$ la base duale di V^* , e $v_1^{**}, \dots, v_n^{**} \in V^{**}$ la base duale della base duale. Si verifica facilmente che $T(v_i) = v_i^{**}$ e quindi T manda una base in una base. Allora è un isomorfismo. $\square$

Possiamo identificare gli spazi V e V^{**} tramite questo isomorfismo canonico. Sia adesso $U \subset V$ un sottospazio. Abbiamo $\text{Ann}(U) \subset V^*$ e quindi $\text{Ann}(\text{Ann}(U)) \subset V^{**} = V$.

Proposizione 4.4.27. *Vale $\text{Ann}(\text{Ann}(U)) = U$.*

Dimostrazione. Sappiamo dall'Esercizio 4.4.24 che $\dim \text{Ann}(\text{Ann}(U)) = \dim U$. Inoltre per costruzione $U \subset \text{Ann}(\text{Ann}(U))$, quindi i due sottospazi coincidono. $\square$

L'annullatore trasforma un k -sottospazio $U \subset V$ in un $(n-k)$ -sottospazio $\text{Ann}(U) \subset V^*$. Otteniamo in questo modo una corrispondenza biunivoca.

Corollario 4.4.28. *L'annullatore induce una corrispondenza biunivoca fra k -sottospazi di V e $(n-k)$ -sottospazi di V^* .*

Dimostrazione. La funzione inversa è sempre l'annullatore: se $W \subset V^*$ è un $(n - k)$ -sottospazio, definiamo $U = \text{Ann}(W) \subset V^{**} = V$ e per la proposizione precedente $\text{Ann}(U) = W$. $\square$

Autovettori e autovalori

In questo capitolo studiamo gli endomorfismi di uno spazio vettoriale V di dimensione n . Sappiamo che qualsiasi endomorfismo $T: V \rightarrow V$ è rappresentabile come una matrice quadrata A dopo aver fissato una base $\mathcal{B}$. La domanda a cui tentiamo di rispondere è la seguente: fra le infinite basi $\mathcal{B}$ possibili, ne esistono alcune che sono più convenienti, ad esempio perché la matrice associata A risulta essere più semplice da gestire?

La risposta a questa domanda è che le matrici più semplici da gestire sono quelle diagonali e le basi $\mathcal{B}$ che forniscono queste matrici sono chiamate *basi di autovettori*. Non tutti gli endomorfismi hanno però delle basi così comode: introdurremo degli algoritmi che ci permetteranno di capire se un endomorfismo abbia una base di autovettori e, in caso affermativo, di determinarla.

5.1. Definizioni

Introduciamo gli strumenti basilari di questo capitolo: gli *autovettori*, gli *autovalori* e infine il *polinomio caratteristico* di una matrice quadrata e di un endomorfismo.

5.1.1. Autovettori e autovalori. Sia $T: V \rightarrow V$ un endomorfismo di uno spazio vettoriale V definito su un campo $\mathbb{K}$. Un *autovettore* di T è un vettore $v \neq 0$ in V per cui

$$T(v) = \lambda v$$

per qualche scalare $\lambda \in \mathbb{K}$ che chiameremo *autovalore di T relativo a v* . Notiamo che λ può essere qualsiasi scalare, anche zero. D'altro canto, rimarchiamo che l'autovettore v *non* può essere zero per definizione. In parole semplici, un autovettore è un vettore (diverso da zero) che viene mandato da T in un multiplo di se stesso. Facciamo alcuni esempi.

Esempio 5.1.1. Consideriamo l'endomorfismo $L_A: \mathbb{R}^2 \rightarrow \mathbb{R}^2$ con

$$A = \begin{pmatrix} 3 & 4 \\ 0 & 2 \end{pmatrix}.$$

Poiché $L_A(e_1) = 3e_1$, il vettore e_1 è autovettore di L_A con autovalore 3. Inoltre notiamo che $L_A\begin{pmatrix} -4 \\ 1 \end{pmatrix} = \begin{pmatrix} -8 \\ 2 \end{pmatrix} = 2\begin{pmatrix} -4 \\ 1 \end{pmatrix}$ e quindi il vettore $\begin{pmatrix} -4 \\ 1 \end{pmatrix}$ è autovettore con autovalore 2.

Esempio 5.1.2. Sia V uno spazio vettoriale su $\mathbb{K}$. Fissiamo $\lambda \in \mathbb{K}$ e definiamo l'endomorfismo $f: V \rightarrow V$, $f(v) = \lambda v$. Tutti i vettori non nulli di V sono autovettori con lo stesso autovalore λ .

Esempio 5.1.3. Sia $L_A: \mathbb{R}^2 \rightarrow \mathbb{R}^2$ con $A = \text{Rot}_\theta$ una rotazione di angolo $\theta \neq 0, \pi$. Ciascun vettore $v \in \mathbb{R}^2$ diverso da zero viene ruotato di un angolo $\theta \neq 0, \pi$ e quindi chiaramente la sua immagine $L_A(v)$ non può essere un multiplo di v . L'endomorfismo L_A non ha autovettori.

Abbiamo visto tre esempi qualitativamente molto differenti: nel primo endomorfismo alcuni vettori dello spazio vettoriale sono autovettori, nel secondo tutti (eccetto lo 0) e nel terzo nessuno.

Facciamo adesso alcune osservazioni semplici ma cruciali.

Osservazione 5.1.4. Sia $f: V \rightarrow V$ un endomorfismo. Se $v \in V$ è autovettore per f con autovalore λ , allora qualsiasi multiplo $w = \mu v$ di v con $\mu \neq 0$ è anch'esso autovettore con lo stesso autovalore λ , infatti

$$f(\mu v) = \mu f(v) = \mu \lambda v = \lambda(\mu v).$$

Se $v \in V$ è autovettore, tutti i vettori non nulli nella retta $\text{Span}(v)$ sono anche loro autovettori con lo stesso autovalore λ .

Gli autovettori con autovalore 0 e 1 sono un po' speciali.

Osservazione 5.1.5. Sia $f: V \rightarrow V$ un endomorfismo. Un vettore $v \neq 0$ è autovettore per f con autovalore 0 se e solo se $f(v) = 0v = 0$, cioè se e solo se $v \in \ker f$.

Un *punto fisso* per un endomorfismo f è un vettore $v \in V$ tale che $f(v) = v$. I punti fissi sono punti che non vengono spostati da f .

Osservazione 5.1.6. Sia $f: V \rightarrow V$ un endomorfismo. Un vettore $v \neq 0$ è autovettore per f con autovalore 1 se e solo se v è un punto fisso.

Esempio 5.1.7. Nella rotazione in $\mathbb{R}^3$ intorno all'asse z descritta nella Sezione 4.4.11, tutti i punti contenuti nell'asse z sono punti fissi. Il vettore e_3 che genera l'asse z è un autovettore con autovalore 1.

Riassumendo:

- Se v è autovettore con autovalore λ , qualsiasi suo multiplo non nullo è anch'esso autovettore con autovalore λ .
- Gli autovettori con autovalore 0 sono i vettori non nulli del nucleo.
- Gli autovettori con autovalore 1 sono i punti fissi non nulli.

Osservazione 5.1.8. Come molte altre cose viste qui, autovettori e autovalori possono essere agevolmente studiati in coordinate rispetto ad una base. Sia $T: V \rightarrow V$ un endomorfismo di uno spazio vettoriale V . Siano $\mathcal{B}$ una base di V e $A = [T]_{\mathcal{B}}^{\mathcal{B}}$ la matrice associata. Sia $v \in V$

un vettore $x = [v]_{\mathcal{B}} \in \mathbb{K}^n$ il vettore delle sue coordinate. Grazie alla Proposizione 4.3.4, troviamo

$$T(v) = \lambda v \iff Ax = \lambda x.$$

Ogni qual volta vediamo l'equazione $T(v) = \lambda v$, possiamo leggerla in coordinate come $Ax = \lambda x$.

5.1.2. Endomorfismi diagonalizzabili. Veniamo adesso al vero motivo per cui abbiamo introdotto autovettori e autovalori.

Definizione 5.1.9. Un endomorfismo $T: V \rightarrow V$ è *diagonalizzabile* se V ha una base $\mathcal{B} = \{v_1, \dots, v_n\}$ composta da autovettori per T .

Il termine “diagonalizzabile” è dovuto al fatto seguente, che è cruciale. Sia $\mathcal{B} = \{v_1, \dots, v_n\}$ una base qualsiasi di V .

Proposizione 5.1.10. La matrice associata $A = [T]_{\mathcal{B}}^{\mathcal{B}}$ è diagonale se e solo se i vettori $v_1, \dots, v_n$ sono tutti autovettori per T .

Dimostrazione. Il vettore v_i è un autovettore $\iff f(v_i) = \lambda_i v_i$ per qualche $\lambda_i \in \mathbb{K} \iff [f(v_i)]_{\mathcal{B}} = \lambda_i e_i$. Questo capita per ogni $i = 1, \dots, n \iff A$ è diagonale, con gli autovalori sulla diagonale principale:

$$A = \begin{pmatrix} \lambda_1 & 0 & \cdots & 0 \\ 0 & \lambda_2 & \cdots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \cdots & \lambda_n \end{pmatrix}.$$

La dimostrazione è conclusa. □

Abbiamo scoperto che un endomorfismo T è diagonalizzabile se e solo se esiste una base $\mathcal{B}$ tale che $A = [T]_{\mathcal{B}}^{\mathcal{B}}$ sia una matrice diagonale. Questo accade precisamente quando $\mathcal{B}$ è una base di autovettori, e gli elementi sulla diagonale principale di A sono precisamente i loro autovalori.

Esempio 5.1.11. L'endomorfismo $L_A: \mathbb{R}^2 \rightarrow \mathbb{R}^2$ con $A = \begin{pmatrix} 3 & 4 \\ 0 & 2 \end{pmatrix}$ già considerato nell'Esempio 5.1.1 è diagonalizzabile, perché come abbiamo notato precedentemente $v_1 = \begin{pmatrix} 1 \\ 0 \end{pmatrix}$ e $v_2 = \begin{pmatrix} -4 \\ 1 \end{pmatrix}$ sono entrambi autovettori e sono indipendenti, quindi formano una base. I loro autovalori sono 3 e 2. Prendendo $\mathcal{B} = \{v_1, v_2\}$ otteniamo quindi

$$[L_A]_{\mathcal{B}}^{\mathcal{B}} = \begin{pmatrix} 3 & 0 \\ 0 & 2 \end{pmatrix}.$$

Esempio 5.1.12. L'endomorfismo $f(v) = \lambda v$ considerato nell'Esempio 5.1.2 è diagonalizzabile, perché *qualsiasi* base di V è una base di autovettori! La matrice associata è sempre λI_n , indipendentemente dalla base scelta.

Esempio 5.1.13. La rotazione di angolo θ considerata nell'Esempio 5.1.3 non è diagonalizzabile per $\theta \neq 0, \pi$ perché non ha autovettori. Per $\theta = 0$ e $\theta = \pi$ la rotazione diventa rispettivamente $f(v) = v$ e $f(v) = -v$ e quindi è diagonalizzabile.

Esempio 5.1.14. L'endomorfismo $T: M(2) \rightarrow M(2)$, $T(A) = {}^tA$ è diagonalizzabile: una base di autovettori è data da

$$A_1 = \begin{pmatrix} 1 & 0 \\ 0 & 0 \end{pmatrix}, \quad A_2 = \begin{pmatrix} 0 & 0 \\ 0 & 1 \end{pmatrix}, \quad A_3 = \begin{pmatrix} 0 & 1 \\ 1 & 0 \end{pmatrix}, \quad A_4 = \begin{pmatrix} 0 & 1 \\ -1 & 0 \end{pmatrix}.$$

Infatti otteniamo

$$T(A_1) = A_1, \quad T(A_2) = A_2, \quad T(A_3) = A_3, \quad T(A_4) = -A_4.$$

Prendendo $\mathcal{B} = \{A_1, A_2, A_3, A_4\}$ otteniamo quindi

$$[T]_{\mathcal{B}}^{\mathcal{B}} = \begin{pmatrix} 1 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 \\ 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & -1 \end{pmatrix}.$$

5.1.3. Matrici diagonali. Un endomorfismo $T: V \rightarrow V$ è diagonalizzabile precisamente quando esiste una base $\mathcal{B}$ rispetto alla quale la matrice associata $A = [T]_{\mathcal{B}}^{\mathcal{B}}$ è diagonale. Perché preferiamo le matrici diagonali? Le preferiamo perché sono molto più maneggevoli.

Notiamo innanzitutto che il prodotto Ax tra una matrice quadrata A ed un vettore x si semplifica molto se A è diagonale:

$$\begin{pmatrix} \lambda_1 & 0 & \dots & 0 \\ 0 & \lambda_2 & \dots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \dots & \lambda_n \end{pmatrix} \begin{pmatrix} x_1 \\ x_2 \\ \vdots \\ x_n \end{pmatrix} = \begin{pmatrix} \lambda_1 x_1 \\ \lambda_2 x_2 \\ \vdots \\ \lambda_n x_n \end{pmatrix}.$$

Questo vuol dire che se $[T]_{\mathcal{B}}^{\mathcal{B}}$ è diagonale è molto facile calcolare le immagini dei vettori. Inoltre il determinante di una matrice diagonale A è semplicemente il prodotto degli elementi sulla diagonale principale:

$$\det A = \lambda_1 \cdots \lambda_n.$$

Anche il prodotto fra matrici diagonali è molto semplice:

$$\begin{pmatrix} \lambda_1 & 0 & \dots & 0 \\ 0 & \lambda_2 & \dots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \dots & \lambda_n \end{pmatrix} \begin{pmatrix} \mu_1 & 0 & \dots & 0 \\ 0 & \mu_2 & \dots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \dots & \mu_n \end{pmatrix} = \begin{pmatrix} \lambda_1 \mu_1 & 0 & \dots & 0 \\ 0 & \lambda_2 \mu_2 & \dots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \dots & \lambda_n \mu_n \end{pmatrix}$$

Questo ci dice che comporre due endomorfismi che si diagonalizzano entrambi con la *stessa* base di autovettori è molto semplice. In particolare,

possiamo calcolare facilmente potenze come A^{100} :

$$A = \begin{pmatrix} \lambda_1 & 0 & \dots & 0 \\ 0 & \lambda_2 & \dots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \dots & \lambda_n \end{pmatrix} \implies A^k = \begin{pmatrix} \lambda_1^k & 0 & \dots & 0 \\ 0 & \lambda_2^k & \dots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \dots & \lambda_n^k \end{pmatrix}.$$

5.1.4. Matrici diagonalizzabili. Abbiamo notato che le matrici diagonali hanno ottime proprietà. Chiamiamo *diagonalizzabili* le matrici che possono essere trasformate in matrici diagonali in un senso opportuno.

Definizione 5.1.15. Una matrice $A \in M(n, \mathbb{K})$ è *diagonalizzabile* se è simile ad una matrice diagonale D .

Quindi A è diagonalizzabile $\iff$ esiste M invertibile tale che

$$D = M^{-1}AM$$

sia diagonale. Il collegamento con la nozione di endomorfismo diagonalizzabile è chiaramente molto forte:

Proposizione 5.1.16. Sia $\mathcal{B}$ una base di V . Un endomorfismo $T: V \rightarrow V$ è *diagonalizzabile* $\iff$ la matrice associata $A = [T]_{\mathcal{B}}^{\mathcal{B}}$ è *diagonalizzabile*.

Dimostrazione. Se T è diagonalizzabile, esiste una base $\mathcal{C}$ di V per cui $D = [T]_{\mathcal{C}}^{\mathcal{C}}$ sia diagonale; se $M = [\text{id}]_{\mathcal{B}}^{\mathcal{C}}$ è la matrice di cambiamento di base da $\mathcal{C}$ a $\mathcal{B}$ otteniamo $D = M^{-1}AM$.

D'altro canto, se $D = M^{-1}AM$ è diagonale per qualche M invertibile, sia $\mathcal{C}$ la base di V formata da vettori le cui coordinate rispetto a $\mathcal{B}$ sono le colonne di M . Allora $M = [\text{id}]_{\mathcal{B}}^{\mathcal{C}}$ e quindi $[T]_{\mathcal{C}}^{\mathcal{C}} = M^{-1}AM = D$. $\square$

Sia $A \in M(n, \mathbb{K})$ una matrice quadrata.

Corollario 5.1.17. A è *diagonalizzabile* $\iff L_A$ è *diagonalizzabile*.

Le matrici diagonalizzabili sono spesso più facili da gestire. Vediamo ad esempio come calcolare la potenza di una matrice diagonalizzabile. Notiamo alcune proprietà algebriche della similitudine fra matrici:

Proposizione 5.1.18. Valgono i fatti seguenti:

$$\begin{aligned} A = M^{-1}BM, \quad C = M^{-1}DM &\implies AC = M^{-1}BDM, \\ A = M^{-1}BM &\implies A^k = M^{-1}B^kM. \end{aligned}$$

per qualsiasi $k \in \mathbb{N}$.

Dimostrazione. Troviamo

$$AC = M^{-1}BMM^{-1}DM = M^{-1}BDM.$$

La seconda implicazione segue dalla prima iterata k volte. $\square$

Esempio 5.1.19. Prendiamo $A = \begin{pmatrix} 3 & 4 \\ 0 & 2 \end{pmatrix}$ e calcoliamo A^{100} . La matrice A non è diagonale, quindi calcolare una potenza di A direttamente è estremamente difficile. Sappiamo però che A è diagonalizzabile: dalla discussione svolta nell'Esempio 5.1.11 ricaviamo $M^{-1}AM = D = \begin{pmatrix} 3 & 0 \\ 0 & 2 \end{pmatrix}$, dove

$$M = [\text{id}]_{\mathcal{C}}^{\mathcal{B}} = \begin{pmatrix} 1 & -4 \\ 0 & 1 \end{pmatrix} \implies M^{-1} = [\text{id}]_{\mathcal{B}}^{\mathcal{C}} = \begin{pmatrix} 1 & 4 \\ 0 & 1 \end{pmatrix}.$$

Qui $\mathcal{B} = \left\{ \begin{pmatrix} 1 \\ 0 \end{pmatrix}, \begin{pmatrix} -4 \\ 1 \end{pmatrix} \right\}$ e $\mathcal{C}$ è la base canonica di $\mathbb{R}^2$. Quindi

$$\begin{aligned} A^{100} &= (MDM^{-1})^{100} = MD^{100}M^{-1} \\ &= \begin{pmatrix} 1 & -4 \\ 0 & 1 \end{pmatrix} \begin{pmatrix} 3^{100} & 0 \\ 0 & 2^{100} \end{pmatrix} \begin{pmatrix} 1 & 4 \\ 0 & 1 \end{pmatrix} \\ &= \begin{pmatrix} 1 & -4 \\ 0 & 1 \end{pmatrix} \begin{pmatrix} 3^{100} & 4 \cdot 3^{100} \\ 0 & 2^{100} \end{pmatrix} = \begin{pmatrix} 3^{100} & 4 \cdot 3^{100} - 2^{102} \\ 0 & 2^{100} \end{pmatrix}. \end{aligned}$$

5.1.5. Rette invarianti. Oltre che per le loro buone proprietà algebriche, gli endomorfismi diagonalizzabili sono preferibili anche per fondati motivi geometrici. Vediamo adesso che una base di autovettori fornisce un utile sistema di *rette invarianti*.

Introduciamo una nozione molto generale.

Definizione 5.1.20. Sia $T: V \rightarrow V$ un endomorfismo. Un *sottospazio T -invariante* (o più semplicemente *invariante*) è un sottospazio vettoriale $U \subset V$ tale che $T(U) \subset U$.

Un sottospazio $U \subset V$ è T -invariante se l'immagine $T(u)$ di ciascun vettore $u \in U$ è ancora un vettore di U . I sottospazi banale $\{0\}$ e totale V sono sempre T -invarianti; trovare altri sottospazi T -invarianti può essere più difficile.

Siamo qui interessati alle rette invarianti. Determinare le rette invarianti di T è di fatto equivalente a trovare gli autovettori per T , in virtù della proposizione seguente.

Proposizione 5.1.21. Sia $T: V \rightarrow V$ un endomorfismo e $v \in V$ un vettore non nullo. La retta $r = \text{Span}(v)$ è T -invariante $\iff v$ è autovettore per T .

Dimostrazione. Se r è T -invariante, $T(v) \in \text{Span}(v)$ e quindi $T(v) = \lambda v$ per qualche $\lambda \in \mathbb{K}$. D'altro canto, se v è un autovettore allora $T(v) \in \text{Span}(v)$ e quindi per linearità $T(\text{Span}(v)) \subset \text{Span}(v)$. $\square$

Per definizione, un endomorfismo $T: V \rightarrow V$ è diagonalizzabile $\iff$ esiste una base $v_1, \dots, v_n$ di autovettori. Per quanto abbiamo appena visto, questo è equivalente all'esistenza di un sistema $\text{Span}(v_1), \dots, \text{Span}(v_n)$ di n rette invarianti, i cui generatori sono una base per V . Questa descrizione ha una forte valenza geometrica nei casi $V = \mathbb{R}^2$ e $\mathbb{R}^3$.

5.1.6. Polinomio caratteristico. Vogliamo adesso tentare di rispondere alle seguenti domande: come facciamo a capire se un dato endomorfismo T sia diagonalizzabile? Come possiamo trovare una base $\mathcal{B}$ che diagonalizzi T , se esiste?

Negli esempi già visti precedentemente gli autovettori erano già dati e dovevamo solo verificare che funzionassero. In generale, come troviamo gli autovettori di un endomorfismo T ? Uno strumento cruciale a questo scopo è il *polinomio caratteristico*.

Definizione 5.1.22. Sia $A \in M(n, \mathbb{K})$. Il *polinomio caratteristico* di $A = (a_{ij})$ è definito nel modo seguente:

$$p_A(\lambda) = \det(A - \lambda I_n) = \det \begin{pmatrix} a_{11} - \lambda & a_{12} & \dots & a_{1n} \\ a_{21} & a_{22} - \lambda & \dots & a_{2n} \\ \vdots & \vdots & \ddots & \vdots \\ a_{n1} & a_{n2} & \dots & a_{nn} - \lambda \end{pmatrix}.$$

Un commento sulla notazione: normalmente un polinomio viene indicato con il simbolo $p(x)$, in cui x rappresenta la variabile; qui generalmente usiamo λ invece di x , perché il simbolo x è spesso usato per indicare un punto di $\mathbb{K}^n$. Il pedice A in p_A indica che il polinomio dipende dalla matrice A e a volte può essere omissso.

Da questa definizione un po' oscura di $p_A(\lambda)$ non risulta affatto chiaro che questo oggetto sia realmente un polinomio. Dimostriamo adesso che è un polinomio e identifichiamo alcuni dei suoi coefficienti.

Proposizione 5.1.23. Il *polinomio caratteristico* $p_A(\lambda)$ è effettivamente un polinomio. Ha grado n e quindi è della forma

$$p_A(\lambda) = a_n \lambda^n + a_{n-1} \lambda^{n-1} + \dots + a_0.$$

Possiamo determinare alcuni dei suoi coefficienti:

- $a_n = (-1)^n$,
- $a_{n-1} = (-1)^{n-1} \text{tr} A$,
- $a_0 = \det A$.

Dimostrazione. Il caso $n = 1$ è molto facile: la matrice è $A = (a)$ e

$$p_A(\lambda) = \det(A - \lambda I) = \det(a - \lambda) = a - \lambda = -\lambda + a.$$

Otteniamo un polinomio di grado 1 con i coefficienti giusti (qui $\det A = \text{tr} A = a$). Il caso $n = 2$ è istruttivo: qui $A = \begin{pmatrix} a & b \\ c & d \end{pmatrix}$ e quindi

$$\begin{aligned} p_A(\lambda) &= \det \begin{pmatrix} a - \lambda & b \\ c & d - \lambda \end{pmatrix} = (a - \lambda)(d - \lambda) - bc \\ &= \lambda^2 - (a + d)\lambda + ad - bc = \lambda^2 - \text{tr} A \lambda + \det A. \end{aligned}$$

Dimostriamo adesso il caso generale. Usando l'interpretazione visiva dell'Osservazione 3.3.1, scriviamo

$$p_A(\lambda) = \det(A - \lambda I_n) = \sum_{c \in \text{Col}(A)} p(c).$$

Qui $p(c)$ è il peso della colorazione c . Il peso $p(c)$ è il prodotto degli elementi nelle caselle colorate di $A - \lambda I_n$ moltiplicato per un segno, quindi è chiaramente un polinomio in λ . L'unica colorazione c che può dare un polinomio di grado n è quella sulla diagonale principale, che ha peso

$$\begin{aligned} p(c) &= (a_{11} - \lambda) \cdots (a_{nn} - \lambda) \\ &= (-1)^n \lambda^n + (-1)^{n-1} (a_{11} + \dots + a_{nn}) \lambda^{n-1} + q(\lambda) \\ &= (-1)^n \lambda^n + (-1)^{n-1} \text{tr} A \lambda^{n-1} + q(\lambda) \end{aligned}$$

dove $q(\lambda)$ è un polinomio di grado $\leq n-2$. Tutte le altre colorazioni c hanno come peso $p(c)$ un polinomio di grado $\leq n-2$. Quindi $p_A(\lambda)$ è un polinomio di grado n e i coefficienti a_n e a_{n-1} sono come indicato.

Resta solo da dimostrare che il termine noto a_0 è il determinante di A . Per fare ciò ricordiamo che $a_0 = p_A(0)$ e scriviamo semplicemente

$$a_0 = p_A(0) = \det(A - 0I) = \det A.$$

La dimostrazione è completa. □

Come il determinante e la traccia, il polinomio caratteristico è invariante per similitudine:

Proposizione 5.1.24. *Se A e B sono simili, allora $p_A(\lambda) = p_B(\lambda)$.*

Dimostrazione. Per ipotesi $A = M^{-1}BM$ con M invertibile. Notiamo che $\lambda I_n = \lambda M^{-1}M = M^{-1}(\lambda I_n)M$. Otteniamo quindi

$$\begin{aligned} p_A(\lambda) &= \det(A - \lambda I_n) = \det(M^{-1}BM - M^{-1}\lambda I_n M) \\ &= \det(M^{-1}(B - \lambda I_n)M) = \det(M^{-1}) \det(B - \lambda I_n) \det M \\ &= \det(B - \lambda I_n) = p_B(\lambda) \end{aligned}$$

grazie al Teorema di Binet. □

Sia adesso $T: V \rightarrow V$ un endomorfismo di uno spazio vettoriale V di dimensione n . Definiamo il polinomio caratteristico $p_T(\lambda)$ di T come il polinomio caratteristico $p_A(\lambda)$ della matrice associata $A = [T]_{\mathcal{B}}^{\mathcal{B}}$ rispetto ad una qualsiasi base $\mathcal{B}$ di V . La definizione non dipende dalla base scelta perché il polinomio caratteristico è invariante per similitudine.

Osservazione 5.1.25. Possiamo definire il polinomio caratteristico di un endomorfismo $T: V \rightarrow V$ direttamente scrivendo $p_T(\lambda) = \det(T - \lambda \text{id}_V)$. Qui $T - \lambda \text{id}_V$ è l'endomorfismo che manda v in $T(v) - \lambda v$.

Esempio 5.1.26. Consideriamo l'endomorfismo $T: V \rightarrow V$ dato da $f(v) = \lambda_0 v$ con $\lambda_0 \in \mathbb{K}$ fissato. La matrice associata a T rispetto a qualsiasi base è $\lambda_0 I_n$. Quindi

$$p_T(\lambda) = \det(\lambda_0 I_n - \lambda I_n) = (\lambda_0 - \lambda)^n.$$

Esempio 5.1.27. Come abbiamo già rimarcato, il polinomio caratteristico di una matrice $A \in M(2, \mathbb{K})$ è

$$p_A(\lambda) = \lambda^2 - \operatorname{tr} A \lambda + \det A.$$

Il polinomio caratteristico di una matrice $A \in M(3, \mathbb{K})$ è

$$p_A(\lambda) = -\lambda^3 + \operatorname{tr} A \lambda^2 + a_1 \lambda + \det A.$$

Tutti i coefficienti di $p_A(\lambda)$ sono invarianti per similitudine: il coefficiente a_1 è un nuovo numero invariante per similitudine, diverso dalla traccia e dal determinante.

Vedremo successivamente altri esempi; prima spieghiamo il motivo per cui abbiamo introdotto il polinomio caratteristico.

5.1.7. Le radici del polinomio caratteristico. Il polinomio caratteristico è uno strumento utile a trovare autovalori e autovettori di un endomorfismo. Sia V uno spazio di dimensione n e $T: V \rightarrow V$ un endomorfismo.

Proposizione 5.1.28. Gli autovalori di T sono precisamente le radici del polinomio caratteristico $p_T(\lambda)$.

Dimostrazione. Scegliamo una base $\mathcal{B}$ e scriviamo $A = [T]_{\mathcal{B}}^{\mathcal{B}}$. Per l'Osservazione 5.1.8, uno scalare $\lambda \in \mathbb{K}$ è autovalore per T se e solo se esiste un $x \in \mathbb{K}^n$ non nullo tale che $Ax = \lambda x$. I fatti seguenti sono tutti equivalenti:

$$\begin{aligned} \exists x \neq 0 \text{ tale che } Ax &= \lambda x \iff \\ \exists x \neq 0 \text{ tale che } Ax - \lambda x &= 0 \iff \\ \exists x \neq 0 \text{ tale che } (A - \lambda I_n)x &= 0 \iff \\ \exists x \neq 0 \text{ tale che } x \in \ker(A - \lambda I_n) &\iff \\ \det(A - \lambda I_n) = 0 &\iff p_A(\lambda) = 0. \end{aligned}$$

La dimostrazione è completa. □

Abbiamo scoperto un metodo per identificare gli autovalori di T : dobbiamo solo (si fa per dire...) trovare le radici del polinomio caratteristico $p_T(\lambda)$. Questo ovviamente non è un problema semplice da risolvere in generale: non esiste una formula per trovare le radici di un polinomio di grado qualsiasi. Per i polinomi di grado due, e quindi per le matrici 2×2 , il problema è però sempre risolvibile. Vediamo alcuni esempi 2×2 .

Esempio 5.1.29. Consideriamo l'endomorfismo $L_A: \mathbb{R}^2 \rightarrow \mathbb{R}^2$ con

$$A = \begin{pmatrix} -1 & 2 \\ -4 & 5 \end{pmatrix}.$$

Il suo polinomio caratteristico è

$$p_A(\lambda) = \lambda^2 - \text{tr}A \lambda + \det A = \lambda^2 - 4\lambda + 3.$$

Le sue radici sono $\lambda = 1$ e $\lambda = 3$. Quindi gli autovalori di A sono 1 e 3. Per ciascun autovalore $\lambda = 1$ e $\lambda = 3$, gli autovettori relativi sono precisamente le soluzioni non banali dei sistemi $A \begin{pmatrix} x \\ y \end{pmatrix} = \begin{pmatrix} x \\ y \end{pmatrix}$ e $A \begin{pmatrix} x \\ y \end{pmatrix} = 3 \begin{pmatrix} x \\ y \end{pmatrix}$. Questi sono:

$$\begin{cases} -x + 2y = x \\ -4x + 5y = y \end{cases}, \quad \begin{cases} -x + 2y = 3x \\ -4x + 5y = 3y \end{cases}$$

Questi sono due sistemi omogenei e le soluzioni rispettive sono due rette

$$\text{Span} \begin{pmatrix} 1 \\ 1 \end{pmatrix}, \quad \text{Span} \begin{pmatrix} 1 \\ 2 \end{pmatrix}.$$

Due generatori qualsiasi di queste rette (ad esempio $v_1 = \begin{pmatrix} 1 \\ 1 \end{pmatrix}$ e $v_2 = \begin{pmatrix} 1 \\ 2 \end{pmatrix}$) formano una base di autovettori per L_A . La matrice A è quindi diagonalizzabile, ed è simile alla matrice $D = \begin{pmatrix} 1 & 0 \\ 0 & 3 \end{pmatrix}$. Con $M = \begin{pmatrix} 1 & 1 \\ 1 & 2 \end{pmatrix}$ otteniamo $D = M^{-1}AM$.

Esempio 5.1.30. Consideriamo una rotazione $L_A: \mathbb{R}^2 \rightarrow \mathbb{R}^2$ antioraria di un angolo $\theta \neq 0, \pi$. La rotazione è descritta dalla matrice

$$A = \text{Rot}_\theta = \begin{pmatrix} \cos \theta & -\sin \theta \\ \sin \theta & \cos \theta \end{pmatrix}.$$

Ragionando geometricamente, abbiamo già dimostrato che L_A non ha autovettori. Possiamo verificarlo anche algebricamente guardando il polinomio caratteristico $p(\lambda) = \lambda^2 - 2\cos \theta + 1$ e notando che ha due radici complesse non reali, perché $\Delta = 4(\cos^2 \theta - 1) < 0$. Il polinomio caratteristico non ha radici reali: quindi non ci sono autovalori, e di conseguenza neppure autovettori.

Il prossimo esempio descrive un fenomeno rilevante: la diagonalizzabilità di una matrice A può dipendere dal campo $\mathbb{K}$ che stiamo considerando.

Esempio 5.1.31. Consideriamo l'endomorfismo $L_A: \mathbb{R}^2 \rightarrow \mathbb{R}^2$ con

$$A = \begin{pmatrix} 0 & -1 \\ 1 & 0 \end{pmatrix}.$$

Questo è una rotazione di $\frac{\pi}{2}$ e quindi sappiamo già che non è diagonalizzabile su $\mathbb{R}$. Se però consideriamo l'endomorfismo $L_A: \mathbb{C}^2 \rightarrow \mathbb{C}^2$ con la stessa matrice A la musica cambia: il polinomio caratteristico è infatti

$p(\lambda) = \lambda^2 + 1$ ed ha soluzioni $\lambda = \pm i$. Studiamo i sistemi $A \begin{pmatrix} x \\ y \end{pmatrix} = i \begin{pmatrix} x \\ y \end{pmatrix}$ e $A \begin{pmatrix} x \\ y \end{pmatrix} = -i \begin{pmatrix} x \\ y \end{pmatrix}$:

$$\begin{cases} -y = ix \\ x = iy \end{cases}, \quad \begin{cases} -y = -ix \\ x = -iy \end{cases}$$

Le soluzioni sono le rette

$$\text{Span} \begin{pmatrix} i \\ 1 \end{pmatrix}, \quad \text{Span} \begin{pmatrix} -i \\ 1 \end{pmatrix}.$$

I vettori $v_1 = \begin{pmatrix} i \\ 1 \end{pmatrix}$ e $v_2 = \begin{pmatrix} -i \\ 1 \end{pmatrix}$ sono autovettori di $L_A: \mathbb{C}^2 \rightarrow \mathbb{C}^2$ con autovalori i e $-i$. Quindi la matrice A non è diagonalizzabile su $\mathbb{R}$, ma lo è su $\mathbb{C}$. Prendendo $\mathcal{B} = \{v_1, v_2\}$ otteniamo

$$[L_A]_{\mathcal{B}}^{\mathcal{B}} = \begin{pmatrix} i & 0 \\ 0 & -i \end{pmatrix}.$$

In virtù di questo esempio, quando parliamo di diagonalizzabilità di matrici dovremo sempre specificare chiaramente con quale campo $\mathbb{K}$ stiamo lavorando. Nel prossimo esempio mostriamo una matrice che non è diagonalizzabile né su $\mathbb{R}$ né su $\mathbb{C}$:

Esempio 5.1.32. Consideriamo l'endomorfismo $L_A: \mathbb{R}^2 \rightarrow \mathbb{R}^2$ con

$$A = \begin{pmatrix} 1 & 1 \\ 0 & 1 \end{pmatrix}$$

Il polinomio caratteristico è $p(\lambda) = \lambda^2 - 2\lambda + 1 = (\lambda - 1)^2$. Il polinomio ha una radice $\lambda = 1$ con molteplicità due. Quindi $\lambda = 1$ è l'unico autovalore. Risolvendo il sistema $A \begin{pmatrix} x \\ y \end{pmatrix} = \begin{pmatrix} x \\ y \end{pmatrix}$ si trova che lo spazio delle soluzioni è la retta $\text{Span}(e_1)$. Troviamo un autovettore per L_A , ma non ne troviamo due indipendenti: quindi L_A non è diagonalizzabile. Con lo stesso ragionamento vediamo che neppure $L_A: \mathbb{C}^2 \rightarrow \mathbb{C}^2$ è diagonalizzabile. La matrice A non è diagonalizzabile su nessun campo, né $\mathbb{R}$ né $\mathbb{C}$.

Riassumendo, abbiamo esaminato vari tipi di esempi:

- La matrice $\begin{pmatrix} 1 & 0 \\ 0 & 1 \end{pmatrix}$ è già diagonale e quindi è diagonalizzabile sia su $\mathbb{C}$ che su $\mathbb{R}$. Ha un solo autovalore 1.
- La matrice $\begin{pmatrix} -1 & 2 \\ -4 & 5 \end{pmatrix}$ è diagonalizzabile sia su $\mathbb{C}$ che su $\mathbb{R}$. Ha due autovalori 1 e 3.
- La matrice $\begin{pmatrix} 0 & -1 \\ 1 & 0 \end{pmatrix}$ è diagonalizzabile su $\mathbb{C}$ ma non su $\mathbb{R}$. Ha due autovalori $\pm i$ su $\mathbb{C}$.
- La matrice $\begin{pmatrix} 1 & 1 \\ 0 & 1 \end{pmatrix}$ non è diagonalizzabile né su $\mathbb{C}$ né su $\mathbb{R}$. Ha un solo autovalore 1.

Esercizio 5.1.33. Per ciascuna delle matrici seguenti, determinare autovalori e autovettori in $\mathbb{R}^3$ e dire se è diagonalizzabile:

$$\begin{pmatrix} 1 & 0 & 0 \\ 0 & 1 & 1 \\ 0 & 0 & 1 \end{pmatrix}, \quad \begin{pmatrix} 2 & 1 & 1 \\ 0 & 2 & 1 \\ 0 & 0 & 2 \end{pmatrix}, \quad \begin{pmatrix} 1 & 1 & 1 \\ 0 & 1 & 0 \\ 1 & 1 & 1 \end{pmatrix}.$$

5.1.8. Matrici triangolari. Ricordiamo che una matrice quadrata A è *triangolare superiore* se è della forma

$$A = \begin{pmatrix} a_{11} & a_{12} & \dots & a_{1n} \\ 0 & a_{22} & \dots & a_{2n} \\ \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \dots & a_{nn} \end{pmatrix}$$

Il determinante di una matrice triangolare A è semplicemente $\det A = a_{11} \cdots a_{nn}$. Anche gli autovalori si determinano immediatamente:

Proposizione 5.1.34. *Gli autovalori di una matrice triangolare A sono gli elementi $a_{11}, \dots, a_{nn}$ sulla diagonale principale.*

Dimostrazione. Vale

$$\begin{aligned} p_A(\lambda) &= \det(A - \lambda I) = \det \begin{pmatrix} a_{11} - \lambda & a_{12} & \dots & a_{1n} \\ 0 & a_{22} - \lambda & \dots & a_{2n} \\ \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \dots & a_{nn} - \lambda \end{pmatrix} \\ &= (a_{11} - \lambda) \cdots (a_{nn} - \lambda). \end{aligned}$$

Le radici di $p_A(\lambda)$ sono quindi $a_{11}, \dots, a_{nn}$. Il caso in cui A sia triangolare inferiore è analogo. $\square$

Notiamo che gli elementi $a_{11}, \dots, a_{nn}$ possono anche ripetersi. Esaminiamo anche il caso in cui una matrice sia descritta a blocchi, con il blocco in basso a sinistra nullo.

Proposizione 5.1.35. *Sia $A \in M(n)$ del tipo*

$$A = \begin{pmatrix} B & C \\ 0 & D \end{pmatrix}$$

con $B \in M(k)$ e $D \in M(n-k)$ per qualche k . Vale la relazione

$$p_A(\lambda) = p_B(\lambda) \cdot p_D(\lambda).$$

Dimostrazione. Scriviamo

$$\begin{aligned} p_A(\lambda) &= \det(A - \lambda I) = \det \begin{pmatrix} B - \lambda I & C \\ 0 & D - \lambda I \end{pmatrix} \\ &= \det(B - \lambda I) \cdot \det(D - \lambda I) = p_B(\lambda) \cdot p_D(\lambda). \end{aligned}$$

Nella terza uguaglianza abbiamo usato l'Esercizio 3.13. $\square$

Quindi gli autovalori di A sono quelli di B uniti a quelli di D .

5.1.9. Esistenza di autovettori. Concludiamo questa sezione con un paio di proposizioni che mostrano l'esistenza di autovettori in alcune condizioni.

Proposizione 5.1.36. Un endomorfismo $T: V \rightarrow V$ di uno spazio vettoriale complesso V ha sempre almeno un autovettore.

Dimostrazione. Il polinomio caratteristico ha sempre almeno una radice (su $\mathbb{C}$). Quindi T ha almeno un autovalore λ , e quindi anche un autovettore. $\square$

Questo risultato non vale su $\mathbb{R}$. Ad esempio, abbiamo visto che una rotazione di un angolo $\theta \neq 0, \pi$ è una trasformazione di $\mathbb{R}^2$ che non ha autovettori. Possiamo comunque ottenere lo stesso risultato in dimensione dispari:

Proposizione 5.1.37. Sia V uno spazio vettoriale reale di dimensione dispari. Un endomorfismo $T: V \rightarrow V$ ha sempre almeno un autovettore.

Dimostrazione. Il polinomio caratteristico ha grado dispari e quindi ha sempre una radice in $\mathbb{R}$ per la Proposizione 1.4.13. $\square$

Da tutte queste considerazioni algebriche deduciamo un corollario geometrico che riguarda lo spazio tridimensionale.

Corollario 5.1.38. Un endomorfismo di $\mathbb{R}^3$ ha sempre una retta invariante.

5.2. Teorema di diagonalizzabilità

Nelle pagine precedenti abbiamo descritto un metodo per trovare gli autovalori di un endomorfismo tramite il suo polinomio caratteristico. In questa sezione ci proponiamo di determinare e studiare gli autovettori. Concludiamo infine il capitolo con il *teorema di diagonalizzabilità*, che fornisce un criterio completo per capire se un endomorfismo sia diagonalizzabile.

5.2.1. Autovettori con autovalori distinti. Dimostriamo alcune proprietà algebriche degli autovettori. Sia $T: V \rightarrow V$ un endomorfismo.

Proposizione 5.2.1. Se $v_1, \dots, v_k \in V$ sono autovettori per T con autovalori $\lambda_1, \dots, \lambda_k$ distinti, allora sono indipendenti.

Dimostrazione. Procediamo per induzione su k . Se $k = 1$, il vettore v_1 è indipendente semplicemente perché non è nullo (per definizione, un autovettore non è mai nullo).

Diamo per buono il caso $k - 1$ e mostriamo il caso k . Supponiamo di avere una combinazione lineare nulla

$$(6) \quad \alpha_1 v_1 + \dots + \alpha_k v_k = 0.$$

Dobbiamo dimostrare che $\alpha_i = 0$ per ogni i . Applicando T otteniamo

$$(7) \quad 0 = \alpha_1 T(v_1) + \dots + \alpha_k T(v_k) = \alpha_1 \lambda_1 v_1 + \dots + \alpha_k \lambda_k v_k.$$

Moltiplicando l'equazione (6) per λ_k ricaviamo

$$\alpha_1 \lambda_k v_1 + \dots + \alpha_k \lambda_k v_k = 0$$

e sottraendo questa equazione dalla (7) deduciamo che

$$\alpha_1(\lambda_1 - \lambda_k)v_1 + \dots + \alpha_{k-1}(\lambda_{k-1} - \lambda_k)v_{k-1} = 0.$$

Questa è una combinazione lineare nulla di $k-1$ autovettori con autovalori distinti: per l'ipotesi induttiva tutti i coefficienti $\alpha_i(\lambda_i - \lambda_k)$ devono essere nulli. Poiché $\lambda_i \neq \lambda_k$, ne deduciamo che $\alpha_i = 0$ per ogni $i = 1, \dots, k-1$ e usando (6) otteniamo anche $\alpha_k = 0$. La dimostrazione è completa. $\square$

Sia V uno spazio vettoriale su $\mathbb{K}$ di dimensione n e $T: V \rightarrow V$ un endomorfismo. Indichiamo come sempre con $p_T(\lambda)$ il suo polinomio caratteristico.

Corollario 5.2.2. *Se il polinomio caratteristico $p_T(\lambda)$ ha n radici distinte in $\mathbb{K}$, l'endomorfismo T è diagonalizzabile.*

Dimostrazione. Le radici sono $\lambda_1, \dots, \lambda_n$ e ciascun λ_i è autovalore di qualche autovettore v_i . Gli autovettori $v_1, \dots, v_n$ sono indipendenti per la Proposizione 5.2.1, quindi sono una base di V . Quindi V ha una base formata da autovettori, cioè è diagonalizzabile. $\square$

Rimarchiamo l'importanza del campo $\mathbb{K}$ nell'enunciato del corollario: il criterio funziona se il polinomio caratteristico ha n radici distinte nello stesso campo $\mathbb{K}$ su cui è definito V . Il criterio si applica anche alle matrici quadrate A , considerate sempre nello stesso campo $\mathbb{K}$.

Esempio 5.2.3. Come notato nell'Esempio 5.1.31, la matrice $\begin{pmatrix} 0 & -1 \\ 1 & 0 \end{pmatrix}$ non è diagonalizzabile su $\mathbb{R}$. Su $\mathbb{C}$ la stessa matrice ha due autovalori distinti $\pm i$ e quindi è diagonalizzabile per il Corollario 5.2.2.

Esempio 5.2.4. Consideriamo una matrice triangolare A . Gli autovalori di A sono gli elementi $a_{11}, \dots, a_{nn}$ sulla diagonale principale. Se gli autovalori $a_{11}, \dots, a_{nn}$ sono tutti distinti, allora A è diagonalizzabile per il Corollario 5.2.2.

Abbiamo ottenuto un criterio di diagonalizzabilità molto economico: possiamo dimostrare in alcuni casi che una matrice è diagonalizzabile senza cercare una esplicita base di autovettori. Ad esempio, la matrice

$$A = \begin{pmatrix} 1 & 7 & -1 \\ 0 & 2 & 8 \\ 0 & 0 & 3 \end{pmatrix}$$

è diagonalizzabile perché ha tre autovalori distinti 1, 2, 3. Notiamo che trovare una base esplicita di autovettori non è immediato: per ciascun

$\lambda = 1, 2, 3$ dobbiamo risolvere l'equazione $Av = \lambda v$ (il lettore è invitato a farlo).

Il teorema di diagonalizzabilità che vogliamo dimostrare sarà una generalizzazione di questo criterio. Per enunciare questo teorema dobbiamo introdurre ancora una nozione.

5.2.2. Autospatio. Sia $T: V \rightarrow V$ un endomorfismo. Per ogni autovalore λ di T definiamo l'*autospatio* $V_\lambda \subset V$ nel modo seguente:

$$V_\lambda = \{v \in V \mid T(v) = \lambda v\}.$$

L'autospatio V_λ consiste precisamente di tutti gli autovettori v con autovalore λ , più l'origine $0 \in V$ (ricordiamo che $0 \in V$ non è autovettore per definizione).

Per ogni autovalore λ di T , ha senso considerare l'endomorfismo $T - \lambda \text{id}$ di V che manda v in $T(v) - \lambda v$ e la proposizione seguente lo mette in relazione con l'autospatio V_λ .

Proposizione 5.2.5. *Per ogni autovalore λ di T abbiamo*

$$V_\lambda = \ker(T - \lambda \text{id}).$$

Dimostrazione. Troviamo

$$\begin{aligned} V_\lambda &= \{v \in V \mid T(v) = \lambda v\} = \{v \in V \mid T(v) - \lambda v = 0\} \\ &= \{v \in V \mid (T - \lambda \text{id})v = 0\} = \ker(T - \lambda \text{id}). \end{aligned}$$

La dimostrazione è conclusa. □

Corollario 5.2.6. *L'autospatio V_λ è un sottospazio vettoriale di V .*

Ricordiamo la somma diretta di più sottospazi definita nella Sezione 2.3.12. Sia V uno spazio vettoriale di dimensione finita. Mostriamo adesso che gli autospatzi di un endomorfismo sono sempre in somma diretta.

Proposizione 5.2.7. *Sia $T: V \rightarrow V$ un endomorfismo e siano $\lambda_1, \dots, \lambda_k$ i suoi autovalori. I corrispondenti autospatzi sono sempre in somma diretta:*

$$V_{\lambda_1} \oplus \dots \oplus V_{\lambda_k}.$$

Dimostrazione. Dobbiamo mostrare che

$$V_{\lambda_i} \cap (V_{\lambda_1} + \dots + V_{\lambda_{i-1}} + V_{\lambda_{i+1}} + \dots + V_{\lambda_k}) = \{0\} \quad \forall i = 1, \dots, k.$$

Sia per assurdo $v \neq 0$ un vettore nell'intersezione. Per ipotesi $v \in V_{\lambda_i}$ e

$$v = v_1 + \dots + v_{i-1} + v_{i+1} + \dots + v_k$$

con $v_j \in V_{\lambda_j}$. Questa è una relazione di dipendenza lineare fra autovettori con autovalori distinti, un assurdo per la Proposizione 5.2.1. □

Corollario 5.2.8. *L'endomorfismo T è diagonalizzabile se e solo se*

$$V = V_{\lambda_1} \oplus \dots \oplus V_{\lambda_k}.$$

Dimostrazione. Sappiamo già che gli autospazi sono in somma diretta, quindi dobbiamo mostrare che $V = V_{\lambda_1} + \cdots + V_{\lambda_k} \iff$ esiste una base di autovettori per T .

($\Rightarrow$) Unendo le basi dei V_{λ_i} otteniamo una base di autovettori per T .

($\Leftarrow$) Ciascun vettore v è combinazione lineare di autovettori e quindi chiaramente $v \in V_{\lambda_1} + \cdots + V_{\lambda_k}$. $\square$

5.2.3. Molteplicità algebrica e geometrica. Per enunciare il teorema di diagonalizzabilità abbiamo bisogno di introdurre due ultime definizioni. In tutta questa sezione V è uno spazio vettoriale di dimensione n .

Sia $T: V \rightarrow V$ un endomorfismo e sia λ un autovalore per T . La *molteplicità algebrica* $m_a(\lambda)$ è la molteplicità di λ come radice del polinomio caratteristico p_T . La *molteplicità geometrica* $m_g(\lambda)$ è la dimensione dell'autospazio associato a λ , cioè

$$m_g(\lambda) = \dim V_\lambda.$$

Esempio 5.2.9. Sia $L_A: \mathbb{R}^2 \rightarrow \mathbb{R}^2$ l'endomorfismo dato da

$$A = \begin{pmatrix} 1 & 1 \\ 0 & 1 \end{pmatrix}.$$

Il polinomio caratteristico è $p_A(\lambda) = \lambda^2 - 2\lambda + 1 = (\lambda - 1)^2$. Troviamo un solo autovalore $\lambda_1 = 1$, con molteplicità algebrica $m_a(1) = 2$. Inoltre

$$m_g(1) = \dim V_1 = \dim \ker(A - I_2) = \dim \ker \begin{pmatrix} 0 & 1 \\ 0 & 0 \end{pmatrix} = 1.$$

Abbiamo quindi trovato in questo caso

$$m_a(1) = 2, \quad m_g(1) = 1.$$

L'esempio ci mostra che le due molteplicità possono essere diverse. In generale, la molteplicità geometrica è un numero che sta sempre tra 1 e la molteplicità algebrica:

Proposizione 5.2.10. *Sia $T: V \rightarrow V$ un endomorfismo. Per ogni autovalore λ_0 di T valgono le disuguaglianze*

$$1 \leq m_g(\lambda_0) \leq m_a(\lambda_0).$$

Dimostrazione. Se λ_0 è un autovalore, esiste almeno un autovettore $v \neq 0$ tale che $T(v) = \lambda_0 v$ e quindi l'autospazio V_{λ_0} ha dimensione almeno uno. Questo dimostra la prima disuguaglianza.

Per dimostrare la seconda, prendiamo una base $\{v_1, \dots, v_k\}$ di V_{λ_0} , formata da $k = m_g(\lambda_0) = \dim V_{\lambda_0}$ vettori, e la completiamo a base $\mathcal{B} = \{v_1, \dots, v_n\}$ di V . Per i primi k vettori della base abbiamo $T(v_i) = \lambda_0 v_i$. Quindi la matrice $[T]_{\mathcal{B}}^{\mathcal{B}}$ è della forma

$$[T]_{\mathcal{B}}^{\mathcal{B}} = \begin{pmatrix} \lambda_0 I_k & C \\ 0 & D \end{pmatrix}$$

dove $\lambda_0 I_k \in M(k)$ e $D \in M(n-k)$. Per il Corollario 5.1.35 il polinomio caratteristico di T è del tipo

$$p_T(\lambda) = p_{\lambda_0 I_k}(\lambda) \cdot p_D(\lambda) = (\lambda_0 - \lambda)^k \cdot p_D(\lambda).$$

Ne segue che λ_0 ha molteplicità almeno $k = m_g(\lambda_0)$ in $p_T(\lambda)$. In altre parole, vale la disuguaglianza $m_a(\lambda_0) \geq m_g(\lambda_0)$. $\square$

5.2.4. Matrici simili. Sia $A \in M(n, \mathbb{K})$ una matrice quadrata. Gli *autovettori* e *autovalori* di A sono per definizione gli autovettori e autovalori di $L_A: \mathbb{K}^n \rightarrow \mathbb{K}^n$.

Notiamo che le molteplicità algebrica e geometrica sono entrambe invarianti per similitudine di una matrice.

Proposizione 5.2.11. *Se A e B sono simili, hanno gli stessi autovalori con le stesse molteplicità algebriche e geometriche.*

Dimostrazione. Le matrici A e B hanno gli stessi autovalori con le stesse molteplicità algebriche perché hanno lo stesso polinomio caratteristico. Sia $m_g^X(\lambda)$ la molteplicità geometrica di λ nella matrice X . Vale $A = M^{-1}BM$ e

$$\begin{aligned} m_g^A(\lambda) &= \dim \ker(A - \lambda I_n) = \dim \ker(M^{-1}BM - M^{-1}\lambda I_n M) \\ &= \dim \ker(M^{-1}(B - \lambda I_n)M) = \dim \ker(B - \lambda I_n) = m_g^B(\lambda). \end{aligned}$$

Nella penultima uguaglianza abbiamo usato che due matrici simili hanno lo stesso rango e quindi anche la stessa dimensione del nucleo. $\square$

D'altro canto, due matrici simili A e B *non* hanno gli stessi autovettori.

5.2.5. Teorema di diagonalizzabilità. Possiamo finalmente enunciare il teorema seguente. Sia V spazio vettoriale su $\mathbb{K}$ di dimensione n .

Teorema 5.2.12 (Teorema di diagonalizzabilità). *Un endomorfismo $T: V \rightarrow V$ è diagonalizzabile se e solo se valgono entrambi i fatti seguenti:*

- (1) $p_T(\lambda)$ ha n radici in $\mathbb{K}$, contate con molteplicità.
- (2) $m_a(\lambda) = m_g(\lambda)$ per ogni autovalore λ di T .

Dimostrazione. Sappiamo che gli autospazi sono in somma diretta e definiamo il sottospazio $W \subset V$ come la loro somma

$$W = V_{\lambda_1} \oplus \cdots \oplus V_{\lambda_k}.$$

Sappiamo che T è diagonalizzabile $\iff W = V \iff \dim W = n$. Inoltre

$$\begin{aligned} \dim W &= \dim V_{\lambda_1} + \cdots + \dim V_{\lambda_k} = m_g(\lambda_1) + \cdots + m_g(\lambda_k) \\ &\leq m_a(\lambda_1) + \cdots + m_a(\lambda_k) \leq n. \end{aligned}$$

Nella prima disuguaglianza abbiamo usato che $m_g(\lambda_i) \leq m_a(\lambda_i)$ e nella seconda che $\lambda_1, \dots, \lambda_k$ sono radici del polinomio caratteristico, che ha grado n e quindi ha al più n radici contate con molteplicità per il Teorema 1.3.7.

Chiaramente $\dim W = n$ se e solo se entrambe le disuguaglianze sono uguaglianze. La prima disuguaglianza è un'uguaglianza precisamente quando $m_g(\lambda_i) = m_a(\lambda_i)$ per ogni i , e la seconda precisamente quando $p_T(\lambda)$ ha n radici in $\mathbb{K}$, contate con molteplicità. $\square$

5.2.6. Esempi. Mostriamo alcuni esempi in cui applichiamo il teorema di diagonalizzabilità.

Esempio 5.2.13. Studiamo la diagonalizzabilità su $\mathbb{R}$ della matrice

$$A = \begin{pmatrix} 3 & 0 & 0 \\ -4 & -1 & -8 \\ 0 & 0 & 3 \end{pmatrix}.$$

Si vede facilmente che il polinomio caratteristico è

$$p_A(\lambda) = (3 - \lambda)(-1 - \lambda)(3 - \lambda),$$

quindi ha radici $\lambda_1 = 3$ con $m_a(\lambda_1) = 2$ e $\lambda_2 = -1$ con $m_a(\lambda_2) = 1$.

Tutte le radici di $p_A(\lambda)$ sono reali, quindi A è diagonalizzabile se e solo se le molteplicità algebriche e geometriche di ciascun autovalore coincidono. Per il secondo autovalore λ_2 è facile: la Proposizione 5.2.10 implica che $m_g(\lambda_2) = m_a(\lambda_2) = 1$ e quindi siamo a posto.

Dobbiamo concentrarci solo sull'autovalore λ_1 che ha $m_a(\lambda_1) = 2$. La Proposizione 5.2.10 ci dice che $m_g(\lambda_1)$ può essere 1 oppure 2: nel primo caso A non è diagonalizzabile, nel secondo sì. Facciamo i conti e troviamo

$$m_g(3) = \dim V_3 = \dim \ker(A - 3I) = 3 - \text{rk}(A - 3I).$$

Nell'ultima uguaglianza abbiamo usato il teorema della dimensione. Quindi

$$\begin{aligned} m_g(3) &= 3 - \text{rk} \begin{pmatrix} 3-3 & 0 & 0 \\ -4 & -1-3 & -8 \\ 0 & 0 & 3-3 \end{pmatrix} \\ &= 3 - \text{rk} \begin{pmatrix} 0 & 0 & 0 \\ -4 & -4 & -8 \\ 0 & 0 & 0 \end{pmatrix} = 3 - 1 = 2. \end{aligned}$$

Abbiamo scoperto che $m_a(\lambda_1) = m_g(\lambda_1) = 2$ e quindi A è diagonalizzabile.

Esempio 5.2.14. Studiamo la diagonalizzabilità su $\mathbb{R}$ della matrice

$$A = \begin{pmatrix} 3 & t+4 & 1 \\ -1 & -3 & -1 \\ 0 & 0 & 2 \end{pmatrix}.$$

al variare del parametro $t \in \mathbb{R}$. Si vede facilmente che il polinomio caratteristico è

$$p_A(\lambda) = (2 - \lambda)(\lambda^2 + t - 5).$$

Se $t > 5$, il polinomio $\lambda^2 + t - 5$ non ha radici reali, quindi $p_A(\lambda)$ ha una radice sola e A non è diagonalizzabile. Se $t \leq 5$, il polinomio caratteristico

ha tre radici reali

$$\lambda_1 = 2, \quad \lambda_2 = \sqrt{5-t}, \quad \lambda_3 = -\sqrt{5-t}.$$

Se le tre radici sono distinte, la matrice A è diagonalizzabile. Restano da considerare i casi in cui le tre radici non sono distinte, e cioè i casi $t = 1$ e $t = 5$. Questi due casi vanno analizzati separatamente con le tecniche mostrate nell'esempio precedente.

Se $t = 1$ gli autovalori sono $\lambda_1 = 2, \lambda_2 = 2$ e $\lambda_3 = -2$ e la matrice è

$$A = \begin{pmatrix} 3 & 5 & 1 \\ -1 & -3 & -1 \\ 0 & 0 & 2 \end{pmatrix}.$$

Calcoliamo la molteplicità geometrica dell'autovalore 2:

$$m_g(2) = 3 - \text{rk} \begin{pmatrix} 1 & 5 & 1 \\ -1 & -5 & -1 \\ 0 & 0 & 0 \end{pmatrix} = 3 - 1 = 2.$$

Otteniamo $m_g(2) = 2 = m_a(2)$ e quindi A è diagonalizzabile.

Se $t = 5$ gli autovalori sono $\lambda_1 = 2, \lambda_2 = 0$ e $\lambda_3 = 0$ e la matrice A è

$$A = \begin{pmatrix} 3 & 9 & 1 \\ -1 & -3 & -1 \\ 0 & 0 & 2 \end{pmatrix}.$$

La molteplicità geometrica dell'autovalore 0 è $m_g(0) = 3 - \text{rk}(A) = 3 - 2 = 1 \neq 2 = m_a(0)$. Quindi A non è diagonalizzabile.

Riassumendo, la matrice A è diagonalizzabile se e solo se $t < 5$.

5.2.7. Traccia, determinante e autovalori. Concludiamo con una osservazione che può essere utile. Sia $A \in M(n, \mathbb{K})$ una matrice quadrata.

Proposizione 5.2.15. *Se $p_A(\lambda)$ ha tutte le radici $\lambda_1, \dots, \lambda_n$ in $\mathbb{K}$, allora*

$$\text{tr} A = \lambda_1 + \dots + \lambda_n,$$

$$\det A = \lambda_1 \cdots \lambda_n.$$

Dimostrazione. Entrambe le uguaglianze sono ovvie se A è diagonale, e quindi sono valide anche se A è diagonalizzabile, perché traccia, determinante e autovalori non cambiano per similitudine. Questo però non basta: dobbiamo dimostrare la proposizione per tutte le matrici, non solo quelle diagonalizzabili.

Sappiamo dalla Proposizione 5.1.23 che

$$p_A(\lambda) = (-1)^n \lambda^n + (-1)^{n-1} \text{tr} A \lambda^{n-1} + \dots + \det A.$$

Poiché $p_A(\lambda)$ ha tutte le radici $\lambda_1, \dots, \lambda_n \in \mathbb{K}$, vale anche

$$p_A(\lambda) = (-1)^n (\lambda - \lambda_1) \cdots (\lambda - \lambda_n)$$

$$= (-1)^n \lambda^n + (-1)^{n-1} (\lambda_1 + \dots + \lambda_n) \lambda^{n-1} + \dots + \lambda_1 \cdots \lambda_n.$$

Eguagliando i coefficienti nelle due scritture di $p_A(\lambda)$ si ottiene la tesi. $\square$

Esercizi

Esercizio 5.1. Quali delle matrici seguenti sono diagonalizzabili su $\mathbb{R}$?

$$\begin{pmatrix} 2 & -1 & 1 \\ 6 & -3 & 4 \\ 3 & -2 & 3 \end{pmatrix}, \quad \begin{pmatrix} 1 & 1 & 1 \\ 1 & 1 & 1 \\ -2 & -2 & -2 \end{pmatrix}, \quad \begin{pmatrix} a & b & c \\ 0 & a & d \\ 0 & 0 & a \end{pmatrix},$$

$$\begin{pmatrix} 1 & k & 0 \\ 1 & k & 1 \\ 0 & -k & 1 \end{pmatrix}, \quad \begin{pmatrix} k & 0 & 0 \\ k+1 & -1 & 0 \\ k-1 & k & k^2 \end{pmatrix}, \quad \begin{pmatrix} k-2 & 0 & -4 & 0 \\ 0 & k & 2 & 0 \\ 0 & 0 & k+2 & 0 \\ 1 & 1 & k & 2-k \end{pmatrix},$$

$$\begin{pmatrix} 1 & 1 & 1 & 1 & 1 & 1 \\ 1 & 1 & 1 & 1 & 1 & 1 \\ 1 & 1 & 1 & 1 & 1 & 1 \\ 1 & 1 & 1 & 1 & 1 & 1 \\ 1 & 1 & 1 & 1 & 1 & 1 \\ 1 & 1 & 1 & 1 & 1 & 1 \end{pmatrix}, \quad \begin{pmatrix} 1 & -1 & 1 & -1 & 1 & -1 & 1 \\ 3 & -1 & -1 & -1 & -1 & -1 & 3 \\ 1 & -1 & 1 & -1 & 1 & -1 & 1 \\ 3 & -1 & -1 & -1 & -1 & -1 & 3 \\ 1 & -1 & 1 & -1 & 1 & -1 & 1 \\ 3 & -1 & -1 & -1 & -1 & -1 & 3 \\ 1 & -1 & 1 & -1 & 1 & -1 & 1 \end{pmatrix}.$$

Quelle che non sono diagonalizzabili su $\mathbb{R}$, lo sono su $\mathbb{C}$? Le risposte dipendono dai parametri reali presenti nella matrice.

Esercizio 5.2. Per quali $t \in \mathbb{C}$ la matrice A studiata nell'Esempio 5.2.14 è diagonalizzabile su $\mathbb{C}$?

Esercizio 5.3. Trova gli autovalori in $\mathbb{C}$ e gli autovettori in $\mathbb{C}^2$ della matrice complessa $\begin{pmatrix} 1 & 2 \\ 1+i & i \end{pmatrix}$.

Esercizio 5.4. Data la matrice $A = \begin{pmatrix} 23 & -84 \\ 6 & -22 \end{pmatrix}$, determina A^{10} .

Esercizio 5.5. Sia $A \in M(2, \mathbb{R})$ una matrice fissata. Considera l'endomorfismo $L_A: M(2, \mathbb{R}) \rightarrow M(2, \mathbb{R})$ dato da $L_A(X) = AX$.

- Sia $A = \begin{pmatrix} 2 & 1 \\ 1 & 1 \end{pmatrix}$. Scrivi la matrice associata a L_A rispetto ad una base di $M(2, \mathbb{R})$ a tua scelta.
- Sia $A = \begin{pmatrix} 2 & 1 \\ 1 & 1 \end{pmatrix}$. L'endomorfismo L_A è diagonalizzabile?
- Dimostra che in generale l'endomorfismo L_A è diagonalizzabile $\iff A$ è diagonalizzabile.
- Scrivi una base di autovettori per L_A nel caso $A = \begin{pmatrix} 0 & 1 \\ 1 & 0 \end{pmatrix}$.

Esercizio 5.6. Considera l'endomorfismo $T: M(2) \rightarrow M(2)$ dato da $T(X) = {}^tX$.

- Scrivi la matrice associata a T rispetto alla base canonica di $M(2)$.
- L'endomorfismo T è diagonalizzabile?

Esercizio 5.7. Costruisci un endomorfismo di $\mathbb{R}^4$ senza autovettori. Più in generale, per ogni $n \geq 1$ costruisci un endomorfismo di $\mathbb{R}^{2n}$ senza autovettori.

Esercizio 5.8. Sia $T: V \rightarrow V$ un endomorfismo invertibile. Mostra che T è diagonalizzabile se e solo se T^{-1} è diagonalizzabile.

Esercizio 5.9. È sempre vero che la composizione di due endomorfismi diagonalizzabili è diagonalizzabile? Se sì, dimostrarlo. Se no, esibisci un controesempio.

Esercizio 5.10. Mostra che se $T: V \rightarrow V$ è diagonalizzabile allora V si decompone in somma diretta come $V = \ker T \oplus \text{Im } T$.

Esercizio 5.11. Mostra che una matrice $A \in M(n, \mathbb{K})$ è diagonalizzabile se e solo se lo è la sua trasposta tA .

Esercizio 5.12. Mostra che una matrice $A \in M(n, \mathbb{K})$ di rango 1 è diagonalizzabile se e solo se $\ker L_A \oplus \operatorname{Im} L_A = \mathbb{K}^n$.

Esercizio 5.13. Considera l'endomorfismo $T: \mathbb{R}_n[x] \rightarrow \mathbb{R}_n[x]$ dato da

$$T(p(x)) = (x+1)p'(x)$$

dove p' è la derivata di p . L'endomorfismo T è diagonalizzabile?

Esercizio 5.14. Sia $T: V \rightarrow V$ un endomorfismo tale che $T \circ T = T$.

- (1) Mostra che $V = \ker T \oplus \operatorname{Im} T$.
- (2) Sia $V_\lambda = \{v \in V \mid T(v) = \lambda v\}$. Mostra che $V_0 = \ker T$ e $V_1 = \operatorname{Im} T$.
- (3) Concludi che T è diagonalizzabile.

Esercizio 5.15. Sia $T: V \rightarrow V$ un endomorfismo tale che $T \circ T = \operatorname{id}$.

- (1) Sia $V_\lambda = \{v \in V \mid T(v) = \lambda v\}$. Mostra che $V = V_1 \oplus V_{-1}$. Il trucco è scrivere ogni vettore $v \in V$ come somma

$$v = \frac{v + T(v)}{2} + \frac{v - T(v)}{2}.$$

- (2) Concludi che T è diagonalizzabile.

Esercizio 5.16. Sia $T: V \rightarrow V$ un endomorfismo tale che $T \circ T = 0$.

- (1) Mostra che l'unico autovalore per T è 0.
- (2) Mostra che T è diagonalizzabile se e solo se $T = 0$.

Esercizio 5.17. Considera una matrice a blocchi

$$A = \begin{pmatrix} B & C \\ 0 & D \end{pmatrix}.$$

in cui A , B e D sono matrici quadrate. Usando l'Esercizio 3.14, mostra che per ogni autovalore λ di A vale

$$m_g^A(\lambda) \leq m_g^B(\lambda) + m_g^D(\lambda)$$

dove indichiamo con $m_g^X(\lambda)$ la molteplicità geometrica di λ nella matrice X .

Esercizio 5.18. Deduci dall'esercizio precedente che se A è diagonalizzabile allora B e D sono entrambe diagonalizzabili. È vero l'opposto?

Esercizio 5.19. Sia $T: V \rightarrow V$ un endomorfismo e $W \subset V$ un sottospazio T -invariante. Usando l'esercizio precedente, mostra che se T è diagonalizzabile allora anche la restrizione

$$T|_W: W \rightarrow W$$

è diagonalizzabile.

Forma di Jordan

In questo capitolo affronteremo il problema seguente: sappiamo che un endomorfismo è più facile da studiare se è diagonalizzabile; cosa possiamo fare nel caso in cui non lo sia? Vedremo che potremo comunque scegliere una matrice relativamente semplice, che è “quasi” diagonale, eccetto che per alcuni valori sopra la diagonale principale: questo tipo di matrice particolare e utile è detto *matrice di Jordan*.

In tutto questo capitolo il campo base $\mathbb{K}$ sarà sempre quello dei numeri complessi $\mathbb{K} = \mathbb{C}$. In questo modo potremo sfruttare il teorema fondamentale dell'algebra – cioè il fatto che ogni polinomio di grado n abbia n radici contate con molteplicità. La teoria nel caso reale è un po' più complicata ed è illustrata brevemente (senza dimostrazioni) nei complementi.

6.1. Forma di Jordan

Una *matrice di Jordan* è un tipo di matrice “quasi” diagonale in cui compaiono solo alcuni valori non nulli sopra la diagonale principale. Mostriamo che qualsiasi matrice A , anche non diagonalizzabile, è simile ad un'unica matrice di Jordan, detta *forma di Jordan* di A .

La forma di Jordan caratterizza completamente la classe di similitudine di A e può essere usata per determinare se due matrici sono simili.

6.1.1. Blocco di Jordan. In tutto questo capitolo lavoriamo sempre con il campo $\mathbb{K} = \mathbb{C}$ dei numeri complessi. Un *blocco di Jordan* è una matrice quadrata $n \times n$ del tipo

$$B_{\lambda,n} = \begin{pmatrix} \lambda & 1 & 0 & \dots & 0 & 0 \\ 0 & \lambda & 1 & \dots & 0 & 0 \\ 0 & 0 & \lambda & \dots & 0 & 0 \\ \vdots & \vdots & \vdots & \ddots & \vdots & \vdots \\ 0 & 0 & 0 & \dots & \lambda & 1 \\ 0 & 0 & 0 & \dots & 0 & \lambda \end{pmatrix}.$$

Qui $\lambda \in \mathbb{C}$ è uno scalare complesso fissato qualsiasi. Un blocco di Jordan $B_{\lambda,n}$ è una matrice $n \times n$ triangolare superiore in cui tutte le caselle della diagonale principale hanno lo stesso numero $\lambda \in \mathbb{C}$ e tutte le caselle sulla diagonale immediatamente superiore hanno lo stesso numero 1. Le altre

diagonali hanno solo zeri. Ad esempio questi sono blocchi di Jordan:

$$(3), \quad \begin{pmatrix} i & 1 \\ 0 & i \end{pmatrix}, \quad \begin{pmatrix} 0 & 1 & 0 \\ 0 & 0 & 1 \\ 0 & 0 & 0 \end{pmatrix}, \quad \begin{pmatrix} -1 & 1 & 0 & 0 \\ 0 & -1 & 1 & 0 \\ 0 & 0 & -1 & 1 \\ 0 & 0 & 0 & -1 \end{pmatrix}.$$

Notiamo che una matrice 1×1 è sempre un blocco di Jordan. Osserviamo anche che un blocco di Jordan $n \times n$ è triangolare e quindi i suoi autovalori sono sulla diagonale: ne segue che la matrice ha un solo autovalore λ . Calcoliamo le sue molteplicità algebrica e geometrica.

Proposizione 6.1.1. *Le molteplicità dell'autovalore λ sono:*

$$m_a(\lambda) = n, \quad m_g(\lambda) = 1.$$

In particolare se $n \geq 2$ un blocco di Jordan non è mai diagonalizzabile.

Dimostrazione. La molteplicità algebrica è chiara perché la matrice è triangolare superiore. Per la molteplicità geometrica, notiamo che

$$m_g(\lambda) = n - \text{rk} \begin{pmatrix} 0 & 1 & 0 & \dots & 0 & 0 \\ 0 & 0 & 1 & \dots & 0 & 0 \\ 0 & 0 & 0 & \dots & 0 & 0 \\ \vdots & \vdots & \vdots & \ddots & \vdots & \vdots \\ 0 & 0 & 0 & \dots & 0 & 1 \\ 0 & 0 & 0 & \dots & 0 & 0 \end{pmatrix} = n - (n - 1) = 1.$$

La matrice descritta ha effettivamente rango $n - 1$ perché, dopo aver tolto la prima colonna, le colonne rimanenti sono chiaramente indipendenti. $\square$

6.1.2. Matrice di Jordan. Una *matrice di Jordan* è una matrice a blocchi

$$J = \begin{pmatrix} B_1 & 0 & \dots & 0 \\ 0 & B_2 & \dots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \dots & B_k \end{pmatrix}$$

dove $B_1, \dots, B_k$ sono blocchi di Jordan di taglia arbitraria. Notiamo che tutti i blocchi che non stanno sulla diagonale principale sono nulli. Ad esempio, queste sono matrici di Jordan:

$$\begin{pmatrix} -i & 1 & 0 \\ 0 & -i & 0 \\ 0 & 0 & 3 \end{pmatrix}, \quad \begin{pmatrix} 1 & 0 & 0 \\ 0 & 1 & 1 \\ 0 & 0 & 1 \end{pmatrix}, \quad \begin{pmatrix} -1 & 0 & 0 \\ 0 & -1 & 0 \\ 0 & 0 & 1+i \end{pmatrix}, \quad \begin{pmatrix} 2 & 1 & 0 \\ 0 & 2 & 1 \\ 0 & 0 & 2 \end{pmatrix}.$$

Le prime due matrici hanno due blocchi di Jordan, la terza ne ha tre, l'ultima ne ha uno solo (è essa stessa un blocco di Jordan). Una matrice diagonale $n \times n$ è una matrice di Jordan con n blocchi di taglia 1×1 .

Notiamo un fatto semplice ma importante: una matrice di Jordan è sempre triangolare superiore e quindi i suoi autovalori sono sulla diagonale.

Proposizione 6.1.2. *Se permutiamo i blocchi di una matrice di Jordan J , otteniamo un'altra matrice di Jordan J' simile a J .*

Dimostrazione. È chiaro che J' è sempre di Jordan. La similitudine fra J e J' discende da un fatto generale: se abbiamo una matrice a blocchi

$$A = \begin{pmatrix} B_1 & 0 & \cdots & 0 \\ 0 & B_2 & \cdots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \cdots & B_k \end{pmatrix}$$

e permutiamo arbitrariamente i blocchi, otteniamo una nuova matrice A' simile a A . Basta usare l'Esercizio 4.15 per mostrare che la permutazione è realizzata da una corrispondente permutazione dei vettori della base canonica di $\mathbb{C}^n$. $\square$

6.1.3. Teorema di Jordan. Enunciamo adesso il teorema più importante di questo capitolo.

Teorema 6.1.3. *Qualsiasi matrice $A \in M(n, \mathbb{C})$ è simile ad una matrice di Jordan J . La matrice di Jordan J è unica a meno di permutare i blocchi.*

Una dimostrazione di questo teorema sarà fornita nella Sezione 6.1.6. È importante notare come la matrice di Jordan J sia unica a meno di permutare i blocchi: per questo J viene chiamata la *forma canonica di Jordan* di A .

Corollario 6.1.4. *Due matrici A e A' sono simili $\iff$ le loro matrici di Jordan J e J' hanno gli stessi blocchi.*

Esempio 6.1.5. Consideriamo le matrici seguenti:

$$\begin{pmatrix} 2 & 0 & 0 \\ 0 & 2 & 0 \\ 0 & 0 & 2 \end{pmatrix}, \begin{pmatrix} 2 & 1 & 0 \\ 0 & 2 & 0 \\ 0 & 0 & 2 \end{pmatrix}, \begin{pmatrix} 2 & 0 & 0 \\ 0 & 2 & 1 \\ 0 & 0 & 2 \end{pmatrix}, \begin{pmatrix} 2 & 1 & 0 \\ 0 & 2 & 1 \\ 0 & 0 & 2 \end{pmatrix}, \begin{pmatrix} 2 & 1 & 1 \\ 0 & 2 & 1 \\ 0 & 0 & 2 \end{pmatrix}.$$

Quali di queste sono simili fra loro? Le prime quattro matrici sono matrici di Jordan. La seconda e la terza sono simili, perché hanno gli stessi blocchi, solo permutati in modo diverso. La prima, la seconda e la quarta sono però tutte non simili fra loro, perché hanno blocchi diversi. L'ultima matrice non è in forma di Jordan e quindi non possiamo dire immediatamente se sia simile ad una delle precedenti: c'è del lavoro da fare.

Il Teorema 6.1.3 è molto potente, ma osserviamo che enunciato in questo modo non fornisce un metodo per costruire J a partire da A . Esiste un algoritmo per costruire J da A che descriveremo alla fine di questa sezione, ma non è semplicissimo e per adesso ci limitiamo ad usare il teorema: nei casi più semplici questo sarà sufficiente.

Esempio 6.1.6. Due matrici simili A e B hanno lo stesso determinante, la stessa traccia, lo stesso polinomio caratteristico, gli stessi autovalori, e per ciascun autovalore hanno anche le stesse molteplicità algebrica e geometrica (queste cose sono state dimostrate nel capitolo precedente). Usando tutte queste proprietà possiamo spesso riconoscere due matrici che non sono simili.

Ad esempio, le matrici dell'Esempio 6.1.5 hanno tutte lo stesso polinomio caratteristico (e quindi stessa traccia, determinante e autovalori). Però le molteplicità geometriche dell'autovalore 2 non sono tutte uguali: facendo un conto vediamo che queste sono rispettivamente

$$3, \quad 2, \quad 2, \quad 1, \quad 1.$$

La quinta matrice quindi può essere simile solo alla quarta. Ma lo è veramente? Usando il Teorema 6.1.3 possiamo rispondere affermativamente nel modo seguente. Ci chiediamo: quale è la forma di Jordan dell'ultima matrice? Siccome ha solo l'autovalore 2, deve essere una delle precedenti: si verifica facilmente infatti che queste sono tutte le possibili matrici di Jordan 3×3 con un solo autovalore 2. Di queste, solo la quarta può essere simile a lei, quindi sono simili.

Esempio 6.1.7. Due matrici A e B che hanno lo stesso polinomio caratteristico e le stesse molteplicità algebriche e geometriche degli autovalori, sono sempre simili? La risposta è negativa: le matrici seguenti

$$\left(\begin{array}{cc|cc} \boxed{1} & \boxed{1} & 0 & 0 \\ \boxed{0} & \boxed{1} & 0 & 0 \\ \hline 0 & 0 & \boxed{1} & \boxed{1} \\ 0 & 0 & \boxed{0} & \boxed{1} \end{array} \right), \quad \left(\begin{array}{ccc|c} \boxed{1} & \boxed{1} & \boxed{0} & 0 \\ \boxed{0} & \boxed{1} & \boxed{1} & 0 \\ \boxed{0} & \boxed{0} & \boxed{1} & 0 \\ \hline 0 & 0 & 0 & \boxed{1} \end{array} \right)$$

hanno un solo autovalore 1 ed in entrambi i casi la molteplicità geometrica di 1 è 2, come si vede facilmente con un conto. Però si decompongono in blocchi di Jordan diversi: quella di sinistra in due blocchi di ordine 2, quella di destra in due blocchi di ordine 3 e 1. Quindi per il Teorema 6.1.3 non sono simili.

6.1.4. Matrici nilpotenti. Una matrice $A \in M(n, \mathbb{C})$ è *nilpotente* se esiste un $k \geq 0$ per cui $A^k = 0$. Ovviamente, se questo accade, allora $A^h = 0$ anche per ogni $h \geq k$. L'*indice di nilpotenza* di A è il più piccolo numero intero $k \geq 1$ per cui $A^k = 0$.

Esempio 6.1.8. Una matrice A è nilpotente con indice di nilpotenza 1 se e solo se $A = 0$. Quindi solo la matrice nulla è nilpotente con indice 1.

Ad esempio, consideriamo il blocco di Jordan $B = B_{0,2}$. Otteniamo

$$B = \begin{pmatrix} 0 & 1 \\ 0 & 0 \end{pmatrix}, \quad B^2 = \begin{pmatrix} 0 & 0 \\ 0 & 0 \end{pmatrix}.$$

Quindi B è nilpotente con indice di nilpotenza 2. Con $B = B_{0,3}$ otteniamo

$$B = \begin{pmatrix} 0 & 1 & 0 \\ 0 & 0 & 1 \\ 0 & 0 & 0 \end{pmatrix}, \quad B^2 = \begin{pmatrix} 0 & 0 & 1 \\ 0 & 0 & 0 \\ 0 & 0 & 0 \end{pmatrix}, \quad B^3 = \begin{pmatrix} 0 & 0 & 0 \\ 0 & 0 & 0 \\ 0 & 0 & 0 \end{pmatrix}.$$

Questo ci dice che B è nilpotente con indice di nilpotenza 3. Con la matrice $B = B_{0,4}$ otteniamo

$$B = \begin{pmatrix} 0 & 1 & 0 & 0 \\ 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 1 \\ 0 & 0 & 0 & 0 \end{pmatrix}, \quad B^2 = \begin{pmatrix} 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 1 \\ 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 \end{pmatrix}, \quad B^3 = \begin{pmatrix} 0 & 0 & 0 & 1 \\ 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 \end{pmatrix}$$

e infine $B^4 = 0$, quindi B ha indice di nilpotenza 4. Più in generale, abbiamo capito che ogni volta che moltiplichiamo $B_{0,n}$ per se stesso spostiamo i numeri 1 sulla diagonale nella diagonale immediatamente superiore. In questo modo dimostriamo il fatto seguente.

Proposizione 6.1.9. *La matrice $B_{0,n}$ è nilpotente con indice di nilpotenza n .*

Dimostrazione. Si dimostra facilmente per induzione su i che la matrice $B_{0,n}^k$ ha zeri ovunque tranne nelle caselle $i, i+k$ in cui c'è 1. $\square$

Esistono molte altre matrici nilpotenti oltre alle $B_{0,n}$. Ad esempio, una qualsiasi matrice simile a $B_{0,n}$ è anch'essa nilpotente:

Esercizio 6.1.10. Siano A e B due matrici simili. Se A è nilpotente con indice k , allora anche B è nilpotente con indice k .

Esercizio 6.1.11. Le matrici 2×2 nilpotenti sono precisamente le

$$\begin{pmatrix} a & b \\ c & -a \end{pmatrix}$$

con $a, b, c \in \mathbb{C}$ tali che $bc = -a^2$. Queste sono precisamente le matrici 2×2 con traccia e determinante entrambi nulli.

La semplice caratterizzazione con traccia e determinante nulli però funziona solo per le matrici 2×2 , non per le 3×3 .

Esercizio 6.1.12. Scrivi una matrice $A \in M(3, \mathbb{C})$ con $\text{tr} A = \det A = 0$ che non sia nilpotente.

Una matrice di Jordan con solo l'autovalore zero è nilpotente. Infatti:

$$J = \begin{pmatrix} B_{0,m_1} & 0 & \dots & 0 \\ 0 & B_{0,m_2} & \dots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \dots & B_{0,m_k} \end{pmatrix} \Rightarrow J^i = \begin{pmatrix} B_{0,m_1}^i & 0 & \dots & 0 \\ 0 & B_{0,m_2}^i & \dots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \dots & B_{0,m_k}^i \end{pmatrix}$$

e otteniamo $J^i = 0$ se $i = \max\{m_1, \dots, m_k\}$. Usando la forma di Jordan, possiamo fornire una caratterizzazione completa delle matrici nilpotenti:

Proposizione 6.1.13. *Una matrice $A \in M(n, \mathbb{C})$ è nilpotente $\iff p_A(\lambda) = \lambda^n \iff A$ ha solo l'autovalore zero.*

Dimostrazione. Se A ha un autovalore $\lambda \neq 0$, allora esiste un $x \in \mathbb{C}^n$ non nullo tale che $Ax = \lambda x$ e allora $A^k x = \lambda^k x \neq 0$ per ogni k . Quindi $A^k \neq 0$ per ogni k e A non è nilpotente.

Se A ha solo l'autovalore zero, la sua forma di Jordan J è fatta di blocchi di tipo $B_{0,m}$ e quindi è nilpotente. Siccome $A \sim J$, anche A è nilpotente. $\square$

6.1.5. Un algoritmo per determinare J . Come possiamo determinare la matrice di Jordan J associata ad una data matrice complessa A ? In molti casi è sufficiente andare per tentativi, come abbiamo visto sopra: si elencano le possibili matrici di Jordan e si usa la molteplicità geometrica per eliminarne alcune, sperando di restare alla fine con una sola. A volte però questo metodo non funziona: esponiamo ora un algoritmo più complesso che funziona sempre.

Sia λ un autovalore di A . Per ciascun numero intero $j = 1, 2, \dots$ definiamo il j -esimo autospazio generalizzato

$$V_\lambda^j = \ker(A - \lambda I_n)^j$$

e la j -esima molteplicità geometrica generalizzata

$$m_j = m_g(\lambda)_j = \dim V_\lambda^j.$$

La molteplicità m_1 è l'usuale molteplicità geometrica di λ .

Proposizione 6.1.14. *Le molteplicità generalizzate formano una successione non decrescente:*

$$m_1 \leq m_2 \leq m_3 \leq \dots$$

Dimostrazione. Scriviamo $B = A - \lambda I_n$. Per ogni $j \geq 1$, se $B^j x = 0$ per qualche $x \in \mathbb{C}^n$, allora anche $B^{j+1} x = 0$. Quindi $\ker B^j \subset \ker B^{j+1}$. $\square$

Proposizione 6.1.15. *Due matrici simili A e B hanno le stesse molteplicità geometriche generalizzate per ciascun autovalore.*

Dimostrazione. Se $A = M^{-1}BM$, allora

$$(A - \lambda I)^j = (M^{-1}BM - \lambda M^{-1}M)^j = (M^{-1}(B - \lambda I)M)^j = M^{-1}(B - \lambda I)^j M.$$

Quindi le matrici $(A - \lambda I)^j$ e $(B - \lambda I)^j$ sono anch'esse simili; hanno lo stesso rango, e quindi anche la stessa dimensione del nucleo. $\square$

Sia J la matrice di Jordan associata ad A . Poiché J e A sono simili, hanno le stesse molteplicità geometriche generalizzate. Il punto fondamentale è che queste determinano univocamente J , come mostrato nella proposizione seguente. Poniamo $m_0 = 0$.

Proposizione 6.1.16. *Per ogni $i \geq 1$, il numero $m_i - m_{i-1}$ è il numero di blocchi di J con autovalore λ e di taglia $\geq i$. In particolare la molteplicità geometrica m_1 è il numero di blocchi con autovalore λ .*

Dimostrazione. Ricordiamo che $B_{\lambda,n}$ indica un blocco di Jordan di ordine n con autovalore λ . A meno di permutare i blocchi, possiamo mettere tutti i blocchi di J con autovalore λ all'inizio, e scrivere

$$J = \begin{pmatrix} B_{\lambda,n_1} & \dots & 0 & 0 \\ \vdots & \ddots & \vdots & \vdots \\ 0 & \dots & B_{\lambda,n_a} & 0 \\ 0 & \dots & 0 & J' \end{pmatrix}.$$

dove $n_1, \dots, n_a$ sono gli ordini dei blocchi con autovalore λ . La matrice di Jordan J' contiene i blocchi con gli altri autovalori. Notiamo che

$$m_1 = \dim \ker(J - \lambda I) = \dim \ker \begin{pmatrix} B_{0,n_1} & \dots & 0 & 0 \\ \vdots & \ddots & \vdots & \vdots \\ 0 & \dots & B_{0,n_a} & 0 \\ 0 & \dots & 0 & J' - \lambda I \end{pmatrix}.$$

La matrice $J' - \lambda I$ è invertibile (perché è triangolare superiore con elementi tutti non nulli sulla diagonale), e m_1 è il numero di colonne nulle della matrice, perché le altre colonne non nulle sono chiaramente indipendenti. C'è una colonna nulla in ogni blocco $B_{0,n}$, quindi $m_1 = a$ è il numero di tali blocchi.

Calcoliamo ora m_2 e notiamo che

$$m_2 = \dim \ker(J - \lambda I)^2 = \dim \ker \begin{pmatrix} B_{0,n_1}^2 & \dots & 0 & 0 \\ \vdots & \ddots & \vdots & \vdots \\ 0 & \dots & B_{0,n_a}^2 & 0 \\ 0 & \dots & 0 & (J' - \lambda I)^2 \end{pmatrix}.$$

La matrice $(J' - \lambda I)^2$ è sempre invertibile, e m_2 è sempre il numero di colonne nulle (perché le colonne non nulle sono indipendenti). Le colonne nulle sono le m_1 colonne di prima, più $m_2 - m_1$ nuove colonne, una per ogni blocco $B_{0,n}$ di taglia almeno due. Quindi $m_2 - m_1$ è proprio il numero di blocchi di taglia ≥ 2 . I passi successivi della dimostrazione sono analoghi. La dimostrazione è completa. $\square$

Esempio 6.1.17. L'esempio numerico seguente dovrebbe chiarire meglio cosa succede ed aiutare anche a capire la dimostrazione appena conclusa. Consideriamo

$$J = \begin{pmatrix} \boxed{2} & \boxed{1} & \boxed{0} & 0 & 0 & 0 & 0 & 0 \\ 0 & 2 & 1 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 2 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & \boxed{2} & \boxed{1} & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 2 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & \boxed{2} & \boxed{1} & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 2 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & \boxed{2} \end{pmatrix}.$$

Questa ha un blocco di ordine 3, due blocchi di ordine 2 e un blocco di ordine 1. Otteniamo

$$m_1 = \dim \ker(J - 2I) = \dim \ker \begin{pmatrix} \boxed{0} & \boxed{1} & \boxed{0} & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 1 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & \boxed{0} & \boxed{1} & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & \boxed{0} & \boxed{0} & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & \boxed{0} & \boxed{1} \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & \boxed{0} \end{pmatrix} = 4$$

perché sono comparse 4 colonne di zeri, una per ogni blocco: infatti la molteplicità geometrica m_1 è il numero di blocchi. Continuiamo e troviamo

$$m_2 = \dim \ker(J - 2I)^2 = \dim \ker \begin{pmatrix} \boxed{0} & \boxed{0} & \boxed{1} & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & \boxed{0} & \boxed{0} & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & \boxed{0} & \boxed{0} & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & \boxed{0} & \boxed{0} \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & \boxed{0} \end{pmatrix} = 7.$$

Al secondo passaggio, si sono manifestate 3 nuove colonne di zeri: una per ogni blocco di taglia ≥ 2 (i blocchi di taglia 1 erano già spariti). I blocchi di taglia ≥ 2 sono effettivamente $m_2 - m_1 = 7 - 4 = 3$.

Infine $(J - 2I)^3 = 0$ ci dice che $m_3 = 8$ e infatti i blocchi di taglia ≥ 3 sono esattamente $m_3 - m_2 = 8 - 7 = 1$.

Dalla Proposizione 6.1.16 possiamo ricavare alcune informazioni più precise sulla successione non decrescente $m_1 \leq m_2 \leq \dots$ delle molteplicità geometriche generalizzate. Sia come sopra J una matrice di Jordan e λ un autovalore per J . Ricordiamo che $m_g(\lambda)$ e $m_a(\lambda)$ indicano le molteplicità geometrica e algebrica di λ .

Corollario 6.1.18. *Sia h la massima taglia di un blocco di Jordan in J con autovalore λ . La successione delle molteplicità geometriche generalizzate di λ è strettamente crescente fino a $m_h = m_a(\lambda)$ e poi si stabilizza:*

$$m_1 = m_g(\lambda) < m_2 < \cdots < m_h = m_a(\lambda) = m_{h+1} = \cdots$$

In particolare notiamo che $m_g(\lambda) = m_a(\lambda)$ precisamente quando $h = 1$, cioè quando tutti i blocchi di Jordan relativi a λ hanno taglia 1.

Esempio 6.1.19. Usiamo l'algoritmo per determinare la forma di Jordan J della matrice

$$A = \begin{pmatrix} 2 & 2 & 0 & -1 \\ 0 & 0 & 0 & 1 \\ 1 & 5 & 2 & -1 \\ 0 & -4 & 0 & 4 \end{pmatrix}.$$

Il polinomio caratteristico è $p_A(\lambda) = (\lambda - 2)^4$. Per determinare la forma di Jordan dobbiamo calcolare le molteplicità geometriche generalizzate dell'autovalore 2. Iniziamo:

$$m_1 = \dim \ker(A - 2I) = \dim \ker \begin{pmatrix} 0 & 2 & 0 & -1 \\ 0 & -2 & 0 & 1 \\ 1 & 5 & 0 & -1 \\ 0 & -4 & 0 & 2 \end{pmatrix} = 2.$$

Scopriamo che la molteplicità geometrica è 2, quindi il numero di blocchi di J è 2. Questo però non è sufficiente a determinare J , perché ci sono due casi possibili, cioè

$$\begin{pmatrix} \boxed{\begin{matrix} 2 & 1 \\ 0 & 2 \end{matrix}} & 0 & 0 \\ 0 & 0 & \boxed{\begin{matrix} 2 & 1 \\ 0 & 2 \end{matrix}} \end{pmatrix} \quad \text{oppure} \quad \begin{pmatrix} \boxed{\begin{matrix} 2 & 1 & 0 \\ 0 & 2 & 1 \\ 0 & 0 & 2 \end{matrix}} & 0 \\ 0 & 0 & 0 & \boxed{2} \end{pmatrix}.$$

Allora continuiamo a calcolare le molteplicità geometriche generalizzate:

$$m_2 = \dim \ker(A - 2I)^2 = \dim \ker \begin{pmatrix} 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 \\ 0 & -4 & 0 & 2 \\ 0 & 0 & 0 & 0 \end{pmatrix} = 3.$$

Quindi il numero di blocchi di J di taglia ≥ 2 è $m_2 - m_1 = 3 - 2 = 1$. Questo esclude il primo dei due casi possibili per J elencati sopra, quindi troviamo

$$J = \begin{pmatrix} \boxed{\begin{matrix} 2 & 1 & 0 \\ 0 & 2 & 1 \\ 0 & 0 & 2 \end{matrix}} & 0 \\ 0 & 0 & 0 & \boxed{2} \end{pmatrix}.$$

Le molteplicità geometriche generalizzate caratterizzano completamente la classe di similitudine di una matrice:

Proposizione 6.1.20. *Due matrici $A, B \in M(n, \mathbb{C})$ sono simili $\iff$ hanno gli stessi autovalori con le stesse molteplicità geometriche generalizzate.*

Dimostrazione. Due matrici sono simili $\iff$ hanno la stessa matrice di Jordan $\iff$ hanno gli stessi autovalori con le stesse molteplicità geometriche generalizzate per la Proposizione 6.1.16. $\square$

6.1.6. Autospazi generalizzati massimali. Adesso forniamo una dimostrazione del Teorema di Jordan. Non è necessario conoscere la dimostrazione per usare il teorema e per comprendere gli argomenti successivi: questa parte può essere saltata e si può continuare la lettura dalla Sezione 6.2.

Ne approfittiamo per studiare il problema dal punto di vista degli endomorfismi, senza usare le matrici. Sia V uno spazio vettoriale complesso di dimensione n e $T: V \rightarrow V$ un endomorfismo.

Per ogni autovalore λ di T e $j \geq 1$ consideriamo l'autospazio generalizzato

$$V_{\lambda}^j = \ker(T - \lambda \text{id})^j.$$

È facile verificare che ogni autospazio generalizzato è contenuto nel successivo:

$$V_{\lambda}^1 \subset V_{\lambda}^2 \subset V_{\lambda}^3 \subset \dots$$

In questa successione di spazi, ogni volta che abbiamo una inclusione stretta la dimensione aumenta, e poiché non può salire più di n questa successione prima o poi si stabilizza: esiste quindi un h per cui $V_{\lambda}^h = V_{\lambda}^{h+j}$ per ogni $j > 0$. Chiamiamo questo V_{λ}^h l'*autospazio generalizzato massimale* relativo all'autovalore λ e lo indichiamo con $V_{\lambda}^{\max}$.

Siano $\lambda_1, \dots, \lambda_k$ gli autovalori di T . Sappiamo che se T è diagonalizzabile lo spazio V è somma diretta degli autospazi $V_{\lambda_1}, \dots, V_{\lambda_k}$. Se T non è diagonalizzabile questo non è più vero, ma V è comunque somma diretta dei suoi autospazi generalizzati massimali.

Proposizione 6.1.21. *Lo spazio vettoriale V si decompone in somma diretta negli autospazi generalizzati massimali di T :*

$$V = V_{\lambda_1}^{\max} \oplus \dots \oplus V_{\lambda_k}^{\max}.$$

Dimostrazione. Sia h per cui $V_{\lambda_1}^h = V_{\lambda_1}^{\max}$. Mostriamo che

$$V = \ker(T - \lambda_1 \text{id})^h \oplus \text{Im}(T - \lambda_1 \text{id})^h.$$

La somma delle dimensioni dei due sottospazi è n per il teorema della dimensione, quindi è sufficiente mostrare che sono in somma diretta. Se $v \in \ker(T - \lambda_1 \text{id})^h \cap \text{Im}(T - \lambda_1 \text{id})^h$, allora $v = (T - \lambda_1 \text{id})^h(w)$ per qualche $w \in V$ e quindi $w \in V_{\lambda_1}^{2h}$. Però $V_{\lambda_1}^{2h} = V_{\lambda_1}^h$ e quindi $w \in \ker(T - \lambda_1 \text{id})^h \Rightarrow v = 0$.

I sottospazi $\ker(T - \lambda_1 \text{id})^h$ e $\text{Im}(T - \lambda_1 \text{id})^h$ sono entrambi T -invarianti (esercizio). La restrizione di T a $\text{Im}(T - \lambda_1 \text{id})^h$ non ha l'autovalore λ_1 ,

quindi ha solo gli autovettori $\lambda_2, \dots, \lambda_k$. Si conclude per induzione sul numero di autovettori per T . $\square$

Notiamo che ciascun $V_{\lambda_i}^{\max}$ è T -invariante e la restrizione di T a $V_{\lambda_i}^{\max}$ è un endomorfismo con il solo autovalore λ_i . Abbiamo quindi trovato un modo per decomporre ogni endomorfismo in endomorfismi aventi ciascuno un solo autovalore.

6.1.7. Dimostrazione del Teorema di Jordan. Sia $T: V \rightarrow V$ un endomorfismo. Una *base di Jordan* per T è una base $\mathcal{B}$ tale che la matrice associata $J = [T]_{\mathcal{B}}^{\mathcal{B}}$ sia una matrice di Jordan. Vediamo adesso come determinare una base di Jordan per T .

Possiamo supporre per semplicità che T abbia un solo autovalore λ , perché V si decompone in somma diretta nei $V_{\lambda_i}^{\max}$ e la restrizione di T a ciascun $V_{\lambda_i}^{\max}$ ha un solo autovalore λ_i . Sarà sufficiente trovare una base di Jordan per ciascuna restrizione: la base di Jordan per T sarà l'unione di queste.

Come è fatta una base di Jordan $\mathcal{B}$? A ciascun blocco di Jordan $B_{\lambda, k}$ di J corrispondono k vettori $v_1, \dots, v_k$ di $\mathcal{B}$ con la proprietà che

$$T(v_1) = \lambda v_1, \quad T(v_2) = \lambda v_2 + v_1, \quad \dots \quad T(v_k) = \lambda v_k + v_{k-1}.$$

È più utile scrivere le uguaglianze in questo modo:

$$(T - \lambda \text{id})(v_1) = 0, \quad (T - \lambda \text{id})(v_2) = v_1, \quad \dots \quad (T - \lambda \text{id})(v_k) = v_{k-1}.$$

Chiamiamo k vettori $v_1, \dots, v_k$ fatti a questo modo una *catena* di vettori. Possiamo visualizzarla così:

$$0 \xleftarrow{T - \lambda \text{id}} v_1 \xleftarrow{T - \lambda \text{id}} v_2 \xleftarrow{T - \lambda \text{id}} \dots \xleftarrow{T - \lambda \text{id}} v_k.$$

Una base di Jordan per T è precisamente una base che è unione disgiunta di alcune catene di vettori (una per ogni blocco di Jordan).

Teorema 6.1.22. *Ogni endomorfismo $T: V \rightarrow V$ ha una base di Jordan.*

Dimostrazione. Come abbiamo detto, possiamo supporre che T abbia un autovalore solo λ . Procediamo per induzione su $n = \dim V$. Se $n = 1$, qualsiasi generatore di V è una base di Jordan. Consideriamo il caso generico n dando per buono i casi $\leq n - 1$.

L'autospazio $V_{\lambda} = \ker(T - \lambda \text{id})$ ha una certa dimensione $m_g(\lambda) \geq 1$. Per il teorema della dimensione, l'immagine $U = \text{Im}(T - \lambda \text{id})$ ha dimensione $n - m_g(\lambda) \leq n - 1$ ed è T -invariante (esercizio). Possiamo quindi considerare l'endomorfismo $T|_U$ ristretto ad U . Siccome $\dim U \leq n - 1$, per l'ipotesi induttiva esiste una base di Jordan per $T|_U$. Questa base è unione di h catene

$$v_1^1, \dots, v_{k_1}^1, \quad \dots, \quad v_1^h, \dots, v_{k_h}^h.$$

Poiché questi vettori stanno tutti in $U = \text{Im}(T - \lambda \text{id})$, per ciascun $i = 1, \dots, h$ esiste un $v_{k_i+1}^i \in V$ tale che

$$(T - \lambda \text{id})(v_{k_i+1}^i) = v_{k_i}^i.$$

Aggiungiamo $v_{k_i+1}^i$ alla fine della i -esima catena, per ogni i . In questo modo abbiamo aggiunto h nuovi vettori alla base. Oltre a questi, ne aggiungiamo altri. Consideriamo i primi vettori delle catene $v_1^1, \dots, v_1^h$. Questi sono vettori indipendenti contenuti in $\ker(T - \lambda \text{id})$. Completiamoli a base di $\ker(T - \lambda \text{id})$ aggiungendo altri vettori $v_1^{h+1}, \dots, v_1^{h+s}$. Ciascuno di questi è una nuova catena formata da un solo vettore.

A questo punto dobbiamo solo dimostrare che i vettori così ottenuti

$$v_1^1, \dots, v_{k_1+1}^1, \quad \dots, \quad v_1^h, \dots, v_{k_h+1}^h, \quad v_1^{h+1}, \quad \dots, \quad v_1^{h+s}$$

formano una base di V . Mostriamo che sono linearmente indipendenti. Supponiamo di avere una combinazione lineare che li annulla:

$$\sum_{i=1}^h (\alpha_1^i v_1^i + \dots + \alpha_{k_i+1}^i v_{k_i+1}^i) + \alpha_1^{h+1} v_1^{h+1} + \dots + \alpha_1^{h+s} v_1^{h+s} = 0.$$

Applicando $T - \lambda \text{id}$ otteniamo

$$\sum_{i=1}^h (\alpha_2^i v_1^i + \dots + \alpha_{k_i+1}^i v_{k_i}^i) = 0.$$

Questi vettori sono una base di U e quindi $\alpha_j^i = 0 \forall j \geq 2$. Rimaniamo con

$$\alpha_1^1 v_1^1 + \dots + \alpha_1^{h+s} v_1^{h+s} = 0.$$

Questi vettori sono una base di $\ker(T - \lambda \text{id})$ e quindi anche questi coefficienti sono tutti zero.

Per mostrare che i vettori costruiti generano V , mostriamo che sono esattamente $n = \dim V$. Dimostriamo cioè che $k_1 + \dots + k_h + h + s = n$. Sappiamo che $\dim \text{Im}(T - \lambda \text{id}) = k_1 + \dots + k_h$ e $\dim \ker(T - \lambda \text{id}) = h + s$, quindi questo segue dal teorema della dimensione. $\square$

Il Teorema 6.1.3 segue dal teorema appena dimostrato. L'esistenza di una forma di Jordan per qualsiasi matrice $A \in M(n, \mathbb{C})$ si ottiene applicando il Teorema 6.1.22 all'endomorfismo $L_A: \mathbb{C}^n \rightarrow \mathbb{C}^n$. L'unicità a meno di permutare i blocchi segue dal fatto che il numero di blocchi di ciascuna taglia è determinato dalle molteplicità geometriche generalizzate, come avevamo già visto nella Proposizione 6.1.16.

6.1.8. Come determinare una base di Jordan. Come troviamo una base di Jordan $\mathcal{B}$ per un dato endomorfismo $T: V \rightarrow V$? L'algoritmo non è semplicissimo: seguiamo la dimostrazione del Teorema 6.1.22.

Come prima cosa determiniamo gli autovalori di T e gli autospazi generalizzati massimali $V_\lambda^{\max}$ per ciascun autovalore λ . A questo scopo notiamo che la successione di autospazi generalizzati

$$V_\lambda^1 \subset V_\lambda^2 \subset V_\lambda^3 \subset \dots$$

appena si stabilizza, lo fa per sempre: intendiamo che se $V_\lambda^i = V_\lambda^{i+1}$ per un certo i allora $V_\lambda^i = V_\lambda^{\max}$ è l'autospazio generalizzato massimale (questo fatto segue dal Teorema di Jordan: esercizio).

L'endomorfismo T si spezza in endomorfismi su questi autospaazi generalizzati che hanno ciascuno un solo autovalore. Lavoriamo separatamente su ciascuna restrizione.

