

Università di Roma Tor Vergata
Corso di Laurea in
Scienza dei Media e della Comunicazione

Algebra lineare,
elementi di geometria analitica
ed aspetti matematici della
prospettiva

Massimo A. Picardello

László Zsidó

BOZZA 1.10.2011 23:39

Indice

Parte 1. Algebra lineare	1
Capitolo 1. Introduzione	3
1.1. Numeri interi, razionali e reali	3
1.2. Alcune notazioni	7
1.3. Appendice: numeri naturali, interi, razionali e reali	9
Capitolo 2. Spazi vettoriali	23
2.1. Definizione di spazio vettoriale	24
2.2. Proprietà della somma	25
2.3. Altre proprietà delle operazioni sui vettori	28
2.4. Combinazioni lineare di vettori	31
2.5. Sistemi di vettori linearmente dipendenti o indipendenti	32
2.6. Sottospazi vettoriali ed insiemi di generatori	36
2.7. Base di uno spazio vettoriale	38
Capitolo 3. Applicazioni lineari e matrici	47
3.1. Definizione di applicazione lineare	47
3.2. Immagine e nucleo di una applicazione lineare	48
3.3. Spazi vettoriali di polinomi ed applicazioni lineari	52
3.4. Esercizi sulle applicazioni lineari su spazi di polinomi	54
3.5. Applicazioni lineari fra $\mathbb{R}^n$ e $\mathbb{R}^m$ e calcolo con matrici	55
3.6. Spazi vettoriali di matrici	57
3.7. Soluzione di sistemi lineari: eliminazione di Gauss	61
3.7.1. Come trovare basi tramite eliminazione di Gauss	65
3.8. Rango di una matrice	70
3.9. Applicazioni lineari e matrici invertibili	73
3.10. Il calcolo della matrice inversa	74
3.11. Minori ed orli di una matrice ed invertibilità	80

3.12. Esercizi sulle trasformazioni lineari	86
3.13. Isomorfismi fra spazi vettoriali	93
Capitolo 4. Cambiamento di base	95
4.1. Trasformazione di coordinate sotto cambiamento di base	95
4.2. Matrice di un'applicazione lineare e cambiamento di basi	100
4.3. Cenni introduttivi sulla diagonalizzazione	108
4.4. Esercizi sulla diagonalizzazione	109
Capitolo 5. Determinante di matrici	115
Capitolo 6. Prodotti scalari e ortogonalità	119
6.1. Prodotto scalare euclideo nel piano ed in $\mathbb{R}^n$	119
6.2. Spazi vettoriali su $\mathbb{C}$	121
6.3. * La definizione generale di prodotto scalare	122
6.4. Matrici complesse, autoaggiunte e simmetriche	124
6.5. * Norma e prodotti scalari definiti positivi	125
6.6. Ortogonalità	127
6.7. Procedimento di ortogonalizzazione di Gram-Schmidt	133
6.8. Matrici ortogonali e matrici unitarie	138
6.9. * Matrice associata ad un prodotto scalare	141
Capitolo 7. Autovalori, autovettori e diagonalizzabilità	147
7.1. Triangolarizzazione e diagonalizzazione	147
7.2. Autovalori, autovettori e diagonalizzazione	148
7.3. Ulteriori esercizi sulla diagonalizzazione sul campo $\mathbb{R}$	157
7.4. Autovalori complessi e diagonalizzazione in $M_n^{\mathbb{C}}$	173
7.5. Diagonalizzabilità di matrici simmetriche o autoaggiunte	174
7.6. Esercizi sulla diagonalizzazione di matrici simmetriche	178
7.7. Qualche applicazione degli autovettori	180
7.7.1. Dinamica delle popolazioni	180
7.7.2. Sistemi dinamici	182
Capitolo 8. Somma diretta di spazi vettoriali	187
Capitolo 9. Triangolarizzazione e forma canonica di Jordan	189
9.1. * Polinomio minimo	189
9.2. Forma canonica di Jordan	189

Capitolo 10. Spazi normati, funzionali lineari e dualità	191
10.1. La disuguaglianza di Cauchy-Schwartz	191
10.2. Operatori lineari e funzionali lineari su spazi normati	192
10.3. * Duale di somme dirette e di complementi ortogonali	195
Parte 2. Geometria analitica e proiettiva	199
Capitolo 11. Vettori e geometria euclidea ed analitica nel piano	201
11.1. L'equazione della retta nel piano	201
11.2. Cerchi nel piano	205
11.3. Esercizi di geometria analitica nel piano	205
11.3.1. Rette	206
11.3.2. Triangoli	209
11.3.3. Parallelogrammi	210
11.3.4. Cerchi	214
Capitolo 12. Vettori e geometria euclidea ed analitica in $\mathbb{R}^3$	221
12.1. L'equazione del piano in $\mathbb{R}^3$	221
12.2. L'equazione della retta in $\mathbb{R}^3$	223
12.3. Proiezioni e distanze fra piani	226
12.4. Sfere	229
12.5. Esercizi di geometria analitica in $\mathbb{R}^3$	230
12.6. Raggi riflessi e rifratti	236
Capitolo 13. * Spazi quozienti, spazi proiettivi	241
13.1. Spazio quoziente e dualità	241
13.2. Lo spazio proiettivo	244
13.3. (*) $\mathbb{P}^n$ come compattificazione di $\mathbb{R}^n$	248
13.4. Trasformazioni proiettive	249
Capitolo 14. * Trasformazioni affini	253
14.1. Moti rigidi in $\mathbb{R}^n$ immersi in trasformazioni lineari di $\mathbb{R}^{n+1}$	253
14.2. Esempio: spostamento di una macchina da ripresa	258
Capitolo 15. * Quaternioni e matrici di rotazione	263
15.1. Espressione delle rotazioni in forma assiale	263
15.2. Rotazioni in $\mathbb{R}^2$, numeri complessi ed estensione a tre dimensioni	264
15.3. Quaternioni	265

15.3.1. Proprietà e definizioni	266
Definizioni	266
Proprietà	267
15.4. Rotazioni in $\mathbb{R}^3$ e coniugazione di quaternioni	268
15.4.1. Composizione di rotazioni e prodotto di quaternioni	270
15.4.2. Matrice di rotazione in termini di quaternioni	271
Parte 3. Matematica della prospettiva	275
Capitolo 16. * Trasformazioni prospettiche	277
16.1. Prospettiva centrale, proiezione standard	278
16.2. Proiezione prospettica ortogonale (o ortografica)	287
16.3. Un'unica matrice per prospettiva centrale e ortogonale	289
16.4. Forma matriciale generale della prospettiva centrale	291
16.5. Punti di fuga della prospettiva centrale	297
16.6. Prospettive parallele	320
16.6.1. Proiezione parallela	321
16.6.2. Proiezione obliqua	324
Proiezioni cavaliera e cabinet	329
16.6.3. Vari tipi di proiezioni ortogonali	331
16.7. Applicazione: rimozione di linee nascoste in grafici 3D	331
16.7.1. Rimozione di linee prospetticamente nascoste	332
16.7.2. Algoritmo veloce di rimozione di linee nascoste in grafici 3D: allineamento della griglia	332
16.7.3. Appendice: pseudocodice per la rimozione di linee nascoste	344
Parte 4. Appendice: norme, prodotti scalari, forme bilineari	363
Capitolo 17. * Appendice: norme, prodotti scalari e forme bilineari	365
17.1. Completezza e compattezza negli spazi metrici	365
17.2. Norme su spazi vettoriali reali	373
17.3. Prodotti scalari su spazi vettoriali reali	381
17.4. Norme e prodotti scalari su spazi vettoriali complessi	389
17.5. Basi ortogonali e proiezione ortogonale	397
Bibliografia	405

Parte 1

Algebra lineare

CAPITOLO 1

Introduzione

Questo capitolo preliminare stabilisce la terminologia sui tipi di numeri ed operazioni aritmetiche che usiamo in seguito. Per questi concetti è sufficiente fare appello all'intuizione del lettore, che, nel caso ritenga utile una maggior precisione, può leggere l'Appendice 1.3 di questo Capitolo oppure far riferimento a [1] (un riferimento bibliografico utile per tutto il presente libro).

1.1. Numeri interi, razionali e reali

Con il termine *algebra* si intende il calcolo, e metodi di calcolo, di numeri naturali, interi, razionali, reali e complessi.

Il primo concetto basilare è quello di *insieme*, che intuitivamente è una collezione di elementi descritti da una qualche proprietà. Per esempio l'insieme dei miei vestiti blu è formato dai vestiti che possiedo nel mio armadio e che sono di colore blu.

Osserviamo che il numero di elementi di un insieme può essere finito (come i miei vestiti blu) o infinito (come gli insiemi di numeri su cui lavoriamo).

L'insieme dei numeri *naturali*, che si indica con $\mathbb{N}$, è formato dai numeri che si possono contare:

$$0, 1, 2, 3, 4, 5, \dots$$

In questo corso elementare si può forse procedere contando che i lettori abbiano già una intuizione precisa dei numeri naturali (con l'operazione di somma), dei numeri interi, in cui c'è anche l'operazione di differenza, e dei numeri razionali, nei quali si introducono anche le operazioni di prodotto e di quoziente. In ogni caso, presentiamo un cenno delle definizioni e costruzioni rigorose nell'Appendice 1.3.

L'insieme dei numeri *naturali*, che si indica con $\mathbb{N}$, è formato da numeri che si possono ordinare consecutivamente:

$$0, 1, 2, 3, 4, 5, \dots$$

I numeri *interi*, il cui insieme è denotato con $\mathbb{Z}$, sono i numeri naturali e i loro *opposti*:

$$0, 1, -1, 2, -2, 3, -3, 4, -4, \dots$$

I numeri naturali e interi si possono sommare e moltiplicare. Si può fare la differenza di due numeri interi, ma non in generale di due numeri naturali: per esempio $1 - 2$ è l'intero -1 che non è un numero naturale.

Lo stesso problema si presenta con i numeri interi per la divisione: il numero

$$1 : 2 = \frac{1}{2}$$

non è intero. Si introducono allora i numeri *razionali*, il cui insieme si denota con $\mathbb{Q}$, che è formato dai quozienti di due numeri interi (con denominatore diverso da zero).

Prima di scrivere la definizione rigorosa di numero razionale, ricordiamo che in insiemistica si usano di solito i simboli di appartenenza $\in$ e di sottoinsieme $\subset$. Per esempio il fatto che l'insieme dei numeri naturali è sottoinsieme dei numeri interi (cioè ogni numero naturale è anche un intero), che a sua volta è sottoinsieme dei numeri razionali, si scrive:

$$\mathbb{N} \subset \mathbb{Z} \subset \mathbb{Q}.$$

Il fatto che $1/2$ non è un numero intero, ma è razionale, si scrive:

$$\frac{1}{2} \notin \mathbb{Z}, \quad \frac{1}{2} \in \mathbb{Q}.$$

DEFINIZIONE 1.1.1. L'insieme dei numeri razionali è

$$\mathbb{Q} = \left\{ \frac{p}{q} \mid p, q \in \mathbb{Z}, q > 0 \right\},$$

dove possiamo supporre p e q *primi tra loro*, cioè senza divisori comuni. Si può supporre $q > 0$, perché se fosse $q < 0$, allora si potrebbe moltiplicare numeratore e denominatore per -1 , ottenendo così una frazione con denominatore positivo.

Per esempio:

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

Con i numeri razionali si può fare la divisione (per un numero diverso da zero):

$$\frac{p}{q} : \frac{m}{n} = \frac{p}{q} \cdot \frac{n}{m} = \frac{pn}{qm}$$

dove il denominatore qm è diverso da zero, perché sia q che m sono diversi da zero.

È conveniente rappresentare i numeri naturali, interi e razionali (e come vedremo anche quelli reali) su una retta:

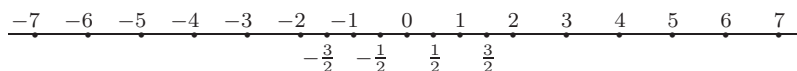


FIGURA 1. Rappresentazione grafica dei numeri naturali, interi e razionali

Ma anche i numeri razionali non sono sufficienti per misurare gli oggetti che troviamo in natura. Consideriamo per esempio la diagonale di un quadrato, come in Figura 6.

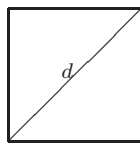


FIGURA 2. La diagonale d di un quadrato

Se il lato è lungo 1 cm, allora per il teorema di Pitagora la diagonale è:

$$d = \sqrt{2} = 0,4142 \dots \text{ cm},$$

cioè d è un numero tale che $d^2 = 2$.

PROPOSIZIONE 1.1.2. *Il numero $d = \sqrt{2}$ non è razionale.*

DIMOSTRAZIONE. Supponiamo per assurdo che $\sqrt{2} = p/q$, con p e q numeri interi primi tra loro. Elevando entrambi i membri al quadrato si ottiene:

$$2 = \frac{p^2}{q^2}, \quad \text{cioè} \quad p^2 = 2q^2.$$

Ne segue che p deve essere un numero pari, quindi $p = 2m$, per un certo intero m . Sostituendo p con $2m$ nella formula precedente, si trova che:

$$p^2 = 4m^2 = 2q^2, \quad \text{perciò} \quad 2m^2 = q^2.$$

Ma allora anche q deve essere un numero pari, divisibile per 2, e questo contraddice l'ipotesi che p e q non abbiano fattori in comune. Abbiamo trovato così una contraddizione, quindi l'ipotesi che $\sqrt{2}$ fosse un numero razionale non può essere vera. $\square$

Si considerano allora anche i numeri *reali*, il cui insieme si indica con $\mathbb{R}$, che si possono scrivere come numeri interi seguiti, dopo la virgola, da infinite cifre decimali. I numeri reali si possono rappresentare sulla stessa retta di Figura 7 e possiamo immaginare che ogni punto della retta corrisponde ad un numero reale.

Possiamo pensare anche i punti di un *piano* come elementi di un insieme su cui poter fare operazioni, come per esempio la somma. Infatti possiamo associare ad ogni punto del piano, determinato da una coppia di numeri reali, un *vettore*, di cui i numeri reali sono le coordinate. I vettori si possono sommare e moltiplicare per uno scalare e formano così uno *spazio vettoriale*, che definiremo dettagliatamente nel prossimo capitolo 2.

Più avanti vedremo anche che ad ogni coppia di numeri reali (ovvero, ad ogni punto del piano) possono essere associati i cosiddetti numeri *complessi*, il cui insieme si denota con $\mathbb{C}$. Il numero complesso associato alla coppia (a, b) si denota con

$$a + bi$$

dove i è l'unità *immaginaria*, che è un numero complesso tale che

$$i^2 = -1, \quad \text{ovvero} \quad i = \sqrt{-1}. \quad (1.1.1)$$

Noi però ci occuperemo principalmente dei numeri reali e non di quelli complessi. Fra gli insiemi di numeri che abbiamo introdotto valgono le seguenti relazione di inclusione:

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

All'interno del corso però i numeri complessi avranno un ruolo abbastanza marginale e quasi sempre avremo a che fare solo con i numeri reali.

Abbiamo detto che con il termine “algebra” si intende il calcolo di operazioni quali la somma e il prodotto di numeri.

Con il termine “algebra lineare”, che è il contenuto di questo corso, si intende lo studio e la risoluzione dei *sistemi di equazioni lineari*, come per esempio:

$$\begin{cases} 2x + 3y = 1 \\ -x + 5y = -2 \end{cases} \quad (1.1.2)$$

cioè di un numero finito di equazioni in cui compaiono *variabili lineari*, ovvero le incognite compaiono nelle espressioni solo con *grado uno* (il grado è l'esponente dell'incognita, che, essendo sempre e solo 1, di solito si traslascia).

Lo strumento per risolvere tali sistemi saranno i *vettori* e le *matrici*. Per esempio il sistema lineare di due equazioni (1.1.2) verrà scritto nel modo seguente:

$$\begin{pmatrix} 2 & 3 \\ -1 & 5 \end{pmatrix} \cdot \begin{pmatrix} x \\ y \end{pmatrix} = \begin{pmatrix} 1 \\ -2 \end{pmatrix},$$

dove il secondo membro è il vettore dei termini noti e al primo membro

$$\begin{pmatrix} 2 & 3 \\ -1 & 5 \end{pmatrix}$$

è la matrice quadrata di ordine 2×2 dei coefficienti delle incognite. Per risolvere tale sistema lineare si useranno allora le proprietà dei vettori e delle matrici che vedremo nel seguito.

1.2. Alcune notazioni

Gli intervalli di numeri reali vengono denotati nel modo seguente:

$$[1, 2] = \{ x \in \mathbb{R} \mid 1 \leq x \leq 2 \},$$

$$[1, 2) = \{ x \in \mathbb{R} \mid 1 \leq x < 2 \},$$

$$(1, +\infty) = \{ x \in \mathbb{R} \mid 1 < x \}.$$

Ricordiamo che possiamo descrivere un insieme elencando tutti gli elementi o indicando una proprietà. Per esempio indichiamo l'insieme dei numeri naturali dispari così:

$$\{ n \in \mathbb{N} \mid n \text{ dispari} \},$$

oppure nel modo seguente:

$$\{ n \in \mathbb{N} \mid n = 2m + 1, \text{ con } m \in \mathbb{N} \}.$$

Una costruzione che useremo spesso è la seguente:

DEFINIZIONE 1.2.1. Consideriamo due insiemi S e T . Il *prodotto cartesiano* di S e T è:

$$S \times T = \{ (s, t) \mid s \in S, t \in T \}.$$

ESEMPIO 1.2.2. Il prodotto cartesiano che considereremo molto spesso è $\mathbb{R} \times \mathbb{R} = \mathbb{R}^2$, detto *piano reale*, che è l'insieme formato dalle coppie (a, b) di numeri reali. Vedremo più avanti che è conveniente scrivere queste coppie in verticale, cioè

$$\begin{pmatrix} a \\ b \end{pmatrix}$$

invece che in orizzontale. Come si può fare la somma di due numeri reali, così si possono sommare le coppie di numeri reali, facendo la somma componente per componente: se $\begin{pmatrix} a_1 \\ a_2 \end{pmatrix}$ e $\begin{pmatrix} b_1 \\ b_2 \end{pmatrix}$ sono due coppie qualsiasi di numeri reali, allora si definisce

$$\begin{pmatrix} a_1 \\ a_2 \end{pmatrix} + \begin{pmatrix} b_1 \\ b_2 \end{pmatrix} = \begin{pmatrix} a_1 + b_1 \\ a_2 + b_2 \end{pmatrix}.$$

□

ESEMPIO 1.2.3. Consideriamo i punti $\begin{pmatrix} 1 \\ -1 \end{pmatrix}$ e $\begin{pmatrix} 2 \\ 0 \end{pmatrix}$ di $\mathbb{R}^2$. Allora la loro somma è:

$$\begin{pmatrix} 1 \\ -1 \end{pmatrix} + \begin{pmatrix} 2 \\ 0 \end{pmatrix} = \begin{pmatrix} 3 \\ -1 \end{pmatrix}.$$

□

Come già sapete, si è soliti rappresentare gli elementi di $\mathbb{R}^2$ usando gli assi cartesiani. Nel prossimo capitolo 2 torneremo subito su questo esempio.

Più in generale si può considerare il prodotto cartesiano di $\mathbb{R}$ per se stesso un numero finito n di volte:

$$\underbrace{\mathbb{R} \times \mathbb{R} \times \cdots \times \mathbb{R}}_{n \text{ fattori}} = \mathbb{R}^n$$

che è un insieme i cui elementi sono le n -uple di numeri reali. Per $n = 3$, si ottiene $\mathbb{R}^3$ che è chiamato *spazio reale*. Per ogni n , si potrà associare ad una n -pla di numeri reali un *vettore*, ottenendo così uno spazio vettoriale di *dimensione* n . In questi spazi vettoriali si potranno fare le stesse operazioni di somma e moltiplicazioni per scalari come nel piano.

1.3. Appendice: numeri naturali, interi, razionali e reali

Partiamo da un solo concetto *primitivo*, cioè non definito, quello di insieme. Un sottoinsieme E_- di un insieme E_+ è un insieme i cui elementi sono tutti contenuti in E_+ ; se E_- non coincide con E_+ si dice che è un *sottoinsieme proprio*.

Definizione assiomatica dei numeri naturali: assiomi di Peano. I numeri naturali sono un insieme $\mathbb{N}$ che verifica i cinque assiomi seguenti, introdotti da Peano:

- P1. Esiste un elemento in $\mathbb{N}$ che denotiamo con 1.
- P2. Per ogni elemento n di $\mathbb{N}$ esiste un altro elemento n' di $\mathbb{N}$ che designiamo come il *successivo* di n .
- P3. 1 non è il successivo di nessun altro numero naturale.
- P4. Se $n \neq m$ allora $n' \neq m'$.
- P5. (**Assioma di induzione.**) Se una proprietà è verificata dal numero 1, ed è tale che, se è verificata da n allora è verificata dal suo successivo n' , allora essa è verificata da tutti gli interi.

Grazie all'assioma di induzione, possiamo definire due operazioni su $\mathbb{N}$, come segue.

DEFINIZIONE 1.3.1. (*L'operazione di somma su $\mathbb{N}$.*) Definiamo $n + 1 = n'$, e poi induttivamente $n + m' = (n + m)'$. Per l'assioma di induzione, questo definisce la somma per ogni coppia di naturali.

DEFINIZIONE 1.3.2. (*L'operazione di moltiplicazione su $\mathbb{N}$.*) Definiamo $n \cdot 1 = n$, e poi induttivamente $n \cdot m' = (n \cdot m) + 1$ (in altre parole, $n \cdot (m + 1) = (n \cdot m) + 1$). Per l'assioma di induzione, questo definisce la moltiplicazione per ogni coppia di naturali.

È facile verificare che valgono le proprietà seguenti.

A1. (*Proprietà associativa della somma:* $m + (n + k) = (m + n) + k$ per ogni $m, n, k \in \mathbb{N}$.)

A2. (*Proprietà commutativa della somma:* $m + n = n + m$ per ogni $m, n \in \mathbb{N}$.)

M1. (*Proprietà associativa della moltiplicazione:* $m \cdot (n \cdot k) = (m \cdot n) \cdot k$ per ogni $m, n, k \in \mathbb{N}$.)

M2. (*Proprietà commutativa della moltiplicazione:* $m \cdot n = n \cdot m$ per ogni $m, n \in \mathbb{N}$.)

M3. (*Esistenza dell'elemento neutro per la moltiplicazione.*) $n \cdot 1 = n$ per ogni $n \in \mathbb{N}$.

D. (*Proprietà distributiva della somma rispetto alla moltiplicazione:* $n \cdot (m + k) = (n \cdot m) + (n \cdot k)$ per ogni $m, n, k \in \mathbb{N}$.)

Vogliamo costruire in maniera concreta una copia degli interi. Per questo scopo dobbiamo esporre vari argomenti preliminari.

DEFINIZIONE 1.3.3. (*Prodotto cartesiano.*) Dati due insiemi A e B , il *prodotto cartesiano* $A \times B$ è l'insieme delle coppie ordinate (a, b) : $a \in A, b \in B$.

Si noti che il prodotto cartesiano $\mathbb{R} \times \mathbb{R}$ viene in tal modo visualizzato geometricamente come un piano: ciascun punto del piano corrisponde infatti, in modo unico, ad una coppia di coordinate. Analogamente, $\mathbb{Z} \times \mathbb{Z}$ è il reticolo dei punti con entrambe le coordinate intere. Prendendo il prodotto cartesiano di tre copie di $\mathbb{R}$ si ottiene lo spazio tridimensionale, ed analogamente per il reticolo tridimensionale degli interi.

NOTA 1.3.4. È ovvio dalla definizione che il prodotto cartesiano è associativo, ma non commutativo. $\square$

DEFINIZIONE 1.3.5. (**Relazione di equivalenza.**) Dato un insieme E , un sottoinsieme $R \subseteq E \times E$ si dice una *relazione di equivalenza* se valgono le seguenti tre proprietà:

- (i) **riflessiva:** $(a, a) \in R$ per ogni $a \in E$
- (ii) **simmetrica:** $(a, b) \in R$ se $(b, a) \in R$
- (iii) **transitiva:** se $(a, b) \in R$ e $(b, c) \in R$ allora anche $(a, c) \in R$

Se $(a, b) \in R$ scriviamo $a \sim b$. Con questa notazione le tre proprietà sopra elencate diventano:

- (i) $a \sim a \quad \forall a \in E$
- (ii) **simmetrica:** $a \sim b \Rightarrow b \sim a$
- (iii) **transitiva:** $a \sim b$ e $b \sim c \Rightarrow a \sim c$

I sottoinsiemi di E consistenti di elementi mutuamente equivalenti (questa è una definizione ben posta grazie alla proprietà transitiva!) si chiamano le *classi di equivalenza* di R . Gli elementi di una classe di equivalenza si chiamano *rappresentanti* della classe.

DEFINIZIONE 1.3.6. (**Funzioni.**) Una funzione (o *mappa*) f da un insieme A ad un insieme B è un sottoinsieme di $A \times B$ tale che non ci siano in f due coppie con lo stesso primo elemento.

Invece di scrivere $(a, b) \in f$ è consuetudine scrivere $f = f(a)$. Questa notazione ispira l'interpretazione abituale, che identifica la funzione f con una *legge* che ad ogni *variabile* $a \in A$ associa un *valore* $b \in B$ (ed uno solo: questo è ciò

‘o che vuol dire che non ci sono due coppie diverse con lo stesso primo elemento). In questa interpretazione, e rammentando la visualizzazione data poco sopra del prodotto cartesiano $A \times B$ come un *piano* generato dagli *assi* A e B , la nozione da noi introdotta di funzione si identifica con quella del suo grafico. Il fatto che non ci siano due coppie con lo stesso primo elemento significa allora che per ogni variabile la funzione ha al più un solo valore, cioè che il grafico non interseca più di una volta nessuna *retta verticale*.

ESEMPIO 1.3.7. (**Equipotenza.**) Introduciamo sulla famiglia di tutti gli insiemi una relazione di equivalenza, che chiamiamo *equipotenza*, nel modo seguente: due insiemi A e B sono equipotenti se esiste una funzione da A a B biunivoca, cioè tale che per ogni $a \in A$ esiste un $b \in B$ (e necessariamente uno solo) tale che $f(a) = b$, e per ogni $b \in B$ esiste un $a \in A$ (ed uno solo) tale che $b = f(a)$.

Lasciamo al lettore la facile verifica del fatto che valgono le proprietà della Definizione 1.3.5. $\square$

NOTA 1.3.8. Il lettore che ha già sviluppato un'intuizione sicura dell'insieme dei numeri interi, che fra poco definiremo in maniera rigorosa, può osservare che gli interi sono equipotenti ad alcuni loro sottoinsiemi propri! In effetti, chiamiamo $\mathbb{N}$ l'insieme degli interi e con $\mathbb{P}$ il sottoinsieme proprio degli interi pari: allora la mappa $n \mapsto 2n$ è una equipotenza fra $\mathbb{N}$ e $\mathbb{P}$. Questo fatto motiva la prossima definizione. $\square$

DEFINIZIONE 1.3.9. (**Insiemi finiti.**) Un insieme si dice *finito* se non è equipotente ad alcun suo sottoinsieme proprio.

Ora introduciamo la *costruzione dei naturali*.

DEFINIZIONE 1.3.10. I numeri naturali sono le classi di equivalenza degli insiemi finiti rispetto alla relazione di equipotenza. L'insieme dei numeri naturali si indica con $\mathbb{N}$. Il numero naturale n dato dalla classe di equivalenza di un insieme A si chiama la *cardinalità* di A .

Questa definizione chiarisce cosa sia l'insieme degli interi, ma non ci dice come costruire un singolo intero. Essa dice che i numeri naturali possono essere considerati come la proprietà che hanno in comune tutti gli insiemi con lo stesso numero di elementi, ma non avrebbe senso dire che un dato numero naturale n può essere considerato come la proprietà che hanno in comune tutti gli insiemi con lo stesso numero n di elementi: questa è una tautologia, perché per spigare cosa sia n fa riferimento a n stesso.

La costruzione rigorosa è la seguente. Il numero 0 è la classe di equivalenza dell'insieme vuoto (quello che non contiene elementi: si noti

che questa classe contiene il solo insieme vuoto). Un insieme si chiama un *singleton* se l'unico suo sottoinsieme proprio è l'insieme vuoto. Tutti i singleton sono nella stessa classe di equivalenza (esercizio!), e tale classe è il numero 1.

L'operazione di somma su $\mathbb{N}$. Consideriamo ora un insieme finito A e sia n la sua cardinalità: se B è un singleton disgiunto da A , diciamo che la cardinalità di $A \cup B$ è $n + 1$. Questo definisce la somma fra un intero ed il numero 1. C'è una naturale estensione di questo concetto a tutte le coppie di interi (cioè una funzione da $\mathbb{N} \times \mathbb{N} \mapsto \mathbb{N}$), ed è l'unica definizione di somma che rispetta la proprietà associativa: $n + 2 = n + (1 + 1) = (n + 1) + 1$, e così via. È facile ed intuitivo vedere che la somma verifica anche la proprietà commutativa. Per una esposizione più rigorosa di questa e delle successive costruzioni dovremmo utilizzare l'assioma di induzione: lasciamo la dimostrazione rigorosa come esercizio.

L'operazione di moltiplicazione su $\mathbb{N}$. La moltiplicazione è una funzione $\mathbb{N} \times \mathbb{N} \mapsto \mathbb{N}$ che verifica la proprietà distributiva del prodotto rispetto alla somma. la sua costruzione induttiva è la seguente: $2n = (1 + 1)n = n + n = n \cdot 2$, $3n = (1 + 1 + 1)n = n + n + n = n \cdot 3$, e così via. È facile vedere che la moltiplicazione è commutativa, associativa e distributiva rispetto alla somma.

DEFINIZIONE 1.3.11. (L'ordinamento di $\mathbb{N}$.) Diciamo che due numeri naturali n, m verificano la disuguaglianza $n < m$ se esiste $k \in \mathbb{N}$ tale che $m = n + k$.

È facile vedere che la somma ed il prodotto rispettano l'ordinamento su $\mathbb{N}$:

O1. Se $n < m$ allora $n + k < m + k$ per ogni $k \in \mathbb{N}$.

O2. Se $n < m$ allora $n \cdot k < m \cdot k$ per ogni $k \in \mathbb{N}$.

Più in generale introduciamo la seguente definizione.

DEFINIZIONE 1.3.12. (Relazione d'ordine.) Dato un insieme E , un sottoinsieme $R \subseteq E \times E$ si dice una *relazione d'ordine*, o un *ordinamento*, se valgono le seguenti tre proprietà:

(i) **riflessiva:** $(a, a) \in R$ per ogni $a \in E$

(ii) **antisimmetrica:** se $(a, b) \in R$ e $(b, a) \in R$ allora $a = b$

(iii) **transitiva:** se $(a, b) \in R$ e $(b, c) \in R$ allora anche $(a, c) \in R$

Se $(a, b) \in R$ scriviamo $a \preceq b$. Con questa notazione le tre proprietà sopra elencate diventano:

(i) $a \preceq a \forall a \in E$

(ii) $a \preceq b$ e $b \preceq a \Rightarrow a = b$

(iii) $a \preceq b$ e $b \preceq c \Rightarrow a \preceq c$

Ovviamente, l'ordinamento dei numeri naturali (ed in seguito degli interi, o dei razionali, o dei reali) è un esempio.

Abbiamo definito la somma di due numeri naturali. Non possiamo invece definire la differenza fra due naturali m e n altro che se $n \leq m$. Sotto questa ipotesi, segue direttamente dalla definizione di ordinamento che esiste un naturale k (ed uno solo) tale che $n + k = m$: il numero k si chiama la differenza fra m e n , e si scrive $k = m - n$. È anche facile verificare che la differenza, nei casi in cui esiste, verifica le proprietà associativa, commutativa e distributiva rispetto alla moltiplicazione. Ma come osservato, essa non esiste per tutte le coppie in $\mathbb{N}$.

Introduciamo una estensione $\mathbb{Z}$ di $\mathbb{N}$ a cui sia possibile estendere l'operazione di differenza in modo che si applichi a tutte le coppie: $\mathbb{Z}$ si chiama l'insieme degli interi (o interi con segno, o interi relativi).

Visualizzazione geometrica dell'operazione di differenza.

Per prima cosa diamo una visualizzazione geometrica della differenza, in termini del prodotto cartesiano $\mathbb{N} \times \mathbb{N}$, che come osservato è un reticolo. Fissato $k \in \mathbb{N}$, per quali coppie di naturali (n, m) , cioè per quali punti del reticolo, si ha che $n - m = k$? Esattamente per i punti che giacciono sulla (semi)retta di pendenza 1 che interseca l'asse orizzontale al punto k . Anzi, richiedere che la differenza $n - m$ sia la stessa per diverse coppie di naturali è una relazione di equivalenza, le cui classi di equivalenza sono i tali rette. Allora è immediato immergere $\mathbb{N}$ in questo insieme di rette: ogni k corrisponde alla retta a pendenza 1 che passa per $(0, k)$. L'operazione di somma si trasporta a queste classi di equivalenza. Infatti possiamo definire la somma di due rette come la retta che si ottiene sommando come vettori i loro punti, cioè l'insieme di tutte le coppie $(n_1 + n_2, m_1 + m_2)$ dove (n_1, m_1) appartiene alla prima

retta e (n_2, m_2) alla seconda. Quest'insieme è ancora una retta perché, se $m_1 - n_1 = k_1$ e $m_2 - n_2 = k_2$, allora $(m_1 + m_2) - (n_1 + n_2) = k_1 + k_2$ non dipende dalla scelta del singolo punto sulla prima o sulla seconda retta, ma solo da k_1 e k_2 . In altre parole, la differenza (ed analogamente la somma) sono invarianti sulle classi di equivalenza date dalle rette, cioè indipendenti dai rappresentanti della classe di equivalenza (diciamo che sono *ben definite* sulle classi).

NOTA 1.3.13. Una presentazione più elegante, ma meno trasparente, avrebbe evitato l'uso del segno meno che abbiamo utilizzato nello scrivere la relazione di equivalenza $(n_1, m_1) \sim (n_2, m_2) \Leftrightarrow m_1 - n_1 = m_2 - n_2$. Avremmo infatti potuto, equivalentemente, scrivere $(n_1, m_1) \sim (n_2, m_2) \Leftrightarrow m_1 + n_2 = n_1 + m_2$. Le classi di equivalenza della relazione si scrivono anch'esse senza ricorrere al segno meno, in questo modo: la classe di equivalenza del numero naturale k è la semiretta $\{k + n, n\}$.

□

Estensione dai naturali $\mathbb{N}$ agli interi $\mathbb{Z}$. A questo punto è facile estendere $\mathbb{N}$ ad un insieme $\mathbb{Z}$ dove la differenza sia definita per ogni coppia di elementi. Abbiamo visto che $\mathbb{N}$ è realizzabile come l'insieme delle (semi)rette del reticolo $V \times \mathbb{N}$ a pendenza 1 che tagliano le ascisse in punti non a sinistra di 0, cioè delle semirette che giacciono nel primo quadrante. Consideriamo l'insieme, più grande, di tutte le rette, o semirette, semirette a pendenza 1. e operazioni di somma e di differenza si estendono come prima, e come prima dipendono solo dalle rette ma non dalla scelta dei singoli punti in esse. Questo insieme di rette lo chiamiamo $\mathbb{Z}$. Se una semiretta non giace nel primo quadrante, essa interseca l'asse orizzontale in un punto a sinistra dell'origine. Il primo di tali punti lo indichiamo con -1 , il secondo con -2 e così via: osserviamo che i punti (n, m) della retta che passa per $-k$ sono quelli per cui $n - m = k$. Osserviamo anche che le rette passanti per $-k$ e per k hanno per somma la retta passante per 0, perché se $n_1 - m_1 = k$ e $m_2 - n_2 = k$ allora $(m_1 + m_2) - (n_1 + n_2) = 0$. In questo stesso modo vediamo che tutte le proprietà della somma si estendono a $\mathbb{Z}$.

Anche la moltiplicazione è ben definita sulle classi di equivalenza

date dalla semirette: osserviamo che, se $m_1 - n_1 = k_1$ e $m_2 - n_2 = k_2$, allora $k_1 k_2 = (m_1 - n_1)(m_2 - n_2) = m_1 m_2 + n_1 n_2 - m_1 n_2 - n_1 m_2$. Quindi definiamo la moltiplicazione delle due rette di pendenza 1 che contengono rispettivamente i punti (n_1, m_1) e (n_2, m_2) come la retta che contiene il punto $(m_1 m_2 + n_1 n_2, m_1 n_2 + n_1 m_2)$. Si verifica facilmente che questa definizione di moltiplicazione è ben posta: essa dipende solo dalle classi di equivalenza, le rette, e non dalla scelta dei rappresentanti usati per formularla (esercizio).

Ora è facile, anche se lungo, verificare che le proprietà: associativa e commutativa della somma e della moltiplicazione, l'esistenza dell'elemento neutro della moltiplicazione, la proprietà distributiva della somma rispetto alla moltiplicazione e la prima proprietà di ordinamento **O1** valgono in $\mathbb{Z}$; invece la seconda proprietà di ordinamento **O2** si estende in questo modo:

O2. Se $n < m$ allora $n \cdot k < m \cdot k$ per ogni $k > 0 \in \mathbb{Z}$, $m \cdot k < n \cdot k$ per ogni $k < 0 \in \mathbb{Z}$, e $n \cdot 0 = 0$ per ogni $n \in \mathbb{Z}$.

In particolare, segue dalla proprietà **O2** che 0 *non ha divisori*: se $nm = 0$ allora uno fra n e m deve essere zero, perché altrimenti il loro prodotto sarebbe o strettamente positivo o strettamente negativo.

Inoltre, è facile vedere che valgono altre due proprietà della somma:

A3. (Esistenza dell'elemento neutro per la somma.) Esiste un numero $0 \in \mathbb{Z}$ tale che $n + 0 = n$ per ogni $n \in \mathbb{Z}$.

A4. (Esistenza dell'opposto.) Per ogni $n \in \mathbb{Z}$ esiste un $m \in \mathbb{Z}$ tale che $n + m = 0$ (si scrive $m = -n$).

Infatti, l'elemento 0 corrisponde alla classe di equivalenza della retta bisettrice: $\{(n, n), n \in \mathbb{Z}\}$, perché se (k, l) sta in un'altra retta, allora $(k + n, l + n) \sim (k, l)$ perché $(k + n) + l = (l + n) + k$ per le proprietà associativa e commutativa della somma. Questo prova la proprietà **A3**. Per provare **A4**, data una coppia (n, m) corrispondente a qualche $k \in \mathbb{Z}$, una coppia appartenente alla retta corrispondente a $-k$ è (m, n) , dal momento che $(n, m) + (m, n) = (n + m, n + m)$ e l'ultima coppia appartiene alla semiretta associata a 0.

NOTA 1.3.14. L'elemento 0 è unico. Infatti, se ce ne fossero due (chiamiamoli 0_1 e 0_2), allora si avrebbe $0_1 = 0_1 + 0_2 = 0_2 + 0_1 = 0_2$, per la

proprietà commutativa della somma. $\square$

Analogamente, l'opposto è unico: più in generale, vale la legge di cancellazione seguente.

COROLLARIO 1.3.15. *Per ogni $n, m \in \mathbb{Z}$ esiste un unico $k \in \mathbb{Z}$ tale che $n + k = m$.*

DIMOSTRAZIONE. È facile verificare che k esiste: basta porre $k = m - n$, la verifica è immediata a partire dalle quattro proprietà della somma. Se ci fossero due diversi k_1 e k_2 con la proprietà dell'enunciato, avremmo $n + k_1 = m = n + k_2$, e sommando $-n$ ad entrambi i membri si trova, di nuovo dalle proprietà associative e di esistenza degli opposti, che $k_1 = k_2$. $\square$

Un insieme con una operazione di somma che verifica le proprietà **A1**, **A2**, **A3** ed **A4** si chiama un *gruppo commutativo*. La regola di cancellazione vale, per questo argomento, in tutti i gruppi commutativi.

Usando le proprietà che abbiamo introdotto si possono provare i risultati cruciali dell'aritmetica: ad esempio l'esistenza e l'unicità della scomposizione in fattori primi (*Teorema fondamentale dell'aritmetica*), le proprietà del massimo comun divisore e del minimo comune multiplo, il teorema della divisione con resto (per ogni $m, n \in \mathbb{N}$ esistono $q \in \mathbb{N}$ e $0 \leq r < m$ tali che $n = qm + r$, e sono unici), l'algoritmo euclideo per la determinazione del massimo comun divisore tramite applicazioni iterate della divisione con resto, e le proprietà delle congruenze modulo un naturale N .

L'operazione di divisione ed i numeri razionali. Abbiamo introdotto in $\mathbb{Z}$ una moltiplicazione, ma non l'operazione inversa, la divisione. È chiaro che la struttura ordinata e sequenziale degli interi, derivante dall'assioma di Peano P5, impedisce di trovare in $\mathbb{Z}$ il reciproco di ogni elemento, e quindi di definire la divisione: infatti i reciproci dei numeri maggiori di 1 dovrebbero essere compresi fra 0 e 1, ma abbiamo visto che non esistono interi con questa proprietà.

Per definire la divisione dobbiamo estendere gli interi ad un insieme più grande, quello dei razionali $\mathbb{Q}$. Vogliamo definire il reciproco di qualunque numero intero non nullo, cioè vogliamo definire le frazioni

$\frac{m}{n}$ se $n \neq 0$. A questo fine definiamo una nuova relazione di equivalenza sulle coppie (n, m) con $m \neq 0$: due tali coppie sono equivalenti *per dilatazione* se esiste $k \in \mathbb{Z}$, $k \neq 0$, tali che $(n_2, m_2) = (kn_1, km_1)$. Preferiamo riformulare questa relazione nel seguente modo equivalente: $(n_1, m_1) \approx (n_2, m_2)$ se $n_1 m_2 = m_1 n_2$. È facile verificare che si tratta di una relazione di equivalenza. Le classi di equivalenza sono i punti del reticolo $\mathbb{Z} \times \mathbb{Z}$ allineati radialmente rispetto all'origine: in altre parole, una classe di equivalenza corrisponde ai punti del reticolo sulla stessa retta passante per l'origine. Osserviamo che in ogni classe di equivalenza esiste una coppia (n, m) di distanza minima dall'origine (questa coppia è unica se si sceglie il suo secondo elemento $m > 0$): essa è caratterizzata dalla proprietà che n e m sono relativamente primi, cioè senza fattori comuni. Quindi, quando si scrivono i razionali come frazioni, questa rappresentazione senza fattori primi è unica (assumendo il denominatore positivo); essa si chiama *rappresentazione ridotta*.

ESERCIZIO 1.3.16. Sebbene gli interi si immergano nei razionali come sottoinsieme, nondimeno gli interi ed i razionali sono equipotenti, ossia esiste fra questi due insiemi una corrispondenza biunivoca. Trovarne una. $\square$

Le operazioni di somma e di moltiplicazione vengono definite, come sempre, sui rappresentanti delle classi di equivalenza. Definiamo queste operazioni in modo che ricalchino le proprietà consuete che vogliamo avere sulle frazioni:

- **Somma in $\mathbb{Q}$** $\equiv \mathbb{N} \times \mathbb{N} \setminus \{0\}$: $(n_1, m_1) \pm (n_2, m_2) \approx (n_1 m_2 \pm m_1 n_2, n_1 m_2)$
- **Moltiplicazione in $\mathbb{Q}$** : $(n_1, m_1) \cdot (n_2, m_2) \approx (n_1 n_2, m_1 m_2)$

È elementare verificare che i rappresentanti del numero 1 sono le coppie (n, n) con $n \neq 0$ (cioè la bisettrice), e quelli del numero 0 sono le coppie $0, n$ con $n \neq 0$ (cioè l'asse orizzontale). È altrettanto immediato verificare che tutte le proprietà aritmetiche di $\mathbb{Z}$ si estendono a $\mathbb{Q}$. Inoltre, ogni razionale non nullo ha un reciproco:

M4. **Esistenza del reciproco in $\mathbb{Q}$.** Per ogni $n, m \in \mathbb{Q}$ con $m \neq 0$ esiste uno ed un solo $k \in \mathbb{Q}$ tale che $n = mk$.

Per dimostrare M4 basta verificare che, se n è rappresentato da (j_1, l_1) e m da (j_2, l_2) , allora il razionale k rappresentato da $(j_1 l_2, j_2 l_1)$ verifica $n = mk$. L'unicità è ovvia: se $mk_1 = n = mk_2$ allora $m \cdot (k_1 - k_2) = 0$, ma l'unico divisore di 0 è 0 e quindi $k_1 - k_2 = 0$.

Un insieme dotato delle operazioni di somma e moltiplicazione con le proprietà elencate più sopra e tale che ogni elemento non nullo ha un reciproco si chiama un *campo*.

Immersione isomorfa degli interi nei razionali. Gli interi sono stati costruiti a partire dal reticolo $\mathbb{N} \times \mathbb{N}$ dei naturali come classi di equivalenza della relazione di equivalenza delle rette a pendenza 1, mentre i razionali sono stati costruiti a partire dal reticolo $\mathbb{Z} \times \mathbb{Z}$ degli interi (che estende il precedente) come classi di equivalenza della relazione di equivalenza della dilatazione. Le due relazioni sono diverse, e quindi non è ovvio che $\mathbb{Z}$ si immerga in $\mathbb{Q}$ preservando le operazioni aritmetiche: se questo avviene, l'immersione ϕ si chiama un *isomorfismo* di $\mathbb{Z}$ su $\phi(\mathbb{Z})$.

Esiste una naturale immersione isomorfa. Questa immersione è data da $\phi(n) = (n, 1)$, $n \in \mathbb{Z}$. È chiaro che ϕ è iniettiva, e si verifica immediatamente che $\phi(n_1 \pm n_2) = \phi(n_1) \pm \phi(n_2)$.

L'immersione preserva l'ordinamento. La relazione d'ordine di $\mathbb{Z}$ si estende a $\mathbb{Q}$. L'ordinamento su $\mathbb{Q}$ si definisce così: se $r \in \mathbb{Q}$ è rappresentato dalla coppia (n, m) ($m \neq 0$), allora si dice che $r > 0$ se n e m sono entrambi positivi o entrambi negativi, e invece $r < 0$ se i segni di n e m sono opposti (ovviamente c'è solo un caso residuo, quello in cui $n = 0$, nel qual caso si ha $r = 0$). È facile vedere che due numeri interi verificano la relazione d'ordine $n < m$ rispetto all'ordinamento di $\mathbb{Z}$ se e solo se le loro immagini in $\mathbb{Q}$ verificano la stessa relazione rispetto all'ordinamento di $\mathbb{Q}$: $\phi(n) < \phi(m)$. Quindi l'immersione di $\mathbb{Z}$ in $\mathbb{Q}$ preserva l'ordinamento.

Risoluzione di equazioni quadratiche e sezioni di Dedekind in $\mathbb{Q}$. Abbiamo dimostrato nella Proposizione 1.1.2 che l'equazione $x^2 = 2$ non ammette soluzioni in $\mathbb{Q}$. Per risolverla dobbiamo estendere il campo $\mathbb{Q}$ dei razionali ad un campo più grande, quello dei numeri

reali, che si denota con $\mathbb{R}$. Vediamo come. Poiché la funzione x^2 è strettamente crescente per $x \geq 0$, e $1^2 = 1 < 2 < 2^2 = 4$, in qualsiasi estensione di $\mathbb{Q}$ che preservi l'ordinamento deve valere $1 < \sqrt{2} < 2$. Ora consideriamo le frazioni $\frac{11}{10}, \frac{12}{10}, \dots, \frac{19}{10}$. Poiché $11^2 < 12^2 < 13^2 < 14^2 = 196 < 200 < 15^2 = 225$, deve valere $1.4 = \frac{14}{10} < 2 < \frac{15}{10} = 1.5$. Continuando così abbiamo $1.41 < 2 < 1.42$, $1.414 < 2 < 1.415$, e così via. In questo modo, in ogni campo che estende $\mathbb{Q}$ preservandone l'ordinamento e nel quale esiste la radice quadrata di 2, costruiamo una successione crescente di razionali minori di $\sqrt{2}$ ed una decrescente di razionali maggiori di $\sqrt{2}$ in cui la differenza fra il maggiorante ed il minorante di indice n è inferiore a 10^{-n} . Ora dobbiamo fare una ulteriore ipotesi, che a questo punto è chiaramente equivalente alla risolubilità dell'equazione $x^2 = 2$: **Proprietà archimedeas**: nell'estensione di $\mathbb{Q}$ non esiste alcun numero positivo minore di 10^{-n} simultaneamente per tutti gli n . Oppure, equivalentemente, riformuliamo l'ipotesi nel modo seguente:

Assioma delle sezioni di Dedekind. Chiamiamo *classe maggiorante* in $\mathbb{Q}$ un sottoinsieme proprio J_+ di $\mathbb{Q}$ tale che, per ogni $r \in J_+$, tutti i $q \in \mathbb{Q}$ con $q > r$ appartengono a J_+ . In maniera simmetrica si definiscono le classi minoranti. Il complementare di J_+ è una classe minorante che indichiamo con J_- . Si osservi che abbiamo appena visto che, se $J_+ = \{r \in \mathbb{Q} : r^2 > 2\}$, allora J_+ non ha minimo in $\mathbb{Q}$ e J_- non ha massimo, ma naturalmente ogni elemento di J_- è minore di tutti gli elementi di J_+ ed ogni elemento di J_+ maggiore di tutti gli elementi di J_- . Una coppia di insiemi J_+, J_- in cui ciascun elemento del primo insieme maggiore di ogni elemento del secondo e viceversa si chiama una *sezione di Dedekind*. Chiamiamo *insieme $\mathbb{R}$ dei numeri reali* l'insieme delle sezioni di Dedekind dei numeri razionali. Estendiamo a $\mathbb{R}$ l'ordinamento di $\mathbb{Q}$ nel modo seguente: il numero reale corrispondente ad una sezione di Dedekind J_+, J_- è minore o uguale di ogni razionale in J_- e maggiore o uguale di ogni razionale in J_+ (esso si chiama *l'elemento separatore* della sezione). È chiaro che l'immersione di $\mathbb{Q}$ in $\mathbb{R}$ rispetta l'ordinamento. Data una sezione di Dedekind J_+, J_- , un numero reale minore o uguale di ogni $r \in J_+$ e maggiore o uguale di ogni $s \in J_-$ si chiama un elemento separatore della sezione. L'assioma di Dedekind

dice che per ogni sezione esiste un unico elemento separatore in $\mathbb{R}$. (Per esercizio si dimostri che, in particolare, J_+ ha minimo in $\mathbb{R}$ oppure J_- ha massimo). Quindi i numeri reali sono gli elementi separatori delle sezioni dei razionali.

Resta solo da estendere a $\mathbb{R}$ le operazioni di somma e di moltiplicazione. La somma di due reali x_1 (associato alla sezione J_+, J_-) e x_2 (associato alla sezione I_+, I_-) è l'elemento separatore della sezione $J_+ + I_+, J_- + I_-$ (qui l'insieme somma $J_+ + I_+$ è definito come $\{x + y : x \in J_+, y \in I_+\}$). Abbiamo visto, dall'esempio di $\sqrt{2}$, che una sezione di Dedekind (cioè un numero reale) corrisponde a due successioni di approssimanti razionali, in cui tutti i numeri della prima sono minoranti di tutti quelli della seconda e tutti quelli della seconda sono maggioranti di tutti quelli della prima, e con l'ulteriore proprietà che i numeri minoranti si avvicinano a quelli maggioranti a meno di una precisione arbitrariamente piccola. Allora è naturale definire la somma trasportandola da questi approssimanti razionali, i quali, avendo un numero finito di cifre decimali oppure uno sviluppo periodico, ci permettono di calcolare la somma colonna per colonna (con eventuale riporto). Quindi la somma di due numeri reali si approssima numericamente troncando i due reali ad approssimanti razionali con lo stesso numero di cifre decimali, sommando tali razionali e poi facendo migliorare l'approssimazione col ripetere il calcolo con via via più cifre decimali. È facile verificare che la somma così definita su $\mathbb{R}$ ha tutte le proprietà precedentemente dimostrate su $\mathbb{Q}$.

La moltiplicazione si estende in modo analogo, prendendo la sezione prodotto $J_+ \cdot I_+, J_- \cdot I_-$. Qui però c'è una difficoltà tecnica che ci accenniamo. Se due numeri reali r_1 e r_2 sono entrambi non negativi, allora il prodotto $J_+ \cdot I_+$ delle loro classi maggioranti è la classe maggiorante del prodotto, e tutto procede come nel caso della somma. Ma se $r_1 < 0 < r_2$ allora la classe maggiorante è $J_+ \cdot I_-$, e se infine $r_1, r_2 < 0$ allora la classe maggiorante è $J_- \cdot I_-$. Nei vari casi la moltiplicazione si esegue sugli approssimanti razionali e poi si trasporta per approssimazioni successive, ma quali siano gli approssimanti dal di sotto e quali quelli dal di sopra dipende dalla casistica che abbiamo elencato.

CAPITOLO 2

Spazi vettoriali

Ad ogni punto $P = (x, y)$ del piano reale $\mathbb{R}^2$ possiamo associare un *vettore*, cioè un segmento orientato che parte dall'origine $O = (0, 0)$ del piano e arriva al punto fissato P , come in Figura 7. Denotiamo questo vettore con $\vec{OP}$.

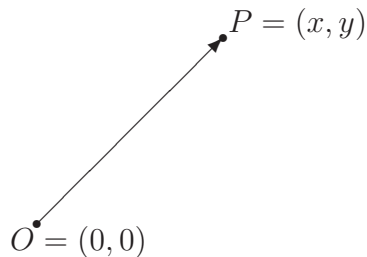


FIGURA 1. Vettore con punto iniziale O e punto finale P .

Ricordiamo che un vettore è determinato dalla sua *lunghezza* (o *modulo*), dalla sua *direzione* e dal suo *verso*. Come è ben noto, si possono sommare due vettori con la cosiddetta *regola del parallelogramma*: infatti si possono pensare i due vettori v, w come lati di un parallelogramma e la loro somma $v + w$ corrisponde alla diagonale del parallelogramma, come mostrato in Figura 8.

Dato un vettore, si può considerare il suo *opposto*, che è il vettore che ha lo stesso modulo, la stessa direzione di quello dato, ma verso opposto. Inoltre si può moltiplicare un vettore per un numero reale $k > 0$, che è il vettore con la stessa direzione e lo stesso verso del vettore dato, ma con la lunghezza moltiplicata per k .

È conveniente identificare un vettore $\vec{OP}$ con il punto del piano P , cioè con la coppia di numeri reali (x, y) che sono le coordinate di P nel piano; quindi consideriamo le proprietà dei vettori in base alle loro

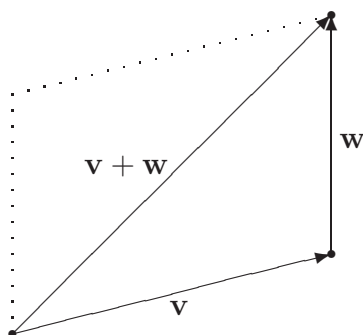


FIGURA 2. Somma di due vettori con la regola del parallelogramma.

coordinate, e non direttamente in base alla lunghezza, direzione e verso del vettore.

2.1. Definizione di spazio vettoriale

DEFINIZIONE 2.1.1. Uno **spazio vettoriale** è un insieme X provvisto di due operazioni:

- la **somma**, che associa ad ogni coppia di elementi di X un terzo elemento di X

$$x, y \in X \mapsto x + y \in X,$$

chiamato *somma* di x e y ;

- la **moltiplicazione per scalari**, che ad ogni elemento di X e ad ogni numero reale associa un altro elemento di X

$$x \in X, \lambda \in \mathbb{R} \mapsto \lambda \cdot x \in X,$$

detto *moltiplicazione di x per lo scalare λ* ;

che devono soddisfare alcune proprietà che vedremo in dettaglio fra poco. Talvolta si dice spazio *lineare*, al posto di spazio *vettoriale*.

ESEMPIO 2.1.2. Nel prodotto cartesiano $\mathbb{R} \times \mathbb{R} = \mathbb{R}^2$, detto *piano reale*, formato dalle coppie (a, b) di numeri reali, possiamo definire le due operazioni di somma e moltiplicazione per uno scalare $\lambda \in \mathbb{R}$ nel

seguinte modo:

$$\begin{pmatrix} a_1 \\ a_2 \end{pmatrix} + \begin{pmatrix} b_1 \\ b_2 \end{pmatrix} := \begin{pmatrix} a_1 + b_1 \\ a_2 + b_2 \end{pmatrix}, \quad (2.1.1)$$

$$\lambda \cdot \begin{pmatrix} a_1 \\ a_2 \end{pmatrix} := \begin{pmatrix} \lambda a_1 \\ \lambda a_2 \end{pmatrix}, \quad (2.1.2)$$

cioè componente per componente.

Associamo ad ogni elemento $P = (a, b)$ di $\mathbb{R}^2$ il vettore *applicato* $\vec{OP}$, cioè il vettore con punto iniziale l'origine O del piano e con punto finale P , come appena visto a pag. 23. Si verifica subito che la somma definita dalla formula (2.1.1) corrisponde proprio alla somma di due vettori con la regola del parallelogramma e che la moltiplicazione per lo scalare $\lambda \in \mathbb{R}$ definita dalla formula (2.1.2) corrisponde alla moltiplicazione della lunghezza del vettore per λ , se λ è positivo, oppure alla moltiplicazione della lunghezza del vettore opposto per $-\lambda$, se λ è negativo. $\square$

Questa corrispondenza giustifica il termine di *spazio vettoriale* per X e la seguente:

DEFINIZIONE 2.1.3. Gli elementi di uno spazio vettoriale X sono detti *vettori*.

2.2. Proprietà della somma

DEFINIZIONE 2.2.1. La somma di due elementi di uno spazio vettoriale X deve verificare le seguenti quattro proprietà:

- (1) la **commutatività**, cioè per ogni $x, y \in X$ deve valere

$$x + y = y + x;$$

- (2) l'**associatività**, cioè per ogni $x, y, z \in X$ deve valere

$$(x + y) + z = x + (y + z);$$

- (3) esistenza dell'elemento **neutro**;

- (4) esistenza dell'**opposto**.

Queste proprietà sono molto naturali perché sono valide per tutti gli insiemi di numeri che già conoscete: interi, razionali e reali. Vedremo in futuro, però, esempi di operazioni su insiemi che non sono commutative, come per esempio la moltiplicazione di due matrici.

La proprietà commutativa ci dice semplicemente che cambiando l'ordine degli addendi, la somma non cambia.

La proprietà associativa invece ci dice che possiamo scrivere la somma di tre vettori senza parentesi, cioè possiamo scrivere:

$$x + y + z$$

senza alcuna ambiguità, invece di scrivere $(x+y)+z$, oppure $x+(y+z)$. Usando la proprietà associativa si può dimostrare che possiamo scrivere anche la somma di n vettori senza indicare alcuna parentesi:

$$x_1 + x_2 + x_3 + \cdots + x_n.$$

Per esempio, la somma di $n = 4$ vettori può essere fatta in molti modi:

$$\begin{aligned} x_1 + x_2 + x_3 + x_4 &= (x_1 + x_2) + (x_3 + x_4) = ((x_1 + x_2) + x_3) + x_4 = \\ &= (x_1 + (x_2 + x_3)) + x_4 = x_1 + ((x_2 + x_3) + x_4) = \\ &= x_1 + (x_2 + (x_3 + x_4)), \end{aligned}$$

ma il risultato è sempre lo stesso, quindi è inutile distinguere con le parentesi in quale ordine eseguire l'addizione. Ricordiamo che per la proprietà commutativa possiamo anche scambiare di posto i vettori.

Quando si considerano un certo numero di vettori qualsiasi, diciamo cinque vettori, si è soliti scrivere:

$$x_1 \quad x_2 \quad x_3 \quad x_4 \quad x_5$$

dove la lettera x indica che i vettori sono indeterminati e i numeri 1, 2, 3, 4 e 5 sono detti *indici*. Per esempio 4 è l'indice dell'elemento x_4 .

Per indicare la somma dei cinque vettori si può scrivere

$$x_1 + x_2 + x_3 + x_4 + x_5$$

oppure, in maniera più compatta, si può scrivere

$$\sum_{i=1}^5 x_i,$$

dove la lettera greca $\sum$ (sigma maiuscola) indica proprio la somma (spesso viene detta anche *sommatoria*) e i viene detto *indice* della sommatoria. Un altro modo per indicare la medesima somma è:

$$\sum_{1 \leq j \leq 5} x_j.$$

Si noti che non ha importanza quale lettera viene usata per denotare l'indice della sommatoria, anche se spesso viene usata la lettera i .

Se invece volessimo indicare, fra dodici vettori $x_1, \dots, x_{12}$ dati, l'addizione dei vettori con indice pari, potremmo scrivere così:

$$\sum_{\substack{1 \leq i \leq 12 \\ i \text{ pari}}} x_i.$$

Torniamo alle proprietà della somma in uno spazio vettoriale e vediamo cosa vogliono dire la terza e la quarta proprietà della definizione **2.2.1**.

L'esistenza dell'elemento neutro significa che esiste un elemento, che indicheremo con 0_X , tale che:

$$0_X + x = x, \quad \text{per ogni } x \in X. \quad (2.2.1)$$

Si dice allora che 0_X è l'elemento **neutro** di X . L'esistenza dell'opposto significa che per ogni elemento x di X esiste un elemento y di X tale che

$$x + y = 0_X. \quad (2.2.2)$$

Si dice che y è l'**opposto di** x e si denota di solito con $-x$.

Naturalmente, se consideriamo l'insieme dei vettori nel piano, l'elemento neutro è il vettore di lunghezza nulla, mentre l'opposto di un vettore dato, come abbiamo già accennato all'inizio del capitolo, è il vettore con stessa lunghezza, stessa direzione e verso opposto del vettore fissato.

ESEMPIO 2.2.2. Nel piano reale $\mathbb{R}^2$, considerato con le operazioni definite con le formule **(2.1.1)** e **(2.1.2)**, l'elemento neutro è $(0, 0) \in \mathbb{R}^2$, mentre l'opposto di (a_1, a_2) è

$$(-a_1, -a_2),$$

perché $(a_1 + a_2) + (-a_1, -a_2) = (a_1 - a_1, a_2 - a_2) = (0, 0) = 0_X$.

□

ESERCIZIO 2.2.3. Dimostrare che l'elemento neutro per la somma in uno spazio vettoriale è *unico*.

SVOLGIMENTO. Supponiamo che y e z siano due elementi dello spazio vettoriale X che soddisfino la definizione di elemento neutro (2.2.1). Allora si ha che $y = y + z$, considerando y come elemento neutro, e che $z = z + y$, considerando invece z come elemento neutro. Ne segue allora che

$$y = y + z = z + y = z$$

come volevasi dimostrare.

□

ESERCIZIO 2.2.4. Sia x un elemento qualsiasi di uno spazio vettoriale X . Dimostrare che l'opposto di x è *unico*.

□

SOLUZIONE. Supponiamo che y e z siano due opposti di x . Per definizione di 0_X si ha che $y = y + 0_X$. Per ipotesi z è opposto di x , quindi possiamo scrivere $0_X = x + z$. Allora si ha che:

$$y = y + 0_X = y + (x + z) = (y + x) + z = 0_X + z = z$$

dove la terza uguaglianza (da sinistra) segue dalla proprietà associativa, la quarta uguaglianza dall'ipotesi che y è opposto di x e infine la quinta e ultima segue dalla definizione di 0_X .

□

OSSERVAZIONE 2.2.5. Si noti che vale $0_X + 0_X = 0_X$, per definizione dell'elemento neutro 0_X applicata ad 0_X stesso, quindi

$$-0_X = 0_X,$$

per definizione di opposto (2.2.2).

2.3. Altre proprietà delle operazioni sui vettori

DEFINIZIONE 2.3.1. La moltiplicazione di un vettore per uno scalare in uno spazio vettoriale X deve verificare le seguenti due proprietà:

(1) l'*associatività*, cioè:

$$\beta \cdot (\alpha \cdot x) = (\beta \cdot \alpha) \cdot x, \quad \text{per ogni } \alpha, \beta \in \mathbb{R}, x \in X;$$

(2) esistenza dell'elemento neutro, che denotiamo con 1:

$$1 \cdot x = x, \quad \text{per ogni } x \in X.$$

Inoltre devono essere verificate altre due proprietà della moltiplicazione per scalari rispetto alla somma, le quali sono dette *proprietà distributive*:

- $\alpha \cdot (x + y) = \alpha \cdot x + \alpha \cdot y$, per ogni $\alpha \in \mathbb{R}$ e $x, y \in X$;
- $(\alpha + \beta) \cdot x = \alpha \cdot x + \beta \cdot x$, per ogni $\alpha, \beta \in \mathbb{R}$ e $x \in X$.

La prima è detta proprietà distributiva della moltiplicazione per scalari rispetto alla somma, mentre la seconda è la proprietà distributiva della somma rispetto alla moltiplicazione per scalari.

Anche se la formulazione delle proprietà può apparire non immediatamente chiara, in verità traduce in formule proprietà che intuitivamente sono evidenti. Per esempio, se $\alpha = 2$, la proprietà distributiva della moltiplicazione per scalari rispetto alla somma ci dice semplicemente che $2(x + y) = 2x + 2y$, come è ragionevole immaginare.

ESEMPIO 2.3.2. Consideriamo il piano reale $\mathbb{R}^2$ come negli esempi 2.1.2 e 2.2.2. Allora è facile verificare che le operazioni somma e moltiplicazione per scalari definite dalle formule (2.1.1) e (2.1.2) soddisfano tutte le proprietà che abbiamo elencato per gli spazi vettoriali. Si può affermare quindi che $\mathbb{R}^2$ è uno spazio vettoriale con tali operazioni.

□

Nei prossimi tre esercizi mostriamo alcune proprietà riguardanti le operazioni di somma e prodotto per scalari che seguono facilmente dalle definizioni e dalle altre proprietà già viste, come sarà chiaro dalle dimostrazioni.

ESERCIZIO 2.3.3. Dimostrare che, per ogni x, y, z elementi di uno spazio vettoriale X , se

$$x + y = x + z \tag{2.3.1}$$

allora $y = z$.

□

SOLUZIONE. Ricordiamo che per ogni $x \in X$ esiste l'opposto di x , che indichiamo con $-x$. Sommando membro a membro $-x$, la formula (2.3.1) è equivalente alla seguente:

$$-x + x + y = -x + x + z. \quad (2.3.2)$$

Il primo membro diventa:

$$-x + (x + y) = (-x + x) + y = 0_X + y = y,$$

perché la prima uguaglianza segue dalla proprietà associativa dell'addizione, la seconda uguaglianza segue dalla definizione di opposto, mentre l'ultima uguaglianza segue dalla definizione di elemento neutro 0_X .

Similmente il secondo membro di (2.3.2) diventa:

$$-x + (x + z) = (-x + x) + z = 0_X + z = z.$$

Concludiamo quindi che $y = z$, come richiesto. $\square$

ESERCIZIO 2.3.4. Dimostrare che $0 \cdot x = 0_X$, per ogni $x \in X$. $\square$

SOLUZIONE. Ricordiamo che possiamo scrivere il numero reale 0 come $0 = 0 + 0$, quindi

$$0 \cdot x = (0 + 0) \cdot x = 0 \cdot x + 0 \cdot x \quad (2.3.3)$$

dove la seconda uguaglianza segue dalla proprietà distributiva dell'addizione rispetto alla moltiplicazione per scalari. Possiamo riscrivere il primo membro così:

$$0 \cdot x = 0_X + 0 \cdot x,$$

per definizione di elemento neutro, così la formula (2.3.3) diventa:

$$0_X + 0 \cdot x = 0 \cdot x + 0 \cdot x$$

e si conclude applicando l'esercizio 2.3.3. $\square$

ESERCIZIO 2.3.5. Dimostrare che $(-1) \cdot x = -x$, per ogni $x \in X$.

$\square$

SOLUZIONE. Ricordando che $x = 1 \cdot x$ e la proprietà distributiva, si ha che:

$$x + (-1) \cdot x = 1 \cdot x + (-1) \cdot x = (1 - 1) \cdot x = 0 \cdot x = 0_X$$

dove l'ultima uguaglianza segue dall'esercizio 2.3.4. Si conclude allora per l'unicità dell'opposto di x . $\square$

2.4. Combinazioni lineari di vettori

DEFINIZIONE 2.4.1. Nello spazio vettoriale X consideriamo due vettori x_1 e x_2 . Dati due scalari $\alpha_1, \alpha_2 \in \mathbb{R}$, si dice che:

$$\alpha_1 x_1 + \alpha_2 x_2,$$

che si può scrivere anche così:

$$\sum_{i=1}^2 \alpha_i x_i,$$

è **combinazione lineare** dei due vettori dati. Gli scalari α_1 e α_2 sono detti **coefficienti** di x_1 e x_2 , rispettivamente.

Per esempio se $\alpha_1 = 2$ e $\alpha_2 = 3$, allora

$$2x_1 + 3x_2$$

è una combinazione lineare di x_1 e x_2 .

ESEMPIO 2.4.2. Siano $\begin{pmatrix} 2 \\ -1 \end{pmatrix}$ e $\begin{pmatrix} 0 \\ 1 \end{pmatrix}$ elementi del piano reale $\mathbb{R}^2$. Calcoliamo la combinazione lineare di questi due vettori con coefficienti rispettivamente -1 e -2 :

$$-1 \begin{pmatrix} 2 \\ -1 \end{pmatrix} - 2 \begin{pmatrix} 0 \\ 1 \end{pmatrix} = \begin{pmatrix} -1 \cdot 2 - 2 \cdot 0 \\ -1 \cdot (-1) - 2 \cdot 1 \end{pmatrix} = \begin{pmatrix} -2 \\ -1 \end{pmatrix}.$$

$\square$

Le combinazioni lineari si possono fare anche per tre o più vettori:

DEFINIZIONE 2.4.3. Se $x_1, \dots, x_n$ sono n vettori dati e $\alpha_1, \dots, \alpha_n$ sono n scalari (tanti quanti i vettori), allora

$$\alpha_1 x_1 + \alpha_2 x_2 + \dots + \alpha_n x_n$$

è la **combinazione lineare** di $x_1, \dots, x_n$ con coefficienti $\alpha_1, \dots, \alpha_n$, rispettivamente, che possiamo scrivere anche:

$$\sum_{i=1}^n \alpha_i x_i.$$

Se consideriamo un solo vettore x e un qualsiasi scalare α , allora il *multiplo* αx di x è anch'esso detto combinazione lineare di x .

ESEMPIO 2.4.4. Consideriamo il vettore $\begin{pmatrix} 1 \\ -1 \end{pmatrix}$.

$$\alpha \begin{pmatrix} 1 \\ -1 \end{pmatrix} = \begin{pmatrix} \alpha \\ -\alpha \end{pmatrix}$$

è un multiplo di $\begin{pmatrix} 1 \\ -1 \end{pmatrix}$ per ogni $\alpha \in \mathbb{R}$. □

ESEMPIO 2.4.5. Consideriamo i seguenti tre elementi di $\mathbb{R}^2$:

$$\begin{pmatrix} -3 \\ 0 \end{pmatrix}, \quad \begin{pmatrix} -2 \\ 1 \end{pmatrix}, \quad \begin{pmatrix} 1 \\ 1 \end{pmatrix}.$$

La combinazione lineare di questi tre elementi con coefficienti rispettivamente 0, 1 e -1 è

$$0 \begin{pmatrix} -3 \\ 0 \end{pmatrix} + \begin{pmatrix} -2 \\ 1 \end{pmatrix} - \begin{pmatrix} 1 \\ 1 \end{pmatrix} = \begin{pmatrix} -3 \\ 0 \end{pmatrix}.$$

□

2.5. Sistemi di vettori linearmente dipendenti o indipendenti

Consideriamo un vettore non nullo $x \in X$. Nella definizione 2.4.3 abbiamo visto che i multipli αx di x sono combinazioni lineari di x .

DEFINIZIONE 2.5.1. Fissato un vettore non nullo $x \in X$, un vettore $y \in X$ che non si può scrivere nella forma $y = \alpha x$, per nessun $\alpha \in \mathbb{R}$, è detto **linearmente indipendente** da x .

Per esempio il vettore $\mathbf{y} = \begin{pmatrix} -3 \\ 3 \end{pmatrix} \in \mathbb{R}^2$ non è linearmente indipendente dal vettore $\mathbf{x} = \begin{pmatrix} 1 \\ -1 \end{pmatrix}$, perché $\mathbf{y} = -3\mathbf{x}$ (si veda l'Esempio 2.4.4).

Il vettore $\mathbf{z} = \begin{pmatrix} 2 \\ 0 \end{pmatrix}$, invece, è linearmente indipendente da $\mathbf{x}$ perché non vale l'uguaglianza $\mathbf{z} = \alpha\mathbf{x}$ per nessun $\alpha \in \mathbb{R}$. Infatti, supponiamo per assurdo che esista $\alpha \in \mathbb{R}$ tale che $\mathbf{z} = \alpha\mathbf{x}$: allora dovrebbe valere

$$\begin{pmatrix} 2 \\ 0 \end{pmatrix} = \begin{pmatrix} \alpha \\ -\alpha \end{pmatrix},$$

cioè $\alpha = 2$ per avere l'uguaglianza della prima coordinata ma $-\alpha = 0$ per l'uguaglianza della seconda. Troviamo così una contraddizione (si veda di nuovo l'Esempio 2.4.4).

OSSERVAZIONE 2.5.2. Se rappresentiamo i vettori nel piano reale come usuale, i multipli di un vettore dato $\mathbf{x}$ stanno sulla retta con la stessa direzione del vettore $\mathbf{x}$. Quindi un vettore $\mathbf{y}$ linearmente indipendente da $\mathbf{x}$ è un vettore che non sta su tale retta.

DEFINIZIONE 2.5.3. Siano $\mathbf{x}$ e $\mathbf{y}$ due vettori di uno spazio vettoriale $\mathbf{X}$. Si dice che $\mathbf{x}$ e $\mathbf{y}$ sono *linearmente dipendenti* se non sono linearmente indipendenti, cioè se $\mathbf{x}$ è un multiplo di $\mathbf{y}$, o equivalentemente se $\mathbf{y}$ è multiplo di $\mathbf{x}$.

Se $\mathbf{y}$ è il vettore nullo $\mathbf{0}_X$, allora per ogni $\mathbf{x} \in \mathbf{X}$ i due vettori $\mathbf{x}$ e $\mathbf{y}$ sono linearmente dipendenti, perché possiamo scrivere $\mathbf{y} = \mathbf{0}_X = 0\mathbf{x}$. In altri termini, il vettore nullo è multiplo di qualsiasi altro vettore.

OSSERVAZIONE 2.5.4. Se x e y sono vettori non nulli, allora x è multiplo di y se e solo se y è multiplo di x . Infatti, se x è multiplo di y , allora esiste $\alpha \in \mathbb{R}$ tale che $x = \alpha y$, con $\alpha \neq 0$, perché $x \neq 0$. Allora anche y è multiplo di x , perché dividendo per α troviamo che $y = \frac{1}{\alpha}x$.

DEFINIZIONE 2.5.5. Siano x e y due vettori dello spazio vettoriale X . Si dice che un vettore z è *linearmente indipendente* da x e y se non

esistono $\alpha, \beta \in \mathbb{R}$ tali che

$$z = \alpha x + \beta y. \quad (2.5.1)$$

In caso contrario, cioè se esistono siffatti α e β , si dice che z è *linearmente dipendente* da x e y .

Osserviamo che l'equazione (2.5.1) si può riscrivere nel modo seguente:

$$\alpha x + \beta y - z = 0_X,$$

cioè abbiamo trovato una combinazione lineare di x , y e z , con coefficienti non tutti nulli, la cui somma è il vettore nullo. Viceversa, se esiste una combinazione lineare di x , y e z

$$\delta x + \xi y + \gamma z = 0, \quad \text{tale che } \gamma \neq 1$$

allora, dividendo per γ , si trova che

$$\frac{\delta}{\gamma} x + \frac{\xi}{\gamma} y + z = 0$$

che possiamo riscrivere come:

$$z = \alpha x + \beta y, \quad \text{dove } \alpha = -\frac{\delta}{\gamma}, \quad \beta = -\frac{\xi}{\gamma}$$

cioè z è combinazione lineare di x e y e quindi è linearmente dipendente da essi.

ESEMPIO 2.5.6. Siano $x = \begin{pmatrix} 2 \\ -1 \end{pmatrix}$ e $y = \begin{pmatrix} -1 \\ 1 \end{pmatrix}$ due elementi di $\mathbb{R}^2$. Il vettore $z = \begin{pmatrix} 1 \\ 0 \end{pmatrix}$ è linearmente dipendente da x e y perché $z = x + y$.

□

Più in generale, la nozione di dipendenza lineare si può dare anche per più di due vettori:

DEFINIZIONE 2.5.7. Consideriamo n vettori $x_1, \dots, x_n$ di uno spazio vettoriale X . Si dice che $x_1, \dots, x_n$ sono **linearmente indipendenti** se non esiste una combinazione lineare di $x_1, \dots, x_n$, a coefficienti non tutti nulli, tale che la loro somma sia il vettore nullo 0_X . In altri termini, $x_1, \dots, x_n$ sono linearmente indipendenti se

$$\lambda_1 x_1 + \dots + \lambda_n x_n = 0_X \implies \lambda_1 = \dots = \lambda_n = 0.$$

Con la prossima proposizione verifichiamo che la definizione 2.5.7 è concorde con le precedenti definizioni 2.5.1 e 2.5.5.

PROPOSIZIONE 2.5.8. *Siano $x_1, \dots, x_n$ vettori di uno spazio vettoriale X . Allora $x_1, \dots, x_n$ sono linearmente indipendenti se e solo se non esiste alcun x_k , con $1 \leq k \leq n$, che sia combinazione lineare dei rimanenti.*

DIMOSTRAZIONE. Supponiamo che esista x_k che sia combinazione lineare dei rimanenti. Per semplicità, poniamo che sia $k = n$. Allora esistono $\alpha_1, \dots, \alpha_{n-1} \in \mathbb{R}$ tali che

$$x_n = \alpha_1 x_1 + \alpha_2 x_2 + \dots + \alpha_{n-1} x_{n-1}.$$

Riscrivendo la formula precedente nel modo seguente:

$$\alpha_1 x_1 + \alpha_2 x_2 + \dots + \alpha_{n-1} x_{n-1} - x_n = 0,$$

cioè abbiamo trovato una combinazione lineare con coefficienti non tutti nulli che dà 0_X , quindi $x_1, \dots, x_n$ sono linearmente dipendenti secondo la definizione (2.5.7).

Viceversa, se $x_1, \dots, x_n$ sono linearmente dipendenti, esiste almeno una combinazione lineare

$$\alpha_1 x_1 + \alpha_2 x_2 + \dots + \alpha_{n-1} x_{n-1} + \alpha_n x_n = 0,$$

con almeno un coefficiente non nullo, che supponiamo per semplicità essere α_n . Allora possiamo dividere tutto per α_n e portare x_n dall'altro membro:

$$x_n = -\frac{\alpha_1}{\alpha_n} x_1 - \frac{\alpha_2}{\alpha_n} x_2 + \dots - \frac{\alpha_{n-1}}{\alpha_n} x_{n-1},$$

cioè abbiamo scritto x_n come combinazione lineare dei rimanenti.

□

Riordinando i vettori se necessario possiamo riscrivere l'enunciato precedente in questo modo:

COROLLARIO 2.5.9. *Siano $x_1, \dots, x_n$ vettori di uno spazio vettoriale X . Allora $x_1, \dots, x_n$ sono linearmente indipendenti se e solo se non esiste alcun x_k , con $1 \leq k \leq n$, che sia combinazione lineare dei vettori precedenti $x_1, \dots, x_{k-1}$.*

ESERCIZIO 2.5.10. Dimostrare che i vettori $\begin{pmatrix} 1 \\ -1 \end{pmatrix}$ e $\begin{pmatrix} 0 \\ 1 \end{pmatrix}$ di $\mathbb{R}^2$ sono linearmente indipendenti usando la definizione 2.5.7.

SVOLGIMENTO. Consideriamo una combinazione lineare dei due vettori fissati che dia 0_X :

$$\lambda_1 \begin{pmatrix} 1 \\ -1 \end{pmatrix} + \lambda_2 \begin{pmatrix} 0 \\ 1 \end{pmatrix} = \begin{pmatrix} 0 \\ 0 \end{pmatrix}$$

Allora deve essere

$$\begin{pmatrix} 1\lambda_1 + 0\lambda_2 \\ -1\lambda_1 + 1\lambda_2 \end{pmatrix} = \begin{pmatrix} \lambda_1 \\ -\lambda_1 + \lambda_2 \end{pmatrix} = \begin{pmatrix} 0 \\ 0 \end{pmatrix}$$

quindi l'uguaglianza della prima coordinata implica che $\lambda_1 = 0$ e quella della seconda che $0 = -\lambda_1 + \lambda_2 = \lambda_2$. Abbiamo così dimostrato che l'unica combinazione lineare dei due vettori fissati che dà 0_X è quella con tutti i coefficienti nulli, quindi i vettori dati sono linearmente indipendenti per la definizione 2.5.7. $\square$

2.6. Sottospazi vettoriali ed insiemi di generatori

DEFINIZIONE 2.6.1. Sia X uno spazio vettoriale. Un sottoinsieme Y di X si dice un **sottospazio vettoriale** di X se valgono le seguenti due condizioni:

- (1) per ogni $y, z \in Y$, si ha $y + z \in Y$;
- (2) per ogni $y \in Y$ e $\lambda \in \mathbb{R}$, si ha $\lambda y \in Y$.

In particolare dalla seconda condizione segue che $0_X \in Y$.

In altri termini, un sottospazio vettoriale Y di X è un sottoinsieme di X in cui si possono fare le operazioni di somma e moltiplicazione per scalari definite in X senza uscire da Y .

ESEMPIO 2.6.2. Consideriamo il vettore $\begin{pmatrix} 1 \\ -1 \end{pmatrix} \in \mathbb{R}^2$ e tutti i suoi multipli $\alpha \begin{pmatrix} 1 \\ -1 \end{pmatrix} = \begin{pmatrix} \alpha \\ -\alpha \end{pmatrix}$, per $\alpha \in \mathbb{R}$, come nell'esempio 2.4.4. Allora l'insieme

$$\left\{ \begin{pmatrix} \alpha \\ -\alpha \end{pmatrix} : \alpha \in \mathbb{R} \right\} \quad (2.6.1)$$

è un sottospazio vettoriale di $\mathbb{R}^2$.

Infatti, sommando due elementi dell'insieme (2.6.1) troviamo:

$$\begin{pmatrix} \alpha \\ -\alpha \end{pmatrix} + \begin{pmatrix} \beta \\ -\beta \end{pmatrix} = \begin{pmatrix} \alpha + \beta \\ -\alpha - \beta \end{pmatrix}$$

che è ancora un elemento dell'insieme. In modo simile si verifica anche la proprietà (2) della definizione 2.6.1. $\square$

DEFINIZIONE 2.6.3. Si dice che l'insieme (2.6.1) è **generato** dal vettore $\begin{pmatrix} 1 \\ -1 \end{pmatrix}$.

Sia X uno spazio vettoriale e consideriamo dei vettori $x_1, \dots, x_n$ di X , per un certo $n \geq 1$.

PROPOSIZIONE 2.6.4. *L'insieme di vettori che sono combinazioni lineari di $x_1, \dots, x_n$, cioè*

$$\{x = \lambda_1 x_1 + \lambda_2 x_2 + \dots + \lambda_n x_n : \lambda_1, \lambda_2, \dots, \lambda_n \in \mathbb{R}\}, \quad (2.6.2)$$

è un sottospazio lineare di X .

DEFINIZIONE 2.6.5. Si dice che (2.6.2) è il sottospazio vettoriale di X **generato da** $x_1, \dots, x_n$.

Osserviamo che 0_X appartiene al sottospazio generato da $x_1, \dots, x_n$ perché

$$0_X = 0 \cdot x_1 + 0 \cdot x_2 + \dots + 0 \cdot x_n.$$

DEFINIZIONE 2.6.6. Diciamo che $\{y_1, \dots, y_k\}$ è un **sistema di generatori** di uno spazio vettoriale X se ogni vettore in X è combinazione lineare di $y_1, \dots, y_k$, cioè se per ogni $x \in X$ esistono scalari $\alpha_1, \dots, \alpha_k \in \mathbb{R}$ tali che

$$x = \sum_{i=1}^k \alpha_i y_i.$$

ESEMPIO 2.6.7. Consideriamo lo spazio reale a tre dimensioni $\mathbb{R}^3$. I tre vettori

$$\begin{pmatrix} 1 \\ 0 \\ 0 \end{pmatrix}, \quad \begin{pmatrix} 0 \\ 1 \\ 0 \end{pmatrix}, \quad \begin{pmatrix} 0 \\ 0 \\ 1 \end{pmatrix},$$

formano un sistema di generatori, infatti ogni vettore $(\alpha_1, \alpha_2, \alpha_3)$ di $\mathbb{R}^3$ si può scrivere come combinazione lineare di questi tre vettori:

$$\begin{pmatrix} \alpha_1 \\ \alpha_2 \\ \alpha_3 \end{pmatrix} = \alpha_1 \begin{pmatrix} 1 \\ 0 \\ 0 \end{pmatrix} + \alpha_2 \begin{pmatrix} 0 \\ 1 \\ 0 \end{pmatrix} + \alpha_3 \begin{pmatrix} 0 \\ 0 \\ 1 \end{pmatrix} = \begin{pmatrix} \alpha_1 \\ 0 \\ 0 \end{pmatrix} + \begin{pmatrix} 0 \\ \alpha_2 \\ 0 \end{pmatrix} + \begin{pmatrix} 0 \\ 0 \\ \alpha_3 \end{pmatrix}.$$

□

2.7. Base di uno spazio vettoriale

Il nostro principale obiettivo ora è provare che il numero di vettori linearmente indipendenti non può essere maggiore del numero di elementi di un sistema di generatori. Questo fatto ci servirà per dimostrare che ogni *base* di uno spazio vettoriale, cioè ogni sistema di generatori linearmente indipendenti, ha lo stesso numero di elementi. Proviamo prima un enunciato ausiliario.

PROPOSIZIONE 2.7.1. *Sia X uno spazio vettoriale. Supponiamo di avere due sistemi linearmente indipendenti,*

$$S = \{x_1, \dots, x_n\} \text{ e } T = \{y_1, \dots, y_k\}.$$

Se gli n vettori $x_1, \dots, x_n$ formano un sistema di generatori per X , allora $k \leq n$.

DIMOSTRAZIONE. Si consideri la famiglia $T_1 = \{y_1, x_1, x_2, \dots, x_n\}$ di vettori. Poiché i vettori $x_1, \dots, x_n$ formano un sistema di generatori, y_1 è combinazione lineare di essi e quindi la famiglia T_1 è linearmente dipendente. Per il Corollario 2.5.9, esiste almeno un x_j che dipende linearmente dai precedenti vettori in T_1 : scartiamo questo vettore x_j e ci riduciamo alla famiglia $S_1 = \{y_1, x_1, x_2, \dots, x_{j-1}, x_{j+1}, \dots, x_n\}$, che continua a generare tutto X perché il vettore che abbiamo scartato è combinazione lineare di quelli in S_1 . Aggiungiamo ora il secondo vettore di T , cioè formiamo la famiglia $T_2 = \{y_2, y_1, x_1, x_2, \dots, x_{j-1}, x_{j+1}, \dots, x_n\}$. Come prima, questa famiglia è linearmente dipendente, ed uno dei suoi vettori dipende linearmente dai precedenti. Questo vettore non può essere y_1 , perché $y_1, \dots, y_k$ sono indipendenti, quindi deve essere qualche x_i con $i \neq j$. Scartiamo anche x_i e continuiamo in questo modo. Ad ogni passo otteniamo una famiglia linearmente dipendente che genera

tutto X , da cui scartare un vettore (uno degli x , non degli y) senza che la sottofamiglia risultante smetta di essere un sistema di generatori. Se i vettori in T fossero meno di quelli in S , alla fine ci ritroveremmo con una famiglia linearmente dipendente di vettori tutti contenuti in $T = \{y_1, \dots, y_k\}$. Ma allora T non è linearmente indipendente, ed abbiamo una contraddizione. $\square$

ESERCIZIO 2.7.2. Con un argomento simmetrico rispetto a quello della Proposizione 2.7.1, provare l'enunciato seguente:

Sia X uno spazio vettoriale. Supponiamo di avere due sistemi di generatori,

$$S = \{y_1, \dots, y_k\} \text{ e } T = \{x_1, \dots, x_n\}.$$

Se i k vettori $y_1, \dots, y_n$ formano un sistema linearmente indipendente, allora $k \leq n$.

SVOLGIMENTO. La dimostrazione è simmetrica rispetto a quella precedente. Consideriamo la famiglia T_1 di vettori $\{y_1, x_1, x_2, \dots, x_n\}$. Esattamente come prima, y_1 è combinazione lineare di $x_1, \dots, x_n$ e quindi la famiglia T_1 è linearmente dipendente. Per il Corollario 2.5.9, esiste almeno un x_j che dipende linearmente dai precedenti vettori in T_1 : scartiamo questo vettore x_j e ci riduciamo alla famiglia $S_1 = \{y_1, x_1, x_2, \dots, x_{j-1}, x_{j+1}, \dots, x_n\}$, che come prima continua a generare tutto X . Aggiungiamo allora il secondo vettore di S per formare la famiglia $T_2 = \{y_2, y_1, x_1, x_2, \dots, x_{j-1}, x_{j+1}, \dots, x_n\}$. Di nuovo, questa famiglia è un sistema di generatori linearmente dipendente, e possiamo scartare uno dei suoi vettori senza che questo scarto le faccia perdere la proprietà di generare tutto X . Questo vettore che possiamo scartare non è però y_1 perché S è un sistema linearmente indipendente. Continuando così, se il numero dei vettori in S fosse maggiore di quelli di T , ci ridurremmo ad una sottofamiglia linearmente dipendente di vettori tutti in S : questo contraddirebbe l'ipotesi che S sia linearmente indipendente. $\square$

Dalla Proposizione precedenti segue facilmente la conclusione desiderata:

TEOREMA 2.7.3. *Sia X uno spazio vettoriale. Supponiamo di avere un sistema T di n vettori non nulli linearmente indipendenti $x_1, \dots, x_n$, ed un sistema S di k generatori non nulli $\{y_1, \dots, y_k\}$. Allora $n \leq k$.*

DIMOSTRAZIONE. Se il sistema S di generatori è anche linearmente indipendente, l'enunciato coincide con la Proposizione 2.7.1. Assumiamo quindi che non lo sia, e scartiamo il primo vettore che è combinazione lineare dei precedenti, riducendolo quindi ad una sottofamiglia S_1 consistente di $k-1$ vettori la quale, come prima, è ancora un sistema di generatori. Se adesso S_1 risulta essere linearmente indipendente, allora, sempre per la Proposizione 2.7.1, si deve avere $n \leq k-1 < k$ ed il teorema è dimostrato. Se invece S_1 è linearmente dipendente, continuiamo a scartare vettori finché non lo diventa (prima o poi deve diventarlo, perché una famiglia consistente di un solo vettore non nullo è linearmente indipendente).

In questo modo abbiamo provato che si può estrarre da S una sottofamiglia S' di generatori linearmente indipendenti, ed il numero di vettori di S' è non inferiore a quello di T . $\square$

DEFINIZIONE 2.7.4. Una **base** di uno spazio vettoriale è un sistema di generatori linearmente indipendenti.

Sia X uno spazio vettoriale e $\{x_1, \dots, x_n\}$ una sua base. Siccome $x_1, \dots, x_n$ è un sistema di generatori di X , ogni elemento x di X si può scrivere come combinazione lineare di $x_1, \dots, x_n$. Il fatto che $x_1, \dots, x_n$ siano anche linearmente indipendenti, implicano che la scrittura di x come combinazione lineare di $x_1, \dots, x_n$ è *unica*, come mostra la seguente:

PROPOSIZIONE 2.7.5. *Sia $\{x_1, \dots, x_n\}$ una base di uno spazio vettoriale X . Sia x un vettore qualsiasi di X . Allora esistono e sono univocamente determinati $\lambda_1, \dots, \lambda_n \in \mathbb{R}$ tali che*

$$x = \lambda_1 x_1 + \lambda_2 x_2 + \dots + \lambda_n x_n. \quad (2.7.1)$$

DIMOSTRAZIONE. Come abbiamo già osservato prima dell'enunciato, $x_1, \dots, x_n$ è un sistema di generatori. Quindi, per ogni $x \in X$, esiste sicuramente la combinazione lineare (2.7.1). Supponiamo che esistano

anche $\lambda'_1, \dots, \lambda'_n$ tali che $x = \sum_{i=1}^n \lambda'_i x_i$. Allora avremmo che:

$$\lambda_1 x_1 + \lambda_2 x_2 + \dots + \lambda_n x_n = \lambda'_1 x_1 + \lambda'_2 x_2 + \dots + \lambda'_n x_n$$

e quindi, portando tutto al primo membro, che:

$$(\lambda_1 - \lambda'_1)x_1 + (\lambda_2 - \lambda'_2)x_2 + \dots + (\lambda_n - \lambda'_n)x_n = 0_X.$$

Ma $x_1, \dots, x_n$ sono linearmente indipendenti e l'equazione precedente è una combinazione lineare che dà 0_X . Perciò tale combinazione lineare deve avere tutti i coefficienti nulli, cioè deve essere:

$$\lambda_1 = \lambda'_1, \quad \lambda_2 = \lambda'_2, \quad \dots \quad \lambda_n = \lambda'_n,$$

come volevasi dimostrare. $\square$

Il prossimo teorema ci mostra una proprietà fondamentale che hanno tutte le basi di uno stesso spazio vettoriale.

TEOREMA 2.7.6. *Due basi di uno spazio vettoriale hanno lo stesso numero di elementi.*

DIMOSTRAZIONE. Siano $\{x_1, \dots, x_n\}$ e $\{y_1, \dots, y_m\}$ due basi dello stesso spazio vettoriale X . In particolare $x_1, \dots, x_n$ sono linearmente indipendenti e $\{y_1, \dots, y_m\}$ è un sistema di generatori, quindi il teorema 2.7.3 implica che $n \leq m$. D'altra parte, è vero che anche $y_1, \dots, y_m$ sono linearmente indipendenti e $\{x_1, \dots, x_n\}$ è un sistema di generatori, perciò ancora lo stesso teorema 2.7.3 implica pure che $m \leq n$. Ne segue allora che $n = m$, cioè due basi qualsiasi hanno lo stesso numero di elementi, come volevasi dimostrare. $\square$

Possiamo quindi dare la seguente:

DEFINIZIONE 2.7.7. La **dimensione** di uno spazio vettoriale è il numero di elementi di una base, che per il teorema precedente non dipende dalla base scelta. Se n è la dimensione dello spazio vettoriale X , allora scriviamo:

$$\dim X = n.$$

ESEMPIO 2.7.8. Consideriamo i seguenti n vettori di $\mathbb{R}^n$:

$$\begin{pmatrix} 1 \\ 0 \\ \vdots \\ 0 \end{pmatrix}, \quad \begin{pmatrix} 0 \\ 1 \\ \vdots \\ 0 \end{pmatrix}, \quad \dots, \quad \begin{pmatrix} 0 \\ 0 \\ \vdots \\ 1 \end{pmatrix}.$$

Come per $\mathbb{R}^2$ o $\mathbb{R}^3$, è facile verificare che questi n vettori sono linearmente indipendenti e che generano tutto $\mathbb{R}^n$. $\square$

Concludiamo questa sezione e questo capitolo mostrando come costruire delle basi di spazio vettoriale.

Supponiamo di avere un sistema di generatori $\{x_1, \dots, x_5\}$ di X . Se questi elementi fossero linearmente indipendenti, sarebbero una base. Ma se non sono linearmente indipendenti, allora bisogna trovare una base.

Per esempio, supponiamo che x_5 sia una combinazione lineare degli altri, cioè di x_1, x_2, x_3 e x_4 :

$$x_5 = \alpha_1 x_1 + \alpha_2 x_2 + \alpha_3 x_3 + \alpha_4 x_4 \quad (2.7.2)$$

con $\alpha_1, \alpha_2, \alpha_3, \alpha_4 \in \mathbb{R}$. Allora x_1, x_2, x_3 e x_4 generano X , perché sappiamo che ogni vettore $x \in X$ è combinazione lineare di $x_1, \dots, x_5$, quindi esistono $\lambda_1, \dots, \lambda_5 \in \mathbb{R}$ tali che

$$\begin{aligned} x &= \lambda_1 x_1 + \dots + \lambda_4 x_4 + \lambda_5 x_5 = \\ &= \lambda_1 x_1 + \dots + \lambda_4 x_4 + \lambda_5 (\alpha_1 x_1 + \alpha_2 x_2 + \alpha_3 x_3 + \alpha_4 x_4) \end{aligned}$$

dove l'ultima uguaglianza segue da (2.7.2), perciò x è combinazione lineare di $x_1, \dots, x_4$.

Generalizzando il ragionamento precedente ad un sistema di n generatori $\{x_1, \dots, x_n\}$ di uno spazio vettoriale X si dimostra la seguente:

PROPOSIZIONE 2.7.9. *Sia $\{x_1, \dots, x_n\}$ un sistema di generatori di uno spazio vettoriale X . Allora esiste un sottoinsieme $\{x_{i_1}, \dots, x_{i_k}\}$ del sistema di generatori che è una base di X .*

DIMOSTRAZIONE. Se $x_1, \dots, x_n$ sono linearmente indipendenti, allora abbiamo già una base di X con $k = n$ e $i_1 = 1, i_2 = 2, \dots, i_k = n$.

n . Se invece $x_1, \dots, x_n$ sono linearmente dipendenti, allora esiste una combinazione lineare:

$$\lambda_1 x_1 + \lambda_2 x_2 + \dots + \lambda_n x_n = 0_X$$

dove i coefficienti $\lambda_1, \dots, \lambda_n$ non sono tutti nulli. Allora esiste i tale che $\lambda_i \neq 0$ e possiamo scrivere x_i in funzione dei rimanenti:

$$x_i = -\frac{\lambda_1}{\lambda_i} x_1 - \dots - \frac{\lambda_{i-1}}{\lambda_i} x_{i-1} - \frac{\lambda_{i+1}}{\lambda_i} x_{i+1} - \dots - \frac{\lambda_n}{\lambda_i} x_n.$$

Ne segue che $\{x_1, \dots, x_{i-1}, x_{i+1}, \dots, x_n\}$ è ancora un sistema di generatori di X . Se ora $x_1, \dots, x_{i-1}, x_{i+1}, \dots, x_n$ sono linearmente indipendenti, abbiamo trovato una base di X e abbiamo finito. Se invece sono linearmente dipendenti, allora esiste una combinazione lineare con coefficienti non tutti nulli che è 0_X , da cui possiamo ricavare uno dei vettori x_j in funzione dei rimanenti. Ripetendo lo stesso ragionamento a questi vettori rimasti, dopo un numero finito di volte troveremo un sistema di generatori di X formato da vettori linearmente indipendenti, che quindi è una base di X . Per costruzione, i vettori della base saranno appartenenti al sistema di generatori dato in partenza. $\square$

Vediamo ora un altro metodo per trovare una base di uno spazio vettoriale X . Sia x_1 un vettore non nullo di X . Se x_1 genera X , allora x_1 è una base. Altrimenti, se x_1 non genera X , esiste un altro elemento x_2 linearmente indipendente da x_1 , cioè che non è multiplo di x_1 . Se x_1 e x_2 generano X , allora sono una base, perché sono linearmente indipendenti. Altrimenti esiste un vettore x_3 che non è combinazione lineare di x_1 e x_2 . Si procede allo stesso modo finché non si trova un sistema di generatori che sono linearmente indipendenti per costruzione, che quindi formano una base.

Si può formalizzare questa idea con la seguente proposizione:

PROPOSIZIONE 2.7.10. *Siano $x_1, \dots, x_m$ vettori linearmente indipendenti di uno spazio vettoriale X (può essere $m = 1$, nel qual caso x_1 è un vettore non nullo qualsiasi). Se la dimensione di X è $n > m$, allora esistono dei vettori $x_{m+1}, \dots, x_n$ tali che $\{x_1, \dots, x_n\}$ è una base di X .*

DIMOSTRAZIONE. Consideriamo il sottospazio vettoriale generato da $x_1, \dots, x_m$. Scegliamo un vettore qualsiasi x_{m+1} non appartenente a

questo sottospazio. Allora $x_1, \dots, x_{m+1}$ sono linearmente indipendenti. Infatti, se non lo fossero, esisterebbero $\lambda_1, \dots, \lambda_{m+1}$ non tutti nulli tali che:

$$\lambda_1 x_1 + \dots + \lambda_m x_m + \lambda_{m+1} x_{m+1} = 0_X;$$

ora ci sono due possibilità: o $\lambda_{m+1} = 0$, ma allora $x_1, \dots, x_m$ sarebbero linearmente dipendenti, in contraddizione con l'ipotesi; oppure $\lambda_{m+1} \neq 0$, ma allora potremmo scrivere x_{m+1} come combinazione lineare di $x_1, \dots, x_m$, contraddicendo l'ipotesi di aver scelto x_m non appartenente al sottospazio generato da $x_1, \dots, x_n$.

A questo punto possiamo ripetere il ragionamento ai vettori linearmente indipendenti $x_1, \dots, x_{m+1}$ e scegliere un vettore qualsiasi x_{m+2} non appartenente al sottospazio vettoriale da essi generato. Con la stessa dimostrazione appena fatta, si vede che $x_1, \dots, x_{m+2}$ sono linearmente indipendenti.

Continuando così, si costruisce una base $x_1, \dots, x_n$ come volevasi dimostrare. $\square$

ESEMPIO 2.7.11. Consideriamo il sistema di generatori $\{x_1, x_2, x_3\}$, dello spazio vettoriale $X = \mathbb{R}^2$, dove:

$$x_1 = \begin{pmatrix} 2 \\ -1 \end{pmatrix}, \quad x_2 = \begin{pmatrix} 0 \\ 1 \end{pmatrix}, \quad x_3 = \begin{pmatrix} 1 \\ 1 \end{pmatrix}.$$

Vediamo come trovare una base di X seguendo la proposizione 2.7.9.

I tre vettori x_1, x_2 e x_3 sono linearmente dipendenti, infatti

$$\begin{aligned} x_1 + 3x_2 - 2x_3 &= \begin{pmatrix} 2 \\ -1 \end{pmatrix} + 3 \begin{pmatrix} 0 \\ 1 \end{pmatrix} - 2 \begin{pmatrix} 1 \\ 1 \end{pmatrix} = \begin{pmatrix} 2 + 0 \cdot 3 - 2 \cdot 1 \\ -1 + 3 \cdot 1 - 2 \cdot 1 \end{pmatrix} \\ &= \begin{pmatrix} 0 \\ 0 \end{pmatrix}. \end{aligned}$$

Allora si ha che:

$$x_1 = 2x_3 - 3x_2$$

e quindi $\{x_2, x_3\}$ è ancora un sistema di generatori di X . Siccome $X = \mathbb{R}^2$ ha dimensione 2, allora $\{x_2, x_3\}$ è una base di X . $\square$

ESEMPIO 2.7.12. Consideriamo il vettore

$$x_1 = \begin{pmatrix} 1 \\ 1 \end{pmatrix}$$

dello spazio vettoriale $X = \mathbb{R}^2$. Vediamo come costruire una base di X seguendo la proposizione 2.7.10.

Il vettore x_1 genera il sottospazio vettoriale di X formato dai vettori

$$\begin{pmatrix} \alpha \\ \alpha \end{pmatrix}$$

dove α è un numero reale qualsiasi. Scegliamo un vettore x_2 non appartenente a tale sottospazio, per esempio

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

Allora x_1 e x_2 sono sicuramente linearmente indipendenti e quindi $\{x_1, x_2\}$ è una base di X perché $\dim X = 2$. $\square$

Si noti che nell'esempio precedente avremmo potuto scegliere un altro vettore x_2 , come

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

o anche

$$\begin{pmatrix} 1 \\ -1 \end{pmatrix}$$

o in altri infiniti modi, perché ci sono infiniti vettori non appartenenti al sottospazio vettoriale generato da x_1 .

CAPITOLO 3

Applicazioni lineari e matrici

3.1. Definizione di applicazione lineare

DEFINIZIONE 3.1.1. Siano $\mathbf{X}$ e $\mathbf{Y}$ spazi vettoriali. Una applicazione $T : \mathbf{X} \longrightarrow \mathbf{Y}$ si chiama *lineare* se rispetta le operazioni di spazio vettoriale di $\mathbf{X}$ e $\mathbf{Y}$, cioè

$$T(\mathbf{x}_1 + \mathbf{x}_2) = T(\mathbf{x}_1) + T(\mathbf{x}_2), \quad \forall \mathbf{x}_1, \mathbf{x}_2 \in \mathbf{X},$$

$T(\lambda \mathbf{x}) = \lambda T(\mathbf{x}), \quad \forall \mathbf{x} \in \mathbf{X}, \forall \lambda \in \mathbb{R}.$ Una applicazione lineare si chiama anche *trasformazione lineare*, ovvero *operatore lineare*.

NOTA 3.1.2. L'immagine tramite T di una combinazione lineare di vettori in $\mathbf{X}$ è la combinazione lineare degli immagini di questi vettori con i coefficienti corrispondenti :

$$T(\lambda_1 \mathbf{x}_1 + \dots + \lambda_k \mathbf{x}_k) = \lambda_1 T(\mathbf{x}_1) + \dots + \lambda_k T(\mathbf{x}_k).$$

□

ESEMPIO 3.1.3. Siano $X = Y = \mathbb{R}^2$ e consideriamo le applicazioni $T : \mathbb{R}^2 \rightarrow \mathbb{R}^2$. Allora l'applicazione:

$$T \begin{pmatrix} x \\ y \end{pmatrix} = \begin{pmatrix} 2x - 3y \\ -x + y \end{pmatrix}$$

è lineare, infatti soddisfa le due proprietà della definizione 3.1.1, come si può verificare direttamente. □

Ma vediamo anche un esempio di applicazione non lineare.

ESEMPIO 3.1.4. Consideriamo ancora le applicazioni $T : \mathbb{R}^2 \rightarrow \mathbb{R}^2$.
L'applicazione:

$$T \begin{pmatrix} x \\ y \end{pmatrix} = \begin{pmatrix} x^2 + y \\ y \end{pmatrix}$$

non è lineare, perché compare un esponente 2 in una delle incognite.
Da ciò segue per esempio che:

$$T \begin{pmatrix} 3 \\ 0 \end{pmatrix} = \begin{pmatrix} 9 \\ 0 \end{pmatrix} \neq \begin{pmatrix} 5 \\ 0 \end{pmatrix} = T \begin{pmatrix} 1 \\ 0 \end{pmatrix} + T \begin{pmatrix} 2 \\ 0 \end{pmatrix}$$

□

Si verifica facilmente che, se $\mathbf{X}, \mathbf{Y}, \mathbf{Z}$ sono spazi vettoriali e $T : \mathbf{X} \rightarrow \mathbf{Y}, S : \mathbf{Y} \rightarrow \mathbf{Z}$ sono applicazioni lineari, allora la *composizione*

$$S \circ T : \mathbf{X} \ni \mathbf{x} \mapsto S(T(\mathbf{x})) \in \mathbf{Z}$$

è pure una applicazione lineare.

3.2. Immagine e nucleo di una applicazione lineare

Sia $T : \mathbf{X} \rightarrow \mathbf{Y}$ una applicazione lineare. Allora

$$\text{Ker}(T) := \{\mathbf{x} \in \mathbf{X} ; T(\mathbf{x}) = \mathbf{0}_{\mathbf{Y}}\} \subset \mathbf{X},$$

chiamato il *nucleo di T* , è un sottospazio lineare di $\mathbf{X}$. Infatti, $\text{Ker}(T)$ contiene $\mathbf{0}_{\mathbf{X}}$ perché $T(\mathbf{0}_{\mathbf{X}}) = \mathbf{0}_{\mathbf{Y}}$,

$$\begin{aligned} \mathbf{x}_1, \mathbf{x}_2 \in \text{Ker}(T) &\implies T(\mathbf{x}_1 + \mathbf{x}_2) = T(\mathbf{x}_1) + T(\mathbf{x}_2) = \mathbf{0}_{\mathbf{Y}} \\ &\implies \mathbf{x}_1 + \mathbf{x}_2 \in \text{Ker}(T) \end{aligned}$$

e

$$\begin{aligned} \mathbf{x} \in \text{Ker}(T), \lambda \in \mathbb{R} &\implies T(\lambda \mathbf{x}) = \lambda T(\mathbf{x}) = \mathbf{0}_{\mathbf{Y}} \\ &\implies \lambda \mathbf{x} \in \text{Ker}(T). \end{aligned}$$

Poi,

$$T(\mathbf{X}) := \{T(\mathbf{x}) ; \mathbf{x} \in \mathbf{X}\} \subset \mathbf{Y},$$

chiamato l'*immagine di T* , è un sottospazio lineare di $\mathbf{Y}$. Infatti, $T(\mathbf{X})$ contiene $\mathbf{0}_{\mathbf{Y}} = T(\mathbf{0}_{\mathbf{X}})$,

$$T(\mathbf{x}_1) + T(\mathbf{x}_2) = T(\mathbf{x}_1 + \mathbf{x}_2) \in T(\mathbf{X})$$

e

$$\lambda T(\mathbf{x}) = T(\lambda \mathbf{x}) \in T(\mathbf{X}).$$

Esiste una relazione notevole fra la dimensione del nucleo e la dimensione dell'immagine di una applicazione lineare:

TEOREMA 3.2.1. (Teorema della dimensione.) *Se $\mathbf{X}$ e $\mathbf{Y}$ sono spazi vettoriali di dimensione finita e $T : \mathbf{X} \longrightarrow \mathbf{Y}$ è una applicazione lineare, allora*

$$\dim(\text{Ker}(T)) + \dim(T(\mathbf{X})) = \dim(\mathbf{X}).$$

* DIMOSTRAZIONE. Sia $\mathbf{v}_1, \dots, \mathbf{v}_p$ una base di $\text{ker}(T)$ e $\mathbf{f}_1, \dots, \mathbf{f}_q$ una base di $T(\mathbf{X})$. Allora esistono $\mathbf{x}_1, \dots, \mathbf{x}_q \in \mathbf{X}$ tale che

$$\mathbf{f}_1 = T(\mathbf{x}_1), \dots, \mathbf{f}_q = T(\mathbf{x}_q).$$

Verifichiamo che i vettori $\mathbf{v}_1, \dots, \mathbf{v}_p, \mathbf{x}_1, \dots, \mathbf{x}_q \in \mathbf{X}$ sono linearmente indipendenti. Infatti, se

$$\alpha_1 \mathbf{v}_1 + \dots + \alpha_p \mathbf{v}_p + \beta_1 \mathbf{x}_1 + \dots + \beta_q \mathbf{x}_q = \mathbf{0}_{\mathbf{X}}, \quad (3.2.1)$$

allora

$$\begin{aligned} \beta_1 \mathbf{f}_1 + \dots + \beta_q \mathbf{f}_q &= \beta_1 T(\mathbf{x}_1) + \dots + \beta_q T(\mathbf{x}_q) \\ &= T(\beta_1 \mathbf{x}_1 + \dots + \beta_q \mathbf{x}_q) \\ &= T(\underbrace{\alpha_1 \mathbf{v}_1 + \dots + \alpha_p \mathbf{v}_p}_{\in \text{Ker}(T)} + \beta_1 \mathbf{x}_1 + \dots + \beta_q \mathbf{x}_q) \\ &= T(\mathbf{0}_{\mathbf{X}}) = \mathbf{0}_{\mathbf{Y}} \end{aligned}$$

e dall'indipendenza lineare dei vettori $\mathbf{f}_1, \dots, \mathbf{f}_q$ risulta che tutti i coefficienti $\beta_1, \dots, \beta_q$ si annullano. Ora da (3.2.1) risulta che $\alpha_1 \mathbf{v}_1 + \dots + \alpha_p \mathbf{v}_p = \mathbf{0}_{\mathbf{X}}$ e l'indipendenza lineare del sistema $\mathbf{v}_1, \dots, \mathbf{v}_p$ implica che anche i coefficienti $\alpha_1, \dots, \alpha_p$ si annullano.

Verifichiamo anche che $\mathbf{v}_1, \dots, \mathbf{v}_p, \mathbf{x}_1, \dots, \mathbf{x}_q \in \mathbf{X}$ generano lo spazio vettoriale $\mathbf{X}$. Sia $\mathbf{x} \in \mathbf{X}$ arbitrario. Poiché $\mathbf{f}_1, \dots, \mathbf{f}_q$ generano $T(\mathbf{X})$, esistono dei scalari $\mu_1, \dots, \mu_q$ tale che

$$T(\mathbf{x}) = \mu_1 \mathbf{f}_1 + \dots + \mu_q \mathbf{f}_q = T(\mu_1 \mathbf{x}_1 + \dots + \mu_q \mathbf{x}_q),$$

cioè

$$\mathbf{x} - (\mu_1 \mathbf{x}_1 + \dots + \mu_q \mathbf{x}_q) \in \text{Ker}(T).$$

Ma $\mathbf{v}_1, \dots, \mathbf{v}_p$ generano $\text{ker}(T)$, perciò esistono sei scalari $\lambda_1, \dots, \lambda_p$ tale che

$$\mathbf{x} - (\mu_1 \mathbf{x}_1 + \dots + \mu_q \mathbf{x}_q) = \lambda_1 \mathbf{v}_1 + \dots + \lambda_p \mathbf{v}_p$$

e così

$$\mathbf{x} = \lambda_1 \mathbf{v}_1 + \dots + \lambda_p \mathbf{v}_p + \mu_1 \mathbf{x}_1 + \dots + \mu_q \mathbf{x}_q.$$

Concludiamo che $\mathbf{v}_1, \dots, \mathbf{v}_p, \mathbf{x}_1, \dots, \mathbf{x}_q$ è una base di $\mathbf{X}$ e così $\dim(\mathbf{X})$ è uguale a $p + q = \dim(\text{Ker}(T)) + \dim(T(\mathbf{X}))$.

□

NOTA 3.2.2. Una applicazione lineare $T : \mathbf{X} \longrightarrow \mathbf{Y}$ è *iniettiva*, cioè applica vettori diversi in vettori diversi, ovvero

$$\mathbf{x}_1, \mathbf{x}_2 \in \mathbf{X}, T(\mathbf{x}_1) = T(\mathbf{x}_2) \implies \mathbf{x}_1 = \mathbf{x}_2, \quad (3.2.2)$$

se e solo se $\text{ker}(T) = \{\mathbf{0}_{\mathbf{X}}\}$, ovvero

$$\mathbf{x} \in \mathbf{X}, T(\mathbf{x}) = \mathbf{0}_{\mathbf{Y}} \implies \mathbf{x} = \mathbf{0}_{\mathbf{X}}. \quad (3.2.3)$$

Infatti, (3.2.2) implica ovviamente (3.2.3), perciò se T è iniettiva, allora $\text{ker}(T) = \{\mathbf{0}_{\mathbf{X}}\}$. Viceversa, se $\text{ker}(T) = \{\mathbf{0}_{\mathbf{X}}\}$ e $\mathbf{x}_1, \mathbf{x}_2 \in \mathbf{X}$ sono tali che $T(\mathbf{x}_1) = T(\mathbf{x}_2)$, allora $T(\mathbf{x}_1 - \mathbf{x}_2) = \mathbf{0}_{\mathbf{Y}}$ e, grazie a (3.2.3), risulta che $\mathbf{x}_1 - \mathbf{x}_2 = \mathbf{0}_{\mathbf{X}}$, cioè $\mathbf{x}_1 = \mathbf{x}_2$. □

Se $\mathbf{X}$ e $\mathbf{Y}$ sono di dimensione finita, allora, per il teorema della dimensione 3.2.1,

$$T \text{ iniettiva} \iff \dim(\text{Ker}(T)) = 0 \iff \dim(T(\mathbf{X})) = \dim(\mathbf{X}).$$

Poiché $T(\mathbf{X})$ è un sottospazio lineare di $\mathbf{Y}$ e così $\dim(T(\mathbf{X})) \leq \dim(\mathbf{Y})$, per l'iniettività di T è necessario che $\dim(\mathbf{X}) \leq \dim(\mathbf{Y})$.

Ricordiamo che una applicazione lineare $T : \mathbf{X} \longrightarrow \mathbf{Y}$ si chiama *surgettiva* se $T(\mathbf{X}) = \mathbf{Y}$. Se $\mathbf{X}$ e $\mathbf{Y}$ sono di dimensione finita, allora,

tenendo conto del teorema della dimensione 3.2.1,

$$\begin{aligned} T \text{ surgettiva} &\iff \dim(T(\mathbf{X})) = \dim(\mathbf{Y}) \\ &\iff \dim(\text{Ker}(T)) = \dim(\mathbf{X}) - \dim(\mathbf{Y}). \end{aligned}$$

In particolare, per la surgettività di T è necessario che $\dim(\mathbf{X}) - \dim(\mathbf{Y})$ sia ≥ 0 , cioè che $\dim(\mathbf{X}) \geq \dim(\mathbf{Y})$.

Se l'applicazione lineare $T : \mathbf{X} \longrightarrow \mathbf{Y}$ è *bigettiva*, cioè simultaneamente iniettiva e surgettiva, allora dobbiamo avere $\dim(\mathbf{X}) = \dim(\mathbf{Y})$. Se $\mathbf{X}$ e $\mathbf{Y}$ sono di dimensione finita uguale, allora dal teorema della dimensione 3.2.1 risulta subito che

$$\dim(\text{Ker}(T)) = 0 \iff \dim(T(\mathbf{X})) = \dim(\mathbf{Y}).$$

Quindi, in questo caso,

$$\begin{aligned} T \text{ bigettiva} &\iff \dim(\mathbf{X}) = \dim(\mathbf{Y}) \ \& \ T \text{ iniettiva} \\ &\iff \dim(\mathbf{X}) = \dim(\mathbf{Y}) \ \& \ T \text{ surgettiva.} \end{aligned} \quad (3.2.4)$$

Se T è bigettiva, allora, associando ad ogni $\mathbf{y} \in \mathbf{Y}$ l'unico $\mathbf{x} \in \mathbf{X}$ che soddisfa $\mathbf{y} = T(\mathbf{x})$, otteniamo l'applicazione inversa $T^{-1} : \mathbf{Y} \longrightarrow \mathbf{X}$ di T , per cui abbiamo

$$T^{-1} \circ T = \mathbb{I}_{\mathbf{X}}, \quad T \circ T^{-1} = \mathbb{I}_{\mathbf{Y}},$$

ove $\mathbb{I}_{\mathbf{X}} : \mathbf{X} \ni \mathbf{x} \longmapsto \mathbf{x} \in \mathbf{X}$ e $\mathbb{I}_{\mathbf{Y}} : \mathbf{Y} \ni \mathbf{y} \longmapsto \mathbf{y} \in \mathbf{Y}$ indicano le rispettive applicazioni identiche.

Si verifica facilmente che, se $T : \mathbf{X} \longrightarrow \mathbf{Y}$, $S : \mathbf{Y} \longrightarrow \mathbf{Z}$ sono applicazioni lineari, allora

$$T, S \text{ iniettive} \implies S \circ T \text{ iniettiva} \implies T \text{ iniettiva}, \quad (3.2.5)$$

$$T, S \text{ surgettive} \implies S \circ T \text{ surgettiva} \implies S \text{ surgettiva}. \quad (3.2.6)$$

Tenendo conto di (3.2.4), risulta che se $\mathbf{X}$, $\mathbf{Y}$ hanno la stessa dimensione finita, allora

$$S \circ T \text{ iniettiva} \implies T \text{ bigettiva},$$

mentre se $\mathbf{Y}$, $\mathbf{Z}$ hanno la stessa dimensione finita, allora

$$S \circ T \text{ surgettiva} \implies S \text{ bigettiva}.$$

3.3. Spazi vettoriali di polinomi ed applicazioni lineari

DEFINIZIONE 3.3.1. (**Polinomi.**) Un *polinomio* (con coefficienti reali) è una espressione $P(x)$ della forma

$$a_0 + a_1 x + \dots + a_n x^n \quad \text{ossia} \quad \sum_{k=0}^n a_k x^k,$$

ove $a_0, a_1, \dots, a_n$ sono *coefficienti* reali, $n \geq 0$ è un intero ed x è una *variabile*. Due polinomi

$$\sum_{k=0}^n a_k x^k, \quad \sum_{k=0}^m b_k x^k$$

con $m \geq n$ sono considerati uguali se $a_k = b_k$ per $k \leq n$ e $b_k = 0$ per $k > n$. L'insieme $\mathcal{P}$ di tutti i polinomi è uno spazio vettoriale rispetto all'addizione

$$\sum_{k=0}^n a_k x^k + \sum_{k=0}^n b_k x^k := \sum_{k=0}^n (a_k + b_k) x^k$$

ed alla moltiplicazione con scalari

$$\lambda \sum_{k=0}^n a_k x^k := \sum_{k=0}^n (\lambda a_k) x^k.$$

Il *grado* di un polinomio non zero $P(x) = \sum_{k=0}^n a_k x^k$ è il più grande intero $0 \leq d \leq n$ per quale

$$a_d \neq 0, \quad a_k = 0 \quad \text{per } k > d$$

e si indica con $\deg(P)$. Il grado del polinomio zero è per definizione $-\infty$. Definiamo anche il prodotto di due polinomi

$$P(x) = \sum_{k=0}^n a_k x^k, \quad Q(x) = \sum_{k=0}^m b_k x^k$$

come segue:

$$(PQ)(x) := \sum_{j=0}^{n+m} \left(\sum_{\substack{0 \leq k_1 \leq n \\ 0 \leq k_2 \leq m \\ k_1 + k_2 = j}} a_{k_1} b_{k_2} \right) x^j,$$

ove $a_{k_1} = 0$ per $k_1 > n$ e $b_{k_2} = 0$ per $k_2 > m$. Il grado di PQ è uguale alla somma $\deg(P) + \deg(Q)$.

La moltiplicazione dei polinomi è associativa e commutativa. Inoltre vale la proprietà di distributività:

$$(P + Q)R = PR + QR.$$

Se $P(x) = \sum_{k=0}^n a_k x^k$ è un polinomio ed r è un numero reale, allora $P(r)$ è il numero che si ottiene sostituendo $x = r$ in $P(x)$, cioè $\sum_{k=0}^n a_k r^k$.

È noto che per tutti i polinomi P e D , $D \neq 0$, possiamo dividere P con D con resto: esistono polinomi Q e R , e sono unici, con la proprietà che

$$P = DQ + R, \quad \deg(R) < \deg(D).$$

Q si chiama il *quoziente* della divisione e verifica $\deg(D) + \deg(Q) = \deg(P)$. R si chiama il *resto* della divisione.

Se $D(x) = d_0 + d_1 x$, allora R è la costante $P\left(-\frac{d_0}{d_1}\right)$.

Per ogni intero $n \geq 0$ denotiamo con $\mathcal{P}_n$ l'insieme di tutti i polinomi di grado al più $\leq n$. Allora $\mathcal{P}_n$ è un sottospazio lineare di $\mathcal{P}$, di dimensione $n + 1$, avendo come base

$$1, x, \dots, x^n.$$

La derivata $P'(x)$ di un polinomio $P(x) = \sum_{k=0}^n a_k x^k$ è

$$a_1 x + 2a_2 x + \dots + n a_n x^{n-1} = \sum_{k=1}^n k a_k x^{k-1}.$$

Per il grado della derivata di un polinomio P non costante abbiamo

$$\deg(P') = \deg(P) - 1.$$

La derivata di un polinomio costante è 0.

Regole di calcolo per la derivazione:

$$(P + Q)' = P' + Q', \quad (\lambda P)' = \lambda P', \quad (PQ)' = P'Q + PQ'.$$

3.4. Esercizi sulle applicazioni lineari su spazi di polinomi

ESERCIZIO 3.4.1. Si trovi il nucleo e l'immagine dell'applicazione lineare $T : \mathcal{P}_3 \longrightarrow \mathcal{P}_4$ che assegna ad ogni $P \in \mathcal{P}_3$ il polinomio $(x + 2)P(x)$.

SVOLGIMENTO. $\ker(T) = \{0\}$, cioè T è iniettiva. Infatti, ogni $0 \neq P \in \mathcal{P}_3$ ha un grado ≥ 0 ed allora il grado di $T(P)$, cioè di $(x + 2)P(x)$, è uguale ad $1 + \deg(P) \geq 1$.

$T(\mathcal{P}_3)$ consiste da tutti i polinomi di grado ≤ 4 che sono divisibili con $x - 2$, cioè quelli si annullano in -2 . In particolare, T non è surgettiva. $\square$

ESERCIZIO 3.4.2. Si trovi il nucleo e l'immagine dell'applicazione lineare $T : \mathcal{P}_5 \longrightarrow \mathcal{P}_1$ che assegna ad ogni $P \in \mathcal{P}_5$ il resto nella divisione di P con $x^2 + 1$.

SVOLGIMENTO. $\ker(T)$ consiste da tutti i polinomi di grado ≤ 5 che sono divisibili con $x^2 + 1$. In particolare, T non è iniettiva.

L'immagine di T è uguale a $\mathcal{P}_1$, cioè T è surgettiva. Infatti, il resto della divisione con $x^2 + 1$ di un qualsiasi $P \in \mathcal{P}_1$ è uguale a P stesso. $\square$

ESERCIZIO 3.4.3. Si trovi il nucleo e l'immagine dell'applicazione lineare $T : \mathcal{P}_4 \longrightarrow \mathcal{P}_3$ che assegna ad ogni $P \in \mathcal{P}_4$ la sua derivata.

SVOLGIMENTO. $\ker(T)$ consiste di tutti i polinomi costanti. In particolare, T non è iniettiva.

D'altro canto, T è surgettiva. Infatti, il polinomio generico

$$P(x) = a_0 + a_1 x + a_2 x^2 + a_3 x^3,$$

è la derivata di

$$a_0 x + \frac{a_1}{2} x^2 + \frac{a_2}{3} x^3 + \frac{a_3}{4} x^4.$$

$\square$

ESERCIZIO 3.4.4. Si trovino il nucleo e l'immagine dell'applicazione lineare $T : \mathcal{P}_3 \longrightarrow \mathcal{P}_3$ che assegna ad ogni $P \in \mathcal{P}_3$ la derivata di $(x - 3)P(x)$.

SVOLGIMENTO. T è iniettiva. Infatti, se $P \in \mathcal{P}_3$ è tale che $T(P) = 0$, allora il polinomio $(x-3)P(x)$ è costante, che è possibile solo se $P = 0$.

Poiché T è una applicazione lineare fra spazi vettoriali della stessa dimensione finita, dall'injectività di T risulta la sua surgettività.

□

3.5. Applicazioni lineari fra $\mathbb{R}^n$ e $\mathbb{R}^m$ e calcolo con matrici

Abbiamo visto che l'immagine tramite T di una combinazione lineare di vettori in $\mathbb{R}^n$ è la combinazione lineare degli immagini di questi vettori con i coefficienti corrispondenti :

$$T(\lambda_1 \mathbf{x}_1 + \dots + \lambda_k \mathbf{x}_k) = \lambda_1 T(\mathbf{x}_1) + \dots + \lambda_k T(\mathbf{x}_k).$$

Quindi una applicazione lineare $T : \mathbb{R}^n \longrightarrow \mathbb{R}^m$ è completamente determinata dai valori assunti negli elementi della base naturale

$$\mathbf{e}_1 = \begin{pmatrix} 1 \\ 0 \\ 0 \\ \vdots \\ 0 \end{pmatrix}, \mathbf{e}_2 = \begin{pmatrix} 0 \\ 1 \\ 0 \\ \vdots \\ 0 \end{pmatrix}, \mathbf{e}_3 = \begin{pmatrix} 0 \\ 0 \\ 1 \\ \vdots \\ 0 \end{pmatrix}, \dots, \mathbf{e}_n = \begin{pmatrix} 0 \\ 0 \\ 0 \\ \vdots \\ 1 \end{pmatrix}$$

di $\mathbb{R}^n$. Infatti, per ogni vettore in $\mathbb{R}^n$

$$\mathbf{x} = \begin{pmatrix} x_1 \\ x_2 \\ x_3 \\ \vdots \\ x_n \end{pmatrix} = \sum_{j=1}^n x_j \mathbf{v}_j$$

abbiamo

$$T(\mathbf{x}) = T\left(\sum_{j=1}^n x_j \mathbf{v}_j\right) = \sum_{j=1}^n x_j T(\mathbf{e}_j). \quad (3.5.1)$$

Pertanto, se i vettori $T(\mathbf{e}_1), \dots, T(\mathbf{e}_n) \in \mathbb{R}^m$, scritti esplicitamente, sono

$$T(\mathbf{e}_1) = \begin{pmatrix} \alpha_{11} \\ \alpha_{21} \\ \alpha_{31} \\ \vdots \\ \alpha_{m1} \end{pmatrix}, T(\mathbf{e}_2) = \begin{pmatrix} \alpha_{12} \\ \alpha_{22} \\ \alpha_{32} \\ \vdots \\ \alpha_{m2} \end{pmatrix}, \dots, T(\mathbf{e}_n) = \begin{pmatrix} \alpha_{1n} \\ \alpha_{2n} \\ \alpha_{3n} \\ \vdots \\ \alpha_{mn} \end{pmatrix},$$

possiamo riscrivere (3.5.1) come segue :

$$\begin{aligned} T(\mathbf{x}) &= x_1 \begin{pmatrix} \alpha_{11} \\ \alpha_{21} \\ \alpha_{31} \\ \vdots \\ \alpha_{m1} \end{pmatrix} + x_2 \begin{pmatrix} \alpha_{12} \\ \alpha_{22} \\ \alpha_{32} \\ \vdots \\ \alpha_{m2} \end{pmatrix} + \dots + x_n \begin{pmatrix} \alpha_{1n} \\ \alpha_{2n} \\ \alpha_{3n} \\ \vdots \\ \alpha_{mn} \end{pmatrix} \\ &\stackrel{!}{=} \underbrace{\begin{pmatrix} \alpha_{11} & \alpha_{12} & \dots & \alpha_{1n} \\ \alpha_{21} & \alpha_{22} & \dots & \alpha_{2n} \\ \alpha_{31} & \alpha_{32} & \dots & \alpha_{3n} \\ \vdots & \vdots & \ddots & \vdots \\ \alpha_{m1} & \alpha_{m2} & \dots & \alpha_{mn} \end{pmatrix}}_{=: A} \underbrace{\begin{pmatrix} x_1 \\ x_2 \\ x_3 \\ \vdots \\ x_n \end{pmatrix}}_{=\mathbf{x}}, \end{aligned}$$

ove il *prodotto* della matrice A di m righe ed n colonne, detta matrice $m \times n$, con il vettore $\mathbf{x}$ di n componenti è il vettore di m componenti

$$\begin{pmatrix} \alpha_{11} x_1 + \alpha_{12} x_2 + \alpha_{13} x_3 + \dots + \alpha_{1n} x_n \\ \alpha_{21} x_1 + \alpha_{22} x_2 + \alpha_{23} x_3 + \dots + \alpha_{2n} x_n \\ \alpha_{31} x_1 + \alpha_{32} x_2 + \alpha_{33} x_3 + \dots + \alpha_{3n} x_n \\ \vdots \\ \alpha_{m1} x_1 + \alpha_{m2} x_2 + \alpha_{m3} x_3 + \dots + \alpha_{mn} x_n \end{pmatrix}.$$

La matrice $m \times n$ ottenuta si chiama *la matrice associata a T* (rispetto alle basi naturali di $\mathbb{R}^n$ e $\mathbb{R}^m$).

Viceversa, se A è una matrice $m \times n$, allora l'applicazione

$$T_A : \mathbb{R}^n \ni \mathbf{x} \longmapsto A\mathbf{x} \in \mathbb{R}^m,$$

ove il prodotto $A\mathbf{x}$ è definito come di cui sopra, è una applicazione lineare e la matrice associata a T_A è A . Perciò

$$A \longmapsto T_A$$

è una corrispondenza biunivoca tra le matrici $m \times n$ e le applicazioni lineari $\mathbb{R}^n \rightarrow \mathbb{R}^m$. ricordiamo che la colonna j -esima di A è uguale all'immagine dell'elemento j -esimo della base naturale di $\mathbb{R}^n$ tramite T_A .

3.6. Spazi vettoriali di matrici

Cominciamo con due notazioni. La matrice dell'applicazione identicamente zero $\mathbb{R}^n \ni \mathbf{x} \mapsto \mathbf{0}_m \in \mathbb{R}^m$ è

$$\mathbf{0}_{mn} := \begin{pmatrix} 0 & 0 & \dots & 0 \\ 0 & 0 & \dots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \dots & 0 \end{pmatrix} \in M_{mn} \text{ (}\mathbf{0}_n \text{ se } m = n\text{)},$$

mentre la matrice dell'applicazione identica $\mathbb{R}^n \ni \mathbf{x} \mapsto \mathbf{x} \in \mathbb{R}^n$ è la *matrice unità* (o *matrice identità*) d'ordine n :

$$\mathbb{I}_n := \begin{pmatrix} 1 & 0 & \dots & 0 \\ 0 & 1 & \dots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \dots & 1 \end{pmatrix} \in M_n.$$

NOTAZIONE 3.6.1. (Applicazioni lineari da $\mathbb{R}^n$ a $\mathbb{R}^m$.) Indichiamo con $\mathcal{L}(\mathbb{R}^n, \mathbb{R}^m)$ l'insieme di tutte le applicazioni lineari $\mathbb{R}^n \rightarrow \mathbb{R}^m$, e con M_{mn} l'insieme di tutte le matrici $m \times n$. Poniamo anche $\mathcal{L}(\mathbb{R}^n)$ per $\mathcal{L}(\mathbb{R}^n, \mathbb{R}^n)$ ed M_n per M_{nn} .

NOTA 3.6.2. L'insieme $\mathcal{L}(\mathbb{R}^n, \mathbb{R}^m)$ è uno spazio vettoriale rispetto alle operazioni seguenti:

- la somma $T_1 + T_2$ di $T_1, T_2 \in \mathcal{L}(\mathbb{R}^n, \mathbb{R}^m)$ è definita da $(T_1 + T_2)(\mathbf{x}) := T_1(\mathbf{x}) + T_2(\mathbf{x})$,
- il prodotto λT di $T \in \mathcal{L}(\mathbb{R}^n, \mathbb{R}^m)$ è definita da $(\lambda T)(\mathbf{x}) := \lambda T(\mathbf{x})$.

Infatti, se a $T_1, T_2 \in \mathcal{L}(\mathbb{R}^n, \mathbb{R}^m)$ corrispondono le matrici

$$A_1 = \left(\alpha_{jk} \right)_{\substack{1 \leq j \leq m \\ 1 \leq k \leq n}}, A_2 = \left(\beta_{jk} \right)_{\substack{1 \leq j \leq m \\ 1 \leq k \leq n}},$$

allora a $T_1 + T_2$ corrisponde

$$A_1 + A_2 := \left(\alpha_{jk} + \beta_{jk} \right)_{\substack{1 \leq j \leq m \\ 1 \leq k \leq n}},$$

mentre se a $T \in \mathcal{L}(\mathbb{R}^n, \mathbb{R}^m)$ corrisponde

$$A = \left(\gamma_{jk} \right)_{\substack{1 \leq j \leq m \\ 1 \leq k \leq n}},$$

allora a λT corrisponde

$$\lambda A := \left(\lambda \gamma_{jk} \right)_{\substack{1 \leq j \leq m \\ 1 \leq k \leq n}}.$$

□

Così abbiamo definito la somma di due matrici dello stesso tipo, e la moltiplicazione di una matrice per un scalare reale. Ora definiamo il prodotto di due matrici (un'operazione in più rispetto a quelle necessarie per definire gli spazi vettoriali).

DEFINIZIONE 3.6.3. (*Prodotto di matrici.*) Siano $T : \mathbb{R}^n \rightarrow \mathbb{R}^m$ e $S : \mathbb{R}^m \rightarrow \mathbb{R}^p$ applicazioni lineari tali che il dominio di S sia uguale allo spazio vettoriale immagine di T . Allora possiamo considerare la composizione $S \circ T : \mathbb{R}^n \rightarrow \mathbb{R}^p$, che è anch'essa lineare. Se $A = (\alpha_{jk})_{j,k} \in M_{mn}$ e $B = (\beta_{hl})_{h,l} \in M_{pm}$ sono le matrici associate a T ed S , allora la *matrice prodotto* BA associata ad $S \circ T$ è la matrice di tipo $p \times n$ data dalla formula di *moltiplicazione righe per colonne* seguente:

$$BA := \left((\beta_{h,1} \ \dots \ \beta_{h,m}) \cdot \begin{pmatrix} \alpha_{1,k} \\ \vdots \\ \alpha_{m,k} \end{pmatrix} \right)_{\substack{1 \leq h \leq p \\ 1 \leq k \leq n}} = \left(\sum_{j=1}^m \beta_{h,j} \alpha_{j,k} \right)_{\substack{1 \leq h \leq p \\ 1 \leq k \leq n}}.$$

Perciò il prodotto BA di due matrici B ed A esiste solo se B ha tante colonne quante sono le righe di A . In questo caso BA ha tante righe quante ne ha B , e tante colonne quante ne ha A . Se la h -esima riga di

B e la k -esima colonna di A sono

$$\text{--- } \mathbf{r}_h \text{ ---} \quad \text{rispettivamente} \quad \begin{array}{c} | \\ \mathbf{c}_k \\ | \end{array},$$

allora $\mathbf{r}_h$ e $\mathbf{c}_k$ hanno lo stesso numero di componenti e l'elemento di BA all'incrocio della h -esima riga e della k -esima colonna è

$$\begin{aligned} & (\text{primo componente di } \mathbf{r}_h) \cdot (\text{primo componente di } \mathbf{c}_k) \\ & + (\text{secondo componente di } \mathbf{r}_h) \cdot (\text{secondo componente di } \mathbf{c}_k) \\ & \quad \vdots \\ & + (\text{l'ultimo componente di } \mathbf{r}_h) \cdot (\text{l'ultimo componente di } \mathbf{c}_k). \end{aligned}$$

Osserviamo che, per $A \in M_{mn}$ e $\mathbf{x} \in \mathbb{R}^n$, il vettore $A\mathbf{x}$, considerato come matrice $m \times 1$, è il prodotto della matrice $m \times n$ A e la matrice $n \times 1$ $\mathbf{x}$.

Se $T : \mathbb{R}^n \longrightarrow \mathbb{R}^m$ è una applicazione lineare, allora

$$\begin{aligned} \ker(T) &:= \{\mathbf{x} \in \mathbb{R}^n ; T(\mathbf{x}) = \mathbf{0}_m\} \subset \mathbb{R}^n \text{ (il nucleo di } T) \\ T(\mathbb{R}^n) &:= \{T(\mathbf{x}) ; \mathbf{x} \in \mathbb{R}^n\} \subset \mathbb{R}^m \text{ (l'immagine di } T) \end{aligned}$$

sono sottospazi lineari. Ricordiamo cosa dice il teorema della dimensione 3.2.1 in questo caso:

$$\dim(\ker(T)) + \dim(T(\mathbb{R}^n)) = \dim(\mathbb{R}^n) = n.$$

Ricordiamo (Nota 3.2.2) che T è *iniettiva* se e solo se $\ker(T) = \{\mathbf{0}_n\}$. Di conseguenza, usando anche il teorema della dimensione 3.2.1, risulta che T è iniettiva se e solo se $\dim(T(\mathbb{R}^n)) = n$. Poiché $\dim(T(\mathbb{R}^n)) \leq \dim(\mathbb{R}^m) = m$, per l'iniettività di T è necessario che $n \leq m$.

D'altro canto, dal Teorema della dimensione 3.2.1 risulta che T è *surgettiva*, cioè $T(\mathbb{R}^n) = \mathbb{R}^m$, se e solo se $\dim(\ker(T)) = n - m$. Quindi per la surgettività di T è necessario che $n \geq m$.

Sia ora

$$A = \begin{pmatrix} \alpha_{11} & \alpha_{12} & \dots & \alpha_{1n} \\ \alpha_{21} & \alpha_{22} & \dots & \alpha_{2n} \\ \vdots & \vdots & \ddots & \vdots \\ \alpha_{m1} & \alpha_{m2} & \dots & \alpha_{mn} \end{pmatrix}$$

una matrice $m \times n$. Allora $\ker(T_A) = \{\mathbf{x} \in \mathbb{R}^n; A\mathbf{x} = \mathbf{0}_m\}$ è l'insieme di tutte le soluzioni

$$\mathbf{x} = \begin{pmatrix} x_1 \\ x_2 \\ \vdots \\ x_n \end{pmatrix}$$

del sistema omogeneo $A\mathbf{x} = \mathbf{0}_m$, ovvero

$$\begin{cases} \alpha_{11}x_1 + \alpha_{12}x_2 + \dots + \alpha_{1n}x_n = 0 \\ \alpha_{21}x_1 + \alpha_{22}x_2 + \dots + \alpha_{2n}x_n = 0 \\ \vdots \\ \alpha_{m1}x_1 + \alpha_{m2}x_2 + \dots + \alpha_{mn}x_n = 0 \end{cases}.$$

Perciò T_A è iniettiva esattamente quando questo sistema ha la sola soluzione $\mathbf{x} = \mathbf{0}_n$. Poiché il sistema precedente si scrive !)

$$A\mathbf{x} = \sum_{j=1}^n x_j \mathbf{c}_j \text{ ove } \mathbf{c}_1, \dots, \mathbf{c}_n \text{ sono le colonne di } A, \quad (3.6.1)$$

l'iniettività di T_A è equivalente all'indipendenza lineare delle sue colonne.

ESEMPIO 3.6.4. (*Il prodotto di matrici non è commutativo.*)

Consideriamo le matrici quadrate di ordine $n \times n$. Siano A e B due matrici di tale ordine. Allora possiamo sempre fare il prodotto $A \cdot B$ e $B \cdot A$, ma in generale otteniamo due risultati diversi. Ecco un esempio.

Siano A e B le seguenti matrici quadrate di ordine 2×2 :

$$A = \begin{pmatrix} 1 & 0 \\ 0 & 0 \end{pmatrix}, \quad B = \begin{pmatrix} 0 & 1 \\ 0 & 0 \end{pmatrix}.$$

Allora il prodotto $A \cdot B$ è

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

che è diverso dal prodotto $B \cdot A$

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

Ciò dimostra che il prodotto di matrici *non* è commutativo, cioè in generale $AB \neq BA$. $\square$

3.7. Soluzione di sistemi lineari: eliminazione di Gauss

Il metodo di eliminazione di Gauss consiste nella trasformazione del sistema dato di equazioni lineari in un sistema triangolare superiore che ha le stesse soluzioni. Questo si fa tramite le seguenti operazioni sulle righe della tabella (o matrice) dei coefficienti (che si chiamano *operazioni di riga*):

- scambio di due righe,
- addizione di un multiplo scalare di una riga ad un'altra,
- moltiplicazione di una riga con un scalare non nullo.

Si noti che eseguire una di queste operazioni di riga non cambia le soluzioni del sistema (il nuovo sistema impone le stesse condizioni di quello originale: si dice che è equivalente). Si procede nel modo seguente. Si comincia, se necessario, con lo scambio della prima riga con una che ha il primo elemento diverso da zero. Poi, sommando successivamente opportuni multipli scalari della prima riga alle altre righe, si azzerano i coefficienti nella prima colonna C_1 .

Ora consideriamo la prima colonna C_2 che ha un elemento non nullo diverso dal primo elemento. Scambiamo, se necessario, la seconda riga con una riga *successiva* che ha un elemento non nullo appartenente a C_2 (se scambiassimo con una riga precedente, in questo caso necessariamente la prima, perderemmo il fatto di aver già reso nulli i coefficienti iniziali di quest'ultima (in questo caso, *il* coefficiente iniziale)). Continuiamo nello stesso modo: sommando successivamente multipli scalari adatti della seconda riga alle righe dalla terza in poi, si azzerano i coefficienti sotto l'elemento (non nullo) dell'intersezione della seconda riga con C_2 .

Procediamo in questo modo finché non arriviamo alla barra oltre alla quale si trovano i termini noti.

ESERCIZIO 3.7.1. Si trovino tutte le soluzioni del sistema di equazioni lineari

$$\begin{cases} -x_1 + 2x_2 + 3x_3 + 2x_4 = 1 \\ 2x_1 - 4x_2 + x_3 + x_4 = -1 \\ 3x_1 + x_2 - 3x_3 + x_4 = -6 \\ 4x_1 - x_2 + x_3 + 4x_4 = -6 \end{cases}$$

SOLUZIONE. In questo caso si ottiene:

$$\begin{array}{cccc|c}
 -1 & 2 & 3 & 2 & 1 \\
 2 & -4 & 1 & 1 & -1 \\
 3 & 1 & -3 & 1 & -6 \\
 4 & -1 & 1 & 4 & -6
 \end{array}
 \quad
 \begin{array}{cccc|c}
 -1 & 2 & 3 & 2 & 1 \\
 0 & 0 & 7 & 5 & 1 \\
 0 & 7 & 6 & 7 & -3 \\
 0 & 7 & 13 & 12 & -2
 \end{array}$$

$$\begin{array}{cccc|c}
 -1 & 2 & 3 & 2 & 1 \\
 0 & 7 & 6 & 7 & -3 \\
 0 & 0 & 7 & 5 & 1 \\
 0 & 7 & 13 & 12 & -2
 \end{array}
 \quad
 \begin{array}{cccc|c}
 -1 & 2 & 3 & 2 & 1 \\
 0 & 7 & 6 & 7 & -3 \\
 0 & 0 & 7 & 5 & 1 \\
 0 & 0 & 7 & 5 & 1
 \end{array}$$

$$\begin{array}{cccc|c}
 -1 & 2 & 3 & 2 & 1 \\
 0 & 7 & 6 & 7 & -3 \\
 0 & 0 & 7 & 5 & 1 \\
 0 & 0 & 0 & 0 & 0
 \end{array}
 .$$

Risulta che la soluzione è

$$\begin{aligned}
 x_4 &= \text{arbitrario}, \\
 x_3 &= \frac{1 - 5x_4}{7}, \\
 x_2 &= \frac{-6x_3 - 7x_4 - 3}{7} = \frac{-19x_4 - 27}{49}, \\
 x_1 &= 2x_2 + 3x_3 + 2x_4 - 1 = \frac{-45x_4 - 82}{49}. \quad \square
 \end{aligned}$$

ESERCIZIO 3.7.2. Si trovino tutte le soluzioni del sistema di equazioni lineari

$$\begin{cases}
 -x_1 + 2x_2 + 3x_3 + 2x_4 = 1 \\
 2x_1 - 4x_2 + x_3 + x_4 = -1 \\
 3x_1 + x_2 - 3x_3 + x_4 = -6 \\
 4x_1 - x_2 + x_3 + 3x_4 = -6
 \end{cases}$$

SOLUZIONE. Usando il metodo di eliminazione di Gauss, trasformiamo il sistema dato in un sistema triangolare superiore equivalente:

$$\begin{array}{cccc|c}
 -1 & 2 & 3 & 2 & 1 \\
 2 & -4 & 1 & 1 & -1 \\
 3 & 1 & -3 & 1 & -6 \\
 4 & -1 & 1 & 3 & -6 \\
 \hline
 -1 & 2 & 3 & 2 & 1 \\
 0 & 7 & 6 & 7 & -3 \\
 0 & 0 & 7 & 5 & 1 \\
 0 & 7 & 13 & 11 & -2 \\
 \hline
 -1 & 2 & 3 & 2 & 1 \\
 0 & 7 & 6 & 7 & -3 \\
 0 & 0 & 7 & 5 & 1 \\
 0 & 0 & 7 & 4 & 1 \\
 \hline
 -1 & 2 & 3 & 2 & 1 \\
 0 & 7 & 6 & 7 & -3 \\
 0 & 0 & 7 & 5 & 1 \\
 0 & 0 & 0 & 1 & 0
 \end{array} .$$

Risulta che l'unica soluzione è

$$\begin{aligned}
 x_4 &= 0, \\
 x_3 &= \frac{1 - 5x_4}{7} = \frac{1}{7}, \\
 x_2 &= \frac{-6x_3 - 7x_4 - 3}{7} = -\frac{27}{49}, \\
 x_1 &= 2x_2 + 3x_3 + 2x_4 - 1 = -\frac{82}{49}.
 \end{aligned}$$

□

ESERCIZIO 3.7.3. Si trovino tutte le soluzioni del sistema di equazioni lineari

$$\begin{cases}
 -x_1 + 2x_2 + 3x_3 + 2x_4 = 1 \\
 2x_1 - 4x_2 + x_3 + x_4 = -1 \\
 3x_1 + x_2 - 3x_3 + x_4 = -6 \\
 4x_1 - x_2 + x_3 + 4x_4 = -5
 \end{cases}$$

SOLUZIONE. Usando il metodo di eliminazione di Gauss, trasformiamo il sistema dato in un sistema triangolare superiore equivalente:

$$\begin{array}{cccc|c}
 -1 & 2 & 3 & 2 & 1 \\
 2 & -4 & 1 & 1 & -1 \\
 3 & 1 & -3 & 1 & -6 \\
 4 & -1 & 1 & 4 & -5
 \end{array}
 \quad
 \begin{array}{cccc|c}
 -1 & 2 & 3 & 2 & 1 \\
 0 & 0 & 7 & 5 & 1 \\
 0 & 7 & 6 & 7 & -3 \\
 0 & 7 & 13 & 12 & -1
 \end{array}$$

$$\begin{array}{cccc|c}
 -1 & 2 & 3 & 2 & 1 \\
 0 & 7 & 6 & 7 & -3 \\
 0 & 0 & 7 & 5 & 1 \\
 0 & 7 & 13 & 12 & -1
 \end{array}
 \quad
 \begin{array}{cccc|c}
 -1 & 2 & 3 & 2 & 1 \\
 0 & 7 & 6 & 7 & -3 \\
 0 & 0 & 7 & 5 & 1 \\
 0 & 0 & 7 & 5 & 2
 \end{array}$$

$$\begin{array}{cccc|c}
 -1 & 2 & 3 & 2 & 1 \\
 0 & 7 & 6 & 7 & -3 \\
 0 & 0 & 7 & 5 & 1 \\
 0 & 0 & 0 & 0 & 1
 \end{array}
 .$$

L'ultima equazione del sistema triangolare superiore ottenuto è

$$0x_1 + 0x_2 + 0x_3 + 0x_4 = 1,$$

che non ha soluzione. Perciò il sistema dato non ha nessuna soluzione.

□

ESERCIZIO 3.7.4. Si trovino tutte le soluzioni del sistema di equazioni lineari

$$\begin{cases}
 2x_1 - x_2 + x_3 + 6x_4 - x_5 = -2 \\
 x_1 - 2x_2 + x_3 + 3x_4 - x_5 = -1 \\
 x_1 + 4x_2 - x_3 + 2x_4 + x_5 = -4 \\
 2x_1 + 2x_2 + 3x_4 = -11
 \end{cases}$$

SOLUZIONE. Usando il metodo di eliminazione di Gauss, trasformiamo il sistema dato in un sistema triangolare superiore equivalente:

$$\begin{array}{ccccc|c}
 2 & -1 & 1 & 6 & -1 & -2 \\
 1 & -2 & 1 & 3 & -1 & -1 \\
 1 & 4 & -1 & 2 & 1 & -4 \\
 2 & 2 & 0 & 3 & 0 & -11
 \end{array}
 \quad
 \begin{array}{ccccc|c}
 1 & -2 & 1 & 3 & -1 & -1 \\
 2 & -1 & 1 & 6 & -1 & -2 \\
 1 & 4 & -1 & 2 & 1 & -4 \\
 2 & 2 & 0 & 3 & 0 & -11
 \end{array}$$

$$\begin{array}{ccccc|c}
1 & -2 & 1 & 3 & -1 & -1 \\
0 & 3 & -1 & 0 & 1 & 0 \\
0 & 6 & -2 & -1 & 2 & -3 \\
0 & 6 & -2 & -3 & 2 & -9
\end{array}
\qquad
\begin{array}{ccccc|c}
1 & -2 & 1 & 3 & -1 & -1 \\
0 & 3 & -1 & 0 & 1 & 0 \\
0 & 0 & 0 & -1 & 0 & -3 \\
0 & 0 & 0 & -3 & 0 & -9
\end{array}$$

$$\begin{array}{ccccc|c}
1 & -2 & 1 & 3 & -1 & -1 \\
0 & 3 & -1 & 0 & 1 & 0 \\
0 & 0 & 0 & -1 & 0 & -3 \\
0 & 0 & 0 & 0 & 0 & 0
\end{array}
.$$

Risulta che

$$x_4 = 3,$$

e poi

$$x_2 = \text{arbitrario},$$

$$x_3 = \text{arbitrario},$$

$$x_5 = -3x_2 + x_3,$$

$$x_1 = -1 + 2x_2 - x_3 - 3x_4 + x_5 = -10 - x_2. \quad \square$$

3.7.1. Come trovare basi tramite eliminazione di Gauss.

Un sistema di vettori $\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_n$ è linearmente indipendente esattamente quando il sistema lineare

$$\lambda_1 \mathbf{x}_1 + \lambda_2 \mathbf{x}_2 + \dots + \lambda_n \mathbf{x}_n = \mathbf{0}$$

nelle variabili $\lambda_1, \dots, \lambda_n$ ha solo la soluzione nulla. Quindi il problema dell'indipendenza lineare può essere trattato mediante la risoluzione di un sistema lineare, e quindi tramite l'eliminazione di Gauss.

ESERCIZIO 3.7.5. Quali vettori tra

$$\mathbf{a} = \begin{pmatrix} -1 \\ 3 \\ 2 \end{pmatrix}, \quad \mathbf{b} = \begin{pmatrix} 1 \\ 2 \\ 1 \end{pmatrix}, \quad \mathbf{c} = \begin{pmatrix} 2 \\ 0 \\ -1 \end{pmatrix}$$

sono combinazioni lineari di

$$\begin{pmatrix} 1 \\ 3 \\ 1 \end{pmatrix}, \begin{pmatrix} 3 \\ 1 \\ -1 \end{pmatrix}, \begin{pmatrix} -1 \\ 1 \\ 1 \end{pmatrix} \quad ?$$

SOLUZIONE. La relazione

$$\lambda_1 \begin{pmatrix} 1 \\ 3 \\ 1 \end{pmatrix} + \lambda_2 \begin{pmatrix} 3 \\ 1 \\ -1 \end{pmatrix} + \lambda_3 \begin{pmatrix} -1 \\ 1 \\ 1 \end{pmatrix} = \mathbf{a}$$

è equivalente al sistema di equazioni lineari

$$\begin{cases} \lambda_1 + 3\lambda_2 - \lambda_3 = -1 \\ 3\lambda_1 + \lambda_2 + \lambda_3 = 3 \\ \lambda_1 - \lambda_2 + \lambda_3 = 2 \end{cases}$$

Usando il metodo di eliminazione di Gauss, si trova che questo sistema ha soluzioni, perciò $\mathbf{a}$ è combinazione lineare dei vettori dati.

In maniera analoga si trova che anche $\mathbf{c}$ è combinazione lineare di questi, ma $\mathbf{b}$ no. $\square$

ESERCIZIO 3.7.6. Si dica se i vettori

$$\begin{pmatrix} 1 \\ -1 \\ 3 \\ 2 \end{pmatrix}, \begin{pmatrix} -1 \\ 2 \\ 0 \\ 3 \end{pmatrix}, \begin{pmatrix} 1 \\ -7 \\ 9 \\ 0 \end{pmatrix}$$

sono linearmente indipendenti o no.

SOLUZIONE. Si trova che il sistema omogeneo di equazioni lineari

$$\begin{cases} \lambda_1 - \lambda_2 + \lambda_3 = 0 \\ -\lambda_1 + 2\lambda_2 - 7\lambda_3 = 0 \\ 3\lambda_1 + 9\lambda_3 = 0 \\ 2\lambda_1 + 3\lambda_2 = 0 \end{cases}$$

ha la sola soluzione $\lambda_1 = \lambda_2 = \lambda_3 = 0$, perciò i vettori dati sono linearmente indipendenti. $\square$

ESERCIZIO 3.7.7. Per quali numeri reali x e y sono i vettori

$$\begin{pmatrix} x \\ y \\ y \end{pmatrix}, \begin{pmatrix} x \\ -y \\ y \end{pmatrix} \in \mathbb{R}^3$$

linearmente indipendenti?

SOLUZIONE. Con calcoli analoghi a prima si trova la soluzione x arbitrario ed $y \neq 0$. $\square$

ESERCIZIO 3.7.8. Si trovi una base per il sottospazio lineare di $\mathbb{R}^4$ generato dai vettori

$$\begin{pmatrix} 1 \\ -1 \\ -1 \\ 1 \end{pmatrix}, \begin{pmatrix} 1 \\ 2 \\ -1 \\ -1 \end{pmatrix}, \begin{pmatrix} -1 \\ 1 \\ 1 \\ 1 \end{pmatrix}, \begin{pmatrix} 2 \\ 1 \\ -2 \\ 2 \end{pmatrix}.$$

SOLUZIONE. Per vedere se i quattro vettori dati sono linearmente indipendenti o no, troviamo le soluzioni del sistema omogeneo

$$\begin{cases} x_1 + x_2 - x_3 + 2x_4 = 0 \\ -x_1 + 2x_2 + x_3 + x_4 = 0 \\ -x_1 - x_2 + x_3 - 2x_4 = 0 \\ x_1 - x_2 + x_3 + 2x_4 = 0 \end{cases}$$

usando eliminazione di Gauss:

$$\begin{array}{cccc|c} 1 & 1 & -1 & 2 & 0 \\ -1 & 2 & 1 & 1 & 0 \\ -1 & -1 & 1 & -2 & 0 \\ 1 & -1 & 1 & 2 & 0 \end{array} \quad \begin{array}{cccc|c} 1 & 1 & -1 & 2 & 0 \\ 0 & 3 & 0 & 3 & 0 \\ 0 & 0 & 0 & 0 & 0 \\ 0 & -2 & 2 & 0 & 0 \end{array}$$

$$\begin{array}{cccc|c} 1 & 1 & -1 & 2 & 0 \\ 0 & 1 & 0 & 1 & 0 \\ 0 & -1 & 1 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 \end{array} \quad \begin{array}{cccc|c} 1 & 1 & -1 & 2 & 0 \\ 0 & 1 & 0 & 1 & 0 \\ 0 & 0 & 1 & 1 & 0 \\ 0 & 0 & 0 & 0 & 0 \end{array}.$$

Risulta che il sistema omogeneo di cui sopra ammette soluzione per $x_4 \neq 0$, quindi il quarto vettore è combinazione lineare dei primi tre.

Di conseguenza il sottospazio lineare X di $\mathbb{R}^4$ generato dai quattro vettori dati è già generato dai primi tre.

Per vedere se i primi tre vettori sono linearmente indipendenti o no, dobbiamo trovare le soluzioni del sistema omogeneo

$$\begin{cases} x_1 + x_2 - x_3 = 0 \\ -x_1 + 2x_2 + x_3 = 0 \\ -x_1 - x_2 + x_3 = 0 \\ x_1 - x_2 + x_3 = 0 \end{cases},$$

cioè le soluzioni del sistema precedente con $x_4 = 0$. Ma i calcoli di cui sopra ci mostrano che se $x_4 = 0$ allora anche $x_1 = x_2 = x_3 = 0$. Perciò i tre vettori

$$\begin{pmatrix} 1 \\ -1 \\ -1 \\ 1 \end{pmatrix}, \begin{pmatrix} 1 \\ 2 \\ -1 \\ -1 \end{pmatrix}, \begin{pmatrix} -1 \\ 1 \\ 1 \\ 1 \end{pmatrix}$$

sono linearmente indipendenti e quindi costituiscono una base di X .

Soluzione alternativa. Invece di scrivere i quattro vettori come colonne e risolvere il sistema omogeneo ottenuto, per vedere se c'è l'indipendenza lineare o no e, nel caso di dipendenza lineare, quale dei quattro vettori sia combinazione lineare degli altri tre, scriviamo i quattro vettori come righe di una matrice rettangolare:

$$\begin{array}{c} \text{----- } \mathbf{r}_1 \text{ -----} \\ \text{----- } \mathbf{r}_2 \text{ -----} \\ \text{----- } \mathbf{r}_3 \text{ -----} \\ \text{----- } \mathbf{r}_4 \text{ -----} \end{array}$$

Sommando il multiplo per un fattore $\lambda \in \mathbb{R}$ di una delle righe ad un'altra, per esempio, $\lambda \mathbf{r}_1$ a $\mathbf{r}_3$, le righe così ottenute, cioè

$$\begin{array}{c} \text{----- } \mathbf{r}_1 \text{ -----} \\ \text{----- } \mathbf{r}_2 \text{ -----} \\ \text{----- } \mathbf{r}_3 + \lambda \mathbf{r}_1 \text{ -----} \\ \text{----- } \mathbf{r}_4 \text{ -----} \end{array},$$

generano lo stesso sottospazio lineare di $\mathbb{R}^4$ delle righe $\mathbf{r}_1, \mathbf{r}_2, \mathbf{r}_3, \mathbf{r}_4$. Infatti,

- ogni combinazione lineare di $\mathbf{r}_1, \mathbf{r}_2, \mathbf{r}_3 + \lambda \mathbf{r}_1, \mathbf{r}_4$ è combinazione lineare di $\mathbf{r}_1, \mathbf{r}_2, \mathbf{r}_3, \mathbf{r}_4$, e
- ogni combinazione lineare di $\mathbf{r}_1, \mathbf{r}_2, \mathbf{r}_3 = (\mathbf{r}_3 + \lambda \mathbf{r}_1) - \lambda \mathbf{r}_1, \mathbf{r}_4$ è combinazione lineare di $\mathbf{r}_1, \mathbf{r}_2, \mathbf{r}_3 + \lambda \mathbf{r}_1, \mathbf{r}_4$.