Ci siamo ricondotti al caso in cui T abbia un solo autovalore λ . Determiniamo la successione di sottospazi T -invarianti:

$$V \supsetneq \operatorname{Im}(T - \lambda \operatorname{id}) \supsetneq \cdots \supsetneq \operatorname{Im}(T - \lambda \operatorname{id})^h = \{0\}.$$

Costruiamo iterativamente (da destra a sinistra) una successione

$$\mathcal{B}_0 \supsetneq \mathcal{B}_1 \supsetneq \cdots \supsetneq \mathcal{B}_{h-1} \supsetneq \mathcal{B}_h = \emptyset$$

dove $\mathcal{B}_j$ è una base di Jordan della restrizione di T al sottospazio $\operatorname{Im}(T - \lambda \operatorname{id})^j$. La base finale $\mathcal{B}_0$ è quindi una base di Jordan per T .

La fabbricazione delle basi funziona iterativamente da destra a sinistra nel modo seguente. Costruiamo $\mathcal{B}_{j-1}$ aggiungendo i seguenti vettori a $\mathcal{B}_j$:

- (1) Per ogni vettore $v \in \mathcal{B}_j \setminus \mathcal{B}_{j+1}$, prendiamo un $w \in \operatorname{Im}(T - \lambda \operatorname{id})^{j-1}$ tale che $(T - \lambda \operatorname{id})(w) = v$ e lo aggiungiamo.
- (2) Se i vettori in $\mathcal{B}_{j-1}$ trovati finora generano $\operatorname{Im}(T - \lambda \operatorname{id})^{j-1}$ siamo a posto. Altrimenti li completiamo a base di $\operatorname{Im}(T - \lambda \operatorname{id})^{j-1}$ usando solo vettori di $\ker(T - \lambda \operatorname{id})$. È sempre possibile fare ciò grazie al teorema della dimensione.

Questo è lo schema:

0	$\xleftarrow{T-\lambda \operatorname{id}}$	v_1^1	$\xleftarrow{T-\lambda \operatorname{id}}$	v_2^1	$\xleftarrow{T-\lambda \operatorname{id}}$	v_3^1	$\xleftarrow{T-\lambda \operatorname{id}}$	$\dots$
$\vdots$		$\vdots$		$\vdots$		$\vdots$		
0	$\xleftarrow{T-\lambda \operatorname{id}}$	v_1^a	$\xleftarrow{T-\lambda \operatorname{id}}$	v_2^a	$\xleftarrow{T-\lambda \operatorname{id}}$	v_3^a	$\xleftarrow{T-\lambda \operatorname{id}}$	$\dots$
		0	$\xleftarrow{T-\lambda \operatorname{id}}$	v_1^{a+1}	$\xleftarrow{T-\lambda \operatorname{id}}$	v_2^{a+1}	$\xleftarrow{T-\lambda \operatorname{id}}$	$\dots$
		$\vdots$		$\vdots$		$\vdots$		
		0	$\xleftarrow{T-\lambda \operatorname{id}}$	v_1^b	$\xleftarrow{T-\lambda \operatorname{id}}$	v_2^b	$\xleftarrow{T-\lambda \operatorname{id}}$	$\dots$
				0	$\xleftarrow{T-\lambda \operatorname{id}}$	v_1^{b+1}	$\xleftarrow{T-\lambda \operatorname{id}}$	$\dots$
				$\vdots$		$\vdots$		
				0	$\xleftarrow{T-\lambda \operatorname{id}}$	v_1^c	$\xleftarrow{T-\lambda \operatorname{id}}$	$\dots$
						0	$\xleftarrow{T-\lambda \operatorname{id}}$	$\dots$
						$\vdots$		$\ddots$
		$\downarrow$		$\downarrow$		$\downarrow$		
		$\mathcal{B}_{h-1}$		$\mathcal{B}_{h-2} \setminus \mathcal{B}_{h-1}$		$\mathcal{B}_{h-3} \setminus \mathcal{B}_{h-2}$		$\dots$

Esempio 6.1.23. Cerchiamo una base di Jordan per $L_A: \mathbb{C}^4 \rightarrow \mathbb{C}^4$ dove A è la matrice dell'Esempio 6.1.19:

$$A = \begin{pmatrix} 2 & 2 & 0 & -1 \\ 0 & 0 & 0 & 1 \\ 1 & 5 & 2 & -1 \\ 0 & -4 & 0 & 4 \end{pmatrix}.$$

La matrice ha come unico autovalore 2. Notiamo che

$$A - 2I = \begin{pmatrix} 0 & 2 & 0 & -1 \\ 0 & -2 & 0 & 1 \\ 1 & 5 & 0 & -1 \\ 0 & -4 & 0 & 2 \end{pmatrix}, \quad (A - 2I)^2 = \begin{pmatrix} 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 \\ 0 & -4 & 0 & 2 \\ 0 & 0 & 0 & 0 \end{pmatrix}.$$

e $(A - 2I)^3 = 0$. Abbiamo

$$\mathbb{C}^4 \supsetneq \text{Im}(A - 2I) \supsetneq \text{Im}(A - 2I)^2 \supsetneq \text{Im}(A - 2I)^3 = \{0\}.$$

e costruiamo partendo da destra la successione di basi:

$$\mathcal{B}_0 \supsetneq \mathcal{B}_1 \supsetneq \mathcal{B}_2 \supsetneq \mathcal{B}_3 = \emptyset.$$

Possiamo prendere ad esempio:

$$\mathcal{B}_2 = \left\{ v_1^1 = \begin{pmatrix} 0 \\ 0 \\ 1 \\ 0 \end{pmatrix} \right\}, \quad \mathcal{B}_1 = \left\{ v_1^1 = \begin{pmatrix} 0 \\ 0 \\ 1 \\ 0 \end{pmatrix}, v_2^1 = \frac{1}{2} \begin{pmatrix} -1 \\ 1 \\ 0 \\ 2 \end{pmatrix} \right\}.$$

Notiamo che effettivamente $(A - 2I)(v_1^1) = 0$ e $(A - 2I)(v_2^1) = v_1^1$. Infine

$$\mathcal{B}_0 = \left\{ v_1^1 = \begin{pmatrix} 0 \\ 0 \\ 1 \\ 0 \end{pmatrix}, v_2^1 = \frac{1}{2} \begin{pmatrix} -1 \\ 1 \\ 0 \\ 2 \end{pmatrix}, v_3^1 = \frac{1}{2} \begin{pmatrix} 1 \\ 0 \\ 0 \\ 1 \end{pmatrix}, v_4^1 = \begin{pmatrix} -3 \\ 1 \\ 0 \\ 2 \end{pmatrix} \right\}.$$

Effettivamente $(A - 2I)(v_3^1) = v_2^1$ e $(A - 2I)(v_4^1) = 0$. Questa è una base di Jordan per L_A e quindi

$$[L_A]_{\mathcal{B}_0}^{\mathcal{B}_0} = \begin{pmatrix} \boxed{\begin{pmatrix} 2 & 1 & 0 \\ 0 & 2 & 1 \\ 0 & 0 & 2 \end{pmatrix}} & \begin{pmatrix} 0 \\ 0 \\ 0 \end{pmatrix} \\ \begin{pmatrix} 0 & 0 & 0 \end{pmatrix} & \boxed{2} \end{pmatrix}.$$

6.2. Teorema di Cayley – Hamilton

Le matrici quadrate possono essere sommate e moltiplicate fra loro. Se A è una matrice quadrata, ha senso scrivere espressioni del tipo $A^4 - A^3 - 2A$. In altre parole, ha senso prendere un polinomio $x^4 - x^3 - 2x$ e quindi sostituire x con una matrice A . Come per le radici di un polinomio, ci possiamo chiedere quando $A^4 - A^3 - 2A = 0$.

In questa sezione studiamo come i polinomi si combinano con le matrici.

6.2.1. Polinomi di matrici. Prendiamo una matrice $A \in M(n, \mathbb{C})$ e un polinomio $p(x) \in \mathbb{C}[x]$ di grado k , entrambi a coefficienti complessi. Abbiamo

$$p(x) = a_k x^k + \dots + a_1 x + a_0.$$

Definiamo adesso una nuova matrice $p(A)$ nel modo seguente:

$$p(A) = a_k A^k + \dots + a_1 A + a_0 I_n.$$

Diciamo che il polinomio $p(x)$ *annulla* la matrice A se¹

$$p(A) = 0.$$

In questa equazione “0” è la matrice $n \times n$ nulla. Vogliamo studiare come sono fatti i polinomi che annullano una matrice data A . Ovviamente, fra questi c’è sempre il *polinomio banale*, cioè la costante 0, che annulla qualsiasi cosa. Prima domanda: c’è dell’altro? La risposta è affermativa:

Proposizione 6.2.1. *Esistono polinomi non banali che annullano A .*

Dimostrazione. Consideriamo le matrici

$$I_n, A, A^2, \dots, A^{n^2}.$$

Queste sono $n^2 + 1$ elementi nello spazio $M(n, \mathbb{C})$ di tutte le matrici $n \times n$. Siccome questo spazio ha dimensione n^2 , questi $n^2 + 1$ elementi non possono essere indipendenti: c’è quindi una combinazione non banale che li annulla

$$a_0 I + a_1 A + a_2 A^2 + \dots + a_{n^2} A^{n^2} = 0$$

con $a_i \in \mathbb{C}$. In altre parole, se definiamo

$$p(x) = a_{n^2} x^{n^2} + \dots + a_1 x + a_0$$

otteniamo $p(A) = 0$. □

La dimostrazione fornisce un polinomio $p(x)$ non banale di grado al massimo n^2 che annulli A . Vogliamo adesso migliorare questo risultato, mostrando che il polinomio caratteristico di A , che ha grado solo n , annulla A . Per ottenere questo profondo risultato noto come *Teorema di Cayley – Hamilton* abbiamo bisogno di alcuni fatti preliminari.

6.2.2. Matrici diagonali a blocchi. In questo capitolo, una *matrice diagonale a blocchi* è una matrice del tipo

$$A = \begin{pmatrix} B_1 & 0 & \dots & 0 \\ 0 & B_2 & \dots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \dots & B_k \end{pmatrix}$$

¹Quando si studiano i polinomi, generalmente se ne fissa uno $p(x)$ e si cercano quei numeri x_0 reali o complessi (chiamati *radici*) per cui $p(x_0) = 0$. Qui la prospettiva è opposta: si fissa una matrice A e si cercano i polinomi $p(x)$ per cui $p(A) = 0$.

dove $B_1, \dots, B_k$ sono matrici quadrate, di taglia arbitraria. A parte i blocchi $B_1, \dots, B_k$ che stanno sulla diagonale, tutti gli altri blocchi della matrice sono nulli. Un esempio importante di matrice diagonale a blocchi è ovviamente una matrice di Jordan.

Proposizione 6.2.2. *Se A è una matrice diagonale a blocchi e $p(x)$ è un polinomio, allora otteniamo*

$$p(A) = \begin{pmatrix} p(B_1) & 0 & \dots & 0 \\ 0 & p(B_2) & \dots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \dots & p(B_k) \end{pmatrix}.$$

In altre parole, per applicare un polinomio p su A basta farlo su ogni blocco.

Dimostrazione. Sia

$$p(x) = a_m x^m + \dots + a_1 x + a_0.$$

Otteniamo

$$p(A) = a_m A^m + \dots + a_1 A + a_0 I_n$$

che può essere scritto come

$$a_m \begin{pmatrix} B_1 & 0 & \dots & 0 \\ 0 & B_2 & \dots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \dots & B_k \end{pmatrix}^m + \dots + a_0 \begin{pmatrix} I & 0 & \dots & 0 \\ 0 & I & \dots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \dots & I \end{pmatrix}$$

dove le I nell'ultima matrice sono tutte matrici identità della taglia opportuna. Questo è equivalente a

$$\begin{pmatrix} a_m B_1^m & 0 & \dots & 0 \\ 0 & a_m B_2^m & \dots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \dots & a_m B_k^m \end{pmatrix} + \dots + \begin{pmatrix} a_0 I & 0 & \dots & 0 \\ 0 & a_0 I & \dots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \dots & a_0 I \end{pmatrix}$$

e sommando le matrici otteniamo

$$\begin{pmatrix} p(B_1) & 0 & \dots & 0 \\ 0 & p(B_2) & \dots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \dots & p(B_k) \end{pmatrix}.$$

La dimostrazione è completa. □

6.2.3. Matrici simili. Come sappiamo, due matrici simili hanno molte cose in comune. La proposizione seguente conferma questo fatto generale.

Proposizione 6.2.3. *Se A e B sono simili e $p(x)$ è un polinomio, allora*

$$p(A) = 0 \iff p(B) = 0.$$

Dimostrazione. Sia

$$p(x) = a_m x^m + \dots + a_1 x + a_0.$$

Per ipotesi $A = MBM^{-1}$ per qualche matrice invertibile M . Ricordiamo che

$$A^k = (MBM^{-1})^k = MB^k M^{-1}.$$

Otteniamo

$$\begin{aligned} p(A) &= a_m (MBM^{-1})^m + \dots + a_1 MBM^{-1} + a_0 I \\ &= Ma_m B^m M^{-1} + \dots + Ma_1 B M^{-1} + Ma_0 I M^{-1} \\ &= Mp(B)M^{-1}. \end{aligned}$$

Nell'ultima uguaglianza abbiamo raccolto M a sinistra e M^{-1} a destra. Segue che $p(B) = 0 \implies p(A) = 0$. Invertendo i ruoli di A e B si ottiene anche la freccia opposta. $\square$

Abbiamo scoperto che due matrici simili sono annullate precisamente dagli stessi polinomi.

6.2.4. Teorema di Cayley – Hamilton. Abbiamo infine tutti gli strumenti per enunciare e dimostrare il seguente teorema.

Teorema 6.2.4 (Cayley – Hamilton). *Il polinomio caratteristico $p_A(\lambda)$ di A annulla A . In altre parole, vale*

$$p_A(A) = 0.$$

Dimostrazione. Consideriamo prima il caso in cui $A = B_{\lambda_0, n}$ sia un blocco di Jordan con autovalore λ_0 . Otteniamo $p_A(\lambda) = (\lambda_0 - \lambda)^n$. Quindi

$$p_A(A) = (\lambda_0 I_n - A)^n = \pm(A - \lambda_0 I_n)^n = \pm B_{0, n}^n = 0.$$

Ricordiamo infatti che $B_{\lambda_0, n} - \lambda_0 I_n = B_{0, n}$ è nilpotente di indice n . Il teorema in questo caso è dimostrato.

Passiamo ora al caso in cui $A = J$ sia una matrice di Jordan, con blocchi di Jordan $B_1, \dots, B_k$. Per il Corollario 5.1.35 abbiamo

$$p_J(x) = p_{B_1}(x) \cdots p_{B_k}(x).$$

Per il caso precedente, sappiamo che $p_{B_i}(B_i) = 0$ e quindi anche $p_J(B_i) = 0$, perché $p_{B_i}(x)$ divide $p_J(x)$. A questo punto la Proposizione 6.2.2 ci dà

$$p_J(J) = \begin{pmatrix} p_J(B_1) & 0 & \dots & 0 \\ 0 & p_J(B_2) & \dots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \dots & p_J(B_k) \end{pmatrix} = 0.$$

Il teorema è quindi verificato anche in questo caso.

Per il Teorema 6.1.3 una generica matrice A è sempre simile ad una matrice di Jordan J . Queste hanno lo stesso polinomio caratteristico $p_A(x) = p_J(x)$, e se $p_J(J) = p_A(J) = 0$ allora anche $p_A(A) = 0$ per la Proposizione 6.2.3. La dimostrazione è conclusa. $\square$

Esempio 6.2.5. Esaminiamo la matrice

$$A = \begin{pmatrix} 0 & -1 \\ 1 & 0 \end{pmatrix}.$$

Sappiamo che $p_A(\lambda) = \lambda^2 + 1$. Per il Teorema di Cayley – Hamilton

$$p_A(A) = A^2 + I = 0.$$

Effettivamente si verifica che $A^2 = -I$.

Esempio 6.2.6. Esaminiamo la matrice

$$A = \begin{pmatrix} 1 & 0 \\ 0 & 0 \end{pmatrix}.$$

Sappiamo che $p_A(\lambda) = \lambda^2 - \lambda$. Per il Teorema di Cayley – Hamilton abbiamo

$$p_A(A) = A^2 - A = 0.$$

Effettivamente si verifica che $A^2 = A$.

Esempio 6.2.7. Più in generale, il polinomio caratteristico di una matrice $A \in M(2, \mathbb{C})$ qualsiasi è $\lambda^2 - \text{tr}A \lambda + \det A$. Quindi otteniamo

$$A^2 = \text{tr}A \cdot A - \det A \cdot I_2.$$

6.3. Polinomio minimo

Il Teorema di Cayley – Hamilton dice che il polinomio caratteristico $p_A(\lambda)$ annulla la matrice A . Se A è una matrice $n \times n$, allora p_A ha grado n . Ci poniamo adesso la domanda seguente: esistono polinomi di grado minore di n che annullano A ? Questa domanda ci porta alla definizione del *polinomio minimo* di una matrice quadrata.

6.3.1. Polinomio minimo. Ricordiamo che un polinomio $p(x)$ è *monico* se il termine di grado più alto di $p(x)$ ha coefficiente 1. In altre parole, un polinomio $p(x)$ di grado k è monico se è del tipo

$$p(x) = x^k + a_{k-1}x^{k-1} + \dots + a_1x + a_0.$$

Notiamo che qualsiasi polinomio non banale può essere trasformato in un polinomio monico dividendolo per il coefficiente a_k del termine di grado più alto. Sia ora $A \in M(n, \mathbb{C})$ una matrice quadrata a coefficienti complessi.

Definizione 6.3.1. Il *polinomio minimo* di A è il polinomio monico $m_A(x)$ di grado minore che annulla A .

Proposizione 6.3.2. *La definizione è ben posta.*

Dimostrazione. Nella definizione abbiamo dato per scontato che ci sia un unico polinomio $m_A(x)$ con quelle proprietà, e adesso dobbiamo dimostrarlo. Supponiamo di avere due polinomi $m_1(x)$ e $m_2(x)$ monici distinti, entrambi di grado minimo k che annullano A . Siccome sono entrambi monici, la differenza $q(x) = m_1(x) - m_2(x)$ è un polinomio non banale di grado strettamente minore di k . Otteniamo

$$q(A) = m_1(A) - m_2(A) = 0 - 0 = 0.$$

Quindi $q(x)$ è un polinomio non banale di grado $< k$ che annulla A , ma questo è assurdo. $\square$

Esempio 6.3.3. Se $A = 0$ è la matrice nulla, allora il suo polinomio minimo è $m_A(x) = x$. Infatti abbiamo

$$m_A(A) = A = 0$$

e quindi $m_A(x)$ annulla A . Questo è sicuramente il polinomio monico di grado minimo fra quelli che annullano A perché ha grado uno.

Se $A = I_n$ allora il suo polinomio minimo è $m_A(x) = x - 1$. Infatti

$$m_A(A) = A - I_n = I_n - I_n = 0$$

e concludiamo come nel caso precedente. Più in generale, le uniche matrici A aventi polinomio minimo di grado uno, cioè con $m_A(x) = x - c$ per qualche $c \in \mathbb{C}$, sono le matrici del tipo $A = cI_n$.

6.3.2. Il polinomio minimo divide tutti gli altri. Sia $A \in M(n, \mathbb{C})$ una matrice quadrata $n \times n$ a coefficienti complessi. Mostriamo un'importante proprietà algebrica del polinomio minimo.

Proposizione 6.3.4. *Il polinomio minimo $m_A(x)$ divide qualsiasi polinomio in $\mathbb{C}[x]$ che annulli la matrice A .*

Dimostrazione. Sia $a(x)$ un polinomio che annulla A . Dividendo $a(x)$ per $m_A(x)$ otteniamo

$$a(x) = m_A(x)q(x) + r(x)$$

per qualche resto $r(x)$ di grado minore di quello di $m_A(x)$. Ne segue che

$$r(A) = a(A) - m_A(A)q(A) = 0 - 0q(A) = 0$$

e allora anche $r(x)$ annulla A . Poiché $r(x)$ ha grado minore di $m_A(x)$, il polinomio $r(x)$ deve essere nullo. Quindi $m_A(x)$ divide $a(x)$. $\square$

Corollario 6.3.5. *Il polinomio minimo $m_A(x)$ divide sempre il polinomio caratteristico $p_A(x)$.*

Dimostrazione. Per il Teorema di Cayley – Hamilton, il polinomio caratteristico $p_A(x)$ annulla A . $\square$

Esempio 6.3.6. Sapendo che $m_A(x)$ divide $p_A(x)$, è possibile determinare $m_A(x)$ per tentativi. Ad esempio, consideriamo

$$A = \begin{pmatrix} 1 & 1 \\ 0 & 1 \end{pmatrix}.$$

Vediamo che il polinomio caratteristico è $p_A(x) = (x-1)^2$. Siccome $m_A(x)$ divide $p_A(x)$, il polinomio $m_A(x)$ è uno dei due polinomi seguenti:

$$x - 1, \quad (x - 1)^2.$$

Procediamo per tentativi, iniziando con il primo: se $q(x) = x - 1$, otteniamo $q(A) = A - I_2 \neq 0$. Quindi il primo non va bene. Allora $m_A(x) = (x - 1)^2$ è il secondo, per esclusione. La lettrice può verificare che infatti $(A - I_2)^2 = 0$.

Esempio 6.3.7. Calcoliamo il polinomio minimo della matrice

$$A = \begin{pmatrix} i & 1 & 0 \\ 0 & i & 1 \\ 0 & 0 & i \end{pmatrix}.$$

Troviamo $p_A(x) = -(x - i)^3$. Quindi il polinomio minimo è uno dei seguenti:

$$x - i, \quad (x - i)^2, \quad (x - i)^3.$$

Se $q(x) = x - i$, otteniamo $q(A) = A - iI_3 \neq 0$, quindi il primo non va bene. Se $q(x) = (x - i)^2$, si verifica che $q(A) = (A - iI_3)^2 \neq 0$, quindi neanche il secondo va bene. Quindi $m_A(x) = (x - i)^3$.

6.3.3. Blocchi di Jordan. Calcoliamo il polinomio minimo di un blocco di Jordan. Ricordiamo che $B_{\lambda,n}$ è il blocco di Jordan $n \times n$ con autovalore λ .

Proposizione 6.3.8. *Il polinomio minimo di $B = B_{\lambda_0,n}$ è uguale al suo polinomio caratteristico:*

$$m_B(\lambda) = p_B(\lambda) = (\lambda - \lambda_0)^n.$$

Dimostrazione. Sappiamo che $m_B(\lambda)$ divide $p_B(\lambda)$ e quindi $m_B(\lambda) = (\lambda - \lambda_0)^k$ per qualche $1 \leq k \leq n$. Dobbiamo determinare il k giusto. Se prendiamo $q(\lambda) = (\lambda - \lambda_0)^k$, otteniamo

$$q(B) = (B - \lambda_0 I_n)^k = B_{0,n}^k$$

dove $B_{0,n}$ è la matrice nilpotente considerata nella Proposizione 6.1.9, che sappiamo già avere ordine di nilpotenza n . Quindi $q(B) = 0$ solo se $k = n$ e ne ricaviamo che $m_A(\lambda) = p_A(\lambda) = (\lambda - \lambda_0)^n$. $\square$

6.3.4. Matrici diagonali a blocchi. Consideriamo una matrice diagonale a blocchi

$$A = \begin{pmatrix} B & 0 \\ 0 & C \end{pmatrix}.$$

Per i polinomi caratteristici, sappiamo dalla Proposizione 5.1.35 che

$$p_A(x) = p_B(x)p_C(x).$$

Il polinomio minimo si comporta diversamente.

Definizione 6.3.9. Dati due polinomi monici $a(x)$ e $b(x)$, il *minimo comune multiplo* di $a(x)$ e $b(x)$ è il più piccolo polinomio monico $m(x)$ che è diviso sia da $a(x)$ che da $b(x)$.

Esercizio 6.3.10. La definizione è ben posta.

Esempio 6.3.11. Il minimo comune multiplo fra $(x-1)(x-2)$ e $(x-1)^2$ è $(x-1)^2(x-2)$.

Esercizio 6.3.12. Se A è una matrice diagonale a blocchi come sopra, allora $m_A(x)$ è il minimo comune multiplo fra $m_B(x)$ e $m_C(x)$.

6.3.5. Da Jordan a polinomio minimo. Conoscendo già la forma di Jordan, è molto facile determinare il polinomio minimo.

Sia come sempre $A \in M(n, \mathbb{C})$ una matrice quadrata. Siano $\lambda_1, \dots, \lambda_k$ gli autovalori di A . Il polinomio caratteristico di A è

$$p_A(\lambda) = \pm(\lambda - \lambda_1)^{m_1} \cdots (\lambda - \lambda_k)^{m_k}$$

dove $m_i = m_a(\lambda_i)$ è la molteplicità algebrica di λ_i . Il polinomio minimo $m_A(\lambda)$ di A divide $p_A(\lambda)$ e quindi è del tipo

$$m_A(\lambda) = (\lambda - \lambda_1)^{d_1} \cdots (\lambda - \lambda_k)^{d_k}$$

dove $d_1, \dots, d_k$ sono dei numeri da determinare. Per adesso sappiamo solo che $d_i \leq m_i$ per ogni i . Sia J la matrice di Jordan associata a A .

Proposizione 6.3.13. Il numero d_i è la massima taglia dei blocchi di Jordan in J con autovalore λ_i .

Dimostrazione. Per un singolo blocco di Jordan $B = B_{\lambda_0, n}$ otteniamo $m_B(\lambda) = (\lambda - \lambda_0)^n$ ed n è effettivamente la taglia di B , quindi ci siamo. Se J è fatto di vari blocchi di Jordan, usando più volte l'Esercizio 6.3.12 troviamo che $m_J(\lambda)$ è il minimo comune multiplo di tutti i polinomi $(\lambda - \lambda_i)^{d_i}$

al variare di tutti i blocchi di Jordan con autovalore λ_i e di taglia variabile d . Quindi otteniamo proprio

$$m_A(\lambda) = (\lambda - \lambda_1)^{d_1} \cdots (\lambda - \lambda_k)^{d_k}$$

dove effettivamente d_i è la massima taglia dei blocchi con autovalore λ_i . $\square$

In particolare, troviamo che $d_i \geq 1$. Quindi per ogni $i = 1, \dots, k$

$$1 \leq d_i \leq m_i = m_a(\lambda_i).$$

Osservazione 6.3.14. Le stesse disuguaglianze $1 \leq m_g(\lambda_i) \leq m_a(\lambda_i)$ valgono per la molteplicità geometrica $m_g(\lambda_i)$, ma attenzione a non fare confusione! Qui d_i è la massima taglia dei blocchi con autovalore λ_i , mentre $m_g(\lambda_i)$ è il numero di questi blocchi. Queste due quantità non sono correlate e sono spesso differenti.

Esempio 6.3.15. Il polinomio minimo di

$$\begin{pmatrix} \boxed{2} & 0 & 0 & 0 & 0 \\ 0 & \boxed{2} & 0 & 0 & 0 \\ 0 & 0 & \boxed{3} & 1 & 0 \\ 0 & 0 & 0 & 3 & 1 \\ 0 & 0 & 0 & 0 & 3 \end{pmatrix}$$

è $(x - 2)(x - 3)^3$.

Enunciamo un importante corollario.

Corollario 6.3.16. *Una matrice A è diagonalizzabile se e solo se il polinomio minimo non ha radici ripetute, cioè è del tipo*

$$m_A(\lambda) = (\lambda - \lambda_1) \cdots (\lambda - \lambda_n)$$

dove $\lambda_1, \dots, \lambda_n$ sono gli autovalori (distinti) di A .

Dimostrazione. Una matrice A è diagonalizzabile se e solo se la sua forma di Jordan è diagonale, e questo accade precisamente quando tutti i blocchi hanno taglia 1. $\square$

6.3.6. Applicazioni. Il Corollario 6.3.16 fornisce un criterio molto potente di diagonalizzabilità. Notiamo una differenza fra polinomio caratteristico e minimo: il primo si calcola subito, ma non ci dice se A sia diagonalizzabile; il secondo è più difficile da calcolare, ma ci dice subito se A sia diagonalizzabile oppure no. Mostriamo qualche applicazione di questo criterio.

Proposizione 6.3.17. *Se $A \in M(n, \mathbb{C})$ è tale che $A^2 = A$, allora A è simile ad una matrice diagonale con solo valori 1 e 0 sulla diagonale.*

Dimostrazione. Per ipotesi $A^2 - A = 0$ e quindi il polinomio $q(x) = x^2 - x = x(x-1)$ annulla A . Il polinomio minimo $m_A(x)$ deve dividere $q(x)$ ed è quindi uno di questi:

$$x(x-1), \quad x, \quad x-1.$$

In tutti i casi $m_A(x)$ non ha radici ripetute e quindi A è diagonalizzabile. Ne segue che A è simile ad una matrice diagonale D , con gli autovalori sulla diagonale. Gli autovalori sono le radici di $m_A(x)$ e quindi possono essere solo 0 e 1. $\square$

L'esercizio seguente si risolve in modo analogo.

Esercizio 6.3.18. Se $A \in M(n, \mathbb{C})$ è tale che $A^2 = I$, allora A è simile ad una matrice diagonale con valori 1 e -1 sulla diagonale.

Attenzione però al caso seguente:

Esercizio 6.3.19. Se $A \in M(n, \mathbb{C})$ è tale che $A^2 = 0$, allora A è diagonalizzabile se e solo se $A = 0$.

Esercizi

Esercizio 6.1. Determina le forme di Jordan delle matrici seguenti:

$$\begin{pmatrix} 5 & -4 \\ 9 & -7 \end{pmatrix}, \quad \begin{pmatrix} 4 & 0 & 0 \\ 2 & 1 & 3 \\ 5 & 0 & 4 \end{pmatrix}, \quad \begin{pmatrix} 1 & 0 & 0 \\ 0 & 1 & -1 \\ -1 & 1 & -1 \end{pmatrix},$$

$$\begin{pmatrix} 5 & 4 & 3 \\ -1 & 0 & -3 \\ 1 & -2 & 1 \end{pmatrix}, \quad \begin{pmatrix} 9 & 7 & 3 \\ -9 & -7 & -4 \\ 4 & 4 & 4 \end{pmatrix}, \quad \begin{pmatrix} 2 & 2 & -1 \\ -1 & -1 & 1 \\ -1 & -2 & 2 \end{pmatrix}.$$

Esercizio 6.2. Elenca tutte le forme di Jordan che hanno polinomio caratteristico $(\lambda - 1)^2(\lambda + 2)^3$. Tra queste, determina quella che ha polinomio minimo $(\lambda - 1)(\lambda + 2)^2$.

Esercizio 6.3. Mostra che l'applicazione $f: \mathbb{C}_n[x] \rightarrow \mathbb{C}_n[x]$ che manda un polinomio $p(x)$ nella sua derivata $p'(x)$ è nilpotente di ordine $n+1$. Determina la sua forma di Jordan.

Esercizio 6.4. Mostra che le matrici complesse $n \times n$ nilpotenti di indice di nilpotenza n sono tutte simili fra loro.

Esercizio 6.5. Determina tutte le forme di Jordan 4×4 con un solo autovalore λ_0 e per ciascuna scrivi il suo polinomio minimo (sono 5.)

Esercizio 6.6. Calcola il polinomio minimo delle matrici seguenti:

$$\begin{pmatrix} 1 & 1 & 0 & 0 \\ -1 & -1 & 0 & 0 \\ -2 & -2 & 2 & 1 \\ 1 & 1 & -1 & 0 \end{pmatrix}, \quad \begin{pmatrix} 2 & 2 & 0 & -1 \\ 0 & 0 & 0 & 1 \\ 1 & 5 & 2 & -1 \\ 0 & -4 & 0 & -4 \end{pmatrix}, \quad \begin{pmatrix} 3 & -1 & 3 & -1 \\ 9 & -3 & 9 & -3 \\ 0 & 0 & 3 & -1 \\ 0 & 0 & 9 & -3 \end{pmatrix}$$

Esercizio 6.7. Trova due matrici $A, B \in M(2)$ che abbiano lo stesso polinomio caratteristico ma polinomi minimi diversi.

Esercizio 6.8. Costruisci due matrici complesse A e A' con lo stesso polinomio caratteristico e lo stesso polinomio minimo, che però non sono simili.

Esercizio 6.9. Costruisci due matrici complesse A e A' con lo stesso polinomio caratteristico, lo stesso polinomio minimo, e le stesse molteplicità geometriche di tutti gli autovalori, che però non sono simili.

Esercizio 6.10. Sia A una matrice complessa 7×7 tale che $(A - I)^3 = 0$ e $\text{rk}(A - I)^2 = 2$. Qual è la forma di Jordan di A ?

Esercizio 6.11. Sia A una matrice complessa quadrata tale che $A^n = I$ per qualche numero intero $n \geq 1$. Mostra che A è sempre diagonalizzabile e determina quali possano essere i suoi autovalori.

Esercizio 6.12. Sia A una matrice complessa tale che $A^2 - 2A + I = 0$. Mostra che A è diagonalizzabile se e solo se $A = I$.

Esercizio 6.13. Dimostra l'Esercizio 5.19 usando il polinomio minimo e il Corollario 6.3.16.

Esercizio 6.14. Sia $V = V_1 \oplus \dots \oplus V_k$ uno spazio vettoriale complesso e $T: V \rightarrow V$ un endomorfismo tale che $T(V_i) \subset V_i$ per ogni i . Mostra che T è diagonalizzabile se e solo se ciascuna restrizione $T|_{V_i}$ è diagonalizzabile.

Complementi

6.1. Forma di Jordan reale. In questo capitolo abbiamo considerato soltanto matrici a coefficienti complessi. Cosa possiamo dire a proposito delle matrici a coefficienti reali?

Un primo fatto da sottolineare è che tutta la teoria esposta nel capitolo si applica anche alle matrici reali $A \in M(n, \mathbb{R})$ che abbiano tutti gli autovalori in $\mathbb{R}$. Queste matrici hanno un'unica forma di Jordan e una base di Jordan in $\mathbb{R}^n$, che possono essere determinate con le stesse tecniche esposte. La dimostrazione del Teorema 6.1.22 si applica anche su $\mathbb{R}$, purché il polinomio caratteristico dell'endomorfismo abbia le radici in $\mathbb{R}$.

Per le matrici reali più generali esiste una versione più complicata della forma canonica di Jordan. L'ingrediente fondamentale è la matrice

$$E_{a,b} = \begin{pmatrix} a & -b \\ b & a \end{pmatrix}$$

definita per ogni coppia $a, b \in \mathbb{R}$. Possiamo anche scrivere

$$E_{a,b} = \rho \begin{pmatrix} \cos \theta & -\sin \theta \\ \sin \theta & \cos \theta \end{pmatrix} = \rho \text{Rot}_\theta$$

prendendo

$$\rho = a^2 + b^2, \quad \cos \theta = \frac{a}{a^2 + b^2}, \quad \sin \theta = \frac{b}{a^2 + b^2}.$$

Geometricamente $E_{a,b}$ è la composizione di una rotazione di angolo θ e di una omotetia di centro l'origine di fattore ρ . Supponiamo sempre $b \neq 0$, cioè $\theta \neq 0, \pi$. Il polinomio caratteristico di $E_{a,b}$ è $p(\lambda) = \lambda^2 - 2a\lambda + a^2 + b^2$ e gli autovalori $\lambda = a \pm bi$ non sono reali.

Oltre ai blocchi di Jordan $B_{\lambda,n}$ già definiti per ogni $\lambda \in \mathbb{R}$, dobbiamo considerare dei blocchi aggiuntivi di taglia $2n$ di questo tipo:

$$B_{a,b,n} = \begin{pmatrix} E_{a,b} & I_2 & 0 & \cdots & 0 \\ 0 & E_{a,b} & I_2 & \cdots & 0 \\ 0 & 0 & E_{a,b} & \cdots & 0 \\ \vdots & \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & 0 & \cdots & E_{a,b} \end{pmatrix}.$$

Notiamo che $B_{a,b,n}$ ha polinomio caratteristico $p(\lambda) = (\lambda^2 - 2a\lambda + a^2 + b^2)^n$ e nessun autovalore reale, quindi nessun autovettore. Ad esempio:

$$B_{a,b,2} = \begin{pmatrix} a & -b & 1 & 0 \\ b & a & 0 & 1 \\ 0 & 0 & a & -b \\ 0 & 0 & b & a \end{pmatrix}.$$

Una *matrice di Jordan reale* è una matrice a blocchi del tipo

$$J = \begin{pmatrix} B_1 & 0 & \cdots & 0 \\ 0 & B_2 & \cdots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \cdots & B_k \end{pmatrix}$$

dove $B_1, \dots, B_k$ sono blocchi di Jordan di taglia arbitraria, di entrambi i tipi $B_{\lambda,n}$ oppure $B_{a,b,n}$. Vale il fatto seguente, che enunciamo senza dimostrazione.

Teorema 6.3.20. *Qualsiasi matrice quadrata reale è simile ad una matrice di Jordan reale, unica a meno di permutare i blocchi.*

Queste sono le possibili forme di Jordan reali di una matrice 2×2 :

$$\begin{pmatrix} \lambda_1 & 0 \\ 0 & \lambda_2 \end{pmatrix}, \quad \begin{pmatrix} \lambda_1 & 1 \\ 0 & \lambda_1 \end{pmatrix}, \quad \rho \begin{pmatrix} \cos \theta & -\sin \theta \\ \sin \theta & \cos \theta \end{pmatrix}.$$

Ogni matrice reale 2×2 è simile ad una di queste.

Prodotti scalari

Nei capitoli precedenti abbiamo introdotto due concetti geometrici fondamentali: gli spazi vettoriali e le applicazioni lineari. Da un punto di vista più algebrico, abbiamo notato che entrambe le nozioni sono strettamente correlate con i sistemi di equazioni lineari – cioè di primo grado in più variabili.

Adesso introduciamo un'altra operazione geometrica fondamentale, che avrà un ruolo predominante da qui fino alla fine del libro: il *prodotto scalare*. Questa nozione, ampiamente usata in matematica e in fisica, da un punto di vista più algebrico è correlata con i polinomi e le equazioni quadratiche – cioè di secondo grado in più variabili.

7.1. Introduzione

Un endomorfismo di uno spazio vettoriale V è una funzione che prende come *input* un vettore di V e restituisce come *output* un vettore di V . Un prodotto scalare invece è una funzione che prende *due* vettori di V e restituisce uno scalare in $\mathbb{R}$. Questa funzione deve ovviamente soddisfare degli assiomi.

In tutto questo capitolo – a differenza dei precedenti – supporremo sempre che il campo sia $\mathbb{K} = \mathbb{R}$. I motivi per escludere i campi $\mathbb{C}$ e $\mathbb{Q}$ sono i seguenti: in molti punti sarà importante parlare di numeri *positivi* e *negativi*, cosa che possiamo fare con $\mathbb{R}$ e $\mathbb{Q}$ ma non con $\mathbb{C}$; sarà anche importante fare la radice quadrata di un numero positivo, cosa che possiamo fare con $\mathbb{R}$ ma non con $\mathbb{Q}$.

7.1.1. Definizioni. Sia V uno spazio vettoriale reale.

Definizione 7.1.1. Un *prodotto scalare* su V è una applicazione

$$\begin{aligned} V \times V &\longrightarrow \mathbb{R} \\ (v, w) &\longmapsto \langle v, w \rangle \end{aligned}$$

che soddisfa i seguenti assiomi:

- (1) $\langle v + v', w \rangle = \langle v, w \rangle + \langle v', w \rangle$
- (2) $\langle \lambda v, w \rangle = \lambda \langle v, w \rangle$
- (3) $\langle v, w + w' \rangle = \langle v, w \rangle + \langle v, w' \rangle$
- (4) $\langle v, \lambda w \rangle = \lambda \langle v, w \rangle$

$$(5) \langle v, w \rangle = \langle w, v \rangle$$

per ogni $v, v', w, w' \in V$ e ogni $\lambda \in \mathbb{R}$.

Il prodotto scalare è una applicazione che prende come *input* due vettori $v, w \in V$ e restituisce come *output* uno scalare, che indichiamo con $\langle v, w \rangle$. In alcuni aspetti il prodotto scalare si comporta come il prodotto fra numeri: gli assiomi (1) e (3) della definizione indicano che il prodotto è distributivo rispetto alla somma, e l'assioma (5) indica che è commutativo. L'aggettivo "scalare" sta però ad indicare una differenza importante: il prodotto scalare di due vettori non è un vettore, ma "solo" uno scalare.

Un po' di terminologia: gli assiomi (1)–(4) dicono che il prodotto scalare è *bilineare*, cioè lineare "sia a sinistra che a destra"; l'assioma (5) dice che il prodotto scalare è *simmetrico*. Notiamo che gli assiomi (1), (2) e (5) implicano (3) e (4), infatti:

$$\langle v, w + w' \rangle \stackrel{(5)}{=} \langle w + w', v \rangle \stackrel{(1)}{=} \langle w, v \rangle + \langle w', v \rangle \stackrel{(5)}{=} \langle v, w \rangle + \langle v, w' \rangle,$$

$$\langle v, \lambda w \rangle \stackrel{(5)}{=} \langle \lambda w, v \rangle \stackrel{(2)}{=} \lambda \langle w, v \rangle \stackrel{(5)}{=} \lambda \langle v, w \rangle.$$

Quindi in realtà gli assiomi (1), (2) e (5) sarebbero sufficienti per definire un prodotto scalare. Come prima applicazione degli assiomi, notiamo che $\langle v, 0 \rangle = 0$ per ogni $v \in V$. Infatti scrivendo 0 come $0 + 0$ troviamo che

$$\langle v, 0 \rangle = \langle v, 0 + 0 \rangle = \langle v, 0 \rangle + \langle v, 0 \rangle$$

e quindi $\langle v, 0 \rangle = 0$. In particolare otteniamo $\langle 0, 0 \rangle = 0$.

Quando vogliamo dare un nome al prodotto scalare, lo indichiamo con una lettera $g: V \times V \rightarrow \mathbb{R}$ e scriviamo $g(v, w)$ invece di $\langle v, w \rangle$.

7.1.2. Prodotto scalare degenere e definito positivo. Introduciamo ora due importanti assiomi aggiuntivi.

Definizione 7.1.2. Un prodotto scalare su V è:

- *degenere* se esiste $v \neq 0$ tale che $\langle v, w \rangle = 0$ per ogni $w \in V$;
- *definito positivo* se $\langle v, v \rangle > 0$ per ogni $v \in V$ non nullo.

I due assiomi appena introdotti sono mutualmente esclusivi:

Proposizione 7.1.3. *Se un prodotto scalare è definito positivo, allora non è degenere.*

Dimostrazione. Ragioniamo per assurdo. Se fosse degenere, esisterebbe $v \in V$ non nullo con $\langle v, w \rangle = 0$ per ogni $w \in V$ ed in particolare $\langle v, v \rangle = 0$, contraddicendo l'ipotesi che $\langle v, v \rangle > 0 \forall v \neq 0$. $\square$

7.1.3. Prodotto scalare euclideo. Introduciamo un esempio importante di prodotto scalare.

Definizione 7.1.4. Il *prodotto scalare euclideo* su $\mathbb{R}^n$ è definito come

$$\langle x, y \rangle = {}^t x \cdot y.$$

Nell'espressione ${}^t x \cdot y$ il vettore ${}^t x$ è un vettore riga e $\cdot$ indica il prodotto fra matrici. In altre parole, se

$$x = \begin{pmatrix} x_1 \\ \vdots \\ x_n \end{pmatrix}, \quad y = \begin{pmatrix} y_1 \\ \vdots \\ y_n \end{pmatrix}$$

allora

$$\langle x, y \rangle = (x_1, \dots, x_n) \cdot \begin{pmatrix} y_1 \\ \vdots \\ y_n \end{pmatrix} = x_1 y_1 + \dots + x_n y_n.$$

Ad esempio, il prodotto scalare fra i vettori $\begin{pmatrix} 1 \\ 3 \end{pmatrix}$ e $\begin{pmatrix} -2 \\ 1 \end{pmatrix}$ di $\mathbb{R}^2$ è

$$\begin{pmatrix} 1 & 3 \end{pmatrix} \begin{pmatrix} -2 \\ 1 \end{pmatrix} = 1 \cdot (-2) + 3 \cdot 1 = 1.$$

Proposizione 7.1.5. *Il prodotto scalare euclideo è un prodotto scalare definito positivo su $\mathbb{R}^n$.*

Dimostrazione. Gli assiomi (1)–(4) seguono dalle proprietà del prodotto fra matrici elencate nella Proposizione 3.4.2. Proviamo (1):

$$\langle x + x', y \rangle = {}^t(x + x') \cdot y = ({}^t x + {}^t x') \cdot y = {}^t x \cdot y + {}^t x' \cdot y = \langle x, y \rangle + \langle x', y \rangle.$$

Gli altri assiomi di bilinearità si mostrano in modo simile. Per la simmetria, notiamo che l'espressione $\langle x, y \rangle = x_1 y_1 + \dots + x_n y_n$ è chiaramente simmetrica in x e y . Il prodotto scalare euclideo è definito positivo, perché se $x \neq 0$ almeno una coordinata x_i di x è diversa da zero e allora

$$\langle x, x \rangle = x_1^2 + \dots + x_n^2 > 0.$$

La dimostrazione è completa. □

Il prodotto scalare euclideo su $\mathbb{R}^2$ e $\mathbb{R}^3$ coincide con l'usuale prodotto scalare fra vettori usato in fisica.

7.1.4. Matrici simmetriche. Abbiamo visto nel Capitolo 4 che una matrice quadrata $A \in M(n, \mathbb{R})$ determina un endomorfismo $L_A: \mathbb{R}^n \rightarrow \mathbb{R}^n$, e che tutti gli endomorfismi di $\mathbb{R}^n$ sono di questo tipo.

Analogamente, vediamo ora che una matrice *simmetrica* $S \in M(n, \mathbb{R})$ determina un prodotto scalare g_S su $\mathbb{R}^n$. Dimosteremo in seguito che tutti i prodotti scalari di $\mathbb{R}^n$ sono di questo tipo.

Ricordiamo che una matrice $S \in M(n)$ è *simmetrica* se ${}^t S = S$. Rammentiamo anche la formula

$${}^t(AB) = {}^t B {}^t A$$

che si applica per ogni prodotto di matrici AB . Più in generale, vale

$${}^t(A_1 \cdots A_k) = {}^t A_k \cdots {}^t A_1.$$

Proposizione 7.1.6. *Una matrice simmetrica S definisce un prodotto scalare g_S su $\mathbb{R}^n$ nel modo seguente:*

$$g_S(x, y) = {}^t x \cdot S \cdot y.$$

Dimostrazione. Notiamo innanzitutto che le tre matrici ${}^t x$, S e y sono di taglia $1 \times n$, $n \times n$ e $n \times 1$, quindi il prodotto ${}^t x S y$ si può effettivamente fare e ha come risultato una matrice 1×1 , cioè un numero. Quindi g_S è effettivamente una funzione $g_S: \mathbb{R}^n \times \mathbb{R}^n \rightarrow \mathbb{R}$. Gli assiomi (1)–(4) seguono dalle proprietà del prodotto fra matrici elencate nella Proposizione 3.4.2. Ad esempio l'assioma (1) di prodotto scalare si dimostra così:

$$g_S(x + x', y) = {}^t(x + x') S y = {}^t x S y + {}^t x' S y = g_S(x, y) + g_S(x', y).$$

Gli assiomi (2)–(4) si verificano in modo analogo. La simmetria (5) segue dal fatto che la matrice S è simmetrica:

$$g_S(x, y) = {}^t x S y = {}^t({}^t x S y) = {}^t y {}^t S {}^t({}^t x) = {}^t y S x = g_S(y, x).$$

La seconda uguaglianza è vera perché ${}^t x S y$ è un numero e quindi è uguale al suo trasposto. Nella quarta abbiamo usato che $S = {}^t S$. $\square$

Ogni matrice simmetrica S determina quindi un prodotto scalare g_S su $\mathbb{R}^n$. La notazione come prodotto fra matrici è efficiente perché è molto compatta. A volte è comunque utile sviluppare il prodotto:

Proposizione 7.1.7. *Vale l'uguaglianza*

$$g_S(x, y) = {}^t x S y = \sum_{i,j=1}^n x_i S_{ij} y_j.$$

Dimostrazione. Basta calcolare due volte il prodotto fra matrici:

$${}^t x S y = \sum_{i=1}^n x_i (S y)_i = \sum_{i=1}^n x_i \sum_{j=1}^n S_{ij} y_j = \sum_{i,j=1}^n x_i S_{ij} y_j.$$

La dimostrazione è conclusa. $\square$

Ricordiamo che $e_1, \dots, e_n$ è la base canonica di $\mathbb{R}^n$.

Corollario 7.1.8. *Vale la relazione seguente:*

$$g_S(e_i, e_j) = {}^t e_i S e_j = S_{ij}.$$

Abbiamo scoperto che l'elemento S_{ij} è precisamente il prodotto scalare $g_S(e_i, e_j)$ fra i vettori e_i ed e_j della base canonica di $\mathbb{R}^n$.

Esempio 7.1.9. Se come matrice simmetrica prendiamo la matrice identità, otteniamo il prodotto scalare euclideo:

$$g_{I_n}(x, y) = {}^t x I_n y = {}^t x y = x_1 y_1 + \dots + x_n y_n.$$

Ad esempio su $\mathbb{R}^2$ con $I_2 = \begin{pmatrix} 1 & 0 \\ 0 & 1 \end{pmatrix}$ otteniamo

$$g_{I_2}(x, y) = x_1 y_1 + x_2 y_2.$$

Esempio 7.1.10. La matrice $S = \begin{pmatrix} 2 & 1 \\ 1 & 1 \end{pmatrix}$ definisce il prodotto scalare su $\mathbb{R}^2$:

$$g_S(x, y) = \begin{pmatrix} x_1 & x_2 \end{pmatrix} \begin{pmatrix} 2 & 1 \\ 1 & 1 \end{pmatrix} \begin{pmatrix} y_1 \\ y_2 \end{pmatrix} = 2x_1y_1 + x_1y_2 + x_2y_1 + x_2y_2.$$

7.1.5. Forme quadratiche. Un polinomio $p(x)$ nelle variabili $x_1, \dots, x_n$ è *omogeneo* se tutti i suoi monomi hanno lo stesso grado. Ad esempio,

$$x_1 + x_2 - 3x_3, \quad 2x_1x_2 - x_3^2 + x_1x_3, \quad x_1^3 - x_2x_3^2$$

sono polinomi omogenei rispettivamente di grado 1, 2, 3 nelle variabili x_1, x_2, x_3 . I polinomi omogenei di grado 1 sono stati studiati quando abbiamo introdotto i sistemi lineari. Così come i polinomi omogenei di grado 1 sono collegati con le applicazioni lineari, quelli di grado 2 sono in relazione con i prodotti scalari. Diamo intanto loro un nome.

Definizione 7.1.11. Una *forma quadratica* è un polinomio omogeneo di grado 2 nelle variabili $x_1, \dots, x_n$.

Il collegamento con i prodotti scalari è immediato.

Proposizione 7.1.12. *Ogni forma quadratica si scrive come*

$$g_S(x, x) = \sum_{i,j=1}^n x_i S_{ij} x_j$$

per un'opportuna (e unica) matrice simmetrica S .

Dimostrazione. Una generica forma quadratica si scrive

$$q(x) = \sum_{1 \leq i \leq j \leq n} a_{ij} x_i x_j.$$

La condizione $i \leq j$ è presente per evitare che lo stesso monomio $x_i x_j$ e $x_j x_i$ compaia due volte. Definendo

$$S_{ij} = S_{ji} = \frac{1}{2} a_{ij}$$

per ogni $1 \leq i, j \leq n$, possiamo sostituire $a_{ij} x_i x_j$ con $S_{ij} x_i x_j + S_{ji} x_j x_i$ e ottenere

$$q(x) = \sum_{i,j=1}^n S_{ij} x_i x_j = g_S(x, x).$$

La matrice S è simmetrica. La dimostrazione è conclusa. $\square$

Ogni forma quadratica $q(x)$ è descritta da una matrice simmetrica. Ad esempio, le matrici simmetriche

$$\begin{pmatrix} 1 & -3 \\ -3 & 0 \end{pmatrix}, \quad \begin{pmatrix} 1 & 0 & 0 \\ 0 & 1 & 0 \\ 0 & 0 & -1 \end{pmatrix}, \quad \begin{pmatrix} 0 & 1 & 1 \\ 1 & 0 & 1 \\ 1 & 1 & 0 \end{pmatrix}$$

definiscono le forme quadratiche

$$x_1^2 - 6x_1x_2, \quad x_1^2 + x_2^2 - x_3^2, \quad 2x_1x_2 + 2x_2x_3 + 2x_3x_1.$$

Notiamo che se $i \neq j$ gli elementi $S_{ij} = S_{ji}$ contribuiscono entrambi al monomio $x_i x_j$, che ha quindi coefficiente pari a $2S_{ij}$. Di converso,

$$q(x) = x_1^2 + 4x_1x_2 - x_2^2 + 4x_3^2$$

è descritta dalla matrice

$$S = \begin{pmatrix} 1 & 2 & 0 \\ 2 & -1 & 0 \\ 0 & 0 & 4 \end{pmatrix}.$$

Qui il monomio $4x_1x_2$ viene spezzato in due monomi $2x_1x_2 + 2x_2x_1$. La matrice identità I_n definisce la forma quadratica

$$x_1^2 + \cdots + x_n^2.$$

Lo studio approfondito dei prodotti scalari che faremo in questo capitolo ci servirà anche a capire le forme quadratiche successivamente nel Capitolo 13. Indichiamo con

$$q_S(x) = g_S(x, x)$$

la forma quadratica definita dalla matrice simmetrica S . Notiamo che per definizione il prodotto scalare g_S è definito positivo precisamente quando la corrispondente forma quadratica è positiva $\forall x \neq 0$:

$$g_S \text{ definito positivo} \iff q_S(x) > 0 \quad \forall x \neq 0.$$

7.1.6. Matrici diagonali. Le matrici diagonali sono matrici simmetriche molto semplici e facili da studiare. Consideriamo una matrice diagonale

$$D = \begin{pmatrix} d_1 & 0 & \cdots & 0 \\ 0 & d_2 & \cdots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \cdots & d_n \end{pmatrix}.$$

Il prodotto scalare $g_D(x, y) = {}^t x D y$ su $\mathbb{R}^n$ è semplicemente:

$$g_D(x, y) = d_1 x_1 y_1 + \cdots + d_n x_n y_n.$$

La corrispondente forma quadratica è

$$q_D(x) = d_1 x_1^2 + \cdots + d_n x_n^2.$$

In particolare è facile capire quando il prodotto scalare è definito positivo:

Proposizione 7.1.13. *Il prodotto scalare g_D è definito positivo $\iff d_i > 0$ per ogni $i = 1, \dots, n$.*

Dimostrazione. Se $d_i > 0$ per ogni i allora otteniamo

$$q_D(x) = x_1^2 d_1 + \cdots + x_n^2 d_n > 0$$

per ogni vettore non nullo x . Quindi il prodotto scalare è definito positivo. D'altra parte, se esiste un $d_i \leq 0$, allora

$$q_D(e_i) = d_i \leq 0$$

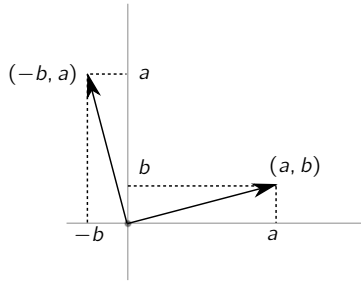


Figura 7.1. Due vettori ortogonali in $\mathbb{R}^2$.

implica che il prodotto scalare non è definito positivo. □

Analogamente è facile capire quando il prodotto scalare è degenere:

Proposizione 7.1.14. *Il prodotto scalare g_D è non degenere $\iff d_i \neq 0$ per ogni $i = 1, \dots, n$.*

Dimostrazione. Notiamo che per ogni $x \in \mathbb{R}^n$ abbiamo $g_D(x, e_i) = d_i x_i$. Se esiste un i tale che $d_i = 0$, allora $g_D(x, e_i) = 0$ per ogni $x \in \mathbb{R}^n$ e quindi il prodotto scalare è degenere.

Se $d_i \neq 0$ per ogni i , il prodotto scalare non è degenere: per ogni $x \in \mathbb{R}^n$ non nullo esiste $y \in \mathbb{R}^n$ tale che $g_D(x, y) \neq 0$. Infatti, siccome $x \neq 0$, deve valere $x_i \neq 0$ per qualche i , e allora $g_D(x_i, e_i) = d_i \neq 0$. □

Ad esempio, se D è la matrice diagonale

$$\begin{pmatrix} 1 & 0 \\ 0 & -3 \end{pmatrix}, \quad \begin{pmatrix} 0 & 0 \\ 0 & 1 \end{pmatrix}, \quad \text{oppure} \quad \begin{pmatrix} 5 & 0 \\ 0 & 1 \end{pmatrix}$$

il prodotto scalare g_D su $\mathbb{R}^2$ è degenere solo nel secondo caso ed è definito positivo solo nel terzo caso. Ricordiamo che

definito positivo $\implies$ non degenere.

Non vale però l'implicazione opposta! La prima matrice $\begin{pmatrix} 1 & 0 \\ 0 & -3 \end{pmatrix}$ definisce un prodotto scalare non degenere che non è però definito positivo.

Ricapitoliamo quanto visto per una matrice diagonale D :

- Il prodotto scalare g_D è definito positivo $\iff d_i > 0 \ \forall i$.
- Il prodotto scalare g_D è non degenere $\iff d_i \neq 0 \ \forall i$.

7.1.7. Vettori ortogonali. Introduciamo alcune definizioni. Sia V munito di un prodotto scalare. Due vettori $v, w \in V$ sono *ortogonali* se

$$\langle v, w \rangle = 0.$$

Esempio 7.1.15. Consideriamo il prodotto scalare euclideo su $\mathbb{R}^2$. I vettori ortogonali ad un vettore dato $\begin{pmatrix} a \\ b \end{pmatrix} \neq 0$ sono quei vettori $\begin{pmatrix} x \\ y \end{pmatrix}$ tali che

$$ax + by = 0.$$

Risolvendo il sistema, ne deduciamo che i vettori ortogonali ad $\begin{pmatrix} a \\ b \end{pmatrix}$ sono precisamente i vettori che stanno nella retta

$$\text{Span} \begin{pmatrix} -b \\ a \end{pmatrix}.$$

Il vettore $\begin{pmatrix} -b \\ a \end{pmatrix}$ forma effettivamente un angolo retto con $\begin{pmatrix} a \\ b \end{pmatrix}$, come mostrato chiaramente dalla Figura 7.1. Il prodotto scalare euclideo sul piano cartesiano $\mathbb{R}^2$ e su $\mathbb{R}^3$ corrisponde alla familiare geometria euclidea; se però scegliamo un altro prodotto scalare su $\mathbb{R}^2$ o $\mathbb{R}^3$ la nozione di ortogonalità può cambiare considerevolmente.

Esempio 7.1.16. Rispetto al prodotto scalare euclideo su $\mathbb{R}^n$, due vettori e_i ed e_j della base canonica con $i \neq j$ sono sempre ortogonali fra loro. Infatti $\langle e_i, e_j \rangle = {}^t e_i \cdot e_j = 0$.

Osservazione 7.1.17. Ricordiamo che un prodotto scalare su V è degenere se esiste un $v \in V$ non nullo tale che $\langle v, w \rangle = 0$ per ogni $w \in V$. Questa condizione può essere espressa dicendo che esiste un vettore $v \in V$ ortogonale a *tutti* i vettori di V .

7.1.8. Vettori isotropi. In un prodotto scalare, può accadere che un vettore sia ortogonale a se stesso! Sia V uno spazio vettoriale munito di un prodotto scalare. Un vettore $v \in V$ tale che $\langle v, v \rangle = 0$ è detto *isotropo*.

Se il prodotto scalare è definito positivo, come ad esempio l'usuale prodotto scalare euclideo su $\mathbb{R}^2$ o su $\mathbb{R}^n$, l'unico vettore isotropo è quello nullo. Prodotti scalari più generali possono però contenere vettori isotropi non banali.

Esempio 7.1.18. Se V è degenere, esiste $v \neq 0$ ortogonale a tutti i vettori di V . In particolare, v è ortogonale a se stesso e quindi è isotropo.

Abbiamo capito che se il prodotto scalare è degenere ci sono vettori isotropi non nulli. È però importante notare che, come mostra il prossimo esempio, ci possono essere dei vettori isotropi non nulli anche in prodotti scalari non degeneri.

Osservazione 7.1.19. Consideriamo il prodotto scalare g_S su $\mathbb{R}^n$ definito da una matrice simmetrica $S \in M(n, \mathbb{R})$. I vettori isotropi di g_S sono precisamente i vettori $x \in \mathbb{R}^n$ che annullano la forma quadratica q_S , cioè le soluzioni dell'equazione di secondo grado in più variabili:

$$q_S(x) = 0.$$

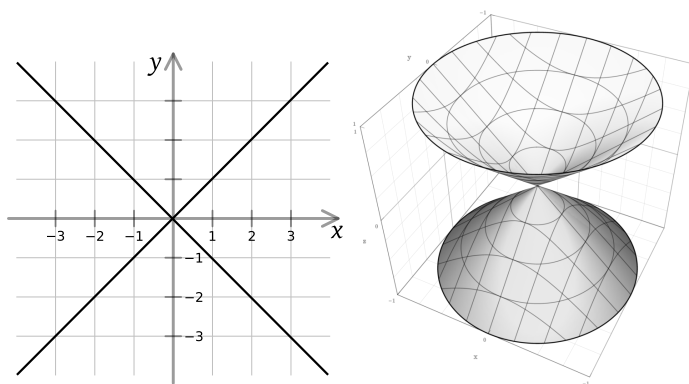


Figura 7.2. Le soluzioni dell'equazione $x_1^2 - x_2^2 = 0$ in $\mathbb{R}^2$ formano le due bisettrici dei quadranti (sinistra). Le soluzioni di $x_1^2 + x_2^2 - x_3^2 = 0$ in $\mathbb{R}^3$ formano un doppio cono (destra).

Esempio 7.1.20. Consideriamo la matrice simmetrica

$$S = \begin{pmatrix} 1 & 0 \\ 0 & -1 \end{pmatrix}.$$

Per quanto visto nella sezione precedente, il prodotto scalare g_S è non degenere e non è definito positivo. La forma quadratica corrispondente è

$$q_S(x) = x_1^2 - x_2^2.$$

I vettori isotropi sono precisamente i vettori $x = \begin{pmatrix} x_1 \\ x_2 \end{pmatrix}$ tali che $x_1^2 - x_2^2 = 0$, cioè tali che $x_1 = \pm x_2$. I vettori isotropi formano quindi le due bisettrici dei quattro quadranti del piano cartesiano come nella Figura 7.2-(sinistra). Il prodotto scalare è non degenere, ma ci sono vettori isotropi non banali.

Esempio 7.1.21. Consideriamo la matrice simmetrica

$$S = \begin{pmatrix} 1 & 0 & 0 \\ 0 & 1 & 0 \\ 0 & 0 & -1 \end{pmatrix}.$$

La forma quadratica corrispondente è

$$q_S(x) = x_1^2 + x_2^2 - x_3^2.$$

I vettori isotropi di g_S sono le soluzioni in $\mathbb{R}^3$ dell'equazione

$$x_1^2 + x_2^2 - x_3^2 = 0.$$

Questi formano il doppio cono mostrato nella Figura 7.2-(destra).

7.1.9. Radicale. Introduciamo adesso una nozione utile a gestire i prodotti scalari degeneri: il *radicale*.

Sia V uno spazio vettoriale munito di un prodotto scalare. Il *radicale* è l'insieme dei vettori $v \in V$ che sono ortogonali a tutti i vettori di V . In altre parole, un vettore v è nel radicale se $\langle v, w \rangle = 0$ per *qualsiasi* vettore $w \in V$.

Il radicale di V è generalmente indicato con $V^\perp$. In simboli:

$$V^\perp = \{v \in V \mid \langle v, w \rangle = 0 \forall w \in V\}.$$

Notiamo subito che se $v \in V^\perp$ allora v è anche isotropo, ma non è vero il viceversa! Un vettore v isotropo non è necessariamente contenuto nel radicale: un vettore v può essere ortogonale a se stesso senza però essere ortogonale anche a tutti gli altri. Osserviamo inoltre che

$$\text{il prodotto scalare è non degenere} \iff V^\perp = \{0\}.$$

Proposizione 7.1.22. *Il radicale $V^\perp$ è un sottospazio vettoriale di V .*

Dimostrazione. Mostriamo che valgono i tre assiomi di sottospazio:

- (1) $0 \in V^\perp$ infatti $\langle 0, w \rangle = 0 \forall w \in V$.
- (2) Se $v, v' \in V^\perp$, allora per ogni $w \in V$ otteniamo

$$\langle v + v', w \rangle = \langle v, w \rangle + \langle v', w \rangle = 0 + 0 = 0$$

e quindi anche $v + v' \in V^\perp$.

- (3) La moltiplicazione per scalare è analoga.

La dimostrazione è conclusa. □

Come primo esempio esaminiamo i prodotti scalari g_S su $\mathbb{R}^n$. Sia $S \in M(n, \mathbb{R})$ una matrice simmetrica.

Proposizione 7.1.23. *Il radicale di g_S è il sottospazio $\ker S$. Quindi il prodotto scalare g_S è degenere $\iff \det S = 0$.*

Dimostrazione. Se $y \in \ker S$, allora per ogni $x \in \mathbb{R}^n$ abbiamo

$$g_S(x, y) = {}^t x S y = 0$$

perché $Sy = 0$. D'altra parte, se un vettore $y \in \mathbb{R}^n$ è tale che ${}^t x S y = 0$ per ogni $x \in \mathbb{R}^n$, allora scriviamo $y' = Sy$ e otteniamo

$${}^t x y' = 0$$

per ogni $x \in \mathbb{R}^n$. Questo implica facilmente che $y' = 0$, cioè che $y \in \ker S$. □

Esempio 7.1.24. Il prodotto scalare su $\mathbb{R}^2$ definito dalla matrice $\begin{pmatrix} 1 & 0 \\ 0 & -1 \end{pmatrix}$ considerato nell'Esempio 7.1.20 è un esempio importante da tenere a mente. Non è degenere perché $\det \begin{pmatrix} 1 & 0 \\ 0 & -1 \end{pmatrix} \neq 0$. Non è neppure definito positivo per la Proposizione 7.1.13. Il radicale è banale (perché è non degenere), ma ci sono comunque dei vettori isotropi (le due bisettrici dei quadranti).

Il vettore $\begin{pmatrix} 1 \\ i \end{pmatrix}$ è ortogonale a se stesso ma non a tutti gli altri (cioè il vettore è isotropo ma non sta nel radicale).

7.1.10. Altri esempi. Introduciamo finalmente alcuni prodotti scalari in spazi vettoriali V diversi da $\mathbb{R}^n$.

Esempio 7.1.25. Un prodotto scalare molto importante usato in analisi nello studio delle funzioni è il seguente. Sia $[a, b] \subset \mathbb{R}$ un intervallo e sia $V = C([a, b])$ lo spazio vettoriale delle funzioni $f: [a, b] \rightarrow \mathbb{R}$ continue. Un prodotto scalare su V è definito nel modo seguente:

$$\langle f, g \rangle = \int_a^b f(t)g(t)dt.$$

La bilinearità si verifica facilmente e la simmetria è immediata. Il prodotto scalare è definito positivo perché

$$\langle f, f \rangle = \int_a^b f(t)^2 dt > 0$$

per ogni funzione continua $f(x)$ non nulla.

È a volte utile definire un prodotto scalare anche sugli spazi di matrici:

Esercizio 7.1.26. Consideriamo il prodotto scalare su $M(n, \mathbb{R})$

$$\langle A, B \rangle = \text{tr}({}^tAB)$$

Mostra che questo è effettivamente un prodotto scalare ed è definito positivo.

Idem per gli spazi di polinomi:

Esempio 7.1.27. Consideriamo sullo spazio $\mathbb{R}_2[x]$ dei polinomi a coefficienti reali di grado ≤ 2 il prodotto scalare

$$\langle p, q \rangle = p(0)q(0) + p(1)q(1) + p(2)q(2).$$

Si verifica facilmente che questo è effettivamente un prodotto scalare. Inoltre è definito positivo: infatti

$$\langle p, p \rangle = p(0)^2 + p(1)^2 + p(2)^2 > 0$$

per qualsiasi polinomio $p(x)$ non nullo di grado ≤ 2 perché un tale polinomio non può annullarsi in tre valori diversi 0, 1, 2 per il Teorema 1.3.7. D'altra parte, il prodotto scalare

$$\langle p, q \rangle = p(0)q(0) + p(1)q(1)$$

è degenere: se prendiamo $p(x) = x(1-x)$ otteniamo $\langle p, q \rangle = 0$ per qualsiasi polinomio $q(x) \in \mathbb{R}_2[x]$. Infine, il prodotto scalare

$$\langle p, q \rangle = p(0)q(0) + p(1)q(1) - p(2)q(2)$$

non è degenere (esercizio: da risolvere dopo aver letto la prossima sezione) ma contiene dei vettori isotropi non banali, ad esempio $p(x) = x-1$, infatti

$$\langle p, p \rangle = (-1)(-1) + 0 \cdot 0 - 1 \cdot 1 = 0.$$

7.1.11. Altri tipi di prodotti scalari. Oltre a *degenere* e *definito positivo*, esistono altri aggettivi che descrivono alcune proprietà che può avere un dato prodotto scalare.

Definizione 7.1.28. Un prodotto scalare su V è

- *semi-definito positivo* se $\langle v, v \rangle \geq 0$ per ogni $v \in V$;
- *definito negativo* se $\langle v, v \rangle < 0$ per ogni $v \in V$ non nullo;
- *semi-definito negativo* se $\langle v, v \rangle \leq 0$ per ogni $v \in V$;
- *indefinito* se esistono v e w per cui $\langle v, v \rangle > 0$ e $\langle w, w \rangle < 0$.

Esercizio 7.1.29. Consideriamo una matrice diagonale D , con elementi $d_1, \dots, d_n$ sulla diagonale. Il prodotto scalare g_D è

- semi-definito positivo $\iff d_i \geq 0$ per ogni i ;
- definito negativo $\iff d_i < 0$ per ogni i ;
- semi-definito negativo $\iff d_i \leq 0$ per ogni i ;
- indefinito $\iff$ esistono i e j con $d_i > 0$ e $d_j < 0$.

7.2. Matrice associata

Come per le applicazioni lineari, è possibile codificare i prodotti scalari in modo efficiente usando le matrici, dopo aver fissato delle basi.

7.2.1. Definizione. Sia V uno spazio vettoriale reale e

$$g: V \times V \longrightarrow \mathbb{R}$$

un prodotto scalare. Sia $\mathcal{B} = \{v_1, \dots, v_n\}$ una base per V .

Definizione 7.2.1. La *matrice associata* al prodotto scalare g nella base $\mathcal{B}$ è la matrice simmetrica S il cui elemento S_{ij} è dato da

$$S_{ij} = g(v_i, v_j).$$

La matrice S è indicata a volte con $[g]_{\mathcal{B}}$.

Esempio 7.2.2. Sia $S \in M(n, \mathbb{R})$ una matrice simmetrica. La matrice associata al prodotto scalare g_S su $\mathbb{R}^n$ rispetto alla base canonica $\mathcal{C}$ è la matrice S stessa, cioè

$$[g_S]_{\mathcal{C}} = S.$$

Infatti per il Corollario 7.1.8 abbiamo $S_{ij} = g_S(e_i, e_j)$.

Esempio 7.2.3. Consideriamo il prodotto scalare euclideo g su $\mathbb{R}^2$. Scegliamo la base

$$\mathcal{B} = \left\{ v_1 = \begin{pmatrix} 1 \\ 0 \end{pmatrix}, v_2 = \begin{pmatrix} 1 \\ 1 \end{pmatrix} \right\}$$

e calcoliamo la matrice associata a g rispetto a $\mathcal{B}$:

$$S = [g]_{\mathcal{B}} = \begin{pmatrix} g(v_1, v_1) & g(v_1, v_2) \\ g(v_2, v_1) & g(v_2, v_2) \end{pmatrix} = \begin{pmatrix} 1 & 1 \\ 1 & 2 \end{pmatrix}.$$

Esercizio 7.2.4. Scrivi la matrice associata al prodotto scalare su $\mathbb{R}_2[x]$

$$g(p, q) = p(-1)q(-1) + p(0)q(0) + p(1)q(1)$$

rispetto alla base canonica $\mathcal{C} = \{1, x, x^2\}$.

Perché nella matrice associata scegliamo di mettere i numeri $S_{ij} = g(v_i, v_j)$? Il motivo principale è che questi n^2 numeri sono sufficienti per calcolare il prodotto scalare $g(v, w)$ di due vettori qualsiasi. Siano

$$v = \lambda_1 v_1 + \dots + \lambda_n v_n, \quad w = \mu_1 v_1 + \dots + \mu_n v_n$$

due vettori arbitrari di V , espressi usando la base $\mathcal{B}$.

Proposizione 7.2.5. *Vale*

$$g(v, w) = \sum_{i=1}^n \sum_{j=1}^n \lambda_i \mu_j g(v_i, v_j).$$

Dimostrazione. Scriviamo

$$v = \lambda_1 v_1 + \dots + \lambda_n v_n, \quad w = \mu_1 v_1 + \dots + \mu_n v_n.$$

Usando la bilinearità di g otteniamo

$$\begin{aligned} g(v, w) &= g(\lambda_1 v_1 + \dots + \lambda_n v_n, \mu_1 v_1 + \dots + \mu_n v_n) \\ &= \sum_{i=1}^n g(\lambda_i v_i, \mu_1 v_1 + \dots + \mu_n v_n) = \sum_{i=1}^n \sum_{j=1}^n g(\lambda_i v_i, \mu_j v_j) \\ &= \sum_{i=1}^n \sum_{j=1}^n \lambda_i \mu_j g(v_i, v_j). \end{aligned}$$

La dimostrazione è conclusa. □

Corollario 7.2.6. *Per ogni coppia di vettori $v, w \in V$ vale l'uguaglianza*

$$g(v, w) = {}^t[v]_{\mathcal{B}} \cdot S \cdot [w]_{\mathcal{B}}.$$

Dimostrazione. Per la Proposizione 7.2.5 abbiamo

$$g(v, w) = \sum_{i=1}^n \sum_{j=1}^n \lambda_i \mu_j S_{ij}$$

Notiamo inoltre che

$$[v]_{\mathcal{B}} = \begin{pmatrix} \lambda_1 \\ \vdots \\ \lambda_n \end{pmatrix}, \quad [w]_{\mathcal{B}} = \begin{pmatrix} \mu_1 \\ \vdots \\ \mu_n \end{pmatrix}.$$

Operando con il prodotto riga per colonna si verifica che effettivamente

$$(\lambda_1 \quad \dots \quad \lambda_n) \cdot S \cdot \begin{pmatrix} \mu_1 \\ \vdots \\ \mu_n \end{pmatrix} = \sum_{i=1}^n \lambda_i \left(\sum_{j=1}^n S_{ij} \mu_j \right) = \sum_{i=1}^n \sum_{j=1}^n \lambda_i \mu_j S_{ij}.$$

Questo conclude la dimostrazione. □

Osservazione 7.2.7. Se scriviamo $x = [v]_B$ e $y = [w]_B$, vediamo che

$$g(v, w) = {}^t x S y.$$

Abbiamo scoperto un fatto importante, del tutto analogo all'Osservazione 4.3.5 per le applicazioni lineari: dopo aver scelto una base, qualsiasi prodotto scalare g su qualsiasi spazio vettoriale V di dimensione n può essere interpretato come un prodotto scalare su $\mathbb{R}^n$ del tipo $g_S(x, y) = {}^t x S y$. Basta sostituire i vettori v e w con le loro coordinate x e y , ed usare la matrice associata S .

Dimostriamo adesso un fatto a cui abbiamo già accennato: i prodotti scalari g_S determinati da matrici simmetriche $S \in M(n, \mathbb{R})$ sono tutti e soli i prodotti scalari su $\mathbb{R}^n$.

Proposizione 7.2.8. *Per ogni prodotto scalare g su $\mathbb{R}^n$ esiste un'unica matrice simmetrica $S \in M(n, \mathbb{R})$ tale che $g = g_S$.*

Dimostrazione. Dato un prodotto scalare g , definiamo $S = [g]_C$ dove C è la base canonica di $\mathbb{R}^n$. Per il corollario precedente abbiamo

$$g(x, y) = {}^t [x]_C S [y]_C = {}^t x S y = g_S(x, y).$$

La dimostrazione è conclusa. $\square$

7.2.2. Cambiamento di base. Come abbiamo fatto per le applicazioni lineari, studiamo adesso come cambia la matrice associata ad un prodotto scalare se cambiamo base. Questo ci porterà a considerare una nuova relazione di equivalenza fra matrici simmetriche.

Siano $\mathcal{B} = \{v_1, \dots, v_n\}$ e $\mathcal{B}' = \{v'_1, \dots, v'_n\}$ due basi di V . Ricordiamo che $M = [\text{id}]_{\mathcal{B}}^{\mathcal{B}'}$ è la matrice di cambiamento di base da $\mathcal{B}'$ a $\mathcal{B}$. Siano S e S' le matrici associate al prodotto scalare g nelle basi $\mathcal{B}$ e $\mathcal{B}'$.

Proposizione 7.2.9. *Vale la relazione seguente:*

$$S' = {}^t M S M.$$

Dimostrazione. Per il Corollario 7.2.6 abbiamo

$$S'_{ij} = g(v'_i, v'_j) = {}^t [v'_i]_{\mathcal{B}} \cdot S \cdot [v'_j]_{\mathcal{B}} = {}^t M^i \cdot S \cdot M^j.$$

Questo vale per ogni i, j , quindi $S' = {}^t M S M$. $\square$

Esempio 7.2.10. Abbiamo già notato nell'Esempio 7.2.3 che la matrice associata al prodotto scalare euclideo di $\mathbb{R}^2$ rispetto alla base

$$\mathcal{B} = \left\{ \begin{pmatrix} 1 \\ 0 \end{pmatrix}, \begin{pmatrix} 1 \\ 1 \end{pmatrix} \right\}$$

è $S = \begin{pmatrix} 1 & 1 \\ 0 & 2 \end{pmatrix}$. La matrice di cambiamento di base è $M = [\text{id}]_C^{\mathcal{B}} = \begin{pmatrix} 1 & 1 \\ 0 & 1 \end{pmatrix}$. La Proposizione 7.2.9 ci dice che $S = {}^t M I_2 M$, e infatti verifichiamo facilmente che vale l'uguaglianza $\begin{pmatrix} 1 & 1 \\ 0 & 2 \end{pmatrix} = \begin{pmatrix} 1 & 0 \\ 0 & 1 \end{pmatrix} \begin{pmatrix} 1 & 1 \\ 0 & 1 \end{pmatrix} \begin{pmatrix} 1 & 1 \\ 0 & 1 \end{pmatrix}$.

7.2.3. Matrici congruenti. L'operazione di cambiamento di base per gli endomorfismi ha portato nella Sezione 4.4.3 ad una relazione di equivalenza fra matrici quadrate chiamata *similitudine*. Con i prodotti scalari, il cambiamento di base porta ad una diversa relazione di equivalenza detta *congruenza*, riservata alle matrici simmetriche.

Definizione 7.2.11. Due matrici $n \times n$ simmetriche S e S' sono *congruenti* se esiste una matrice $n \times n$ invertibile M per cui $S = {}^tMS'M$.

Esercizio 7.2.12. La congruenza fra matrici simmetriche è una relazione di equivalenza.

Si deve stare attenti a non confondere questa relazione di equivalenza con la similitudine, dove al posto della trasposta tM sta l'inversa M^{-1} .

Esempio 7.2.13. Le matrici

$$S = \begin{pmatrix} 1 & 0 \\ 0 & 1 \end{pmatrix}, \quad S' = \begin{pmatrix} 2 & 0 \\ 0 & 1 \end{pmatrix}$$

non sono simili perché hanno determinanti diversi, ma sono congruenti: prendiamo $M = \begin{pmatrix} \sqrt{2} & 0 \\ 0 & 1 \end{pmatrix}$ e verifichiamo facilmente che $S' = {}^tMSM$.

Due matrici simili S e S' hanno sempre lo stesso determinante, due matrici congruenti no! Possiamo comunque dimostrare quanto segue.

Proposizione 7.2.14. Se S e S' sono congruenti, allora $\det S$ e $\det S'$ hanno lo stesso segno, cioè $\det S$ è positivo, nullo, o negativo $\iff \det S'$ è positivo, nullo, o negativo rispettivamente.

Dimostrazione. Per ipotesi esiste una M con $\det M \neq 0$ tale che $S' = {}^tMSM$ e quindi usando il teorema di Binet troviamo

$$\det S' = \det({}^tM) \det S \det M = \det M \det S \det M = (\det M)^2 \det S.$$

La proposizione è dimostrata perché $(\det M)^2 > 0$. $\square$

Studieremo approfonditamente la relazione di congruenza fra matrici quando classificheremo i prodotti scalari, nella Sezione 7.4.

7.3. Sottospazio ortogonale

Introduciamo finalmente un concetto geometrico fondamentale, quello di *ortogonalità* fra sottospazi vettoriali.

7.3.1. Sottospazio ortogonale. Sia V uno spazio vettoriale dotato di un prodotto scalare. Sia $W \subset V$ un sottospazio.

Definizione 7.3.1. Il *sottospazio ortogonale* $W^\perp$ è l'insieme

$$W^\perp = \{v \in V \mid \langle v, w \rangle = 0 \ \forall w \in W\}.$$

L'insieme $W^\perp$ è formato da quei vettori $v \in V$ che sono ortogonali a *tutti* i vettori di W . Notiamo che questa definizione è coerente con la notazione $V^\perp$ già usata per definire il radicale, che come sappiamo è l'insieme dei vettori ortogonali a qualsiasi vettore di tutto lo spazio V .

Proposizione 7.3.2. *L'insieme $W^\perp$ è un sottospazio di V .*

Dimostrazione. Come di consueto dobbiamo mostrare tre cose:

- (1) $0 \in W^\perp$, infatti $\langle 0, w \rangle = 0$ sempre;
- (2) se $v, v' \in W^\perp$ allora $v + v' \in W^\perp$, infatti
$$\langle v + v', w \rangle = \langle v, w \rangle + \langle v', w \rangle = 0 + 0 = 0 \quad \forall w \in W;$$
- (3) se $v \in W^\perp$ allora $\lambda v \in W^\perp$, infatti
$$\langle \lambda v, w \rangle = \lambda \langle v, w \rangle = \lambda 0 = 0 \quad \forall w \in W.$$

La dimostrazione è completa. □

Per determinare $W^\perp$ concretamente è molto utile il fatto seguente.

Proposizione 7.3.3. *Se $W = \text{Span}(w_1, \dots, w_k)$, allora*

$$W^\perp = \{v \in V \mid \langle v, w_i \rangle = 0 \quad \forall i = 1, \dots, k\}.$$

Dimostrazione. Verifichiamo la doppia inclusione fra i due insiemi.

($\subset$). Se $v \in W^\perp$ allora $\langle v, w \rangle = 0$ per ogni $w \in W$ e quindi in particolare $\langle v, w_i \rangle = 0$ per ogni $i = 1, \dots, k$.

($\supset$). Se $\langle v, w_i \rangle = 0$ per ogni $i = 1, \dots, k$ allora per ogni $w \in W$ possiamo scrivere $w = \lambda_1 w_1 + \dots + \lambda_k w_k$ e quindi otteniamo

$$\langle v, w \rangle = \langle v, \lambda_1 w_1 + \dots + \lambda_k w_k \rangle = \lambda_1 \langle v, w_1 \rangle + \dots + \lambda_k \langle v, w_k \rangle = 0.$$

Quindi $v \in W^\perp$. □

Concretamente, l'insieme $W^\perp$ può essere descritto come l'insieme dei vettori v che soddisfano le k relazioni $\langle v, w_i \rangle = 0$ con $i = 1, \dots, k$.

7.3.2. Esempi. Facciamo alcuni esempi di calcolo del sottospazio ortogonale $W^\perp$ in vari contesti. Useremo spesso la Proposizione 7.3.3.

Esempio 7.3.4. Consideriamo $\mathbb{R}^2$ con il prodotto scalare euclideo. Il sottospazio ortogonale ad una retta $W = \text{Span}\begin{pmatrix} a \\ b \end{pmatrix}$ è la retta $W^\perp = \text{Span}\begin{pmatrix} -b \\ a \end{pmatrix}$, si veda l'Esempio 7.1.15 e la Figura 7.1.

Infatti $W^\perp$ consta di tutti vettori $\begin{pmatrix} x \\ y \end{pmatrix}$ che sono ortogonali al generatore $\begin{pmatrix} a \\ b \end{pmatrix}$ di W . Quindi è lo spazio delle soluzioni dell'equazione $ax + by = 0$, cioè la retta $W^\perp = \text{Span}\begin{pmatrix} -b \\ a \end{pmatrix}$.

Esempio 7.3.5. Consideriamo $\mathbb{R}^3$ con il prodotto scalare euclideo. Consideriamo una retta generica

$$W = \text{Span} \begin{pmatrix} a \\ b \\ c \end{pmatrix}.$$

Il sottospazio ortogonale $W^\perp$ consta di tutti quei vettori che sono ortogonali al generatore della retta. Un vettore generico è ortogonale al generatore se soddisfa l'equazione

$$ax + by + cz = 0.$$

Quindi $W^\perp$ è il piano descritto proprio da questa equazione:

$$W^\perp = \{ax + by + cz = 0\}.$$

Esempio 7.3.6. Consideriamo sempre $\mathbb{R}^3$ con il prodotto scalare euclideo. Sia adesso $W = \{ax + by + cz = 0\}$ un piano generico. Il sottospazio $W^\perp$ consta di quei vettori che sono ortogonali a tutti i vettori di W . Uno di questi è certamente

$$v = \begin{pmatrix} a \\ b \\ c \end{pmatrix}$$

perché l'equazione stessa che definisce W dice che tutti i vettori di W sono ortogonali a v . Quindi il sottospazio $W^\perp$ contiene la retta generata da v . Con un argomento di dimensione che descriveremo nelle prossime pagine potremo dedurre che $W^\perp$ è proprio questa retta, cioè

$$W^\perp = \text{Span}(v).$$

Introduciamo un metodo generale per calcolare il sottospazio ortogonale in presenza di un prodotto scalare su $\mathbb{R}^n$. Sia g_S un prodotto scalare in $\mathbb{R}^n$ determinato da una matrice simmetrica S . Sia $W = \text{Span}(v_1, \dots, v_k)$ un sottospazio vettoriale di $\mathbb{R}^n$.

Proposizione 7.3.7. *Il sottospazio ortogonale $W^\perp$ è espresso in forma cartesiana come soluzione del sistema lineare omogeneo con k equazioni:*

$$W^\perp = \{x \in \mathbb{R}^n \mid {}^t v_i S x = 0 \ \forall i = 1, \dots, k\}.$$

Dimostrazione. Le k equazioni esprimono la condizione che x sia ortogonale ai vettori $v_1, \dots, v_k$. $\square$

Esempio 7.3.8. Consideriamo $\mathbb{R}^2$ con il prodotto scalare g_S definito dalla matrice simmetrica $S = \begin{pmatrix} 1 & 0 \\ 0 & -1 \end{pmatrix}$. Se $r = \text{Span}\begin{pmatrix} a \\ b \end{pmatrix}$, allora

$$r^\perp = \{ax - by = 0\}.$$

Risolvendo troviamo che $r^\perp$ è la retta $r^\perp = \text{Span}\begin{pmatrix} b \\ a \end{pmatrix}$. Notiamo qui un paio di fatti interessanti. Il primo è che, siccome il prodotto scalare *non* è quello euclideo, le due rette $r = \text{Span}\begin{pmatrix} a \\ b \end{pmatrix}$ e $r^\perp = \text{Span}\begin{pmatrix} b \\ a \end{pmatrix}$ in generale non sono ortogonali rispetto al prodotto scalare euclideo: si veda la Figura 7.3-(destra). Il secondo, più sorprendente, è che se $a = b$ le due rette r e $r^\perp$ in realtà coincidono! Questo è collegato al fatto che se $a = b$ il vettore $\begin{pmatrix} a \\ a \end{pmatrix}$ è isotropo ed è quindi ortogonale a se stesso.

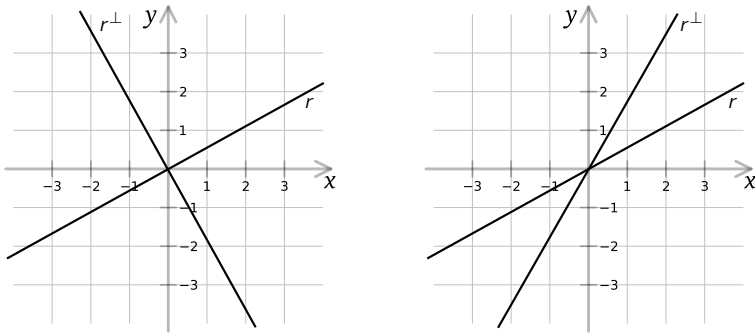


Figura 7.3. La retta $r^\perp$ ortogonale ad r con il prodotto scalare euclideo (sinistra) e con il prodotto scalare dell'Esempio 7.3.8 (destra).

Esempio 7.3.9. Consideriamo $\mathbb{R}^2$ con il prodotto scalare g_S definito dalla matrice simmetrica $S = \begin{pmatrix} 1 & 0 \\ 0 & 0 \end{pmatrix}$. Se $W = \text{Span}\begin{pmatrix} a \\ b \end{pmatrix}$, allora

$$W^\perp = \{ax = 0\}.$$

Notiamo che si presentano solo due casi: se $a \neq 0$, allora $W^\perp$ è la retta verticale $x = 0$. Se invece $a = 0$, il sottospazio $W^\perp$ è tutto il piano $\mathbb{R}^2$.

Gli esempi precedenti mostrano i fatti seguenti:

- (1) Se il prodotto scalare non è definito positivo, può accadere che $W^\perp$ e W non siano in somma diretta.
- (2) Se il prodotto scalare è degenere, può accadere che $W^\perp$ abbia dimensione maggiore di quanto ci aspettavamo.

Nelle prossime pagine studieremo entrambi questi fenomeni.

7.3.3. Restrizioni. Introduciamo adesso una semplice operazione che consiste nel restringere un prodotto scalare ad un sottospazio. Ne avremo bisogno successivamente per studiare le relazioni fra un sottospazio W ed il suo ortogonale $W^\perp$.

Sia V dotato di un prodotto scalare g . Restringendo il prodotto scalare g ai soli vettori di un sottospazio $W \subset V$, si ottiene un prodotto scalare per lo spazio W che possiamo indicare con $g|_W$. Il prodotto scalare $g|_W$ è la *restrizione* di g a W .

La relazione fra g e $g|_W$ può essere sottile: non è detto che una buona proprietà di g si conservi sulla restrizione $g|_W$. Se g è definito positivo questa proprietà si conserva, se g è solo non degenere può non conservarsi:

Proposizione 7.3.10. *Notiamo che:*

- (1) se g è definito positivo, anche $g|_W$ è definito positivo;
- (2) se g è non degenere, non è detto che $g|_W$ sia non degenere!

Dimostrazione. (1). Se $g(v, v) > 0$ per ogni $v \in V$ non nullo, lo stesso vale per ogni $v \in W$ non nullo.

(2) Qui dobbiamo fornire un controesempio. Prendiamo $S = \begin{pmatrix} 1 & 0 \\ 0 & -1 \end{pmatrix}$ ed il prodotto scalare $g = g_S$. Sappiamo che g è non degenere. Definiamo quindi W come la retta $W = \text{Span}(v)$ con $v = \begin{pmatrix} 1 \\ 1 \end{pmatrix}$. Il prodotto scalare $g|_W$ è il prodotto scalare nullo, perché v è isotropo: per ogni coppia di vettori $\lambda v, \mu v \in W$ otteniamo

$$g(\lambda v, \mu v) = \lambda \mu g(v, v) = 0.$$

In particolare la restrizione $g|_W$ è degenere. □

Alla fine della dimostrazione abbiamo anche notato il fatto seguente.

Proposizione 7.3.11. *Se $W = \text{Span}(v)$, la restrizione $g|_W$ è degenere $\iff v$ è un vettore isotropo.*

7.3.4. Dimensione del sottospazio ortogonale. Dimostriamo adesso un teorema che chiarisce, a seconda del tipo di prodotto scalare che abbiamo, che relazione ci sia fra un sottospazio e il suo ortogonale.

Sia V uno spazio vettoriale di dimensione n munito di un prodotto scalare g . Sia $W \subset V$ un sottospazio vettoriale.

Teorema 7.3.12. *Valgono i fatti seguenti:*

- (1) $\dim W^\perp \geq n - \dim W$;
- (2) se g è non-degenere, allora $\dim W^\perp = n - \dim W$;
- (3) se $g|_W$ è non-degenere, allora $V = W \oplus W^\perp$;
- (4) se g è definito positivo, allora $V = W \oplus W^\perp$.

Dimostrazione. Sia $w_1, \dots, w_k$ una base di W . Un vettore $v \in V$ è ortogonale ad ogni vettore di W se (e solo se) è ortogonale a ciascun vettore della base $w_1, \dots, w_k$, cioè se e solo se $g(v, w_i) = 0$ per ogni i . Consideriamo l'applicazione lineare $T: V \rightarrow \mathbb{R}^k$ definita nel modo seguente:

$$T(v) = \begin{pmatrix} g(v, w_1) \\ \vdots \\ g(v, w_k) \end{pmatrix}.$$

Lo spazio $W^\perp$ è precisamente il nucleo di T . Per il teorema della dimensione

$$\dim W^\perp = \dim \ker T = n - \dim \text{Im } T \geq n - k$$

infatti $\dim \text{Im } T \leq k$, perché $\text{Im } T \subset \mathbb{R}^k$. Questo mostra il punto (1), quindi ora passiamo a mostrare (2). Completiamo $w_1, \dots, w_k$ a base $w_1, \dots, w_n$ di V e consideriamo come sopra l'applicazione $T': V \rightarrow \mathbb{R}^n$ data da

$$T'(v) = \begin{pmatrix} g(v, w_1) \\ \vdots \\ g(v, w_n) \end{pmatrix}.$$

Siccome g non è degenere, la mappa T' è iniettiva. Poiché $\dim V = \dim \mathbb{R}^n = n$, la mappa T' è anche suriettiva. Questo implica facilmente che anche T è suriettiva. Quindi $\dim \operatorname{Im} T = k$ e $\dim W^\perp = n - k$. Il punto (2) è dimostrato.

Passiamo a dimostrare il punto (3). Se $g|_W$ è non degenere, nessun vettore non nullo di W è ortogonale a tutti i vettori di W e quindi $W \cap W^\perp = \{0\}$. Per il punto (1) sappiamo che $\dim W + \dim W^\perp \geq n$, quindi per la formula di Grassmann deduciamo che $\dim(W + W^\perp) = \dim W + \dim W^\perp \geq n$, quindi $W + W^\perp = V$ e allora $V = W \oplus W^\perp$.

Se g è definito positivo, ogni restrizione $g|_W$ è definita positiva ed in particolare non è degenere, quindi (3) implica (4). $\square$

Una decomposizione del tipo $V = W \oplus W^\perp$ si chiama *somma diretta ortogonale*. Per quanto appena visto, questa esiste se $g|_W$ è non degenere.

Esercizio 7.3.13. Sia V dotato di un prodotto scalare g e $U, W \subset V$ due sottospazi. Valgono i fatti seguenti:

- se $U \subset W$ allora $W^\perp \subset U^\perp$;
- $(U^\perp)^\perp \supset U$, con l'uguaglianza se g è non-degenere;
- $(U + W)^\perp = U^\perp \cap W^\perp$.

7.3.5. Prodotto scalare euclideo. Mostriamo come determinare il sottospazio ortogonale nello spazio $\mathbb{R}^n$ dotato del prodotto scalare euclideo. La procedura è sorprendentemente semplice.

Se $W \subset \mathbb{R}^n$ è un sottospazio, sappiamo che

$$W \oplus W^\perp = \mathbb{R}^n$$

perché il prodotto scalare euclideo è definito positivo. Il sottospazio $W \subset \mathbb{R}^n$ è descritto in forma parametrica come $W = \operatorname{Span}(B^1, \dots, B^k)$, dove $B^1, \dots, B^k$ sono le colonne di una matrice $B \in M(n, k)$, oppure in forma cartesiana come $W = \{Ax = 0\}$ per qualche $A \in M(m, n)$. In entrambi i casi è facile descrivere lo spazio ortogonale $W^\perp$.

Proposizione 7.3.14. *Valgono i fatti seguenti.*

$$\begin{aligned} W = \operatorname{Span}(B^1, \dots, B^k) &\implies W^\perp = \{ {}^t Bx = 0 \} \\ W = \{ Ax = 0 \} &\implies W^\perp = \operatorname{Span}({}^t A_1, \dots, {}^t A_m). \end{aligned}$$

Da una descrizione di W in forma parametrica (rispettivamente, cartesiana), si ottiene una descrizione di $W^\perp$ in forma cartesiana (rispettivamente, parametrica) facendo semplicemente la trasposta della matrice.

Dimostrazione. In entrambi i casi, per costruzione $U = \{ {}^t Bx = 0 \}$ oppure $U = \operatorname{Span}({}^t A_1, \dots, {}^t A_m)$ è un sottospazio che contiene solo vettori ortogonali a W . Quindi $U \subset W^\perp$. Si vede anche che $\dim U = n - \dim W = \dim W^\perp$ e quindi $U = W^\perp$. $\square$

Esempio 7.3.15. Consideriamo $\mathbb{R}^5$ con il prodotto scalare euclideo. Sia $W \subset \mathbb{R}^5$ il piano

$$W = \text{Span} \left(\begin{pmatrix} 1 \\ 2 \\ -11 \\ 0 \\ 5 \end{pmatrix}, \begin{pmatrix} -1 \\ 0 \\ 9 \\ 1 \\ -7 \end{pmatrix} \right).$$

Il suo ortogonale è il sottospazio tridimensionale

$$W^\perp = \begin{cases} x_1 + 2x_2 - 11x_3 + 5x_5 = 0, \\ -x_1 + 9x_3 + x_4 - 7x_5 = 0. \end{cases}$$

Esempio 7.3.16. Consideriamo $\mathbb{R}^3$ con il prodotto scalare euclideo. Sia $W \subset \mathbb{R}^3$ la retta descritta in forma cartesiana come

$$\begin{cases} 2x - 3y + z = 0 \\ x + y = 0 \end{cases}$$

Il suo ortogonale è il piano

$$W = \text{Span} \left(\begin{pmatrix} 2 \\ -3 \\ 1 \end{pmatrix}, \begin{pmatrix} 1 \\ 1 \\ 0 \end{pmatrix} \right).$$

7.4. Classificazione dei prodotti scalari

Lo scopo principale di questa sezione è ottenere una classificazione completa dei prodotti scalari su spazi vettoriali di dimensione finita.

7.4.1. Basi ortogonali e ortonormali. In presenza di un prodotto scalare, fra le tante basi disponibili è spesso preferibile scegliere una base in cui i vettori sono tutti ortogonali fra loro. Una base di questo tipo è detta *ortogonale*.

Sia g un prodotto scalare su uno spazio vettoriale V di dimensione n .

Definizione 7.4.1. Una base $\mathcal{B} = \{v_1, \dots, v_n\}$ dello spazio V è *ortogonale* se $g(v_i, v_j) = 0$ per ogni $i \neq j$. In altre parole, una base $\mathcal{B}$ è ortogonale se due vettori distinti di $\mathcal{B}$ sono sempre ortogonali fra loro.

Una base ortogonale è *ortonormale* se per ciascun i il numero $g(v_i, v_i)$ è 1, 0 oppure -1 .

Sia $S = [g]_{\mathcal{B}}$ la matrice associata a g nella base $\mathcal{B}$. Segue immediatamente dalla definizione che

$\mathcal{B}$ è ortogonale $\iff S = [g]_{\mathcal{B}}$ è una matrice diagonale.

Ricordiamo infatti che $S_{ij} = g(v_i, v_j)$, quindi chiedere che $g(v_i, v_j) = 0$ per ogni $i \neq j$ è come chiedere che S sia diagonale.

D'altra parte $\mathcal{B}$ è ortonormale $\iff$ la matrice S è diagonale e inoltre gli elementi sulla diagonale sono soltanto di tre tipi: 1, 0 e -1 .

Esempio 7.4.2. La base canonica è ortonormale per il prodotto scalare euclideo. Infatti $\langle e_i, e_j \rangle = 0$ per ogni $i \neq j$ e inoltre $\langle e_i, e_i \rangle = 1$ per ogni i .

Esempio 7.4.3. Le basi seguenti sono basi ortogonali di $\mathbb{R}^2$ e $\mathbb{R}^3$ rispetto al prodotto scalare euclideo:

$$\mathcal{B}_1 = \left\{ \begin{pmatrix} 1 \\ 1 \end{pmatrix}, \begin{pmatrix} -1 \\ 1 \end{pmatrix} \right\}, \quad \mathcal{B}_2 = \left\{ \begin{pmatrix} 2 \\ 2 \\ 1 \end{pmatrix}, \begin{pmatrix} -1 \\ 2 \\ -2 \end{pmatrix}, \begin{pmatrix} 2 \\ -1 \\ -2 \end{pmatrix} \right\}.$$

Le basi non sono però ortonormali. Possiamo però trasformarle in basi ortonormali moltiplicando ciascun vettore per un opportuno scalare; ad esempio, le basi seguenti sono ortonormali:

$$\mathcal{C}_1 = \left\{ \begin{pmatrix} \frac{\sqrt{2}}{2} \\ \frac{\sqrt{2}}{2} \end{pmatrix}, \begin{pmatrix} -\frac{\sqrt{2}}{2} \\ \frac{\sqrt{2}}{2} \end{pmatrix} \right\}, \quad \mathcal{C}_2 = \left\{ \begin{pmatrix} \frac{2}{3} \\ \frac{2}{3} \\ \frac{1}{3} \end{pmatrix}, \begin{pmatrix} -\frac{1}{3} \\ \frac{2}{3} \\ -\frac{2}{3} \end{pmatrix}, \begin{pmatrix} \frac{2}{3} \\ -\frac{1}{3} \\ -\frac{2}{3} \end{pmatrix} \right\}.$$

Questo processo di trasformazione di una base ortogonale in ortonormale è detto *normalizzazione* e verrà descritto nella Sezione 7.4.3.

Esempio 7.4.4. Consideriamo il prodotto scalare $g = g_5$ su $\mathbb{R}^2$ definito dalla matrice $S = \begin{pmatrix} 0 & 1 \\ 1 & 0 \end{pmatrix}$. Abbiamo

$$g(x, y) = \begin{pmatrix} x_1 & x_2 \end{pmatrix} \begin{pmatrix} 0 & 1 \\ 1 & 0 \end{pmatrix} \begin{pmatrix} y_1 \\ y_2 \end{pmatrix} = x_2 y_1 + x_1 y_2.$$

La base $v_1 = \begin{pmatrix} 1 \\ 1 \end{pmatrix}$, $v_2 = \begin{pmatrix} 1 \\ -1 \end{pmatrix}$ è una base ortogonale per il prodotto scalare, perché $g(v_1, v_2) = 0$. Questa ovviamente non è l'unica base ortogonale possibile, ce ne sono molte altre: ad esempio anche $w_1 = \begin{pmatrix} 1 \\ 2 \end{pmatrix}$, $w_2 = \begin{pmatrix} 1 \\ -2 \end{pmatrix}$ è una base ortogonale, sempre perché $g(w_1, w_2) = 0$.

7.4.2. Esistenza di basi ortogonali. Mostriamo che è sempre possibile trovare basi ortogonali. Sia V uno spazio vettoriale di dimensione n con un prodotto scalare g .

Proposizione 7.4.5. *Esiste sempre una base ortogonale $\mathcal{B}$ per V .*

Dimostrazione. Lavoriamo per induzione su $n = \dim V$. Se $n = 1$, una base di V è formata da un vettore solo ed è sempre ortogonale. Supponiamo che l'enunciato sia valido per $n - 1$ e lo mostriamo per n .

Supponiamo che esista almeno un vettore $v \in V$ non isotropo. Consideriamo la retta vettoriale $W = \text{Span}(v)$. La restrizione di g a W è non-degenere perché $g(v, v) \neq 0$. Quindi il Teorema 7.3.12-(3) ci dice che $V = W \oplus W^\perp$.

Lo spazio $W^\perp$ ha dimensione $n - 1$ e quindi possiamo applicare l'ipotesi induttiva al prodotto scalare $g|_{W^\perp}$ ottenuto restringendo g a $W^\perp$. Per l'ipotesi induttiva esiste una base ortogonale $v_2, \dots, v_n$ per $W^\perp$. Aggiungendo v alla sequenza otteniamo una base $\mathcal{B} = \{v, v_2, \dots, v_n\}$ per V che è ancora ortogonale, infatti $v_i \in W^\perp$ implica $g(v, v_i) = 0$ per ogni i .

Resta da considerare il caso in cui tutti i vettori in V siano isotropi, cioè $g(v, v) = 0$ per ogni $v \in V$. La facile uguaglianza

$$g(v, w) = \frac{g(v + w, v + w) - g(v, v) - g(w, w)}{2} = \frac{0 - 0 - 0}{2} = 0$$

ci mostra che $g(v, w) = 0$ per ogni $v, w \in V$ e quindi qualsiasi base di V è ortogonale. $\square$

Ricordiamo che se $\mathcal{B}$ è una base ortogonale allora la matrice associata a g nella base $\mathcal{B}$ è diagonale. Ne ricaviamo il corollario seguente.

Corollario 7.4.6. *Qualsiasi matrice simmetrica S è congruente ad una matrice diagonale.*

Dimostrazione. Consideriamo il prodotto scalare $g = g_S$ su $\mathbb{R}^n$. Sappiamo che $S = [g]_{\mathcal{C}}$ dove $\mathcal{C}$ è la base canonica. Esiste sempre una base ortogonale $\mathcal{B}$ per g , quindi $D = [g]_{\mathcal{B}}$ è diagonale. Allora

$$S = {}^tMDM$$

con $M = [\text{id}]_{\mathcal{B}}^{\mathcal{C}}$. Quindi S è congruente a D . $\square$

7.4.3. Normalizzazione. Descriviamo adesso un semplice algoritmo di *normalizzazione*, che trasforma una base ortogonale in una base ortonormale.

Se $\mathcal{B} = \{v_1, \dots, v_n\}$ è una base ortogonale per un certo prodotto scalare, possiamo sostituire ciascun vettore v_i con un nuovo vettore w_i definito nel modo seguente:

$$w_i = \begin{cases} \frac{v_i}{\sqrt{\langle v_i, v_i \rangle}} & \text{se } \langle v_i, v_i \rangle > 0, \\ v_i & \text{se } \langle v_i, v_i \rangle = 0, \\ \frac{v_i}{\sqrt{-\langle v_i, v_i \rangle}} & \text{se } \langle v_i, v_i \rangle < 0. \end{cases}$$

Il risultato è una nuova base $\{w_1, \dots, w_n\}$ in cui il numero $\langle w_i, w_i \rangle$ può essere solo 1, -1 , oppure 0, per ogni i . Il lettore è invitato a verificare questo fatto.