Analogamente, uno scambio di due righe o la moltiplicazione di una delle righe con un scalare non nullo non cambia il sottospazio lineare X di $\mathbb{R}^4$ generato dalle righe. In altre parole, le operazioni di riga non cambiano il sottospazio vettoriale generato dalle righe. Perciò, svolgendo l'eliminazione di Gauss, otteniamo una matrice

$$\begin{array}{c} \text{----- } \mathbf{r}'_1 \text{ -----} \\ \text{----- } \mathbf{r}'_2 \text{ -----} \\ \text{----- } \mathbf{r}'_3 \text{ -----} \\ \text{----- } \mathbf{r}'_4 \text{ -----} \end{array}$$

tale che $\mathbf{r}'_1, \mathbf{r}'_2, \mathbf{r}'_3, \mathbf{r}'_4$ generano X . Ma le righe non nulle di questa matrice sono linearmente indipendenti, perciò costituiscono una base per X .

Questo metodo ha il vantaggio di ottenere una base tramite una sola applicazione dell'eliminazione di Gauss. Per di più, gli elementi della base trovata hanno molte componenti uguali a 0. Però, di solito, la base trovata non è sottoinsieme del sistema dato di generatori.

Nel caso presente:

$$\begin{array}{cccc|cccc} 1 & -1 & -1 & 1 & 1 & -1 & -1 & 1 \\ 1 & 2 & -1 & -1 & 0 & 3 & 0 & -2 \\ -1 & 1 & 1 & 1 & 0 & 0 & 0 & 2 \\ 2 & 1 & -2 & 2 & 0 & 3 & 0 & 0 \\ \hline 1 & -1 & -1 & 1 & 1 & -1 & -1 & 1 \\ 0 & 1 & 0 & 0 & 0 & 1 & 0 & 0 \\ 0 & 0 & 0 & 1 & 0 & 0 & 0 & 1 \\ 0 & 3 & 0 & -2 & 0 & 0 & 0 & 0 \end{array}$$

Cosicché una base per X è

$$\begin{pmatrix} 1 \\ -1 \\ -1 \\ 1 \end{pmatrix}, \begin{pmatrix} 0 \\ 1 \\ 0 \\ 0 \end{pmatrix}, \begin{pmatrix} 0 \\ 0 \\ 0 \\ 1 \end{pmatrix}.$$

□

3.8. Rango di una matrice

Osserviamo che, a causa di (3.6.1), l'immagine di T_A è il sottospazio lineare di $\mathbb{R}^m$ generato dalle colonne di A . Perciò il *rango* di A , definito come la dimensione dell'immagine dell'applicazione lineare $T_A : \mathbb{R}^n \ni \mathbf{x} \mapsto A\mathbf{x} \in \mathbb{R}^m$, è uguale alla dimensione del sottospazio lineare di $\mathbb{R}^m$ generato dalle colonne di A . Poiché ogni sistema di generatori contiene una base, possiamo dire che il rango di A è il numero massimo di colonne linearmente indipendenti. Indichiamo il rango di A con $\text{rg}(A)$ o $\text{rg}A$. Poiché A ha n colonne, $\text{rg}(A) \leq n$. D'altro canto, poiché le colonne di A appartengono ad $\mathbb{R}^m$ e la dimensione di ogni sottospazio lineare di $\mathbb{R}^m$ è minore o uguale alla dimensione m di $\mathbb{R}^m$, abbiamo anche $\text{rg}(A) \leq m$. Quindi:

$$\text{rg}(A) \leq \max\{m, n\}, \quad A \in M_{mn}.$$

NOTA 3.8.1. (**Invertibilità di una matrice.**) L'iniettività di T_A è equivalente a $\text{rg}(A) = n$, mentre la surgettività di T_A è equivalente a $\text{rg}(A) = m$. Perciò T_A è bigettiva se e solo se A è quadrata e $\text{rg}(A) = m = n$. In questo caso A si chiama *invertibile*. □

DEFINIZIONE 3.8.2. (**Matrice trasposta.**) La *trasposta* di una matrice $A \in M_{mn}$ è la matrice $A^\top \in M_{nm}$ che si ottiene da A cambiando le colonne in righe (ed analogamente le righe in colonne):

$$\left(\begin{array}{cccc} | & | & & | \\ \mathbf{c}_1 & \mathbf{c}_2 & \dots & \mathbf{c}_n \\ | & | & & | \end{array} \right)^\top = \left(\begin{array}{ccc} - & \mathbf{c}_1 & - \\ - & \mathbf{c}_2 & - \\ & \vdots & \\ - & \mathbf{c}_n & - \end{array} \right),$$

cioè

$$\begin{pmatrix} \alpha_{11} & \alpha_{12} & \dots & \alpha_{1n} \\ \alpha_{21} & \alpha_{22} & \dots & \alpha_{2n} \\ \vdots & \vdots & \ddots & \vdots \\ \alpha_{m1} & \alpha_{m2} & \dots & \alpha_{mn} \end{pmatrix}^{\top} = \begin{pmatrix} \alpha_{11} & \alpha_{21} & \dots & \alpha_{m1} \\ \alpha_{12} & \alpha_{22} & \dots & \alpha_{m2} \\ \vdots & \vdots & \ddots & \vdots \\ \alpha_{1n} & \alpha_{2n} & \dots & \alpha_{mn} \end{pmatrix}.$$

NOTA 3.8.3. (**Regole di calcolo per la matrice trasposta.**)

$$\begin{aligned} (A + B)^{\top} &= A^{\top} + B^{\top} \text{ per } A, B \in M_{mn}, \\ (\lambda A)^{\top} &= \lambda A^{\top}, \\ (BA)^{\top} &= A^{\top} B^{\top} \text{ per } A \in M_{mn}, B \in M_{pm}, \\ (A^{\top})^{\top} &= A. \end{aligned}$$

□

PROPOSIZIONE 3.8.4.

$$\operatorname{rg}(A) = \operatorname{rg}(A^{\top}), \quad A \in M_{mn}. \quad (3.8.1)$$

* DIMOSTRAZIONE. Osserviamo che la trasposta $\mathbf{x}^{\top}$ di ogni $\mathbf{x} \in \mathbb{R}^n$, considerata come matrice $n \times 1$, è una matrice $1 \times n$, cioè un vettore riga. Usando le regole di calcolo di cui sopra, otteniamo una uguaglianza importante:

$$\mathbf{y}^{\top} A \mathbf{x} = (A^{\top} \mathbf{y})^{\top} \mathbf{x}, \quad A \in M_{mn}, \mathbf{x} \in \mathbb{R}^n, \mathbf{y} \in \mathbb{R}^m \quad (3.8.2)$$

(tutti i prodotti di matrici in (3.8.2) hanno senso perché il numero di colonne di ciascuna corrisponde al numero di righe della successiva).

Ne segue che

$$\{\mathbf{x} \in \mathbb{R}^n; A \mathbf{x} = \mathbf{0}_m\} \cap A^{\top} \mathbb{R}^m = \{\mathbf{0}_n\}, \quad A \in M_{mn}. \quad (3.8.3)$$

Infatti, se

$$\mathbf{x} = \begin{pmatrix} x_1 \\ \vdots \\ x_n \end{pmatrix} \in \mathbb{R}^n$$

è tale che nello stesso tempo $A \mathbf{x} = \mathbf{0}_m$ e $\mathbf{x} = A^{\top} \mathbf{y}$ per almeno un $\mathbf{y} \in \mathbb{R}^m$, allora (3.8.2) implica che

$$\mathbf{x}^{\top} \mathbf{x} = (A^{\top} \mathbf{y})^{\top} \mathbf{x} = \mathbf{y}^{\top} A \mathbf{x} = \mathbf{y}^{\top} \mathbf{0}_m = (0).$$

Ma l'unico elemento di $\mathbf{x}^\top \mathbf{x}$ è uguale a $x_1^2 + \dots + x_n^2$, perciò dobbiamo avere $x_1 = \dots = x_n = 0$, cioè $\mathbf{x} = \mathbf{0}_n$.

(Vedremo più avanti che l'applicazione

$$\mathbb{R}^n \times \mathbb{R}^n \ni (\mathbf{x}, \mathbf{y}) \longmapsto \mathbf{y}^\top \mathbf{x} \in \mathbb{R},$$

chiamata il *prodotto scalare (naturale)* su $\mathbb{R}^n$, è di importanza fondamentale anche per altri propositi.)

Ora siamo pronti per dimostrare che

$$\text{rg}(A^\top) \leq \text{rg}(A) \quad \text{per ogni } A \in M_{mn}. \quad (3.8.4)$$

Indichiamo con r il rango di $\text{rg}(A^\top)$. Allora esistono r colonne di $\text{rg}(A^\top)$, diciamo $\mathbf{x}_1, \dots, \mathbf{x}_r \in \mathbb{R}^n$, linearmente indipendenti. Verifichiamo che i vettori $A\mathbf{x}_1, \dots, A\mathbf{x}_r$ sono ancora linearmente indipendenti:

Se $\sum_{j=1}^r \lambda_j A\mathbf{x}_j = \mathbf{0}_m$, cioè $A\left(\sum_{j=1}^r \lambda_j \mathbf{x}_j\right) = \mathbf{0}_m$, allora $\sum_{j=1}^r \lambda_j \mathbf{x}_j$ appartiene al nucleo di $\mathbb{R}^n \ni \mathbf{x} \rightarrow A\mathbf{x} \in \mathbb{R}^m$. Ma questa somma si trova anche in $A^\top \mathbb{R}^m$, perciò (3.8.3) implica che $\sum_{j=1}^r \lambda_j \mathbf{x}_j = \mathbf{0}_n$. Siccome i vettori $\mathbf{x}_1, \dots, \mathbf{x}_r$ sono linearmente indipendenti, concludiamo che tutti i coefficienti $\lambda_1, \dots, \lambda_r$ si annullano.

Ne segue che la dimensione dell'immagine di $\mathbb{R}^n \ni \mathbf{x} \rightarrow A\mathbf{x} \in \mathbb{R}^m$, cioè il rango di A , è maggiore o uguale di r . Questo prova (3.8.4).

Infine, applicando la disuguaglianza dimostrata per $A^\top$, otteniamo anche la disuguaglianza inversa: $\text{rg}(A) = \text{rg}((A^\top)^\top) \leq \text{rg}(A^\top)$.

□

NOTA 3.8.5. Per (3.8.1) il rango di una matrice è uguale anche alla dimensione dello spazio lineare generato dalle sue righe, ovvero al numero massimo di righe linearmente indipendenti. □

COROLLARIO 3.8.6. (**Calcolo del rango.**) *Per calcolare il rango di una matrice A , basta applicare l'eliminazione di Gauss ad A e contare il numero delle righe non zero nella tabella ottenuto.*

DIMOSTRAZIONE. I passaggi dell'eliminazione di Gauss non cambiano lo spazio lineare generato dalle righe e le righe non zero nella tabella finale costituiscono una base di questo spazio lineare. Possiamo calcolare il rango di A anche applicando l'eliminazione di Gauss alle righe della matrice trasposta, cioè alle colonne di A .

ESERCIZIO 3.8.7. Si determini il rango della matrice

$$\begin{pmatrix} 3 & 1 & 1 & 4 \\ 0 & 4 & 10 & 1 \\ 1 & 7 & 17 & 3 \\ 2 & 2 & 4 & 3 \end{pmatrix}.$$

SOLUZIONE. Per semplificare i calcoli fin dall'inizio, sottraiamo prima l'ultima riga dalla prima. Poi svolgiamo l'eliminazione di Gauss usuale, ottenendo

$$\begin{aligned} \operatorname{rg} \begin{pmatrix} 3 & 1 & 1 & 4 \\ 0 & 4 & 10 & 1 \\ 1 & 7 & 17 & 3 \\ 2 & 2 & 4 & 3 \end{pmatrix} &= \operatorname{rg} \begin{pmatrix} 1 & -1 & -3 & 1 \\ 0 & 4 & 10 & 1 \\ 1 & 7 & 17 & 3 \\ 2 & 2 & 4 & 3 \end{pmatrix} \\ &= \operatorname{rg} \begin{pmatrix} 1 & -1 & -3 & 1 \\ 0 & 4 & 10 & 1 \\ 0 & 8 & 20 & 2 \\ 0 & 4 & 10 & 1 \end{pmatrix} = \operatorname{rg} \begin{pmatrix} 1 & -1 & -3 & 1 \\ 0 & 4 & 10 & 1 \\ 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 \end{pmatrix} \\ &= 2. \end{aligned} \quad \square$$

3.9. Applicazioni lineari e matrici invertibili

Abbiamo introdotto la nozione di matrice invertibile nella Nota [3.8.1](#). Perché una matrice A sia invertibile, per prima cosa A dev'essere quadrata. In tal caso, $A \in M_n$ è invertibile se e solo se $\operatorname{rg}(A)$ è uguale ad n , cioè assume il valore massimo possibile.

Sia $A \in M_n$. Se A è invertibile, allora l'applicazione lineare $T_A : \mathbb{R}^n \longrightarrow \mathbb{R}^n$ ha una inversa $T_A^{-1} : \mathbb{R}^n \longrightarrow \mathbb{R}^n$, che è anch'essa lineare (verificare per esercizio). La matrice associata a T_A^{-1} si chiama *l'inversa* di A ed è

indicata con A^{-1} . Poiché $T_A^{-1} \circ T_A$ e $T_A \circ T_A^{-1}$ sono uguali all'applicazione identica su $\mathbb{R}^n$, per le matrici associate abbiamo

$$A A^{-1} = A^{-1} A = \mathbb{I}_n.$$

PROPOSIZIONE 3.9.1. *Se $A, B \in M_n$*

$$A B = \mathbb{I}_n \iff B A = \mathbb{I}_n \implies A \text{ è invertibile ed } A^{-1} = B. \quad (\text{I})$$

* DIMOSTRAZIONE. Supponiamo dapprima che $A B = \mathbb{I}_n$. Allora $T_A \circ T_B = T_{AB} = T_{\mathbb{I}_n} = \mathbb{I}_{\mathbb{R}^n}$. Per (3.2.6) T_A è surgettiva e poi (3.2.4) implica che T_A è di fatto bigettiva. Perciò A è invertibile e

$$B = \mathbb{I}_n B = (A^{-1} A) B = A^{-1} (A B) = A^{-1} \mathbb{I}_n = A^{-1}.$$

Analogamente, se $B A = \mathbb{I}_n$, allora $T_B \circ T_A = T_{BA} = T_{\mathbb{I}_n} = \mathbb{I}_{\mathbb{R}^n}$ implica l'iniettività, e poi la bigettività di T_A . Di conseguenza, A è anche in questo caso invertibile e

$$B = B \mathbb{I}_n = B (A A^{-1}) = (B A) A^{-1} = \mathbb{I}_n A^{-1} = A^{-1}.$$

□

NOTA 3.9.2. La composizione di due operatori lineari invertibili è ovviamente ancora invertibile. Ne segue che il prodotto righe per colonne di due matrici $A, B \in M_n$ è una matrice invertibile (ed infatti si verifica immediatamente che l'inversa è $B^{-1} A^{-1}$). Da questo segue che le matrici quadrate $n \times n$ invertibili formano un *gruppo*, con la terminologia dell'Appendice 1.3 □

NOTAZIONE 3.9.3. (**Il gruppo lineare.**) *Il gruppo delle matrici invertibili reali a dimensione n si denota con $GL_n(\mathbb{R})$; il gruppo delle matrici invertibili a coefficienti complessi si denota con $GL_n(\mathbb{C})$.*

3.10. Il calcolo della matrice inversa

Usiamo l'uguaglianza (I) per calcolare la matrice inversa.

Sia infatti

$$A = \begin{pmatrix} \alpha_{11} & \alpha_{12} & \dots & \alpha_{1n} \\ \alpha_{21} & \alpha_{22} & \dots & \alpha_{2n} \\ \vdots & \vdots & \ddots & \vdots \\ \alpha_{n1} & \alpha_{n2} & \dots & \alpha_{nn} \end{pmatrix}$$

una matrice quadrata di ordine n e poniamoci il problema di verificare l'invertibilità di A e, nel caso di invertibilità, di calcolare la matrice inversa A^{-1} .

Grazie ad (I), il problema posto è equivalente a verificare l'esistenza di una matrice

$$B = \begin{pmatrix} \beta_{11} & \beta_{12} & \dots & \beta_{1n} \\ \beta_{21} & \beta_{22} & \dots & \beta_{2n} \\ \vdots & \vdots & \ddots & \vdots \\ \beta_{n1} & \beta_{n2} & \dots & \beta_{nn} \end{pmatrix} = \begin{pmatrix} | & | & & | \\ \mathbf{b}_1 & \mathbf{b}_2 & \dots & \mathbf{b}_n \\ | & | & & | \end{pmatrix}$$

con $AB = \mathbb{I}_n$ e, nel caso di esistenza, a calcolarla. Ma $AB = \mathbb{I}_n$ significa

$$\begin{aligned} & \begin{pmatrix} \alpha_{11} & \alpha_{12} & \dots & \alpha_{1n} \\ \alpha_{21} & \alpha_{22} & \dots & \alpha_{2n} \\ \vdots & \vdots & \ddots & \vdots \\ \alpha_{n1} & \alpha_{n2} & \dots & \alpha_{nn} \end{pmatrix} \begin{pmatrix} | & | & & | \\ \mathbf{b}_1 & \mathbf{b}_2 & \dots & \mathbf{b}_n \\ | & | & & | \end{pmatrix} = \\ & = \begin{pmatrix} | & | & & | \\ \mathbf{e}_1 & \mathbf{e}_2 & \dots & \mathbf{e}_n \\ | & | & & | \end{pmatrix}, \end{aligned}$$

ove $\mathbf{e}_1, \mathbf{e}_2, \dots, \mathbf{e}_n$ è la base naturale di $\mathbb{R}^n$, cioè i n sistemi di equazioni

lineari

$$\begin{pmatrix} & & & \\ & A & & \\ & & & \end{pmatrix} \begin{pmatrix} | \\ \mathbf{b}_1 \\ | \end{pmatrix} = \begin{pmatrix} | \\ \mathbf{e}_1 \\ | \end{pmatrix},$$

.....

(3.10.1)

$$\begin{pmatrix} & & & \\ & A & & \\ & & & \end{pmatrix} \begin{pmatrix} | \\ \mathbf{b}_n \\ | \end{pmatrix} = \begin{pmatrix} | \\ \mathbf{e}_n \\ | \end{pmatrix}.$$

Per la soluzione di questi n sistemi usiamo il metodo di eliminazione di Gauss. Poiché gli coefficienti sono i stessi, possiamo svolgere l'eliminazione di Gauss *simultaneamente per tutti i sistemi*.

Partiamo quindi con la tabella

$$\begin{array}{cccc|cccc} \alpha_{11} & \alpha_{12} & \dots & \alpha_{1n} & 1 & 0 & \dots & 0 \\ \alpha_{21} & \alpha_{22} & \dots & \alpha_{2n} & 0 & 1 & \dots & 0 \\ \vdots & \vdots & \ddots & \vdots & \vdots & \vdots & \ddots & \vdots \\ \alpha_{n1} & \alpha_{n2} & \dots & \alpha_{nn} & 0 & 0 & \dots & 1 \end{array}.$$

Tramite il processo di eliminazione di Gauss otteniamo una tabella con la matrice dei coefficienti triangolare superiore, cioè n sistemi triangolari superiori:

$$\begin{array}{cccc|cccc} \alpha'_{11} & \alpha'_{12} & \dots & \alpha'_{1n} & \gamma'_{11} & \gamma'_{12} & \dots & \gamma'_{1n} \\ 0 & \alpha'_{22} & \dots & \alpha'_{2n} & \gamma'_{21} & \gamma'_{22} & \dots & \gamma'_{2n} \\ \vdots & \vdots & \ddots & \vdots & \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \dots & \alpha'_{nn} & \gamma'_{n1} & \gamma'_{n2} & \dots & \gamma'_{nn} \end{array}. \quad (3.10.2)$$

Se uno degli elementi diagonali $\alpha'_{11}, \alpha'_{22}, \dots, \alpha'_{nn}$ si annulla, allora il rango di A è $< n$ e quindi A non è invertibile. Infatti, se per esempio α'_{33} è il primo elemento zero sulla diagonale, allora la terza colonna di coefficienti

$$\begin{pmatrix} \gamma'_{13} \\ \gamma'_{23} \\ \gamma'_{33} \\ \vdots \\ 0 \end{pmatrix} = \begin{pmatrix} \gamma'_{13} \\ \gamma'_{23} \\ 0 \\ \vdots \\ 0 \end{pmatrix}$$

appartiene al sottospazio lineare bidimensionale

$$\left\{ \begin{pmatrix} x_1 \\ x_2 \\ 0 \\ \vdots \\ 0 \end{pmatrix} ; x_1, x_2 \in \mathbb{R} \right\}$$

di $\mathbb{R}^n$, generato dai vettori linearmente indipendenti

$$\begin{pmatrix} \gamma'_{11} \\ 0 \\ 0 \\ \vdots \\ 0 \end{pmatrix}, \begin{pmatrix} \gamma'_{12} \\ \gamma'_{22} \\ 0 \\ \vdots \\ 0 \end{pmatrix}, \quad \gamma'_{11}, \gamma'_{22} \neq 0.$$

Se invece ogni elemento sulla diagonale è non zero, allora A è invertibile. Per trovare la matrice inversa, cioè risolvere gli n sistemi (3.10.1) che sono equivalenti agli n sistemi (3.10.2), applichiamo alla tabella (3.10.2) il processo di eliminazione di Gauss, partendo dal coefficiente *pivot* γ'_{nn} nell'angolo destro inferiore e procedendo verso alto ed a sinistra:

- Sommando alla $(n-1)$ -esima riga il multiplo con $-\frac{\alpha'_{n-1,n}}{\alpha'_{nn}}$ dell'ultima riga, azzeriamo l'elemento sopra α'_{nn} . Poi, sommando alla $(n-2)$ -esima riga il multiplo con $-\frac{\alpha'_{n-2,n}}{\alpha'_{nn}}$ dell'ultima riga, azzeriamo anche il secondo elemento sopra α'_{nn} . Dopo $n-1$ passi la matrice dei coefficienti diventa

$$\begin{array}{cccccc}
\alpha'_{11} & \alpha''_{12} & \cdots & \alpha''_{1,n-1} & 0 & \\
0 & \alpha'_{22} & \cdots & \alpha''_{2,n-1} & 0 & \\
\vdots & \vdots & \ddots & \vdots & \vdots & \cdot \\
0 & 0 & \cdots & \alpha'_{n-1,n-1} & 0 & \\
0 & 0 & \cdots & 0 & \alpha'_{nn} &
\end{array}$$

- Ora, sommando successivamente multipli scalari appropriati della $(n-1)$ -esima riga a tutte le righe precedenti, azzeriamo tutti gli elementi della $(n-1)$ -esima colonna che si trovano sopra α'_{nn} . Procediamo in questo modo fino all'azzeramento del coefficiente sopra α'_{22} .

Il risultato è una tabella che rappresenta n sistemi diagonali equivalenti a (3.10.1) :

$$\begin{array}{ccccc|cccc}
\alpha'_{11} & 0 & \cdots & 0 & 0 & \gamma''_{11} & \gamma''_{12} & \cdots & \gamma''_{1n} \\
0 & \alpha'_{22} & \cdots & 0 & 0 & \gamma''_{21} & \gamma''_{22} & \cdots & \gamma''_{2n} \\
\vdots & \vdots & \ddots & \vdots & \vdots & \vdots & \vdots & \ddots & \vdots \\
0 & 0 & \cdots & \alpha'_{n-1,n-1} & 0 & \gamma''_{n-1,1} & \gamma''_{n-1,2} & \cdots & \gamma''_{n-1,n-1} \\
0 & 0 & \cdots & 0 & \alpha'_{nn} & \gamma''_{n1} & \gamma''_{n2} & \cdots & \gamma''_{nn}
\end{array}$$

Poiché tutti gli elementi sulla diagonale sono non zero, possiamo dividere ogni k -esima riga con l'elemento diagonale α'_{kk} su di essa, ottenendo

$$\begin{array}{ccccc|cccc}
1 & 0 & \cdots & 0 & 0 & \gamma''_{11} & \gamma''_{12} & \cdots & \gamma''_{1n} \\
0 & 1 & \cdots & 0 & 0 & \gamma''_{21} & \gamma''_{22} & \cdots & \gamma''_{2n} \\
\vdots & \vdots & \ddots & \vdots & \vdots & \vdots & \vdots & \ddots & \vdots \\
0 & 0 & \cdots & 1 & 0 & \gamma''_{n-1,1} & \gamma''_{n-1,2} & \cdots & \gamma''_{n-1,n-1} \\
0 & 0 & \cdots & 0 & 1 & \gamma''_{n1} & \gamma''_{n2} & \cdots & \gamma''_{nn}
\end{array} \quad (3.10.3)$$

Ora da (3.10.3) leggiamo che la soluzione del primo sistema in (3.10.1), cioè la prima colonna della matrice B inversa di A , è

$$\begin{pmatrix} | \\ \mathbf{b}_1 \\ | \end{pmatrix} = \begin{pmatrix} \gamma'_{11} \\ \vdots \\ \gamma'_{n1} \end{pmatrix}.$$

Poi, la soluzione del secondo sistema in (3.10.1), cioè la seconda colonna della matrice B inversa di A , è

$$\begin{pmatrix} | \\ \mathbf{b}_2 \\ | \end{pmatrix} = \begin{pmatrix} \gamma'_{12} \\ \vdots \\ \gamma'_{n2} \end{pmatrix}$$

e così via fino all'uguaglianza

$$\begin{pmatrix} | \\ \mathbf{b}_n \\ | \end{pmatrix} = \begin{pmatrix} \gamma'_{1n} \\ \vdots \\ \gamma'_{nn} \end{pmatrix}.$$

Perciò la matrice inversa di A è la matrice dei termini noti in (3.10.3):

$$\begin{pmatrix} \alpha_{11} & \alpha_{12} & \cdots & \alpha_{1n} \\ \alpha_{21} & \alpha_{22} & \cdots & \alpha_{2n} \\ \vdots & \vdots & \ddots & \vdots \\ \alpha_{n1} & \alpha_{n2} & \cdots & \alpha_{nn} \end{pmatrix}^{-1} = \begin{pmatrix} \gamma''_{11} & \gamma''_{12} & \cdots & \gamma''_{1n} \\ \gamma''_{21} & \gamma''_{22} & \cdots & \gamma''_{2n} \\ \vdots & \vdots & \ddots & \vdots \\ \gamma''_{n1} & \gamma''_{n2} & \cdots & \gamma''_{nn} \end{pmatrix}. \quad \square$$

ESEMPIO 3.10.1. Vediamo se la matrice

$$A = \begin{pmatrix} 0 & 1 & 2 \\ -1 & 3 & 0 \\ 1 & -2 & 1 \end{pmatrix}$$

è invertibile o no e nel caso di invertibilità calcoliamone l'inversa.

SVOLGIMENTO. Tramite l'eliminazione di Gauss riduciamo A a forma triangolare superiore:

$$\begin{array}{ccc|ccc} 0 & 1 & 2 & 1 & 0 & 0 & 1 & -2 & 1 & 0 & 0 & 1 \\ -1 & 3 & 0 & 0 & 1 & 0 & 0 & 1 & 2 & 1 & 0 & 0 \\ 1 & -2 & 1 & 0 & 0 & 1 & -1 & 3 & 0 & 0 & 1 & 0 \\ \\ 1 & -2 & 1 & 0 & 0 & 1 & 1 & -2 & 1 & 0 & 0 & 1 \\ 0 & 1 & 2 & 1 & 0 & 0 & 0 & 1 & 2 & 1 & 0 & 0 \\ 0 & 1 & 1 & 0 & 1 & 1 & 0 & 0 & -1 & -1 & 1 & 1 \end{array}.$$

Poiché tutti gli elementi sulla diagonale sono non zero, A è invertibile. Per calcolare la matrice inversa, tramite operazioni elementari sulle righe azzeriamo anche i coefficienti sopra la diagonale:

$$\begin{array}{ccc|ccc} 1 & -2 & 0 & -1 & 1 & 2 & 1 & 0 & 0 & -3 & 5 & 6 \\ 0 & 1 & 0 & -1 & 2 & 2 & 0 & 1 & 0 & -1 & 2 & 2 \\ 0 & 0 & -1 & -1 & 1 & 1 & 0 & 0 & -1 & -1 & 1 & 1 \end{array} ,$$

Ora normalizziamo ogni riga dividendone tutti gli elementi per il suo elemento diagonale: in tal modo otteniamo come matrice di coefficienti la matrice unità:

$$\begin{array}{ccc|ccc} 1 & 0 & 0 & -3 & 5 & 6 \\ 0 & 1 & 0 & -1 & 2 & 2 \\ 0 & 0 & 1 & 1 & -1 & -1 \end{array} .$$

Concludiamo che

$$A^{-1} = \begin{pmatrix} -3 & 5 & 6 \\ -1 & 2 & 2 \\ 1 & -1 & -1 \end{pmatrix} . \quad \square$$

3.11. Minori ed orli di una matrice ed invertibilità

DEFINIZIONE 3.11.1. (*Minori, ovvero sottomatrici.*) Sia

$$A = \begin{pmatrix} \alpha_{11} & \alpha_{12} & \dots & \alpha_{1n} \\ \alpha_{21} & \alpha_{22} & \dots & \alpha_{2n} \\ \vdots & \vdots & \ddots & \vdots \\ \alpha_{m1} & \alpha_{m2} & \dots & \alpha_{mn} \end{pmatrix}$$

una matrice $m \times n$. Una *sottomatrice* (detta anche un *minore*) di A è una matrice che si ottiene cancellando alcune righe e colonne di A :

$$B = \begin{pmatrix} \alpha_{j_1 k_1} & \alpha_{j_1 k_2} & \dots & \alpha_{j_1 k_q} \\ \alpha_{j_2 k_1} & \alpha_{j_2 k_2} & \dots & \alpha_{j_2 k_q} \\ \vdots & \vdots & \ddots & \vdots \\ \alpha_{j_p k_1} & \alpha_{j_p k_2} & \dots & \alpha_{j_p k_q} \end{pmatrix}$$

è la sottomatrice ottenuta tramite la cancellazione di tutte le righe di A tranne quelle con indici $j_1 < j_2 < \dots < j_p$, e di tutte le colonne tranne quelle con indici $k_1 < k_2 < \dots < k_q$, ove $1 \leq p \leq n$ e $1 \leq q \leq m$.

NOTA 3.11.2. Se B è una sottomatrice di A , allora $\text{rg}(B) \leq \text{rg}(A)$. $\square$

DEFINIZIONE 3.11.3. (*Orli di una matrice.*) Sia B una sottomatrice di A . Una sottomatrice B' di A si ottiene *orlando* B se B' si ottiene cancellando una riga ed una colonna di meno di quante si cancellano per ottenere B .

Per esempio, togliendo le righe di indici 1 e 2, nonché le colonne di indici 2, 3 e 5 della matrice

$$A = \begin{pmatrix} 7 & -1 & 0 & 2 & 3 \\ -1 & 2 & 3 & 1 & -1 \\ -2 & 5 & -4 & 7 & 0 \\ 0 & -1 & 1 & 0 & 1 \end{pmatrix},$$

si ottiene la sottomatrice

$$B = \begin{pmatrix} -2 & 7 \\ 0 & 0 \end{pmatrix}.$$

Una sottomatrice di A che si ottiene orlando B è, per esempio,

$$\begin{pmatrix} -1 & 1 & -1 \\ -2 & 7 & 0 \\ 0 & 0 & 1 \end{pmatrix},$$

che si ottiene cancellando solo la prima riga e le colonne di indici 2 e 3 di A (abbiamo orlato B con la seconda riga e la quinta colonna di A). Un altro esempio si ottiene orlando B con la prima riga e la terza colonna di A :

$$\begin{pmatrix} 7 & 0 & 2 \\ -2 & -4 & 7 \\ 0 & 1 & 0 \end{pmatrix}.$$

Il rango di una matrice qualsiasi può essere espresso per mezzo di sottomatrici quadrate:

TEOREMA 3.11.4. (**Teorema degli orli.**) *Sia A una matrice $m \times n$. Allora $\text{rg}(A) = r$ se e solo se esiste una sottomatrice quadrata B di ordine r di A , tale che*

- B è invertibile e
- ogni sottomatrice quadrata di ordine $r + 1$ di A , che si ottiene orlando B , è non invertibile.

* DIMOSTRAZIONE. Sia $\text{rg}(A) = r$. Allora esistono r colonne linearmente indipendenti, diciamo quelli con indici

$$k_1 < k_2 < \dots < k_r,$$

di A . Se cancelliamo tutte le altre colonne, otteniamo una sottomatrice A' di A con m righe ed r colonne, tale che $\text{rg}(A') = r$. Ora, per (3.8.1), esistono r righe linearmente indipendenti, diciamo quelle con indici

$$j_1 < j_2 < \dots < j_r,$$

di A' . Se cancelliamo tutte le altre righe di A' otteniamo una sottomatrice quadrata B di A' di ordine r con tutte le righe linearmente indipendenti, quindi B è invertibile. Ovviamente, B è una sottomatrice anche di A e si verifica facilmente che ogni sottomatrice quadrata di ordine $r + 1$ di A che si ottiene orlando B è non invertibile.

Supponiamo adesso che esista una sottomatrice quadrata B di ordine r di A tale che le due condizioni nell'enunciato siano soddisfatte. Siano

$$j_1 < j_2 < \dots < j_r,$$

gli indici delle righe di A e

$$k_1 < k_2 < \dots < k_r,$$

gli indici delle colonne di A che non si cancellano per ottenere B . Allora le colonne di A con indici $k_1, k_2, \dots, k_r$ sono linearmente indipendenti. Verifichiamo che tutte le altre colonne di A sono combinazioni lineari di queste, dimostrando quindi che il rango di A è uguale ad r .

Supponiamo che una colonna di A , il cui indice indichiamo con k' , non sia combinazione lineare delle colonne con indici $k_1, k_2, \dots, k_r$. Allora le colonne con indici $k', k_1, k_2, \dots, k_r$ sono linearmente indipendenti. Pertanto il rango della sottomatrice A'' di A che si ottiene

cancellando tutte le colonne con indici diversi da $k', k_1, k_2, \dots, k_r$ è uguale a $r + 1$.

Per (3.8.1) sappiamo che la dimensione del sottospazio lineare di $\mathbb{R}^m$ generato dalle righe di A'' è uguale ad $r + 1$. Ma le righe con indici $j_1, j_2, \dots, j_r$, cioè le righe di B , sono linearmente indipendenti. Se tutte le altre righe di A'' fossero combinazioni lineari di queste, allora il rango di A'' sarebbe r , perciò esiste almeno una riga di A'' , diciamo quella con indice j' , che non è combinazione lineare delle righe con indici $j_1, j_2, \dots, j_r$. Allora le righe di A'' con indici $j'', j_1, j_2, \dots, j_r$ sono linearmente indipendenti e quindi la sottomatrice quadrata B'' di ordine $r + 1$ di A'' che si ottiene cancellando tutte le righe con indice diverso da $j'', j_1, j_2, \dots, j_r$ è invertibile. Ma B'' è una sottomatrice di A ottenuta orlando B e quindi otteniamo una contraddizione all'ipotesi fatta su B . $\square$

Il rango interviene nel seguente criterio per la risolubilità di un sistema di equazioni lineari:

TEOREMA 3.11.5. (**Rouché-Capelli.**) *Siano*

$$A = \begin{pmatrix} \alpha_{11} & \alpha_{12} & \dots & \alpha_{1n} \\ \alpha_{21} & \alpha_{22} & \dots & \alpha_{2n} \\ \vdots & \vdots & \ddots & \vdots \\ \alpha_{m1} & \alpha_{m2} & \dots & \alpha_{mn} \end{pmatrix}$$

una matrice $m \times n$ e

$$\mathbf{b} = \begin{pmatrix} b_1 \\ b_2 \\ \vdots \\ b_m \end{pmatrix} \in \mathbb{R}^m.$$

Allora il sistema $A\mathbf{x} = \mathbf{b}$ di m equazioni lineari in n incognite ammette

(almeno) una soluzione se e solo se

$$\begin{aligned} \operatorname{rg} \begin{pmatrix} \alpha_{11} & \alpha_{12} & \dots & \alpha_{1n} \\ \alpha_{21} & \alpha_{22} & \dots & \alpha_{2n} \\ \vdots & \vdots & \ddots & \vdots \\ \alpha_{m1} & \alpha_{m2} & \dots & \alpha_{mn} \end{pmatrix} &= \\ = \operatorname{rg} \begin{pmatrix} \alpha_{11} & \alpha_{12} & \dots & \alpha_{1n} & b_1 \\ \alpha_{21} & \alpha_{22} & \dots & \alpha_{2n} & b_2 \\ \vdots & \vdots & \ddots & \vdots & \vdots \\ \alpha_{m1} & \alpha_{m2} & \dots & \alpha_{mn} & b_m \end{pmatrix} . \end{aligned}$$

$\underbrace{\hspace{10em}}_{=A} \quad \underbrace{\hspace{1em}}_{=\mathbf{b}}$

* DIMOSTRAZIONE. Siano $\mathbf{c}_1, \dots, \mathbf{c}_n$ le colonne di A e supponiamo che il sistema $A\mathbf{x} = \mathbf{b}$ ammetta una soluzione

$$\mathbf{x} = \begin{pmatrix} x_1 \\ x_2 \\ \vdots \\ x_n \end{pmatrix} \in \mathbb{R}^n.$$

Allora dalla formula (3.6.1) segue

$$\mathbf{b} = A\mathbf{x} = \sum_{j=1}^n x_j \mathbf{c}_j.$$

Perciò i vettori $\mathbf{c}_1, \dots, \mathbf{c}_n, \mathbf{b}$ generano lo stesso sottospazio lineare di $\mathbb{R}^m$ di $\mathbf{c}_1, \dots, \mathbf{c}_n$, la cui dimensione è uguale a

$$\operatorname{rg} \begin{pmatrix} | & | & & | \\ \mathbf{c}_1 & \mathbf{c}_2 & \dots & \mathbf{c}_n \\ | & | & & | \end{pmatrix} = \operatorname{rg} \begin{pmatrix} | & | & & | & | \\ \mathbf{c}_1 & \mathbf{c}_2 & \dots & \mathbf{c}_n & \mathbf{b} \\ | & | & & | & | \end{pmatrix}. \quad (3.11.1)$$

$\underbrace{\hspace{10em}}_{=A}$

Supponiamo viceversa che valga (3.11.1). Se $\mathbf{X}$ ed $\mathbf{X}_{\mathbf{b}}$ indicano i sottospazi lineari di $\mathbb{R}^m$ generati rispettivamente

- (1) dalle colonne di A e
- (2) dalle colonne di A e dal termine noto $\mathbf{b}$,

allora (3.11.1) significa che $\dim \mathbf{X} = \dim \mathbf{X}_{\mathbf{b}}$ e quindi implica l'uguaglianza $\mathbf{X} = \mathbf{X}_{\mathbf{b}}$. Ora, se il rango di A , è r , esistono r colonne $\mathbf{c}_{k_1}, \dots, \mathbf{c}_{k_r}$ di A che costituiscono una base per $\mathbf{X}$. Poiché $\mathbf{b} \in \mathbf{X}$, segue che $\mathbf{b}$

è combinazione lineare dei vettori $\mathbf{c}_{k_1}, \dots, \mathbf{c}_{k_r}$, e quindi, a maggior ragione, di tutte le colonne di A . Ora, per (3.6.1), questo significa che il sistema $A\mathbf{x} = \mathbf{b}$ ammette soluzioni. $\square$

ESEMPIO 3.11.6. Troviamo tutti i valori del parametro reale a per i quali il sistema di equazioni lineari

$$\begin{cases} 2x_1 - x_2 - x_3 = 1 \\ -2x_1 + a^2x_2 + (1-a)x_3 = -1 \\ x_1 + 3x_2 + x_3 = 2 \end{cases}$$

non ha soluzione.

SVOLGIMENTO. Per il teorema di Rouché-Capelli 3.11.5 il sistema ha soluzione se e soltanto se

$$\operatorname{rg} \begin{pmatrix} 2 & -1 & -1 \\ -2 & a^2 & 1-a \\ 1 & 3 & 1 \end{pmatrix} = \operatorname{rg} \begin{pmatrix} 2 & -1 & -1 & 1 \\ -2 & a^2 & 1-a & -1 \\ 1 & 3 & 1 & 2 \end{pmatrix}.$$

Tramite operazioni elementari sulle righe si ottiene

$$\begin{aligned} \operatorname{rg} \begin{pmatrix} 2 & -1 & -1 \\ -2 & a^2 & 1-a \\ 1 & 3 & 1 \end{pmatrix} &= \operatorname{rg} \begin{pmatrix} 1 & 3 & 1 \\ 2 & -1 & -1 \\ -2 & a^2 & 1-a \end{pmatrix} = \\ &= \operatorname{rg} \begin{pmatrix} 1 & 3 & 1 \\ 0 & -7 & -3 \\ 0 & 6+a^2 & 3-a \end{pmatrix} = \operatorname{rg} \begin{pmatrix} 1 & 3 & 1 \\ 0 & 1 & \frac{3}{7} \\ 0 & 6+a^2 & 3-a \end{pmatrix} = \\ &= \operatorname{rg} \begin{pmatrix} 1 & 3 & 1 \\ 0 & 1 & \frac{3}{7} \\ 0 & 0 & \frac{1}{7}(3-7a-3a^2) \end{pmatrix} = \begin{cases} 3 & \text{se } 3-7a-3a^2 \neq 0, \\ 2 & \text{se } 3-7a-3a^2 = 0. \end{cases} \end{aligned}$$

Ne segue che nel caso $3 - 7a - 3a^2 \neq 0$ il sistema ha soluzione. Se invece $3 - 7a - 3a^2 = 0$, allora abbiamo

$$\begin{aligned} & \operatorname{rg} \begin{pmatrix} 2 & -1 & -1 & 1 \\ -2 & a^2 & 1-a & -1 \\ 1 & 3 & 1 & 2 \end{pmatrix} = \operatorname{rg} \begin{pmatrix} 1 & 3 & 1 & 2 \\ 2 & -1 & -1 & 1 \\ -2 & a^2 & 1-a & -1 \end{pmatrix} = \\ & = \operatorname{rg} \begin{pmatrix} 1 & 3 & 1 & 2 \\ 0 & -7 & -3 & -3 \\ 0 & 6+a^2 & 3-a & 3 \end{pmatrix} = \operatorname{rg} \begin{pmatrix} 1 & 3 & 1 & 2 \\ 0 & 1 & \frac{3}{7} & \frac{3}{7} \\ 0 & 6+a^2 & 3-a & 3 \end{pmatrix} = \\ & = \operatorname{rg} \begin{pmatrix} 1 & 3 & 1 & 2 \\ 0 & 1 & \frac{3}{7} & \frac{3}{7} \\ 0 & 0 & \frac{1}{7}(3-7a-3a^2) & \frac{1}{7}(3-3a^2) \end{pmatrix} = \\ & = \operatorname{rg} \begin{pmatrix} 1 & 3 & 1 & 2 \\ 0 & 1 & \frac{3}{7} & \frac{3}{7} \\ 0 & 0 & 0 & a \end{pmatrix} = 3 \end{aligned}$$

perché $3 - 7a - 3a^2 = 0 \Rightarrow a \neq 0$. Perciò il sistema non ha soluzione esattamente quando $3 - 7a - 3a^2 = 0$, ossia $a = \frac{-7 \pm \sqrt{85}}{6}$. $\square$

3.12. Esercizi sulle trasformazioni lineari

ESERCIZIO 3.12.1. Si verifichi che l'applicazione

$$\mathbb{R}^3 \ni \begin{pmatrix} x_1 \\ x_2 \\ x_3 \end{pmatrix} \mapsto \begin{pmatrix} x_1 + x_2 \\ x_2 + x_3 \end{pmatrix} \in \mathbb{R}^2$$

è lineare e si trovi la matrice associata.

SOLUZIONE. Gli elementi della base naturale di $\mathbb{R}^*$ sono mandati in

$$\begin{pmatrix} 1 \\ 0 \end{pmatrix}, \begin{pmatrix} 1 \\ 1 \end{pmatrix}, \begin{pmatrix} 0 \\ 1 \end{pmatrix}.$$

Pertanto la matrice associata è

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

e l'applicazione può essere scritta nella forma

$$\mathbb{R}^3 \ni \begin{pmatrix} x_1 \\ x_2 \\ x_3 \end{pmatrix} \mapsto \begin{pmatrix} 1 & 1 & 0 \\ 0 & 1 & 1 \end{pmatrix} \begin{pmatrix} x_1 \\ x_2 \\ x_3 \end{pmatrix} \in \mathbb{R}^2.$$

□

ESERCIZIO 3.12.2. Si trovi una matrice A di tipo 4×3 tale che, per ogni $\lambda_1, \lambda_2, \lambda_3 \in \mathbb{R}$, la combinazione lineare

$$\lambda_1 \begin{pmatrix} 2 \\ -3 \\ 0 \\ -1 \end{pmatrix} + \lambda_2 \begin{pmatrix} -2 \\ 1 \\ 1 \\ 0 \end{pmatrix} + \lambda_3 \begin{pmatrix} 1 \\ 5 \\ 2 \\ -3 \end{pmatrix}$$

sia uguale al prodotto

$$A \begin{pmatrix} \lambda_1 \\ \lambda_2 \\ \lambda_3 \end{pmatrix}.$$

SOLUZIONE. Si usa l'uguaglianza (3.6.1) :

$$A = \begin{pmatrix} 2 & -2 & 1 \\ -3 & 1 & 5 \\ 0 & 1 & 2 \\ -1 & 0 & -3 \end{pmatrix}. \quad \square$$

ESERCIZIO 3.12.3. Si trovi la matrice dell'applicazione lineare $T : \mathbb{R}^3 \longrightarrow \mathbb{R}^3$ che manda

$$\begin{pmatrix} 1 \\ 0 \\ 0 \end{pmatrix} \mapsto \begin{pmatrix} 2 \\ 1 \\ 3 \end{pmatrix}, \quad \begin{pmatrix} 0 \\ 1 \\ 0 \end{pmatrix} \mapsto \begin{pmatrix} -1 \\ 2 \\ 1 \end{pmatrix}, \quad \begin{pmatrix} 0 \\ 0 \\ 1 \end{pmatrix} \mapsto \begin{pmatrix} 3 \\ 4 \\ -2 \end{pmatrix}.$$

SOLUZIONE. Le colonne della matrice A associata a T sono le immagini tramite T degli elementi della base naturale di $\mathbb{R}^3$, perciò

$$A = \begin{pmatrix} 2 & -1 & 3 \\ 1 & 2 & 4 \\ 3 & 1 & -2 \end{pmatrix}.$$

□

ESERCIZIO 3.12.4. Esiste una applicazione lineare $T : \mathbb{R}^3 \longrightarrow \mathbb{R}^3$ che manda

$$\begin{pmatrix} 0 \\ -1 \\ 1 \end{pmatrix} \mapsto \begin{pmatrix} 1 \\ 0 \\ 0 \end{pmatrix}, \quad \begin{pmatrix} 1 \\ 3 \\ -2 \end{pmatrix} \mapsto \begin{pmatrix} 0 \\ 1 \\ 0 \end{pmatrix}, \quad \begin{pmatrix} 2 \\ 0 \\ 1 \end{pmatrix} \mapsto \begin{pmatrix} 0 \\ 0 \\ 1 \end{pmatrix} ?$$

Se esistono tali applicazioni, se ne trovi una.

SOLUZIONE. Se T esistesse, sarebbe surgettiva, quindi, per (I), bigettiva. Inoltre T^{-1} manderebbe

$$\begin{pmatrix} 1 \\ 0 \\ 0 \end{pmatrix} \mapsto \begin{pmatrix} 0 \\ -1 \\ 1 \end{pmatrix}, \quad \begin{pmatrix} 0 \\ 1 \\ 0 \end{pmatrix} \mapsto \begin{pmatrix} 1 \\ 3 \\ -2 \end{pmatrix}, \quad \begin{pmatrix} 0 \\ 0 \\ 1 \end{pmatrix} \mapsto \begin{pmatrix} 2 \\ 0 \\ 1 \end{pmatrix},$$

perciò, come nell'Esercizio 3.12.2, la matrice associata a T^{-1} dovrebbe essere

$$\begin{pmatrix} 0 & 1 & 2 \\ -1 & 3 & 0 \\ 1 & -2 & 1 \end{pmatrix}.$$

Ma questa matrice è invertibile, avendo l'inversa

$$\begin{pmatrix} 0 & 1 & 2 \\ -1 & 3 & 0 \\ 1 & -2 & 1 \end{pmatrix}^{-1} = \begin{pmatrix} -3 & 5 & 6 \\ -1 & 2 & 2 \\ 1 & -1 & -1 \end{pmatrix}$$

(si veda l'Esempio 3.10.1). Ne concludiamo che

$$T : \mathbb{R}^3 \ni \begin{pmatrix} x_1 \\ x_2 \\ x_3 \end{pmatrix} \longmapsto \begin{pmatrix} -3 & 5 & 6 \\ -1 & 2 & 2 \\ 1 & -1 & -1 \end{pmatrix} \begin{pmatrix} x_1 \\ x_2 \\ x_3 \end{pmatrix} \in \mathbb{R}^3$$

è l'unica applicazione lineare con le proprietà desiderate.

□

ESERCIZIO 3.12.5. Si trovino il nucleo e l'immagine dell'applicazione lineare

$$T : \mathbb{R}^3 \ni \begin{pmatrix} x_1 \\ x_2 \\ x_3 \end{pmatrix} \mapsto \begin{pmatrix} 2 & -1 & 1 \\ 1 & 1 & 5 \\ 2 & 0 & 4 \end{pmatrix} \begin{pmatrix} x_1 \\ x_2 \\ x_3 \end{pmatrix} \in \mathbb{R}^3 .$$

SOLUZIONE. Il nucleo di T consiste da tutte le soluzioni del sistema omogeneo

$$\begin{cases} 2x_1 - x_2 + x_3 = 0 \\ x_1 + x_2 + 5x_3 = 0 \\ 2x_1 \quad \quad + 4x_3 = 0 \end{cases}$$

che ora risolviamo.

$$\begin{array}{ccc|c} 2 & -1 & 1 & 0 \\ 1 & 1 & 5 & 0 \\ 2 & 0 & 4 & 0 \end{array} , \quad \begin{array}{ccc|c} 1 & 1 & 5 & 0 \\ 2 & -1 & 1 & 0 \\ 2 & 0 & 4 & 0 \end{array} , \quad \begin{array}{ccc|c} 1 & 1 & 5 & 0 \\ 0 & -3 & -9 & 0 \\ 0 & -2 & -6 & 0 \end{array}$$

$$\begin{array}{ccc|c} 1 & 1 & 5 & 0 \\ 0 & 1 & 3 & 0 \\ 0 & 0 & 0 & 0 \end{array} ,$$

$x_3 =$ arbitrario

$$\begin{array}{l} x_2 = -3x_3 \\ x_1 = -x_2 - 5x_3 = -2x_3 \end{array} , \text{ cioè } \begin{pmatrix} x_1 \\ x_2 \\ x_3 \end{pmatrix} = \begin{pmatrix} -2x_3 \\ -3x_3 \\ x_3 \end{pmatrix} = x_3 \begin{pmatrix} -2 \\ -3 \\ 1 \end{pmatrix} .$$

Concludiamo che $\text{Ker}(T)$ consiste da tutti i multipli scalari di $\begin{pmatrix} -2 \\ -3 \\ 1 \end{pmatrix}$:

$$\text{Ker}(T) = \left\{ \lambda \begin{pmatrix} -2 \\ -3 \\ 1 \end{pmatrix} ; \lambda \in \mathbb{R} \right\} .$$

D'altra parte, l'immagine di T è il sottospazio lineare di $\mathbb{R}^3$ generato dalle colonne della matrice associata, cioè dai vettori

$$\begin{pmatrix} 2 \\ 1 \\ 2 \end{pmatrix} , \quad \begin{pmatrix} -1 \\ 1 \\ 0 \end{pmatrix} , \quad \begin{pmatrix} 1 \\ 5 \\ 4 \end{pmatrix} .$$

Volendo, possiamo trovare una base per questo sottospazio, scrivendo i vettori di cui sopra come le righe di una matrice, svolgendo l'eliminazione di Gauss e poi considerando le righe non zero della matrice risultante:

$$\begin{array}{cccc} 2 & 1 & 2 & -1 & 1 & 0 & -1 & 1 & 0 & -1 & 1 & 0 \\ -1 & 1 & 0 & 2 & 1 & 2 & 0 & 3 & 2 & 0 & 3 & 2 \\ 1 & 5 & 4 & 1 & 5 & 4 & 0 & 6 & 4 & 0 & 0 & 0 \end{array}$$

Così otteniamo la base $\begin{pmatrix} -1 \\ 1 \\ 0 \end{pmatrix}, \begin{pmatrix} 0 \\ 3 \\ 2 \end{pmatrix}$ per $T(\mathbb{R}^3)$:

$$T(\mathbb{R}^3) = \left\{ \lambda_1 \begin{pmatrix} -1 \\ 1 \\ 0 \end{pmatrix} + \lambda_2 \begin{pmatrix} 0 \\ 3 \\ 2 \end{pmatrix} ; \lambda_1, \lambda_2 \in \mathbb{R} \right\}.$$

□

ESERCIZIO 3.12.6. Si determini il rango della matrice

$$\begin{pmatrix} 2 & 3 & 1 & 1 \\ 1 & -1 & 5 & 3 \\ 5 & 5 & 3 & 5 \\ 3 & 7 & -7 & -1 \end{pmatrix}.$$

SOLUZIONE. Tramite operazioni elementari sulle righe si ottiene

$$\begin{aligned} \text{rg} \begin{pmatrix} 2 & 3 & 1 & 1 \\ 1 & -1 & 5 & 3 \\ 5 & 5 & 3 & 5 \\ 3 & 7 & -7 & -1 \end{pmatrix} &= \text{rg} \begin{pmatrix} 1 & -1 & 5 & 3 \\ 2 & 3 & 1 & 1 \\ 2 & -2 & 10 & 6 \\ 3 & 7 & -7 & -1 \end{pmatrix} = \\ &= \text{rg} \begin{pmatrix} 1 & -1 & 5 & 3 \\ 0 & 5 & -9 & -5 \\ 0 & 0 & 0 & 0 \\ 0 & 10 & -22 & -10 \end{pmatrix} = \text{rg} \begin{pmatrix} 1 & -1 & 5 & 3 \\ 0 & 5 & -9 & -5 \\ 0 & 0 & 0 & 0 \\ 0 & 0 & -4 & 0 \end{pmatrix} = \\ &= 3. \end{aligned}$$

□

ESERCIZIO 3.12.7. Si dica se la matrice

$$A = \begin{pmatrix} 2 & 7 & 3 \\ 3 & 9 & 4 \\ 1 & 5 & 3 \end{pmatrix}$$

è invertibile o no. Nel caso di invertibilità si calcoli la sua matrice inversa.

SOLUZIONE. Tramite operazioni elementari sulle righe si azzerano i coefficienti sotto la diagonale:

$$\begin{array}{ccc|ccc} 2 & 7 & 3 & 1 & 0 & 0 \\ 3 & 9 & 4 & 0 & 1 & 0 \\ 1 & 5 & 3 & 0 & 0 & 1 \end{array}, \quad \begin{array}{ccc|ccc} 1 & 5 & 3 & 0 & 0 & 1 \\ 2 & 7 & 3 & 1 & 0 & 0 \\ 3 & 9 & 4 & 0 & 1 & 0 \end{array},$$

$$\begin{array}{ccc|ccc} 1 & 5 & 3 & 0 & 0 & 1 \\ 0 & -3 & -3 & 1 & 0 & -2 \\ 0 & -6 & -5 & 0 & 1 & -3 \end{array}, \quad \begin{array}{ccc|ccc} 1 & 5 & 3 & 0 & 0 & 1 \\ 0 & 1 & 1 & -1/3 & 0 & 2/3 \\ 0 & 0 & 1 & -2 & 1 & 1 \end{array},$$

Poiché tutti gli elementi del diagonale sono non zero, la matrice A è diagonalizzabile. Usando di nuovo operazioni elementari sulle righe, si azzerano anche i coefficienti sopra la diagonale:

$$\begin{array}{ccc|ccc} 1 & 5 & 0 & 6 & -3 & -2 \\ 0 & 1 & 0 & 5/3 & -1 & -1/3 \\ 0 & 0 & 1 & -2 & 1 & 1 \end{array}, \quad \begin{array}{ccc|ccc} 1 & 0 & 0 & -7/3 & 2 & -1/3 \\ 0 & 1 & 0 & 5/3 & -1 & -1/3 \\ 0 & 0 & 1 & -2 & 1 & 1 \end{array}.$$

Così

$$A^{-1} = \begin{pmatrix} -7/3 & 2 & -1/3 \\ 5/3 & -1 & -1/3 \\ -2 & 1 & 1 \end{pmatrix}.$$

□

ESERCIZIO 3.12.8. Si dica se la matrice

$$A = \begin{pmatrix} 2 & 1 & -1 & 0 \\ 0 & -2 & 0 & 1 \\ -3 & 0 & 1 & 0 \\ 0 & 2 & 1 & -2 \end{pmatrix}$$

è invertibile o no. Nel caso di invertibilità si calcoli la sua matrice inversa.

SOLUZIONE. Tramite operazioni elementari sulle righe si azzerano i coefficienti sotto la diagonale:

$$\begin{array}{cccc|cccc|cccc} 2 & 1 & -1 & 0 & 1 & 0 & 0 & 0 & -1 & 1 & 0 & 0 & 1 & 0 & 1 & 0 \\ 0 & -2 & 0 & 1 & 0 & 1 & 0 & 0 & 0 & -2 & 0 & 1 & 0 & 1 & 0 & 0 \\ -3 & 0 & 1 & 0 & 0 & 0 & 1 & 0 & -3 & 0 & 1 & 0 & 0 & 0 & 1 & 0 \\ 0 & 2 & 1 & -2 & 0 & 0 & 0 & 1 & 0 & 2 & 1 & -2 & 0 & 0 & 0 & 1 \end{array}$$

$$\begin{array}{cccc|cccc} -1 & 1 & 0 & 0 & 1 & 0 & 1 & 0 \\ 0 & -2 & 0 & 1 & 0 & 1 & 0 & 0 \\ 0 & -3 & 1 & 0 & -3 & 0 & -2 & 0 \\ 0 & 2 & 1 & -2 & 0 & 0 & 0 & 1 \end{array}$$

$$\begin{array}{cccc|cccc} -1 & 1 & 0 & 0 & 1 & 0 & 1 & 0 \\ 0 & -2 & 0 & 1 & 0 & 1 & 0 & 0 \\ 0 & -1 & 1 & -1 & -3 & -1 & -2 & 0 \\ 0 & 0 & 1 & -1 & 0 & 1 & 0 & 1 \end{array}$$

$$\begin{array}{cccc|cccc|cccc|cccc} -1 & 1 & 0 & 0 & 1 & 0 & 1 & 0 & -1 & 1 & 0 & 0 & 1 & 0 & 1 & 0 \\ 0 & 1 & -1 & 1 & 3 & 1 & 2 & 0 & 0 & 1 & -1 & 1 & 3 & 1 & 2 & 0 \\ 0 & -2 & 0 & 1 & 0 & 1 & 0 & 0 & 0 & 0 & -2 & 3 & 6 & 3 & 4 & 0 \\ 0 & 0 & 1 & -1 & 0 & 1 & 0 & 1 & 0 & 0 & 1 & -1 & 0 & 1 & 0 & 1 \end{array}$$

$$\begin{array}{cccc|cccc|cccc} -1 & 1 & 0 & 0 & 1 & 0 & 1 & 0 & -1 & 1 & 0 & 0 & 1 & 0 & 1 & 0 \\ 0 & 1 & -1 & 1 & 3 & 1 & 2 & 0 & 0 & 1 & -1 & 1 & 3 & 1 & 2 & 0 \\ 0 & 0 & 1 & -1 & 0 & 1 & 0 & 1 & 0 & 0 & 1 & -1 & 0 & 1 & 0 & 1 \\ 0 & 0 & -2 & 3 & 6 & 3 & 4 & 0 & 0 & 0 & 0 & 1 & 6 & 5 & 4 & 2 \end{array} \cdot$$

Poiché tutti gli elementi della diagonale sono non nulli, la matrice A è diagonalizzabile. Usando di nuovo operazioni elementari sulle righe, si azzeriamo anche i coefficienti sopra la diagonale:

$$\begin{array}{cccc|cccc|cccc|cccc} 1 & -1 & 0 & 0 & -1 & 0 & -1 & 0 & 1 & 0 & 0 & 0 & 2 & 2 & 1 & 1 \\ 0 & 1 & 0 & 0 & 3 & 2 & 2 & 1 & 0 & 1 & 0 & 0 & 3 & 2 & 2 & 1 \\ 0 & 0 & 1 & -1 & 0 & 1 & 0 & 1 & 0 & 0 & 1 & 0 & 6 & 6 & 4 & 3 \\ 0 & 0 & 0 & 1 & 6 & 5 & 4 & 2 & 0 & 0 & 0 & 1 & 6 & 5 & 4 & 2 \end{array} \cdot$$

Pertanto

$$A^{-1} = \begin{pmatrix} 2 & 2 & 1 & 1 \\ 3 & 2 & 2 & 1 \\ 6 & 6 & 4 & 3 \\ 6 & 5 & 4 & 2 \end{pmatrix}.$$

□

3.13. Isomorfismi fra spazi vettoriali

NOTAZIONE 3.13.1. Una applicazione lineare da uno spazio vettoriale $\mathbf{X}$ ad uno spazio $\mathbf{Y}$ si chiama un omomorfismo del primo spazio nel secondo.

DEFINIZIONE 3.13.2. Un omomorfismo da $\mathbf{X}$ a $\mathbf{Y}$ biunivoco, cioè iniettivo e surgettivo, si chiama un isomorfismo fra $\mathbf{X}$ e $\mathbf{Y}$.

PROPOSIZIONE 3.13.3. Sia $T : \mathbf{X} \longrightarrow \mathbf{Y}$ un isomorfismo. Allora $\dim \mathbf{X} = \dim \mathbf{Y}$.

DIMOSTRAZIONE. Poiché T è iniettivo, si ha $\text{Ker}(T) = \mathbf{0}$ e quindi segue dal teorema della dimensione (Teorema 3.2.1) che $\dim \mathbf{X} = \dim T\mathbf{Y}$. Poiché T è surgettivo, l'enunciato è dimostrato.

NOTA 3.13.4. Un omomorfismo iniettivo e surgettivo è invertibile, e l'inversa è ancora un omomorfismo; inoltre è ancora iniettiva e surgettiva, quindi un isomorfismo, che indichiamo con *l'isomorfismo inverso*.

□

PROPOSIZIONE 3.13.5. *Data una base $\mathcal{V}$ in uno spazio vettoriale $\mathbf{X}$ di dimensione n , la mappa delle coordinate $\mathcal{F}_{\mathcal{V}}$ (si veda la Notazione 4.1.1) è un isomorfismo fra $\mathbf{X}$ e $\mathbb{R}^n$.*

DIMOSTRAZIONE. È chiaro che $\mathcal{F}_{\mathcal{V}}$ è un omomorfismo surgettivo; l'injectività equivale al fatto che gli elementi della base sono linearmente indipendenti.

NOTA 3.13.6. Per il momento stiamo considerando spazi vettoriali sul campo reale, ma in seguito estenderemo lo studio anche agli spazi vettoriali sul campo complesso. La Proposizione 3.13.5 si estende agli spazi complessi, con identica dimostrazione: ogni spazio vettoriale di dimensione n sui complessi è isomorfo a $\mathbb{C}^n$.

□

COROLLARIO 3.13.7. *Due spazi vettoriali della stessa dimensione (finita) sono isomorfi.*

CAPITOLO 4

Cambiamento di base

4.1. Trasformazione di coordinate sotto cambiamento di base

Sia $\mathbf{X}$ uno spazio vettoriale e $\mathcal{V} = \{\mathbf{v}_1, \dots, \mathbf{v}_n\}$ una base di $\mathbf{X}$. Allora

$$\mathbb{R}^n \ni \begin{pmatrix} \lambda_1 \\ \vdots \\ \lambda_n \end{pmatrix} \mapsto \lambda_1 \mathbf{v}_1 + \dots + \lambda_n \mathbf{v}_n \in \mathbf{X}$$

è una applicazione lineare bigettiva e la sua applicazione inversa è

$$F_{\mathcal{V}} : \mathbf{X} \ni \mathbf{x} \mapsto \begin{pmatrix} \lambda_1^{(\mathcal{V})}(\mathbf{x}) \\ \vdots \\ \lambda_n^{(\mathcal{V})}(\mathbf{x}) \end{pmatrix} \in \mathbb{R}^n,$$

ove $\lambda_1^{(\mathcal{V})}(\mathbf{x}), \dots, \lambda_n^{(\mathcal{V})}(\mathbf{x})$ sono i coefficienti di $\mathbf{v}_1, \dots, \mathbf{v}_n$ nello sviluppo di $\mathbf{x}$ quale combinazione lineare degli elementi della base $\mathcal{V}$:

$$\mathbf{x} = \lambda_1^{(\mathcal{V})}(\mathbf{x}) \mathbf{v}_1 + \dots + \lambda_n^{(\mathcal{V})}(\mathbf{x}) \mathbf{v}_n.$$

$F_{\mathcal{V}}$ applica $\mathbf{v}_k$ nel k -esimo elemento della base naturale di $\mathbb{R}^n$, quindi $\mathcal{V}$ nella base naturale di $\mathbb{R}^n$. Perciò $F_{\mathcal{V}}$ ci offre una identificazione dello spazio vettoriale $\mathbf{X}$ con $\mathbb{R}^n$, tale che ad $\mathcal{V}$ corrisponda la base naturale di $\mathbb{R}^n$.

NOTAZIONE 4.1.1. *La mappa $F_{\mathcal{V}}$ introdotto qui sopra è la mappa delle coordinate, che per ogni vettore restituisce le sue coordinate rispetto alla base scelta.*

Facciamo un esempio.

ESEMPIO 4.1.2. Siano $\mathcal{E} = \{\mathbf{e}_1, \mathbf{e}_2, mbe_3\}$ la base canonica

$$\mathcal{E} = \mathcal{E}_3 = \left\{ \begin{pmatrix} 1 \\ 0 \\ 0 \end{pmatrix}, \dots, \begin{pmatrix} 0 \\ 0 \\ 1 \end{pmatrix} \right\}$$

di $\mathbb{R}^3$ e $\mathcal{V} = \{\mathbf{v}_1, \mathbf{v}_2, mbv_3\}$ la base seguente: $\mathbf{v}_1 = \mathbf{e}_1 + \mathbf{e}_2 = \begin{pmatrix} 1 \\ 1 \\ 0 \end{pmatrix}$,

$$\mathbf{v}_2 = \mathbf{e}_1 - \mathbf{e}_2 = \begin{pmatrix} 1 \\ -1 \\ 0 \end{pmatrix}, \mathbf{v}_3 = \mathbf{e}_3 = \begin{pmatrix} 0 \\ 0 \\ 1 \end{pmatrix}. \text{ Sia } \mathbf{x} = \begin{pmatrix} x_1 \\ x_2 \\ x_3 \end{pmatrix} : \text{quindi}$$

x_i è la i -sima coordinata (o componente) di $\mathbf{x}$ nella base canonica.

Troviamo le componenti y_i dello stesso vettore $\mathbf{x}$ nella base $\mathcal{V}$.

SVOLGIMENTO. Per prima cosa scriviamo la matrice M (nella base canonica!) della applicazione lineare U che manda la base canonica nella base $\mathcal{V}$ (nell'ordine: $U(\mathbf{e}_i) = (v)_i$, $i = 1, 2, 3$). Questo è facile: la i -sima colonna di M_U è il vettore $\mathbf{v}_i$, cioè

$$M = \begin{pmatrix} 1 & 1 & 0 \\ 1 & -1 & 0 \\ 0 & 0 & 1 \end{pmatrix}.$$

Un facile calcolo basato sull'eliminazione di Gauss, che lasciamo per esercizio, mostra che la matrice inversa M^{-1} , che indichiamo con $C = C_{\mathcal{V}, \mathcal{E}}$ è

$$C = \begin{pmatrix} \frac{1}{2} & \frac{1}{2} & 0 \\ \frac{1}{2} & -\frac{1}{2} & 0 \\ 0 & 0 & 1 \end{pmatrix}.$$

(si arriva allo stesso risultato risolvendo per sostituzione il sistema lineare

$$A \begin{pmatrix} x_1 \\ x_2 \\ x_3 \end{pmatrix} = \begin{pmatrix} y_1 \\ y_2 \\ y_3 \end{pmatrix}$$

il che è facile perché la terza equazione del sistema è $x_3 = y_3$, e quindi in pratica il sistema lineare è bidimensionale invece che tridimensionale).

4.1. TRASFORMAZIONE DI COORDINATE SOTTO CAMBIAMENTO DI BASE 17

Abbiamo $\mathbf{x} = \sum_{i=1}^3 x_i \mathbf{e}_i$ e vogliamo trovare i numeri y_i tali che $\mathbf{x} = \sum_{j=1}^3 y_j \mathbf{v}_j$. D'altra parte, $U^{-1} \mathbf{v}_i = \mathbf{e}_i$, quindi

$$\mathbf{x} = \sum_{i=1}^3 x_i \mathbf{e}_i = \sum_{i=1}^3 x_i U^{-1} \mathbf{v}_i = \sum_{i=1}^3 x_i \sum_{j=1}^3 A_{ij} \mathbf{v}_j.$$

Scambiamo ora l'ordine delle due somme:

$$\sum_{j=1}^3 y_j \mathbf{v}_j = \sum_{j=1}^3 \sum_{i=1}^3 x_i A_{ij} \mathbf{v}_j.$$

Poiché i vettori $\mathbf{v}_j$, $j = 1, 2, 3$ sono una base i coefficienti dell'ultima uguaglianza devono essere uguali a sinistra e a destra: quindi per $j = 1, 2, 3$

$$y_j = \sum_{i=1}^3 A_{ij}.$$

L'ultima uguaglianza si può scrivere come prodotto righe per colonne

$$\begin{pmatrix} y_1 \\ y_2 \\ y_3 \end{pmatrix} = (x_1, x_2, x_3) A,$$

dove però il vettore $\mathbf{x}$ ora è il fattore di sinistra nel prodotto, ed è quindi scritto come vettore riga anziché colonna (altrimenti il prodotto sarebbe stato colonne per righe!).

In questo esempio, svolgendo il prodotto, si trova la seguente regola di trasformazione di coordinate:

$$\begin{pmatrix} y_1 \\ y_2 \\ y_3 \end{pmatrix} = \begin{pmatrix} \frac{1}{2}x_1 + \frac{1}{2}x_2 \\ \frac{1}{2}x_1 - \frac{1}{2}x_2 \\ x_3 \end{pmatrix}.$$

□

NOTA 4.1.3. Cautela: la matrice che realizza la trasformazione delle coordinate nel passaggio da una base ad un'altra è *l'inversa* della matrice del cambiamento di base. Infatti, sebbene la matrice che trasforma i vettori $\mathbf{e}_i$ nei $\mathbf{v}_i$ sia M_U , nella trasformazione delle coordinate dei vettori sotto questo cambiamento di base interviene invece la matrice

inversa $C = M_U^{-1}$. Si dice che le coordinate si trasformano in maniera controgradiente rispetto ai vettori di base. $\square$

Riassumendo quanto esposto nell'Esempio 4.1.2, abbiamo mostrato:

PROPOSIZIONE 4.1.4. (Trasformazione di coordinate per cambiamento di riferimento dalla base canonica ad un'altra base.) *Se indichiamo con $\mathbf{e}_i$ i vettori della base canonica e con $\mathbf{v}_i$ ($i = 1, \dots, n$) quelli di un'altra base in $\mathbb{R}^n$, e con U la trasformazione lineare che verifica $U\mathbf{e}_i = \mathbf{v}_i$, allora la trasformazione delle coordinate dello stesso vettore $\mathbf{x} = \sum_{i=1}^n x_i \mathbf{e}_i = \sum_{i=1}^n y_i \mathbf{v}_i$ sotto questo cambiamento di base è data da*

$$y_j = \sum_{i=1}^n x_i A_{ij}, \quad (4.1.1)$$

dove $A := M_U^{-1}$ è l'inversa della matrice M_U che esprime la trasformazione lineare U nella base canonica.

NOTAZIONE 4.1.5. (Matrice di passaggio dalla base $\mathcal{F}$ alla base $\mathcal{V}$.) *A causa del suo ruolo nella regola di trasformazione di coordinate, la matrice $C = C_{\mathcal{V}, \mathcal{F}} := M_U^{-1}$ si chiama la matrice di passaggio (o del cambiamento di base) dalla base canonica*

$$\mathcal{E}_n = \left\{ \begin{pmatrix} 1 \\ \vdots \\ 0 \end{pmatrix}, \dots, \begin{pmatrix} 0 \\ \vdots \\ 1 \end{pmatrix} \right\}$$

alla base $\mathcal{V}$. Essa è quindi la matrice che realizza la seguente trasformazione $F_{\mathcal{V}} \circ F_{\mathcal{E}}^{-1}$:

$$\mathbb{R}^n \xrightarrow{F_{\mathcal{E}}^{-1}} (X) \xrightarrow{F_{\mathcal{V}}} \mathbb{R}^n$$

Rammentiamo che la colonna k -sima di tale matrice consiste dei coefficienti dello sviluppo rispetto alla base $\mathcal{V}$ del k -simo vettore della base canonica (per maggiori dettagli si veda la successiva Nota 4.1.7).

Più in generale, nel passare da una base $\mathcal{F}$ alla base $\mathcal{V}$ la matrice $C = C_{\mathcal{V}, \mathcal{F}}$ che realizza la corrispondente trasformazione di coordinate $F_{\mathcal{V}} \circ \mathcal{F}^{-1}$

$$\mathbb{R}^n \xrightarrow{F_{\mathcal{F}}^{-1}} (X) \xrightarrow{F_{\mathcal{V}}} \mathbb{R}^n$$

è quella associata alla trasformazione lineare $F_{\mathcal{V}} \circ F_{\mathcal{F}}^{-1}$

NOTA 4.1.6. È immediato che $C_{\mathcal{V},\mathcal{F}} = C_{\mathcal{V},\mathcal{E}}C_{\mathcal{E},\mathcal{F}}$, perché $F_{\mathcal{V}} \circ \mathcal{F}^{-1} = F_{\mathcal{V}} \circ \mathcal{E}^{-1} \circ F_{\mathcal{E}} \circ \mathcal{F}^{-1}$. Quindi tutte le matrici di passaggio si esprimono in termini di matrici di passaggio dalla base canonica e loro inverse.

□

NOTA 4.1.7. (**Forma della matrice di cambiamento di base.**)

Estendendo l'osservazione fatta per il passaggio dalla base canonica alla base $\mathcal{V}$, osserviamo ora che la colonna k -esima di $C_{\mathcal{F},\mathcal{E}}$ consiste dei coefficienti dello sviluppo di $\mathbf{e}_k$ rispetto alla base $\mathcal{F}$: se

$$\mathbf{e}_k = \sum_{j=1}^n c_{jk} \mathbf{f}_j,$$

allora la colonna k -esima di $C_{\mathcal{F},\mathcal{E}}$ è

$$\begin{pmatrix} c_{1k} \\ c_{2k} \\ \vdots \\ c_{nk} \end{pmatrix},$$

quindi

$$C_{\mathcal{F},\mathcal{E}} = \begin{pmatrix} c_{11} & c_{12} & \dots & c_{1n} \\ c_{21} & c_{22} & \dots & c_{2n} \\ \vdots & \vdots & \ddots & \vdots \\ c_{n1} & c_{n2} & \dots & c_{nn} \end{pmatrix}. \quad (4.1.2)$$

Nello spazio vettoriale $\mathbb{R}^n$ c'è un modo semplice per scrivere la matrice di passaggio da una base ad un'altra.

Calcoliamo prima la matrice di passaggio da una base arbitraria $\mathcal{F} = \{\mathbf{f}_1, \dots, \mathbf{f}_n\}$ di $\mathbb{R}^n$ alla base canonica

$$\mathcal{E}_n = \left\{ \begin{pmatrix} 1 \\ \vdots \\ 0 \end{pmatrix}, \dots, \begin{pmatrix} 0 \\ \vdots \\ 1 \end{pmatrix} \right\}$$

e viceversa. Poiché i coefficienti dello sviluppo di $\mathbf{f}_k$ rispetto ad $\mathcal{E}_n$ sono le componenti del vettore $\mathbf{f}_k$, la matrice di passaggio da $\mathcal{F}$ a $\mathcal{E}_n$ è

$$C_{\mathcal{E}_n, \mathcal{F}} = \begin{pmatrix} | & & | \\ \mathbf{f}_1 & \dots & \mathbf{f}_n \\ | & & | \end{pmatrix}.$$

Ora possiamo calcolare anche la matrice di passaggio da $\mathcal{E}_n$ a $\mathcal{F}$:

$$C_{\mathcal{F}, \mathcal{E}_n} = C_{\mathcal{E}_n, \mathcal{F}}^{-1} = \begin{pmatrix} | & & | \\ \mathbf{f}_1 & \dots & \mathbf{f}_n \\ | & & | \end{pmatrix}^{-1}.$$

Se $\mathcal{F} = \{\mathbf{f}_1, \dots, \mathbf{f}_n\}$ e $\mathcal{V} = \{\mathbf{v}_1, \dots, \mathbf{v}_n\}$ sono due basi qualsiasi di $\mathbb{R}^n$, allora calcoliamo la matrice di passaggio da $\mathcal{F}$ a $\mathcal{V}$ quale il prodotto della matrice di passaggio da $\mathcal{F}$ a $\mathcal{E}_n$ con la matrice di passaggio da $\mathcal{E}_n$ a $\mathcal{V}$:

$$C_{\mathcal{V}, \mathcal{F}} = C_{\mathcal{V}, \mathcal{E}_n} C_{\mathcal{E}_n, \mathcal{F}} = \begin{pmatrix} | & & | \\ \mathbf{v}_1 & \dots & \mathbf{v}_n \\ | & & | \end{pmatrix}^{-1} \begin{pmatrix} | & & | \\ \mathbf{f}_1 & \dots & \mathbf{f}_n \\ | & & | \end{pmatrix}.$$

□

4.2. Matrice di un'applicazione lineare e cambiamento di basi

DEFINIZIONE 4.2.1. (*Matrice associata ad una trasformazione lineare in una data coppia di basi.*) Siano $\mathbf{X}, \mathbf{Y}$ spazi vettoriali e $T : \mathbf{X} \longrightarrow \mathbf{Y}$ una applicazione lineare. Indichiamo con $\mathcal{V} = \{\mathbf{v}_1, \dots, \mathbf{v}_n\}$ una base di $\mathbf{X}$ e con $\mathcal{F} = \{\mathbf{f}_1, \dots, \mathbf{f}_n\}$ una base di $\mathbf{Y}$. Identificando $\mathbf{X}$ con $\mathbb{R}^n$ ed $\mathbf{Y}$ con $\mathbb{R}^m$ tramite $F_{\mathcal{V}}$ ed $F_{\mathcal{F}}$, T diventa una applicazione lineare $\mathbb{R}^n \longrightarrow \mathbb{R}^m$. Più precisamente, a $T : \mathbf{X} \longrightarrow \mathbf{Y}$ corrisponde univocamente (una volta scelte le basi) $T_{\mathcal{F}, \mathcal{V}} : \mathbb{R}^n \longrightarrow \mathbb{R}^m$ definita da $T_{\mathcal{F}, \mathcal{V}} := F_{\mathcal{F}} \circ T \circ F_{\mathcal{V}}^{-1}$:

$$\begin{array}{ccc} \mathbf{X} & \xrightarrow{T} & \mathbf{Y} \\ F_{\mathcal{V}}^{-1} \uparrow & & \downarrow F_{\mathcal{F}} \\ \mathbb{R}^n & \xrightarrow{T_{\mathcal{F}, \mathcal{V}}} & \mathbb{R}^m \end{array}$$

4.2. MATRICE DI UN'APPLICAZIONE LINEARE E CAMBIAMENTO DI BASE

La matrice $A_{\mathcal{F},\mathcal{V}} = A_{\mathcal{F},\mathcal{V}}(T) \in M_{mn}$ associata a $T_{\mathcal{F},\mathcal{V}}$ si chiama la *matrice associata a T rispetto alle basi $\mathcal{V}$ e $\mathcal{F}$* . Se $\mathbf{X} = \mathbf{Y}$ e $\mathcal{V} = \mathcal{F}$, allora questa matrice, associata alla trasformazione $F_{\mathcal{V}} \circ T \circ F_{\mathcal{V}}^{-1}$, viene chiamata la *matrice associata a T rispetto alla base $\mathcal{V}$* e la si indica semplicemente con $A_{\mathcal{V}}(T)$.