Notiamo inoltre che possiamo anche riordinare gli elementi di questa base in modo che prima ci siano i vettori w_i con $\langle w_i, w_i \rangle = 1$, poi quelli con $\langle w_i, w_i \rangle = -1$, ed infine quelli con $\langle w_i, w_i \rangle = 0$. La matrice associata $S = [g]_{\mathcal{B}}$ a questo punto diventa del tipo

$$S = \begin{pmatrix} I_k & 0 & 0 \\ 0 & -I_h & 0 \\ 0 & 0 & 0 \end{pmatrix}.$$

Riassumiamo questa indagine con un enunciato. Sia come sempre V uno spazio vettoriale reale di dimensione n e g un prodotto scalare su V .

Proposizione 7.4.7. *Esiste sempre una base ortonormale $\mathcal{B}$ per V tale che la matrice associata $S = [g]_{\mathcal{B}}$ sia del tipo*

$$(8) \quad S = \begin{pmatrix} I_k & 0 & 0 \\ 0 & -I_h & 0 \\ 0 & 0 & 0 \end{pmatrix}.$$

Dimostrazione. Esiste una base ortogonale, che può essere trasformata in ortonormale tramite la normalizzazione e quindi riordinata come descritto sopra. $\square$

Corollario 7.4.8. *Qualsiasi matrice simmetrica è congruente ad una matrice diagonale S come in (8).*

Corollario 7.4.9. *Per qualsiasi forma quadratica $q(x) = {}^t x S x$ esiste un cambiamento di variabili $x = Mx'$ per cui $q(Mx')$ è uguale a*

$$(x'_1)^2 + \cdots + (x'_k)^2 - (x'_{k+1})^2 - \cdots - (x'_{k+h})^2.$$

Questo fatto sarà utile nel Capitolo 13.

7.4.4. Segnatura. Abbiamo scoperto che esistono sempre delle basi ortogonali. Adesso definiamo una terna di numeri, detta *segnatura*, che caratterizza completamente un prodotto scalare.

Sia V uno spazio vettoriale di dimensione n munito di un prodotto scalare g . La *segnatura* è una terna di numeri naturali

$$(i_+, i_-, i_0)$$

definiti nel modo seguente:

- $i_+ = \max\{\dim W \mid W \subset V \text{ sottospazio con } g|_W \text{ definito positivo}\}$
- $i_- = \max\{\dim W \mid W \subset V \text{ sottospazio con } g|_W \text{ definito negativo}\}$
- $i_0 = \dim V^\perp$

I numeri i_+, i_-, i_0 sono chiamati *indice di positività*, *indice di negatività* e *indice di nullità* di g . L'indice di positività (rispettivamente, negatività) è la massima dimensione di un sottospazio $W \subset V$ su cui la restrizione $g|_W$ è definita positiva (rispettivamente, negativa). L'indice di nullità è più semplicemente la dimensione del radicale.

La definizione appena data suona abbastanza astratta, ed è quindi importante adesso fornire degli strumenti di calcolo. Il più importante è il seguente.

Teorema 7.4.10 (Teorema di Sylvester). *Sia $\mathcal{B} = \{v_1, \dots, v_n\}$ una base ortogonale per V . Valgono i fatti seguenti:*

- (1) i_+ è il numero di vettori v_i della base con $\langle v_i, v_i \rangle > 0$.
- (2) i_- è il numero di vettori v_i della base con $\langle v_i, v_i \rangle < 0$.
- (3) i_0 è il numero di vettori v_i della base con $\langle v_i, v_i \rangle = 0$.

Dimostrazione. Per semplicità ordiniamo i vettori in modo che

- (1) $v_1, \dots, v_k$ siano tali che $\langle v_i, v_i \rangle > 0$,

(2) $v_{k+1}, \dots, v_{k+h}$ siano tali che $\langle v_i, v_i \rangle < 0$,

(3) $v_{k+h+1}, \dots, v_n$ siano tali che $\langle v_i, v_i \rangle = 0$.

Dobbiamo dimostrare che

$$k = i_+, \quad h = i_-, \quad n - (k + h) = i_0.$$

Consideriamo i sottospazi

$$W_+ = \text{Span}(v_1, \dots, v_k), \quad W_- = \text{Span}(v_{k+1}, \dots, v_{k+h}),$$

$$W_0 = \text{Span}(v_{k+h+1}, \dots, v_n).$$

Notiamo che

$$V = W_+ \oplus W_- \oplus W_0.$$

La Proposizione 7.1.23 ci dice che W_0 è proprio il radicale, quindi effettivamente $i_0 = n - (k + h)$.

Le restrizioni $g|_{W_+}$ e $g|_{W_-}$ sono definite positive e negative rispettivamente per la Proposizione 7.1.13, quindi $i_+ \geq k$ e $i_- \geq h$. Supponiamo per assurdo che non valgano le uguaglianze, cioè ad esempio che $i_+ > k$. Allora esiste un altro sottospazio W di dimensione $i_+ > k$ su cui la restrizione $g|_W$ è definita positiva. Notiamo che

$$\dim W + \dim(W_- \oplus W_0) = i_+ + (n - k) > n$$

e quindi per la formula di Grassman l'intersezione $W \cap (W_- \oplus W_0)$ ha dimensione almeno uno. Quindi esiste un $v \neq 0$ con $v \in W$ (e quindi $\langle v, v \rangle > 0$) e $v \in W_- \oplus W_0$ (e quindi $\langle v, v \rangle \leq 0$): assurdo. $\square$

Dalla proposizione possiamo dedurre questa proprietà della segnatura, che non era affatto ovvia all'inizio.

Corollario 7.4.11. *La somma dei tre indici della segnatura è sempre*

$$i_+ + i_- + i_0 = n = \dim V.$$

7.4.5. Proprietà della segnatura. La segnatura (i_+, i_-, i_0) caratterizza il prodotto scalare g . Molte delle definizioni che abbiamo dato possono essere reinterpretate con la segnatura:

- g è definito positivo $\iff (i_+, i_-, i_0) = (n, 0, 0)$;
- g è definito negativo $\iff (i_+, i_-, i_0) = (0, n, 0)$;
- g è degenere $\iff i_0 > 0$;
- g è semi-definito positivo $\iff i_- = 0$;
- g è semi-definito negativo $\iff i_+ = 0$;
- g è indefinito $\iff i_+ > 0$ e $i_- > 0$.

La lettrice è invitata a dimostrare tutte queste facili equivalenze.

La segnatura caratterizza i prodotti scalari e quindi anche le matrici simmetriche. Definiamo la segnatura (i_+, i_-, i_0) di una matrice simmetrica S come la segnatura del prodotto scalare g_S su $\mathbb{R}^n$.

Ricordiamo che se g è un prodotto scalare e $\mathcal{B}$ è una base di V , allora g su V viene rappresentato in coordinate come g_S su $\mathbb{R}^n$. I prodotti scalari g e g_S hanno la stessa segnatura, che è per definizione la segnatura di S .

Da un punto di vista algebrico, la segnatura è utile a classificare completamente le matrici a meno di congruenza.

Proposizione 7.4.12. *Due matrici simmetriche sono congruenti se e solo se hanno la stessa segnatura.*

Dimostrazione. Due matrici simmetriche congruenti rappresentano lo stesso prodotto scalare in basi diverse e quindi hanno la stessa segnatura. Due matrici simmetriche con la stessa segnatura sono congruenti alla stessa matrice diagonale e quindi sono congruenti per la proprietà transitiva. $\square$

La segnatura (i_+, i_-, i_0) di una matrice diagonale D si calcola semplicemente contando gli elementi positivi, negativi e nulli sulla diagonale di D . Per una matrice simmetrica S non diagonale il calcolo della segnatura può essere più complesso.

7.4.6. Calcolo della segnatura. Abbiamo dimostrato che la segnatura (i_+, i_-, i_0) caratterizza un prodotto scalare, quindi è naturale adesso porsi il problema seguente: come si calcola la segnatura di un prodotto scalare dato?

Un metodo consiste nel determinare una base ortogonale seguendo la dimostrazione della Proposizione 7.4.5 e quindi usare il Teorema di Sylvester. Il problema con questo approccio è che la costruzione della base ortogonale può essere laboriosa. Cerchiamo quindi un modo di determinare la segnatura senza dover costruire una base ortogonale.

In seguito descriveremo alcuni metodi generali per risolvere questo problema; per adesso ci accontentiamo di fare alcune considerazioni. Sia g un prodotto scalare su uno spazio vettoriale V e S la matrice associata a g rispetto ad una base $\mathcal{B} = \{v_1, \dots, v_n\}$ qualsiasi. Abbiamo visto nella Proposizione 7.2.14 che il segno di $\det S$ non dipende dalla base $\mathcal{B}$. Poiché S è congruente ad una matrice diagonale del tipo (8), ne deduciamo che:

- $\det S = 0 \iff i_0 > 0$;
- se $\det S > 0$, allora i_- è pari,
- se $\det S < 0$, allora i_- è dispari.

Notiamo anche un'altra cosa. Se $W \subset V$ è un sottospazio su cui la restrizione $g|_W$ è definita positiva, allora $i_+ \geq \dim W$. In particolare, se $S_{ii} > 0$ per qualche i , allora $\langle v_i, v_i \rangle = S_{ii} > 0$ e quindi la restrizione $g|_{\text{Span}(v_i)}$ sulla retta $\text{Span}(v_i)$ è definita positiva: ne segue che $i_+ > 0$. Riassumendo:

- se esiste i con $S_{ii} > 0$, allora $i_+ > 0$;
- se esiste i con $S_{ii} < 0$, allora $i_- > 0$.

Queste informazioni sono sufficienti per trattare completamente il caso bidimensionale. Consideriamo il prodotto scalare g_S su $\mathbb{R}^2$ definito da

$$S = \begin{pmatrix} a & b \\ b & c \end{pmatrix}.$$

Notiamo che $\det S = ac - b^2$. Questa è la casistica:

- se $S = 0$, la segnatura è $(0, 0, 2)$;
- se $\det S = 0$ e $a > 0$ oppure $c > 0$, la segnatura è $(1, 0, 1)$;
- se $\det S = 0$ e $a < 0$ oppure $c < 0$, la segnatura è $(0, 1, 1)$;
- se $\det S > 0$ e $a > 0$ oppure $c > 0$, la segnatura è $(2, 0, 0)$;
- se $\det S > 0$ e $a < 0$ oppure $c < 0$, la segnatura è $(0, 2, 0)$;
- se $\det S < 0$, la segnatura è $(1, 1, 0)$.

Il lettore è invitato a dimostrare i vari casi usando i criteri elencati sopra.

Osservazione 7.4.13. Attenzione: un elemento positivo sulla diagonale di S implica $i_+ \geq 1$, ma due elementi positivi non implicano $i_+ \geq 2$. Ad esempio la segnatura di $\begin{pmatrix} 1 & 2 \\ 2 & 1 \end{pmatrix}$ è $(1, 1, 0)$.

Esempio 7.4.14. Calcoliamo la segnatura della matrice

$$S = \begin{pmatrix} 2 & 2 & 1 & 1 \\ 2 & 2 & 1 & 1 \\ 1 & 1 & 0 & 0 \\ 1 & 1 & 0 & 0 \end{pmatrix}.$$

Notiamo che $\text{rk} S = 2$ e quindi $i_0 = 4 - 2 = 2$. Le segnature possibili sono:

$$(2, 0, 2), \quad (1, 1, 2), \quad (0, 2, 2).$$

Ci sono elementi positivi sulla diagonale, quindi $i_+ > 0$ e la terza possibilità è esclusa. Ci resta da capire quale delle prime due segnature sia quella giusta. Se consideriamo il piano $W = \text{Span}(e_2, e_3)$, notiamo che la matrice associata alla restrizione $g|_W$ nella base e_2, e_3 è la sottomatrice quadrata

$$\begin{pmatrix} 2 & 1 \\ 1 & 0 \end{pmatrix}.$$

Questa ha determinante negativo: la segnatura di $g|_W$ è $(1, 1, 0)$. Quindi W contiene una retta su cui g è definito positivo e una su cui g è definito negativo. Quindi $i_- \geq 1$ e l'unica segnatura possibile è $(1, 1, 2)$.

Osservazione 7.4.15. Sia V uno spazio vettoriale con prodotto scalare g avente segnatura (i_+, i_-, i_0) . Se $W \subset V$ è un sottospazio, indichiamo con (i_+^W, i_-^W, i_0^W) la segnatura della restrizione $g|_W$. Dalle definizioni otteniamo

$$i_+^W \leq i_+, \quad i_-^W \leq i_-.$$

Non è vero che $i_0^W \leq i_0$ in generale (esercizio: costruisci un esempio).

7.4.7. Criterio di Jacobi. Introduciamo un criterio semplice per determinare se una matrice $n \times n$ simmetrica S sia definita positiva. Per ogni $i = 1, \dots, n$ sia d_i il determinante del minore S_i di S ottenuto considerando solo le prime i righe e le prime i colonne di S .

Lemma 7.4.16 (Criterio di Jacobi). *Una matrice S è definita positiva se e solo se $d_i > 0$ per ogni i .*

Dimostrazione. La matrice S definisce il prodotto scalare g_S su $\mathbb{R}^n$ e S_i è la matrice associata alla restrizione di g sul sottospazio coordinato $V_i = \text{Span}(e_1, \dots, e_i)$ rispetto alla base $e_1, \dots, e_i$. Se S è definita positiva, anche $g|_{V_i}$ lo è, ed in particolare deve valere $d_i = \det S_i > 0$ per ogni i .

Mostriamo ora l'implicazione inversa: supponiamo che $d_i > 0$ per ogni i e deduciamo che g è definito positivo. Dimostriamo il lemma per induzione su n . Se $n = 1$, la matrice è un numero positivo d_1 e siamo a posto. Supponiamo il lemma vero per $n - 1$ e lo dimostriamo per n .

Per l'ipotesi induttiva la restrizione $g|_{V_{n-1}}$ è definita positiva. Consideriamo la segnatura (i_+, i_-, i_0) di g . Siccome la restrizione $g|_{V_{n-1}}$ è definita positiva, abbiamo $i_+ \geq n - 1$ e le possibilità per la segnatura di g sono

$$(n, 0, 0), \quad (n - 1, 1, 0), \quad (n - 1, 0, 1).$$

Siccome $d_n > 0$, l'unica possibile è la prima. $\square$

Corollario 7.4.17. *Una matrice S è definita negativa se e solo se $d_i < 0$ per ogni i dispari e $d_i > 0$ per ogni i pari.*

Dimostrazione. La matrice S è definita negativa se e solo se $-S$ è definita positiva. A questo punto basta notare che $\det(-A) = (-1)^k \det A$ per ogni matrice A di taglia $k \times k$. $\square$

Con una dimostrazione analoga si dimostra anche questo criterio:

Lemma 7.4.18. *Se $d_i \neq 0$ per ogni i , allora la segnatura di S è $(i_+, i_-, 0)$ dove i_+ è il numero di permanenze di segno nella successione*

$$1, d_1, \dots, d_n$$

e i_- è il numero di cambiamenti di segno.

Dimostrazione. La dimostrazione è per induzione su n come quella del Lemma 7.4.16 ed è lasciata per esercizio. $\square$

Notiamo che il criterio non si applica se $d_i = 0$ per qualche i .

Esercizio 7.4.19. Verifica il criterio appena enunciato quando S è una matrice diagonale invertibile (cioè senza zeri sulla diagonale).

Esempio 7.4.20. Studiamo la segnatura della matrice

$$\begin{pmatrix} 1 & 0 & a \\ 0 & 2 & 0 \\ a & 0 & 3 \end{pmatrix}$$

al variare di $a \in \mathbb{R}$. Otteniamo

$$d_1 = 1, \quad d_2 = 2, \quad d_3 = 6 - 2a^2.$$

Se $|a| > \sqrt{3}$ allora $d_3 < 0$ e quindi la segnatura è $(2, 1, 0)$. Se $|a| < \sqrt{3}$ allora $d_3 > 0$ e quindi la segnatura è $(3, 0, 0)$. Se $a = \pm\sqrt{3}$ il criterio di Jacobi non si applica e dobbiamo usare altri metodi. La matrice S ha rango 2 e quindi $i_0 = 3 - 2 = 1$. D'altra parte la restrizione di g al piano

coordinato $W = \text{Span}(e_1, e_2)$ ha come matrice associata $\begin{pmatrix} 1 & 0 \\ 0 & 2 \end{pmatrix}$ che è definita positiva: quindi $g|_W$ è definita positiva e $i_+ \geq 2$. L'unico caso possibile è $(2, 0, 1)$. Riassumendo, la segnatura è

- $(3, 0, 0)$ per $|a| < \sqrt{3}$,
- $(2, 0, 1)$ per $|a| = \sqrt{3}$,
- $(2, 1, 0)$ per $|a| > \sqrt{3}$.

7.5. Isometrie

Ogni volta che in matematica si definiscono delle strutture come gruppi, spazi vettoriali, prodotti scalari, esistono delle opportune trasformazioni che preservano interamente queste strutture: queste trasformazioni sono chiamate *isomorfismi* per i gruppi e per gli spazi vettoriali, e *isometrie* per i prodotti scalari.

7.5.1. Definizione. Siano V e W spazi vettoriali dotati ciascuno di un prodotto scalare. Una *isometria* è un isomorfismo $T: V \rightarrow W$ tale che

$$(9) \quad \langle v, w \rangle = \langle T(v), T(w) \rangle \quad \forall v, w \in V.$$

Brevemente, un'isometria è un isomorfismo che preserva il prodotto scalare.

Esempio 7.5.1. L'identità $\text{id}: V \rightarrow V$ è un'isometria.

Esempio 7.5.2. La riflessione rispetto all'origine

$$r: V \longrightarrow V, \quad r(v) = -v$$

è sempre un'isometria, infatti

$$\langle r(v), r(w) \rangle = \langle -v, -w \rangle = (-1)^2 \langle v, w \rangle = \langle v, w \rangle \quad \forall v, w \in V.$$

Esempio 7.5.3 (Riflessione ortogonale). Sia $W \subset V$ un sottospazio su cui la restrizione $g|_W$ non sia degenere. Per il Teorema 7.3.12 lo spazio V si decompone in una somma diretta ortogonale:

$$V = W \oplus W^\perp.$$

Ricordiamo dall'Esempio 4.1.12 che in presenza di una somma diretta $V = W \oplus W^\perp$ è definita una riflessione $r_W: V \rightarrow V$ che opera nel modo seguente:

$$r_W(w + w') = w - w'$$

per ogni $w \in W$ e $w' \in W^\perp$. In questo contesto la riflessione r_W è detta *riflessione ortogonale* perché i due sottospazi W e $W^\perp$ sono ortogonali.

La riflessione ortogonale r_W è una isometria: se prendiamo due vettori generici $v + v'$ e $w + w'$ di V con $v, w \in W$ e $v', w' \in W^\perp$ otteniamo

$$\langle v + v', w + w' \rangle = \langle v, w \rangle + \langle v', w' \rangle,$$

$$\langle r_W(v + v'), r_W(w + w') \rangle = \langle v - v', w - w' \rangle = \langle v, w \rangle + \langle v', w' \rangle.$$

Abbiamo usato che $\langle v, w' \rangle = \langle v', w \rangle = 0$ perché $v, w \in W$ e $v', w' \in W^\perp$.

Osservazione 7.5.4. Notiamo che l'inversa T^{-1} di un'isometria è anch'essa un'isometria, e la composizione $T \circ T'$ di due isometrie è una isometria. Le isometrie $V \rightarrow V$ da uno spazio vettoriale V munito di prodotto scalare in sé formano un gruppo con l'operazione di composizione.

7.5.2. Con le basi. Per mostrare che un isomorfismo $T: V \rightarrow W$ è una isometria dovremmo verificare la condizione (9) su ogni coppia $v, w \in V$. In realtà è sufficiente farlo per gli elementi di una base fissata. Sia $\mathcal{B} = \{v_1, \dots, v_n\}$ una base di V .

Lemma 7.5.5. *Un isomorfismo T è un'isometria se e solo se*

$$(10) \quad \langle v_i, v_j \rangle = \langle T(v_i), T(v_j) \rangle \quad \forall i, j.$$

Dimostrazione. Chiaramente (9) implica (10), mostriamo adesso il contrario. Per due generici v e w scriviamo

$$v = \lambda_1 v_1 + \dots + \lambda_n v_n, \quad w = \mu_1 v_1 + \dots + \mu_n v_n$$

e usando la bilinearità del prodotto scalare e la linearità di T troviamo

$$\langle v, w \rangle = \sum_{i=1}^n \sum_{j=1}^n \lambda_i \mu_j \langle v_i, v_j \rangle,$$

$$\langle T(v), T(w) \rangle = \sum_{i=1}^n \sum_{j=1}^n \lambda_i \mu_j \langle T(v_i), T(v_j) \rangle.$$

Quindi (10) implica (9). $\square$

Corollario 7.5.6. *Sia $v_1, \dots, v_n$ una base ortonormale per V . Una funzione lineare $T: V \rightarrow W$ è una isometria $\iff$ anche le immagini $T(v_1), \dots, T(v_n)$ formano una base ortonormale per W dello stesso tipo (con questo intendiamo che $\langle v_i, v_i \rangle = \langle T(v_i), T(v_i) \rangle$ per ogni i).*

Possiamo dire quindi che le isometrie sono precisamente le applicazioni lineari che mandano basi ortonormali in basi ortonormali dello stesso tipo.

Esempio 7.5.7. Forniamo un'altra dimostrazione che la riflessione ortogonale r_W dell'Esempio 7.5.3 è una isometria. Siano $v_1, \dots, v_k$ e $v_{k+1}, \dots, v_n$ basi ortonormali di W e $W^\perp$. Assieme formano una base ortonormale $v_1, \dots, v_n$ di V . La riflessione r_W manda questa base nella base

$$v_1, \dots, v_k, -v_{k+1}, \dots, -v_n$$

che è anch'essa ortonormale e dello stesso tipo.

7.5.3. Con le matrici. Cerchiamo adesso di tradurre la definizione di isometria nel linguaggio concreto delle matrici. Siano V e V' spazi vettoriali muniti di prodotti scalari g e g' e di basi $\mathcal{B}$ e $\mathcal{B}'$ rispettivamente. Consideriamo un isomorfismo $T: V \rightarrow V'$. Siano

$$S = [g]_{\mathcal{B}}, \quad S' = [g']_{\mathcal{B}'}, \quad A = [T]_{\mathcal{B}'}^{\mathcal{B}}$$

le matrici associate a tutti gli attori in scena.

Proposizione 7.5.8. *L'isomorfismo T è una isometria se e solo se*

$$S = {}^tAS'A.$$

Dimostrazione. Sappiamo che T è una isometria se e solo se $g(v_i, v_j) = g'(T(v_i), T(v_j))$ per ogni i, j . Sappiamo anche che

$$S_{ij} = g(v_i, v_j).$$

Ricordiamo che la colonna i -esima di A è $A^i = [T(v_i)]_{\mathcal{B}'}$. Quindi

$$({}^tAS'A)_{ij} = ({}^tA^i)S'A^j = g'(T(v_i), T(v_j)).$$

Ne deduciamo che T è una isometria se e solo se $S = {}^tAS'A$. $\square$

In particolare, consideriamo il caso in cui V sia munito di un prodotto scalare g e $T: V \rightarrow V$ sia un isomorfismo. Siano $S = [g]_{\mathcal{B}}$ e $A = [T]_{\mathcal{B}}^{\mathcal{B}}$.

Corollario 7.5.9. *L'isomorfismo T è una isometria se e solo se*

$$S = {}^tASA.$$

Consideriamo adesso un caso ancora più specifico:

Corollario 7.5.10. *Sia g definito positivo e $\mathcal{B}$ una base ortonormale. L'isomorfismo T è una isometria se e solo se*

$${}^tAA = I_n.$$

Dimostrazione. Con queste ipotesi $S = [g]_{\mathcal{B}} = I_n$, quindi la relazione $S = {}^tASA$ diventa $I_n = {}^tAA$. $\square$

7.5.4. Matrici ortogonali. Consideriamo il caso che ci sta più a cuore, quello di $\mathbb{R}^n$ munito del suo prodotto scalare euclideo. Sia $A \in M(n)$ una matrice quadrata. Consideriamo l'endomorfismo $L_A: \mathbb{R}^n \rightarrow \mathbb{R}^n$.

Corollario 7.5.11. *L'endomorfismo L_A è una isometria $\iff {}^tAA = I_n$.*

Dimostrazione. La matrice A è la matrice associata a L_A rispetto alla base canonica, che è ortonormale: quindi segue dal Corollario 7.5.10. $\square$

Una matrice $A \in M(n)$ a coefficienti reali tale che ${}^tAA = I_n$ è detta *ortogonale*. Studieremo le matrici ortogonali più approfonditamente nel prossimo capitolo. Per adesso ci limitiamo a mostrare alcuni esempi.

Esempio 7.5.12. Le rotazioni e riflessioni ortogonali di $\mathbb{R}^2$ descritte nelle Definizioni 4.4.14 e 4.4.16 sono entrambe isometrie. Si verifica facilmente infatti che le matrici

$$\text{Rot}_{\theta} = \begin{pmatrix} \cos \theta & -\sin \theta \\ \sin \theta & \cos \theta \end{pmatrix}, \quad \text{Rif}_{\theta} = \begin{pmatrix} \cos \theta & \sin \theta \\ \sin \theta & -\cos \theta \end{pmatrix}$$

sono entrambe ortogonali. Anche la rotazione intorno all'asse z descritta nella Sezione 4.4.11 è determinata dalla matrice

$$A = \begin{pmatrix} \cos \theta & -\sin \theta & 0 \\ \sin \theta & \cos \theta & 0 \\ 0 & 0 & 1 \end{pmatrix} = \begin{pmatrix} \text{Rot}_{\theta} & 0 \\ 0 & 1 \end{pmatrix}.$$

è una isometria perché A è una matrice ortogonale.

7.5.5. Trasformazioni di Lorentz. Nella relatività di Einstein, l'universo in assenza di materia è modellizzato (localmente) come uno *spazio-tempo di Minkowski*. Questo è lo spazio euclideo $\mathbb{R}^4$ con variabili x, y, z, t , munito però non del prodotto scalare euclideo, ma bensì del *prodotto scalare lorentziano* g_S determinato dalla matrice diagonale

$$S = \begin{pmatrix} 1 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 \\ 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & -c^2 \end{pmatrix}$$

dove c è la velocità della luce. Notiamo che la sua segnatura è $(3, 1, 0)$: il prodotto scalare è non degenere ma non è definito positivo.

Esercizio 7.5.13. Sia $v \in (-c, c)$ un numero. Il *fattore di Lorentz* è

$$\gamma = \frac{1}{\sqrt{1 - \frac{v^2}{c^2}}}.$$

Considera la *matrice di trasformazione di Lorentz*

$$A = \begin{pmatrix} 1 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 \\ 0 & 0 & \gamma & -\gamma v \\ 0 & 0 & -\gamma \frac{v}{c^2} & \gamma \end{pmatrix}.$$

Dimostra che $L_A: \mathbb{R}^4 \rightarrow \mathbb{R}^4$ è isometria per il prodotto scalare lorentziano.

7.5.6. Spazi isometrici. Due spazi vettoriali V e V' muniti di prodotti scalari g e g' sono *isometrici* se esiste una isometria $T: V \rightarrow V'$. La segnatura caratterizza completamente gli spazi isometrici.

Proposizione 7.5.14. *Due spazi vettoriali V e V' con prodotti scalari g e g' sono isometrici $\iff$ hanno la stessa segnatura.*

Dimostrazione. Siano $\mathcal{B}$ e $\mathcal{B}'$ basi per V e V' . Scriviamo $S = [g]_{\mathcal{B}}$ e $S' = [g']_{\mathcal{B}'}$. Se esiste una isometria $T: V \rightarrow V'$, allora con $A = [T]_{\mathcal{B}'}^{\mathcal{B}}$ otteniamo ${}^tAS'A = S$, quindi S e S' sono congruenti, quindi hanno la stessa segnatura.

Anche le implicazioni opposte sono vere: se S e S' hanno la stessa segnatura sono congruenti, quindi esiste un A tale che $AS'A = S$ e possiamo interpretare A come matrice associata ad una isometria. $\square$

Corollario 7.5.15. *Gli spazi vettoriali con un prodotto scalare definito positivo di dimensione n sono tutti isometrici fra loro. In particolare sono isometrici a $\mathbb{R}^n$ con il prodotto scalare euclideo.*

Gli spazi di dimensione n muniti di un prodotto scalare definito positivo sono tutti isometrici fra loro e quindi in un certo senso hanno tutti le stesse proprietà intrinseche. Studieremo più approfonditamente questi spazi nel prossimo capitolo.

Esercizi

Esercizio 7.1. Sia $g(x, y) = {}^t x S y$ il prodotto scalare su $\mathbb{R}^2$ definito da $S = \begin{pmatrix} -1 & 0 \\ 0 & 1 \end{pmatrix}$. Scrivi la matrice associata a g nella base

$$\mathcal{B} = \left\{ \begin{pmatrix} 1 \\ 1 \end{pmatrix}, \begin{pmatrix} 1 \\ -1 \end{pmatrix} \right\}.$$

Verifica che la matrice che ottieni è congruente a S .

Esercizio 7.2. Sia $V = M(n, \mathbb{R})$. Considera le funzioni

$$f: V \times V \longrightarrow \mathbb{R}, \quad f(A, B) = \text{tr}(AB),$$

$$g: V \times V \longrightarrow \mathbb{R}, \quad g(A, B) = \text{tr}({}^t AB).$$

- (1) Mostra che f e g sono prodotti scalari.
- (2) Mostra che g è definito positivo.
- (3) Il prodotto scalare f è definito positivo?
- (4) Determina una base ortonormale per g .
- (5) Trova la segnatura di f .

Prova a rispondere a queste domande nel caso $n = 2$. È anche un utile esercizio scrivere le matrici associate ai prodotti scalari rispetto alla base canonica di $M(2)$.

Esercizio 7.3. Considera il prodotto scalare su $\mathbb{R}_2[x]$ dato da

$$\langle p, q \rangle = p(1)q(1) + p(-1)q(-1).$$

Scrivi la matrice associata al prodotto scalare rispetto alla base canonica. Determina il radicale di questo prodotto scalare.

Esercizio 7.4. Sia g un prodotto scalare su V . Mostra che:

- (1) se tutti i vettori sono isotropi, allora il prodotto scalare è banale, cioè $\langle v, w \rangle = 0$ per ogni $v, w \in V$;
- (2) se g è indefinito, cioè esistono v e w con $g(v, v) > 0$ e $g(w, w) < 0$, allora esiste un vettore isotropo non banale.

Esercizio 7.5. Siano D_1 e D_2 due matrici diagonali, entrambe aventi i termini sulla diagonale tutti strettamente positivi. Mostra che le due matrici D_1 e D_2 sono congruenti. Determina concretamente una M tale che $D_1 = {}^t M D_2 M$.

Esercizio 7.6. Sia V uno spazio vettoriale di dimensione finita.

- Mostra che per ogni $v \in V$ non nullo esiste un prodotto scalare g definito positivo su V tale che $g(v, v) = 1$.
- Mostra che per ogni coppia $v, w \in V$ di vettori indipendenti esiste un prodotto scalare g definito positivo su V tale che $g(v, w) = 0$.

Esercizio 7.7. Considera il prodotto scalare su $\mathbb{R}_2[x]$ dato da

$$\langle p, q \rangle = p(0)q(0) + p'(0)q'(0) + p''(0)q''(0)$$

dove p' e p'' sono i polinomi ottenuti come derivata prima e seconda di p . Considera l'insieme $W \subset \mathbb{R}_2[x]$ formato da tutti i polinomi $p(x)$ che si annullano in $x = 1$, cioè tali che $p(1) = 0$.

- Mostra che W è un sottospazio vettoriale e calcola la sua dimensione.
- Determina il sottospazio ortogonale $W^\perp$.

Esercizio 7.8. Considera su $M(2)$ il prodotto scalare $\langle A, B \rangle = \text{tr}({}^t AB)$ già esaminato in un esercizio precedente.

- Se $W \subset M(2)$ è un sottospazio vettoriale, è sempre vero che $M(2) = W \oplus W^\perp$?
- Calcola $\dim W$ e determina $W^\perp$ nei casi seguenti:
 - (1) W è il sottospazio formato dalle matrici simmetriche.
 - (2) W è il sottospazio formato dalle matrici diagonali.

Esercizio 7.9. Considera il prodotto scalare su $\mathbb{R}_2[x]$ dato da

$$\langle p, q \rangle = p(0)q(0) - p(1)q(1).$$

- (1) Determina il radicale del prodotto scalare.
- (2) Determina un polinomio isotropo che non sia contenuto nel radicale.
- (3) I vettori isotropi formano un sottospazio vettoriale?

Esercizio 7.10. Considera il prodotto scalare g_S su $\mathbb{R}^3$ con

$$S = \begin{pmatrix} 0 & 0 & 1 \\ 0 & 1 & 0 \\ 1 & 0 & 0 \end{pmatrix}$$

- (1) Il prodotto scalare è degenere?
- (2) Esiste un sottospazio vettoriale $V \subset \mathbb{R}^3$ di dimensione 1 tale che $V \subset V^\perp$?
- (3) Determinare i vettori isotropi. Formano un sottospazio vettoriale?
- (4) Determinare un piano vettoriale $W \subset \mathbb{R}^3$ tale che la restrizione del prodotto scalare a W sia definita positiva.
- (5) Determinare un piano vettoriale $W \subset \mathbb{R}^3$ tale che la restrizione del prodotto scalare a W sia degenere.

Esercizio 7.11. Sia V spazio vettoriale di dimensione n , dotato di un prodotto scalare g . Siano $U, W \subset V$ sottospazi vettoriali qualsiasi. Mostra i fatti seguenti:

- (1) Se $U \subset W$, allora $W^\perp \subset U^\perp$.
- (2) $U \subset (U^\perp)^\perp$. Se g è non degenere, allora $U = (U^\perp)^\perp$.

Esercizio 7.12. Trova una base ortogonale per i prodotti scalari g_S su $\mathbb{R}^2$ e $\mathbb{R}^3$ determinati dalle matrici S seguenti:

$$\begin{pmatrix} 1 & 1 \\ 1 & 2 \end{pmatrix}, \quad \begin{pmatrix} 0 & -1 \\ -1 & 0 \end{pmatrix}, \quad \begin{pmatrix} 0 & 0 & 1 \\ 0 & 1 & 0 \\ 1 & 0 & 0 \end{pmatrix}$$

Esercizio 7.13. Determina la segnatura delle matrici, al variare di $t \in \mathbb{R}$:

$$\begin{pmatrix} 1 & 1 & 1 \\ 1 & 1 & 1 \\ 1 & 1 & t \end{pmatrix}, \quad \begin{pmatrix} 0 & 1 & 0 & 0 \\ 1 & 0 & 0 & 0 \\ 0 & 0 & 0 & 1 \\ 0 & 0 & 1 & 0 \end{pmatrix}.$$

Esercizio 7.14. Considera il prodotto scalare su $\mathbb{R}_3[x]$ dato da

$$\langle p, q \rangle = p(1)q(-1) + p(-1)q(1).$$

- Trova la segnatura.
- Determina il radicale.
- Determina $W^\perp$ dove $W = \text{Span}(x + x^2)$.

Esercizio 7.15. Determina quali fra queste matrici sono congruenti:

$$\begin{pmatrix} 7 & 0 & 0 \\ 0 & 0 & -\sqrt{2} \\ 0 & -\sqrt{2} & 0 \end{pmatrix}, \quad \begin{pmatrix} 1 & 2 & 3 \\ 2 & 3 & 2 \\ 3 & 2 & 1 \end{pmatrix}, \quad \begin{pmatrix} 1 & 1 & 1 \\ 1 & 0 & 1 \\ 1 & 1 & 0 \end{pmatrix}, \quad \begin{pmatrix} 0 & 0 & \sqrt{17} \\ 0 & 1 & 0 \\ \sqrt{17} & 0 & 0 \end{pmatrix}.$$

Esercizio 7.16. Calcola la segnatura delle matrici al variare di $\alpha \in \mathbb{R}$:

$$\begin{pmatrix} 1 & 2 & 1 \\ 2 & 1 & 1 \\ 1 & 1 & 1 \end{pmatrix}, \quad \begin{pmatrix} 1 & 0 & 0 \\ 0 & 8 & -4 \\ 0 & -4 & 2 \end{pmatrix},$$

$$\begin{pmatrix} 1 & 2 & 1 \\ 2 & 0 & \alpha \\ 1 & \alpha & \alpha^2 \end{pmatrix}, \quad \begin{pmatrix} \alpha & \alpha+1 & \alpha+2 \\ \alpha+1 & \alpha+2 & \alpha+1 \\ \alpha+2 & \alpha+1 & \alpha \end{pmatrix}.$$

Esercizio 7.17. Considera il prodotto scalare su $\mathbb{R}^3$ dato dalla matrice

$$S = \begin{pmatrix} 1 & 2 & -1 \\ 2 & 3 & 1 \\ -1 & 1 & -8 \end{pmatrix}.$$

Sia inoltre

$$W = \{2x + y = z\}.$$

- (1) Calcola la segnatura di S .
- (2) Determina una base per $W^\perp$.

Esercizio 7.18. Considera il prodotto scalare su $\mathbb{R}^4$ dato dalla matrice

$$S = \begin{pmatrix} 1 & 2 & 0 & 3 \\ 2 & -1 & 0 & 1 \\ 0 & 0 & 2 & -2 \\ 3 & 1 & -2 & 6 \end{pmatrix}$$

Sia inoltre

$$W = \text{Span} \left(\begin{pmatrix} 0 \\ 1 \\ 0 \\ 1 \end{pmatrix}, \begin{pmatrix} 1 \\ 0 \\ -1 \\ 0 \end{pmatrix} \right).$$

- (1) Calcola la segnatura di S .
- (2) Determina una base per $W^\perp$.

Esercizio 7.19. Determina un prodotto scalare g su $\mathbb{R}^2$ tale che $\begin{pmatrix} 0 \\ 1 \end{pmatrix}$ e $\begin{pmatrix} 1 \\ 1 \end{pmatrix}$ formino una base ortonormale. Scrivi la matrice associata a g rispetto alla base canonica di $\mathbb{R}^2$.

Esercizio 7.20. Sia g un prodotto scalare su $\mathbb{R}^3$ con segnatura $(2, 1, 0)$. Dimostra che esistono dei piani $W \subset \mathbb{R}^3$ tali che $g|_W$ abbia segnatura

$$(2, 0, 0), \quad (1, 1, 0), \quad (1, 0, 1)$$

mentre non esistono dei piani $W \subset \mathbb{R}^3$ tali che $g|_W$ abbia segnatura

$$(0, 2, 0), \quad (0, 1, 1), \quad (0, 0, 2).$$

Complementi

7.1. Tensori. I tensori sono oggetti matematici che generalizzano vettori, endomorfismi e prodotti scalari. Sono ampiamente usati in fisica ed hanno un ruolo fondamentale nella relatività generale di Einstein. Questi possono essere introdotti in due modi.

Punto di vista fisico. Il primo approccio, che potremmo chiamare "fisico", consiste nel definire un tensore come una tabella di numeri che può

avere dimensione arbitraria. Questa definizione include i vettori numerici di lunghezza n , le matrici $n \times n$, i cubi $n \times n \times n$, eccetera. Queste tabelle multidimensionali contengono rispettivamente n , n^2 e n^3 coefficienti.

La definizione però non finisce qui (sarebbe troppo facile...) perché in realtà un tensore è un oggetto che dipende dal sistema di riferimento usato per il problema fisico in esame: se modifichiamo il sistema di riferimento, i numeri nella tabella devono cambiare secondo alcune formule fissate.

Più precisamente, sia V uno spazio vettoriale su $\mathbb{K}$ munito di una base $\mathcal{B}$. Un *tensore* di tipo (h, k) per V è una tabella di n^{h+k} numeri

$$T_{j_1 \dots j_k}^{i_1 \dots i_h} \in \mathbb{K}$$

che variano secondo $h+k$ indici $1 \leq i_1, \dots, i_h, j_1, \dots, j_k \leq n$, i primi h scritti in alto e gli altri k in basso. Se usiamo un'altra base $\mathcal{B}'$, i numeri devono cambiare come segue. Siano $M = [\text{id}]_{\mathcal{B}'}^{\mathcal{B}}$, $N = [\text{id}]_{\mathcal{B}}^{\mathcal{B}'}$ le matrici di cambiamento di base, che sono una l'inversa dell'altra, e scriviamo $M_j^i = M_{ij}$, $N_j^i = N_{ij}$. I coefficienti del tensore nella nuova base devono essere i seguenti:

$$(11) \quad \hat{T}_{j_1 \dots j_k}^{i_1 \dots i_h} = \sum M_{i_1}^{i_1} \dots M_{i_h}^{i_h} N_{j_1}^{m_1} \dots N_{j_k}^{m_k} T_{m_1 \dots m_k}^{i_1 \dots i_h}.$$

La somma va fatta su tutti gli indici $1 \leq i_1, \dots, i_h, m_1, \dots, m_k \leq n$, quindi ci sono n^{k+h} addendi. Questa formula sembra complicata perché è esposta nella più totale generalità: nei casi concreti i numeri h e k sono molto bassi e tutto risulta molto più semplice.

Vediamo adesso che effettivamente questa nozione generalizza (con un linguaggio un po' diverso) molte cose già viste in questo libro.

Le coordinate di un vettore $v \in V$ sono un tensore di tipo $(1, 0)$. Queste formano infatti un vettore numerico $T = [v]_{\mathcal{B}}$ che reagisce quando si cambia base esattamente come un tensore di tipo $(1, 0)$. Infatti (11) diventa

$$\hat{T}^i = \sum_{l=1}^n M_j^l T^l$$

e questo equivale a scrivere il prodotto matrice per vettore

$$[v]_{\mathcal{B}'} = M[v]_{\mathcal{B}}$$

come prescritto dalla Proposizione 4.3.13. Passando in coordinate è possibile quindi interpretare i vettori di V come tensori di tipo $(1, 0)$.

Questo è solo l'inizio: le coordinate di un endomorfismo $f: V \rightarrow V$ sono un tensore $T = [f]_{\mathcal{B}}^{\mathcal{B}}$ di tipo $(1, 1)$. Infatti

$$\hat{T}_j^i = \sum_{l,m} M_j^l N_m^m T_m^l = \sum_{l,m} M_j^l T_m^l N_j^m$$

equivale a scrivere il prodotto fra matrici

$$[f]_{\mathcal{B}'}^{\mathcal{B}'} = M[f]_{\mathcal{B}}^{\mathcal{B}} N = N^{-1}[f]_{\mathcal{B}}^{\mathcal{B}} N$$

che descrive proprio come reagisce $[f]_B^B$ quando si cambia base. Le coordinate di un prodotto scalare g su V sono un tensore $T = [g]_B$ di tipo $(0, 2)$, perché

$$\hat{T}_{i,j} = \sum_{l,m} N_l^i N_j^m T_{l,m} = \sum_{l,m} N_l^i T_{l,m} N_j^m$$

equivale al prodotto fra matrici

$$[g]_{B'} = {}^t N [g]_B N.$$

Riassumendo: vettori, endomorfismi e prodotti scalari sono interpretabili (tramite le loro coordinate) come tensori di tipo $(1, 0)$, $(1, 1)$ e $(0, 2)$.

Cos'è un tensore di tipo $(0, 1)$? Sembra un vettore di V , ma in realtà le sue coordinate cambiano con N invece che con M . Si può interpretare un tale oggetto, detto *covettore*, come un vettore dello spazio duale V^* di V introdotto nella Sezione 4.1.

Punto di vista matematico. Introduciamo adesso l'approccio "matematico". Sia V uno spazio vettoriale di dimensione n . Ricordiamo dalla Sezione 4.1 che lo *spazio duale* V^* è lo spazio di tutte le applicazioni lineari $V \rightarrow \mathbb{K}$. Gli elementi di V e V^* sono chiamati rispettivamente vettori e covettori. Lo spazio biduale V^{**} è canonicamente identificato con V .

Abbiamo bisogno di estendere le nozioni di linearità e bilinearità a quella più generale di *multilinearità*. Sia $V_1 \times \cdots \times V_m$ il prodotto cartesiano di alcuni spazi vettoriali. Una funzione

$$f: V_1 \times \cdots \times V_m \longrightarrow \mathbb{K}$$

è *multilineare* se

$$f(v_1, \dots, v_{i-1}, \lambda v + \mu w, v_{i+1}, \dots, v_m) =$$

$$\lambda f(v_1, \dots, v_{i-1}, v, v_{i+1}, \dots, v_m) + \mu f(v_1, \dots, v_{i-1}, w, v_{i+1}, \dots, v_m)$$

per ogni $i = 1, \dots, m$, $v_j \in V_j$, $v, w \in V_i$ e $\lambda, \mu \in \mathbb{K}$. Notiamo che se $m = 1$ oppure $m = 2$ questa definizione è quella usuale di linearità e bilinearità. A parte la pesantezza notazionale, *multilineare* vuol dire semplicemente lineare in ogni componente.

Definizione 7.5.16. Un *tensore* di tipo (h, k) è una funzione multilineare

$$T: \underbrace{V^* \times \cdots \times V^*}_h \times \underbrace{V \times \cdots \times V}_k \longrightarrow \mathbb{K}.$$

Alcuni oggetti già visti in questo libro sono tensori di un tipo preciso:

- Un tensore di tipo $(1, 0)$ è una funzione lineare $V^* \rightarrow \mathbb{K}$, cioè un elemento di $V^{**} = V$. Quindi è un *vettore*.
- Un tensore di tipo $(0, 1)$ è una funzione lineare $V \rightarrow \mathbb{K}$, cioè un elemento di V^* . Quindi è un *covettore*.
- Un tensore di tipo $(0, 2)$ è una funzione bilineare $V \times V \rightarrow \mathbb{K}$. I *prodotti scalari* ne sono un esempio.

- Un tensore di tipo $(1, 1)$ è in realtà un *endomorfismo* di V .

Per giustificare l'ultima frase, notiamo che un tensore T di tipo $(1, 1)$ induce una funzione lineare $L: V \rightarrow V^{**}$ nel modo seguente:

$$L(v)(w) = T(v, w) \quad \forall v \in V, \forall w \in V^*.$$

Identificando V^{**} con V , la funzione L è effettivamente un endomorfismo di V . Con ragionamenti analoghi, vediamo che:

- Un tensore di tipo $(2, 0)$ è una funzione bilineare $V^* \times V^* \rightarrow \mathbb{K}$.
- Un tensore di tipo $(1, 2)$ è in realtà una funzione bilineare $V \times V \rightarrow V$.

Un tensore T di tipo $(1, 2)$ può essere interpretato come una funzione bilineare $V \times V \rightarrow V$. Questa è $L: V \times V \rightarrow V^{**}$ con

$$L(v, v')(w) = T(v, v', w).$$

Il *prodotto vettoriale* in $\mathbb{R}^3$ che introdurremo nella Sezione 9.1 è un esempio di tensore di tipo $(1, 2)$.

Se fissiamo una base, tutti questi concetti molto astratti possono essere letti in coordinate dove diventano delle più concrete tabelle di numeri. Sia $\mathcal{B} = \{v_1, \dots, v_n\}$ una base per V e $\mathcal{B}^* = \{v_1^*, \dots, v_n^*\}$ la base duale di V^* . L'esercizio seguente è una generalizzazione delle Proposizioni 4.1.18 e 7.2.5 e mostra che un tensore è completamente determinato dai valori che assume sulle basi. Sia T un tensore di tipo (h, k) .

Esercizio 7.5.17. Sia

$$T_{j_1 \dots j_k}^{i_1 \dots i_h} \in \mathbb{K}$$

una tabella $(k+h)$ -dimensionale di n^{k+h} scalari dipendenti da indici $i_1, \dots, i_h, j_1, \dots, j_k$ che variano tutti da 1 a n . Esiste un unico tensore T di tipo (h, k) tale che

$$T_{j_1 \dots j_k}^{i_1 \dots i_h} = T(v_{i_1}^*, \dots, v_{i_h}^*, v_{j_1}, \dots, v_{j_k})$$

per ogni $i_1, \dots, i_h, j_1, \dots, j_k$.

Gli scalari $T_{j_1 \dots j_k}^{i_1 \dots i_h}$ sono le *coordinate* di T rispetto alla base $\mathcal{B}$. In coordinate, vettori e covettori diventano vettori numerici, endomorfismi e prodotti scalari diventano matrici, e tensori più complicati diventano tabelle multidimensionali.

Applicazioni. A cosa servono i tensori? In fisica sono ampiamente usati per descrivere i campi elettromagnetici e gravitazionali. Ci ricordiamo dalla fisica delle superiori che un campo elettrico è un *campo vettoriale*: ad ogni punto dello spazio assegniamo un vettore in $\mathbb{R}^3$. Ci ricordiamo anche che le equazioni di Maxwell descrivono l'elettromagnetismo in una maniera unificata che ha però messo in discussione alla fine del XIX secolo lo spazio e il tempo come entità assolute e separate, ed ha aperto le porte verso la relatività ristretta di Einstein.

Nella formulazione adottata nel XX secolo, un campo elettromagnetico può essere descritto come un *campo tensoriale* in cui ad ogni punto

dello spaziotempo $\mathbb{R}^4$ assegniamo un tensore di tipo $(0, 2)$. Nella relatività generale di Einstein la gravità è responsabile della *curvatura* dello spaziotempo, un concetto complesso e profondo che è interamente codificato da un campo tensoriale di tipo $(1, 3)$ detto *tensore di Riemann*.

I campi tensoriali sono ampiamente usati anche in altri ambiti della fisica e dell'ingegneria, ad esempio nella meccanica del continuo.

Prodotti scalari definiti positivi

Nel capitolo precedente abbiamo introdotto i prodotti scalari nella loro generalità; in questo ci concentriamo sui prodotti scalari *definiti positivi*. Introduciamo per la prima volta alcuni concetti familiari della geometria euclidea: quello di *lunghezza* di un vettore, di *angolo* fra due vettori, e di *distanza* fra due punti. Tutte queste nozioni geometriche familiari sono presenti in qualsiasi spazio vettoriale munito di un prodotto scalare definito positivo.

Lo spazio vettoriale che geometricamente ci interessa di più è come sempre lo spazio euclideo n -dimensionale $\mathbb{R}^n$, con un occhio di riguardo per gli spazi $\mathbb{R}^2$ e $\mathbb{R}^3$ che sono utili a descrivere il mondo in cui viviamo. Continuiamo però ad avere una prospettiva unificante: i concetti geometrici appena descritti sono utili anche a studiare spazi più astratti, come ad esempio alcuni spazi di funzioni, ed hanno quindi importanti applicazioni anche in algebra ed in analisi.

8.1. Nozioni geometriche

Nella geometria euclidea e in fisica, il prodotto scalare fra due vettori del piano o dello spazio viene generalmente definito usando angoli e lunghezze: il *prodotto scalare* di due vettori v e w che formano un angolo acuto θ come nella Figura 8.1 è il numero

$$\langle v, w \rangle = \|v\| \cdot \|w\| \cos \theta.$$

In questa formula $\|v\|$ e $\|w\|$ sono le lunghezze di v e w . Nel caso in cui w abbia lunghezza unitaria, la formula diventa $\|v\| \cos \theta$ e può essere interpretata geometricamente come la lunghezza della proiezione ortogonale di v su w come nella Figura 8.1. Se w non è unitario, si deve moltiplicare questa lunghezza per $\|w\|$.

Nella geometria analitica adottata in questo libro preferiamo percorrere questa strada in senso opposto: introduciamo *prima* il prodotto scalare (come abbiamo fatto abbondantemente nel capitolo precedente) e *poi* sulla base di questo definiamo angoli, lunghezze e distanze. Questo approccio è preferibile perché si applica con lo stesso linguaggio ad un'ampia varietà di spazi vettoriali.

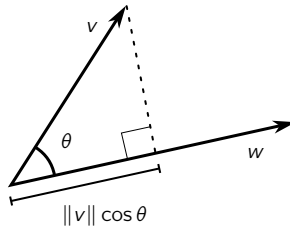


Figura 8.1. Se w ha lunghezza 1, il prodotto scalare di v e w è la lunghezza della proiezione di v su w , pari a $\|v\| \cos \theta$.

In tutto il capitolo V sarà sempre uno spazio vettoriale reale munito di un prodotto scalare definito positivo. L'esempio fondamentale è $\mathbb{R}^n$ con il prodotto scalare euclideo, e siamo particolarmente interessati a $\mathbb{R}^2$ e $\mathbb{R}^3$.

8.1.1. Norma. Sia V uno spazio vettoriale munito di un prodotto scalare definito positivo. Definiamo alcune nozioni geometriche molto naturali. Iniziamo con la *norma* di un vettore.

Definizione 8.1.1. La *norma* di un vettore $v \in V$ è il numero reale

$$\|v\| = \sqrt{\langle v, v \rangle}.$$

La norma di v , detta anche *modulo*, va interpretata come la *lunghezza* del vettore v . Nella definizione usiamo il fatto che il prodotto scalare è definito positivo e quindi ha senso la radice quadrata di $\langle v, v \rangle \geq 0$.

Proposizione 8.1.2. *Valgono le proprietà seguenti:*

- (1) $\|v\| > 0$ se $v \neq 0$ e $\|0\| = 0$,
- (2) $\|\lambda v\| = |\lambda| \|v\|$,
- (3) $|\langle v, w \rangle| \leq \|v\| \|w\|$,
- (4) $\|v + w\| \leq \|v\| + \|w\|$

per ogni $v, w \in V$ e ogni $\lambda \in \mathbb{R}$.

La disuguaglianza (3) è nota come *disuguaglianza di Cauchy-Schwarz* e la (4) è una versione della *disuguaglianza triangolare*.

Dimostrazione. I primi due punti seguono semplicemente dalla definizione $\|v\| = \sqrt{\langle v, v \rangle}$. Il punto (1) è verificato perché il prodotto scalare è definito positivo, il (2) si ottiene così:

$$\|\lambda v\| = \sqrt{\langle \lambda v, \lambda v \rangle} = \sqrt{\lambda^2 \langle v, v \rangle} = |\lambda| \sqrt{\langle v, v \rangle} = |\lambda| \|v\|.$$

Il punto (3) invece non è ovvio. Consideriamo due numeri $a, b \in \mathbb{R}$ e notiamo che per ogni $v, w \in V$ vale la disuguaglianza

$$\begin{aligned} 0 &\leq \|av + bw\|^2 = \langle av + bw, av + bw \rangle \\ &= a^2 \langle v, v \rangle + b^2 \langle w, w \rangle + 2ab \langle v, w \rangle = a^2 \|v\|^2 + b^2 \|w\|^2 + 2ab \langle v, w \rangle. \end{aligned}$$

Sostituendo $a = \|w\|^2$ e $b = -\langle v, w \rangle$ otteniamo

$$0 \leq \|w\|^4 \|v\|^2 + \langle v, w \rangle^2 \|w\|^2 - 2\|w\|^2 \langle v, w \rangle^2.$$

Dividendo per $\|w\|^2$ si ottiene

$$\|w\|^2 \|v\|^2 \geq \langle v, w \rangle^2$$

e la (3) è ottenuta prendendo le radici quadrate. La disuguaglianza (4) è un semplice corollario della (3):

$$\begin{aligned} \|v + w\|^2 &= \langle v + w, v + w \rangle = \|v\|^2 + \|w\|^2 + 2\langle v, w \rangle \\ &\leq \|v\|^2 + \|w\|^2 + 2\|v\|\|w\| = (\|v\| + \|w\|)^2. \end{aligned}$$

Anche qui basta prendere le radici per concludere. $\square$

Esempio 8.1.3. La norma di un vettore $x \in \mathbb{R}^n$ con il prodotto scalare euclideo è

$$\|x\| = \sqrt{x_1^2 + \cdots + x_n^2}.$$

Esempio 8.1.4. Consideriamo il prodotto scalare g_S su $\mathbb{R}^2$ definito dalla matrice $S = \begin{pmatrix} 2 & 1 \\ 1 & 1 \end{pmatrix}$. Notiamo che il prodotto scalare è definito positivo. La norma di un vettore x rispetto a questo prodotto scalare è

$$\|x\| = \sqrt{2x_1^2 + x_2^2 + 2x_1x_2}.$$

8.1.2. Applicazioni. Mostriamo alcune applicazioni importanti.

Osservazione 8.1.5. La disuguaglianza di Cauchy-Schwarz applicata a $\mathbb{R}^n$ con il prodotto scalare euclideo assume questa forma: presi $2n$ numeri reali qualsiasi $x_1, \dots, x_n, y_1, \dots, y_n \in \mathbb{R}$, vale sempre la disuguaglianza

$$\left| \sum_{i=1}^n x_i y_i \right| \leq \sqrt{\sum_{i=1}^n x_i^2 \sum_{j=1}^n y_j^2}.$$

Osservazione 8.1.6. Consideriamo lo spazio $C([a, b])$ delle funzioni continue $f: [a, b] \rightarrow \mathbb{R}$ con il prodotto scalare definito positivo dell'Esempio 7.1.25. La disuguaglianza di Cauchy-Schwarz per $C([a, b])$ assume questa forma: date due funzioni $f, g: [a, b] \rightarrow \mathbb{R}$ continue, vale sempre la disuguaglianza

$$\left| \int_a^b f(t)g(t)dt \right| \leq \sqrt{\int_a^b f(t)^2 dt \int_a^b g(t)^2 dt}.$$

Sia V uno spazio vettoriale munito di un prodotto scalare definito positivo. Elenchiamo un paio di semplici esercizi.

Esercizio 8.1.7. Per ogni $v, w \in V$ otteniamo

$$\langle v, w \rangle = \frac{1}{4}(\|v + w\|^2 - \|v - w\|^2).$$

Il prossimo ha un'immediata applicazione nella geometria del piano.

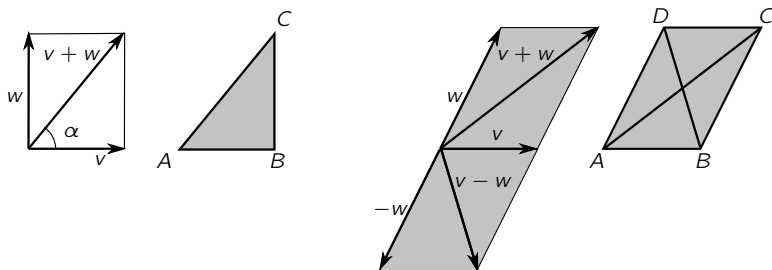


Figura 8.2. Un triangolo rettangolo e la legge del parallelogramma.

Esercizio 8.1.8 (Legge del parallelogramma). Per ogni $v, w \in V$ otteniamo

$$\|v + w\|^2 + \|v - w\|^2 = 2(\|v\|^2 + \|w\|^2).$$

Abbiamo dimostrato la *legge del parallelogramma*: per qualsiasi parallelogramma come nella Figura 8.2-(destra) vale la relazione

$$|AC|^2 + |BD|^2 = 2(|AB|^2 + |AD|^2).$$

Possiamo apprezzare subito i vantaggi di un approccio astratto ai prodotti scalari: con un linguaggio unificato e dei teoremi molto generali, in poche righe abbiamo dimostrato dei fatti non banali di algebra (disuguaglianze numeriche), analisi (integrali di funzioni) e geometria (parallelogrammi nel piano).

8.1.3. Angoli. A questo punto possiamo definire una nozione di angolo fra due vettori non nulli. Sia come sempre V uno spazio vettoriale munito di un prodotto scalare definito positivo.

Definizione 8.1.9. L'*angolo* fra due vettori $v, w \in V$ non nulli è il numero $\theta \in [0, \pi]$ per cui

$$\cos \theta = \frac{\langle v, w \rangle}{\|v\| \|w\|}.$$

Effettivamente per la disuguaglianza di Cauchy-Scharz il membro di destra è un numero $x \in [-1, 1]$ e quindi esiste un unico $\theta \in [0, \pi]$ per cui $\cos \theta = x$. In altre parole, si può usare la funzione *arcocoseno* e scrivere

$$\theta = \arccos \frac{\langle v, w \rangle}{\|v\| \|w\|}.$$

Osservazione 8.1.10. Tutte le funzioni trigonometriche come $\sin x$, $\cos x$, $\tan x$, $\arcsin x$, $\arccos x$, ecc. possono essere definite intrinsecamente come serie di potenze. Sempre usando le serie si dimostra la relazione fondamentale $\sin^2 x + \cos^2 x = 1$. Assumiamo una certa familiarità con queste funzioni trigonometriche, generalmente studiate nel corso di analisi.

Notiamo che θ è acuto, retto o ottuso (cioè minore di $\frac{\pi}{2}$, uguale a $\frac{\pi}{2}$ o maggiore di $\frac{\pi}{2}$) $\iff$ il prodotto scalare $\langle v, w \rangle$ è rispettivamente positivo,

nullo o negativo. Due vettori non nulli sono ortogonali $\iff$ formano un angolo retto.

Se v e w sono ortogonali come nella Figura 8.2-(sinistra) e α è l'angolo tra v e $v + w$, ritroviamo le formule familiari con cui normalmente si definiscono seno e coseno, ed il Teorema di Pitagora:

Proposizione 8.1.11. *Valgono le uguaglianze*

$$\begin{aligned}\|v\| &= \|v + w\| \cos \alpha, & \|w\| &= \|v + w\| \sin \alpha, \\ \|v + w\|^2 &= \|v\|^2 + \|w\|^2.\end{aligned}$$

Dimostrazione. Mostriamo l'ultima. Abbiamo

$$\|v + w\|^2 = \langle v + w, v + w \rangle = \langle v, v \rangle + \langle w, w \rangle + 2\langle v, w \rangle = \|v\|^2 + \|w\|^2.$$

Mostriamo la prima. Dalla definizione di angolo

$$\cos \alpha = \frac{\langle v, v + w \rangle}{\|v\| \|v + w\|} = \frac{\langle v, v \rangle + \langle v, w \rangle}{\|v\| \|v + w\|} = \frac{\|v\|^2}{\|v\| \|v + w\|} = \frac{\|v\|}{\|v + w\|}.$$

La seconda è lasciata per esercizio. $\square$

Abbiamo dimostrato il Teorema di Pitagora: per ogni triangolo rettangolo come nella Figura 8.2-(sinistra) vale la relazione

$$|AC|^2 = |AB|^2 + |BC|^2.$$

8.1.4. Distanze. Quando abbiamo una norma, è immediato definire una nozione di *distanza* fra due punti. Sia come sempre V uno spazio vettoriale reale dotato di un prodotto scalare definito positivo.

Quando interpretiamo V in modo più geometrico, può essere utile indicare i punti di V con le lettere maiuscole P, Q, R , eccetera. Possiamo anche definire il vettore

$$\overrightarrow{PQ} = Q - P.$$

Notiamo in particolare che

$$\overrightarrow{PQ} + \overrightarrow{QR} = \overrightarrow{PR}.$$

Definizione 8.1.12. La *distanza* fra due punti $p, q \in V$ è il numero

$$d(P, Q) = \|\overrightarrow{PQ}\| = \|Q - P\|.$$

La distanza ha tre proprietà basilari.

Proposizione 8.1.13. *Valgono i fatti seguenti per ogni $P, Q, R \in V$:*

- (1) $d(P, Q) > 0$ se $P \neq Q$ e $d(P, P) = 0$;
- (2) $d(P, Q) = d(Q, P)$;
- (3) $d(P, R) \leq d(P, Q) + d(Q, R)$.

Dimostrazione. Le proprietà (1) e (2) sono facili conseguenze delle proprietà della norma. Per la (3) abbiamo

$$d(P, R) = \|\overrightarrow{PR}\| = \|\overrightarrow{PQ} + \overrightarrow{QR}\| \leq \|\overrightarrow{PQ}\| + \|\overrightarrow{QR}\| = d(P, Q) + d(Q, R).$$

La dimostrazione è completa. $\square$

La proprietà (3) è la *disuguaglianza triangolare*; nel triangolo con vertici P, Q, R , la lunghezza di ogni lato è minore o uguale alla somma delle lunghezze degli altri due:

$$|PR| \leq |PQ| + |QR|.$$

8.1.5. Proiezione ortogonale. Una nozione importante in presenza di un prodotto scalare è quella di *ortogonalità*. Mostriamo adesso come sia possibile proiettare ortogonalmente vettori su altri vettori.

Sia come sempre V uno spazio vettoriale munito di un prodotto scalare definito positivo. Sia w un vettore non nullo di V e $U = \text{Span}(w)$ la retta generata da w . Sappiamo dal Teorema 7.3.12 che lo spazio V si decompone in somma diretta ortogonale come

$$V = U \oplus U^\perp.$$

Come descritto nell'Esempio 4.1.11, questa somma diretta induce una proiezione $p_U: V \rightarrow U$ che indichiamo semplicemente così:

$$p_w: V \rightarrow U.$$

La proiezione ortogonale p_w può essere descritta esplicitamente.

Proposizione 8.1.14. *Vale*

$$p_w(v) = \frac{\langle v, w \rangle}{\langle w, w \rangle} w.$$

Dimostrazione. Il vettore v si scrive in modo unico come

$$v = p_w(v) + v'$$

con $v' \in U^\perp$. Poiché $p_w(v) \in \text{Span}(w)$, abbiamo $p_w(v) = kw$ per un $k \in \mathbb{R}$ da determinare. La condizione è che $v' = v - kw$ sia ortogonale a w , cioè

$$\langle v - kw, w \rangle = 0.$$

Questo implica che

$$\langle v, w \rangle - k \langle w, w \rangle = 0 \implies k = \frac{\langle v, w \rangle}{\langle w, w \rangle}.$$

La dimostrazione è conclusa. □

Si veda la Figura 8.3. La proiezione ortogonale $p_w(v)$ è ottenuta moltiplicando w per il *coefficiente di Fourier*

$$\frac{\langle v, w \rangle}{\langle w, w \rangle}.$$

Questo coefficiente comparirà in vari punti nelle pagine seguenti.

Esercizio 8.1.15. Vale l'uguaglianza

$$\|p_w(v)\| = \frac{|\langle v, w \rangle|}{\|w\|}.$$

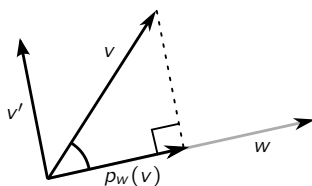


Figura 8.3. Il vettore v si decompone nella sua proiezione ortogonale $p_w(v)$ su w , che è parallela a w , e in un vettore v' che è ortogonale a w .

8.1.6. Coefficienti di Fourier. Mostriamo adesso come i coefficienti di Fourier forniscano immediatamente le coordinate di un vettore rispetto ad una base ortogonale.

Sia come sempre V uno spazio vettoriale dotato di un prodotto scalare definito positivo. Sia inoltre $\mathcal{B} = \{v_1, \dots, v_n\}$ una base ortogonale per V .

Proposizione 8.1.16. Per qualsiasi vettore $v \in V$ vale la relazione

$$v = p_{v_1}(v) + \dots + p_{v_n}(v).$$

In altre parole, ciascun vettore v è somma delle sue proiezioni ortogonali sugli elementi della base. Usando i coefficienti di Fourier scriviamo:

$$(12) \quad v = \frac{\langle v, v_1 \rangle}{\langle v_1, v_1 \rangle} v_1 + \dots + \frac{\langle v, v_n \rangle}{\langle v_n, v_n \rangle} v_n.$$

Dimostrazione. Sappiamo che

$$v = \lambda_1 v_1 + \dots + \lambda_n v_n.$$

Mostriamo adesso che i λ_i sono i coefficienti di Fourier. Facendo il prodotto scalare con v_i di ambo i membri e usando $\langle v_j, v_i \rangle = 0 \forall j \neq i$ otteniamo

$$\langle v, v_i \rangle = \lambda_i \langle v_i, v_i \rangle$$

e quindi effettivamente

$$\lambda_i = \frac{\langle v, v_i \rangle}{\langle v_i, v_i \rangle}.$$

La dimostrazione è conclusa. □

L'equazione (12) assume una forma particolarmente semplice se $\mathcal{B}$ è ortonormale: in questo caso $\langle v_i, v_i \rangle = 1$ per ogni i e l'equazione diventa

$$v = \langle v, v_1 \rangle v_1 + \dots + \langle v, v_n \rangle v_n.$$

La Proposizione 8.1.16 può essere enunciata anche nel modo seguente:

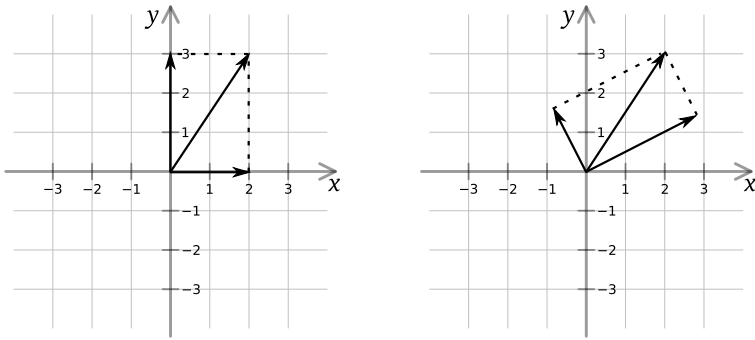


Figura 8.4. Le componenti di un vettore rispetto alla base canonica e ad un'altra base ortogonale. Queste si calcolano facilmente usando i coefficienti di Fourier.

Corollario 8.1.17. *Le coordinate di un vettore v rispetto ad una base ortogonale $\mathcal{B} = \{v_1, \dots, v_n\}$ sono i coefficienti di Fourier. Scriviamo quindi:*

$$[v]_{\mathcal{B}} = \begin{pmatrix} \langle v, v_1 \rangle \\ \langle v_1, v_1 \rangle \\ \vdots \\ \langle v, v_n \rangle \\ \langle v_n, v_n \rangle \end{pmatrix}.$$

Se la base $\mathcal{B}$ è ortonormale, scriviamo

$$[v]_{\mathcal{B}} = \begin{pmatrix} \langle v, v_1 \rangle \\ \vdots \\ \langle v, v_n \rangle \end{pmatrix}.$$

Esempio 8.1.18. Consideriamo $\mathbb{R}^n$ con il prodotto scalare euclideo. La base canonica $\mathcal{C} = \{e_1, \dots, e_n\}$ è ortonormale. Se $x \in \mathbb{R}^n$ è un vettore qualsiasi, il coefficiente di Fourier $\langle x, e_i \rangle$ di x rispetto a $\mathcal{C}$ è semplicemente la sua coordinata x_i . Quindi con il Corollario 8.1.17 ritroviamo semplicemente una uguaglianza che già conoscevamo:

$$[x]_{\mathcal{C}} = \begin{pmatrix} x_1 \\ \vdots \\ x_n \end{pmatrix}.$$

Esempio 8.1.19. Consideriamo $\mathbb{R}^2$ con il prodotto scalare euclideo e prendiamo la base ortogonale $v_1 = \begin{pmatrix} 2 \\ 1 \end{pmatrix}$, $v_2 = \begin{pmatrix} -1 \\ 2 \end{pmatrix}$. Se $v = \begin{pmatrix} 2 \\ 3 \end{pmatrix}$ allora

$$v = \frac{\langle v, v_1 \rangle}{\langle v_1, v_1 \rangle} v_1 + \frac{\langle v, v_2 \rangle}{\langle v_2, v_2 \rangle} v_2 = \frac{7}{5} \begin{pmatrix} 2 \\ 1 \end{pmatrix} + \frac{4}{5} \begin{pmatrix} -1 \\ 2 \end{pmatrix} = \begin{pmatrix} 2 \\ 3 \end{pmatrix}.$$

Abbiamo calcolato le coordinate di $\begin{pmatrix} 2 \\ 3 \end{pmatrix}$ rispetto alla base $\mathcal{B} = \{v_1, v_2\}$:

$$\left[\begin{pmatrix} 2 \\ 3 \end{pmatrix} \right]_{\mathcal{B}} = \frac{1}{5} \begin{pmatrix} 7 \\ 4 \end{pmatrix}.$$

La Figura 8.4 mostra le componenti di v rispetto alla base canonica e rispetto alla base $\mathcal{B}$.

L'aspetto fondamentale di questa costruzione è che le coordinate di un vettore rispetto ad una base ortogonale possono essere calcolate direttamente con i coefficienti di Fourier, senza dover impostare e risolvere un sistema lineare. Questa costruzione funziona ovviamente solo se la base $v_1, \dots, v_n$ è ortogonale. Il prodotto scalare ci è di aiuto per fare dei calcoli più rapidamente.

Notiamo infine che la norma di un vettore dipende dalle norme delle sue proiezioni. Il risultato seguente è una generalizzazione del Teorema di Pitagora in più dimensioni. Sia $\mathcal{B} = \{v_1, \dots, v_n\}$ base ortogonale per V .

Proposizione 8.1.20. *Per qualsiasi vettore $v \in V$ vale l'uguaglianza*

$$\|v\|^2 = \|p_{v_1}(v)\|^2 + \dots + \|p_{v_n}(v)\|^2.$$

Dimostrazione. Sappiamo che

$$v = p_{v_1}(v) + \dots + p_{v_n}(v)$$

e quindi

$$\begin{aligned} \langle v, v \rangle &= \langle p_{v_1}(v) + \dots + p_{v_n}(v), p_{v_1}(v) + \dots + p_{v_n}(v) \rangle \\ &= \langle p_{v_1}(v), p_{v_1}(v) \rangle + \dots + \langle p_{v_n}(v), p_{v_n}(v) \rangle. \end{aligned}$$

Abbiamo usato che $\langle v_i, v_j \rangle = 0$ se $i \neq j$. □

Usando l'Esercizio 8.1.15, possiamo riscrivere l'uguaglianza della proposizione nel modo seguente:

$$(13) \quad \|v\|^2 = \frac{\langle v, v_1 \rangle^2}{\|v_1\|^2} + \dots + \frac{\langle v, v_n \rangle^2}{\|v_n\|^2}.$$

8.1.7. Ortogonalizzazione di Gram-Schmidt. Abbiamo appena scoperto che le basi ortogonali sono molto utili, perché possiamo calcolare le coordinate di qualsiasi vettore tramite i coefficienti di Fourier. Ci poniamo il seguente problema: come facciamo a costruire una base ortogonale?

Sia come sempre V uno spazio vettoriale dotato di un prodotto scalare definito positivo. Descriviamo adesso un algoritmo che prende come *input* dei vettori indipendenti $v_1, \dots, v_k \in V$ e restituisce come *output* dei vettori $w_1, \dots, w_k$ indipendenti e ortogonali. Questo algoritmo è noto come *ortogonalizzazione di Gram-Schmidt*.

L'algoritmo è il seguente. Si costruiscono i vettori $w_1, \dots, w_k$ induttivamente in questo modo:

$$\begin{aligned} w_1 &= v_1, \\ w_2 &= v_2 - p_{w_1}(v_2), \\ w_3 &= v_3 - p_{w_1}(v_3) - p_{w_2}(v_3), \\ &\vdots \\ w_k &= v_k - p_{w_1}(v_k) - \dots - p_{w_{k-1}}(v_k). \end{aligned}$$

In altre parole, ad ogni passo togliamo al vettore v_i le sue proiezioni ortogonali sui vettori $w_1, \dots, w_{i-1}$ già trovati precedentemente, in modo che il nuovo vettore w_i sia ortogonale a questi. Esplicitando i coefficienti di Fourier:

$$\begin{aligned} w_1 &= v_1, \\ w_2 &= v_2 - \frac{\langle v_2, w_1 \rangle}{\langle w_1, w_1 \rangle} w_1, \\ w_3 &= v_3 - \frac{\langle v_3, w_1 \rangle}{\langle w_1, w_1 \rangle} w_1 - \frac{\langle v_3, w_2 \rangle}{\langle w_2, w_2 \rangle} w_2, \\ &\vdots \\ w_k &= v_k - \frac{\langle v_k, w_1 \rangle}{\langle w_1, w_1 \rangle} w_1 - \dots - \frac{\langle v_k, w_{k-1} \rangle}{\langle w_{k-1}, w_{k-1} \rangle} w_{k-1}. \end{aligned}$$

Dimostriamo adesso che l'algoritmo funziona.

Proposizione 8.1.21. *L'algoritmo di Gram-Schmidt produce effettivamente dei vettori $w_1, \dots, w_k$ indipendenti e ortogonali. Inoltre per ogni $j = 1, \dots, k$, i primi j vettori della base generano gli stessi spazi di prima:*

$$(14) \quad \text{Span}(v_1, \dots, v_j) = \text{Span}(w_1, \dots, w_j).$$

Dimostrazione. Dimostriamo tutto per induzione su j . Il caso $j = 1$ è banale, quindi esaminiamo il passo induttivo.

Il vettore w_j è costruito togliendo a v_j le sue proiezioni ortogonali sui vettori precedenti $w_1, \dots, w_{j-1}$ e quindi ciò che resta è ortogonale a questi. Più formalmente, per ogni $i < j$ abbiamo

$$\begin{aligned} \langle w_j, w_i \rangle &= \langle v_j, w_i \rangle - \sum_{h=1}^{j-1} \langle p_{w_h}(v_j), w_i \rangle = \langle v_j, w_i \rangle - \langle p_{w_i}(v_j), w_i \rangle \\ &= \langle v_j, w_i \rangle - \frac{\langle v_j, w_i \rangle}{\langle w_i, w_i \rangle} \langle w_i, w_i \rangle = 0. \end{aligned}$$

Nella seconda uguaglianza abbiamo usato che $\langle w_h, w_i \rangle = 0$ se $h \neq i$. Notiamo inoltre che ciascun w_j è costruito come combinazione lineare dei $v_1, \dots, v_j$, quindi $\text{Span}(w_1, \dots, w_j) \subset \text{Span}(v_1, \dots, v_j)$. D'altra parte, si

può anche scrivere facilmente v_j come combinazione lineare dei $w_1, \dots, w_j$, quindi anche $\text{Span}(w_1, \dots, w_j) \supset \text{Span}(v_1, \dots, v_j)$ e dimostriamo (14).

In particolare $\dim \text{Span}(w_1, \dots, w_k) = k$ e deduciamo che i nuovi vettori $w_1, \dots, w_k$ sono effettivamente indipendenti. $\square$

Esempio 8.1.22. Ortogonalizziamo i vettori

$$v_1 = \begin{pmatrix} 1 \\ 0 \end{pmatrix}, \quad v_2 = \begin{pmatrix} 1 \\ 1 \end{pmatrix}.$$

rispetto al prodotto scalare euclideo di $\mathbb{R}^2$. Otteniamo

$$w_1 = v_1 = \begin{pmatrix} 1 \\ 0 \end{pmatrix},$$

$$w_2 = v_2 - p_{w_1}(v_2) = v_2 - \frac{\langle v_2, w_1 \rangle}{\langle w_1, w_1 \rangle} w_1 = \begin{pmatrix} 1 \\ 1 \end{pmatrix} - \frac{1}{1} \begin{pmatrix} 1 \\ 0 \end{pmatrix} = \begin{pmatrix} 0 \\ 1 \end{pmatrix}.$$

I vettori ottenuti w_1 e w_2 sono effettivamente ortogonali.

Esempio 8.1.23. Ortogonalizziamo i vettori

$$v_1 = \begin{pmatrix} 1 \\ 1 \\ 0 \end{pmatrix}, \quad v_2 = \begin{pmatrix} 0 \\ 1 \\ 1 \end{pmatrix}, \quad v_3 = \begin{pmatrix} 1 \\ 0 \\ 1 \end{pmatrix}$$

rispetto al prodotto scalare euclideo di $\mathbb{R}^3$. Otteniamo

$$w_1 = v_1 = \begin{pmatrix} 1 \\ 1 \\ 0 \end{pmatrix},$$

$$w_2 = v_2 - \frac{\langle v_2, w_1 \rangle}{\langle w_1, w_1 \rangle} w_1 = \begin{pmatrix} 0 \\ 1 \\ 1 \end{pmatrix} - \frac{1}{2} \begin{pmatrix} 1 \\ 1 \\ 0 \end{pmatrix} = \begin{pmatrix} -\frac{1}{2} \\ \frac{1}{2} \\ 1 \end{pmatrix},$$

$$w_3 = v_3 - \frac{\langle v_3, w_1 \rangle}{\langle w_1, w_1 \rangle} w_1 - \frac{\langle v_3, w_2 \rangle}{\langle w_2, w_2 \rangle} w_2 = \begin{pmatrix} 1 \\ 0 \\ 1 \end{pmatrix} - \frac{1}{2} \begin{pmatrix} 1 \\ 1 \\ 0 \end{pmatrix} - \frac{\frac{1}{2}}{\frac{3}{2}} \begin{pmatrix} -\frac{1}{2} \\ \frac{1}{2} \\ 1 \end{pmatrix} = \begin{pmatrix} \frac{2}{3} \\ \frac{2}{3} \\ \frac{2}{3} \end{pmatrix}$$

I vettori ottenuti w_1 , w_2 e w_3 sono effettivamente ortogonali.

8.1.8. Riscaldamento. L'operazione di *riscaldamento* di un vettore $v \neq 0$ consiste nella sostituzione di v con $v' = \lambda v$, per qualche $\lambda \neq 0$. Se $w_1, \dots, w_k$ sono vettori ortogonali, possiamo sostituire ciascun w_i con un suo riscaldamento $w'_i = \lambda_i w_i$, ed i nuovi vettori $w'_1, \dots, w'_k$ sono ancora ortogonali.

Nell'Esempio 8.1.23 può essere utile riscaldare i vettori:

$$w'_1 = w_1, \quad w'_2 = 2w_2, \quad w'_3 = \frac{3}{2}w_3$$

ed ottenere una nuova base di vettori ortogonali, più semplice da descrivere:

$$w'_1 = \begin{pmatrix} 1 \\ 1 \\ 0 \end{pmatrix}, \quad w'_2 = \begin{pmatrix} -1 \\ 1 \\ 2 \end{pmatrix}, \quad w'_3 = \begin{pmatrix} 1 \\ -1 \\ 1 \end{pmatrix}.$$

La *normalizzazione* di un vettore $v \neq 0$ è il riscalamento

$$v' = \frac{v}{\|v\|}.$$

Questa operazione è importante perché

$$\|v'\| = \frac{\|v\|}{\|v\|} = 1.$$

Il vettore normalizzato v' ha sempre norma uno. Per trasformare una base ortogonale in una base ortonormale è sufficiente normalizzare i vettori. Questa operazione ha spesso l'inconveniente di introdurre delle radici quadrate.

Esempio 8.1.24. Normalizzando la base ortogonale trovata nell'Esempio 8.1.23 otteniamo la base ortonormale seguente:

$$w'_1 = \begin{pmatrix} \frac{\sqrt{2}}{2} \\ \frac{\sqrt{2}}{2} \\ 0 \end{pmatrix}, \quad w'_2 = \begin{pmatrix} -\frac{\sqrt{6}}{6} \\ \frac{\sqrt{6}}{6} \\ \frac{\sqrt{6}}{3} \end{pmatrix}, \quad w'_3 = \begin{pmatrix} \frac{\sqrt{3}}{3} \\ -\frac{\sqrt{3}}{3} \\ \frac{\sqrt{3}}{3} \end{pmatrix}.$$

Si verifica facilmente che hanno tutti norma uno.

8.1.9. Ortogonalità. Sia come sempre V uno spazio vettoriale dotato di un prodotto scalare definito positivo. Lavorare con vettori ortogonali ha molti vantaggi. Uno di questi è che l'indipendenza lineare è garantita a priori:

Proposizione 8.1.25. *Se $v_1, \dots, v_k \in V$ sono tutti ortogonali e non nulli, allora sono indipendenti.*

Dimostrazione. Supponiamo che esista una combinazione lineare nulla

$$0 = \lambda_1 v_1 + \dots + \lambda_k v_k.$$

Facendo il prodotto scalare di entrambi i membri con v_i otteniamo

$$0 = \langle 0, v_i \rangle = \langle \lambda_1 v_1 + \dots + \lambda_k v_k, v_i \rangle = \lambda_i \langle v_i, v_i \rangle$$

perché $\langle v_j, v_i \rangle = 0$ se $j \neq i$. A questo punto $\langle v_i, v_i \rangle > 0$ implica $\lambda_i = 0$. Abbiamo scoperto che $\lambda_i = 0 \forall i$ e questo conclude la dimostrazione. $\square$

È sempre possibile completare un insieme di vettori ortogonali a base ortogonale. Supponiamo che $\dim V = n$.

Proposizione 8.1.26. *Se $v_1, \dots, v_k \in V$ sono ortogonali e non nulli, esistono altri vettori $v_{k+1}, \dots, v_n$ tali che $v_1, \dots, v_n$ sia una base ortogonale.*

Dimostrazione. Sia $U = \text{Span}(v_1, \dots, v_k)$. Prendiamo $v_{k+1}, \dots, v_n$ base ortogonale qualsiasi di $U^\perp$. $\square$

Enunciamo una disuguaglianza che ha varie applicazioni in analisi.

Proposizione 8.1.27 (Disuguaglianza di Bessel). Se $v_1, \dots, v_k$ sono ortogonali e non nulli, e $v \in V$ è un vettore qualsiasi, allora

$$\sum_{i=1}^k \frac{\langle v, v_i \rangle^2}{\|v_i\|^2} \leq \|v\|^2.$$

Dimostrazione. Completiamo ad una base ortogonale $v_1, \dots, v_n$ e quindi applichiamo (13). $\square$

8.1.10. Proiezioni su sottospazi. Sia V uno spazio vettoriale con prodotto scalare definito positivo e $U \subset V$ un sottospazio. Per il Teorema 7.3.12 otteniamo la somma diretta ortogonale

$$V = U \oplus U^\perp.$$

Questa definisce come nell'Esempio 4.1.11 una proiezione

$$p_U: V \longrightarrow U$$

detta *proiezione ortogonale*. Possiamo decomporre ciascun vettore $v \in V$ come somma delle sue proiezioni ortogonali su U e $U^\perp$:

$$v = p_U(v) + p_{U^\perp}(v).$$

Mostriamo come determinare concretamente $p_U(v)$. Entrano in gioco ancora una volta i coefficienti di Fourier.

Proposizione 8.1.28. Sia $v_1, \dots, v_k$ una base ortogonale per U . Allora

$$p_U(v) = \frac{\langle v, v_1 \rangle}{\langle v_1, v_1 \rangle} v_1 + \dots + \frac{\langle v, v_k \rangle}{\langle v_k, v_k \rangle} v_k.$$

Dimostrazione. Sia $v_{k+1}, \dots, v_n$ una base ortogonale per $U^\perp$. Quindi $v_1, \dots, v_n$ per V è una base ortogonale per V . Sappiamo che

$$v = \underbrace{\frac{\langle v, v_1 \rangle}{\langle v_1, v_1 \rangle} v_1 + \dots + \frac{\langle v, v_k \rangle}{\langle v_k, v_k \rangle} v_k}_u + \underbrace{\frac{\langle v, v_{k+1} \rangle}{\langle v_{k+1}, v_{k+1} \rangle} v_{k+1} + \dots + \frac{\langle v, v_n \rangle}{\langle v_n, v_n \rangle} v_n}_{u'}.$$

Abbiamo scritto $v = u + u'$ con $u \in U$ e $u' \in U^\perp$. Quindi $u = p_U(v)$. $\square$

Esempio 8.1.29. Consideriamo $U = \{x + y - z = 0\} \subset \mathbb{R}^3$ con il prodotto scalare euclideo. Cerchiamo una base ortogonale per U e troviamo ad esempio

$$v_1 = \begin{pmatrix} 1 \\ 1 \\ 2 \end{pmatrix}, \quad v_2 = \begin{pmatrix} 1 \\ -1 \\ 0 \end{pmatrix}.$$

A questo punto determiniamo la proiezione ortogonale su U di un generico vettore di $\mathbb{R}^3$ usando la proposizione precedente, in questo modo:

$$p_U \begin{pmatrix} x \\ y \\ z \end{pmatrix} = \frac{x + y + 2z}{6} \begin{pmatrix} 1 \\ 1 \\ 2 \end{pmatrix} + \frac{x - y}{2} \begin{pmatrix} 1 \\ -1 \\ 0 \end{pmatrix} = \frac{1}{3} \begin{pmatrix} 2x - y + z \\ -x + 2y + z \\ x + y + 2z \end{pmatrix}$$

Quindi la matrice associata a p_U nella base canonica è

$$\frac{1}{3} \begin{pmatrix} 2 & -1 & 1 \\ -1 & 2 & 1 \\ 1 & 1 & 2 \end{pmatrix}.$$

Usando questa matrice possiamo calcolare la proiezione ortogonale su U di qualsiasi vettore di $\mathbb{R}^3$.

Ricordiamo alcuni fatti generali sulle proiezioni già notati precedentemente. Innanzitutto p_U manda ogni vettore di U in sé stesso, mentre manda ogni vettore di $U^\perp$ in zero. In particolare otteniamo:

$$\text{Im } p_U = U, \quad \ker p_U = U^\perp.$$

Se prendiamo una base $\{v_1, \dots, v_k\}$ di U e una base $\{v_{k+1}, \dots, v_n\}$ di $U^\perp$, assieme formano una base $\mathcal{B} = \{v_1, \dots, v_n\}$ di V rispetto alla quale la matrice associata è molto semplice:

$$[p_U]_{\mathcal{B}}^{\mathcal{B}} = \begin{pmatrix} I_k & 0 \\ 0 & 0 \end{pmatrix}.$$

In particolare p_U è diagonalizzabile con autovalori 0 e 1. Gli autospazi relativi sono precisamente $U^\perp$ e U . Vale infine la relazione

$$p_U \circ p_U = p_U.$$

Abbiamo già incontrato delle trasformazioni che soddisfano questa relazione nella Proposizione 6.3.17.

Esempio 8.1.30. Sia $U = \{x + y + z = 0\}$ un piano in $\mathbb{R}^3$. La base

$$\mathcal{B} = \left\{ \begin{pmatrix} 1 \\ -1 \\ 0 \end{pmatrix}, \begin{pmatrix} 0 \\ 1 \\ -1 \end{pmatrix}, \begin{pmatrix} 1 \\ 1 \\ 1 \end{pmatrix} \right\}$$

ha i primi due vettori in U e l'ultimo in $U^\perp$. Quindi

$$[p_U]_{\mathcal{B}}^{\mathcal{B}} = \begin{pmatrix} 1 & 0 & 0 \\ 0 & 1 & 0 \\ 0 & 0 & 0 \end{pmatrix}.$$

Nella base canonica la proiezione è rappresentata dalla matrice $B = M^{-1}AM$ con $M = [\text{id}]_{\mathcal{B}}^{\mathcal{C}}$. Facendo i conti otteniamo

$$B = \frac{1}{3} \begin{pmatrix} 1 & 0 & 1 \\ -1 & 1 & 1 \\ 0 & -1 & 1 \end{pmatrix} \begin{pmatrix} 1 & 0 & 0 \\ 0 & 1 & 0 \\ 0 & 0 & 0 \end{pmatrix} \begin{pmatrix} 2 & -1 & -1 \\ 1 & 1 & -2 \\ 1 & 1 & 1 \end{pmatrix} = \frac{1}{3} \begin{pmatrix} 2 & -1 & -1 \\ -1 & 2 & -1 \\ -1 & -1 & 2 \end{pmatrix}$$

Il lettore è invitato a controllare che usando la Proposizione 8.1.28 come abbiamo fatto nell'Esempio 8.1.29 otteniamo la stessa matrice.

Concludiamo questa sezione con una osservazione che può tornare utile. Può capitare in molti casi che sia più agevole calcolare la proiezione

$p_{U^\perp}$ sullo spazio ortogonale $U^\perp$ invece di p_U . Si può quindi calcolare prima $p_{U^\perp}$ e quindi ritrovare p_U scrivendo semplicemente

$$p_U(v) = v - p_{U^\perp}(v).$$

Esempio 8.1.31. Studiamo l'Esempio 8.1.30 seguendo questo suggerimento. Lo spazio ortogonale $U^\perp$ è una retta, generata dal vettore

$$w = \begin{pmatrix} 1 \\ 1 \\ 1 \end{pmatrix}.$$

Ne segue che $p_U(v) = v - p_{U^\perp}(v) = v - p_w(v)$. Quindi

$$p_U \begin{pmatrix} x \\ y \\ z \end{pmatrix} = \begin{pmatrix} x \\ y \\ z \end{pmatrix} - \frac{x+y+z}{3} \begin{pmatrix} 1 \\ 1 \\ 1 \end{pmatrix} = \frac{1}{3} \begin{pmatrix} 2x - y - z \\ -x + 2y - z \\ -x - y + 2z \end{pmatrix}.$$

Abbiamo ritrovato la stessa matrice associata dell'Esempio 8.1.30.

8.2. Isometrie

Abbiamo definito le isometrie nella Sezione 7.5. Le studiamo adesso più approfonditamente nel contesto dei prodotti scalari definiti positivi.

8.2.1. Definizioni equivalenti. Letteralmente, il termine “isometria” indica che T preserva la geometria degli oggetti, cioè non li distorce. Quando il prodotto scalare è definito positivo, possiamo definire questa nozione in almeno tre modi equivalenti:

Proposizione 8.2.1. *Sia $T: V \rightarrow W$ un isomorfismo fra spazi dotati di un prodotto scalare definito positivo. I fatti seguenti sono equivalenti:*

- (1) T è una isometria,
- (2) T preserva la norma, cioè $\|T(v)\| = \|v\| \forall v \in V$,
- (3) T preserva la distanza, cioè $d(v, w) = d(T(v), T(w)) \forall v, w \in V$.

Dimostrazione. (1) $\Rightarrow$ (2). Se T è una isometria, allora

$$\|T(v)\|^2 = \langle T(v), T(v) \rangle = \langle v, v \rangle = \|v\|^2.$$

(2) $\Rightarrow$ (3). Vale

$$d(v, w) = \|v - w\| = \|T(v - w)\| = \|T(v) - T(w)\| = d(T(v), T(w)).$$

(3) $\Rightarrow$ (2). Vale

$$\|v\| = d(0, v) = d(T(0), T(v)) = d(0, T(v)) = \|T(v)\|.$$

(2) $\Rightarrow$ (1). Questo è meno facile dei precedenti:

$$\begin{aligned}
 \langle v, w \rangle &= \frac{\langle v+w, v+w \rangle - \langle v, v \rangle - \langle w, w \rangle}{2} = \frac{\|v+w\|^2 - \|v\|^2 - \|w\|^2}{2} \\
 &= \frac{\|T(v+w)\|^2 - \|T(v)\|^2 - \|T(w)\|^2}{2} \\
 &= \frac{\langle T(v+w), T(v+w) \rangle - \langle T(v), T(v) \rangle - \langle T(w), T(w) \rangle}{2} \\
 &= \langle T(v), T(w) \rangle.
 \end{aligned}$$

La dimostrazione è completa. $\square$

8.2.2. Matrici ortogonali. Per comprendere le isometrie di un fissato spazio dobbiamo studiare le matrici ortogonali.

Sia V uno spazio vettoriale munito di un prodotto scalare definito positivo. Sia $\mathcal{B}$ una base ortonormale per V e $T: V \rightarrow V$ un endomorfismo. Per il Corollario 7.5.10,

$$T \text{ è una isometria} \iff A = [T]_{\mathcal{B}}^{\mathcal{B}} \text{ è ortogonale.}$$

Ricordiamo che A è ortogonale se ${}^tAA = I_n$. Studiamo adesso alcune semplici proprietà delle matrici ortogonali.

Proposizione 8.2.2. *Se $A \in M(n)$ è ortogonale, allora*

- $\det A = \pm 1$,
- *i possibili autovalori di A sono ± 1 .*

Dimostrazione. Se ${}^tAA = I_n$, allora per Binet

$$1 = \det I_n = \det({}^tA) \det A = (\det A)^2.$$

Quindi $\det A = \pm 1$. L'endomorfismo L_A è una isometria per $\mathbb{R}^n$ munito del prodotto scalare euclideo. Se λ è un autovalore per A , esiste $v \neq 0$ tale che $Av = \lambda v$. L'isometria L_A preserva la norma e quindi

$$\|v\| = \|Av\| = \|\lambda v\| = |\lambda| \|v\|.$$

Ne segue che $\lambda = \pm 1$. $\square$

Dal fatto che $\det A = \pm 1$ deduciamo che una matrice ortogonale A è sempre invertibile; inoltre l'inversa di A è sorprendentemente semplice da determinare: dalla relazione ${}^tAA = I$ deduciamo infatti che $A^{-1} = {}^tA$. In particolare vale anche $A {}^tA = I$.

Proposizione 8.2.3. *Una matrice $A \in M(n)$ è ortogonale $\iff$ i vettori colonna $A^1, \dots, A^n$ formano una base ortonormale di $\mathbb{R}^n$.*

Dimostrazione. La relazione ${}^tAA = I$ è equivalente alla relazione $\langle A^i, A^j \rangle = 1$ se $i = j$ e 0 se $i \neq j$. Questo conclude la dimostrazione. In alternativa, questa è una conseguenza del Corollario 7.5.6, visto che i vettori colonna sono le immagini della base canonica (che è ortonormale). $\square$

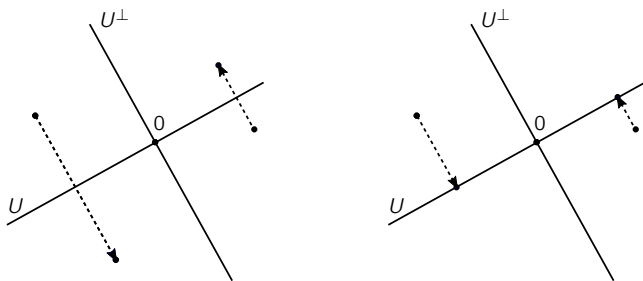


Figura 8.5. La riflessione e la proiezione ortogonale lungo un sottospazio U .

8.2.3. Riflessioni ortogonali. Sia V uno spazio vettoriale munito di un prodotto scalare definito positivo. Sia $U \subset V$ un sottospazio. Nell'Esempio 7.5.3 abbiamo definito la *riflessione ortogonale* $r_U: V \rightarrow V$. Questa può essere riscritta nel modo seguente. Ogni vettore $v \in V$ si decompone nelle proiezioni ortogonali su U e $U^\perp$:

$$v = p_U(v) + p_{U^\perp}(v).$$

La riflessione ortogonale lungo U è l'isometria

$$r_U(v) = p_U(v) - p_{U^\perp}(v) = v - 2p_{U^\perp}(v).$$

Si veda la Figura 8.5. Se $v_1, \dots, v_k$ e $v_{k+1}, \dots, v_n$ sono basi ortonormali di U e $U^\perp$, allora $\mathcal{B} = \{v_1, \dots, v_n\}$ è una base ortonormale di V e la matrice associata a r_U nella base $\mathcal{B}$ è molto semplice:

$$[r_U]_{\mathcal{B}}^{\mathcal{B}} = \begin{pmatrix} I_k & 0 \\ 0 & -I_{n-k} \end{pmatrix}.$$

Questa matrice è ortogonale: grazie al Corollario 7.5.10 confermiamo che r_U sia effettivamente una isometria. Gli autospazi di r_U sono

$$V_1 = U, \quad V_{-1} = U^\perp.$$

Notiamo che U può avere dimensione arbitraria ed il segno di

$$\det r_U = (-1)^{n-k} = (-1)^{n-\dim U}$$

dipende dalla parità di $\dim U$.

Esempio 8.2.4. Nel caso in cui $U = \{0\}$, otteniamo la riflessione r_U rispetto all'origine, che cambia di segno a tutti i vettori:

$$r_U(v) = -v.$$

Il determinante di questa applicazione è $(-1)^{\dim V}$. Se $V = \mathbb{R}^2$, la riflessione rispetto all'origine è una rotazione di angolo π .

Esempio 8.2.5. Nel caso in cui U sia un iperpiano (cioè un sottospazio di dimensione $\dim V - 1$) la riflessione r_U è semplice da scrivere. Sia v_0 un generatore della retta $U^\perp$. Troviamo

$$r_U(v) = v - 2 \frac{\langle v, v_0 \rangle}{\langle v_0, v_0 \rangle} v_0.$$

Questa formula è molto simile a quella che rappresenta la proiezione p_U , con un coefficiente 2 al posto di 1. Ad esempio, se $U \subset \mathbb{R}^3$ è il piano

$$U = \{x + y + z = 0\}$$

otteniamo

$$r_U \begin{pmatrix} x \\ y \\ z \end{pmatrix} = \begin{pmatrix} x \\ y \\ z \end{pmatrix} - 2 \frac{x+y+z}{3} \begin{pmatrix} 1 \\ 1 \\ 1 \end{pmatrix} = \frac{1}{3} \begin{pmatrix} x-2y-2z \\ -2x+y-2z \\ -2x-2y+1z \end{pmatrix}.$$

La riflessione r_U rispetto alla base canonica è rappresentata dalla matrice

$$\frac{1}{3} \begin{pmatrix} 1 & -2 & -2 \\ -2 & 1 & -2 \\ -2 & -2 & 1 \end{pmatrix}$$

che è effettivamente ortogonale.

Nella prossima proposizione vediamo come costruire facilmente una riflessione con un comportamento determinato su un vettore specifico.

Proposizione 8.2.6. *Siano $v, w \in V$ vettori arbitrari con $\|v\| = \|w\|$. Sia $U = \text{Span}(v+w)$. Vale*

$$r_U(v) = w.$$

Dimostrazione. Il vettore v si decompone come

$$v = \frac{v+w}{2} + \frac{v-w}{2}.$$

Il primo addendo è in U , il secondo in $U^\perp$ perché

$$\langle v-w, v+w \rangle = \|v\|^2 - \|w\|^2 = 0.$$

Quindi

$$r_U(v) = \frac{v+w}{2} - \frac{v-w}{2} = w.$$

La dimostrazione è completa. $\square$

Notiamo che $U = \text{Span}(v+w)$ è una retta se $v \neq -w$ ed è l'origine $\{0\}$ se $v = -w$.

8.2.4. Isometrie del piano. Vogliamo adesso classificare completamente le isometrie del piano $\mathbb{R}^2$, dotato del prodotto scalare euclideo. Per quanto appena visto, questo problema è equivalente a classificare le matrici ortogonali 2×2 . Le matrici ortogonali 2×2 possono essere descritte tutte facilmente:

Proposizione 8.2.7. *Le matrici ortogonali in $M(2)$ sono le seguenti:*

$$\text{Rot}_\theta = \begin{pmatrix} \cos \theta & -\sin \theta \\ \sin \theta & \cos \theta \end{pmatrix}, \quad \text{Rif}_\theta = \begin{pmatrix} \cos \theta & \sin \theta \\ \sin \theta & -\cos \theta \end{pmatrix}$$

al variare di $\theta \in [0, 2\pi)$.

Dimostrazione. I vettori colonna A^1 e A^2 di una matrice ortogonale A devono formare una base ortonormale. Scriviamo A^1 come un generico vettore unitario

$$A^1 = \begin{pmatrix} \cos \theta \\ \sin \theta \end{pmatrix}.$$

Il secondo vettore A^2 deve essere ortogonale ad A^1 ed unitario: ci sono solo due possibilità che differiscono solo di un segno, e otteniamo così le matrici Rot_θ e Rif_θ . $\square$

Notiamo che $\det \text{Rot}_\theta = 1$ e $\det \text{Rif}_\theta = -1$. Come sappiamo già fin dalle Sezioni 4.4.8 e 4.4.9, queste matrici rappresentano una rotazione antioraria di angolo θ e una riflessione ortogonale rispetto ad una retta che forma un angolo $\frac{\theta}{2}$ con l'asse x .

Corollario 8.2.8. *Le isometrie di $\mathbb{R}^2$ sono rotazioni e riflessioni.*

8.2.5. Rotazioni nello spazio. Vogliamo adesso studiare le isometrie dello spazio $\mathbb{R}^3$. Avevamo introdotto nella Sezione 4.4.11 la rotazione di un angolo θ intorno all'asse z , rappresentata rispetto alla base canonica dalla matrice

$$A = \begin{pmatrix} \cos \theta & -\sin \theta & 0 \\ \sin \theta & \cos \theta & 0 \\ 0 & 0 & 1 \end{pmatrix} = \begin{pmatrix} \text{Rot}_\theta & 0 \\ 0 & 1 \end{pmatrix}.$$

Generalizziamo questa costruzione scrivendo la rotazione intorno ad un asse orientata r qualsiasi. Prendiamo una base ortonormale positiva

$$\mathcal{B} = \{v_1, v_2, v_3\}$$

dove $v_1 \in r$ è orientato come r . Ricordiamo che la base è positiva se la matrice $(v_1 v_2 v_3)$ ha determinante 1. I vettori v_2, v_3 formano una base ortonormale del piano $r^\perp$. La rotazione intorno ad r di angolo θ è l'endomorfismo $T: \mathbb{R}^3 \rightarrow \mathbb{R}^3$ la cui matrice associata è

$$A = [T]_{\mathcal{B}}^{\mathcal{B}} = \begin{pmatrix} 1 & 0 & 0 \\ 0 & \text{Rot}_\theta \end{pmatrix} = \begin{pmatrix} 1 & 0 & 0 \\ 0 & \cos \theta & -\sin \theta \\ 0 & \sin \theta & \cos \theta \end{pmatrix}.$$

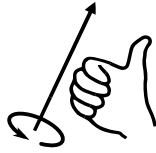


Figura 8.6. La regola della mano destra serve a determinare la direzione di rotazione intorno ad una retta orientata.

Si vede facilmente che ${}^tAA = I$, cioè la matrice A è ortogonale. Siccome $\mathcal{B}$ è una base ortonormale, per il Corollario 7.5.10 l'endomorfismo T è una isometria. Si tratta di una rotazione di angolo θ intorno alla retta orientata r . Il verso della rotazione può essere determinato usando la *regola della mano destra* mostrata nella Figura 8.6.

I casi $\theta = 0$ e $\theta = \pi$ sono un po' particolari:

- una rotazione di angolo $\theta = 0$ è l'identità;
- una rotazione di angolo $\theta = \pi$ è una riflessione rispetto a r .

Ovviamente è sempre possibile scrivere T nella base canonica, cambiando la matrice per similitudine nel modo usuale. Se $M = [\text{id}]_{\mathcal{C}}^{\mathcal{B}}$ è la matrice di cambiamento di base, otteniamo

$$A' = [T]_{\mathcal{C}}^{\mathcal{C}} = MAM^{-1}.$$

La nuova matrice A' può essere però molto complicata. Poiché A' rappresenta una isometria rispetto ad una base $\mathcal{C}$ ortonormale, la matrice A' sarà comunque ortogonale. Anche la matrice M è ortogonale, perché le sue colonne sono una base ortonormale: segue che l'inversa $M^{-1} = {}^tM$ si calcola agevolmente facendo la trasposta.

Esempio 8.2.9. Determiniamo la matrice A' che rappresenta (rispetto alla base canonica $\mathcal{C}$) la rotazione T intorno alla retta

$$r = \text{Span} \begin{pmatrix} 1 \\ 1 \\ 1 \end{pmatrix}$$

di angolo $\theta = \frac{2\pi}{3}$. Cerchiamo una base ortonormale v_1, v_2, v_3 con $v_1 \in r$. Ad esempio prendiamo questi:

$$v_1 = \frac{\sqrt{3}}{3} \begin{pmatrix} 1 \\ 1 \\ 1 \end{pmatrix}, \quad v_2 = \frac{\sqrt{2}}{2} \begin{pmatrix} 1 \\ -1 \\ 0 \end{pmatrix}, \quad v_3 = \frac{\sqrt{6}}{6} \begin{pmatrix} 1 \\ 1 \\ -2 \end{pmatrix}.$$

Otteniamo

$$A = [T]_{\mathcal{B}}^{\mathcal{B}} = \begin{pmatrix} 1 & 0 \\ 0 & \text{Rot}_{\theta} \end{pmatrix}.$$

Quindi $A' = {}^tMAM$, dove ${}^tM = (v_1 v_2 v_3)$. La matrice A' è pari a

$$\frac{1}{6} \begin{pmatrix} 2\sqrt{3} & 3\sqrt{2} & \sqrt{6} \\ 2\sqrt{3} & -3\sqrt{2} & \sqrt{6} \\ 2\sqrt{3} & 0 & -2\sqrt{6} \end{pmatrix} \begin{pmatrix} 1 & 0 & 0 \\ 0 & -\frac{1}{2} & -\frac{\sqrt{3}}{2} \\ 0 & \frac{\sqrt{3}}{2} & -\frac{1}{2} \end{pmatrix} \begin{pmatrix} 2\sqrt{3} & 2\sqrt{3} & 2\sqrt{3} \\ 3\sqrt{2} & -3\sqrt{2} & 0 \\ \sqrt{6} & \sqrt{6} & -2\sqrt{6} \end{pmatrix} \frac{1}{6}$$

Svolgendo un po' di conti si trova infine la matrice

$$A' = \begin{pmatrix} 0 & 0 & 1 \\ 1 & 0 & 0 \\ 0 & 1 & 0 \end{pmatrix}.$$

Come verifica, notiamo che r è effettivamente una retta invariante (è l'autospazio V_1) e che $\text{tr}A = \text{tr}A' = 0$.

8.2.6. Antirotazioni. Introduciamo adesso alcune isometrie di $\mathbb{R}^3$ che invertono l'orientazione dello spazio. Sia r una retta vettoriale in $\mathbb{R}^3$, e sia $U = r^\perp$ il piano vettoriale ad essa ortogonale. Ovviamente $r \cap U = \{0\}$.

Definizione 8.2.10. Una *antirotazione* $T: \mathbb{R}^3 \rightarrow \mathbb{R}^3$ di asse r è la composizione di una rotazione intorno a r e di una riflessione rispetto a U .

Scriviamo la matrice associata ad una antirotazione T rispetto ad una base opportuna. Se prendiamo un vettore unitario $v_1 \in r$ e lo completiamo a base ortonormale $\mathcal{B} = \{v_1, v_2, v_3\}$, otteniamo una base ortonormale $\{v_2, v_3\}$ per U e scriviamo la composizione nel modo seguente:

$$[T]_{\mathcal{B}}^{\mathcal{B}} = \begin{pmatrix} -1 & 0 \\ 0 & I_2 \end{pmatrix} \cdot \begin{pmatrix} 1 & 0 \\ 0 & \text{Rot}_\theta \end{pmatrix} = \begin{pmatrix} -1 & 0 \\ 0 & \text{Rot}_\theta \end{pmatrix}.$$

Come per le rotazioni, i casi $\theta = 0$ e $\theta = \pi$ sono un po' particolari:

- una antirotazione di angolo $\theta = 0$ è una riflessione rispetto a U ;
- una antirotazione di angolo $\theta = \pi$ è una riflessione rispetto all'origine.

Ricapitolando, rotazioni e antirotazioni in $\mathbb{R}^3$ sono isometrie che, rispetto ad una base ortonormale, si rappresentano nel modo seguente:

$$\begin{pmatrix} 1 & 0 \\ 0 & \text{Rot}_\theta \end{pmatrix}, \quad \begin{pmatrix} -1 & 0 \\ 0 & \text{Rot}_\theta \end{pmatrix}.$$

Questo insieme di matrici contiene anche tutte le riflessioni rispetto a piani, rette, e all'origine. Le riflessioni si ottengono con i valori $\theta = 0$ oppure π .

Nella prossima sezione dimostriamo che tutte le isometrie di $\mathbb{R}^3$ sono rotazioni o antirotazioni.

8.2.7. Isometrie dello spazio. Abbiamo classificato precedentemente le isometrie di $\mathbb{R}^2$, e passiamo adesso alle isometrie di $\mathbb{R}^3$. Osserviamo innanzitutto che descrivere esplicitamente tutte le matrici ortogonali 3×3 è possibile ma molto più laborioso che in dimensione due. Ci accontenteremo quindi di dare una descrizione geometrica chiara di tutte le isometrie possibili. Ci servirà questo fatto generale:

Proposizione 8.2.11. *Sia $T: V \rightarrow V$ una isometria di uno spazio V dotato di un prodotto scalare definito positivo. Per ogni sottospazio $U \subset V$*

$$T(U)^\perp = T(U^\perp).$$

Dimostrazione. Se $v \in U^\perp$, allora $\langle v, u \rangle = 0$ per ogni $u \in U$. Siccome T è un'isometria, otteniamo $\langle T(v), T(u) \rangle = 0$ per ogni $u \in U$, in altre parole $T(v) \in T(U)^\perp$. Abbiamo mostrato che $T(U^\perp) \subset T(U)^\perp$. I due sottospazi coincidono perché hanno la stessa dimensione, pari a $\dim V - \dim U$. $\square$

Corollario 8.2.12. *Se $T(U) = U$, allora $T(U^\perp) = U^\perp$. In altre parole se U è T -invariante allora anche $U^\perp$ è T -invariante.*

Il teorema seguente è un risultato importante della geometria dello spazio. La sua dimostrazione è anche interessante perché usa molti degli argomenti trattati nei capitoli precedenti: matrici associate, autovettori e prodotti scalari.

Teorema 8.2.13. *Ogni isometria di $\mathbb{R}^3$ è una rotazione oppure un'antirotazione.*

Dimostrazione. Sia $T: \mathbb{R}^3 \rightarrow \mathbb{R}^3$ una isometria. Siccome $\mathbb{R}^3$ ha dimensione dispari, per la Proposizione 5.1.37 esiste almeno un autovettore $v \neq 0$ per T . L'autovalore di v è ± 1 per la Proposizione 8.2.2. Quindi $T(v) = \pm v$. A meno di rinormalizzare v possiamo supporre che $\|v\| = 1$.

La retta $\text{Span}(v)$ è T -invariante e quindi anche il piano $U = \text{Span}(v)^\perp$ lo è per il Corollario 8.2.12. Prendiamo una base $\mathcal{B} = \{v, v_2, v_3\}$ ortonormale. I vettori v_2, v_3 formano una base ortonormale di U e otteniamo

$$A = [T]_{\mathcal{B}}^{\mathcal{B}} = \begin{pmatrix} 1 & 0 \\ 0 & B \end{pmatrix}$$

dove B è una matrice 2×2 ortogonale che rappresenta la restrizione $T|_U$ nella base ortonormale $\{v_2, v_3\}$. Per la Proposizione 8.2.7, la matrice B è una matrice di rotazione o di riflessione. Se è di rotazione, otteniamo

$$A = \begin{pmatrix} \pm 1 & 0 \\ 0 & \text{Rot}_\theta \end{pmatrix}$$

e quindi A è una rotazione o antirotazione a seconda del segno di ± 1 . Se B è una matrice di riflessione, esiste una base ortonormale $\{v'_2, v'_3\}$ di U rispetto alla quale la riflessione è la matrice $\begin{pmatrix} -1 & 0 \\ 0 & 1 \end{pmatrix}$. Con la base $\mathcal{B}' = \{v, v'_2, v'_3\}$ troviamo

$$A' = [T]_{\mathcal{B}'}^{\mathcal{B}'} = \begin{pmatrix} \pm 1 & 0 & 0 \\ 0 & -1 & 0 \\ 0 & 0 & 1 \end{pmatrix}.$$

Questa è una riflessione (rispetto ad un piano o ad una retta a seconda del segno di ± 1), un caso particolare di rotazione o antirotazione. $\square$

Abbiamo scoperto che ogni isometria di $\mathbb{R}^3$ è una rotazione o una antirotazione. Ci poniamo il problema seguente: concretamente, data una matrice ortogonale A , come facciamo a determinare geometricamente di che tipo di (anti-)rotazione si tratta?

Rispondere a questa domanda è molto facile. Innanzitutto, se $\det A = 1$ si tratta di una rotazione; se invece $\det A = -1$, è una antirotazione. L'angolo di (anti-)rotazione θ è inoltre facilmente determinato (a meno di segno) dalla traccia di A , tramite la formula seguente:

$$\operatorname{tr} A = \pm 1 + 2 \cos \theta$$

dove $\pm 1 = \det A$ dipende se A è una rotazione o una antirotazione. Quindi possiamo scrivere

$$\cos \theta = \frac{\operatorname{tr} A - \det A}{2}.$$

In questo modo si determina immediatamente l'angolo θ a meno di segno. Se $\theta = 0$ oppure π , la matrice A rappresenta in realtà una riflessione (oppure è l'identità), che può essere descritta trovando gli autospazi V_1 e V_{-1} .

Nel caso generico $\theta \neq 0, \pi$, l'autospazio V_1 (rispettivamente, V_{-1}) è precisamente l'asse di rotazione (rispettivamente, antirotazione).

Esempio 8.2.14. Data la matrice ortogonale

$$A = \begin{pmatrix} 0 & 0 & 1 \\ 1 & 0 & 0 \\ 0 & 1 & 0 \end{pmatrix}.$$

vediamo che $\det A = 1$ e quindi L_A è una rotazione; abbiamo $\operatorname{tr} A = 0$ e quindi $\cos \theta = -\frac{1}{2}$, che implica $\theta = \pm \frac{2\pi}{3}$. Infine, l'asse di rotazione è l'autospazio

$$V_1 = \ker(A - I) = \ker \begin{pmatrix} -1 & 0 & 1 \\ 1 & -1 & 0 \\ 0 & 1 & -1 \end{pmatrix} = \operatorname{Span} \begin{pmatrix} 1 \\ 1 \\ 1 \end{pmatrix}.$$

Esercizi

Esercizio 8.1. Considera $\mathbb{R}^3$ con il prodotto scalare euclideo. Siano $U, W \subset \mathbb{R}^3$ due piani vettoriali distinti. Mostra che

- (1) $U \cap W$, $U^\perp$ e $W^\perp$ sono tre rette vettoriali.
- (2) Vale $(U \cap W)^\perp = U^\perp + W^\perp$.
- (3) Esiste una base v_1, v_2, v_3 di $\mathbb{R}^3$ tale che

$$v_1 \in (U \cap W), \quad v_2 \in U^\perp, \quad v_3 \in W^\perp.$$

- (4) Determina i vettori v_1, v_2, v_3 nel caso seguente:

$$U = \{x + y = 0\}, \quad W = \{x + z = 0\}.$$

Dimostra questi fatti in modo algebrico, ma visualizza il problema con un disegno.

Esercizio 8.2. Considera i seguenti vettori di $\mathbb{R}^3$:

$$\begin{pmatrix} 1 \\ -1 \\ 0 \end{pmatrix}, \quad \begin{pmatrix} 2 \\ 0 \\ 1 \end{pmatrix}, \quad \begin{pmatrix} 0 \\ -1 \\ 1 \end{pmatrix}$$

- (1) Mostra che formano una base e ortogonalizzala con Gram-Schmidt.
- (2) Calcola le coordinate del vettore $2e_1 - 5e_2 + e_3$ in questa base.

Esercizio 8.3. Considera il prodotto scalare definito positivo su $\mathbb{R}_2[x]$ dato da

$$\langle p, q \rangle = p(0)q(0) + p(1)q(1) + p(2)q(2).$$

Costruisci una base ortogonale che contenga il vettore $x + x^2$.

Esercizio 8.4. Considera $\mathbb{R}^3$ con il prodotto scalare euclideo. Determina una base ortonormale per ciascuno dei sottospazi seguenti:

$$U = \text{Span} \left(\begin{pmatrix} 1 \\ 2 \\ 0 \end{pmatrix}, \begin{pmatrix} 1 \\ 0 \\ 2 \end{pmatrix} \right),$$

$$W = \{x + y + z = 0\}.$$

Determina infine una base ortonormale v_1, v_2, v_3 di $\mathbb{R}^3$ con queste proprietà:

$$v_1 \in U \cap W, \quad v_2 \in U, \quad v_3 \in W.$$

Quante sono le basi possibili di questo tipo?

Esercizio 8.5. Considera $\mathbb{R}^4$ con il prodotto scalare Euclideo. Determina una base ortonormale del sottospazio

$$W = \{w - x + y + z = 0\}.$$

Completa questa base ad una base ortonormale di $\mathbb{R}^4$.

Esercizio 8.6. Considera il prodotto scalare su $\mathbb{R}^2$ definito dalla matrice

$$\begin{pmatrix} 2 & 1 \\ 1 & 1 \end{pmatrix}.$$

Determina una base ortonormale per questo prodotto scalare.

Esercizio 8.7. Considera il piano

$$U = \text{Span} \left(\begin{pmatrix} 1 \\ 1 \\ 1 \\ 1 \end{pmatrix}, \begin{pmatrix} 1 \\ 0 \\ 1 \\ 0 \end{pmatrix} \right) \subset \mathbb{R}^4.$$

- Determina il piano ortogonale $U^\perp$.
- Trova una base ortogonale per U ed una base ortogonale per $U^\perp$.
- Nota che l'unione di queste due basi è una base $\mathcal{B}$ ortogonale per $\mathbb{R}^4$. Scrivi le coordinate dei vettori della base canonica e_1 e e_3 rispetto a questa base.

Negli esercizi seguenti, quando non è specificato diversamente, il prodotto scalare su $\mathbb{R}^n$ è sempre quello euclideo.

Esercizio 8.8. Considera il piano $U = \{x + y - 2z = 0\} \subset \mathbb{R}^3$ e scrivi le matrici associate alla proiezione ortogonale p_U e alla riflessione ortogonale r_U rispetto alla base canonica di $\mathbb{R}^3$.

Esercizio 8.9. Considera il piano $U = \{x + 2y + z = 0\} \subset \mathbb{R}^3$. Trova una base ortonormale per U e completala a base ortonormale per $\mathbb{R}^3$.

Esercizio 8.10. Considera il piano $U = \{4x - 3y + z = 0\} \subset \mathbb{R}^3$ e indica con p e q le proiezioni ortogonali di $\mathbb{R}^3$ rispettivamente su U e $U^\perp$. Considera la matrice

$$A = \begin{pmatrix} 2 & 1 & -1 \\ 1 & 4 & 2 \\ -1 & 2 & 3 \end{pmatrix}.$$

- (1) Scrivi la matrice associata a p rispetto alla base canonica di $\mathbb{R}^3$.
- (2) L'endomorfismo $f(x) = p(x) - q(x)$ è diagonalizzabile?
- (3) Determina il sottospazio ortogonale a U rispetto al prodotto scalare $\langle x, y \rangle = {}^t x A y$.

Esercizio 8.11. Siano $U, U' \subset \mathbb{R}^3$ due piani distinti. Considera la composizione $f = p_{U'} \circ p_U$ delle proiezioni ortogonali su U e U' . Mostra che f ha un autovalore 0, un autovalore 1, e un terzo autovalore $\lambda \in [0, 1]$.

Suggerimento. Disegna i due piani e determina geometricamente tre autovettori indipendenti per f .

Esercizio 8.12. Trova assi e angoli di rotazione delle isometrie di $\mathbb{R}^3$:

$$\begin{pmatrix} 0 & 1 & 0 \\ 0 & 0 & -1 \\ -1 & 0 & 0 \end{pmatrix}, \quad \begin{pmatrix} 0 & 1 & 0 \\ 1 & 0 & 0 \\ 0 & 0 & -1 \end{pmatrix}, \quad \frac{1}{3} \begin{pmatrix} -1 & 2 & 2 \\ 2 & -1 & 2 \\ 2 & 2 & -1 \end{pmatrix}$$

Esercizio 8.13. Scrivi una matrice $A \in M(3)$ che rappresenti una rotazione intorno all'asse

$$r = \text{Span} \begin{pmatrix} 1 \\ 1 \\ -1 \end{pmatrix}$$

di angolo $\frac{2\pi}{3}$.

Esercizio 8.14. Se $r_U: \mathbb{R}^3 \rightarrow \mathbb{R}^3$ è la riflessione lungo un piano $U \subset \mathbb{R}^3$, che tipo di isometria è $-r_U$?

Esercizio 8.15. Mostra che esistono esattamente 48 matrici ortogonali $A \in M(3)$ i cui coefficienti A_{ij} siano tutti numeri interi. Quante di queste preservano l'orientazione di $\mathbb{R}^3$? Quali sono i loro assi? L'insieme

$$C = \left\{ \begin{pmatrix} x \\ y \\ z \end{pmatrix} \in \mathbb{R}^3 \mid -1 \leq x, y, z \leq +1 \right\}$$

è un cubo centrato nell'origine. Verifica che le 48 matrici ortogonali A così ottenute sono tutte simmetrie del cubo, cioè $L_A(C) = C$.

Esercizio 8.16. Scrivi una matrice $A \in M(3)$ che rappresenti una riflessione rispetto a:

- (1) il piano $U = \{x - y + 3z = 0\}$;
- (2) la retta $U = \text{Span}(e_1 + e_3)$;
- (3) l'origine $U = \{0\}$.

Esercizio 8.17. Scrivi una isometria di $\mathbb{R}^4$ che non ha autovettori. Mostra che in $\mathbb{R}^{2n}$ esistono sempre isometrie senza autovettori.

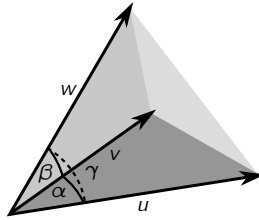


Figura 8.7. Tre vettori non nulli u, v, w nello spazio e gli angoli α, β, γ fra questi.

Esercizio 8.18. Sia $R: \mathbb{R}^3 \rightarrow \mathbb{R}^3$ una rotazione di un angolo θ lungo un asse r . Sia $P: \mathbb{R}^3 \rightarrow \mathbb{R}^3$ la proiezione ortogonale su un piano U che contiene r . Per quali valori di $\theta \in [0, \pi]$ l'endomorfismo $P \circ R$ è diagonalizzabile?

Complementi

8.1. Angoli fra tre vettori nello spazio. Teniamo dritti pollice, indice e medio della mano destra: queste tre dita formano a coppie tre angoli α, β e γ . Usiamo gli strumenti di questo capitolo per mostrare un fatto intuitivo di geometria solida: fra questi tre angoli vale la disuguaglianza triangolare, cioè un angolo non può essere più grande della somma degli altri due, e può essere uguale alla somma degli altri due solo se le tre dita sono complanari. Questo fatto vale in un qualsiasi spazio vettoriale definito positivo.

Sia come sempre V uno spazio vettoriale munito di un prodotto scalare definito positivo. Siano $u, v, w \in V$ tre vettori non nulli e α, β, γ gli angoli fra questi, come mostrato nella Figura 8.7.

Proposizione 8.2.15. *Abbiamo $\gamma \leq \alpha + \beta$. Inoltre $\gamma = \alpha + \beta$ se e solo se vale uno dei fatti seguenti:*

- (1) u e w sono indipendenti e $v = \lambda u + \mu w$ per qualche $\lambda, \mu \geq 0$;
- (2) $u = \lambda w$ per qualche $\lambda < 0$;
- (3) $u = \lambda w$ e $v = \mu w$ per qualche $\lambda, \mu > 0$.

In particolare, se u, v, w sono indipendenti vale sempre $\gamma < \alpha + \beta$. L'uguaglianza $\gamma = \alpha + \beta$ si ottiene solo se i tre vettori sono complanari (cioè contenuti in un piano) e configurati opportunamente (il vettore v deve stare "tra u e w ": questo è il senso geometrico della condizione (1)).

Dimostrazione. Possiamo supporre che u, v, w siano unitari. Abbiamo

$$\cos \alpha = \langle u, v \rangle, \quad \cos \beta = \langle v, w \rangle, \quad \cos \gamma = \langle w, u \rangle.$$

Consideriamo prima il caso in cui u, v, w siano indipendenti. Sia v_1, v_2, v_3 la terna di vettori ortonormali ottenuti da v, w, u tramite l'ortogonalizzazione di Gram-Schmidt. Abbiamo $v_1 = v$, $w \in \text{Span}(v_1, v_2)$ e $u \in$

$\text{Span}(v_1, v_2, v_3)$. Troviamo

$$\begin{aligned} v &= v_1, \\ w &= \cos\beta v_1 + \sin\beta v_2, \\ u &= \cos\alpha v_1 + a v_2 + b v_3 \end{aligned}$$

con $\cos^2\alpha + a^2 + b^2 = 1$ e $b > 0$. Ne deduciamo che $|a| < \sin\alpha$ e allora

$$\begin{aligned} \cos\gamma &= \langle w, u \rangle = \cos\alpha \cos\beta + a \sin\beta \\ &> \cos\alpha \cos\beta - \sin\alpha \sin\beta = \cos(\alpha + \beta). \end{aligned}$$

Quindi $\gamma < \alpha + \beta$. Resta da considerare il caso in cui u, v, w siano dipendenti. Se v e w sono indipendenti, si può usare l'ortogonalizzazione di Gram-Schmidt su v, w , trovare v_1, v_2 e ottenere le stesse uguaglianze di sopra con $b = 0$ e $a = \pm \sin\alpha$. Analogamente si deduce che $\gamma \leq \alpha + \beta$ e inoltre

$$\gamma = \alpha + \beta \iff a = -\sin\alpha \text{ e } \alpha + \beta \leq \pi.$$

Questo accade precisamente quando $a \leq 0$ e $\alpha + \beta \leq \pi$. Questo equivale a chiedere che v giaccia "tra u e w ", cioè $v = \lambda u + \mu w$ per qualche $\lambda, \mu \geq 0$.

Se u e w sono multipli si conclude facilmente. $\square$

Osservazione 8.2.16. Esiste un'altra dimostrazione della disuguaglianza $\gamma < \alpha + \beta$ nel caso in cui i vettori u, v, w siano indipendenti e unitari. Consideriamo lo spazio tridimensionale $U = \text{Span}(u, v, w)$. Sia g il prodotto scalare definito positivo su V . La *matrice di Gram* è la matrice $S = [g|_U]_{\mathcal{B}}$ che rappresenta $g|_U$ nella base $\mathcal{B} = \{u, v, w\}$. Ha la forma seguente:

$$S = \begin{pmatrix} \langle u, u \rangle & \langle u, v \rangle & \langle u, w \rangle \\ \langle v, u \rangle & \langle v, v \rangle & \langle v, w \rangle \\ \langle w, u \rangle & \langle w, v \rangle & \langle w, w \rangle \end{pmatrix} = \begin{pmatrix} 1 & \cos\alpha & \cos\gamma \\ \cos\alpha & 1 & \cos\beta \\ \cos\gamma & \cos\beta & 1 \end{pmatrix}.$$

Poiché $g|_U$ è definito positivo, si deve avere $\det S > 0$. Con un po' di impegno si dimostra che questa condizione implica che $\gamma < \alpha + \beta$.

Lo spazio euclideo

Nei capitoli precedenti abbiamo prodotto una notevole quantità di arnesi atti a costruire e studiare gli spazi vettoriali nella più ampia generalità. In questo capitolo e nel successivo, applichiamo finalmente questi strumenti per esaminare da vicino la geometria dello spazio $\mathbb{R}^n$, munito del prodotto scalare euclideo.

Siamo interessati soprattutto al piano $\mathbb{R}^2$ e allo spazio $\mathbb{R}^3$ che sono ovviamente quelli che ci riguardano più da vicino perché descrivono lo spazio in cui viviamo. Iniziamo introducendo il *prodotto vettoriale*, un'operazione importante fra vettori dello spazio. Quindi studiamo i sottospazi affini nel piano e nello spazio e infine le loro trasformazioni.

9.1. Prodotto vettoriale

Il prodotto scalare è una operazione che prende due vettori e restituisce uno scalare. Introduciamo qui un'altra operazione, il *prodotto vettoriale*, che prende due vettori e restituisce un *vettore*. Il prodotto vettoriale è definito solo su $\mathbb{R}^3$.

9.1.1. Definizione. Consideriamo due vettori $v, w \in \mathbb{R}^3$:

$$v = \begin{pmatrix} v_1 \\ v_2 \\ v_3 \end{pmatrix}, \quad w = \begin{pmatrix} w_1 \\ w_2 \\ w_3 \end{pmatrix}.$$

Il *prodotto vettoriale* fra v e w è il vettore

$$v \times w = \begin{pmatrix} v_2 w_3 - v_3 w_2 \\ v_3 w_1 - v_1 w_3 \\ v_1 w_2 - v_2 w_1 \end{pmatrix}.$$

Notiamo che

$$v \times w = \begin{pmatrix} d_1 \\ -d_2 \\ d_3 \end{pmatrix}$$

dove d_i è il determinante del minore 2×2 ottenuto cancellando la i -esima riga dalla matrice

$$A = \begin{pmatrix} v_1 & w_1 \\ v_2 & w_2 \\ v_3 & w_3 \end{pmatrix}.$$

Un metodo per tenere a mente la definizione di prodotto vettoriale è il seguente. Si ottiene il prodotto vettoriale calcolando formalmente il determinante di questa “matrice”:

$$\begin{aligned} v \times w &= \det \begin{pmatrix} v_1 & w_1 & e_1 \\ v_2 & w_2 & e_2 \\ v_3 & w_3 & e_3 \end{pmatrix} \\ &= \det \begin{pmatrix} v_2 & w_2 \\ v_3 & w_3 \end{pmatrix} e_1 - \det \begin{pmatrix} v_1 & w_1 \\ v_3 & w_3 \end{pmatrix} e_2 + \det \begin{pmatrix} v_1 & w_1 \\ v_2 & w_2 \end{pmatrix} e_3. \end{aligned}$$

Questa è solo una regola mnemonica: quella matrice in realtà non è una vera matrice perché e_1, e_2, e_3 non sono numeri ma i vettori della base canonica.

Esempio 9.1.1. Vale $e_1 \times e_2 = e_3$.

9.1.2. Proprietà. Dimostriamo adesso alcune proprietà del prodotto vettoriale. Siano $v, w \in \mathbb{R}^3$ due vettori qualsiasi.

Proposizione 9.1.2. *Il vettore $v \times w$ è ortogonale sia a v che a w .*

Dimostrazione. Si ottiene

$$\begin{aligned} \langle v \times w, v \rangle &= \det \begin{pmatrix} v_2 & w_2 \\ v_3 & w_3 \end{pmatrix} v_1 - \det \begin{pmatrix} v_1 & w_1 \\ v_3 & w_3 \end{pmatrix} v_2 + \det \begin{pmatrix} v_1 & w_1 \\ v_2 & w_2 \end{pmatrix} v_3 \\ &= \det \begin{pmatrix} v_1 & w_1 & v_1 \\ v_2 & w_2 & v_2 \\ v_3 & w_3 & v_3 \end{pmatrix} = 0. \end{aligned}$$

La seconda uguaglianza è lo sviluppo di Laplace sull'ultima colonna. Il determinante è nullo perché la matrice ha due colonne uguali. Analogamente si trova $\langle v \times w, w \rangle = 0$. $\square$

Proposizione 9.1.3. *Il vettore $v \times w$ è nullo $\iff v$ e w sono dipendenti.*

Dimostrazione. Ricordiamo che

$$v \times w = \begin{pmatrix} d_1 \\ -d_2 \\ d_3 \end{pmatrix}$$

dove d_1, d_2, d_3 sono i determinanti dei minori 2×2 della matrice A avente come colonne v e w . Per la Proposizione 3.3.13 otteniamo

$$v \times w = 0 \iff d_1 = d_2 = d_3 = 0 \iff \text{rk} A \leq 1 \iff v \text{ e } w \text{ sono dipendenti.}$$

La dimostrazione è conclusa. $\square$

Il prodotto vettoriale è un utile strumento per determinare se due vettori v, w siano indipendenti e, in caso affermativo, per costruire un terzo vettore non nullo ortogonale ad entrambi. Vogliamo adesso caratterizzare geometricamente la norma di $v \times w$.

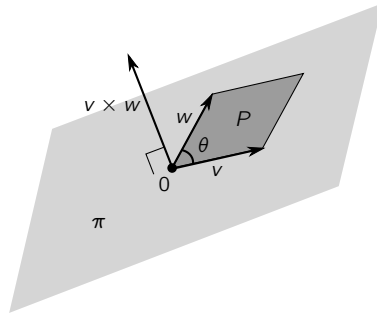


Figura 9.1. Il modulo $\|v \times w\|$ è pari all'area del parallelogramma P con lati v e w nel piano $\pi = \text{Span}(v, w)$.

Proposizione 9.1.4. *Vale l'equazione*

$$\|v \times w\|^2 + \langle v, w \rangle^2 = \|v\|^2 \|w\|^2.$$

Dimostrazione. Troviamo

$$\begin{aligned} \|v \times w\|^2 &= (v_2 w_3 - v_3 w_2)^2 + (v_1 w_3 - v_3 w_1)^2 + (v_1 w_2 - v_2 w_1)^2 \\ &= (v_1^2 + v_2^2 + v_3^2)(w_1^2 + w_2^2 + w_3^2) - (v_1 w_1 + v_2 w_2 + v_3 w_3)^2 \\ &= \|v\|^2 \|w\|^2 - \langle v, w \rangle^2. \end{aligned}$$

La dimostrazione è completa. $\square$

Supponiamo che v e w siano indipendenti. I due vettori sono contenuti nel piano $\pi = \text{Span}(v, w)$. Indichiamo con $P \subset \pi$ il parallelogramma avente lati v e w . Si veda la Figura 9.1.

Proposizione 9.1.5. *L'area di P è*

$$\text{Area}(P) = \|v\| \|w\| \sin \theta.$$

Dimostrazione. Se prendiamo v come base, l'altezza relativa ha lunghezza $\|w\| \sin \theta$. Usando la formula base $\times$ altezza si trova l'area. $\square$

Mostriamo che questa è anche uguale al modulo di $v \times w$.

Corollario 9.1.6. *Il modulo del prodotto vettoriale è*

$$\|v \times w\| = \|v\| \|w\| \sin \theta = \text{Area}(P)$$

dove θ è l'angolo formato da v e w e P è il parallelogramma con lati v e w .

Dimostrazione. Ricordiamo che $\langle v, w \rangle = \|v\| \|w\| \cos \theta$. Troviamo

$$\begin{aligned} \|v \times w\|^2 &= \|v\|^2 \|w\|^2 - \langle v, w \rangle^2 = \|v\|^2 \|w\|^2 (1 - \cos^2 \theta) \\ &= \|v\|^2 \|w\|^2 \sin^2 \theta. \end{aligned}$$

Quindi $\|v \times w\| = \|v\| \|w\| \sin \theta = \text{Area}(P)$. $\square$

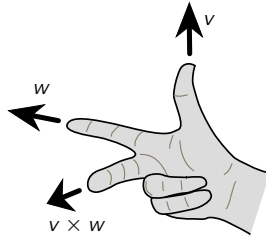


Figura 9.2. La regola della mano destra.

Se v e w sono dipendenti, allora $v \times w = 0$. Se sono indipendenti, allora la direzione di $v \times w$ è ortogonale al piano contenente v e w e il modulo di $v \times w$ è l'area di P . Queste due condizioni determinano il vettore $v \times w$ a meno di segno: per scegliere il “verso giusto” è sufficiente usare la regola della mano destra mostrata nella Figura 9.2 e giustificata dalla proposizione seguente.

Proposizione 9.1.7. *Se v e w sono indipendenti, la terna $v, w, v \times w$ è una base positiva di $\mathbb{R}^3$.*

Dimostrazione. Mostriamo che la matrice avente come colonne $v, w, v \times w$ ha determinante positivo. Sviluppiamo sulla terza colonna:

$$\det \begin{pmatrix} v_1 & w_1 & \det \begin{pmatrix} v_2 & w_2 \\ v_3 & w_3 \end{pmatrix} \\ v_2 & w_2 & -\det \begin{pmatrix} v_1 & w_1 \\ v_3 & w_3 \end{pmatrix} \\ v_3 & w_3 & \det \begin{pmatrix} v_1 & w_1 \\ v_2 & w_2 \end{pmatrix} \end{pmatrix} = \left(\det \begin{pmatrix} v_2 & w_2 \\ v_3 & w_3 \end{pmatrix} \right)^2 + \left(\det \begin{pmatrix} v_1 & w_1 \\ v_3 & w_3 \end{pmatrix} \right)^2 + \left(\det \begin{pmatrix} v_1 & w_1 \\ v_2 & w_2 \end{pmatrix} \right)^2 > 0.$$

La dimostrazione è completa. $\square$

Possiamo quindi definire geometricamente il prodotto vettoriale $v \times w$ di due vettori indipendenti come l'unico vettore ortogonale ad entrambi, di lunghezza pari all'area del parallelogramma con lati v e w , e orientato positivamente rispetto a v e w .

Dalla definizione seguono facilmente anche i fatti seguenti:

- $v \times w = -w \times v$ per ogni $v, w \in \mathbb{R}^3$,
- il prodotto vettoriale $\times: \mathbb{R}^3 \times \mathbb{R}^3 \rightarrow \mathbb{R}^3$ è bilineare.

Come per il prodotto scalare, per bilinearità intendiamo:

$$\begin{aligned} (v + v') \times w &= v \times w + v' \times w, & (\lambda v) \times w &= \lambda(v \times w), \\ v \times (w + w') &= v \times w + v \times w', & v \times (\lambda w) &= \lambda(v \times w). \end{aligned}$$

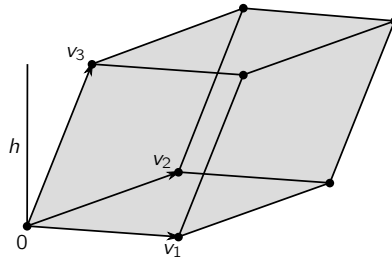


Figura 9.3. Il volume del parallelepipedo è il valore assoluto del determinante $|\det(v_1 v_2 v_3)|$.

Abbiamo visto che il prodotto vettoriale ha molte proprietà. Notiamo che a differenza del prodotto scalare, il prodotto vettoriale non è commutativo, ma *anticommutativo*: vale $v \times w = -w \times v$.

C'è infine una differenza fondamentale fra il prodotto vettoriale e altri prodotti, come quello fra numeri o fra matrici: non è verificata la proprietà associativa! Infatti ad esempio:

$$\begin{aligned}(e_1 \times e_2) \times e_2 &= e_3 \times e_2 = -e_2 \times e_3 = -e_1, \\ e_1 \times (e_2 \times e_2) &= e_1 \times 0 = 0.\end{aligned}$$

Esercizio 9.1.8 (Prodotto triplo). Siano $v_1, v_2, v_3 \in \mathbb{R}^3$. Mostra che

$$\langle v_1 \times v_2, v_3 \rangle = \langle v_1, v_2 \times v_3 \rangle = \det(v_1 | v_2 | v_3).$$

9.1.3. Applicazioni. Come prima applicazione del prodotto vettoriale possiamo finalmente dimostrare la relazione fra determinanti e volumi di parallelepipedi enuncata nella Sezione 3.3.10.

Siano $v_1, v_2, v_3 \in \mathbb{R}^3$ tre vettori indipendenti. Indichiamo con P il parallelepipedo con lati v_1, v_2 e v_3 mostrato nella Figura 9.3.

Proposizione 9.1.9. *Il volume del parallelepipedo P è*

$$\text{Vol}(P) = |\det(v_1 | v_2 | v_3)|.$$

Dimostrazione. Consideriamo come base di P il parallelogramma con lati v_1 e v_2 . Il volume di P è l'area di base per l'altezza h mostrata nella Figura 9.3. L'area di base è $\|v_1 \times v_2\|$. L'altezza si trova facendo la proiezione ortogonale di v_3 sul vettore $v_1 \times v_2$, che è ortogonale alla base. Quindi

$$\text{Vol}(P) = \|v_1 \times v_2\| \|p_{v_1 \times v_2}(v_3)\| = |\langle v_1 \times v_2, v_3 \rangle| = |\det(v_1 | v_2 | v_3)|.$$

Abbiamo usato gli Esercizi 8.1.15 e 9.1.8. □

Il prodotto vettoriale è impiegato innanzitutto per trovare velocemente un vettore ortogonale a due vettori indipendenti dati v_1 e v_2 .

Si può usare il prodotto vettoriale per passare velocemente da coordinate parametriche a cartesiane per qualsiasi piano vettoriale di $\mathbb{R}^3$. Se $W = \text{Span}(v, w)$ è un piano di $\mathbb{R}^3$, si può scrivere in forma cartesiana come

$$W = \{ax + by + cz = 0\}$$

dove a, b, c sono i coefficienti del vettore $v \times w$.

Esempio 9.1.10. Un piano

$$W = \text{Span} \left\{ \begin{pmatrix} a_1 \\ b_1 \\ c_1 \end{pmatrix}, \begin{pmatrix} a_2 \\ b_2 \\ c_2 \end{pmatrix} \right\}$$

ha come vettore ortogonale

$$v = \begin{pmatrix} a_1 \\ b_1 \\ c_1 \end{pmatrix} \times \begin{pmatrix} a_2 \\ b_2 \\ c_2 \end{pmatrix} = \begin{pmatrix} b_1 c_2 - b_2 c_1 \\ a_2 c_1 - a_1 c_2 \\ a_1 b_2 - a_2 b_1 \end{pmatrix}$$

e quindi ha equazioni cartesiane

$$W = \{(b_1 c_2 - b_2 c_1)x + (a_2 c_1 - a_1 c_2)y + (a_1 b_2 - a_2 b_1)z = 0\}.$$

Vedremo nelle pagine seguenti altre applicazioni geometriche del prodotto vettoriale. Il prodotto vettoriale è molto usato in fisica, ad esempio nell'elettromagnetismo.

9.2. Sottospazi affini

Studiamo in questa sezione i sottospazi affini di $\mathbb{R}^n$, con un occhio di riguardo per le rette e i piani in $\mathbb{R}^2$ e $\mathbb{R}^3$.

9.2.1. Forma parametrica e cartesiana. Ricordiamo brevemente la nozione di sottospazio affine introdotta nella Sezione 3.2.2.

Un sottospazio affine di uno spazio vettoriale V è un sottoinsieme del tipo $S = x + W$, dove $W \subset V$ è un sottospazio vettoriale. Per la Proposizione 3.2.4 il sottospazio W è univocamente determinato da S (mentre il punto x non lo è: si tratta di un qualsiasi punto in S), è chiamato la *giacitura* di S ed è indicato con $\text{giac}(S)$. La dimensione di S è per definizione quella di $\text{giac}(S)$.

La giacitura è determinata in modo intrinseco da S nel modo seguente:

Esercizio 9.2.1. Sia S un sottospazio affine. Abbiamo

$$\text{giac}(S) = \{\overrightarrow{PQ} \mid P \in S, Q \in S\}.$$

Siamo particolarmente interessati ai sottospazi affini di $\mathbb{R}^n$. Un tale sottospazio affine può descritto in forma *parametrica* (cioè *esplicita*) come

$$S = x + \text{Span}(v_1, \dots, v_k) = \{x + t_1 v_1 + \dots + t_k v_k \mid t_1, \dots, t_k \in \mathbb{R}\}$$

dove i vettori $v_1, \dots, v_k \in \mathbb{R}^n$ formano una base per la giacitura W . Qui ovviamente $\dim S = k$. In alternativa S può essere descritto in forma *cartesiana* (cioè *implicita*) come insieme di soluzioni di un sistema lineare

$$S = \{x \in \mathbb{R}^n \mid Ax = b\}.$$

Qui $A \in M(m, n)$ e $b \in \mathbb{R}^m$. Sappiamo per il Teorema di Rouché–Capelli che $S \neq \emptyset \iff \text{rk} A = \text{rk}(A \mid b)$ e in questo caso $\dim S = n - \text{rk} A$.

I sottospazi affini di $\mathbb{R}^3$ sono i punti, le rette, i piani, e $\mathbb{R}^3$ stesso.

Esempio 9.2.2. Un piano affine π in $\mathbb{R}^3$ è descritto da una equazione

$$\pi = \{ax + by + cz = d\}.$$

Una retta affine r in $\mathbb{R}^3$ è descritta da due equazioni di questo tipo, oppure più agevolmente in forma parametrica:

$$r = \left\{ \begin{pmatrix} x_0 \\ y_0 \\ z_0 \end{pmatrix} + t \begin{pmatrix} x \\ y \\ z \end{pmatrix} \right\}.$$

Per passare da coordinate cartesiane a parametriche si risolve il sistema $Ax = b$, ad esempio usando l'algoritmo di Gauss–Jordan.

A volte può servire fare l'opposto, cioè passare da coordinate parametriche a cartesiane. Questo passaggio è particolarmente facile per i piani in $\mathbb{R}^3$. Sia

$$\pi = \{P_0 + tv_1 + uv_2\}$$

un piano affine in $\mathbb{R}^3$, dove $v_1, v_2 \in \mathbb{R}^3$ sono vettori indipendenti. Calcoliamo il vettore $u = v_1 \times v_2$. Questo avrà tre coefficienti a, b, c . Il piano π può essere descritto come

$$\pi = \{ax + by + cz = d\}$$

per un certo $d \in \mathbb{R}$ che può essere trovato imponendo che $P_0 \in \pi$.

Esempio 9.2.3. Consideriamo

$$\pi = \left\{ \begin{pmatrix} 1 \\ 2 \\ -3 \end{pmatrix} + t \begin{pmatrix} 1 \\ 0 \\ 1 \end{pmatrix} + u \begin{pmatrix} 2 \\ -1 \\ 0 \end{pmatrix} \right\}.$$

Troviamo

$$\begin{pmatrix} 1 \\ 0 \\ 1 \end{pmatrix} \times \begin{pmatrix} 2 \\ -1 \\ 0 \end{pmatrix} = \begin{pmatrix} 1 \\ 2 \\ -1 \end{pmatrix}$$

Quindi

$$\pi = \{x + 2y - z = d\}$$

per qualche $d \in \mathbb{R}$, che determiniamo imponendo

$$\begin{pmatrix} 1 \\ 2 \\ -3 \end{pmatrix} \in \pi \implies 1 + 4 + 3 = d$$

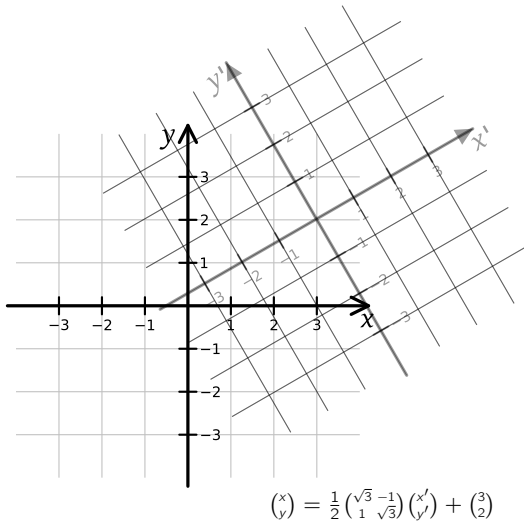


Figura 9.4. Un cambiamento di sistema di riferimento.

e quindi $d = 8$. Abbiamo trovato una equazione cartesiana per π :

$$\pi = \{x + 2y - z = 8\}.$$

9.2.2. Traslare l'origine. Abbiamo visto nei capitoli precedenti che per studiare un determinato problema in $\mathbb{R}^n$ è spesso conveniente non usare la base canonica ma *cambiare base*. Quando studiamo i sottospazi affini, possiamo anche scegliere di *traslare l'origine*. Questo vuol dire fissare un punto $P \in \mathbb{R}^n$ e prendere delle nuove coordinate x' che differiscano dalle precedenti x nel modo seguente:

$$x' = x - P.$$

In altre parole:

$$\begin{pmatrix} x'_1 \\ \vdots \\ x'_n \end{pmatrix} = \begin{pmatrix} x_1 \\ \vdots \\ x_n \end{pmatrix} - \begin{pmatrix} P_1 \\ \vdots \\ P_n \end{pmatrix}.$$

Adesso il punto P è diventato la nuova origine dello spazio. Dopo aver traslato l'origine, il sottospazio affine $P+W$ è descritto nelle nuove coordinate come il sottospazio vettoriale W .

Esempio 9.2.4. Consideriamo la retta $r = \left\{ \begin{pmatrix} x \\ y \end{pmatrix} = \begin{pmatrix} 1 \\ 2 \end{pmatrix} + t \begin{pmatrix} -1 \\ 1 \end{pmatrix} \right\}$ in $\mathbb{R}^2$. Questa è una retta affine. Adesso trasliamo l'origine, prendendo come nuova origine il punto $\begin{pmatrix} 1 \\ 2 \end{pmatrix}$. Quindi $\begin{pmatrix} x' \\ y' \end{pmatrix} = \begin{pmatrix} x \\ y \end{pmatrix} - \begin{pmatrix} 1 \\ 2 \end{pmatrix}$ e $r = \left\{ \begin{pmatrix} x' \\ y' \end{pmatrix} = t \begin{pmatrix} -1 \\ 1 \end{pmatrix} \right\}$ nelle nuove coordinate.

È spesso più utile scrivere x in funzione di x' come

$$x = x' + P.$$

Esempio 9.2.5. Consideriamo la retta r in $\mathbb{R}^2$ avente equazione $x + 2y = 3$. Cambiamo coordinate traslando l'origine nel punto $P = \begin{pmatrix} a \\ b \end{pmatrix}$, scrivendo

$$\begin{pmatrix} x \\ y \end{pmatrix} = \begin{pmatrix} x' \\ y' \end{pmatrix} + \begin{pmatrix} a \\ b \end{pmatrix}.$$

Quindi $x = x' + a$ e $y = y' + b$. Per trovare la nuova equazione per r basta sostituire e troviamo $(x' + a) + 2(y' + b) = 3$, cioè $x' + 2y' = 3 - a - 2b$.

Possiamo anche simultaneamente traslare l'origine e cambiare base di $\mathbb{R}^n$, scegliendo un'altra base ortonormale al posto di quella canonica. Una tale operazione è un *cambiamento di sistema di riferimento*. Si tratta di prendere nuove coordinate x' che differiscano dalle precedenti x nel modo seguente:

$$x = Ax' + P$$

dove $P \in \mathbb{R}^n$ è la nuova origine e $A \in M(n)$ è una matrice ortogonale che avrà l'effetto di modificare le direzioni degli assi. Si veda la Figura 9.4.

9.2.3. Intersezioni. Una peculiarità importante dei sottospazi affini, che li differenzia molto dai sottospazi vettoriali, sta nel fatto che i primi possono non intersecarsi, mentre gli ultimi devono sempre intersecarsi almeno nell'origine. Per questo motivo i sottospazi affini sono più complicati da studiare dei sottospazi vettoriali.

Due sottospazi affini $S, S' \subset \mathbb{R}^n$ sono *incidenti* se $S \cap S' \neq \emptyset$. Se sono incidenti, possiamo prendere $P \in S \cap S'$ e scrivere entrambi come

$$S = P + W, \quad S' = P + W'.$$

In questo modo chiaramente otteniamo

$$S \cap S' = P + W \cap W'.$$

Ne deduciamo in particolare che l'intersezione di due spazi affini è sempre uno spazio affine (se non è vuota).

Esempio 9.2.6. Se S e S' sono entrambi descritti in forma cartesiana, la loro intersezione $S \cap S'$ è descritta in forma cartesiana semplicemente unendo le equazioni di S e S' . Ad esempio se $S = \{x + y = 1\}$ e $S' = \{x - y + z = 3\}$ sono due piani in $\mathbb{R}^3$, la loro intersezione $S \cap S'$ è l'insieme delle soluzioni di

$$\begin{cases} x + y = 1, \\ x - y + z = 3. \end{cases}$$

Se S è descritto in forma cartesiana e S' in forma parametrica, per trovare $S \cap S'$ è sufficiente dare il punto generico di S' in pasto alle equazioni

di S e trovare i parametri che le soddisfano. Ad esempio, se $S = \{x + y - z = 2\}$ è un piano in $\mathbb{R}^3$ e

$$S' = \left\{ \begin{pmatrix} 1 \\ -1 \\ 1 \end{pmatrix} + t \begin{pmatrix} 1 \\ 2 \\ -3 \end{pmatrix} \right\} = \left\{ \begin{pmatrix} 1+t \\ -1+2t \\ 1-3t \end{pmatrix} \right\}$$

è una retta in $\mathbb{R}^3$, per trovare $S \cap S'$ si sostituisce

$$x = 1 + t, \quad y = -1 + 2t, \quad z = 1 - 3t$$

nell'equazione che definisce S e si ottiene

$$1 + t - 1 + 2t - 1 + 3t = 2 \iff t = \frac{1}{2}$$

e quindi l'intersezione $S \cap S'$ è il punto

$$S \cap S' = \left\{ \frac{1}{2} \begin{pmatrix} 3 \\ 0 \\ -1 \end{pmatrix} \right\}.$$

Se entrambi S e S' sono in forma parametrica, si eguaglia il punto generico di S con quello di S' e si trova quali parametri risolvono il sistema. Questo procedimento può essere più laborioso dei precedenti.

Se S e S' sono due sottospazi affini incidenti di $\mathbb{R}^n$, possiamo all'occorrenza scegliere un punto $P \in S \cap S'$ e quindi spostare l'origine in P . Con queste nuove coordinate gli spazi S, S' e $S \cap S'$ coincidono con le loro giaciture W, W' e $W \cap W'$; adesso sono sottospazi vettoriali e possiamo applicare le tecniche dei capitoli precedenti per studiarli.

Se S e S' non sono incidenti, si comportano in modo molto diverso da quanto visto per i sottospazi vettoriali, come vedremo fra poco.

Siano $S, S' \subset \mathbb{R}^n$ due sottospazi affini. Vediamo subito una condizione sulle giaciture che garantisce che i sottospazi si intersechino.

Proposizione 9.2.7. *Se $\text{giac}(S) + \text{giac}(S') = \mathbb{R}^n$, allora i sottospazi S e S' sono incidenti.*

Dimostrazione. Abbiamo

$$S = \{P + t_1 v_1 + \cdots + t_k v_k\}, \quad S' = \{Q + u_1 w_1 + \cdots + u_h w_h\}.$$

I due spazi sono incidenti $\iff$ il sistema

$$P + t_1 v_1 + \cdots + t_k v_k = Q + u_1 w_1 + \cdots + u_h w_h$$

nelle variabili $t_1, \dots, t_k, u_1, \dots, u_h$ ha soluzione. Per ipotesi $v_1, \dots, v_k, w_1, \dots, w_h$ generano $\mathbb{R}^n$, quindi si può scrivere $P - Q$ come loro combinazione lineare, e questa combinazione lineare fornisce una soluzione. $\square$

9.2.4. Sottospazio generato. Sia $X \subset \mathbb{R}^n$ un sottoinsieme qualsiasi. Il *sottospazio affine generato* da X è l'intersezione S di tutti i sottospazi affini contenenti X . L'intersezione S è effettivamente uno spazio affine.

Se X è un insieme finito di punti, il sottospazio generato S può essere descritto esplicitamente.

Proposizione 9.2.8. Se $X = \{P_0, \dots, P_k\}$, allora $S = P_0 + W$ con

$$W = \text{Span}(\overrightarrow{P_0P_1}, \dots, \overrightarrow{P_0P_k}).$$

Dimostrazione. Qualsiasi sottospazio affine S' contenente X è del tipo $S' = P_0 + W'$ per qualche sottospazio W' . Siccome $P_i \in S'$, il vettore $\overrightarrow{P_0P_i}$ appartiene a W' , per ogni i . Quindi $W' \supset W$. Ne segue che $S' \supset S$ e quindi S è l'intersezione di tutti i sottospazi affini contenenti X . $\square$

Notiamo che in questo caso $\dim S = \dim W \leq k$. Diciamo che i punti $P_0, \dots, P_k$ sono *affinemente indipendenti* se $\dim S = k$. Notiamo questa conseguenza.

Corollario 9.2.9. Dati $k + 1$ punti affinemente indipendenti in $\mathbb{R}^n$, esiste un unico sottospazio affine S di dimensione k che li contiene (e nessun sottospazio di dimensione $< k$).

Vediamo alcuni esempi. Due punti $P_0, P_1 \in \mathbb{R}^n$ sono affinemente indipendenti se e solo se sono distinti. Due punti distinti generano una retta r , in forma parametrica

$$r = P_0 + \text{Span}(\overrightarrow{P_0P_1}).$$

Ritroviamo qui una versione di uno dei postulati di Euclide: per due punti distinti di $\mathbb{R}^n$ passa una sola retta.

Esempio 9.2.10. La retta passante per i punti $P_0 = e_1 + e_2$ e $P_1 = e_2 - 2e_3$ in $\mathbb{R}^3$ ha equazione parametrica

$$r = \left\{ \begin{pmatrix} 1 \\ 1 \\ 0 \end{pmatrix} + t \begin{pmatrix} -1 \\ 0 \\ -2 \end{pmatrix} \right\}.$$

Proposizione 9.2.11. La retta in $\mathbb{R}^2$ passante per $P_0 = \begin{pmatrix} x_0 \\ y_0 \end{pmatrix}$ e $P_1 = \begin{pmatrix} x_1 \\ y_1 \end{pmatrix}$ ha equazione cartesiana

$$\det \begin{pmatrix} x - x_0 & y - y_0 \\ x_1 - x_0 & y_1 - y_0 \end{pmatrix} = 0.$$

Dimostrazione. L'equazione è lineare e soddisfatta da P_0 e P_1 . Quindi descrive la retta passante per i due punti. $\square$

Tre punti $P_0, P_1, P_2 \in \mathbb{R}^n$ sono affinemente indipendenti se e solo se non sono *allineati*, cioè contenuti in una retta. Per tre punti non allineati passa un solo piano π . Questo ha equazione parametrica

$$\pi = \{P_0 + t\overrightarrow{P_0P_1} + u\overrightarrow{P_0P_2}\}.$$

Il fatto che i punti non siano allineati garantisce che i vettori $\overrightarrow{P_0P_1}$ e $\overrightarrow{P_0P_2}$ siano indipendenti.

Proposizione 9.2.12. *Consideriamo tre punti $P_0, P_1, P_2 \in \mathbb{R}^3$ non allineati. Siano x_i, y_i, z_i le coordinate di P_i . Il piano in $\mathbb{R}^3$ passante per P_0, P_1 e P_2 ha equazione cartesiana*

$$\det \begin{pmatrix} x - x_0 & y - y_0 & z - z_0 \\ x_1 - x_0 & y_1 - y_0 & z_1 - z_0 \\ x_2 - x_0 & y_2 - y_0 & z_2 - z_0 \end{pmatrix} = 0.$$

Dimostrazione. Stessa dimostrazione della Proposizione 9.2.11. L'equazione è lineare e soddisfatta dai 3 punti. $\square$

La nozione di sottospazio generato è utile anche in presenza di alcuni insiemi X infiniti, ad esempio nel caso in cui X sia unione di un sottospazio affine e di un punto esterno ad esso.

Proposizione 9.2.13. *Sia $S \subset \mathbb{R}^n$ un sottospazio affine e $P \notin S$. Il sottospazio affine S' generato da $S \cup \{P\}$ ha dimensione*

$$\dim S' = \dim S + 1.$$

Dimostrazione. Abbiamo

$$S = Q + \text{Span}(v_1, \dots, v_k)$$

con $v_1, \dots, v_k$ indipendenti. Definiamo

$$S' = P + \text{Span}(\overrightarrow{PQ}, v_1, \dots, v_k).$$

Si vede facilmente che $S \subset S'$ e che qualsiasi sottospazio contenente P e S deve contenere S' . Quindi S' è il sottospazio generato da $S \cup \{P\}$.

Il vettore $\overrightarrow{PQ}$ non è combinazione lineare dei $v_1, \dots, v_k$ perché $P \notin S$. Quindi effettivamente $\dim S' = k + 1 = \dim S + 1$. $\square$

Troviamo come conseguenza un importante fatto geometrico.

Corollario 9.2.14. *Dati una retta $r \subset \mathbb{R}^3$ e un punto $P \notin r$, esiste un unico piano $\pi \subset \mathbb{R}^3$ che li contiene entrambi.*

Concretamente, se $r = Q + \text{Span}(v)$ allora il piano π è

$$\pi = P + \text{Span}(\overrightarrow{PQ}, v).$$

Esempio 9.2.15. Il piano π che contiene il punto $P = e_1 - e_2$ e la retta r dell'Esercizio 9.2.10 è

$$\pi = \left\{ \begin{pmatrix} 1 \\ -1 \\ 0 \end{pmatrix} + t \begin{pmatrix} 0 \\ 2 \\ 0 \end{pmatrix} + u \begin{pmatrix} -1 \\ 0 \\ 2 \end{pmatrix} \right\}.$$

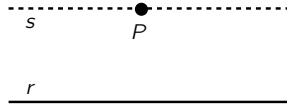


Figura 9.5. Il V Postulato di Euclide: dati r e $P \notin r$, esiste un'unica retta s passante per P parallela ad r .

9.2.5. Parallelismi. Introduciamo per la prima volta in questo libro la nozione di parallelismo. Diciamo che due sottospazi affini $S, S' \subset \mathbb{R}^n$ sono *paralleli* se $\text{giac}(S) \subset \text{giac}(S')$ oppure $\text{giac}(S') \subset \text{giac}(S)$.

Proposizione 9.2.16. Due sottospazi affini S e S' paralleli sono disgiunti oppure uno contenuto nell'altro.

Dimostrazione. Supponiamo che siano incidenti. Dopo aver spostato l'origine in un punto di $S \cap S'$ possiamo supporre che S e S' coincidano con le loro giaciture e quindi sono uno contenuto nell'altro per ipotesi. $\square$

Corollario 9.2.17. Due sottospazi affini paralleli della stessa dimensione coincidono oppure sono disgiunti.

Due sottospazi S e S' che non sono né incidenti né paralleli sono detti *sghebbi*. Ci convinciamo che la geometria che abbiamo definito in $\mathbb{R}^2$ è effettivamente quella euclidea dimostrando il V Postulato di Euclide, descritto nella Figura 9.5.

Proposizione 9.2.18 (V Postulato di Euclide). Data una retta r in $\mathbb{R}^2$ e un punto $P \notin r$, esiste un'unica retta s parallela a r passante per P .

Dimostrazione. La retta s deve passare per P e avere la giacitura di r , quindi è necessariamente $s = P + \text{giac}(r)$. $\square$

9.2.6. Ortogonalità. Definiamo una nozione di ortogonalità fra rette e altri sottospazi affini.

Siano r e S una retta e un sottospazio affine di $\mathbb{R}^n$. Diciamo che r e S sono *ortogonali* se si intersecano in un punto $r \cap S = \{P\}$ e se tutti i vettori di $\text{giac}(r)$ sono ortogonali a tutti i vettori di $\text{giac}(S)$.

Mostriamo un fatto geometrico importante.

Proposizione 9.2.19. Sia $S \subset \mathbb{R}^n$ un sottospazio affine e $P \notin S$ un punto. Esiste un'unica retta r passante per P ortogonale a S .

Dimostrazione. Sia S' il sottospazio generato da P e S . Ricordiamo che $\dim S' = \dim S + 1$. Una tale retta r deve essere contenuta in S' . La sua giacitura deve essere ortogonale alla giacitura di S ma contenuta nella giacitura di S' . Deve quindi valere la somma diretta ortogonale

$$\text{giac}(S') = \text{giac}(S) \oplus \text{giac}(r).$$

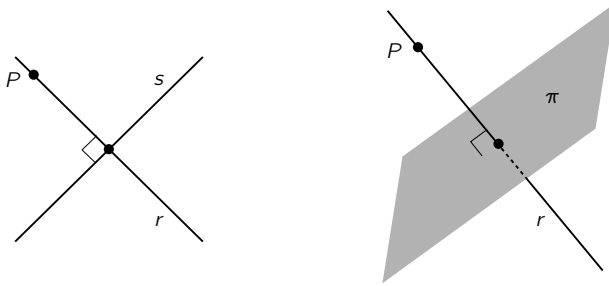


Figura 9.6. Dato un punto P ed una retta s (sinistra) oppure un piano π (destra) disgiunti da P , esiste un'unica retta r passante per P ortogonale a s (sinistra) o π (destra).

In altre parole $\text{giac}(r)$ è la retta ortogonale a $\text{giac}(S)$ dentro $\text{giac}(S')$. Questo identifica $\text{giac}(r)$ e quindi $r = P + \text{giac}(r)$ univocamente. $\square$

Questo fatto generale ha alcune conseguenze importanti nella geometria del piano e dello spazio, mostrate nella Figura 9.6.

Corollario 9.2.20. *Dati una retta s in $\mathbb{R}^2$ oppure $\mathbb{R}^3$ e un punto $P \notin s$, esiste un'unica retta r passante per P ortogonale a s .*

Corollario 9.2.21. *Dati un piano $\pi \subset \mathbb{R}^3$ e un punto $P \notin \pi$, esiste un'unica retta r passante per P ortogonale a π .*

Esempio 9.2.22. Calcoliamo concretamente r in un caso. Siano

$$P = \begin{pmatrix} 1 \\ 2 \\ 3 \end{pmatrix}, \quad s = \left\{ \begin{pmatrix} 1 \\ 1 \\ 1 \end{pmatrix} + t \begin{pmatrix} -1 \\ 1 \\ -1 \end{pmatrix} \right\}.$$

Determiniamo la retta r passante per P ortogonale a s . Scriviamo il piano π contenente P e s in forma parametrica:

$$\pi = \left\{ \begin{pmatrix} 1 \\ 2 \\ 3 \end{pmatrix} + t \begin{pmatrix} 0 \\ -1 \\ -2 \end{pmatrix} + u \begin{pmatrix} -1 \\ 1 \\ -1 \end{pmatrix} \right\}.$$

Dobbiamo determinare la somma diretta ortogonale $\text{giac}(r) \oplus \text{giac}(s) = \text{giac}(\pi)$. A questo scopo ortogonalizziamo la base trovata per $\text{giac}(\pi)$ lasciando fisso il vettore che genera $\text{giac}(s)$ e modificando l'altro. Quindi

$$\text{giac}(\pi) = \text{Span} \left\{ \begin{pmatrix} 1 \\ -4 \\ -5 \end{pmatrix}, \begin{pmatrix} -1 \\ 1 \\ -1 \end{pmatrix} \right\}.$$

La giacitura di r è il nuovo vettore trovato. Quindi

$$r = \left\{ \begin{pmatrix} 1 \\ 2 \\ 3 \end{pmatrix} + t \begin{pmatrix} 1 \\ -4 \\ -5 \end{pmatrix} \right\}.$$

Verifichiamo che effettivamente r sia ortogonale ad s : le due giaciture sono ortogonali e le due rette si intersecano nel punto

$$\frac{1}{3} \begin{pmatrix} 4 \\ 2 \\ 4 \end{pmatrix}.$$

Un metodo alternativo per trovare r passa per l'esercizio seguente.

Esercizio 9.2.23. Dati un punto P e una retta r in $\mathbb{R}^3$, esiste un unico piano π contenente P e ortogonale a r .

Questo è il metodo alternativo:

Esempio 9.2.24. Con i dati dell'Esempio 9.2.22, possiamo determinare il piano π' contenente P e ortogonale a s . Il piano ha equazione $ax + by + cz = d$ con a, b, c coefficienti del vettore che genera $\text{giac}(s)$, quindi

$$\pi' = \{-x + y - z = d\}$$

per un d da determinare imponendo che π' contenga P . Troviamo $d = -1 + 2 - 3 = -2$. A questo punto troviamo

$$Q = \pi' \cap r = \frac{1}{3} \begin{pmatrix} 4 \\ 2 \\ 4 \end{pmatrix}.$$

La retta r è l'unica retta passante per P e Q .

9.2.7. Posizioni. Studiamo adesso in che posizione reciproca possono stare due sottospazi affini in $\mathbb{R}^2$ e $\mathbb{R}^3$. Iniziamo con due rette affini nel piano.

Proposizione 9.2.25. Due rette affini distinte r e r' in $\mathbb{R}^2$ sono:

- incidenti se $\text{giac}(r) \neq \text{giac}(r')$,
- parallele se $\text{giac}(r) = \text{giac}(r')$.

Dimostrazione. Segue dalla Proposizione 9.2.7. □

Passiamo a studiare le posizioni di due rette nello spazio. Due rette affini distinte r e r' in $\mathbb{R}^3$ possono essere:

- incidenti se $r \cap r' \neq \emptyset$;
- parallele se $r \cap r' = \emptyset$ e $\text{giac}(r) = \text{giac}(r')$;
- sghembe se $r \cap r' = \emptyset$ e $\text{giac}(r) \neq \text{giac}(r')$.

Le tre possibilità sono disegnate nella Figura 9.7. Due rette r e r' sono *complanari* se esiste un piano π che le contiene entrambe.

Proposizione 9.2.26. Due rette r e r' in $\mathbb{R}^3$ sono *complanari* $\iff$ sono incidenti o parallele.

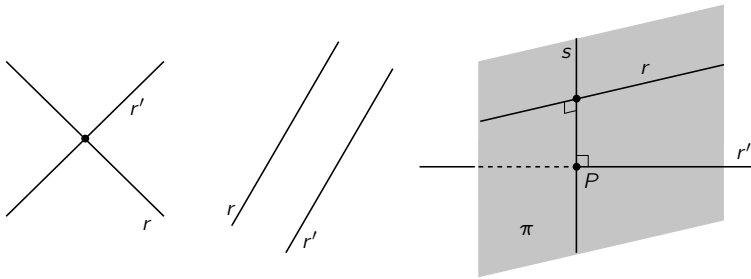


Figura 9.7. Due rette r, r' nello spazio possono essere incidenti (sinistra), parallele (centro) o sghembe (destra). Due rette sghembe hanno un'unica perpendicolare comune s (destra).

Dimostrazione. Se sono contenute in un piano, sono incidenti o parallele per la Proposizione 9.2.25. D'altro canto, se sono incidenti prendiamo $P \in r \cap r'$ e notiamo che r e r' sono contenute nel piano $\pi = P + \text{giac}(r) + \text{giac}(r')$. Se sono parallele, prendiamo $P \in r$ e $P' \in r'$ e notiamo (esercizio) che r e r' sono contenute nel piano

$$\pi = P + \text{giac}(r) + \text{Span}(\overrightarrow{PP'}).$$

La dimostrazione è conclusa. $\square$

Le rette sghembe non sono complanari e quindi sono le più difficili da visualizzare. Ad esempio il risultato seguente non è geometricamente intuitivo come i precedenti e può risultare sorprendente.

Proposizione 9.2.27. *Date due rette sghembe r, r' , esiste un'unica retta s ortogonale ad entrambe.*

Dimostrazione. Per ipotesi $\text{giac}(r)$ e $\text{giac}(r')$ sono rette vettoriali distinte, quindi $W = \text{giac}(r) + \text{giac}(r')$ è un piano vettoriale e $W^\perp$ è l'unica retta vettoriale ortogonale sia a $\text{giac}(r)$ che a $\text{giac}(r')$. Esaminiamo il piano

$$\pi = r + W^\perp = \{p + v \mid p \in r, v \in W^\perp\},$$

disegnato nella Figura 9.7-(destra). Abbiamo $\text{giac}(\pi) = \text{giac}(r) + W^\perp$.

Siccome $\text{giac}(\pi) + \text{giac}(r') = \mathbb{R}^3$, per la Proposizione 9.2.7 i sottospazi π e r' si intersecano in un punto P come mostrato in figura. La retta s cercata è l'unica retta contenente P e ortogonale a r , cioè $s = P + W^\perp$. $\square$

Esempio 9.2.28. Consideriamo le rette

$$r = \left\{ \begin{pmatrix} 1 \\ 1 \\ 1 \end{pmatrix} + t \begin{pmatrix} -1 \\ 2 \\ 1 \end{pmatrix} \right\}, \quad r' = \left\{ \begin{pmatrix} 0 \\ 1 \\ 2 \end{pmatrix} + u \begin{pmatrix} 1 \\ 1 \\ 1 \end{pmatrix} \right\}.$$

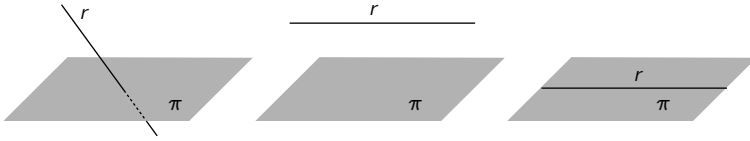


Figura 9.8. Una retta r ed un piano π possono essere incidenti (sinistra), paralleli (centro) o entrambi (destra).

Non sono parallele e si verifica facilmente che non si intersecano, quindi sono sghembe. Determiniamo la retta s ortogonale ad entrambe. Abbiamo $W = \text{giac}(r) + \text{giac}(r')$ e con il prodotto vettoriale si trova

$$W^\perp = \text{Span} \begin{pmatrix} 1 \\ 2 \\ -3 \end{pmatrix}.$$

Il piano $\pi = r + W^\perp$ è

$$\pi = \left\{ \begin{pmatrix} 1 \\ 1 \\ 1 \end{pmatrix} + t \begin{pmatrix} -1 \\ 2 \\ 1 \end{pmatrix} + u \begin{pmatrix} 1 \\ 2 \\ -3 \end{pmatrix} \right\} = \{4x + y + 2z = 7\}.$$

Questo interseca r' nel punto

$$P = \frac{1}{7} \begin{pmatrix} 2 \\ 9 \\ 16 \end{pmatrix}.$$

Quindi la retta cercata s è

$$s = \left\{ \frac{1}{7} \begin{pmatrix} 2 \\ 9 \\ 16 \end{pmatrix} + t \begin{pmatrix} 1 \\ 2 \\ -3 \end{pmatrix} \right\}.$$

Abbiamo studiato le posizioni reciproche di due rette; adesso passiamo a considerare le posizioni reciproche di una retta e un piano.

Un piano π ed una retta r nello spazio sono:

- incidenti se $\text{giac}(\pi) \oplus \text{giac}(r) = \mathbb{R}^3$,
- paralleli se $\text{giac}(r) \subset \text{giac}(\pi)$,
- incidenti e paralleli se $r \subset \pi$.

I tre casi sono mostrati nella Figura 9.8. L'intersezione $r \cap \pi$ è rispettivamente un punto, vuota, e l'intera retta r nei tre casi.

9.2.8. Fasci. Introduciamo uno strumento che, oltre ad avere un valore geometrico intrinseco, è a volte utile per fare i conti più velocemente.

Un *fascio* di rette nel piano o di piani nello spazio è un insieme di rette o piani che soddisfano una semplice condizione geometrica.

Rette. Introduciamo due tipi di fasci di rette nel piano $\mathbb{R}^2$: quelle passanti per un punto fissato, e quelle parallele ad una retta fissata.

Sia $P_0 = \begin{pmatrix} x_0 \\ y_0 \end{pmatrix}$ un punto. Il *fascio di rette passanti per P_0* è l'insieme di tutte le rette che contengono P_0 . Vediamo come descriverle concretamente. Consideriamo le equazioni cartesiane di due rette r_0, r_1 distinte qualsiasi che passano per P_0 :

$$r_0 = \{a_0x + b_0y + c_0 = 0\}, \quad r_1 = \{a_1x + b_1y + c_1 = 0\}.$$

Per ogni coppia $t, u \in \mathbb{R}$ di numeri reali non entrambi nulli, definiamo la retta

$$r_{t,u} = \{t(a_0x + b_0y + c_0) + u(a_1x + b_1y + c_1) = 0\}.$$

In particolare $r_{1,0} = r_0$ e $r_{0,1} = r_1$. Se moltiplichiamo entrambi i parametri t e u per la stessa costante $\lambda \neq 0$ otteniamo in realtà la stessa retta:

$$r_{t,u} = r_{\lambda t, \lambda u} \quad \forall \lambda \neq 0.$$

Proposizione 9.2.29. *Le rette $r_{t,u}$ al variare di t, u sono precisamente tutte le rette passanti per P_0 .*

Dimostrazione. La retta $r_{t,u}$ contiene P_0 perché

$$t(a_0x_0 + b_0y_0 + c_0) + u(a_1x_0 + b_1y_0 + c_1) = t \cdot 0 + u \cdot 0 = 0.$$

D'altro canto, sia r una retta passante per P_0 . Sia Q un altro punto di r diverso da P_0 . Sostituendo Q nell'equazione che definisce $r_{t,u}$ si trovano dei parametri t e u che la risolvono e quindi una retta $r_{t,u}$ che passa per Q . Poiché per due punti P_0 e Q passa una sola retta, deve essere $r_{t,u} = r$. $\square$

Al variare di $t, u \in \mathbb{R}$, con la condizione che non siano entrambi nulli, le rette $r_{t,u}$ descrivono il fascio di rette passanti per P_0 .

Esempio 9.2.30. Il fascio di rette passanti per $P_0 = \begin{pmatrix} 2 \\ 1 \end{pmatrix}$ può essere costruito in questo modo. Ne prendiamo due, ad esempio quelle parallele agli assi:

$$\{x - 2 = 0\}, \quad \{y - 1 = 0\}$$

e si scrive

$$r_{t,u} = \{t(x - 2) + u(y - 1) = 0\} = \{tx + uy - 2t - u = 0\}.$$

È spesso semplice pescare nel fascio una retta che soddisfi delle condizioni richieste. Ad esempio, se cerchiamo la retta del fascio passante per $Q = \begin{pmatrix} 1 \\ 2 \end{pmatrix}$, sostituiamo $x = 1$, $y = 2$ per trovare $-t + u = 0$ con soluzione ad esempio $t = 1$, $u = 1$. Quindi

$$r_{1,1} = \{x + y - 3 = 0\}$$

è la retta che passa per P_0 e Q . Se invece cerchiamo la retta del fascio parallela alla retta $s = \{x + 2y = 18\}$, imponiamo che abbia gli stessi coefficienti (di x e y) di s , troviamo $t = 1$, $u = 2$ e e quindi otteniamo

$$r_{1,2} = \{x + 2y - 4 = 0\}.$$

Questa è la retta parallela a s passante per P_0 . Analogamente, se vogliamo la retta del fascio perpendicolare alla retta $s_0 = \{x - 3y = 2\}$ cerchiamo la retta del fascio con giacitura $\text{Span}\left(\begin{pmatrix} 1 \\ -3 \end{pmatrix}\right)$ e troviamo $r_{3,1}$. Questa è la retta perpendicolare a s_0 passante per P_0 .

Il *fascio di rette parallele* ad una retta data r è l'insieme di tutte le rette parallele a r . Questo è costruito molto semplicemente: data un'equazione cartesiana $\{ax + by = c\}$ di r , si sostituisce c con un parametro t .

Esempio 9.2.31. Il fascio di rette parallele alla retta $r = \{x + 2y = 7\}$ è dato da $\{x + 2y = t\}$.

Piani. Come sopra, il *fascio di piani contenenti una retta r* è l'insieme di tutti i piani contenenti r , ed è costruito prendendo due equazioni cartesiane di due piani qualsiasi contenenti r e facendo le loro combinazioni lineari in t e u . Il fascio di piani paralleli ad un piano con equazione $\{ax + by + cz = d\}$ si ottiene sostituendo d con un parametro t .

L'esempio seguente mostra come utilizzare agevolmente i fasci per risolvere alcuni problemi.

Esempio 9.2.32. Consideriamo la retta ed il punto

$$r = \begin{cases} x - y + 1 = 0, \\ 3x + y - z = 0, \end{cases} \quad P = \begin{pmatrix} 1 \\ 1 \\ 1 \end{pmatrix}.$$

Per trovare il piano π contenente r ed il punto P è sufficiente scrivere il fascio di piani contenenti r

$$\pi_{t,u} = \{t(x - y + 1) + u(3x + y - z) = 0\}$$

ed imporre al piano generico del fascio che contenga P . In questo modo otteniamo $t + 3u = 0$ e una soluzione è $t = 3, u = -1$. Il piano cercato è $\pi_{3,-1} = \{-4y + z + 3 = 0\}$.

9.2.9. Angoli. Abbiamo definito precedentemente una nozione di ortogonalità fra una retta e un sottospazio affine incidenti. Ci concentriamo adesso su $\mathbb{R}^3$ ed estendiamo questo concetto definendo l'angolo fra due sottospazi affini $S, S' \subset \mathbb{R}^3$ incidenti qualsiasi.

Ci sono alcuni casi da considerare separatamente.

Angolo tra rette. Siano r e r' due rette in $\mathbb{R}^3$ che si intersecano in un punto P come nella Figura 9.9. Abbiamo $r = P + \text{Span}(v)$ e $r' = P + \text{Span}(v')$ e definiamo l'angolo fra r e r' come l'angolo θ formato da v e v' . A seconda di come si scelgano i generatori v e v' gli angoli possibili sono due: θ e $\pi - \theta$.

Angolo tra retta e piano. Siano r una retta e π un piano che si intersecano in un punto P come nella Figura 9.10-(sinistra). Definiamo l'angolo θ fra r e π nel modo seguente.

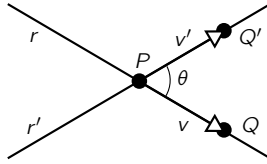


Figura 9.9. L'angolo fra due rette r e r' che si intersecano in un punto P è definito come l'angolo θ fra i vettori $v = Q - P$ e $v' = Q' - P'$. Gli angoli così ottenuti sono in realtà due, θ e $\pi - \theta$, a seconda delle posizioni di Q e Q' .

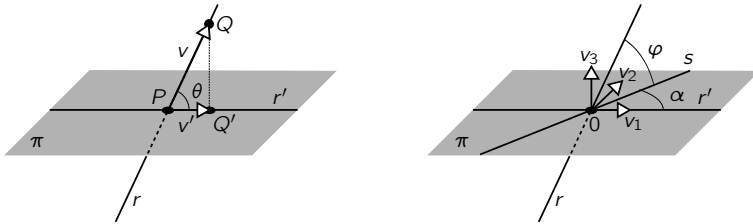


Figura 9.10. L'angolo θ fra una retta r ed un piano π che si intersecano in P è uguale all'angolo fra $v = Q - P$ e $v' = Q' - P'$ dove $Q \in r$ è un punto qualsiasi diverso da P e Q' è la sua proiezione ortogonale su π (sinistra). L'angolo θ fra r e r' è minore di quello φ fra r e s (destra).

Trasliamo l'origine in modo che $P = 0$. A questo punto π è un sottospazio vettoriale e definiamo il vettore $v' = p_\pi(v)$ facendo la proiezione ortogonale di v su π . Se $v' = 0$ poniamo $\theta = \frac{\pi}{2}$, altrimenti definiamo θ come l'angolo fra v e v' . Si veda la Figura 9.10-(sinistra).

In altre parole, se r non è ortogonale a π , proiettando r ortogonalmente su π otteniamo una retta $r' \subset \pi$ come nella Figura 9.10-(sinistra) e definiamo θ come l'angolo acuto fra r e r' .

A questo punto ci chiediamo: cosa succede se prendiamo un'altra retta $s \subset \pi$ passante per $P = 0$ al posto di r' come nella Figura 9.10-(destra)? L'angolo φ fra r e s è maggiore o minore di quello θ fra r e r' ? La lettrice è invitata a tentare di indovinare la risposta usando la propria intuizione geometrica, prima di leggere la proposizione seguente.

Proposizione 9.2.33. *Sia $s \subset \pi$ una retta passante per P diversa da r' . L'angolo φ fra s e r è maggiore di quello θ fra r' e r .*

Dimostrazione. A meno di traslare le coordinate, possiamo supporre che $P = 0$. Scegliamo una base ortonormale $\mathcal{B} = \{v_1, v_2, v_3\}$ come nella Figura 9.10-(destra), tale che $v_1 \in r'$ e $v_1, v_2 \in \pi$. Consideriamo adesso due vettori $v \in r$ e $u \in s$ unitari. Valgono le seguenti uguaglianze:

$$v = \cos \theta v_1 + \sin \theta v_3, \quad u = \cos \alpha v_1 + \sin \alpha v_2$$

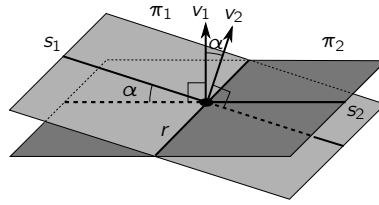


Figura 9.11. L'angolo fra due piani π_1 e π_2 che si intersecano in una retta r è definito come l'angolo α fra due rette $s_1 \subset \pi_1$ e $s_2 \subset \pi_2$ incidenti entrambe ortogonali a r . Questo è anche uguale all'angolo α formato da due vettori v_1, v_2 ortogonali ai piani.

dove $\alpha > 0$ è l'angolo tra r' e s . Otteniamo

$$\cos \varphi = \langle v, u \rangle = \cos \theta \cos \alpha < \cos \theta.$$

La funzione coseno è monotona decrescente: quindi $\varphi > \theta$. $\square$

Angolo tra piani. Siano π_1 e π_2 due piani che si intersecano in una retta $r = \pi_1 \cap \pi_2$. L'*angolo diedrale* fra π_1 e π_2 è definito nel modo seguente: si prendono due rette $s_1 \subset \pi_1$ e $s_2 \subset \pi_2$ incidenti ed entrambe ortogonali a r come nella Figura 9.11; l'angolo diedrale fra π_1 e π_2 è per definizione l'angolo α fra s_1 e s_2 . Si ottengono in realtà due angoli α e $\pi - \alpha$.

Esercizio 9.2.34. Siano v_1 e v_2 due vettori non nulli ortogonali a π_1 e π_2 come nella Figura 9.11. L'angolo α fra π_1 e π_2 è uguale a quello fra v_1 e v_2 .

Usando l'esercizio, è molto facile calcolare l'angolo diedrale fra due piani:

$$\pi_1 = \{a_1x + b_1y + c_1z = d_1\}, \quad \pi_2 = \{a_2x + b_2y + c_2z = d_2\}.$$

Questo è semplicemente l'angolo fra i vettori dei coefficienti.

Esempio 9.2.35. I piani $\pi_1 = \{x + y - z = 3\}$ e $\pi_2 = \{x - y - z = 9\}$ formano un angolo

$$\theta = \arccos \frac{\left\langle \begin{pmatrix} 1 \\ 1 \\ -1 \end{pmatrix}, \begin{pmatrix} 1 \\ -1 \\ -1 \end{pmatrix} \right\rangle}{\left\| \begin{pmatrix} 1 \\ 1 \\ -1 \end{pmatrix} \right\| \left\| \begin{pmatrix} 1 \\ -1 \\ -1 \end{pmatrix} \right\|} = \arccos \frac{1}{\sqrt{3}\sqrt{3}} = \arccos \frac{1}{3}.$$

9.2.10. Distanze. Definiamo infine la *distanza* $d(S, S')$ fra due sottospazi affini S e S' di $\mathbb{R}^3$. Se sono incidenti, la distanza è zero. Se i sottospazi sono disgiunti la distanza è positiva ed è definita in modo diverso caso per caso.

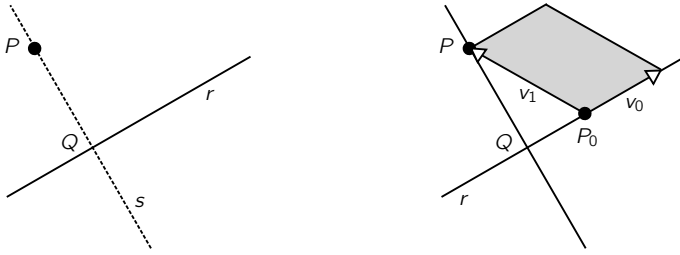


Figura 9.12. La distanza fra un punto P e una retta r è la lunghezza del segmento PQ ortogonale a r .

Distanza fra punti. Come già sappiamo, la distanza $d(P, Q)$ fra due punti $P, Q \in \mathbb{R}^3$ è definita usando la norma:

$$d(P, Q) = \|\overrightarrow{PQ}\| = \|Q - P\|.$$

Distanza fra punto e retta. La distanza $d(P, r)$ fra un punto P ed una retta r nello spazio è definita nel modo seguente. Tracciamo come nella Figura 9.12 la perpendicolare s a r nel punto P e definiamo

$$d(P, r) = d(P, Q)$$

dove $Q = r \cap s$. Se la retta r è espressa in forma parametrica $r = \{P_0 + tv_0\}$ come nella Figura 9.12-(destra), la distanza si calcola facilmente usando il prodotto vettoriale.

Proposizione 9.2.36. *Vale l'uguaglianza*

$$d(P, r) = \frac{\|v_0 \times v_1\|}{\|v_0\|}$$

dove $v_1 = \overrightarrow{P_0P} = P - P_0$.

Dimostrazione. L'area del parallelogramma mostrato nella Figura 9.12-(destra) può essere calcolata in due modi: con il prodotto vettoriale $\|v_0 \times v_1\|$ e come base per altezza $\|v_0\| \cdot d(P, Q)$. L'uguaglianza

$$\|v_0 \times v_1\| = \|v_0\| \cdot d(P, Q)$$

implica la formula per $d(P, r) = d(P, Q)$. □

Questa formula è molto conveniente: non abbiamo bisogno di determinare il punto Q per calcolare la distanza fra P e r .

Esempio 9.2.37. Consideriamo il punto P e la retta r seguenti:

$$P = \begin{pmatrix} 1 \\ -2 \\ 3 \end{pmatrix}, \quad r = \left\{ \begin{pmatrix} -1 \\ 0 \\ 1 \end{pmatrix} + t \begin{pmatrix} 3 \\ 2 \\ 1 \end{pmatrix} \right\}$$

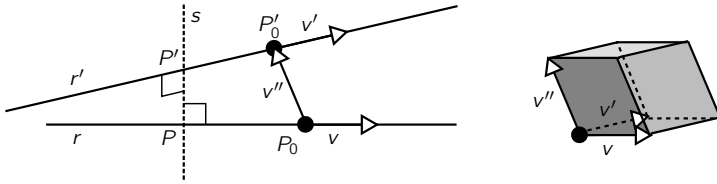


Figura 9.13. La distanza fra due rette sghembe è la lunghezza del segmento PP' ortogonale ad entrambe.

La distanza fra P e r è

$$d(P, r) = \frac{\left\| \begin{pmatrix} 3 \\ 2 \\ 1 \end{pmatrix} \times \begin{pmatrix} 2 \\ -2 \\ 2 \end{pmatrix} \right\|}{\left\| \begin{pmatrix} 3 \\ 2 \\ 1 \end{pmatrix} \right\|} = \frac{\left\| \begin{pmatrix} 6 \\ -4 \\ -10 \end{pmatrix} \right\|}{\sqrt{14}} = \frac{2\sqrt{38}}{\sqrt{14}} = \frac{2}{7}\sqrt{133}.$$

Esercizio 9.2.38. La distanza fra un punto $P = \begin{pmatrix} x_0 \\ y_0 \end{pmatrix}$ e una retta r in $\mathbb{R}^2$ di equazione $ax + by + c = 0$ è

$$d(P, r) = \frac{|ax_0 + by_0 + c|}{\sqrt{a^2 + b^2}}.$$

Dimostreremo una formula analoga nella Proposizione 9.2.42.

Distanza fra rette. La distanza $d(r, r')$ fra due rette r ed r' disgiunte nel piano o nello spazio è definita nel modo seguente. Le due rette r e r' sono parallele o sghembe. In entrambi i casi troviamo una retta s perpendicolare ad entrambe (questa è unica se sono sghembe, non lo è se sono parallele). Definiamo quindi

$$d(r, r') = d(P, P')$$

dove $P = r \cap s$ e $P' = r' \cap s$. Si veda la Figura 9.13-(sinistra) nel caso in cui r e r' siano sghembe.

Anche in questo caso, se $r = \{P_0 + tv\}$ e $r' = \{P'_0 + uv'\}$ sono espresse in forma parametrica la distanza si calcola facilmente. Se sono parallele, la distanza $d(r, r')$ è uguale a $d(P_0, r')$ ed applichiamo la Proposizione 9.2.36. Consideriamo il caso in cui le rette siano sghembe: si veda la Figura 9.13.

Proposizione 9.2.39. Se r e r' sono sghembe, vale la formula

$$d(r, r') = \frac{|\det(v|v'|v'')|}{\|v \times v'\|}$$

dove $v'' = \overrightarrow{P_0P'_0} = P'_0 - P_0$.

Dimostrazione. Ragioniamo come per la Proposizione 9.2.36. La dimostrazione è simile alla precedente. Il volume del parallelepipedo nella

Figura 9.13-(destra) generato dai vettori v , v' e v'' può essere calcolato in due modi: come $|\det(v|v'|v'')|$ e come area di base per altezza, cioè $\|v \times v'\| \cdot d$. L'altezza d è proprio $d = d(P, P') = d(r, r')$ e quindi si ottiene

$$|\det(v|v'|v'')| = d(r, r') \cdot \|v \times v'\|$$

che implica la formula per $d(r, r')$. $\square$

Notiamo anche qui l'efficacia di una formula che consente di calcolare la distanza fra due rette con uno sforzo minimo: non è necessario determinare la retta s , né i punti P e P' .

Esempio 9.2.40. Calcoliamo la distanza fra le rette sghembe

$$r = \left\{ \begin{pmatrix} 2 \\ -5 \\ 1 \end{pmatrix} + t \begin{pmatrix} 2 \\ 0 \\ 1 \end{pmatrix} \right\}, \quad r' = \left\{ \begin{pmatrix} 1 \\ 1 \\ 0 \end{pmatrix} + u \begin{pmatrix} -1 \\ -2 \\ 3 \end{pmatrix} \right\}.$$

Questa è

$$d(r, r') = \frac{\left| \det \begin{pmatrix} 2 & -1 & -1 \\ 0 & -2 & 6 \\ 1 & 3 & -1 \end{pmatrix} \right|}{\left\| \begin{pmatrix} 2 \\ 0 \\ 1 \end{pmatrix} \times \begin{pmatrix} -1 \\ -2 \\ 3 \end{pmatrix} \right\|} = \frac{|2(2-18) + (-6-2)|}{\sqrt{2^2 + (-7)^2 + (-4)^2}} = \frac{40}{\sqrt{69}}.$$

Con un ragionamento analogo troviamo anche un semplice algoritmo che ci permette di capire se due rette date in forma parametrica siano complanari. Siano $r = \{P_0 + tv\}$ e $r' = \{P'_0 + uv'\}$ due rette qualsiasi nello spazio.

Proposizione 9.2.41. *Le rette r e r' sono complanari se e solo se*

$$\det(v|v'|v'') = 0$$

dove $v'' = \overrightarrow{P_0P'_0} = P'_0 - P_0$.

Dimostrazione. Se sono complanari, i vettori v , v' e v'' giacciono sullo stesso piano e quindi $\det(v|v'|v'') = 0$. Se non sono complanari, i vettori v , v' e v'' sono indipendenti e quindi $\det(v|v'|v'') \neq 0$. $\square$

Distanza fra punto e piano. Infine, la distanza fra un punto P_0 ed un piano π nello spazio è definita in modo simile a quanto già visto: si traccia la perpendicolare s a π passante per P_0 e si definisce

$$d(P_0, \pi) = d(P_0, Q)$$

con $Q = s \cap \pi$, si veda la Figura 9.14. Se

$$\pi = \{ax + by + cz = d\}, \quad P_0 = \begin{pmatrix} x_0 \\ y_0 \\ z_0 \end{pmatrix}$$

allora la distanza si calcola facilmente con una formula.

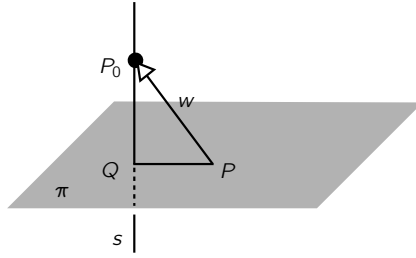


Figura 9.14. La distanza fra un punto P_0 e un piano π è la lunghezza del segmento P_0Q ortogonale a π .

Proposizione 9.2.42. *Vale*

$$d(P_0, \pi) = \frac{|ax_0 + by_0 + cz_0 - d|}{\sqrt{a^2 + b^2 + c^2}}$$

Dimostrazione. Consideriamo

$$v = \begin{pmatrix} a \\ b \\ c \end{pmatrix}, \quad P = \begin{pmatrix} x \\ y \\ z \end{pmatrix}$$

con $P \in \pi$ punto qualsiasi, come nella Figura 9.14. Il vettore v è ortogonale a π . Sia $w = P_0 - P$. La figura mostra che

$$\begin{aligned} d(P_0, Q) &= \|p_v(w)\| = \frac{|\langle v, w \rangle|}{|v|} = \frac{|a(x_0 - x) + b(y_0 - y) + c(z_0 - z)|}{\sqrt{a^2 + b^2 + c^2}} \\ &= \frac{|ax_0 + by_0 + cz_0 - d|}{\sqrt{a^2 + b^2 + c^2}}. \end{aligned}$$

Nell'ultima uguaglianza abbiamo usato che $P \in \pi$ e quindi $ax + by + cz = d$. $\square$

Esempio 9.2.43. Consideriamo il piano ed il punto seguenti in $\mathbb{R}^3$:

$$\pi = \{2x - y + z = 4\}, \quad P_0 = \begin{pmatrix} 2 \\ 1 \\ -2 \end{pmatrix}.$$

Applicando la formula troviamo

$$d(P_0, \pi) = \frac{|2 \cdot 2 - 1 - 2 - 4|}{\sqrt{6}} = \frac{\sqrt{6}}{2}.$$

9.3. Affinità

Tutti gli endomorfismi di $\mathbb{R}^n$ per definizione devono fissare l'origine. Vogliamo adesso considerare alcune trasformazioni di $\mathbb{R}^n$ che spostano l'origine in un punto arbitrario. Queste si chiamano *affinità*.

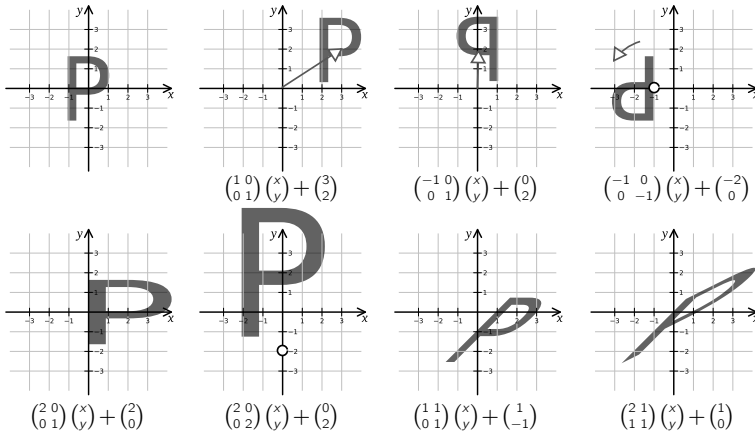


Figura 9.15. Sette trasformazioni affini di $\mathbb{R}^2$, che trasformano la lettera P in alto a sinistra come mostrato. La prima è una traslazione di vettore $(3, 2)$, la seconda una glissoriflessione lungo l'asse y di passo 2, la terza una rotazione di angolo π intorno al punto $(-1, 0)$, la quinta una omotetia di fattore 2 e centro $(0, -2)$.

9.3.1. Definizione. Un'affinità di $\mathbb{R}^n$ è una mappa $f: \mathbb{R}^n \rightarrow \mathbb{R}^n$ del tipo

$$f(x) = Ax + b$$

dove A è una matrice $n \times n$ e $b \in \mathbb{R}^n$ un vettore. Chiamiamo la matrice A la *componente lineare* di f ed il vettore b la *componente traslatoria*. Notiamo che la componente traslatoria è l'immagine del vettore nullo:

$$f(0) = A0 + b = b.$$

Due esempi fondamentali di affinità sono gli endomorfismi

$$f(x) = Ax$$

in cui $b = 0$, e le *traslazioni*

$$f(x) = x + b$$

in cui $A = I_n$. Qualsiasi affinità $f(x) = Ax + b$ è la composizione $f = h \circ g$ di un endomorfismo $g(x) = Ax$ e di una traslazione $h(y) = y + b$.

9.3.2. Isomorfismo affine. Una affinità $f(x) = Ax + b$ è un *isomorfismo affine* (o *trasformazione affine*) se è bigettiva. Alcuni isomorfismi affini di $\mathbb{R}^2$ sono mostrati nella Figura 9.15.

Proposizione 9.3.1. Una affinità $f(x) = Ax + b$ è un *isomorfismo affine* $\iff A$ è invertibile. In questo caso l'inversa è l'*isomorfismo affine*

$$g(y) = A^{-1}y - A^{-1}b.$$

Dimostrazione. Se A non è invertibile, per Rouché-Capelli il sistema $Ax + b = 0$ non può avere una unica soluzione e quindi f non è una bigezione.

Se A è invertibile, moltiplicando a sinistra entrambi i membri dell'equazione $y = Ax + b$ per A^{-1} troviamo $A^{-1}y = x + A^{-1}b$ e quindi

$$x = A^{-1}y - A^{-1}b.$$

Ne segue che f è una bigezione e la sua inversa è l'isomorfismo affine g descritto nell'enunciato. $\square$

Vogliamo adesso esaminare l'effetto di un isomorfismo affine su di un sottospazio. Siano $f(x) = Ax + b$ un isomorfismo affine e $S = P + W$ un sottospazio di $\mathbb{R}^n$. L'isomorfismo affine sposta il sottospazio S in un certo insieme $f(S)$, l'immagine di S secondo f .

Proposizione 9.3.2. *L'immagine $f(S)$ è il sottospazio affine*

$$f(S) = f(P) + L_A(W).$$

In particolare $f(S)$ ha la stessa dimensione di S .

Dimostrazione. Abbiamo

$$S = P + W = \{P + w \mid w \in W\}$$

e quindi

$$\begin{aligned} f(S) &= \{A(P + w) + b \mid w \in W\} = \{AP + b + Aw \mid w \in W\} \\ &= \{f(P) + L_A(w) \mid w \in W\} = f(P) + L_A(W) \end{aligned}$$

Siccome A è invertibile abbiamo $\dim L_A(W) = \dim W$ per l'Esercizio 4.5. $\square$

Abbiamo scoperto che un isomorfismo affine trasforma un sottospazio affine S in un sottospazio affine $f(S)$ della stessa dimensione, la cui giacitura $L_A(W)$ è ottenuta applicando la parte lineare L_A alla giacitura W di S .

In particolare la giacitura $L_A(W)$ dell'immagine di S dipende solo dalla componente lineare A di f e non dalla parte traslatoria; questo fatto non è sorprendente: traslando un sottospazio affine non si cambia la sua giacitura.

Notiamo adesso che gli isomorfismi affini preservano il parallelismo:

Proposizione 9.3.3. *Un isomorfismo affine manda sottospazi paralleli in sottospazi paralleli.*

Dimostrazione. Sia $f(x) = Ax + b$ un isomorfismo affine. Se $S = P + W$ e $S' = P' + W'$ sono sottospazi paralleli con $W \subset W'$, allora $f(S) = f(P) + L_A(W)$ e $f(S') = f(P') + L_A(W')$ sono paralleli con $L_A(W) \subset L_A(W')$. $\square$

Gli isomorfismi affini preservano il parallelismo, ma non necessariamente le distanze fra punti né gli angoli fra sottospazi! Questa non è una novità: sappiamo già che gli endomorfismi lineari L_A in generale non preservano né norma, né distanza, né angoli. Vedremo degli esempi successivamente.

Osservazione 9.3.4. Gli isomorfismi affini formano un gruppo con la composizione: la composizione di due isomorfismi è un isomorfismo e ogni isomorfismo ha un inverso. L'elemento neutro è l'identità $f(x) = x$.

Date due basi di uno spazio vettoriale V , esiste un unico isomorfismo lineare di V che manda la prima base nella seconda. Dimostriamo un risultato analogo per gli isomorfismi affini.

Proposizione 9.3.5. *Siano $P_0, \dots, P_n$ e $Q_0, \dots, Q_n$ due insiemi di $n+1$ punti affinementemente indipendenti in $\mathbb{R}^n$. Esiste un unico isomorfismo affine $f: \mathbb{R}^n \rightarrow \mathbb{R}^n$ tale che $f(P_i) = Q_i$ per ogni i .*

Dimostrazione. Ricordiamo dalla Sezione 9.2.4 che i punti $P_0, \dots, P_n$ sono affinementemente indipendenti precisamente quando i vettori $\overrightarrow{P_0P_1}, \dots, \overrightarrow{P_0P_n}$ sono indipendenti. Questi n vettori sono quindi una base di $\mathbb{R}^n$. Sia $L_A: \mathbb{R}^n \rightarrow \mathbb{R}^n$ l'unico endomorfismo tale che

$$L_A(\overrightarrow{P_0P_i}) = \overrightarrow{Q_0Q_i} \quad \forall i.$$

Questo è un isomorfismo di $\mathbb{R}^n$ perché $Q_0, \dots, Q_n$ sono affinementemente indipendenti e quindi i vettori immagine $\overrightarrow{Q_0Q_1}, \dots, \overrightarrow{Q_0Q_n}$ sono indipendenti. Quindi A è invertibile. Definiamo inoltre

$$b = Q_0 - AP_0.$$

L'isomorfismo affine $f(x) = Ax + b$ manda P_i in Q_i per ogni i , infatti:

$$f(P_i) = f(P_0 + \overrightarrow{P_0P_i}) = AP_0 + A\overrightarrow{P_0P_i} + b = Q_0 + \overrightarrow{Q_0Q_i} = Q_i.$$

Le scelte di A e b sono state forzate da queste condizioni, quindi questo è l'unico isomorfismo affine possibile. $\square$

Esercizio 9.3.6. Determina l'isomorfismo affine di $\mathbb{R}^2$ che manda i punti $\begin{pmatrix} 1 \\ 1 \end{pmatrix}$, $\begin{pmatrix} -2 \\ 0 \end{pmatrix}$, $\begin{pmatrix} 1 \\ -1 \end{pmatrix}$ rispettivamente in $\begin{pmatrix} 0 \\ 1 \end{pmatrix}$, $\begin{pmatrix} 2 \\ 2 \end{pmatrix}$, $\begin{pmatrix} -1 \\ 0 \end{pmatrix}$.

9.3.3. Traslare l'origine. Quando si studiano le affinità è a volte conveniente traslare l'origine in un punto P , ottenendo quindi nuove coordinate tramite la formula

$$x' = x - P.$$

Se scriviamo una affinità nelle nuove coordinate x' come

$$f(x') = Ax' + b,$$

sostituendo $f(x') = y' = y - P$ e $x' = x - P$ troviamo $y - P = A(x - P) + b$ con $y = f(x)$. Quindi la funzione f nelle coordinate iniziali è la seguente:

$$(15) \quad f(x) = A(x - P) + b + P.$$

Possiamo anche scrivere l'affinità f in altri modi:

$$f(x) = Ax - AP + b + P = Ax + (I_n - A)P + b.$$

Notiamo un fatto rilevante: la parte lineare A della affinità è sempre la stessa, mentre la componente traslatoria è cambiata da b a $(I_n - A)P + b$.

Nei prossimi esempi usiamo questo procedimento per definire alcune affinità molto importanti.

Esempio 9.3.7 (Omotetia di centro P). L'*omotetia* di centro l'origine e di fattore $\lambda \neq 0$ è l'isomorfismo affine $f(x) = \lambda x$. La componente lineare è λI_n e quella traslatoria è nulla.

L'*omotetia* di centro un punto arbitrario $P \in \mathbb{R}^n$ è l'affinità ottenuta cambiando coordinate e mettendo l'origine in P , cioè con $x' = x - P$, e prendendo $f(x') = \lambda x'$ nelle nuove coordinate. Nelle vecchie coordinate, l'omotetia di centro P si scrive come

$$f(x) = \lambda(x - P) + P = \lambda x + (1 - \lambda)P.$$

Ad esempio, l'omotetia di centro $\begin{pmatrix} 1 \\ 2 \end{pmatrix} \in \mathbb{R}^2$ e di fattore $\lambda = 2$ è

$$f \begin{pmatrix} x \\ y \end{pmatrix} = 2 \begin{pmatrix} x \\ y \end{pmatrix} - \begin{pmatrix} 1 \\ 2 \end{pmatrix}.$$

Come riprova, si può verificare che effettivamente $f\begin{pmatrix} 1 \\ 2 \end{pmatrix} = \begin{pmatrix} 1 \\ 2 \end{pmatrix}$.

La quinta trasformazione mostrata nella Figura 9.15 è una omotetia di centro $\begin{pmatrix} 0 \\ -2 \end{pmatrix}$ e di fattore 2.

Esempio 9.3.8 (Rotazione di centro P). Ricordiamo che la rotazione di angolo θ intorno all'origine è la mappa L_A con $A = \text{Rot}_\theta$. Come per le omotetie, definiamo la rotazione intorno ad un punto P qualsiasi scrivendo

$$f(x) = \text{Rot}_\theta(x - P) + P = \text{Rot}_\theta x + (I_2 - \text{Rot}_\theta)P.$$

Esplicitamente, se $P = \begin{pmatrix} x_0 \\ y_0 \end{pmatrix}$ otteniamo

$$f \begin{pmatrix} x \\ y \end{pmatrix} = \begin{pmatrix} \cos \theta & -\sin \theta \\ \sin \theta & \cos \theta \end{pmatrix} \begin{pmatrix} x \\ y \end{pmatrix} + \begin{pmatrix} (1 - \cos \theta)x_0 + \sin \theta y_0 \\ -\sin \theta x_0 + (1 - \cos \theta)y_0 \end{pmatrix}.$$

Come controllo, si può verificare che $f(P) = P$. Notiamo che la parte lineare di una rotazione di angolo θ intorno ad un certo punto P è sempre Rot_θ , mentre la parte traslatoria è più complicata da calcolare.

Ad esempio, la rotazione di angolo π intorno a $\begin{pmatrix} -1 \\ 0 \end{pmatrix}$ è

$$f \begin{pmatrix} x \\ y \end{pmatrix} = \begin{pmatrix} -1 & 0 \\ 0 & -1 \end{pmatrix} \begin{pmatrix} x \\ y \end{pmatrix} + \begin{pmatrix} -2 \\ 0 \end{pmatrix}.$$

Questa è la terza l'affinità mostrata nella Figura 9.15. Come riprova, ricordiamo che in qualsiasi affinità la parte traslatoria $b = \begin{pmatrix} -2 \\ 0 \end{pmatrix}$ è l'immagine $b = f(0)$ dell'origine, ed effettivamente torna geometricamente che una rotazione di angolo π intorno a $\begin{pmatrix} -1 \\ 0 \end{pmatrix}$ mandi l'origine in $\begin{pmatrix} -2 \\ 0 \end{pmatrix}$.

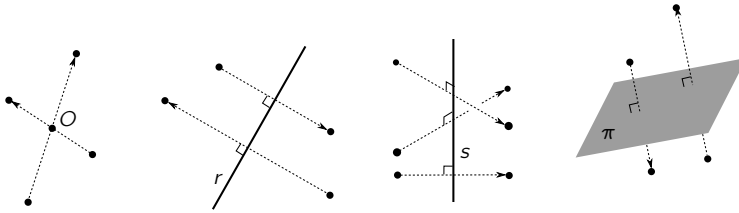


Figura 9.16. Una riflessione ortogonale rispetto ad un punto O , ad una retta r nel piano, ad una retta s nello spazio, e ad un piano π nello spazio.

Esempio 9.3.9 (Riflessione ortogonale). Sia $S \subset \mathbb{R}^n$ un sottospazio affine. La *riflessione ortogonale* rispetto a S è la mappa $r_S: \mathbb{R}^n \rightarrow \mathbb{R}^n$ definita nel modo seguente. Per ogni $P \in \mathbb{R}^n$, per la Proposizione 9.2.19 esiste un'unica retta r ortogonale a S passante per P . Sia $Q = r \cap S$. Definiamo

$$r_S(P) = P + 2\overrightarrow{PQ}.$$

I casi che ci interessano di più sono quelli in $\mathbb{R}^2$ e $\mathbb{R}^3$, in cui S è un punto, una retta o un piano: si veda la Figura 9.16.

Per scrivere r_S è sufficiente scegliere un punto $P \in S$ e spostare l'origine in P con $x' = x - P$. Con queste coordinate S è un sottospazio vettoriale e r_S è l'usuale riflessione ortogonale già vista nell'Esempio 7.5.3.