NOTA 4.2.2. La colonna k -sima di $A_{\mathcal{F},\mathcal{V}}$ è

$$\begin{aligned} & T_{\mathcal{F},\mathcal{V}} \left(\text{l'elemento } k\text{-simo della base naturale di } \mathbb{R}^n \right) \\ &= T_{\mathcal{F},\mathcal{V}}(F_{\mathcal{V}}(\mathbf{v}_k)) = F_{\mathcal{F}}(T(\mathbf{v}_k)), \end{aligned}$$

cioè i coefficienti dello sviluppo di $T(\mathbf{v}_k)$ rispetto alla base $\mathcal{F}$ di $\mathbf{Y}$ considerati come un vettore in $\mathbb{R}^m$. Allora, per la definizione di $A_{\mathcal{F},\mathcal{V}}$, l'immagine di

$$\mathbf{x} = \sum_{k=1}^n \lambda_k \mathbf{v}_k \in \mathbf{X}$$

tramite T è

$$T(\mathbf{x}) = \sum_{j=1}^m \mu_j \mathbf{f}_j \in \mathbf{Y},$$

ove

$$\begin{pmatrix} \mu_1 \\ \vdots \\ \mu_m \end{pmatrix} = A_{\mathcal{F},\mathcal{V}} \begin{pmatrix} \lambda_1 \\ \vdots \\ \lambda_n \end{pmatrix}.$$

□

ESEMPIO 4.2.3. Si calcoli la matrice dell'applicazione lineare T data da

$$\mathbb{R}^2 \ni \begin{pmatrix} x_1 \\ x_2 \end{pmatrix} \mapsto \begin{pmatrix} 1 & 1 \\ 1 & 3 \end{pmatrix} \begin{pmatrix} x_1 \\ x_2 \end{pmatrix} \in \mathbb{R}^2$$

rispetto alla base

$$\mathbf{v}_1 = \begin{pmatrix} 1 \\ 1 \end{pmatrix}, \quad \mathbf{v}_2 = \begin{pmatrix} 1 \\ -1 \end{pmatrix}$$

di $\mathbb{R}^2$.

SVOLGIMENTO. Sia

$$\mathbf{e}_1 = \begin{pmatrix} 1 \\ 0 \end{pmatrix}, \quad \mathbf{e}_2 = \begin{pmatrix} 0 \\ 1 \end{pmatrix}$$

la base canonica in $\mathbb{R}^2$. Osserviamo che

$$\mathbf{v}_1 = \mathbf{e}_1 + \mathbf{e}_2 \quad \text{e} \quad \mathbf{v}_2 = \mathbf{e}_1 - \mathbf{e}_2, \quad (4.2.1)$$

da cui

$$\mathbf{e}_1 = \frac{1}{2}(\mathbf{v}_1 + \mathbf{v}_2) \quad (4.2.2)$$

$$\mathbf{e}_2 = \frac{1}{2}(\mathbf{v}_1 - \mathbf{v}_2) \quad (4.2.3)$$

(si noti che (4.2.2) è il calcolo dell'*inversa della trasformazione di cambio di base* (4.2.1). Pertanto,

$$T(\mathbf{e}_1) = \mathbf{e}_1 + \mathbf{e}_2 = \mathbf{v}_1 \quad (4.2.4)$$

$$T(\mathbf{e}_2) = \mathbf{e}_1 + 3\mathbf{e}_2 = \mathbf{v}_1 + 2\mathbf{e}_2 = \frac{3}{2}\mathbf{v}_1 - \mathbf{v}_2. \quad (4.2.5)$$

Quindi, in seguito alla Nota 4.2.2, nelle basi $\mathcal{E}$ e $\mathcal{V}$ la matrice di T , $A_{\mathcal{V},\mathcal{E}}$, è

$$\begin{pmatrix} 1 & \frac{3}{2} \\ 0 & -1 \end{pmatrix}.$$

Ma la domanda che ci siamo posti è di trovare la matrice $A_{\mathcal{V}} = A_{\mathcal{V},\mathcal{V}}$ di T nella base $\mathcal{V}$. A questo scopo dobbiamo esprimere i vettori $\mathbf{e}_1$ ed $\mathbf{e}_2$ nel primo membro di (4.2.4) in termini della base $\mathbf{v}_1, \mathbf{v}_2$, facendo uso della trasformazione inversa del cambio di base calcolata in (4.2.2). Si ottiene:

$$\begin{aligned} \frac{1}{2}T(\mathbf{v}_1) + \frac{1}{2}T(\mathbf{v}_2) &= \mathbf{v}_1 \\ \frac{1}{2}T(\mathbf{v}_1) - \frac{1}{2}T(\mathbf{v}_2) &= \frac{3}{2}\mathbf{v}_1 - \mathbf{v}_2, \end{aligned}$$

da cui, risolvendo,

$$\begin{aligned} T(\mathbf{v}_1) &= \frac{5}{4}\mathbf{v}_1 - \frac{1}{2}\mathbf{v}_2 \\ T(\mathbf{v}_2) &= -\frac{1}{4}\mathbf{v}_1 + \frac{1}{2}\mathbf{v}_2. \end{aligned}$$

Pertanto

$$A_{\mathcal{V}} = \begin{pmatrix} \frac{5}{4} & -\frac{1}{4} \\ -\frac{1}{2} & \frac{1}{2} \end{pmatrix}.$$

□

ESEMPIO 4.2.4. Si consideri l'applicazione lineare T che manda i vettori della base canonica di $\mathbb{R}^3$, $\mathbf{e}_1$, $\mathbf{e}_2$, $\mathbf{e}_3$, rispettivamente nei vettori $\mathbf{v}_1 := \mathbf{e}_1 + 2\mathbf{e}_2$, $\mathbf{v}_2 = \mathbf{e}_1$, $\mathbf{v}_3 = \mathbf{e}_3$. Si calcoli la matrice $A_{\mathcal{V}}$ di T nella base $\mathcal{V} = \{\mathbf{v}_1, \mathbf{v}_2, \mathbf{v}_3\}$.

SVOLGIMENTO. Esprimiamo l'immagine sotto T dei vettori $\mathbf{v}_i$ in termini ancora dei vettori $\mathbf{v}_i$:

$$T(\mathbf{v}_1) = T(\mathbf{e}_1) + 2T(\mathbf{e}_2) = 2\mathbf{e}_1 + 2\mathbf{e}_2 = \mathbf{v}_1 = \mathbf{v}_2$$

$$T(\mathbf{v}_2) = T(\mathbf{e}_1) = \mathbf{v}_1$$

$$T(\mathbf{v}_3) = T(\mathbf{e}_3) = \mathbf{v}_3.$$

Si noti che nell'ultimo passaggio abbiamo espresso i vettori $\mathbf{e}_i$ in termini dei $\mathbf{v}_i$, quindi anche in questo caso abbiamo invertito la matrice del cambiamento di base che manda gli $\mathbf{e}_i$ nei $\mathbf{v}_i$, $i = 1, 2, 3$.) Pertanto la matrice che cerchiamo è

$$A_{\mathcal{V}} = \begin{pmatrix} 1 & 1 & 0 \\ 1 & 0 & 0 \\ 0 & 0 & 1 \end{pmatrix}.$$

□

Il ruolo della matrice inversa del cambiamento di base in questi due esempi lascia intravedere la forma generale della soluzione del problema del cambiamento della rappresentazione matriciale di un operatore lineare al cambiare delle basi:

TEOREMA 4.2.5. *Siano $\mathbf{X}, \mathbf{Y}$ spazi vettoriali di dimensione rispettivamente n e m , e $T : \mathbf{X} \rightarrow \mathbf{Y}$ una applicazione lineare. Indichiamo con $\mathcal{V} = \{\mathbf{v}_1, \dots, \mathbf{v}_n\}$ e $\mathcal{V}' = \{\mathbf{v}'_1, \dots, \mathbf{v}'_n\}$ due basi di $\mathbf{X}$ e con $\mathcal{F} = \{\mathbf{f}_1, \dots, \mathbf{f}_m\}$ e $\mathcal{F}' = \{\mathbf{f}'_1, \dots, \mathbf{f}'_m\}$ due basi di $\mathbf{Y}$.*

Siano $C_V = C_{\mathcal{V}', \mathcal{V}}$ la matrice (4.1.2) del cambiamento di base dalla base $\mathcal{V}$ alla base $\mathcal{V}'$, e $C_F = C_{\mathcal{F}', \mathcal{F}}$ la matrice (4.1.2) del cambiamento

di base dalla base $\mathcal{F}$ alla base $\mathcal{F}'$. Allora, con la notazione stabilita nella Definizione 4.2.1, si ha

$$A_{\mathcal{F}', \mathcal{V}'} = C_V^{-1} A_{\mathcal{F}, \mathcal{V}} C_F .$$

DIMOSTRAZIONE. Dalla Definizione 4.2.1 sappiamo che la matrice $A_{\mathcal{F}, \mathcal{V}}$ rappresenta (nelle basi canoniche di $\mathbb{R}^n$ e $\mathbb{R}^m$) la trasformazione lineare $T_{\mathcal{F}, \mathcal{V}} := F_{\mathcal{F}} \circ T \circ F_{\mathcal{V}}^{-1}$. Perciò, per ogni vettore $\mathbf{x}$ in $\mathbf{X}$,

$$\begin{aligned} C_V^{-1} A_{\mathcal{F}, \mathcal{V}} C_F \mathbf{x} &= (F_{\mathcal{F}'} \circ F_{\mathcal{F}}^{-1})^{-1} \circ (F_{\mathcal{F}} \circ T \circ F_{\mathcal{V}}^{-1}) \circ (F_{\mathcal{V}} \circ F_{\mathcal{V}'}^{-1}) \mathbf{x} \\ &= F_{\mathcal{F}'} \circ F_{\mathcal{F}}^{-1} \circ F_{\mathcal{F}} \circ T \circ F_{\mathcal{V}}^{-1} \circ F_{\mathcal{V}} \circ F_{\mathcal{V}'}^{-1} \mathbf{x} \\ &= F_{\mathcal{F}'} \circ T \circ F_{\mathcal{V}'}^{-1} \mathbf{x} = A_{\mathcal{F}', \mathcal{V}'} \mathbf{x} . \end{aligned}$$

□

NOTA 4.2.6. Riassumendo,

$$A_{\mathcal{V}', \mathcal{F}'}(T) = C_{\mathcal{V}', \mathcal{V}} A_{\mathcal{V}, \mathcal{F}}(T) C_{\mathcal{F}, \mathcal{F}'} = C_{\mathcal{V}, \mathcal{V}'}^{-1} A_{\mathcal{V}, \mathcal{F}}(T) C_{\mathcal{F}, \mathcal{F}'} .$$

In particolare, nel caso di $A \in M_{mn}$, una base $\mathcal{F} = \{\mathbf{e}_1, \dots, \mathbf{e}_n\}$ di $\mathbb{R}^n$ ed una base $\mathcal{V} = \{\mathbf{f}_1, \dots, \mathbf{f}_m\}$ di $\mathbb{R}^m$, la matrice dell'applicazione lineare

$$T_A : \mathbb{R}^n \ni \mathbf{x} \longmapsto A \mathbf{x} \in \mathbb{R}^m$$

rispetto alle basi $\mathcal{F}$ ed $\mathcal{V}$ è uguale ad

$$\begin{aligned} A_{\mathcal{V}, \mathcal{F}}(T_A) &= C_{\mathcal{N}_m, \mathcal{V}}^{-1} A_{\mathcal{N}_m, \mathcal{N}_n}(T_A) C_{\mathcal{N}_n, \mathcal{F}} = C_{\mathcal{N}_m, \mathcal{V}}^{-1} A C_{\mathcal{N}_n, \mathcal{F}} \\ &= \begin{pmatrix} | & & | \\ \mathbf{f}_1 & \dots & \mathbf{f}_m \\ | & & | \end{pmatrix}^{-1} A \begin{pmatrix} | & & | \\ \mathbf{e}_1 & \dots & \mathbf{e}_n \\ | & & | \end{pmatrix} . \end{aligned}$$

Chiaramente,

$$\begin{aligned} T \text{ iniettiva} &\iff T_{\mathcal{F}, \mathcal{V}} \text{ iniettiva,} \\ T \text{ surgettiva} &\iff T_{\mathcal{F}, \mathcal{V}} \text{ surgettiva,} \end{aligned}$$

e quindi

$$T \text{ bigettiva} \iff T_{\mathcal{F}, \mathcal{V}} \text{ bigettiva} \iff A_{\mathcal{F}, \mathcal{V}} \text{ invertibile.}$$

La matrice della applicazione lineare identicamente zero è la matrice zero rispetto a qualsiasi basi. La matrice della applicazione identica $\mathbb{I}_{\mathbf{X}}$

di uno spazio vettoriale $\mathbf{X}$ rispetto ad una base di $\mathbf{X}$ è chiaramente la matrice unità. Se però $\mathcal{V} = \{\mathbf{v}_1, \dots, \mathbf{v}_n\}$ e $\mathcal{F} = \{\mathbf{f}_1, \dots, \mathbf{f}_n\}$ sono due basi diverse di $\mathbf{X}$, allora la matrice $C_{\mathcal{F}, \mathcal{V}} := A_{\mathcal{F}, \mathcal{V}}(\mathbb{I}_{\mathbf{X}})$ di $\mathbb{I}_{\mathbf{X}}$ rispetto alle basi $\mathcal{V}$ ed $\mathcal{F}$ trasforma le coordinate di un vettore $\mathbf{x} \in \mathbf{X}$ rispetto ad $\mathcal{V}$ nelle coordinate dello stesso $\mathbf{x}$ rispetto a $\mathcal{F}$, e quindi è esattamente la *matrice di passaggio* dalla base $\mathcal{V}$ alla base $\mathcal{F}$ definita nella Notazione 4.1.5. $\square$

ESEMPIO 4.2.7. Sia $\mathcal{P}_2$ lo spazio vettoriale di tutti i polinomi (con coefficienti reali) di grado al più ≤ 2 , cioè tutti i polinomi della forma

$$a_0 + a_1 x + a_2 x^2, \quad a_0, a_1, a_2 \in \mathbb{R}.$$

Chiamiamo *base canonica* di $\mathcal{P}_2$ la base $1, x, x^2$, e consideriamo l'applicazione lineare $T : \mathcal{P}_2 \rightarrow \mathcal{P}_2$ che assegna ad ogni $P \in \mathcal{P}_2$ la derivata di $(x-1)P(x)$. Calcoliamo la matrice associata a T rispetto alla base naturale di $\mathcal{P}_2$.

SVOLGIMENTO. Poiché

$$\begin{aligned} T(1) &= (x-1)' = 1 = 1 + 0x + 0x^2, \\ T(x) &= (x^2 - x)' = 2x - 1 = -1 + 2x + 0x^2, \\ T(x^2) &= (x^3 - x^2)' = 3x^2 - 2x = 0 - 2x + 3x^2, \end{aligned}$$

la matrice A associata a T rispetto alla base naturale di $\mathcal{P}_2$ ha le colonne

$$\begin{pmatrix} 1 \\ 0 \\ 0 \end{pmatrix}, \begin{pmatrix} -1 \\ 2 \\ 0 \end{pmatrix}, \begin{pmatrix} 0 \\ -2 \\ 3 \end{pmatrix}$$

e perciò

$$A = \begin{pmatrix} 1 & -1 & 0 \\ 0 & 2 & -2 \\ 0 & 0 & 3 \end{pmatrix}.$$

Di conseguenza, la derivata del polinomio $(x-1)(a_0 + a_1 x + a_2 x^2)$ è uguale al polinomio $b_0 + b_1 x + b_2 x^2$, ove

$$\begin{pmatrix} b_0 \\ b_1 \\ b_2 \end{pmatrix} = \begin{pmatrix} 1 & -1 & 0 \\ 0 & 2 & -2 \\ 0 & 0 & 3 \end{pmatrix} \begin{pmatrix} a_0 \\ a_1 \\ a_2 \end{pmatrix}.$$

La matrice A è invertibile, perciò T è bigettiva e perciò ad ogni $Q \in \mathcal{P}_2$ corrisponde un unico $P \in \mathcal{P}_2$ tale che Q sia la derivata di $(x-1)P(x)$. $\square$

I seguenti fatti sono ormai ovvi:

PROPOSIZIONE 4.2.8. *Siano $\mathbf{X}, \mathbf{Y}, \mathbf{Z}$ spazi vettoriali,*

$$T : \mathbf{X} \longrightarrow \mathbf{Y}, S : \mathbf{Y} \longrightarrow \mathbf{Z} \text{ applicazioni lineari,}$$

ed

$$\mathcal{V} = \{\mathbf{v}_1, \dots, \mathbf{v}_n\}, \mathcal{F} = \{\mathbf{f}_1, \dots, \mathbf{f}_m\}, \mathcal{G} = \{\mathbf{g}_1, \dots, \mathbf{g}_p\}$$

basi di $\mathbf{X}, \mathbf{Y}, \mathbf{Z}$, rispettivamente.

Allora la matrice associata alla composizione $S \circ T$ rispetto ad $\mathcal{V}$ e $\mathcal{G}$ è il prodotto $A_{\mathcal{G}, \mathcal{F}}(S) A_{\mathcal{F}, \mathcal{V}}(T) \in M_{np}$.

Di conseguenza, se T è bigettiva, allora il prodotto

$$A_{\mathcal{V}, \mathcal{F}}(T^{-1}) A_{\mathcal{F}, \mathcal{V}}(T) = A_{\mathcal{V}, \mathcal{V}}(T^{-1} \circ T) = A_{\mathcal{V}, \mathcal{V}}(\mathbb{I}_{\mathbf{X}})$$

è uguale alla matrice unità. Risulta che $A_{\mathcal{V}, \mathcal{F}}(T^{-1}) = A_{\mathcal{F}, \mathcal{V}}(T)^{-1}$.

In particolare, la matrice di passaggio da una base di $\mathbf{X}$ ad un'altra base è l'inversa della matrice di passaggio dalla seconda base alla prima.

ESEMPIO 4.2.9. Consideriamo di nuovo lo spazio vettoriale $\mathcal{P}_2$ di tutti i polinomi di grado al più ≤ 2 . Troviamo le matrici di passaggio:

- dalla base naturale di $\mathcal{P}_2$, cioè $1, x, x^2$, alla base $1, x-1, (x-1)^2$,
- dalla base $1, x-1, (x-1)^2$ alla base naturale di $\mathcal{P}_2$.

Dopo di che, troviamo la matrice associata all'applicazione lineare $T : \mathcal{P}_2 \longrightarrow \mathcal{P}_2$, che assegna ad ogni $P \in \mathcal{P}_2$ la derivata di $(x-1)P(x)$ ed è stata considerata qui sopra, rispetto alla base $1, x-1, (x-1)^2$. Che connessione esiste fra la matrice trovata e quella calcolata nell'esempio precedente?

SVOLGIMENTO. Calcoliamo prima la matrice di passaggio dalla base $1, x-1, (x-1)^2$ alla base naturale di $\mathcal{P}_2$: poiché

$$\begin{aligned} 1 &= 1 + 0x + 0x^2, \\ x-1 &= -1 + 1x + 0x^2, \end{aligned}$$

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

essa è

$$C = \begin{pmatrix} 1 & -1 & 1 \\ 0 & 1 & -2 \\ 0 & 0 & 1 \end{pmatrix}.$$

Troviamo la matrice di passaggio nel senso inverso invertendo la matrice C , col metodo di eliminazione di Gauss:

$$\begin{array}{ccc|ccc} 1 & -1 & 1 & 1 & 0 & 0 & 1 & -1 & 0 \\ 0 & 1 & -2 & 0 & 1 & 0 & 0 & 1 & 0 \\ 0 & 0 & 1 & 0 & 0 & 1 & 0 & 0 & 1 \end{array},$$

$$\begin{array}{ccc|ccc} 1 & 0 & 0 & 1 & 1 & 1 \\ 0 & 1 & 0 & 0 & 1 & 2 \\ 0 & 0 & 1 & 0 & 0 & 1 \end{array}$$

$$C^{-1} = \begin{pmatrix} 1 & 1 & 1 \\ 0 & 1 & 2 \\ 0 & 0 & 1 \end{pmatrix};$$

ma avremmo potuto trovare il risultato anche direttamente: poiché

$$\begin{aligned} 1 &= 1 + 0(x-1) + 0(x-1)^2, \\ x &= 1 + 1(x-1) + 0(x-1)^2, \\ x^2 &= 1 + 2(x-1) + 1(x-1)^2, \end{aligned}$$

la matrice di passaggio dalla base naturale di $\mathcal{P}_2$ alla base $1, x-1, (x-1)^2$ è quella con le colonne

$$\begin{pmatrix} 1 \\ 0 \\ 0 \end{pmatrix}, \begin{pmatrix} -1 \\ 1 \\ 0 \end{pmatrix}, \begin{pmatrix} 1 \\ -2 \\ 1 \end{pmatrix}.$$

Ora possiamo trovare la matrice della trasformazione T , sia in maniera diretta analoga a quanto fatto nell'Esempio 4.2.7, sia applicando il Teorema 4.2.5. I due metodi sono equivalenti: illustriamo il secondo. Poiché

$$\begin{aligned} T(1) &= (x-1)' = 1, \\ T(x-1) &= ((x-1)^2)' = 2(x-1), \\ T((x-1)^2) &= ((x-1)^3)' = 3(x-1)^2, \end{aligned}$$

la matrice di T rispetto alla base $1, x-1, (x-1)^2$ è la matrice diagonale

$$\begin{pmatrix} 1 & 0 & 0 \\ 0 & 2 & 0 \\ 0 & 0 & 3 \end{pmatrix}.$$

Abbiamo calcolato le matrici di passaggio C e C^{-1} , e grazie al Teorema 4.2.5 ne ricaviamo di nuovo la matrice di T rispetto alla base naturale di $\mathcal{P}_2$, già calcolata nell'esempio precedente:

$$\begin{aligned} C \begin{pmatrix} 1 & 0 & 0 \\ 0 & 2 & 0 \\ 0 & 0 & 3 \end{pmatrix} C^{-1} &= \begin{pmatrix} 1 & -1 & 1 \\ 0 & 1 & -2 \\ 0 & 0 & 1 \end{pmatrix} \begin{pmatrix} 1 & 0 & 0 \\ 0 & 2 & 0 \\ 0 & 0 & 3 \end{pmatrix} \begin{pmatrix} 1 & 1 & 1 \\ 0 & 1 & 2 \\ 0 & 0 & 1 \end{pmatrix} \\ &= \begin{pmatrix} 1 & -1 & 0 \\ 0 & 2 & -2 \\ 0 & 0 & 3 \end{pmatrix}. \quad \square \end{aligned}$$

4.3. Cenni introduttivi sulla diagonalizzazione

Anticipiamo qui, come breve cenno, un argomento che sarà trattato in maggior dettaglio in seguito nel Capitolo 7. Siano $\mathbf{X}, \mathbf{Y}$ spazi vettoriali della stessa dimensione n , $T : \mathbf{X} \rightarrow \mathbf{Y}$ una applicazione lineare, e $\mathcal{V} = \{\mathbf{v}_1, \dots, \mathbf{v}_n\}$, $\mathcal{F} = \{\mathbf{f}_1, \dots, \mathbf{f}_n\}$ basi di $\mathbf{X}$ e $\mathbf{Y}$, rispettivamente. Allora la matrice associata a T rispetto alle basi $\mathcal{V}$ e $\mathcal{F}$ è *diagonale*, cioè ha tutti gli elementi fuori della diagonale uguali a zero, se e solo se

$$T(\mathbf{v}_k) \text{ è un multiplo scalare } \lambda_k \mathbf{f}_k \text{ di } \mathbf{f}_k, \quad 1 \leq k \leq n. \quad (4.3.1)$$

In tal caso si ha:

$$A_{\mathcal{F}, \mathcal{V}} = \begin{pmatrix} \lambda_1 & 0 & \dots & 0 \\ 0 & \lambda_2 & \dots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \dots & 0 \end{pmatrix}.$$

Se $\mathbf{X} = \mathbf{Y}$ e $\mathcal{V} = \mathcal{F}$, allora la condizione (4.3.1) si scrive :

$$T(\mathbf{v}_k) = \lambda_k \mathbf{v}_k \text{ per un opportuno } \lambda_k \in \mathbb{R}, \quad 1 \leq k \leq n.$$

DEFINIZIONE 4.3.1. In generale, se T è una applicazione lineare di uno spazio vettoriale $\mathbf{X}$ in se stesso, un scalare λ si chiama *autovalore* di T se esiste $\mathbf{0}_{\mathbf{X}} \neq \mathbf{x} \in \mathbf{X}$ tale che

$$T(\mathbf{x}) = \lambda \mathbf{x}. \quad (4.3.2)$$

Tutti i vettori *non nulli* $\mathbf{x} \in \mathbf{X}$, che soddisfano (4.3.2), si chiamano *autovettori* di T corrispondenti all'autovalore λ ed il sottospazio lineare

$$\text{Ker}(T - \lambda \mathbb{I}_{\mathbf{X}}) = \{\mathbf{x} \in \mathbf{X}; T(\mathbf{x}) = \lambda \mathbf{x}\}$$

di $\mathbf{X}$ si chiama *l'autospazio* di T corrispondente a λ .

Se A è una matrice quadrata d'ordine n , allora gli autovalori, autovettori ed autospazi di $T_A : \mathbb{R}^n \ni \mathbf{x} \mapsto A\mathbf{x} \in \mathbb{R}^n$ si chiamano rispettivamente autovalori, autovettori ed autospazi della matrice A .

4.4. Esercizi sulla diagonalizzazione

ESERCIZIO 4.4.1. Sia A la matrice 3×3 seguente:

$$A = \begin{pmatrix} 1 & 1 & 1 \\ 1 & 1 & 1 \\ 1 & 1 & 1 \end{pmatrix}.$$

Con un calcolo diretto analogo a quello dell'Esempio 4.2.3, si calcoli come cambia la matrice dell'applicazione lineare indotta da A (nella base canonica), cioè

$$\mathbb{R}^3 \ni \begin{pmatrix} x_1 \\ x_2 \\ x_3 \end{pmatrix} \mapsto A \begin{pmatrix} x_1 \\ x_2 \\ x_3 \end{pmatrix} \in \mathbb{R}^3$$

quando si passa alla base

$$\mathbf{v}_1 = \begin{pmatrix} 1 \\ 1 \\ 1 \end{pmatrix}, \quad \mathbf{v}_2 = \begin{pmatrix} 1 \\ -1 \\ 0 \end{pmatrix}, \quad \mathbf{v}_3 = \begin{pmatrix} 0 \\ -1 \\ 1 \end{pmatrix}$$

di $\mathbb{R}^3$.

SOLUZIONE. Poiché

$$\begin{aligned}
 T(\mathbf{v}_1) &= \begin{pmatrix} 3 \\ 3 \\ 3 \end{pmatrix} = 3\mathbf{v}_1 + 0\mathbf{v}_2 + 0\mathbf{v}_3, \\
 T(\mathbf{v}_2) &= \mathbf{0}_3 = 0\mathbf{v}_1 + 0\mathbf{v}_2 + 0\mathbf{v}_3, \\
 T(\mathbf{v}_3) &= \mathbf{0}_3 = 0\mathbf{v}_1 + 0\mathbf{v}_2 + 0\mathbf{v}_3,
 \end{aligned}$$

la matrice desiderata è

$$\begin{pmatrix} 3 & 0 & 0 \\ 0 & 0 & 0 \\ 0 & 0 & 0 \end{pmatrix}.$$

Osserviamo che la matrice così ottenuta è diagonale. Quindi il cambiamento di base ha diagonalizzato la rappresentazione matriciale della applicazione lineare data: cioè, la nuova base è una base di autovettori. Daremo maggiori dettagli nel Capitolo 7. $\square$

ESERCIZIO 4.4.2. Si calcoli la matrice di passaggio dalla base naturale di $\mathbb{R}^3$ alla base

$$\begin{pmatrix} 1 \\ 1 \\ 1 \end{pmatrix}, \begin{pmatrix} 1 \\ -1 \\ 0 \end{pmatrix}, \begin{pmatrix} 0 \\ -1 \\ 1 \end{pmatrix}$$

SOLUZIONE. La matrice richiesta è

$$B = \begin{pmatrix} 1 & 1 & 0 \\ 1 & -1 & -1 \\ 1 & 0 & 1 \end{pmatrix}^{-1}.$$

La calcoliamo:

$$\begin{array}{ccc|ccc|ccc}
 1 & 1 & 0 & 1 & 0 & 0 & 1 & 1 & 0 & 1 & 0 & 0 \\
 1 & -1 & -1 & 0 & 1 & 0 & 0 & -2 & -1 & -1 & 1 & 0 \\
 1 & 0 & 1 & 0 & 0 & 1 & 0 & -1 & 1 & -1 & 0 & 1 \\
 \\
 1 & 1 & 0 & 1 & 0 & 0 & 1 & 1 & 0 & 1 & 0 & 0 \\
 0 & 1 & -1 & 1 & 0 & -1 & 0 & 1 & -1 & 1 & 0 & -1 \\
 0 & -2 & -1 & -1 & 1 & 0 & 0 & 0 & -3 & 1 & 1 & -2 \\
 \\
 1 & 1 & 0 & 1 & 0 & 0 & 1 & 1 & 0 & 1 & 0 & 0 \\
 0 & 1 & -1 & 1 & 0 & -1 & 0 & 1 & 0 & 2/3 & -1/3 & -1/3 \\
 0 & 0 & 1 & -1/3 & -1/3 & 2/3 & 0 & 0 & 1 & -1/3 & -1/3 & 2/3
 \end{array}$$

$$\begin{array}{ccc|ccc} 1 & 0 & 0 & 1/3 & 1/3 & 1/3 \\ 0 & 1 & 0 & 2/3 & -1/3 & -1/3 \\ 0 & 0 & 1 & -1/3 & -1/3 & 2/3 \end{array} .$$

Concludiamo che

$$B = \begin{pmatrix} 1/3 & 1/3 & 1/3 \\ 2/3 & -1/3 & -1/3 \\ -1/3 & -1/3 & 2/3 \end{pmatrix} .$$

□

ESERCIZIO 4.4.3. Si risolva l'Esercizio 4.4.1 non con il calcolo diretto, bensì applicando il Teorema di cambiamento di base 4.2.5.

SOLUZIONE. Per prima cosa scriviamo la matrice $C = C_{\mathcal{V}, \mathcal{E}}$ di passaggio dalla base canonica $\mathcal{E}$ alla base $\mathcal{V}$. Chiaramente C è la matrice che ha per colonne i vettori di $\mathcal{V}$:

$$C = \begin{pmatrix} 1 & 1 & 0 \\ 1 & -1 & -1 \\ 1 & 0 & 1 \end{pmatrix} .$$

L'inversa $B := C^{-1} = C_{\mathcal{E}, \mathcal{V}}$ è stata calcolata nel precedente Esercizio 4.4.2: essa vale

$$B = \begin{pmatrix} 1/3 & 1/3 & 1/3 \\ 2/3 & -1/3 & -1/3 \\ -1/3 & -1/3 & 2/3 \end{pmatrix} .$$

Ora, per il Teorema 4.2.5, la matrice $A_{\mathcal{V}}$ associata alla trasformazione lineare T nella base $\mathcal{V}$ è

$$A_{\mathcal{V}} = B^{-1}AB ,$$

dove, come già nell'Esercizio 4.4.2, abbiamo indicato con

$$A = \begin{pmatrix} 1 & 1 & 1 \\ 1 & 1 & 1 \\ 1 & 1 & 1 \end{pmatrix}$$

la matrice che rappresenta l'operatore nella base canonica. Svolgendo le moltiplicazioni fra matrici si trova che

$$A_{\mathcal{V}} = B^{-1}AB = \begin{pmatrix} 3 & 0 & 0 \\ 0 & 0 & 0 \\ 0 & 0 & 0 \end{pmatrix}. \quad (4.4.1)$$

Si osservi che è opportuno verificare l'identità (4.4.1), ma farlo è equivalente a verificare che $A = BA_{\mathcal{V}}B^{-1}$, una identità che richiede calcoli molto più semplici perché la matrice $A_{\mathcal{V}} = \begin{pmatrix} 3 & 0 & 0 \\ 0 & 0 & 0 \\ 0 & 0 & 0 \end{pmatrix}$ ha tutti i coefficienti nulli tranne il primo. L'elementare verifica viene lasciata al lettore.

Osserviamo che *il cambiamento di basi, e la corrispondente operazione matriciale $A \longrightarrow B^{-1}AB$, ha diagonalizzato la matrice A di partenza*. Si veda il Capitolo 7 nel seguito. $\square$

ESERCIZIO 4.4.4. Si trovi la matrice quadrata A d'ordine 3 tale che l'applicazione lineare $\mathbb{R}^3 \ni \mathbf{x} \longmapsto A\mathbf{x} \in \mathbb{R}^3$ trasformi

$$\begin{pmatrix} 1 \\ 1 \\ 1 \end{pmatrix} \mapsto 3 \begin{pmatrix} 1 \\ 1 \\ 1 \end{pmatrix}, \quad \begin{pmatrix} 0 \\ 1 \\ 1 \end{pmatrix} \mapsto \begin{pmatrix} 0 \\ 1 \\ 1 \end{pmatrix}, \quad \begin{pmatrix} 0 \\ -1 \\ 1 \end{pmatrix} \mapsto \frac{1}{2} \begin{pmatrix} 0 \\ -1 \\ 1 \end{pmatrix}.$$

SOLUZIONE. Indichiamo

$$\mathbf{v}_1 := \begin{pmatrix} 1 \\ 1 \\ 1 \end{pmatrix}, \quad \mathbf{v}_2 := \begin{pmatrix} 0 \\ 1 \\ 1 \end{pmatrix}, \quad \mathbf{v}_3 := \begin{pmatrix} 0 \\ -1 \\ 1 \end{pmatrix}.$$

Allora $\mathbf{v}_1, \mathbf{v}_2, \mathbf{v}_3$ è una base di $\mathbb{R}^3$ e l'applicazione lineare T di cui sopra è stata definita tramite la sua matrice

$$\begin{pmatrix} 3 & 0 & 0 \\ 0 & 1 & 0 \\ 0 & 0 & 1/2 \end{pmatrix}$$

rispetto a questa base. Poiché la matrice del passaggio dalla base $\mathbf{v}_1, \mathbf{v}_2, \mathbf{v}_3$ alla base canonica di $\mathbb{R}^3$ è

$$C = \begin{pmatrix} 1 & 0 & 0 \\ 1 & 1 & -1 \\ 1 & 1 & 1 \end{pmatrix},$$

la matrice del passaggio dalla base canonica di $\mathbb{R}^3$ alla base $\mathbf{v}_1, \mathbf{v}_2, \mathbf{v}_3$ sarà C^{-1} . La calcoliamo:

$$\begin{array}{ccc|ccc} 1 & 0 & 0 & 1 & 0 & 0 & 1 & 0 & 0 \\ 1 & 1 & -1 & 0 & 1 & 0 & 0 & 1 & -1 \\ 1 & 1 & 1 & 0 & 0 & 1 & 0 & 1 & -1 \end{array}, \quad \begin{array}{ccc|ccc} 1 & 0 & 0 & 1 & 0 & 0 & 1 & 0 & 0 \\ 0 & 1 & -1 & -1 & 1 & 0 & 0 & 1 & 0 \\ 0 & 0 & 2 & 0 & -1 & 1 & 0 & 0 & 1 \end{array}$$

e perciò

$$C^{-1} = \begin{pmatrix} 1 & 0 & 0 \\ -1 & 1/2 & 1/2 \\ 0 & -1/2 & 1/2 \end{pmatrix},$$

la matrice associata a T rispetto alla base canonica di $\mathbb{R}^3$, cioè la matrice A per cui $T = T_A$, è

$$\begin{aligned} & C \begin{pmatrix} 3 & 0 & 0 \\ 0 & 1 & 0 \\ 0 & 0 & 1/2 \end{pmatrix} C^{-1} \\ &= \begin{pmatrix} 1 & 0 & 0 \\ 1 & 1 & -1 \\ 1 & 1 & 1 \end{pmatrix} \begin{pmatrix} 3 & 0 & 0 \\ 0 & 1 & 0 \\ 0 & 0 & 1/2 \end{pmatrix} \begin{pmatrix} 1 & 0 & 0 \\ -1 & 1/2 & 1/2 \\ 0 & -1/2 & 1/2 \end{pmatrix} \\ &= \begin{pmatrix} 3 & 0 & 0 \\ 2 & 3/4 & 1/4 \\ 2 & 1/4 & 3/4 \end{pmatrix}. \end{aligned}$$

□

ESERCIZIO 4.4.5. Si scriva la matrice dell'applicazione lineare

$$\mathbb{R}^2 \ni \begin{pmatrix} x_1 \\ x_2 \end{pmatrix} \mapsto \begin{pmatrix} 2 & 1 \\ -1 & 3 \\ 1 & 0 \end{pmatrix} \begin{pmatrix} x_1 \\ x_2 \end{pmatrix} \in \mathbb{R}^3$$

rispetto alle basi

$$\begin{pmatrix} 1 \\ 1 \end{pmatrix}, \begin{pmatrix} 1 \\ -1 \end{pmatrix} \text{ e } \begin{pmatrix} 1 \\ 1 \\ 1 \end{pmatrix}, \begin{pmatrix} 0 \\ 1 \\ 1 \end{pmatrix}, \begin{pmatrix} 0 \\ -1 \\ 1 \end{pmatrix}.$$

SOLUZIONE.

$$\begin{pmatrix} 1 & 0 & 0 \\ 1 & 1 & -1 \\ 1 & 1 & 1 \end{pmatrix}^{-1} \begin{pmatrix} 2 & 1 \\ -1 & 3 \\ 1 & 0 \end{pmatrix} \begin{pmatrix} 1 & 1 \\ 1 & -1 \end{pmatrix} = \begin{pmatrix} 3 & 1 \\ -3/2 & -5/2 \\ -1/2 & 5/2 \end{pmatrix}.$$

□

CAPITOLO 5

Determinante di matrici

Consideriamo le matrici *quadrate*, cioè le matrici di ordine $n \times n$, per un certo $n \geq 1$. Ad ogni matrice quadrata si può associare un numero reale che in qualche modo ne *determina* alcune proprietà fondamentali.

DEFINIZIONE 5.0.6. Il ***determinante*** è una funzione che associa ad ogni matrice A di ordine $n \times n$ un numero reale, che indichiamo con $\det(A)$, che soddisfa le seguenti proprietà:

- $\det(I) = 1$;
- se A ha due righe uguali, allora $\det(A) = 0$;
- $\det(A)$ è una funzione lineare sulle righe di A .

L'ultima proprietà significa che se moltiplichiamo una riga di A per uno scalare λ , allora anche il determinante viene moltiplicato per λ . Allo stesso modo se sommiamo ad una riga di A un vettore, allora il determinante della matrice ottenuta è la somma di $\det(A)$ e del determinante della matrice con il vettore al posto della riga di A .

Dalle proprietà della definizione 5.0.6, si possono dimostrare altre proprietà della funzione determinante:

PROPOSIZIONE 5.0.7. *Consideriamo le matrici quadrate di ordine $n \times n$ e la funzione determinante su di esse. Allora:*

- se A ha una riga nulla, allora $\det(A) = 0$;
- scambiando due righe qualunque di A , allora $\det(A)$ cambia di segno;
- se le righe di A sono linearmente dipendenti, allora $\det(A) = 0$.

Si può dimostrare che esiste davvero, ed è unica, la funzione determinante che soddisfa la definizione 5.0.6 e le proprietà indicate.

Nel caso di matrici 2×2 , il determinante è facile da calcolare:

$$\det(A) = \begin{vmatrix} a_{11} & a_{12} \\ a_{21} & a_{22} \end{vmatrix} = a_{11}a_{22} - a_{21}a_{12},$$

dove

$$A = \begin{pmatrix} a_{11} & a_{12} \\ a_{21} & a_{22} \end{pmatrix}.$$

ESEMPIO 5.0.8. Consideriamo la matrice A quadrata di ordine 2×2 :

$$\begin{pmatrix} 2 & 1 \\ -1 & 0 \end{pmatrix}$$

allora il determinante di A è:

$$\det(A) = \begin{vmatrix} 2 & 1 \\ -1 & 0 \end{vmatrix} = 2 \cdot 0 - (-1) \cdot 1 = 1.$$

□

OSSERVAZIONE 5.0.9. Il determinante è definito solo per le matrici *quadrate*, cioè di ordine $n \times n$, per un certo intero $n \geq 1$. Se invece una matrice A non è quadrata, cioè è di ordine $n \times m$ con $n \neq m$, allora il determinante *non è definito*.

Il determinante di una matrice di ordine 3×3

$$A = \begin{pmatrix} a_{11} & a_{12} & a_{13} \\ a_{21} & a_{22} & a_{23} \\ a_{31} & a_{32} & a_{33} \end{pmatrix}$$

si può calcolare con la formula di Sarrus:

$$\begin{aligned} \det(A) &= \\ &= a_{11}a_{22}a_{33} + a_{12}a_{23}a_{31} + a_{13}a_{21}a_{32} \\ &\quad - a_{13}a_{22}a_{31} - a_{11}a_{23}a_{32} - a_{12}a_{21}a_{33}. \end{aligned}$$

ESEMPIO 5.0.10. Consideriamo la matrice

$$A = \begin{pmatrix} 1 & -1 & 0 \\ 2 & 1 & 1 \\ 1 & -2 & 2 \end{pmatrix}$$

Allora il determinante di A è:

$$\det(A) = 2 + (-1) + 0 - 0 - (-2) - (-4) = 7.$$

□

In generale, se consideriamo una matrice quadrata di ordine $n \times n$

$$A = \begin{pmatrix} a_{11} & a_{12} & \cdots & a_{1n} \\ a_{21} & a_{22} & \cdots & a_{2n} \\ \vdots & \vdots & \ddots & \vdots \\ a_{n1} & a_{n2} & \cdots & a_{nn} \end{pmatrix} \quad (5.0.2)$$

allora si può calcolare il determinante sviluppandolo lungo una riga o una colonna con il metodo di Laplace. La dimostrazione è per induzione sull'ordine n delle matrici (questa dimostrazione si basa sull'assioma di induzione, assioma **P5** dei numeri interi, (1.3)). Il risultato è il seguente:

PROPOSIZIONE 5.0.11. *Sia A una matrice quadrata di ordine $n \times n$ come nella formula (5.0.2). Scegliendo di sviluppare la i -esima riga, il determinante di A è:*

$$\det(A) = \sum_{j=1}^n (-1)^{i+j} a_{ij} \det(A_{ij}),$$

dove A_{ij} è la matrice quadrata di ordine $(n-1) \times (n-1)$ ottenuta da A eliminando la i -esima riga e la j -esima colonna.

Nella proposizione 5.0.7 abbiamo visto che se le righe della matrice A sono linearmente dipendenti, allora il determinante di A è nullo. Vale anche il viceversa, come afferma la seguente:

PROPOSIZIONE 5.0.12. *Il determinante $\det(A)$ di una matrice A è zero se e solo se le righe di A sono linearmente indipendenti.*

Si può dimostrare anche che una matrice A è invertibile se e solo se il determinante di A è diverso da zero. In tal caso, l'inversa di A è

$$A^{-1} = \frac{1}{\det(A)} \begin{pmatrix} A_{11} & -A_{21} & \cdots & (-1)^{n-1}A_{n1} \\ -A_{12} & A_{22} & \cdots & (-1)^n A_{n2} \\ \vdots & \vdots & \ddots & \vdots \\ (-1)^{n-1}A_{1n} & (-1)^n A_{2n} & \cdots & A_{nn} \end{pmatrix}$$

dove A_{ij} è il determinante della matrice di ordine $(n-1) \times (n-1)$ ottenuta da A eliminando la i -esima riga e la j -esima colonna.

Un'altra proprietà importante del determinante è data dal seguente:

TEOREMA 5.0.13 (Binet). *Siano A e B sono due matrici quadrate di ordine $n \times n$. Allora*

$$\det(A \cdot B) = \det(A) \cdot \det(B). \quad (5.0.3)$$

ESERCIZIO 5.0.14. Consideriamo due matrici quadrate A e B di ordine 2×2 o 3×3 . Verificate che vale la formula (5.0.3). $\square$

CENNO DI SOLUZIONE. Vediamo esplicitamente il caso delle matrici di ordine 2×2 . Lasciamo al lettore il caso di quelle di ordine 3×3 . Siano A e B matrici di ordine 2×2 :

$$A = \begin{pmatrix} a_{11} & a_{12} \\ a_{21} & a_{22} \end{pmatrix}, \quad B = \begin{pmatrix} b_{11} & b_{12} \\ b_{21} & b_{22} \end{pmatrix}.$$

Allora il prodotto AB è la seguente matrice di ordine 2×2 :

$$AB = \begin{pmatrix} a_{11}b_{11} + a_{12}b_{21} & a_{11}b_{12} + a_{12}b_{22} \\ a_{21}b_{11} + a_{22}b_{21} & a_{21}b_{12} + a_{22}b_{22} \end{pmatrix},$$

che ha determinante

$$\begin{aligned} \det(AB) &= \\ &= (a_{11}b_{11} + a_{12}b_{21})(a_{21}b_{12} + a_{22}b_{22}) - (a_{21}b_{11} + a_{22}b_{21})(a_{11}b_{12} + a_{12}b_{22}) \\ &= a_{12}b_{21}a_{21}b_{12} + a_{11}b_{11}a_{22}b_{22} - a_{21}b_{11}a_{12}b_{22} - a_{22}b_{21}a_{11}b_{12} \\ &= (a_{11}a_{22} - a_{21}a_{12})(b_{11}b_{22} - b_{21}b_{12}) = \det(A) \cdot \det(B), \end{aligned}$$

che è proprio ciò che volevamo dimostrare. $\square$

CAPITOLO 6

Prodotti scalari e ortogonalità

6.1. Prodotto scalare euclideo nel piano ed in $\mathbb{R}^n$

Questa sezione ha carattere introduttivo: tutte le definizioni e proprietà qui illustrate verranno riprese nel seguito in veste più generale e rigorosa.

DEFINIZIONE 6.1.1. Il *prodotto scalare naturale, o euclideo* di due vettori $\mathbf{x}$ e $\mathbf{y}$ in $\mathbb{R}^2$ è la lunghezza della proiezione ortogonale di $\mathbf{y}$ su $\mathbf{x}$:

$$\mathbf{x} \cdot \mathbf{y} = \|\mathbf{x}\| \|\mathbf{y}\| \cos \theta$$

dove $\|\cdot\|$ denota la lunghezza di un vettore e θ è l'angolo formato dai vettori $\mathbf{x}$ e $\mathbf{y}$.

Si noti che il prodotto scalare di $\mathbf{x}$ e $\mathbf{y}$ è un numero reale, cioè appunto uno scalare, il che spiega la terminologia.

Vorremmo trovare un metodo per calcolare il prodotto scalare in termini delle coordinate dei vettori, senza dover usare la trigonometria per ricavare il coseno dell'angolo da essi formato: indubbiamente tale metodo è vantaggioso, anche se a prima vista sembra assoggettare il risultato alla scelta di una base.

Consideriamo allora due vettori $\mathbf{x}$ e $\mathbf{y}$ nel piano $\mathbb{R}^2$, diciamo $\mathbf{x} = (x_1, x_2)$ e $\mathbf{y} = (y_1, y_2)$, dove x_1, x_2 e y_1, y_2 sono le coordinate rispetto alla base canonica $e_1 = (1, 0)$ e $e_2 = (0, 1)$ di $\mathbb{R}^2$.

Siano θ_1 e θ_2 gli angoli formati rispettivamente dai vettori $\mathbf{x}$ e $\mathbf{y}$ con la semiretta positiva delle ascisse. Allora l'angolo formato dai vettori $\mathbf{x}$ e $\mathbf{y}$ è

$$\theta = \theta_2 - \theta_1.$$

Ricordiamo le *formule di addizione e sottrazione* dei coseni:

$$\cos(\alpha + \beta) = \cos \alpha \cos \beta - \sin \alpha \sin \beta$$

che ci serviranno fra poco e osserviamo che per definizione si ha che:

$$\cos \theta_1 = \frac{x_1}{\|\mathbf{x}\|}, \quad \sin \theta_1 = \frac{x_2}{\|\mathbf{x}\|}, \quad \cos \theta_2 = \frac{y_1}{\|\mathbf{y}\|}, \quad \sin \theta_2 = \frac{y_2}{\|\mathbf{y}\|}.$$

Dalle due equazioni precedenti segue che:

$$\begin{aligned} \cos \theta &= \cos(\theta_2 - \theta_1) = \cos \theta_1 \cos \theta_2 + \sin \theta_1 \sin \theta_2 \\ &= \frac{x_1 y_1}{\|\mathbf{x}\| \|\mathbf{y}\|} + \frac{x_2 y_2}{\|\mathbf{x}\| \|\mathbf{y}\|} = \frac{x_1 y_1 + x_2 y_2}{\|\mathbf{x}\| \|\mathbf{y}\|}. \end{aligned}$$

Ma allora il prodotto scalare di $\mathbf{x}$ e $\mathbf{y}$ è

$$\mathbf{x} \cdot \mathbf{y} = \|\mathbf{x}\| \|\mathbf{y}\| \cos \theta = x_1 y_1 + x_2 y_2. \quad (6.1.1)$$

Scrivendo i vettori $\mathbf{x}$ e $\mathbf{y}$ come colonne:

$$\mathbf{x} = \begin{pmatrix} x_1 \\ x_2 \end{pmatrix}, \quad \mathbf{y} = \begin{pmatrix} y_1 \\ y_2 \end{pmatrix}$$

il prodotto scalare si può scrivere anche così:

$$\mathbf{x} \cdot \mathbf{y} = x_1 y_1 + x_2 y_2 = \begin{pmatrix} x_1 & x_2 \end{pmatrix} \begin{pmatrix} y_1 \\ y_2 \end{pmatrix} = \mathbf{x}^t \mathbf{y}.$$

OSSERVAZIONE 6.1.2. Il prodotto scalare è commutativo: scambiando di posto i due vettori, il prodotto scalare non cambia:

$$\mathbf{x} \cdot \mathbf{y} = \mathbf{y} \cdot \mathbf{x},$$

come è evidente dalla formula (6.1.1) e dalla definizione intuitiva data all'inizio.

Un prodotto scalare naturale si può definire anche in $\mathbb{R}^n$ analogamente al caso $\mathbb{R}^2$.

DEFINIZIONE 6.1.3. Dati due vettori di $\mathbb{R}^n$

$$\mathbf{x} = \begin{pmatrix} x_1 \\ x_2 \\ \vdots \\ x_n \end{pmatrix}, \quad \mathbf{y} = \begin{pmatrix} y_1 \\ y_2 \\ \vdots \\ y_n \end{pmatrix}$$

definiamo il *prodotto scalare euclideo* di $\mathbf{x}$ e $\mathbf{y}$ come

$$\mathbf{x} \cdot \mathbf{y} = \mathbf{x}^t \mathbf{y} = \begin{pmatrix} x_1 & x_2 & \dots & x_n \end{pmatrix} \begin{pmatrix} y_1 \\ y_2 \\ \vdots \\ y_n \end{pmatrix} = x_1 y_1 + x_2 y_2 + \dots + x_n y_n.$$

ESEMPIO 6.1.4. Consideriamo i seguenti due vettori in $\mathbb{R}^3$:

$$\begin{pmatrix} 2 \\ 1 \\ -1 \end{pmatrix}, \quad \begin{pmatrix} 1 \\ 0 \\ 1 \end{pmatrix}$$

Allora il loro prodotto scalare è

$$2 \cdot 1 + 1 \cdot 0 - 1 \cdot 1 = 1.$$

□

6.2. Spazi vettoriali su $\mathbb{C}$

La definizione di spazio vettoriale sul campo complesso $\mathbb{C}$ differisce da quella degli spazi vettoriali su $\mathbb{R}$, Definizione 2.1.1, solo perché al posto di scalari reali si usano ora scalari complessi.

DEFINIZIONE 6.2.1. Dato uno spazio vettoriale reale $\mathbf{X}_{\mathbb{R}}$, la sua *complessificazione* $\mathbf{X}_{\mathbb{C}}$ è lo spazio vettoriale di tutte le combinazioni lineari a coefficienti complessi dei vettori di una qualsiasi base $\mathbf{x}_1, \dots, \mathbf{x}_n$ di $\mathbf{X}_{\mathbb{R}}$. È chiaro che lo spazio $\mathbf{X}_{\mathbb{C}}$ non dipende dalla scelta della base in $\mathbf{X}_{\mathbb{R}}$ usata per costruirlo, e che quella base è anche una base in $\mathbf{X}_{\mathbb{C}}$: quindi $\mathbf{X}_{\mathbb{C}}$ ha la stessa dimensione su $\mathbb{C}$ che $\mathbf{X}_{\mathbb{R}}$ ha su $\mathbb{R}$ (cautela: poiché il campo complesso $\mathbb{C}$ è uno spazio vettoriale di dimensione due sul campo di scalari $\mathbb{R}$, ogni spazio vettoriale complesso $\mathbf{X}_{\mathbb{C}}$ di dimensione n si può anche considerare come uno spazio vettoriale $\mathbf{X}_{\mathbb{R}}$ su $\mathbb{R}$, ed in tal caso la sua dimensione su $\mathbb{R}$ è $2n$: una base è $\mathbf{x}_1, i\mathbf{x}_1, \mathbf{x}_2, i\mathbf{x}_2, \dots, \mathbf{x}_n, i\mathbf{x}_n$. Se uno spazio vettoriale $\mathbf{X}_{\mathbb{C}}$ su $\mathbb{C}$ è la complessificazione di uno spazio vettoriale $\mathbf{X}_{\mathbb{R}}$ su $\mathbb{R}$, allora $\mathbf{X}_{\mathbb{R}}$ si chiama una *forma reale* di $\mathbf{X}_{\mathbb{C}}$.

6.3. * La definizione generale di prodotto scalare

Il contenuto di questa Sezione viene espanso ed approfondito nell'Appendice 17. Per una prima lettura ci si può limitare a quanto segue.

Estendiamo la nozione di prodotto scalare con le seguenti definizioni.

DEFINIZIONE 6.3.1. (*Forme bilineari e forme hermitiane.*)

- (i) Dato uno spazio vettoriale $\mathbf{X}$ su un campo K (in questo libro $K = \mathbb{R}$ o $\mathbb{C}$, i numeri reali o rispettivamente i numeri complessi), una applicazione $\phi : \mathbf{X} \times \mathbf{X} \longrightarrow K$ si dice una *forma bilineare* se è lineare rispetto ad entrambe le variabili, cioè se

$$(a_1) \quad \phi(\mathbf{x}_1 + \mathbf{x}_2, \mathbf{y}) = \phi(\mathbf{x}_1, \mathbf{y}) + \phi(\mathbf{x}_2, \mathbf{y})$$

$$(a_2) \quad \phi(\mathbf{x}, \mathbf{y}_1 + \mathbf{y}_2) = \phi(\mathbf{x}, \mathbf{y}_1) + \phi(\mathbf{x}, \mathbf{y}_2)$$

$$(b) \quad \phi(\lambda \mathbf{x}, \mathbf{y}) = \phi(\mathbf{x}, \lambda \mathbf{y}) = \lambda \phi(\mathbf{x}, \mathbf{y}) \quad \text{per ogni scalare } \lambda.$$

- (ii) Sia $\mathbf{X}$ uno spazio vettoriale sul campo $\mathbb{C}$. Una applicazione $\phi : \mathbf{X} \times \mathbf{X} \longrightarrow \mathbb{C}$ si dice una *forma hermitiana* se è lineare rispetto alla prima variabile ed *antilineare* rispetto alla seconda, cioè se

$$(a_1) \quad \phi(\mathbf{x}_1 + \mathbf{x}_2, \mathbf{y}) = \phi(\mathbf{x}_1, \mathbf{y}) + \phi(\mathbf{x}_2, \mathbf{y})$$

$$(a_2) \quad \phi(\mathbf{x}, \mathbf{y}_1 + \mathbf{y}_2) = \phi(\mathbf{x}, \mathbf{y}_1) + \phi(\mathbf{x}, \mathbf{y}_2)$$

$$(b_1) \quad \phi(\lambda \mathbf{x}, \mathbf{y}) = \lambda \phi(\mathbf{x}, \mathbf{y}) \quad \text{per ogni scalare } \lambda$$

$$(b_2) \quad \phi(\mathbf{x}, \lambda \mathbf{y}) = \overline{\lambda} \phi(\mathbf{x}, \mathbf{y}) \quad \text{per ogni scalare } \lambda \text{ (qui, come sempre, } \overline{\lambda} \text{ denota il complesso coniugato di } \lambda \text{)}$$

DEFINIZIONE 6.3.2. (*Simmetria di una forma bilineare*)

- (i) Una forma bilineare ϕ sullo spazio vettoriale $\mathbf{X}$ si dice *simmetrica* se

$$\phi(\mathbf{x}, \mathbf{y}) = \phi(\mathbf{y}, \mathbf{x})$$

per ogni $\mathbf{x}, \mathbf{y}$ in $\mathbf{X}$.

- (ii) Se lo spazio $\mathbf{X}$ è complesso, la forma bilineare ϕ si dice *simmetrica coniugata*, o *sesquilineare*, se

$$\phi(\mathbf{x}, \mathbf{y}) = \overline{\phi(\mathbf{y}, \mathbf{x})}$$

per ogni $\mathbf{x}, \mathbf{y}$ in $\mathbf{X}$.

DEFINIZIONE 6.3.3. (*Prodotto scalare.*)

- (i) Un *prodotto scalare sul campo reale* su uno spazio vettoriale reale $\mathbf{X}$ è una forma bilineare $\phi : \mathbf{X} \times \mathbf{X} \longrightarrow \mathbb{R}$ *simmetrica*, nel senso della Definizione 6.3.2 (i). Di solito si scrive $\langle \mathbf{u}, \mathbf{v} \rangle$ invece di $\phi(\mathbf{u}, \mathbf{v})$.
- (ii) Più in generale, un *prodotto scalare sul campo complesso* (detto anche *prodotto hermitiano*) su uno spazio vettoriale complesso $\mathbf{X}$ è una forma bilineare $\phi : \mathbf{X} \times \mathbf{X} \longrightarrow \mathbb{C}$ *simmetrica coniugata*, nel senso della Definizione 6.3.2 (ii).

ESEMPIO 6.3.4. Il prodotto scalare euclideo in $\mathbb{R}^n$, definito nella Sezione 6.1 mediante la formula

$$\mathbf{x} \cdot \mathbf{y} = \sum_{i=1}^n x_i y_i \quad (6.3.1)$$

dove i numeri x_i e y_i sono le coordinate dei vettori $\mathbf{x}$ e $\mathbf{y}$ *rispetto alla base canonica* (oppure rispetto ad una qualsiasi base prefissata) è un esempio di prodotto scalare reale. Si noti che questa definizione dipende dalla scelta di base.

L'esempio corrispondente di prodotto scalare complesso su $\mathbb{C}^n$ è l'analogo prodotto scalare euclideo

$$\mathbf{x} \cdot \mathbf{y} = \sum_{i=1}^n x_i \overline{y_i}, \quad (6.3.2)$$

che dipende anch'esso dalla base scelta. $\square$

ESEMPIO 6.3.5. (***L'integrale come prodotto scalare.***) Sia P_n lo spazio dei polinomi di grado al più n sull'intervallo $[0, 1] \subset \mathbb{R}$, introdotto nella Sezione 3.3. Definiamo un prodotto scalare fra i polinomi $p(x) = \sum_{k=0}^n a_k x^k$ e $q(x) = \sum_{k=0}^n b_k x^k$ come l'integrale

$$\langle p, q \rangle := \int_0^1 p(x) q(x) dx$$

(per la definizione di integrale si veda un libro di testo di Analisi Matematica). Osserviamo che i vettori della base canonica, cioè i monomi $e_i(x) = x^i$, verificano

$$\langle e_i, e_j \rangle = \int_0^1 x^{i+j} dx = \frac{1}{i+j+1}.$$

□

6.4. Matrici complesse, autoaggiunte e simmetriche

Una volta scelte le basi in due spazi vettoriali complessi di dimensioni rispettivamente n e m , ogni applicazione lineare fra i due spazi è rappresentata da una matrice $m \times n$ come nella sezione 3.5, ma questa volta le matrici sono a coefficienti complessi.

NOTAZIONE 6.4.1. *L'insieme delle matrici $m \times n$ a coefficienti complessi è uno spazio vettoriale su $\mathbb{C}$ che si indica con $M_m^{\mathbb{C}}n$. Per chiarezza, quando ci sia adito a dubbio, indicheremo lo spazio delle matrici reali con $M_m^{\mathbb{C}}n$.*

Rammentiamo dal Capitolo 3 la regola con cui si associa una matrice ad una applicazione lineare, riformulandola ora in termini di prodotti scalari:

PROPOSIZIONE 6.4.2. **(Matrice associata ad una applicazione lineare.)**

- (i) Denotiamo con $\{\mathbf{e}_1, \dots, \mathbf{e}_n\}$ la base canonica in $\mathbb{C}^n$. Per ogni applicazione lineare $T : \mathbb{C}^n \rightarrow \mathbb{C}^n$, la matrice $A = A_T$ associata a T nella base canonica ha per coefficienti

$$a_{ij} = T\mathbf{e}_i \cdot \mathbf{e}_j.$$

- (ii) Sia $\mathbf{X}$ uno spazio vettoriale con base $\mathcal{X} = \{\mathbf{x}_1, \dots, \mathbf{x}_n\}$ e sia $T : \mathbf{X} \rightarrow \mathbf{Y}$ una applicazione lineare. La matrice $A = A_{\mathcal{X}, \mathcal{X}}(T)$ associata a T nella base $\mathcal{X}$ ha per coefficienti

$$a_{ij} = T\mathbf{x}_i \cdot \mathbf{x}_j.$$

Da queste espressioni segue immediatamente:

COROLLARIO 6.4.3. *Sia $A \in M_{nn}^{\mathbb{R}}$. Allora, per ogni $\mathbf{x}$ e $\mathbf{y} \in \mathbb{R}^n$,*

$$A\mathbf{x} \cdot \mathbf{y} = \mathbf{x} \cdot A^{\top}\mathbf{y}.$$

DEFINIZIONE 6.4.4. **(Aggiunto.)** Sia $A \in M_{mn}^{\mathbb{C}}$. Si chiama *matrice aggiunta* (o semplicemente *aggiunto*) di A la matrice $A^* \in M_{nm}^{\mathbb{C}}$ data

da

$$\begin{pmatrix} \alpha_{11} & \alpha_{12} & \dots & \alpha_{1n} \\ \alpha_{21} & \alpha_{22} & \dots & \alpha_{2n} \\ \vdots & \vdots & \ddots & \vdots \\ \alpha_{m1} & \alpha_{m2} & \dots & \alpha_{mn} \end{pmatrix}^* = \begin{pmatrix} \overline{\alpha_{11}} & \overline{\alpha_{21}} & \dots & \overline{\alpha_{m1}} \\ \overline{\alpha_{12}} & \overline{\alpha_{22}} & \dots & \overline{\alpha_{m2}} \\ \vdots & \vdots & \ddots & \vdots \\ \overline{\alpha_{1n}} & \overline{\alpha_{2n}} & \dots & \overline{\alpha_{mn}} \end{pmatrix}$$

(di nuovo, la barra indica il complesso coniugato). In altre parole, $A^* = \overline{A}^\top = \overline{A}^\top$.

NOTA 6.4.5. Per la matrice A^* valgono regole di calcolo analoghe a quelle della Nota 3.8.3; inoltre, se A è una matrice quadrata, fra aggiunto e prodotto scalare complesso c'è lo stesso legame che abbiamo già illustrato nel Corollario 6.4.3 fra trasposta e prodotto scalare reale:

$$\langle A\mathbf{x}, \mathbf{y} \rangle = \langle \mathbf{x}, A^*\mathbf{y} \rangle, \quad (6.4.1)$$

ovvero, con la terminologia alternativa introdotta in (6.3.2) per il prodotto scalare euclideo,

$$A\mathbf{x} \cdot \mathbf{y} = \mathbf{x} \cdot A^*\mathbf{y}. \quad (6.4.2)$$

Per linearità basta provare (6.4.2) per i vettori della base canonica: in tal caso, siccome $A\mathbf{e}_i \cdot \mathbf{e}_j = \alpha_{ij}$, (6.4.2) è precisamente la definizione di aggiunto 6.4.4. $\square$

DEFINIZIONE 6.4.6. (**Matrici simmetriche e matrici autoaggiunte.**)

Una matrice $A \in M_{nn}^{\mathbb{R}}$ si dice *simmetrica* se $A = A^\top$. Più in generale, una matrice $A \in M_{nn}^{\mathbb{C}}$ si dice *autoaggiunta* se $A = A^*$. Analogamente, se una applicazione lineare verifica $\langle A\mathbf{x}, \mathbf{x} \rangle = \langle \mathbf{x}, A\mathbf{x} \rangle$, A si dice *simmetrica* se agisce su uno spazio vettoriale reale, ed *autoaggiunta* se agisce su uno spazio vettoriale complesso.

6.5. * Norma e prodotti scalari definiti positivi

Il contenuto di questa Sezione viene espanso ed approfondito nell'Appendice 17. Per una prima lettura ci si può limitare a quanto segue.

DEFINIZIONE 6.5.1. (*Norma.*) Una *norma* su uno spazio vettoriale $\mathbf{X}$ è una funzione $N : \mathbf{X} \longrightarrow [0, +\infty)$ con le seguenti proprietà:

- (i) $N(\mathbf{x}) = 0$ se e solo se $\mathbf{x} = \mathbf{0}$,
- (ii) $N(\lambda\mathbf{x}) = |\lambda| N(\mathbf{x})$ per ogni vettore $\mathbf{x}$ e scalare λ ,
- (iii) (**disuguaglianza triangolare.**) $N(\mathbf{x} + \mathbf{y}) \leq N(\mathbf{x}) + N(\mathbf{y})$ per tutti i vettori $\mathbf{x}$ e $\mathbf{y}$.

NOTAZIONE 6.5.2. *Scriviamo d'ora in poi $\|\mathbf{x}\| := N(\mathbf{x})$.*

DEFINIZIONE 6.5.3. (*Identità del parallelogramma.*) Si dice che una norma verifica l'*identità del parallelogramma* se

$$\|\mathbf{x} + \mathbf{y}\| + \|\mathbf{x} - \mathbf{y}\| = 2\|\mathbf{x}\| + 2\|\mathbf{y}\|. \quad (6.5.1)$$

NOTA 6.5.4. (**Norma euclidea.**) La *lunghezza* di un vettore $\mathbf{x} = (x_1, x_2, \dots, x_n)$ in $\mathbb{R}^n$, data da

$$\|\mathbf{x}\| = \sqrt{\sum_{i=1}^n |x_i|^2}$$

è una norma, che si indica con *norma euclidea*; è da questo esempio di norma che proviene il nome della disuguaglianza triangolare (ogni lato di un triangolo ha lunghezza non superiore alla somma degli altri due). Questa norma verifica l'identità del parallelogramma (6.5.1), la quale equivale al ben noto risultato della geometria elementare che la somma delle lunghezze delle diagonali di un parallelogramma è uguale al suo perimetro (in effetti è da questo fatto che viene il nome dell'identità).

Analogamente, la lunghezza di un vettore nello spazio vettoriale $\mathbb{C}^n$ sul campo complesso è la norma definita esattamente come sopra ma interpretando il modulo nel senso del modulo dei numeri complessi.

□

DEFINIZIONE 6.5.5. (**Prodotto scalare definito positivo.**) Si dice che un prodotto scalare è *semidefinito positivo* se $\langle \mathbf{x}, \mathbf{x} \rangle \geq 0$. Si dice che un prodotto scalare è *definito positivo* se $\langle \mathbf{x}, \mathbf{x} \rangle > 0$ per ogni vettore $\mathbf{x} \neq \mathbf{0}$ (e naturalmente $\langle \mathbf{0}, \mathbf{0} \rangle = 0$).

COROLLARIO 6.5.6. (**Norma indotta da un prodotto scalare definito positivo.**) *Se un prodotto scalare sullo spazio vettoriale $\mathbf{X}$ è definito positivo, l'espressione $\sqrt{\langle \mathbf{x}, \mathbf{x} \rangle}$ è una norma su X che verifica l'identità del parallelogramma (6.5.1).*

DIMOSTRAZIONE. In virtù delle Definizioni 6.5.5 e 6.5.1 tutto è evidente eccetto forse l'identità del parallelogramma, che si ottiene sommando le identità seguenti:

$$\begin{aligned}\langle \mathbf{x} + \mathbf{y}, \mathbf{x} + \mathbf{y} \rangle &= \langle \mathbf{x}, \mathbf{x} \rangle + \langle \mathbf{x}, \mathbf{y} \rangle + \langle \mathbf{y}, \mathbf{x} \rangle + \langle \mathbf{y}, \mathbf{y} \rangle \\ \langle \mathbf{x} - \mathbf{y}, \mathbf{x} - \mathbf{y} \rangle &= \langle \mathbf{x}, \mathbf{x} \rangle - \langle \mathbf{x}, \mathbf{y} \rangle - \langle \mathbf{y}, \mathbf{x} \rangle + \langle \mathbf{y}, \mathbf{y} \rangle .\end{aligned}$$

NOTA 6.5.7. Vedremo in seguito (Corollario 6.6.2) che è vero anche il viceversa: una norma che verifica l'identità del parallelogramma identifica un prodotto scalare definito positivo. $\square$

ESEMPIO 6.5.8. Alla luce della Nota 6.5.4, la norma euclidea in $\mathbb{R}^n$ (rispettivamente $\mathbb{C}^n$) è indotta dal rispettivo prodotto scalare euclideo. Quindi il prodotto scalare euclideo è definito positivo. $\square$

ESEMPIO 6.5.9. Il prodotto scalare sullo spazio dei polinomi, introdotto nell'Esempio 6.3.5, è definito positivo. $\square$

DEFINIZIONE 6.5.10. (**Prodotto scalare non degenere.**) Si dice che un prodotto scalare è *non degenere* se per ciascun vettore $\mathbf{x}$ la condizione $\langle \mathbf{x}, \mathbf{y} \rangle = 0$ per ogni $\mathbf{y}$ implica $\mathbf{x} = \mathbf{0}$, cioè se non esistono vettori non nulli ortogonali a tutti i vettori.

Da queste definizioni risulta ovvio il Corollario seguente:

COROLLARIO 6.5.11. *Ogni prodotto scalare definito positivo è non degenere.*

6.6. Ortogonalità

Il contenuto di questa Sezione viene espanso ed approfondito nell'Appendice 17. Per una prima lettura ci si può limitare a quanto segue.

LEMMA 6.6.1. (**Polarizzazione di forme bilineari.**)

- (i) Sia $\mathbf{X}$ uno spazio vettoriale sul campo reale e $\phi : \mathbf{X} \times \mathbf{X} \longrightarrow \mathbb{R}$ una forma bilineare simmetrica. Allora $\phi(\mathbf{x}, \mathbf{x}) = 0$ per ogni vettore $\mathbf{x}$ se e solo se $\phi(\mathbf{x}, \mathbf{y}) = 0$ per ogni coppia di vettori $\mathbf{x}$ e $\mathbf{y}$.
- (ii) Più in generale, lo stesso risultato vale se $\mathbf{X}$ è uno spazio vettoriale sui complessi e $\phi : \mathbf{X} \times \mathbf{X} \longrightarrow \mathbb{C}$ una forma hermitiana simmetrica coniugata.

DIMOSTRAZIONE. L'implicazione inversa è ovvia. Per l'implicazione diretta, consideriamo dapprima il caso reale (parte (i)). Per ogni $\mathbf{x}, \mathbf{y}$ si formi il vettore $\mathbf{v}_+ = \mathbf{x} + \mathbf{y}$. Per la proprietà di simmetria della forma bilineare (Definizione 6.3.2 (i)) si ha

$$0 = \phi(\mathbf{v}_+, \mathbf{v}_+) = \phi(\mathbf{x}, \mathbf{x}) + \phi(\mathbf{x}, \mathbf{y}) + \phi(\mathbf{y}, \mathbf{x}) + \phi(\mathbf{y}, \mathbf{y}) = 2\phi(\mathbf{x}, \mathbf{y}).$$

Questo prova (i).

La dimostrazione di (ii) (il caso complesso) è simile. Si considerino il vettore $\mathbf{v}_1$ di prima ed il vettore $\mathbf{v}_2 = \mathbf{x} + i\mathbf{y}$. Si usa la sesquilinearità della forma bilineare (Definizione 6.3.2 (ii)). Allora l'argomento di prima, applicato di nuovo al vettore $\mathbf{v}_1$, questa volta dà

$$0 = \phi(\mathbf{v}_1, \mathbf{v}_1) = \phi(\mathbf{x}, \mathbf{y}) + \phi(\mathbf{y}, \mathbf{x}) = 2 \operatorname{Re}(\phi(\mathbf{x}, \mathbf{y})),$$

ed applicato al vettore $\mathbf{v}_2$ dà

$$0 = \phi(\mathbf{v}_2, \mathbf{v}_2) = \phi(\mathbf{x}, i\mathbf{y}) + \phi(i\mathbf{y}, \mathbf{x}) = 2 \operatorname{Re}(-i\phi(\mathbf{x}, \mathbf{y})) = 2 \operatorname{Im}(\phi(\mathbf{x}, \mathbf{y})).$$

Concludiamo che $\phi(\mathbf{x}, \mathbf{y}) = 0$. Questo prova (ii).

COROLLARIO 6.6.2. I valori diagonali $\phi(\mathbf{x}, \mathbf{x})$ di una forma bilineare o hermitiana determinano univocamente la forma. In particolare, ogni norma che soddisfa l'identità del parallelogramma (6.5.1) identifica univocamente il prodotto scalare definito positivo da cui essa è indotta nel senso del Corollario 6.5.6: quindi c'è una corrispondenza biunivoca fra le norme che verificano l'identità del parallelogramma ed i prodotti scalari definiti positivi.

DIMOSTRAZIONE. Se due forme bilineari o hermitiane ϕ e ψ coincidono sulla diagonale, applicando il Lemma 6.6.1 alla forma $\phi - \psi$ vediamo che questa forma è zero ovunque.

DEFINIZIONE 6.6.3. (**Ortogonalità.**) Due vettori x e y di uno spazio vettoriale X si dicono **ortogonali**, o *perpendicolari*, se

$$\langle \mathbf{x}, \mathbf{y} \rangle = 0.$$

DEFINIZIONE 6.6.4. (**Sottospazio ortogonale.**) Scelto un prodotto scalare in $\mathbf{X}$, per ogni sottoinsieme $\mathbf{E} \subset \mathbf{X}$ l'insieme

$$\mathbf{E}^\perp := \{ \mathbf{x} : \langle \mathbf{x}, \mathbf{v} \rangle = 0 \quad \text{per ogni } \mathbf{v} \in \mathbf{E} \}$$

si chiama *ortogonale* di $\mathbf{E}$.

PROPOSIZIONE 6.6.5. (**Complemento ortogonale e proiezione ortogonale.**) In ogni spazio vettoriale $\mathbf{X}$ munito di prodotto scalare si ha:

- (i) L'ortogonale $\mathbf{E}^\perp$ di ogni sottoinsieme (non solo di un sottospazio!) $\mathbf{E} \subset \mathbf{X}$ è un sottospazio vettoriale di $\mathbf{X}$, che d'ora in avanti chiameremo *sottospazio ortogonale* di $\mathbf{E}$.
- (ii) $\mathbf{0}^\perp = \mathbf{X}$.
- (iii) Se il prodotto scalare è non degenere, allora $\mathbf{X}^\perp = \mathbf{0}$.
- (iv) Il prodotto scalare è non degenere se per ogni vettore $\mathbf{v}$ il sottospazio unidimensionale $\mathbf{V}$ che esso genera è strettamente contenuto in $\mathbf{X}$, cioè $\dim \mathbf{V} < \dim \mathbf{X}$. Più in generale, se $\mathbf{V} \neq \mathbf{0}$ è un qualsiasi sottospazio non nullo di $\mathbf{X}$, allora il suo ortogonale $\mathbf{V}^\perp$ è un sottospazio proprio di $\mathbf{X}$.
- (v) Se il prodotto scalare è definito positivo, allora per ogni sottospazio $\mathbf{V}$ si ha $\mathbf{V} \cap \mathbf{V}^\perp = \emptyset$ ed ogni $\mathbf{x} \in \mathbf{X}$ si decompone in modo unico come $\mathbf{x} = \mathbf{v} + \mathbf{w}$ con $\mathbf{v} \in \mathbf{V}$ e $\mathbf{w} \in \mathbf{V}^\perp$. Il vettore $\mathbf{v}$ si chiama la **proiezione ortogonale** di $\mathbf{x}$ su $\mathbf{V}$. In tal caso per ogni sottospazio vettoriale $\mathbf{V} \subset \mathbf{X}$ si ha $\dim \mathbf{V} + \dim \mathbf{V}^\perp = \dim \mathbf{X}$.

DIMOSTRAZIONE. Le parti (i), (ii) e (iii) sono ovvie: ne lasciamo la dimostrazione al lettore per esercizio. Per provare (iv), basta ricordare (Definizione 6.5.10) che il prodotto scalare è non degenere se e solo se nessun vettore è ortogonale a tutto lo spazio. Per provare (v) osserviamo che, se $\mathbf{x} \in \mathbf{V} \cap \mathbf{V}^\perp$, allora $\langle \mathbf{x}, \mathbf{x} \rangle = 0$ e quindi $\mathbf{x} = \mathbf{0}$ perché il prodotto scalare è definito positivo. Rimane solo da dimostrare l'enunciato sulle

dimensioni. Per ogni $\mathbf{x} \in \mathbf{X}$ scriviamo $\mathbf{x} = \mathbf{v} + \mathbf{w}$ come sopra e definiamo l'operatore di proiezione su $\mathbf{V}$ come $P\mathbf{x} = \mathbf{v}$. È evidente che P è una applicazione lineare, $P : \mathbf{X} \longrightarrow \mathbf{X}$, l'immagine di P è il sottospazio $\mathbf{V}$ e $\text{Ker}(P) = \mathbf{V}^\perp$. Ora l'identità $\dim \mathbf{V} + \dim \mathbf{V}^\perp = \dim X$ non è altro che il teorema della dimensione, Teorema 3.2.1.

NOTA 6.6.6. La dimostrazione della parte (v) della Proposizione 6.6.5, cioè dell'esistenza della proiezione ortogonale, qui è formulata senza far ricorso alla nozione di continuità della norma. Il nostro approccio vale pertanto in dimensione finita, ma l'enunciato è più generale, e vale per una classe di spazi vettoriali a dimensione infinita muniti di prodotto scalare (gli *spazi di Hilbert*). Per la dimostrazione nel caso generale rinviando il lettore ad un [libro su questo argomento](#), come ad esempio [9]. $\square$

LEMMA 6.6.7. *Se i vettori $\mathbf{x}_1, \dots, \mathbf{x}_k$ sono non zeri ed a due a due ortogonali rispetto ad un prodotto scalare non degenere, allora essi sono linearmente indipendenti.*

DIMOSTRAZIONE. Siano λ_i scalari tali che $\sum_{i=1}^k \lambda_i \mathbf{x}_i = \mathbf{0}$: dobbiamo mostrare che $\lambda_j = 0$ per ogni j con $1 \leq j \leq k$. Ma per ogni j si ha

$$0 = \left\langle \sum_{i=1}^k \lambda_i \mathbf{x}_i, \mathbf{x}_j \right\rangle = \lambda_j \langle \mathbf{x}_j, \mathbf{x}_j \rangle.$$

Poiché il prodotto scalare è non degenere, $\langle \mathbf{x}_j, \mathbf{x}_j \rangle \neq 0$ per ogni j : ne segue che $\lambda_j = 0$.

NOTAZIONE 6.6.8. *Dato un prodotto scalare, un insieme di vettori si dice un sistema ortogonale rispetto a quel prodotto scalare se essi sono a due a due ortogonali. Se inoltre tutti i vettori sono di norma 1 il sistema si chiama un emphisistema ortonormale.*

TEOREMA 6.6.9. **(Esistenza di basi ortogonali ed ortonormali.)**

- (i) *Sia $\mathbf{X}$ uno spazio vettoriale con un prodotto scalare: allora $\mathbf{X}$ ha una base ortogonale (cautela: non è detto che esista una base ortonormale, perché i vettori di base potrebbero avere*

lunghezza zero rispetto a questo prodotto scalare, e quindi non essere normalizzabili).