Ad esempio, determiniamo la riflessione ortogonale rispetto al piano affine

$$\pi = \{x + y - z = 3\}.$$

A questo scopo scegliamo un punto in π , ad esempio $P = 3e_1$. Spostiamo l'origine in P scrivendo le nuove coordinate

$$x' = x - 3, \quad y' = y, \quad z' = z.$$

Adesso il piano π ha equazione $\pi = \{x' + y' - z' = 0\}$. Come abbiamo visto nella Sezione 8.2.3, la riflessione lungo π si scrive come

$$r_\pi \begin{pmatrix} x' \\ y' \\ z' \end{pmatrix} = \begin{pmatrix} x' \\ y' \\ z' \end{pmatrix} - 2 \frac{x' + y' - z'}{3} \begin{pmatrix} 1 \\ 1 \\ -1 \end{pmatrix} = \frac{1}{3} \begin{pmatrix} 1 & -2 & 2 \\ -2 & 1 & 2 \\ 2 & 2 & 1 \end{pmatrix} \begin{pmatrix} x' \\ y' \\ z' \end{pmatrix}.$$

Usiamo (15) per scrivere la riflessione nelle coordinate originarie:

$$\begin{aligned} r_\pi \begin{pmatrix} x \\ y \\ z \end{pmatrix} &= \frac{1}{3} \begin{pmatrix} 1 & -2 & 2 \\ -2 & 1 & 2 \\ 2 & 2 & 1 \end{pmatrix} \begin{pmatrix} x \\ y \\ z \end{pmatrix} + \frac{1}{3} \begin{pmatrix} 2 & 2 & -2 \\ 2 & 2 & -2 \\ -2 & -2 & 2 \end{pmatrix} \begin{pmatrix} 3 \\ 0 \\ 0 \end{pmatrix} \\ &= \frac{1}{3} \begin{pmatrix} 1 & -2 & 2 \\ -2 & 1 & 2 \\ 2 & 2 & 1 \end{pmatrix} \begin{pmatrix} x \\ y \\ z \end{pmatrix} + \begin{pmatrix} 2 \\ 2 \\ -2 \end{pmatrix}. \end{aligned}$$

Come verifica, possiamo notare che il punto P è fissato da r_π .

Esercizio 9.3.10. Scrivi nella forma $f(x) = Ax + b$ la riflessione di $\mathbb{R}^2$ rispetto alla retta affine $x + y = 1$. Riesci a indovinare chi è b geometricamente prima di fare i conti, sapendo che $b = f(0)$?

Osservazione 9.3.11. Una riflessione rispetto ad un punto $P \in \mathbb{R}^n$ è anche una omotetia intorno a P con fattore -1 . In dimensione $n = 2$, questa è anche una rotazione di angolo π intorno a P .

Esempio 9.3.12 (Rotazione attorno ad un asse). Abbiamo definito nella Sezione 8.2.5 la rotazione intorno ad un asse r di angolo θ , dove r è una qualsiasi retta vettoriale in $\mathbb{R}^3$. Estendiamo questa definizione al caso in cui r sia una retta affine con lo stesso procedimento visto negli esempi precedenti: prendiamo un punto $P \in r$ qualsiasi, cambiamo coordinate con $x' = x - P$ per mettere l'origine in P , e quindi definiamo la rotazione con queste nuove coordinate.

Ad esempio, scriviamo la rotazione di angolo $\frac{2\pi}{3}$ intorno all'asse

$$r = \left\{ \begin{pmatrix} 1 \\ 2 \\ 0 \end{pmatrix} + t \begin{pmatrix} 1 \\ 1 \\ 1 \end{pmatrix} \right\}.$$

Scegliamo $P \in r$, ad esempio $P = e_1 + 2e_2$. Spostiamo l'origine in P scrivendo

$$x' = x - 1, \quad y' = y - 2, \quad z' = z.$$

In queste coordinate r è una retta vettoriale. Abbiamo già visto nell'Esempio 8.2.9 che la rotazione di angolo $\frac{2\pi}{3}$ intorno ad r è l'endomorfismo

$$r' \begin{pmatrix} x' \\ y' \\ z' \end{pmatrix} = \begin{pmatrix} 0 & 0 & 1 \\ 1 & 0 & 0 \\ 0 & 1 & 0 \end{pmatrix} \begin{pmatrix} x' \\ y' \\ z' \end{pmatrix}.$$

Nelle coordinate originarie questa diventa

$$\begin{aligned} r \begin{pmatrix} x \\ y \\ z \end{pmatrix} &= \begin{pmatrix} 0 & 0 & 1 \\ 1 & 0 & 0 \\ 0 & 1 & 0 \end{pmatrix} \begin{pmatrix} x \\ y \\ z \end{pmatrix} + \begin{pmatrix} 1 & 0 & -1 \\ -1 & 1 & 0 \\ 0 & -1 & 1 \end{pmatrix} \begin{pmatrix} 1 \\ 2 \\ 0 \end{pmatrix} \\ &= \begin{pmatrix} 0 & 0 & 1 \\ 1 & 0 & 0 \\ 0 & 1 & 0 \end{pmatrix} \begin{pmatrix} x \\ y \\ z \end{pmatrix} + \begin{pmatrix} 1 \\ 1 \\ -2 \end{pmatrix}. \end{aligned}$$

Esempio 9.3.13 (Antirotazione). Dati una retta orientata r e un piano π perpendicolari fra loro, e dato un angolo θ , si definisce l'*antirotazione* intorno a r di angolo θ con piano invariante π la composizione di una rotazione di angolo θ intorno ad r con una riflessione lungo π . Spostando l'origine nel punto $P = r \cap \pi$, questa può essere scritta come affinità in modo simile a quanto visto negli esempi precedenti. Si veda la Sezione 8.2.6.

9.3.4. Isometrie affini. Passiamo adesso a considerare gli isomorfismi affini che ci interessano di più: quelli che preservano le distanze.

Una *isometria affine* è un isomorfismo affine $f(x) = Ax + b$ che preserva le distanze, cioè tale che

$$d(f(P), f(Q)) = d(P, Q)$$

per ogni coppia di punti $P, Q \in \mathbb{R}^n$.

Esempio 9.3.14. Una traslazione $f(x) = x + b$ è una isometria affine:

$$d(f(P), f(Q)) = d(P+b, Q+b) = \|(P+b) - (Q+b)\| = \|P - Q\| = d(P, Q).$$

Descriviamo subito un semplice criterio per capire se una data affinità sia un'isometria affine.

Proposizione 9.3.15. *Una affinità $f(x) = Ax + b$ è una isometria affine $\iff$ la matrice A è ortogonale.*

Dimostrazione. ($\Leftarrow$) L'affinità $g(x) = Ax$ è una isometria per il Corollario 7.5.10. Anche la traslazione $h(y) = y + b$ è una isometria. Componendo due isometrie si ottiene una isometria, quindi anche la composizione $f = h \circ g$ è una isometria.

($\Rightarrow$) Se $f(x) = Ax + b$ è una isometria, componendola con la traslazione $g(y) = y - b$ si ottiene l'isometria $h(x) = g(f(x)) = Ax$. Siccome h è una isometria, per la Proposizione 8.2.1 la matrice A è ortogonale. $\square$

Molti degli isomorfismi affini descritti in precedenza sono isometrie: oltre alle traslazioni, troviamo le riflessioni ortogonali rispetto ad un sotto-spazio affine, le rotazioni (intorno ad un punto in $\mathbb{R}^2$ o ad un asse in $\mathbb{R}^3$) e le antirotazioni in $\mathbb{R}^3$. Vedremo fra poco altri esempi.

Una isometria affine $f(x) = Ax + b$ ha sempre $\det A = \pm 1$. Se $\det A = 1$, diciamo che l'isometria affine *preserva* l'orientazione dello spazio, mentre se $\det A = -1$ diciamo che la *inverte*. Come per gli endomorfismi, le isometrie affini di $\mathbb{R}^2$ e $\mathbb{R}^3$ che preservano l'orientazione hanno un comportamento più familiare, mentre quelle che la invertono specchiano gli oggetti, trasformando in particolare una mano destra in una mano sinistra e viceversa. Si veda la Sezione 4.4.6.

Osservazione 9.3.16. Le isometrie affini formano un gruppo con la composizione. Infatti l'inversa di una isometria è una isometria, e componendo due isometrie si ottiene una isometria.

9.3.5. Punti fissi. Nello studio delle affinità (ed in particolare delle isometrie affini) giocano un ruolo importante i *punti fissi*. Questi sono i punti che non vengono spostati.

Definizione 9.3.17. Un *punto fisso* per una affinità $f(x) = Ax + b$ è un punto $P \in \mathbb{R}^n$ tale che $f(P) = P$.

Esempio 9.3.18. Se $f: \mathbb{R}^2 \rightarrow \mathbb{R}^2$ è una rotazione intorno a $P \in \mathbb{R}^2$, allora P è un punto fisso. D'altra parte, una traslazione $f(x) = x + b$ con $b \neq 0$ non ha nessun punto fisso.

Se f ha un punto fisso P , è naturale prendere le coordinate $x' = x - P$, perché in queste nuove coordinate il punto fisso P è la nuova origine 0 e quindi l'affinità f letta nelle nuove coordinate diventa un endomorfismo:

$$f(x') = Ax'.$$

Gli endomorfismi sono più facili da studiare delle affinità perché non compare il termine traslatorio b e possiamo usare molti degli strumenti introdotti nei capitoli precedenti come autovettori, polinomio caratteristico, ecc. Per questo motivo è spesso utile capire subito se una affinità ha dei punti fissi, ed eventualmente trovarli. La proposizione seguente fornisce un criterio potente.

Proposizione 9.3.19. *Se 1 non è autovalore per A , l'affinità $f(x) = Ax + b$ ha esattamente un punto fisso.*

Dimostrazione. Un punto $P \in \mathbb{R}^n$ è fisso se e solo se

$$AP + b = P \iff (A - I_n)P = -b.$$

Se 1 non è autovalore per A , la matrice $A - I_n$ è invertibile. Quindi esiste un unico P per cui $(A - I_n)P = -b$. $\square$

Se 1 è autovalore per A , può capitare che non ci siano punti fissi (ad esempio, nelle traslazioni) o che ce ne siano (ad esempio, nelle riflessioni ortogonali). Nel caso in cui ci siano dei punti fissi, sappiamo che sono un sottospazio affine e conosciamo la loro dimensione.

Proposizione 9.3.20. *Sia $f(x) = Ax + b$ una affinità. Se ci sono punti fissi, questi formano un sottospazio affine di dimensione $\dim \ker(A - I_n)$.*

Dimostrazione. Abbiamo visto nella dimostrazione precedente che i punti fissi sono le soluzioni del sistema lineare $(A - I_n)x = -b$. Per il Teorema di Rouché - Capelli, se ci sono soluzioni queste formano un sottospazio affine di dimensione $n - \text{rk}(A - I_n) = \dim \ker(A - I_n)$. $\square$

L'insieme dei punti fissi di f è indicato con $\text{Fix}(f)$. In breve:

- Se 1 non è autovalore per A , allora $\text{Fix}(f)$ è un punto.
- Se 1 è autovalore per A , allora $\text{Fix}(f)$ è vuoto oppure un sottospazio affine di dimensione $m_g(1) = \dim \ker(A - I_n)$.

9.3.6. Isometrie affini del piano che preservano l'orientazione. Vogliamo adesso classificare completamente tutte le isometrie affini di $\mathbb{R}^2$ e $\mathbb{R}^3$. Iniziamo con il piano $\mathbb{R}^2$, e ci limitiamo a considerare le isometrie affini che preservano l'orientazione.

Una isometria affine di $\mathbb{R}^2$ che preserva l'orientazione è del tipo

$$f(x) = Ax + b$$

dove A è ortogonale con $\det A = 1$. Per la Proposizione 8.2.7 abbiamo $A = \text{Rot}_\theta$ per qualche $\theta \in [0, 2\pi)$ e quindi

$$f(x) = \text{Rot}_\theta x + b.$$

Proposizione 9.3.21. *Ci sono tre casi possibili:*

- (1) Se $\theta = 0$ e $b = 0$, allora f è l'identità.
- (2) Se $\theta = 0$ e $b \neq 0$, allora f è una traslazione.
- (3) Se $\theta \neq 0$, allora f è una rotazione intorno al punto

$$P = -(\text{Rot}_\theta - I_2)^{-1}b = -\frac{1}{2(1 - \cos \theta)}(\text{Rot}_{-\theta} - I_2)b.$$

Dimostrazione. Se $\theta = 0$ chiaramente f è l'identità o una traslazione. Se $\theta \neq 0$, la matrice Rot_θ non ha autovalore 1 e quindi f ha un punto fisso P , soluzione del sistema $(\text{Rot}_\theta - I_2)P = -b$ e quindi $P = -(\text{Rot}_\theta - I_2)^{-1}b$. Svolgendo i conti si trova che $(\text{Rot}_\theta - I_2)^{-1} = 1/(2 - 2\cos \theta)(\text{Rot}_{-\theta} - I_2)$.

Se mettiamo P al centro con nuove coordinate $x' = x - P$ troviamo $f(x') = \text{Rot}_\theta x'$ e quindi f è una rotazione intorno a P . $\square$

Abbiamo scoperto che una isometria affine f di $\mathbb{R}^2$ che preserva l'orientazione è l'identità, una traslazione, oppure una rotazione intorno ad un punto $P \in \mathbb{R}^2$. I tre casi sono caratterizzati dall'insieme $\text{Fix}(f)$ dei punti fissi di f :

- se $\text{Fix}(f) = \mathbb{R}^2$ allora f è l'identità;
- se $\text{Fix}(f)$ è un punto allora f è una rotazione;
- se $\text{Fix}(f) = \emptyset$ allora f è una traslazione.

Esercizio 9.3.22. Siano f e g due rotazioni di angolo $\frac{\pi}{2}$ intorno a punti diversi. Mostra che $f \circ g$ è una riflessione rispetto ad un punto di $\mathbb{R}^2$.

Esercizio 9.3.23. Scrivi due rotazioni f e g intorno a punti diversi tali che la composizione $f \circ g$ sia una traslazione.

9.3.7. Isometrie affini del piano che invertono l'orientazione. Passiamo a studiare le isometrie affini del piano che invertono l'orientazione. Sappiamo che le riflessioni ortogonali rispetto alle rette di $\mathbb{R}^2$ appartengono a questa categoria. C'è dell'altro.

Definizione 9.3.24. Una *glissoriflessione* di $\mathbb{R}^2$ è la composizione di una riflessione lungo una retta r con una traslazione non banale (cioè diversa dall'identità) parallela a r . La retta r , orientata nella direzione della traslazione, è detta *asse della glissoriflessione*.

Si veda la Figura 9.17. Notiamo che una glissoriflessione non ha punti fissi. Il *passo* della glissoriflessione è il valore $d(P, f(P))$ per un qualsiasi punto P della retta r . Il passo indica quanto vengono spostati i punti dell'asse r . Geometricamente, una glissoriflessione è determinata dall'asse orientato r e dal passo.

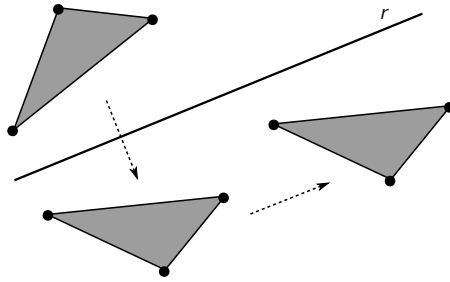


Figura 9.17. Una glissoriflessione.

Esempio 9.3.25. Sia $b_1 > 0$. L'isometria affine

$$f \begin{pmatrix} x \\ y \end{pmatrix} = \begin{pmatrix} 1 & 0 \\ 0 & -1 \end{pmatrix} \begin{pmatrix} x \\ y \end{pmatrix} + \begin{pmatrix} b_1 \\ 0 \end{pmatrix} = \begin{pmatrix} x + b_1 \\ -y \end{pmatrix}$$

è una glissoriflessione ottenuta componendo una riflessione rispetto all'asse x con una traslazione orizzontale di passo b_1 . L'asse della glissoriflessione è l'asse x ed il passo è b_1 .

La seconda affinità mostrata nella Figura 9.15 è una glissoriflessione lungo l'asse y con passo 2.

Per la Proposizione 8.2.7 una isometria affine di $\mathbb{R}^2$ che inverte l'orientazione è del tipo

$$f(x) = \text{Rif}_\theta x + b$$

con $\theta \in [0, 2\pi)$ e $b \in \mathbb{R}^2$.

Proposizione 9.3.26. *Ci sono due casi possibili:*

- se $\text{Fix}(f)$ è una retta, allora f è una riflessione ortogonale.
- se $\text{Fix}(f) = \emptyset$, allora f è una glissoriflessione.

Dimostrazione. Sappiamo che esiste una base ortonormale v_1, v_2 formata da autovettori di Rif_θ rispetto alla quale Rif_θ diventa $\begin{pmatrix} 1 & 0 \\ 0 & -1 \end{pmatrix}$. Cambiamo sistema di riferimento scegliendo v_1, v_2 come base ortonormale. Nelle nuove coordinate l'isometria diventa

$$f \begin{pmatrix} x \\ y \end{pmatrix} = \begin{pmatrix} 1 & 0 \\ 0 & -1 \end{pmatrix} \begin{pmatrix} x \\ y \end{pmatrix} + \begin{pmatrix} b_1 \\ b_2 \end{pmatrix} = \begin{pmatrix} x + b_1 \\ -y + b_2 \end{pmatrix}.$$

Con un ulteriore cambio di coordinate $x' = x$ e $y' = y - b_2/2$ troviamo

$$f' \begin{pmatrix} x' \\ y' \end{pmatrix} = \begin{pmatrix} 1 & 0 \\ 0 & -1 \end{pmatrix} \begin{pmatrix} x' \\ y' \end{pmatrix} + \begin{pmatrix} b_1 \\ b_2 \end{pmatrix} = \begin{pmatrix} x' + b_1 \\ -y' \end{pmatrix}.$$

Questa è una glissoriflessione se $b_1 \neq 0$ e una riflessione se $b_1 = 0$. $\square$

Riassumendo, oltre all'identità ci sono quattro tipi di isometrie affini di $\mathbb{R}^2$, si scrivono come $f(x) = Ax + b$ e sono classificate nella Tabella 9.1 a seconda del segno di $\det A$ e della presenza di punti fissi.

	$\text{Fix}(f) = \emptyset$	$\text{Fix}(f) \neq \emptyset$
$\det A = 1$	traslazioni	rotazioni
$\det A = -1$	glissoriflessioni	riflessioni

Tabella 9.1. Le isometrie affini di $\mathbb{R}^2$ non banali (cioè diverse dall'identità).

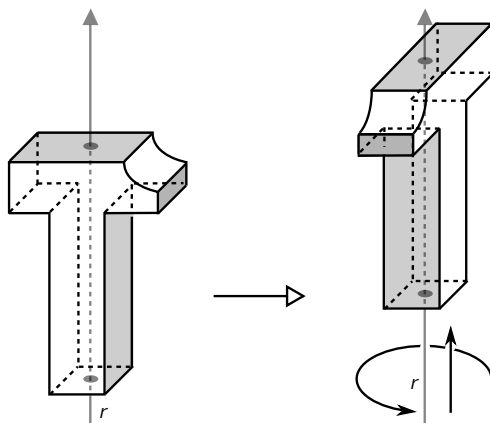


Figura 9.18. Una rototraslazione intorno ad un asse r .

9.3.8. Isometrie dello spazio. Passiamo a classificare le isometrie affini dello spazio. Oltre a quelle già introdotte precedentemente, come riflessioni lungo rette o piani, rotazioni e antirotazioni, ci sono un paio di isometrie nuove.

Definizione 9.3.27. Una *rototraslazione* in $\mathbb{R}^3$ è la composizione di una rotazione intorno ad un asse r e di una traslazione non banale (cioè diversa dall'identità) parallela a r . La retta r orientata nella direzione della traslazione, è detta *asse* della rototraslazione.

Notiamo che una rototraslazione non ha punti fissi. Il *passo* della rototraslazione è il valore $d(P, f(P))$ per un qualsiasi punto P della retta r . Il passo indica di quanto vengono spostati i punti dell'asse. Geometricamente, una rototraslazione è determinata dall'asse orientato r , dall'angolo θ di rotazione e dal passo.

Esempio 9.3.28. Sia $b_1 > 0$. L'isometria affine

$$f \begin{pmatrix} x \\ y \\ z \end{pmatrix} = \begin{pmatrix} 1 & 0 & 0 \\ 0 & \cos \theta & -\sin \theta \\ 0 & \sin \theta & \cos \theta \end{pmatrix} \begin{pmatrix} x \\ y \\ z \end{pmatrix} + \begin{pmatrix} b_1 \\ 0 \\ 0 \end{pmatrix}$$

è una rototraslazione ottenuta componendo una rotazione di angolo θ intorno all'asse x con una traslazione lungo l'asse di passo b_1 . La rototraslazione ha asse x , angolo θ e passo b_1 .

Definizione 9.3.29. Una *glissoriflessione* in $\mathbb{R}^3$ è la composizione di una riflessione rispetto ad un piano π e di una traslazione lungo una direzione parallela a π .

Notiamo che una glissoriflessione non ha punti fissi.

Esempio 9.3.30. L'isometria

$$f \begin{pmatrix} x \\ y \\ z \end{pmatrix} = \begin{pmatrix} -1 & 0 & 0 \\ 0 & 1 & 0 \\ 0 & 0 & 1 \end{pmatrix} \begin{pmatrix} x \\ y \\ z \end{pmatrix} + \begin{pmatrix} 0 \\ b_2 \\ b_3 \end{pmatrix}$$

è una glissoriflessione ottenuta riflettendo rispetto al piano $\pi = \{x = 0\}$ e poi traslando di un vettore b parallelo a π .

Classifichiamo infine tutte le isometrie affini di $\mathbb{R}^3$.

Proposizione 9.3.31. *Ogni isometria affine $f(x) = Ax + b$ di $\mathbb{R}^3$ diversa dall'identità è di uno dei tipi seguenti:*

- se $\det A = 1$, è una traslazione, rotazione o rototraslazione;
- se $\det A = -1$, è una riflessione, antirotazione o glissoriflessione.

Dimostrazione. Sappiamo che A è una matrice ortogonale e quindi per il Teorema 8.2.13 rappresenta l'identità, una rotazione o una antirotazione. Consideriamo questi tre casi separatamente.

Se $A = I_3$ allora f è l'identità se $b = 0$ oppure una traslazione se $b \neq 0$. Se A è una rotazione di angolo $\theta \neq 0$ intorno ad un asse r , allora prendendo come nuovo sistema di riferimento una base ortonormale v_1, v_2, v_3 con $v_1 \in r$ otteniamo

$$f \begin{pmatrix} x \\ y \\ z \end{pmatrix} = \begin{pmatrix} 1 & 0 & 0 \\ 0 & \cos \theta & -\sin \theta \\ 0 & \sin \theta & \cos \theta \end{pmatrix} \begin{pmatrix} x \\ y \\ z \end{pmatrix} + \begin{pmatrix} b_1 \\ b_2 \\ b_3 \end{pmatrix}.$$

Poiché $\theta \neq 0$, la mappa $g: \mathbb{R}^2 \rightarrow \mathbb{R}^2$

$$g \begin{pmatrix} y \\ z \end{pmatrix} = \begin{pmatrix} \cos \theta & -\sin \theta \\ \sin \theta & \cos \theta \end{pmatrix} \begin{pmatrix} y \\ z \end{pmatrix} + \begin{pmatrix} b_2 \\ b_3 \end{pmatrix}$$

è una rotazione intorno ad un punto $\begin{pmatrix} y_0 \\ z_0 \end{pmatrix} \in \mathbb{R}^2$. Cambiando variabili con

$$x' = x, \quad y' = y - y_0, \quad z' = z - z_0$$

otteniamo che la stessa f si scrive come

$$f' \begin{pmatrix} x' \\ y' \\ z' \end{pmatrix} = \begin{pmatrix} 1 & 0 & 0 \\ 0 & \cos \theta & -\sin \theta \\ 0 & \sin \theta & \cos \theta \end{pmatrix} \begin{pmatrix} x' \\ y' \\ z' \end{pmatrix} + \begin{pmatrix} b_1 \\ 0 \\ 0 \end{pmatrix}.$$

$\text{Fix}(f)$	$\emptyset$	punto	retta	piano
$\det A = 1$	(roto-)traslazioni		rotazioni	
$\det A = -1$	glissoriflessioni	antirotazioni		riflessioni

Tabella 9.2. Le isometrie affini di $\mathbb{R}^3$ classificate a seconda del tipo di punti fissi.

Questa è una rotazione intorno all'asse x' se $b_1 = 0$, ed una rototraslazione intorno all'asse x' se $b_1 \neq 0$.

Se A è una antirotazione, ci sono due casi da considerare. Se l'antirotazione ha angolo $\theta = 0$, è in realtà una riflessione rispetto ad un piano π . Prendendo una base ortonormale v_1, v_2, v_3 con $v_2, v_3 \in \pi$ otteniamo

$$f \begin{pmatrix} x \\ y \\ z \end{pmatrix} = \begin{pmatrix} -1 & 0 & 0 \\ 0 & 1 & 0 \\ 0 & 0 & 1 \end{pmatrix} \begin{pmatrix} x \\ y \\ z \end{pmatrix} + \begin{pmatrix} b_1 \\ b_2 \\ b_3 \end{pmatrix}.$$

Dopo un cambio di coordinate $x' = x - b_1/2$ sparisce il termine b_1 e quindi otteniamo una glissoriflessione (se $b_2 \neq 0$ oppure $b_3 \neq 0$) oppure una riflessione (se $b_2 = b_3 = 0$) rispetto a π .

Se l'antirotazione ha angolo $\theta \neq 0$, allora A non ha autovalore 1 e quindi ha esattamente un punto fisso P . Spostando l'origine in P otteniamo che f è una antirotazione $f(x') = Ax'$. $\square$

I tipi di isometrie affini sono elencati nella Tabella 9.2, a seconda del segno di $\det A$ e del tipo di punti fissi.

9.3.9. Similitudini. Le isometrie affini preservano la distanza fra punti. Esiste una classe più ampia di isomorfismi affini che è molto importante in geometria, quella delle *similitudini*: queste non preservano necessariamente la distanza fra punti, ma conservano gli angoli fra vettori.

Definizione 9.3.32. Una *similitudine* è un isomorfismo affine $f: \mathbb{R}^n \rightarrow \mathbb{R}^n$ del tipo seguente:

$$f(x) = \lambda Ax + b$$

dove A è una matrice ortogonale, $\lambda > 0$ e $b \in \mathbb{R}^n$.

La parte lineare λA della similitudine è una composizione di una isometria A e di una omotetia λ .

Esempio 9.3.33. Le isometrie sono similitudini. Le omotetie sono similitudini. Componendo isometrie e omotetie si ottengono altre similitudini. Ad esempio, la similitudine in $\mathbb{R}^2$

$$f \begin{pmatrix} x \\ y \end{pmatrix} = \begin{pmatrix} 2y \\ -2x \end{pmatrix} + \begin{pmatrix} 1 \\ 1 \end{pmatrix}$$

è una rotazione di $\frac{3\pi}{2}$ intorno all'origine composta con una omotetia di fattore 2 avente sempre come centro l'origine e infine con una traslazione

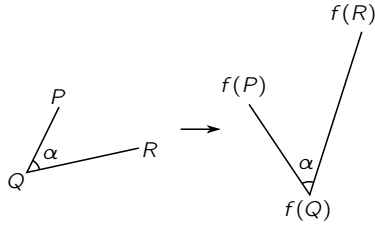


Figura 9.19. Le similitudini preservano gli angoli.

di vettore $\begin{pmatrix} 1 \\ 1 \end{pmatrix}$. Per capire la natura geometrica di questa trasformazione cerchiamo i punti fissi: risolvendo le equazioni $x = 2y + 1$ e $y = -2x + 1$ troviamo un unico punto fisso

$$P = \frac{1}{5} \begin{pmatrix} 3 \\ -1 \end{pmatrix}.$$

Se cambiamo le coordinate in modo da spostare l'origine in P con $\begin{pmatrix} x' \\ y' \end{pmatrix} = \begin{pmatrix} x \\ y \end{pmatrix} - P$ la similitudine diventa

$$f \begin{pmatrix} x' \\ y' \end{pmatrix} = \begin{pmatrix} 2y' \\ -2x' \end{pmatrix}.$$

Quindi f è la composizione di una rotazione di angolo $\frac{3\pi}{2}$ e di una omotetia entrambe con centro P .

Notiamo subito che in una similitudine le distanze fra punti cambiano in modo molto controllato: vengono tutte moltiplicate per λ .

Proposizione 9.3.34. *Sia $f(x) = \lambda Ax + b$ una similitudine, con A matrice ortogonale. Per ogni coppia di punti $P, Q \in \mathbb{R}^n$ vale*

$$d(f(P), f(Q)) = \lambda d(P, Q).$$

Dimostrazione. La funzione f è la composizione di una isometria L_A , di una omotetia $g(x) = \lambda x$ e di una traslazione $h(y) = y + b$. La prima e la terza sono isometrie e non modificano le distanze fra punti. La seconda le modifica in questo modo:

$$d(g(P), g(Q)) = d(\lambda P, \lambda Q) = \|\lambda Q - \lambda P\| = \lambda \|Q - P\| = \lambda d(P, Q).$$

La dimostrazione è completa. $\square$

Le distanze fra punti cambiano tutte dello stesso fattore $\lambda > 0$. Il numero λ è il *fattore di dilatazione* della similitudine f . Una similitudine è una isometria se e solo se $\lambda = 1$.

Come avevamo accennato, le similitudini preservano sempre gli angoli. Dati tre punti $P, Q, R \in \mathbb{R}^n$ con $P, R \neq Q$, indichiamo con $\alpha(P\hat{Q}R)$ l'angolo fra i vettori $\overrightarrow{QP}$ e $\overrightarrow{QR}$. Si veda la Figura 9.19.

Proposizione 9.3.35. *Sia f una similitudine. Vale*

$$\alpha(f(P)\widehat{f(Q)f(R)}) = \alpha(P\widehat{Q}R)$$

per qualsiasi terna $P, Q, R \in \mathbb{R}^n$ di punti distinti.

Dimostrazione. Una similitudine è la composizione di un'isometria lineare, un'omotetia $g(x) = \lambda x$ e una traslazione. Tutte e tre queste trasformazioni affini non cambiano gli angoli. $\square$

Mostriamo infine che le similitudini con fattore di dilatazione $\lambda \neq 1$ hanno sempre esattamente un punto fisso. Questo risultato contrasta molto con quanto accade per le isometrie, che come abbiamo visto possono anche avere nessun punto fisso (ad esempio, le traslazioni) o infiniti (ad esempio, le riflessioni lungo rette o piani).

Proposizione 9.3.36. *Una similitudine $f: \mathbb{R}^n \rightarrow \mathbb{R}^n$ con dilatazione $\lambda \neq 1$ ha esattamente un punto fisso P .*

Dimostrazione. La similitudine è $f(x) = \lambda Ax + b$. La matrice A è ortogonale e quindi può avere come autovalori solo ± 1 . Allora λA può avere come autovalori solo $\pm \lambda$. Siccome $\lambda \neq 1$, la matrice λA non ha l'autovalore 1 e quindi f ha esattamente un punto fisso per la Proposizione 9.3.19. $\square$

Il punto fisso P è il *centro* della similitudine.

9.3.10. Congruenza, similitudine e equivalenza affine. Studieremo nelle prossime pagine alcuni oggetti geometrici contenuti in $\mathbb{R}^n$. A questo scopo introduciamo alcune definizioni.

Definizione 9.3.37. Due sottoinsiemi $S, S' \subset \mathbb{R}^n$ sono *congruenti* o *isometrici* se esiste una isometria affine f di $\mathbb{R}^n$ tale che $f(S) = S'$. I due sottoinsiemi sono *simili* se esiste una similitudine f tale che $f(S) = S'$. Sono *affinemente equivalenti* se esiste un isomorfismo affine f tale che $f(S) = S'$.

Per due sottoinsiemi S e S' valgono sempre queste implicazioni:

congruenti $\implies$ simili $\implies$ affinemente equivalenti.

Le implicazioni opposte non valgono sempre, come vedremo nelle prossime pagine studiando i triangoli: si veda la Figura 9.20.

La congruenza, la similitudine e l'equivalenza affine sono tutte e tre relazioni di equivalenza fra sottoinsiemi di $\mathbb{R}^n$.

Esercizio 9.3.38. Due sottospazi affini S e S' di $\mathbb{R}^n$ sono congruenti $\iff$ sono simili $\iff$ sono affinemente equivalenti $\iff$ hanno la stessa dimensione.

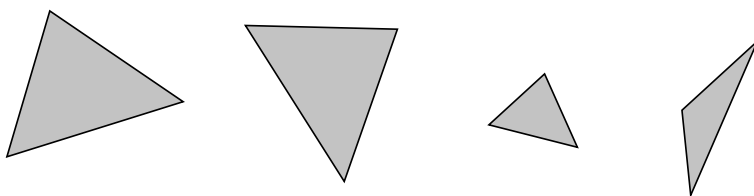


Figura 9.20. Il primo triangolo è congruente al secondo, è simile (ma non congruente) al terzo, ed è affinementemente equivalente (ma non simile) al quarto triangolo.

9.3.11. Aree e volumi. Abbiamo notato nella Sezione 3.3.10 che il valore assoluto del determinante di una matrice 2×2 o 3×3 misura l'area o il volume del parallelogramma o parallelepipedo generato dai vettori colonna. Questa informazione si riflette nelle trasformazioni affini.

Proposizione 9.3.39. Sia $f(x) = Ax + b$ un isomorfismo affine di $\mathbb{R}^2$. Sia $S \subset \mathbb{R}^2$ un oggetto geometrico che ha una certa area $\text{Area}(S)$. Allora

$$\text{Area}(f(S)) = |\det A| \cdot \text{Area}(S).$$

Dimostrazione. L'oggetto geometrico S si decompone in tanti rettangoli. Per quanto visto nella Sezione 3.3.10, ciascun rettangolo viene trasformato in parallelogrammi di area $|\det A|$ volte quella del rettangolo originale. Quindi tutta l'area di $f(S)$ è $|\det A|$ volte quella di S .

L'oggetto S si decompone in realtà in infiniti rettangoli sempre più piccoli: per rendere rigoroso questo argomento avremmo bisogno del calcolo infinitesimale. $\square$

Vale un teorema analogo, con una dimostrazione simile, nello spazio.

Proposizione 9.3.40. Sia $f(x) = Ax + b$ un isomorfismo affine di $\mathbb{R}^3$. Sia $S \subset \mathbb{R}^3$ un oggetto geometrico che ha un certo volume $\text{Vol}(S)$. Allora

$$\text{Vol}(f(S)) = |\det A| \cdot \text{Vol}(S).$$

Ad esempio, le due ultime trasformazioni della Figura 9.15 hanno entrambe $|\det A| = 1$ e quindi preservano le aree degli oggetti: nonostante distorcano molto la P in figura, non ne modificano l'area totale!

Esercizi

Esercizio 9.1. Considera tre punti qualsiasi del piano $\mathbb{R}^2$

$$A = \begin{pmatrix} x_A \\ y_A \end{pmatrix}, \quad B = \begin{pmatrix} x_B \\ y_B \end{pmatrix}, \quad C = \begin{pmatrix} x_C \\ y_C \end{pmatrix}$$

Mostra che i tre punti sono allineati se e solo se

$$\det \begin{pmatrix} x_A & y_A & 1 \\ x_B & y_B & 1 \\ x_C & y_C & 1 \end{pmatrix} = 0.$$

Esercizio 9.2. Sia $u \in \mathbb{R}^3$ unitario e $w \in \mathbb{R}^3$ ortogonale a u . Mostra che

$$u \times (u \times w) = -w.$$

Esercizio 9.3. Considera le rette affini

$$r = \left\{ \begin{pmatrix} 1 \\ 1 \\ 0 \end{pmatrix} + \text{Span} \begin{pmatrix} 0 \\ 0 \\ 1 \end{pmatrix} \right\}, \quad s = \left\{ \begin{pmatrix} 2 \\ 0 \\ 3 \end{pmatrix} + \text{Span} \begin{pmatrix} 0 \\ 1 \\ 0 \end{pmatrix} \right\}.$$

Dimostra che sono sghembe. Determina l'unica retta affine perpendicolare ad entrambe.

Esercizio 9.4. Considera i punti

$$P = \begin{pmatrix} 1 \\ 1 \\ 0 \end{pmatrix}, \quad Q = \begin{pmatrix} -1 \\ 0 \\ 1 \end{pmatrix}, \quad R = \begin{pmatrix} 2 \\ 0 \\ 0 \end{pmatrix}, \quad S = \begin{pmatrix} 0 \\ 0 \\ 1 \end{pmatrix}.$$

- (1) Trova delle equazioni parametriche per la retta r passante per P e Q .
- (2) Trova delle equazioni cartesiane per il piano π ortogonale a r e passante per R .
- (3) Trova delle equazioni cartesiane per il piano π' contenente r e parallelo alla retta passante per R e S .

Esercizio 9.5. Considera le rette

$$r = \begin{cases} x - y + z = 0, \\ 2x + y + z = 1, \end{cases} \quad s = \begin{cases} x + z = 1, \\ y - z = -2. \end{cases}$$

- (1) Determina se sono incidenti, parallele o sghembe. Calcola la loro distanza e trova tutte le rette perpendicolari a entrambe.
- (2) Determina il piano contenente r parallelo a s .

Esercizio 9.6. Siano P e Q due punti distinti di $\mathbb{R}^3$. Mostra che il luogo dei punti equidistanti da P e Q è un piano π .

Esercizio 9.7. Siano r e r' due rette sghembe in $\mathbb{R}^3$ e P un punto disgiunto da r e r' , tale che il piano π passante per P e r non sia parallelo a r' , ed il piano π' passante per P e r' non sia parallelo ad r . Mostra che esiste un'unica retta s passante per P che interseca r e r' .

Suggerimento. Considera l'intersezione fra π e r' .

Esercizio 9.8. Nell'esercizio precedente, determina esplicitamente s nel caso seguente:

$$P = \begin{pmatrix} 2 \\ 1 \\ 0 \end{pmatrix}, \quad r = \left\{ \frac{x+3}{2} = 1 - y = -2z \right\}, \quad r' = \begin{cases} 2x - y - 1 = 0, \\ 3x + 3y - z = 0. \end{cases}$$

Esercizio 9.9. Considera il piano affine $\pi = \{x + y - z = 2\}$ in $\mathbb{R}^3$. Scrivi l'affinità $f(x) = Ax + b$ che rappresenta una riflessione rispetto a π .

Esercizio 9.10. Scrivi l'affinità $f(x) = Ax + b$ di $\mathbb{R}^2$ che rappresenta una rotazione di angolo $\frac{\pi}{2}$ intorno al punto $P = \begin{pmatrix} 3 \\ 1 \end{pmatrix}$.

Esercizio 9.11. Scrivi l'affinità $f(x) = Ax + b$ di $\mathbb{R}^3$ che rappresenta una rotazione di angolo $\frac{2\pi}{3}$ intorno all'asse r seguente:

$$r = \left\{ \begin{pmatrix} 1 \\ 0 \\ 1 \end{pmatrix} + t \begin{pmatrix} 1 \\ 1 \\ 1 \end{pmatrix} \right\}$$

Esercizio 9.12. Siano f e g due rotazioni del piano $\mathbb{R}^2$ di angolo $\frac{\pi}{2}$ intorno a due punti P e Q . Mostra che la composizione $f \circ g$ è una riflessione rispetto ad un qualche punto R .

Esercizio 9.13. Siano f e g due riflessioni di $\mathbb{R}^2$ intorno a due rette affini r e s . Mostra che:

- se r e s sono parallele, la composizione $f \circ g$ è una traslazione;
- se r e s sono incidenti, la composizione $f \circ g$ è una rotazione.

Esercizio 9.14. Scrivi un'affinità $f(x) = Ax + b$ di $\mathbb{R}^2$ tale che $f(r_1) = r_2$, $f(r_2) = r_3$, $f(r_3) = r_1$ dove r_1, r_2, r_3 sono le rette

$$r_1 = \{y = 1\}, \quad r_2 = \{x = 1\}, \quad r_3 = \{x + y = -1\}.$$

Esercizio 9.15. Scrivi una isometria $f(x) = Ax + b$ di $\mathbb{R}^2$ senza punti fissi, tale che $f(r) = r'$ dove r e r' sono le rette

$$r = \{y = x\}, \quad r' = \{y = -x\}.$$

Esercizio 9.16. Scrivi una rototraslazione $f(x) = Ax + b$ in $\mathbb{R}^3$ che abbia asse

$$r = \left\{ \begin{pmatrix} 1 \\ -1 \\ 2 \end{pmatrix} + t \begin{pmatrix} 0 \\ 0 \\ 1 \end{pmatrix} \right\}$$

tale che

$$f \begin{pmatrix} 0 \\ 0 \\ 1 \end{pmatrix} = \begin{pmatrix} 2 \\ -2 \\ 2 \end{pmatrix}.$$

Esercizio 9.17. Considera i piani affini $\pi_1 = \{x = 0\}$ e $\pi_2 = \{y = 0\}$ in $\mathbb{R}^3$. Scrivi una isometria $f(x) = Ax + b$ di $\mathbb{R}^3$ senza punti fissi tale che $f(\pi_1) = \pi_2$.

Esercizio 9.18. Sia $f: \mathbb{R}^3 \rightarrow \mathbb{R}^3$ una rototraslazione con asse r . Mostra che i punti che si spostano di meno sono proprio quelli di r , cioè se $P \in r$ e $Q \notin r$ allora

$$d(P, f(P)) < d(Q, f(Q)).$$

Esercizio 9.19. Considera la matrice ortogonale

$$A = \frac{1}{7} \begin{pmatrix} -6 & 3 & -2 \\ 2 & 6 & 3 \\ 3 & 2 & -6 \end{pmatrix}$$

- (1) Determina che tipo di isometria sia $L_A: \mathbb{R}^3 \rightarrow \mathbb{R}^3$.
- (2) Per quali $b \in \mathbb{R}^3$ l'isometria affine $f_b(x) = Ax + b$ è una rotazione intorno ad un asse?

Esercizio 9.20. Sia r la retta affine passante per l'origine e per $e_1 + e_2$ e s la retta definita dalle equazioni

$$\begin{cases} y + z = 1, \\ x + 2y + 2z = 1 \end{cases}.$$

- (1) Scrivi una isometria affine che porti r in s .
- (2) Scrivi una isometria affine che porti r in s che non abbia punti fissi, e una che abbia almeno un punto fisso.
- (3) Esiste una riflessione ortogonale che porta r in s ?

Esercizio 9.21. Considera in $\mathbb{R}^3$ il piano $\pi = \{z = 1\}$ e la retta

$$r = \left\{ \begin{pmatrix} 2+t \\ t \\ -1 \end{pmatrix} \mid t \in \mathbb{R} \right\}.$$

- (1) Determina la distanza fra π e r .
- (2) Determina delle equazioni cartesiane per il piano π' contenente r e perpendicolare a π .
- (3) Scrivi una isometria $f(x) = Ax + b$ che sposti il piano π in π' .

Poligoni e poliedri

Il piano $\mathbb{R}^2$ e lo spazio $\mathbb{R}^3$ contengono varie figure geometriche, che ci interessano da vicino perché possono essere usate per modellizzare gli oggetti che ci stanno intorno: molecole, manufatti, pianeti, eccetera. Le figure geometriche più semplici sono i *cerchi* ed i *poligoni* nel piano e le *sfere* ed i *poliedri* nello spazio. In questo capitolo utilizziamo gli strumenti di geometria e algebra lineare introdotti nelle pagine precedenti per definire e studiare approfonditamente queste forme geometriche.

10.1. Poligoni

I poligoni sono (assieme ai cerchi) le figure geometriche più semplici che si possono disegnare in una porzione limitata di piano. Fra questi giocano un ruolo molto importante ovviamente i triangoli.

10.1.1. Segmenti. Introduciamo alcune definizioni. Siano A e B due punti distinti di $\mathbb{R}^n$.

Definizione 10.1.1. Il *segmento* con estremi A e B è il sottoinsieme di $\mathbb{R}^n$

$$AB = \{A + t\overrightarrow{AB} \mid t \in [0, 1]\}.$$

Il segmento AB ha una *lunghezza* $|AB|$ pari alla norma di $\overrightarrow{AB}$ e cioè alla distanza fra A e B :

$$|AB| = \|\overrightarrow{AB}\| = d(A, B).$$

Si veda la Figura 10.1-(sinistra).

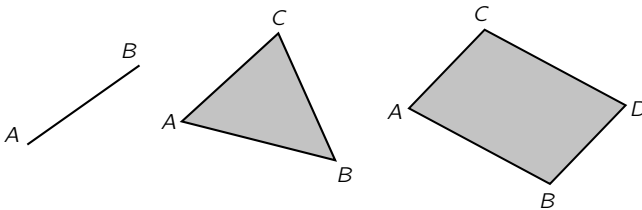


Figura 10.1. Un segmento, un triangolo e un parallelogramma.

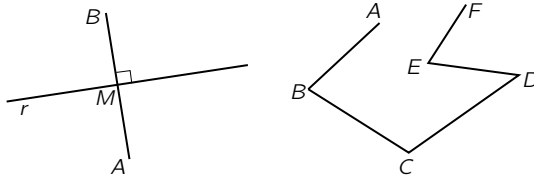


Figura 10.2. L'asse di un segmento AB è la retta r ortogonale ad AB passante per il punto medio M (sinistra). Una linea spezzata con 5 segmenti e 6 vertici (destra).

Esercizio 10.1.2. Valgono le uguaglianze

$$\begin{aligned} AB = BA &= \{tA + (1-t)B \mid t \in [0, 1]\} \\ &= \{tA + uB \mid t, u \geq 0, t + u = 1\}. \end{aligned}$$

Il segmento AB è contenuto nell'unica retta r passante per A e B . Il punto medio del segmento AB è il punto

$$M = \frac{1}{2}(A + B) = A + \frac{1}{2}\overrightarrow{AB}.$$

Il punto medio M divide il segmento AB in due segmenti AM e MB entrambi della stessa lunghezza $\frac{1}{2}|AB|$.

Se A e B sono due punti del piano $\mathbb{R}^2$, l'asse del segmento AB è la retta ortogonale ad AB passante per il punto medio M . Si veda la Figura 10.2-(sinistra).

Proposizione 10.1.3. L'asse del segmento AB è l'insieme dei punti equidistanti dai punti A e B .

Dimostrazione. Scegliamo un sistema di riferimento per $\mathbb{R}^2$ in cui l'origine è il punto medio di AB e l'asse x contiene AB . Con questo sistema troviamo $B = \begin{pmatrix} a \\ 0 \end{pmatrix}$, $A = \begin{pmatrix} -a \\ 0 \end{pmatrix}$ e l'asse è $\{x = 0\}$. La distanza quadra di un punto $\begin{pmatrix} x \\ y \end{pmatrix}$ da A e da B rispettivamente è $(x - a)^2 + y^2$ e $(x + a)^2 + y^2$. Queste sono uguali se e solo se $x = 0$. $\square$

Due segmenti AB e BC in $\mathbb{R}^n$ che condividono lo stesso estremo B sono *concatenati*. Una successione finita di segmenti consecutivamente concatenati come nella Figura 10.2-(destra) è una *linea spezzata*. Gli estremi dei segmenti si chiamano *vertici*. La linea spezzata è *chiusa* se il secondo estremo dell'ultimo segmento coincide con il primo estremo del primo segmento. È *intrecciata* se ci sono altre intersezioni fra i segmenti oltre a quelle ovvie (cioè i vertici in cui si intersecano due segmenti successivi, o l'ultimo ed il primo se è chiusa).

Studiamo adesso quando due segmenti sono simili o congruenti.

Proposizione 10.1.4. I segmenti in $\mathbb{R}^n$ sono tutti simili fra loro. Due segmenti sono congruenti $\iff$ hanno la stessa lunghezza.

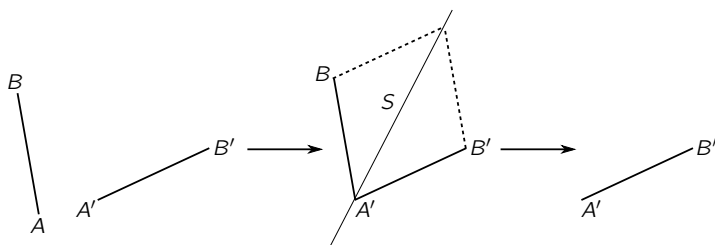


Figura 10.3. Una traslazione e una riflessione sono sufficienti per sovrapporre due segmenti congruenti AB e $A'B'$.

Dimostrazione. Siccome le isometrie preservano le distanze, due segmenti di lunghezza diversa non possono essere congruenti.

D'altro canto, siano AB e $A'B'$ due segmenti della stessa lunghezza. La Figura 10.3 descrive come sia possibile spostare isometricamente AB su $A'B'$ in due passi. Con la traslazione $f(x) = x + A' - A$ possiamo portare A a coincidere con A' . Successivamente, per la Proposizione 8.2.6, possiamo fare una riflessione lungo il sottospazio affine

$$S = A + \text{Span}(\overrightarrow{AB} + \overrightarrow{A'B'})$$

che sposta B in B' .

Se AB e $A'B'$ non hanno la stessa lunghezza, possiamo preventivamente trasformare AB con una omotetia di centro qualsiasi in modo che abbia la stessa lunghezza di $A'B'$ e quindi procedere come sopra. Quindi tutti i segmenti sono simili. $\square$

10.1.2. Angoli. Sia $P \in \mathbb{R}^n$ un punto e $v \in \mathbb{R}^n$ un vettore non nullo. La *semiretta* uscente da P con vettore direzione v è il sottoinsieme di $\mathbb{R}^n$

$$s = \{P + tv \mid t \geq 0\}.$$

La *semiretta opposta* è $\{P - tv \mid t \geq 0\}$ e l'unione delle due semirette forma la retta affine $P + \text{Span}(v)$.

Sia $P \in \mathbb{R}^2$ e siano s e s' due semirette distinte uscenti da P , con vettori direzione v e v' . Le due semirette suddividono il piano in due zone dette *angoli*.

Più precisamente: se v e v' sono indipendenti, l'*angolo convesso* fra s e s' è il sottoinsieme

$$A = \{tv + uv' \mid t, u \geq 0\}$$

mentre l'*angolo concavo* è

$$A' = \{tv + uv' \mid t \leq 0 \text{ oppure } u \leq 0\}.$$

Si veda la Figura 10.4-(sinistra). L'*ampiezza* dell'angolo convesso A è l'angolo fra i vettori v e v' , cioè il numero $\alpha \in (0, \pi)$ definito tramite la

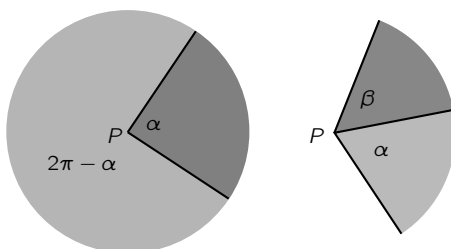


Figura 10.4. Due semirette distinte uscenti da P determinano due angoli, di ampiezza α e $2\pi - \alpha$ (sinistra). Se affianchiamo due angoli con vertice P con una semiretta in comune di ampiezza α e β , otteniamo un angolo di ampiezza $\alpha + \beta$ (destra).

formula

$$\cos \alpha = \frac{\langle v, v' \rangle}{|v||v'|}.$$

L'ampiezza dell'angolo concavo A' è invece $2\pi - \alpha$. Se v e v' non sono indipendenti, hanno verso opposto e $r = s \cup s'$ è una retta affine. Questa divide il piano in due *semipiani* A e A' nel modo seguente:

$$A = \{w \in \mathbb{R}^2 \mid \det(vw) \geq 0\}, \quad A' = \{w \in \mathbb{R}^2 \mid \det(vw) \leq 0\}.$$

Ciascun semipiano è un *angolo piatto* ed ha ampiezza π . In tutti i casi P è il *vertice* di ciascuno dei due angoli. Si vede facilmente che due angoli sono congruenti se e solo se hanno la stessa ampiezza.

Se affianchiamo due angoli di ampiezza α e β lungo una semiretta comune come nella Figura 10.4-(destra), otteniamo un angolo di ampiezza $\alpha + \beta$. Questo è conseguenza della Proposizione 8.2.15.

Osservazione 10.1.5. Può sembrare ovvio che affiancando due angoli le ampiezze si sommino, ma non lo è per come abbiamo definito le ampiezze: questa proprietà discende dalla formula di addizione $\cos(\alpha + \beta) = \cos \alpha \cos \beta - \sin \alpha \sin \beta$ che è stata usata nella dimostrazione della Proposizione 8.2.15. Questo fatto sarebbe stato invece ovvio se avessimo definito l'ampiezza di un angolo come la lunghezza dell'arco di circonferenza unitaria corrispondente: questa strada è forse più naturale ma avrebbe necessitato una definizione rigorosa della lunghezza delle curve nel piano, che non è banale e fa uso del calcolo infinitesimale.

Un angolo di ampiezza minore di $\frac{\pi}{2}$, uguale a $\frac{\pi}{2}$, o compreso fra $\frac{\pi}{2}$ e π è detto *acuto*, *retto* o *ottuso*.

Tre punti A, B, C non allineati del piano determinano un angolo concavo, che indichiamo con $\widehat{ABC}$, di vertice B e con semirette $\{B + t\overrightarrow{BA}\}$ e $\{B + u\overrightarrow{BC}\}$.

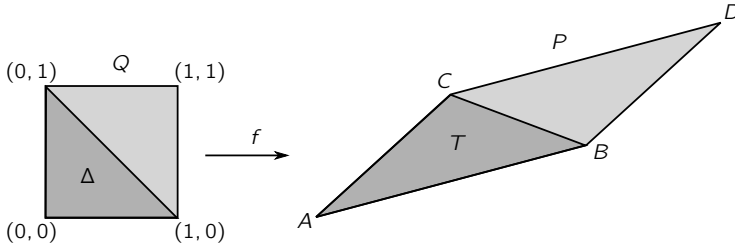


Figura 10.5. L'isomorfismo affine f manda il triangolo Δ nel triangolo T ed il quadrato Q nel parallelogramma P .

10.1.3. Triangoli e parallelogrammi. Definiamo adesso triangoli e parallelogrammi, già mostrati nella Figura 10.1.

Siano A , B e C tre punti non allineati di $\mathbb{R}^n$.

Definizione 10.1.6. Il *triangolo* con vertici A , B e C è il sottoinsieme di $\mathbb{R}^n$

$$T = \{A + t\overrightarrow{AB} + u\overrightarrow{AC} \mid t, u \in [0, 1], t + u \leq 1\}.$$

Questa definizione è molto concreta ma può risultare un po' oscura e necessita di una giustificazione geometrica. Notiamo innanzitutto che lo spazio dei parametri usato nella definizione è

$$\Delta = \{(t, u) \in \mathbb{R}^2 \mid t, u \in [0, 1], t + u \leq 1\}.$$

Questo spazio è effettivamente il triangolo con vertici $(0, 0)$, $(1, 0)$ e $(0, 1)$. Consideriamo adesso l'isomorfismo affine $f: \mathbb{R}^2 \rightarrow \mathbb{R}^2$ dato da

$$f \begin{pmatrix} t \\ u \end{pmatrix} = A + t\overrightarrow{AB} + u\overrightarrow{AC}.$$

Notiamo che $T = f(\Delta)$ è l'immagine del triangolo Δ lungo l'isomorfismo affine f . La mappa f manda i vertici di Δ in A , B , C e gli spigoli di Δ in AB , BC , CA . Si veda la Figura 10.5.

Analogamente definiamo il *parallelogramma* L con lati AB e AC come il sottoinsieme

$$P = \{A + t\overrightarrow{AB} + u\overrightarrow{AC} \mid t, u \in [0, 1]\}.$$

Come sopra, notiamo che $P = f(Q)$ dove Q è il quadrato

$$Q = \{(t, u) \in \mathbb{R}^2 \mid t, u \in [0, 1]\}.$$

Si veda sempre la Figura 10.5. I *vertici* del triangolo sono A , B e C , mentre quelli del parallelogramma sono A , B , C e $D = A + \overrightarrow{AB} + \overrightarrow{AC}$. I *lati* del triangolo sono i segmenti AB , BC e CA , mentre quelli del parallelogramma sono i segmenti AB , BD , DC e CA .

Esercizio 10.1.7. Vale la seguente uguaglianza

$$T = \{sA + tB + uC \mid s, t, u \geq 0, s + t + u = 1\}.$$

Come abbiamo visto nel Corollario 9.1.6, l'area del parallelogramma L è

$$\text{Area}(L) = \|\vec{AB} \times \vec{AC}\|.$$

Poiché il parallelogramma si decompone in due triangoli isometrici a T , ne deduciamo che

$$\text{Area}(T) = \frac{1}{2} \|\vec{AB} \times \vec{AC}\|.$$

10.1.4. Circonferenze e cerchi. La *circonferenza* in $\mathbb{R}^2$ di centro $P \in \mathbb{R}^2$ e di raggio $r > 0$ è l'insieme dei punti Q che hanno distanza r da P :

$$C = \{Q \in \mathbb{R}^2 \mid d(Q, P) = r\}.$$

Se $P = \begin{pmatrix} x_0 \\ y_0 \end{pmatrix}$ e $Q = \begin{pmatrix} x \\ y \end{pmatrix}$, si può scrivere la condizione $d(Q, P)^2 = r^2$ così:

$$(x - x_0)^2 + (y - y_0)^2 = r^2.$$

La circonferenza C è l'insieme delle soluzioni di questa equazione di secondo grado. Sviluppando i termini si ottiene

$$(16) \quad x^2 + y^2 - 2x_0x - 2y_0y + x_0^2 + y_0^2 - r^2 = 0.$$

Una equazione generica del tipo

$$(17) \quad a(x^2 + y^2) + bx + cy + d = 0$$

descrive una circonferenza se le disuguaglianze seguenti sono soddisfatte:

$$a \neq 0, \quad b^2 + c^2 > 4ad.$$

In questo caso l'equazione rappresenta effettivamente una circonferenza di centro $P = -\frac{1}{2a} \begin{pmatrix} b \\ c \end{pmatrix}$ e raggio $r = \frac{1}{2|a|} \sqrt{b^2 + c^2 - 4ad}$. Infatti basta sostituire questi valori in (16) per ottenere (17).

Si può anche descrivere C in forma parametrica:

$$C = \left\{ \begin{pmatrix} x_0 + r \cos \theta \\ y_0 + r \sin \theta \end{pmatrix} \mid \theta \in [0, 2\pi) \right\}.$$

Il *cerchio* D di centro $P = \begin{pmatrix} x_0 \\ y_0 \end{pmatrix}$ e raggio $r > 0$ è l'insieme dei punti che hanno distanza $\leq r$ da P , cioè

$$D = \{Q \in \mathbb{R}^2 \mid d(Q, P) \leq r\}.$$

Questi sono i punti $Q = \begin{pmatrix} x \\ y \end{pmatrix}$ che soddisfano la disuguaglianza

$$(x - x_0)^2 + (y - y_0)^2 \leq r^2.$$

10.1.5. Teoremi su circonferenze e cerchi. Dimostriamo alcuni teoremi su circonferenze e cerchi.

Proposizione 10.1.8. *Per tre punti non allineati*

$$A = \begin{pmatrix} x_A \\ y_A \end{pmatrix}, \quad B = \begin{pmatrix} x_B \\ y_B \end{pmatrix}, \quad C = \begin{pmatrix} x_C \\ y_C \end{pmatrix}$$

passa un'unica circonferenza, di equazione

$$\det \begin{pmatrix} x^2 + y^2 & x & y & 1 \\ x_A^2 + y_A^2 & x_A & y_A & 1 \\ x_B^2 + y_B^2 & x_B & y_B & 1 \\ x_C^2 + y_C^2 & x_C & y_C & 1 \end{pmatrix} = 0.$$

Dimostrazione. L'equazione generica di una circonferenza è

$$a(x^2 + y^2) + bx + cy + d = 0$$

con le condizioni $a \neq 0$ e $b^2 + c^2 > 4ad$. Determiniamo le variabili a, b, c, d imponendo il passaggio per A tramite l'equazione

$$a(x_A^2 + y_A^2) + bx_A + cy_A + d = 0$$

ed il passaggio per B e C con altre due equazioni simili. Otteniamo il sistema

$$\begin{pmatrix} x_A^2 + y_A^2 & x_A & y_A & 1 \\ x_B^2 + y_B^2 & x_B & y_B & 1 \\ x_C^2 + y_C^2 & x_C & y_C & 1 \end{pmatrix} \begin{pmatrix} a \\ b \\ c \\ d \end{pmatrix} = 0.$$

Sia d_i il determinante del minore 3×3 della matrice dei coefficienti ottenuto togliendo la i -esima colonna, per $i = 1, \dots, 4$. Abbiamo $d_1 \neq 0$ perché i punti non sono allineati (Esercizio 9.1). Per la Proposizione 3.4.15 le soluzioni sono

$$a = \lambda d_1, \quad b = -\lambda d_2, \quad c = \lambda d_3, \quad d = -\lambda d_4$$

al variare di $\lambda \in \mathbb{R}$. L'unica circonferenza passante per A, B, C ha equazione

$$(18) \quad d_1(x^2 + y^2) - d_2x + d_3y - d_4 = 0.$$

La condizione $b^2 + c^2 > 4ad$ è soddisfatta perché se non lo fosse l'equazione (18) avrebbe una o nessuna soluzione e noi già sappiamo che ne ha almeno tre A, B, C . L'equazione (18) può essere scritta come nell'enunciato, sviluppando il determinante lungo la prima riga. $\square$

È possibile usare questa equazione per ricavare le coordinate del centro della circonferenza e il raggio in funzione delle coordinate di A, B e C , ma le formule risultanti sono un po' complicate.

Studiamo adesso le intersezioni fra una circonferenza e una retta.

Proposizione 10.1.9. *Una circonferenza e una retta si possono intersecare in 0, 1 oppure 2 punti.*

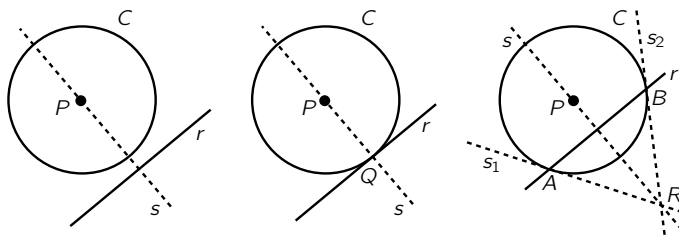


Figura 10.6. Una circonferenza C e una retta r si intersecano in 0, 1 o 2 punti. La retta s è ortogonale a r e passa per il centro P .

Dimostrazione. Una retta ha equazione del tipo $y = mx + n$ oppure $x = my + n$. Sostituendo la y (oppure la x) nell'equazione della circonferenza si ottiene un'equazione di secondo grado in x (oppure in y) che può avere 0, 1 oppure 2 soluzioni. $\square$

Si veda la Figura 10.6. Una retta che interseca una circonferenza in 1 o 2 punti è detta rispettivamente *tangente* e *secante*. Una *corda* è l'intersezione di un cerchio e di una retta secante. Se la corda contiene il centro del cerchio è detta *diametro*. Il centro divide il diametro in due segmenti detti *raggi*. Un raggio è un qualsiasi segmento con estremi nel centro e in un punto della circonferenza.

Sia C una circonferenza di centro P e raggio r .

Proposizione 10.1.10. *Per ogni $Q \in C$ esiste un'unica retta r tangente a C in Q , ed è quella ortogonale alla retta s passante per P e Q .*

Si veda la Figura 10.6-(centro).

Dimostrazione. Scegliamo un sistema di riferimento in cui l'origine è il centro di C e Q è il punto $\begin{pmatrix} 0 \\ r \end{pmatrix}$. Ora $C = \{x^2 + y^2 = r^2\}$ e $r = \left\{ \begin{pmatrix} 0 \\ r \end{pmatrix} + t \begin{pmatrix} a \\ b \end{pmatrix} \right\}$ è una retta generica passante per Q . Si veda la Figura 10.7-(sinistra). L'intersezione $C \cap r$ è data dai parametri $t \in \mathbb{R}$ che soddisfano l'equazione

$$(r + ta)^2 + (tb)^2 = r^2 \iff (a^2 + b^2)t^2 + 2art = 0.$$

Questa equazione in t ha sempre due soluzioni distinte, eccetto quando $a = 0$. Quindi r è tangente a $C \iff a = 0 \iff r$ è verticale, cioè ortogonale a s , che in questo sistema di riferimento è l'asse x . $\square$

Proposizione 10.1.11. *Per ogni punto R esterno alla circonferenza C esistono esattamente due rette s_1 e s_2 tangenti a C passanti per R .*

Dimostrazione. Scegliamo un sistema di riferimento in cui l'origine è il centro di C e R è il punto $\begin{pmatrix} c \\ 0 \end{pmatrix}$ con $c > r$. Ora $l = \left\{ \begin{pmatrix} c \\ 0 \end{pmatrix} + t \begin{pmatrix} a \\ b \end{pmatrix} \right\}$ è una retta generica passante per il punto R e la circonferenza C ha equazione $C = \{x^2 + y^2 = r^2\}$. Si veda la Figura 10.7-(destra). Possiamo supporre

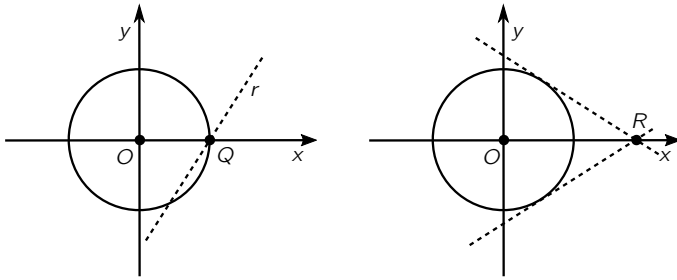


Figura 10.7. Dimostrazione delle Proposizioni 10.1.10 e 10.1.11.

che il vettore direttore di l sia unitario, cioè $a^2 + b^2 = 1$. L'intersezione $C \cap l$ è data dai parametri $t \in \mathbb{R}$ che soddisfano l'equazione

$$(c + ta)^2 + (tb)^2 = r^2 \iff t^2 + 2act + c^2 - r^2 = 0.$$

Questa equazione in t ha una sola soluzione precisamente quando

$$\frac{\Delta}{2} = a^2 c^2 - c^2 + r^2 = 0 \iff a = \pm \frac{\sqrt{c^2 - r^2}}{c}.$$

Queste determinano le due rette tangenti passanti per R . □

Passiamo a studiare le circonferenze a meno di similitudine e congruenza.

Proposizione 10.1.12. *Tutte le circonferenze del piano sono simili. Due circonferenze sono congruenti $\iff$ hanno lo stesso raggio.*

Dimostrazione. Date due circonferenze C e C' con centri in P e P' e con raggi r e r' , è possibile spostare C in C' con la similitudine

$$f(x) = \frac{r'}{r}x + P' - \frac{r'}{r}P.$$

Notiamo infatti che $f(P) = P'$ e inoltre f ha coefficiente di dilatazione $\frac{r'}{r}$, quindi manda i punti a distanza r da P nei punti a distanza r' da P' .

Se $r = r'$, la f appena definita è una isometria. Se $r \neq r'$, non esiste nessuna isometria che sposti C in C' perché le isometrie preservano le distanze e quindi le lunghezze dei raggi. □

10.1.6. Teoremi sui triangoli. Dimostriamo alcuni noti teoremi sui triangoli. Dato un triangolo con vertici A, B, C , indichiamo con α, β, γ i suoi angoli interni, come nella Figura 10.8.

Teorema 10.1.13 (Teorema del coseno). *Vale*

$$|BC|^2 = |AB|^2 + |AC|^2 - 2|AB| \cdot |AC| \cos \alpha.$$

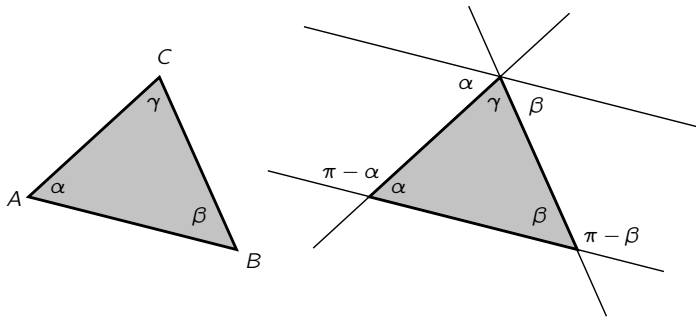


Figura 10.8. Un triangolo (sinistra). Dimostrazione che $\alpha + \beta + \gamma = \pi$ (destra).

Dimostrazione. Con il prodotto scalare troviamo

$$\begin{aligned} |BC|^2 &= \langle \vec{BC}, \vec{BC} \rangle = \langle \vec{AC} - \vec{AB}, \vec{AC} - \vec{AB} \rangle = |\vec{AB}|^2 + |\vec{AC}|^2 - 2\langle \vec{AB}, \vec{AC} \rangle \\ &= |\vec{AB}|^2 + |\vec{AC}|^2 - 2|\vec{AB}| \cdot |\vec{AC}| \cos \alpha. \end{aligned}$$

La dimostrazione è conclusa. □

Corollario 10.1.14 (Teorema di Pitagora). Se $\alpha = \frac{\pi}{2}$ vale

$$|BC|^2 = |\vec{AB}|^2 + |\vec{AC}|^2.$$

Scriviamo per semplicità

$$a = |BC|, \quad b = |CA|, \quad c = |AB|, \quad p = \frac{a+b+c}{2}.$$

La formula seguente è utile in alcuni casi perché esprime l'area del triangolo usando solo le lunghezze dei suoi lati.

Teorema 10.1.15 (Formula di Erone). Vale

$$\text{Area}(T) = \sqrt{p(p-a)(p-b)(p-c)}.$$

Dimostrazione. Per il Teorema del coseno abbiamo

$$\cos \alpha = \frac{b^2 + c^2 - a^2}{2bc}.$$

Quindi

$$\begin{aligned} \text{Area}(T) &= \frac{1}{2}bc \sin \alpha = \frac{1}{2}bc \sqrt{1 - \cos^2 \alpha} \\ &= \frac{1}{2}bc \sqrt{\frac{4b^2c^2 - (b^2 + c^2 - a^2)^2}{4b^2c^2}} \\ &= \frac{1}{4} \sqrt{4b^2c^2 - (b^2 + c^2 - a^2)^2} = \sqrt{p(p-a)(p-b)(p-c)}. \end{aligned}$$

L'ultima uguaglianza è lasciata per esercizio. □

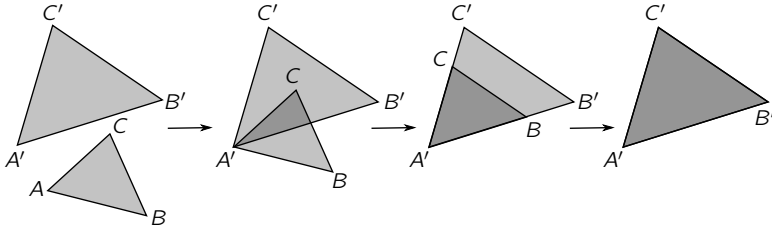


Figura 10.9. Due triangoli con gli stessi angoli interni sono simili.

Teorema 10.1.16 (Teorema dei seni). *Valgono le uguaglianze*

$$\frac{a}{\sin \alpha} = \frac{b}{\sin \beta} = \frac{c}{\sin \gamma} = \frac{abc}{2\text{Area}(T)}.$$

Dimostrazione. Abbiamo $\text{Area}(T) = \frac{1}{2}ab\sin\gamma$. Moltiplicando per c e scrivendo questa uguaglianza con a, b, c permutati si ottiene la tesi. $\square$

Abbiamo già dimostrato la *disuguaglianza triangolare*:

$$|AB| \leq |BC| + |CA|.$$

Mostriamo un altro fatto ben noto.

Proposizione 10.1.17. *La somma degli angoli interni è $\alpha + \beta + \gamma = \pi$.*

Dimostrazione. Come nella Figura 10.8, disegniamo la retta contenente AB e la sua parallela passante per C . Dalla figura segue che $\alpha + \beta + \gamma$ è un angolo piatto. $\square$

Un triangolo T è

- *equilatero* se $a = b = c$,
- *isoscele* se (a meno di permutare i lati) $a = b \neq c$,
- *scaleno* se a, b, c sono numeri distinti.

Dal teorema dei seni, a lunghezze di lati uguali corrispondono angoli uguali. Quindi in un triangolo equilatero abbiamo $\alpha = \beta = \gamma$, in un isoscele $\alpha = \beta \neq \gamma$ e in uno scaleno gli angoli α, β, γ sono tutti distinti.

Inoltre un triangolo è

- *acutangolo* se ha tutti gli angoli interni acuti,
- *rettangolo* se ha un angolo interno retto,
- *ottusangolo* se ha un angolo interno ottuso.

Verifichiamo adesso che gli angoli interni α, β, γ caratterizzano completamente il triangolo a meno di similitudine.

Proposizione 10.1.18. *Due triangoli sono simili $\iff$ hanno gli stessi angoli interni α, β, γ .*

Dimostrazione. ($\Rightarrow$) Le similitudini preservano gli angoli, quindi due triangoli simili hanno gli stessi angoli interni.

($\Leftarrow$) Siano T e T' due triangoli di vertici A, B, C e A', B', C' , entrambi con angoli interni α, β e γ . Costruiamo nella Figura 10.9 in alcuni passi una similitudine f che sposti T in T' .

La figura spiega già tutto; in ogni caso, questi sono i dettagli. Supponiamo che i vertici A', B', C' di T' siano ordinati in senso antiorario. Come prima cosa, se i vertici A, B, C non sono anche loro ordinati in senso antiorario, specchiamo T con una riflessione rispetto ad una retta qualsiasi in modo che lo diventino. Quindi con una traslazione spostiamo A su A' e poi con una rotazione intorno ad A facciamo in modo che i segmenti AB e AB' siano sovrapposti. Infine, con una omotetia di centro A trasformiamo T in modo che B si sovrapponga a B' . Siccome gli angoli in $A = A'$ e $B = B'$ dei due triangoli sono gli stessi, i due triangoli così costruiti coincidono, come mostrato in figura. $\square$

D'altro canto, le lunghezze dei lati caratterizzano completamente il triangolo a meno di congruenza.

Proposizione 10.1.19. *Due triangoli sono congruenti $\iff$ i lati hanno le stesse lunghezze a, b, c .*

Dimostrazione. Per la Formula di Erone i due triangoli T e T' hanno la stessa area e per il Teorema dei seni hanno anche gli stessi angoli interni. Quindi T e T' sono simili. La congruenza che sposta T in T' è costruita come nella dimostrazione precedente, senza usare l'omotetia. $\square$

Siano T e T' due triangoli di vertici A, B, C e A', B', C' , con angoli interni relativi α, β, γ e α', β', γ' .

Esercizio 10.1.20 (Criteri di congruenza). I due triangoli T e T' sono congruenti se vale una qualsiasi delle condizioni seguenti:

- (1) $|AB| = |A'B'|$, $|AC| = |A'C'|$ e $\alpha = \alpha'$,
- (2) $|AB| = |A'B'|$, $\alpha = \alpha'$ e $\beta = \beta'$,
- (3) $|AB| = |A'B'|$, $|AC| = |A'C'|$ e $|BC| = |B'C'|$.

Diamo una formula per l'area del triangolo T in $\mathbb{R}^2$ con vertici

$$A = \begin{pmatrix} x_A \\ y_A \end{pmatrix}, \quad B = \begin{pmatrix} x_B \\ y_B \end{pmatrix}, \quad C = \begin{pmatrix} x_C \\ y_C \end{pmatrix}.$$

Proposizione 10.1.21 (Area del triangolo). *Vale*

$$\text{Area}(T) = \frac{1}{2} \left| \det \begin{pmatrix} x_A & y_A & 1 \\ x_B & y_B & 1 \\ x_C & y_C & 1 \end{pmatrix} \right|.$$

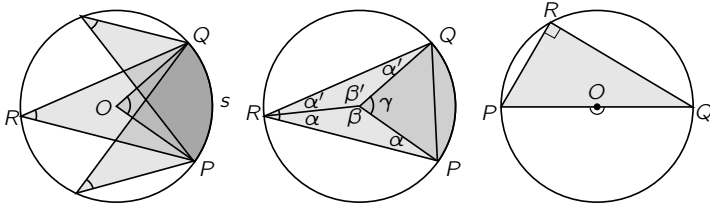


Figura 10.10. Un angolo al centro e alcuni angoli alla circonferenza che insistono sullo stesso arco s delimitato dai punti P e Q (sinistra). Dimostrazione del fatto che l'angolo al centro ha ampiezza doppia di quello alla circonferenza (centro). Un triangolo con vertici sulla circonferenza con il diametro come lato ha l'angolo opposto retto (destra).

Dimostrazione. Con un paio di mosse di Gauss sulle righe troviamo

$$d = \det \begin{pmatrix} x_A & y_A & 1 \\ x_B & y_B & 1 \\ x_C & y_C & 1 \end{pmatrix} = \det \begin{pmatrix} x_B - x_A & y_B - y_A \\ x_C - x_A & y_C - y_A \end{pmatrix}.$$

Considerando il piano dentro $\mathbb{R}^3$, troviamo

$$\overrightarrow{BA} = \begin{pmatrix} x_B - x_A \\ y_B - y_A \\ 0 \end{pmatrix}, \quad \overrightarrow{CA} = \begin{pmatrix} x_C - x_A \\ y_C - y_A \\ 0 \end{pmatrix} \implies \overrightarrow{BA} \times \overrightarrow{CA} = \begin{pmatrix} 0 \\ 0 \\ d \end{pmatrix}.$$

Dal Corollario 9.1.6 otteniamo $\text{Area}(T) = \frac{1}{2}|d|$. □

La formula è coerente con l'Esercizio 9.1, che sostiene che i tre punti A, B, C sono allineati precisamente quando questo determinante si annulla. Notiamo infine il fatto seguente.

Proposizione 10.1.22. *Tutti i triangoli sono affinemente equivalenti.*

Dimostrazione. Consideriamo due triangoli T, T' con vertici A, B, C e A', B', C' . Per la Proposizione 9.3.5 esiste una affinità che manda la terna A, B, C in A', B', C' , e quindi T in T' . □

10.1.7. Ancora sulle circonferenze. Usiamo quanto abbiamo visto con i triangoli per dimostrare alcuni fatti noti sulle circonferenze.

Sia C una circonferenza di centro O . Un *angolo al centro* è un qualsiasi angolo A con centro O . Un *arco di circonferenza* è l'intersezione $s = C \cap A$ fra C e un angolo al centro A . Si tratta di una porzione di circonferenza delimitata da due punti P e Q , si veda la Figura 10.10-(sinistra).

Un *angolo alla circonferenza* è un angolo A con centro un punto R della circonferenza C , delimitato da due semirette secanti, o da una secante e una tangente. Si veda la Figura 10.10-(sinistra).

Se A è un angolo al centro o alla circonferenza, l'intersezione $s = A \cap C$ è un arco di circonferenza e diciamo che A *insiste* su s .

Proposizione 10.1.23. *Siano A e A' un angolo al centro e alla circonferenza che insistono sullo stesso arco s . L'ampiezza di A' è doppia di quella di A .*

Dimostrazione. Il caso in cui l'angolo al centro sia convesso e interamente contenuto nell'angolo alla circonferenza è mostrato nella Figura 10.10-(centro). I due triangoli chiari sono isosceli (due loro lati sono raggi), quindi gli angoli α e α' si ripetono come mostrato. Abbiamo

$$2\alpha + \beta = 2\alpha' + \beta' = \pi, \quad \beta + \beta' + \gamma = 2\pi$$

e questo implica facilmente che $\gamma = 2(\alpha + \alpha')$. Gli altri casi, in cui l'angolo al centro è concavo oppure non è contenuto in quello alla circonferenza, sono lasciati per esercizio. $\square$

Questa proposizione ha come conseguenza che gli angoli alla circonferenza che insistono sullo stesso arco come nella Figura 10.10-(sinistra) sono tutti congruenti. Un altro corollario è il fatto che un triangolo PQR come nella Figura 10.10-(destra), con i vertici in C e avente il diametro PQ come lato, è sempre rettangolo (perché l'angolo in R è metà dell'angolo piatto in O).

10.1.8. Punti notevoli del triangolo. Il triangolo è una figura apparentemente molto semplice, ma in realtà foriera di un'enorme quantità di costruzioni geometriche raffinate. Se da una parte la circonferenza contiene essenzialmente un solo punto privilegiato, il suo centro, dall'altra il triangolo ne contiene una miriade, ciascuno con le sue proprietà: questi punti sono detti *punti notevoli*. In questa sezione ne vediamo solo quattro, generalmente già studiati nelle scuole superiori.

Definizione 10.1.24. Sia T un triangolo di vertici A, B e C .

- La *mediana* di vertice A è il segmento AM con M punto medio del lato opposto BC .
- L'*altezza* di vertice A è il segmento AH ortogonale alla retta s che contiene BC , con $H \in s$.
- La *bisettrice* di vertice A è il segmento AN con $N \in BC$ tale che gli angoli $\hat{B}AN$ e $\hat{N}AC$ siano congruenti.
- L'*asse* del lato BC è la retta r ortogonale a BC passante per il punto medio M di BC .

La Figura 10.11 mostra la mediana, l'altezza, la bisettrice e l'asse relativi al lato BC . Se il triangolo è isoscele con $|AB| = |AC|$, si vede facilmente che $H = N = M$ e quindi mediana, altezza e bisettrice coincidono, e l'asse è la retta che le contiene.

Come accennato, definiamo adesso quattro punti notevoli del triangolo: il *baricentro*, l'*incentro*, il *circocentro* e l'*ortocentro*. Ciascuno di questi risulta essere il punto di incontro di mediane, bisettrici, assi e altezze.

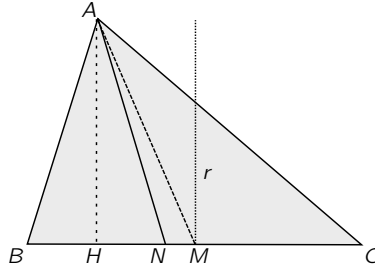


Figura 10.11. La mediana AM , l'altezza AH , la bisettrice AN e l'asse r .

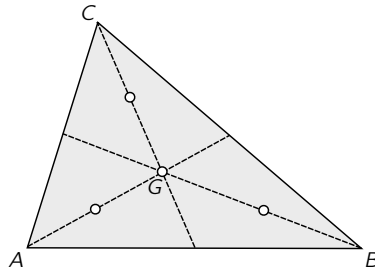


Figura 10.12. Il baricentro G di un triangolo taglia ciascuna mediana in due segmenti, uno lungo il doppio dell'altro.

In tutta questa sezione T è un triangolo di vertici A , B e C . Iniziamo definendo il *baricentro* come il punto

$$G = \frac{1}{3}(A + B + C).$$

Si veda un esempio nella Figura 10.12.

Proposizione 10.1.25. *Le tre mediane di T si intersecano nel baricentro. Il baricentro taglia ciascuna mediana in due segmenti, uno doppio dell'altro.*

Dimostrazione. Mostriamo che il baricentro $G = \frac{1}{3}(A + B + C)$ è contenuto in ciascuna mediana. Consideriamo ad esempio la mediana AM dove $M = \frac{1}{2}(B + C)$ è il punto medio di BC . Troviamo effettivamente

$$G = \frac{1}{3}(A + B + C) = A + \frac{2}{3} \left(\frac{B + C}{2} - A \right) = A + \frac{2}{3} \overrightarrow{AM}.$$

Questo ci dice che G appartiene alla mediana, e che $|AG| = 2|GM|$ perché il baricentro G sta esattamente a $2/3$ del percorso da A a M . $\square$

Da un punto di vista fisico, il baricentro è il centro di massa di un sistema formato da tre palline di uguale peso posizionate nei punti A , B e C . L'esercizio seguente è un semplice conto.

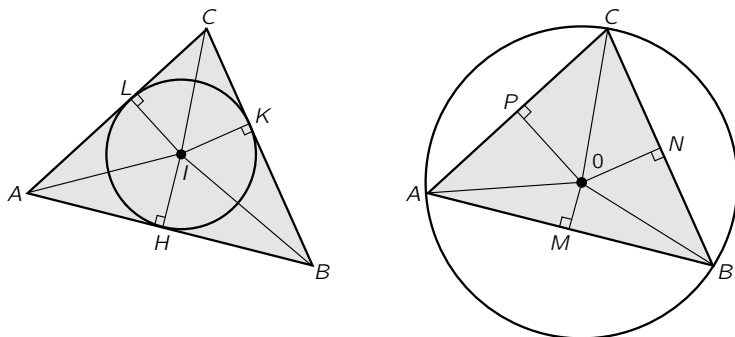


Figura 10.13. L'incentro è centro del cerchio inscritto nel triangolo (sinistra). Il circocentro è centro del cerchio circoscritto (destra).

Esercizio 10.1.26. Vale la relazione

$$\vec{GA} + \vec{GB} + \vec{GC} = 0.$$

Passiamo ad un altro punto notevole. Ricordiamo che a, b, c indicano le lunghezze dei lati BC, CA e AB . Definiamo l'*incentro* come il punto

$$I = \frac{aA + bB + cC}{a + b + c}.$$

Un cerchio contenuto nel triangolo è *inscritto* se è tangente a tutti i lati.

Proposizione 10.1.27. *Le tre bisettrici di T si intersecano nell'incentro. L'incentro è il centro dell'unico cerchio inscritto nel triangolo.*

Dimostrazione. Mostriamo che I è contenuto in ciascuna bisettrice. Scriviamo la bisettrice uscente da A . La bisettrice di due vettori *unitari* v e u è sempre $\frac{1}{2}(v + u)$. Quindi la bisettrice uscente da A ha equazione parametrica

$$A + t \left(\frac{\vec{AB}}{|\vec{AB}|} + \frac{\vec{AC}}{|\vec{AC}|} \right).$$

Per $t = \frac{bc}{a+b+c}$ otteniamo proprio l'incentro:

$$\begin{aligned} A + \frac{bc}{a+b+c} \left(\frac{\vec{AB}}{c} + \frac{\vec{AC}}{b} \right) &= A + \frac{b(B-A) + c(C-A)}{a+b+c} \\ &= \frac{aA + bB + cC}{a+b+c} = I. \end{aligned}$$

Analogamente I è contenuto nelle bisettrici uscenti da B e da C . Mostriamo adesso geometricamente che I è il centro dell'unico cerchio inscritto in T .

Disegniamo nella Figura 10.13 l'incentro I e i tre segmenti perpendicolari uscenti da I sui lati del triangolo. Per costruzione gli angoli $\hat{H}\hat{A}I$ e

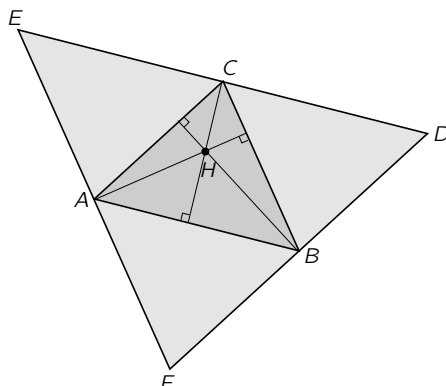


Figura 10.14. Le tre altezze di ABC sono porzioni degli assi del triangolo più grande DEF e quindi si intersecano in un punto H , che è il circocentro di DEF e l'ortocentro di ABC .

$\hat{L}\hat{A}\hat{I}$ sono congruenti, quindi anche i triangoli relativi HAI e LAI lo sono per l'Esercizio 10.1.20-(2). Ne segue che $|IH| = |IL|$ e lavorando con B troviamo anche $|IH| = |IK|$. Quindi la circonferenza di centro I e raggio $|IH|$ è contenuta nel triangolo e tangente ai suoi lati.

La circonferenza iscritta è unica: sia C' un'altra circonferenza inscritta di centro I' , tangente ai lati in alcuni punti H' , K' e L' . Abbiamo $|H'I| = |L'I|$. Usando il Teorema di Pitagora troviamo $|AH'| = |AL'|$. Per l'Esercizio 10.1.20-(3) i triangoli $H'A'I'$ e $L'A'I'$ sono congruenti e quindi $A'I'$ è la bisettrice di $\hat{B}\hat{A}\hat{C}$. Ragionando analogamente con B e C troviamo che I' è contenuto nelle tre bisettrici, quindi $I' = I$ e allora $C' = C$. $\square$

Un cerchio è *circoscritto* al triangolo se la circonferenza ne contiene i vertici. Sappiamo dalla Proposizione 10.1.8 che esiste un unico cerchio circoscritto al triangolo e chiamiamo *circocentro* il suo centro O .

Proposizione 10.1.28. *Le tre assi di T si intersecano nel circocentro.*

Dimostrazione. Disegniamo il cerchio circoscritto come nella Figura 10.13-(destra). Tracciamo dal centro O le perpendicolari sui tre lati. Notiamo che AOB è isoscele e quindi $|AM| = |MB|$. Allora OM è un asse e contiene O . Analogamente ON e OP sono assi che contengono O .

In alternativa, con argomenti simili si può verificare direttamente che i tre assi si intersecano in un punto e quindi ridimostrare l'esistenza e unicità del cerchio circoscritto, senza usare la Proposizione 10.1.8. $\square$

Passiamo infine a considerare le altezze.

Proposizione 10.1.29. *Le tre altezze di T si intersecano in un punto H detto ortocentro.*

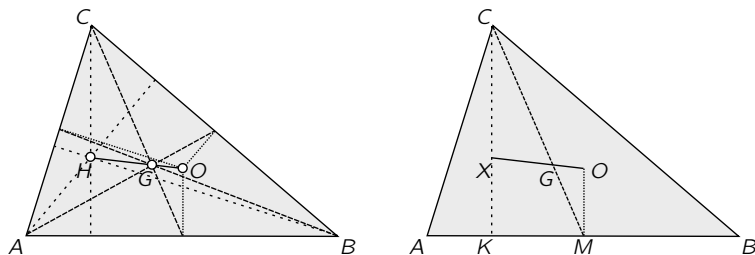


Figura 10.15. La linea di Eulero contiene l'ortocentro H , il baricentro G e il circocentro O . Vale $|HG| = 2|GO|$ (sinistra). Dimostrazione dell'esistenza della linea di Eulero (destra).

Dimostrazione. Come mostrato nella Figura 10.14, tracciamo le parallele dei tre lati dai vertici opposti per costruire un triangolo DEF più grande. Le tre altezze di ABC sono porzioni degli assi di DEF e quindi si incontrano in un punto H , che è il circocentro di DEF e l'ortocentro di ABC . $\square$

Abbiamo definito quattro punti notevoli: il baricentro G , l'incentro I , il circocentro O e l'ortocentro H . Se il triangolo è equilatero, questi coincidono tutti. Il teorema seguente è stato dimostrato da Eulero nel 1765.

Teorema 10.1.30 (Linea di Eulero). *Se il triangolo non è equilatero, il circocentro O , il baricentro G e l'ortocentro H sono distinti e allineati, in quest'ordine. Inoltre*

$$|HG| = 2|GO|.$$

Si veda la Figura 10.15-(sinistra).

Dimostrazione. Consideriamo il baricentro G e l'ortocentro O . Lasciamo come esercizio il fatto che se $G = O$ allora il triangolo è equilatero; quindi supponiamo $G \neq O$.

Come nella Figura 10.15-(destra), identifichiamo un punto X tracciando il segmento OG e proseguendo finché non otteniamo

$$(19) \quad |XG| = 2|GO|.$$

Il nostro scopo è mostrare che X è l'ortocentro del triangolo. Tracciamo la mediana CM , che passa per G , e l'asse MO relativa al lato AB come in figura. Ricordiamo dalla Proposizione 10.1.25 che

$$(20) \quad |CG| = 2|GM|.$$

Le relazioni (19) e (20) implicano che i triangoli CXG e MOG sono simili. Quindi OM e CX sono paralleli, e allora CX è perpendicolare ad AB . Quindi X è contenuto nell'altezza CK relativa ad AB . Lavorando

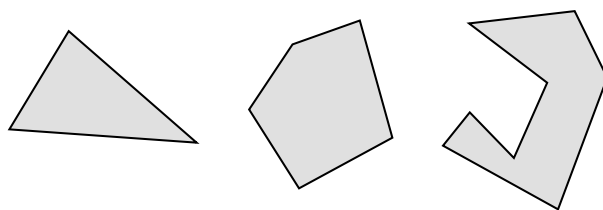


Figura 10.16. Alcuni poligoni. I primi due sono convessi, il terzo no.

analogamente con tutti e tre i lati troviamo che X è contenuto in tutte e tre le altezze e quindi $X = H$ è l'ortocentro. $\square$

La linea che contiene O , G e H è detta *linea di Eulero*.

Corollario 10.1.31 (Problema di Sylvester). *Per qualsiasi triangolo vale*

$$\overrightarrow{OH} = \overrightarrow{OA} + \overrightarrow{OB} + \overrightarrow{OC}.$$

Dimostrazione. Abbiamo

$$\begin{aligned}\overrightarrow{OH} &= 3\overrightarrow{OG} = \overrightarrow{OG} + \overrightarrow{OG} + \overrightarrow{OG} \\ &= \overrightarrow{OG} + \overrightarrow{GA} + \overrightarrow{OG} + \overrightarrow{GB} + \overrightarrow{OG} + \overrightarrow{GC} = \overrightarrow{OA} + \overrightarrow{OB} + \overrightarrow{OC}.\end{aligned}$$

Nella terza uguaglianza abbiamo usato l'Esercizio 10.1.26. $\square$

Una conseguenza di questo corollario è la seguente: se disegniamo una circonferenza di centro l'origine O di $\mathbb{R}^2$ e prendiamo tre punti A, B, C su di essa, l'ortocentro H del triangolo con vertici A, B, C si ottiene semplicemente facendo la somma vettoriale $H = A + B + C$ dei vertici.

10.1.9. Poligoni. I triangoli sono elementi importanti di una classe più ampia di oggetti geometrici, i *poligoni*.

Definizione 10.1.32. Un *poligono* P è un sottoinsieme di $\mathbb{R}^2$ delimitato da una linea poligonale chiusa non intrecciata.

Alcuni esempi sono mostrati nella Figura 10.16. Questa definizione sfrutta implicitamente un teorema profondo, il *Teorema della curva di Jordan*, che garantisce che una linea chiusa non intrecciata effettivamente separi il piano $\mathbb{R}^2$ in due zone, una limitata interna e l'altra illimitata esterna. Questo teorema può sembrare ovvio per linee chiuse semplici come quelle nella Figura 10.16, ma non lo è se la linea chiusa è più complessa, come nella Figura 10.17: prima di guardare la parte destra della figura, riesci ad identificare con i tuoi occhi la zona interna e quella esterna?

La linea spezzata che definisce il poligono è il suo *contorno* e la sua lunghezza è il *perimetro*. Ciascuno spigolo della linea è un *lato* del poligono

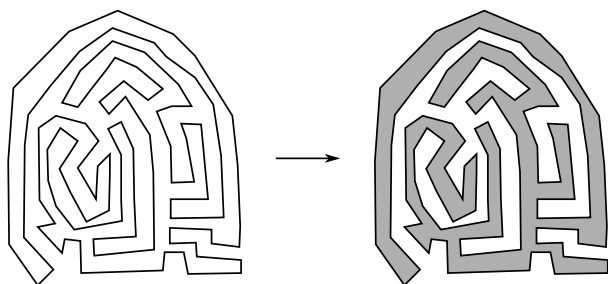


Figura 10.17. Per il Teorema della curva di Jordan, qualsiasi curva poligonale chiusa non intrecciata divide il piano in due zone: una interna (colorata di grigio a destra) e una esterna.

e ciascun estremo è un *vertice*. Un poligono P ha lo stesso numero $n \geq 3$ di lati e di vertici.

Ciascun vertice di P ha un *angolo interno*, delimitato dai due spigoli adiacenti e contenente la parte del poligono vicina al vertice. Il poligono è *convesso* se tutti i suoi angoli interni sono convessi o piatti. Si veda la Figura 10.16.

Una *diagonale* di P è un segmento che congiunge due vertici non consecutivi.

Esercizio 10.1.33. Un poligono P con $n \geq 3$ lati ha $\frac{n(n-3)}{2}$ diagonali.

In particolare un triangolo non ha diagonali e un quadrilatero ne ha due. Una diagonale è *interna* se è interamente contenuta in P . Notiamo che una diagonale può non essere interna: questo accade quando il poligono non è convesso, come nella Figura 10.16-(destra).

Esercizio 10.1.34. Un poligono P con $n \geq 4$ lati ha sempre almeno una diagonale interna.

Traccia. Dimostra intanto che P ha sempre almeno un vertice A con angolo interno convesso. Siano B e C i vertici adiacenti ad A . Se BC è una diagonale interna, siamo a posto. Altrimenti, il triangolo ABC contiene altri vertici di P . Fra questi, prendi quello D più vicino ad A e dimostra che AD è una diagonale interna. $\square$

Questo esercizio ha delle conseguenze interessanti.

Proposizione 10.1.35. *Un poligono P con n lati si decompone sempre in $n - 2$ triangoli lungo alcune sue diagonali.*

Si veda la Figura 10.18.

Dimostrazione. Dimostriamo la proposizione per induzione su $n \geq 3$. Per $n = 3$ il poligono P è già un triangolo e siamo a posto. Se $n > 3$, per

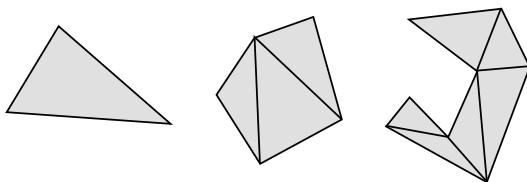


Figura 10.18. I poligoni della Figura 10.16 decomposti in triangoli lungo alcune diagonali.

l'esercizio precedente esiste sempre una diagonale interna. Spezzando P lungo la diagonale interna otteniamo due poligoni di $m \geq 3$ e $k \geq 3$ lati, con $m + k = n + 2$. Poiché $m < n$ e $k < n$, possiamo applicare l'ipotesi induttiva su entrambi i poligoni: questi si suddividono in triangoli, e quindi anche P . $\square$

Corollario 10.1.36. *La somma degli angoli interni in un poligono di n lati è $(n - 2)\pi$.*

Dimostrazione. Il poligono si decompone in $n - 2$ triangoli e ciascuno contribuisce con π . $\square$

Quindi la somma degli angoli interni di un quadrilatero è sempre 2π .

10.1.10. Sottoinsiemi convessi. La nozione di convessità in realtà non si applica solo ai poligoni: si tratta di un concetto geometrico molto più ampio valido per sottoinsiemi arbitrari di $\mathbb{R}^n$, cui accenniamo brevemente.

Definizione 10.1.37. Un sottoinsieme $S \subset \mathbb{R}^n$ è *convesso* se per ogni coppia di punti $P, Q \in S$ il segmento PQ è interamente contenuto in S .

Esempio 10.1.38. I sottospazi affini $S \subset \mathbb{R}^n$ sono convessi. I triangoli sono convessi. Un semipiano in $\mathbb{R}^2$ è sempre convesso.

Un aspetto fondamentale della teoria è che questa definizione si mantiene per intersezione:

Proposizione 10.1.39. *Se S_1 e S_2 sono convessi, anche $S_1 \cap S_2$ lo è.*

Dimostrazione. Presi due punti $A, B \in S_1 \cap S_2$, il segmento AB è contenuto in S_1 e S_2 perché sono convessi, quindi è contenuto nell'intersezione. $\square$

Abbiamo introdotto due nozioni diverse di convessità per i poligoni e dobbiamo dimostrare che coincidono. Lasciamo questo fatto per esercizio.

Esercizio 10.1.40. Un poligono $P \subset \mathbb{R}^2$ è convesso (nel senso della Definizione 10.1.37) se e solo se i suoi angoli interni sono piatti o convessi.

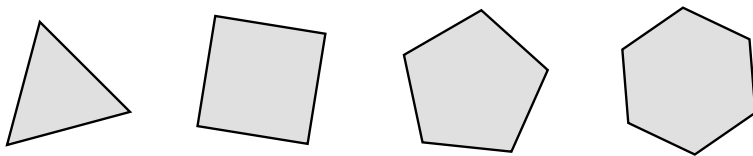


Figura 10.19. Alcuni poligoni regolari.

10.1.11. Poligoni regolari. Un poligono convesso con n lati è *regolare* se ha tutti gli angoli della stessa ampiezza e tutti i lati della stessa lunghezza. Alcuni esempi sono mostrati nella Figura 10.19.

Un poligono regolare con n lati ha tutti gli angoli interni pari a

$$\frac{n-2}{n}\pi.$$

Si possono usare i numeri complessi per descrivere convenientemente un poligono regolare. Ad esempio, i numeri complessi

$$e^{\frac{2k\pi i}{n}}$$

per $k = 0, 1, \dots, n-1$ descrivono i vertici di un poligono regolare con n lati centrato nell'origine e con tutti i vertici nella circonferenza unitaria.

10.2. Poliedri

I poliedri sono (assieme alle sfere) le figure geometriche più semplici che si possono definire e studiare in una porzione limitata di spazio. Sono l'analogo tridimensionale dei poligoni.

10.2.1. Definizione. La definizione di poliedro è simile a quella di poligono data precedentemente.

Definizione 10.2.1. Un *poliedro* è un sottoinsieme di $\mathbb{R}^3$ delimitato da un numero finito di poligoni che si intersecano a coppie lungo i loro lati.

Come un poligono è una porzione di piano delimitata da alcuni segmenti che si intersecano a coppie lungo i loro estremi, così il poliedro è una porzione di spazio delimitata da alcuni poligoni che si intersecano a coppie lungo i loro lati. Si suppone implicitamente che i poligoni si intersechino solo lungo i lati (niente intrecciamento) e che formino un blocco solo. Anche in questo caso esistono dei teoremi che garantiscono che effettivamente un poliedro suddivide sempre lo spazio in una zona interna limitata ed una zona esterna illimitata.

I poligoni che delimitano il poliedro P sono le *facce* di P . I lati delle facce sono gli *spigoli* di P e i vertici delle facce sono i *vertici* di P . Quindi un poliedro P ha facce, spigoli e vertici.

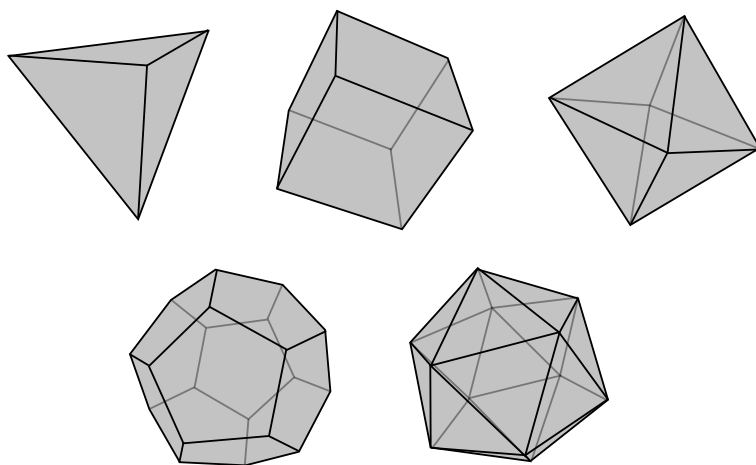


Figura 10.20. I cinque poliedri regolari, chiamati anche solidi platonici.

10.2.2. Poliedri regolari. Abbiamo definito precedentemente un poligono regolare come un poligono che ha tutti i lati e tutti gli angoli uguali. Adesso vorremmo fare la stessa cosa nello spazio: definire un *poliedro regolare* come un poliedro che ha vertici, spigoli e facce tutti uguali. Questa condizione può essere espressa rigorosamente in vari modi. Quello più semplice è il seguente.

Definizione 10.2.2. Un poliedro P è *regolare* se le sue facce sono poligoni regolari congruenti che si incontrano lo stesso numero di volte ad ogni vertice.

La combinatoria di un poliedro regolare è caratterizzata da due numeri, tradizionalmente indicati tra parentesi graffe come $\{p, q\}$. Questo simbolo indica un poliedro regolare in cui le facce sono tutte poligoni regolari con p lati, tale che q facce si incontrano ad ogni vertice. Il simbolo $\{p, q\}$ si chiama *simbolo di Schläfli*, dal matematico svizzero Schläfli che nel XIX secolo studiò i poliedri regolari in dimensione arbitraria. Qui $p, q \geq 3$.

Come è noto fin dall'antichità, ci sono precisamente cinque poliedri regolari a meno di similitudine. Questi sono il *tetraedro*, il *cubo*, l'*ottaedro*, il *dodecaedro* e l'*icosaedro* e sono mostrati nella Figura 10.20. La lettrice può convincersi guardando la figura che i loro simboli di Schläfli sono rispettivamente

$$\{3, 3\}, \quad \{4, 3\}, \quad \{3, 4\}, \quad \{5, 3\}, \quad \{3, 5\}.$$

Ad esempio il tetraedro corrisponde al simbolo $\{3, 3\}$ perché ha facce triangolari che si incontrano a tre su ogni vertice, il cubo è $\{4, 3\}$ perché ha facce quadrate che si incontrano a tre, mentre l'icosaedro ha simbolo $\{3, 5\}$

	tetraedro	cubo	ottaedro	dodecaedro	icosaedro
F	4	6	8	12	20
S	6	12	12	30	30
V	4	8	6	20	12

Tabella 10.1. Il numero di facce F , di spigoli S e di vertici V di ciascun solido platonico.

perché ha facce triangolari che si incontrano a cinque su ogni vertice. I cinque poliedri sono anche chiamati *solidi platonici*.

Ci si può convincere del fatto che questi cinque siano gli unici poliedri regolari tendando di costruirli con dei pezzi di carta: se ad esempio si prendono dei triangoli equilateri congruenti e si tenta di affiancarli lungo un vertice, ci si rende conto che è possibile mettere 3, 4, oppure 5 triangoli intorno ad un vertice, ma non 6: se ne mettiamo 6 il vertice si appiattisce. Continuando quindi ad assemblare triangoli mantenendo lo stesso numero ad ogni vertice (che è 3, 4 oppure 5) si finisce magicamente con un poliedro che si chiude perfettamente, che sarà il tetraedro, l'ottaedro o l'icosaedro a seconda che si sia scelto di affiancarne 3, 4 o 5 ad ogni vertice. Analogamente partendo da quadrati o pentagoni ci si accorge concretamente che non è possibile affiancarne più di 3, e procedendo in questo modo si ottiene un cubo o un dodecaedro. Infine, si verifica concretamente che non è possibile affiancare neppure 3 poligoni regolari con $p \geq 6$ lati lungo un vertice.

10.2.3. Relazione di Eulero. Un poliedro P ha un certo numero F di facce, un numero S di spigoli e un numero V di vertici. I numeri F , S e V per i cinque solidi platonici sono mostrati nella Tabella 10.1. Notiamo che per tutti e 5 i solidi vale la relazione

$$F - S + V = 2.$$

Questa formula, detta *relazione di Eulero*, è valida per una classe ben più ampia di poliedri.

Teorema 10.2.3 (Relazione di Eulero). *Per qualsiasi poliedro convesso P vale la formula seguente:*

$$F - S + V = 2.$$

Dimostrazione. Come nella Figura 10.21, togliamo una faccia da P , e poi successivamente togliamo tutte le facce una ad una, stando attenti a non spezzare mai in due l'oggetto durante il processo. Come suggerito in figura, alla fine rimaniamo con una faccia sola.