(ii) *Se il prodotto scalare è definito positivo, allora $\mathbf{X}$ ha una base ortonormale.*

DIMOSTRAZIONE.

(i) Procediamo per induzione su $n = \dim \mathbf{X}$ (la validità di questo metodo dimostrativo segue dall'assioma di induzione, assioma **P5** dei numeri interi, (1.3)). Se $n = 1$ le basi hanno un solo elemento: basta allora scegliere un qualsiasi vettore non nullo. Supponiamo $n > 1$. Ci sono due possibilità:

- o $\langle \mathbf{x}, \mathbf{x} \rangle = 0$ per ogni $\mathbf{x}$, oppure
- esiste un vettore $\mathbf{x}_1$ tale che $\langle \mathbf{x}_1, \mathbf{x}_1 \rangle \neq 0$.

Nel primo caso, il Lemma di polarizzazione 6.6.1 ci dice che il prodotto scalare è identicamente zero, cioè $\langle \mathbf{x}, \mathbf{y} \rangle = 0$ per ogni $\mathbf{x}, \mathbf{y}$. In questo caso banale ogni base è ortogonale (non ortonormale, perché tutti i vettori hanno norma zero!), e l'enunciato vale.

Nel secondo caso, indichiamo con $\mathbf{X}_1$ il sottospazio generato da $\mathbf{x}_1$. Sappiamo che la norma $\langle \mathbf{x}_1, \mathbf{x}_1 \rangle$ di $\mathbf{x}_1$ è non nulla, e questo equivale a dire che $\mathbf{X}_1$ non è contenuto nel suo ortogonale $\mathbf{X}_1^\perp$. Quindi, anche se ora, a differenza della Proposizione 6.6.5 (v), non sappiamo se $\dim \mathbf{X}_1 + \dim \mathbf{X}_1^\perp = n$, certamente sappiamo che $\dim \mathbf{X}_1^\perp < n$.

Ogni vettore $\mathbf{x} \in \mathbf{X}$ è la somma di un vettore in $\mathbf{X}_1$ ed un altro vettore nel sottospazio ortogonale $\mathbf{X}_1^\perp$: non abbiamo l'unicità della decomposizione come nella Proposizione 6.6.5 (v), perché ora non stiamo supponendo che il prodotto scalare sia definito positivo, ma l'esistenza sì, perché

$$\mathbf{x} = \left(\mathbf{x} - \frac{\langle \mathbf{x}, \mathbf{x}_1 \rangle}{\langle \mathbf{x}_1, \mathbf{x}_1 \rangle} \mathbf{x}_1 \right) + \frac{\langle \mathbf{x}, \mathbf{x}_1 \rangle}{\langle \mathbf{x}_1, \mathbf{x}_1 \rangle} \mathbf{x}_1. \quad (6.6.1)$$

Poiché $\dim \mathbf{X}_1^\perp < n$, per ipotesi di induzione possiamo assumere che esista una base ortogonale $\{\mathbf{x}_2, \mathbf{x}_3, \dots, \mathbf{x}_n\}$ in $\mathbf{X}_1^\perp$. Allora tutti questi vettori sono ortogonali a $\mathbf{x}_1$, e quindi

$\{\mathbf{x}_1, \mathbf{x}_2, \mathbf{x}_3, \dots, \mathbf{x}_n\}$ è una base ortogonale in $\mathbf{X}$. Questo prova (i).

(ii) Anche in questa parte della dimostrazione, se la dimensione $n = 1$ il risultato è ovvio: si prende un vettore non nullo e si osserva che, poiché il prodotto scalare è definito positivo, la norma di quel vettore è positiva, quindi normalizzandolo si ottiene una base ortonormale (consistente in un solo vettore, in questo caso).

Se invece $n > 1$ consideriamo un sottospazio $\mathbf{X}_1$ di dimensione 1, generato da un vettore $\mathbf{x}_1$ che come sopra possiamo assumere di norma 1. Il sottospazio ortogonale $\mathbf{X}_1^\perp$ ha dimensione $n - 1$ (Proposizione 6.6.5 (v)), e quindi, per ipotesi di induzione, ha una base ortonormale che indichiamo con $\{\mathbf{x}_2, \mathbf{x}_3, \dots, \mathbf{x}_n\}$. Allora, analogamente alla parte (i) della dimostrazione, $\{\mathbf{x}_1, \mathbf{x}_2, \mathbf{x}_3, \dots, \mathbf{x}_n\}$ è una base ortonormale in $\mathbf{X}$.

NOTA 6.6.10. La proprietà del Teorema 6.6.9 (ii) non vale se il prodotto scalare è soltanto non degenere. Infatti, sotto tali ipotesi, possono esistere vettori di norma zero, cioè ortogonali a sé stessi. Ad esempio, il prodotto scalare dell'Esempio 6.9.2 (iii), diciamo associato alla matrice diagonale

$$D = \begin{pmatrix} 1 & & & \\ & -1 & & \\ & & 0 & \\ & & & \ddots \\ & & & & 0 \end{pmatrix}$$

si espande come

$$\mathbf{x} \cdot \mathbf{y} = x_1 \overline{y_1} - x_2 \overline{y_2}$$

e quindi il vettore $(1, -1, 0, \dots, 0)$ ha norma nulla, cioè è ortogonale a sé stesso, ma il prodotto scalare è non degenere come osservato nell'Esempio 6.9.2 (iii) e 6.9.2 (iv). $\square$

6.7. Procedimento di ortogonalizzazione di Gram-Schmidt

In questa sezione ci proponiamo di fornire una dimostrazione costruttiva dell'esistenza di basi ortonormali in spazi vettoriali muniti di un prodotto scalare definito positivo (Teorema 6.6.9 (ii), basandoci sulla proiezione ortogonale (Proposizione 6.6.5 (v)) ed applicando iterativamente l'identità (6.6.1).

Consideriamo una base $\{\mathbf{x}_1, \dots, \mathbf{x}_n\}$ di uno spazio vettoriale $\mathbf{X}$. Scegliamo in $\mathbf{X}$ un prodotto scalare definito positivo, in modo che il prodotto scalare di ogni vettore non nullo con sé stesso sia positivo (una norma). Naturalmente il caso tipico è quello del prodotto scalare euclideo.

Vogliamo costruire esplicitamente una base *ortonormale* $\{\mathbf{x}'_1, \dots, \mathbf{x}'_n\}$. Il metodo che stiamo per sviluppare si chiama il *procedimento di ortogonalizzazione di Gram-Schmidt*.

Per prima cosa poniamo $\mathbf{x}'_1 = \mathbf{x}_1$. Poi definiamo $\mathbf{x}'_2$ come in (6.6.1):

$$\mathbf{x}'_2 = \mathbf{x}_2 - \frac{\langle \mathbf{x}_2, \mathbf{x}_1 \rangle}{\langle \mathbf{x}_1, \mathbf{x}_1 \rangle} \mathbf{x}_1.$$

In tal modo, come osservato in (6.6.1), si ha $\langle \mathbf{x}'_2, \mathbf{x}'_1 \rangle = 0$ (come del resto si verifica anche immediatamente). Si osservi che $\mathbf{x}'_2$ si ottiene da $\mathbf{x}_2$ sottraendogli la componente lungo il normalizzato di $\mathbf{x}'_1$:

$$\mathbf{x}'_2 = \mathbf{x}_2 - \left\langle \mathbf{x}_2, \frac{\mathbf{x}_1}{\sqrt{\langle \mathbf{x}_1, \mathbf{x}_1 \rangle}} \right\rangle \frac{\mathbf{x}_1}{\sqrt{\langle \mathbf{x}_1, \mathbf{x}_1 \rangle}}.$$

La dimostrazione della Proposizione 6.6.5 (v) rivela che stiamo sottraendo a $\mathbf{x}_2$ la sua proiezione ortogonale lungo $\mathbf{x}_1$. Analogamente, definiamo $\mathbf{x}'_3$ in maniera che sia ortogonale a $\mathbf{x}'_1$ e $\mathbf{x}'_2$: per questo si deve sottrarre a $\mathbf{x}_3$ la sua proiezione ortogonale sullo spazio bidimensionale generato da $\mathbf{x}'_1$ e $\mathbf{x}'_2$: quindi dobbiamo sottrargli le sue componenti lungo i vettori che si ottengono normalizzando $\mathbf{x}'_1$ e $\mathbf{x}'_2$, cioè

$$\mathbf{x}'_3 = \mathbf{x}_3 - \frac{\langle \mathbf{x}_3, \mathbf{x}'_1 \rangle}{\langle \mathbf{x}'_1, \mathbf{x}'_1 \rangle} \mathbf{x}'_1 - \frac{\langle \mathbf{x}_3, \mathbf{x}'_2 \rangle}{\langle \mathbf{x}'_2, \mathbf{x}'_2 \rangle} \mathbf{x}'_2.$$

Procediamo così per $i = 2, \dots, n$:

$$\mathbf{x}'_i = \mathbf{x}_i - \sum_{j=1}^{i-1} \frac{\langle \mathbf{x}_i, \mathbf{x}'_j \rangle}{\langle \mathbf{x}'_j, \mathbf{x}'_j \rangle} \mathbf{x}'_j. \quad (6.7.1)$$

In tal modo abbiamo dimostrato:

PROPOSIZIONE 6.7.1. (Ortogonalizzazione di Gram-Schmidt.) *La famiglia $\{\mathbf{x}'_1, \dots, \mathbf{x}'_n\}$ costruita iterativamente in (6.7.1) è una base ortogonale di $\mathbf{X}$.*

Si noti che, per costruzione, il sottospazio vettoriale generato da $\{\mathbf{x}'_1, \dots, \mathbf{x}'_i\}$ coincide per ogni i con il sottospazio vettoriale generato da $\{\mathbf{x}_1, \dots, \mathbf{x}_i\}$.

ESEMPIO 6.7.2. Consideriamo i seguenti tre vettori in $\mathbb{R}^4$:

$$x_1 = \begin{pmatrix} 1 \\ -1 \\ 1 \\ -1 \end{pmatrix}, \quad x_2 = \begin{pmatrix} 1 \\ 0 \\ 1 \\ 0 \end{pmatrix}, \quad x_3 = \begin{pmatrix} 1 \\ 0 \\ 0 \\ 0 \end{pmatrix}$$

Vogliamo applicare il procedimento di ortogonalizzazione di Gram-Schmidt per trovare una base ortogonale del sottospazio di $\mathbb{R}^4$ generato da x_1, x_2, x_3 , rispetto al consueto prodotto scalare euclideo.

SVOLGIMENTO. Innanzitutto poniamo

$$x'_1 = x_1 = \begin{pmatrix} 1 \\ -1 \\ 1 \\ -1 \end{pmatrix}.$$

Dopo si definisce:

$$x'_2 = x_2 - \frac{x_2 \cdot x'_1}{x'_1 \cdot x'_1} x'_1 = \begin{pmatrix} 1 \\ 0 \\ 1 \\ 0 \end{pmatrix} - \frac{1}{2} \begin{pmatrix} 1 \\ -1 \\ 1 \\ -1 \end{pmatrix} = \frac{1}{2} \begin{pmatrix} 1 \\ 1 \\ 1 \\ 1 \end{pmatrix}$$

Infine si definisce:

$$x'_3 = x_3 - \frac{x_3 \cdot x'_2}{x'_2 \cdot x'_2} x'_2 - \frac{x_3 \cdot x'_1}{x'_1 \cdot x'_1} x'_1 = \begin{pmatrix} 1 \\ 0 \\ 0 \\ 0 \end{pmatrix} - \frac{1}{4} \begin{pmatrix} 1 \\ 1 \\ 1 \\ 1 \end{pmatrix} - \frac{1}{4} \begin{pmatrix} 1 \\ -1 \\ 1 \\ -1 \end{pmatrix} = \frac{1}{2} \begin{pmatrix} 1 \\ 0 \\ -1 \\ 0 \end{pmatrix}$$

Abbiamo trovato così la seguente base ortogonale:

$$x'_1 = \begin{pmatrix} 1 \\ -1 \\ 1 \\ -1 \end{pmatrix}, \quad x'_2 = \frac{1}{2} \begin{pmatrix} 1 \\ 1 \\ 1 \\ 1 \end{pmatrix}, \quad x'_3 = \frac{1}{2} \begin{pmatrix} 1 \\ 0 \\ -1 \\ 0 \end{pmatrix}.$$

□

ESEMPIO 6.7.3. Troviamo una base ortogonale per il sottospazio lineare di $\mathbb{R}^4$ generato dai vettori

$$\begin{pmatrix} 1 \\ -1 \\ -1 \\ 1 \end{pmatrix}, \quad \begin{pmatrix} 1 \\ 2 \\ -1 \\ -1 \end{pmatrix}, \quad \begin{pmatrix} -1 \\ 1 \\ 1 \\ 1 \end{pmatrix}, \quad \begin{pmatrix} 2 \\ 1 \\ -2 \\ 2 \end{pmatrix}.$$

SVOLGIMENTO. Per prima cosa vediamo se i quattro vettori dati sono linearmente indipendenti o no, trovando le soluzioni del sistema omogeneo

$$\begin{cases} x_1 + x_2 - x_3 + 2x_4 = 0 \\ -x_1 + 2x_2 + x_3 + x_4 = 0 \\ -x_1 - x_2 + x_3 - 2x_4 = 0 \\ x_1 - x_2 + x_3 + 2x_4 = 0 \end{cases}$$

mediante eliminazione di Gauss (Sezione 3.7):

$$\begin{array}{cccc|c} 1 & 1 & -1 & 2 & 0 \\ -1 & 2 & 1 & 1 & 0 \\ -1 & -1 & 1 & -2 & 0 \\ 1 & -1 & 1 & 2 & 0 \end{array} \quad \begin{array}{cccc|c} 1 & 1 & -1 & 2 & 0 \\ 0 & 3 & 0 & 3 & 0 \\ 0 & 0 & 0 & 0 & 0 \\ 0 & -2 & 2 & 0 & 0 \end{array}$$

$$\begin{array}{cccc|cccc|c}
1 & 1 & -1 & 2 & 0 & 1 & 1 & -1 & 2 & 0 \\
0 & 1 & 0 & 1 & 0 & 0 & 1 & 0 & 1 & 0 \\
0 & -1 & 1 & 0 & 0 & 0 & 0 & 1 & 1 & 0 \\
0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0
\end{array}$$

Ne segue che il sistema omogeneo di cui sopra ammette soluzione con $x_4 \neq 0$, quindi il quarto vettore è combinazione lineare dei primi tre. Di conseguenza il sottospazio lineare $\mathbf{X}$ di $\mathbb{R}^4$ generato dai quattro vettori dati è già generato dai primi tre.

Per vedere se i primi tre vettori sono linearmente indipendenti o no, dobbiamo trovare le soluzioni del sistema omogeneo

$$\begin{cases}
x_1 + x_2 - x_3 = 0 \\
-x_1 + 2x_2 + x_3 = 0 \\
-x_1 - x_2 + x_3 = 0 \\
x_1 - x_2 + x_3 = 0
\end{cases},$$

cioè le soluzioni del sistema precedente con $x_4 = 0$. Ma i calcoli di cui sopra ci mostrano che, se $x_4 = 0$, allora anche $x_1 = x_2 = x_3 = 0$. Perciò i tre vettori

$$\mathbf{a}_1 = \begin{pmatrix} 1 \\ -1 \\ -1 \\ 1 \end{pmatrix}, \quad \mathbf{a}_2 = \begin{pmatrix} 1 \\ 2 \\ -1 \\ -1 \end{pmatrix}, \quad \mathbf{a}_3 = \begin{pmatrix} -1 \\ 1 \\ 1 \\ 1 \end{pmatrix}$$

sono linearmente indipendenti e quindi costituiscono una base di $\mathbf{X}$.

Per ottenere una base ortogonale $\mathbf{b}_1, \mathbf{b}_2, \mathbf{b}_3$ di $\mathbf{X}$, ortogonalizziamo i vettori $\mathbf{a}_1, \mathbf{a}_2, \mathbf{a}_3$:

$$\mathbf{b}_1 = \mathbf{a}_1 = \begin{pmatrix} 1 \\ -1 \\ -1 \\ 1 \end{pmatrix}.$$

Per $\mathbf{b}_2$ possiamo prendere qualsiasi multiplo non zero di

$$\begin{aligned} \mathbf{a}_2 - \frac{\mathbf{a}_2 \cdot \mathbf{b}_1}{|\mathbf{b}_1|^2} \mathbf{b}_1 &= \\ &= \begin{pmatrix} 1 \\ 2 \\ -1 \\ -1 \end{pmatrix} - \frac{\begin{pmatrix} 1 \\ 2 \\ -1 \\ -1 \end{pmatrix} \cdot \begin{pmatrix} 1 \\ -1 \\ -1 \\ 1 \end{pmatrix}}{\left| \begin{pmatrix} 1 \\ -1 \\ -1 \\ 1 \end{pmatrix} \right|^2} \begin{pmatrix} 1 \\ -1 \\ -1 \\ 1 \end{pmatrix} = \frac{1}{4} \begin{pmatrix} 5 \\ 7 \\ -5 \\ -3 \end{pmatrix} \end{aligned}$$

e ci conviene scegliere

$$\mathbf{b}_2 = \begin{pmatrix} 5 \\ 7 \\ -5 \\ -3 \end{pmatrix}.$$

Infine, per $\mathbf{b}_3$ scegliamo un multiplo non zero di

$$\begin{aligned} \mathbf{a}_3 - \frac{\mathbf{a}_3 \cdot \mathbf{b}_1}{|\mathbf{b}_1|^2} \mathbf{b}_1 - \frac{\mathbf{a}_3 \cdot \mathbf{b}_2}{|\mathbf{b}_2|^2} \mathbf{b}_2 &= \\ &= \begin{pmatrix} -1 \\ 1 \\ 1 \\ 1 \end{pmatrix} - \frac{-2}{4} \begin{pmatrix} 1 \\ -1 \\ -1 \\ 1 \end{pmatrix} - \frac{-6}{108} \begin{pmatrix} 5 \\ 7 \\ -5 \\ -3 \end{pmatrix} = \frac{2}{9} \begin{pmatrix} -1 \\ 4 \\ 1 \\ 6 \end{pmatrix} \end{aligned}$$

ed è naturale scegliere

$$\mathbf{b}_3 = \begin{pmatrix} -1 \\ 4 \\ 1 \\ 6 \end{pmatrix}.$$

Concludiamo che i tre vettori

$$\mathbf{b}_1 = \begin{pmatrix} 1 \\ -1 \\ -1 \\ 1 \end{pmatrix}, \quad \mathbf{b}_2 = \begin{pmatrix} 5 \\ 7 \\ -5 \\ -3 \end{pmatrix}, \quad \mathbf{b}_3 = \begin{pmatrix} -1 \\ 4 \\ 1 \\ 6 \end{pmatrix}$$

costituiscono una base ortogonale di $\mathbf{X}$.

□

6.8. Matrici ortogonali e matrici unitarie

LEMMA 6.8.1. (Preservazione di norme e di angoli.) *Siamo $\mathbf{x}$ e $\mathbf{y}$ due vettori in uno spazio vettoriale complesso $X_{\mathbb{C}}$, ed A una applicazione lineare di $X_{\mathbb{C}}$ in sé. Allora $\langle A\mathbf{x}, A\mathbf{x} \rangle = \langle \mathbf{x}, \mathbf{x} \rangle$ per ogni vettore $\mathbf{x}$ se e solo se $\langle A\mathbf{x}, A\mathbf{y} \rangle = \langle \mathbf{x}, \mathbf{y} \rangle$ per ogni coppia di vettori $\mathbf{x}, \mathbf{y}$.*

DIMOSTRAZIONE. Costruiamo la forma hermitiana

$$\phi(\mathbf{x}, \mathbf{y}) = \langle A\mathbf{x}, A\mathbf{y} \rangle - \langle \mathbf{x}, \mathbf{y} \rangle.$$

L'enunciato segue immediatamente applicando a questa forma hermitiana il Lemma 6.6.1.

DEFINIZIONE 6.8.2. (Matrici ortogonali e matrici unitarie.) Una matrice quadrata $A \in M_n^{\mathbb{R}}$ si dice *ortogonale* se preserva la norma dei vettori (che abbiamo anche chiamato lunghezza), cioè se $\|A\mathbf{x}\| = \langle A\mathbf{x}, A\mathbf{x} \rangle = \langle \mathbf{x}, \mathbf{x} \rangle = \|\mathbf{x}\|^2$ per ogni vettore $\mathbf{x}$. Analogamente, una matrice quadrata $A \in M_n^{\mathbb{C}}$ si dice *unitaria* se preserva la norma in $\mathbb{C}^n$.

COROLLARIO 6.8.3. *Il prodotto righe per colonne di matrici unitarie (rispettivamente, ortogonali) è una matrice unitaria (rispettivamente, ortogonale).*

DIMOSTRAZIONE. Ogni matrice definisce una applicazione lineare, ed il prodotto righe per colonne corrisponde alla composizione delle applicazioni (Definizione 3.6.3). Se due applicazioni preservano la norma, anche il prodotto la preserva.

NOTA 6.8.4. Se A è una matrice ortogonale o unitaria, allora, grazie al Lemma 6.6.1, per ogni coppia di vettori $\mathbf{x}, \mathbf{y}$ in $\mathbb{R}^n$ o $\mathbb{C}^n$ rispettivamente si ha $A\mathbf{x} \cdot A\mathbf{y} = \mathbf{x} \cdot \mathbf{y}$. In particolare, ricordando che se due vettori in $\mathbb{R}^n$ formano un angolo θ allora $\mathbf{x} \cdot \mathbf{y} = \|\mathbf{x}\| \|\mathbf{y}\| \cos \theta$, ed osservando che $\|A\mathbf{x}\| = \|\mathbf{x}\|$ per la Definizione 6.8.2, concludiamo che l'azione di una matrice ortogonale A su $\mathbb{R}^n$ preserva gli angoli fra i vettori. □

ESEMPIO 6.8.5. (*Matrici di rotazione.*) La Nota 6.8.4 mostra che le matrici ortogonali corrispondono agli operatori di rotazione su $\mathbb{R}^n$: sono quelle che ruotano una base ortonormale in un'altra base ortonormale, quindi le loro colonne sono vettori ortonormali. Ad esempio, in $\mathbb{R}^2$, le loro colonne sono i vettori del tipo $(\cos \theta, \sin \theta)$, se chiamiamo θ l'angolo di rotazione (si veda più in generale il successivo Teorema 6.8.7). In particolare, la più generale matrice reale ortogonale a dimensione 2 è

$$\begin{pmatrix} \cos \theta & \sin \theta \\ -\sin \theta & \cos \theta \end{pmatrix}$$

che manda il vettore $\mathbf{e}_1 = \begin{pmatrix} 1 \\ 0 \end{pmatrix}$ in $\begin{pmatrix} \cos \theta \\ -\sin \theta \end{pmatrix}$, e quindi è una rotazione di un angolo θ in senso antiorario. $\square$

COROLLARIO 6.8.6. (i) Una matrice $A \in M_{\mathbb{R}}^n$ è ortogonale se e solo se $A^{-1} = A^T$

(ii) Una matrice $A \in M_{\mathbb{C}}^n$ è unitaria se e solo se $A^{-1} = A^*$

DIMOSTRAZIONE. Basta provare la proprietà (ii), perché essa implica la (i).

Per la precedente Nota 6.8.4, A è unitaria se e solo se, per ogni $\mathbf{x}, \mathbf{y}$ in $\mathbb{C}^n$, si ha $\mathbf{x} \cdot A^* A \mathbf{y} = A \mathbf{x} \cdot A \mathbf{y} = \mathbf{x} \cdot \mathbf{y}$. Quindi $\mathbf{x} \cdot (A^* A - \mathbb{I}) \mathbf{y} = 0$. Poiché il prodotto scalare è non degenere, questo equivale a dire che $A^* A = \mathbb{I}$, cioè $A^* = A^{-1}$. Si noti che, a sua volta, quest'ultima asserzione equivale a dire che $AA^* = \mathbb{I}$.)

TEOREMA 6.8.7. Le seguenti condizioni sono equivalenti:

- (i) una matrice $A \in M_{\mathbb{C}}^n$ è unitaria;
- (ii) le colonne di A sono vettori ortonormali rispetto al prodotto scalare euclideo in $\mathbb{C}^n$,
- (iii) le righe di A sono vettori ortonormali rispetto al prodotto scalare euclideo in $\mathbb{C}^n$,
- (iv) A^* è unitaria,
- (v) A manda una base ortonormale in una base ortonormale (rispetto al prodotto scalare euclideo).

In particolare, una matrice reale è ortogonale se e solo se le sue colonne sono ortonormali rispetto al prodotto scalare in $\mathbb{R}^n$, e se e solo se lo sono le sue righe, e se e solo se lo è $A^\top$, e se e solo se manda una base ortonormale di $\mathbb{R}^n$ in una base ortonormale di $\mathbb{R}^n$.

DIMOSTRAZIONE. Sia $A \in M_n^{\mathbb{C}}$ unitaria. Indichiamo come sempre con $\{\mathbf{e}_i\}$ i vettori della base canonica in $\mathbb{C}^n$. Utilizziamo la seguente notazione abituale: il simbolo δ_{ij} , detto *simbolo di Krönecker*, vale 1 se $i = j$ e 0 altrimenti. Segue dalla Nota 6.8.4 che

$$A\mathbf{e}_i \cdot A\mathbf{e}_j = \mathbf{e}_i \cdot \mathbf{e}_j = \delta_{ij}. \quad (6.8.1)$$

Pertanto le colonne di A , cioè i vettori $A\mathbf{e}_i$, sono una famiglia ortonormale, e quindi (i) implica (ii). Viceversa, se le colonne sono ortonormali, cioè se vale (6.8.1), allora, scrivendo $\mathbf{x} = \sum_{i=1}^n x_i \mathbf{e}_i$, otteniamo per linearità che $A\mathbf{x} \cdot A\mathbf{x} = \sum_{i,j=1}^n x_i \overline{x_j} \delta_{ij} = \sum_{i=1}^n |x_i|^2 = \|\mathbf{x}\|^2$, e quindi A è unitaria. Pertanto (ii) implica (i).

Ora, se A è unitaria, grazie a (ii) essa è invertibile, quindi lo è anche A^* , perché $\det A^* = \det \overline{A^\top} = \overline{\det A^\top} = \overline{\det A}$. Perciò ogni vettore $\mathbf{y} \in \mathbb{C}^n$ si può scrivere (in modo unico) come $\mathbf{y} = A^*\mathbf{x}$ per qualche vettore $\mathbf{x}$. Pertanto, per definizione di aggiunto (Definizione 6.4.4),

$$A^*\mathbf{x} \cdot A^*\mathbf{x} = \mathbf{y} \cdot \mathbf{y} = A\mathbf{y} \cdot A\mathbf{y} = A^{*-1}\mathbf{y} \cdot A^{*-1}\mathbf{y} = \mathbf{x} \cdot \mathbf{x}$$

quindi A^* è unitaria e (i) implica (iv). Poiché $A^{**} = A$, anche (iv) implica (i). Allora, in base a (ii), il fatto che A^* sia unitaria equivale ad asserire che le sue colonne formino una famiglia ortonormale di vettori: ma queste colonne sono il complesso coniugato delle righe di A (Definizione 6.4.4), e quindi le righe di A sono una famiglia ortonormale, e (iv) equivale a (iii). Infine, il fatto che le colonne di A siano un sistema ortonormale chiaramente equivale a dire che A manda la base canonica in una base ortonormale. Pertanto la matrice che manda una base ortonormale nella base canonica è l'inversa A^{-1} di una matrice unitaria: ma per una matrice unitaria si ha $A^{-1} = A^*$, e A^* è ancora unitaria grazie alla parte (iv) del teorema. Da questo segue che una matrice che manda una base ortonormale $\mathcal{F}$ in una base ortonormale $\mathcal{V}$ è il prodotto di due matrici unitarie (quella che manda $\mathcal{F}$ nella base canonica e quella che manda la base canonica in $\mathcal{V}$). Per il Corollario

6.8.3 questo prodotto è ancora una matrice unitaria. Questo prova la parte (v).

6.9. * Matrice associata ad un prodotto scalare

L'Esempio 6.3.4 sottolinea che, quando il prodotto scalare è espresso in termini di coordinate, esso evidentemente dipende dalla scelta della base. Chiariamo ora questa dipendenza.

PROPOSIZIONE 6.9.1. (Matrice associata ad una forma bilineare o hermitiana.) *Esiste una corrispondenza biunivoca fra le forme bilineari o hermitiane ϕ su $\mathbb{R}^n$ (rispettivamente $\mathbb{C}^n$) e le matrici M $n \times n$, che verifica la regola seguente: la matrice M dà origine alla forma ϕ_M definita da*

$$\phi_M(\mathbf{x}, \mathbf{y}) = \mathbf{x} \cdot M\mathbf{y}$$

ovvero, facendo uso del prodotto righe per colonne,

$$\phi_M(\mathbf{x}, \mathbf{y}) = \mathbf{x}^\top M\mathbf{y}. \quad (6.9.1)$$

In altre parole, in $\mathbb{R}^n$

$$\phi_M(\mathbf{x}, \mathbf{y}) = \sum_{i,j=1}^n m_{ij}x_iy_j$$

(dove i numeri x_i e y_i sono le coordinate nella base canonica) e più in generale in $\mathbb{C}^n$

$$\phi_M(\mathbf{x}, \mathbf{y}) = \sum_{i,j=1}^n m_{ij}x_i\overline{y_j}.$$

Inoltre:

- ϕ_M è un prodotto scalare reale se e solo se M è simmetrica (nel senso della Definizione 6.4.6, $M = M^\top$;
- ϕ_M è un prodotto scalare complesso (ovvero prodotto hermitiano) se e solo se M è autoaggiunta (nel senso della Definizione 6.4.6, $M = M^* := \overline{M}^\top$;
- se ϕ_M è un prodotto scalare, allora esso è non degenere se e solo se $\det M \neq 0$;

- se ϕ_M è un prodotto scalare, allora esso è definito positivo se e solo se M è definita positiva, nel senso che $\sum_{i,j=1}^n m_{ij} x_i \overline{x_j} \geq 0$ per ogni n -pla di numeri complessi $\{x_1, \dots, x_n\}$.

DIMOSTRAZIONE. Data una matrice M a dimensione n , è chiaro che ϕ_M è una forma bilineare se si usa il prodotto scalare reale su $\mathbb{R}^n$, o hermitiana si usa il prodotto scalare complesso su $\mathbb{C}^n$, nel senso dell'Esempio 6.3.4. Viceversa, data una forma bilineare (rispettivamente hermitiana) ϕ , consideriamo la matrice M i cui coefficienti sono

$$m_{ij} = \phi(\mathbf{e}_i, \mathbf{e}_j)$$

dove i vettori $\mathbf{e}_i$ sono i vettori della base canonica. Grazie alla bilinearità si ha

$$\phi(\mathbf{x}, \mathbf{y}) = \sum_{i,j=1}^n m_{ij} x_i \overline{y_j}$$

dove i numeri x_i e y_i sono le coordinate nella base canonica. Ne segue che M è la matrice associata alla forma ϕ , cioè che $\phi = \phi_M$, e quindi la corrispondenza è biunivoca.

Delle restanti proprietà, le prime due sono conseguenze ovvie delle formule della Proposizione 6.4.2. Proviamo la terza per un prodotto scalare reale. La forma bilineare ϕ_M dà luogo ad un prodotto scalare degenerare se e solo se esiste un vettore $\mathbf{y} \neq \mathbf{0}$ tale che $\phi_M(\mathbf{x}, \mathbf{y}) = \mathbf{x} \cdot M\mathbf{y} = 0$ per ogni $\mathbf{x}$. In particolare, $M\mathbf{y} \cdot M\mathbf{y} = 0$. Ma il prodotto scalare euclideo è definito positivo (Nota 6.5.8), quindi $M\mathbf{y} = \mathbf{0}$ e M non è iniettiva. Viceversa, se $\text{Ker}(M) \neq \mathbf{0}$, sia $\mathbf{y}$ un vettore non nullo in $\text{Ker}(M)$: allora $\phi_M(\mathbf{x}, \mathbf{y}) = \mathbf{x} \cdot M\mathbf{y} = 0$ per ogni $\mathbf{x}$ ed il prodotto scalare ϕ_M è degenerare. Infine, dire che ϕ_M è un prodotto scalare definito positivo equivale a dire che $\sum_{i,j=1}^n m_{ij} x_i \overline{x_j} = \phi_M(\mathbf{x}, \mathbf{x}) \geq 0$ per ogni $\mathbf{x}$.

ESEMPIO 6.9.2. (**Prodotti scalari associati a matrici diagonali.**)

- (i) Il prodotto scalare euclideo è associato alla matrice identità: $\mathbf{x} \cdot \mathbf{y} = \mathbf{x} \cdot \mathbb{I}\mathbf{y}$.
- (ii) Un prodotto scalare ϕ_D associato ad una matrice diagonale D è definito positivo se e solo se tutti i termini diagonali $\{d_i, \quad i = 1, \dots, n\}$ della matrice sono positivi. Infatti un tale prodotto

scalare si espande come $\phi_D(\mathbf{x}, \mathbf{x}) = \sum_{i=1}^n d_i |x_i|^2$. Perciò, se i d_i sono tutti positivi, allora $\phi_D(\mathbf{x}, \mathbf{x}) = \sum_{i=1}^n d_i |x_i|^2 \geq 0$, invece se esiste un termine negativo, diciamo d_1 , allora $\phi_D(\mathbf{e}_1, \mathbf{e}_1) = d_1 < 0$.

- (iii) Un prodotto scalare ϕ_D associato ad una matrice diagonale D che possiede due termini diagonali di segno opposto, diciamo $d_1 < 0 < d_2$ non è quindi definito positivo, ed in particolare l'espressione $N(\mathbf{x}) = \phi_D(\mathbf{x}, \mathbf{x})$ non è una norma perché esistono vettori non nulli con $N(\mathbf{x}) = 0$. Ad esempio, consideriamo i vettori $\mathbf{x} = (x_1, x_2, 0, 0, \dots, 0)$ con tutte le componenti nulle dopo la seconda: per questi vettori si ha $\phi_D(\mathbf{x}, \mathbf{x}) = d_1 |x_1|^2 - d_2 |x_2|^2$, ed il secondo membro si annulla per opportuni x_1 e x_2 non nulli. Nonostante questo, il prodotto $\mathbf{x}$ tale che $\phi_D(\mathbf{x}, \mathbf{y}) = 0$ per ogni vettore $\mathbf{y}$, potremmo prendere al posto di $\mathbf{y}$ lo i -simo scalare è non degenerare a meno che qualcuno dei termini diagonali sia nullo. Infatti, se esistesse un vettore non nullo vettore della base canonica, $\mathbf{e}_i$, ed otterremmo $d_i x_i = \phi_D(\mathbf{x}, \mathbf{e}_i) = 0$, da cui $x_i = 0$ visto che $d_i \neq 0$. Ma allora tutte le componenti x_i di $\mathbf{x}$ sono nulle, il che contraddice l'ipotesi $\mathbf{x} \neq \mathbf{0}$.
- (iv) Se invece la matrice diagonale D ha un termine diagonale nullo, diciamo $d_1 = 0$, allora ϕ_D è degenerare, perché $\phi_D(\mathbf{e}_1, \mathbf{y}) = d_1 |y_1|^2 = 0$ per ogni vettore $\mathbf{y}$.

□

NOTAZIONE 6.9.3. (Prodotti scalari in forma diagonale.) *Quando un prodotto scalare è associato ad una matrice diagonale, come nel precedente Esempio 6.9.2, diciamo che esso è espresso in forma diagonale, o anche senza termini misti. Se il prodotto scalare non è espresso in forma diagonale ma lo diventa in una base diversa, allora diciamo che esso è riducibile a forma diagonale tramite un cambiamento di base.*

Ora finalmente possiamo discutere cosa succede quando si fissa un

prodotto scalare, ad esempio il prodotto scalare euclideo nella base canonica, e poi si cambia base.

PROPOSIZIONE 6.9.4. (Prodotti scalari e cambio di base.)

- (i) Sia M la matrice associata ad un prodotto scalare in una data base $\mathcal{F}$ nel senso della Proposizione 6.9.1, e sia $\mathcal{F}'$ una nuova base. Sia $C = C_{\mathcal{F}, \mathcal{F}'}$ la matrice di passaggio dalla base $\mathcal{F}'$ alla base $\mathcal{F}$, come nella Notazione 4.1.5. Allora la matrice associata al prodotto scalare nella base $\mathcal{F}'$ è $C^\top MC$.
- (ii) Se le basi $\mathcal{F}$ e $\mathcal{F}'$ sono basi ortonormali, allora la matrice associata al prodotto scalare nella nuova base $\mathcal{F}'$ è la matrice $C^{-1}MC$.
- (iii) La base canonica è ortogonale per un prodotto scalare se e solo se la matrice associata al prodotto scalare nella base canonica è diagonale; analogamente, una base è ortogonale per un prodotto scalare se e solo se la matrice ad esso associata in quella base è una matrice diagonale.
- (iv) Ogni prodotto scalare è riducibile a forma diagonale tramite un cambiamento di base. Esiste quindi una matrice invertibile C tale che $C^\top MC$ è una matrice diagonale.
- (v) Se il prodotto scalare è definito positivo, la matrice di cambiamento di base che realizza la diagonalizzazione è unitaria nel caso complesso, e ortogonale nel caso reale (come definite nella Definizione 6.8.2). Data la matrice M associata al prodotto scalare, esiste quindi una matrice unitaria (od ortogonale) C tale che $C^\top MC = C^{-1}MC$ è una matrice diagonale.

DIMOSTRAZIONE. Abbiamo visto (6.9.1) che il legame fra il prodotto scalare la matrice M è il seguente: $\langle \mathbf{x}, \mathbf{y} \rangle = \mathbf{x}^\top M \mathbf{y}$. Per la bilinearità (proprietà (i) – (a₁) e (i) – (a₂) della Definizione 6.3.1, e analoghe nel caso hermitiano), il prodotto scalare è univocamente identificato dai coefficienti di matrice $m_{ij} = \mathbf{f}_j^\top M \mathbf{f}_i$ (nella base $\mathcal{F}$). D'altra parte, abbiamo osservato nella Nota 4.1.3 che, se $A = C_{\mathcal{F}', \mathcal{F}}$, allora $A^{-1} \mathbf{f}_i = \mathbf{f}'_i$ per ogni $i = 1, \dots, n$; poiché $A^{-1} = C_{\mathcal{F}, \mathcal{F}'} := C$ questo equivale a dire $C^{-1} \mathbf{f}_i = \mathbf{f}'_i$, cioè $C \mathbf{f}'_i = \mathbf{f}_i$. Allora nella base $\mathcal{F}'$ il prodotto scalare è associato alla matrice $m'_{ij} := \mathbf{f}'_j^\top M \mathbf{f}'_i$, cioè alla matrice $M' = C^\top MC$.

Questo prova (i).

Per provare (ii), osserviamo che, se le due basi sono ortonormali, allora la matrice C di cambiamento di base è una matrice unitaria (o, nel caso reale, ortogonale) per il Teorema 6.8.7. Pertanto $C^\top = C^{-1}$, e la parte (ii) segue dalla parte (i).

La base canonica è ortogonale rispetto al prodotto scalare associato in essa ad una matrice M se e solo se $m_{ij} = \langle \mathbf{e}_i, \mathbf{e}_j \rangle = 0$ per $i \neq j$, cioè se e solo se la matrice M è diagonale. Identico ragionamento vale in qualunque altra base. Questo prova (iii).

Proviamo ora la parte (iv). Sia M la matrice associata al prodotto scalare in una data base. Sappiamo che M è simmetrica nel caso reale, e più in generale autoaggiunta nel caso complesso. In entrambi i casi, per la parte (iii) di questo teorema, M è diagonale se e solo se la base è ortogonale rispetto al prodotto scalare. D'altro canto, grazie al Teorema 6.6.9, una base ortogonale esiste. Quindi, in questa base, M si riduce a forma diagonale, ed il prodotto scalare diventa espresso in forma diagonale.

In maggior dettaglio, supponiamo dapprima che lo spazio vettoriale si $\mathbb{R}^n$ e che le matrici siano reali. Sia C la matrice che implementa il cambiamento di base: se indichiamo con $\mathbf{e}_i$ i vettori della base canonica e con $\mathbf{f}_i$ quelli della base ortogonale, abbiamo $C\mathbf{e}_i = \mathbf{f}_i$. Consideriamo un nuovo prodotto scalare

$$\langle \mathbf{x}, \mathbf{y} \rangle' := \langle C\mathbf{x}, C\mathbf{y} \rangle = (C\mathbf{x})^\top M C\mathbf{y} = \mathbf{x}^\top C^\top M C \mathbf{x}.$$

Questo nuovo prodotto scalare è quello ottenuto passando alla nuova base, e cioè quello trovato nella parte (i) della dimostrazione, la cui matrice associata è $C^\top M C$. Osserviamo che la base canonica è ortogonale rispetto a questo nuovo prodotto scalare:

$$\langle \mathbf{e}_i, \mathbf{e}_j \rangle' = \langle \mathbf{f}_i, \mathbf{f}_j \rangle = 0 \quad \text{se } i \neq j$$

perché la base dei vettori $\mathbf{f}_i$ è ortogonale rispetto al prodotto scalare originale. Quindi la matrice associata al nuovo prodotto scalare, cioè $C^\top M C$, è diagonale grazie alla parte (iii).

Nel caso dello spazio complesso $\mathbb{C}^n$ l'argomento è analogo. Questo prova (iv).

La parte (v) equivale a dimostrare che per ogni matrice A autoaggiunta (rispettivamente, simmetrica) esiste una matrice U unitaria (rispettivamente, ortogonale) tale che $U^T A U = U^{-1} A U$ è diagonale. Questo risultato verrà dimostrato in seguito nel Corollario 7.5.4.

CAPITOLO 7

Autovalori, autovettori e diagonalizzabilità

In questo capitolo analizziamo in maggiore profondità la nozione di *diagonalizzazione* accennata nella Sezione 4.3 ed elaborata nell'Esercizio 4.4.3.

7.1. Triangolarizzazione e diagonalizzazione

Siano $\mathbf{X}, \mathbf{Y}$ due spazi vettoriali e $T : \mathbf{X} \longrightarrow \mathbf{Y}$ una applicazione lineare. Siano anche $\mathcal{V} = \{\mathbf{v}_1, \dots, \mathbf{v}_n\}$ e $\mathcal{V}' = \{\mathbf{v}'_1, \dots, \mathbf{v}'_n\}$ due basi di $\mathbf{X}$, mentre $\mathcal{F} = \{\mathbf{f}_1, \dots, \mathbf{f}_m\}$ e $\mathcal{F}' = \{\mathbf{f}'_1, \dots, \mathbf{f}'_m\}$ sono due basi di $\mathbf{Y}$. Il Teorema 4.2.5 esprime la matrice di T rispetto ad $\mathcal{V}'$ ed $\mathcal{F}'$ in termini della matrice di T rispetto ad $\mathcal{V}$ ed $\mathcal{F}$ e le matrici di cambiamenti di base nel modo seguente:

$$A_{\mathcal{F}', \mathcal{V}'}(T) = C_{\mathcal{F}', \mathcal{F}} A_{\mathcal{F}, \mathcal{V}}(T) C_{\mathcal{V}, \mathcal{V}'} = C_{\mathcal{F}, \mathcal{F}'}^{-1} A_{\mathcal{F}, \mathcal{V}}(T) C_{\mathcal{V}, \mathcal{V}'}.$$

In particolare, nel caso di $A \in M_{mn}$, una base $\mathcal{V} = \{\mathbf{v}_1, \dots, \mathbf{v}_n\}$ di $\mathbb{R}^n$ ed una base $\mathcal{F} = \{\mathbf{f}_1, \dots, \mathbf{f}_m\}$ di $\mathbb{R}^m$, la matrice dell'applicazione lineare

$$T_A : \mathbb{R}^n \ni \mathbf{x} \longmapsto A \mathbf{x} \in \mathbb{R}^m$$

rispetto alle basi $\mathcal{V}$ ed $\mathcal{F}$ è uguale ad

$$\begin{aligned} A_{\mathcal{F}, \mathcal{V}}(T_A) &= C_{\mathcal{N}_m, \mathcal{F}}^{-1} A_{\mathcal{N}_m, \mathcal{N}_n}(T_A) C_{\mathcal{N}_n, \mathcal{V}} = C_{\mathcal{N}_m, \mathcal{F}}^{-1} A C_{\mathcal{N}_n, \mathcal{V}} \\ &= \begin{pmatrix} | & & | \\ \mathbf{f}_1 & \dots & \mathbf{f}_m \\ | & & | \end{pmatrix}^{-1} A \begin{pmatrix} | & & | \\ \mathbf{v}_1 & \dots & \mathbf{v}_n \\ | & & | \end{pmatrix}. \end{aligned}$$

Pertanto, se possiamo trovare delle basi $\mathcal{V}$ ed $\mathcal{F}$ tale che $A_{\mathcal{F}, \mathcal{V}}(T_A)$ abbia una forma particolare, per esempio

- **triangolare** (*superiore*), cioè con tutti gli elementi sotto la diagonale uguali a zero, o

- **diagonale**, cioè con tutti gli elementi fuori della diagonale uguali a zero,

allora A si esprime tramite una matrice di questo tipo particolare mediante la formula

$$A = \begin{pmatrix} | & & | \\ \mathbf{f}_1 & \dots & \mathbf{f}_m \\ | & & | \end{pmatrix} A_{\mathcal{F}, \mathcal{V}}(T_A) \begin{pmatrix} | & & | \\ \mathbf{v}_1 & \dots & \mathbf{v}_n \\ | & & | \end{pmatrix}^{-1}. \quad (7.1.1)$$

Particolarmente interessante è il caso in cui $n = m$ e $\mathcal{V} = \mathcal{F}$. Allora la matrice in (7.1.1) è quadrata A d'ordine n :

$$A = \begin{pmatrix} | & & | \\ \mathbf{v}_1 & \dots & \mathbf{v}_n \\ | & & | \end{pmatrix} A_{\mathcal{V}, \mathcal{V}}(T_A) \begin{pmatrix} | & & | \\ \mathbf{v}_1 & \dots & \mathbf{v}_n \\ | & & | \end{pmatrix}^{-1}. \quad (7.1.2)$$

Se $A_{\mathcal{V}, \mathcal{V}}(T_A)$ è triangolare, diciamo che A è stata *triangolarizzata*, mentre se $A_{\mathcal{V}, \mathcal{V}}(T_A)$ è diagonale, diciamo che A è stata *diagonalizzata* (la stessa terminologia è stata introdotta nella Sezione 4.3). Nel seguito rivolgiamo l'attenzione al problema di trovare come diagonalizzare una matrice quadrata. Un esempio è stato già presentato nell'Esercizio 4.4.3.

7.2. Autovalori, autovettori e diagonalizzazione

Siano $\mathbf{X}, \mathbf{Y}$ spazi vettoriali della stessa dimensione n , $T : \mathbf{X} \rightarrow \mathbf{Y}$ una applicazione lineare, ed $\mathcal{V} = \{\mathbf{v}_1, \dots, \mathbf{v}_n\}$, $\mathcal{F} = \{\mathbf{f}_1, \dots, \mathbf{f}_n\}$ basi di $\mathbf{X}$ e $\mathbf{Y}$, rispettivamente. Allora la matrice associata a T rispetto ad $\mathcal{V}$ e $\mathcal{F}$ è diagonale se e solo se

$$T(\mathbf{v}_k) \text{ è un multiplo scalare } \lambda_k \mathbf{f}_k \text{ di } \mathbf{f}_k, \quad 1 \leq k \leq n, \quad (7.2.1)$$

ed in tal caso abbiamo:

$$A_{\mathcal{F}, \mathcal{V}} = \begin{pmatrix} \lambda_1 & 0 & \dots & 0 \\ 0 & \lambda_2 & \dots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \dots & 0 \end{pmatrix}.$$

Se $\mathbf{X} = \mathbf{Y}$ e $\mathcal{V} = \mathcal{F}$, allora la condizione (7.2.1) si scrive :

$$T(\mathbf{v}_k) = \lambda_k \mathbf{v}_k \text{ per un opportuno } \lambda_k \in \mathbb{R}, \quad 1 \leq k \leq n.$$

DEFINIZIONE 7.2.1. (**Autovalori, autovettori ed autospazi.**) Se T è una applicazione lineare di uno spazio vettoriale $\mathbf{X}$ in se stesso, un scalare λ si chiama *autovalore* di T se esiste $\mathbf{0}_{\mathbf{X}} \neq \mathbf{x} \in \mathbf{X}$ tale che

$$T(\mathbf{x}) = \lambda \mathbf{x}. \quad (7.2.2)$$

Tutti i vettori non nulli $\mathbf{x} \in \mathbf{X}$ che soddisfano (7.2.2) si chiamano *autovettori* di T corrispondenti all'autovalore λ ed il sottospazio lineare

$$\text{Ker}(T - \lambda \mathbb{I}_{\mathbf{X}}) = \{\mathbf{x} \in \mathbf{X}; T(\mathbf{x}) = \lambda \mathbf{x}\}$$

di $\mathbf{X}$ si chiama *l'autospazio* di T corrispondente a λ . La dimensione di un'autospazio $\text{Ker}(T - \lambda \mathbb{I}_{\mathbf{X}})$ si chiama la *molteplicità geometrica* di λ .

NOTA 7.2.2. Osserviamo che il multiplo di un autovettore per un scalare non nullo è ancora un autovettore corrispondente allo stesso autovalore:

$$T(\mathbf{x}) = \lambda \mathbf{x} \implies T(\alpha \mathbf{x}) = \lambda (\alpha \mathbf{x}).$$

□

Con queste definizioni possiamo dire che la matrice di una applicazione lineare $T : \mathbf{X} \longrightarrow \mathbf{X}$ rispetto ad una base $\mathcal{V} = \{\mathbf{v}_1, \dots, \mathbf{v}_n\}$ di $\mathbf{X}$ è diagonale se e solo se tutti i vettori $\mathbf{v}_k$ della base sono autovettori di T .

Il prossimo risultato dimostra la proprietà cruciale che autovettori corrispondenti ad autovalori diversi sono linearmente indipendenti:

TEOREMA 7.2.3. (**Indipendenza lineare degli autovettori.**) *Siano $\mathbf{X}$ uno spazio vettoriale e $T : \mathbf{X} \longrightarrow \mathbf{X}$ una applicazione lineare. Se*

$\lambda_1, \dots, \lambda_k$ sono autovalori diversi di T e

$\mathbf{x}_j$ è un autovettore di T corrispondente a λ_j , $1 \leq j \leq k$,

allora i vettori $\mathbf{x}_1, \dots, \mathbf{x}_k$ sono linearmente indipendenti.

DIMOSTRAZIONE. (Facoltativa.)

Dobbiamo verificare l'implicazione

$$\left. \begin{array}{l} \lambda_1, \dots, \lambda_k \text{ autovalori diversi di } T \\ \mathbf{0}_X \neq \mathbf{x}_j \in \text{Ker}(T - \lambda_j \mathbb{I}_X), 1 \leq j \leq k \\ \sum_{j=1}^k \alpha_j \mathbf{x}_j = \mathbf{0}_X \end{array} \right\} \implies \alpha_1 = \dots = \alpha_k = 0. \quad (\mathbb{I}_k)$$

L'implicazione $(\mathbb{I}_1)$ è chiara.

Supponiamo adesso che l'implicazione $(\mathbb{I}_k)$ non sia sempre vera e sia $k \geq 2$ il più piccolo numero naturale per il quale essa non vale. Allora esistono

autovalori diversi $\lambda_1, \dots, \lambda_k$ di T ,
 autovettori $\mathbf{x}_1 \in \text{Ker}(T - \lambda_1 \mathbb{I}_X), \dots, \mathbf{x}_k \in \text{Ker}(T - \lambda_k \mathbb{I}_X)$,
 scalari $\alpha_1, \dots, \alpha_k$ non tutti nulli

tali che

$$\sum_{j=1}^k \alpha_j \mathbf{x}_j = \mathbf{0}_X \quad \text{e quindi anche} \quad \sum_{j=1}^k \alpha_j \lambda_j \mathbf{x}_j = T\left(\sum_{j=1}^k \alpha_j \mathbf{x}_j\right) = \mathbf{0}_X.$$

Allora

$$\sum_{j=1}^{k-1} \alpha_j (\lambda_k - \lambda_j) \mathbf{x}_j = \lambda_k \sum_{j=1}^k \alpha_j \mathbf{x}_j - \sum_{j=1}^k \alpha_j \lambda_j \mathbf{x}_j = \mathbf{0}_X.$$

Per il modo in cui è stato scelto k , i $k-1$ autovettori $\mathbf{x}_1, \dots, \mathbf{x}_{k-1}$ sono linearmente indipendenti: perciò dobbiamo avere $\alpha_j (\underbrace{\lambda_k - \lambda_j}_{\neq 0}) = 0$, e

così $\alpha_j = 0$ per ogni $1 \leq j \leq k-1$. Ma allora l'uguaglianza $\sum_{j=1}^k \alpha_j \mathbf{x}_j = \mathbf{0}_X$ diventa $\alpha_k \mathbf{x}_k = \mathbf{0}_X$ e pertanto anche $\alpha_k = 0$, in contraddizione con l'ipotesi che non tutti i scalari $\alpha_1, \dots, \alpha_k$ siano nulli. $\square$

Grazie a questo Teorema 7.2.3, se $\lambda_1, \dots, \lambda_k$ sono autovalori diversi di una applicazione lineare $T : X \rightarrow X$ e, per ogni $1 \leq j \leq k$, $\mathbf{v}_1^{(j)}, \dots, \mathbf{v}_{d_j}^{(j)}$ è una base dell'autospazio $\text{Ker}(T - \lambda_j \mathbb{I}_X)$, allora

$$\underbrace{\mathbf{v}_1^{(1)}, \dots, \mathbf{v}_{d_1}^{(1)}}_{d_1}, \dots, \underbrace{\mathbf{v}_1^{(k)}, \dots, \mathbf{v}_{d_k}^{(k)}}_{d_k}$$

è un sistema linearmente indipendente e pertanto è una base del sottospazio lineare di X generato da tutti gli autovettori di T .

Infatti, se $\underbrace{\sum_{p=1}^{d_1} \alpha_p^{(1)} \mathbf{v}_p^{(1)}}_{=: \mathbf{x}_1} + \dots + \underbrace{\sum_{p=1}^{d_k} \alpha_p^{(k)} \mathbf{v}_p^{(k)}}_{=: \mathbf{x}_k} = \mathbf{0}_{\mathbf{X}}$, allora per il teo-

rema sull'indipendenza lineare degli autovettori $\mathbf{x}_1 = \dots = \mathbf{x}_k = \mathbf{0}_{\mathbf{X}}$. Siccome, per ogni $1 \leq j \leq k$, $\mathbf{v}_1^{(j)}, \dots, \mathbf{v}_{d_j}^{(j)}$ sono linearmente indipendenti, risulta che $\alpha_1^{(j)} = \dots = \alpha_{d_j}^{(j)} = 0$.

Abbiamo così dimostrato il risultato seguente:

PROPOSIZIONE 7.2.4. *Esiste una base di $\mathbf{X}$ consistente di soli autovettori di una applicazione lineare $T : \mathbf{X} \rightarrow \mathbf{X}$, ossia una base rispetto a quale la matrice di T è diagonale, se e solo se la somma delle molteplicità geometriche di tutti gli autovalori di T è uguale alla dimensione di $\mathbf{X}$. In particolare, se T ha tanti autovalori diversi quanta è la dimensione di $\mathbf{X}$, allora otteniamo una tale base scegliendo per ogni autovalore un corrispondente autovettore.*

Se A è una matrice quadrata d'ordine n , allora gli autovalori, autovettori ed autospazi di $T_A : \mathbb{R}^n \ni \mathbf{x} \mapsto A\mathbf{x} \in \mathbb{R}^n$ si chiamano rispettivamente autovalori, autovettori ed autospazi della matrice A . Per diagonalizzare A dobbiamo trovare una base di $\mathbb{R}^n$ consistente di soli autovettori di A . Questo è possibile se e solo se la somma delle molteplicità geometriche di tutti gli autovalori di A è uguale a n .

NOTA 7.2.5. (Procedimento per la diagonalizzazione.) Per diagonalizzare una matrice quadrata A di ordine n , dobbiamo prima calcolare tutti gli autovalori di A , cioè

- tutti i $\lambda \in \mathbb{R}$ per i quali il sistema omogeneo $(A - \lambda \mathbb{I}_n) \mathbf{x} = \mathbf{0}_n$ ha almeno una soluzione non nulla.

Se $\lambda_1, \dots, \lambda_k$ sono tutti gli autovalori diversi di A , allora per ogni $1 \leq j \leq k$ dobbiamo trovare una base per il relativo autospazio, cioè

- una base $\mathbf{v}_1^{(j)}, \dots, \mathbf{v}_{d_j}^{(j)}$ per lo spazio vettoriale di tutte le soluzioni del sistema omogeneo $(A - \lambda \mathbb{I}_n) \mathbf{x} = \mathbf{0}_n$.

Se $d_1 + \dots + d_k = n$, allora A è diagonalizzabile,

$$\underbrace{\mathbf{v}_1^{(1)}, \dots, \mathbf{v}_{d_1}^{(1)}}_{d_1}, \dots, \underbrace{\mathbf{v}_1^{(k)}, \dots, \mathbf{v}_{d_k}^{(k)}}_{d_k}$$

è una base di $\mathbb{R}^n$ consistente di soli autovettori di A e la formula (7.1.2) di diagonalizzazione diventa:

$$A = \left(\underbrace{\begin{array}{c|c|c} | & & | \\ \mathbf{v}_1^{(1)} & \dots & \mathbf{v}_{d_1}^{(1)} \\ | & & | \end{array}}_{d_1} \dots \underbrace{\begin{array}{c|c|c} | & & | \\ \mathbf{v}_1^{(k)} & \dots & \mathbf{v}_{d_k}^{(k)} \\ | & & | \end{array}}_{d_k} \right) \cdot \begin{pmatrix} \lambda_1 & \dots & 0 & & \\ \vdots & \ddots & \vdots & & 0 \\ 0 & \dots & \lambda_1 & & \\ & & & \ddots & \\ & & & & \lambda_k & \dots & 0 \\ & 0 & & & \vdots & \ddots & \vdots \\ & & & & 0 & \dots & \lambda_k \end{pmatrix} \cdot \left(\underbrace{\begin{array}{c|c|c} | & & | \\ \mathbf{v}_1^{(1)} & \dots & \mathbf{v}_{d_1}^{(1)} \\ | & & | \end{array}}_{d_1} \dots \underbrace{\begin{array}{c|c|c} | & & | \\ \mathbf{v}_1^{(k)} & \dots & \mathbf{v}_{d_k}^{(k)} \\ | & & | \end{array}}_{d_k} \right)^{-1}.$$

□

Il passo più difficile è trovare tutti gli autovalori di A . Questo si fa trovando i valori di λ per i quali $A - \lambda\mathbb{I}$ non è invertibile. Diamo due esempi nei quale determiniamo la perdita dell'invertibilità dapprima in maniera artigianale mediante il metodo di eliminazione di Gauss, e poi con il calcolo del determinante.

ESEMPIO 7.2.6. Diagonalizziamo le matrici

$$A = \begin{pmatrix} 2 & -1 \\ -1 & 2 \end{pmatrix}, \quad B = \begin{pmatrix} 0 & 1 \\ 0 & 0 \end{pmatrix}.$$

SVOLGIMENTO. Gli autovalori di A sono i numeri reali λ per quali il sistema omogeneo

$$\begin{pmatrix} 2-\lambda & -1 \\ -1 & 2-\lambda \end{pmatrix} \begin{pmatrix} x_1 \\ x_2 \end{pmatrix} = \begin{pmatrix} 0 \\ 0 \end{pmatrix},$$

cioè

$$\begin{cases} (2-\lambda)x_1 - x_2 = 0 \\ -x_1 + (2-\lambda)x_2 = 0 \end{cases},$$

ammette una soluzione non banale. Risolviamo questo sistema dapprima usando il metodo di eliminazione di Gauss: dai calcoli

$$\begin{array}{cc|c} 2-\lambda & -1 & 0 \\ -1 & 2-\lambda & 0 \end{array}, \quad \begin{array}{cc|c} -1 & 2-\lambda & 0 \\ 2-\lambda & -1 & 0 \end{array}$$

$$\begin{array}{cc|c} -1 & 2-\lambda & 0 \\ 0 & (2-\lambda)^2 - 1 & 0 \end{array}$$

risulta che il sistema ha soluzione non banale se e solo se $(2-\lambda)^2 - 1 = 0$, cioè $2-\lambda = \pm 1$, $\lambda = 2 \mp 1$. Pertanto gli autovalori di A sono $\lambda_1 = 1$ e $\lambda_2 = 3$.

A questo punto sappiamo già che A è diagonalizzabile: è una matrice 2×2 che ha due autovalori distinti.

Troviamo un autovettore corrispondente a $\lambda_1 = 1$: il sistema

$$\begin{array}{cc|c} -1 & 1 & 0 \\ 0 & 0 & 0 \end{array}$$

ha la soluzione non banale $x_1 = 1$, $x_2 = 1$. Poi troviamo un autovettore corrispondente anche a $\lambda_2 = 3$:

$$\begin{array}{cc|c} -1 & -1 & 0 \\ 0 & 0 & 0 \end{array}$$

ha la soluzione non banale $x_1 = 1$, $x_2 = -1$. Perciò gli autovettori

$$\begin{pmatrix} 1 \\ 1 \end{pmatrix} \text{ (per } \lambda_1 = 1), \quad \begin{pmatrix} 1 \\ -1 \end{pmatrix} \text{ (per } \lambda_2 = 3)$$

costituiscono una base di $\mathbb{R}^2$ ed una diagonalizzazione di A è

$$A = \begin{pmatrix} 1 & 1 \\ 1 & -1 \end{pmatrix} \begin{pmatrix} 1 & 0 \\ 0 & 3 \end{pmatrix} \begin{pmatrix} 1 & 1 \\ 1 & -1 \end{pmatrix}^{-1}$$

$$= \begin{pmatrix} 1 & 1 \\ 1 & -1 \end{pmatrix} \begin{pmatrix} 1 & 0 \\ 0 & 3 \end{pmatrix} \begin{pmatrix} 1/2 & 1/2 \\ 1/2 & -1/2 \end{pmatrix}. \text{ Per trovare}$$

gli autovalori di A abbiamo calcolato per quali λ la matrice $A - \lambda \mathbb{I}$ non è invertibile. Questo passaggio si poteva svolgere in un colpo solo per tutti gli autovalori, semplicemente imponendo che $\det(A - \lambda \mathbb{I}) = 0$ (come visto nel Capitolo 5). La funzione della variabili λ data da $\det(A - \lambda \mathbb{I})$ si chiama *il polinomio caratteristico* della matrice A , ed è un polinomio dello stesso grado della dimensione di A , nel caso presente di grado 2, quindi risolvibile per radicali. Le sue radici sono tutti gli autovalori. (Quando $n > 2$ non sempre si sanno trovare le radici dell'equazione, e quindi il metodo del determinante porta alla soluzione solo se si riescono a trovare per ispezione diretta abbastanza radici).

Nel caso presente si ottiene:

$$P(\lambda) = \begin{vmatrix} 2 - \lambda & -1 \\ -1 & 2 - \lambda \end{vmatrix} = (2 - \lambda)^2 - 1 = \lambda^2 - 4\lambda + 3$$

le cui radici sono appunto i due autovalori $\lambda = 1$ e $\lambda = 3$ precedentemente trovati.

In maniera analoga si procede con la matrice B . Si trova che il solo autovalore di B è 0 e che l'autospazio di B corrispondente a questo autovalore consiste dai multipli scalari del vettore $\begin{pmatrix} 1 \\ 0 \end{pmatrix}$. Perciò non esiste una base di $\mathbb{R}^2$ consistente da soli autovettori di B , ossia B non è diagonalizzabile. $\square$

ESERCIZIO 7.2.7. Si diagonalizzi la matrice

$$A = \begin{pmatrix} -1 & 1 & 0 \\ 4 & 0 & 4 \\ 0 & 1 & 1 \end{pmatrix}.$$

SOLUZIONE. Calcoliamo anche questa volta gli autovalori dapprima in maniera diretta, con il metodo di eliminazione di Gauss, e poi in maniera più rapida tramite le radici del *polinomio caratteristico* introdotto nel precedente Esempio 7.2.6.

Gli autovalori di A sono gli scalari λ per quali il sistema omogeneo

$$\begin{pmatrix} -1-\lambda & 1 & 0 \\ 4 & -\lambda & 4 \\ 0 & 1 & 1-\lambda \end{pmatrix} \begin{pmatrix} x_1 \\ x_2 \\ x_3 \end{pmatrix} = \begin{pmatrix} 0 \\ 0 \\ 0 \end{pmatrix},$$

cioè

$$\begin{cases} (-1-\lambda)x_1 + x_2 = 0 \\ 4x_1 - \lambda x_2 + 4x_3 = 0 \\ x_2 + (1-\lambda)x_3 = 0 \end{cases},$$

ammette soluzione non zero. Applichiamo il metodo di eliminazione di Gauss:

$$\begin{array}{ccc|ccc|c} -1-\lambda & 1 & 0 & 0 & 1 & -\frac{\lambda}{4} & 1 & 0 \\ 4 & -\lambda & 4 & 0 & 0 & 1 & 1-\lambda & 0 \\ 0 & 1 & 1-\lambda & 0 & -1-\lambda & 1 & 0 & 0 \\ \hline 1 & -\frac{\lambda}{4} & 1 & 0 & 1 & -\frac{\lambda}{4} & 1 & 0 \\ 0 & 1 & 1-\lambda & 0 & 0 & 1 & 1-\lambda & 0 \\ 0 & 1-\frac{\lambda}{4}-\frac{\lambda^2}{4} & 1+\lambda & 0 & 0 & 0 & \frac{9\lambda}{4}-\frac{\lambda^3}{4} & 0 \end{array},$$

mostra che il sistema ha soluzione non zero se e solo se $\frac{9\lambda}{4} - \frac{\lambda^3}{4} = 0$. Perciò gli autovalori di A sono $\lambda_1 = 0$, $\lambda_2 = 3$, $\lambda_3 = -3$.

Avremmo trovato più velocemente questi tre autovalori se avessimo calcolato le radici del polinomio caratteristico

$$\begin{vmatrix} -1-\lambda & 1 & 0 \\ 4 & -\lambda & 4 \\ 0 & 1 & 1-\lambda \end{vmatrix},$$

che, dopo varie semplificazioni, si trova essere un multiplo di $9\lambda - \lambda^3$. Si noti che questo calcolo, persino nel presente caso di dimensione piccola ($n = 3$), è laborioso. In effetti, il determinante (o meglio un suo multiplo) si calcola più facilmente proprio se si esegue l'eliminazione di Gauss! Si osservi infatti che il determinante non cambia sotto le operazioni di riga del metodo di Gauss tranne per la moltiplicazione per costanti (-1 quando si scambiano due righe adiacenti, e α quando si moltiplica una riga per la costante α): quindi, già in questo esempio a dimensione soltanto 3, il modo più agevole di calcolare il determinante è di svolgere l'eliminazione di gauss come abbiamo fatto sopra e poi

calcolare il determinante dell'ultima riduzione, quella in cui la matrice è diventata triangolare, e quindi ha per determinante il prodotto dei termini diagonali.

Abbiamo trovato i tre autovalori di A . Poiché A è una matrice 3×3 ed ha tre autovalori distinti, essa è diagonalizzabile. Per trovare una diagonalizzazione di A , troviamo un autovettore ad ogni autovalore:

- Autovettore corrispondente a $\lambda_1 = 0$:

$$\begin{array}{ccc|c} -1 & 1 & 0 & 0 \\ 4 & 0 & 4 & 0 \\ 0 & 1 & 1 & 0 \end{array}, \quad \begin{array}{ccc|c} -1 & 1 & 0 & 0 \\ 0 & 4 & 4 & 0 \\ 0 & 1 & 1 & 0 \end{array}, \quad \begin{array}{ccc|c} -1 & 1 & 0 & 0 \\ 0 & 1 & 1 & 0 \\ 0 & 0 & 0 & 0 \end{array},$$

$$\mathbf{x} = \begin{pmatrix} 1 \\ 1 \\ -1 \end{pmatrix}.$$

- Autovettore corrispondente a $\lambda_2 = 3$:

$$\begin{array}{ccc|c} -4 & 1 & 0 & 0 \\ 4 & -3 & 4 & 0 \\ 0 & 1 & -2 & 0 \end{array}, \quad \begin{array}{ccc|c} -4 & 1 & 0 & 0 \\ 0 & -2 & 4 & 0 \\ 0 & 1 & -2 & 0 \end{array}, \quad \begin{array}{ccc|c} -4 & 1 & 0 & 0 \\ 0 & 1 & -2 & 0 \\ 0 & 0 & 0 & 0 \end{array},$$

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

- Autovettore corrispondente a $\lambda_3 = -3$:

$$\begin{array}{ccc|c} 2 & 1 & 0 & 0 \\ 4 & 3 & 4 & 0 \\ 0 & 1 & 4 & 0 \end{array}, \quad \begin{array}{ccc|c} 2 & 1 & 0 & 0 \\ 0 & 1 & 4 & 0 \\ 0 & 1 & 4 & 0 \end{array}, \quad \begin{array}{ccc|c} 2 & 1 & 0 & 0 \\ 0 & 1 & 4 & 0 \\ 0 & 0 & 0 & 0 \end{array},$$

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

Abbiamo così ottenuto la diagonalizzazione

$$\begin{pmatrix} -1 & 1 & 0 \\ 4 & 0 & 4 \\ 0 & 1 & 1 \end{pmatrix} = \begin{pmatrix} 1 & 1 & 2 \\ 1 & 4 & -4 \\ -1 & 2 & 1 \end{pmatrix} \cdot \begin{pmatrix} 0 & 0 & 0 \\ 0 & 3 & 0 \\ 0 & 0 & -3 \end{pmatrix} \cdot \begin{pmatrix} 1 & 1 & 2 \\ 1 & 4 & -4 \\ -1 & 2 & 1 \end{pmatrix}^{-1}.$$

□

7.3. Ulteriori esercizi sulla diagonalizzazione sul campo $\mathbb{R}$

ESERCIZIO 7.3.1. Si dica se la matrice

$$\begin{pmatrix} 7 & 6 \\ 6 & 2 \end{pmatrix}$$

è diagonalizzabile o no e se lo è si trovi una sua diagonalizzazione.

SOLUZIONE. Il polinomio caratteristico della matrice è

$$P(\lambda) = \begin{vmatrix} 7-\lambda & 6 \\ 6 & 2-\lambda \end{vmatrix} = (7-\lambda)(2-\lambda) - 36 = \lambda^2 - 9\lambda - 22$$

e risolvendo l'equazione $P(\lambda) = 0$ risultano gli autovalori $\lambda_1 = 11$ e $\lambda_2 = -2$. Poiché $\lambda_1 \neq \lambda_2$, la matrice è diagonalizzabile.

- Autovettore corrispondente a $\lambda_1 = 11$:

$$\begin{array}{ccc} \begin{array}{cc|c} -4 & 6 & 0 \\ 6 & -9 & 0 \end{array} & , & \begin{array}{cc|c} -2 & 3 & 0 \\ 2 & -3 & 0 \end{array} & , & \begin{array}{cc|c} -2 & 3 & 0 \\ 0 & 0 & 0 \end{array} \end{array}$$

$$\mathbf{x} = \begin{pmatrix} 3 \\ 2 \end{pmatrix}.$$

- Autovettore corrispondente a $\lambda_2 = -2$:

$$\begin{array}{ccc} \begin{array}{cc|c} 9 & 6 & 0 \\ 6 & 4 & 0 \end{array} & , & \begin{array}{cc|c} 3 & 2 & 0 \\ 3 & 2 & 0 \end{array} & , & \begin{array}{cc|c} 3 & 2 & 0 \\ 0 & 0 & 0 \end{array} \end{array}$$

$$\mathbf{y} = \begin{pmatrix} -2 \\ 3 \end{pmatrix}.$$

Si ottiene la diagonalizzazione

$$\begin{aligned} \begin{pmatrix} 7 & 6 \\ 6 & 2 \end{pmatrix} &= \begin{pmatrix} 3 & -2 \\ 2 & 3 \end{pmatrix} \cdot \begin{pmatrix} 11 & 0 \\ 0 & -2 \end{pmatrix} \cdot \begin{pmatrix} 3 & -2 \\ 2 & 3 \end{pmatrix}^{-1} \\ &= \begin{pmatrix} 3 & -2 \\ 2 & 3 \end{pmatrix} \cdot \begin{pmatrix} 11 & 0 \\ 0 & -2 \end{pmatrix} \cdot \begin{pmatrix} \frac{3}{13} & \frac{2}{13} \\ -\frac{2}{13} & \frac{3}{13} \end{pmatrix}. \quad \square \end{aligned}$$

ESERCIZIO 7.3.2. Si discuta la diagonalizzazione della matrice

$$\begin{pmatrix} -1 & 1 & 3 \\ 0 & 1 & 0 \\ 3 & 1 & -1 \end{pmatrix}.$$

SOLUZIONE. Il polinomio caratteristico della matrice è

$$\begin{vmatrix} -1-\lambda & 1 & 3 \\ 0 & 1-\lambda & 0 \\ 3 & 1 & -1-\lambda \end{vmatrix},$$

che si calcola usando sviluppo rispetto alla seconda riga:

$$(1-\lambda) \begin{vmatrix} -1-\lambda & 3 \\ 3 & -1-\lambda \end{vmatrix} = (1-\lambda)(\lambda^2 + 2\lambda - 8).$$

Risultano tre autovalori diversi: 1 , -4 , 2 , quindi la matrice è diagonalizzabile. Gli autovettori corrispondenti all'autovalore λ sono le soluzioni non nulle del sistema omogeneo

$$\begin{pmatrix} -1-\lambda & 1 & 3 \\ 0 & 1-\lambda & 0 \\ 3 & 1 & -1-\lambda \end{pmatrix} \begin{pmatrix} x_1 \\ x_2 \\ x_3 \end{pmatrix} = \begin{pmatrix} 0 \\ 0 \\ 0 \end{pmatrix}$$

Si ottiene

$$\begin{pmatrix} 1 \\ -1 \\ 1 \end{pmatrix} \text{ per } 1, \quad \begin{pmatrix} 1 \\ 0 \\ -1 \end{pmatrix} \text{ per } -4, \quad \begin{pmatrix} 1 \\ 0 \\ 1 \end{pmatrix} \text{ per } 2.$$

Si ha quindi la seguente diagonalizzazione della matrice:

$$\begin{aligned}
& \begin{pmatrix} -1 & 1 & 3 \\ 0 & 1 & 0 \\ 3 & 1 & -1 \end{pmatrix} = \\
& = \begin{pmatrix} 1 & 1 & 1 \\ -1 & 0 & 0 \\ 1 & -1 & 1 \end{pmatrix} \cdot \begin{pmatrix} 1 & 0 & 0 \\ 0 & -4 & 0 \\ 0 & 0 & 2 \end{pmatrix} \cdot \begin{pmatrix} 1 & 1 & 1 \\ -1 & 0 & 0 \\ 1 & -1 & 1 \end{pmatrix}^{-1} = \\
& = \begin{pmatrix} 1 & 1 & 1 \\ -1 & 0 & 0 \\ 1 & -1 & 1 \end{pmatrix} \cdot \begin{pmatrix} 1 & 0 & 0 \\ 0 & -4 & 0 \\ 0 & 0 & 2 \end{pmatrix} \cdot \begin{pmatrix} 0 & -1 & 0 \\ 1/2 & 0 & -1/2 \\ 1/2 & 1 & 1/2 \end{pmatrix}.
\end{aligned}$$

□

ESERCIZIO 7.3.3. Si discuta la diagonalizzazione della matrice

$$\begin{pmatrix} 5 & -7 & 8 \\ -7 & 5 & 8 \\ 8 & 8 & -10 \end{pmatrix}.$$

SOLUZIONE. Il polinomio caratteristico della matrice è

$$P(\lambda) = \begin{vmatrix} 5 - \lambda & -7 & 8 \\ -7 & 5 - \lambda & 8 \\ 8 & 8 & -10 - \lambda \end{vmatrix},$$

che si calcola sottraendo la seconda riga alla prima ed osservando che

dopo questa operazione la prima riga diventa divisibile per $12 - \lambda$:

$$\begin{aligned}
 P(\lambda) &= \begin{vmatrix} 12 - \lambda & \lambda - 12 & 0 \\ -7 & 5 - \lambda & 8 \\ 8 & 8 & -10 - \lambda \end{vmatrix} = \\
 &= (12 - \lambda) \begin{vmatrix} 1 & -1 & 0 \\ -7 & 5 - \lambda & 8 \\ 8 & 8 & -10 - \lambda \end{vmatrix} = \\
 &= (12 - \lambda) \begin{vmatrix} 1 & 0 & 0 \\ -7 & -2 - \lambda & 8 \\ 8 & 16 & -10 - \lambda \end{vmatrix} = \\
 &= (12 - \lambda) \begin{vmatrix} -2 - \lambda & 8 \\ 16 & -10 - \lambda \end{vmatrix} = \\
 &= (12 - \lambda) (\lambda^2 + 12\lambda - 108).
 \end{aligned}$$

Si ottengono quindi gli autovalori $\lambda_1 = 12$ e

$$\lambda_{2,3} = -6 \pm \sqrt{36 + 108} = -6 \pm 12, \text{ cioè } \lambda_2 = 6, \lambda_3 = -18,$$

- Autovettore corrispondente a $\lambda_1 = 12$:

$$\begin{array}{ccc|c} -7 & -7 & 8 & 0 \\ -7 & -7 & 8 & 0 \\ 8 & 8 & -22 & 0 \end{array} \quad , \quad \begin{array}{ccc|c} -7 & -7 & 8 & 0 \\ 0 & 0 & 0 & 0 \\ 0 & 0 & -90 & 0 \end{array}$$

$$\mathbf{x} = \begin{pmatrix} 1 \\ -1 \\ 0 \end{pmatrix}.$$

- Autovettore corrispondente a $\lambda_2 = 6$:

$$\begin{array}{ccc|c} -1 & -7 & 8 & 0 \\ -7 & -1 & 8 & 0 \\ 8 & 8 & -16 & 0 \end{array} \quad , \quad \begin{array}{ccc|c} -1 & -7 & 8 & 0 \\ -7 & -1 & 8 & 0 \\ 1 & 1 & -2 & 0 \end{array}$$

$$\begin{array}{ccc|c} -1 & -7 & 8 & 0 \\ 0 & 48 & -48 & 0 \\ 0 & -6 & 6 & 0 \end{array} \quad , \quad \begin{array}{ccc|c} -1 & -7 & 8 & 0 \\ 0 & 1 & -1 & 0 \\ 0 & 0 & 0 & 0 \end{array}$$

$$\mathbf{y} = \begin{pmatrix} 1 \\ 1 \\ 1 \end{pmatrix}.$$

- Autovettore corrispondente a $\lambda_3 = -18$:

$$\begin{array}{ccc|c}, & 1 & 1 & 1 & | & 0 \\ 23 & -7 & 8 & | & 0 \\ -7 & 23 & 8 & | & 0 \\ 8 & 8 & 8 & | & 0 \end{array}, \quad \begin{array}{ccc|c}, & 1 & 1 & 1 & | & 0 \\ 1 & 1 & 1 & | & 0 \\ 0 & -30 & -15 & | & 0 \\ 0 & 30 & 15 & | & 0 \end{array},$$

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

Si ottiene la diagonalizzazione

$$\begin{aligned} & \begin{pmatrix} 5 & -7 & 8 \\ -7 & 5 & 8 \\ 8 & 8 & -10 \end{pmatrix} = \\ & = \begin{pmatrix} 1 & 1 & 1 \\ -1 & 1 & 1 \\ 0 & 1 & -2 \end{pmatrix} \cdot \begin{pmatrix} 12 & 0 & 0 \\ 0 & 6 & 0 \\ 0 & 0 & -18 \end{pmatrix} \cdot \begin{pmatrix} 1 & 1 & 1 \\ -1 & 1 & 1 \\ 0 & 1 & -2 \end{pmatrix}^{-1} = \\ & = \begin{pmatrix} 1 & 1 & 1 \\ -1 & 1 & 1 \\ 0 & 1 & -2 \end{pmatrix} \cdot \begin{pmatrix} 12 & 0 & 0 \\ 0 & 6 & 0 \\ 0 & 0 & -18 \end{pmatrix} \cdot \begin{pmatrix} 1/2 & -1/2 & 0 \\ 1/3 & 1/3 & 1/3 \\ 1/6 & 1/6 & -1/3 \end{pmatrix}. \end{aligned}$$

□

ESERCIZIO 7.3.4. Si discuta la diagonalizzazione della matrice

$$\begin{pmatrix} 16 & -12 & 0 \\ -12 & 9 & 0 \\ 0 & 0 & 25 \end{pmatrix}.$$

SOLUZIONE. Anzitutto, la matrice è simmetrica e perciò certamente diagonalizzabile, come dimostreremo in un prossimo teorema

Il polinomio caratteristico della matrice è

$$\begin{aligned} P(\lambda) &= \begin{vmatrix} 16-\lambda & -12 & 0 \\ -12 & 9-\lambda & 0 \\ 0 & 0 & 25-\lambda \end{vmatrix} = (25-\lambda) \begin{vmatrix} 16-\lambda & -12 \\ -12 & 9-\lambda \end{vmatrix} = \\ &= (25-\lambda) ((16-\lambda)(9-\lambda) - 144) = -\lambda(\lambda-25)^2, \end{aligned}$$

perciò gli autovalori della matrice sono $\lambda_1 = 0$ e $\lambda_2 = 25$.

• Autovettori corrispondenti a $\lambda_1 = 0$:

$$\begin{array}{ccc|c} 16 & -12 & 0 & 0 \\ -12 & 9 & 0 & 0 \\ 0 & 0 & 25 & 0 \end{array}, \quad \begin{array}{ccc|c} 4 & -3 & 0 & 0 \\ -4 & 3 & 0 & 0 \\ 0 & 0 & 1 & 0 \end{array}, \quad \begin{array}{ccc|c} 4 & -3 & 0 & 0 \\ 0 & 0 & 0 & 0 \\ 0 & 0 & 1 & 0 \end{array},$$

$$\mathbf{x} = \begin{pmatrix} 3 \\ 4 \\ 0 \end{pmatrix}.$$

• Autovettori corrispondenti a $\lambda_2 = 25$: dai calcoli

$$\begin{array}{ccc|c} -9 & -12 & 0 & 0 \\ -12 & -16 & 0 & 0 \\ 0 & 0 & 0 & 0 \end{array}, \quad \begin{array}{ccc|c} -3 & -4 & 0 & 0 \\ -3 & -4 & 0 & 0 \\ 0 & 0 & 0 & 0 \end{array}, \quad \begin{array}{ccc|c} 3 & 4 & 0 & 0 \\ 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 \end{array}$$

risulta che la forma generale di questi autovettori è

$$\begin{pmatrix} -\frac{4}{3}x_2 \\ x_2 \\ x_3 \end{pmatrix} = \frac{1}{3}x_2 \begin{pmatrix} -4 \\ 3 \\ 0 \end{pmatrix} + x_3 \begin{pmatrix} 0 \\ 0 \\ 1 \end{pmatrix},$$

perciò abbiamo due autovettori linearmente indipendenti corrispondenti a $\lambda_2 = 25$:

$$\mathbf{y} = \begin{pmatrix} -4 \\ 3 \\ 0 \end{pmatrix}, \quad \mathbf{z} = \begin{pmatrix} 0 \\ 0 \\ 1 \end{pmatrix}.$$

Si ottiene la diagonalizzazione:

$$\begin{aligned}
 & \begin{pmatrix} 16 & -12 & 0 \\ -12 & 9 & 0 \\ 0 & 0 & 25 \end{pmatrix} \\
 &= \begin{pmatrix} 3 & -4 & 0 \\ 4 & 3 & 0 \\ 0 & 0 & 1 \end{pmatrix} \cdot \begin{pmatrix} 0 & 0 & 0 \\ 0 & 25 & 0 \\ 0 & 0 & 25 \end{pmatrix} \cdot \begin{pmatrix} 3 & -4 & 0 \\ 4 & 3 & 0 \\ 0 & 0 & 1 \end{pmatrix}^{-1} = \\
 &= \begin{pmatrix} 3 & -4 & 0 \\ 4 & 3 & 0 \\ 0 & 0 & 1 \end{pmatrix} \cdot \begin{pmatrix} 0 & 0 & 0 \\ 0 & 25 & 0 \\ 0 & 0 & 25 \end{pmatrix} \cdot \begin{pmatrix} 3/25 & 4/25 & 0 \\ -4/25 & 3/25 & 0 \\ 0 & 0 & 1 \end{pmatrix}.
 \end{aligned}$$

□

ESERCIZIO 7.3.5. Si discuta la diagonalizzabilità della matrice

$$\begin{pmatrix} 3 & 1 & 2 \\ -2 & 0 & -2 \\ -1 & 0 & -1 \end{pmatrix}.$$

SOLUZIONE. Il polinomio caratteristico della matrice è

$$P(\lambda) = \begin{vmatrix} 3-\lambda & 1 & 2 \\ -2 & -\lambda & -2 \\ -1 & 0 & -1-\lambda \end{vmatrix},$$

che si calcola sommando la seconda riga alla prima ed osservando che dopo questa operazione la prima riga diventa divisibile per $1-\lambda$:

$$\begin{aligned}
 P(\lambda) &= \begin{vmatrix} 1-\lambda & 1-\lambda & 0 \\ -2 & -\lambda & -2 \\ -1 & 0 & -1-\lambda \end{vmatrix} = (1-\lambda) \begin{vmatrix} 1 & 1 & 0 \\ -2 & -\lambda & -2 \\ -1 & 0 & -1-\lambda \end{vmatrix} = \\
 &= (1-\lambda) \begin{vmatrix} 1 & 0 & 0 \\ -2 & 2-\lambda & -2 \\ -1 & 1 & -1-\lambda \end{vmatrix} = (1-\lambda) \begin{vmatrix} 2-\lambda & -2 \\ 1 & -1-\lambda \end{vmatrix} = \\
 &= (1-\lambda)(\lambda^2 - \lambda) = -\lambda(\lambda - 1)^2.
 \end{aligned}$$

Pertanto si trovano gli autovalori $\lambda_1 = 0$, $\lambda_2 = 1$.

Risolvendo il sistema omogeneo

$$\begin{pmatrix} 3-\lambda & 1 & 2 \\ -2 & -\lambda & -2 \\ -1 & 0 & -1-\lambda \end{pmatrix} \begin{pmatrix} x_1 \\ x_2 \\ x_3 \end{pmatrix} = \begin{pmatrix} 0 \\ 0 \\ 0 \end{pmatrix}$$

per $\lambda = 0$, si trovano come soluzioni i multipli scalari del vettore

$$\begin{pmatrix} -1 \\ 1 \\ 1 \end{pmatrix},$$

mentre risolvendolo per $\lambda = 1$, si trovano i multipli scalari del vettore

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

Di conseguenza il sottospazio lineare di $\mathbb{R}^3$ generato da tutti gli autovettori della matrice è il sottospazio generato dai due vettori di cui sopra e pertanto non è uguale a $\mathbb{R}^3$. Concludiamo che la matrice non è diagonalizzabile. $\square$

ESERCIZIO 7.3.6. Si trovino gli autovalori della matrice

$$\begin{pmatrix} 5 & -2 & 0 & 1 \\ -2 & 5 & 1 & 0 \\ 0 & 1 & 5 & -2 \\ 1 & 0 & -2 & 5 \end{pmatrix}.$$

SOLUZIONE. Il polinomio caratteristico della matrice è

$$P(\lambda) = \begin{vmatrix} 5-\lambda & -2 & 0 & 1 \\ -2 & 5-\lambda & 1 & 0 \\ 0 & 1 & 5-\lambda & -2 \\ 1 & 0 & -2 & 5-\lambda \end{vmatrix},$$

che si comincia a calcolare sommando la seconda, terza e quarta riga alla prima ed osservando che dopo questa operazione la prima riga diventa divisibile per $4-\lambda$:

$$\begin{aligned}
P(\lambda) &= \begin{vmatrix} 4-\lambda & 4-\lambda & 4-\lambda & 4-\lambda \\ -2 & 5-\lambda & 1 & 0 \\ 0 & 1 & 5-\lambda & -2 \\ 1 & 0 & -2 & 5-\lambda \end{vmatrix} = \\
&= (4-\lambda) \begin{vmatrix} 1 & 1 & 1 & 1 \\ -2 & 5-\lambda & 1 & 0 \\ 0 & 1 & 5-\lambda & -2 \\ 1 & 0 & -2 & 5-\lambda \end{vmatrix} = \\
&= (4-\lambda) \begin{vmatrix} 1 & 1 & 1 & 1 \\ 0 & 7-\lambda & 3 & 2 \\ 0 & 1 & 5-\lambda & -2 \\ 0 & -1 & -3 & 4-\lambda \end{vmatrix} = \\
&= (4-\lambda) \begin{vmatrix} 7-\lambda & 3 & 2 \\ 1 & 5-\lambda & -2 \\ -1 & -3 & 4-\lambda \end{vmatrix}.
\end{aligned}$$

Ora si sottraggono alla prima colonna la seconda e la terza e si osserva che dopo questa operazione la prima colonna diventa divisibile per $2-\lambda$:

$$\begin{aligned}
P(\lambda) &= (4-\lambda) \begin{vmatrix} 2-\lambda & 3 & 2 \\ -2+\lambda & 5-\lambda & -2 \\ -2+\lambda & -3 & 4-\lambda \end{vmatrix} = \\
&= (4-\lambda)(2-\lambda) \begin{vmatrix} 1 & 3 & 2 \\ -1 & 5-\lambda & -2 \\ -1 & -3 & 4-\lambda \end{vmatrix} = \\
&= (4-\lambda)(2-\lambda) \begin{vmatrix} 1 & 3 & 2 \\ 0 & 8-\lambda & 0 \\ 0 & 0 & 6-\lambda \end{vmatrix} = \\
&= (4-\lambda)(2-\lambda)(8-\lambda)(6-\lambda).
\end{aligned}$$

Risultano gli autovalori $\lambda_1 = 2$, $\lambda_2 = 4$, $\lambda_3 = 6$, $\lambda_4 = 8$.

□

ESERCIZIO 7.3.7. Consideriamo la matrice

$$A = \begin{pmatrix} 2 & 6 & -9 \\ -1 & -5 & 9 \\ 0 & p & 8 \end{pmatrix},$$

dove p è un parametro reale.

- (i) Si calcoli il polinomio caratteristico di A .
- (ii) Per che valori di p tutti gli autovalori di A sono reali?
- (iii) Per che valori di p tutti gli autovalori di A sono reali e distinti?
- (iv) Per che valori di p la matrice A è diagonalizzabile?

SOLUZIONE.

- (i) Il polinomio caratteristico di A è

$$\begin{aligned} \begin{vmatrix} 2-\lambda & 6 & -9 \\ -1 & -5-\lambda & 9 \\ 0 & p & 8-\lambda \end{vmatrix} &= \begin{vmatrix} 1-\lambda & 1-\lambda & 0 \\ -1 & -5-\lambda & 9 \\ 0 & p & 8-\lambda \end{vmatrix} \\ &= (1-\lambda) \begin{vmatrix} 1 & 1 & 0 \\ -1 & -5-\lambda & 9 \\ 0 & p & 8-\lambda \end{vmatrix} \\ &= (1-\lambda) \begin{vmatrix} 1 & 1 & 0 \\ 0 & -4-\lambda & 9 \\ 0 & p & 8-\lambda \end{vmatrix} \\ &= (1-\lambda) (\lambda^2 - 4\lambda - 32 - 9p). \end{aligned}$$

- (ii) Gli autovalori di A sono $\lambda_1 = 1$ e $\lambda_{2,3} = 2 \pm 3\sqrt{4+p}$, quindi sono tutti reali se e soltanto se $p \geq -4$.

- (iii) Risulta

$$\lambda_2 = \lambda_3 \quad \text{per} \quad p = -4$$

e

$$\lambda_1 = \lambda_3 \quad \text{per} \quad p = -4 + \frac{1}{9} = -\frac{35}{9},$$

perciò A risulta avere tutti gli autovalori reali e distinti esattamente per $-\frac{35}{9} \neq p > -4$.

- (iv) Calcolando tutti gli autovettori esplicitamente nei casi $p = -4$ (autovalori 1 e 2) e $p = -\frac{35}{9}$ (autovalori 1 e 2), si trova che essi non generano $\mathbb{R}^3$ in nessuno di questi casi. Sapendo poi che

A è diagonalizzabile quando tutti gli autovalori sono reali e distinti, risulta che A è diagonalizzabile esattamente in questo caso, cioè quando $-\frac{35}{9} \neq p > -4$.

□

ESERCIZIO 7.3.8. Per che valori reali a la matrice

$$\begin{pmatrix} 1 & 1 & 0 \\ 0 & 2 & a \\ 0 & 0 & 2 \end{pmatrix}$$

è diagonalizzabile? Nei casi in cui lo è si trovi una diagonalizzazione della matrice.

SOLUZIONE. Il polinomio caratteristico della matrice si calcola facilmente, perché è il determinante di una matrice triangolare, e quindi è il prodotto dei termini diagonali:

$$\begin{vmatrix} 1-\lambda & 1 & 0 \\ 0 & 2-\lambda & a \\ 0 & 0 & 2-\lambda \end{vmatrix} = (1-\lambda)(2-\lambda)^2.$$

Abbiamo quindi gli autovalori $\lambda_1 = 1$, $\lambda_2 = 2$. Ora calcoliamo gli autovettori.

Troviamo il nucleo di $A - \lambda\mathbb{I}$ per $\lambda_1 = 1$:

$$\begin{pmatrix} 0 & 1 & 0 \\ 0 & 1 & a \\ 0 & 0 & 1 \end{pmatrix} \begin{pmatrix} x_1 \\ x_2 \\ x_3 \end{pmatrix} = \begin{pmatrix} 0 \\ 0 \\ 0 \end{pmatrix} \implies x_1 \text{ arbitrario, } x_2 = x_3 = 0,$$

quindi gli autovettori corrispondenti a $\lambda_1 = 1$ sono i multipli non nulli di

$$\mathbf{x} = \begin{pmatrix} 1 \\ 0 \\ 0 \end{pmatrix}.$$

Si osservi che x_1 è arbitrario perché il sistema $(A - \mathbb{I})\mathbf{x} = \mathbf{0}$ non impone su x_1 alcuna condizione, perché la matrice dei coefficienti, $A - \mathbb{I}$, ha la prima *colonna* nulla.

Per $\lambda_2 = 2$:

$$\begin{pmatrix} -1 & 1 & 0 \\ 0 & 0 & a \\ 0 & 0 & 0 \end{pmatrix} \begin{pmatrix} x_1 \\ x_2 \\ x_3 \end{pmatrix} = \begin{pmatrix} 0 \\ 0 \\ 0 \end{pmatrix} \implies x_1 - x_2 = 0, \quad ax_3 = 0,$$

quindi per la terza componente si trova

$$x_3 = \begin{cases} 0, & \text{se } a \neq 0 \\ \text{arbitrario}, & \text{se } a = 0 \end{cases}.$$

Si osservi che, anche qui, x_3 è indeterminato quando $a = 0$ perché *la terza colonna* è nulla, non perché lo sia *la terza riga*. Il fatto che la terza riga sia zero significa semplicemente che la terza riga non impone alcuna condizione a nessuna variabile, non solo alla terza.

Pertanto, nel caso di $a \neq 0$ gli autovettori corrispondenti a $\lambda_2 = 2$ sono i multipli non nulli di

$$\mathbf{y} = \begin{pmatrix} 1 \\ 1 \\ 0 \end{pmatrix},$$

mentre nel caso $a = 0$ possiamo trovare due autovettori linearmente indipendenti:

$$\mathbf{y} = \begin{pmatrix} 1 \\ 1 \\ 0 \end{pmatrix}, \quad \mathbf{z} = \begin{pmatrix} 1 \\ 1 \\ 1 \end{pmatrix}.$$

Di conseguenza, per $a \neq 0$ gli autovettori della matrice data generano un sottospazio di dimensione 2 di $\mathbb{R}^3$ e quindi la matrice non è diagonalizzabile. Per $a = 0$ si ottiene invece la diagonalizzazione

$$\begin{aligned} & \begin{pmatrix} 1 & 1 & 0 \\ 0 & 2 & 0 \\ 0 & 0 & 2 \end{pmatrix} = \\ &= \begin{pmatrix} 1 & 1 & 1 \\ 0 & 1 & 1 \\ 0 & 0 & 1 \end{pmatrix} \cdot \begin{pmatrix} 1 & 0 & 0 \\ 0 & 2 & 0 \\ 0 & 0 & 2 \end{pmatrix} \cdot \begin{pmatrix} 1 & 1 & 1 \\ 0 & 1 & 1 \\ 0 & 0 & 1 \end{pmatrix}^{-1} = \\ &= \begin{pmatrix} 1 & 1 & 1 \\ 0 & 1 & 1 \\ 0 & 0 & 1 \end{pmatrix} \cdot \begin{pmatrix} 1 & 0 & 0 \\ 0 & 2 & 0 \\ 0 & 0 & 2 \end{pmatrix} \cdot \begin{pmatrix} 1 & -1 & 0 \\ 0 & 1 & -1 \\ 0 & 0 & 1 \end{pmatrix}. \quad \square \end{aligned}$$

NOTA 7.3.9. Osserviamo che nel precedente Esercizio 7.3.8 la diagonalizzabilità si perde quando $a \neq 0$, come conseguenza del coefficiente a non nullo sopra la diagonale nel blocco diagonale 2×2 che corrisponde all'autovalore 2, non del coefficiente 1 che si trova sopra la diagonale, il quale non pregiudica affatto la diagonalizzabilità nel caso in cui a è nullo. Questo interessante esempio indica che *la impossibilità di diagonalizzare si collega al fatto che le matrici si possano portare a forma triangolare nella quale sopra la diagonale ci siano termini non nulli nei blocchi corrispondenti ad autovalori ripetuti*. Qui ci limitiamo ad osservare che una matrice con coefficienti uguali sulla diagonale, diciamo λ , e zeri al di sotto (o al di sopra) non può essere diagonalizzabile, perché se lo fosse la forma diagonale sarebbe $\lambda \mathbb{I}$, ma la matrice identità rimane la stessa in qualunque base, perché manda ogni vettore di ogni base in sé stesso, e la stessa cosa ovviamente accade ai multipli dell'identità. Una situazione analoga si ha se una matrice ha un blocco triangolare (diciamo superiore) con coefficienti uguali sulla diagonale e qualche numero non nullo al di sopra. La parte non banale della riduzione a forma triangolare di questo tipo (*forma canonica di Jordan*) consiste nel dimostrare che tutte le matrici si possono ridurre a questa forma. Questo fatto verrà approfondito e dimostrato nella Sezione 9.2 sulla forma canonica di Jordan. $\square$

ESERCIZIO 7.3.10. Per che valori reali a la matrice

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

è diagonalizzabile? Si trovino gli autovettori della matrice.

SOLUZIONE. La matrice è triangolare, e quindi gli autovalori sono i termini diagonali:

$$\begin{vmatrix} 1 - \lambda & 1 & 0 \\ 0 & 2 - \lambda & a \\ 0 & 0 & 2 - \lambda \end{vmatrix} = (1 - \lambda)(2 - \lambda)^2.$$

Abbiamo quindi gli autovalori $\lambda_1 = 1$, $\lambda_2 = a$, $\lambda_3 = -a$. Per $a \neq 0$ i tre autovalori sono diversi e quindi la matrice è diagonalizzabile.

Invece per $a = 0, -1$ e 1 due dei tre autovalori sono uguali. Per $a = 0$ non c'è dubbio sulla diagonalizzabilità perché la matrice è diagonale. Dall'esempio precedente ci accorgiamo che la matrice è diagonalizzabile per $a = 1$, perché il blocco dei due coefficienti diagonale uguali non ha un numero non nullo sopra la diagonale. Invece questo succede per $a = -1$, e quindi in tale caso il risultato annunciato prima sulla forma canonica di Jordan esclude la diagonalizzabilità. Ora calcoliamo gli autovettori.

Poiché la matrice è diagonale per $a = 0$, d'ora in poi assumiamo $a \neq 0$. Troviamo il nucleo di $A - \lambda \mathbb{I}$ per $\lambda = 1$:

$$\begin{pmatrix} a-1 & 0 & 0 \\ 0 & 0 & a \\ 0 & 0 & -(1+a) \end{pmatrix} \begin{pmatrix} x_1 \\ x_2 \\ x_3 \end{pmatrix} = \begin{pmatrix} 0 \\ 0 \\ 0 \end{pmatrix}$$

$$\Rightarrow \begin{cases} \text{se } a = 1, & x_1, x_2 \text{ arbitrari, } x_3 = 0 \\ \text{se } a = -1, & x_2 \text{ arbitrario, } x_1 = x_3 = 0 \\ \text{se } a \neq 1, -1, & x_2 \text{ arbitrario, } x_1 = x_3 = 0 \end{cases}.$$

Quindi per $a = 1$ l'autovalore $\lambda = 1$ ha molteplicità 2 ed il suo autospazio ha dimensione 2 (è generato da

$$\mathbf{e}_1 = \begin{pmatrix} 1 \\ 0 \\ 0 \end{pmatrix} \quad \text{e} \quad \mathbf{e}_2 = \begin{pmatrix} 0 \\ 1 \\ 0 \end{pmatrix}.$$

Aggiungendo a questi due autovettori l'autovettore relativo all'altro autovalore $\lambda = -a = -1$, che risulta essere $\mathbf{e}_2 - \mathbf{e}_3$, si ottiene una base di autovettori, e quindi la diagonalizzabilità. Invece per $a = -1$ si ha ancora autovalore 1 con molteplicità 2, ma l'autospazio è solo unidimensionale, generato da

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

e quindi si perde la diagonalizzabilità. In maniera analoga, si trova che, per $a \neq 1, -1$ gli autovettori sono: $\mathbf{e}_2$ per $\lambda = 1$, $\mathbf{e}_1$ per $\lambda = a$, $a\mathbf{e}_2 + (1-a)\mathbf{e}_3$ per $\lambda = -a$ (il lettore è invitato a verificare per esercizio).

Ovviamente per tutti questi valori di a i tre autovettori trovati sono una base e la matrice è diagonalizzabile. $\square$

ESERCIZIO 7.3.11. Per che valore reale a la matrice

$$\begin{pmatrix} a & a-1 & 0 \\ 0 & a^2 & a \\ 0 & 0 & 1 \end{pmatrix}$$

è diagonalizzabile? Nei casi in cui lo è si trovi una diagonalizzazione della matrice. $\square$

ESERCIZIO 7.3.12. Per che valori reali a la matrice

$$\begin{pmatrix} 1 & a & 0 & 0 \\ a & 1 & 0 & 0 \\ 0 & 0 & 2a-1 & a \\ 0 & 0 & 0 & a^2 \end{pmatrix}$$

è diagonalizzabile? Nei casi in cui lo è si trovi una diagonalizzazione della matrice.

SUGGERIMENTO. Ci limitiamo ad osservare che la matrice si decompone in due blocchi 2×2 e quindi il determinante si spezza come prodotto. Il primo blocco porta ad autovalori che sono soluzione dell'equazione $(\lambda-a)^2-1=0$, e cioè $\lambda = a \pm 1$. Il secondo blocco è triangolare e i suoi autovalori sono i coefficienti diagonali $2a-1$ e a^2 . Si osservi che questi due autovalori sono uguali solo $a^2-2a+1=(a-1)^2=0$, cioè solo se $a=1$. Quando questi due autovalori sono uguali il secondo blocco non è diagonalizzabile, a causa del coefficiente a sopra la diagonale, che non è nullo in quanto $a=1$. Bisogna però osservare che l'autovalore 1 viene anche dall'altro blocco, e quindi ha molteplicità 3: l'autospazio potrebbe avere uno o due autovettori indipendenti, ma è facile vedere che ne ha due, per il seguente fatto interessante, la cui verifica lasciamo al lettore:

NOTA 7.3.13. Se una matrice A di dimensione n si decompone a blocchi, nel senso che per $i=1, \dots, k$ esistono matrici A_i di dimensione n_i con $\sum_{i=1}^k n_i = n$ ed i coefficienti di A sono nulli al di fuori di k blocchi quadrati disgiunti disposti sulla diagonale ed uguali alle sottomatrici

A_i , allora gli autovalori di A_i sono chiaramente anche autovettori di A (per la diagonalità dei blocchi e la corrispondente fattorizzazione del determinante, ed ogni autovettore $\mathbf{x}$ di A_i in $\mathbb{C}^{n_i}$ induce un autovettore $\bar{\mathbf{x}}$ di A in $\mathbb{C}^n$ nel modo seguente:

$$\bar{\mathbf{x}} = \begin{pmatrix} \mathbf{0}_{\mathbb{C}^{n_1}} \\ \mathbf{0}_{\mathbb{C}^{n_2}} \\ \vdots \\ \mathbf{0}_{\mathbb{C}^{n_{i-1}}} \\ \mathbf{x} \\ \mathbf{0}_{\mathbb{C}^{n_{i+1}}} \\ \vdots \\ \mathbf{0}_{\mathbb{C}^{n_k}} \end{pmatrix}.$$

È chiaro che autovettori provenienti da blocchi diversi sono ortogonali, così come autovettori indotti da autovalori diversi dello stesso blocco sono linearmente indipendenti. Per un trattamento di esempi di questo tipo sulla base di una notazione più adeguata rinviamo il lettore al capitolo 8 sulla somma diretta di spazi vettoriali. $\square$

$\square$

ESERCIZIO 7.3.14. Per che valori reali a la matrice

$$\begin{pmatrix} 1 & a & 0 & 0 \\ 0 & 1 & 0 & 0 \\ 0 & 0 & 1+a & a \\ 0 & 0 & 0 & a^2 \end{pmatrix}$$

è diagonalizzabile? Nei casi in cui lo è si trovi una diagonalizzazione della matrice.

SUGGERIMENTO. Rispetto al precedente Esercizio 7.3.12 ora la matrice è triangolare. È ovvio che gli autovalori sono $1, 1, 1+a, a$, e l'autovalore 1 ed il primo blocco diagonale, con autovalore 1 di molteplicità 2 , porta solo un autovettore se $a \neq 0$ a causa del coefficiente non nullo a sopra la diagonale. Invece se $a = 0$ questo blocco è diagonale e

dà luogo, come nel succitato Esercizio, a due autovettori (che ovviamente sono $\mathbf{e}_1$ ed $\mathbf{e}_2$). In realtà se $a = 0$ allora uno degli autovalori del secondo blocco, $1 + a$, vale anch'esso 1, ma siccome il corrispondente autovettore è indotto da un altro blocco (si veda la Nota 7.3.13), allora esso è ortogonale al precedente, e l'autovalore 1 ha molteplicità 3 ma autospazio bidimensionale. $\square$

7.4. Autovalori complessi e diagonalizzazione in $M_n^{\mathbb{C}}$

Il polinomio caratteristico di ogni matrice $A \in M_n^{\mathbb{C}}$ è un polinomio di grado n a coefficienti complessi. Un celebre risultato, il Teorema Fondamentale dell'Algebra, la cui dimostrazione qui viene omessa, afferma che ogni polinomio complesso di grado n ha sempre n radici complesse, se le si conta con la loro molteplicità. Quindi, a differenza del caso reale, ogni matrice complessa di grado n ha sempre n autovalori complessi, ma alcuni possono essere ripetuti. Se gli autovalori sono distinti, anche se magari non reali, i corrispondenti autovettori sono indipendenti (il Teorema di indipendenza 7.2.3 continua a valere nel caso complesso, con la stessa dimostrazione). Se invece un autovalore è ripetuto, diciamo con molteplicità m , allora non è detto che il suo autospazio contenga m autovettori linearmente indipendenti, e se questo non accade non si può avere una base di autovettori e la matrice non è diagonalizzabile. (Per avere esempi, si veda in seguito il Capitolo 9).

Si osservi che, se la matrice A è reale, nondimeno il suo polinomio caratteristico, che è a coefficienti reali, può avere qualche radice λ complessa. In tal caso questi autovalori complessi devono avere autovettori con qualche componente complessa, altrimenti l'equazione $A\mathbf{x} = \lambda\mathbf{x}$ non potrebbe valere per $\mathbf{x} \neq \mathbf{00}$. Pertanto A non è diagonalizzabile in $M_n^{\mathbb{R}}$ (si dice che A non è *diagonalizzabile sui reali*). Però potrebbe esserlo in $M_n^{\mathbb{C}}$. Ecco un esempio:

ESEMPIO 7.4.1. (*Diagonalizzazione sui complessi di matrici di*

rotazione.) Come nell'Esempio 6.8.5, sia A_θ la matrice reale di rotazione 2×2 data da

$$A_\theta = \begin{pmatrix} \cos \theta & \sin \theta \\ -\sin \theta & \cos \theta \end{pmatrix}.$$

È ovvio che, a meno che $\theta = 0$ o π (nei quali casi $A_\theta = \pm \mathbb{I}$ lascia invariante la direzione di ogni vettore nel piano), A_θ non può avere autovettori reali, perché ogni vettore non nullo in $\mathbb{R}^2$ viene ruotato di un angolo diverso da 0 e da π e quindi non viene mandato in un multiplo di se stesso. Vedremo infatti che la matrice ha autovettori complessi invece che reali, e quindi non è diagonalizzabile sui reali, ma anche che i due autovettori, complessi coniugati, sono diversi (certo! se no sarebbero reali!), e quindi A_θ è *diagonalizzabile sui complessi*.

In effetti, si ha

$$\begin{aligned} \det(A_\theta - \lambda \mathbb{I}) &= \begin{vmatrix} \cos \theta - \lambda & \sin \theta \\ -\sin \theta & \cos \theta - \lambda \end{vmatrix} \\ &= (\cos \theta - \lambda)^2 + \sin^2 \theta = \lambda^2 - 2\lambda \cos \theta + 1, \end{aligned}$$

le cui radici sono $\lambda_\pm = \cos \theta \pm i \sin \theta = e^{\pm i\theta}$. La forma diagonale di A_θ , su $\mathbb{C}$, è quindi

$$\begin{pmatrix} e^{i\theta} & 0 \\ 0 & e^{-i\theta} \end{pmatrix}.$$

Gli autovettori complessi $\mathbf{x}_\pm$ sono i vettori nel nucleo di $A - \lambda_\pm \mathbb{I}$: il metodo di eliminazione di Gauss ci permette di trovarli. Lasciamo i calcoli al lettore. $\square$

7.5. Diagonalizzabilità di matrici simmetriche o autoaggiunte

La definizione di matrici simmetriche o autoaggiunte è stata data in 6.4.6.

TEOREMA 7.5.1. (Autovalori ed autovettori di matrici simmetriche o autoaggiunte.) *Sia A una matrice simmetrica (o più in generale autoaggiunta). Allora*

- (i) *Tutti gli autovalori di A sono numeri reali.*

- (ii) *Autovettori di A corrispondenti ad autovalori diversi sono ortogonali.*
- (iii) *Esiste almeno un autovalore reale di A .*

DIMOSTRAZIONE.

- (i) Se la matrice A è autoaggiunta essa agisce sullo spazio complesso $\mathbb{C}^n$; se è simmetrica, agisce sullo spazio reale $\mathbb{R}^3$ ma più in generale si può considerare la sua azione sullo spazio complesso $\mathbb{C}^3$. In entrambi i casi i suoi autovalori, a priori, potrebbero essere complessi (ovviamente, quando A è una matrice a coefficienti reali, se essa ha un autovalore complesso il corrispondente autovettore deve essere a coefficienti complessi, e quindi non esiste in $\mathbb{R}^3$, ma - al più - in $\mathbb{C}^3$). Nella dimostrazione consideriamo il caso più generale dell'azione della matrice autoaggiunta A sullo spazio complesso $\mathbb{C}^n$, munito quindi del prodotto scalare sui complessi

$$\langle \mathbf{x}, \mathbf{y} \rangle = \sum_{i=1}^n x_i \overline{y_i}$$

introdotto in (6.3.2).

Sia $\mathbf{x}$ un autovettore di A con autovalore λ . Poiché A è a coefficienti reali abbiamo $A^* = A^T = A$, e quindi

$$A\mathbf{x} \cdot \mathbf{x} = \mathbf{x} \cdot A\mathbf{x}$$

grazie a (6.4.2).

Pertanto

$$\lambda \mathbf{x} \cdot \mathbf{x} = A\mathbf{x} \cdot \mathbf{x} = \mathbf{x} \cdot A\mathbf{x} = \overline{\lambda} \mathbf{x} \cdot \mathbf{x}.$$

Dal momento che l'autovettore $\mathbf{x}$ deve essere non nullo, la sua norma $\|\mathbf{x}\| = \mathbf{x} \cdot \mathbf{x}$ è diversa da zero, e quindi l'identità precedente implica $\lambda \in \mathbb{R}$.

- (ii) Siano $\mathbf{x}_1, \mathbf{x}_2$ autovettori di A con rispettivi autovalori $\lambda_1 \neq \lambda_2$. Allora

$$\lambda_1 \mathbf{x}_1 \cdot \mathbf{x}_2 = A\mathbf{x}_1 \cdot \mathbf{x}_2 = \mathbf{x}_1 \cdot A\mathbf{x}_2 = \lambda_2 \mathbf{x}_1 \cdot \mathbf{x}_2$$

(nell'ultima uguaglianza abbiamo scritto λ_2 invece di $\overline{\lambda_2}$ perché l'autovalore è reale grazie alla parte (i) del teorema). Poiché $\lambda_1 \neq \lambda_2$ ne segue $\mathbf{x}_1 \cdot \mathbf{x}_2 = 0$.

- (iii) Per il Teorema Fondamentale dell'algebra, il polinomio caratteristico di A ha almeno una radice, e quindi A ha almeno un autovalore; per la parte (i) questo autovalore è reale.

Quando si sa che due autovettori sono ortogonali, li si può scegliere ortonormali: basta normalizzarli, cioè dividerli per la loro norma (lunghezza). Quindi il precedente Teorema 7.5.1 porta a questa conseguenza:

COROLLARIO 7.5.2. *Se tutti gli autovalori di una matrice autoaggiunta (in particolare simmetrica) sono distinti, allora la matrice è diagonalizzabile, ed esiste una base ortonormale di autovettori. Equivalentemente, se tutti gli autovalori di un operatore autoaggiunto (in particolare simmetrico) sono distinti, allora esiste una base ortonormale di autovettori.*

TEOREMA 7.5.3. (Diagonalizzabilità di operatori autoaggiunti.)

- (i) *Ogni matrice autoaggiunta è diagonalizzabile sui complessi, e la matrice di cambiamento di base che la diagonalizza è una matrice unitaria. Equivalentemente, ogni applicazione lineare autoaggiunta (noi diremo ora operatore autoaggiunto) ammette una base ortonormale di autovettori.*
- (ii) *Ogni matrice simmetrica è diagonalizzabile sui reali, e la matrice di cambiamento di base che la diagonalizza è una matrice ortogonale. Equivalentemente, ogni operatore simmetrico su uno spazio vettoriale reale ammette una base ortonormale di autovettori.*

DIMOSTRAZIONE. Dimostriamo solo la parte (i): la dimostrazione della parte (ii) è identica.

Gli autovettori con autovalori distinti si possono scegliere ortonormali per il Corollario 7.5.2. Quindi basta provare che per ogni autovalore si possono trovare tanti autovettori quanta è la molteplicità dell'autovalore (se non sono fra loro ortogonali li si può ortogonalizzare all'interno dell'autospazio con il procedimento di Gram-Schmidt illustrato

nella Sezione 6.7, che rimpiazza i vettori con loro combinazioni lineari opportune, le quali quindi restano autovettori.

Equivalentemente, ora mostriamo la seguente asserzione: *ogni operatore autoaggiunto A ammette una base ortonormale di autovettori*. La dimostrazione ricalca da vicino quella del Teorema 6.6.9 (cioè proprio del teorema di proiezione ortogonale che dà origine alla ortogonalizzazione di Gram-Schmidt: in effetti le due dimostrazioni sono essenzialmente equivalenti, visto che ora stiamo usando il prodotto scalare euclideo, che è definito positivo).

Procediamo di nuovo per induzione sulla dimensione n dello spazio vettoriale su cui A opera (ancora una volta stiamo usando l'assioma di induzione degli interi, (1.3)). Se $n = 1$ l'asserzione è ovvia, perché qualsiasi vettore non nullo forma una base, e si può normalizzare. Assumiamo ora $n > 1$. Per il Teorema 7.5.1 (iii) l'operatore A ha un autovalore reale λ ; indichiamo con $\mathbf{x}_1$ un autovettore corrispondente, e sia $\mathbf{X}_1$ lo spazio vettoriale che esso genera. Mostriamo che il complemento ortogonale $\mathbf{X}_1^\perp$ è *invariante* sotto l'operatore A , cioè che $A\mathbf{X}_1^\perp \subset \mathbf{X}_1^\perp$. Questo è vero perché, se $\mathbf{x} \in \mathbf{X}_1^\perp$, cioè se $\mathbf{x} \cdot \mathbf{x}_1 = 0$, allora

$$A\mathbf{x} \cdot \mathbf{x}_1 = \mathbf{x} \cdot A\mathbf{x}_1 = 0 = \lambda\mathbf{x} \cdot A\mathbf{x}_1 = 0.$$

D'altra parte, sempre perché il prodotto scalare euclideo è definito positivo, si ha $\dim \mathbf{X}_1^\perp = n - 1$ (Proposizione 6.6.5 (v); qui basterebbe usare il fatto che il prodotto scalare euclideo è non degenerare e ricorrere alla Proposizione 6.6.5 (iv)). Perciò, per ipotesi di induzione, $\mathbf{X}_1^\perp$ ha una base ortonormale: aggiungendo a questa base il vettore $\mathbf{x}_1$, come nella dimostrazione del Teorema 6.6.9, si ottiene una base ortonormale per $\mathbf{X}$. L'asserzione è dimostrata.

La matrice che porta dalla base canonica alla base ortonormale, cioè la matrice del cambiamento di base che diagonalizza A , è unitaria (nel caso reale, ortogonale) per il Teorema 6.8.7, perché manda la base canonica, ovviamente ortonormale rispetto al prodotto scalare euclideo, in una base ortonormale.

COROLLARIO 7.5.4. *Per ogni matrice A autoaggiunta (rispettivamente,*

simmetrica) esiste una matrice U unitaria (rispettivamente, ortogonale) tale che $U^\top AU = U^{-1}AU$ è diagonale.

7.6. Esercizi sulla diagonalizzazione di matrici simmetriche

ESERCIZIO 7.6.1. La matrice

$$A = \begin{pmatrix} 0 & 2 & 2 \\ 2 & 0 & 2 \\ 2 & 2 & 0 \end{pmatrix}$$

è diagonalizzabile? Se lo è, si trovi una diagonalizzazione della matrice tramite passaggio ad una base ortonormale, cioè tramite una matrice ortogonale di cambio di base.

SVOLGIMENTO. La matrice è simmetrica e reale, quindi per il Corollario 7.5.4 essa è diagonalizzabile ed esiste una base ortonormale in cui si diagonalizza. Raccogliendo i termini si vede facilmente che il polinomio caratteristico è $P(\lambda) \equiv (\lambda + 2)^2(\lambda - 4)$. L'autovalore $\lambda = 4$ ha molteplicità 1 e quindi dà luogo ad un autospazio di dimensione 1. Invece l'autovalore $\lambda = -2$ ha molteplicità 2, ma il calcolo del nucleo di $A - \lambda\mathbb{I} = A + \mathbb{I}$ porta ad un autospazio bidimensionale generato dai vettori

$$\mathbf{x}_1 = \begin{pmatrix} 1 \\ -1 \\ 0 \end{pmatrix} \quad \text{e} \quad \mathbf{x}_2 = \begin{pmatrix} 1 \\ 0 \\ -1 \end{pmatrix}$$

cioè $\mathbf{x}_1 = \mathbf{e}_1 - \mathbf{e}_2$ e $\mathbf{x}_2 = \mathbf{e}_1 - \mathbf{e}_3$. Il procedimento di ortogonalizzazione di Gram-Schmidt della Sezione 6.7 ci permette di trovare in questo autospazio due autovettori ortogonali, che risultano essere $\mathbf{x}_1 = \mathbf{e}_1 - \mathbf{e}_3$ e $\mathbf{x}'_2 = \frac{1}{2}\mathbf{e}_1 + \frac{1}{2}\mathbf{e}_2 - \mathbf{e}_3$. Ora, normalizzando, si trovano due autovettori ortonormali: $\mathbf{y}_1 = \frac{1}{\sqrt{2}}(\mathbf{e}_1 + \mathbf{e}_2)$ e $\mathbf{y}_2 = \frac{1}{\sqrt{6}}(\mathbf{e}_1 + \mathbf{e}_2 - 2\mathbf{e}_3)$. Il terzo autovettore, quello di autovalore 1, che indichiamo con $\mathbf{x}_3$, si determina calcolando il nucleo di $A - \mathbb{I}$. Si trova $\mathbf{x}_3 = \mathbf{e}_1 + \mathbf{e}_2 + \mathbf{e}_3$, e questo autovettore, nel rispetto del Teorema 7.5.1 (ii), è ortogonale agli altri due. Normalizzando si ottiene il terzo autovettore della base ortonormale: $\mathbf{y}_3 = \frac{1}{\sqrt{3}}(\mathbf{e}_1 + \mathbf{e}_2 + \mathbf{e}_3)$. Allora la matrice ortonormale del cambiamento

di base è

$$O = \begin{pmatrix} \frac{1}{\sqrt{2}} & \frac{1}{\sqrt{6}} & \frac{1}{\sqrt{3}} \\ -\frac{1}{\sqrt{2}} & \frac{1}{\sqrt{6}} & \frac{1}{\sqrt{3}} \\ 0 & \frac{-2}{\sqrt{6}} & \frac{1}{\sqrt{3}} \end{pmatrix}$$

e si ha

$$O^{-1}AO = O^{\top}AO = \begin{pmatrix} -2 & 0 & 0 \\ 0 & -2 & 0 \\ 0 & 0 & 4 \end{pmatrix}.$$

□

ESERCIZIO 7.6.2. La matrice

$$\begin{pmatrix} 1 & 2 & 0 & 0 \\ 2 & 1 & 0 & 0 \\ 0 & 0 & 1 & 2 \\ 0 & 0 & 2 & 1 \end{pmatrix}$$

è diagonalizzabile? Se lo è si trovi una diagonalizzazione tramite una matrice ortogonale.

SVOLGIMENTO. La matrice è diagonale a blocchi, e per di più i due blocchi sono identici. Gli autovalori che ciascun blocco genera sono -1 e 3 , quindi ciascuno dei due ha molteplicità due ma A è diagonalizzabile perché è simmetrica. Gli autovettori si trovano immediatamente grazie alla Nota 7.3.13: per l'autovalore -1 , due autovettori indipendenti (ed ovviamente ortogonali, perché provenienti da blocchi diversi) sono $\mathbf{e}_1 - \mathbf{e}_2$ e $\mathbf{e}_3 - \mathbf{e}_4$, mentre per l'autovalore 3 abbiamo ad esempio i seguenti: $\mathbf{e}_1 + \mathbf{e}_2$ ed $\mathbf{e}_3 + \mathbf{e}_4$. Per normalizzare basta dividere questi autovettori per $\sqrt{2}$. La base ortonormale così ottenuta porta alla diagonalizzazione

$$O^{-1}AO = O^{\top}AO = \begin{pmatrix} -1 & 0 & 0 & 0 \\ 0 & -1 & 0 & 0 \\ 0 & 0 & 3 & 0 \\ 0 & 0 & 0 & 3 \end{pmatrix},$$

dove

$$O = \begin{pmatrix} \frac{1}{\sqrt{2}} & 0 & \frac{1}{\sqrt{2}} & 0 \\ -\frac{1}{\sqrt{2}} & 0 & \frac{1}{\sqrt{2}} & 0 \\ 0 & \frac{1}{\sqrt{2}} & 0 & \frac{1}{\sqrt{2}} \\ 0 & -\frac{1}{\sqrt{2}} & 0 & \frac{1}{\sqrt{2}} \end{pmatrix}.$$

□

7.7. Qualche applicazione degli autovettori

Ci sono innumerevoli impieghi degli autovettori nella matematica, ed altrettanto innumerevoli sono i casi in cui gli autovettori compaiono nelle applicazioni della matematica alla modellazione di fenomeni fisici, chimici, biologici, economici, ingegneristici. Ne illustriamo qui due tipiche: *dinamica delle popolazioni* e *sistemi dinamici*. La presentazione della prima è tratta da [2].

7.7.1. Dinamica delle popolazioni. Si consideri una specie animale. Ogni anno, oppure in ogni ciclo generazionale (la cui lunghezza chiameremo comunque *anno*, ma l'esempio funziona con qualsiasi altro ciclo riproduttivo) alcuni animali muoiono, ne nascono di nuovi e tutti gli animali che restano in vita invecchiano di un anno. Questo dà luogo ad una distribuzione d'età della popolazione di quella specie, che può variare di anno in anno. Certe distribuzioni potrebbero restare stabili. Vogliamo trovare se ci sono e quali sono le distribuzioni di età stabili nel tempo. Per semplicità supponiamo che ci sia una proporzione fissa fra maschi e femmine. Questo è sensato in molte specie animali. Ad esempio, nei mammiferi la probabilità che nasca un maschio o una femmina è la stessa. Questo di per sé non implica che ci siano in vita lo stesso numero di maschi e femmine, perché la lunghezza media della vita potrebbe essere diversa fra i due generi: ma è molto ragionevole supporre che ci sia una proporzione fissa stabile nel tempo (e non lontana dall'uguaglianza se la vita media è lunga).

Sotto questa ipotesi, per stimare la distribuzione d'età della popolazione basta stimare la distribuzione d'età delle femmine (o dei maschi, ma le femmine sono più opportune da considerare perché la maternità

di un cucciolo è molto più documentabile della sua paternità). Allora modelliamo il problema come segue.

Supponiamo che la durata massima della vita sia n anni. Indichiamo con

- $x_i^{(k)}$ il numero di femmine di età i anni che sono vive all'anno k (con $1 \leq k \leq n$);
- f_i la percentuale di femmine di età i anni che rimangono vive un anno dopo (con $0 \leq i \leq n$);
- b_i il numero medio di cuccioli femmine generati da una femmina di età i (con $0 \leq i \leq n$).

Questi dati si raccolgono sperimentalmente. I coefficienti b_i sono nulli per i valori di i prima dell'età prepuberale, la quale però è zero se l'unità di tempo è il ciclo riproduttivo; se invece l'unità di tempo è fissata indipendentemente, ad esempio un anno, comunque molte specie animali raggiungono la pubertà fin dal primo anno, le femmine hanno cuccioli prima. Per elevati valori di i la fertilità di solito diminuisce e quindi i b_i diventano piccoli, e così pure naturalmente gli f_i , ma questo fatto nello sviluppo matematico è inessenziale.

Le nuove femmine (cioè di età zero anni) all'anno $k + 1$ sono generate fra l'anno k e l'anno $k + 1$ in proporzione ai pesi b_i da madri di varie fasce d'età in vita all'anno k : pertanto

$$x_0^{(k+1)} = \sum_{i=0}^n b_i x_i^{(k)}.$$

Invece il numero di femmine di età $i > 0$ vive all'anno $k + 1$ è quello delle femmine che all'anno k avevano età $i - 1$ moltiplicato per il fattore di sopravvivenza f_i :

$$x_i^{(k+1)} = f_{i-1} x_{i-1}^{(k)}.$$

Scrivendo la distribuzione d'età come un vettore

$$\mathbf{x}^{(k)} = \left(x_0^{(k)}, x_1^{(k)}, \dots, x_n^{(k)} \right),$$

otteniamo quindi il sistema lineare

$$\mathbf{x}^{(k+1)} = A\mathbf{x}^{(k)},$$

dove la matrice A della dinamica della popolazione è

$$A = \begin{pmatrix} b_0 & b_1 & b_2 & \cdots & b_{n-1} & b_n \\ f_0 & 0 & 0 & \cdots & 0 & 0 \\ 0 & f_1 & 0 & \cdots & 0 & 0 \\ \vdots & \vdots & \vdots & & \vdots & \vdots \\ 0 & 0 & 0 & \cdots & f_{n-1} & 0 \end{pmatrix}.$$

Stiamo cercando le distribuzioni d'età che rimangono stabili all'aumentare del tempo k , cioè tali che

$$\mathbf{x}^{(k)} = \mathbf{x}^{(k+1)} = A\mathbf{x}^{(k)},$$

e quindi stiamo cercando nient'altro che gli autovettori della matrice A con autovalore 1. Ad esempio, se una specie vive al massimo due

anni, diventa pubere all'ultimo anno di vita nel quale mediamente ogni femmina ha 6 cuccioli, e la sopravvivenza è il 50% fra la nascita ed il primo anno ed il 30% fra il primo ed il secondo anno, la matrice della dinamica della popolazione è

$$A = \begin{pmatrix} 0 & 0 & 6 \\ \frac{1}{2} & 0 & 0 \\ 0 & \frac{1}{3} & 0 \end{pmatrix},$$

e (a meno di multipli) si trova un solo autovettore con autovalore 1, il vettore $(6, 3, 1)$. Quindi, normalizzando, si ottiene la distribuzione stabile di età della popolazione: 60% nella fascia fra zero ed un anno, 30% nella fascia fra uno e due anni, 10% nella fascia sopra i due anni. Si noti che in questo caso l'autospazio di autovalore 1 ha dimensione 1, e quindi c'è un'unica distribuzione di età stabile nel tempo, ma in altri casi potrebbero essercene due o anche tre, oppure nessuna se la matrice A non ha autovalori reali.

7.7.2. Sistemi dinamici. Un operatore lineare agisce su uno spazio vettoriale fissandone l'origine, ma in generale spostando gli altri punti. Fissata una base, la sua azione si visualizza come l'azione della matrice associata sui vettori, con l'operazione di prodotto righe per colonna.

DEFINIZIONE 7.7.1. Si chiama *orbita* del punto $\mathbf{x} \in \mathbb{R}^n$ sotto l'azione della matrice $A \in M_{nn}$ la successione di punti $\{A^n \mathbf{x}, n \in \mathbb{N}\}$.

Siamo interessati a studiare l'andamento asintotico delle orbite. Ovviamente l'origine coincide con la sua stessa orbita, poiché un punto fisso di ogni operatore lineare. Cosa accade alle orbite degli altri punti per grandi valori di n ? A seconda del dato iniziale e della scelta dell'operatore, l'orbita può tendere a zero, allontanarsi illimitatamente o restare a distanza costante dall'origine; analoghe situazioni si hanno per spazi vettoriali sul campo complesso.

Per studiare l'andamento asintotico indipendentemente dal fatto che i punti dell'orbita siano vettori che diventano piccoli oppure grandi, ci interesseremo particolarmente alla loro *direzionale asintotica*, cioè al loro valore normalizzato $A^n \mathbf{x} / \|A^n \mathbf{x}\|$.

ESEMPIO 7.7.2. (*Esempi di orbite in $\mathbb{R}^2$.*)

(i) L'orbita di $\mathbf{x}$ sotto la matrice di rotazione

$$A_\theta = \begin{pmatrix} \cos \theta & \sin \theta \\ -\sin \theta & \cos \theta \end{pmatrix}$$

è contenuta nella circonferenza di raggio $\|\mathbf{x}\|$. In effetti, l'orbita si ripete periodicamente dopo N passi se θ è del tipo $\frac{2\pi}{N}$ con N intero, ed altrimenti è densa ovunque nella circonferenza, nel senso che per ogni punto della circonferenza c'è un punto dell'orbita arbitrariamente vicino.

(ii) Se

$$B = \begin{pmatrix} a & 0 \\ 0 & b \end{pmatrix}$$

con $a > 1 > b > 0$, allora l'orbita di $\mathbf{x} = (x_1, x_2)$ è

$$\{(a^n x_1, b^n x_2), n = 1, 2, \dots\}.$$

Se ad esempio $\mathbf{x}$ è nel primo quadrante, cioè $x_1, x_2 > 0$, allora l'orbita rimane nel primo quadrante, ma la prima componente $a^n x_1$ tende a $+\infty$, mentre la seconda tende a zero. Le direzioni invece tendono al vettore $\mathbf{e}_1 = (1, 0)$, perché a_n tende ad infinito mentre b^n tende a zero, e quindi $\frac{a^n}{\sqrt{a^{2n} + b^{2n}}}$ tende a uno,

mentre $\frac{a^n}{\sqrt{a^{2n}+b^{2n}}}$ tende a zero. Si osservi che $\mathbf{e}_1$ è un autovettore corrispondente al più grande dei due autovalori, l'autovalore a .

(iii) Se

$$C = \begin{pmatrix} a & 0 \\ 0 & b \end{pmatrix}$$

con $a > b > 1$, allora entrambe le componenti di $C^n \mathbf{x}$ divergono, ma la prima diverge più rapidamente della seconda, ed anche in questo caso la direzione asintotica dei punti dell'orbita tende ad $\mathbf{e}_1$, cioè all'autovettore corrispondente al massimo autovalore.

(iv) Se

$$D = \begin{pmatrix} a & 0 \\ 0 & b \end{pmatrix}$$

con $0 < b < a < 1$, allora entrambe le componenti di $D^n \mathbf{x}$ tendono a zero, ma la prima tende a zero più lentamente della seconda, ed anche in questo caso la direzione asintotica dei punti dell'orbita tende ad $\mathbf{e}_1$, cioè all'autovettore corrispondente al massimo autovalore.

(v) Se

$$E = \begin{pmatrix} a & 0 \\ 0 & 1 \end{pmatrix}$$

con $a > 1$, allora la prima componente di $E^n \mathbf{x}$ diverge e la seconda rimane costante: anche in questo caso la direzione asintotica dei punti dell'orbita tende ad $\mathbf{e}_1$, cioè all'autovettore corrispondente al massimo autovalore.

(vi) Se

$$F = \begin{pmatrix} 1 & 0 \\ 0 & b \end{pmatrix}$$

con $b < 1$, allora la prima componente di $F^n \mathbf{x}$ rimane costante e la seconda tende a zero: anche in questo caso la direzione asintotica dei punti dell'orbita tende ad $\mathbf{e}_1$, cioè all'autovettore corrispondente al massimo autovalore.

- (vi) Infine, se negli esempi precedenti si cambia base, con un cambiamento di base che manda gli autovettori $\mathbf{e}_1, \mathbf{e}_2$ in due nuovi vettori, rispettivamente $\mathbf{v}_1, \mathbf{v}_2$, allora il cambiamento di base trasforma le matrici dei punti precedenti in matrici che non sono più diagonali, e che hanno per autovettori $\mathbf{v}_1, \mathbf{v}_2$, rispettivamente con autovalore a e b . Le orbite cambiano in maniera prevedibile, restando ora non nel primo quadrante, bensì nel settore angolare delimitato dai vettori $\mathbf{v}_1, \mathbf{v}_2$; i punti dell'orbita, corrispondentemente, si avvicinano in direzione, asintoticamente, alla direzione del trasformato di $\mathbf{e}_1$, e cioè $\mathbf{v}_1$, che è l'autovettore della matrice trasformata corrispondente al massimo autovalore.
- (vii) Tutti gli esempi qui sopra cambiano senza differenze essenziali se i termini diagonali a e b non sono più positivi (in valore assoluto non cambia nulla) o non più reale (prendendo il modulo dei numeri complessi si ritorna agli esempi dei punti precedenti). C'è però una differenza sottile. Consideriamo il caso di due autovalori $a = b = 1$. Se la matrice è diagonalizzabile sui reali, allora è la matrice identità ed ogni orbita è un punto (e viceversa). Però, se la matrice non è diagonalizzabile sui reali, può esserlo sui complessi. Nel caso la matrice sia a coefficienti reali ma diagonalizzabile sui complessi, allora essa è una matrice di rotazione come A_θ nella parte (i) di questo Esempio, i punti delle orbite girano dentro circonferenze e le direzioni ruotano senza avvicinarsi ad una direzione limite. Infine, se la matrice non è diagonalizzabile neppure sui complessi, allora è riducibile, con un opportuno cambiamento di base, ad una matrice triangolare, cioè del tipo

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

(si veda la Sezione 9.2 sulla forma canonica di Jordan). In tal caso l'unico autovettore è $(1, 0)$ (coordinate rispetto a questa base dove la matrice si triangolarizza) e l'orbita di $\mathbf{x} = (x_1, x_2)$, rispetto a questa nuova base, consiste dei punti

$$G^n \mathbf{x} = (x_1 + nx_2, x_2), \quad n = 1, 2, \dots$$

(si verifichi per esercizio). Quindi, se $\mathbf{x}$ non giace sul primo asse, cioè se $x_2 \neq 0$, allora i punti dell'orbita si avvicinano in direzione a $(1, 0)$, cioè all'unico autovettore. Nel caso rimanente, $x_2 = 0$, l'orbita giace nel primo asse e quindi è gioco forza che si avvicini in direzione al versore di tale asse (la direzione non cambia da un punto dell'orbita al successivo).

□

PROPOSIZIONE 7.7.3. *L'esempio precedente, come l'analisi che abbiamo fatto rivela, rispecchia la situazione generale. la direzione asintotica dei punti dell'orbita generica si avvicina a quella dell'autovalore di autovettore massimo (se ci sono autovettori di modulo diverso da 1 o da 0). Questo fatto è vero, per lo stesso argomento, non solo in dimensione 2 ma anche in dimensione n .*

Questo esempio mostra anche la rilevanza dello *spazio delle direzioni*, detto anche *spazio proiettivo*, che riprenderemo in esame nella Sezione **13.2**.

Somma diretta di spazi vettoriali

DEFINIZIONE 8.0.4. (*Somma diretta.*)

- (i) Siano $\mathbf{V}$, $\mathbf{W}$ sottospazi vettoriali di uno spazio vettoriale $\mathbf{X}$. Diciamo che $\mathbf{X}$ è *somma diretta* di $\mathbf{V}$ e $\mathbf{W}$, e scriviamo $\mathbf{X} = \mathbf{V} \oplus \mathbf{W}$, se
- (a) $\mathbf{V} \cap \mathbf{W} = \emptyset$
 - (b) $\mathbf{V} \cup \mathbf{W}$ genera $\mathbf{X}$ come spazio vettoriale: in altre parole, $\mathbf{X} = \mathbf{V} + \mathbf{W}$ è lo spazio delle combinazioni lineari di vettori in $\mathbf{V}$ e $\mathbf{W}$, rispettivamente.
- (ii) Più in generale, dati due spazi vettoriali $\mathbf{V}$, $\mathbf{W}$ sullo stesso campo, ma non necessariamente sottospazi di uno stesso spazio vettoriale, si costruisce uno spazio vettoriale $\mathbf{Y}$ *somma diretta di $\mathbf{V}$ e $\mathbf{W}$* , nel quale $\mathbf{V}$, $\mathbf{W}$ si immergono come sottospazi, nel modo seguente: $\mathbf{Y}$ è l'insieme prodotto cartesiano $\mathbf{V} \times \mathbf{W}$, la somma di due coppie è definita da $(\mathbf{v}_1, \mathbf{w}_1) + (\mathbf{v}_2, \mathbf{w}_2) = (\mathbf{v}_1 + \mathbf{v}_2, \mathbf{w}_1 + \mathbf{w}_2)$, ed il prodotto per scalari è dato da $\lambda(\mathbf{v}, \mathbf{w}) = (\lambda\mathbf{v}, \lambda\mathbf{w})$. È evidente che $\mathbf{V}$, $\mathbf{W}$ si immergono in $\mathbf{X}$ come sottospazi vettoriali, che le combinazioni lineari dei loro vettori generano $\mathbf{X}$ e che la loro intersezione come sottospazi di $\mathbf{X}$ consiste solo del sottospazio banale $\{\mathbf{0}\}$

NOTA 8.0.5. È chiaro che gli omomorfismi di $\mathbf{V}$ (rispettivamente, $\mathbf{W}$) in $\mathbf{X}$ definiti da $\mathbf{v} \mapsto (\mathbf{v}, \mathbf{0})$ (rispettivamente, $\mathbf{w} \mapsto (\mathbf{0}, \mathbf{w})$) sono iniettivi. Essi si chiamano le *iniezioni canoniche* dei fattori nella somma diretta. Viceversa, gli omomorfismi da $\mathbf{Y}$ a $\mathbf{V}$ (rispettivamente, $\mathbf{W}$) dati da $(\mathbf{v}, \mathbf{w}) \mapsto \mathbf{v}$ (rispettivamente, $(\mathbf{v}, \mathbf{w}) \mapsto \mathbf{w}$) sono surgettivi sui fattori: si chiamano le *proiezioni canoniche*, e si indicano con P_1 e P_2 , rispettivamente. □

PROPOSIZIONE 8.0.6. *Date una base $\mathcal{V}$ nello spazio vettoriale $\mathbf{V}$ ed una base $\mathcal{W}$ in $\mathbf{W}$, la loro unione $\mathcal{V} \cup \mathcal{W}$ è una base in $\mathbf{V} \oplus \mathbf{W}$.*

DIMOSTRAZIONE. $\mathcal{V} \cup \mathcal{W}$ genera lo spazio vettoriale $\mathbf{Y} = \mathbf{V} + \mathbf{W}$ in base alla proprietà (ii) della Definizione 8.0.4. Per provare che $\mathcal{V} \cup \mathcal{W}$ è una base, supponiamo che sia nulla una combinazione lineare di questi vettori, diciamo $\sum_{i=1}^{n_1} \lambda_i \mathbf{v}_i + \sum_{j=1}^{n_2} \mu_j \mathbf{w}_j = \mathbf{0}$. Allora le due sommatorie in questa combinazione lineare devono essere separatamente nulle, perché se non fosse così il vettore $\sum_{i=1}^{n_1} \lambda_i \mathbf{v}_i \in \mathbf{V}$ coinciderebbe con il vettore $-\sum_{j=1}^{n_2} \mu_j \mathbf{w}_j \in \mathbf{W}$, contrariamente alla condizione (i) della Definizione 8.0.4. Ma poiché $\mathcal{V}$ e $\mathcal{W}$ sono basi dei rispettivi spazi vettoriali, tutti i coefficienti λ_i e μ_j devono essere nulli. Questo prova l'indipendenza lineare.

NOTA 8.0.7. Abbiamo già incontrato due esempi di somma diretta.

(i) L'enunciato della Proposizione 6.6.5 (v) equivale a:

PROPOSIZIONE 8.0.8. *Se $\mathbf{X}$ è uno spazio vettoriale munito di un prodotto scalare definito positivo, allora per ogni sottospazio $\mathbf{V}$ si ha*

$$\mathbf{X} = \mathbf{V} \oplus \mathbf{V}^\perp. \quad (8.0.1)$$

come si verifica immediatamente rivedendo la dimostrazione.

(ii) La Nota 7.3.13 illustra il fatto che, se lo spazio vettoriale $\mathbf{X}$ si decompone come $\mathbf{X} = \bigoplus_{i=1}^k \mathbf{V}_i$, con $\dim \mathbf{V}_i = n_i$ e $\sum_{i=1}^k n_i = n$, ed A è un operatore lineare su $\mathbf{X}$ che preserva i blocchi, cioè tale che $A\mathbf{V}_i \subset \mathbf{V}_i$ per ogni i , allora in una base ciascuno dei cui elementi appartiene ad uno dei sottospazi $\mathbf{V}_i$ la matrice di A si decompone a blocchi, ed ogni blocco corrisponde al riquadro non nullo della matrice $P_i A P_i$, dove P_i è l'operatore di proiezione ortogonale su $\mathbf{V}_i$, la cui matrice è diagonale con autovalore 1 sui vettori della base canonica appartenenti a $\mathbf{V}_i$ e zero altrove.

□

CAPITOLO 9

Triangolarizzazione e forma canonica di Jordan

9.1. * Polinomio minimo

9.2. Forma canonica di Jordan

Spazi normati, funzionali lineari e dualità

10.1. La disuguaglianza di Cauchy-Schwartz

PROPOSIZIONE 10.1.1. *Se V è uno spazio vettoriale munito di un prodotto scalare e $v, w \in V$ allora*

$$|(v, w)| \leq \|v\|^2 \|w\|^2$$

DIMOSTRAZIONE. Osserviamo dapprima che la disuguaglianza è vera in ogni spazio V di dimensione 2. Questo segue dalla Definizione 6.1.1: $(v, w) = \|v\| \|w\| \cos(\theta)$ dove θ è l'angolo formato dai vettori v e w ; oppure, in termini di coordinate nella base canonica, dalle seguenti uguaglianze:

se $v = v_1 e_1 + v_2 e_2$ e $w = w_1 e_1 + w_2 e_2$, allora

$$\begin{aligned} |(v, w)|^2 &= |v_1 \bar{w}_1 + v_2 \bar{w}_2|^2 \\ &= |v_1 \bar{w}_1|^2 + |v_2 \bar{w}_2|^2 + 2 \operatorname{Re}(v_1 \bar{w}_1 - \bar{v}_2 w_2) \end{aligned}$$

mentre

$$\begin{aligned} \|v\|^2 \|w\|^2 &= (|v_1|^2 + |v_2|^2)(|w_1|^2 + |w_2|^2) \\ &= |v_1 w_1|^2 + |v_2 w_2|^2 + |v_1|^2 |w_2|^2 + |w_1|^2 |v_2|^2. \end{aligned}$$

Infatti da qui segue che la disuguaglianza dell'enunciato equivale a dire che $\forall v_1, v_2, w_1, w_2$

$$2 \operatorname{Re}(v_1 \bar{w}_1 - \bar{v}_2 w_2) \leq |v_1|^2 |w_2|^2 + |w_1|^2 |v_2|^2$$

cioè $\forall a, b \in \mathbb{C}$

$$2 \operatorname{Re}(a \bar{b}) \leq |a|^2 + |b|^2.$$

Questo è vero perché

$$|a|^2 + |b|^2 - 2 \operatorname{Re}(a\bar{b}) = |a - b|^2 \geq 0.$$

Questo dimostra la Proposizione se la dimensione di V è uguale a 2. Ma se V è arbitrario, una volta scelti v, w applichiamo la disuguaglianza appena provata allo spazio bidimensionale generato da v e w . $\square$

10.2. Operatori lineari e funzionali lineari su spazi normati

In questa sezione denotiamo con *spazio normato* uno spazio vettoriale munito di una norma, nel senso della Definizione 6.5.1. In tutti i risultati presumiamo che lo spazio vettoriale sia a dimensione finita, ma una variante dei risultati vale per operatori lineari *continui* su spazi normati a dimensione infinita: per questo rammenteremo volta per volta che la dimensione è finita. Per il caso generale, si consulti un [libro appropriato](#), ad esempio [9].

DEFINIZIONE 10.2.1. Se V, W sono spazi normati e $T : V \rightarrow W$ è un operatore lineare, allora si chiama *norma* di T il numero

$$\|T\| = \inf \{C > 0 : \|Tv\|_W \leq C\|v\|_V \ \forall v \in V\}$$

COROLLARIO 10.2.2. *Dati due operatori lineari $A : V \rightarrow W$ e $B : U \rightarrow V$, indichiamo con $AB : U \rightarrow W$ l'operatore composto $A \circ B$. Allora*

$$\|AB\| \leq \|A\|\|B\|.$$

DIMOSTRAZIONE. . Segue immediatamente dalla associatività, grazie alla quale per ogni vettore $u \in U$ si ha

$$\|(AB)u\|_W = \|A(Bu)\|_W \leq \|A\|\|Bu\|_V \leq \|A\|\|B\|\|u\|_U$$

da cui l'enunciato.

PROPOSIZIONE 10.2.3. *Se V ha dimensione finita, $W = V$ e T è diagonale, allora*

$$\|T\| = \max \{|\lambda| : \lambda \text{ autovalore di } T\}.$$

Se invece T è diagonalizzabile, cioè se esiste una base di $v_1, \dots, v_n$ di V tale che T manda ogni vettore v_i in un suo multiplo $\lambda_i v_i$, allora

l'operatore C di cambiamento di base che manda la base canonica nella base dei v_i diagonalizza T , nel senso che $C^{-1}TC$ è diagonale, e

$$\|T\| \leq \|C\| \|C^{-1}\| \max \{|\lambda| : \lambda \text{ autovalore di } T\} .$$

DIMOSTRAZIONE. . Supponiamo che T sia diagonale (nella base canonica): cioè che gli autovettori siano $\mathbf{e}_1, \dots, \mathbf{e}_n$, e che gli autovalori siano ordinati in modo che $|\lambda_1| \geq |\lambda_2| \geq \dots \geq |\lambda_n|$. Allora

$$Tv = T\left(\sum_{i=1}^n v_i e_i\right) = \sum_{i=1}^n \lambda_i v_i e_i,$$

e quindi

$$\|Tv\|^2 = \sum_{i=1}^n |\lambda_i|^2 |v_i|^2 \leq |\lambda_1|^2 \sum_{i=1}^n |v_i|^2 = |\lambda_1|^2 \|v\|^2.$$

Perciò $\|T\| \leq |\lambda_1|$. D'altra parte, $T\mathbf{e}_1 = \lambda_1 \mathbf{e}_1$, e quindi $\|T\| \geq |\lambda_1|$. Ne segue che

$$\|T\| = |\lambda_1| = \max \{|\lambda| : \lambda \text{ autovalore di } T\} .$$

Questo prova la prima parte dell'enunciato. La seconda parte segue immediatamente dalla prima e dal Corollario 10.2.2.

DEFINIZIONE 10.2.4. Un funzionale lineare F su uno spazio vettoriale normato V (sul campo complesso) è un operatore lineare $F : V \rightarrow \mathbb{C}$.

In base alla Definizione 10.2.1, la norma di F è data da:

$$\|F\| = \inf \{C > 0 : |Fv| \leq C\|v\| \ \forall v \in V\} .$$

Ora studiamo i funzionali lineari su uno spazio vettoriale a dimensione finita (enunciati e dimostrazioni si possono estendere a spazi normati a dimensione infinita, completi rispetto alla norma: si consulti un [libro di testo sugli spazi di Hilbert](#), ad esempio [9]).

TEOREMA 10.2.5. (**Teorema di rappresentazione di Riesz.**) *Se H è uno spazio normato a dimensione finita, i funzionali lineari su H formano uno spazio vettoriale isomorfo a H , e l'isomorfismo preserva la norma. In particolare, per ogni $u \in H$ il funzionale lineare F_u definito da*

$$F_u(v) = (v, u)_H \quad \forall v \in H$$

verifica $\|F_u\| = \|u\|$. Viceversa ogni funzionale su H è del tipo F_u per qualche $u \in H$, e questo u è unico.

NOTAZIONE 10.2.6. Lo spazio dei funzionali lineari su uno spazio normato V si chiama il duale di V e si indica con V' .

DIMOSTRAZIONE. E' ovvio che H' è uno spazio vettoriale. Sia $u \in H$ e $F_u(v) = (v, u)$. Allora F_u è un funzionale lineare e per la disuguaglianza di Cauchy-Schwartz (Proposizione 10.1.1) si ha

$$|F_u(v)| = |(v, u)| \leq \|v\| \|u\|.$$

Quindi, per la Proposizione 10.1.1, $\|F_u\| \leq \|u\|$. Ma $F_u(u) = (u, u) = \|u\|^2$, e perciò

$$\|F_u\| = \inf \{C > 0 : |F_u(v)| \leq C\|v\| \ \forall v \in H\} \geq \|u\|.$$

Quindi $\|F_u\| = \|u\|$.

Viceversa, mostriamo che $\forall F \in H'$ esiste $u \in H$ tale che $F = F_u$. Se $F = 0$, allora questo è vero per $u = 0$. Altrimenti sia

$$M = \text{Ker } F = \{v \in H : F(v) = 0\}.$$

Osserviamo che M è un sottospazio vettoriale di H , perché F è lineare. Poiché $F \neq 0$ si ha $M \neq H$. Perciò possiamo trovare $v_0 \in H$ con $v_0 \notin M$. Ma allora esiste $w \in H$, $w \perp M$: basta prendere $w = v_0 - P_M(v_0)$, dove P_M è la proiezione ortogonale su M definita nella Proposizione 6.6.5 (v). Riscalando, otteniamo un vettore w_0 tale che $w_0 \perp M$ e $F(w_0) = 1$.

Per ogni $h \in H$, ponendo $\alpha = F(h)$, osserviamo che

$$F(h - \alpha w_0) = F(h) - \alpha F(w_0) = \alpha - \alpha = 0.$$

Quindi $h - \alpha w_0 \in M$, e perciò

$$0 = (h - \alpha w_0, w_0) = (h, w_0) - \alpha \|w_0\|^2 = (h, w_0) - F(h) \|w_0\|^2.$$

Poniamo $u = \frac{w_0}{\|w_0\|^2}$: allora ne segue che, $\forall h \in H$, $F(h) = (h, u)$. Quindi $F = F_u$.

Questo vettore u è unico, perché se $F = F_u = F_{u'}$, allora

$$(h, u) = (h, u') \quad \forall h \in H,$$

e quindi $u - u' \perp H$. Ma allora $u = u'$. □

NOTA 10.2.7. Il Teorema di rappresentazione di Riesz 10.2.5 asserisce che ogni spazio vettoriale di dimensione finita V , è isomorfo al suo duale. Si osservi che, se $e_i = (0, 0, \dots, 1, 0, \dots, 0)$ con 1 al posto i è un vettore della base canonica, allora il funzionale lineare associato è

$$F_{e_i}(v) = (v, e_i) = v_i,$$

cioè il suo valore è la i -esima coordinata del vettore a cui F_{e_i} si applica. Quindi i funzionali che leggono le singole coordinate formano una base in V' che chiameremo base canonica di V' (una volta fissata una base canonica in V). $\square$

ESERCIZIO 10.2.8. Sia H uno spazio normato e sia $F \in H'$ un funzionale lineare continuo su H . Dimostrare che se $K = \text{Ker } F = \{x : F(x) = 0\}$, allora $\dim \text{Ker}^\perp = 1$.

Svolgimento. $y \in \text{Ker}^\perp \Leftrightarrow (y, x) = 0$ se $F(x) = 0$. Se y_1, y_2 sono due vettori linearmente indipendenti in $\text{Ker}^\perp$, allora

$$(y_1 - y_2, x) = 0 \quad \forall x \in K \Rightarrow y_1 - y_2 \in \text{Ker}^\perp.$$

Scegliamo y_1 e y_2 in modo tale che $F(y_1) = F(y_2) = 1$ (ciò è sempre possibile dividendo y_i per $F(y_i)$, $i = 1, 2$).

Allora

$$F(y_1 - y_2) = F(y_1) - F(y_2) = 1 - 1 = 0 \Rightarrow y_1 - y_2 \in \text{Ker}^\perp.$$

Ma $\text{Ker}^\perp \cap \text{Ker}\{F=0\}$ perché il prodotto scalare è non degenere (Definizione 6.5.10) e quindi $y_1 - y_2 = 0$, cioè $y_1 = y_2$. $\square$

NOTA 10.2.9. L'esercizio può essere risolto anche utilizzando il Teorema di rappresentazione di Riesz 10.2.5. $\square$

10.3. * Duale di somme dirette e di complementi ortogonali

DEFINIZIONE 10.3.1. Ogni prodotto scalare su uno spazio $\mathbf{X}$ induce un prodotto scalare sul suo duale $\mathbf{X}'$ nel modo seguente. Sia ϕ l'isomorfismo da $\mathbf{X}$ a $\mathbf{X}'$ presentato nella Nota 10.2.7: per semplicità scriviamo x' invece di $\phi(x)$. Allora

$$\langle x', y' \rangle := \langle x, y \rangle.$$

DEFINIZIONE 10.3.2. (*Immersione del duale di un sottospazio nel duale dello spazio.*) Immergiamo in $\mathbf{X}'$ il duale $\mathbf{V}'$ del sottospazio $\mathbf{V}$: per ogni $\mathbf{x} \in \mathbf{X}$, $\mathbf{v}' \in \mathbf{V}'$ poniamo

$$\langle \mathbf{v}', \mathbf{x} \rangle := \langle \mathbf{v}', P_{\mathbf{V}} \mathbf{x} \rangle$$

dove $P_{\mathbf{V}}$ è la proiezione canonica di $\mathbf{X}$ su $\mathbf{V}$ introdotta nella Definizione 8.0.4 (ii) e nella Nota 8.0.5. In altre parole, l'immersione che stiamo definendo è $\mathbf{v}' \mapsto \mathbf{v}' \circ P_{\mathbf{V}}$, ma con abuso di notazione continuiamo ad indicare con $\mathbf{v}'$ l'immagine di $\mathbf{v}'$ in $\mathbf{X}'$ sotto questa immersione.

NOTA 10.3.3. (* *Inversi destro e sinistro dell'immersione da un sottospazio e proiezione per restrizione.*)

Ricapitolando, osserviamo che l'immersione $\psi : \mathbf{V}' \longrightarrow \mathbf{X}'$ introdotta nella Definizione 10.3.2 è data da

$$\psi(\mathbf{v}') = \mathbf{v}' \circ P_{\mathbf{V}}$$

dove $P_{\mathbf{V}}(x) = v$ se $\mathbf{x}$ si decompone come $\mathbf{x} = \mathbf{v} + \mathbf{w}$ con $\mathbf{v} \in \mathbf{V}$, $\mathbf{w} \in \mathbf{V}^{\perp}$.

Ora costruiamo un inverso sinistro di ψ . Sia $P_{\mathbf{V}}^{\sharp}$ la *proiezione per restrizione* di $\mathbf{X}'$ su $\mathbf{V}'$, definita da

$$P_{\mathbf{V}}^{\sharp} \mathbf{x}' = \mathbf{x}'|_{\mathbf{V}}. \quad (10.3.1)$$

Analogamente a prima, scriviamo la decomposizione ortogonale $\mathbf{x} = \mathbf{v}_{\mathbf{x}} + \mathbf{w}_{\mathbf{x}}$ con $\mathbf{v}_{\mathbf{x}} \in \mathbf{V}$, $\mathbf{w}_{\mathbf{x}} \in \mathbf{V}^{\perp}$. Allora $P_{\mathbf{V}}^{\sharp} \psi(\mathbf{v}') = P_{\mathbf{V}}^{\sharp} \circ \mathbf{v}' \circ P_{\mathbf{V}}$ è il funzionale su $\mathbf{V}$ che soddisfa

$$P_{\mathbf{V}}^{\sharp}(\psi(\mathbf{v}'))(\mathbf{v}_{\mathbf{x}}) = \mathbf{v}'(P_{\mathbf{V}} \mathbf{x})|_{\mathbf{V}} = \mathbf{v}'(\mathbf{v}_{\mathbf{x}}).$$

Pertanto

$$P_{\mathbf{V}}^{\sharp}(\psi(\mathbf{v}'))(\mathbf{v}_{\mathbf{x}}) = \mathbf{v}'(\mathbf{v}_{\mathbf{x}})$$

cioè $P_{\mathbf{V}}^{\sharp} \circ \psi$ è l'applicazione identità.

Inoltre, $P_{\mathbf{V}}^{\sharp}$ fornisce anche, nell'unico senso possibile, un inverso destro dell'immersione ψ . Infatti, sia $\mathbf{x}' \in \mathbf{X}'$ e sia $\mathbf{z}' = \mathbf{z}'_{\mathbf{x}'}$ la sua *proiezione su $\mathbf{V}$* , cioè il funzionale su $\mathbf{X}$ che coincide con $\mathbf{x}'$ su $\mathbf{V}$ e

vale zero su $\mathbf{V}^\perp$. Allora $\psi(P_{\mathbf{V}}^\# \mathbf{x}')$ è il funzionale in $\mathbf{X}'$ tale che, per ogni vettore $\mathbf{y} \in \mathbf{X}$,

$$\psi(P_{\mathbf{V}}^\# \mathbf{x}')(\mathbf{y}) = \psi(\mathbf{x}'|_{\mathbf{V}})(\mathbf{y}) = \mathbf{z}'(\mathbf{y}).$$

Quindi, se $\mathbf{x}'$ coincide con la propria proiezione su $\mathbf{V}$, cioè se è nullo su $\mathbf{V}^\perp$, allora su $\mathbf{x}'$ l'applicazione $P_{\mathbf{V}}^\#$ è l'inverso destro di ψ ; in generale invece $P_{\mathbf{V}}^\#$ coincide con l'inverso destro di ψ solo modulo l'immagine della proiezione $P_{\mathbf{V}}^\#$, cioè modulo il sottospazio dei funzionali lineari su $\mathbf{X}'$ che si annullano su $\mathbf{V}$ (in altre parole, sulla restrizione a $\mathbf{V}$).

□

PROPOSIZIONE 10.3.4. *Sia $\mathbf{X} = \mathbf{V} \oplus \mathbf{W}$. Allora $\mathbf{X}' = \mathbf{V}' \oplus \mathbf{W}'$.*

DIMOSTRAZIONE. Anche in questa dimostrazione consideriamo un funzionale lineare $\mathbf{v}' \in \mathbf{V}'$ come esteso ad un funzionale su $\mathbf{X}$, seguendo la terminologia (inaccurata ma chiara) della Definizione 10.3.2.