Analizziamo come cambia il numero intero $F - S + V$ ad ogni passaggio. Nel togliere la prima faccia, il numero F di facce scende di uno mentre spigoli e vertici restano inalterati: quindi il numero $F - S + V$ scende

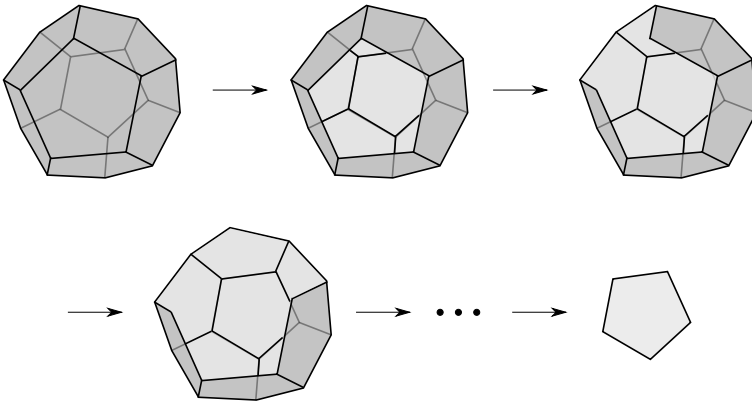


Figura 10.21. Dimostrazione della relazione di Eulero.

di uno. Al secondo passo scendono di uno sia F che S mentre V resta inalterato: quindi il numero $F - S + V$ non cambia. Si vede facilmente che il numero $F - S + V$ non cambia in nessun passo successivo. Alla fine troviamo un solo poligono con un certo numero n di lati; quindi alla fine $F - S + V = 1 - n + n = 1$. Poiché l'intero $F - S + V$ è cambiato solo al primo passo scendendo di uno, allora con i valori iniziali avevamo $F - S + V = 2$. $\square$

Osservazione 10.2.4. Usando la relazione di Eulero è possibile ridimostrare che non esistono poliedri regolari con simboli di Schläfli diversi dai $\{p, q\}$ elencati precedentemente. Consideriamo un poliedro regolare con simbolo di Schläfli $\{p, q\}$. Ogni faccia ha p spigoli: contiamo pF spigoli in tutto, però ogni spigolo è stato contato due volte perché adiacente a due facce; quindi

$$S = \frac{pF}{2} \implies F = \frac{2S}{p}.$$

Analogamente, ogni spigolo tocca due vertici e quindi contiamo $2S$ vertici, però ciascun vertice è stato contato q volte e quindi

$$V = \frac{2S}{q}.$$

Possiamo quindi esprimere F e V in funzione di S . Sostituendo nella relazione di Eulero troviamo

$$2 = F - S + V = \frac{2S}{p} - S + \frac{2S}{q} = S \left(\frac{2}{p} + \frac{2}{q} - 1 \right)$$

che implica

$$\frac{1}{S} = \frac{1}{p} + \frac{1}{q} - \frac{1}{2}.$$

Questo numero deve essere positivo. Ricordando che $p, q \geq 3$, questo capita solo per le coppie $\{p, q\}$ seguenti:

$$\{3, 3\}, \quad \{3, 4\}, \quad \{3, 5\}, \quad \{4, 3\}, \quad \{5, 3\}.$$

Questi sono esattamente i simboli dei 5 solidi platonici.

10.2.4. Sfere. Introduciamo adesso un oggetto geometrico fondamentale, l'analogo tridimensionale della circonferenza nel piano.

Definizione 10.2.5. La *sfera* in $\mathbb{R}^3$ di centro $P \in \mathbb{R}^3$ e di raggio $r > 0$ è l'insieme S dei punti che hanno distanza r da P .

La sfera S è definita dall'equazione:

$$(x - x_0)^2 + (y - y_0)^2 + (z - z_0)^2 = r^2$$

dove x_0, y_0, z_0 sono le coordinate del centro P . È possibile rappresentare in modo parametrico i punti della sfera usando le *coordinate sferiche*, cioè la longitudine $\varphi \in [0, 2\pi)$ e un parametro $\theta \in [0, \pi]$ analogo alla latitudine, già usate per il globo terrestre:

$$\begin{cases} x = x_0 + r \sin \theta \cos \varphi, \\ y = y_0 + r \sin \theta \sin \varphi, \\ z = z_0 + r \cos \theta. \end{cases}$$

I punti con $\theta = 0, \pi$ sono i poli: in questi la longitudine φ non è definita.

La proposizione seguente si dimostra in modo del tutto analogo alla Proposizione 10.1.8.

Proposizione 10.2.6. *Per quattro punti non complanari dello spazio*

$$A = \begin{pmatrix} x_A \\ y_A \\ z_A \end{pmatrix}, \quad B = \begin{pmatrix} x_B \\ y_B \\ z_B \end{pmatrix}, \quad C = \begin{pmatrix} x_C \\ y_C \\ z_C \end{pmatrix}, \quad D = \begin{pmatrix} x_D \\ y_D \\ z_D \end{pmatrix}$$

passa un'unica sfera, di equazione

$$\det \begin{pmatrix} x^2 + y^2 + z^2 & x & y & z & 1 \\ x_A^2 + y_A^2 + z_A^2 & x_A & y_A & z_A & 1 \\ x_B^2 + y_B^2 + z_B^2 & x_B & y_B & z_B & 1 \\ x_C^2 + y_C^2 + z_C^2 & x_C & y_C & z_C & 1 \\ x_D^2 + y_D^2 + z_D^2 & x_D & y_D & z_D & 1 \end{pmatrix} = 0.$$

La *sfera solida* di centro P e raggio $r > 0$ è l'insieme Z dei punti che hanno distanza $\leq r$ da P , cioè dei punti che soddisfano la disequazione

$$(x - x_0)^2 + (y - y_0)^2 + (z - z_0)^2 \leq r^2.$$

La sfera solida è la parte interna della sfera di centro P e raggio r .

Come per le circonferenze, si dimostra che una sfera S ed una retta r possono intersecarsi in 0, 1 oppure 2 punti. Nel caso si intersechino in un punto, diciamo che r è *tangente* ad S . Nel caso in cui la retta r passi per il centro di S , la sua intersezione con la sfera solida delimitata da S è un

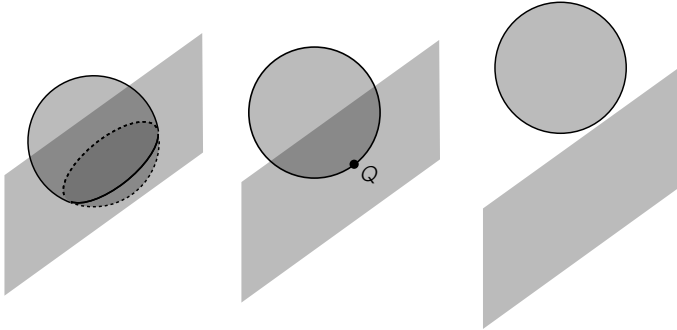


Figura 10.22. Una sfera ed un piano possono intersecarsi in una circonferenza, in un punto di tangenza Q , o da nessuna parte.

diametro di S , tagliato in due *raggi* dal centro. Un raggio è un qualsiasi segmento che congiunge il centro di S con un punto di S .

Un piano π può intersecare una sfera in una circonferenza, in un punto, oppure in nessun punto, si veda la Figura 10.22. Nel caso in cui π intersechi S in un punto Q , diciamo che π è *tangente* ad S in Q . Come per la circonferenza, il piano π tangente ad S in un punto $Q \in S$ è precisamente il piano ortogonale al raggio PQ .

Esempio 10.2.7. Si consideri la sfera S di centro P con coordinate $(0, 1, -1)$ e di raggio 3. Questa ha equazione $x^2 + (y - 1)^2 + (z + 1)^2 = 9$. Il punto Q di coordinate $(1, 3, 1)$ appartiene alla sfera. Il piano π tangente ad S nel punto Q è ortogonale al vettore

$$\overrightarrow{PQ} = Q - P = \begin{pmatrix} 1 \\ 3 \\ 1 \end{pmatrix} - \begin{pmatrix} 0 \\ 1 \\ -1 \end{pmatrix} = \begin{pmatrix} 1 \\ 2 \\ 2 \end{pmatrix}$$

e deve passare per Q . Queste due condizioni lo determinano:

$$\pi = \{x + 2y + 2z = 9\}.$$

Le rette tangenti a S nel punto Q sono precisamente tutte le rette in π che passano per Q .

10.2.5. Tetraedri e parallelepipedi. Così come la sfera è l'analogo tridimensionale della circonferenza, possiamo dire che il tetraedro è l'analogo tridimensionale del triangolo. Un *tetraedro* è un poliedro con quattro facce. Equivalentemente, può essere definito nel modo seguente:

Definizione 10.2.8. Il *tetraedro* con vertici A, B, C e D è l'insieme

$$T = \{A + s\overrightarrow{AB} + t\overrightarrow{AC} + u\overrightarrow{AD} \mid s, t, u \in [0, 1], s + t + u \leq 1\} \subset \mathbb{R}^3.$$

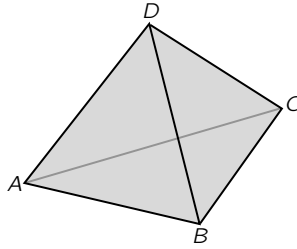


Figura 10.23. Un tetraedro con vertici A, B, C, D .

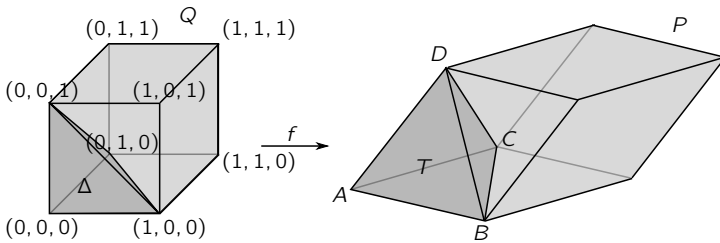


Figura 10.24. L'isomorfismo affine f manda il tetraedro Δ nel tetraedro T ed il cubo Q nel parallelepipedo P .

Si veda la Figura 10.23. Come nella Sezione 10.1.3, possiamo capire meglio questa scrittura notando che $T = f(\Delta)$ dove

$$\Delta = \{(s, t, u) \in \mathbb{R}^3 \mid s, t, u \in [0, 1], s + t + u \leq 1\}$$

è un tetraedro con vertici $(0, 0, 0)$, $(1, 0, 0)$, $(0, 1, 0)$, $(0, 0, 1)$ e $f: \mathbb{R}^3 \rightarrow \mathbb{R}^3$ è l'isomorfismo affine

$$f \begin{pmatrix} s \\ t \\ u \end{pmatrix} = A + s\overrightarrow{AB} + t\overrightarrow{AC} + u\overrightarrow{AD}$$

che manda i vertici di Δ nei vertici di T . La trasformazione affine f è descritta nella Figura 10.24.

Un tetraedro ha 4 facce triangolari ABC , ABD , ACD e BCD , 6 spigoli AB , AC , AD , BC , BD e CD e 4 vertici A, B, C, D .

Analogamente definiamo il *parallelepipedo* P con lati AB, AC e AD come il sottoinsieme

$$P = \{A + s\overrightarrow{AB} + t\overrightarrow{AC} + u\overrightarrow{AD} \mid s, t, u \in [0, 1]\}.$$

Come sopra, notiamo che $P = f(Q)$ dove Q è il cubo

$$Q = \{(s, t, u) \in \mathbb{R}^3 \mid s, t, u \in [0, 1]\}.$$

Si veda sempre la Figura 10.24. Il parallelepipedo ha come facce 6 parallelogrammi; ha 12 spigoli e 8 vertici. Per la Proposizione 9.1.9, ha

volume

$$\text{Vol}(P) = |\det(\overrightarrow{AB}|\overrightarrow{AC}|\overrightarrow{AD})|.$$

Proposizione 10.2.9. *Il volume del tetraedro T è*

$$\text{Vol}(T) = \frac{1}{6} |\det(\overrightarrow{AB}|\overrightarrow{AC}|\overrightarrow{AD})|.$$

Dimostrazione. Usiamo la formula “area di base $\times$ altezza diviso 3”, notando che l’area di base è metà di quella del parallelepipedo e l’altezza è la stessa: quindi viene un sesto del volume del parallelepipedo. $\square$

Possiamo scrivere una formula ancora più esplicita. Siano

$$A = \begin{pmatrix} x_A \\ y_A \\ z_A \end{pmatrix}, \quad B = \begin{pmatrix} x_B \\ y_B \\ z_B \end{pmatrix}, \quad C = \begin{pmatrix} x_C \\ y_C \\ z_C \end{pmatrix}, \quad D = \begin{pmatrix} x_D \\ y_D \\ z_D \end{pmatrix}.$$

La formula seguente è del tutto analoga alla Proposizione 10.1.21.

Corollario 10.2.10. *Vale*

$$\text{Vol}(T) = \frac{1}{6} \left| \det \begin{pmatrix} x_A & y_A & z_A & 1 \\ x_B & y_B & z_B & 1 \\ x_C & y_C & z_C & 1 \\ x_D & y_D & z_D & 1 \end{pmatrix} \right|.$$

Dimostrazione. Troviamo

$$\begin{aligned} \text{Vol}(T) &= \frac{1}{6} \left| \det \begin{pmatrix} x_B - x_A & x_C - x_A & x_D - x_A \\ y_B - y_A & y_C - y_A & y_D - y_A \\ z_B - z_A & z_C - z_A & z_D - z_A \end{pmatrix} \right| \\ &= \frac{1}{6} \left| \det \begin{pmatrix} x_B - x_A & x_C - x_A & x_D - x_A & x_A \\ y_B - y_A & y_C - y_A & y_D - y_A & y_A \\ z_B - z_A & z_C - z_A & z_D - z_A & z_A \\ 0 & 0 & 0 & 1 \end{pmatrix} \right| \\ &= \frac{1}{6} \left| \det \begin{pmatrix} x_B & x_C & x_D & x_A \\ y_B & y_C & y_D & y_A \\ z_B & z_C & z_D & z_A \\ 1 & 1 & 1 & 1 \end{pmatrix} \right|. \end{aligned}$$

Si conclude quindi permutando le colonne e facendo la trasposta. $\square$

La formula di Erone esprime l’area di un triangolo in funzione delle lunghezze dei suoi lati. Esiste un’analoga formula che esprime il volume di un tetraedro in funzione delle lunghezze dei suoi spigoli, scoperta da Piero della Francesca nel XV secolo. Nel linguaggio del XXI secolo la formula può essere scritta come determinante di una opportuna matrice simmetrica che ha lunghezze degli spigoli fra i suoi coefficienti.

Teorema 10.2.11 (Formula di Piero della Francesca). *Sia T un tetraedro con vertici A, B, C, D . Vale*

$$\text{Vol}(T)^2 = \frac{1}{288} \det \begin{pmatrix} 0 & 1 & 1 & 1 & 1 \\ 1 & 0 & |AB|^2 & |AC|^2 & |AD|^2 \\ 1 & |AB|^2 & 0 & |BC|^2 & |BD|^2 \\ 1 & |AC|^2 & |BC|^2 & 0 & |CD|^2 \\ 1 & |AD|^2 & |BD|^2 & |CD|^2 & 0 \end{pmatrix}.$$

Dimostrazione. Dal Corollario 10.2.10, usando $\det M = \det {}^t M$ e il teorema di Binet troviamo

$$\begin{aligned} \text{Vol}(T)^2 &= \frac{1}{6} \det \begin{pmatrix} 1 & x_A & y_A & z_A \\ 1 & x_B & y_B & z_B \\ 1 & x_C & y_C & z_C \\ 1 & x_D & y_D & z_D \end{pmatrix} \cdot \frac{1}{6} \det \begin{pmatrix} 1 & 1 & 1 & 1 \\ x_A & x_B & x_C & x_D \\ y_A & y_B & y_C & y_D \\ z_A & z_B & z_C & z_D \end{pmatrix} \\ &= \frac{1}{36} \det \begin{pmatrix} 1 + \|A\|^2 & 1 + \langle A, B \rangle & 1 + \langle A, C \rangle & 1 + \langle A, D \rangle \\ 1 + \langle B, A \rangle & 1 + \|B\|^2 & 1 + \langle B, C \rangle & 1 + \langle B, D \rangle \\ 1 + \langle C, A \rangle & 1 + \langle C, B \rangle & 1 + \|C\|^2 & 1 + \langle C, D \rangle \\ 1 + \langle D, A \rangle & 1 + \langle D, B \rangle & 1 + \langle D, C \rangle & 1 + \|D\|^2 \end{pmatrix} \end{aligned}$$

Usando lo sviluppo di Laplace e le mosse di Gauss sulle righe troviamo

$$\begin{aligned} \text{Vol}(T)^2 &= \frac{1}{36} \det \begin{pmatrix} 1 & 1 & 1 & 1 & 1 \\ 0 & 1 + \|A\|^2 & 1 + \langle A, B \rangle & 1 + \langle A, C \rangle & 1 + \langle A, D \rangle \\ 0 & 1 + \langle B, A \rangle & 1 + \|B\|^2 & 1 + \langle B, C \rangle & 1 + \langle B, D \rangle \\ 0 & 1 + \langle C, A \rangle & 1 + \langle C, B \rangle & 1 + \|C\|^2 & 1 + \langle C, D \rangle \\ 0 & 1 + \langle D, A \rangle & 1 + \langle D, B \rangle & 1 + \langle D, C \rangle & 1 + \|D\|^2 \end{pmatrix} \\ &= \frac{1}{36} \det \begin{pmatrix} 1 & 1 & 1 & 1 & 1 \\ -1 & \|A\|^2 & \langle A, B \rangle & \langle A, C \rangle & \langle A, D \rangle \\ -1 & \langle B, A \rangle & \|B\|^2 & \langle B, C \rangle & \langle B, D \rangle \\ -1 & \langle C, A \rangle & \langle C, B \rangle & \|C\|^2 & \langle C, D \rangle \\ -1 & \langle D, A \rangle & \langle D, B \rangle & \langle D, C \rangle & \|D\|^2 \end{pmatrix} \end{aligned}$$

Il minore 4×4 in basso a destra M ha determinante nullo perché è prodotto di matrici con determinante nullo:

$$M = \begin{pmatrix} 0 & x_A & y_A & z_A \\ 0 & x_B & y_B & z_B \\ 0 & x_C & y_C & z_C \\ 0 & x_D & y_D & z_D \end{pmatrix} \begin{pmatrix} 0 & 0 & 0 & 0 \\ x_A & x_B & x_C & x_D \\ y_A & y_B & y_C & y_D \\ z_A & z_B & z_C & z_D \end{pmatrix}.$$

Usando lo sviluppo di Laplace sulla prima colonna e $\det M = 0$, troviamo

$$\begin{aligned} \text{Vol}(T)^2 &= \frac{1}{36} \det \begin{pmatrix} 0 & 1 & 1 & 1 & 1 \\ -1 & \|A\|^2 & \langle A, B \rangle & \langle A, C \rangle & \langle A, D \rangle \\ -1 & \langle B, A \rangle & \|B\|^2 & \langle B, C \rangle & \langle B, D \rangle \\ -1 & \langle C, A \rangle & \langle C, B \rangle & \|C\|^2 & \langle C, D \rangle \\ -1 & \langle D, A \rangle & \langle D, B \rangle & \langle D, C \rangle & \|D\|^2 \end{pmatrix} \\ &= \frac{1}{288} \det \begin{pmatrix} 0 & 1 & 1 & 1 & 1 \\ 1 & -2\|A\|^2 & -2\langle A, B \rangle & -2\langle A, C \rangle & -2\langle A, D \rangle \\ 1 & -2\langle B, A \rangle & -2\|B\|^2 & -2\langle B, C \rangle & -2\langle B, D \rangle \\ 1 & -2\langle C, A \rangle & -2\langle C, B \rangle & -2\|C\|^2 & -2\langle C, D \rangle \\ 1 & -2\langle D, A \rangle & -2\langle D, B \rangle & -2\langle D, C \rangle & -2\|D\|^2 \end{pmatrix}. \end{aligned}$$

Nella seconda uguaglianza abbiamo moltiplicato per $-1, -2, -2, -2, -2$ le colonne e per $-\frac{1}{2}$ la prima riga. Aggiungendo multipli della prima riga alle altre righe e della prima colonna alle altre colonne otteniamo

$$\text{Vol}(T)^2 = \begin{pmatrix} 0 & 1 & 1 & 1 & \dots \\ 1 & \|A\|^2 - 2\|A\|^2 + \|A\|^2 & \|A\|^2 - 2\langle A, B \rangle + \|B\|^2 & \dots & \dots \\ 1 & \|B\|^2 - 2\langle B, A \rangle + \|A\|^2 & \|B\|^2 - 2\|B\|^2 + \|B\|^2 & \dots & \dots \\ 1 & \|C\|^2 - 2\langle C, A \rangle + \|A\|^2 & \|C\|^2 - 2\langle C, B \rangle + \|B\|^2 & \dots & \dots \\ 1 & \|D\|^2 - 2\langle D, A \rangle + \|A\|^2 & \|D\|^2 - 2\langle D, B \rangle + \|B\|^2 & \dots & \dots \end{pmatrix}$$

e quindi la tesi. $\square$

Possiamo esprimere la formula di Erone per un triangolo ABC con una scrittura simile (e dimostrarla esattamente allo stesso modo):

$$\text{Area}(T) = -\frac{1}{16} \det \begin{pmatrix} 0 & 1 & 1 & 1 \\ 1 & 0 & |AB|^2 & |AC|^2 \\ 1 & |AB|^2 & 0 & |BC|^2 \\ 1 & |AC|^2 & |BC|^2 & 0 \end{pmatrix}.$$

Esercizi

Esercizio 10.1 (Cerchio di Feuerbach). Sia ABC un triangolo. Mostra che la circonferenza che passa per i tre punti medi dei tre lati passa anche per i piedi delle tre altezze e per i punti medi dei segmenti compresi tra l'ortocentro e i vertici. Si veda la Figura 10.25.

Esercizio 10.2. Mostra che i parallelogrammi sono tutti affinemente equivalenti. Mostra che due trapezi sono affinemente equivalenti se e solo se hanno lo stesso rapporto di lunghezze fra base maggiore e minore.

Esercizio 10.3. Sia $ABCD$ un quadrilatero convesso. Mostra che

$$\text{Area}(ABCD) = \frac{|AC| \cdot |BD| \sin \alpha}{2}$$

dove α è l'angolo formato dalle sue due diagonali.

Un quadrilatero è *ciclico* se i suoi vertici giacciono su una circonferenza.

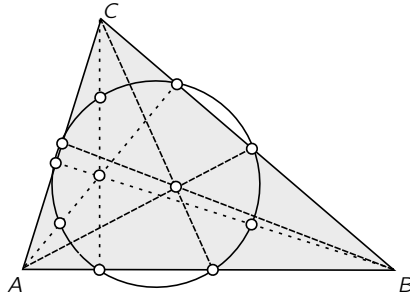


Figura 10.25. Il cerchio di Feuerbach di un triangolo passa per i tre punti medi dei lati, i tre piedi delle altezze, e i tre punti medi dei segmenti compresi tra ortocentro e vertici.

Esercizio 10.4. Sia $ABCD$ un quadrilatero convesso. Mostra che il quadrilatero è ciclico se e solo se $\alpha + \gamma = \beta + \delta$, dove $\alpha, \beta, \gamma, \delta$ sono le ampiezze degli angoli in A, B, C e D .

Esercizio 10.5 (Teorema di Tolomeo). Se $ABCD$ è un quadrilatero ciclico,

$$|AC| \cdot |BD| = |AB| \cdot |CD| + |AD| \cdot |BC|.$$

Esercizio 10.6. Sia $ABCD$ un quadrilatero convesso. Mostra che esiste un cerchio inscritto nel quadrilatero (cioè tangente ai quattro lati) se e solo se $|AB| + |CD| = |DA| + |BC|$.

Complementi

10.1. Coordinate baricentriche. Abbiamo definito nel piano $\mathbb{R}^2$ due tipi di coordinate: quelle cartesiane e quelle polari. Nello spazio $\mathbb{R}^3$ abbiamo costruito le coordinate cartesiane, sferiche e cilindriche. Introduciamo qui un nuovo tipo di coordinate, dette *coordinate baricentriche*, che funzionano in ogni dimensione. Così come le coordinate cartesiane in $\mathbb{R}^n$ dipendono dalla scelta di un sistema di riferimento, le coordinate baricentriche dipendono dalla scelta di $n + 1$ punti affinementemente indipendenti.

Fissiamo $n + 1$ punti affinementemente indipendenti $A_0, \dots, A_n$ in $\mathbb{R}^n$. Ci riferiamo alla Sezione 9.2.4 per le definizioni.

Siano adesso $t_0, \dots, t_n$ numeri reali, con $t_0 + \dots + t_n \neq 0$. Definiamo il *punto di coordinate baricentriche* $(t_0, \dots, t_n)$ come il punto

$$P = \frac{t_0 A_0 + \dots + t_n A_n}{t_0 + \dots + t_n} \in \mathbb{R}^n.$$

Notiamo subito alcuni fatti. Innanzitutto i punti di coordinate

$$(1, 0, \dots, 0), \quad (0, 1, \dots, 0), \quad \dots \quad (0, 0, \dots, 1)$$

sono precisamente i punti $A_0, \dots, A_n$. Inoltre le coordinate di un punto non sono univocamente determinate: se $\lambda \neq 0$, le coordinate baricentriche $(t_0, \dots, t_n)$ e $(\lambda t_0, \dots, \lambda t_n)$ identificano in realtà lo stesso punto. Affinché le coordinate di un punto siano univocamente determinate, possiamo

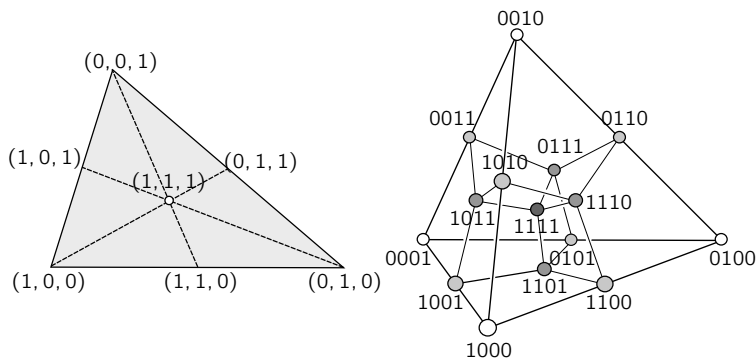


Figura 10.26. Le coordinate baricentriche dei vertici, dei punti medi dei segmenti e del baricentro di un triangolo di vertici $A_0 = (1, 0, 0)$, $A_1 = (0, 1, 0)$ e $A_2 = (0, 0, 1)$ (sinistra). Le coordinate baricentriche dei vertici, dei punti medi dei lati, dei baricentri delle facce, e del baricentro in un tetraedro di vertici $A_0 = (1, 0, 0, 0)$, $A_1 = (0, 1, 0, 0)$, $A_2 = (0, 0, 1, 0)$ e $A_3 = (0, 0, 0, 1)$. Nel disegno scriviamo $abcd$ invece di (a, b, c, d) per semplicità (destra).

normalizzarle in modo che la loro somma sia uno, cioè $t_0 + \dots + t_n = 1$. Dimostriamo questo fatto:

Proposizione 10.2.12. *Qualsiasi $P \in \mathbb{R}^n$ si scrive in modo unico come*

$$P = t_0 A_0 + \dots + t_n A_n$$

con $t_0, \dots, t_n \in \mathbb{R}$ tali che $t_0 + \dots + t_n = 1$.

Dimostrazione. I vettori $\overrightarrow{A_0 A_1}, \dots, \overrightarrow{A_0 A_n}$ sono una base di $\mathbb{R}^n$. Quindi P si scrive in modo unico come

$$P = A_0 + \lambda_1 \overrightarrow{A_0 A_1} + \dots + \lambda_n \overrightarrow{A_0 A_n} = (1 - \lambda_1 - \dots - \lambda_n) A_0 + \lambda_1 A_1 + \dots + \lambda_n A_n.$$

Prendiamo $t_0 = 1 - \lambda_1 - \dots - \lambda_n$, $t_1 = \lambda_1, \dots, t_n = \lambda_n$. □

Gli spazi euclidei che ci interessano di più sono, come sempre, il piano $\mathbb{R}^2$ e lo spazio $\mathbb{R}^3$. Nella Figura 10.26 sono disegnati un triangolo in $\mathbb{R}^2$ ed un tetraedro in $\mathbb{R}^3$. I vertici A_0, A_1, A_2 del triangolo e A_0, A_1, A_2, A_3 del tetraedro sono punti indipendenti e definiscono un sistema di coordinate baricentriche, rispettivamente in $\mathbb{R}^2$ e in $\mathbb{R}^3$. Sono mostrate le coordinate baricentriche dei vertici, dei punti medi dei lati, dei baricentri del triangolo, delle facce e del tetraedro. Le coordinate nella figura non sono normalizzate: la loro somma può essere diversa da 1.

Dalla figura si capisce subito perché le coordinate baricentriche si chiamano in questo modo: il baricentro del triangolo o del tetraedro è il punto con coordinate tutte uguali. Notiamo inoltre che, per l'Esercizio 10.1.7,

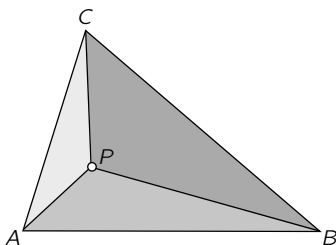


Figura 10.27. I rapporti fra le coordinate baricentriche s , t e u di un punto P del triangolo sono uguali ai rapporti fra le aree dei triangoli corrispondenti BCP , CAP e ABP .

i punti del triangolo ABC sono precisamente tutti i punti del piano aventi coordinate baricentriche (t_0, t_1, t_2) tutte non negative, cioè con $t_i \geq 0$ per ogni i . Analogamente i punti del tetraedro $ABCD$ sono precisamente tutti i punti (t_0, t_1, t_2, t_3) dello spazio con $t_i \geq 0$ per ogni i . Le coordinate baricentriche descrivono in modo efficiente i punti interni di un triangolo o di un tetraedro.

Forniamo adesso due interpretazioni delle coordinate baricentriche, una fisica e l'altra geometrica. Ci limitiamo a studiare il piano, ma le stesse interpretazioni sono valide anche nello spazio con opportune modifiche. Sia T un triangolo in $\mathbb{R}^2$ di vertici A, B e C . Sia P un punto del triangolo di coordinate baricentriche (s, t, u) rispetto ai punti A, B, C . Si veda la Figura 10.27.

Da un punto di vista fisico, il punto P è il baricentro di un sistema formato da 3 palline centrate nei punti A, B e C aventi massa s, t e u . Notiamo che se s è molto più grande di t e u , allora il punto P è effettivamente molto più vicino ad A che agli altri punti B e C .

Da un punto di vista geometrico, le coordinate baricentriche misurano i rapporti fra le aree dei tre triangoli in cui viene diviso T , come mostrato nella Figura 10.27.

Proposizione 10.2.13. *I rapporti fra le coordinate s, t e u sono uguali ai rapporti fra le aree dei triangoli corrispondenti BCP, CAP e ABP , cioè:*

$$\frac{s}{\text{Area}(BCP)} = \frac{t}{\text{Area}(CAP)} = \frac{u}{\text{Area}(ABP)}.$$

Dimostrazione. Possiamo supporre che le coordinate siano normalizzate, cioè che $s + t + u = 1$. Quindi

$$P = sA + tB + uC = \begin{pmatrix} s x_A + t x_B + u x_C \\ s y_A + t y_B + u y_C \end{pmatrix}.$$

Punto	Coordinate baricentriche
Baricentro G	$(1, 1, 1)$
Incentro I	(a, b, c)
Ortocentro H	$(\tan \alpha, \tan \beta, \tan \gamma)$
Circocentro O	$(\sin 2\alpha, \sin 2\beta, \sin 2\gamma)$

Tabella 10.2. Le coordinate baricentriche (non normalizzate) dei 4 punti notevoli più importanti del triangolo.

Per le Proposizioni 10.1.21 e 3.3.8 l'area $R = \text{Area}(ABP)$ è uguale a

$$\begin{aligned}
 R &= \frac{1}{2} \left| \det \begin{pmatrix} x_A & y_A & 1 \\ x_B & y_B & 1 \\ sx_A + tx_B + ux_C & sy_A + ty_B + uy_C & s + t + u \end{pmatrix} \right| \\
 &= \frac{1}{2} \left| \det \begin{pmatrix} x_A & y_A & 1 \\ x_B & y_B & 1 \\ sx_A & sy_A & s \end{pmatrix} + \det \begin{pmatrix} x_A & y_A & 1 \\ x_B & y_B & 1 \\ tx_B & ty_B & t \end{pmatrix} + \det \begin{pmatrix} x_0 & y_0 & 1 \\ x_1 & y_1 & 1 \\ ux_C & uy_C & u \end{pmatrix} \right| \\
 &= \frac{1}{2} \left| 0 + 0 + u \det \begin{pmatrix} x_A & y_A & 1 \\ x_B & y_B & 1 \\ x_C & y_C & 1 \end{pmatrix} \right| = u \text{Area}(T).
 \end{aligned}$$

Abbiamo dimostrato che

$$\text{Area}(ABP) = u \text{Area}(ABC).$$

Una uguaglianza analoga vale per le aree di BCP e CAP . Quindi i rapporti fra le aree dei tre triangoli sono come i rapporti fra s , t e u . $\square$

Esaminando la dimostrazione della proposizione notiamo che l'interpretazione geometrica funziona anche se P è un punto esterno al triangolo; alcune coordinate in s , t , u possono essere negative e i triangoli corrispondenti sono esterni al triangolo: in questo caso la loro area va intesa come area negativa.

Descriviamo adesso le coordinate baricentriche dei quattro punti notevoli del triangolo T che abbiamo studiato nelle pagine precedenti. Indichiamo come di consueto con a, b, c le lunghezze dei lati $|BC|, |CA|, |AB|$ e con α, β, γ gli angoli ai vertici.

Proposizione 10.2.14. *Le coordinate baricentriche di baricentro, incentro, circocentro e ortocentro sono elencate nella Tabella 10.2.*

Dimostrazione. Le coordinate del baricentro e dell'incentro seguono dalle loro definizioni. L'ortocentro è descritto nella Figura 10.28. Si verifica

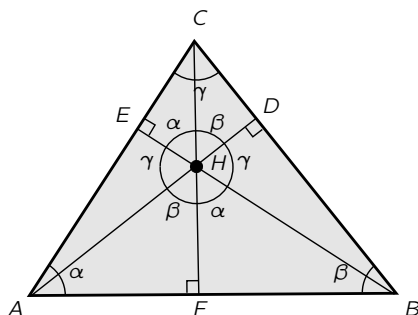


Figura 10.28. Calcolo delle coordinate baricentriche dell'ortocentro.

facilmente che gli angoli intorno a H sono α, β, γ come descritto. Quindi

$$\begin{aligned} \text{Area}(ABH) &= \frac{1}{2}c|HF| = \frac{1}{2}c|AF| \cot \beta = \frac{1}{2}cb \cos \alpha \cot \beta \\ &= \text{Area}(T) \cot \alpha \cot \beta = K \cdot \text{Area}(T) \tan \gamma \end{aligned}$$

con $K = \cot \alpha \cot \beta \cot \gamma$. I rapporti fra le aree sono quindi come mostrato nella Tabella 10.2. Anche le coordinate del circocentro si trovano calcolando i rapporti fra aree: questo conto è più semplice del precedente ed è lasciato per esercizio. $\square$

Teorema spettrale

Nei capitoli precedenti abbiamo studiato approfonditamente due aspetti riguardanti gli spazi vettoriali: gli endomorfismi ed i prodotti scalari. In questo capitolo descriviamo un enunciato, noto come *teorema spettrale*, che collega questi due argomenti.

Il teorema spettrale descrive completamente per quali endomorfismi sia possibile trovare una base che sia contemporaneamente formata da autovettori e ortonormale. Questo teorema centrale dell'algebra lineare ha molte conseguenze in vari campi della matematica e della scienza; lo useremo nel Capitolo 13 per classificare coniche e quadriche.

11.1. Prodotti hermitiani

In tutto il Capitolo 7 abbiamo introdotto i prodotti scalari soltanto per spazi vettoriali reali. Mostriamo adesso brevemente come sia possibile definire una nozione analoga anche per gli spazi vettoriali complessi: questa estensione è nota con il nome di *prodotto hermitiano*.

11.1.1. Definizione. Sia V uno spazio vettoriale complesso.

Definizione 11.1.1. Un *prodotto hermitiano* su V è una applicazione

$$\begin{aligned} V \times V &\longrightarrow \mathbb{C} \\ (v, w) &\longmapsto \langle v, w \rangle \end{aligned}$$

che soddisfa i seguenti assiomi:

- (1) $\langle v + v', w \rangle = \langle v, w \rangle + \langle v', w \rangle$,
- (2) $\langle \lambda v, w \rangle = \lambda \langle v, w \rangle$,
- (3) $\langle v, w \rangle = \overline{\langle w, v \rangle}$

per ogni $v, v', w \in V$ e ogni $\lambda \in \mathbb{C}$.

Notiamo immediatamente che l'assioma (3) differisce da quello analogo per i prodotti scalari per la presenza di un coniugio. Il motivo di questo cambiamento sarà presto chiaro. Intanto da questi assiomi deduciamo facilmente altre proprietà:

- (4) $\langle v, w + w' \rangle = \langle v, w \rangle + \langle v, w' \rangle$,
- (5) $\langle v, \lambda w \rangle = \overline{\lambda} \langle v, w \rangle$,
- (6) $\langle 0, w \rangle = \langle v, 0 \rangle = 0$

per ogni $v, w, w' \in V$ e $\lambda \in \mathbb{C}$. Ad esempio, la (5) si mostra così:

$$\langle v, \lambda w \rangle = \overline{\langle \lambda w, v \rangle} = \overline{\lambda \langle w, v \rangle} = \bar{\lambda} \overline{\langle w, v \rangle} = \bar{\lambda} \langle v, w \rangle.$$

Perché abbiamo inserito il coniugio in (3), che per adesso ha come unico effetto quello di rendere lievemente più complicati i passaggi algebrici? Lo abbiamo fatto essenzialmente per ottenere la proprietà seguente:

(7) $\langle v, v \rangle$ è un numero reale, per ogni $v \in V$.

Per dimostrare ciò, notiamo che l'assioma (3) implica che

$$\langle v, v \rangle = \overline{\langle v, v \rangle}$$

e quindi effettivamente $\langle v, v \rangle \in \mathbb{R}$. Il coniugio in (3) è un trucco per ottenere che $\langle v, v \rangle$ sia un numero reale.

Osservazione 11.1.2. Gli assiomi (1), (2) e (4) sono gli stessi dei prodotti scalari, mentre le proprietà (5) è differente. Diciamo che il prodotto hermitiano è lineare a sinistra e *antilineare* a destra. Unendo queste due proprietà, diciamo anche che un prodotto hermitiano è *sesquilineare* invece che *bilineare*, dal latino *sesqui* che denota il numero 1,5. In un certo senso, un prodotto hermitiano è lineare una volta e mezzo ma non due.

11.1.2. Prodotto hermitiano definito positivo. In un prodotto hermitiano, è effettivamente cruciale che $\langle v, v \rangle$ sia un numero reale: in questo modo è perfettamente sensato chiedere che questo numero sia positivo o negativo, a seconda delle necessità. Possiamo quindi introdurre questa definizione, come abbiamo fatto nel caso reale.

Definizione 11.1.3. Un prodotto hermitiano è *definito positivo* se $\langle v, v \rangle > 0$ per ogni $v \neq 0$.

In un prodotto hermitiano definito positivo si definiscono in modo del tutto analogo al caso reale molti dei concetti già visti precedentemente. La *norma* di un vettore è

$$\|v\| = \sqrt{\langle v, v \rangle}.$$

Si possono quindi introdurre le nozioni di *spazio ortogonale* e di *basse ortonormale* ed il procedimento di Gram – Schmidt funziona anche in questo contesto.

11.1.3. Prodotto hermitiano euclideo. Se $x \in \mathbb{C}^n$ è un vettore e più in generale $A \in M(m, n, \mathbb{C})$ è una matrice a coefficienti complessi, indichiamo con $\bar{x}$ e $\bar{A}$ il vettore o matrice ottenuto da x oppure A facendo il coniugio di ogni singolo elemento. Il *prodotto hermitiano euclideo* su $\mathbb{C}^n$ è dato da

$$\langle x, y \rangle = {}^t x \bar{y}.$$

Questo è effettivamente un prodotto hermitiano, analogo al prodotto scalare euclideo: notiamo il coniugio presente sulla variabile y , grazie al quale risulta valido l'assioma (3), infatti

$$\langle y, x \rangle = {}^t y \bar{x} = {}^t \bar{x} y = \overline{{}^t x \bar{y}} = \overline{\langle x, y \rangle}.$$

Come nel caso reale, il prodotto hermitiano euclideo è definito positivo, perché per ogni $x \in \mathbb{C}^n$ non nullo troviamo

$$\langle x, x \rangle = {}^t x \bar{x} = x_1 \bar{x}_1 + \cdots + x_n \bar{x}_n = |x_1|^2 + \cdots + |x_n|^2 > 0.$$

11.1.4. Matrici hermitiane. Le matrici hermitiane hanno lo stesso ruolo delle matrici simmetriche nei prodotti scalari.

Una *matrice hermitiana* è una matrice quadrata complessa H per cui

$${}^t H = \bar{H}.$$

In altre parole vale $H_{ij} = \bar{H}_{ji}$ per ogni i, j . Ad esempio, la matrice seguente è una matrice hermitiana:

$$\begin{pmatrix} 2 & 1+i \\ 1-i & 1 \end{pmatrix}.$$

Notiamo che in una matrice hermitiana gli elementi H_{ii} sulla diagonale devono necessariamente essere reali, perché $H_{ii} = \bar{H}_{ii}$ per ogni i .

Osservazione 11.1.4. Una matrice quadrata a coefficienti reali è hermitiana se e solo se è simmetrica.

Così come una matrice simmetrica S definisce un prodotto scalare g_S su $\mathbb{R}^n$, una matrice hermitiana H determina un prodotto hermitiano g_H su $\mathbb{C}^n$ nel modo seguente:

$$g_H(x, y) = {}^t x H \bar{y}.$$

Notiamo che effettivamente l'assioma (3) è soddisfatto:

$$g_H(x, y) = {}^t x H \bar{y} = {}^t ({}^t x H \bar{y}) = {}^t \bar{y} {}^t H x = {}^t \bar{y} \bar{H} x = \overline{{}^t y H x} = \overline{g_H(y, x)}.$$

La verifica degli assiomi (1) e (2) è identica a quando visto per i prodotti scalari. Come per i prodotti scalari dimostriamo che

$$g_H(e_i, e_j) = {}^t e_i H \bar{e}_j = {}^t e_i H e_j = H_{ij}$$

dove $e_1, \dots, e_n$ è la base canonica di $\mathbb{C}^n$.

11.1.5. Altri esempi. I numeri complessi sono presenti in numerosi ambiti della scienza moderna e in alcuni contesti i prodotti hermitiani sono più frequenti dei prodotti scalari. Ci limitiamo a fare un esempio di prodotto hermitiano utile in analisi.

Esempio 11.1.5. Sia V lo spazio vettoriale complesso formato da tutte le funzioni $f: [a, b] \rightarrow \mathbb{C}$ continue. Definiamo su V un prodotto hermitiano nel modo seguente.

$$\langle f, g \rangle = \int_a^b f(t) \overline{g(t)} dt.$$

Questo prodotto hermitiano è definito positivo perché

$$\langle f, f \rangle = \int_a^b f(t) \overline{f(t)} dt = \int_a^b |f(t)|^2 dt > 0$$

per ogni funzione continua non nulla f .

11.1.6. Matrice associata. Come nel caso reale, se V è un prodotto hermitiano e $\mathcal{B} = \{v_1, \dots, v_n\}$ è una base di V , possiamo definire la *matrice associata* H nel modo seguente:

$$H_{ij} = \langle v_i, v_j \rangle$$

per ogni i, j . Abbiamo scelto la lettera H invece di S perché questa matrice non è simmetrica ma bensì hermitiana: infatti per l'assioma (3) abbiamo

$$H_{ij} = \bar{H}_{ji}.$$

Per ogni coppia di vettori $v, w \in V$ possiamo sempre scrivere

$$(21) \quad \langle v, w \rangle = {}^t[v]_{\mathcal{B}} \cdot H \cdot \overline{[w]_{\mathcal{B}}}.$$

Questo fatto si dimostra esattamente come il Corollario 7.2.6, stando attenti alla sesquilinearità, cioè al coniugio nella proprietà (5).

In modo simile alla Proposizione 7.2.8 si mostra che qualsiasi prodotto hermitiano su $\mathbb{C}^n$ è della forma g_H per una matrice hermitiana H .

11.2. Endomorfismi autoaggiunti

Introduciamo in questo capitolo una importante classe di endomorfismi. In tutta questa sezione indichiamo con V uno spazio vettoriale reale dotato di un prodotto scalare definito positivo, oppure uno spazio vettoriale complesso dotato di un prodotto hermitiano definito positivo.

11.2.1. Definizione. Un endomorfismo $T: V \rightarrow V$ è *autoaggiunto* se

$$(22) \quad \langle T(v), w \rangle = \langle v, T(w) \rangle$$

per ogni $v, w \in V$.

Notiamo subito che questa definizione ricorda un po' quella di isometria, in cui si chiede che $\langle v, w \rangle = \langle T(v), T(w) \rangle$. Ci sono effettivamente delle analogie fra le due definizioni, ma anche delle importanti differenze.

11.2.2. Con le matrici. Come con le isometrie, cerchiamo subito di tradurre la definizione di endomorfismo autoaggiunto in una condizione concreta sulle matrici associate.

Scegliamo una base ortonormale $\mathcal{B}$ di V . Sia $T: V \rightarrow V$ un endomorfismo e $A = [T]_{\mathcal{B}}^{\mathcal{B}}$ la sua matrice associata. Nella proposizione seguente è importante che la base $\mathcal{B}$ sia ortonormale.

Proposizione 11.2.1. *L'endomorfismo T è autoaggiunto se e solo se la matrice A è hermitiana.*

Dimostrazione. Per bilinearità (o sesquilinearità nel caso complesso), per verificare se T è autoaggiunto è sufficiente testare la condizione (22) sugli elementi v_i, v_j della base $\mathcal{B}$. Nel caso complesso, dall'equazione (21) ricaviamo

$$\langle T(v_i), v_j \rangle = {}^t[T(v_i)]_{\mathcal{B}} \cdot I_n \cdot \overline{[v_j]_{\mathcal{B}}} = {}^t A^i e_j = A_{ji},$$

$$\langle v_i, T(v_j) \rangle = {}^t[v_i]_{\mathcal{B}} \cdot I_n \cdot \overline{[T(v_j)]_{\mathcal{B}}} = {}^t e_i \bar{A}^j = \bar{A}_{ij}.$$

Quindi T è autoaggiunto se e solo se

$$A_{ji} = \bar{A}_{ij}$$

per ogni i, j , in altre parole se e solo se A è hermitiana. Il caso reale è del tutto analogo (senza i coniugi). $\square$

In particolare, otteniamo il fatto seguente.

Corollario 11.2.2. *Sia A una matrice $n \times n$. L'endomorfismo $L_A: \mathbb{C}^n \rightarrow \mathbb{C}^n$ è autoaggiunto rispetto al prodotto hermitiano euclideo di $\mathbb{C}^n$ se e solo se la matrice A è hermitiana.*

Analogamente, se A è reale, l'endomorfismo $L_A: \mathbb{R}^n \rightarrow \mathbb{R}^n$ è autoaggiunto rispetto al prodotto scalare euclideo di $\mathbb{R}^n$ se e solo se la matrice A è simmetrica (ricordiamo che una matrice reale è hermitiana $\iff$ è simmetrica).

Notiamo che la Proposizione 11.2.1 è valida solo se prendiamo una base $\mathcal{B}$ ortonormale, come mostra questo esempio.

Esempio 11.2.3. Consideriamo $\mathbb{R}^2$ con il prodotto scalare euclideo. L'endomorfismo L_A definito dalla matrice $A = \begin{pmatrix} 2 & 1 \\ 1 & 1 \end{pmatrix}$ è autoaggiunto perché la matrice A è simmetrica e stiamo usando la base canonica che è ortonormale.

Se scriviamo L_A rispetto ad un'altra base ortonormale, ad esempio $\mathcal{B} = \left\{ \begin{pmatrix} 0 \\ 1 \end{pmatrix}, \begin{pmatrix} -1 \\ 0 \end{pmatrix} \right\}$, otteniamo una nuova matrice simmetrica:

$$A' = \begin{pmatrix} 0 & 1 \\ -1 & 0 \end{pmatrix} \begin{pmatrix} 2 & 1 \\ 1 & 1 \end{pmatrix} \begin{pmatrix} 0 & -1 \\ 1 & 0 \end{pmatrix} = \begin{pmatrix} 1 & -1 \\ -1 & 2 \end{pmatrix}.$$

Questo è coerente con la Proposizione 11.2.1. Se però scriviamo A rispetto ad una base *non* ortonormale, ad esempio $\mathcal{B} = \left\{ \begin{pmatrix} -1 \\ 1 \end{pmatrix}, \begin{pmatrix} 1 \\ 0 \end{pmatrix} \right\}$, la nuova matrice associata può non essere simmetrica. In questo caso otteniamo

$$A' = \begin{pmatrix} 0 & 1 \\ 1 & 1 \end{pmatrix} \begin{pmatrix} 2 & 1 \\ 1 & 1 \end{pmatrix} \begin{pmatrix} -1 & 1 \\ 1 & 0 \end{pmatrix} = \begin{pmatrix} 0 & 1 \\ -1 & 3 \end{pmatrix}.$$

Osservazione 11.2.4. Le matrici hermitiane (o simmetriche) hanno un doppio ruolo in questo capitolo: sia come matrici che rappresentano un prodotto hermitiano (o scalare), che come matrici che rappresentano un endomorfismo autoaggiunto (rispetto ad una base ortonormale). È importante non fare confusione fra questi due concetti ben distinti.

11.2.3. Sottospazi invarianti. Mostriamo che gli endomorfismi autoaggiunti si comportano con i sottospazi invarianti in modo simile alle isometrie. Questa proposizione è analoga al Corollario 8.2.12.

Proposizione 11.2.5. *Sia $T: V \rightarrow V$ un endomorfismo autoaggiunto e $U \subset V$ un sottospazio. Se $T(U) \subset U$ allora $T(U^\perp) \subset U^\perp$.*

Dimostrazione. Prendiamo un vettore qualsiasi $v \in U^\perp$ e dobbiamo dimostrare che $T(v) \in U^\perp$. Per ogni vettore $u \in U$ otteniamo

$$\langle T(v), u \rangle = \langle v, T(u) \rangle = 0$$

perché $T(u) \in U$ per ipotesi. Quindi $T(v)$ è ortogonale a tutti i vettori $u \in U$, in altre parole $T(v) \in U^\perp$. $\square$

11.3. Il teorema

Introduciamo adesso un risultato noto come *teorema spettrale*. Si tratta di uno dei enunciati più importanti dell'algebra lineare, con varie implicazioni in diversi rami della scienza moderna. Il teorema collega due importanti argomenti studiati in questo libro: la diagonalizzabilità degli endomorfismi e i prodotti scalari definiti positivi.

La motivazione principale che sta dietro al teorema è la seguente considerazione: abbiamo visto che gli endomorfismi più semplici da studiare sono quelli diagonalizzabili, cioè quelli che hanno una base di autovettori. In quali casi possiamo chiedere che questa base di autovettori sia anche ortonormale? Il teorema spettrale risponde completamente a questa domanda: questo è possibile precisamente per gli endomorfismi autoaggiunti.

11.3.1. Enunciato. Sia come sempre V uno spazio vettoriale complesso dotato di un prodotto hermitiano definito positivo, oppure uno spazio reale dotato di un prodotto scalare definito positivo.

Teorema 11.3.1 (Teorema spettrale). *Un endomorfismo $T: V \rightarrow V$ è autoaggiunto $\iff$ ha una base ortonormale di autovettori e tutti i suoi autovalori sono in $\mathbb{R}$.*

Dimostrazione. ($\Leftarrow$) Se T ha una base $\mathcal{B}$ ortonormale di autovettori, la matrice $A = [T]_{\mathcal{B}}^{\mathcal{B}}$ è diagonale; poiché gli autovalori sono reali, la matrice A è diagonale con numeri reali sulla diagonale e quindi è chiaramente hermitiana. Per la Proposizione 11.2.1 l'endomorfismo T è autoaggiunto.

($\Rightarrow$) Dimostriamo il teorema innanzitutto nel caso complesso. Iniziamo verificando che tutti gli autovalori di T sono reali: se λ è autovalore, allora c'è un autovettore v e $T(v) = \lambda v$, quindi

$$\lambda \langle v, v \rangle = \langle \lambda v, v \rangle = \langle T(v), v \rangle = \langle v, T(v) \rangle = \langle v, \lambda v \rangle = \bar{\lambda} \langle v, v \rangle.$$

Poiché $\langle v, v \rangle > 0$, deduciamo che $\lambda = \bar{\lambda}$ e quindi $\lambda \in \mathbb{R}$.

Dimostriamo adesso per induzione sulla dimensione n di V che T ha una base ortonormale di autovettori. Il caso $n = 1$ è banale: tutti gli endomorfismi di uno spazio di dimensione uno hanno una base ortonormale di autovettori, formata da un elemento solo di norma uno.

Supponiamo il fatto vero per la dimensione $n-1$ e lo mostriamo per n . Poiché siamo sui complessi, esiste sempre un autovettore $v \in V$. La retta $\text{Span}(v)$ è chiaramente T -invariante e quindi anche lo spazio ortogonale

$U = \text{Span}(v)^\perp$ è T -invariante per la Proposizione 11.2.5, cioè vale $T(U) \subset U$. La restrizione

$$T|_U: U \rightarrow U$$

è un endomorfismo autoaggiunto di uno spazio vettoriale U di dimensione $n-1$. Per l'ipotesi induttiva, esiste una base $v_2, \dots, v_n$ ortonormale formata da autovettori per $T|_U$. Aggiungendo il vettore iniziale v , rinormalizzato in modo che abbia norma uno, otteniamo una base $\mathcal{B} = \{v, v_2, \dots, v_n\}$ ortonormale formata da autovettori di T .

Abbiamo dimostrato il teorema spettrale nel caso complesso. Come importante conseguenza, notiamo che una matrice reale simmetrica S ha sempre tutti gli autovalori reali: infatti l'endomorfismo $L_S: \mathbb{C}^n \rightarrow \mathbb{C}^n$ è autoaggiunto (rispetto al prodotto hermitiano euclideo) per la Proposizione 11.2.1 e quindi L_S ha gli autovalori tutti reali e una base ortonormale di autovettori in $\mathbb{C}^n$.

Resta infine da dimostrare il teorema spettrale nel caso reale. Tutta la dimostrazione appena descritta si applica se riusciamo a dimostrare che T ha solo autovalori reali. Prendiamo una base ortonormale $\mathcal{B}$ di V e notiamo che $S = [T]_{\mathcal{B}}^{\mathcal{B}}$ è hermitiana, cioè simmetrica, per la Proposizione 11.2.1. Per quanto appena detto una matrice simmetrica reale ha tutti gli autovalori reali e quindi siamo a posto. $\square$

Come molti teoremi importanti, il teorema spettrale può essere enunciato in vari modi diversi. Il corollario che segue è una versione del teorema spettrale nel caso reale, scritta con il linguaggio delle matrici. Consideriamo lo spazio $\mathbb{R}^n$ con l'usuale prodotto scalare euclideo.

Corollario 11.3.2. *Sia A una matrice $n \times n$ reale. Sono fatti equivalenti:*

- (1) A è simmetrica;
- (2) L_A ha una base ortonormale di autovettori;
- (3) esiste una matrice ortogonale M tale che

$${}^tMAM = M^{-1}AM = D$$

sia una matrice diagonale.

Dimostrazione. L'equivalenza (1) $\Leftrightarrow$ (2) è il teorema spettrale. L'esistenza di una base $\mathcal{B}$ di autovettori per L_A è equivalente all'esistenza di una matrice M tale che $M^{-1}AM = D$ sia diagonale. Inoltre $\mathcal{B}$ è ortonormale se e solo se M è ortogonale, quindi (2) $\Leftrightarrow$ (3). Notiamo che $M^{-1} = {}^tM$ in questo caso. $\square$

Riassumendo, abbiamo scoperto due fatti non banali: il primo è che ogni matrice simmetrica reale è diagonalizzabile; il secondo è che è sempre possibile scegliere una base di autovettori che sia anche ortonormale. Solo le matrici simmetriche hanno entrambe queste proprietà.

Esempio 11.3.3. Sia come sempre V uno spazio vettoriale complesso dotato di un prodotto hermitiano definito positivo, oppure uno spazio reale

dotato di un prodotto scalare definito positivo. Sia $U \subset V$ un sottospazio vettoriale. La riflessione ortogonale r_U e la proiezione ortogonale p_U lungo U , definite nell'Esempio 7.5.3 e nella Sezione 8.1.10, sono entrambi esempi importanti di endomorfismi autoaggiunti.

Per verificare questo fatto, anziché usare la definizione di endomorfismo autoaggiunto, si può invocare il teorema spettrale e notare che entrambe r_U e p_U hanno una base ortonormale di autovettori, con autovalori reali 1, -1 oppure 0, ottenuta unendo due qualsiasi basi ortonormali di U e di $U^\perp$.

Applicando ancora il teorema spettrale, ne deduciamo che proiezioni e riflessioni ortogonali in $\mathbb{R}^n$ (con il prodotto scalare euclideo) sono rappresentate (rispetto alla base canonica) sempre da matrici simmetriche. Abbiamo visto un caso nell'Esempio 8.2.5.

11.3.2. Segnatura. Il teorema spettrale ha varie conseguenze inaspettate: una di queste è un metodo definitivo (ma laborioso) per calcolare la segnatura (i_+, i_-, i_0) di una matrice simmetrica S .

Proposizione 11.3.4. *Sia S una matrice simmetrica e (i_+, i_-, i_0) la sua segnatura. I numeri i_+ , i_- e i_0 sono pari al numero di autovalori positivi, negativi e nulli di S .*

Dimostrazione. Sappiamo dal Corollario 11.3.2 che esiste una matrice ortogonale M per cui

$${}^tMSM = M^{-1}SM = D$$

sia diagonale. Le matrici S e D sono simultaneamente simili e congruenti! La matrice D ha la stessa segnatura di S perché sono congruenti, e ha gli stessi autovalori di S perché sono simili. Quindi è sufficiente mostrare la proposizione per D , ma per le diagonali la proposizione è nota (si veda la Sezione 7.4.5). $\square$

Abbiamo apparentemente trovato un metodo infallibile per determinare la segnatura di una matrice simmetrica S : è sufficiente determinare i suoi autovalori. Ma come troviamo gli autovalori di S ? Questi sono le radici del polinomio caratteristico $p_S(x)$, però sfortunatamente non esistono formule risolutive per trovare le radici di un polinomio quando questo ha grado ≥ 5 , e di fatto l'unica formula semplice è quella in grado 2 che si impara alle superiori.

Per fortuna non è necessario determinare precisamente gli autovalori, è sufficiente il loro *segno*. A questo scopo possiamo usare il *criterio di Cartesio*:

Teorema 11.3.5 (Criterio di Cartesio). *Sia $p(x) = a_n x^n + \dots + a_m x^m$ un polinomio con n radici reali, scritto in modo che il primo e l'ultimo coefficiente a_n e a_m siano entrambi non nulli. Valgono i fatti seguenti:*

- (1) *il numero di radici nulle di $p(x)$ è m ;*

- (2) *il numero di radici positive di $p(x)$ è pari al numero di cambiamenti di segno nella sequenza $a_n, \dots, a_m$ dei coefficienti, da cui siano stati rimossi i coefficienti nulli.*

Ad esempio, il polinomio $p(x) = x^4 - x^2$ ha due radici nulle (perché $m = 2$) e una radice positiva perché la sequenza dei coefficienti è $1, -1$ e contiene una variazione. Effettivamente sappiamo che le radici sono $0, 0, 1, -1$.

Il numero di radici negative è n meno il numero di radici nulle e positive (attenzione: questo non è necessariamente il numero di permanenze di segno nella successione dei coefficienti!)

Con il criterio di Cartesio, è sempre possibile determinare la segnatura di una matrice esaminando i segni dei coefficienti del suo polinomio caratteristico. Notiamo però che il calcolo esplicito del polinomio caratteristico di una matrice $n \times n$ può essere molto dispendioso: quando funziona, il criterio di Jacobi spiegato nella Sezione 7.4.7 è più immediato.

Forniamo una dimostrazione del criterio nei complementi al capitolo.

Esercizi

Esercizio 11.1. Verifica che la prima matrice (simmetrica) ha una base ortonormale di autovettori mentre la seconda no:

$$\begin{pmatrix} 2 & 1 \\ 1 & 1 \end{pmatrix}, \quad \begin{pmatrix} 1 & 1 \\ 0 & 2 \end{pmatrix}.$$

Esercizio 11.2. Determina una base ortonormale di autovettori (rispetto al prodotto hermitiano euclideo di $\mathbb{C}^2$) per l'endomorfismo autoaggiunto definito dalla matrice hermitiana seguente:

$$\begin{pmatrix} 1 & i \\ -i & 1 \end{pmatrix}.$$

Nei prossimi due esercizi V è sempre uno spazio vettoriale reale con un prodotto scalare definito positivo.

Esercizio 11.3. Sia $T: V \rightarrow V$ endomorfismo. Supponiamo che T sia contemporaneamente un endomorfismo autoaggiunto e una isometria. Mostra che T è una riflessione ortogonale lungo un sottospazio $U \subset V$. Chi è U ?

Esercizio 11.4. Sia $T: V \rightarrow V$ un endomorfismo autoaggiunto. Mostra che se $T^2 = T$ allora T è la proiezione ortogonale lungo un sottospazio $U \subset V$. Chi è U ?

Esercizio 11.5. Sia A una matrice quadrata reale diagonalizzabile, con tutti gli autovalori positivi. Mostra che esiste una radice quadrata B di A , cioè esiste una matrice quadrata B tale che $A = B^2$.

Se A è simmetrica, mostra che puoi chiedere che B sia simmetrica.

Esercizio 11.6. Mostra che $-I_n$ ha una radice quadrata se e solo se n è pari.

Esercizio 11.7. Siano g e h due prodotti scalari sullo stesso spazio vettoriale reale V . Sia g definito positivo. Mostra che esiste una base che sia simultaneamente ortogonale per entrambi i prodotti scalari.

Esercizio 11.8. Considera lo spazio vettoriale reale $V = C[a, b]$ delle funzioni continue su $[a, b]$ a valori in $\mathbb{R}$, dotato del prodotto scalare definito positivo

$$\langle f, g \rangle = \int_a^b f(t)g(t)dt.$$

Sia $h \in C[a, b]$ una funzione fissata. Considera l'endomorfismo

$$T: V \rightarrow V, \quad T(f) = hf$$

che moltiplica qualsiasi funzione f per h , trasformando $f(x)$ in $h(x)f(x)$. Mostra che T è un endomorfismo autoaggiunto per V .

Esercizio 11.9. Considera lo spazio $V = C_c^\infty(\mathbb{R})$ formato da tutte le funzioni $f: \mathbb{R} \rightarrow \mathbb{R}$ che sono derivabili infinite volte e che abbiano *supporto compatto*, cioè tali che esista un $M > 0$ (dipendente dalla funzione f) per cui $f(x) = 0$ per ogni $|x| > M$. Lo spazio V è uno spazio vettoriale, dotato del prodotto scalare definito positivo

$$\langle f, g \rangle = \int_{\mathbb{R}} f(t)g(t)dt.$$

L'integrale ha senso (ed il risultato è un numero finito) perché f e g sono entrambe a supporto compatto. Considera l'endomorfismo

$$T: V \rightarrow V, \quad T(f) = f''$$

che trasforma la funzione f nella sua derivata seconda f'' . Mostra che T è un endomorfismo autoaggiunto per V .

Suggerimento. Integra per parti.

Complementi

11.1. Dimostrazione del criterio di Cartesio. Dimostriamo il Teorema 11.3.5. Faremo uso di alcuni teoremi di analisi. Iniziamo dimostrando un paio di fatti sui polinomi. Tutti i polinomi in questa sezione sono a coefficienti reali.

Lemma 11.3.6. *Consideriamo un polinomio*

$$p(x) = a_n x^n + \cdots + a_0$$

con $a_n, a_0 \neq 0$ aventi tutte le radici reali. Il numero di radici positive (contate con molteplicità) di $p(x)$ ed il numero di cambiamenti di segno nei coefficienti hanno la stessa parità.

Dimostrazione. Siano $x_1, \dots, x_n$ le radici di $p(x)$, tutte diverse da zero perché $a_0 \neq 0$. Per il Corollario 1.4.10,

$$p(x) = a_n(x - x_1) \cdots (x - x_n).$$

Quindi $a_0 = a_n(-x_1) \cdots (-x_n)$. Ne deduciamo che $p(x)$ ha un numero pari di radici positive $\iff a_n$ e a_0 hanno lo stesso segno $\iff$ nella successione $a_n, \dots, a_0$ il segno cambia un numero pari di volte. $\square$

Il lemma seguente mostra che se un polinomio ha tutte le radici reali, anche la sua derivata le ha, ed il numero di radici positive può scendere al massimo di uno. Durante la dimostrazione vediamo inoltre che le radici della derivata devono stare in zone ben localizzate.

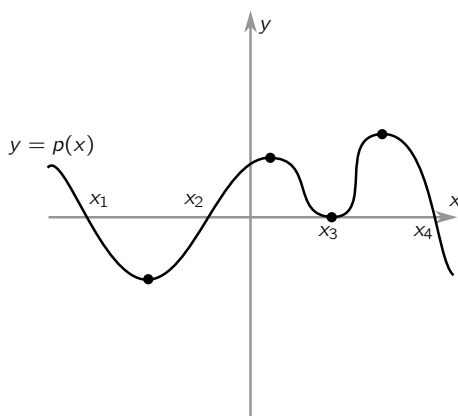


Figura 11.1. Se $p(x)$ è un polinomio con tutte radici reali, le radici della derivata $p'(x)$ sono localizzate in zone precise: la derivata $p'(x)$ ha le radici multiple di $p(x)$, ciascuna con una molteplicità in meno, ed esattamente una radice semplice tra due radici consecutive di $p(x)$.

Lemma 11.3.7. *Sia $p(x)$ un polinomio di grado n con tutte le radici reali e con k radici positive. La sua derivata $p'(x)$ ha anch'essa tutte le radici reali e le sue radici positive sono k oppure $k - 1$.*

Dimostrazione. Consideriamo le radici $x_1 < \dots < x_h$ di $p(x)$, ciascuna con una certa molteplicità $m_1, \dots, m_h$. Per ipotesi $m_1 + \dots + m_h = n$.

Per l'Esercizio 1.12, se $m_i \geq 2$ allora x_i è anche radice di $p'(x)$ con molteplicità $m_i - 1$. Per il Teorema di Rolle, tra due radici consecutive x_i e x_{i+1} c'è un punto di massimo o di minimo per $p(x)$ e quindi c'è almeno un'altra radice di $p'(x)$. Con queste due osservazioni abbiamo trovato che $p'(x)$ ha almeno

$$\sum_{i=1}^h (m_i - 1) + h - 1 = \sum_{i=1}^h m_i - 1 = n - 1$$

radici. Siccome $p'(x)$ ha grado $n - 1$, ne ha esattamente $n - 1$.

Inoltre sappiamo localizzare tutte le radici di $p'(x)$. Sono quelle multiple di $p(x)$, ciascuna con una molteplicità in meno, più una radice semplice all'interno di ogni intervallo (x_i, x_{i+1}) . Si veda la Figura 11.1.

Siano $x_j < \dots < x_h$ le radici positive di $p(x)$. Applicando lo stesso conto fatto sopra a queste radici, si deduce facilmente che $p'(x)$ ha $k - 1$ oppure k radici positive: la risposta precisa dipende dalla posizione della radice di $p'(x)$ che sta nell'intervallo (x_{j-1}, x_j) , che può essere positiva oppure no. $\square$

Dimostriamo adesso il criterio di Cartesio, il Teorema 11.3.5.

Dimostrazione. Il punto (1) è ovvio, quindi mostriamo il punto (2) per induzione sul grado n di $p(x)$. Se $n = 1$ allora $p(x) = a_1x + a_0$ e chiaramente esiste una radice positiva precisamente quando a_1 e a_0 hanno segno opposto.

Assumiamo il caso $n - 1$ e dimostriamo il criterio per n . Abbiamo

$$p(x) = a_nx^n + \cdots + a_0.$$

Se $a_0 = 0$ allora $p(x) = xq(x)$. Il polinomio $q(x)$ ha grado $n-1$, quindi per ipotesi induttiva soddisfa il punto (2). Poiché $p(x)$ e $q(x)$ hanno le stesse radici positive e gli stessi coefficienti, si deduce che anche $p(x)$ lo soddisfa.

Supponiamo $a_0 \neq 0$ e consideriamo la derivata

$$p'(x) = na_nx^{n-1} + \cdots + a_1.$$

Notiamo che i segni dei coefficienti $p'(x)$ sono proprio gli stessi di $p(x)$, meno l'ultimo, quello di a_0 . Per l'ipotesi induttiva il numero t di radici positive di $p'(x)$ è pari al numero t di cambiamenti di segno nella sequenza $a_n, \dots, a_1$. Per il lemma precedente $p(x)$ ha t oppure $t + 1$ radici positive. Inoltre ci sono t oppure $t + 1$ cambiamenti di segno nella sequenza $a_n, \dots, a_0$.

Il numero di radici positive in $p(x)$ ed il numero di cambiamenti di segno nella sequenza $a_n, \dots, a_0$ sono ciascuno t oppure $t + 1$, quindi sono uguali o differiscono di 1. Non possono però differire di un numero dispari per il Lemma 11.3.6. Quindi sono uguali (entrambi t oppure entrambi $t + 1$). $\square$

Geometria proiettiva

Introduciamo in questo capitolo una nuova geometria, chiamata *geometria proiettiva*. La geometria proiettiva è un arricchimento della geometria affine, realizzato aggiungendo allo spazio euclideo $\mathbb{R}^n$ alcuni “punti all’infinito”. Questo arricchimento presenta numerosi vantaggi, sia algebrici che geometrici. Renderemo in particolare rigoroso l’assunto informale che due rette parallele “si incontrano all’infinito”.

La geometria proiettiva sarà utilizzata nel capitolo seguente per esaminare coniche e quadriche secondo una prospettiva unificante. Vedremo ad esempio che ellissi, parabole e iperboli sono in realtà interpretabili come coniche aventi tutte lo stesso “tipo proiettivo”, ma posizionate in modo differente rispetto ai punti all’infinito del piano. Notiamo che questa interpretazione è utile ma non è strettamente necessaria per lo studio di coniche e quadriche ed è possibile saltare il presente capitolo e passare alla lettura del Capitolo 13.

12.1. Lo spazio proiettivo

Definiamo lo spazio proiettivo. Si tratta di una costruzione che può essere fatta a partire da qualsiasi spazio vettoriale V e che produce uno spazio con una sua geometria, con punti, rette, piani e più in generale sottospazi di dimensione arbitraria.

La costruzione dello spazio proiettivo richiede un certo livello di astrazione perché fa uso dell’insieme quoziente introdotto nella Sezione 1.1.10. Può essere utile la lettura preventiva della Sezione 10.1 sulle coordinate baricentriche, che sono molto simili alle *coordinate omogenee* che impiegheremo qui. L’idea di fondo è quella di aggiungere allo spazio euclideo $\mathbb{R}^n$ una coordinata in più: quando questa coordinata è nulla, otterremo dei punti nuovi, che stanno “all’infinito”. Per motivi teorici, la definizione precisa si discosta però molto da questa idea, che ritroveremo solo dopo alcune pagine.

12.1.1. Definizione. Sia V uno spazio vettoriale su un certo campo $\mathbb{K}$. Lo *spazio proiettivo su V* è l’insieme quoziente

$$\mathbb{P}(V) = (V \setminus \{0\})/\sim$$

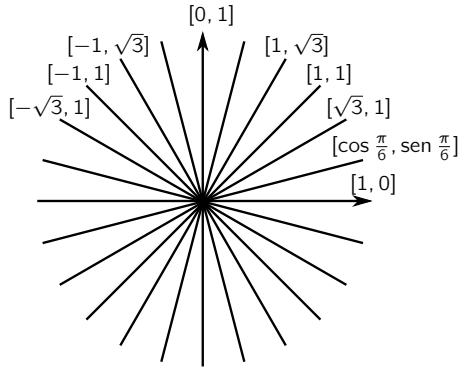


Figura 12.1. Ciascuna retta vettoriale in $\mathbb{R}^2$ è un punto in $\mathbb{P}(\mathbb{R}^2)$. Se la retta è del tipo $\text{Span}(v)$ con $v = (x_1, x_2)$, il punto corrispondente di $\mathbb{P}(\mathbb{R}^2)$ ha coordinate omogenee $[x_1, x_2]$. Alcune di queste sono mostrate in figura.

dove $\sim$ è la relazione di equivalenza che identifica due vettori non nulli precisamente quando sono multipli:

$$v \sim w \iff \exists \lambda \in \mathbb{K}, \lambda \neq 0 \text{ tale che } w = \lambda v.$$

Ricordiamo dalla Sezione 1.1.10 che una relazione di equivalenza $\sim$ induce un insieme quoziente, i cui elementi sono le classi di equivalenza di $\sim$. Un elemento di $\mathbb{P}(V)$ è una classe di equivalenza di vettori non nulli di V . Notiamo che abbiamo escluso l'origine di V nella definizione di $\mathbb{P}(V)$.

Esiste un'altra interpretazione di $\mathbb{P}(V)$. Ciascun vettore $v \in V$ non nullo genera una retta $r = \text{Span}(v)$. Due vettori non nulli v, w generano la stessa retta precisamente quando sono multipli, cioè quando $v \sim w$. Quindi lo spazio $\mathbb{P}(V)$ può essere interpretato semplicemente come l'insieme delle rette vettoriali di V . Si veda la Figura 12.1 per $V = \mathbb{R}^2$.

La *dimensione* dello spazio proiettivo $\mathbb{P}(V)$ è definita come

$$\dim \mathbb{P}(V) = \dim V - 1.$$

Osservazione 12.1.1. Se $\dim V = 1$, allora V è una retta e $\mathbb{P}(V)$ consta di un punto solo. Questo è coerente con il fatto che $\dim \mathbb{P}(V) = 1 - 1 = 0$.

Lo spazio $\mathbb{P}(\mathbb{R}^2)$ può essere interpretato come l'insieme delle rette vettoriali in $\mathbb{R}^2$, si veda la Figura 12.1. Intuitivamente, queste rette sono determinate da un solo parametro (ad esempio l'angolo θ che formano con l'asse x) e formano uno spazio unidimensionale. Infatti $\dim \mathbb{P}(\mathbb{R}^2) = 2 - 1 = 1$.

12.1.2. Coordinate omogenee. Cerchiamo adesso di dare un nome concreto agli elementi di $\mathbb{P}(V)$. Ciascun vettore non nullo $v \in V$ rappresenta un elemento di $\mathbb{P}(V)$ che indichiamo con $[v]$. Ricordiamo che per

definizione

$$[v] = [w] \iff \exists \lambda \in \mathbb{K}, \lambda \neq 0 \text{ tale che } w = \lambda v.$$

Sia $v_1, \dots, v_n$ una base di V . Dati $x_1, \dots, x_n \in \mathbb{K}$ non tutti nulli, possiamo costruire il vettore non nullo $v = x_1 v_1 + \dots + x_n v_n$ e conseguentemente il punto

$$P = [v] = [x_1 v_1 + \dots + x_n v_n] \in \mathbb{P}(V).$$

Diciamo che $x_1, \dots, x_n$ sono le *coordinate omogenee* del punto P rispetto alla base $v_1, \dots, v_n$. Generalmente le coordinate omogenee sono indicate fra parentesi quadre:

$$[x_1, \dots, x_n].$$

Dobbiamo fare adesso un'osservazione cruciale: le coordinate omogenee di P non sono uniche. Due coordinate omogenee

$$[x_1, \dots, x_n], \quad [y_1, \dots, y_n]$$

identificano lo stesso punto $P \in \mathbb{P}(V)$ se e solo se sono multiple, cioè se esiste un $\lambda \neq 0$ tale che $y_i = \lambda x_i$ per ogni $i = 1, \dots, n$. Questo è dovuto al fatto che v e λv identificano lo stesso punto $[v] = [\lambda v]$ in $\mathbb{P}(V)$. Lavorando con gli spazi proiettivi dobbiamo convivere con questa ambiguità: coordinate diverse possono identificare lo stesso punto. Le coordinate proiettive sono simili alle coordinate baricentriche introdotte nella Sezione 10.1.

Ricordiamo anche che per ipotesi le coordinate omogenee $[x_1, \dots, x_{n+1}]$ non sono mai tutte nulle.

12.1.3. Lo spazio proiettivo $\mathbb{P}^n(\mathbb{K})$. Così come $\mathbb{K}^n$ è il prototipo di spazio vettoriale n -dimensionale, analogamente esiste un modello di spazio proiettivo n -dimensionale, che indichiamo con $\mathbb{P}^n(\mathbb{K})$. La definizione è la seguente:

$$\mathbb{P}^n(\mathbb{K}) = \mathbb{P}(\mathbb{K}^{n+1}).$$

Notiamo che secondo la nostra definizione

$$\dim \mathbb{P}^n(\mathbb{K}) = \dim \mathbb{K}^{n+1} - 1 = n + 1 - 1 = n.$$

Ciascun punto di $\mathbb{P}^n(\mathbb{K})$ è descritto tramite coordinate omogenee

$$[x_1, \dots, x_{n+1}]$$

rispetto alla base canonica di $\mathbb{K}^{n+1}$. Per ogni $x \in \mathbb{K}^{n+1}$ non nullo, le coordinate omogenee di $[x] \in \mathbb{P}^n(\mathbb{K})$ sono semplicemente $[x_1, \dots, x_{n+1}]$. Ricordiamo ancora che le coordinate $[x_1, \dots, x_{n+1}]$ e $[\lambda x_1, \dots, \lambda x_{n+1}]$ indicano lo stesso punto di $\mathbb{P}^n(\mathbb{K})$, per ogni $\lambda \neq 0$.

Gli spazi $\mathbb{P}^1(\mathbb{K})$, $\mathbb{P}^2(\mathbb{K})$ e $\mathbb{P}^3(\mathbb{K})$ sono particolarmente importanti per noi e sono chiamati la *retta proiettiva*, il *piano proiettivo* e lo *spazio proiettivo*. Siamo ovviamente interessati soprattutto al campo $\mathbb{K} = \mathbb{R}$ reale. Per semplicità indicheremo a volte $\mathbb{P}^n(\mathbb{R})$ con $\mathbb{P}^n$.

Nella retta proiettiva $\mathbb{P}^1(\mathbb{K})$ ciascun punto ha due coordinate omogenee $[x_1, x_2]$. Si veda la Figura 12.1 nel caso $\mathbb{K} = \mathbb{R}$. Ad esempio, $[1, -1]$ e $[-2, 2]$ identificano lo stesso punto mentre $[1, -1]$ e $[1, 0]$ sono chiaramente punti diversi perché le coordinate non sono multiple.

Nel piano proiettivo $\mathbb{P}^2(\mathbb{K})$ ogni punto ha tre coordinate omogenee $[x_1, x_2, x_3]$, e le coordinate $[1, -3, 2]$ e $[5, -15, 10]$ indicano in realtà lo stesso punto.

12.1.4. Sottospazio proiettivo. Sia V uno spazio vettoriale e $\mathbb{P}(V)$ il corrispondente spazio proiettivo. Ciascun sottospazio vettoriale $W \subset V$ definisce un *sottospazio proiettivo* $\mathbb{P}(W) \subset \mathbb{P}(V)$ nel modo ovvio: il sottospazio $\mathbb{P}(W)$ è l'insieme delle classi $[w]$ dei vettori $w \in W$.

Notiamo che $\dim \mathbb{P}(W) = \dim(W) - 1$ e $\dim \mathbb{P}(V) = \dim V - 1$. Quindi

$$\dim \mathbb{P}(V) - \dim \mathbb{P}(W) = \dim V - \dim W.$$

Questo numero è chiamato la *codimensione* di $\mathbb{P}(V)$ in $\mathbb{P}(W)$. Un *iperpiano* è un sottospazio proiettivo di codimensione 1.

Consideriamo lo spazio $\mathbb{P}^n(\mathbb{K})$. Un sottospazio proiettivo di $\mathbb{P}^n(\mathbb{K})$ è del tipo $\mathbb{P}(W)$ dove $W \subset \mathbb{K}^n$ è un sottospazio vettoriale. Se il sottospazio $W \subset \mathbb{K}^n$ è descritto in forma cartesiana come luogo di zeri di un sistema di equazioni omogenee

$$Ax = 0$$

allora $\mathbb{P}(W)$ è descritto semplicemente dalle stesse equazioni $Ax = 0$ in coordinate omogenee. Così $\mathbb{P}(W)$ è descritto in *forma cartesiana*.

Notiamo che è sempre possibile prendere un numero di equazioni pari alla codimensione di $\mathbb{P}(W)$. In particolare un iperpiano proiettivo è descritto da una sola equazione. Ad esempio, nel piano proiettivo $\mathbb{P}^2(\mathbb{K})$ un punto ha coordinate omogenee $[x_1, x_2, x_3]$ ed una retta in $\mathbb{P}^2(\mathbb{K})$ è descritta da una equazione omogenea

$$a_1x_1 + a_2x_2 + a_3x_3 = 0.$$

Analogamente nello spazio proiettivo $\mathbb{P}^3(\mathbb{K})$ i punti hanno coordinate omogenee $[x_1, x_2, x_3, x_4]$ ed un piano è descritto da una equazione omogenea

$$a_1x_1 + a_2x_2 + a_3x_3 + a_4x_4 = 0.$$

Una retta proiettiva in $\mathbb{P}^3(\mathbb{K})$ è data da due equazioni indipendenti

$$\begin{cases} a_{11}x_1 + a_{12}x_2 + a_{13}x_3 + a_{14}x_4 = 0, \\ a_{21}x_1 + a_{22}x_2 + a_{23}x_3 + a_{24}x_4 = 0. \end{cases}$$

Esempio 12.1.2. In geometria proiettiva si tenta di disegnare per quanto possibile i punti, le rette e i piani come nella usuale geometria affine. I disegni aiutano a mettere in moto la nostra intuizione geometrica; il collegamento fra spazi euclidei e spazi proiettivi sarà più chiaro fra poco, quando interpreteremo $\mathbb{P}^n$ come lo spazio euclideo $\mathbb{R}^n$ con dei punti aggiuntivi all'infinito.

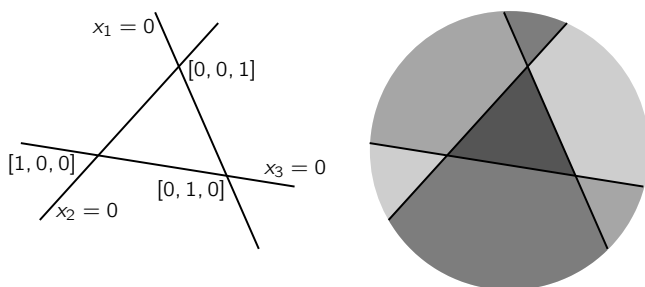


Figura 12.2. Tre punti e tre rette in $\mathbb{P}^2$ (sinistra). Le rette tagliano il piano proiettivo in quattro triangoli: le zone con lo stesso colore sono in realtà parte dello stesso triangolo che "attraversa l'infinito" (destra).

Ad esempio, abbiamo disegnato nella Figura 12.2-(sinistra) le rette

$$\{x_1 = 0\}, \quad \{x_2 = 0\}, \quad \{x_3 = 0\}$$

nel piano proiettivo reale $\mathbb{P}^2$. Queste si intersecano a coppie nei punti $[1, 0, 0]$, $[0, 1, 0]$ e $[0, 0, 1]$ come mostrato in figura.

Osservazione 12.1.3. Come abbiamo detto, le coordinate omogenee sono molto simili alle coordinate baricentriche descritte nella Sezione 10.1. Ad una attenta analisi si trovano però delle stranezze: in realtà le tre rette disegnate nella Figura 12.2-(sinistra) dividono il piano proiettivo $\mathbb{P}^2$ in quattro triangoli, quello centrale ed altri tre esterni come suggerito nella Figura 12.2-(destra). Come vedremo fra poco, il piano proiettivo $\mathbb{P}^2$ è come il piano euclideo $\mathbb{R}^2$ a cui siano stati aggiunti opportunamente dei "punti all'infinito" che permettono agli oggetti di attraversare l'infinito ribucando dall'altra parte dello spazio: in questo modo le due porzioni aventi lo stesso colore in figura sono in realtà pezzi di un unico triangolo.

Osservazione 12.1.4. Le equazioni che descrivono i sottospazi proiettivi di $\mathbb{P}^n(\mathbb{K})$ in forma cartesiana sono tutte omogenee. È importante notare che una equazione non omogenea *non* ha senso nello spazio proiettivo! Nel piano $\mathbb{P}^2$, l'equazione omogenea $x_1 + x_2 - x_3 = 0$ ha senso ed identifica una retta, mentre l'equazione non omogenea $x_1 + x_2 - x_3 = 3$ non ha senso, perché ambigua: il punto $[1, 1, -1]$ la soddisfa mentre $[2, 2, -2]$ no, però sono coordinate diverse che indicano lo stesso punto! Per questo motivo con lo spazio proiettivo tutte le equazioni che scriveremo saranno omogenee.

12.1.5. Intersezioni. Consideriamo uno spazio proiettivo $\mathbb{P}(V)$ e due sottospazi proiettivi $\mathbb{P}(U)$ e $\mathbb{P}(W)$. Studiamo adesso la loro intersezione.

Ricordiamo che per definizione U e W sono sottospazi vettoriali di V . Ci sono due casi possibili da considerare. Se U e W sono in somma diretta,

cioè se $U \cap W = \{0\}$, allora l'intersezione fra i sottospazi proiettivi è vuota:

$$\mathbb{P}(U) \cap \mathbb{P}(W) = \emptyset.$$

In questo caso si dice che i due sottospazi proiettivi sono *sghembi* o semplicemente *disgiunti*.

Se invece $U \cap W$ è un sottospazio vettoriale di dimensione almeno uno, allora i due sottospazi proiettivi si intersecano, e la loro intersezione è sempre un sottospazio proiettivo:

$$\mathbb{P}(U) \cap \mathbb{P}(W) = \mathbb{P}(U \cap W).$$

In questo caso diciamo che i due sottospazi proiettivi sono *incidenti*.

Esempio 12.1.5. I due piani $\{x_1 + x_2 = 0\}$ e $\{x_2 + x_4 = 0\}$ di $\mathbb{P}^3$ si intersecano in una retta r , descritta in forma cartesiana dal sistema

$$\begin{cases} x_1 + x_2 = 0, \\ x_2 + x_4 = 0. \end{cases}$$

Notiamo ad esempio che i punti $[0, 0, 1, 0]$ e $[1, -1, 0, 1]$ appartengono entrambi alla retta r .

D'altra parte, le due rette proiettive r e s in $\mathbb{P}^3$ descritte dai sistemi

$$r = \begin{cases} x_1 + x_2 = 0, \\ x_2 + x_4 = 0, \end{cases} \quad s = \begin{cases} x_1 = 0, \\ x_3 = 0 \end{cases}$$

non si intersecano e quindi sono sghembe. Infatti affiancando i due sistemi si ottiene come unica soluzione quella banale $x_1 = x_2 = x_3 = x_4 = 0$, ma ricordiamo che un punto con coordinate omogenee tutte nulle non esiste.

La lettrice avrà notato che due sottospazi possono essere solo incidenti o sghembi: non è prevista, come nel caso affine, la possibilità che siano paralleli. Effettivamente una peculiarità importante della geometria proiettiva è la totale assenza di parallelismi: due sottospazi non sono mai paralleli, in nessun senso. In particolare, due rette nel piano si intersecano sempre.

Proposizione 12.1.6. *Due rette proiettive distinte in $\mathbb{P}^2(\mathbb{K})$ si intersecano sempre in un punto.*

Dimostrazione. Ricordiamo che $\mathbb{P}^2(\mathbb{K}) = \mathbb{P}(\mathbb{K}^3)$. L'enunciato è conseguenza del fatto che due piani vettoriali distinti in $\mathbb{K}^3$ si intersecano sempre in una retta vettoriale. $\square$

Concretamente, se le due rette in $\mathbb{P}^2(\mathbb{K})$ sono descritte dalle equazioni

$$a_1x_1 + a_2x_2 + a_3x_3 = 0, \quad b_1x_1 + b_2x_2 + b_3x_3 = 0$$

allora la loro intersezione è descritta dal sistema

$$\begin{cases} a_1x_1 + a_2x_2 + a_3x_3 = 0, \\ b_1x_1 + b_2x_2 + b_3x_3 = 0. \end{cases}$$

I coefficienti formano una matrice 2×3 di rango due: siccome le rette sono distinte, le due righe sono indipendenti. Per la Proposizione 3.4.15, le due rette si intersecano nel punto

$$[a_2b_3 - a_3b_2, a_3b_1 - a_1b_3, a_1b_2 - a_2b_1].$$

Osservazione 12.1.7. Nel piano proiettivo $\mathbb{P}^2$ non vale il V Postulato di Euclide: dati una retta $r \subset \mathbb{P}^2$ ed un punto $P \in \mathbb{P}^2$ esterno ad essa, non esiste nessuna retta s passante per P che sia disgiunta da r . La geometria proiettiva è una geometria non euclidea.

Analogamente si dimostra il fatto seguente.

Proposizione 12.1.8. *Due piani proiettivi distinti in $\mathbb{P}^3(\mathbb{K})$ si intersecano sempre in una retta.*

Con le stesse tecniche possiamo studiare le possibili posizioni di una retta e un piano nello spazio

Esercizio 12.1.9. Siano π e r un piano ed una retta proiettiva in $\mathbb{P}^3(\mathbb{K})$. Si possono presentare due casi:

- $r \subset \pi$,
- $r \not\subset \pi$ e quindi r e π si intersecano in un punto.

Osservazione 12.1.10. Notiamo anche qui l'assenza di parallelismi. I teoremi nella geometria proiettiva sono spesso *più semplici* che nella geometria affine: nella geometria affine la possibilità che alcuni sottospazi siano paralleli ha tipicamente l'effetto di complicare gli enunciati, che devono tenere conto di un numero maggiore di configurazioni particolari.

12.1.6. Sottospazio generato. Come nell'ambito vettoriale e in quello affine, anche in geometria proiettiva è possibile definire un sottospazio in forma *parametrica*, cioè come insieme generato da alcuni elementi e descritto da parametri variabili in un certo insieme.

Sia V uno spazio vettoriale e $\mathbb{P}(V)$ lo spazio proiettivo associato. Sia $X \subset \mathbb{P}(V)$ un sottoinsieme qualsiasi. Il *sottospazio proiettivo generato* da X è l'intersezione S di tutti i sottospazi proiettivi contenenti X . L'intersezione S è effettivamente un sottospazio proiettivo.

Se X è un numero finito di punti $P_0, \dots, P_k$, il sottospazio da loro generato può essere descritto esplicitamente. Ricordiamo che ciascun P_i è rappresentato da un vettore $v_i \in V$, unico a meno di riscalamento, e scriviamo $P_i = [v_i]$.

Proposizione 12.1.11. *Se $X = \{P_0, \dots, P_k\}$ e $P_i = [v_i] \forall i$, allora*

$$S = \mathbb{P}(\text{Span}(v_0, \dots, v_k))$$

$$= \{[t_0v_0 + \dots + t_kv_k] \mid t_0, \dots, t_k \in \mathbb{K} \text{ non tutti nulli}\}.$$

Dimostrazione. Per definizione, abbiamo $S = \mathbb{P}(W)$ dove W è l'intersezione di tutti i sottospazi di V che contengono i vettori $v_0, \dots, v_k$. Allora $W = \text{Span}(v_0, \dots, v_k)$. $\square$

Notiamo che in questo caso

$$\dim S = \dim(\text{Span}(v_0, \dots, v_k)) - 1 \leq k + 1 - 1 = k.$$

Lo spazio S generato da $k + 1$ punti ha dimensione al massimo k . Se $\dim S = k$ allora diciamo che i $k + 1$ punti sono *proiettivamente indipendenti*. Questo accade se e solo se i rappresentanti $v_0, \dots, v_k$ sono vettori indipendenti in V .

Corollario 12.1.12. *Dati $k + 1$ punti proiettivamente indipendenti in $\mathbb{P}(V)$ esiste un unico sottospazio proiettivo S di dimensione k che li contiene (e nessun sottospazio di dimensione $< k$).*

Vediamo alcuni esempi. Due punti $P, Q \in \mathbb{P}^n(\mathbb{K})$ sono proiettivamente indipendenti se e solo se sono distinti. Due punti P, Q distinti generano una retta r , descritta in forma parametrica nel modo seguente. Se

$$P = [x_1, \dots, x_{n+1}], \quad Q = [y_1, \dots, y_{n+1}]$$

allora

$$r = \{[tx_1 + uy_1, \dots, tx_{n+1} + uy_{n+1}] \mid t, u \in \mathbb{K} \text{ non entrambi nulli}\}.$$

Notiamo che le coppie (t, u) e $(\lambda t, \lambda u)$ parametrizzano lo stesso punto: solo le coordinate omogenee $[t, u]$ importano veramente e quindi possiamo scrivere

$$r = \{[tx_1 + uy_1, \dots, tx_{n+1} + uy_{n+1}] \mid [t, u] \in \mathbb{P}^1(\mathbb{K})\}.$$

In questo modo la retta proiettiva r è parametrizzata dalla retta proiettiva standard $\mathbb{P}^1(\mathbb{K})$, come è giusto che sia.

In modo analogo si descrivono in forma parametrica dei sottospazi di dimensione più alta. Tre punti in $\mathbb{P}^3(\mathbb{K})$ sono proiettivamente indipendenti se e solo se non sono contenuti in una retta, cioè se non sono *allineati*. Come in geometria affine, si può passare da una descrizione in forma cartesiana ad una in forma parametrica e viceversa usando gli strumenti dell'algebra lineare.

Esempio 12.1.13. La retta r in $\mathbb{P}^2$ passante per $P = [1, 1, 0]$ e $Q = [0, 1, -1]$ è data dai punti

$$r = \{[t, t + u, -u] \mid [t, u] \in \mathbb{P}^1\}.$$

Si passa facilmente ad una scrittura cartesiana notando che

$$r = \{-x_1 + x_2 + x_3 = 0\}.$$

Il piano π in $\mathbb{P}^3$ passante per i tre punti $P = [1, -1, 0, 0]$, $Q = [1, 0, 1, 0]$ e $R = [0, 0, 0, 1]$ in forma parametrica è la seguente

$$\pi = \{[s + t, -s, t, u] \mid [s, t, u] \in \mathbb{P}^2\}.$$

In forma cartesiana diventa

$$\pi = \{x_1 + x_2 - x_3 = 0\}.$$

12.1.7. Formula di Grassmann. Nel passare dalla geometria vettoriale a quella affine abbiamo perso la formula di Grassmann: questa formula infatti vale per i sottospazi vettoriali ma *non* vale per i sottospazi affini quando questi sono disgiunti; questa perdita è essenzialmente dovuta alla presenza di fenomeni di parallelismo, che complicano un po' le cose. Vediamo adesso che la formula torna ad essere valida in geometria proiettiva, perché il parallelismo qui è stato eliminato: questo è uno dei suoi pregi.

Per enunciare la formula di Grassmann dobbiamo definire la somma di due sottospazi proiettivi. Sia V uno spazio vettoriale e $\mathbb{P}(V)$ lo spazio proiettivo associato. Siano S_1 e S_2 due sottospazi proiettivi di $\mathbb{P}(V)$. Il *sottospazio somma* $S_1 + S_2$ è il sottospazio proiettivo generato dall'unione $S_1 \cup S_2$. Ad esempio, se S_1 e S_2 sono due punti, allora $S_1 + S_2$ è la retta che li contiene; se S_1 è un punto e S_2 è una retta in $\mathbb{P}^3(\mathbb{K})$ che non contiene S_1 , allora $S_1 + S_2$ è l'unico piano che contiene S_1 e S_2 .

Diciamo per comodità che l'insieme vuoto ha dimensione -1 . In questo modo, se $S_1 \cap S_2 = \emptyset$ allora $\dim(S_1 \cap S_2) = -1$.

Proposizione 12.1.14 (Formula di Grassmann). *Per qualsiasi coppia di sottospazi proiettivi S_1, S_2 in $\mathbb{P}(V)$ vale*

$$\dim(S_1 + S_2) + \dim(S_1 \cap S_2) = \dim S_1 + \dim S_2.$$

Dimostrazione. Segue facilmente dall'analogo vettoriale. Abbiamo $S_1 = \mathbb{P}(W_1)$, $S_2 = \mathbb{P}(W_2)$, $S_1 \cap S_2 = \mathbb{P}(W_1 \cap W_2)$, $S_1 + S_2 = \mathbb{P}(W_1 + W_2)$ quindi l'enunciato è conseguenza della formula di Grassmann per W_1 e W_2

$$\dim(W_1 + W_2) + \dim(W_1 \cap W_2) = \dim W_1 + \dim W_2.$$

Aggiungendo -1 a ciascun addendo si ottiene la formula proiettiva. $\square$

12.2. Completamento proiettivo

Abbiamo accennato al fatto che lo spazio proiettivo possa essere ottenuto aggiungendo i "punti all'infinito" allo spazio euclideo. Spieghiamo adesso rigorosamente questo fenomeno.

12.2.1. Le carte affini. Consideriamo lo spazio vettoriale $\mathbb{K}^n$ e lo spazio proiettivo $\mathbb{P}^n(\mathbb{K})$. Mostriamo come sia possibile vedere il primo come un sottoinsieme del secondo.

Fissiamo un $i \in \{1, \dots, n+1\}$. Definiamo il sottoinsieme di $\mathbb{P}^n(\mathbb{K})$

$$E_i = \{[x_1, \dots, x_{n+1}] \mid x_i \neq 0\}.$$

Notiamo che la condizione $x_i \neq 0$ è ben posta, perché non cambia se moltiplichiamo le coordinate omogenee per una costante non nulla (ricordiamo che le coordinate omogenee sono definite solo a meno di una costante moltiplicativa non nulla). D'altro canto, consideriamo l'iperpiano

$$H_i = \{[x_1, \dots, x_{n+1}] \mid x_i = 0\}.$$

Otteniamo una partizione dello spazio $\mathbb{P}^n(\mathbb{K})$ in due insiemi disgiunti:

$$(23) \quad \mathbb{P}^n(\mathbb{K}) = E_i \cup H_i.$$

Consideriamo adesso la mappa

$$f_i: \mathbb{K}^n \longrightarrow E_i$$

definita nel modo seguente:

$$f_i(x_1, \dots, x_n) = [x_1, \dots, x_{i-1}, 1, x_i, \dots, x_n].$$

La mappa inserisce una coordinata addizionale costantemente uguale a 1 prima della posizione i -esima.

Proposizione 12.2.1. *La mappa f_i è una bigezione. La sua inversa è*

$$f_i^{-1}[y_1, \dots, y_{n+1}] = \left(\frac{y_1}{y_i}, \dots, \frac{y_{i-1}}{y_i}, \frac{y_{i+1}}{y_i}, \dots, \frac{y_{n+1}}{y_i} \right).$$

Dimostrazione. Notiamo che effettivamente $f_i^{-1}: E_i \rightarrow \mathbb{K}^n$ è ben definita perché $y_i \neq 0$ in E_i . Si verifica facilmente che f_i^{-1} è l'inversa di f_i e quindi f_i è una bigezione. $\square$

Siccome f_i è una bigezione, possiamo identificare gli insiemi $\mathbb{K}^n$ e E_i tramite f_i . Concretamente, il punto $(x_1, \dots, x_n)$ di $\mathbb{K}^n$ viene identificato con la sua immagine $[x_1, \dots, x_{i-1}, 1, x_i, \dots, x_n]$. La mappa f_i è chiamata la *i -esima carta affine*. Il termine "carta" è in analogia con le carte geografiche, che sono mappe bigettive da $\mathbb{R}^2$ (o da una porzione di $\mathbb{R}^2$) su una parte della superficie del globo terrestre.

Ora che abbiamo identificato $\mathbb{K}^n$ e E_i , possiamo riscrivere (23) così:

$$\mathbb{P}^n(\mathbb{K}) = \mathbb{K}^n \cup H_i.$$

Come promesso, in questo modo possiamo interpretare lo spazio proiettivo $\mathbb{P}^n(\mathbb{K})$ come lo spazio vettoriale $\mathbb{K}^n$ dotato di alcuni punti aggiuntivi, quelli dell'iperpiano H_i . I punti aggiuntivi di H_i sono chiamati i *punti all'infinito* e l'iperpiano H_i è l'*iperpiano all'infinito*.

Esaminiamo adesso queste nozioni geometricamente nei casi che ci interessano. Tutta la costruzione dipende da i , e generalmente per semplicità si sceglie $i = n + 1$. Scriviamo quindi

$$\mathbb{P}^n(\mathbb{K}) = \mathbb{K}^n \cup H$$

con $H = H_{n+1} = \{x_{n+1} = 0\}$. Il punto $(x_1, \dots, x_n) \in \mathbb{K}^n$ viene identificato con $[x_1, \dots, x_n, 1] \in \mathbb{P}^n(\mathbb{K})$. L'iperpiano all'infinito $H = \{x_{n+1} = 0\}$ contiene tutti i punti del tipo $[x_1, \dots, x_n, 0]$. Ricordiamo che un iperpiano è un sottospazio proiettivo di codimensione 1.

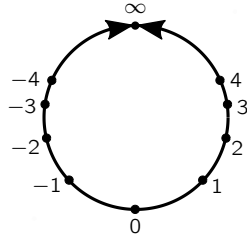


Figura 12.3. La retta proiettiva $\mathbb{P}^1$ è ottenuta aggiungendo un punto all'infinito alla retta reale $\mathbb{R}$.

12.2.2. La retta proiettiva. Studiamo la retta proiettiva reale $\mathbb{P}^1$. Per quanto appena visto, la retta proiettiva si decompone come

$$\mathbb{P}^1 = \mathbb{R} \cup H$$

dove $x \in \mathbb{R}$ è identificato con $[x, 1] \in \mathbb{P}^1$. L'iperpiano H all'infinito è un sottospazio in $\mathbb{P}^1$ di codimensione 1, cioè un punto:

$$H = \{[x_1, 0] \in \mathbb{P}^1 \mid x_1 \neq 0\} = \{[1, 0]\}.$$

La retta proiettiva $\mathbb{P}^1$ è quindi formata da punti del tipo $[x, 1]$, che possono essere identificati con la retta reale $\mathbb{R}$, più un punto aggiuntivo $[1, 0]$, il *punto all'infinito*, che può essere indicato con il simbolo ∞ . La retta proiettiva ha in realtà la forma di una circonferenza come nella Figura 12.3.

12.2.3. Il piano proiettivo. Consideriamo adesso il piano proiettivo reale $\mathbb{P}^2$. Questo si decompone come

$$\mathbb{P}^2 = \mathbb{R}^2 \cup H$$

dove $(x_1, x_2) \in \mathbb{R}^2$ è identificato con $[x_1, x_2, 1] \in \mathbb{P}^2$ e H è la retta all'infinito

$$H = \{[x_1, x_2, 0] \in \mathbb{P}^2\} = \{x_3 = 0\}.$$

Il piano proiettivo $\mathbb{P}^2$ è quindi formato dai punti del tipo $[x_1, x_2, 1]$, che possono essere identificati con $\mathbb{R}^2$, più una *retta all'infinito* aggiuntiva, formata da punti del tipo $[x_1, x_2, 0]$ al variare di $[x_1, x_2] \in \mathbb{P}^1$.

A differenza di $\mathbb{P}^1$, nel piano proiettivo non abbiamo un solo punto all'infinito, ma ne abbiamo infiniti, e formano a loro volta una retta proiettiva. Ad esempio, i punti seguenti

$$[0, 1, 0], \quad [1, 1, 0], \quad [1, 2, 0], \quad [1, 0, 0]$$

sono quattro punti distinti, tutti all'infinito. Qual è il loro significato geometrico? Cerchiamo di rispondere a questa domanda. Geometricamente, esiste un punto all'infinito per ogni retta passante per l'origine: alla retta $r = \text{Span}(v)$ con $v = (x_1, x_2)$ corrisponde il punto all'infinito $[x_1, x_2, 0]$. Il

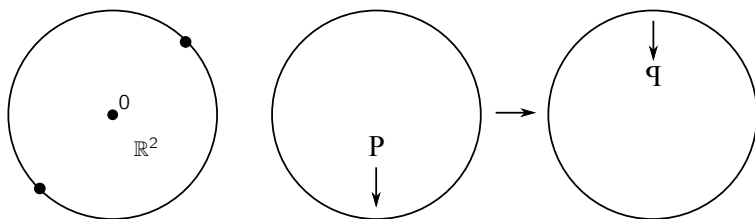


Figura 12.4. Il piano proiettivo $\mathbb{P}^2$ è ottenuto aggiungendo i “punti all’infinito” a $\mathbb{R}^2$. Precisamente, viene aggiunto un punto ad ogni retta passante per l’origine. Possiamo pensare a $\mathbb{R}^2$ come ad un cerchio senza contorno, a cui abbiamo aggiunto dei punti al contorno: due punti opposti sono però in realtà lo stesso punto (sinistra). È possibile per un oggetto attraversare l’infinito e trovarsi quindi nella parte opposta del disco: attenzione però che l’attraversamento causa uno specchiamento dell’oggetto! Questa descrizione spiega perché nella Figura 12.2-(destra) il piano proiettivo è diviso in 4 triangoli (destra).

piano proiettivo $\mathbb{P}^2$ è ottenuto dal piano euclideo $\mathbb{R}^2$ aggiungendo un punto all’infinito per ogni retta passante per l’origine. Si veda la Figura 12.4.

12.2.4. Lo spazio proiettivo. Come nei casi precedenti, abbiamo

$$\mathbb{P}^3 = \mathbb{R}^3 \cup H$$

dove $(x_1, x_2, x_3) \in \mathbb{R}^3$ è identificato con $[x_1, x_2, x_3, 1]$ e $H = \{x_4 = 0\}$ è il piano dei punti all’infinito. Anche qui i punti all’infinito corrispondono alle rette di $\mathbb{R}^3$ passanti per l’origine: per ogni retta ce n’è uno.

Si può rappresentare la forma di $\mathbb{P}^3$ alla stessa maniera della Figura 12.4, disegnando $\mathbb{R}^3$ come una sfera solida tridimensionale senza contorno e aggiungendo i punti al contorno, in cui però due punti antipodali risultano identificati. Si può verificare con un po’ di pazienza che in questo caso l’attraversamento dell’infinito *non* causa uno specchiamento dell’oggetto!

12.2.5. Completamento di un sottospazio affine. Mostriamo adesso che i sottospazi affini di $\mathbb{K}^n$ diventano sottospazi proiettivi di $\mathbb{P}^n(\mathbb{K})$ tramite un procedimento di *completamento*, che consiste nell’aggiunta dei loro punti all’infinito.