Dobbiamo provare che $\mathbf{V}' + \mathbf{W}' = \mathbf{X}'$, e che ogni funzionale in $\mathbf{X}'$ si decompone in un unico modo come somma di un termine in $\mathbf{V}'$ ed uno in $\mathbf{W}'$. Per questo basta ricordare che ogni elemento $\mathbf{x}' \in \mathbf{X}'$ è indotto in modo unico da un vettore $\mathbf{x} \in \mathbf{X}$ tramite l'isomorfismo considerato nell'Nota 10.2.7. D'altra parte, ogni $\mathbf{x} \in \mathbf{X}$ si decompone *in modo unico* come $\mathbf{x} = \mathbf{v} + \mathbf{w}$ per qualche $\mathbf{v} \in \mathbf{V}$, $\mathbf{w} \in \mathbf{W}$, per la Definizione 8.0.4 (ii) ed il fatto che $\mathbf{V}$ e $\mathbf{W}$ sono sottospazi disgiunti di $\mathbf{X}$. Perciò il funzionale $\mathbf{x}'$ su $\mathbf{X}$ indotto da $\mathbf{x}$ tramite l'isomorfismo della Nota 10.2.7 si spezza come $\mathbf{x}' = \mathbf{v}' + \mathbf{w}'$, e lo spezzamento è unico.

Ora consideriamo il duale di complementi ortogonali.

PROPOSIZIONE 10.3.5. *Se $\mathbf{V}$ è un sottospazio vettoriale di uno spazio $\mathbf{X}$ munito di prodotto scalare, allora $(\mathbf{V}^\perp)' = (\mathbf{V}')^\perp$.*

DIMOSTRAZIONE. Utilizziamo il prodotto scalare indotto sul duale introdotto nella Definizione 10.3.1, con lo stesso abuso di notazione spiegato nella Definizione 10.3.2:

$$\begin{aligned} (\mathbf{V}')^\perp &= \{\mathbf{x}' \in \mathbf{X}' : \langle \mathbf{x}', \mathbf{v}' \rangle = 0 \ \forall \ \mathbf{v}' \in \mathbf{V}'\} \\ &= \{\mathbf{x}' \in \mathbf{X}' : \langle \mathbf{x}, \mathbf{v} \rangle = 0 \ \forall \ \mathbf{v} \in \mathbf{V}\} = (\mathbf{V}^\perp)'. \end{aligned}$$

Dall'identità (8.0.1) e dalla precedente Proposizione segue

COROLLARIO 10.3.6. *Per ogni sottospazio $\mathbf{V} \subset \mathbf{X}$, si ha $\mathbf{X}' = \mathbf{V}' \oplus (\mathbf{V}')^\perp$.*

NOTA 10.3.7. Per uso futuro, conviene ora osservare che, dato un prodotto scalare su $\mathbf{X}$, l'immersione di uno spazio vettoriale nel suo duale che esso induce nel senso della Definizione 10.3.2 manda il complemento ortogonale di un sottospazio $\mathbf{V}$ in

$$\{\mathbf{x}' : \ker(\mathbf{x}') \supset \mathbf{V}\} = \{\mathbf{x}' : \mathbf{x}'(\mathbf{v}) = 0 \forall \mathbf{v} \in \mathbf{V}\}.$$

Infatti

$$\mathbf{V}^\perp = \{\mathbf{x} : \langle \mathbf{x}, \mathbf{v} \rangle = 0 \forall \mathbf{v} \in \mathbf{V}\} = \{\mathbf{x}' : \mathbf{x}'(\mathbf{v}) = 0 \forall \mathbf{v} \in \mathbf{V}\}$$

(l'ultima identità segue dalla forma dell'isomorfismo nella Definizione 10.3.2). $\square$

Parte 2

Geometria analitica e proiettiva

CAPITOLO 11

Vettori e geometria euclidea ed analitica nel piano

11.1. L'equazione della retta nel piano

Ci sono due modi di scrivere l'equazione di una retta:

- *Equazione parametrica:* la retta che passa per il punto $\mathbf{p}$ ed ha la direzione del vettore $\mathbf{v}$ ha equazione

$$\mathbf{r}(t) = \mathbf{p} + t\mathbf{v}$$

con $t \in \mathbb{R}$.

- *Equazione cartesiana:* consideriamo due punti distinti $\mathbf{p} = (p_1, p_2)$ e $\mathbf{q} = (q_1, q_2)$ non verticalmente allineati, cioè tali che $p_1 \neq q_1$. La retta che passa per questi due punti ha pendenza $m = \frac{q_2 - p_2}{q_1 - p_1}$, e quindi i suoi punti (x_1, x_2) , per proporzionalità, soddisfano l'equazione $x_2 - p_2 = m(x_1 - p_1)$, cioè

$$x_2 = p_2 - m(x_1 - p_1). \quad (11.1.1)$$

Se i due punti sono allineati verticalmente, allora si scambia il ruolo di ascisse ed ordinate e si ottiene l'equazione

$$x_1 = p_1 - m(x_2 - p_2), \quad (11.1.2)$$

dove ora il denominatore è necessariamente non nullo.

Si noti che l'equazione cartesiana è del tipo

$$ax_1 + bx_2 = d \quad (11.1.3)$$

dove a , b e d sono costanti reali. Questa formulazione ha il vantaggio che non contiene denominatori, e quindi non la si deve cambiare se la retta è verticale.

NOTA 11.1.1. (*Equazione della retta in forma normale.*) Si osservi che, se poniamo $\mathbf{n} = (a, b)$ e $\mathbf{x} = (x_1, x_2)$, l'equazione (11.1.3) si scrive $\mathbf{n} \cdot \mathbf{x} = d$. Se $d = 0$, allora la retta è quella perpendicolare a $\mathbf{n}$

che passa per l'origine (perché se $d = 0$ l'origine soddisfa l'equazione). Se invece $d \neq 0$, allora la retta è costituita dai punti la cui proiezione sulla direzione di $\mathbf{n}$ vale $d \frac{\mathbf{n}}{\|\mathbf{n}\|}$ (lo si dimostri per esercizio), ossia è la parallela a quella precedente a distanza $\frac{|d|}{\|\mathbf{n}\|}$ in direzione $\mathbf{n}$ (quindi a distanza $\frac{|d|}{\|\mathbf{n}\|}$ dall'origine).

In particolare, se si normalizza il vettore $\mathbf{n}$, il coefficiente al lato di destra dell'equazione della retta rappresenta la distanza dall'origine; l'equazione così normalizzata si chiama *equazione normale*:

$$\frac{\mathbf{n}}{\|\mathbf{n}\|} \cdot \mathbf{x} = \frac{d}{\|\mathbf{n}\|}.$$

In forma parametrica, per equazione normale si intende quella in cui il vettore direzionale è normalizzato ed il punto di partenze (per $t = 0$) è un vettore perpendicolare alla retta:

$$\mathbf{r}(t) = \mathbf{p} - P_{\mathbf{v}}(\mathbf{p}) + t \frac{\mathbf{v}}{\|\mathbf{v}\|} = \mathbf{p} - \frac{\mathbf{p} \cdot \mathbf{v}}{\|\mathbf{v}\|^2} \mathbf{v} + t \frac{\mathbf{v}}{\|\mathbf{v}\|}$$

dove $t \in \mathbb{R}$ e $P_{\mathbf{v}}(\mathbf{p}) = \frac{\mathbf{p} \cdot \mathbf{v}}{\|\mathbf{v}\|^2} \mathbf{v}$ è la proiezione del vettore $\mathbf{p}$ lungo $\mathbf{v}$. La distanza della retta dall'origine è la proiezione di un suo punto qualsiasi, ad esempio $\mathbf{p}$, perpendicolarmente a $\mathbf{v}$, e quindi è

$$\left\| \mathbf{p} - \frac{\mathbf{p} \cdot \mathbf{v}}{\|\mathbf{v}\|^2} \mathbf{v} \right\|.$$

□

ESEMPIO 11.1.2. (*Proiezione ortogonale di un punto su una retta.*) Calcoliamo la proiezione ortogonale $\mathbf{p}$ del punto $\mathbf{x}$ sulla retta $\mathbf{R}$ di equazione $ax + by = d$. Sappiamo dalla Nota 11.1.1 che il versore $\mathbf{n} = \frac{1}{\sqrt{a^2+b^2}}(a, b)$ è il versore normale a $\mathbf{R}$. Possiamo allora trovare $\mathbf{p}$ in due modi equivalenti.

Il primo consiste nel considerare la retta uscente da $\mathbf{x}$ in direzione $\mathbf{n}$, cioè di equazione parametrica $\mathbf{r}(t) = \mathbf{p} + t\mathbf{n}$, e trovare il punto di intersezione con $\mathbf{R}$: questo significa imporre che il punto $\mathbf{r}(t)$ soddisfi l'equazione $ax + by = d$, trovare il valore t_0 di t (positivo o negativo a seconda del verso del versore normale) per cui questo succede e ricavare $\mathbf{p} = \mathbf{r}(t_0)$.

Il secondo modo consiste nello scomporre $\mathbf{x}$ rispetto ad una base ortonormale in cui uno dei vettori sia parallelo a $\mathbf{R}$. Per semplicità

rinormalizziamo i coefficienti dell'equazione della retta $ax + by = d$ in maniera da scriverla come $ax + by = d'$, dove ora $a^2 + b^2 = 1$ e $d' = d/\sqrt{a^2 + b^2}$. Trattiamo prima il caso in cui $\mathbf{R}$ passi per l'origine, cioè sia $d = d' = 0$. Allora il versore normale a $\mathbf{R}$ è $\mathbf{n}_1 = (a, b)$. Un versore perpendicolare a $\mathbf{n}_1$, e quindi giacente lungo $\mathbf{R}$, è $\mathbf{n}_2 = (-b, a)$. Da (6.6.1) abbiamo la decomposizione ortogonale $\mathbf{x} = \mathbf{x} \cdot \mathbf{n}_1 \mathbf{n}_1 + \mathbf{x} \cdot \mathbf{n}_2 \mathbf{n}_2$. Allora $\mathbf{p} := \mathbf{x} - \mathbf{x} \cdot \mathbf{n}_1 \mathbf{n}_1$ è la proiezione ortogonale cercata: infatti $\mathbf{p}$ giace in $\mathbf{R}$ ed è tale che $\mathbf{x} - \mathbf{p}$ è proporzionale a $\mathbf{n}_1$ e quindi perpendicolare a $\mathbf{R}$.

Nel caso generale in cui $d' \neq 0$ il procedimento precedente ci dà la proiezione $\mathbf{p}'$ di $\mathbf{x}$ non su $\mathbf{R}$ ma sulla sua parallela $\mathbf{R}'$ che passa per l'origine. Ora la proiezione $\mathbf{p}$ su $\mathbf{R}$ si ottiene spostandosi perpendicolarmente a $\mathbf{R}'$ di una distanza $|d'|$:

$$\mathbf{p} = \mathbf{p}' + d' \mathbf{n}_1$$

(nel secondo membro abbiamo preso la somma invece della differenza perché, se ad esempio $a, b > 0$ allora la retta $\mathbf{R}$ è a destra di $\mathbf{R}'$ se $d' > 0$ (cioè $d > 0$) e a sinistra altrimenti, quindi per spostarsi da $\mathbf{R}$ a $\mathbf{R}'$ bisogna muoversi nella direzione di $\mathbf{n}_1$ se $d' > 0$ e al contrario se $d' < 0$).

In particolare, se una retta ha equazione $\mathbf{n} \cdot \mathbf{x} = d$ (e quindi equazione normale $\mathbf{n} \cdot \mathbf{x} / \|\mathbf{n}\| = d / \|\mathbf{n}\|$), la distanza di un punto $\mathbf{p}$ dalla retta è

$$\left| \frac{\mathbf{n} \cdot \mathbf{p} - d}{\|\mathbf{n}\|} \right|. \quad (11.1.4)$$

Infatti, per quanto osservato sopra, questa è la differenza fra le distanze dall'origine della retta data e della sua parallela che passa per $\mathbf{p}$.

□

ESEMPIO 11.1.3. (*Base e altezza di un triangolo.*)

Dati i tre vertici $\mathbf{x}_i$, con $i = 1, 2, 3$, consideriamo la base del triangolo data dal segmento che unisce $\mathbf{x}_1$ a $\mathbf{x}_2$. Per tracciare l'altezza dal vertice $\mathbf{x}_3$ basta applicare il procedimento del precedente Esempio 11.1.2. La base giace sulla retta di equazione parametrica

$$\mathbf{r}(t) = \mathbf{x}_1 + t(\mathbf{x}_2 - \mathbf{x}_1)$$

con $t \in \mathbb{R}$, e (se abbiamo scelto $\mathbf{x}_1 = (x', y')$ e $\mathbf{x}_2 = (x'', y'')$ non allineati verticalmente, cioè $x' \neq x''$) l'equazione cartesiana (11.1.1)

$$y - y' = m(x - x')$$

con coefficiente angolare

$$m = \frac{y'' - y'}{x'' - x'}$$

(se i due punti sono allineati verticalmente si deve scambiare il ruolo delle coordinate x e y in queste uguaglianze, come facemmo in (11.1.2)).

In ciascuna di queste formulazioni applichiamo il corrispondente metodo dell'Esempio precedente per trovare il piede dell'altezza.

□

ESEMPIO 11.1.4. (*Equazioni delle bisettrici di due rette.*) Grazie a (11.1.4), date due rette $\mathbf{n}_1 \cdot \mathbf{x} = d_1$ e $\mathbf{n}_2 \cdot \mathbf{x} = d_2$, le loro due bisettrici hanno equazione

$$\left| \frac{\mathbf{n}_1 \cdot \mathbf{x} - d_1}{\|\mathbf{n}_1\|} \right| = \left| \frac{\mathbf{n}_2 \cdot \mathbf{x} - d_2}{\|\mathbf{n}_2\|} \right| \quad (11.1.5)$$

e cioè

$$\frac{\mathbf{n}_1 \cdot \mathbf{x} - d_1}{\|\mathbf{n}_1\|} = \frac{\mathbf{n}_2 \cdot \mathbf{x} - d_2}{\|\mathbf{n}_2\|}$$

e

$$\frac{\mathbf{n}_1 \cdot \mathbf{x} - d_1}{\|\mathbf{n}_1\|} = -\frac{\mathbf{n}_2 \cdot \mathbf{x} - d_2}{\|\mathbf{n}_2\|}.$$

□

ESERCIZIO 11.1.5. Scrivere le equazioni delle bisettrici della seguente coppia di rette:

$$\begin{aligned} 2x + 3y &= 1 \\ 3x + 2y &= 1. \end{aligned}$$

□

11.2. Cerchi nel piano

Il cerchio di centro $\mathbf{c}$ e raggio $r > 0$ è l'insieme dei punti $\mathbf{x}$ a distanza da $\mathbf{c}$ uguale a r , quindi la sua equazione è

$$\|\mathbf{x} - \mathbf{c}\| = r$$

cioè

$$(x_1 - c_1)^2 + (x_2 - c_2)^2 = r^2. \quad (11.2.1)$$

L'equazione è quadratica, non lineare, e quindi esula dal dominio dell'algebra lineare, ma certamente non da quello della geometria: pertanto la studiamo brevemente qui.

Sviluppando i quadrati otteniamo:

$$x^2 + y^2 + 2c_1x + 2c_2y + c_1^2 + c_2^2 - r^2 = 0.$$

Viceversa, consideriamo più generale l'equazione quadratica

$$x^2 + y^2 + a_1x + a_2y + b = 0.$$

Allora, ponendo

$$\begin{aligned} c_i &= \frac{a_i}{2} \\ r^2 &= c_1^2 + c_2^2 - b, \end{aligned} \quad (11.2.2)$$

questa equazione quadratica si riconduce a (11.2.1) se $b \leq c_1^2 + c_2^2$ (nel caso si abbia l'uguaglianza il cerchio consiste di un solo punto), ed invece non ha soluzioni reali se $b > c_1^2 + c_2^2$.

L'equazione parametrica del cerchio è

$$\begin{aligned} x &= c_1 + r \cos t \\ y &= c_2 + r \sin t \end{aligned} \quad (11.2.3)$$

dove $0 \leq t < 2\pi$.

11.3. Esercizi di geometria analitica nel piano

In vari punti di questa Sezione useremo il seguente risultato, troppo banale per essere elencato come esercizio:

NOTA 11.3.1. Dati due vettori $\mathbf{a}$ e $\mathbf{b}$, il *punto medio* del segmento da essi sotteso è la loro media aritmetica $\mathbf{c} = \mathbf{a} + \mathbf{b}$. Analogamente, per ripartire il suddetto segmento in N parti, si considerano i punti intermedi $\mathbf{c}_1, \dots, \mathbf{c}_{N-1}$ definiti da $\mathbf{c}_i = \mathbf{a} + \frac{i}{N}(\mathbf{b} - \mathbf{a})$ (per $i = 0$ si ritrova $\mathbf{a}$ e per $i = N$ si ritrova $\mathbf{b}$, per i valori intermedi di i si ottengono i vettori $\mathbf{c}_i$). $\square$

11.3.1. Rette.

ESERCIZIO 11.3.2. (*Asse di un segmento, ovvero retta equidistante da due punti dati.*) Trovare l'equazione della retta equidistante dai punti $(1, 2)$ e $(2, 4)$.

SUGGERIMENTO. La retta è quella che passa per il punto medio del segmento sotteso da $(1, 2)$ e $(2, 4)$ e con vettore direzionale perpendicolare a questo segmento. $\square$

ESERCIZIO 11.3.3. Trovare a e c tali che le equazioni lineari

$$ax + y + 2 = 0$$

$$4x + 2y + c = 0$$

abbiano le stesse soluzioni. Qual è il significato geometrico della soluzione?

SVOLGIMENTO. Basta prendere $a = 2$ e $c = 4$ e le due equazioni diventano multiple l'una dell'altra, quindi coincidenti. Le equazioni rappresentano rette nel piano: con la scelta $a = 2$ e $c = 4$ le due rette coincidono. $\square$

ESERCIZIO 11.3.4. Trovare a e c tali che il sistema lineare

$$ax + y + 2 = 0$$

$$4x + 2y + c = 0$$

sia impossibile, e spiegare il significato geometrico della soluzione.

SVOLGIMENTO. Se si prende $a = 2$ e $c \neq 4$ e le parti lineari delle due equazioni diventano multiple l'una dell'altra, ma i coefficienti costanti non rispettano la stessa proporzione, quindi il sistema è impossibile. Geometricamente, le due equazioni rappresentano due rette parallele

ma non coincidenti, che quindi non si intersecano: le soluzioni del sistema sono proprio le intersezioni delle due rette, quindi in questo caso non esistono. $\square$

ESERCIZIO 11.3.5. Sia $\mathbf{r}$ la retta di equazione $3x + y = 1$. Trovare:

- (i) l'equazione della retta simmetrica di $\mathbf{r}$ rispetto all'asse delle x ;
- (ii) l'equazione della retta simmetrica di $\mathbf{r}$ rispetto all'asse delle y ;
- (iii) l'equazione della retta simmetrica di $\mathbf{r}$ rispetto all'origine.

SVOLGIMENTO.

- (i) L'equazione della retta simmetrica rispetto all'asse x deve ottenersi cambiando x in $-x$ nell'equazione di $\mathbf{r}$, quindi è $-3x + y = 1$.

Un modo più diretto (ma più lungo) di procedere è il seguente. Per prima cosa si scrive $\mathbf{r}$ in forma parametrica. Un suo vettore normale è $(3, 1)$, quindi un suo vettore direzionale è $(1, -3)$, e $\mathbf{r}$ passa per $(0, 1)$, quindi l'equazione è $\mathbf{r}(t) = (0, 1) + t(1, -3)$. A questo punto per ottenere l'equazione parametrica della retta riflessa basta cambiare il segno delle componenti x dei vettori. Dall'equazione parametrica si ricava facilmente quella cartesiana.

Un terzo modo è di determinare due punti per cui passa $\mathbf{r}$, ad esempio le intersezioni con gli assi $(0, 1)$ e $(\frac{1}{3}, 0)$, ribaltare questi rispetto all'asse x in $(0, -1)$ e $(\frac{1}{3}, 0)$ e trovare l'equazione della retta per questi nuovi due punti.

- (ii) Ragionando come prima si trova l'equazione $3x - y = 1$.
- (iii) $-3x - y = -1$, cioè $3x + y = 1$. Si noti che questa è proprio la retta parallela alla precedente che passa alla stessa distanza dall'origine ma dall'altro lato.

$\square$

ESERCIZIO 11.3.6. Trovare due rette perpendicolari che passano per il punto $(2, 3)$ e determinano sull'asse x un segmento

- (i) di lunghezza 4;

(ii) di lunghezza minima.

Trovare la distanza minima di un punto (x_0, y_0) dall'asse x per la quale esistono due rette ortogonali che passano per (x_0, y_0) e staccano sull'asse x un segmento di lunghezza 4.

SVOLGIMENTO. Le rette che passano per $(2, 3)$ hanno equazione $ax + by = 2a + 3b$, con a, b fissati. Per ciascuna di tali rette, la famiglia di rette ad essa ortogonali si ottiene a partire da un vettore ortogonale al suo vettore normale (a, b) : le equazioni di queste rette ortogonali sono $bx - ay = d$, con $d \in \mathbb{R}$. Fra queste, quella che passa per $(2, 3)$ è $bx - ay = 2b - 3a$. Le intersezioni delle due rette ortogonali con l'asse x si ottengono ponendo $y = 0$ nelle loro equazioni, e sono $x_1 = \frac{2a+3b}{a}$ e $x_2 = \frac{2b-3a}{b}$ (se a oppure b sono zero una delle due rette è parallela all'asse x e quindi non c'è intersezione: possiamo quindi assumere che entrambi i denominatori siano non nulli). Pertanto la lunghezza del segmento tagliato dalle due rette sull'asse x è

$$|x_2 - x_1| = \left| \frac{2b-3a}{b} - \frac{2a+3b}{a} \right| = 3 \frac{a^2 + b^2}{|ab|}. \quad (11.3.1)$$

Fissato un valore di a , questo segmento ha lunghezza 4 per il valore di b che risolve l'equazione $3(a^2 + b^2) - 4ab = 0$, la quale ha due soluzioni reali perché il discriminante $16a^2 - 12a^2$ è positivo: questo risolve la parte (i) dell'esercizio. Per la parte (ii), osserviamo che la lunghezza minima del segmento si ricava prendendo il minimo rispetto ad a e b del lato di destra di (11.3.1). Ora, $a^2 + b^2 = (a - b)^2 + 2ab$, e quindi il lato di destra di (11.3.1) si riscrive come $2 + \frac{(a-b)^2}{|ab|}$: evidentemente questa espressione ha minimo per $b = a$.

Infine, le due rette ortogonali che passano per (x, y) si scrivono come prima: $ax + by = ax_0 + by_0$ e $bx - ay = bx_0 - ay_0$. La lunghezza del segmento staccato sull'asse x da queste due rette si ricava come in (11.3.1), e vale

$$\left| \frac{bx_0 - ay_0}{b} - \frac{ax_0 + by_0}{a} \right| = |y_0| \frac{a^2 + b^2}{|ab|}$$

ed il minimo di queste distanze si trova imponendo che l'equazione quadratica $|y_0|(a^2 + b^2) = 4|ab|$, per un dato a , abbia una sola soluzione in b . Questo equivale ad imporre che il discriminante sia nullo. Possiamo

assumere $y_0 > 0$ (ci sono due soluzioni simmetriche, una con $y_0 > 0$ e l'altra data da $-y_0$). Lasciamo i calcoli al lettore. $\square$

ESERCIZIO 11.3.7. Trovare la retta che passa per il punto $(1, 1)$ e determina con i semiassi x, y positivi un triangolo di area minima.

SVOLGIMENTO. L'equazione delle rette attraverso il punto $(1, 1)$ è $ax + by = a + b$, con a, b costanti reali. Le intersezioni con gli assi sono rispettivamente $x = \frac{a+b}{a}$ e $y = \frac{a+b}{b}$ (possiamo assumere che a e b siano non nulli, perché altrimenti la retta è parallela ad uno degli assi, e quindi non determina con essi un triangolo). Quindi l'area del triangolo è $\frac{(a+b)^2}{2|ab|}$. Come dimostrato alla fine del precedente Esercizio 11.3.6, il minimo si ha quando $a = b$, cioè per la retta $x + y = 2$. $\square$

ESERCIZIO 11.3.8. Trovare i punti della retta $y = 2x$ che sono a distanza 3 dal punto $(1, 1)$.

SVOLGIMENTO. Sia x l'ascissa del punto cercato; allora la condizione di appartenere alla retta impone che il punto abbia coordinate $(2x, x)$. La condizione di essere a distanza 3 da $(1, 1)$ equivale a richiedere la condizione $(2x-1)^2 + (x-1)^2 = 9$. Sviluppando il primo membro troviamo una equazione quadratica con due soluzioni reali, che individuano i due punti cercati.

Questo procedimento equivale a mettere a sistema l'equazione (lineare) della retta con l'equazione (quadratica) della sfera di raggio 3 e centro $(1, 1)$. $\square$

11.3.2. Triangoli.

ESERCIZIO 11.3.9. (*Triangoli equilateri.*) Trovare i triangoli equilateri che hanno due vertici nei punti $(1, 2)$ e $(2, 4)$.

SUGGERIMENTO. Il terzo vertice deve trovarsi sull'asse del segmento sotteso da $(1, 2)$ e $(2, 4)$. Una volta scritta l'equazione di questa retta (come nel precedente Esercizio 11.3.2), basta determinare i punti su di essa a distanza da, diciamo, $(1, 2)$ uguale alla lunghezza del suddetto segmento (cioè $\sqrt{5}$). Ci sono due tali punti, e quindi due soluzioni del presente problema. $\square$

ESERCIZIO 11.3.10. (Mediane di un triangolo.) Dato il triangolo di vertici $(0, 0)$, $(-1, 3)$ e $(4, 4)$ trovare le tre mediane, cioè i segmenti che congiungono ciascun vertice col punto medio del lato opposto. Mostrare che le tre mediane si incontrano in un punto, determinando tale punto.

□

ESERCIZIO 11.3.11. (*Bisettrici di un triangolo.*) Rispondere alle domande dell'esercizio precedente per le bisettrici anziché per le mediane.

SUGGERIMENTO. Utilizzare le equazioni (11.1.5).

□

ESERCIZIO 11.3.12. (Altezze di un triangolo.) Dato il triangolo di vertici $(0, 0)$, $(-1, 3)$ e $(4, 4)$ trovare le tre altezze, cioè i segmenti da ciascun vertice che incrociano perpendicolarmente il lato opposto. Mostrare che le tre altezze si incontrano in un punto, determinando tale punto.

SUGGERIMENTO. Per ogni vertice, si scriva l'equazione della retta che contiene il lato opposto, del tipo $ax + by + c = 0$. Il vettore normale a questa retta è (a, b) ; un vettore perpendicolare a questo è $(b - a)$. Perciò la famiglia di rette perpendicolari al lato opposto è data da $bx - ay + d = 0$ al variare di d in $\mathbb{R}$. Trovare il valore di d per una di cui queste rette passa per il vertice dato.

□

11.3.3. Parallelogrammi.

ESERCIZIO 11.3.13. (*Parallelogrammi.*) (*Punti medi della diagonale di un parallelogramma.*)

- (i) Mostrare che le due diagonali di un parallelogramma si incontrano nel loro punto medio, e più in generale che un quadrilatero è un parallelogramma se e solo se le diagonali hanno lo stesso punto medio.
- (ii) Mostrare che il quadrilatero di vertici $(0, 0)$, $(-1, 3)$, $(4, 4)$ e $(3, -1)$ non è un parallelogramma.
- (iii) Mostrare che se i vertici $(0, 0)$ e $(4, 4)$ sono vertici opposti di un parallelogramma, gli altri due vertici devono essere $(t, 4 - t)$

e $(4-t, t)$ per qualche $t \in \mathbb{R}$ (nel caso $t = 2$ il parallelogramma degenera in un segmento).

SVOLGIMENTO. Proviamo con metodi di geometria analitica, cioè facendo uso delle coordinate, questo celebre enunciato che in geometria euclidea segue dalle proprietà dei triangoli simili individuati da ciascuna diagonale. Senza perdita di generalità supponiamo che il primo vertice $\mathbf{v}_0$ di un parallelogramma sia l'origine ed il primo lato giaccia sull'asse x : il vertice $\mathbf{v}_1$ adiacente all'origine ha coordinate che indichiamo con $(x, 0)$. Sia $\mathbf{v}_2 = (x+h, y)$ il vertice successivo. Allora il punto medio della diagonale da $\mathbf{v}_0$ a $\mathbf{v}_2$ è $\mathbf{c} = (\frac{x+h}{2}, \frac{y}{2})$. Il quarto vertice ha la stessa ordinata y del terzo, visto che il parallelogramma ha il lato opposto sull'asse x ; la sua ascissa è uguale alla differenza fra le ascisse di $\mathbf{v}_2$ e $\mathbf{v}_1$, perché il lato sotteso da questi due vertici deve essere parallelo a quello dall'origine a $\mathbf{v}_3$. Quindi $\mathbf{v}_3 = (h, y)$, ed il punto medio della diagonale da $\mathbf{v}_1$ a $\mathbf{v}_3$ è ancora $\mathbf{c} = (\frac{x+h}{2}, \frac{y}{2})$.

Viceversa, se si parte con un quadrilatero col primo vertice nell'origine ed il primo lato sull'asse x , chiamiamo il secondo vertice $\mathbf{v}_1 = (x, 0)$, il terzo $\mathbf{v}_2 = (x+h, y)$ e sia $\mathbf{c} = (\frac{x+h}{2}, \frac{y}{2})$ il punto medio della diagonale da $\mathbf{v}_0$ a $\mathbf{v}_2$. Se questo punto è anche il punto medio dell'altra diagonale, allora il quarto vertice $\mathbf{v}_3$ deve soddisfare l'identità $\mathbf{v}_3 = \mathbf{v}_1 + 2(\mathbf{c} - \mathbf{v}_1) = 2\mathbf{c} - \mathbf{v}_1 = (h, y)$, e quindi il quadrilatero è un parallelogramma.

Le altre due proprietà sono dirette conseguenze della prima.

□

ESERCIZIO 11.3.14. I punti $(1, 5)$ e $(4, 2)$ sono due vertici adiacenti di un parallelogramma; il punto di incontro delle diagonali è $(3, 2)$. Trovare gli altri due vertici. □

ESERCIZIO 11.3.15. (*I parallelogrammi con diagonali perpendicolari sono i rombi.*) Mostrare che un parallelogramma è un rombo se e solo se le diagonali si intersecano ad angolo retto (in particolare, un rettangolo è un quadrato se e solo se vale tale condizione di ortogonalità).

SVOLGIMENTO. Questo esercizio si può dimostrare elegantemente con la geometria euclidea. Un verso è chiaro. Infatti, se il parallelogramma è un rombo, ogni tre vertici consecutivi formano un triangolo isoscele, dove il terzo lato è una diagonale, e la mediana rispetto a tale lato giace sull'altra diagonale. Per simmetria questa mediana è anche un'altezza, cioè incontra la base (l'altra diagonale) ad un angolo retto. Proviamo quindi il viceversa. Chiamiamo i vertici, presi in ordine consecutivo, $\mathbf{v}_0, \mathbf{v}_1, \mathbf{v}_2, \mathbf{v}_3$. Sia $\mathbf{c}$ il punto medio delle due diagonali (è lo stesso punto per entrambe perché stiamo considerando un parallelogramma). Se le diagonali si incrociano ad angolo retto, applichiamo il teorema di Pitagora dapprima al triangolo $\mathbf{v}_0, \mathbf{c}, \mathbf{v}_3$ ed otteniamo che la lunghezza del lato da $\mathbf{v}_0$ a $\mathbf{v}_3$ è la radice quadrata della somma dei quadrati delle lunghezze delle due semidiagonali. Lo stesso ragionamento vale per gli altri tre lati, ma le lunghezze delle semidiagonali sono le stesse in tutti e quattro i casi, e quindi i quattro lati sono uguali.

Svolgiamo la seconda parte dell'esercizio con metodi di geometria analitica. Senza perdita di generalità supponiamo che il primo vertice $\mathbf{v}_0$ sia l'origine ed il primo lato giaccia sull'asse x . Come nell'Esercizio 11.3.13, scriviamo gli altri tre vertici come $\mathbf{v}_1 = (x, 0)$, $\mathbf{v}_2 = (x + h, y)$ e $\mathbf{v}_3 = (h, y)$. Il punto medio della diagonale da $\mathbf{v}_0$ a $\mathbf{v}_2$ è $\mathbf{c} = (\frac{x+h}{2}, \frac{y}{2})$. Pertanto il segmento da $\mathbf{v}_0$ a $\mathbf{c}$ è dato dal vettore $(\frac{x+h}{2}, \frac{y}{2})$, mentre il segmento da $\mathbf{c}$ a $\mathbf{v}_3$ è dato dal vettore $(h, y) - (\frac{x+h}{2}, \frac{y}{2}) = (\frac{h-x}{2}, \frac{y}{2})$. Ora, applicando il teorema di Pitagora come sopra, troviamo che il quadrato della lunghezza del lato da $\mathbf{v}_0$ a $\mathbf{v}_3$ è

$$h^2 + y^2 = [(x+h)^2/4 + y^2/4] + [(x-h)^2/4 + y^2/4] \quad (11.3.2)$$

$$= \frac{1}{4} [(x-h)^2 + (x+h)^2 + 2y^2] . \quad (11.3.3)$$

Svolgendo i quadrati in questa relazione arriviamo all'identità $h^2 + y^2 = x^2$, cioè $\sqrt{h^2 + y^2} = |x|$. Quindi il lato da $\mathbf{v}_0$ a $\mathbf{v}_3$ è lungo quanto il lato da $\mathbf{v}_0$ a $\mathbf{v}_1$, e perciò il parallelogramma ha due lati adiacenti della stessa lunghezza: allora, per similitudine, sono tutti e quattro della stessa lunghezza, ed il parallelogramma è un rombo.

Questa dimostrazione analitica è più facile se si suppone che il parallelogramma sia un rettangolo e si vuol dimostrare che l'ortogonalità delle diagonali lo forza ad essere un quadrato: in questo caso $h = 0$ e l'identità (11.3.2) dà $|x| = |y|$, quindi il parallelogramma ha tutti i lati uguali. $\square$

ESERCIZIO 11.3.16. (*Lunghezze di lati e diagonali di un parallelogramma.*) Mostrare che la somma dei quadrati delle lunghezze dei lati di un parallelogramma è uguale alla somma dei quadrati delle lunghezze delle diagonali.

SVOLGIMENTO. Come nel precedente Esercizio 11.3.15, supponiamo che il primo vertice sia l'origine ed il primo lato sia sull'asse x , con vertice finale in $(x, 0)$. Chiamiamo di nuovo $(x + h, y)$ le coordinate del secondo vertice. Abbiamo visto che il punto medio della diagonale dall'origine al vertice opposto è $\mathbf{c} = (\frac{x+h}{2}, \frac{y}{2})$, ed il quarto vertice è (h, y) .

Ora, il quadrato della diagonale dal primo al terzo vertice vale $(x + h)^2 + y^2$, e quello della diagonale dal secondo al quarto vertice vale $(x - h)^2 + y^2$, quindi la somma dei quadrati delle diagonali è $2x^2 + 2h^2 + 2y^2$. La somma dei quadrati dei quattro lati è il doppio della somma dei quadrati di due lati adiacenti, e quindi vale $2(x^2 + h^2 + y^2)$: quindi la somma dei quadrati delle diagonali coincide con quella dei lati. $\square$

ESERCIZIO 11.3.17. Mostrare che, se le diagonali di un parallelogramma hanno la stessa lunghezza, il parallelogramma è un rettangolo.

SVOLGIMENTO. Come nell'Esercizio 11.3.15, supponiamo che il primo vertice sia l'origine ed il primo lato sia sull'asse x , con vertice finale in $(x, 0)$, secondo vertice in $(x + h, y)$: come in quell'Esercizio, segue che il quarto vertice è (h, y) . Dobbiamo mostrare che $h = 0$.

Come visto nell'Esercizio 11.3.16, la lunghezza della diagonale dal primo al terzo vertice vale $(x + h)^2 + y^2$, e quello della diagonale dal secondo al quarto vertice vale $(x - h)^2 + y^2$. Perciò le due diagonali hanno la

stessa lunghezza se e solo se $(x+h)^2 = (x-h)^2$, ovvero $xh = 0$. Poiché x è la lunghezza del primo lato, si ha $x \neq 0$, e quindi $h = 0$. $\square$

11.3.4. Cerchi.

ESERCIZIO 11.3.18. Trovare l'equazione del cerchio che passa per i punti $(1, 2)$, $(0, 0)$ e $(3, 3)$.

SUGGERIMENTO. Si considerano due delle rette equidistanti da coppie di questi punti, come nel precedente Esercizio 11.3.2. Intersecando queste due rette troviamo il centro del cerchio; allora il raggio è uguale alla distanza fra questo centro ed uno qualsiasi dei vertici.

Presentiamo un metodo alternativo geometricamente meno elegante, ma che talvolta richiede meno passaggi. Individuiamo il centro (x, y) della sfera imponendo che sia equidistante dai tre punti assegnati. nel caso presente, questo significa imporre le condizioni

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

dalle quali i termini quadratici si cancellano dopo che si sviluppano i quadrati, e quello che resta è il sistema lineare

$$-2x - 4y + 2 = 0 = -6x - 6y + 18$$

cioè

$$x + 2y = 1$$

$$x + y = 3.$$

Il centro del cerchio è la soluzione di questo sistema lineare, data da $x = 5$, $y = -2$. Geometricamente questa soluzione rappresenta il punto di intersezione delle due rette $x + 2y = 1$ e $x + y = 3$. È facile verificare che la prima di queste due rette è proprio quella equidistante dai punti $(1, 2)$ e $(0, 0)$, e la seconda è equidistante da $(0, 0)$ e $(1, 2)$. Perciò questo secondo metodo in effetti equivale a quello precedente. $\square$

ESERCIZIO 11.3.19. (*Triangoli equilateri.*) Trovare gli altri due vertici del triangolo equilatero con un vertice in $(1, 1)$ e lato opposto sulla retta di equazione $x + 3y = 1$.

SUGGERIMENTO. Il modo più semplice è di calcolare il piede $\mathbf{p}$ della proiezione ortogonale del punto $(1, 1)$ sulla retta, e poi trovare i due punti della retta la cui distanza da $\mathbf{p}$ è la metà del segmento della proiezione (che è l'altezza del triangolo). Un altro modo consiste nello scrivere l'equazione della circonferenza con centro $(1, 1)$ e raggio $r > 0$ arbitrario, calcolare le intersezioni di questa circonferenza con la retta data, e trovare r tale che la lunghezza della corda da esse determinata sia proprio r . $\square$

ESERCIZIO 11.3.20. (*Triangoli inscritti in un cerchio.*) I punti $\mathbf{v}_1 = (2, 0)$ e $\mathbf{v}_2 = (-2, 4)$ giacciono nel cerchio di centro $\mathbf{c} = (0, 2)$ e raggio 2 (infatti entrambi sono a distanza 2 da $\mathbf{c}$). Trovare il punto $\mathbf{v} = (x, y)$ su questo cerchio tale che il triangolo così formato abbia area massima, e mostrare che è isoscele rispetto al terzo vertice.

SVOLGIMENTO. Dall'equazione parametrica del cerchio(11.2.3) otteniamo $\mathbf{v} = (2 \cos t, 2 + 2 \sin t)$. Calcoliamo la proiezione ortogonale di questo punto sul segmento da $\mathbf{v}_2$ a $\mathbf{v}_1$. Il vettore direzionale di questo segmento è $\mathbf{v}_1 - \mathbf{v}_2 = (4, -4)$. I versori ortonormali sono $\pm \mathbf{n} = \pm(\frac{1}{\sqrt{2}}, \frac{1}{\sqrt{2}})$. Quindi la proiezione ortogonale cercata è

$$\begin{aligned} \mathbf{h} &= \mathbf{v} - (\mathbf{v} \cdot \mathbf{n})\mathbf{n} = \mathbf{v} - \sqrt{2}(\cos t + \sin t + 1)\mathbf{n} \\ &= (\cos t - \sin t - 1, -\cos(t) + \sin t + 1). \end{aligned}$$

Perciò l'altezza del triangolo relativa al vertice $\mathbf{v}$ è il segmento $\mathbf{h} - \mathbf{v} = -(\cos t + \sin t + 1)$, la cui lunghezza è $\sqrt{2}|1 + \cos t + \sin t|$. Invece la lunghezza della base è $\|\mathbf{v}_1 - \mathbf{v}_2\| = \|(4, -4)\| = 4\sqrt{2}$. Pertanto l'area del triangolo è proporzionale al modulo di $1 + \cos t + \sin t$, che ha massimo quando $\cos t + \sin t$ ha massimo o minimo. Ma $\cos t + \sin t = \frac{1}{2} \sin(2t)$ ha massimo e minimo in $t = \frac{\pi}{4}$ e $t = \frac{3\pi}{4}$, rispettivamente (e così via per periodicità). Si calcolino le aree dei due triangoli così ottenuti e si trovi la più grande delle due (è quella del triangolo il cui vertice $\mathbf{v}$ sta dall'altro lato di $\mathbf{v}_1$ e $\mathbf{v}_2$ rispetto al diametro parallelo a questi due vertici).

Si verifichi che il triangolo così ottenuto è isoscele. $\square$

ESERCIZIO 11.3.21. Si dimostri che il triangolo di area massima iscritto in un cerchio è equilatero.

SVOLGIMENTO. Si fissino arbitrariamente due dei tre vertici: il precedente Esercizio 11.3.20 mostra che il triangolo è isoscele rispetto al terzo vertice. $\square$

ESERCIZIO 11.3.22. Trovare le tangenti al cerchio $x^2 + y^2 - 2x + 2y - 1 = 0$ che si incontrano nel punto $(4, 4)$.

SUGGERIMENTO. Scrivere l'equazione della generica retta che passa per $(4, 4)$ e trovare le intersezioni con il cerchio (mettendo a sistema le due equazioni). I punti di intersezione sono le soluzioni dell'equazione quadratica così ricavata. A seconda del fatto che il discriminante di questa equazione sia positivo, nullo o negativo si ottengono due, una o nessuna soluzione. Le rette tangenti sono quelle che corrispondono ad una sola soluzione: se ne trovano esattamente due.

Un metodo alternativo (ma di fatto equivalente) consiste nello scrivere in forma parametrica, grazie a (11.2.3), i punti del cerchio, che nel caso presente, in base a (11.2.2), è $\mathbf{r}t = (1 + \cos t, 1 + \sin t)$, ed il vettore dal centro del cerchio al punto generico su di esso (che evidentemente è $(\cos t, \sin t)$, e trovare i valori del parametro t per cui questo vettore sia perpendicolare al vettore dal punto $(4, 4)$ al punto del cerchio $\mathbf{r}t$, cioè il vettore $(-3 + \cos t, -3 + \sin t)$. $\square$

ESERCIZIO 11.3.23. Trovare l'equazione della famiglia di cerchi tangenti simultaneamente alle rette $x = y$ e $x = -y$.

SVOLGIMENTO. I cerchi hanno centro sulle bisettrici delle due rette, che in generale si ricavano da (11.1.5) ed in questo caso sono evidentemente $x = 0$ e $y = 0$. Quindi i centri sono parametrizzati da $(t, 0)$ e $(0, s)$, con $t, s \in \mathbb{R}$. La distanza di $(t, 0)$ dalle rette $x = \pm y$ è $|t|/\sqrt{2}$, come si vede subito dal teorema di Pitagora (se le rette non fossero le bisettrici occorrerebbe calcolare il piede della proiezione ortogonale su di esse del punto $(t, 0)$). Quindi l'equazione di questa famiglia di cerchi è $(x - t)^2 + y^2 = t^2/2$, al variare del parametro $t \in \mathbb{R}$. Analogamente,

per i cerchi tangenti alle due bisettrici e centrati sull'asse y si trova l'equazione $x^2 + (y - s)^2 = s^2/2$. $\square$

ESERCIZIO 11.3.24. Scrivere l'equazione del cerchio tangente alla retta $x = y$ nel punto $\mathbf{t} = (1, 1)$ ed anche alla retta $x = 0$.

SVOLGIMENTO. I cerchi tangenti alla retta $x = y$ nel punto $(1, 1)$ hanno centro sulla perpendicolare di questa retta che passa per $(1, 1)$, e quindi sulla retta $y = -x + 2$, ovvero, in forma parametrica, $\mathbf{r}(t) = (t, 2 - t)$. La proiezione ortogonale dei punti della retta $\mathbf{r}$ sulla retta $x = y$ è il punto $(1, 1)$. La distanza di $(t, 2 - t)$ dalla retta $x = y$ è la distanza di $(t, 2 - t)$ da $(1, 1)$, cioè $\sqrt{(t - 1)^2 + (1 - t)^2} = \sqrt{2}|t - 1|$. Invece la distanza dalla retta $x = 0$ evidentemente è $|t|$ (in un caso meno evidente occorrerebbe calcolare la proiezione ortogonale). Quindi il centro del cerchio verifica l'equazione $|t| = \sqrt{2}|t - 1|$, cioè $t^2 = 2(t - 1)^2$, ovvero $t^2 - 4t + 2 = 0$, ed il raggio è $|t|$. Ci sono due soluzioni per ogni $t \neq 0$: una è data da cerchi tangenti a $x = 0$ in punti con y positivo, l'altra con y negativo (in entrambi i casi il cerchio si trova nel semipiano $x \geq 0$). $\square$

ESERCIZIO 11.3.25. Scrivere l'equazione del cerchio simultaneamente tangente alle rette $x = 0$, $y = 0$ e $x + y = 2$.

SUGGERIMENTO. Il centro è il punto di incontro delle bisettrici delle tre rette. In questo caso le bisettrici delle prime due sono le rette $x = \pm y$. I punti della retta $x = y$ sono parametrizzati da (t, t) con $t \in \mathbb{R}$: la loro distanza dalle rette $x = 0$ e $y = 0$ è $|t|$. La proiezione ortogonale di ciascuno di questi punti sulla retta $x + y = 2$ si ottiene spostandosi lungo il versore ortogonale a questa retta, che è $\mathbf{n} = (\frac{1}{\sqrt{2}}, \frac{1}{\sqrt{2}})$, fino ad incontrare la retta $x + y = 2$, e quindi è il punto $(1, 1)$. Quindi la distanza di (t, t) dalla retta $x + y = 1$ è $\sqrt{(1 - t)^2 + (1 - t)^2} = \sqrt{2}|t - 1|$, e la condizione di tangenza simultanea diventa l'equazione $|t| = \sqrt{2}|t - 1|$, ovvero $t^2 - 4t + 2 = 0$.

Invece, se il centro è sulla retta $x = -y$, allora è del tipo $(s, -s)$ con $s \in \mathbb{R}$. La retta $x = -y$ è parallela alla retta $x + y = 2$, ed esattamente a distanza 2 da essa. Quindi in questo caso il centro soddisfa l'equazione $|s| = 2$.

Si noti che, grazie ai valori assoluti, si trovano quattro soluzioni: un cerchio tangente per ciascuno dei tre settori illimitati delimitati dalle tre rette, ed uno nel settore limitato da esse racchiuso. $\square$

ESERCIZIO 11.3.26. Scrivere l'equazione della retta tangente nell'origine al cerchio con centro in $(2, 2)$ e passante per l'origine.

SVOLGIMENTO. La retta tangente è quella perpendicolare al raggio, cioè al segmento dal centro $(1, 12)$ al punto di tangenza $(0, 0)$. Questo raggio giace lungo la retta $y = x$, quindi la tangente è la sua perpendicolare $y = -x$. $\square$

ESERCIZIO 11.3.27. Scrivere l'equazione della tangente nell'origine al cerchio $(x - 1)^2 + (y - 2)^2 = 5$.

SVOLGIMENTO. Il raggio dal centro $(1, 2)$ all'origine è $\mathbf{r} = -(1, 2)$. La retta tangente è ortogonale al raggio e passa per l'origine, quindi il suo vettore normale è $\mathbf{r}$ e l'equazione è $2x - y = 0$. $\square$

ESERCIZIO 11.3.28. Scrivere l'equazione del cerchio passante per l'origine e tangente alla retta $x + 3y = 4$ nel punto $(1, 1)$.

SUGGERIMENTO. Il centro del cerchio si trova sulla retta perpendicolare a $x + 3y = 4$ nel punto $(1, 1)$, cioè sulla retta $3y - x = 2$. Il quadrato della distanza dei punti di questa retta da $(1, 1)$ è $(x - 1)^2 + (y - 1)^2$; il fatto che il cerchio passi anche per l'origine equivale a richiedere che tale distanza sia uguale a quella dall'origine, cioè che si abbia $(x - 1)^2 + (y - 1)^2 = x^2 + y^2$. Il centro del cerchio verifica quindi questa equazione quadratica e l'equazione lineare $3y - x = 2$. $\square$

ESERCIZIO 11.3.29. Scrivere l'equazione del cerchio passante per il punto $(3, 4)$ e tangente alle rette $x + 2y = 1$ e $3x - y = 2$.

SUGGERIMENTO. Il centro del cerchio è equidistante dalle due rette e dal punto. Quindi lo si trova calcolando le proiezioni ortogonali di un generico punto (x, y) sulle due rette, trovandone le rispettive distanze ed uguagliandole fra loro ed alla distanza da $(3, 4)$, cioè alla radice quadrata di $(x - 3)^2 + (y - 4)^2$. $\square$

ESERCIZIO 11.3.30. Dimostrare che i due cerchi $x^2 + y^2 - 2ax - 2by + c = 0$ e $x^2 + y^2 - 2a'x - 2b'y + c' = 0$ si intersecano perpendicolarmente se e solo se $2aa' + 2bb' = c + c'$.

SUGGERIMENTO. I due cerchi si incontrano perpendicolarmente se e solo se il triangolo formato dai due centri e dal punto di intersezione è rettangolo. Siano r, r' i raggi e d la distanza fra i due centri: per il teorema di Pitagora, questo equivale alla condizione $r^2 + r'^2 = d^2$. Si ricavano da (11.2.2) i valori dei raggi e dei centri, e quindi la distanza d . □

CAPITOLO 12

Vettori e geometria euclidea ed analitica in $\mathbb{R}^3$

Alcuni esempi geometrici nello spazio tridimensionale sono analoghi a quelli dati nel piano nel Capitolo 11, ma, anche quando la trattazione è pressoché identica parola per parola, per completezza la riportiamo di nuovo nel presente contesto.

12.1. L'equazione del piano in $\mathbb{R}^3$

- *L'equazione cartesiana* di un piano in tre dimensioni reali è

$$ax_1 + bx_2 + cx_3 = d \quad (12.1.1)$$

dove a, b, c e d sono costanti reali.

Se poniamo $\mathbf{n} = (a, b, c)$ e $\mathbf{x} = (x_1, x_2, x_3)$, l'equazione del piano si scrive

$$\mathbf{n} \cdot \mathbf{x} = d. \quad (12.1.2)$$

In questa forma l'equazione vale anche in qualsiasi dimensione, cioè per iperpiani di codimensione 1 in $\mathbb{R}^n$. Se $d = 0$, allora il piano passa per l'origine, ed il vettore $\mathbf{n}$ è perpendicolare ai punti del piano. Se invece $d \neq 0$ allora il piano è formato dai vettori la cui proiezione sulla direzione di $\mathbf{n}$ vale $d \frac{\mathbf{n}}{\|\mathbf{n}\|}$, cioè è il piano parallelo a quello precedente a distanza $\frac{d}{\|\mathbf{n}\|}$ in direzione $\mathbf{n}$.

In particolare, se tre punti $\mathbf{u}$, $\mathbf{q}$ e $\mathbf{r}$ non sono allineati, cioè non multipli uno dell'altro, allora sostituendo le loro componenti nell'equazione $ax_1 + bx_2 + cx_3 = d$ otteniamo un sistema lineare di tre equazioni nelle incognite a, b, c e d . Questo sistema ha un grado di incertezza, essendo un sistema con una incognita in più del numero di equazioni, ma per ogni valore di d ha un'unica soluzione. Quindi, *a meno di multipli*, c'è un'unica soluzione, che ci dà l'equazione cartesiana del piano. È

del resto ovvio che, per ogni costante α non nulla, le equazioni $ax_1 + bx_2 + cx_3 = d$ e $\alpha ax_1 + \alpha bx_2 + \alpha cx_3 = \alpha d$ descrivono lo stesso piano (si noti anche che i due vettori normali sono multipli uno dell'altro).

- *L'equazione normale* del piano si ottiene normalizzando il vettore perpendicolare in (12.1.2): $\mathbf{n} \cdot \mathbf{x} / \|\mathbf{n}\| = d / \|\mathbf{n}\|$. La distanza dall'origine è quindi $d / \|\mathbf{n}\|$.
- *L'equazione parametrica* del piano in tre dimensioni che contiene le due rette passanti per un punto $\mathbf{u}$ e dirette lungo i vettori rispettivamente $\mathbf{v}$ e $\mathbf{w}$ è

$$\mathbf{p}(s, t) = \mathbf{u} + s\mathbf{v} + t\mathbf{w}$$

dove s, t sono due parametri reali. In particolare, se tre punti $\mathbf{u}, \mathbf{q}$ e $\mathbf{r}$ non sono allineati, allora $\mathbf{q} - \mathbf{u}$ e $\mathbf{r} - \mathbf{u}$ non sono multipli uno dell'altro, e l'equazione parametrica

$$\mathbf{p}(s, t) = \mathbf{u} + s(\mathbf{q} - \mathbf{u}) + t(\mathbf{r} - \mathbf{u})$$

determina un piano, l'unico piano che contiene i tre punti.

ESEMPIO 12.1.1. (*Il piano bisettore di un segmento, cioè equidistante da due punti.*) Siano $\mathbf{a}$ e $\mathbf{b}$ due vettori distinti, e $\mathbf{m} = \frac{1}{2}(\mathbf{a} + \mathbf{b})$ il loro punto medio, come nella Nota 11.3.1. Il piano equidistante da $\mathbf{a}$ e $\mathbf{b}$ è quello che taglia perpendicolarmente a metà il segmento che li congiunge, cioè è ortogonale a $\mathbf{b} - \mathbf{a}$ e passa per $\mathbf{m}$. Quindi questo piano ha per vettore normale il vettore $\mathbf{n} = \mathbf{b} - \mathbf{a}$: la sua equazione è

$$\mathbf{x} \cdot \mathbf{n} = \mathbf{m} \cdot \mathbf{n} = \frac{1}{2}(\mathbf{a} + \mathbf{b}) \cdot (\mathbf{b} - \mathbf{a}) = \frac{1}{2}(\|\mathbf{a}\|^2 - \|\mathbf{b}\|^2).$$

□

ESERCIZIO 12.1.2. Scrivere l'equazione (cartesiana e parametrica) del piano che passa per i punti $(0, 1, 2)$, $(3, -1, 1)$, $(1, 1, 3)$. □

ESERCIZIO 12.1.3. Scrivere l'equazione del piano bisettore del segmento da $(1, 0, 0)$ a $(0, 0, 1)$. □

ESERCIZIO 12.1.4. In quali punti il piano $2x + y + 3z = 4$ taglia gli assi coordinati?

SUGGERIMENTO. Per trovare l'intersezione con l'asse x basta porre $y = z = 0$ nell'equazione del piano, ed analogamente per le intersezioni con gli altri due assi. $\square$

12.2. L'equazione della retta in $\mathbb{R}^3$

Anche qui, come prima, abbiamo due modi di scrivere l'equazione di una retta:

- *Equazione parametrica:* questo modo è identico a quello visto in due dimensioni. La retta che passa per il punto $\mathbf{p}$ ed ha la direzione del vettore $\mathbf{v}$ ha equazione

$$\mathbf{r}(t) = \mathbf{p} + t\mathbf{v}$$

con $t \in \mathbb{R}$.

Osserviamo che, se la retta passa per due punti $\mathbf{p}$ e $\mathbf{q}$, allora la sua equazione parametrica è

$$\mathbf{r}(t) = \mathbf{p} + t(\mathbf{q} - \mathbf{p})$$

con $t \in \mathbb{R}$.

- *Equazione cartesiana:* Si noti che l'analogo tridimensionale dell'equazione cartesiana (11.1.3) della retta in $\mathbb{R}^2$ dovrebbe essere

$$ax_1 + bx_2 + cx_3 = d$$

dove a, b, c e d sono costanti reali, ma questa non è l'equazione di una retta in $\mathbb{R}^3$, bensì di un piano, come visto sopra in (12.1.1). Per determinare l'equazione cartesiana di una retta basta osservare che, in tre dimensioni, ogni retta si può interpretare come l'intersezione di due piani (non paralleli, e non univocamente determinati), e quindi invece di un'equazione come (12.1.1) abbiamo un sistema di due equazioni:

$$\begin{aligned} a_1x_1 + b_1x_2 + c_1x_3 &= d_1 \\ a_2x_1 + b_2x_2 + c_2x_3 &= d_2. \end{aligned}$$

Quindi l'equazione della retta è un sistema di due equazioni lineari, cioè una equazione vettoriale:

$$\begin{pmatrix} a_1 & b_1 & c_1 \\ a_2 & b_2 & c_2 \end{pmatrix} \begin{pmatrix} x_1 \\ x_2 \\ x_3 \end{pmatrix} = \begin{pmatrix} d_1 \\ d_2 \end{pmatrix} . \quad (12.2.1)$$

I vettori (a_1, b_1, c_1) e (a_2, b_2, c_2) sono perpendicolari ai due piani rispettivi, e poiché tali piani non sono paralleli i due vettori determinano un unico (a meno di multipli) terzo vettore perpendicolare ad entrambi, che quindi è il vettore di direzione della retta.

ESEMPIO 12.2.1. (*La retta intersezione di due piani.*) Ogni coppia di piani non paralleli definisce una ed una sola retta di intersezione. Se le equazioni dei due piani sono $a_1x + b_1y + c_1z = d_1$ e $a_2x + b_2y + c_2z = d_2$ allora essi sono paralleli esattamente quando i loro vettori normali (a_1, b_1, c_1) e (a_2, b_2, c_2) sono multipli uno dell'altro. In tal caso il sistema delle due equazioni dei piani è l'equazione cartesiana della retta.

Per scrivere l'equazione parametrica della retta occorre trovare il suo vettore direzionale. Come osservato più sopra, questo vettore $\mathbf{v}$ è, a meno di multipli, quello che completa una terna ortogonale i cui primi due vettori sono i vettori perpendicolari ai due piani, quindi si ottiene dal procedimento di ortogonalizzazione di Gram-Schmidt (Sezione 6.7), o equivalentemente ma più velocemente calcolando il prodotto vettoriale dei due vettori (Sezione ??). A questo punto per trovare l'equazione parametrica della retta basta trovare un qualsiasi suo punto $\mathbf{p}$: l'equazione della retta è $\mathbf{r}(t) = \mathbf{p} + t\mathbf{v}$. Per trovare $\mathbf{p}$ occorre trovare almeno una soluzione del sistema (12.2.1), e quindi, ad esempio, tramite il metodo di eliminazione di Gauss (Sezione 3.7).

È interessante il caso in cui i due piani sono assegnati tramite due terne di punti che vi appartengono. Abbiamo visto come determinare le equazioni dei due piani, e quindi un vettore direzionale $\mathbf{v}$ della retta. Il caso più comodo in questo contesto è quello in cui le due terne di punti ne contengono uno in comune: allora questo punto $\mathbf{p}$ appartiene

ad entrambi i piani e l'equazione parametrica della retta è $\mathbf{r}(t) = \mathbf{p} + t\mathbf{v}$.

□

ESERCIZIO 12.2.2. (*Il fascio di piani che contengono una retta assegnata.*)

(i) Sia $\mathbf{r}$ una retta espressa in forma parametrica, $\mathbf{r}(t) = \mathbf{p} + t\mathbf{v}$.

Come si scrivono le equazioni dei piani π che la contengono?

(ii) Ora esprimiamo $\mathbf{r}$ in forma cartesiana, come intersezione di due piani, $\pi_0 : \{(x_1, x_2, x_3) : a_0x_1 + b_0x_2 + c_0x_3 = d_0\}$ e $\pi_1 : \{(x_1, x_2, x_3) : a_1x_1 + b_1x_2 + c_1x_3 = d_1\}$. Come si scrive ora il fascio di piani che contiene questa retta?

SVOLGIMENTO. La prima parte è facile: i piani che contengono la retta sono quelli che contengono un punto della retta, ad esempio $\mathbf{p}$, ed hanno vettore normale $\mathbf{n}$ perpendicolare al vettore direzionale della retta, ossia $\mathbf{v}$. Scrivendo $\mathbf{n} = (n_1, n_2, n_3)$, $\mathbf{p} = (p_1, p_2, p_3)$ e $\mathbf{x} = (x_1, x_2, x_3)$ troviamo l'equazione $\langle \mathbf{n}, \mathbf{x} \rangle = \langle \mathbf{n}, \mathbf{p} \rangle$, ovvero

$$n_1x_1 + n_2x_2 + n_3x_3 = n_1p_1 + n_2p_2 + n_3p_3. \quad (12.2.2)$$

Per la seconda parte, basta prendere i vettori normali ai due piani, diciamo $\mathbf{n}_0 = (a_0, b_0, c_0)$ e $\mathbf{n}_1 = (a_1, b_1, c_1)$, e considerare i piani con vettore normale dato da una combinazione lineare di $\mathbf{n}_0$ e $\mathbf{n}_1$: infatti i vettori normali dei piani appartenenti al fascio di piani che contengono $\mathbf{r}$ generano il piano perpendicolare a $\mathbf{r}$ e quindi sono combinazioni lineari di due qualsiasi indipendenti di essi. Allora, sia $\lambda = (\lambda_0, \lambda_1)$ e consideriamo il vettore normale $\mathbf{n}_\lambda = \lambda_0\mathbf{n}_0 + \lambda_1\mathbf{n}_1$. Poniamo $d_\lambda = \lambda_0d_0 + \lambda_1d_1$ e $\mathbf{x} = (x_1, x_2, x_3)$. Allora il piano che passa per $\mathbf{r}$ ed ha $\mathbf{n}_\lambda$ come vettore normale è quello con equazione $\langle \mathbf{n}_\lambda, \mathbf{x} \rangle = d_\lambda$. Infatti questo piano contiene la retta data, perché il sistema delle due equazioni lineari $\langle \mathbf{n}_\lambda, \mathbf{x} \rangle = d_\lambda$ e $\langle \mathbf{n}_1, \mathbf{x} \rangle = d_1$ è equivalente al sistema originale $\langle \mathbf{n}_0, \mathbf{x} \rangle = d_0$, $\langle \mathbf{n}_1, \mathbf{x} \rangle = d_1$, purché $\mathbf{n}_0$ e $\mathbf{n}_1$ siano linearmente indipendenti (come si vede subito sottraendo la seconda equazione dalla prima).

□

ESERCIZIO 12.2.3. (*Il piano che contiene una retta assegnata ed è perpendicolare ad un'altra retta data.*) Date due rette in forma parametrica, $\mathbf{r}(t) = \mathbf{p} + t\mathbf{v}$ e $\mathbf{s}(t) = \mathbf{q} + t\mathbf{u}$, mostrare che esiste un piano che contiene $\mathbf{r}$ ed è perpendicolare a $\mathbf{s}$ se e solo se i vettori direzionali $\mathbf{v}$ e $\mathbf{u}$ sono ortogonali.

SVOLGIMENTO. Senza perdita di generalità possiamo supporre, a meno di una traslazione, che la retta $\mathbf{r}$ passi per l'origine, ossia $\mathbf{p} = \mathbf{o}$: infatti la traslazione non cambia il vettore direzionale della retta. Esattamente per lo stesso motivo possiamo supporre che anche $\mathbf{s}$ passi per l'origine: basta rimpiazzare $\mathbf{s}$ con la sua parallela che passa per l'origine ed osservare che se un piano è perpendicolare a questa retta parallela è anche perpendicolare alla retta originale. Ora le due rette si incontrano nell'origine. Se i loro vettori direzionali $\mathbf{v}$ e $\mathbf{u}$ sono ortogonali, allora il piano per l'origine con vettore normale $\mathbf{u}$ contiene il vettore $\mathbf{v}$ e quindi la retta $\mathbf{r}(t) = t\mathbf{v}$. Viceversa, ogni piano che contiene $\mathbf{v}$ e quindi la retta $\mathbf{r}(t) = t\mathbf{v}$ deve avere vettore normale ortogonale a $\mathbf{v}$. Un tale piano è perpendicolare a $\mathbf{s}(t) = t\mathbf{u}$ se e solo se il suo vettore normale è multiplo di $\mathbf{u}$, e pertanto $\mathbf{v}$ deve essere ortogonale a $\mathbf{u}$. $\square$

ESERCIZIO 12.2.4. I vettori $\mathbf{v} = (1, 0, 0)$ e $\mathbf{u} = (0, 1, 1)$ sono ortogonali. Trovare l'equazione del piano π che contiene la retta $\mathbf{r}(t) = (0, 1, 0) + t\mathbf{v}$ ed è perpendicolare alla retta $\mathbf{s}(t) = (0, 0, 1) + t\mathbf{u}$.

SVOLGIMENTO. In base al precedente Esercizio 12.2.3, il vettore normale a π è $\mathbf{u}$. Inoltre π deve contenere $\mathbf{r}$, ovvero deve contenere il punto $\mathbf{p}$ (poiché π ha per vettore normale un vettore ortogonale a $\mathbf{v}$, e quindi contiene $\mathbf{v}$, quest'ultima asserzione equivale a dire che π contiene tutta la retta $\mathbf{r}$, in base alla parte (i) dell'Esercizio 12.2.2). Pertanto conosciamo un punto contenuto in π ed il vettore normale di π : ora l'equazione cartesiana di π si trova come in (12.2.2). $\square$

12.3. Proiezioni e distanze fra piani

ESEMPIO 12.3.1. (*Proiezione ortogonale di un punto su un piano.*) Il problema della proiezione ortogonale di un punto $\mathbf{x} \in \mathbb{R}^3$ su

un piano π si risolve in maniera analoga al caso bidimensionale trattato nell'Esempio 11.1.2. Ci sono quindi due modi:

- *Tramite l'equazione parametrica della retta perpendicolare.* Sia $ax_1 + bx_2 + cx_3 = d$ l'equazione del piano π . Allora un vettore normale al piano è $\mathbf{n} = (a, b, c)$, e l'equazione parametrica della retta che passa per $\mathbf{x}$ ed è perpendicolare a π è $\mathbf{r}(t) = \mathbf{x} + t\mathbf{n}$. Imponendo che $\mathbf{r}(t)$ soddisfi l'equazione $ax_1 + bx_2 + cx_3 = d$ si trova il valore t_0 del parametro t per cui $\mathbf{r}(t)$ giace in π , quindi il piede $\mathbf{p} := \mathbf{r}(t_0)$ della proiezione ortogonale di $\mathbf{x}$ sul piano π .
- *Tramite la decomposizione ortogonale di Gram-Schmidt.* Come prima, si considera il vettore normale $\mathbf{n} = (a, b, c)$, ed ora lo si normalizza (per semplicità supponiamo che sia già di norma 1). Supponiamo dapprima che il piano passi per l'origine, cioè che si abbia $d = 0$. Allora questo piano è un sottospazio bidimensionale: in esso si costruisce una base ortonormale $\mathbf{v}_1, \mathbf{v}_2$ tramite il procedimento di Gram-Schmidt. La terna $\mathbf{v}_1, \mathbf{v}_2, \mathbf{n}$ è una base ortonormale in $\mathbb{R}^3$. Ora la proiezione ortogonale $\mathbf{p}$ di $\mathbf{x}$ su π si ottiene da (6.6.1), sottraendo le componenti lungo π :

$$\mathbf{p} = \mathbf{x} - \mathbf{x} \cdot \mathbf{v}_1 \mathbf{v}_1 - \mathbf{x} \cdot \mathbf{v}_2 \mathbf{v}_2.$$

Si osservi che non c'è bisogno di prendere una base ortonormale in π : qualunque base va bene, però le componenti di $\mathbf{x}$ lungo questi vettori di base si scrivono in maniera più complicata, come nell'uguaglianza (??).

Ora veniamo al caso generale in cui il piano π non passa per l'origine, cioè $d \neq 0$: π passa a distanza $|d|$ dall'origine. In questo caso il procedimento precedente ci dà la proiezione $\mathbf{p}'$ di $\mathbf{x}$ non su π ma sul piano π' ad esso parallelo che passa per l'origine. Ora la proiezione $\mathbf{p}$ su π si ottiene spostandosi perpendicolarmente a π' di una distanza $|d|$:

$$\mathbf{p} = \mathbf{p}' + d\mathbf{n}$$

(nel secondo membro abbiamo preso la somma invece della differenza perché, se ad esempio $a, b, c > 0$ allora il piano π è spostato verso l'ottante positivo rispetto a π' se $d > 0$ e verso quello negativo (cioè con tutte e tre le coordinate negative) altrimenti. Quindi per spostarsi da π a π' bisogna muoversi nella direzione di $\mathbf{n}$ se $d > 0$ e al contrario se $d < 0$.

In particolare, se un piano ha equazione $\mathbf{n} \cdot \mathbf{x} = d$ (e quindi equazione normale $\mathbf{n} \cdot \mathbf{x} / \|\mathbf{n}\| = d / \|\mathbf{n}\|$), la distanza di un punto $\mathbf{p}$ dal piano è

$$\left| \frac{\mathbf{n} \cdot \mathbf{p} - d}{\|\mathbf{n}\|} \right|. \quad (12.3.1)$$

Infatti, per quanto osservato sopra, questa è la differenza fra le distanze dall'origine del piano dato e del suo parallelo che passa per $\mathbf{p}$.

In particolare, l'equazione del piano con vettore normale $\mathbf{n}$ che contiene un dato punto $\mathbf{p}$ è

$$\mathbf{n} \cdot \mathbf{x} = \mathbf{n} \cdot \mathbf{p}. \quad (12.3.2)$$

□

ESERCIZIO 12.3.2. Trovare la distanza fra il punto $(1, 1, 1)$ ed il piano $x + 3y - z = 2$.

□

ESERCIZIO 12.3.3. Trovare la distanza fra i due piani paralleli seguenti:

$$\begin{aligned} x + 3y - 2z &= 1 \\ x + 3y - 2z &= 4. \end{aligned}$$

□

ESEMPIO 12.3.4. (*Proiezione ortogonale di un punto su una retta in $\mathbb{R}^3$.*) Anche questo problema si può risolvere in due modi equivalenti.

- Tramite l'equazione parametrica della retta perpendicolare.
- Tramite la decomposizione ortogonale di Gram-Schmidt.

□

ESERCIZIO 12.3.5. Trovare la distanza del punto $(1, 2, 3)$ dalla retta intersezione dei due piani $2x + y - z = 0$ e $x - 2y + z = 1$. $\square$

ESEMPIO 12.3.6. (*Equazioni dei piani bisettori di due piani dati.*) Grazie a (12.3.1), dati due piani $\mathbf{n}_1 \cdot \mathbf{x} = d_1$ e $\mathbf{n}_2 \cdot \mathbf{x} = d_2$, i loro due piani bisettori hanno equazione

$$\left| \frac{\mathbf{n}_1 \cdot \mathbf{x} - d_1}{\|\mathbf{n}_1\|} \right| = \left| \frac{\mathbf{n}_2 \cdot \mathbf{x} - d_2}{\|\mathbf{n}_2\|} \right|$$

e cioè

$$\frac{\mathbf{n}_1 \cdot \mathbf{x} - d_1}{\|\mathbf{n}_1\|} = \frac{\mathbf{n}_2 \cdot \mathbf{x} - d_2}{\|\mathbf{n}_2\|}$$

e

$$\frac{\mathbf{n}_1 \cdot \mathbf{x} - d_1}{\|\mathbf{n}_1\|} = -\frac{\mathbf{n}_2 \cdot \mathbf{x} - d_2}{\|\mathbf{n}_2\|}.$$

$\square$

12.4. Sfere

Il modo di descrivere sfere tridimensionali tramite equazioni quadratiche è identico a quello già presentato per i cerchi nel piano (Sezione 11.2). Lo ripetiamo in questo ambiente più generale.

La sfera di centro $\mathbf{c}$ e raggio $r > 0$ è l'insieme dei punti $\mathbf{x}$ a distanza da $\mathbf{c}$ uguale a r , quindi l'equazione della sfera è

$$\|\mathbf{x} - \mathbf{c}\| = r$$

cioè

$$\sum_{i=1}^3 (x_i - c_i)^2 = r^2. \quad (12.4.1)$$

Sviluppando i quadrati nell'equazione della sfera (12.4.1) otteniamo:

$$\sum_{i=1}^3 x_i^2 + 2 \sum_{i=1}^3 c_i x_i + \sum_{i=1}^3 c_i^2 - r^2 = 0.$$

Viceversa, consideriamo più generale l'equazione quadratica

$$\sum_{i=1}^3 x_i^2 + \sum_{i=1}^3 a_i x_i + b = 0.$$

Allora, ponendo $c_i = \frac{a_i}{2}$ e $r^2 = \sum_{i=1}^3 c_i^2 - b$, questa equazione quadratica si riconduce a (12.4.1) se $b \leq \sum_{i=1}^3 c_i^2$ (nel caso si abbia l'uguaglianza la sfera consiste di un solo punto), ed invece non ha soluzioni reali se $b > \sum_{i=1}^3 c_i^2$.

Se invece di equazioni del tipo $\sum_{i=1}^3 (x_i - c_i)^2 = r^2$ scegliamo tre coefficienti positivi b_i e consideriamo l'equazione

$$\sum_{i=1}^3 b_i x_i^2 + 2 \sum_{i=1}^3 c_i x_i + \sum_{i=1}^3 c_i^2 - r^2 = 0,$$

possiamo riportare questa equazione alla precedente tramite un cambiamento di scala su ciascun asse. Quindi le soluzioni di questa equazione sono ellissi con gli assi principali paralleli agli assi coordinati. Tramite una rotazione ci riportiamo al caso generale di ellissi con gli assi principali disposti in altre orientazioni (ma ovviamente consistenti di una terna ortogonale, anche se non si incrociano nell'origine).

12.5. Esercizi di geometria analitica in $\mathbb{R}^3$

ESERCIZIO 12.5.1. Siano $\mathbf{p} = (0, 0, 1)$ e $\mathbf{q} = (2, 2, 3)$. Per quali valori delle costanti reali a e d la retta che passa per $\mathbf{p}$ e $\mathbf{q}$ è parallela al piano $ax + y + z + d = 0$?

SVOLGIMENTO. Basta imporre che il vettore normale al piano sia ortogonale al vettore di direzione della retta, che è $(2, 2, 2)$. Quindi l'equazione del piano è un multiplo di $2x + 2y + 2z + d = 0$ (la costante d non viene determinata, perché la condizione sul piano riguarda solo il suo vettore normale, non la distanza dall'origine: cambiando d il piano si trasforma in un altro piano parallelo al primo.) Pertanto la risposta è $a = 1$, d arbitrario.

In alternativa, ma con molti più calcoli non necessari, si può scrivere l'equazione della retta e metterla a sistema con quella del piano, poi imporre che non ci siano soluzioni. $\square$

ESERCIZIO 12.5.2. Scrivere l'equazione del piano parallelo all'asse z che interseca l'asse x nel punto $(2, 0, 0)$ e l'asse y in $(0, 1, 0)$.

SVOLGIMENTO. Il piano è parallelo all'asse z se e solo se il suo vettore normale è perpendicolare a $\mathbf{e}_3 = (0, 0, 1)$, e quindi del tipo $(a, b, 0)$.

Pertanto l'equazione del piano è della forma $ax + by + d = 0$. L'intersezione con l'asse x deve quindi essere in un punto $(x, 0, 0)$ che soddisfa $ax + d = 0$: poiché si richiede che questa intersezione sia in $(2, 0, 0)$ dobbiamo avere $2a + d = 0$. Con lo stesso ragionamento per l'intersezione con l'asse y si ottiene $b + d = 0$. Da queste due equazioni ricaviamo $2a = b = -d$: quindi il piano ha equazione $ax + 2ay - 2a = 0$. Poiché il piano non cambia se dividiamo tutti i coefficienti dell'equazione per la stessa costante, possiamo scrivere per il piano l'equazione $x + 2y - 2 = 0$.

□

ESERCIZIO 12.5.3. (*Piani che contengono una retta assegnata.*)

Scrivere l'equazione dei piani che contengono la retta $x = 2t, y = t, z = t + 1$, dove t è una variabile reale.

SVOLGIMENTO. Il vettore direzionale della retta è $\mathbf{d} = (2, 1, 1)$. I piani che contengono la retta hanno vettore normale ortogonale a $\mathbf{d}$, quindi sono quelli del tipo $ax + by + cz + d = 0$ dove $(a, b, c) \perp (2, 1, 1)$ (cioè $2a + b + c = 0$) che contengono un punto della retta, ad esempio il punto $(0, 0, 1)$ (ottenuto ponendo $t = 0$ nell'equazione parametrica della retta). Dall'ultima condizione segue $c + d = 0$, cioè $d = -c$. Pertanto i piani richiesti soddisfano l'equazione $ax + by + cz - c = 0$ con la condizione $2a + b + c = 0$: una famiglia a due parametri di piani.

□

ESERCIZIO 12.5.4. (*Piani che contengono una retta data come intersezione di due piani fissati.*) Scrivere l'equazione dei piani che contengono la retta intersezione dei piani $x + y + z = 0$ e $x + 2y + 3z - 1 = 0$.

SVOLGIMENTO. Fissiamo un punto nell'intersezione dei due piani, cioè nella retta, diciamo $\mathbf{p} = (-1, 1, 0)$ (ottenuto scegliendo $z = 0$ nelle equazioni dei due piani, e mettendole a sistema).

Non occorre trovare l'equazione parametrica della retta: il suo vettore direzionale $\mathbf{d}$ è perpendicolare ai vettori normali dei due piani assegnati, cioè ai vettori $\mathbf{n}_1 = (1, 1, 1)$ e $\mathbf{n}_2 = (1, 2, 3)$. Ora, come nel precedente Esercizio 12.5.3, i piani che contengono la retta hanno vettore normale

$\mathbf{n}$ ortogonale a $\mathbf{d}$, e quindi appartenente al piano generato da $\mathbf{n}_1$ e $\mathbf{n}_2$ (perché in $\mathbb{R}^3$ questo piano è il complemento ortogonale di $\mathbf{d}$: *ma come si risolverebbe lo stesso esercizio in $\mathbb{R}^n$?*). Pertanto $\mathbf{n} = s\mathbf{n}_1 + t\mathbf{n}_2 = (s+t, s+2t, s+3t)$ per qualche $s, t \in \mathbb{R}$. Quindi i piani cercati sono la famiglia a due parametri di equazione

$$\mathbf{n} \cdot (x, y, z) = (s+t)x + (s+2t)y + (s+3t)z = \mathbf{n} \cdot \mathbf{p} = t.$$

Ponendo $s+t = u$ riscriviamo la famiglia a due parametri di equazioni nel seguente modo leggermente più semplice:

$$ux + (u+t)y + (u+2t)z - t = 0.$$

□

ESERCIZIO 12.5.5. (*Piano che contiene un punto assegnato ed una retta data come intersezione di due piani fissati.*) Scrivere l'equazione del piano che contiene il punto $\mathbf{p} = (1, 1, -2)$ e la retta intersezione dei piani σ di equazione $x + y + z = 0$ e τ di equazione $x + 2y + 3z - 1 = 0$. Si risolva lo stesso problema nel caso in cui, invece di $\mathbf{p}$, si prende il punto $\mathbf{q} = (-1, 1, 0)$.

SVOLGIMENTO. Dal precedente Esercizio 12.5.4 sappiamo che i piani che passano per la retta data verificano l'equazione

$$ux + (u+t)y + (u+2t)z - t = 0$$

dove u e t sono due parametri reali. Fra questi piani, quindi, quello che passa anche per il punto $(1, 1, -2)$ soddisfa $u + (u+t) - 2(u+2t) - t = 0$, cioè $-5t = 0$, ossia $t = 0$. L'equazione è $ux + uy + uz = 0$: poiché il piano rimane invariato se riscaldiamo i coefficienti dell'equazione, essa si può riscrivere come $z + y + z = 0$.

Si noti che il punto $\mathbf{p}$ appartiene già al piano σ , e quindi il piano cercato è proprio σ . Invece, come già osservato nel precedente Esercizio 12.5.4, il punto $\mathbf{q} = (-1, 1, 0)$ appartiene ad entrambi i piani σ e τ , e quindi alla retta data dalla loro intersezione. Perciò imporre che fra i vari piani che contengono tale retta se ne scelga uno che passa per $\mathbf{q}$ non comporta nessuna restrizione: tutti lo fanno. Ed infatti, se si ripete il ragionamento precedente, imponendo ora che u e t siano tali

che l'equazione $ux + (u + t)y + (u + 2t)z - t = 0$ sia soddisfatta dal punto $(x, y, z) = (-1, 1, 0)$, otteniamo l'identità $0 = 0$ invece che una equazione in u e t . $\square$

ESERCIZIO 12.5.6. Trovare l'equazione del piano che passa per la retta di equazione parametrica $\mathbf{r}(t) = (1 + t, 2t, t - 1)$ ed è parallelo alla retta $\mathbf{s}(t) = (2 + t, t, 3t + 2)$.