Consideriamo la partizione

$$\mathbb{P}^n(\mathbb{K}) = \mathbb{K}^n \cup H.$$

Ricordiamo che un punto $(x_1, \dots, x_n) \in \mathbb{K}^n$ è identificato con $[x_1, \dots, x_n, 1] \in \mathbb{P}^n(\mathbb{K})$ e $H = \{x_{n+1} = 0\}$ è l’iperpiano all’infinito.

Sia $S \subset \mathbb{K}^n$ uno sottospazio affine, descritto in forma cartesiana come

$$(24) \quad \begin{cases} a_{11}x_1 + \cdots + a_{1n}x_n + b_1 = 0, \\ \vdots \\ a_{k1}x_1 + \cdots + a_{kn}x_n + b_k = 0. \end{cases}$$

Come sottolineato nell'Osservazione 12.1.4, queste equazioni non sono omogenee e quindi non hanno senso in $\mathbb{P}^n(\mathbb{K})$. Rimediamo facilmente, trasformandole in equazioni omogenee nel modo più naïf: aggiungiamo una nuova coordinata x_{n+1} e la assegniamo ai termini noti b_i . Il *completamento proiettivo* di S in $\mathbb{P}^n(\mathbb{K})$ è il sottospazio proiettivo $\bar{S}$ di $\mathbb{P}^n(\mathbb{K})$ definito in forma cartesiana dalle equazioni omogenee

$$(25) \quad \begin{cases} a_{11}x_1 + \cdots + a_{1n}x_n + b_1x_{n+1} = 0, \\ \vdots \\ a_{k1}x_1 + \cdots + a_{kn}x_n + b_kx_{n+1} = 0. \end{cases}$$

Il processo con cui abbiamo trasformato le equazioni da (24) a (25) si chiama *omogeneizzazione*. Ricordiamo che $\mathbb{P}^n(\mathbb{K}) = \mathbb{K}^n \cup H$. Studiamo adesso le intersezioni di $\bar{S}$ con $\mathbb{K}^n$ e H .

Proposizione 12.2.2. *Valgono i fatti seguenti:*

- $\bar{S} \cap \mathbb{K}^n = S$.
- $\bar{S} \cap H$ è il sottospazio proiettivo di H descritto dal sistema

$$(26) \quad \begin{cases} a_{11}x_1 + \cdots + a_{1n}x_n = 0, \\ \vdots \\ a_{k1}x_1 + \cdots + a_{kn}x_n = 0, \\ x_{n+1} = 0. \end{cases}$$

Dimostrazione. I punti di $\bar{S} \cap \mathbb{K}^n$ sono i punti $[x_1, \dots, x_n, 1]$ che soddisfano il sistema (25), e questi sono precisamente i punti $(x_1, \dots, x_n)$ che soddisfano (24), cioè S .

Il sistema (26) descrive l'intersezione fra $\bar{S}$ e H . □

Il completamento proiettivo $\bar{S}$ consiste quindi nel sottospazio affine originario S , più alcuni punti aggiuntivi che stanno all'infinito: questi sono i "punti all'infinito di S " e sono descritti dal sistema (26). Possiamo scrivere

$$\bar{S} = S \cup S_\infty$$

dove $S_\infty = \bar{S} \cap H$ è il sottospazio proiettivo dell'iperpiano all'infinito H descritto dal sistema (26). Calcolando le dimensioni con Rouché - Capelli, troviamo

$$\dim \bar{S} = \dim S, \quad \dim S_\infty = \dim S - 1.$$

Se S è una retta affine i suoi punti all'infinito S_∞ sono in realtà un punto solo, mentre se S è un piano affine i suoi punti all'infinito S_∞ formano una retta proiettiva.

Esempio 12.2.3. Sia $r \subset \mathbb{R}^2$ una retta affine descritta dall'equazione $a_1x_1 + a_2x_2 + b = 0$. Il suo completamento $\bar{r} \subset \mathbb{P}^2$ è la retta proiettiva

$$a_1x_1 + a_2x_2 + bx_3 = 0.$$

Il completamento $\bar{r}$ aggiunge ad r il suo punto all'infinito r_∞ , soluzione di

$$a_1x_1 + a_2x_2 = 0, \quad x_3 = 0.$$

Quindi $r_\infty = [a_2, -a_1, 0]$. Abbiamo trovato

$$\bar{r} = r \cup \{[a_2, -a_1, 0]\}.$$

Possiamo finalmente dare sostanza all'idea intuitiva secondo cui due rette parallele si incontrano all'infinito. Due rette affini parallele hanno equazioni

$$a_1x_1 + a_2x_2 + b_1 = 0, \quad a_1x_1 + a_2x_2 + b_2 = 0$$

e sono distinte se $b_1 \neq b_2$. Queste due rette hanno chiaramente lo stesso punto all'infinito $[a_2, -a_1, 0]$. Due rette affini parallele distinte non si intersecano in $\mathbb{R}^2$, ma i loro completamenti si intersecano all'infinito.

Esempio 12.2.4. Sia $\pi \subset \mathbb{R}^3$ un piano affine descritto dall'equazione $a_1x_1 + a_2x_2 + a_3x_3 + b = 0$. Il suo completamento $\bar{\pi} \subset \mathbb{P}^3$ è il piano proiettivo

$$a_1x_1 + a_2x_2 + a_3x_3 + bx_4 = 0.$$

Il completamento $\bar{\pi}$ aggiunge a π una retta all'infinito, di equazioni

$$a_1x_1 + a_2x_2 + a_3x_3 = 0, \quad x_4 = 0.$$

Ad esempio, il piano affine $\{x_1 - 1 = 0\}$ in $\mathbb{R}^3$ si completa al piano proiettivo $\{x_1 - x_4 = 0\}$ in $\mathbb{P}^3$. I punti all'infinito che abbiamo aggiunto sono la retta

$$\{x_1 - x_4 = x_4 = 0\} = \{x_1 = 0, x_4 = 0\} = \{[0, t, u, 0] \mid [t, u] \in \mathbb{P}^1\}.$$

Questa è una retta contenuta nel piano all'infinito. Come sopra, due piani paralleli condividono la stessa retta all'infinito. Due piani paralleli distinti in $\mathbb{R}^3$ non si intersecano, ma i loro completamenti si intersecano in una retta all'infinito.

12.3. Proiettività

Ogni spazio che si rispetti ha le sue funzioni privilegiate, che in qualche senso rispettano la struttura dello spazio, e lo spazio proiettivo non fa eccezione. Queste si chiamano *proiettività*.

12.3.1. Definizione. Le affinità sono trasformazioni particolari dello spazio affine. Analogamente, le *proiettività* sono trasformazioni particolari dello spazio proiettivo.

Sia $f: V \rightarrow W$ un isomorfismo fra spazi vettoriali. Questo induce una applicazione fra i relativi spazi proiettivi $f_*: \mathbb{P}(V) \rightarrow \mathbb{P}(W)$ nel modo ovvio:

$$f_*([v]) = [f(v)].$$

Proposizione 12.3.1. *La funzione f_* è ben definita.*

Dimostrazione. Poiché f è iniettiva, abbiamo $v \neq 0 \Rightarrow f(v) \neq 0$ e quindi $[f(v)] \in \mathbb{P}(W)$ ha senso per ogni $v \neq 0$.

Inoltre, se rappresentiamo il punto $[v]$ come $[\lambda v]$ con $\lambda \neq 0$ otteniamo come immagine lo stesso punto $f_*([\lambda v]) = [\lambda f(v)] = [f(v)]$ grazie alla linearità della funzione f . $\square$

Una funzione $f_*: \mathbb{P}(V) \rightarrow \mathbb{P}(W)$ costruita in questo modo è chiamata una *proiettività*. La funzione f_* è biunivoca: la sua inversa è la proiettività $(f^{-1})_*$ indotta dall'inversa f^{-1} di f .

Siamo particolarmente interessati alle proiettività $\mathbb{P}^n(\mathbb{K}) \rightarrow \mathbb{P}^n(\mathbb{K})$. Queste sono per definizione indotte da mappe lineari invertibili $L_A: \mathbb{K}^{n+1} \rightarrow \mathbb{K}^{n+1}$, determinate a loro volta da matrici quadrate invertibili $A \in M(n+1, \mathbb{K})$. Concretamente, la proiettività $(L_A)_*: \mathbb{P}^n(\mathbb{K}) \rightarrow \mathbb{P}^n(\mathbb{K})$ manda il punto di coordinate $[x_1, \dots, x_{n+1}]$ nel punto

$$[a_{1,1}x_1 + \dots + a_{1,n+1}x_{n+1}, \dots, a_{n+1,1}x_1 + \dots + a_{n+1,n+1}x_{n+1}].$$

Esempio 12.3.2. La matrice

$$A = \begin{pmatrix} 1 & 2 & 0 \\ 0 & 1 & 0 \\ 0 & 0 & -2 \end{pmatrix}$$

descrive la proiettività

$$[x_1, x_2, x_3] \mapsto [x_1 + 2x_2, x_2, -2x_3]$$

del piano proiettivo $\mathbb{P}^2$.

Le proiettività di $\mathbb{P}^n(\mathbb{K})$ formano un gruppo con la composizione. Notiamo che due matrici A e λA con $\lambda \neq 0$ descrivono in realtà la stessa proiettività. Notiamo anche che una proiettività manda un sottospazio proiettivo in un sottospazio proiettivo della stessa dimensione.

Osservazione 12.3.3. In geometria affine, abbiamo prima definito una affinità come una qualsiasi mappa $f(x) = Ax + b$, e poi abbiamo chiamato *isomorfismi affini* le affinità invertibili, quelle con $\det A \neq 0$. Qui invece inseriamo la condizione di invertibilità direttamente nella nozione di proiettività.

12.3.2. Riferimenti proiettivi. In geometria affine, un isomorfismo affine di $\mathbb{R}^n$ è determinato dai suoi valori su un qualsiasi insieme di $n + 1$ punti affinemente indipendenti. Ci potremmo aspettare un analogo risultato per gli isomorfismi proiettivi, ma questo non è possibile, come mostra l'esempio seguente.

Esempio 12.3.4. Sia $a \neq 0$. Consideriamo la proiettività di $\mathbb{P}^n$

$$f_a([x_1, \dots, x_{n+1}]) = [x_1, \dots, x_n, ax_{n+1}].$$

Notiamo che f_a manda ciascuno dei punti $[1, 0, \dots, 0], \dots, [0, \dots, 0, 1]$ in se stesso, però manda $[1, \dots, 1]$ nel punto $[1, \dots, 1, a]$, che dipende da a . I punti $[1, 0, \dots, 0], \dots, [0, \dots, 0, 1]$ sono $n + 1$ punti proiettivamente indipendenti, ma le loro immagini non determinano l'intera funzione f_a .

Per determinare univocamente una proiettività $n + 1$ punti non sono quindi sufficienti: vedremo che è necessario prenderne $n + 2$. Sia V uno spazio vettoriale con $\dim V = n + 1$ e quindi $\dim \mathbb{P}(V) = n$.

Definizione 12.3.5. Un *riferimento proiettivo* per $\mathbb{P}(V)$ è un insieme di $n + 2$ punti $P_0, \dots, P_{n+1}$, tali che qualsiasi sottoinsieme di $n + 1$ punti sia formato da punti proiettivamente indipendenti.

Proposizione 12.3.6. *Un insieme di punti $P_0, \dots, P_{n+1} \in \mathbb{P}(V)$ è un riferimento proiettivo se e solo se esiste una base $v_0, \dots, v_n$ di V rispetto alla quale i punti hanno coordinate*

$$P_0 = [1, 0, \dots, 0], \quad \dots \quad P_n = [0, \dots, 0, 1], \quad P_{n+1} = [1, \dots, 1].$$

Dimostrazione. Se esiste una tale base, è facile vedere che i punti formano un riferimento proiettivo. D'altro canto, supponiamo che i punti formino un riferimento proiettivo. Prendiamo dei rappresentanti $P_0 = [w_0], \dots, P_{n+1} = [w_{n+1}]$. Per ipotesi, togliendo un qualsiasi vettore da $w_0, \dots, w_{n+1}$ si ottiene una base per V . In particolare $w_0, \dots, w_n$ è una base e quindi

$$w_{n+1} = \lambda_0 w_0 + \dots + \lambda_n w_n.$$

Se $\lambda_i = 0$ ottengo una dipendenza lineare fra i vettori $w_0, \dots, w_{i-1}, w_{i+1}, \dots, w_{n+1}$, che è escluso. Quindi $\lambda_i \neq 0$ per ogni i . Prendiamo $v_i = \lambda_i w_i$ per ogni $i \leq n$ e otteniamo le coordinate omogenee cercate per i P_i . $\square$

Proposizione 12.3.7. *Siano $P_0, \dots, P_{n+1}$ e $Q_0, \dots, Q_{n+1}$ due riferimenti proiettivi per $\mathbb{P}(V)$. Esiste un'unica proiettività $f: \mathbb{P}(V) \rightarrow \mathbb{P}(V)$ tale che $f(P_i) = Q_i$ per ogni i .*

Dimostrazione. Siano $v_0, \dots, v_n$ e $w_0, \dots, w_n$ basi di V rispetto a cui i punti P_i e Q_j hanno le coordinate descritte nella proposizione precedente. L'isomorfismo $g: V \rightarrow V$ tale che $g(v_i) = w_j$ per ogni i induce la proiettività $f = g_*: \mathbb{P}(V) \rightarrow \mathbb{P}(V)$ cercata. Si vede facilmente che g è

univocamente determinato a meno di uno scalare moltiplicativo, quindi f è unica. $\square$

In un certo senso abbiamo più libertà nel definire una proiettività di quanta ne avevamo nel definire un'affinità. Informalmente, il motivo è che in una proiettività possiamo spostare un punto dall'infinito al finito e viceversa, cosa chiaramente preclusa agli isomorfismi affini. Chiariamo un po' questo aspetto nel paragrafo seguente.

12.3.3. Da isomorfismo affine a proiettività. Mostriamo come estendere qualsiasi isomorfismo affine di $\mathbb{R}^n$ ad una proiettività di $\mathbb{P}^n$.

Consideriamo un isomorfismo affine di $\mathbb{R}^n$, del tipo

$$f(x) = Ax + b$$

con $A \in M(n, \mathbb{R})$ invertibile e $b \in \mathbb{R}^n$. L'isomorfismo affine f induce una proiettività $\bar{f}: \mathbb{P}^n \rightarrow \mathbb{P}^n$, definita dalla matrice $(n+1) \times (n+1)$ invertibile

$$\bar{A} = \begin{pmatrix} A & b \\ 0 & 1 \end{pmatrix}.$$

Notiamo che per ogni $x \in \mathbb{R}^n$ otteniamo

$$\bar{A} \begin{pmatrix} x \\ 1 \end{pmatrix} = \begin{pmatrix} A & b \\ 0 & 1 \end{pmatrix} \begin{pmatrix} x \\ 1 \end{pmatrix} = \begin{pmatrix} Ax + b \\ 1 \end{pmatrix}.$$

$$\bar{A} \begin{pmatrix} x \\ 0 \end{pmatrix} = \begin{pmatrix} A & b \\ 0 & 1 \end{pmatrix} \begin{pmatrix} x \\ 0 \end{pmatrix} = \begin{pmatrix} Ax \\ 0 \end{pmatrix}.$$

Ricordiamoci della partizione $\mathbb{P}^n = \mathbb{R}^n \cup H$. Ne deduciamo che $\bar{f}$ si comporta esattamente come f sui punti di $\mathbb{R}^n$, ed in più sposta i punti all'infinito in H con una proiettività determinata dalla matrice A .

Tutti gli isomorfismi affini f di $\mathbb{R}^n$ si estendono a proiettività $\bar{f}$ di $\mathbb{P}^n$ in questo modo.

Esempio 12.3.8. La traslazione $f(x) = x + b$ si estende alla proiettività $\bar{f}$ determinata dalla matrice $\begin{pmatrix} I_n & b \\ 0 & 1 \end{pmatrix}$, cioè

$$\bar{f}([x_1, \dots, x_{n+1}]) = [x_1 + b_1 x_{n+1}, \dots, x_n + b_n x_{n+1}, x_{n+1}].$$

Poiché $A = I_n$, i punti all'infinito sono tutti punti fissi:

$$\bar{f}([x_1, \dots, x_n, 0]) = [x_1, \dots, x_n, 0].$$

Esempio 12.3.9. Consideriamo la rotazione $f\left(\begin{smallmatrix} x_1 \\ x_2 \end{smallmatrix}\right) = \text{Rot}_\theta\left(\begin{smallmatrix} x_1 \\ x_2 \end{smallmatrix}\right)$ di $\mathbb{R}^2$ di un angolo θ intorno all'origine. Questo isomorfismo affine si estende alla proiettività $\bar{f}$ determinata dalla matrice

$$\bar{A} = \begin{pmatrix} \text{Rot}_\theta & 0 \\ 0 & 1 \end{pmatrix}.$$

La proiettività è la seguente:

$$\bar{f}([x_1, x_2, x_3]) = [x_1 \cos \theta - x_2 \sin \theta, x_1 \sin \theta + x_2 \cos \theta, x_3].$$

Notiamo che i punti all'infinito vengono ruotati anche loro.

Esempio 12.3.10. Consideriamo una rototraslazione in $\mathbb{R}^3$ con asse x_1 e passo t :

$$f \begin{pmatrix} x_1 \\ x_2 \\ x_3 \end{pmatrix} = \begin{pmatrix} 1 & 0 & 0 \\ 0 & \cos \theta & -\sin \theta \\ 0 & \sin \theta & \cos \theta \end{pmatrix} \begin{pmatrix} x_1 \\ x_2 \\ x_3 \end{pmatrix} + \begin{pmatrix} t \\ 0 \\ 0 \end{pmatrix}$$

La sua estensione $\bar{f}$ è

$$\bar{f}([x_1, x_2, x_3, x_4]) = [x_1 + tx_4, x_2 \cos \theta - x_3 \sin \theta, x_2 \sin \theta + x_3 \cos \theta, x_4].$$

Anche qui i punti all'infinito vengono in qualche senso ruotati intorno al punto $[1, 0, 0, 0]$, che è il punto all'infinito dell'asse x_1 e viene fissato da $\bar{f}$.

12.3.4. Punti fissi. Abbiamo notato negli esempi precedenti che la traslazione $f(x) = x + b$ con $b \neq 0$, pur non avendo punti fissi in $\mathbb{R}^n$, fissa in realtà tutti i punti all'infinito. Ci chiediamo quindi come studiare in generale i punti fissi di una proiettività. Scopriamo subito che questo problema si riconduce ad uno studio di autovettori.

Proposizione 12.3.11. *Sia $f: V \rightarrow V$ un isomorfismo che induce una proiettività $f_*: \mathbb{P}(V) \rightarrow \mathbb{P}(V)$. Un punto $P = [v] \in \mathbb{P}(V)$ è un punto fisso per $f_* \iff v$ è autovettore per f .*

Dimostrazione. Il punto P è fisso se e solo se $[f(v)] = [v]$, cioè se e solo se $f(v) = \lambda v$ per qualche $\lambda \neq 0$. $\square$

Esempio 12.3.12. La matrice A descritta nell'Esempio 12.3.2 ha due autovalori 1 e -2 , con autospazi

$$V_1 = \text{Span} \begin{pmatrix} 1 \\ 0 \\ 0 \end{pmatrix}, \quad V_{-2} = \text{Span} \begin{pmatrix} 0 \\ 0 \\ 1 \end{pmatrix}.$$

Quindi i punti fissi in $\mathbb{P}^2$ della proiettività descritta da A sono due: $[1, 0, 0]$ e $[0, 0, 1]$. D'altra parte, se

$$B = \begin{pmatrix} 1 & 0 & 0 \\ 0 & 1 & 0 \\ 0 & 0 & -2 \end{pmatrix}$$

allora gli autospazi di B sono il piano e la retta vettoriali:

$$V_1 = \text{Span}(e_1, e_2), \quad V_{-2} = \text{Span}(e_3).$$

I punti fissi in $\mathbb{P}^3$ della proiettività descritta da B sono corrispondentemente tutti quelli della retta $r = \{x_3 = 0\}$ ed il punto $[0, 0, 1]$.

Dalla teoria degli autovettori che abbiamo sviluppato nei capitoli precedenti possiamo dedurre facilmente un fatto non banale.

Proposizione 12.3.13. *Ogni proiettività di $\mathbb{P}^2$ ha almeno un punto fisso.*

Dimostrazione. Ogni matrice 3×3 ha almeno un autovettore. $\square$

La proposizione è abbastanza sorprendente perché non è valida in geometria affine: le traslazioni sono affinità di $\mathbb{R}^2$ senza punti fissi (ma, come abbiamo visto, queste fissano tutti i punti all'infinito).

Esercizi

Esercizio 12.1. Scrivi l'equazione omogenea che descrive la retta in $\mathbb{P}^2$ che passa per i punti

$$[1, 1, 0], \quad [0, 0, 1].$$

Esercizio 12.2. Scrivi l'equazione omogenea che descrive il piano proiettivo $\pi \subset \mathbb{P}^3$ che passa per i punti

$$[1, 2, 1, 0], \quad [1, 0, 1, 0], \quad [0, 0, 1, 1].$$

Scrivi in forma parametrica la retta $r = \pi \cap \pi'$, intersezione dei piani proiettivi π e $\pi' = \{x_1 + 2x_3 - x_4 = 0\}$. Determina l'equazione omogenea che descrive il piano che passa per r e il punto $[1, 1, 0, 1]$.

Esercizio 12.3. Determina per quali valori $k \in \mathbb{R}$ le rette di equazione

$$kx_1 + x_2 = 0, \quad x_1 + x_3 = 0, \quad 3x_1 + x_2 + 2x_3$$

si intersecano tutte e tre in un punto.

Esercizio 12.4. Determina la proiettività di $\mathbb{P}^1$ che fissa $[1, 0]$ e $[0, 1]$ e manda $[2, 1]$ in $[1, 3]$.

Esercizio 12.5. Determina la proiettività di $\mathbb{P}^2$ che manda i punti $[1, 0, 1]$, $[2, 1, -1]$, $[0, 1, 1]$, $[0, 1, 0]$ rispettivamente in $[0, 2, 2]$, $[4, 7, 1]$, $[0, 3, 1]$, $[1, 3, 0]$.

Esercizio 12.6. Scrivi una proiettività di $\mathbb{P}^3$ senza punti fissi.

Esercizio 12.7. Per quali numeri naturali n esiste una proiettività di $\mathbb{P}^n$ con esattamente n punti fissi? (Se c'è fornisci un esempio, se non c'è dimostrarlo).

Esercizio 12.8. Siano $r, r' \subset \mathbb{P}^3$ due rette proiettive disgiunte e $P \in \mathbb{P}^3$ un punto disgiunto da entrambe. Mostra che esiste un'unica retta $s \subset \mathbb{P}^3$ che interseca r , r' e P . (Questa è la versione proiettiva dell'Esercizio 9.7. Nota che qui non abbiamo bisogno di imporre delle condizioni di non-parallelismo: gli enunciati proiettivi sono spesso più semplici.)

Esercizio 12.9. Con le notazioni dell'esercizio precedente, siano

$$r = \begin{cases} x_1 + x_2 - x_3 - x_4 = 0 \\ x_1 + x_3 = 0 \end{cases}, \quad r' = \begin{cases} x_1 + 2x_2 - x_4 = 0 \\ x_2 - x_3 = 0 \end{cases}, \quad P = [1, 1, 1, 0].$$

Mostra che r e r' sono disgiunte. Determina s .

Complementi

12.1. I Teoremi di Desargues e di Pappo. In geometria proiettiva molti teoremi noti di geometria euclidea come quello di Pitagora non sono validi semplicemente perché non hanno senso: nel piano proiettivo $\mathbb{P}^2$ non esistono i concetti di distanza fra punti né di angolo fra segmenti. Come abbiamo visto, non esiste neppure la nozione di parallelismo.

In geometria proiettiva esistono innanzitutto i sottospazi e le relazioni di incidenza fra questi. Abbiamo già visto in questo capitolo dei teoremi

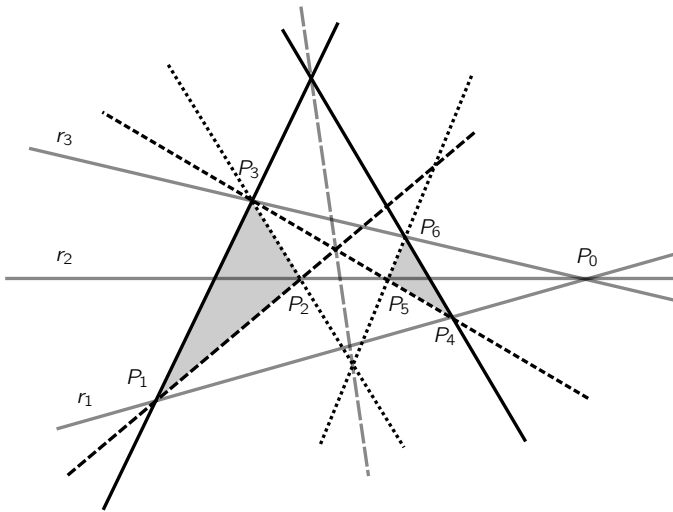


Figura 12.5. Il Teorema di Desargues.

generali, che si distinguono per la loro semplicità e generalità: per due punti passa una sola retta; due piani distinti in $\mathbb{P}^3$ si intersecano in una retta; eccetera. In questa sezione presentiamo due teoremi più elaborati, di Desargues e di Pappo.

Sia $\mathbb{P}(V)$ un piano proiettivo, cioè uno spazio proiettivo arbitrario di dimensione 2. Dati due punti distinti P, Q , indichiamo con $L(P, Q)$ la retta passante per P e Q .

Teorema 12.3.14 (Teorema di Desargues). *Siano r_1, r_2, r_3 tre rette incidenti in un punto P_0 in $\mathbb{P}(V)$. Siano*

$$P_1, P_4 \in r_1, \quad P_2, P_5 \in r_2, \quad P_3, P_6 \in r_3$$

dei punti arbitrari tutti distinti e diversi da P_0 . I tre punti

$$L(P_1, P_3) \cap L(P_4, P_6), \quad L(P_2, P_3) \cap L(P_5, P_6), \quad L(P_1, P_2) \cap L(P_4, P_5)$$

sono allineati.

Si veda la Figura 12.5.

Dimostrazione. Siano $v_0, \dots, v_6 \in V$ con $P_i = [v_i]$. Le condizioni di allineamento fra punti ci dicono che

$$v_0 = \lambda_1 v_1 + \lambda_4 v_4 = \lambda_2 v_2 + \lambda_5 v_5 = \lambda_3 v_3 + \lambda_6 v_6$$

per opportuni scalari $\lambda_1, \dots, \lambda_6$ non nulli (perché $P_0 \neq P_i \forall i \geq 1$). Il vettore

$$\lambda_1 v_1 - \lambda_3 v_3 = \lambda_6 v_6 - \lambda_4 v_4$$

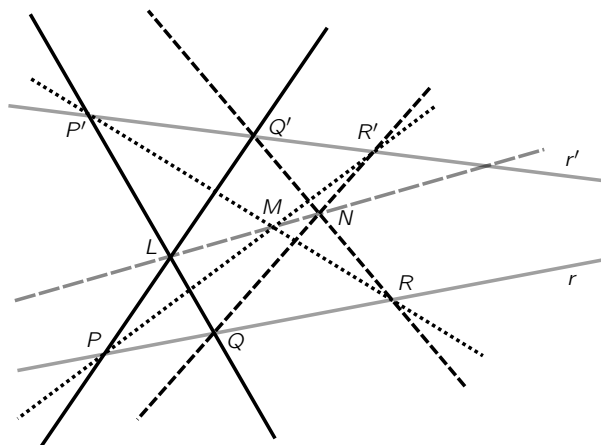


Figura 12.6. Il Teorema di Pappo.

rappresenta il punto $L(P_1, P_3) \cap L(P_4, P_6)$. Analogamente i vettori

$$\lambda_3 v_3 - \lambda_2 v_2 = \lambda_5 v_5 - \lambda_6 v_6$$

$$\lambda_2 v_2 - \lambda_1 v_1 = \lambda_4 v_4 - \lambda_5 v_5$$

rappresentano i punti $L(P_2, P_3) \cap L(P_5, P_6)$ e $L(P_1, P_2) \cap L(P_4, P_5)$. I tre vettori appena descritti sono dipendenti perché la loro somma è nulla. Quindi i tre punti corrispondenti in $\mathbb{P}(V)$ sono allineati. $\square$

Sia sempre $\mathbb{P}(V)$ un piano proiettivo.

Teorema 12.3.15 (Teorema di Pappo). *Siano r, r' due rette distinte in $\mathbb{P}(V)$. Siano*

$$P, Q, R \in r, \quad P', Q', R' \in r'$$

punti distinti e diversi da $O = r \cap r'$. I tre punti

$$L(P, Q') \cap L(P', Q), \quad L(Q, R') \cap L(Q', R), \quad L(P, R') \cap L(P', R)$$

sono allineati.

Si veda la Figura 12.6.

Dimostrazione. Dobbiamo mostrare che i tre punti L, M, N nella Figura 12.6 sono allineati. I punti P, P', Q, Q' formano un riferimento proiettivo. Fissiamo una base di V per cui questi abbiano coordinate omogenee

$$P = [1, 0, 0], \quad P' = [0, 1, 0], \quad Q = [0, 0, 1], \quad Q' = [1, 1, 1].$$

Quindi

$$R = [a, 0, b], \quad R' = [c, c + d, c]$$

per qualche $a, b, c, d \neq 0$. Si vede facilmente che

$$L = [0, 1, 1], \quad M = [ac, b(c+d), bc], \quad N = [ac, a(c+d), a(c+d) - bd]$$

sono allineati. $\square$

12.II. Dualità. Oltre ad eliminare i parallelismi, la geometria proiettiva ci permette di interpretare le rette di $\mathbb{P}^2$ (e più in generale gli iperpiani di $\mathbb{P}^n$) come punti di un altro spazio vettoriale *duale*. Questa proprietà è peculiare della geometria proiettiva e non è presente né nella geometria vettoriale né in quella affine. Usando questo strumento è possibile “dualizzare” i teoremi trovando nuovi teoremi. Descriviamo questo fenomeno.

Sia V uno spazio vettoriale su $\mathbb{K}$ di dimensione $n+1$. Nella Sezione 4.I abbiamo definito lo spazio duale V^* come lo spazio di tutte le applicazioni lineari $V \rightarrow \mathbb{K}$. Ricordiamo che $\dim V^* = n+1$.

Definizione 12.3.16. Lo spazio proiettivo *duale* a $\mathbb{P}(V)$ è $\mathbb{P}(V^*)$.

Per definizione, un punto di $\mathbb{P}(V^*)$ è un elemento $[f]$ dove $f: V \rightarrow \mathbb{K}$ è un'applicazione lineare non nulla. Vale $[f] = [g]$ se e solo se $g = \lambda f$ con $\lambda \neq 0$. Notiamo un fatto geometricamente importante: se $f: V \rightarrow \mathbb{K}$ è una funzione lineare non nulla, il suo nucleo

$$U = \ker f$$

è un iperpiano di V che non cambia se sostituiamo f con λf , quindi è ben definito a partire dal punto $[f]$. L'iperpiano $U = \ker f$ definisce a sua volta un iperpiano in $\mathbb{P}(V)$. Abbiamo appena scoperto che un punto in $\mathbb{P}(V^*)$ definisce un iperpiano in $\mathbb{P}(V)$.

Proposizione 12.3.17. Abbiamo costruito una corrispondenza biunivoca fra i punti di $\mathbb{P}(V^*)$ e gli iperpiani di $\mathbb{P}(V)$.

Dimostrazione. Per verificare che questa è una corrispondenza biunivoca dobbiamo mostrare due fatti di algebra lineare che lasciamo come esercizio.

Suriettività: per ogni iperpiano $U \subset V$ esiste una funzione lineare $f: V \rightarrow \mathbb{K}$ con $U = \ker f$.

Iniettività: se due funzioni lineari $f, g: V \rightarrow \mathbb{K}$ hanno $\ker f = \ker g$, allora $f = \lambda g$ per qualche $\lambda \neq 0$. $\square$

Possiamo interpretare i punti di $\mathbb{P}(V^*)$ come gli iperpiani in $\mathbb{P}(V)$. Di converso, abbiamo scoperto che gli iperpiani in $\mathbb{P}(V)$ formano in modo naturale uno spazio proiettivo $\mathbb{P}(V^*)$ e quindi possiamo usare il linguaggio della geometria proiettiva per studiarli. Questa è una proprietà specifica della geometria proiettiva: gli iperpiani vettoriali di V e affini di $\mathbb{R}^n$ non formano né uno spazio vettoriale né affine.

Ad esempio, possiamo dire che k iperpiani in $\mathbb{P}(V)$ sono *proiettivamente indipendenti* se lo sono visti come punti in $\mathbb{P}(V^*)$.

Lo spazio proiettivo duale di $\mathbb{P}^n(\mathbb{K})$ è indicato con $\mathbb{P}^n(\mathbb{K})^*$. Siamo interessati soprattutto al caso reale e indichiamo con $(\mathbb{P}^n)^*$ il duale di $\mathbb{P}^n$.

Osservazione 12.3.18. La dualità tra $\mathbb{P}^n(\mathbb{K})$ e $\mathbb{P}^n(\mathbb{K})^*$ può essere descritta più concretamente. Gli spazi $\mathbb{K}^{n+1}$ e $(\mathbb{K}^{n+1})^*$ sono rispettivamente i vettori colonna ed i vettori riga, ed un vettore riga $(a_1, \dots, a_{n+1}) \in (\mathbb{K}^{n+1})^*$ non nullo definisce l'iperpiano

$$(27) \quad a_1 x_1 + \dots + a_{n+1} x_{n+1} = 0$$

in $\mathbb{K}^{n+1}$. Nel proiettivo, il vettore $[a_1, \dots, a_{n+1}] \in \mathbb{P}^n(\mathbb{K})^*$ determina l'iperpiano in $\mathbb{P}^n(\mathbb{K})$ definito dalla stessa equazione (27). Ad esempio, il punto $[1, 0, -1] \in (\mathbb{P}^2)^*$ descrive l'iperpiano $x_1 - x_3 = 0$ in $\mathbb{P}^2$.

Sia $S \subset \mathbb{P}(V)$ un sottospazio proiettivo di dimensione k . Indichiamo con $\hat{S} \subset \mathbb{P}(V^*)$ l'insieme degli iperpiani che contengono S .

Proposizione 12.3.19. *Il sottoinsieme $\hat{S} \subset \mathbb{P}(V^*)$ è un sottospazio proiettivo di dimensione $n - k - 1$. Trasformando S in $\hat{S}$ otteniamo una corrispondenza biunivoca fra i k -sottospazi di $\mathbb{P}(V)$ e gli $(n - k - 1)$ -sottospazi di $\mathbb{P}(V^*)$.*

Dimostrazione. Il sottospazio proiettivo S corrisponde ad un sottospazio vettoriale $U \subset V$ e per costruzione $\hat{S}$ corrisponde al suo annullatore $\text{Ann}(U) \subset V^*$. L'enunciato è una conseguenza del Corollario 4.4.28. $\square$

Se $S \subset \mathbb{P}(V)$ è un iperpiano, allora $\hat{S} \subset \mathbb{P}(V^*)$ è semplicemente il punto che indica quell'iperpiano. Se S ha codimensione 2, allora $\hat{S}$ è una retta di $\mathbb{P}(V^*)$, una "retta di iperpiani" detta *fascio*. Ritroviamo qui la nozione di fascio già esplorata in geometria affine nella Sezione 9.2.8.

Esempio 12.3.20. Consideriamo la retta $r = \{[t + u, t, u, t - u]\}$ in $\mathbb{P}^3$. Il sottospazio duale $\hat{r} \subset (\mathbb{P}^3)^*$ è il fascio di piani che contengono la retta r . Determiniamo concretamente i punti di $\hat{r}$.

È sufficiente trovare due piani distinti che contengono r , ad esempio i piani $\{x_1 - x_2 - x_3 = 0\}$ e $\{x_2 - x_3 - x_4 = 0\}$, che corrispondono ai punti $[1, -1, -1, 0]$ e $[0, 1, -1, -1]$ in $(\mathbb{P}^3)^*$. La retta $\hat{r}$ è quindi $\{[s, r - s, -r - s, -r]\}$.

Abbiamo definito una dualità che trasforma k -sottospazi $S \subset \mathbb{P}(V)$ in $(n - k - 1)$ -sottospazi $\hat{S} \subset \mathbb{P}(V^*)$. Questa trasformazione scambia le inclusioni: se $S \subset S'$, allora $\hat{S} \supset \hat{S}'$. Notiamo un altro fatto interessante. Siano $S, S' \subset \mathbb{P}(V)$ due sottospazi proiettivi.

Proposizione 12.3.21. *La dualità trasforma $S \cap S'$ in $\hat{S} + \hat{S}'$.*

Dimostrazione. Sia $X \subset \mathbb{P}(V^*)$ il sottospazio duale a $S \cap S'$. Gli iperpiani che contengono S oppure S' contengono anche $S \cap S'$, quindi X contiene $\hat{S}$ e $\hat{S}'$, e quindi contiene $\hat{S} + \hat{S}'$. Con la formula di Grassmann otteniamo (esercizio) che $\dim X = \dim(\hat{S} + \hat{S}')$ e quindi $X = \hat{S} + \hat{S}'$. $\square$

Analogamente si dimostra che la dualità trasforma $S + S'$ in $\hat{S} \cap \hat{S}'$. Questi fatti hanno una conseguenza profonda sulla geometria degli spazi

proiettivi. Supponiamo che P sia una proposizione su $\mathbb{P}^n$ che riguarda i suoi sottospazi e le loro incidenze. Ad esempio:

Due punti distinti in $\mathbb{P}^2$ sono contenuti entrambi in un'unica retta.

La *proposizione duale* P^* è ottenuta scambiando k -sottospazi con $(n - k - 1)$ -sottospazi e i relativi contenimenti. Ad esempio, in questo caso P^* è

Due rette distinte in $\mathbb{P}^2$ contengono entrambe un solo punto.

Teorema 12.3.22 (Principio di dualità). P è vera $\iff P^*$ è vera.

Dimostrazione. La dualità trasforma k -sottospazi di $\mathbb{P}(V)$ in $(n - k - 1)$ -sottospazi di $\mathbb{P}(V^*)$ scambiando le inclusioni. Visto che gli spazi proiettivi $\mathbb{P}^n$ e $(\mathbb{P}^n)^*$ hanno la stessa dimensione, e sono quindi proiettivamente isomorfi, tutti i teoremi generali validi per $\mathbb{P}^n$ lo sono anche per $(\mathbb{P}^n)^*$. Quindi la dualità deve trasformare proposizioni vere in proposizioni vere. $\square$

Notiamo che $P^{**} = P$, cioè la proposizione duale a P^* è la P di partenza. Per il principio di dualità, ogni proposizione vera ha una proposizione duale che è anch'essa vera. Questo principio può essere usato per dedurre nuovi teoremi da teoremi già noti.

Esempio 12.3.23. Mostriamo alcuni esempi di proposizioni vere e delle loro duali (anch'esse vere, per il principio di dualità). Se P è la proposizione

*k punti indipendenti in $\mathbb{P}^n$ sono contenuti
in un unico $(k - 1)$ -sottospazio*

allora la sua duale P^* è

*k iperpiani indipendenti in $\mathbb{P}^n$ si intersecano
in un unico $(n - k)$ -sottospazio.*

Se P è la proposizione

*Un punto Q ed una retta r in $\mathbb{P}^3$ con $Q \notin r$
sono contenuti in un unico piano*

allora la sua duale P^* è

*Un piano π ed una retta r in $\mathbb{P}^3$ con $\pi \not\supset r$
si intersecano in un unico punto.*

Una proposizione P è *autoduale* se coincide con la sua duale P^* . Un esempio di proposizione autoduale è la seguente:

Due rette distinte in $\mathbb{P}^3$ si intersecano $\iff$ sono complanari.

Esercizio 12.3.24. Scrivi la proposizione duale a quella dell'Esercizio 12.8 e dimostrarla direttamente; quindi deduci l'esercizio come corollario usando il principio di dualità.

Esercizio 12.3.25. Scrivi le proposizioni duali ai Teoremi di Desargues e di Pappo. Nota che si ottengono in entrambi i casi le stesse configurazioni di punti e rette!

Quadriche

L'ultimo capitolo del libro è dedicato allo studio delle *quadriche*. Una quadrica è il luogo di zeri di una equazione di secondo grado in $\mathbb{R}^n$. Siamo ovviamente interessati soprattutto alle quadriche in $\mathbb{R}^2$ e in $\mathbb{R}^3$.

Le quadriche del piano $\mathbb{R}^2$ sono curve familiari, chiamate *coniche*, costruite fin dall'antichità intersecando un cono con un piano e già studiate alle superiori: fra queste troviamo ellissi, parabole e iperboli.

Le quadriche dello spazio $\mathbb{R}^3$ sono superfici di un tipo relativamente semplice, ampiamente usate dall'essere umano nella costruzione di manufatti. Fra queste troviamo ellissoidi, paraboloidi e iperboloidi.

Per studiare coniche e quadriche useremo abbondantemente gli strumenti affinati nei capitoli precedenti: le quadriche sono strettamente collegate alle forme quadratiche e quindi ai prodotti scalari, le classificheremo con la segnatura ed il teorema spettrale, le completeremo con la geometria proiettiva aggiungendo i loro punti all'infinito.

13.1. Introduzione

Nei capitoli precedenti abbiamo studiato molte forme geometriche *lineari*, cioè definite usando equazioni di primo grado; in questo capitolo alziamo il tiro ed esaminiamo alcuni oggetti *quadratici*, cioè definiti usando equazioni di secondo grado. Questi oggetti comprendono le *coniche* nel piano, quali le ellissi, iperboli e parabole già viste nelle scuole superiori.

13.1.1. Quadriche affini. Introduciamo subito gli oggetti principali che studieremo in questo capitolo.

Definizione 13.1.1. Una *quadrica* in $\mathbb{R}^n$ è il luogo di zeri

$$Q = \{x \in \mathbb{R}^n \mid p(x_1, \dots, x_n) = 0\}$$

di un polinomio $p(x_1, \dots, x_n)$ di secondo grado nelle variabili $x_1, \dots, x_n$.

Sono esempi di quadriche le circonferenze in $\mathbb{R}^2$ e le sfere in $\mathbb{R}^3$ definite nel Capitolo 10. Una *conica* è una quadrica in $\mathbb{R}^2$.

Usiamo a volte il simbolo $p(x)$ al posto di $p(x_1, \dots, x_n)$ per semplicità. Scriveremo un generico polinomio $p(x)$ di secondo grado nella forma

seguente:

$$(28) \quad \sum_{i,j=1}^n a_{ij}x_i x_j + 2 \sum_{i=1}^n b_i x_i + c$$

con l'ipotesi aggiuntiva che $a_{ij} = a_{ji}$ per ogni i, j . È effettivamente sempre possibile scrivere un polinomio di secondo grado in questo modo, raccogliendo prima i monomi di secondo grado, poi quelli di primo, e infine il termine noto c . Il primo addendo è una forma quadratica, codificata dalla matrice simmetrica $A = (a_{ij})$, si veda la Proposizione 7.1.12. Possiamo anche scrivere $p(x)$ così:

$$(29) \quad {}^t x A x + 2 {}^t b x + c$$

dove $b \in \mathbb{R}^n$ è il vettore di coordinate $b_1, \dots, b_n$. Si può anche usare il linguaggio dei prodotti scalari e delle forme quadratiche e scrivere $p(x)$ in questa maniera:

$$q_A(x) + 2\langle b, x \rangle + c$$

dove $q_A(x) = {}^t x A x$ è la forma quadratica associata alla matrice simmetrica A e $\langle b, x \rangle = {}^t b x$ è il prodotto scalare euclideo dei vettori b e x . Infine, possiamo anche definire la matrice $\bar{A} \in M(n+1, \mathbb{R})$ ed il vettore $\bar{x} \in \mathbb{R}^{n+1}$

$$\bar{A} = \begin{pmatrix} A & b \\ {}^t b & c \end{pmatrix}, \quad \bar{x} = \begin{pmatrix} x \\ 1 \end{pmatrix}$$

e scrivere (29) in forma compatta così:

$$(30) \quad {}^t \bar{x} \bar{A} \bar{x}.$$

Svolgendo il prodotto si vede che (30) è effettivamente equivalente a (29):

$${}^t \bar{x} \bar{A} \bar{x} = ({}^t x \quad 1) \begin{pmatrix} A & b \\ {}^t b & c \end{pmatrix} \begin{pmatrix} x \\ 1 \end{pmatrix} = {}^t x A x + {}^t x b + {}^t b x + c = {}^t x A x + 2 {}^t b x + c.$$

Il fattore 2 davanti ai monomi di primo grado in (28) è stato messo proprio per ottenere la forma compatta (30).

Esempio 13.1.2. Siamo interessati soprattutto agli spazi $\mathbb{R}^2$ e $\mathbb{R}^3$. In questi, useremo come di consueto le variabili x, y, z invece di x_1, x_2, x_3 . Il polinomio di secondo grado in due variabili

$$x^2 - 2xy + 4y^2 - x + 6y + 8$$

si può scrivere come

$$(x, y) \begin{pmatrix} 1 & -1 \\ -1 & 4 \end{pmatrix} \begin{pmatrix} x \\ y \end{pmatrix} + 2 \left(-\frac{1}{2}, 3\right) \begin{pmatrix} x \\ y \end{pmatrix} + 8$$

oppure in forma compatta come

$$(x, y, 1) \begin{pmatrix} 1 & -1 & -\frac{1}{2} \\ -1 & 4 & 3 \\ -\frac{1}{2} & 3 & 8 \end{pmatrix} \begin{pmatrix} x \\ y \\ 1 \end{pmatrix}.$$

Una quadrica Q è descritta dall'equazione

$${}^t\bar{x}\bar{A}\bar{x} = 0.$$

Osservazione 13.1.3. Se sostituiamo $\bar{A}$ con $\lambda\bar{A}$ per qualche scalare $\lambda \neq 0$, otteniamo l'equazione ${}^t\bar{x}\lambda\bar{A}\bar{x} = 0$, che definisce la stessa quadrica Q . Quindi è sempre possibile moltiplicare $\bar{A}$ arbitrariamente per uno scalare $\lambda \neq 0$ senza alterare Q .

Introduciamo una definizione importante.

Definizione 13.1.4. La quadrica ${}^t\bar{x}\bar{A}\bar{x} = 0$ è *degenere* se $\det \bar{A} = 0$.

Vedremo successivamente che le quadriche non degeneri sono generalmente più interessanti di quelle degeneri. Studieremo approfonditamente le proprietà geometriche delle coniche nella Sezione 13.2 e delle quadriche in $\mathbb{R}^3$ nelle sezioni successive. In quelle pagine analizzeremo alcuni esempi concreti di coniche e quadriche; per adesso ci limitiamo a dimostrare alcuni teoremi algebrici generali che porteranno successivamente alla loro completa classificazione.

13.1.2. Cambiamento di coordinate. Siamo interessati a vedere come muta il polinomio che descrive una quadrica se cambiamo sistema di riferimento. Questo processo sarà fondamentale in seguito quando classificheremo le quadriche in $\mathbb{R}^2$ e $\mathbb{R}^3$.

Consideriamo una quadrica $Q \subset \mathbb{R}^n$, descritta da un'equazione

$$(31) \quad {}^t\bar{x}\bar{A}\bar{x} = 0.$$

Stiamo scrivendo il polinomio in forma compatta come in (30). Supponiamo di spostare il sistema di riferimento, con un cambio di variabili

$$x = Mx' + P$$

dove M è una matrice ortogonale e $P \in \mathbb{R}^n$ un punto. Ricordiamo che con M cambiamo base ortonormale e con P trasliamo l'origine in P . Definiamo

$$\bar{M} = \begin{pmatrix} M & P \\ 0 & 1 \end{pmatrix}.$$

Possiamo scrivere il cambio di variabili in forma compatta come

$$\bar{x} = \bar{M}\bar{x}'$$

infatti svolgendo i conti si ottiene

$$\bar{M}\bar{x}' = \begin{pmatrix} M & P \\ 0 & 1 \end{pmatrix} \begin{pmatrix} x' \\ 1 \end{pmatrix} = \begin{pmatrix} Mx' + P \\ 1 \end{pmatrix} = \begin{pmatrix} x \\ 1 \end{pmatrix} = \bar{x}.$$

Proposizione 13.1.5. *Nelle coordinate x' la quadrica Q ha equazione*

$${}^t\bar{x}'\bar{A}'\bar{x}' = 0$$

con

$$\bar{A}' = {}^t\bar{M}\bar{A}\bar{M}.$$

Dimostrazione. Sostituiamo $\bar{x} = \bar{M}\bar{x}'$ e troviamo

$${}^t\bar{x}\bar{A}\bar{x} = {}^t(\bar{M}\bar{x}')\bar{A}\bar{M}\bar{x}' = {}^t\bar{x}'({}^t\bar{M}\bar{A}\bar{M})\bar{x}'.$$

La dimostrazione è conclusa. $\square$

Abbiamo scoperto che se cambiamo sistema di riferimento la matrice $\bar{A}$ cambia per congruenza tramite la matrice $\bar{M}$. Questo mostra in particolare che la condizione di essere degenere, definita dalla condizione $\det \bar{A} = 0$, non dipende dal sistema di riferimento scelto, perché $\det \bar{A} = 0 \iff \det \bar{A}' = 0$.

Vediamo più nel dettaglio come cambiano i coefficienti del polinomio che descrive Q . Scriviamo più esplicitamente l'equazione (31) come

$${}^tAx + 2 {}^tbx + c = 0.$$

Vediamo come cambiano A , b e c .

Corollario 13.1.6. *Nelle coordinate x' la quadrica Q ha equazione*

$${}^t_{x'}A'x' + 2 {}^t_{b'}x' + c' = 0$$

con

$$\begin{aligned} A' &= {}^tMAM, \\ b' &= {}^tM(AP + b), \\ c' &= {}^tPAP + 2 {}^tbP + c. \end{aligned}$$

Dimostrazione. Svolgiamo i conti:

$$\begin{aligned} \bar{A}' &= {}^t\bar{M}\bar{A}\bar{M} = \begin{pmatrix} {}^tM & 0 \\ {}^tP & 1 \end{pmatrix} \begin{pmatrix} A & b \\ b & c \end{pmatrix} \begin{pmatrix} M & P \\ 0 & 1 \end{pmatrix} = \begin{pmatrix} {}^tM & 0 \\ {}^tP & 1 \end{pmatrix} \begin{pmatrix} AM & AP + b \\ bM & bP + c \end{pmatrix} \\ &= \begin{pmatrix} {}^tMAM & {}^tMAP + {}^tMb \\ {}^tPAM + {}^tbM & {}^tPAP + {}^tPb + {}^tbP + c \end{pmatrix} \\ &= \begin{pmatrix} {}^tMAM & {}^tM(AP + b) \\ {}^t(AP + b)M & {}^tPAP + 2 {}^tbP + c \end{pmatrix} \end{aligned}$$

La dimostrazione è conclusa. $\square$

Notiamo che a fronte di un cambio di coordinate le matrici A e $\bar{A}$ mutano entrambe per congruenza, con la differenza importante che A cambia per congruenza rispetto ad una matrice M ortogonale mentre $\bar{A}$ cambia tramite una matrice $\bar{M}$ che è ortogonale se e solo se $P = 0$.

Dal Corollario 13.1.6 segue che se cambiamo coordinate con una traslazione $x = x' + P$ la matrice A resta immutata. Se invece cambiamo coordinate lasciando fissa l'origine, cioè con $x = Mx'$, allora il termine noto c resta immutato.

13.1.3. Diagonalizzazione di A . Mostriamo adesso che, a meno di cambiare sistema di riferimento, possiamo sempre supporre che la matrice simmetrica A sia in realtà diagonale. Questa è una applicazione diretta del teorema spettrale.

Proposizione 13.1.7. *Per ogni quadrica Q esiste un sistema di riferimento in cui Q abbia equazione*

$${}^t x A x + 2 {}^t b x + c = 0$$

con A matrice diagonale.

Dimostrazione. Partiamo da un sistema di riferimento qualsiasi. La quadrica Q ha equazione ${}^t x A x + 2 {}^t b x + c = 0$. La matrice A è simmetrica. Se A non è diagonale, per il teorema spettrale esiste una M ortogonale per cui $A' = {}^t M A M$ sia diagonale. Basta quindi prendere delle nuove coordinate x' con $x = M x'$ ed applicare il Corollario 13.1.6 per sostituire A con A' . $\square$

L'equazione con A diagonale può essere scritta esplicitamente così:

$$a_1 x_1^2 + \cdots + a_n x_n^2 + 2b_1 x_1 + \cdots + 2b_n x_n + c = 0.$$

Il coefficiente a_i è l'elemento a_{ii} della matrice diagonale A .

Esempio 13.1.8. Consideriamo la conica $C \subset \mathbb{R}^2$ di equazione

$$34x^2 + 41y^2 + 24xy + 20x - 10y + 4 = 0.$$

Otteniamo

$$\bar{A} = \begin{pmatrix} 34 & 12 & 10 \\ 12 & 41 & -5 \\ 10 & -5 & 4 \end{pmatrix}, \quad A = \begin{pmatrix} 34 & 12 \\ 12 & 41 \end{pmatrix}, \quad b = \begin{pmatrix} 10 \\ -5 \end{pmatrix}, \quad c = 4.$$

La matrice simmetrica A ha una base ortonormale v_1, v_2 di autovettori. Facendo i conti troviamo che gli autovalori sono $\lambda_1 = 25, \lambda_2 = 50$ e gli autovettori di norma unitaria sono

$$v_1 = \frac{1}{5} \begin{pmatrix} 4 \\ -3 \end{pmatrix}, \quad v_2 = \frac{1}{5} \begin{pmatrix} 3 \\ 4 \end{pmatrix}.$$

Scegliamo un sistema di riferimento in cui questi vettori diventano i generatori degli assi. Quindi poniamo $\begin{pmatrix} x \\ y \end{pmatrix} = M \begin{pmatrix} x' \\ y' \end{pmatrix}$ con

$$M = (v_1 \ v_2) = \frac{1}{5} \begin{pmatrix} 4 & 3 \\ -3 & 4 \end{pmatrix}.$$

Dal Corollario 13.1.6 deduciamo che la conica C nelle nuove coordinate è descritta dalle matrici

$$A' = {}^t M A M = \begin{pmatrix} 25 & 0 \\ 0 & 50 \end{pmatrix}, \quad b' = {}^t M b = \begin{pmatrix} 11 \\ 2 \end{pmatrix}, \quad c' = c = 4.$$

Nelle nuove coordinate la conica C ha equazione

$$25(x')^2 + 50(y')^2 + 22x' + 4y' + 4 = 0.$$

13.1.4. Centro di simmetria di una quadrica. Molte quadriche (ma non tutte) hanno una simmetria centrale. Analizziamo questo fenomeno.

Sia $Q \subset \mathbb{R}^n$ una quadrica. Sia $P \in \mathbb{R}^n$ un punto e $r_P: \mathbb{R}^n \rightarrow \mathbb{R}^n$ la riflessione rispetto a P . Diciamo che P è un *centro di simmetria* (più semplicemente, *centro*) per la quadrica Q se

$$r_P(Q) = Q.$$

Vediamo subito un criterio efficace per determinare un centro per Q .

Proposizione 13.1.9. *Sia $Q \subset \mathbb{R}^n$ una quadrica di equazione*

$${}^t x A x + 2 {}^t b x + c = 0.$$

Se $P \in \mathbb{R}^n$ è tale che $AP + b = 0$, allora P è un centro di simmetria per Q .

Dimostrazione. Per il Corollario 13.1.6, se spostiamo il centro delle coordinate in P con il cambio di variabili $x = x' + P$ otteniamo $b' = AP + b = 0$, quindi i termini di primo grado nel polinomio spariscono. Nelle nuove coordinate troviamo un polinomio $p(x') = {}^t x' A x' + c'$ che ha solo monomi di grado pari e quindi $p(-x') = p(x') \forall x' \in \mathbb{R}^n$. Ne segue che

$$x' \in Q \iff p(x') = 0 \iff p(-x') = 0 \iff -x' \in Q$$

e allora l'origine (che in queste coordinate è il punto P) è un centro per Q . $\square$

Esistono quadriche che hanno centri di simmetria e quadriche che non ne hanno. Se A è invertibile un centro c'è sempre e sappiamo come identificarlo.

Corollario 13.1.10. *Se A è invertibile, il punto $P = -A^{-1}b$ è un centro.*

Esempio 13.1.11. Continuiamo ad analizzare la conica C dell'Esempio 13.1.8. Dopo aver ruotato gli assi, il polinomio che definisce C ha coefficienti

$$A = \begin{pmatrix} 25 & 0 \\ 0 & 50 \end{pmatrix}, \quad b = \begin{pmatrix} 11 \\ 2 \end{pmatrix}, \quad c = 4.$$

La matrice A è invertibile e quindi troviamo un centro

$$P = -A^{-1}b = -\frac{1}{50} \begin{pmatrix} 2 & 0 \\ 0 & 1 \end{pmatrix} \begin{pmatrix} 11 \\ 2 \end{pmatrix} = -\frac{1}{25} \begin{pmatrix} 11 \\ 1 \end{pmatrix}.$$

Trasliamo le coordinate con $\begin{pmatrix} x \\ y \end{pmatrix} = \begin{pmatrix} x' \\ y' \end{pmatrix} + P$. Per il Corollario 13.1.6 il nuovo polinomio che definisce C ha coefficienti

$$A' = A = \begin{pmatrix} 25 & 0 \\ 0 & 50 \end{pmatrix}, \quad b' = 0, \quad c' = {}^t P A P + 2 {}^t b P + c = -\frac{23}{25}.$$

Nelle nuove coordinate la conica C ha equazione

$$25(x')^2 + 50(y')^2 - \frac{23}{25} = 0.$$

Con le tecniche che descriveremo nella Sezione 13.2 potremo dedurre che C è un'ellisse e studiarla geometricamente.

13.1.5. Il caso A invertibile. Assemblando i risultati dimostrati nelle ultime pagine, possiamo descrivere in modo soddisfacente tutte le quadriche in cui la matrice A sia invertibile.

Sia Q una quadrica in $\mathbb{R}^n$, luogo di zeri di un polinomio

$${}^t x A x + 2 {}^t b x + c.$$

Proposizione 13.1.12. *Se A è invertibile, esiste un sistema di coordinate x' in cui Q abbia equazione*

$$a'_1(x'_1)^2 + \cdots + a'_n(x'_n)^2 + c' = 0$$

con $a'_1, \dots, a'_n \neq 0$.

Dimostrazione. Con la Proposizione 13.1.7 possiamo sostituire A con una matrice diagonale congruente ad A . Siccome A è invertibile, la quadrica ha un centro P e spostando successivamente l'origine in P eliminiamo b lasciando inalterato A . Il risultato è un'equazione della forma descritta. $\square$

Se A è invertibile esiste un sistema di coordinate in cui $\bar{A}$ diventa diagonale. Abbiamo analizzato un caso concreto negli Esempi 13.1.8 e 13.1.11.

13.2. Coniche

Studiamo in questa sezione le *coniche*, cioè le quadriche in $\mathbb{R}^2$. Definiamo ellissi, iperboli e parabole in modo geometrico e dimostriamo successivamente che queste sono luoghi di zeri di equazioni di secondo grado. Quindi studiamo le loro proprietà e le classifichiamo completamente.

13.2.1. Circonferenza. Abbiamo già definito nella Sezione 10.1.4 la *circonferenza* di centro $P \in \mathbb{R}^2$ e di raggio $r > 0$ come l'insieme

$$C = \{Q \in \mathbb{R}^2 \mid d(Q, P) = r\}.$$

Se $P = \begin{pmatrix} x_0 \\ y_0 \end{pmatrix}$ e $Q = \begin{pmatrix} x \\ y \end{pmatrix}$, abbiamo riscritto questa condizione così:

$$(x - x_0)^2 + (y - y_0)^2 = r^2.$$

13.2.2. Ellisse. Siano F_1 e F_2 due punti di $\mathbb{R}^2$. L'*ellisse* con *fuochi* F_1, F_2 e *semiasse* $a > 0$ è l'insieme

$$E = \{Q \in \mathbb{R}^2 \mid d(F_1, Q) + d(F_2, Q) = 2a\}.$$

Si veda la Figura 13.1-(sinistra). Chiediamo sempre che $2a$ sia strettamente maggiore della distanza $2c$ fra i due fuochi, perché i casi $2a \leq 2c$ non sono interessanti: se $2a = 2c$ allora E coincide con il segmento $F_1 F_2$ e se $2a < 2c$ l'insieme E è vuoto (per la disuguaglianza triangolare).

Notiamo che se $F_1 = F_2$ l'ellisse è una circonferenza di centro F_1 e raggio a . D'ora in poi supponiamo che $F_1 \neq F_2$. È conveniente fissare

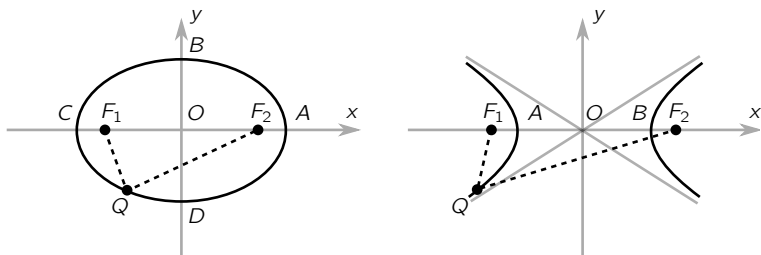


Figura 13.1. L'ellisse è il luogo di punti Q del piano con $|F_1Q| + |F_2Q| = 2a$ (sinistra). L'iperbole è il luogo di punti Q del piano con $||F_1Q| - |F_2Q|| = 2a$ (destra).

un sistema di riferimento in modo che l'origine sia il punto medio O del segmento F_1F_2 e l'asse x contenga questo segmento, come nella Figura 13.1-(sinistra). Dimostriamo che l'ellisse è effettivamente una conica.

Proposizione 13.2.1. *L'ellisse è la conica di equazione*

$$\frac{x^2}{a^2} + \frac{y^2}{b^2} = 1$$

dove $b^2 = a^2 - c^2$.

Dimostrazione. Abbiamo $F_1 = \begin{pmatrix} -c \\ 0 \end{pmatrix}$ e $F_2 = \begin{pmatrix} c \\ 0 \end{pmatrix}$ e per definizione l'ellisse è l'insieme dei punti che soddisfano l'equazione

$$\sqrt{(x+c)^2 + y^2} + \sqrt{(x-c)^2 + y^2} = 2a.$$

Possiamo elevare al quadrato entrambi i membri, dividere per due e trovare

$$x^2 + y^2 + c^2 - 2a^2 = -\sqrt{(x+c)^2 + y^2} \sqrt{(x-c)^2 + y^2}.$$

Eleviamo ancora al quadrato e facciamo un po' di conti fino ad ottenere

$$(a^2 - c^2)x^2 + a^2y^2 = a^2(a^2 - c^2).$$

Con $b^2 = a^2 - c^2$ questa equazione è equivalente a quella dell'enunciato. $\square$

Il punto medio O fra i due fuochi F_1 e F_2 è un centro di simmetria per l'ellisse. Dall'equazione ricaviamo che i punti A, B, C, D nella Figura 13.1-(sinistra) hanno coordinate

$$A = \begin{pmatrix} a \\ 0 \end{pmatrix}, \quad B = \begin{pmatrix} 0 \\ b \end{pmatrix}, \quad C = \begin{pmatrix} -a \\ 0 \end{pmatrix}, \quad D = \begin{pmatrix} 0 \\ -b \end{pmatrix}$$

In particolare $2a$ e $2b$ sono le lunghezze degli assi AC e BD . Notiamo che

$$|OF_2| = c, \quad |OB| = b \implies |BF_2| = \sqrt{b^2 + c^2} = a = |OA|.$$

Una descrizione parametrica dell'ellisse è la seguente:

$$E = \left\{ \begin{pmatrix} a \cos \theta \\ b \sin \theta \end{pmatrix} \mid \theta \in [0, 2\pi) \right\}.$$

13.2.3. Iperbole. Siano F_1 e F_2 due punti distinti di $\mathbb{R}^2$. L'*iperbole* con fuochi $F_1, F_2 \in \mathbb{R}^2$ e semiasse $a > 0$ è l'insieme

$$I = \{Q \in \mathbb{R}^2 \mid |d(F_1, Q) - d(F_2, Q)| = 2a\}.$$

Si veda la Figura 13.1-(destra). Chiediamo sempre che $2a$ sia strettamente minore della distanza $2c$ fra i due fuochi, perché i casi $2a \geq 2c$ non sono interessanti: se $2a = 2c$ allora I coincide con le due semirette uscenti dai fuochi ciascuna in direzione opposta all'altro fuoco e se $2a > 2c$ allora $I = \emptyset$ (per la disuguaglianza triangolare).

Come con l'ellisse, è utile fissare un sistema di riferimento in cui l'origine sia il punto medio O del segmento F_1F_2 e l'asse x contenga questo segmento come nella Figura 13.1-(destra). Dimostriamo che l'iperbole è effettivamente una conica.

Proposizione 13.2.2. *L'iperbole è la conica di equazione*

$$\frac{x^2}{a^2} - \frac{y^2}{b^2} = 1$$

dove $b^2 = c^2 - a^2$.

Dimostrazione. Abbiamo $F_1 = \begin{pmatrix} -c \\ 0 \end{pmatrix}$ e $F_2 = \begin{pmatrix} c \\ 0 \end{pmatrix}$ e per definizione l'iperbole è l'insieme dei punti che soddisfano l'equazione

$$|\sqrt{(x+c)^2 + y^2} - \sqrt{(x-c)^2 + y^2}| = 2a.$$

Possiamo elevare al quadrato entrambi i membri, dividere per due e trovare

$$x^2 + y^2 + c^2 - 2ax = \sqrt{(x+c)^2 + y^2} \sqrt{(x-c)^2 + y^2}.$$

Gli stessi conti dell'ellisse portano alla stessa equazione di prima

$$(a^2 - c^2)x^2 + a^2y^2 = a^2(a^2 - c^2).$$

Con $b^2 = c^2 - a^2$ questa equazione è equivalente a quella dell'enunciato. $\square$

Il punto medio O fra i due fuochi F_1 e F_2 è un centro di simmetria per l'iperbole. Dall'equazione ricaviamo che i punti A, B nella Figura 13.1-(destra) hanno coordinate

$$A = \begin{pmatrix} -a \\ 0 \end{pmatrix}, \quad B = \begin{pmatrix} a \\ 0 \end{pmatrix}.$$

Per scrivere l'iperbole in forma parametrica dobbiamo introdurre alcune funzioni trigonometriche.

Definizione 13.2.3. Le funzioni *seno iperbolico* $\sinh: \mathbb{R} \rightarrow \mathbb{R}$, *coseno iperbolico* $\cosh: \mathbb{R} \rightarrow \mathbb{R}$ e *tangente iperbolica* $\tanh: \mathbb{R} \rightarrow \mathbb{R}$ sono le seguenti:

$$\sinh(x) = \frac{e^x - e^{-x}}{2}, \quad \cosh(x) = \frac{e^x + e^{-x}}{2}, \quad \tanh(x) = \frac{\sinh(x)}{\cosh(x)}.$$

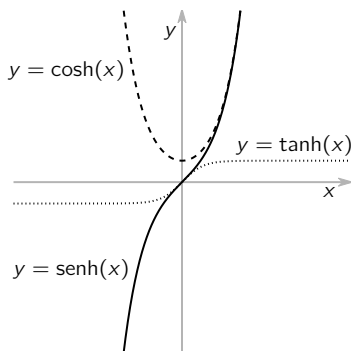


Figura 13.2. I grafici delle funzioni $\sinh(x)$, $\cosh(x)$ e $\tanh(x)$.

I grafici delle funzioni sono mostrati nella Figura 13.2. Queste nuove funzioni trigonometriche hanno alcune somiglianze con le usuali seno, coseno e tangente. Intanto notiamo che $\sinh(x)$ è una funzione dispari, cioè $\sinh(-x) = -\sinh(x)$, mentre $\cosh(x)$ è una funzione pari, cioè $\cosh(-x) = \cosh(x)$. Inoltre è facile vedere che

$$\cosh^2(x) - \sinh^2(x) = 1 \quad \forall x \in \mathbb{R}.$$

Una descrizione parametrica dell'iperbole è la seguente:

$$(32) \quad I = \left\{ \begin{pmatrix} \pm a \cosh \theta \\ b \sinh \theta \end{pmatrix} \mid \theta \in \mathbb{R} \right\}.$$

L'iperbole ha due componenti (dette *rami*) ed il segno $\pm$ dipende da quale componente stiamo parametrizzando. Le due rette

$$y = \pm \frac{b}{a}x$$

mostrate nella Figura 13.1-(destra) sono gli *asintoti* dell'iperbole. Gli asintoti sono caratterizzati dalla proprietà seguente, che usa la parametrizzazione (32).

Proposizione 13.2.4. *Quando $\theta \rightarrow \pm\infty$, la distanza del punto dell'iperbole da uno degli asintoti tende a zero.*

Dimostrazione. Consideriamo i punti

$$P(\theta) = \begin{pmatrix} \pm a \cosh \theta \\ b \sinh \theta \end{pmatrix}, \quad Q(\theta) = \begin{pmatrix} \pm a \cosh \theta \\ b \cosh \theta \end{pmatrix}.$$

Abbiamo $P(\theta) \in I$ mentre $Q(\theta)$ è contenuto in uno dei due asintoti. La distanza fra i due punti è

$$d(P(\theta), Q(\theta)) = b(\cosh \theta - \sinh \theta) = be^{-\theta}$$

e tende a zero se $\theta \rightarrow +\infty$. Per $\theta \rightarrow -\infty$ si sceglie $Q(\theta)$ con seconda coordinata $-b \cosh \theta$ e si ragiona in modo analogo. $\square$

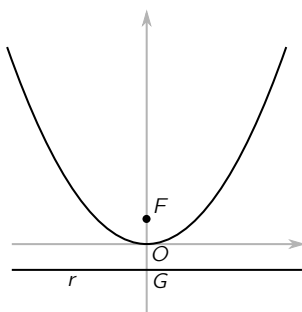


Figura 13.3. Una parabola.

13.2.4. Parabola. Siano F e r un punto ed una retta disgiunti in $\mathbb{R}^2$. La parabola con fuoco F e direttrice r è l'insieme

$$P = \{Q \in \mathbb{R}^2 \mid d(Q, F) = d(Q, r)\}.$$

Si veda la Figura 13.3. Sia $2a$ la distanza fra F e r . Siano G la proiezione ortogonale di F su r ed O il punto medio del segmento FG . Fissiamo un sistema di riferimento con l'origine in O e l'asse x parallela a r , come in figura. Dimostriamo che la parabola è effettivamente una conica.

Proposizione 13.2.5. La parabola è la conica di equazione

$$y = \frac{x^2}{4a}.$$

Dimostrazione. Abbiamo $F = \begin{pmatrix} 0 \\ a \end{pmatrix}$, $G = \begin{pmatrix} 0 \\ -a \end{pmatrix}$ e $r = \{y = -a\}$. Per definizione la parabola è l'insieme dei punti che soddisfano l'equazione

$$\sqrt{x^2 + (y - a)^2} = |y + a|$$

Elevando al quadrato si ottiene semplicemente

$$x^2 - 4ay = 0$$

che è equivalente all'equazione scritta nell'enunciato. $\square$

Notiamo che una parabola *non* ha un centro di simmetria. A differenza dell'ellissi e dell'iperbole, la parabola è il grafico di una funzione e quindi può essere parametrizzata semplicemente come

$$P = \left\{ \begin{pmatrix} t \\ \frac{t^2}{4a} \end{pmatrix} \mid t \in \mathbb{R} \right\}.$$

13.2.5. Eccentricità. Vedremo successivamente che ellissi, iperboli e parabole sono molto collegate fra loro. Per adesso notiamo che sono tutte e tre risultato di una medesima costruzione geometrica, il cui risultato differisce per un numero positivo chiamato *eccentricità*.

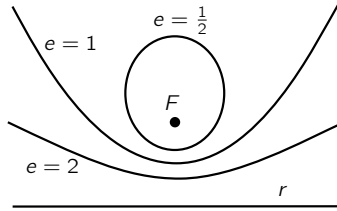


Figura 13.4. A seconda dell'eccentricità $e > 0$ otteniamo una iperbole, parabola o ellisse.

Siano F e r un punto ed una retta disgiunti in $\mathbb{R}^2$. Sia $e > 0$ un numero, chiamato *eccentricità*. Studiamo l'insieme

$$C = \{Q \in \mathbb{R}^2 \mid d(Q, F) = ed(Q, r)\}.$$

Si veda la Figura 13.4. La proposizione seguente mostra che C è sempre una conica, il cui tipo dipende dall'eccentricità e .

Proposizione 13.2.6. *L'insieme C è un'ellisse, una parabola o un'iperbole, a seconda che sia $e < 1$, $e = 1$ oppure $e > 1$.*

Dimostrazione. Se $e = 1$ siamo a posto, quindi supponiamo $e \neq 1$. Sia $d = d(F, r)$. Fissiamo un sistema di riferimento in cui

$$F = \left(\frac{e^2 d}{|1 - e^2|}, 0 \right), \quad r = \left\{ x = \frac{d}{|1 - e^2|} \right\}.$$

Notiamo che effettivamente

$$d(F, r) = \frac{|e^2 d - d|}{|1 - e^2|} = d.$$

I punti di C sono quelli che soddisfano l'equazione

$$\sqrt{\left(x - \frac{e^2 d}{|1 - e^2|}\right)^2 + y^2} = e \left|x - \frac{d}{|1 - e^2|}\right|.$$

Eleviamo al quadrato e troviamo

$$\left(x - \frac{e^2 d}{|1 - e^2|}\right)^2 + y^2 = e^2 \left(x - \frac{d}{|1 - e^2|}\right)^2$$

che diventa

$$(1 - e^2)x^2 + y^2 = \frac{e^2 d^2}{1 - e^2}.$$

Dividendo per $e^2 d^2 / (1 - e^2)$ troviamo l'equazione di un'ellisse se $e < 1$ e di un'iperbole se $e > 1$. $\square$

Il punto F e la retta r sono il *fuoco* e la *direttrice* della conica C .

13.2.6. Classificazione delle coniche non degeneri. Vogliamo classificare completamente le coniche: ci concentriamo soprattutto su quelle non degeneri, che come vedremo sono le più interessanti.

Sia C una conica in $\mathbb{R}^2$, di equazione

$${}^t x A x + 2 {}^t b x + c = 0.$$

Ricordiamo che per definizione C è non degenera se la matrice simmetrica

$$\bar{A} = \begin{pmatrix} A & b \\ {}^t b & c \end{pmatrix}$$

è invertibile. Consideriamo intanto questo caso. Poiché $\bar{A}$ è simmetrica ed invertibile, le sue possibili segnature sono

$$(3, 0, 0), \quad (2, 1, 0), \quad (1, 2, 0), \quad (0, 3, 0).$$

Per l'Osservazione 13.1.3, possiamo cambiare $\bar{A}$ con $-\bar{A}$ e in questo modo cambiare la segnatura da $(p, m, 0)$ a $(m, p, 0)$ senza alterare la conica C . Quindi in realtà i casi possibili che ci interessano sono due: diciamo che la matrice $\bar{A}$ è *definita* se ha segnatura $(3, 0, 0)$ o $(0, 3, 0)$ (corrispondenti ad un prodotto scalare definito positivo o negativo) e *indefinita* negli altri casi $(2, 1, 0)$ o $(1, 2, 0)$ (corrispondenti ad un prodotto scalare indefinito).

Teorema 13.2.7. *Se C è non degenera, si presentano i casi seguenti:*

- (1) *Se $\bar{A}$ è definita, allora la conica C è vuota, cioè $C = \emptyset$.*
- (2) *Se $\bar{A}$ è indefinita, allora la conica C è*
 - *un'iperbole se $\det A < 0$,*
 - *una parabola se $\det A = 0$,*
 - *un'ellisse se $\det A > 0$.*

Dimostrazione. Per la Proposizione 13.1.7 possiamo supporre che A sia diagonale. Consideriamo prima il caso $\det A \neq 0$. Per la Proposizione 13.1.12 possiamo supporre che anche $\bar{A} = \begin{pmatrix} A & 0 \\ 0 & c \end{pmatrix}$ sia diagonale. Siccome $\bar{A}$ è invertibile, abbiamo $c \neq 0$. Possiamo quindi dividere tutto per c e ottenere che C abbia equazione

$$a_{11}x^2 + a_{22}y^2 + 1 = 0$$

con $a_{11}, a_{22} \neq 0$. Se a_{11}, a_{22} sono entrambi positivi otteniamo l'insieme vuoto, se sono entrambi negativi otteniamo un'ellisse e se hanno segni opposti un'iperbole. Si presenta il primo caso se $\bar{A}$ è definita. Gli altri due se $\bar{A}$ è indefinita e $\det A$ è positivo o negativo, rispettivamente.

Resta da considerare il caso in cui $\det A = 0$. A meno di scambiare le variabili x e y otteniamo

$$\bar{A} = \begin{pmatrix} a_{11} & 0 & b_1 \\ 0 & 0 & b_2 \\ b_1 & b_2 & c \end{pmatrix}.$$

Qui $\det \bar{A} = -b_2^2 a_{11} \neq 0$ e quindi $a_{11}, b_2 \neq 0$. La conica C ha equazione

$$a_{11}x^2 + 2b_1x + 2b_2y + c = 0.$$

Dopo aver diviso tutto per $2b_2$ otteniamo un'equazione del tipo

$$y = ax^2 + bx + c$$

con opportuni coefficienti a, b, c e $a \neq 0$. Una traslazione opportuna (esercizio) trasforma l'equazione nella forma $y = ax^2$. Quindi C è una parabola. $\square$

13.2.7. Classificazione delle coniche degeneri. Classifichiamo adesso le coniche degeneri. Durante la classificazione scopriremo che effettivamente queste coniche sono meno interessanti delle precedenti.

Sia C una conica degenera, luogo di zeri di un'equazione

$${}^t x A x + 2 {}^t b x + c.$$

Per ipotesi la matrice simmetrica

$$\bar{A} = \begin{pmatrix} A & b \\ {}^t b & c \end{pmatrix}$$

ha rango ≤ 2 . Nel caso in cui il rango sia 2 le signature possibili sono

$$(2, 0, 1), \quad (1, 1, 1), \quad (0, 2, 1).$$

In analogia con il caso non-degenere diciamo che $\bar{A}$ è *semidefinita* nel primo e terzo caso e *indefinita* nel secondo.

Teorema 13.2.8. *Se C è degenera, si presentano i casi seguenti:*

- (1) Se $\text{rk} \bar{A} = 2$ e $\bar{A}$ è semidefinita, allora C è
 - un punto se $\det A \neq 0$,
 - l'insieme vuoto se $\det A = 0$.
- (2) Se $\text{rk} \bar{A} = 2$ e $\bar{A}$ è indefinita, allora C è
 - l'unione di due rette distinte incidenti se $\det A \neq 0$,
 - l'unione di due rette distinte parallele se $\det A = 0$.
- (3) Se $\text{rk} \bar{A} = 1$ allora C è una retta.

Dimostrazione. Per la Proposizione 13.1.7 possiamo supporre che A sia diagonale. Consideriamo prima il caso in cui A sia invertibile. Per la Proposizione 13.1.12 possiamo supporre che $\bar{A} = \begin{pmatrix} A & 0 \\ 0 & c \end{pmatrix}$ sia diagonale. Poiché $\bar{A}$ non è invertibile, abbiamo $c = 0$. Quindi C è descritta da un'equazione

$$a_{11}x^2 + a_{22}y^2 = 0$$

con $a_{11}, a_{22} \neq 0$. Se $\bar{A}$ è semidefinita, allora a_{11} e a_{22} hanno lo stesso segno e l'unica soluzione è il punto $\begin{pmatrix} 0 \\ 0 \end{pmatrix}$. Se $\bar{A}$ è indefinita hanno segno opposto e a meno di scambiare e rinominare le variabili scriviamo l'equazione come

$$a^2x^2 - b^2y^2 = (ax + by)(ax - by) = 0$$

con $a, b \neq 0$. Le soluzioni di questa equazione sono i punti delle due rette r e s di equazioni rispettivamente

$$ax + by = 0, \quad ax - by = 0.$$

Le due rette sono distinte e si intersecano nell'origine. Abbiamo $C = r \cup s$.

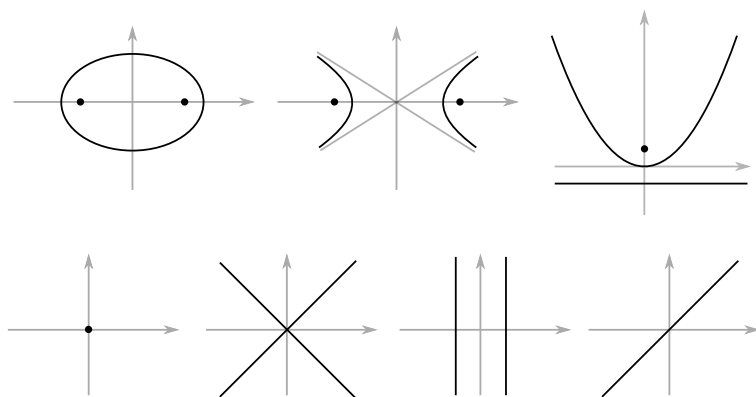


Figura 13.5. Le coniche non degeneri (ellisse, iperbole e parabola) e quelle degeneri (un punto, due rette incidenti, due rette parallele e una retta doppia). Tutte le coniche tranne la parabola hanno almeno un centro di simmetria. Il centro di simmetria è unico per l'ellissi, l'iperbole, il punto e le rette incidenti. Le ultime due coniche degeneri hanno invece più centri di simmetria, che formano una retta.

Ci resta da considerare il caso $\det A = 0$. A meno di scambiare le variabili x e y otteniamo

$$\bar{A} = \begin{pmatrix} a_{11} & 0 & b_1 \\ 0 & 0 & b_2 \\ b_1 & b_2 & c \end{pmatrix}$$

con $a_{11} \neq 0$. Poiché $\det \bar{A} = -a_{11}b_2^2$ si annulla, abbiamo $b_2 = 0$. La conica C ha quindi equazione

$$a_{11}x^2 + 2b_1x + c = 0.$$

Questa equazione di secondo grado ha 0, 1 oppure 2 soluzioni, e C è rispettivamente l'insieme vuoto, una retta verticale, o due rette verticali. Questo dipende ovviamente da $\Delta = 4(b_1^2 - a_{11}c)$ che è negativo, nullo o positivo a seconda che $\bar{A}$ sia semidefinito, di rango 1, oppure indefinito. $\square$

Nel caso in cui C sia una sola retta, questa viene chiamata *retta doppia*, perché è in un certo senso una retta con “molteplicità due”.

Adesso sappiamo identificare il tipo di conica esaminando le matrici $\bar{A}$ e A . Tutti i casi in cui la conica non è vuota sono riassunti nella Figura 13.5.

Esempio 13.2.9. Consideriamo la famiglia di coniche

$$C_t = \{(2t - 1)x^2 + 6txy + ty^2 + 2x = 0\}$$

dipendenti da un parametro $t \in \mathbb{R}$. Cerchiamo di capire che tipo di conica è C_t al variare di t . Scriviamo la matrice

$$\bar{A} = \begin{pmatrix} 2t-1 & 3t & 1 \\ 3t & t & 0 \\ 1 & 0 & 0 \end{pmatrix}.$$

Il determinante è $\det \bar{A} = -t$ e quindi la conica è degenere solo per $t = 0$. Per questo valore $C_0 = \{-x^2 + 2x = 0\} = \{x(-x + 2) = 0\}$ è l'unione di due rette parallele $\{x = 0\}$ e $\{x = 2\}$.

Se $t \neq 0$ la conica è non degenere. Per il Lemma 7.4.16 e il Corollario 7.4.17, la matrice $\bar{A}$ non è definita (né positiva né negativa) e quindi è indefinita per ogni t . Ne segue che C_t non è mai vuota.

Calcoliamo $\det A = t(2t - 1) - 9t^2 = -7t^2 - t$ e deduciamo che C_t è una ellisse per $t \in (-\frac{1}{7}, 0)$, una parabola per $t = -\frac{1}{7}$ e una iperbole per $t \in (-\infty, -\frac{1}{7}) \cup (0, +\infty)$.

Dalla classificazione delle coniche possiamo dedurre il fatto seguente.

Proposizione 13.2.10. *Sia C una conica non vuota di equazione ${}^t\bar{x}\bar{A}\bar{x} = 0$. I centri di simmetria di C sono precisamente i punti P tali che $AP + b = 0$.*

Dimostrazione. Sappiamo che tali punti sono centri di simmetria, e analizzando le coniche nella Figura 13.5 si vede che non ce ne sono altri. $\square$

13.2.8. Fascio di coniche. Ricordiamo dalla Sezione 9.2.8 che un fascio di rette è una famiglia di rette dipendente da un paio di parametri omogenei. Un'analoga definizione funziona per le coniche.

Siano C_1 e C_2 due coniche distinte, descritte da due polinomi $p_1(x)$ e $p_2(x)$ di secondo grado non multipli fra loro. Il *fascio di coniche* generato da $p_1(x)$ e $p_2(x)$ è l'insieme di coniche $C_{t,u}$ determinate dall'equazione

$$C_{t,u} = \{x \in \mathbb{R}^2 \mid tp_1(x) + up_2(x) = 0\}.$$

Qui t, u sono due numeri reali arbitrari, non entrambi nulli. Abbiamo $C_{1,0} = C_1$ e $C_{0,1} = C_2$. Notiamo che $C_{t,u} = C_{\lambda t, \lambda u}$ perché le due coniche sono descritte da polinomi multipli.

Esempio 13.2.11. L'insieme di coniche descritto nell'Esempio 13.2.9 può essere completato ad un fascio di coniche nel modo seguente:

$$C_{t,u} = \{(2t - u)x^2 + 6txy + ty^2 + 2ux = 0\}$$

Questo è il fascio generato dai polinomi

$$p_1(x) = 2x^2 + 6xy + y^2, \quad p_2(x) = -x^2 + 2x.$$

Osservazione 13.2.12. Un polinomio di grado ≤ 2 in due variabili

$$a_{11}x^2 + 2a_{12}xy + a_{22}y^2 + 2b_1x + 2b_2y + c$$

è determinato dal punto $(a_{11}, a_{12}, a_{22}, b_1, b_2, c) \in \mathbb{R}^6$. I polinomi di grado ≤ 2 sono parametrizzati da punti in $\mathbb{R}^6$. Se consideriamo i polinomi a meno di moltiplicazione per scalare, come è naturale fare nel contesto delle coniche, questi formano lo spazio proiettivo $\mathbb{P}(\mathbb{R}^6) = \mathbb{P}^5$.

Con questa interpretazione, un fascio $C_{t,u}$ generato da $p_1(x)$ e $p_2(x)$ è semplicemente la retta proiettiva in $\mathbb{P}^5$ che passa per i due punti $p_1(x)$ e $p_2(x)$. Fasci di coniche e rette proiettive in $\mathbb{P}^5$ sono la stessa cosa.

I fasci sono uno strumento potente nello studio delle coniche. Mostriamo intanto che un fascio di coniche copre tutti i punti del piano.

Proposizione 13.2.13. *Sia $C_{t,u}$ un fascio di coniche. Per ogni $P \in \mathbb{R}^2$ esiste una conica $C_{t,u}$ del fascio tale che $P \in C_{t,u}$.*

Dimostrazione. Abbiamo $C_{t,u} = \{tp_1(x) + up_2(x) = 0\}$. Imponiamo che $C_{t,u}$ passi per P scrivendo $tp_1(P) + up_2(P) = 0$. Questa equazione ha sempre una soluzione $(t, u) \neq (0, 0)$. Infatti: se $p_1(P) = p_2(P) = 0$ allora qualsiasi (t, u) è soluzione, altrimenti $t = -p_2(P)$ e $u = p_1(P)$ funzionano. $\square$

Esempio 13.2.14. Cerchiamo nel fascio dell'Esempio 13.2.11 la conica passante per $P = \begin{pmatrix} 3 \\ 2 \end{pmatrix}$. Imponiamo

$$(2t - u)3^2 + 6t3 \cdot 2 + t2^2 + 2u3 = 0 \iff 58t - 3u = 0$$

e troviamo la soluzione $t = 3, u = 58$. La conica $C_{3,58}$ del fascio passa per P .

Notiamo anche il fatto seguente.

Proposizione 13.2.15. *Se due coniche distinte di un fascio passano per P , allora tutte le coniche del fascio passano per P .*

Dimostrazione. Siano $p_1(x)$ e $p_2(x)$ i polinomi relativi a queste due coniche: possiamo prendere loro come generatori del fascio. Se $p_1(P) = 0$ e $p_2(P) = 0$, allora $tp_1(P) + up_2(P) = 0$ per ogni t, u . $\square$

Osservazione 13.2.16. Dalla nostra definizione di fascio $C_{t,u}$, può capitare in casi molto particolari che il polinomio $tp_1(x) + up_2(x)$ sia in realtà di grado 1 o 0 per qualche valore di t e u .

13.2.9. Coniche e rette. Studiamo in questo paragrafo le possibili posizioni reciproche di una conica C e una retta r nel piano. Iniziamo valutando la loro intersezione.

Proposizione 13.2.17. *L'intersezione $C \cap r$ fra una conica C e una retta r consta di 0, 1 o 2 punti, oppure è l'intera retta r .*

Dimostrazione. La conica è luogo di zeri di un polinomio di secondo grado

$$C = \{ {}^t x A x + 2 {}^t b x + c = 0 \}.$$

La retta r è in forma parametrica

$$r = \{P + tv \mid t \in \mathbb{R}\}.$$

Per studiare le intersezioni fra C e r , sostituiamo la forma parametrica $x = P + tv$ nella prima equazione e troviamo

$${}^t(P + tv)A(P + tv) + 2 {}^t b(P + tv) + c = 0.$$

Questa è un'equazione polinomiale in t di grado al massimo due:

$$(33) \quad ({}^t v A v) t^2 + 2({}^t P A v + {}^t b v) t + {}^t P A P + 2 {}^t b P + c = 0.$$

Questa equazione ha al massimo due soluzioni, tranne nel caso molto particolare in cui questo polinomio si annulla completamente e allora tutti i t sono soluzioni. $\square$

È possibile che C contenga r nel caso in cui C sia degenera.

Corollario 13.2.18. *Una conica C che passa per tre punti allineati è una conica degenera che contiene la retta r passante per i tre punti.*

Dimostrazione. Poiché $C \cap r$ contiene almeno 3 punti, è tutto r . $\square$

13.2.10. Conica per 5 punti dati. Sappiamo che per 3 punti non allineati passa esattamente una circonferenza: dimostriamo adesso un fatto analogo per le coniche. Diciamo che $k \geq 3$ punti nel piano sono in *posizione generale* se 3 qualsiasi di questi non sono mai allineati. Mostriamo che per 5 punti in posizione generale passa esattamente una conica.

Proposizione 13.2.19. *Siano $P_1, \dots, P_5 \in \mathbb{R}^2$ punti in posizione generale. Esiste un'unica conica C che li contenga tutti. La conica C è non degenera.*

Dimostrazione. Ogni P_i ha coordinate $\begin{pmatrix} x_i \\ y_i \end{pmatrix}$. Scriviamo una conica generica come soluzione dell'equazione

$$(34) \quad a_{11}x^2 + 2a_{12}xy + a_{22}y^2 + 2b_1x + 2b_2y + c = 0.$$

Imponiamo il passaggio per $P_1, \dots, P_5$ e otteniamo un sistema lineare omogeneo di 5 equazioni dove le variabili sono i 6 coefficienti $a_{11}, a_{21}, a_{22}, b_1, b_2, c$

$$\begin{cases} a_{11}x_1^2 + 2a_{12}x_1y_1 + a_{22}y_1^2 + 2b_1x_1 + 2b_2y_1 + c = 0, \\ \vdots \\ a_{11}x_5^2 + 2a_{12}x_5y_5 + a_{22}y_5^2 + 2b_1x_5 + 2b_2y_5 + c = 0. \end{cases}$$

Questo sistema ha almeno una retta vettoriale di soluzioni. Quindi esiste una soluzione non nulla, ovvero una conica C che passa per i 5 punti.

La conica C è non degenera: se fosse degenera, sarebbe un'unione di una o due rette, e una di queste rette dovrebbe contenere almeno 3 dei 5 punti, un assurdo perché i 5 punti sono a 3 non allineati.

Dobbiamo infine dimostrare che C è unica. Supponiamo per assurdo che ci siano due coniche distinte C_1, C_2 che passano per i 5 punti, descritte da due polinomi $p_1(x, y)$ e $p_2(x, y)$ con coefficienti non multipli fra loro.

Tutte le coniche del fascio $C_{t,u} = \{tp_1(x,y) + up_2(x,y)\}$ passano per i 5 punti e non sono degeneri (per lo stesso argomento già usato). Sia P_6 un punto qualsiasi allineato con P_1 e P_2 . Per la Proposizione 13.2.13 esiste una conica $C_{t,u}$ del fascio che passa anche per P_6 . Questa conica contiene tre punti allineati P_1, P_2, P_6 e quindi è degenera per il Corollario 13.2.18, un assurdo. $\square$

Per determinare concretamente la conica che passa per 5 punti, un metodo consiste nel risolvere il sistema (34) con 5 equazioni e 6 incognite. Introduciamo adesso un modo alternativo meno dispendioso, che fa uso ancora una volta dei fasci.

13.2.11. Fascio di coniche per 4 punti. Abbiamo dimostrato che per 5 punti in posizione generale passa una sola conica. È naturale adesso chiedersi cosa succeda se consideriamo solo 4 punti.

Proposizione 13.2.20. Le coniche che passano per 4 punti in posizione generale formano un fascio.

Dimostrazione. Siano P_1, P_2, P_3, P_4 i quattro punti in posizione generale e scriviamo $P_i = \begin{pmatrix} x_i \\ y_i \end{pmatrix}$ per ogni i . Le coniche passanti per i 4 punti sono quelle i cui coefficienti soddisfanno le equazioni

$$\begin{cases} a_{11}x_1^2 + 2a_{12}x_1y_1 + a_{22}y_1^2 + 2b_1x_1 + 2b_2y_1 + c = 0, \\ \vdots \\ a_{11}x_4^2 + 2a_{12}x_4y_4 + a_{22}x_4^2 + 2b_1x_4 + 2b_2y_4 + c = 0. \end{cases}$$

Lo spazio delle soluzioni ha dimensione almeno $6 - 4 = 2$. Mostriamo che deve essere esattamente 2. Sia P_5 un punto tale che $P_1, \dots, P_5$ siano ancora in posizione generale. Se lo spazio delle soluzioni fosse almeno 3, aggiungendo come condizione il passaggio per P_5 otterrei uno spazio di dimensione almeno 2, ma per 5 punti in posizione generale passa una sola conica, assurdo.

Lo spazio delle soluzioni ha dimensione 2 e forma quindi un fascio, determinato da due generatori qualsiasi $p_1(x)$ e $p_2(x)$. $\square$

Nel fascio di coniche che passano per 4 punti in posizione generale ci sono tre coniche degeneri: queste sono i tre tipi di coppie di rette che passano per i 4 punti mostrate nella Figura 13.6. Possiamo scrivere agevolmente il fascio prendendo due di queste coniche degeneri come generatori.

Esempio 13.2.21. Descriviamo il fascio di coniche passante per i punti

$$P_1 = \begin{pmatrix} 1 \\ 0 \end{pmatrix}, \quad P_2 = \begin{pmatrix} 1 \\ -1 \end{pmatrix}, \quad P_3 = \begin{pmatrix} 2 \\ 1 \end{pmatrix}, \quad P_4 = \begin{pmatrix} 3 \\ 0 \end{pmatrix}.$$

Per fare ciò determiniamo due delle tre coniche degeneri appena descritte; le coniche degeneri C_1 e C_2 disegnate nella Figura 13.6-(destra) hanno

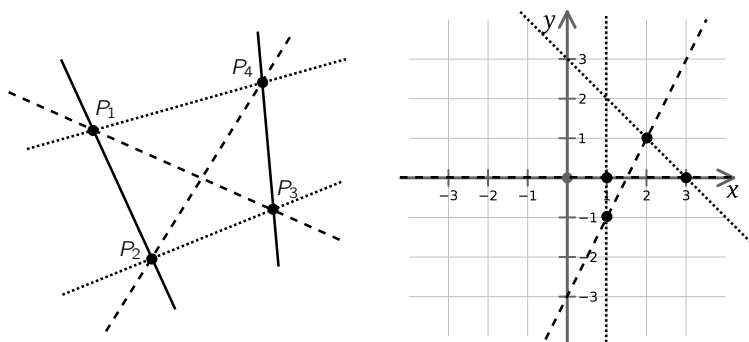


Figura 13.6. Il fascio di coniche passante per quattro punti in posizione generale contiene tre coniche degeneri: ciascuna di queste è l'unione di due rette (sinistra). Come determinare rapidamente il fascio di coniche per quattro punti (destra).

equazioni

$$(x-1)(x+y-3)=0, \quad y(2x-y-3)=0.$$

Quindi il fascio è semplicemente

$$\begin{aligned} C_{t,u} &= \{t(x-1)(x+y-3) + uy(2x-y-3) = 0\} \\ &= \{tx^2 - uy^2 + (t+2u)xy - 4tx + (-t-3u)y + 3t = 0\}. \end{aligned}$$

I fasci possono essere usati per determinare agevolmente la conica passante per 5 punti in posizione generale.

Esempio 13.2.22. Determiniamo la conica passante per i 5 punti

$$P_1 = \begin{pmatrix} 1 \\ 0 \end{pmatrix}, \quad P_2 = \begin{pmatrix} 1 \\ -1 \end{pmatrix}, \quad P_3 = \begin{pmatrix} 2 \\ 2 \end{pmatrix}, \quad P_4 = \begin{pmatrix} 3 \\ 0 \end{pmatrix}, \quad P_5 = \begin{pmatrix} -1 \\ -1 \end{pmatrix}.$$

Abbiamo già determinato il fascio di coniche passanti per i primi quattro:

$$C_{t,u} = \{t(x-1)(x+y-3) + uy(2x-y-3) = 0\}.$$

A questo punto imponiamo semplicemente il passaggio per P_5 e otteniamo

$$t(-2)(-5) + u(-1)(-4) = 0 \iff 10t + 4u = 0.$$

Una soluzione è ad esempio $t = 2, u = -5$. Quindi la conica passante per i 5 punti ha equazione

$$\begin{aligned} C_{2,-5} &= \{2(x-1)(x+y-3) - 5y(2x-y-3) = 0\} \\ &= \{2x^2 + 5y^2 - 8xy - 8x + 13y + 6 = 0\}. \end{aligned}$$

13.2.12. Polare. Introduciamo adesso una relazione fra punti e rette detta *polare*. Useremo questa nozione per definire e studiare le condizioni di tangenza fra coniche e rette.

Sia C una conica non vuota in $\mathbb{R}^2$, descritta da un'equazione

$${}^t\bar{x}\bar{A}\bar{x} = 0.$$

Sia $P \in \mathbb{R}^2$ un punto qualsiasi. La *polare* di P rispetto a C è il luogo di punti

$$r = \{x \in \mathbb{R}^2 \mid {}^t\bar{P}\bar{A}\bar{x} = 0\}.$$

Usiamo il simbolo r per indicare la polare di P perché questa è effettivamente quasi sempre una retta. Ricordiamo dalla Proposizione 13.2.10 che un punto P è un centro di simmetria per C se e solo se $AP + b = 0$.

Proposizione 13.2.23. *Valgono i fatti seguenti:*

- Se P non è un centro per C , allora r è la retta di equazione

$$(35) \quad {}^t(AP + b)x + {}^tPb + c = 0.$$

- Se P è un centro e $P \in C$, allora $r = \mathbb{R}^2$
- Se P è un centro e $P \notin C$, allora $r = \emptyset$.

Dimostrazione. Sviluppando si ottiene

$$\begin{aligned} r &= \left\{ \begin{pmatrix} {}^tP & 1 \end{pmatrix} \begin{pmatrix} A & b \\ {}^tb & c \end{pmatrix} \begin{pmatrix} x \\ 1 \end{pmatrix} = 0 \right\} \\ &= \{({}^tPA + {}^tb)x + {}^tPb + c = 0 \}. \end{aligned}$$

Se P non è un centro, allora $AP + b \neq 0$, quindi ${}^tPA + {}^tb \neq 0$ ed effettivamente l'equazione descrive una retta.

Se P è un centro, allora $AP + b = 0$ e $r = \{ {}^tPb + c = 0 \}$. Notiamo che

$$P \in C \iff {}^tPAP + 2 {}^tPb + c = 0 \iff {}^tPb + c = 0.$$

Quindi $r = \mathbb{R}^2$ se $P \in C$ e $r = \emptyset$ se $P \notin C$. □

Esempio 13.2.24. Sia C la conica di equazione

$$x^2 - 2xy - y^2 - 2x + 3 = 0.$$

Si verifica facilmente che la conica è un'iperbole con centro in $O = \frac{1}{2} \begin{pmatrix} 1 \\ -1 \end{pmatrix}$. Scriviamo

$$\bar{A} = \begin{pmatrix} 1 & -1 & -1 \\ -1 & -1 & 0 \\ -1 & 0 & 3 \end{pmatrix}, \quad \bar{x} = \begin{pmatrix} x \\ y \\ 1 \end{pmatrix}, \quad \bar{P} = \begin{pmatrix} x_0 \\ y_0 \\ 1 \end{pmatrix}.$$

La polare di P ha equazione

$$\begin{pmatrix} x_0 & y_0 & 1 \end{pmatrix} \begin{pmatrix} 1 & -1 & -1 \\ -1 & -1 & 0 \\ -1 & 0 & 3 \end{pmatrix} \begin{pmatrix} x \\ y \\ 1 \end{pmatrix} = 0$$

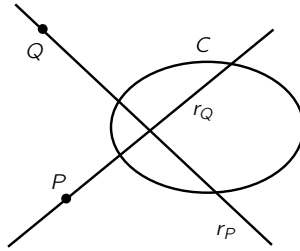


Figura 13.7. La legge di reciprocità: un punto Q è nella polare r_P di P se e solo se P è nella polare r_Q di Q .

cioè

$$r = \{(x_0 - y_0 - 1)x - (x_0 + y_0)y - x_0 + 3 = 0\}.$$

Questa è una retta per ogni $P \neq O$. Per $P = O$ è l'insieme vuoto.

Se la conica è non degenera, la polare è un'operazione iniettiva:

Proposizione 13.2.25. *Se C è non degenera, punti non centrali differenti hanno rette polari differenti.*

Dimostrazione. La matrice $\bar{A}$ è invertibile perché C è non degenera. Punti diversi P, P' danno vettori $\bar{P}, \bar{P}' \in \mathbb{R}^3$ indipendenti. Poiché $\bar{A}$ è invertibile, anche $\bar{A}\bar{P}$ e $\bar{A}\bar{P}'$ sono vettori indipendenti di $\mathbb{R}^3$, e quindi ${}^t\bar{P}\bar{A}$ e ${}^t\bar{P}'\bar{A}$ sono i coefficienti di due rette affini distinte. $\square$

Sia come sopra $C \subset \mathbb{R}^2$ una conica non vuota di equazione ${}^t\bar{x}\bar{A}\bar{x} = 0$. La polare ha molte proprietà geometriche che ci aiutano a capire meglio la conica C . Mostriamo innanzitutto che i punti di C sono precisamente quelli contenuti nella propria polare. Sia $P \in \mathbb{R}^2$ un punto qualsiasi.

Proposizione 13.2.26. *La polare r di P contiene $P \iff P \in C$.*

Dimostrazione. Entrambe le condizioni sono equivalenti a ${}^t\bar{P}\bar{A}\bar{P} = 0$. $\square$

La costruzione della polare soddisfa una proprietà semplice ma cruciale, detta *legge di reciprocità*, mostrata in Figura 13.7. Siano $P, Q \in \mathbb{R}^2$ due punti qualsiasi del piano.

Proposizione 13.2.27 (Legge di reciprocità). *Vale il fatto seguente:*

$$Q \text{ sta nella polare di } P \iff P \text{ sta nella polare di } Q.$$

Dimostrazione. Entrambe le condizioni sono equivalenti a ${}^t\bar{P}\bar{A}\bar{Q} = 0$. $\square$

Applicando la Proposizione 13.2.23, troviamo una prima conseguenza.

Corollario 13.2.28. *Ogni centro Q di C è contenuto nelle polari di tutti i punti $P \in \mathbb{R}^2$ se $Q \in C$ e di nessun punto se $Q \notin C$.*

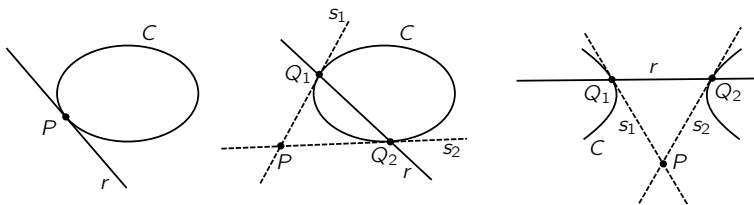


Figura 13.8. Se $P \in C$ non è centrale, la tangente r alla conica C in P è per definizione la polare di P . Se $P \notin C$, la sua polare r interseca la conica precisamente nei punti Q per cui la retta s contenente il segmento PQ è tangente a C .

Notiamo infine un fatto importante. La costruzione della polare si basa sulla matrice $\bar{A}$ e quindi potrebbe *a priori* dipendere non solo dalla conica C ma anche dal sistema di riferimento scelto. Questo per fortuna non è il caso: la costruzione dipende solo dalla conica.

Proposizione 13.2.29. *La polare non dipende dal sistema di riferimento.*

Dimostrazione. Sia r la polare di P . Se cambiamo sistema di riferimento con $\bar{x} = \bar{M}\bar{x}'$, otteniamo $\bar{A}' = {}^t\bar{M}\bar{A}\bar{M}$ e $\bar{P} = \bar{M}\bar{P}'$. Quindi

$${}^t\bar{P}\bar{A}\bar{x} = 0 \iff {}^t\bar{P}'{}^t\bar{M}\bar{A}\bar{M}\bar{x}' = 0 \iff {}^t\bar{P}'\bar{A}'\bar{x}' = 0.$$

La dimostrazione è conclusa. $\square$

13.2.13. Tangenza fra retta e conica. A cosa serve la polare? Principalmente a definire e studiare una nozione di tangenza fra rette e coniche.

Una conica C non degenera può intersecare una retta s in 0, 1 oppure 2 punti; diciamo rispettivamente che C ed s sono *disgiunte*, *tangenti* o *secanti*. Quindi C e s sono tangenti in un punto P precisamente se $C \cap s = \{P\}$.

Nel caso in cui C sia degenera la definizione di tangenza è diversa. Forniamo adesso una definizione generale che funziona per qualsiasi C . Vedremo successivamente che nel caso non degenera questa è equivalente a quella già data precedentemente.

Definizione 13.2.30. Siano C una conica, s una retta e P un punto contenuto in entrambe. Diciamo che s e C sono *tangenti in P* se la retta s è contenuta nella polare r di P .

Ci sono essenzialmente due casi da considerare:

- (1) Se $P \in C$ non è un centro, la polare r è una retta. Quindi l'unica retta s passante per P tangente a C è la polare $s = r$, di equazione

$${}^t(AP + b)x + {}^tPb + c = 0.$$

Si veda la Figura 13.8-(sinistra).

- (2) Se $P \in C$ è un centro, la polare è $r = \mathbb{R}^2$ e quindi per definizione *tutte* (sic!) le rette s passanti per P sono tangenti a C .

Esempio 13.2.31. Consideriamo l'ellisse standard di equazione

$$\frac{x^2}{a^2} + \frac{y^2}{b^2} = 1.$$

Si consideri un punto dell'ellisse

$$P = \begin{pmatrix} a \cos \theta \\ b \sin \theta \end{pmatrix}.$$

Determiniamo la retta r tangente all'ellisse in P . Troviamo

$$\bar{A} = \begin{pmatrix} \frac{1}{a^2} & 0 & 0 \\ 0 & \frac{1}{b^2} & 0 \\ 0 & 0 & -1 \end{pmatrix}, \quad r = \{ {}^t \bar{P} \bar{A} \bar{x} = 0 \} = \left\{ \frac{\cos \theta}{a} x + \frac{\sin \theta}{b} y - 1 = 0 \right\}.$$

In forma parametrica:

$$r = \left\{ \begin{pmatrix} a \cos \theta \\ b \sin \theta \end{pmatrix} + t \begin{pmatrix} a \sin \theta \\ -b \cos \theta \end{pmatrix} \right\}.$$

Sia $P \in C$ un punto non centrale di una conica C . Ricordiamo che questa configurazione si può presentare in due casi:

- (1) C è non degenera e $P \in C$ è un punto qualsiasi, oppure
- (2) C è unione di due rette distinte e P non è la loro intersezione.

Studiamo la retta tangente a P in entrambi i casi.

Proposizione 13.2.32. *Sia $P \in C$ un punto non centrale e r la retta tangente a C in P . Valgono i fatti seguenti.*

- (1) Se C è non degenera, r interseca C solo in P .
- (2) Se C consta di due rette distinte, allora r è quella che contiene P .

Dimostrazione. (1). Supponiamo per assurdo che ci sia un'altra intersezione Q . Siccome $Q \in r$, per la legge di reciprocità il punto P è contenuto nella polare s di Q . La polare s contiene anche Q perché $Q \in C$. Entrambe le rette r e s contengono P e Q , quindi $r = s$, assurdo per la Proposizione 13.2.25 perché C è non degenera.

(2). Se le due rette sono incidenti, si intersecano in un centro Q ; allora per il Corollario 13.2.28 la retta r contiene sia P che Q e quindi è la retta di C contenente P . Se le due rette sono parallele, in un opportuno sistema di riferimento $C = \{a^2 x^2 = 1\}$ con $a \neq 0$ e r si calcola (esercizio). $\square$

I punti centrali P contenuti in una conica C sono molto particolari: qualsiasi retta passante per P è tangente a C . Questo fenomeno si presenta solo in questi casi:

- (1) C è unione di due rette incidenti e P è la loro intersezione, oppure
- (2) C è una retta doppia e P è un suo punto qualsiasi, oppure
- (3) C è un punto e $P = C$.

La polare di un punto P non contenuto nella conica C è anch'essa uno strumento utile, perché ci permette di identificare immediatamente tutte le rette passanti per P tangenti a C . Si veda sempre la Figura 13.8.

Proposizione 13.2.33. *Se $P \notin C$, la sua polare r interseca C precisamente nei punti Q per cui la retta s contenente PQ è tangente a C .*

Dimostrazione. La retta s contenente PQ è tangente a $C \iff P$ è contenuto nella polare di $Q \iff Q$ è contenuto nella polare r di P . $\square$

Esempio 13.2.34. Consideriamo la conica $C = \{2x^2 - 2xy + y^2 - 5 = 0\}$ ed il punto $P = (1, 4)$. Notiamo che $P \notin C$ e determiniamo tutte le tangenti a C passanti per P . La polare r in P ha equazione

$$\begin{pmatrix} 1 & 4 & 1 \end{pmatrix} \begin{pmatrix} 2 & -1 & 0 \\ -1 & 1 & 0 \\ 0 & 0 & -5 \end{pmatrix} \begin{pmatrix} x \\ y \\ 1 \end{pmatrix} = 0.$$

Quindi $r = \{-2x + 3y - 5 = 0\}$. Determiniamo l'intersezione $r \cap C$ così:

$$\begin{cases} 2x^2 - 2xy + y^2 - 5 = 0, \\ -2x + 3y - 5 = 0. \end{cases}$$

Si trovano facilmente due soluzioni $A = (2, 3)$ e $B = (-1, 1)$. Quindi le tangenti a C passanti per P sono le rette che contengono rispettivamente PA e PB , e cioè rispettivamente

$$\{x + y = 5\}, \quad \{3x - 2y = -5\}.$$

13.2.14. Altri fasci. Oltre al fascio di coniche passanti per 4 punti in posizione generale esistono altri fasci utili che introduciamo brevemente.

La filosofia che sta dietro alla costruzione dei fasci è la seguente: una conica generica C è del tipo

$$a_{11}x^2 + 2a_{12}xy + a_{22}y^2 + 2b_1x + 2b_2y + c = 0.$$

La conica è determinata da sei parametri, importanti solo a meno di riscalamento (lo spazio dei parametri è $\mathbb{P}^5(\mathbb{R})$). Il passaggio per un punto P fissato si traduce in una equazione lineare omogenea nei coefficienti.

Oltre al passaggio per punti fissati, ci sono altre condizioni geometriche che si traducono in equazioni lineari omogenee nei coefficienti di C . Oltre a contenere P , possiamo ad esempio chiedere che la conica C sia tangente in P ad una retta r fissata. Se $r = \{P + tv\}$, questa condizione si traduce facilmente nel chiedere che

$${}^t v(AP + b) = 0.$$

Anche questa è un'equazione lineare omogenea nei coefficienti di C . In modo simile a come abbiamo fatto nella Proposizione 13.2.20, usando questa condizione di tangenza possiamo formare un altro paio di tipi di fasci.

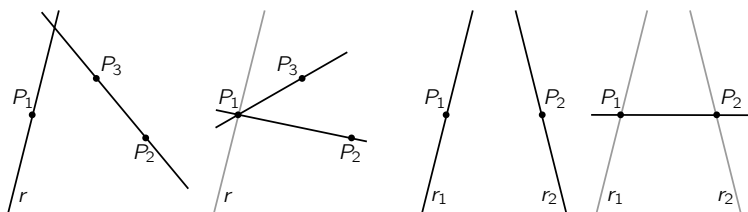


Figura 13.9. Due generatori per il fascio delle coniche passanti per P_1, P_2, P_3 e tangenti ad r (sinistra) e per il fascio delle coniche passanti per P_1, P_2 e tangenti a r_1, r_2 (destra).

Proposizione 13.2.35. *Formano un fascio:*

- (1) le coniche passanti per tre punti non allineati P_1, P_2, P_3 e tangenti in P_1 ad una retta fissata r ;
- (2) le coniche passanti per due punti distinti P_1, P_2 e tangenti in questi a due rette r_1, r_2 fissate.

Dimostrazione. La dimostrazione è simile alla Proposizione 13.2.20. Le condizioni imposte sono sempre 4 e lineari, quindi otteniamo uno spazio di soluzioni di dimensione almeno $6 - 4 = 2$. Per mostrare che la dimensione non può essere ≥ 3 , aggiungiamo un'altra condizione e mostriamo che c'è una sola conica che la soddisfa.

Nel primo caso, possiamo aggiungere il passaggio per un altro punto P_4 allineato con P_2 e P_3 . L'unica conica che soddisfa queste nuove restrizioni è l'unione $r \cup s$ dove s è la retta contenente P_2, P_3, P_4 . Nel secondo caso, aggiungiamo un punto P_3 allineato con P_1 e P_2 . L'unica conica è la retta doppia passante per i tre punti. $\square$

La costruzione di entrambi questi fasci è molto semplice: come per le coniche che passano per 4 punti, il trucco è usare coniche degeneri come generatori. I generatori da usare in entrambi i fasci sono descritti nella Figura 13.9. Nei primi tre casi sono coppie di rette, nell'ultimo una retta doppia. La lettrice si può convincere che in tutti e 4 i generatori le condizioni di tangenza sono rispettate: ricordiamo che tutte le rette passanti per un centro sono tangenti, e che in una retta doppia tutti i punti sono centri.

Entrambi i fasci appena descritti possono essere quindi usati per determinare agevolmente una conica del fascio che soddisfi una condizione aggiuntiva.

Esempio 13.2.36. Determiniamo la conica C che passa per i punti

$$P_1 = \begin{pmatrix} 1 \\ -1 \end{pmatrix}, \quad P_2 = \begin{pmatrix} 1 \\ 1 \end{pmatrix}, \quad P_3 = \begin{pmatrix} 2 \\ 2 \end{pmatrix}, \quad P_4 = \begin{pmatrix} 3 \\ 1 \end{pmatrix}$$

ed è tangente in P_1 alla retta $r = \{y = -x\}$. A questo scopo scriviamo il fascio delle coniche passanti per P_1, P_2, P_3 e tangenti ad r in P_1 . Come suggerito dalla Figura 13.9, due generatori sono le coniche di equazioni

$$\{(x+y)(x-y) = 0\}, \quad \{(x-1)(3x-y-4) = 0\}.$$

Il fascio è quindi

$$C_{t,u} = \{t(x+y)(x-y) + u(x-1)(3x-y-4) = 0\}.$$

Imponiamo il passaggio per P_4 :

$$t \cdot 4 \cdot 2 + u \cdot 2 \cdot 4 = 0.$$

Una soluzione è $t = 1, u = -1$. Quindi la conica cercata è

$$\begin{aligned} C_{1,-1} &= \{(x+y)(x-y) - (x-1)(3x-y-4) = 0\} \\ &= \{-2x^2 - y^2 + xy + 7x - y - 4 = 0\}. \end{aligned}$$

13.3. Quadriche proiettive

Introduciamo le quadriche proiettive in $\mathbb{P}^n$. La lettura di questa sezione è opzionale e può essere saltata. La sezione successiva, che contiene la classificazione delle quadriche nello spazio $\mathbb{R}^3$, può essere letta anche senza conoscere la geometria proiettiva (a cui faremo riferimento solo sporadicamente).

Così come lo studio dei sottospazi proiettivi è più semplice di quello dei sottospazi affini, analogamente vedremo che le quadriche proiettive si classificano più facilmente di quelle in $\mathbb{R}^n$.

Il punto di vista proiettivo sarà quindi utile per capire meglio $\mathbb{R}^n$. Studiando i punti all'infinito comprenderemo meglio le coniche in $\mathbb{R}^2$ e le quadriche in $\mathbb{R}^3$. Potremo ad esempio definire gli *asintoti* di un'iperbole come le tangenti ai suoi punti all'infinito.

13.3.1. Polinomi omogenei. Ricordiamo che un polinomio $p(x)$ nelle variabili $x_1, \dots, x_n$ è omogeneo se tutti i suoi monomi hanno lo stesso grado. Ad esempio, il polinomio

$$p(x) = 2x_1x_2 - x_3^2 + x_1x_3$$

è un polinomio omogeneo di grado 2 nelle variabili x_1, x_2, x_3 . Una proprietà cruciale dei polinomi omogenei è la seguente.

Proposizione 13.3.1. *Sia $p(x)$ un polinomio omogeneo di grado d . Vale*

$$p(\lambda x) = \lambda^d p(x) \quad \forall x \in \mathbb{R}^n, \forall \lambda \in \mathbb{R}.$$

Dimostrazione. Un monomio $m(x)$ di grado d è del tipo

$$m(x) = ax_1^{i_1} \cdots x_k^{i_k}$$

con $i_1 + \cdots + i_k = d$. Quindi

$$m(\lambda x) = a(\lambda x_1)^{i_1} \cdots (\lambda x_k)^{i_k} = \lambda^{i_1 + \cdots + i_k} ax_1^{i_1} \cdots x_k^{i_k} = \lambda^d m(x).$$

Poiché $p(x)$ è somma di monomi di grado d , otteniamo la tesi. $\square$

Corollario 13.3.2. *Sia $p(x)$ un polinomio omogeneo di grado d . Vale*

$$p(x) = 0 \iff p(\lambda x) = 0 \quad \forall x \in \mathbb{R}^n, \forall \lambda \neq 0.$$

Per questa proprietà, i polinomi omogenei sono particolarmente utili in geometria proiettiva.

13.3.2. Quadriche proiettive. Le proprietà dei polinomi omogenei ci consentono di definire le quadriche proiettive nel modo seguente.

Definizione 13.3.3. Una *quadrica proiettiva* in $\mathbb{P}^n$ è il luogo di zeri

$$Q = \{[x_1, \dots, x_{n+1}] \in \mathbb{P}^n \mid p(x_1, \dots, x_{n+1}) = 0\}$$

di un polinomio $p(x_1, \dots, x_{n+1})$ omogeneo di secondo grado.

È essenziale che $p(x)$ sia omogeneo affinché abbia senso l'espressione $p(x_1, \dots, x_{n+1}) = 0$ per un punto x di coordinate $[x_1, \dots, x_{n+1}]$. Questa uguaglianza deve infatti continuare ad essere verificata se moltiplichiamo tutte le variabili $x_1, \dots, x_{n+1}$ per $\lambda \neq 0$, si veda l'Osservazione 12.1.4.

Scriveremo un generico polinomio omogeneo di grado due come

$$p(x) = {}^t x A x$$

dove A è una matrice simmetrica $(n+1) \times (n+1)$. Se cambiamo A con λA per qualche $\lambda \neq 0$ otteniamo un polinomio che definisce la stessa quadrica Q .

Una quadrica in $\mathbb{P}^2$ è una *conica proiettiva*.

Definizione 13.3.4. La quadrica ${}^t x A x = 0$ è *degenere* se $\det A = 0$.

Esempio 13.3.5. L'equazione

$$x_1^2 + x_2^2 - x_3^2 + 4x_1x_2 - 2x_1x_3 = 0$$

definisce una conica proiettiva. La matrice A è

$$A = \begin{pmatrix} 1 & 2 & -1 \\ 2 & 1 & 0 \\ -1 & 0 & -1 \end{pmatrix}.$$

Poiché $\det A \neq 0$, la conica non è degenere.

13.3.3. Cambiamento di coordinate. Come nel caso affine, siamo interessati a vedere come muta il polinomio che definisce una quadrica proiettiva se cambiamo coordinate.

Sia $x = Mx'$ un cambiamento di coordinate per $\mathbb{R}^{n+1}$. Qui M è una qualsiasi matrice invertibile $(n+1) \times (n+1)$. Otteniamo un analogo cambiamento di coordinate per i punti di $\mathbb{P}^n$. Se $Q \subset \mathbb{P}^n$ è una quadrica di equazione

$${}^t x A x = 0,$$

nelle nuove coordinate $x = Mx'$ l'equazione si trasforma in

$$0 = {}^t(Mx')AMx' = {}^t x'({}^tMAM)x'.$$

Quindi nelle nuove coordinate la quadrica Q ha equazione

$${}^t x' A' x' = 0$$

con $A' = {}^tMAM$. La matrice dei coefficienti A cambia per congruenza, similmente a quanto visto nel caso affine.

13.3.4. Classificazione delle quadriche proiettive. La classificazione delle quadriche proiettive è molto semplice.

Teorema 13.3.6 (Classificazione delle quadriche proiettive). *Per ogni quadrica proiettiva $Q \subset \mathbb{P}^n$ esiste un sistema di coordinate in cui Q abbia equazione*

$$(36) \quad x_1^2 + \dots + x_k^2 - x_{k+1}^2 - \dots - x_{k+h}^2 = 0$$

per qualche $k, h \geq 0$ con $1 \leq k + h \leq n + 1$.

Dimostrazione. Questa è una conseguenza del Corollario 7.4.9. $\square$

L'equazione (36) è la *forma canonica* della quadrica proiettiva ${}^t x A x = 0$ ed è determinata direttamente da A nel modo seguente: la terna

$$(k, h, n - (k + h))$$

è la segnatura della matrice simmetrica A . Si possono quindi usare tutte le tecniche studiate nel Capitolo 7 per calcolare la segnatura per identificare la forma canonica di una quadrica proiettiva. La quadrica è non degenere precisamente quando $k + h = n$.

Osservazione 13.3.7. Se moltiplichiamo tutta l'equazione (36) per -1 otteniamo un'altra forma canonica che descrive in realtà la stessa quadrica con k e h invertiti. Quindi le forme canoniche con segnatura

$$(k, h, n - (k + h)), \quad (h, k, n - (h + k))$$

descrivono in realtà la stessa quadrica.

Passiamo ora a studiare più da vicino le quadriche in $\mathbb{P}^1$, $\mathbb{P}^2$ e $\mathbb{P}^3$.

13.3.5. Quadriche in $\mathbb{P}^1$. Studiare una quadrica ${}^t x A x = 0$ in $\mathbb{P}^1$ è molto simile a studiare le soluzioni di un polinomio di secondo grado in una variabile: praticamente è la stessa cosa, l'unica differenza è che qui fra le possibili soluzioni c'è anche il punto all'infinito. Otteniamo 0, 1 oppure 2 punti, a seconda del determinante di A che gioca lo stesso ruolo del Δ in una equazione di secondo grado.

Sia Q una quadrica proiettiva in $\mathbb{P}^1$ di equazione ${}^t x A x = 0$. Notiamo che A è una matrice simmetrica 2×2 .

Proposizione 13.3.8. *La quadrica Q è:*

- (1) *vuota se $\det A > 0$;*

- (2) un punto se $\det A = 0$;
 (3) due punti se $\det A < 0$.

Dimostrazione. Tenendo in conto l'Osservazione 13.3.7, le possibili segnature di A da considerare sono

$$(2, 0, 0), \quad (1, 1, 0), \quad (1, 0, 1)$$

corrispondenti a $\det A$ positivo, negativo e nullo. Le forme canoniche sono

$$x_1^2 + x_2^2 = 0, \quad x_1^2 - x_2^2 = 0, \quad x_1^2 = 0.$$

Per la prima non ci sono soluzioni, con la seconda ce ne sono due $[1, 1]$ e $[1, -1]$, nella terza otteniamo solo $[0, 1]$. $\square$

13.3.6. Coniche proiettive. Usiamo il Teorema 13.3.6 per dedurre una classificazione completa delle coniche proiettive.

Sia C una conica proiettiva in $\mathbb{P}^2$ di equazione ${}^t x A x = 0$. Notiamo che A è una matrice simmetrica 3×3 . Dal Teorema 13.3.6, combinato con l'Osservazione 13.3.7, ricaviamo questa classificazione:

Proposizione 13.3.9. *Esiste un sistema di coordinate rispetto a cui C è una delle coniche seguenti:*

- (1) $x_1^2 + x_2^2 + x_3^2 = 0$ (insieme vuoto) se $\text{rk} A = 3$ e A è definita,
- (2) $x_1^2 + x_2^2 - x_3^2 = 0$ (circonferenza) se $\text{rk} A = 3$ e A è indefinita,
- (3) $x_1^2 + x_2^2 = 0$ (un punto) se $\text{rk} A = 2$ e A è semidefinita,
- (4) $x_1^2 - x_2^2 = 0$ (due rette incidenti) se $\text{rk} A = 2$ e A è indefinita,
- (5) $x_1^2 = 0$ (retta doppia) se $\text{rk} A = 1$.

Studiamo un po' più nel dettaglio ciascun tipo di conica proiettiva, giustificando le brevi descrizioni date fra parentesi nell'enunciato.

- (1) $C = \{x_1^2 + x_2^2 + x_3^2 = 0\}$ è l'insieme vuoto.
- (2) $C = \{x_1^2 + x_2^2 - x_3^2 = 0\}$ ha effettivamente la forma di una circonferenza, perché può essere parametrizzata come

$$C = \{[\cos \theta, \sin \theta, 1] \mid \theta \in [0, 2\pi)\}.$$

- (3) $C = \{x_1^2 + x_2^2 = 0\} = \{[0, 0, 1]\}$ è un punto.
- (4) $C = \{x_1^2 - x_2^2 = 0\} = \{(x_1 + x_2)(x_1 - x_2) = 0\}$ è unione delle rette proiettive $x_1 + x_2 = 0$ e $x_1 - x_2 = 0$, che si intersecano in $[0, 0, 1]$.
- (5) $C = \{x_1^2 = 0\}$ è la retta $x_1 = 0$, da intendersi con "molteplicità due", cioè una retta doppia.

In particolare una conica proiettiva non degenerare è vuota oppure una circonferenza.

Esempio 13.3.10. Consideriamo la famiglia di coniche proiettive

$$C_k = \{kx_0^2 + x_1^2 + 5x_2^2 + 4x_0x_1 - 4x_0x_2 - 2x_1x_2 = 0\}$$

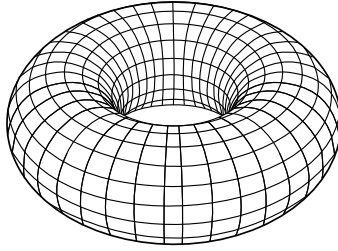


Figura 13.10. Un toro. La superficie proiettiva di equazione $x_1^2 + x_2^2 - x_3^2 - x_4^2 = 0$ ha questa forma.

dipendenti da un parametro $k \in \mathbb{R}$. La matrice dei coefficienti è

$$A = \begin{pmatrix} k & 2 & -2 \\ 2 & 1 & -1 \\ -2 & -1 & 5 \end{pmatrix}.$$

Il determinante è $\det A = 4k - 16$. Quindi:

- se $k > 4$ la segnatura è $(3, 0, 0)$ e quindi $C_k = \emptyset$;
- se $k = 4$ la segnatura è $(2, 0, 1)$ e quindi C_4 è un punto;
- se $k < 4$ la segnatura è $(2, 1, 0)$ e quindi C_k è una circonferenza.

Passiamo adesso dal piano allo spazio.

13.3.7. Quadriche in $\mathbb{P}^3$. Usiamo il Teorema 13.3.6 per classificare le quadriche proiettive in $\mathbb{P}^3$.

Sia Q una quadrica proiettiva in $\mathbb{P}^3$ di equazione ${}^t x A x = 0$. Notiamo che A è una matrice simmetrica 4×4 . Dal Teorema 13.3.6, combinato con l'Osservazione 13.3.7, ricaviamo questa classificazione:

Corollario 13.3.11. *Esiste un sistema di coordinate rispetto a cui Q è una delle quadriche seguenti:*

- (1) $x_1^2 + x_2^2 + x_3^2 + x_4^2 = 0$ (*insieme vuoto*),
- (2) $x_1^2 + x_2^2 + x_3^2 - x_4^2 = 0$ (*sfera*),
- (3) $x_1^2 + x_2^2 - x_3^2 - x_4^2 = 0$ (*toro*),
- (4) $x_1^2 + x_2^2 + x_3^2 = 0$ (*un punto*),
- (5) $x_1^2 + x_2^2 - x_3^2 = 0$ (*cono*),
- (6) $x_1^2 + x_2^2 = 0$ (*una retta*),
- (7) $x_1^2 - x_2^2 = 0$ (*due piani incidenti*),
- (8) $x_1^2 = 0$ (*piano doppio*).

Gli otto possibili casi sono determinati come sempre dalla segnatura di A . Studiamo un po' più nel dettaglio ciascuna quadrica proiettiva, giustificando le brevi descrizioni date fra parentesi nell'enunciato.

- (1) $Q = \{x_1^2 + x_2^2 + x_3^2 + x_4^2 = 0\}$ è l'insieme vuoto.

- (2) $Q = \{x_1^2 + x_2^2 + x_3^2 - x_4^2 = 0\}$ ha effettivamente la forma di una sfera, perché può essere parametrizzata come

$$Q = \{[x_1, x_2, x_3, 1] \mid x_1^2 + x_2^2 + x_3^2 = 1\}$$

e ricordiamo che i punti di $\mathbb{R}^3$ con $x_1^2 + x_2^2 + x_3^2 = 1$ formano una sfera di centro l'origine.

- (3) $Q = \{x_1^2 + x_2^2 - x_3^2 - x_4^2 = 0\}$ può essere parametrizzata come

$$Q = \{[\cos \theta, \sin \theta, \cos \varphi, \sin \varphi] \mid \theta, \varphi \in [0, 2\pi)\}.$$

Questa superficie ha la forma di un *toro*, mostrato nella Figura 13.10.

- (4) $Q = \{x_1^2 + x_2^2 + x_3^2 = 0\} = \{[0, 0, 0, 1]\}$ è un punto.

- (5) $Q = \{x_1^2 + x_2^2 - x_3^2 = 0\}$ può essere parametrizzata come

$$Q = \{[\cos \theta, \sin \theta, 1, t] \mid \theta \in [0, 2\pi), t \in \mathbb{R}\} \cup \{[0, 0, 0, 1]\}.$$

Questa superficie ha una figura un po' più complicata, perché contiene un punto $[0, 0, 0, 1]$ di tipo "singolare" in cui non è liscia, come nel cono nella Figura 7.2-(destra).

- (6) $Q = \{x_1^2 + x_2^2 = 0\} = \{x_1 = x_2 = 0\}$ è una retta.

- (7) $Q = \{x_1^2 - x_2^2 = 0\} = \{(x_1 + x_2)(x_1 - x_2) = 0\}$ è unione dei piani proiettivi $x_1 + x_2 = 0$ e $x_1 - x_2 = 0$, che si intersecano nella retta $x_1 = x_2 = 0$.

- (8) $Q = \{x_1^2 = 0\}$ è il piano $x_1 = 0$, da intendersi con "molteplicità due", cioè un piano doppio.

13.3.8. Completamento proiettivo di una quadrica. Come abbiamo accennato, le quadriche proiettive ci aiutano a capire le quadriche di $\mathbb{R}^n$. Abbiamo visto nel Capitolo 12 che un sottospazio affine di $\mathbb{R}^n$ può essere completato ad un sottospazio proiettivo di $\mathbb{P}^n$ aggiungendo i suoi punti all'infinito. Mostriamo adesso che la stessa costruzione funziona per le quadriche.

Come nel Capitolo 12, scriviamo

$$\mathbb{P}^n = \mathbb{R}^n \cup H$$

dove un punto $x = {}^t(x_1, \dots, x_n) \in \mathbb{R}^n$ è identificato con $[x_1, \dots, x_n, 1] \in \mathbb{P}^n$ e H è l'iperpiano $\{x_{n+1} = 0\}$ dei punti all'infinito.

Sia $Q \subset \mathbb{R}^n$ una quadrica di equazione

$${}^t\bar{x}\bar{A}\bar{x} = 0.$$

Ricordiamo che

$$\bar{x} = \begin{pmatrix} x \\ 1 \end{pmatrix}, \quad \bar{A} = \begin{pmatrix} A & b \\ {}^tb & c \end{pmatrix}.$$

Il *completamento proiettivo* di Q è la quadrica proiettiva $\bar{Q} \subset \mathbb{P}^n$ descritta dall'equazione

$${}^t\bar{x}\bar{A}\bar{x} = 0$$

dove questa volta $x = {}^t(x_1, \dots, x_{n+1})$. Questa equazione differisce dalla precedente per il fatto che il polinomio a sinistra dell'uguale adesso è omogeneo e con una variabile in più x_{n+1} . Più concretamente, se $Q \subset \mathbb{R}^n$ è definita da un'equazione di secondo grado

$$\sum_{i,j=1}^n a_{ij}x_i x_j + 2 \sum_{i=1}^n b_i x_i + c = 0$$

allora $\bar{Q} \subset \mathbb{P}^n$ è definita dall'equazione omogenea

$$\sum_{i,j=1}^n a_{ij}x_i x_j + 2 \sum_{i=1}^n b_i x_i x_{n+1} + c x_{n+1}^2 = 0.$$

Abbiamo trasformato il polinomio che definisce Q in un polinomio omogeneo inserendo la variabile aggiuntiva x_{n+1} . Questo procedimento è identico a quello usato nella Sezione 12.2.5 per definire il completamento di un sottospazio affine.

Proposizione 13.3.12. *Valgono i fatti seguenti:*

- $\bar{Q} \cap \mathbb{R}^n = Q$.
- $\bar{Q} \cap H$ è la quadrica proiettiva descritta dall'equazione

$${}^t x A x = 0, \quad x_{n+1} = 0$$

$$\text{con } x = {}^t(x_1, \dots, x_n).$$

Dimostrazione. La dimostrazione è la stessa della Proposizione 12.2.2. $\square$

Riassumendo, il completamento $\bar{Q}$ della quadrica affine Q è ottenuto aggiungendo a Q una quadrica proiettiva nell'iperpiano all'infinito.

Esempio 13.3.13. Usiamo le variabili x, y, z invece di x_1, x_2, x_3 . Se $C = \{x^2 + y^2 = 1\}$ è una circonferenza in $\mathbb{R}^2$, il completamento $\bar{C} = \{x^2 + y^2 = z^2\}$ è ancora una circonferenza in $\mathbb{P}^2$: nessun punto all'infinito è stato aggiunto.

Se $C = \{y = x^2\}$ è una parabola in $\mathbb{R}^2$, il completamento $\bar{C} = \{yz = x^2\} = \{x^2 - yz = 0\}$ è una conica proiettiva descritta dalla matrice

$$\bar{A} = \begin{pmatrix} 1 & 0 & 0 \\ 0 & 0 & -\frac{1}{2} \\ 0 & -\frac{1}{2} & 0 \end{pmatrix}.$$

La matrice ha segnatura $(2, 1, 0)$ e quindi $\bar{C}$ è una circonferenza. Abbiamo ottenuto una circonferenza $\bar{C}$ aggiungendo un punto all'infinito alla parabola C . È stato aggiunto il punto all'infinito $[0, 1, 0]$.

Se $C = \{xy = 1\}$ è una iperbole in $\mathbb{R}^2$, il completamento $\bar{C} = \{xy = z^2\}$ è ancora una circonferenza in $\mathbb{P}^2$. Questa volta sono stati aggiunti due punti all'infinito $[1, 0, 0]$ e $[0, 1, 0]$.

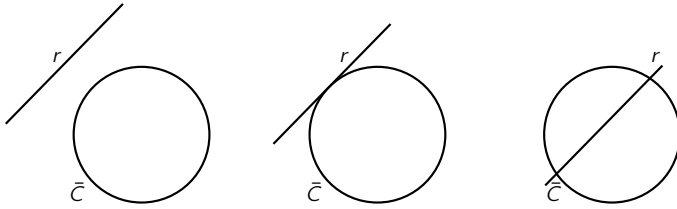


Figura 13.11. Il completamento proiettivo $\bar{C}$ di una conica non degenera C è sempre una circonferenza $\bar{C}$. Questa interseca la retta r all'infinito in 0, 1 oppure 2 punti, e la conica C è corrispondentemente un'ellisse (sinistra), una parabola (centro) o un'iperbole (destra).

13.3.9. Completamento di una conica affine. Nell'Esempio 13.3.13 abbiamo descritto il completamento di alcune coniche. Descriviamo adesso questo fenomeno in generale.

Proposizione 13.3.14. *Sia $C \subset \mathbb{R}^2$ una conica non vuota e $\bar{C} \subset \mathbb{P}^2$ il suo completamento. Valgono i fatti seguenti:*

- (1) C è un'ellisse, parabola o iperbole $\iff \bar{C}$ è una circonferenza.
- (2) C è un punto $\iff \bar{C}$ è un punto.
- (3) C è due rette parallele o incidenti $\iff \bar{C}$ è due rette incidenti.
- (4) C è una retta doppia $\iff \bar{C}$ è una retta doppia.

In tutti i casi $\bar{C}$ è ottenuto aggiungendo 0, 1 oppure 2 punti all'infinito a C . Si aggiungono due punti se C è un'iperbole o due rette incidenti, uno se è una parabola o una retta doppia, zero negli altri casi.

Dimostrazione. La prima parte è una conseguenza dei Teoremi 13.2.7 e 13.2.8 e della Proposizione 13.3.9. I punti all'infinito possono essere intuiti geometricamente (c'è un punto all'infinito per ogni ramo illimitato della conica) oppure algebricamente: i punti all'infinito sono descritti dall'equazione ${}^t x A x = 0$ e sono 2, 1, 0 a seconda di $\det A$, si veda la Proposizione 13.3.8. $\square$

Si veda la Figura 13.11.

13.3.10. Retta polare e tangente. È possibile definire la polare per le coniche proiettive, in modo analogo a quanto fatto nella Sezione 13.2.12. Usando la polare si definisce una nozione di tangenza fra retta e conica.

Sia C una conica proiettiva in $\mathbb{P}^2$, di equazione

$${}^t x A x = 0.$$

Sia $P \in \mathbb{P}^2$ un punto qualsiasi. La *polare* di P è il luogo di punti

$$r = \{x \in \mathbb{P}^2 \mid {}^t P A x = 0\}.$$

Consideriamo per semplicità solo il caso in cui C sia non degenera. In questo caso A è invertibile e quindi r è sempre una retta: la costruzione polare induce una corrispondenza biunivoca fra i punti P e le rette r di $\mathbb{P}^2$.

Come nella Sezione 13.2.12 si dimostra che la polare r di P contiene P se e solo se $P \in C$, ed in questo caso diciamo che r è la *retta tangente* a C in P . Vale inoltre sempre la legge di reciprocità:

Q sta nella polare di $P \iff P$ sta nella polare di Q .

Esempio 13.3.15. Consideriamo la conica proiettiva non degenera

$$C = \{x_1^2 + x_2^2 - x_3^2 + 4x_1x_2 - 2x_1x_3 = 0\}$$

già studiata nell'Esempio 13.3.5. Sia $P = [0, 1, 1] \in C$. La retta r tangente a C in P ha equazione

$$(0 \quad 1 \quad 1) \begin{pmatrix} 1 & 2 & -1 \\ 2 & 1 & 0 \\ -1 & 0 & -1 \end{pmatrix} \begin{pmatrix} x_1 \\ x_2 \\ x_3 \end{pmatrix} = 0$$

e quindi

$$r = \{x_1 + x_2 - x_3 = 0\}.$$

Sia $C \subset \mathbb{R}^2$ una conica non degenera, di centro O (quindi C è una ellisse o una iperbole). Ricordiamo che la polare in $\mathbb{R}^2$ del centro O è vuota. In geometria proiettiva la nozione di centro non ha più senso e ogni punto ha una sua retta polare. Possiamo quindi chiederci chi sia la polare di O rispetto al completamento $\bar{C}$ di C in $\mathbb{P}^2 \supset \mathbb{R}^2$. C'è una sola risposta plausibile:

Proposizione 13.3.16. *La polare di O rispetto a $\bar{C}$ in $\mathbb{P}^2$ è la retta all'infinito.*

Dimostrazione. Il punto O soddisfa l'equazione $AO + b = 0$. Quindi

$$\bar{A}\bar{O} = \begin{pmatrix} A & b \\ {}^t b & c \end{pmatrix} \begin{pmatrix} O \\ 1 \end{pmatrix} = \begin{pmatrix} AO + b \\ {}^t bO + c \end{pmatrix} = \begin{pmatrix} 0 \\ {}^t bO + c \end{pmatrix}.$$

Ne segue che la polare $r = \{{}^t \bar{O} \bar{A} \bar{x} = 0\}$ ha equazione

$$r = \{({}^t bO + c)x_3 = 0\} = \{x_3 = 0\}$$

ed è quindi la retta all'infinito. □

In geometria proiettiva è possibile definire ed usare i fasci come nelle Sezioni 13.2.11 e 13.2.14. Una conica proiettiva è descritta da 6 parametri, visti a meno di moltiplicazione per una costante: questi parametri formano uno spazio $\mathbb{P}^6$ e un fascio è una retta proiettiva in $\mathbb{P}^6$. Per 5 punti in $\mathbb{P}^2$ in posizione generale passa una sola conica proiettiva.

13.3.11. Asintoti. Grazie all'introduzione dei punti all'infinito possiamo infine definire e studiare gli asintoti di un'iperbole.

Sia $C \subset \mathbb{R}^2$ un'iperbole. Il suo completamento $\bar{C}$ è ottenuto aggiungendo a C due punti all'infinito.

Definizione 13.3.17. Le rette tangenti a $\bar{C}$ in $\mathbb{P}^2$ nei suoi due punti all'infinito sono gli *asintoti* dell'iperbole C .

Proposizione 13.3.18. *I due asintoti si intersecano nel centro dell'iperbole.*

Dimostrazione. Sappiamo dalla Proposizione 13.3.16 che la polare del centro è la retta all'infinito; per la legge di reciprocità, la polare di un qualsiasi punto all'infinito contiene il centro. In particolare gli asintoti contengono il centro. $\square$

Esempio 13.3.19. Determiniamo gli asintoti dell'iperbole

$$C = \{x^2 - 2xy - y^2 - 2x + 3 = 0\}$$

introdotta nell'Esempio 13.2.24. Il suo centro è $O = \frac{1}{2} \begin{pmatrix} 1 \\ -1 \end{pmatrix}$. I due punti all'infinito si trovano risolvendo l'equazione ${}^t x A x = 0$, cioè

$$x^2 - 2xy - y^2 = 0.$$

Ponendo $y = 1$ troviamo le soluzioni $x = 1 \pm \sqrt{2}$. I punti all'infinito sono

$$Q_1 = [1 - \sqrt{2}, 1, 0], \quad Q_2 = [1 + \sqrt{2}, 1, 0].$$

Gli asintoti sono le rette passanti per O nelle direzioni Q_1 e Q_2 , cioè

$$r_1 = \left\{ \frac{1}{2} \begin{pmatrix} 1 \\ -1 \end{pmatrix} + t \begin{pmatrix} 1 - \sqrt{2} \\ 1 \end{pmatrix} \right\}, \quad r_2 = \left\{ \frac{1}{2} \begin{pmatrix} 1 \\ -1 \end{pmatrix} + t \begin{pmatrix} 1 + \sqrt{2} \\ 1 \end{pmatrix} \right\}.$$

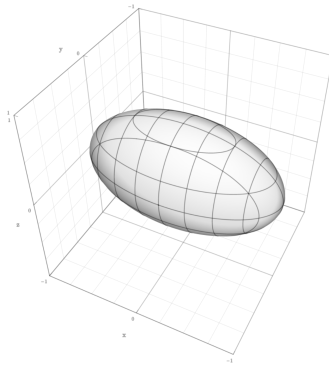
13.4. Quadriche in $\mathbb{R}^3$

Concludiamo il capitolo con lo studio e la classificazione delle quadriche in $\mathbb{R}^3$. Come per le coniche, la classificazione delle quadriche in $\mathbb{R}^3$ è più complicata di quella proiettiva, perché ci sono molti più casi da considerare.

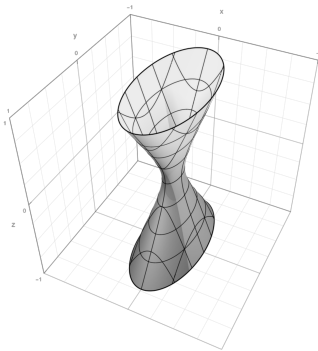
13.4.1. Esempi di quadriche non degeneri. Costruiamo ora cinque tipi di quadriche in $\mathbb{R}^3$, disegnate nella Figura 13.12. In seguito mostreremo che questi esempi esauriscono tutte le quadriche non degeneri in $\mathbb{R}^3$.

Ellissoide. L'*ellissoide* di assi x, y, z e con semiassi $a, b, c > 0$ è la quadrica in $\mathbb{R}^3$ di equazione

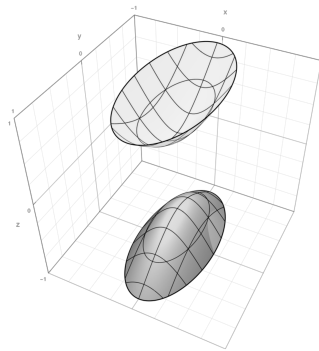
$$\frac{x^2}{a^2} + \frac{y^2}{b^2} + \frac{z^2}{c^2} = 1.$$



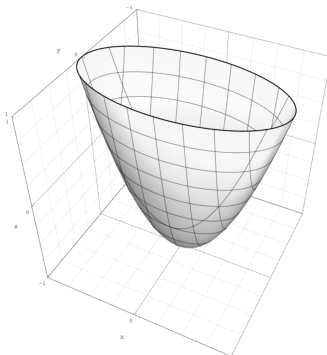
ellissoide



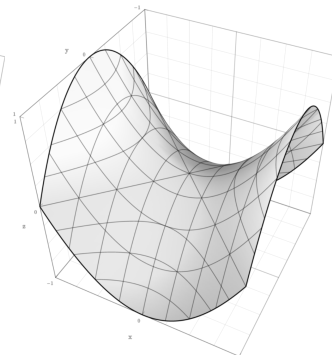
iperboloide a una falda



iperboloide a due falde



paraboloide ellittico



paraboloide iperbolico

Figura 13.12. I cinque tipi di quadriche non degeneri in $\mathbb{R}^3$. Per ciascuna quadrica sono disegnate le intersezioni con alcuni piani verticali ed orizzontali: queste intersezioni sono sempre delle coniche.

Si veda la Figura 13.12. Per convincersi che questa superficie ha la forma disegnata in figura, notiamo che se intersechiamo l'ellissoide con un piano orizzontale $z = h$ otteniamo la conica

$$\frac{x^2}{a^2} + \frac{y^2}{b^2} = 1 - \frac{h^2}{c^2}.$$

Questa è un'ellisse se $h^2 < c^2$, un punto se $h^2 = c^2$ e l'insieme vuoto se $h^2 > c^2$. Otteniamo analogamente ellissi anche intersecando con i piani $x = k$ e $y = l$ se $k^2 < a^2$ e $l^2 < b^2$. Se $a = b = c$ l'ellissoide è una sfera di centro l'origine e raggio a .

I punti dell'ellissoide possono essere descritti in forma parametrica usando una variante delle coordinate sferiche:

$$\begin{aligned}x &= a \cos \theta \cos \varphi, \\y &= b \cos \theta \sin \varphi, \\z &= c \sin \theta,\end{aligned}$$

con $\theta \in [0, \pi]$ e $\varphi \in [0, 2\pi)$.

Il completamento proiettivo dell'ellissoide è la quadrica proiettiva

$$\frac{x^2}{a^2} + \frac{y^2}{b^2} + \frac{z^2}{c^2} - w^2 = 0$$

Usiamo le variabili x, y, z, w anziché x_1, x_2, x_3, x_4 . Questa quadrica in $\mathbb{P}^3$ è una sfera. L'ellissoide non ha punti all'infinito, perché il sistema

$$\frac{x^2}{a^2} + \frac{y^2}{b^2} + \frac{z^2}{c^2} = 0, \quad w = 0$$

non ha soluzioni in $\mathbb{P}^3$.

Iperboloide ad una falda. L'*iperboloide ad una falda* di assi x, y, z e con semiassi $a, b, c > 0$ è la quadrica in $\mathbb{R}^3$ di equazione

$$\frac{x^2}{a^2} + \frac{y^2}{b^2} - \frac{z^2}{c^2} = 1.$$

Si veda la Figura 13.12. Per convincersi che questa superficie ha la forma disegnata in figura, notiamo che se la intersechiamo con un piano orizzontale $z = h$ otteniamo ellissi tutte simili fra loro, la cui grandezza cresce con h^2 . Invece se la intersechiamo con un piano $x = k$ o $y = l$ otteniamo delle iperboli.

L'iperboloide ad una falda ha la peculiarità di essere una *superficie rigata*, cioè è una unione di rette. La Figura 13.13-(sinistra) mostra come l'iperboloide sia effettivamente una unione di rette in due modi diversi. Le rette sono

$$r_{\alpha}^{\pm} = \left\{ \begin{pmatrix} a \cos \alpha \\ b \sin \alpha \\ 0 \end{pmatrix} + t \begin{pmatrix} -a \sin \alpha \\ b \cos \alpha \\ \pm c \end{pmatrix} \right\}.$$

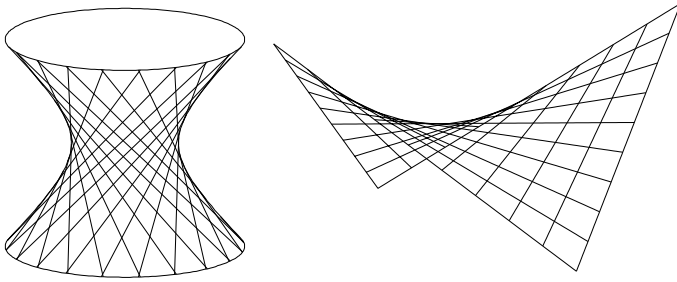


Figura 13.13. L'iperboloide ad una falda ed il paraboloide iperbolico sono superfici rigate.

Le famiglie r_α^+ e r_α^- ricoprono l'iperboloide al variare di $\alpha \in [0, 2\pi)$. Per questa proprietà geometrica l'iperboloide ad una falda è spesso usato in architettura.

Il completamento proiettivo dell'iperboloide ad una falda è la quadrica

$$\frac{x^2}{a^2} + \frac{y^2}{b^2} - \frac{z^2}{c^2} - w^2 = 0.$$

Questa quadrica proiettiva in $\mathbb{P}^3$ è un toro. I punti all'infinito dell'iperboloide sono le soluzioni del sistema

$$\frac{x^2}{a^2} + \frac{y^2}{b^2} - \frac{z^2}{c^2} = 0, \quad w = 0$$

e formano una circonferenza. Aggiungendo una circonferenza all'infinito, l'iperboloide ad una falda si chiude e diventa un toro.

Iperboloide a due falde. L'*iperboloide a due falde* di assi x, y, z e con semiassi $a, b, c > 0$ è la quadrica in $\mathbb{R}^3$ di equazione

$$\frac{x^2}{a^2} + \frac{y^2}{b^2} - \frac{z^2}{c^2} = -1.$$

Si veda la Figura 13.12. Per convincersi che questa superficie ha la forma disegnata in figura, notiamo che se la intersechiamo con un piano orizzontale $z = h$ otteniamo ellissi precisamente quando $h^2 > c^2$, mentre non otteniamo nulla se $h^2 < c^2$. Se invece la intersechiamo con un piano verticale $x = k$ o $y = l$ otteniamo delle iperboli.

Il completamento dell'iperboloide a due falde è la quadrica proiettiva

$$\frac{x^2}{a^2} + \frac{y^2}{b^2} - \frac{z^2}{c^2} + w^2 = 0.$$

Questa quadrica proiettiva in $\mathbb{P}^3$ è una sfera. I punti all'infinito dell'iperboloide a due falde formano una circonferenza come per l'iperboloide ad una falda. Aggiungendo una circonferenza all'infinito, le due falde dell'iperboloide di uniscono a formare una sfera.

Paraboloide ellittico. Il *paraboloide ellittico* di assi x, y, z e con semiassi $a, b > 0$ è la quadrica in $\mathbb{R}^3$ di equazione

$$z = \frac{x^2}{a^2} + \frac{y^2}{b^2}.$$

Si veda la Figura 13.12. Per convincersi che questa superficie ha la forma disegnata in figura, notiamo che se la intersechiamo con un piano orizzontale $z = h$ otteniamo ellissi precisamente quando $h > 0$, mentre non otteniamo nulla se $h < 0$. Invece se la intersechiamo con un piano $x = k$ o $y = l$ otteniamo delle parabole.

Se $a = b$, il paraboloide ellittico è *regolare*. Questa è la superficie usata per costruire antenne paraboliche.

Il completamento del paraboloide ellittico è la quadrica proiettiva

$$\frac{x^2}{a^2} + \frac{y^2}{b^2} - zw = 0.$$

Esaminando la segnatura della forma quadratica si vede che questa quadrica proiettiva è una sfera in $\mathbb{P}^3$. Il paraboloide ellittico ha un unico punto all'infinito, il punto $[0, 0, 1, 0]$ soluzione del sistema

$$\frac{x^2}{a^2} + \frac{y^2}{b^2} = 0, \quad w = 0.$$

Aggiungendo un punto all'infinito, il paraboloide ellittico si chiude e diventa una sfera.

Paraboloide iperbolico. Il *paraboloide iperbolico* di assi x, y, z e con semiassi $a, b > 0$ è la quadrica in $\mathbb{R}^3$ di equazione

$$z = \frac{x^2}{a^2} - \frac{y^2}{b^2}.$$

Si veda la Figura 13.12. Per convincersi che questa superficie ha la forma disegnata in figura, notiamo che se la intersechiamo con un piano orizzontale $z = h$ otteniamo un'iperbole se $h \neq 0$ e due rette incidenti se $h = 0$. Se invece la intersechiamo con un piano $x = k$ o $y = l$ otteniamo delle parabole.

La forma del paraboloide vicino all'origine ricorda la sella di un cavallo: la sua intersezione con il piano verticale xz è una parabola rivolta verso l'alto, mentre quella con il piano verticale yz è una parabola rivolta verso il basso.

Il paraboloide iperbolico è una superficie rigata. La Figura 13.13- (destra) mostra come la superficie sia effettivamente una unione di rette in due modi diversi. Le rette sono

$$r_u^\pm = \left\{ \begin{pmatrix} au \\ 0 \\ u^2 \end{pmatrix} + t \begin{pmatrix} a \\ \pm b \\ 2u \end{pmatrix} \right\}$$

Le famiglie r_u^+ e r_u^- ricoprono il paraboloide al variare di $u \in \mathbb{R}$. Per questa proprietà geometrica il paraboloide iperbolico è spesso usato in architettura.

Il completamento del paraboloide iperbolico è la quadrica proiettiva

$$\frac{x^2}{a^2} - \frac{y^2}{b^2} - zw = 0.$$

Esaminando la segnatura della forma quadratica si vede che questa quadrica proiettiva è un toro in $\mathbb{P}^3$. I punti all'infinito del paraboloide iperbolico sono le soluzioni del sistema

$$\frac{x^2}{a^2} - \frac{y^2}{b^2} = 0, \quad w = 0$$

e formano due rette incidenti. Aggiungendo due rette incidenti all'infinito, il paraboloide iperbolico si chiude e diventa un toro.

In generale, diciamo che una quadrica Q è un ellissoide, iperboloide (ad una o due falde) o paraboloide (ellittico o iperbolico) se assume la forma appena descritta in un opportuno sistema di coordinate.

13.4.2. Classificazione delle quadriche non degeneri. Con tecniche simili a quelle usate per le coniche possiamo classificare le quadriche in $\mathbb{R}^3$.

Sia Q una quadrica in $\mathbb{R}^3$, luogo di zeri di un polinomio

$${}^t x A x + {}^t b x + c = 0.$$

Ricordiamo che Q è non degenera se la matrice simmetrica

$$\bar{A} = \begin{pmatrix} A & b \\ {}^t b & c \end{pmatrix}$$

è invertibile. Consideriamo questo caso. Le possibili segnature per $\bar{A}$ sono

$$(4, 0, 0), \quad (3, 1, 0), \quad (2, 2, 0), \quad (1, 3, 0), \quad (0, 4, 0).$$

La matrice $\bar{A}$ è *definita* nel primo e l'ultimo caso e *indefinita* negli altri. Notiamo che l'indice di negatività è pari precisamente quando $\det \bar{A}$ è positivo. In virtù dell'Osservazione 13.3.7, ci sono solo 3 tipi di segnature da considerare:

- (1) $(4, 0, 0)$ oppure $(0, 0, 4)$, cioè $\bar{A}$ definita;
- (2) $(3, 1, 0)$ oppure $(1, 3, 0)$, cioè $\bar{A}$ indefinita e $\det \bar{A} < 0$;
- (3) $(2, 2, 0)$, cioè $\bar{A}$ indefinita e $\det \bar{A} > 0$.

Teorema 13.4.1. *Se Q è non degenera, si presentano i casi seguenti:*

- (1) *Se $\bar{A}$ è definita, allora la quadrica è vuota.*
- (2) *Se $\bar{A}$ è indefinita e $\det \bar{A} < 0$, allora la quadrica è*
 - *un ellissoide se $\det A \neq 0$ e A è definita,*
 - *un iperboloide a due falde se $\det A \neq 0$ e A è indefinita,*
 - *un paraboloide ellittico se $\det A = 0$.*
- (3) *Se $\bar{A}$ è indefinita e $\det \bar{A} > 0$, allora la quadrica è*
 - *un iperboloide ad una falda se $\det A \neq 0$,*
 - *un paraboloide iperbolico se $\det A = 0$.*

$\bar{A}$	$\det \bar{A} < 0$			$\det \bar{A} > 0$	
$\bar{Q}$	sfera			toro	
Q	ellissoide	iperboloide a due falde	paraboloide ellittico	iperboloide a una falda	paraboloide iperbolico
A	$\det A \neq 0$ definita	$\det A \neq 0$ indefinita	$\det A = 0$	$\det A \neq 0$	$\det A = 0$

Tabella 13.1. I cinque tipi di quadriche non degeneri in $\mathbb{R}^3$. Si suppone che $\bar{A}$ sia indefinita (altrimenti la quadrica è vuota). Qui Q è la quadrica e $\bar{Q}$ il suo completamento proiettivo.

Dimostrazione. Per la Proposizione 13.1.7 possiamo supporre che A sia diagonale. Consideriamo prima il caso in cui A sia invertibile. Per la Proposizione 13.1.12 possiamo supporre che $\bar{A} = \begin{pmatrix} A & 0 \\ 0 & c \end{pmatrix}$ sia diagonale. Siccome $\bar{A}$ è invertibile, abbiamo $c \neq 0$. Possiamo quindi dividere tutto per c e ottenere che Q è luogo di zeri di

$$ax^2 + by^2 + cz^2 + 1 = 0$$

con opportuni coefficienti $a, b, c \neq 0$. A seconda dei segni di a, b, c otteniamo l'insieme vuoto, un ellissoide o un iperboloide a una o due falde. Il tipo è determinato dalle segnature di A e $\bar{A}$ come scritto nell'enunciato.

Resta da considerare il caso in cui $\det A = 0$. A meno di scambiare le variabili x, y e z otteniamo

$$\bar{A} = \begin{pmatrix} a_{11} & 0 & 0 & b_1 \\ 0 & a_{22} & 0 & b_2 \\ 0 & 0 & 0 & b_3 \\ b_1 & b_2 & b_3 & c \end{pmatrix}.$$

Abbiamo $\det \bar{A} = -b_3^2 a_{11} a_{22} \neq 0$, quindi $a_{11}, a_{22}, b_3 \neq 0$. La conica Q ha equazione

$$a_{11}x^2 + a_{22}y^2 + 2b_1x + 2b_2y + 2b_3z + c = 0.$$

Dopo aver diviso tutto per $2b_3$ otteniamo un'equazione del tipo

$$z = ax^2 + by^2 + cx + dy + e$$

con opportuni coefficienti a, b, c, d, e e con $a, b \neq 0$. Una traslazione opportuna (esercizio) trasforma l'equazione nella forma $z = ax^2 + by^2$. Quindi Q è un paraboloide, ellittico o iperbolico a seconda di $\det \bar{A}$. $\square$

I tipi di quadriche sono riassunti nella Tabella 13.1.

13.4.3. Alcune quadriche degeneri. Descriviamo alcune quadriche degeneri, disegnate nella Figura 13.14.

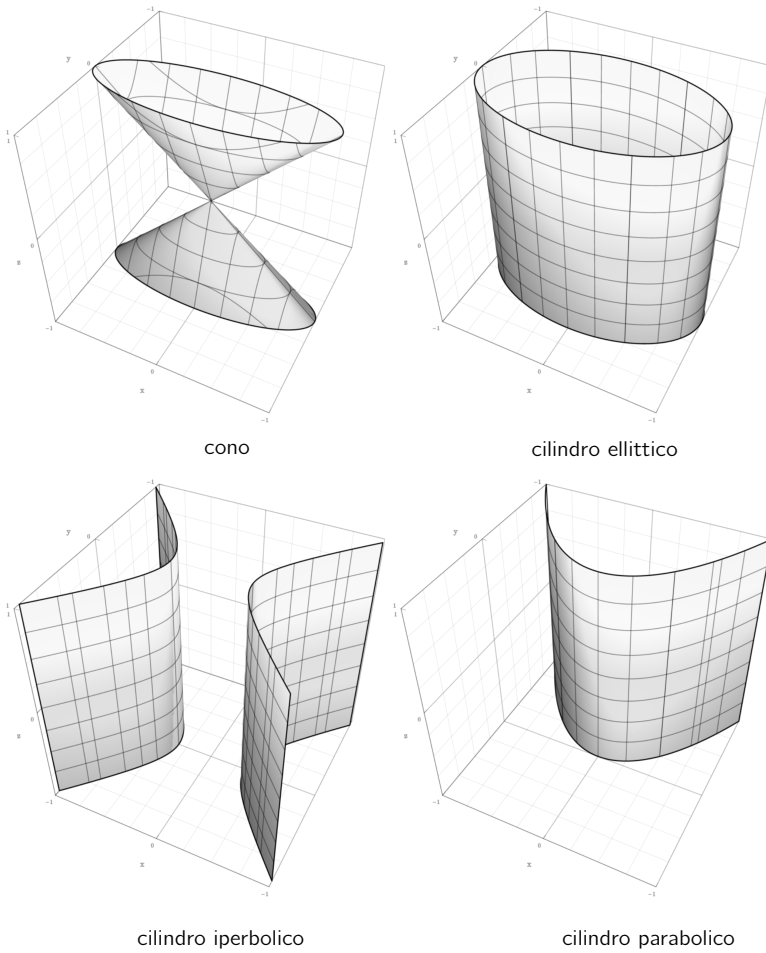


Figura 13.14. Alcune quadriche degeneri.

Cono. Il cono di assi x, y, z e con semiassi $a, b > 0$ è la quadrica in $\mathbb{R}^3$ di equazione

$$z^2 = \frac{x^2}{a^2} + \frac{y^2}{b^2}.$$

Si veda la Figura 13.14. L'intersezione con un piano orizzontale $z = h$ è una ellisse se $h \neq 0$ e un punto se $h = 0$. L'intersezione con un piano verticale $x = k$ è un'iperbole se $k \neq 0$ e due rette incidenti per $k = 0$. L'intersezione con il piano $z = y/b + k$ per $k \neq 0$ è una parabola.

Intersecando il cono con dei piani otteniamo tutti i tipi di coniche, ed è per questo che vengono chiamate in questo modo sin dall'antichità.

Il cono è una superficie rigata, unione delle rette passanti per l'origine

$$r_\theta = \text{Span} \begin{pmatrix} a \cos \theta \\ b \sin \theta \\ 1 \end{pmatrix}$$

al variare di $\theta \in [0, 2\pi)$.

Cilindri. Data una conica C nel piano $\mathbb{R}^2$, di equazione

$$a_{11}x^2 + a_{22}y^2 + 2a_{12}xy + 2b_1x + 2b_2y + c = 0$$

è possibile definire una conica Q in $\mathbb{R}^3$ usando la stessa equazione. La conica Q è il *cilindro* con base C . In altre parole, un cilindro è una conica definita da un'equazione in cui non compare la variabile z .

Nella Figura 13.14 sono mostrati dei cilindri con base ellittica, iperbolica e parabolica. I cilindri sono ovviamente delle superfici rigate, unioni di rette verticali.

13.4.4. Classificazione delle quadriche degeneri. Concludiamo il capitolo con la classificazione delle quadriche degeneri.

Sia $Q \subset \mathbb{R}^3$ una quadrica degenera, luogo ai zeri di un'equazione

$${}^t x A x + {}^t b x + c = 0.$$

Per ipotesi la matrice simmetrica

$$\bar{A} = \begin{pmatrix} A & b \\ {}^t b & c \end{pmatrix}$$

ha rango ≤ 3 . Nel caso in cui il rango sia 3 le segnature possibili sono

$$(3, 0, 1), \quad (2, 1, 1), \quad (1, 2, 1), \quad (0, 3, 1).$$

Diciamo che $\bar{A}$ è *semidefinita* nel primo e ultimo caso e *indefinita* altrimenti.

Nel caso in cui il rango sia 2 le segnature possibili sono

$$(2, 0, 2), \quad (1, 1, 2), \quad (0, 2, 2)$$

e $\bar{A}$ è *semidefinita* nel primo e terzo caso e *indefinita* nel secondo.

Teorema 13.4.2. Se Q è degenera, si presentano i casi seguenti:

- (1) Se $\text{rk} \bar{A} = 3$ e $\bar{A}$ è *semidefinita*, allora Q è
 - un punto se $\det A \neq 0$,
 - l'insieme vuoto se $\det A = 0$.
- (2) Se $\text{rk} \bar{A} = 3$ e $\bar{A}$ è *indefinita*, allora Q è
 - un cono se $\det A \neq 0$,
 - un cilindro ellittico se $\text{rk} A = 2$ e A è *semidefinita*,
 - un cilindro iperbolico se $\text{rk} A = 2$ e A è *indefinita*,
 - un cilindro parabolico se $\text{rk} A = 1$.
- (3) Se $\text{rk} \bar{A} = 2$ e $\bar{A}$ è *semidefinita*, allora Q è
 - una retta se $\text{rk} A = 2$,

- *l'insieme vuoto se $\text{rk}A = 1$.*
- (4) *Se $\text{rk}\bar{A} = 2$ e $\bar{A}$ è indefinita, allora Q è*
 - *due piani incidenti se $\text{rk}A = 2$,*
 - *due piani paralleli se $\text{rk}A = 1$.*
- (5) *Se $\text{rk}\bar{A} = 1$ allora Q è un piano.*

Dimostrazione. Per la Proposizione 13.1.7 possiamo supporre che A sia diagonale. Consideriamo prima il caso in cui A sia invertibile. Per la Proposizione 13.1.12 possiamo supporre che $\bar{A} = \begin{pmatrix} A & 0 \\ 0 & c \end{pmatrix}$ sia diagonale. Poiché $\bar{A}$ non è invertibile, abbiamo $c = 0$. Quindi Q è descritta da un'equazione

$$ax^2 + by^2 + cz^2 = 0$$

con opportuni coefficienti $a, b, c \neq 0$. La quadrica è un punto o un cono a seconda che $\bar{A}$ sia semidefinita o indefinita.

Consideriamo il caso $\det A = 0$. A meno di scambiare le variabili x, y e z otteniamo

$$\bar{A} = \begin{pmatrix} a_{11} & 0 & 0 & b_1 \\ 0 & a_{22} & 0 & b_2 \\ 0 & 0 & 0 & b_3 \\ b_1 & b_2 & b_3 & c \end{pmatrix}$$

Sappiamo che $\det \bar{A} = -a_{11}a_{22}b_3^2 = 0$. Consideriamo i vari casi. Se $b_3 = 0$, la variabile z non è presente nell'equazione della quadrica e quindi Q è un cilindro, la cui base è la conica C descritta dalla matrice

$$\begin{pmatrix} a_{11} & 0 & b_1 \\ 0 & a_{22} & b_2 \\ b_1 & b_2 & c \end{pmatrix}.$$

Dalla classificazione delle coniche otteniamo i vari casi come nell'enunciato (se C è un punto, una retta, due rette, allora Q è rispettivamente una retta, un piano, due piani). Infine, c'è da considerare il caso $b_3 \neq 0$ e $a_{22} = 0, a_{11} \neq 0$. Qui $\text{rk}\bar{A} = 3$ e l'equazione diventa

$$a_{11}x^2 + 2b_1x + 2b_2y + 2b_3z + c = 0.$$

Con un opportuno cambio di variabili (esercizio) si fa sparire una variabile e quindi si ottiene ancora un cilindro parabolico. $\square$

Esercizi

Esercizio 13.1. Considera al variare di $h \in \mathbb{R}$ la conica

$$C_h = \left\{ x^2 - 2hxy + \frac{1}{2}y^2 - 2hx + \frac{1}{2} = 0 \right\} \subset \mathbb{R}^2.$$

Determina per ogni $h \in \mathbb{R}$ il tipo di conica (ellisse, parabola, ecc.) ed i suoi centri quando esistono. Se la conica è un'iperbole, determina gli asintoti.

Esercizio 13.2. Considera al variare di $t \in \mathbb{R}$ la conica

$$C_t = \{ x^2 + (1-t)y^2 + 2tx - 2(1-t)y + 2 - t = 0 \}.$$

Determina per ogni $t \in \mathbb{R}$ il tipo di conica ed i suoi centri quando esistono. Per quali $t \in \mathbb{R}$ la conica è una circonferenza? Se la conica è un'iperbole, determina gli asintoti.

Esercizio 13.3. Considera al variare di $k \in \mathbb{R}$ la conica

$$C_k = \{x^2 - 6xy + ky^2 + 2x + k = 0\}.$$

Determina per ogni $k \in \mathbb{R}$ il tipo di conica ed i suoi centri quando esistono. Se la conica è un'iperbole, determina gli asintoti.

Esercizio 13.4. Determina la conica passante per i punti

$$P_1 = \begin{pmatrix} 1 \\ 0 \end{pmatrix}, \quad P_2 = \begin{pmatrix} 1 \\ 1 \end{pmatrix}, \quad P_3 = \begin{pmatrix} -1 \\ -1 \end{pmatrix}, \quad P_4 = \begin{pmatrix} -1 \\ 0 \end{pmatrix}$$

e tangente in P_4 alla retta $y = x + 1$.

Esercizio 13.5. Determina la conica tangente alle rette

$$r = \{x + y + 3 = 0\}, \quad r' = \{x - y - 2 = 0\}$$

nei punti $P = \begin{pmatrix} -2 \\ -1 \end{pmatrix}$, $P' = \begin{pmatrix} 1 \\ -1 \end{pmatrix}$ e passante per l'origine.

Esercizio 13.6. Determina il fascio di coniche proiettive passanti per i punti

$$[1, 1, 0], \quad [1, 0, 1], \quad [0, 1, 1], \quad [0, 0, 1].$$

Determina tutte le coniche degeneri del fascio.

Esercizio 13.7. Determina il tipo delle quadriche seguenti:

$$6xy + 8xz - 5y = 0,$$

$$x^2 + 2y^2 + z^2 + 2xz - 4 = 0,$$

$$z^2 = xy,$$

$$5y^2 + 12yz + 2x + 4 = 0,$$

$$x^2 - 2xy + 2y^2 + 2yz + z^2 - 2y + 2 = 0.$$

Esercizio 13.8. Considera al variare di $k \in \mathbb{R}$ la conica

$$C_k = \{x^2 + 2kxy + y^2 + 2kx + 2ky + 2k - 2 = 0\}.$$

Determina per ogni $k \in \mathbb{R}$ il tipo di conica ed i suoi centri quando esistono. Per tutti i k per cui C_k è una ellisse, determina il rapporto fra l'asse maggiore e quello minore al variare di k . Per tutti i k per cui C_k è una iperbole, determina gli asintoti.

Esercizio 13.9. Scrivi l'equazione della parabola passante per $(0, 2)$ che ha vertice in $(-2, 2)$ e direttrice parallela a $x + y = 0$.

Esercizio 13.10. Date due rette incidenti r e s in $\mathbb{R}^2$ scritte in forma cartesiana, costruisci un metodo per scrivere il fascio di tutte le iperboli aventi r e s come asintoti.

Esercizio 13.11. Scrivi il fascio di coniche tangenti nell'origine all'asse y e passanti per i punti $(1, 0)$ e $(1, 2)$. Determina l'unica parabola del fascio.

Esercizio 13.12. Scrivi il fascio di coniche tangenti nel punto $(1, 0)$ alla retta $x + y - 1 = 0$ e tangenti nel punto $(3, 0)$ alla retta $x - y - 3 = 0$. Determina l'unica parabola del fascio.

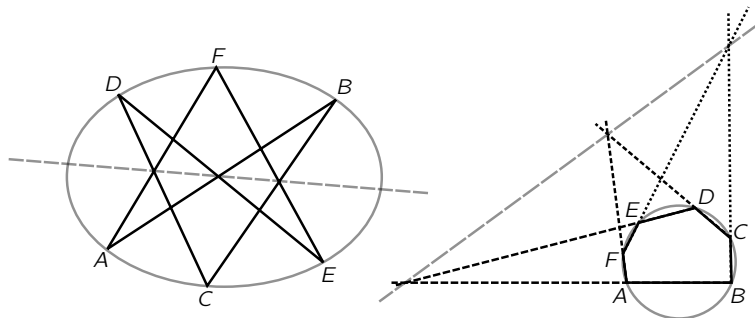


Figura 13.15. Il Teorema di Pascal in due configurazioni diverse.

Complementi

13.1. Il Teorema di Pascal. Nel 1639, all'età di 16 anni, Pascal enunciò un teorema, noto anche come *Hexagrammum mysticum*, che generalizza il Teorema di Pappo introdotto nella Sezione 12.1. Indichiamo con $L(P, Q)$ la retta contenente i due punti P e Q .

Teorema 13.4.3 (Teorema di Pascal). *Siano $A, B, C, D, E, F \in \mathbb{P}^2$ sei punti distinti che giacciono tutti su una conica non degenera. I tre punti*

$$L(A, B) \cap L(D, E), \quad L(B, C) \cap L(E, F), \quad L(C, D) \cap L(F, A)$$

sono allineati.

La configurazione dei 6 punti distinti A, B, C, D, E, F nella conica è totalmente arbitraria. Due esempi sono mostrati nella Figura 13.15.

Dimostrazione. Poiché la conica è non degenera, i punti A, C, E non sono allineati; quindi possiamo fissare un sistema di riferimento in cui

$$A = [1, 0, 0], \quad C = [0, 1, 0], \quad E = [0, 0, 1].$$

La conica ha una certa equazione

$$ax^2 + by^2 + cz^2 + dxy + eyz + fzx = 0.$$

Il passaggio per A, C, E impone che $a = b = c = 0$. Quindi l'equazione è

$$(37) \quad dxy + eyz + fzx = 0.$$

Poniamo

$$B = [x_1, y_1, z_1], \quad D = [x_2, y_2, z_2], \quad F = [x_3, y_3, z_3].$$

Abbiamo $x_i, y_i, z_i \neq 0$ per ogni i perché non ci sono triple di punti allineati fra i sei A, B, C, D, E, F . Dividendo (37) per xyz otteniamo

$$(38) \quad \frac{d}{z_1} + \frac{e}{x_1} + \frac{f}{y_1} = 0$$

per ogni $i = 1, 2, 3$. Otteniamo facilmente che

$$L(A, B) \cap L(D, E) = [x_2 y_1, y_2 y_1, y_2 z_1] = [x_2/y_2, 1, z_1/y_1],$$

$$L(B, C) \cap L(E, F) = [x_1 x_3, x_1 y_3, z_1 x_3] = [1, y_3/x_3, z_1/x_1],$$

$$L(C, D) \cap L(F, A) = [x_2 z_3, z_2 y_3, z_2 z_3] = [x_2/z_2, y_3/z_3, 1].$$

I tre punti sono allineati perché

$$\det \begin{pmatrix} 1 & y_3/x_3 & z_1/x_1 \\ x_2/y_2 & 1 & z_1/y_1 \\ x_2/z_2 & y_3/z_3 & 1 \end{pmatrix} = x_2 y_3 z_1 \det \begin{pmatrix} 1/x_2 & 1/x_3 & 1/x_1 \\ 1/y_2 & 1/y_3 & 1/y_1 \\ 1/z_2 & 1/z_3 & 1/z_1 \end{pmatrix} = 0.$$

Il determinante dell'ultima matrice è zero perché le tre righe sono dipendenti: la loro combinazione con coefficienti e, f, d è nulla per (38). $\square$

Il Teorema di Pascal è un teorema di geometria proiettiva che è quindi valido anche nel piano cartesiano, su qualsiasi conica non degenera: dati sei punti distinti A, B, C, D, E, F in una ellisse, parabola o iperbole di $\mathbb{R}^2$, il teorema si applica sempre (stando attenti alla possibilità che qualche punto di intersezione finisca all'infinito). Si tratta quindi di un risultato molto generale che, ottenuto con una sola dimostrazione, si applica in una quantità notevole di casi geometrici differenti.

Il Teorema di Pascal può essere effettivamente interpretato come una generalizzazione del Teorema di Pappo, dal caso di una conica degenera (due rette incidenti) a quello di una conica non degenera.

Soluzioni di alcuni esercizi

Capitolo 1

Esercizio 1.1. $\sqrt{n} = p/q \iff n^2 = p^2/q^2$. Supponiamo che la frazione p/q sia ridotta ai minimi termini, cioè p e q non hanno fattori in comune. Allora p^2/q^2 è un numero naturale se e solo se $q = \pm 1$.

Esercizio 1.2. Effettivamente $p_1 \cdots p_n + 1$ non è divisibile per p_i , perché dividendo $p_1 \cdots p_n + 1$ per p_i otteniamo come resto 1.

Esercizio 1.8. La terza.

Esercizio 1.11. L'insieme X contiene n elementi. Se f è iniettiva, elementi diversi vanno in elementi diversi, e quindi l'immagine contiene n elementi. Poiché ci sono n elementi nel codominio, f è suriettiva. Se invece f non è iniettiva, ci sono almeno due elementi che vanno nello stesso e quindi l'immagine contiene al più $n - 1$ elementi, quindi non può essere suriettiva.

Con gli insiemi infiniti non è vero, ad esempio $f: \mathbb{N} \rightarrow \mathbb{N}, f(x) = 2x$ è iniettiva ma non suriettiva.

Esercizio 1.12. Scriviamo $p(x) = (x - a)^m q(x)$ con $q(a) \neq 0$. Usando la formula della derivata di un prodotto di funzioni, otteniamo

$$p'(x) = m(x - a)^{m-1}q(x) + (x - a)^m q'(x) = (x - a)^{m-1}(mq(x) + (x - a)q'(x)).$$

Inoltre a non è radice del secondo fattore $h(x) = mq(x) + (x - a)q'(x)$ perché $h(a) = mq(a) \neq 0$.

Esercizio 1.13. $\sqrt{2} + \sqrt{2}i, -\sqrt{2} + \sqrt{2}i, -\sqrt{2} - \sqrt{2}i, \sqrt{2} - \sqrt{2}i$.

Esercizio 1.14. $e^{2k\pi i/7}$ per $k = 0, \dots, 6$.

Capitolo 2

Esercizio 2.14. $\begin{pmatrix} 1/3 \\ -2/3 \end{pmatrix}, \begin{pmatrix} 1/3 \\ 1/3 \end{pmatrix}, \begin{pmatrix} 3 \\ 1 \end{pmatrix}$.

Esercizio 2.20. Sappiamo che esiste una base $v_1, \dots, v_n$ per V . Prendiamo $W = \text{Span}(v_1, \dots, v_k)$.

Esercizio 2.21. Prendiamo una base $v_1, \dots, v_k$ per U . Sappiamo che questa si può completare a base $v_1, \dots, v_n$ per V . Definiamo $W = \text{Span}(v_{k+1}, \dots, v_n)$.

Esercizio 2.24. Per ogni $t \neq -1$.

Esercizio 2.26. Per ipotesi esiste base $v_1, \dots, v_n$. Basta prendere $V_j = \text{Span}(v_j)$.

Esercizio 2.27. (5) A meno di permutare a, b, c otteniamo:

$$(3, 3, 3, 3), \quad (3, 2, 2, 2), \quad (2, 2, 2, 2), \quad (2, 2, 2, 1).$$

Capitolo 3

Esercizio 3.1.

- $x_1 = 1, x_2 = 8, x_3 = 3, x_4 = 6.$
- $x_1 = 1 - 3t_1 - 4t_2, x_2 = 2 - t_1 - 3t_2, x_3 = t_1, x_4 = t_2, x_5 = 3.$

Esercizio 3.2.

- Se $\alpha = 0$, allora $x = 1 - \beta, y = t, z = 1.$
- Se $\alpha \neq 0$ e $\alpha + \beta \neq 0$, allora $x = \frac{\alpha + \beta - \beta^2}{\alpha + \beta}, y = 0, z = \frac{\beta}{\alpha + \beta}.$
- Se $\alpha \neq 0$ e $\alpha + \beta = 0$, non ci sono soluzioni.

Esercizio 3.3. 5, 2160, 0.

Esercizio 3.4.

- Una retta per $c = 22$, vuoto per $c \neq 22.$
- Un punto per $c = -2$, vuoto per $c \neq -2.$
- Un punto per $c \neq 0, 1$, una retta affine per $c = 0$, vuoto per $c = 1.$
- Un punto per $c \neq 0, 1$, una retta affine per $c = 0$, vuoto per $c = 1.$

Esercizio 3.5.

$$\frac{1}{2} \begin{pmatrix} 1 & i \\ -i & -1 \end{pmatrix}, \quad \frac{1}{24} \begin{pmatrix} 1 & 11 & -6 \\ -1 & 13 & -18 \\ -4 & 4 & 0 \end{pmatrix}, \quad \begin{pmatrix} -2 & 6 & -3 & 1 \\ 0 & -3 & 1 & 0 \\ -1 & 1 & 0 & 0 \\ 1 & 0 & 0 & 0 \end{pmatrix}.$$

Esercizio 3.7. 2, 2, 1, 3.

Esercizio 3.8.

- (1) $3k = 0, 1$ per $k \neq 0.$
 (2) ${}^t(-1 \ 1 \ 0 \ 0), {}^t(0 \ 0 \ 4 \ 5)$

Esercizio 3.11. ${}^tA = -A$ quindi $\det A = \det({}^tA) = \det(-A)$. Usando l'esercizio precedente $\det(-A) = (-1)^{17} \det A = -\det A$. Quindi $\det A = -\det A \Rightarrow \det A = 0$.

Capitolo 4

Esercizio 4.1. $\ker T = \{0\}, \operatorname{Im} T = \operatorname{Span}(x, x^2, x^3).$

Esercizio 4.3. $\ker f = \{0\}, \operatorname{Im} f = \operatorname{Span}({}^t(2 \ 1 \ 3), {}^t(1 \ -1 \ 2)).$

$$\begin{pmatrix} 2 & 1 \\ 1 & -1 \\ 3 & 2 \end{pmatrix}, \quad \begin{pmatrix} 1 & -1 \\ 1 & 1 \\ 0 & 4 \end{pmatrix}.$$

Esercizio 4.4. $\frac{1}{2} \begin{pmatrix} 1 & -1 & 1 \\ -3 & 3 & -1 \\ 1 & 1 & 1 \end{pmatrix}.$

Esercizio 4.6. $A = [T]_{\mathcal{C}}^{\mathcal{C}} = \begin{pmatrix} 0 & 1 & 0 & \cdots & 0 \\ 0 & 1 & 2 & \cdots & 0 \\ 0 & 0 & 2 & \cdots & 0 \\ \vdots & \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & 0 & \cdots & n \end{pmatrix},$ $\ker T$ sono i polinomi

costanti, $\operatorname{Im} T$ sono i polinomi divisibili per $x + 1$, cioè quelli che si annullano in -1 , $\operatorname{tr} A = \frac{(n+1)n}{2}, \det A = 0$.

Esercizio 4.7.

(1) $\ker L_A = \{0\}$, $\operatorname{Im} L_A = M(2)$.

(2) $\begin{pmatrix} 1 & 0 & 2 & 0 \\ 0 & 1 & 0 & 2 \\ 3 & 0 & 4 & 0 \\ 0 & 3 & 0 & 4 \end{pmatrix}$.

(3) 4.

(4) Prendendo $B = \{e_{11}, e_{21}, e_{12}, e_{22}\}$ si ottiene $[L_A]_B^B = \begin{pmatrix} A & 0 \\ 0 & A \end{pmatrix}$ che ha determinante $(\det A)^2$ ed è invertibile $\iff \det A \neq 0$.

Esercizio 4.8. (3) Il nucleo contiene sempre I_2 e A . Se I e A sono dipendenti, cioè $A = \lambda I_2$, allora T è banale. In ogni caso $\dim \ker T \geq 2$.

Esercizio 4.9. Ad esempio $\begin{pmatrix} 2 & 1 & -3 \\ 1 & 0 & -1 \\ 0 & -1 & 1 \end{pmatrix}$.

Esercizio 4.11. $\begin{pmatrix} 1 & 0 & 0 \\ -1 & 0 & 0 \\ -1 & -1 & 1 \end{pmatrix}, \begin{pmatrix} 0 & 0 & 0 \\ 1 & 1 & 0 \\ 1 & 1 & 0 \end{pmatrix}$.

Esercizio 4.12. Sia $B = \{v_1, \dots, v_n\}$ base di V tale che gli ultimi $n - r$ vettori siano una base di $\ker f$. Siano $w_1 = f(v_1), \dots, w_r = f(v_r)$. Prendi un completamento $C = \{w_1, \dots, w_m\}$ a base di W qualsiasi.

Esercizio 4.14. T è iniettiva perché un polinomio di grado $\leq n$ non nullo non può annullarsi negli $n + 1$ punti $x_0, \dots, x_n$. Siccome $\dim \mathbb{K}_n[x] = \dim \mathbb{K}^{n+1}$, T è un isomorfismo. La matrice associata rispetto alle basi canoniche è quella di Vandermonde. Come conseguenza, la matrice di Vandermonde è invertibile.

Esercizio 4.16. Prendi base $v_1, \dots, v_k$ di U e completa a base $v_1, \dots, v_m$ di V . Prendi base $w_1, \dots, w_h$ di Z e completa a base $w_1, \dots, w_n$ di W . Usa queste basi per definire un isomorfismo $T: \operatorname{Hom}(V, W) \rightarrow M(n, m)$. Nota che $T(S)$ è l'insieme di tutte le matrici $\begin{pmatrix} A & B \\ 0 & C \end{pmatrix}$ con $A \in M(h, k)$, $B \in M(h, m - k)$ e $C \in M(n - h, m - k)$. Questo insieme è un sottospazio di $M(n, m)$ di dimensione

$$hk + h(m - k) + (n - h)(m - k) = mn + hk - kn.$$

Capitolo 5

Esercizio 5.1. Sì. No. Per $b = c = d = 0$. Per $k = 0$. Per $k \neq 1$. Per $k \neq 1, 2$.

Sì: la matrice ha rango 1. Una base di autovettori è data prendendo il vettore ${}^t(1 \ 1 \ 1 \ 1 \ 1 \ 1)$ e 5 generatori del nucleo.

No: la matrice ha rango 2 (ci sono solo due tipi di righe) quindi l'autovalore 0 ha molteplicità geometrica 5. Il vettore ${}^t(1 \ 1 \ 1 \ 1 \ 1 \ 1)$ è autovettore con autovalore 1. Quindi le radici del polinomio caratteristico sono $0, 0, 0, 0, 0, 1, \lambda_0$ con λ_0 da determinare. La traccia è sempre la somma $1 + \lambda_0$ degli autovalori, quindi $1 + \lambda_0 = 1 \Rightarrow \lambda_0 = 0$. Quindi 0 ha molteplicità algebrica 5 e geometrica 6: non è diagonalizzabile.

Le matrici di questi esempi non diagonalizzabili su $\mathbb{R}$ non lo sono neppure su $\mathbb{C}$.

Esercizio 5.2. Per ogni $t \in \mathbb{C}$ diverso da 5.

Esercizio 5.3. $-1, 2 + i, \begin{pmatrix} -1 \\ 1 \end{pmatrix}, \begin{pmatrix} 1-i \\ 1 \end{pmatrix}$.

Esercizio 5.4. $\begin{pmatrix} 2^{13} - 7 & -7 \cdot 2^{12} + 28 \\ 2^{11} - 2 & -7 \cdot 2^{10} + 8 \end{pmatrix}$.

Esercizio 5.5. Se si sceglie $e_{11}, e_{21}, e_{12}, e_{22}$ la matrice associata a L_A è semplicemente $\begin{pmatrix} A & 0 \\ 0 & A \end{pmatrix}$ che è diagonalizzabile se e solo se A lo è (esercizio).

Esercizio 5.8. Se T è diagonalizzabile, esiste $v_1, \dots, v_n$ base di autovettori per T . Questa è anche una base di autovettori per T^{-1} .

Esercizio 5.10. Sia $v_1, \dots, v_n$ base di autovettori, con i primi $v_1, \dots, v_k$ che hanno autovalore 0. Allora $\ker T = \text{Span}(v_1, \dots, v_k)$, $\text{Im } T = \text{Span}(v_{k+1}, \dots, v_n)$.

Capitolo 6

Esercizio 6.1. $\begin{pmatrix} -1 & 1 \\ 0 & -1 \end{pmatrix}, \begin{pmatrix} 1 & 0 & 0 \\ 0 & 4 & 1 \\ 0 & 0 & 4 \end{pmatrix}, \begin{pmatrix} 0 & 1 & 0 \\ 0 & 0 & 0 \\ 0 & 0 & 1 \end{pmatrix},$
 $\begin{pmatrix} -2 & 0 & 0 \\ 0 & 4 & 1 \\ 0 & 0 & 4 \end{pmatrix}, \begin{pmatrix} 2 & 1 & 0 \\ 0 & 2 & 1 \\ 0 & 0 & 2 \end{pmatrix}, \begin{pmatrix} 1 & 0 & 0 \\ 0 & 1 & 1 \\ 0 & 0 & 1 \end{pmatrix}.$

Esercizio 6.4. La loro forma di Jordan è necessariamente $B_{0,n}$. Sono tutte simili a $B_{0,n}$ e quindi simili tra loro.

Esercizio 6.6. $x^2(x-1)^2, (x+2)^2(x-2)^2, x^2.$

Esercizio 6.9.

$$\begin{pmatrix} \boxed{\begin{matrix} 1 & 1 & 0 \\ 0 & 1 & 1 \\ 0 & 0 & 1 \end{matrix}} & \begin{matrix} 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 \end{matrix} \\ \begin{matrix} 0 & 0 & 0 & 1 \\ 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 \end{matrix} & \boxed{\begin{matrix} 1 & 1 & 0 & 0 \\ 0 & 1 & 1 & 0 \\ 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 1 \end{matrix}} \end{pmatrix}, \begin{pmatrix} \boxed{\begin{matrix} 1 & 1 & 0 \\ 0 & 1 & 1 \\ 0 & 0 & 1 \end{matrix}} & \begin{matrix} 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 \end{matrix} \\ \begin{matrix} 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 \end{matrix} & \boxed{\begin{matrix} 1 & 1 \\ 0 & 1 \end{matrix}} \begin{matrix} 0 & 0 \\ 0 & 0 \end{matrix} \\ \begin{matrix} 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 \end{matrix} & \boxed{\begin{matrix} 1 & 1 \\ 0 & 1 \end{matrix}} \end{pmatrix}.$$

Capitolo 7

Esercizio 7.1. $\begin{pmatrix} 0 & -2 \\ -2 & 0 \end{pmatrix}.$

Esercizio 7.2.

(2) $g(A, A) = \text{tr}({}^tAA) = \sum_{i,j=1}^n A_{ij}^2 > 0$ se $A \neq 0$.

(3) No. Esistono matrici $A \neq 0$ con $\text{tr}(A^2) = 0$.

(4) La base canonica, ad esempio.

(5) $\left(\frac{n(n+1)}{2}, \frac{n(n-1)}{2}, 0\right).$

Esercizio 7.3. Il radicale è $\text{Span}(x^2 - 1)$.

Esercizio 7.7. $W = \text{Span}(x - 1, x^2 - x). W^\perp = \text{Span}(x^2 + 4x + 4).$

Esercizio 7.8.

(1) Le antisimmetriche.

(2) Le matrici del tipo $\begin{pmatrix} 0 & a \\ b & 0 \end{pmatrix}.$

Esercizio 7.10.

(1) No.

(2) Sì.

(3) $\{y^2 + 2xz = 0\}$. No.

Esercizio 7.13. $t < 1$: $(1, 1, 1)$, $t = 1$: $(1, 0, 2)$, $t > 1$: $(2, 0, 1)$.

Esercizio 7.15. La prima, la seconda e la quarta sono congruenti.

Esercizio 7.16. $(2, 1, 0)$. $(2, 0, 1)$. Per $\alpha \in (0, 4/5)$: $(1, 2, 0)$; per $\alpha \in (-\infty, 0) \cup (4/5, \infty)$: $(2, 1, 0)$; per $\alpha \in \{0, 4/5\}$: $(1, 1, 1)$. Per $\alpha < -1$: $(1, 2, 0)$; per $\alpha = -1$: $(1, 1, 0)$; per $\alpha > -1$: $(2, 1, 0)$.

Capitolo 8

Esercizio 8.2. $\begin{pmatrix} 1 \\ -1 \\ 0 \end{pmatrix}, \begin{pmatrix} 1 \\ 1 \\ 1 \end{pmatrix}, \begin{pmatrix} -1 \\ -\frac{1}{2} \\ 1 \end{pmatrix}$. Coordinate: $\begin{pmatrix} \frac{7}{6} \\ \frac{2}{3} \\ \frac{1}{6} \end{pmatrix}$.

Esercizio 8.8. $\frac{1}{3} \begin{pmatrix} 2 & -1 & 2 \\ -1 & 2 & 2 \\ 2 & 2 & -1 \end{pmatrix}, \frac{1}{6} \begin{pmatrix} 5 & -1 & 2 \\ -1 & 5 & 2 \\ 2 & 2 & 2 \end{pmatrix}$.

Esercizio 8.10.

$$(1) \frac{1}{26} \begin{pmatrix} 10 & 12 & -4 \\ 12 & 17 & 3 \\ -4 & 3 & 25 \end{pmatrix}$$

(2) Sì, perché se prendiamo una base v_1, v_2 di U e un generatore v_3 di $U^\perp$, e usiamo la base $B = \{v_1, v_2, v_3\}$ di $\mathbb{R}^3$, otteniamo

$$[f]_B^B = [\rho]_B^B - [q]_B^B = \begin{pmatrix} 1 & 0 & 0 \\ 0 & 1 & 0 \\ 0 & 0 & 0 \end{pmatrix} - \begin{pmatrix} 0 & 0 & 0 \\ 0 & 0 & 0 \\ 0 & 0 & 1 \end{pmatrix} = \begin{pmatrix} 1 & 0 & 0 \\ 0 & 1 & 0 \\ 0 & 0 & -1 \end{pmatrix}$$

che è diagonale.

$$(3) U^\perp = \text{Span}\left(\begin{pmatrix} 53 & -40 & 46 \end{pmatrix}\right).$$

$$\text{Esercizio 8.13. } \begin{pmatrix} 0 & 0 & -1 \\ 1 & 0 & 0 \\ 0 & -1 & 0 \end{pmatrix}.$$

Esercizio 8.14. Una riflessione $r_{U^\perp}$ rispetto alla retta $U^\perp$.

Esercizio 8.15. Gli unici vettori di norma uno con coefficienti interi sono questi: $e_1, e_2, e_3, -e_1, -e_2, -e_3$. Usando solo questi vettori si possono costruire 48 basi diverse, quindi 48 matrici ortogonali diverse. Di queste, metà, cioè 24, hanno determinante positivo. Una è l'identità, le altre 23 sono rotazioni. Più precisamente ci sono:

- rotazioni di angolo $\frac{\pi}{2}, \pi, \frac{3\pi}{2}$ intorno ai 3 assi coordinati x, y, z , quindi $3 \cdot 3 = 9$ rotazioni in tutto.
- rotazioni di angolo $\frac{\pi}{3}$ oppure $\frac{2\pi}{3}$ intorno ai 4 assi $\text{Span}(e_1 \pm e_2 \pm e_3)$, quindi $2 \cdot 4 = 8$ rotazioni in tutto.
- rotazioni di angolo π intorno ai 6 assi $\text{Span}(e_1 \pm e_2), \text{Span}(e_2 \pm e_3), \text{Span}(e_3 \pm e_1)$. Quindi 6 rotazioni in tutto.

Il totale è effettivamente $9+8+6 = 23$ e si verifica che queste sono tutte simmetrie del cubo.

Esercizio 8.18. Per ogni $\theta \neq \frac{\pi}{2}$.

Capitolo 9

$$\text{Esercizio 9.3. } \begin{pmatrix} 0 \\ 1 \\ 3 \end{pmatrix} + \text{Span} \begin{pmatrix} 1 \\ 0 \\ 0 \end{pmatrix}.$$

Esercizio 9.4.

$$(1) \begin{pmatrix} 1 \\ 1 \\ 0 \end{pmatrix} + \text{Span} \begin{pmatrix} -2 \\ -1 \\ 1 \end{pmatrix}.$$

$$(2) 2x + y - z = 4.$$

$$(3) x + 2z = 1.$$

Esercizio 9.5.

(1) Sono sghembe e a distanza $\frac{\sqrt{6}}{6}$. La retta perpendicolare ad entrambe

$$\text{è } \begin{pmatrix} -7/3 \\ 4/3 \\ 10/3 \end{pmatrix} + \text{Span} \begin{pmatrix} 2 \\ 1 \\ 1 \end{pmatrix}.$$

$$\text{Esercizio 9.8. } s = \begin{pmatrix} 2 \\ 1 \\ 0 \end{pmatrix} + \text{Span} \begin{pmatrix} 7 \\ 6 \\ 3 \end{pmatrix}.$$

Esercizio 9.14. Calcola i punti $P = r_1 \cap r_2$, $Q = r_2 \cap r_3$ e $R = r_3 \cap r_1$ e costruisci f in modo che $f(P) = Q$, $f(Q) = R$, $f(R) = P$.

Esercizio 9.15. Ad esempio $f \begin{pmatrix} x \\ y \end{pmatrix} = \begin{pmatrix} 1 & 0 \\ 0 & -1 \end{pmatrix} \begin{pmatrix} x \\ y \end{pmatrix} + \begin{pmatrix} 1 \\ -1 \end{pmatrix}$.

$$\text{Esercizio 9.16. } f \begin{pmatrix} x \\ y \\ z \end{pmatrix} = \begin{pmatrix} -1 & 0 & 0 \\ 0 & -1 & 0 \\ 0 & 0 & 1 \end{pmatrix} \begin{pmatrix} x \\ y \\ z \end{pmatrix} + \begin{pmatrix} 2 \\ -2 \\ 1 \end{pmatrix}.$$

Esercizio 9.19.

- (1) Una rotazione di asse $r = {}^t(1 \quad 5 \quad 1)$ e di angolo $\pm \arccos \frac{13}{14}$.
 (2) Per tutti i $b \in r^\perp$.

Capitolo 11

Esercizio 11.3. L'endomorfismo ha una base ortonormale di autovettori con autovalori che possono essere solo 1 e -1 . Quindi $V = V_1 \oplus V_{-1}$ è una somma diretta ortogonale di autospazi e l'endomorfismo è una riflessione ortogonale lungo V_1 .

Esercizio 11.7. Sia $\mathcal{B}$ base ortonormale per g . Abbiamo $[g]_{\mathcal{B}} = I_n$ e $[f]_{\mathcal{B}} = S$ per qualche matrice simmetrica S . Per il Corollario 11.3.2 esiste una M ortogonale tale che ${}^tMSM = D$ sia diagonale. Sia $\mathcal{B}'$ base tale che $M = [\text{id}]_{\mathcal{B}}^{\mathcal{B}'}$. Otteniamo $[g]_{\mathcal{B}'} = I_n$ e $[f]_{\mathcal{B}'} = D$.

Esercizio 11.8. $\langle T(f), g \rangle = \langle hf, g \rangle = \int h(t)f(t)g(t)dt = \int f(t)h(t)g(t)dt = \langle f, hg \rangle = \langle f, T(g) \rangle$.

Esercizio 11.9. $\langle T(f), g \rangle = \langle f'', g \rangle = \int f''(t)g(t)dt = - \int f'(t)g'(t)dt$ che è uguale a $\int f(t)g''(t)dt = \langle f, g'' \rangle = \langle f, T(g) \rangle$ integrando due volte per parti (qui $\lim_{t \rightarrow \infty} f'(t)g(t) = \lim_{t \rightarrow \infty} f(t)g'(t) = 0$ perché f e g hanno supporto compatto).

Capitolo 12

Esercizio 12.1. $x_1 - x_2 = 0$.

Esercizio 12.2. $x_1 - x_3 + x_4 = 0$. $\{-t, u, 2t, 3t\}$. $x_1 + 2x_3 - x_4$.

Esercizio 12.3. $k = 1$.

Esercizio 12.4. $\begin{pmatrix} 1 & 0 \\ 0 & 0 \end{pmatrix}$.

$$\text{Esercizio 12.5. } \begin{pmatrix} 1 & 1 & -1 \\ 2 & 3 & 0 \\ 1 & 0 & 1 \end{pmatrix}.$$

Esercizio 12.9. $s = \{[5t + u, -t + u, -t + u, 3t]\}$

Capitolo 13

Esercizio 13.1. $|h| < 1/2$: insieme vuoto; $|h| = 1/2$: un punto; $|h| \in (1/2, \sqrt{2}/2)$: ellisse; $|h| = \sqrt{2}/2$: parabola; $|h| > \sqrt{2}/2$: iperbole.

Il centro esiste per $|h| \geq 1/2$ e $|h| \neq \sqrt{2}/2$ ed è $\frac{h}{1-2h^2}(1, 2h)$.

Gli asintoti esistono per $|h| > \sqrt{2}/2$ e sono le rette

$$\frac{h}{1-2h^2} \begin{pmatrix} 1 \\ 2h \end{pmatrix} + \text{Span} \left(\begin{pmatrix} h \pm \sqrt{h^2 - 1/2} \\ 1 \end{pmatrix} \right).$$

Esercizio 13.2. $t < -1$: ellisse; $t = -1$: punto; $t \in (-1, 1)$: insieme vuoto; $t = 1$: retta; $t > 1$: iperbole.

È una circonferenza per $t = 0$.

Per $t > 1$ gli asintoti sono le rette

$$\begin{pmatrix} -t \\ 1 \end{pmatrix} + \text{Span} \left(\begin{pmatrix} \pm \sqrt{t-1} \\ 1 \end{pmatrix} \right).$$

Esercizio 13.4. $x^2 + 2y^2 - 2xy - 1 = 0$.

Esercizio 13.5. $x^2 - y^2 + x + y = 0$.

Esercizio 13.7. Iperboloide ad una falda. Cilindro ellittico. Cono. Paraboloide iperbolico. Paraboloide ellittico.

Esercizio 13.8. Per $|k| < 1$ ellisse, per $k = -1$ parabola, per $k = 1$ due rette parallele, per $|k| > 1$ iperbole. Per $k \neq \pm 1$ il centro è unico ed è $\frac{1}{k+1} \begin{pmatrix} k \\ k \end{pmatrix}$. Per $k = -1$ non c'è centro, per $k = 1$ i centri formano la retta $x + y + 1 = 0$. Il rapporto fra gli assi per $|k| < 1$ è $\sqrt{(1+k)/(1-k)}$. Per $|k| > 1$ gli asintoti sono

$$\left\{ -\frac{1}{k+1} \begin{pmatrix} k \\ k \end{pmatrix} + t \begin{pmatrix} -k \pm \sqrt{k^2 - 1} \\ 1 \end{pmatrix} \right\}.$$

Esercizio 13.9. $x^2 - 2xy + y^2 + 6x - 10y + 16 = 0$.

Esercizio 13.10. Le iperboli con asintoti fissati formano effettivamente un fascio: sono le coniche che hanno centro in $r \cap s$ e che intersecano la retta all'infinito in due punti r_∞ e s_∞ fissati; la prima condizione è espressa con due equazioni lineari sui coefficienti, con le altre due otteniamo 4 condizioni lineari in tutto. Se le rette hanno equazione $ax + by + c = 0$ e $dx + ey + f = 0$, allora il fascio è $(ax + by + c)(dx + ey + f) = k$ al variare di k .

Esercizio 13.11. $kx^2 + 2xy - y^2 - kx = 0$. $(x-y)^2 - x = 0$.

Esercizio 13.12. $x^2 + (k-1)y^2 - 4x - 2y + 3 = 0$. $y = \frac{1}{2}(x^2 - 4x + 3)$.

Indice analitico

- $M(m, n, \mathbb{K})$, 50
- $\text{Hom}(V, W)$, 134
- $\text{Im } f$, 124
- $\ker f$, 123
- $\mathbb{K}[x]$, 48
- $\mathbb{K}_n[x]$, 54
- affinità, 291
- algoritmo
 - di completamento a base, 67
 - di estrazione di una base, 68
 - di Gauss, 80
 - di Gauss - Jordan, 82
- ampiezza di un angolo, 314
- anello, 35
- angolo, 313
 - acuto, 314
 - al centro, 323
 - alla circonferenza, 323
 - concavo, 314
 - convesso, 313
 - diedrale, 287
 - fra sottospazi affini incidenti, 285
 - fra vettori, 242
 - ottuso, 314
 - piatto, 314
 - retto, 314
- annullatore, 148
- antirotazione, 297
- applicazione lineare, 115
- arco di circonferenza, 323
- asintoto, 392, 418
- asse
 - di un segmento, 312
 - immaginaria, 26
 - reale, 26
- autospazio, 165
 - generalizzato, 178
 - massimale, 182
- autovalore, 151
- autovettore, 151
- baricentro, 325
- base, 62
 - canonica
 - di $\mathbb{K}^n$, 62
 - di $\mathbb{K}_n[x]$, 63
 - di $M(m, n, \mathbb{K})$, 63
 - di Jordan, 183
 - duale, 148
 - ortogonale, 219
 - ortonormale, 219
- bigezione, 15
- blocco di Jordan, 173
- campo, 36
- carta affine, 368
- cerchio, 316
- circocentro, 327
- circonferenza, 316
- coefficiente di Fourier, 244
- combinazione lineare, 52
- congruenza fra matrici, 213
- conica, 383
- coniugio di un numero complesso, 26
- controimmagine di una funzione, 13
- coordinate
 - baricentriche, 342
 - cilindriche, 145
 - di un vettore rispetto ad una base, 64
 - omogenee, 360
 - polari, 27
 - sferiche, 336
- corrispondenza biunivoca, 15
- covettore, 147
- criterio
 - di Cartesio, 354
 - di Jacobi, 225
- determinante, 93
- dimensione
 - di un sottospazio affine, 88
 - di uno spazio vettoriale, 65
- dimostrazione
 - per assurdo, 4
 - per induzione, 9
- distanza fra punti, 243
- disuguaglianza
 - di Bessel, 251
 - di Cauchy-Schwarz, 240

- triangolare, 240, 244, 321
- divisione fra polinomi, 21
- eccentricità, 393
- ellisse, 389
- endomorfismo, 137
 - diagonalizzabile, 153
- equazione lineare, 51
- fascio, 283
 - di coniche, 398
 - di piani, 285
 - di rette, 283
- forma
 - cartesiana, 54, 272, 362
 - parametrica, 54, 272
 - quadratica, 203
- formula
 - di Erone, 320
 - di Grassmann, 70, 367
- funzione, 12
 - bigettiva, 14
 - composta, 16
 - iniettiva, 14
 - inversa, 15
 - lineare, 115
 - suriettiva, 14
- giacitura, 88
- gruppo, 34
 - simmetrico, 34
- immagine di una funzione, 13
- incentro, 326
- indipendenza
 - affine, 277
 - lineare, 60
 - proiettiva, 366
- insieme
 - numerabile, 38
 - vuoto, 5
- iperbole, 391
- iperpiano, 362
- isometria, 227, 253
 - affine, 298
- isomorfismo, 127
 - affine, 292
- legge
 - del parallelogramma, 242
 - di reciprocità, 404
- leggi di De Morgan, 6
- linea
 - di Eulero, 328
 - spezzata, 312
- matrice, 49
 - antisimmetrica, 56, 73
 - associata
 - ad un prodotto hermitiano, 350
 - ad un prodotto scalare, 210
 - ad una applicazione lineare, 130, 137
 - dei cofattori, 109
 - di cambiamento di base, 135
 - di Gram, 265
 - di Jordan, 174
 - reale, 197
 - di Vandermonde, 113
 - diagonale, 56, 154
 - a blocchi, 187
 - diagonalizzabile, 155
 - hermitiana, 349
 - identità, 96
 - nilpotente, 176
 - quadrata, 56
 - simmetrica, 56, 73
 - trasposta, 73
 - triangolare, 56
- matrici
 - congruenti, 213
 - simili, 139
- minore, 100
- modulo
 - di un numero complesso, 26
 - di un vettore, 240
- molteplicità di una radice, 23
- molteplicità geometrica di un
 - autovalore, 166
- mossa di Gauss, 80
- norma di un vettore, 240, 348
- normalizzazione di un vettore, 250
- nucleo, 123
- numero
 - complesso, 25
 - intero, 4
 - naturale, 4
 - razionale, 4
 - reale, 4
- omotetia, 295
- orientazione, 298
- ortocentro, 327
- ortogonalizzazione di Gram-Schmidt, 247, 348
- parabola, 393
- parallelogramma, 315
- permutazione, 17
- piano
 - complesso, 26
 - proiettivo, 361
- pivot, 80
- polarità, 403
- poliedro, 332
- polinomio
 - caratteristico, 157
 - monico, 23
- prodotto
 - cartesiano, 8
 - fra matrici, 104
 - hermitiano, 347
 - definito positivo, 348
 - euclideo, 348
 - scalare, 199
 - definito negativo, 210

- definito positivo, 200
- degenere, 200
- euclideo, 200
- indefinito, 210
- lorentziano, 230
- semi-definito negativo, 210
- semi-definito positivo, 210
- triplo, 271
- vettoriale, 267
- proiettività, 373
- proiezione, 119
 - ortogonale, 244, 251
- proprietà
 - di Archimede, 42
- punti indipendenti, 277
- punto
 - fisso, 298
 - medio, 312
- quadrica
 - affine, 383
 - degenere, 385
 - proiettiva, 410
- quantificatore, 7
- radicale di un prodotto scalare, 208
- radice di un polinomio, 22
- rango, 89
- regola
 - della mano destra, 258
 - della mano destra, 270
 - di Cramer, 110
- relazione
 - di equivalenza, 10
 - di Eulero, 334
- restrizione
 - di un prodotto scalare, 216
 - di una funzione, 121
- retta
 - invariante, 156
 - polare, 403
 - proiettiva, 361
- rette complanari, 281
- riflessione, 120
 - ortogonale, 143, 227, 255, 296
- rotazione, 143, 295
- segmento, 311
- segnatura, 222
- segno di una permutazione, 19
- semipiano, 314
- semiretta, 313
- sesquilinearità, 348
- simbolo di Schläfli, 333
- similitudine
 - affine, 304
 - fra matrici, 139
- sistema lineare, 51
 - omogeneo, 51
 - associato, 85
- somma
 - di sottospazi, 58
 - diretta, 71, 74
 - ortogonale, 218
- sottoinsieme, 5
- sottomatrice, 92
- sottospazi affini
 - incidenti, 275
 - paralleli, 279
 - sghebbi, 364
- sottospazio
 - affine, 87, 272
 - generato, 53, 277
 - invariante, 156
 - ortogonale, 213
 - proiettivo, 362
 - somma, 367
 - vettoriale, 50
- spazio
 - cartesiano, 9
 - euclideo, 43
 - ortogonale, 348
 - proiettivo, 359, 361
 - duale, 380
 - vettoriale, 47
 - duale, 147
- spaziotempo di Minkowski, 230
- superficie rigata, 420
- sviluppo di Laplace, 96
- tensore, 233
- teorema
 - del coseno, 319
 - della curva di Jordan, 329
 - della dimensione, 124
 - di Binet, 107
 - di Cayley - Hamilton, 189
 - di Desargues, 378
 - di Jordan, 175
 - di Pappo, 379
 - di Pitagora, 243, 320
 - di Rouché - Capelli, 89
 - di Sylvester, 222
 - fondamentale dell'algebra, 31
 - spetttrale, 352
- traccia di una matrice, 141
- trasformazione
 - del piano, 142
 - lineare, 141
- traslazione, 292
- trasposizione, 18
- triangolo, 315
 - equilatero, 321
 - isoscele, 321
 - scaleno, 321
- V postulato di Euclide, 279
- vertice, 314
- vettore isotropo, 206
- vettori ortogonali, 205