SVOLGIMENTO. Il vettore direzionale della retta $\mathbf{r}$ è $\mathbf{d}_1 = (1, 2, 1)$; quello di $\mathbf{s}$ è $\mathbf{d}_2 = (1, 1, 3)$. Poiché il piano cercato è parallelo ad entrambe le rette, il suo vettore normale $\mathbf{n}$ è ortogonale sia a $\mathbf{d}_1$ sia a $\mathbf{d}_2$: esso si può ricavare dal prodotto vettoriale di $\mathbf{d}_1$ e $\mathbf{d}_2$ (Sezione ??), oppure direttamente imponendo che $\mathbf{n} = (a, b, c)$ verifichi $\mathbf{n} \cdot \mathbf{d}_1 = a - c = 0$ e $\mathbf{n} \cdot \mathbf{d}_2 = a + b + 3c = 0$. Il sistema lineare di queste due equazioni nelle incognite a, b e c ha per soluzioni $(a, -4a, a)$: riscalando otteniamo un vettore normale $\mathbf{n} = (1, -4, 1)$.

Scegliamo un punto $\mathbf{p}$ nella retta $\mathbf{r}$, ad esempio $\mathbf{p} = (1, 0, -1)$, e scriviamo il punto generico (x, y, z) come $\mathbf{x}$. Allora l'equazione (12.3.2) del piano è $\mathbf{n} \cdot \mathbf{x} = \mathbf{n} \cdot \mathbf{p}$, ossia $x - 4y + z = 0$. $\square$

ESERCIZIO 12.5.7. Trovare l'equazione del piano che passa per la retta di equazione parametrica $\mathbf{r}(t) = (1 + t, 2t, t - 1)$ ed è perpendicolare al piano $x + y + 2z - 1 = 0$

SVOLGIMENTO. Come nel precedente Esercizio 12.5.6, imponiamo che il vettore normale $\mathbf{n} := (a, b, c)$ del piano cercato sia ortogonale al vettore direzionale $\mathbf{d} = (1, 2, 1)$ della retta, e che passi per un suo punto, diciamo $\mathbf{p} = (1, 0, -1)$. Inoltre, ora richiediamo che $\mathbf{n}$ sia anche ortogonale al vettore normale del piano $x + y + 2z - 1 = 0$, cioè a $(1, 1, 2)$. Le condizioni imposte determinano $\mathbf{n}$ come la soluzione (a meno di multipli) del sistema $a + 2b + c = 0$, $a + b + 2c = 0$, cioè $a = c$, $b = -c$. Quindi un vettore normale al piano cercato è $\mathbf{n} = (1, -1, 1)$, e l'equazione (12.3.2) di questo piano diventa $x - y + z = 0$. $\square$

ESERCIZIO 12.5.8. Trovare l'equazione della retta $\mathbf{r}$ che interseca le rette $\mathbf{r}_1(t) = (1, t, 2t)$, $\mathbf{r}_2(t) = (1, 2t, t)$ e passa per il punto $\mathbf{p} = (1, 1, 1)$.

SUGGERIMENTO. Il modo più rapido è di trovare l'equazione del piano che passa per la retta $\mathbf{r}_1$ ed il punto $\mathbf{p}$, e da essa ricavare il punto $\mathbf{q}$ di intersezione di questo piano con la retta $\mathbf{r}_2$. Ora la retta $\mathbf{r}$ desiderata è quella che passa per i punti $\mathbf{p}$ e $\mathbf{q}$.

Assai più pedissequamente si può scrivere l'equazione della retta che passa per il generico punto $(1, t, 2t)$ di $\mathbf{r}_1$ ed il generico punto $(1, 2s, s)$ di $\mathbf{r}_2(t)$ (basta usare il vettore direzionale $\mathbf{d}_t s = (1, t, 2t) - (1, 2s, s)$), ed imporre che passi anche per $\mathbf{p} = (1, 1, 1)$. $\square$

ESERCIZIO 12.5.9. Trovare l'equazione della retta $\mathbf{r}$ che interseca le rette $\mathbf{r}_1(t) = (1, t, 2t)$, $\mathbf{r}_2(t) = (1, 2t, t)$ ed è parallela alla retta $\mathbf{r}_3(t) = (t, -t, 0)$.

SUGGERIMENTO. Il fatto che $\mathbf{r}$ e $\mathbf{r}_3$ siano parallele equivale a dire che hanno lo stesso vettore direzionale, cioè $\mathbf{d} = (1, -1, 0)$ (o equivalentemente un suo multiplo). Scegliamo un punto $\mathbf{p} = \mathbf{r}_1(t_0)$ sulla retta $\mathbf{r}_1$, e consideriamo la retta $\mathbf{v}$ di equazione parametrica $\mathbf{v}(s) = \mathbf{p} + s\mathbf{d}$. Per tutti i valori di t_0 per cui questa retta $\mathbf{v}$ interseca $\mathbf{r}_2$, essa dà una soluzione del problema proposto. Per trovare quando questo accade si procede in maniera analoga al precedente Esercizio 12.5.8: si determina l'equazione del piano che contiene $\mathbf{r}_1$ e $\mathbf{v}$ e si vede se questo piano ha intersezione con $\mathbf{r}_2$. $\square$

ESERCIZIO 12.5.10. Trovare l'equazione della retta uscente dall'origine che interseca la retta $\mathbf{r}(t) = (1, t, 2t)$ ed è parallela al piano $2x + z = 0$.

SUGGERIMENTO. La condizione di parallelismo al piano equivale a richiedere che il vettore direzionale $\mathbf{d}$ della retta cercata sia ortogonale al vettore normale del piano $\mathbf{n} = (2, 0, 1)$. Quindi abbiamo un sottospazio bidimensionale di tali vettori (l'ortogonale di $\mathbf{n}$), che corrispondono ad un piano di rette uscenti dall'origine: la retta desiderata è quella fra queste che interseca $\mathbf{r}$, cioè quella il cui vettore $\mathbf{d}$ è diretto verso il punto di intersezione di questo piano con $\mathbf{r}_2$ (cioè è un multiplo di tale punto). $\square$

ESERCIZIO 12.5.11. Trovare l'equazione della retta passante per $(1, 3, 2)$ e perpendicolare al piano $2x - y + z - 2 = 0$.

SVOLGIMENTO. Il vettore direzionale della retta deve essere ortogonale al piano, quindi multiplo del suo vettore normale $(2, -1, 1)$. Pertanto la retta ha equazione parametrica $\mathbf{r}(t) = (1, 3, 2) + t(2, -1, 1) = (1 + 2t, 3 - t, 2 + t)$. $\square$

ESERCIZIO 12.5.12. Trovare l'equazione della retta passante per $(1, 3, 2)$ e perpendicolare al piano $2x - y + z - 2 = 0$.

SUGGERIMENTO. Ci si riporta al precedente Esercizio 12.5.11 osservando che la retta cercata è quella passante per $(1, 3, 2)$ e parallela alla retta di intersezione dei due piani (Esempio 12.2.1). $\square$

ESERCIZIO 12.5.13. Trovare l'equazione della retta uscente dall'origine che interseca perpendicolarmente la retta $\mathbf{r}(t) = (1 + t, 2t, -t)$.

SUGGERIMENTO. La retta cercata ha vettore direzionale perpendicolare a quello di $\mathbf{r}$, cioè a $(1, 2, -1)$. Questi vettori direzionali giacciono quindi nel piano ortogonale a $(1, 2, -1)$. Fra questi, quello desiderato è quello che punta verso l'intersezione della retta $\mathbf{r}$ con questo piano, esattamente come nell'Esercizio 12.5.11. $\square$

ESERCIZIO 12.5.14. Trovare la distanza minima fra le rette $\mathbf{r}(t) = (0, 0, t)$ e $\mathbf{q}(u) = (u, u + 1, u + 2)$.

SVOLGIMENTO. Il segmento che connette un punto generico della prima retta ad un punto della seconda ha lunghezza minima se e solo se è perpendicolare ad entrambe le rette. Allora prendiamo in esame la prima retta e consideriamo il piano ad essa ortogonale che passa per il suo punto generico $\mathbf{r}(t)$. Questo piano è ortogonale al suo vettore direzionale $(0, 0, 1)$, e quindi la sua equazione (12.3.2) è $x = (0, 0, 1) \cdot \mathbf{r}(t)$, cioè $z = t$. Intersechiamo tale piano con la retta $\mathbf{q}$ e troviamo il punto di intersezione $\mathbf{q}(u) = (u, u + 1, u + 2)$ con $u + 2 = t$, cioè il punto $\mathbf{q}(t - 2) = (t - 2, t - 1, t)$. Il segmento che unisce $\mathbf{r}(t)$ con $\mathbf{q}(t - 2)$ è diretto come $\mathbf{s}(t) := \mathbf{q}(t - 2) - \mathbf{r}(t) = (t - 2, t - 1, 0)$. Ora basta imporre che questo vettore sia ortogonale a $\mathbf{q}$, cioè al suo vettore direzionale $\mathbf{d} = (1, 1, 1)$. Ma $\mathbf{s}(t) \cdot \mathbf{d} = 2t - 3$, da cui si ricava $t = \frac{3}{2}$. La distanza minima richiesta quindi è $\|\mathbf{s}(\frac{3}{2})\| = \frac{1}{\sqrt{2}}$. $\square$

ESERCIZIO 12.5.15. Trovare l'equazione della sfera con centro in $\mathbf{c} = (1, 2, 3)$ e tangente al piano $x + y + 2z - 1 = 0$.

SUGGERIMENTO. Il raggio dal centro della sfera al punto di tangenza deve essere ortogonale al piano, quindi diretto come il suo vettore normale $\mathbf{n} = (1, 1, 2)$. Perciò questo raggio giace lungo la retta $\mathbf{r}(t) = \mathbf{c} + t\mathbf{n}$. Intersecando la retta $\mathbf{r}$ con il piano tangente troviamo ora il punto di tangenza $\mathbf{t}$, e la norma $\|\mathbf{t} - \mathbf{c}\|$, ossia la sua distanza dal centro $\mathbf{c}$, è il raggio della sfera. $\square$

ESERCIZIO 12.5.16. Trovare l'equazione della sfera tangente nell'origine al piano σ di equazione $x + y + z = 0$ e tangente anche al piano τ di equazione $x + y + 2z - 1 = 0$.

SUGGERIMENTO. Il centro della sfera deve giacere sulla retta $\mathbf{r}$ uscente dall'origine e perpendicolare a σ , cioè diretta come il suo vettore normale $\mathbf{p} = (1, 1, 1)$. Normalizziamo $\mathbf{p}$ ed otteniamo il versore normale $\mathbf{n} = \left(\frac{1}{\sqrt{3}}, \frac{1}{\sqrt{3}}, \frac{1}{\sqrt{3}}\right)$. Quindi abbiamo l'equazione parametrica $\mathbf{r}(t) = t\mathbf{n}$. Poiché $\mathbf{n}$ ha norma 1, la distanza di $\mathbf{r}(t)$ dall'origine è $|t|$. calcoliamo allora, come nell'Esempio 12.3.1, la proiezione ortogonale $\mathbf{p}(t)$ del punto $\mathbf{r}(t)$ sul piano τ . La sfera è tangente a τ se e solo se la distanza dal suo centro al punto di tangenza su τ è la stessa di quella dal centro al punto di tangenza su σ (ossia l'origine). Pertanto, per trovare il centro della sfera, basta imporre che il centro sia il punto $\mathbf{r}(t)$ tale che $\|\mathbf{r}(t) - \mathbf{p}(t)\| = \|\mathbf{r}(t)\| = |t|$. Dopo di che, il raggio della sfera è $\|\mathbf{r}(t)\| = |t|$. $\square$

12.6. Raggi riflessi e rifratti

Varie procedure in Computer Graphics richiedono di tracciare raggi (cioè rette parametriche) nello spazio e farli riflettere in maniera speculare su superficie. Il problema, espresso in termini del vettore direzionale $\mathbf{p}$ del raggio, si riduce al seguente. Dato un versore $\mathbf{p}$ ed un piano π che passa per l'origine con versore normale $\mathbf{n}$, come si determina il versore $\mathbf{v}$ riflesso speculare di $\mathbf{p}$ rispetto a $\mathbf{n}$? Per ragioni di semplicità di programmazione e velocità di esecuzione del codice, vorremmo

trovare la risposta solo in termini di prodotti scalari (somme di prodotti di coordinate, rapidi da calcolare) invece che tramite espressioni trigonometriche.

La proiezione del vettore $\mathbf{p}$ sul vettore normale $\mathbf{n}$ vale $\cos\theta\mathbf{n}$, dove θ è l'angolo di incidenza del raggio sul piano π . Perciò la componente di $\mathbf{p}$ ortogonale a $\mathbf{n}$ vale $\mathbf{s} = \mathbf{p} - \cos\theta\mathbf{n}$. Si osservi che questo vettore $\mathbf{s}$ è quello che verifica l'identità $\mathbf{p} + \mathbf{s} = \cos\theta\mathbf{n} = (\mathbf{p} \cdot \mathbf{n})\mathbf{n}/|\mathbf{n}|$, cioè porta da $\mathbf{p}$ al piede della sua proiezione lungo $\mathbf{n}$.

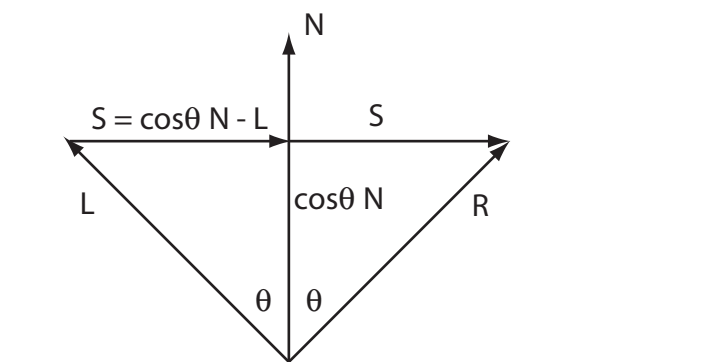


FIGURA 1. Vettore riflesso

È quindi chiaro che si ha

$$\mathbf{v} = \mathbf{p} + 2\mathbf{s} = 2\cos\theta\mathbf{n} - \mathbf{p} = 2(\mathbf{p} \cdot \mathbf{n})\mathbf{n} - \mathbf{p}.$$

Talvolta i raggi di luce, anziché riflettersi sulla superficie di un materiale, vi si rifrangono dentro. Determiniamo l'espressione del vettore rifratto dalla legge di Snell della rifrazione, che asserisce che il raggio incidente viene deviato in maniera da soddisfare la seguente relazione fra l'angolo θ_i che esso forma con la normale alla superficie (nel nostro caso un piano π e l'analogo angolo θ_t formato dal raggio rifratto, cioè trasmesso al secondo materiale:

$$\eta_i \sin \theta_i = \eta_t \sin \theta_t,$$

dove η_i e η_t sono i rispettivi indici di rifrazione dei due materiali, definiti come il rapporto tra la velocità della luce nel vuoto e nel materiale: essi variano con la lunghezza d'onda della luce. Il vuoto ha quindi indice di rifrazione 1; i materiali hanno indice di rifrazione più elevato.

Calcoliamo il versore direzionale $\mathbf{t}$ del raggio rifratto. Per prima cosa, determiniamo il versore $\mathbf{m}$ perpendicolare al versore normale $\mathbf{n}$ che giace nel piano generato da $\mathbf{n}$ e dal versore della direzione di incidenza $\mathbf{i}$, e punta verso il lato opposto di $\mathbf{i}$ rispetto a $\mathbf{n}$. La proiezione di $\mathbf{i}$ lungo $\mathbf{n}$ vale $\cos \theta_i \mathbf{n}$. Perciò il vettore $\cos \theta_i \mathbf{n} - \mathbf{i}$ è diretto parallelamente al piano π , e quindi è perpendicolare a $\mathbf{n}$. Poiché $\mathbf{n}$ e $\mathbf{i}$ hanno norma 1 (sono versori!) e $\cos \theta_i \mathbf{n}$ è ortogonale a $\mathbf{n}$ e quindi anche a $\cos \theta_i \mathbf{n}$, per il teorema di Pitagora la norma di $\cos \theta_i \mathbf{n} - \mathbf{i}$ vale $\sqrt{1 - \cos^2 \theta_i} = \sin \theta_i$ (il seno è positivo perché l'angolo θ_i è compreso fra 0 e $\pi/2$). Quindi, normalizzando, troviamo un vettore $\mathbf{m} = (\cos \theta_i \mathbf{n} - \mathbf{i}) / \sin \theta_i$.

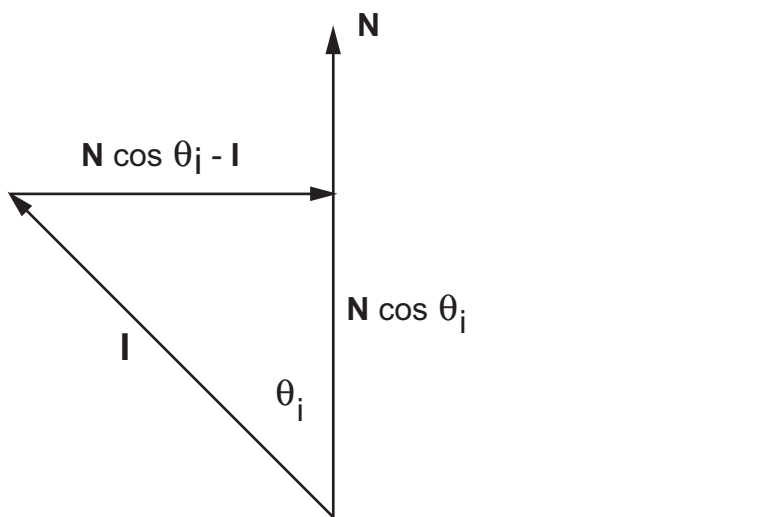


FIGURA 2. Offset trasversale del raggio incidente

Ora, per la legge di Snell, il versore del raggio trasmesso è

$$\mathbf{t} = \sin \theta_t \mathbf{m} - \cos \theta_t \mathbf{n} = \frac{\sin \theta_t}{\sin \theta_i} (\cos \theta_i \mathbf{n} - \mathbf{i}) - \cos \theta_t \mathbf{n}.$$

D'altra parte, sempre per la legge di Snell, si ha $\sin \theta_t / \sin \theta_i = \eta_i / \eta_t$. Quindi, ponendo $\eta_\tau = \eta_i / \eta_t$, otteniamo

$$\mathbf{t} = (\eta_\tau \cos \theta_i - \cos \theta_t) \mathbf{n} - \eta_\tau \mathbf{i}.$$

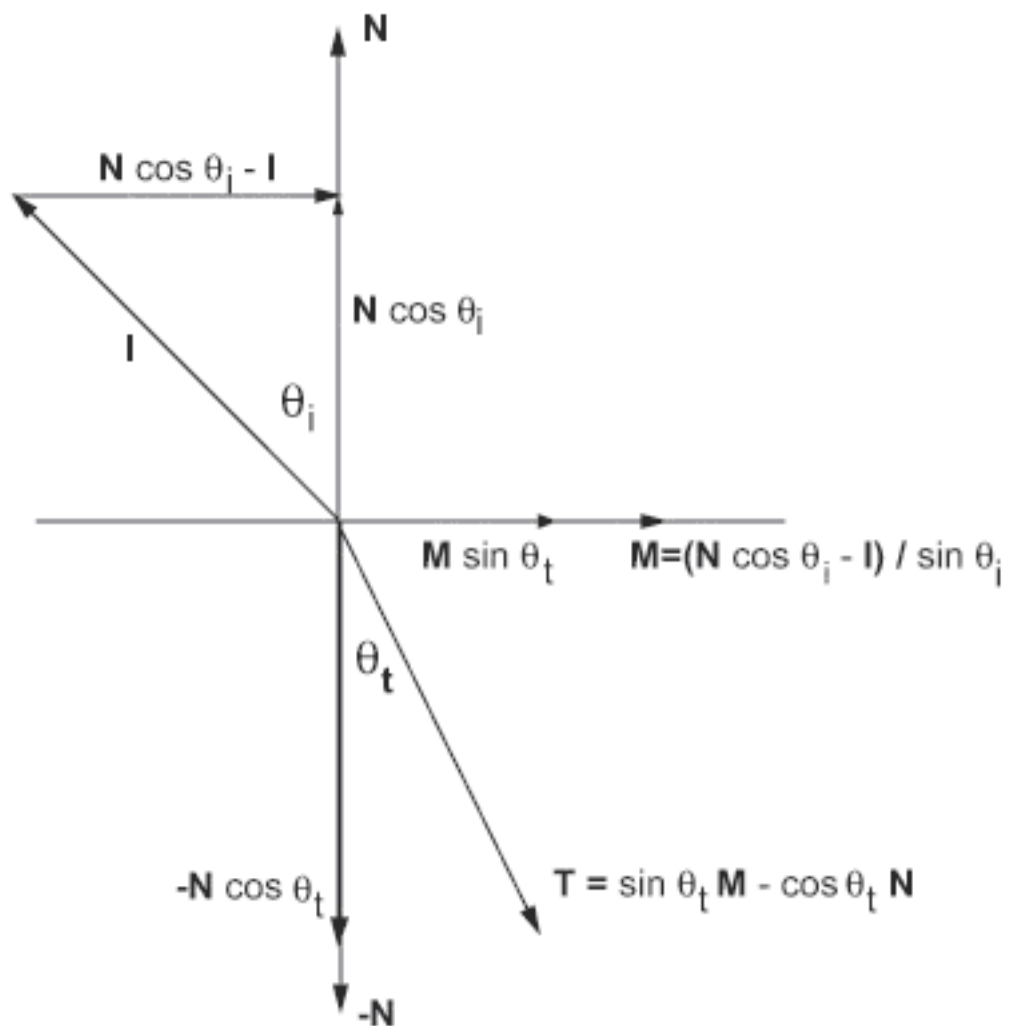


FIGURA 3. Il versore rifratto

D'altra parte $\cos \theta_i = \mathbf{n} \cdot \mathbf{i}$, e

$$\begin{aligned} \cos \theta_t &= \sqrt{1 - \sin^2 \theta_t} = \sqrt{1 - \eta_r^2 \sin^2 \theta_i} \\ &= \sqrt{1 - \eta_r^2 (1 - \cos^2 \theta_i)} = \sqrt{1 - \eta_r^2 (1 - (\mathbf{n} \cdot \mathbf{i})^2)}. \end{aligned}$$

Abbiamo così ricavato $\mathbf{t}$ unicamente in termini di prodotti scalari:

$$\mathbf{t} = \left(\eta_\tau \mathbf{n} \cdot \mathbf{i} - \sqrt{1 - \eta_\tau^2 (1 - (\mathbf{n} \cdot \mathbf{i})^2)} \right) \mathbf{n} - \eta_\tau \mathbf{i}.$$

CAPITOLO 13

* Spazi quozienti, spazi proiettivi

13.1. Spazio quoziente e dualità

DEFINIZIONE 13.1.1. (*Spazio quoziente.*) Sia $\mathbf{X}$ uno spazio vettoriale (a dimensione finita) e $\mathbf{V}$ un suo sottospazio. Il sottospazio $\mathbf{V}$ induce una relazione di equivalenza su $\mathbf{X}$, nel senso della Definizione 1.3.5: la classe laterale con rappresentante $\mathbf{x} \in \mathbf{X}$ è

$$\mathbf{x} + \mathbf{V} = \{\mathbf{x} + \mathbf{v} : \mathbf{v} \in \mathbf{V}\}.$$

Il quoziente di $\mathbf{X}$ rispetto a questa relazione di equivalenza si chiama lo *spazio quoziente* $\mathbf{X}/\mathbf{V}$.

NOTA 13.1.2. È immediato verificare che, se $\mathbf{X} = \mathbf{V} \oplus \mathbf{W}$, allora si ha l'isomorfismo $\mathbf{X}/\mathbf{V} \simeq \mathbf{W}$. $\square$

DEFINIZIONE 13.1.3. (*Norma e prodotto scalare nello spazio quoziente.*)

Sia $\mathbf{X}$ uno spazio normato (a dimensione finita), e quindi munito di un prodotto scalare definito positivo indotto dalla norma (Corollario 6.5.6), e $\mathbf{V}$ un sottospazio di $\mathbf{X}$. La norma in $\mathbf{X}$ proietta una norma nello spazio quoziente $\mathbf{X}/\mathbf{V}$ nel modo seguente:

$$\|\tilde{\mathbf{x}}\| = \|\mathbf{x} + \mathbf{V}\| = \min\{\|\mathbf{x} + \mathbf{v}\| : \mathbf{v} \in \mathbf{V}\}. \quad (13.1.1)$$

Grazie al Corollario 6.6.2, questa norma quoziente è univocamente associata all'unico prodotto scalare sullo spazio quoziente da cui essa è indotta, necessariamente definito positivo. Quindi il prodotto scalare su $\mathbf{X}$ proietta un (unico) prodotto scalare su $\mathbf{X}/\mathbf{V}$.

ESEMPIO 13.1.4. Si consideri lo spazio $\mathbb{R}^n$ col consueto prodotto scalare euclideo (6.3.1). Sia $n = m + k$, con $m, k > 0$, e quindi $\mathbb{R}^n = \mathbb{R}^m \oplus \mathbb{R}^k$. Il quoziente $\mathbb{R}^n/\mathbb{R}^k$ è isomorfo a $\mathbb{R}^m$. La norma quoziente indotta su

$\mathbb{R}^m$ dalla norma euclidea di $\mathbb{R}^n$ è data da

$$\|\tilde{\mathbf{x}}\|_{\mathbb{R}^m} = \min\{\|\mathbf{x} + \mathbf{v}\|_{\mathbb{R}^n} : \mathbf{v} \in \mathbb{R}^k\} = \|P_{\mathbb{R}^m}\mathbf{x}\|_{\mathbb{R}^n},$$

dove $P_{\mathbb{R}^m}$ è la proiezione canonica sul primo addendo della somma diretta (Nota 8.0.5), e quindi $P_{\mathbb{R}^m}\mathbf{x} = (x_1, x_2, \dots, x_m, 0, 0, \dots, o)$. In altre parole, la norma indotta sul quoziente $\mathbb{R}^m$ dalla norma euclidea sullo spazio $\mathbb{R}^n$ è precisamente la norma euclidea di $\mathbb{R}^m$. Identico argomento vale in $\mathbb{C}^n$.

Poiché la norma euclidea è quella indotta dal prodotto scalare euclideo (Nota 6.5.8), il fatto che la norma euclidea di $\mathbb{R}^n$ proiettata sullo spazio quoziente $\mathbb{R}^m$ coincida con la norma euclidea del quoziente equivale a dire che il prodotto scalare euclideo proiettato sul quoziente coincide con il prodotto scalare euclideo in $\mathbb{R}^m$. $\square$

LEMMA 13.1.5. *Per ogni sottospazio $\mathbf{V}$ di $\mathbf{X}$, il duale dello spazio quoziente $\mathbf{X}/\mathbf{V}$ consiste dei funzionali in $\mathbf{X}'$ che sono costanti sulle classi laterali di $\mathbf{V}$ in $\mathbf{X}$ (cioè che hanno lo stesso valore sui rappresentanti di ciascuna classe $\mathbf{x} + \mathbf{V}$).*

DIMOSTRAZIONE. Sia $\mathbf{x}' \in \mathbf{X}'$ tale che $\mathbf{x}'$ si annulla su $\mathbf{V}$, ossia $\mathbf{x}'(\mathbf{y} + \mathbf{v}) = \mathbf{x}'(\mathbf{y})$ per ogni $\mathbf{v} \in \mathbf{V}$. Allora $\mathbf{x}'$ induce un funzionale lineare sullo spazio quoziente $\mathbf{X}/\mathbf{V}$, che indichiamo con $\tilde{P}_{\mathbf{V}}(\mathbf{x}')$ (oppure con $\tilde{P}_{\mathbf{V}}\mathbf{x}'$ quando non c'è adito a confusione), definito dalla regola naturale:

$$\tilde{P}_{\mathbf{V}}(\mathbf{x}')(\mathbf{y} + \mathbf{V}) = \mathbf{x}'(\mathbf{y})$$

per ogni $\mathbf{y} \in \mathbf{X}$. Questo funzionale è ben definito sullo spazio quoziente perché per ipotesi i valori che assume non dipendono dalla scelta del rappresentante delle classi laterali del quoziente, anche se la definizione che abbiamo scritto fa riferimento al rappresentante $\mathbf{y}$.

Viceversa, ogni funzionale lineare $\mathbf{f}$ sul quoziente $\mathbf{X}/\mathbf{V}$ induce un funzionale lineare $\tilde{\mathbf{f}}$ su $\mathbf{X}$ costante sulle classi laterali di $\mathbf{V}$ secondo la regola naturale simmetrica della precedente:

$$\mathbf{f}(\mathbf{y}) = \mathbf{f}(\mathbf{y} + \mathbf{V}).$$

NOTA 13.1.6. L'operatore lineare

$$\tilde{P}_{\mathbf{V}} : \{\mathbf{x}' \in \mathbf{X}'; \mathbf{x}' \text{ si annulla su } \mathbf{V}\} \longrightarrow (\mathbf{X}/\mathbf{V})'$$

introdotto nella dimostrazione del precedente Lemma 13.1.5 è un isomorfismo e verifica $\tilde{P}_{\mathbf{V}}\tilde{\mathbf{f}} = \mathbf{f}$. $\square$

LEMMA 13.1.7. *Un funzionale lineare $\mathbf{x}' \in \mathbf{X}'$ è costante sulle classi laterali del sottospazio $\mathbf{V} \subset \mathbf{X}$*

- *se e soltanto se esiste un funzionale lineare $\mathbf{f} \in (\mathbf{X}/\mathbf{V})'$ tale che $\mathbf{x}' = \tilde{P}_{\mathbf{V}}^{-1}\mathbf{f}$;*
- *se e soltanto se $P_{\mathbf{V}}^{\sharp}\mathbf{f} = \mathbf{f} \mid_{\mathbf{V}} = \mathbf{0}_{\mathbf{V}'}$*

DIMOSTRAZIONE. La prima parte è chiara. Per la seconda basta osservare che, grazie alla linearità, un funzionale lineare $\mathbf{x}'$ è costante sulle classi laterali di $\mathbf{V}$ se e solo se assume il valore 0 sulla classe laterale $\mathbf{0} + \mathbf{V} = \mathbf{V}$, cioè se si annulla su $\mathbf{V}$, quindi se e solo se è nel nucleo della proiezione di restrizione $P_{\mathbf{V}}^{\sharp}$.

Si noti invece che la caratterizzazione dei funzionali che *coincidono fra loro* sulle classi laterali del sottospazio è la seguente, come si ricava facilmente facendo uso della linearità:

* ESERCIZIO 13.1.8. Due funzionali lineari $\mathbf{x}'_1, \mathbf{x}'_2$ su $\mathbf{X}$ coincidono sulle classi laterali del sottospazio $\mathbf{V} \subset \mathbf{X}$

- se e solo se $\mathbf{x}'_1 - \mathbf{x}'_2$ vale zero su $\mathbf{V}$,
- se e solo se $\mathbf{x}'_1 - \mathbf{x}'_2 \in \ker P_{\mathbf{V}}^{\sharp}$ (dove $P_{\mathbf{V}}^{\sharp}$ è la proiezione per restrizione introdotta in (10.3.1) della Nota 10.3.3),
- se e solo se $\mathbf{x}'_1$ e $\mathbf{x}'_2$ sono equivalenti secondo la relazione di equivalenza data dal fare il quoziente per $\ker P_{\mathbf{V}}^{\sharp}$,
- se e solo se $\mathbf{x}'_1$ e $\mathbf{x}'_2$ appartengono alla stessa classe laterale in $\mathbf{X}/\ker P_{\mathbf{V}}^{\sharp}$.

LEMMA 13.1.9. $\ker P_{\mathbf{V}}^{\sharp} = \mathbf{V}^{\perp}$, dove ora lo spazio ortogonale $\mathbf{V}^{\perp}$ è inteso come un sottospazio del duale $\mathbf{X}'$, nel senso della Nota 10.3.7, tramite l'immersione di uno spazio vettoriale nel suo duale introdotta nella Definizione 10.3.2.

DIMOSTRAZIONE. Grazie alla Nota 10.3.7, si ha

$$\ker P_V^\# = \{\mathbf{x}' \in \mathbf{X}' : \mathbf{x}'|_{\mathbf{V}} = \mathbf{0}\} = \{\mathbf{x}' : \mathbf{x}'(\mathbf{v}) = 0 \forall \mathbf{v} \in \mathbf{V}\} = \mathbf{V}^\perp.$$

In seguito ai Lemmi 13.1.5, 13.1.7 e 13.1.9 ora concludiamo che:

PROPOSIZIONE 13.1.10. *Per ogni sottospazio $\mathbf{V}$ di $\mathbf{X}$, il duale dello spazio quoziente $\mathbf{X}/\mathbf{V}$ è isomorfo al quoziente $\mathbf{X}'/\mathbf{V}^\perp$, tramite l'isomorfismo della Definizione 13.1.3.*

ESEMPIO 13.1.11. (**Polinomi omogenei.**) Questo esempio si può considerare in un numero arbitrario di variabili, ma è piuttosto banale in una variabile. Per semplicità di notazione lo formuliamo nel caso di due variabili x e y , ma la costruzione vale in generale.

Consideriamo lo spazio vettoriale $P_n[x, y]$ dei polinomi di grado al più n in due variabili, i cui elementi sono

$$\sum_{i+j \leq n} a_{ij} x^i y^j$$

con a_{ij} reali (o più in generale complessi). Osserviamo che ogni polinomio si scompone in termini polinomiali omogenei:

$$\sum_{i+j \leq n} a_{ij} x^i y^j = \sum_{k=0}^n \sum_{i+j=k} a_{ij} x^i y^j = \sum_{k=0}^n \sum_{i=0}^k c_i^{(k)} x^i y^{k-i}.$$

Allora lo spazio quoziente $Q_n[x, y] := P_n[x, y]/P_{n-1}[x, y]$ è isomorfo allo spazio vettoriale dei polinomi omogenei di grado n , $\sum_{k=0}^n \sum_{i+j=k} a_{ij} x^i y^j$.

□

13.2. Lo spazio proiettivo

Abbiamo visto nella Definizione 13.1.1 come, a partire da uno spazio vettoriale $\mathbf{X}$, si può introdurre lo spazio quoziente rispetto a (la relazione di equivalenza generata da) un sottospazio $\mathbf{X}$. Più in generale, possiamo considerare relazioni di equivalenza $\sim$ su sottoinsiemi $\mathbf{X}$ di uno spazio vettoriale, che non rispettano necessariamente le operazioni algebriche di somma e di moltiplicazione per scalari. In tal caso, definiamo il quoziente $\mathbf{X}/\sim$ non è uno spazio vettoriale. Ad

esempio, l'*equivalenza per dilatazione* identifica due vettori non nulli in $\mathbb{R}^n$ se sono multipli uno dell'altro. Le classi di equivalenza sono le rette passanti per l'origine tolta l'origine stessa, che possiamo pensare parametrizzati dai versori modulo il segno, cioè dai punti della sfera unitaria con punti antipodali identificati (le direzioni dall'origine verso l'infinito). Se si preferisce separare i segni, cioè avere una relazione di equivalenza le cui classi siano parametrizzate dai punti della sfera, si devono considerare equivalenti solo vettori multipli l'uno dell'altro per scalari positivi.

Sommando i vettori di una retta con quelli di un'altra otteniamo una figura piana, non una retta, quindi la relazione di equivalenza non rispetta la somma; ovviamente non rispetta la moltiplicazione per scalari, perché moltiplicando per 0 usciamo dalle classi di equivalenza. Però rispetta la moltiplicazione per scalari non nulli.

NOTA 13.2.1. (**Proiezione stereografica.**) Una parte delle classi di equivalenza della relazione di equivalenza per dilatazione (cioè le rette uscenti dall'origine, scavate dell'origine stessa) sono in corrispondenza biunivoca con i punti di una sfera tangente all'origine in $\mathbb{R}^n$, detta sfera stereografica. Infatti, se indichiamo con σ l'iperpiano (cioè il sottospazio a dimensione $n-1$) di tangenza, le classi di equivalenza date dalle rette che non giacciono su σ sono in corrispondenza biunivoca con i punti della sfera eccetto il punto di tangenza, ed anche con i punti dell'altro iperpiano τ tangente a questa sfera e parallelo a σ . Ad esempio, se la sfera ha diametro di lunghezza 1 e se per σ scegliamo l'iperpiano $\{\mathbf{x} : x_n = 0\}$, si ha $\tau = \{\mathbf{x} : x_n = 1\}$. In tal modo si crea una corrispondenza fra i punti di quest'ultimo iperpiano (che è una copia di $\mathbb{R}^{n-1}$) e quelli della sfera stereografica tolta l'origine: questa corrispondenza si chiama *proiezione stereografica*. Si noti che nella proiezione stereografica punti della sfera che si muovono verso l'origine sono associati a punti di τ che si muovono verso l'infinito di $\mathbb{R}^{n-1}$. Il caso a dimensione due è illustrato nella Figura 1.

Si osservi che, in dimensione superiore a 2, le rette *orizzontali*, cioè che giacciono in σ , costituiscono (tolta l'origine) differenti classi di equivalenza per dilatazione, che identificano differenti direzioni versi

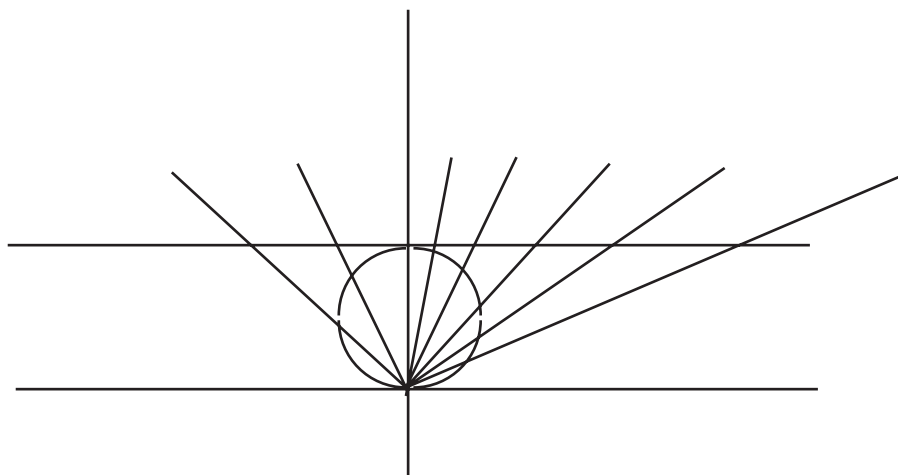


FIGURA 1. Proiezione stereografica in due dimensioni

l'infinito, e si parametrizzano ovviamente con una sfera a dimensione $n - 1$ (nel caso di $\mathbb{R}^3$, con un cerchio, il cerchio di raggio, ad esempio, 1 nel piano di base $\sigma = \{\mathbf{x} : x_n = 0\}$). Invece, sulla sfera stereografica, c'è un unico punto sul piano di base σ , e quindi la proiezione stereografica non separa le direzioni verso l'infinito sul piano di tangenza. $\square$

A differenza dell'esempio precedente, siamo interessati ad equivalenze che separino le direzioni all'infinito. Chiariamo meglio questo concetto, prendendo ad esempio il piano. Consideriamo le rette nel piano, e dichiariamo due di esse equivalenti se sono parallele: possiamo pensare che le classi di equivalenza individuino le *direzioni all'infinito*. Questo significa pensare la direzione del vettore $\mathbf{x} = (x_1, x_2)$ come la pendenza della retta formata da tutti i multipli di $\mathbf{x}$. Questa pendenza è data dal rapporto $\frac{x_1}{x_2}$ (che può essere infinito, il che accade quando la retta è l'asse $x_1 = 0$).

DEFINIZIONE 13.2.2. (*Spazio proiettivo.*) Vogliamo aggiungere ai punti di $\mathbb{R}^2$ un ulteriore insieme di *punti all'infinito*, o forse dovremmo dire *direzioni all'infinito*, parametrizzato dai punti di un cerchio, come spiegato nella Nota 13.2.1. A questo scopo prendiamo in esame una variante dell'ambiente geometrico tridimensionale introdotto in quella Nota per la proiezione stereografica, con il piano di base $\sigma = \{\mathbf{x} : x_0 =$

$0\}$ ed il suo traslato $\tau = \{\mathbf{x} : x_0 = 1\}$. Introduciamo su $\mathbb{R}^3 \setminus \{0\}$ una relazione di equivalenza che identifichi i punti di τ con i punti di uno spazio vettoriale bidimensionale $\mathbb{R}^2$, ed i punti del cerchio unitario (o di un qualsiasi raggio non nullo) con centro l'origine in σ con i punti all'infinito di $\mathbb{R}^2$.

Questa relazione di equivalenza è nient'altro che l'equivalenza per dilatazioni non nulle già introdotta all'inizio di questa Sezione: due punti in $\mathbb{R}^3$ diversi dall'origine, diciamo (x_0, x_1, x_2) e (x'_0, x'_1, x'_2) sono equivalenti se esiste uno scalare $\lambda \neq 0$ tale che

$$(x'_0, x'_1, x'_2) = (\lambda x_0, \lambda x_1, \lambda x_2) .$$

In particolare, se $x_0 \neq 0$, cioè se il vettore non si trova nel piano di base σ , allora ogni classe di equivalenza ha un rappresentante nel piano τ , ed uno solo, dato dal punto $(1, \frac{x_1}{x_0}, \frac{x_2}{x_0})$; invece per la classe di equivalenza di un punto del piano di base, cioè del tipo $\mathbf{v} = (0, x_1, x_2)$, si può scegliere il rappresentante dato dal punto del cerchio unitario in σ con la stessa direzione di $\mathbf{v}$, cioè dal versore $\mathbf{v}/\|\mathbf{v}\|$.

Lo spazio quoziente di $\mathbb{R}^3 \setminus \{0\}$ rispetto a questa relazione di equivalenza, cioè l'insieme di tutte le classi di equivalenza si indica con $\mathbb{P}^2$. Come già osservato, non costituisce uno spazio vettoriale, ed infatti l'operazione di somma non passa al quoziente, perché dipende dalla scelta dei rappresentanti, e quindi non è ben definita sulle classi di equivalenza.

Più precisamente:

ESEMPIO 13.2.3. La classe di equivalenza per dilatazione del vettore $s(x_0, x_1, x_2) + t(x'_0, x'_1, x'_2)$ dipende dalla scelta degli scalari s e t , a meno che (x_0, x_1, x_2) e (x'_0, x'_1, x'_2) siano multipli non nulli l'uno dell'altro. Infatti, basta osservare che per $s = 0$ la classe di equivalenza è quella di (x'_0, x'_1, x'_2) , mentre per $t = 0$ è quella di (x_0, x_1, x_2) . $\square$

In tal modo lo spazio proiettivo $\mathbb{P}^2$ parametrizza i punti al finito ed all'infinito grazie a tre coordinate (x_0, x_1, x_2) , dette *coordinate omogenee* perché in esse le equazioni delle rette, che in $\mathbb{R}^2$ sono del tipo $a_1x_1 + a_2x_2 + a_0 = 0$, diventano (grazie all'equivalenza per dilatazioni

non nulle) equazioni omogenee, del tipo $a_0x_0 + a_1x_1 + a_2x_2 = 0$. I punti all'infinito, come abbiamo visto, sono geometricamente in corrispondenza biunivoca con quelli del cerchio unitario nel piano $x_0 = 0$, ma (sempre grazie all'equivalenza per dilatazioni non nulle) sono anche in corrispondenza con i punti (x_0, x_1, x_2) che verificano l'equazione $x_0 = 0$, la quale è l'equazione di un piano in tre dimensioni (il piano $x_0 = 0$, appunto), ovvero di una retta in due dimensioni.

Analogamente si tratta il caso della geometria proiettiva in $n + 1$ dimensioni reali ($\mathbb{P}^n(\mathbb{R})$), o anche, più in generale, complesse ($\mathbb{P}^n(\mathbb{C})$).

Per comodità nell'uso degli indici, d'ora in avanti permuteremo ciclicamente le coordinate ed indicheremo i punti al finito con $(x_1, x_2, \dots, x_{n-1}, 1)$ e quelli all'infinito con $(x_1, x_2, \dots, x_{n-1}, 0)$. Quindi l'iperpiano dei punti al finito in $\mathbb{P}^{n-1}$ ora ha equazione $x_n = 1$ e quello dei punti all'infinito $x_n = 0$.

13.3. (*) $\mathbb{P}^n$ come compattificazione di $\mathbb{R}^n$

In $\mathbb{R}^{n+1}$ due punti sono vicini a meno di, diciamo, $\varepsilon > 0$ se il secondo giace in una palla (aperta) di raggio ε centrata nel primo. Diciamo che queste palle generano la topologia di $\mathbb{R}^{n+1}$, cioè danno luogo alla consueta nozione di convergenza di successioni. La nozione corrispondente in $\mathbb{P}^n$ deve essere basata su insiemi aperti costituiti da classi di equivalenza per dilatazione non nulla (semirette scavate dell'origine). Ad ogni palla aperta B in $\mathbb{R}^{n+1}$ associamo un *aperto* C_B in $\mathbb{P}^n$ nel modo seguente: C_B è l'insieme delle rette in $\mathbb{R}^{n+1}$ uscenti dall'origine (tolta l'origine stessa) che intersecano B : quindi un cono scavato dell'origine. Ciascuno dei suddetti coni in $\mathbb{R}^{n+1}$ interseca l'iperpiano $(x_1, x_2, \dots, x_n, 1)$ dei punti al finito di $\mathbb{P}^n$ in una ellisse. Poiché ciascuna di queste ellissi è iscritta e circoscritta a palle sferiche dell'iperpiano, è evidente che la nozione di convergenza così indotta sul sottoinsieme dei punti al finito di $\mathbb{P}^n$ coincide con quella usuale di $\mathbb{R}^n$. Invece, sull'iperpiano $(x_1, x_2, \dots, x_n, 0)$, i coni $n + 1$ -dimensionali di cui sopra hanno per intersezioni ancora coni (a dimensione n). Quindi, sul sottoinsieme dei punti all'infinito, la nozione di convergenza riguarda le direzioni all'infinito: una successione di classi converge ad una classe

data se e solo se, scelto per ciascuna classe un vettore per rappresentante, i rappresentanti delle classi della successione hanno direzioni che convergono a quella della classe limite (la convergenza quindi avviene sulla sfera unitaria).

Consideriamo allora una successione di punti al finito in $\mathbb{P}^n$, cioè punti dell'iperpiano $x_{n+1} = 1$. In questo iperpiano la nozione di convergenza di $\mathbb{P}^n$, come abbiamo visto, è quella usuale. Allora, se questi punti formano una successione che tende all'infinito, essi si allontanano dall'origine, e quindi le rette uscenti dall'origine che li contengono diventano via via più orizzontali. Poiché la topologia di $\mathbb{P}^n$ è data da coni come descritto sopra, una successione di punti al finito tende ad un punto all'infinito esattamente quando, dato un qualunque cono aperto che ha per asse la retta orizzontale corrispondente a quel punto all'infinito, le rette che essi descrivono giacciono, da un dato indice in poi, dentro tale cono. È facile verificare che questa nozione di convergenza equivale a dire che le direzioni dei punti al finito tendono a quella del punto limite all'infinito (nella naturale nozione di convergenza della sfera unitaria data dai versori in $\mathbb{R}^{n+1}$).

Con le definizioni e gli argomenti sviluppati nell'Appendice (Capitolo 17), Sezione 17.1, si vede che ogni successione di punti di $\mathbb{P}^n$ ammette una sottosuccessione convergente. Se la successione di punti in $\mathbb{P}^n$ è una successione di punti al finito che, vista nell'iperpiano $x_{n+1} = 1$, è illimitata, allora essa ammette una sottosuccessione che converge ad un punto all'infinito nella nozione di convergenza di $\mathbb{P}^n$ (e cioè nel senso che convergono le direzioni). Adottando di nuovo la terminologia della Sezione 17.1 dell'Appendice, diciamo che lo spazio proiettivo $\mathbb{P}^n$ è una *compattificazione* di $\mathbb{R}^n$.

13.4. Trasformazioni proiettive

Abbiamo già introdotto il gruppo lineare $GL_n(\mathbb{R})$, cioè delle matrici reali invertibili a dimensione n . (Notazione 3.9.3). Ora vogliamo considerare la forma matriciale degli operatori lineari sullo spazio proiettivo $\mathbb{P}^{n-1}$. Per questo scopo fissiamo una scelta privilegiata di rappresentanti nelle classi di equivalenza, quella legata alla rappresentazione

stereografica introdotta nella precedente Sezione 13.2: per i rappresentanti delle classi al finito scegliamo $(x_1, x_2, \dots, x_{n-1}, 1)$, mentre per le classi all'infinito $(x_1, x_2, \dots, x_{n-1}, 0)$.

Come visto nella Sezione 13.2, per le classi in $\mathbb{P}^{n-1}$ usiamo n coordinate *omogenee*, definite univocamente solo a meno di dilatazioni, rammentandoci però del fatto che in $\mathbb{P}^{n-1}$ bastano $n - 1$ parametri per identificare le classi. Infatti l'insieme delle classi al finito è isomorfo a $\mathbb{R}^{n-1}$, e ad esse si devono aggiungere quelle all'infinito, che formano una circonferenza in $\mathbb{R}^{n-1}$, o se si preferisce un iperpiano: l'iperpiano $x_0 = 0$. Questo iperpiano, si parametrizza con soli $n - 2$ parametri (ad esempio $n - 2$ variabili cartesiane $x_1, \dots, x_{n-2}$, oppure, se lo si pensa come una circonferenza, $n - 2$ angoli di Eulero). Per evitare ambiguità notazionali, in questa Sezione denotiamo in termini di coordinate il vettore $\mathbf{x} \in \mathbb{R}^n$ con

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

e la sua classe di equivalenza per dilatazione, cioè l'elemento di $\mathbb{P}^{n-1}$ di cui $\mathbf{x}$ è rappresentante, con

$$\begin{bmatrix} x_1 \\ \vdots \\ x_n \end{bmatrix}.$$

Le matrici in $GL_n(\mathbb{R})$ agiscono sui vettori, invece che sulle classi di equivalenza. Però, per linearità, se due vettori $\mathbf{x}$ e $\mathbf{y}$ sono multipli uno dell'altro, $\mathbf{x} = \lambda \mathbf{y}$, ogni matrice M in $GL_n(\mathbb{R})$ verifica $M\mathbf{x} = \lambda M\mathbf{y}$. Quindi M preserva le classi di equivalenza, e perciò l'azione di $GL_n(\mathbb{R})$ su $\mathbb{R}^n$ induce una azione ben definita di $GL_n(\mathbb{R})$ sullo spazio proiettivo $\mathbb{P}^{n-1}$. Però l'azione di matrici diverse può essere identica: per via del fatto che la scelta del rappresentante di una classe è determinata solo a meno di multipli, per ogni $\lambda \neq 0$ reale l'azione di M coincide con quella di λM . Pertanto, viste come operatori sullo spazio proiettivo, M e λM si identificano. Questo porta a definire correttamente il gruppo di operatori lineari che agiscono sullo spazio proiettivo in termini di equivalenza sotto dilatazione, come segue.

DEFINIZIONE 13.4.1. (**Gruppo lineare proiettivo.**) Il gruppo lineare proiettivo $PGL_n(\mathbb{R})$ è il gruppo i cui elementi sono le classi di equivalenza di $GL_n(\mathbb{R})$ sotto la relazione di equivalenza che identifica due matrici M_1 e M_2 se esiste $\lambda \in \mathbb{R}$, $\lambda \neq 0$ tale che $M_1 = \lambda M_2$.

Le operazioni di prodotto e di inverso in $GL_n(\mathbb{R})$ (prodotto righe per colonne) rispettano le classi di equivalenza, perché il prodotto di multipli di due matrici è un multiplo della matrice prodotto, e quindi le operazioni di gruppo su $GL_n(\mathbb{R})$ inducono operazioni analoghe ben definite su $PGL_n(\mathbb{R})$.

DEFINIZIONE 13.4.2. (**Quoziente sotto dilatazione dello spazio delle matrici reali.**) Analogamente alla Definizione 13.4.1, introduciamo lo spazio vettoriale $M_n(\mathbb{R})$ delle matrici $n \times n$ a coefficienti reali, ewd il suo quoziente $PM_n(\mathbb{R})$ modulo l'equivalenza per dilatazione positiva, cioè modulo il sottoinsieme $\{\lambda \mathbb{I}\}$, con $\lambda > 0$ e $\mathbb{I}$ la matrice identità.

NOTAZIONE 13.4.3. In analogia a quanto fatto per $\mathbb{P}^n$, indichiamo nel modo seguente gli elementi di $PGL_n(\mathbb{R})$ e $PM_n(\mathbb{R})$ a partire dalle matrici in $GL_n(\mathbb{R})$ e $PM_n(\mathbb{R})$ che li rappresentano: l'elemento di $PGL_n(\mathbb{R})$ o $PM_n(\mathbb{R})$ associato alla matrice

$$\begin{pmatrix} m_{11} & \cdots & m_{1n} \\ & \ddots & \\ m_{n1} & \cdots & m_{nn} \end{pmatrix}$$

si scrive con le parentesi quadre:

$$\begin{pmatrix} m_{11} & \cdots & m_{1n} \\ & \ddots & \\ m_{n1} & \cdots & m_{nn} \end{pmatrix}$$

ESEMPIO 13.4.4. Per semplificare la notazione, restringiamo in questo esempio la nostra attenzione al caso dello spazio proiettivo tridimensionale $\mathbb{P}^3$, con coordinate omogenee x, y, z, w (questo è il caso di interesse nelle applicazioni della geometria proiettiva alla Computer Graphics, che vedremo in seguito nella Sezione 14.2. Consideriamo una classe con un rappresentante al finito, cioè del tipo $(x, y, z, 1)$. Sia

$M \in PM_4(\mathbb{R})$, ed indichiamo con lettere maiuscole le coordinate dei rappresentanti della classe in cui M manda $(x, y, z, 1)$. La notazione quindi diventa:

$$\begin{bmatrix} X \\ Y \\ Z \\ W \end{bmatrix} = \begin{pmatrix} m_{11} & \cdots & m_{1n} \\ & \ddots & \\ m_{n1} & \cdots & m_{nn} \end{pmatrix} \begin{bmatrix} x \\ y \\ z \\ 1 \end{bmatrix}. \quad (13.4.1)$$

Si osservi che $M \in PM_4(\mathbb{R})$ può mandare una classe al finito, come $[x, y, z, 1]$ sia in un'altra classe al finito, sia in una classe all'infinito, del tipo $[x, y, z, 0]$. Ad esempio, la matrice identità manda ciascuna classe in sé stessa, mentre la (classe di equivalenza dell) matrice di proiezione sui punti all'infinito,

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

manda $[x, y, z, 1]$ in $[x, y, z, 0]$, ed infine, se $\lambda \neq 0$,

$$\begin{pmatrix} 1 & 0 & 0 & 0 \\ & \ddots & & \\ 0 & \cdots & 1 & 0 \\ 0 & 0 & \lambda & 0 \end{pmatrix} \quad (13.4.2)$$

manda $[x, y, z, 0]$ in $[x, y, z, \lambda z]$, che, almeno per $z \neq 0$, è una classe al finito. $\square$

CAPITOLO 14

* Trasformazioni affini

14.1. Moti rigidi in $\mathbb{R}^n$ immersi in trasformazioni lineari di $\mathbb{R}^{n+1}$

L'insieme delle trasformazioni lineari di $\mathbb{R}$ in sé è l'insieme dei funzionali lineari su $\mathbb{R}$, e quindi, per il Teorema 10.2.5, è uno spazio vettoriale isomorfo a $\mathbb{R}$: ogni $a \in \mathbb{R}$ corrisponde alla trasformazione lineare su $\mathbb{R}$ data dalla dilatazione $x \mapsto ax$ (una costante a è infatti la più generale matrice 1×1).

Consideriamo la famiglia (anzi il *gruppo*, nel senso della Appendice 1.3) data da tutte le composizioni di queste trasformazioni lineari (le dilatazioni) e le *traslazioni* di $\mathbb{R}$, cioè le applicazioni (non lineari perché non preservano l'origine!) di $\mathbb{R}$ in sé definite da $x \mapsto x + b$, dove $b \in \mathbb{R}$. È facile costruire esempi che mostrano che questo è un gruppo non commutativo. Esso non consiste quindi di matrici 1×1 : però ora mostriamo che i suoi elementi possono essere rappresentati come matrici 2×2 , cioè di applicazioni lineari su $\mathbb{R}^2$.

Infatti:

LEMMA 14.1.1. *Immergiamo la retta reale come una retta nello spazio bidimensionale $\mathbb{R}^2$, ad esempio come la retta $\{x_2 = 1\}$, parallela all'asse delle ascisse ad ordinata 1. Allora esiste una applicazione lineare di $\mathbb{R}^2$ in sé che preserva questa retta, e su di essa agisce traslandone l'origine (cioè il punto $(0, 1)$ e dilatando gli intervalli di un fattore a .*

DIMOSTRAZIONE. Per prima cosa troviamo quali operatori lineari su $\mathbb{R}^2$ preservano la retta $\{x_2 = 1\}$. Scriviamo questa condizione come una condizione sulla matrice A di un tale operatore nella base canonica. La

condizione è che per ogni $s \in \mathbb{R}$ esista $t \in \mathbb{R}$ tale che

$$A \begin{pmatrix} s \\ 1 \end{pmatrix} = \begin{pmatrix} t \\ 1 \end{pmatrix}.$$

In particolare, se si sceglie $s = 0$, si vede che A manda il secondo vettore della base canonica in un vettore con ordinata 1, e quindi la seconda colonna di A è del tipo $(b, 1)$ per qualche $b \in \mathbb{R}$. D'altra parte, la differenza fra due vettori sulla retta $\{x_2 = 1\}$ è un vettore proporzionale a $\mathbf{e}_1 = (1, 0)$, e quindi questo vettore viene mandato in un multiplo $(a, 0)$ di $\mathbf{e}_1$. Pertanto la prima colonna della matrice A è del tipo $(a, 0)$ per qualche $a \in \mathbb{R}$. Riassumendo si ha:

$$A = \begin{pmatrix} a & b \\ 0 & 1 \end{pmatrix}.$$

Osserviamo che questa matrice manda $(s, 1)$ in $(as+b, 1)$, e quindi trasla l'origine della retta $\{x_2 = 1\}$ di una quantità b , e dilata gli intervalli di questa retta di un fattore a .

In tre dimensioni il problema corrispondente è quello di quali operatori lineari agiscono preservando il piano $x_3 = 1$ e traslandone l'origine di un vettore *orizzontale* $(v_1, v_2, 0)$. Esattamente lo stesso argomento di prima porta ora alla conclusione seguente.

PROPOSIZIONE 14.1.2. *Le trasformazioni lineari di $\mathbb{R}^3$ in $\mathbb{R}^3$ che lasciano invariante il piano $\{x_3 = 1\}$ e traslano il punto $(0, 0, 1)$ (l'origine del piano affine) nel punto $(v_1, v_2, 1)$ sono precisamente quelle la cui espressione matriciale (nella base canonica) è triangolare a blocchi, del tipo*

$$A = \begin{pmatrix} a & b & v_1 \\ c & d & v_2 \\ 0 & 0 & 1 \end{pmatrix}$$

con $a, b, c, d \in \mathbb{R}$.

Ora riscriviamo queste considerazioni per $\mathbb{R}^n$.

DEFINIZIONE 14.1.3. (*Trasformazioni affini.*) Una trasformazione di $\mathbb{R}^n$ in sé che si ottiene componendo una traslazione con una trasformazione lineare si chiama una *trasformazione affine*.

PROPOSIZIONE 14.1.4. *Se identifichiamo $\mathbb{R}^n$ con l'iperpiano di equazione $x_{n+1} = 1$ in $\mathbb{R}^{n+1}$, le trasformazioni affini di $\mathbb{R}^n$ corrispondono alle trasformazioni lineari di $\mathbb{R}^{n+1}$ la cui matrice (nella base canonica) ha forma triangolare a blocchi del tipo*

$$A_{n+1} = \begin{pmatrix} a_{11} & \dots & a_{1n} & v_1 \\ & \ddots & & \\ a_{n1} & \dots & a_{nn} & v_n \\ 0 & \dots & 0 & 1 \end{pmatrix}$$

dove il vettore $(v_1, v_2, \dots, v_n)$ è il punto in cui viene traslata l'origine di $\mathbb{R}^n$.

Una trasformazione affine T è identificata da dove manda l'origine e dove manda i vettori di una base. Supponiamo che mandi l'origine in $\mathbf{v}$ ed il vettore $\mathbf{e}_j$ della base canonica in $\mathbf{u}_j$, per $j = 1, \dots, n$. Allora T è identificata anche dal vettore $\mathbf{v}$ immagine dell'origine e dai *vettori applicati* a $\mathbf{v}$ dati da $\mathbf{v}_j = \mathbf{e}_j - \mathbf{v}$. Più precisamente:

DEFINIZIONE 14.1.5. (**Spazio affine, vettori applicati e trasformazioni affini.**) Lo spazio affine reale $\mathbf{X}$ a dimensione n è lo spazio $\mathbb{R}^n$ con una scelta di origine $\mathbf{o}$ da cui calcolare le coordinate, cioè è la coppia $\mathbf{o} \times \mathbb{R}^n \subset \mathbb{R}^{2n}$; il vettore $\mathbf{o} \in \mathbb{R}^n$ si chiama il punto di applicazione dei vettori. Il generico elemento $(\mathbf{o}, \mathbf{v})$ di $\mathbf{X}$ si chiama vettore applicato al punto $\mathbf{o}$. Per evitare ambiguità fra vettori e vettori applicati, denotiamo il vettore applicato $(\mathbf{o}, \mathbf{v})$ con $\overrightarrow{\mathbf{o}\mathbf{v}}$.

COROLLARIO 14.1.6. *Una trasformazione affine T su $\mathbf{X}$ sposta l'origine in una nuova origine $\mathbf{o}'$, ed i vettori applicati a $\mathbf{o}$ in vettori applicati a $\mathbf{o}'$; essa è univocamente determinata dalla nuova origine $\mathbf{o}'$ (che esprime la traslazione che è stata applicata all'origine iniziale $\mathbf{o}$) e dalle immagini dei vettori di una base applicati a $\mathbf{o}$, che sono vettori applicati a $\mathbf{o}'$.*

NOTA 14.1.7. **Cautela:** poichè una trasformazione affine sposta l'origine, essa manda vettori applicati all'origine in vettori applicati altrove (lo spostamento essendo dato dalla componente di traslazione). Perciò

non è più vero che la matrice che la rappresenta ha per colonne l'immagine dei vettori canonici di base applicati all'origine (Nota 4.2.2): bisogna traslarli perchè siano applicati alla nuova origine. La forma che la matrice assume viene presentata nel prossimo enunciato.

□

COROLLARIO 14.1.8. (i) *La traslazione di $\mathbb{R}^n$ che manda l'origine in $\mathbf{v}$ (e quindi che sposta gli assi coordinati in assi a questi paralleli, ossia manda i vettori $\mathbf{e}_i$ della base canonica in $\mathbf{e}_i + \mathbf{v}$) ha per matrice*

$$A = \begin{pmatrix} 1 & 0 & \dots & 0 & v_1 \\ & & \ddots & & \\ 0 & 0 & \dots & 1 & v_n \\ 0 & 0 & \dots & 0 & 1 \end{pmatrix}$$

Osserviamo che, quando consideriamo lo spazio affine come immerso in $\mathbb{R}^{n+1}$, l' i -simo vettore della base canonica diventa $\mathbf{e}_i = (0, \dots, 0, 1, 0, \dots, 0, 1)$ con il primo 1 alla componente i -sima, e la matrice A manda questo vettore in $\mathbf{w}_i := \mathbf{e}_i + \mathbf{v} = (v_1, v_2, \dots, v_{i-1}, v_i + 1, v_{i+1}, \dots, v_n, 1)$. La colonna i -sima della matrice è quindi il vettore $\mathbf{e}_i = \mathbf{w}_i - \mathbf{v}$.

(ii) *Più in generale, la trasformazione affine che manda l'origine in $\mathbf{o}' \equiv \mathbf{v}$ ed il vettore $\mathbf{e}_j$ della base canonica in $\mathbf{w}_j$, cioè il vettore applicato $\overrightarrow{\mathbf{o}\mathbf{e}_j}$ nel vettore applicato $\overrightarrow{\mathbf{v}\mathbf{w}_j} = \mathbf{w}_j - \mathbf{v}$ (per $j = 1, \dots, n$) è*

$$A = \begin{pmatrix} w_{11} - v_1 & w_{21} - v_1 & \dots & w_{n1} - v_1 & v_1 \\ & & \ddots & & \\ w_{1n} - v_n & w_{2n} - v_n & \dots & w_{nn} - v_n & v_n \\ 0 & 0 & \dots & 0 & 1 \end{pmatrix}$$

dove abbiamo posto $w_{ij} = (\mathbf{w}_i)_j$, la j -sima componente del vettore $\mathbf{w}_i$. Si noti che le colonne della sottomatrice $n \times n$ in alto a sinistra di A sono date dalle coordinate dei vettori applicati $\overrightarrow{\mathbf{v}\mathbf{w}_j}$ immagini dei vettori canonici di base (che si

possono pensare applicati all'origine) $\mathbf{e}_j = \overrightarrow{\mathbf{o}\mathbf{e}_j}$. La colonna i -sima della matrice è quindi il vettore $\mathbf{w}_i - \mathbf{v}$.

DIMOSTRAZIONE. La cui verifica è immediata: basta applicare la trasformazione affine ai vettori canonici di base $\mathbf{e}_i$.

Come nella Nota 4.1.3 ora abbiamo:

COROLLARIO 14.1.9. (**Cambio di riferimento affine.**) Sia $\{\overrightarrow{\mathbf{o}\mathbf{e}_j}, j = 1, \dots, n\}$ una base di vettori applicati a $\mathbf{o}$, e $\{\overrightarrow{\mathbf{v}\mathbf{w}_j}, j = 1, \dots, n\}$ una base di vettori applicati a $\mathbf{v} \equiv \mathbf{o}'$. Indichiamo con x_i le componenti di un vettore $\overrightarrow{\mathbf{o}\mathbf{x}}$ applicato ad $\mathbf{o}$ nella prima base, e con x'_i le componenti del vettore applicato a $\mathbf{v}$ che si ottiene dallo stesso punto spaziale $\mathbf{x}$, cioè il vettore applicato $\overrightarrow{\mathbf{v}\mathbf{x}}$: in altre parole, sia

$$\overrightarrow{\mathbf{o}\mathbf{x}} = \sum_{i=1}^n x_i \overrightarrow{\mathbf{o}\mathbf{e}_i}$$

e

$$\overrightarrow{\mathbf{o}'\mathbf{x}} = \sum_{i=1}^n x'_i \overrightarrow{\mathbf{o}'\mathbf{w}_i}$$

Allora le nuove coordinate affini $(x'_1, \dots, x'_n)$ sono l'immagine delle coordinate $(x_1, \dots, x_n)$ sotto la matrice inversa di quella del Corollario 14.1.8, cioè si ha

$$\begin{pmatrix} x_1 \\ x_2 \\ \vdots \\ x_n \\ 1 \end{pmatrix} = \begin{pmatrix} w_{11} - v_1 & w_{21} - v_1 & \dots & w_{n1} - v_1 & v_1 \\ & & \ddots & & \\ w_{1n} - v_n & w_{2n} - v_n & \dots & w_{nn} - v_n & v_n \\ 0 & 0 & \dots & 0 & 1 \end{pmatrix} \begin{pmatrix} x'_1 \\ x'_2 \\ \vdots \\ x'_n \\ 1 \end{pmatrix}.$$

La matrice inversa che realizza la trasformazione di coordinate si chiama la matrice del cambiamento di riferimento affine.

NOTA 14.1.10. Poiché le trasformazioni lineari mandano parallelogrammi in parallelogrammi, e parallelepipedi in parallelepipedi, lo stesso è vero per ogni trasformazione affine, in quanto essa differisce da una trasformazione lineare per una traslazione. Ad esempio, una trasformazione del piano che manda il quadrato Q di vertici $(0,0)$, $(1,0)$,

$(0, 1)$, $(1, 1)$ nel quadrilatero di vertici $(1, 3)$, $(2, 2)$, $(-1, -1)$, $(3, 2)$ non è affine, perché il quadrilatero non è un parallelogramma.

Viceversa, dato un parallelogramma P_2 nel piano o un parallelepipedo P_3 in $\mathbb{R}^3$, esiste sempre una trasformazione lineare che manda il cubo unitario del rispettivo spazio (rispettivamente $\mathbb{R}^2$ o $\mathbb{R}^3$) in un traslato di P_2 (rispettivamente, P_3) che ha un vertice nell'origine (infatti i vertici del cubo contigui all'origine formano una base nel rispettivo spazio, per cui esiste una trasformazione lineare che li manda nei vertici contigui all'origine in P_2 (rispettivamente P_3), e la regola del parallelogramma asserisce che i vertici restanti vengono mandati da questa trasformazione nei corrispondenti altri vertici del parallelogramma o parallelepipedo.

Quindi esiste sempre una trasformazione affine che manda il cubo unitario in un qualsiasi parallelogramma (in $\mathbb{R}^2$) o parallelepipedo (in $\mathbb{R}^3$).

□

14.2. Esempio: spostamento di una macchina da ripresa

Per praticità, e per il suo significato geometrico, formuliamo questo esempio in tre dimensioni. Vogliamo trovare la forma matriciale della rototraslazione in tre dimensioni senza deflessione laterale del piano verticale (o *tilt*): questa trasformazione è quella che determina il cambiamento degli assi del sistema di riferimento quando un osservatore si sposta senza inclinare lateralmente il piano verticale della macchina da ripresa, e quindi è essenziale in Computer Graphics. In effetti, sarebbe naturale identificare l'asse laterale con l'asse x , l'asse di osservazione con l'asse y , e come terzo asse (di elevazione) l'asse z . In realtà, è tradizione in Computer Graphics scegliere come asse y quello verticale ed invece riservare l'asse z è per la profondità (per una motivazione di questa scelta apparentemente stravagante di coordinate si veda nel seguito la Sezione 16.4). Quindi qui ci conformiamo a questa abitudine: il versore della base canonica che viene spostato nel nuovo versore di osservazione è il terzo versore canonico $\mathbf{e}_3$. Cominciamo con il determinare la rotazione richiesta: questo è il termine lineare 3×3

della trasformazione affine che vogliamo trovare. dopo incorporeremo anche la traslazione per ottenere l'intera trasformazione affine 4×4 .

LEMMA 14.2.1. *Denotiamo come sempre con $\mathbf{e}_1, \mathbf{e}_2, \mathbf{e}_3$ i vettori della base canonica in $\mathbb{R}^3$. Per ogni $\mathbf{w} \neq \mathbf{0}$ prefissato in $\mathbb{R}^3$, la trasformazione lineare di $\mathbb{R}^3$ che manda $\mathbf{e}_3$ in un versore $\mathbf{w}$ ed $\mathbf{e}_1$ in un vettore che giace sul piano di base $\{x_3 = 0\}$ ha per matrice la matrice ortonormale (nel senso della Definizione 6.8.2)*

$$A = \begin{pmatrix} a & b & w_1 \\ c & d & w_2 \\ 0 & e & w_3 \end{pmatrix}$$

dove $a, b, c, d, e \in \mathbb{R}$. In particolare, se la trasformazione manda la base canonica in una base ortogonale $\mathbf{u}, \mathbf{v}, \mathbf{w}$ destrorsa (ossia con con orientamento concorde a quello della base canonica) con il primo versore sul piano di base, le colonne della matrice sono proporzionali a quelle della matrice ortogonale

$$B = \begin{pmatrix} -w_2 & -w_1w_3 & w_1 \\ w_1 & -w_2w_3 & w_2 \\ 0 & w_1^2 + w_2^2 & w_3 \end{pmatrix}$$

Normalizzando i vettori colonna della matrice B e rammentando che $\mathbf{w}$ è già scelto di norma 1, si ottiene la forma esplicita della matrice ortogonale richiesta:

$$A = \begin{pmatrix} -\frac{w_2}{\sqrt{w_1^2 + w_2^2}} & -\frac{w_1w_3}{\sqrt{w_1^2 + w_2^2}} & w_1 \\ \frac{w_1}{\sqrt{w_1^2 + w_2^2}} & -\frac{w_2w_3}{\sqrt{w_1^2 + w_2^2}} & w_2 \\ 0 & \sqrt{w_1^2 + w_2^2} & w_3 \end{pmatrix}$$

DIMOSTRAZIONE. La prima espressione è immediata perché le colonne della matrice sono le immagini dei vettori canonici di base, ed $\mathbf{e}_1$ viene mandato in un vettore con terza componente nulla. La seconda segue dal fatto che un vettore del tipo $(a, b, 0)$ ortogonale a (w_1, w_2, w_3) è multiplo di $(-w_2, w_1, 0)$, e che questi due vettori perpendicolari determinano, a meno di multipli, un solo terzo vettore perpendicolare

ad entrambi, che è $(-w_1w_3, -w_2w_3, w_1^2 + w_2^2)$ (si veda la Sezione ? sul prodotto vettoriale). $\square$

Se la base $\mathbf{u}, \mathbf{v}, \mathbf{w}$ nel Lemma 14.2.1 è ortonormale, la matrice ortogonale che ne risulta si esprime più chiaramente in termini di angoli anziché di coordinate. A questo scopo introduciamo gli angoli di Eulero come segue.

DEFINIZIONE 14.2.2. (*Angoli di Eulero.*) Ogni vettore $\mathbf{x} \in \mathbb{R}^3$ di norma $\rho > 0$ si esprime in termini degli angoli ϕ di latitudine e θ di longitudine (detti *angoli di Eulero*) mediante la seguente trasformazione di coordinate:

$$x = \rho \cos \theta \sin \phi$$

$$y = \rho \sin \theta \sin \phi$$

$$z = \rho \cos \phi,$$

con $0 \leq \theta < 2\pi$, $0 \leq \phi \leq \pi$.

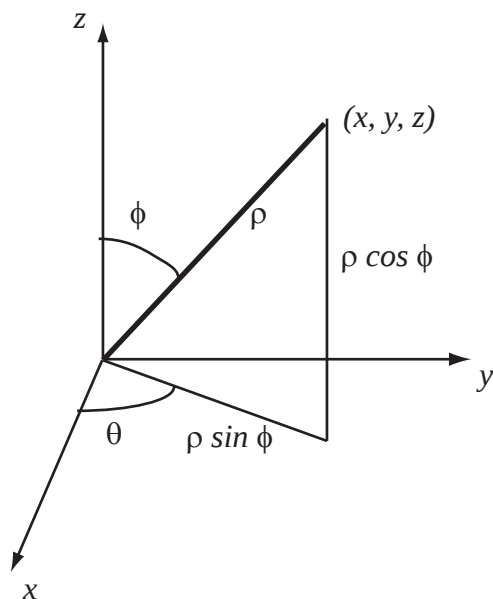


FIGURA 1. Angoli di Eulero

COROLLARIO 14.2.3. Sia $\mathbf{u}$, $\mathbf{v}$, $\mathbf{w}$ una base ortonormale destrorsa in $\mathbb{R}^3$, e scriviamo il vettore $\mathbf{w}$ in termini degli angoli di Eulero: $\mathbf{w} = (\cos \theta \sin \varphi, \sin \theta \sin \varphi, \cos \varphi)$. Allora la matrice ortonormale A del Lemma 14.2.1 diventa

$$A = \begin{pmatrix} -\sin \theta & \cos \theta \cos \varphi & \cos \theta \sin \varphi \\ \cos \theta & \sin \theta \cos \varphi & \sin \theta \sin \varphi \\ 0 & -\sin \varphi & \cos \varphi \end{pmatrix}$$

oppure la matrice che si ottiene da questa cambiando i segni dei coefficienti delle prime due colonne (il che corrisponde a ruotare la base ortonormale di un angolo π intorno all'asse $\mathbf{w}$).

DIMOSTRAZIONE. Il vettore $\mathbf{u}$ deve giacere nel piano $\{z = 0\}$, quindi la condizione di ortogonalità con $\mathbf{w} = (\cos \theta \sin \varphi, \sin \theta \sin \varphi, \cos \varphi)$ equivale a richiedere che $\mathbf{u}$ sia proporzionale a $(-\sin \theta \sin \varphi, \cos \theta \sin \varphi, 0)$. Osserviamo che la norma di quest'ultimo vettore è $\sqrt{\sin^2 \varphi} = |\sin \varphi|$, e $\sin \varphi \geq 0$ perché $0 \leq \varphi \leq \pi$. Ne segue che $\mathbf{u} = \pm(-\sin \theta, \cos \theta, 0)$. (**Nota:** anche senza applicare la definizione di prodotto scalare euclideo, il fatto che il vettore $\mathbf{u}$ dipenda solo da θ e non da φ è ovvio. Infatti esso è perpendicolare a $\mathbf{w}$ e giace sul piano di base, quindi se φ varia il vettore $\mathbf{w}$ ruota intorno all'asse generato da $\mathbf{u}$: perciò $\mathbf{u}$ rimane invariante sotto questa rotazione, cioè non dipende da φ . Inoltre esso giace nel piano di base ed è perpendicolare a $\mathbf{w}$ la cui proiezione sul piano di base, data da $(\cos \theta \sin \varphi, 0)$, una volta normalizzata diventa $(\cos \theta, \sin \theta, 0)$. Quindi è chiaro che deve essere

$$\mathbf{u} = \pm(-\sin \theta, \cos \theta, 0). \quad (14.2.1)$$

Questo mostra quanto sia geometricamente più evidente ragionare con gli angoli di Eulero invece che con le coordinate cartesiane.)

A questo punto, il terzo vettore $\mathbf{v}$ della terna ortonormale destrorsa si ottiene, come prima, calcolando il prodotto vettore $\mathbf{w} \times \mathbf{u}$. Il risultato è

$$\mathbf{v} = \pm(\cos \theta \cos \varphi, \sin \theta \cos \varphi, -\sin \varphi).$$

dove, affinché la terna sia ortonormale destrorsa, il segno deve essere lo stesso che in (14.2.1). (**Nota:** anche questo risultato è geometricamente evidente, perché $\mathbf{v}$, essendo perpendicolare a $\mathbf{u}$ che è equatoriale, deve

giacere nello stesso piano verticale in cui giace $\mathbf{w}$, e quindi essere del tipo $\mathbf{v} = (\cos \theta \sin \varphi', \sin \theta \sin \varphi', \cos \varphi')$ per qualche angolo $\varphi' \in [0, \pi]$. Ma siccome $\mathbf{v}$ è anche perpendicolare a $\mathbf{w}$, che ha una inclinazione φ rispetto al piano equatoriale, la sua inclinazione rispetto al piano equatoriale deve essere $\varphi' = \varphi \pm \frac{\pi}{2}$, da cui il risultato.) $\square$

A questo punto la matrice della trasformazione affine dello spostamento della macchina da ripresa si ottiene immediatamente dal Corollario 14.1.8 (ii), ed è la seguente:

PROPOSIZIONE 14.2.4. *La trasformazione affine di $\mathbb{R}^3$ che trasla l'origine in $\mathbf{v}$ e manda la base canonica in una base ortonormale di versori applicati $\overrightarrow{\mathbf{v}\mathbf{u}}$, $\overrightarrow{\mathbf{v}\mathbf{v}}$, $\overrightarrow{\mathbf{v}\mathbf{w}}$, con primo versore parallelo al piano di base, agisce su $\mathbb{R}^3$ come la restrizione all'iperpiano $\{x_4 = 1\}$ in $\mathbb{R}^4$ della trasformazione lineare su $\mathbb{R}^4$ la cui matrice è*

$$A = \begin{pmatrix} -\frac{w_2}{\sqrt{w_1^2 + w_2^2}} - v_1 & -\frac{w_1 w_3}{\sqrt{w_1^2 + w_2^2}} - v_2 & w_1 - v_3 & v_1 \\ \frac{w_1}{\sqrt{w_1^2 + w_2^2}} - v_1 & -\frac{w_2 w_3}{\sqrt{w_1^2 + w_2^2}} - v_2 & w_2 - v_3 & v_2 \\ -v_1 & \sqrt{w_1^2 + w_2^2} - v_2 & w_3 - v_3 & v_3 \\ 0 & 0 & 0 & 1 \end{pmatrix}.$$

Se questa matrice viene espressa in termini di angoli di Eulero, come nel Corollario 14.2.3, essa diventa

$$A = \begin{pmatrix} \sin \theta - v_1 & \cos \theta \cos \varphi - v_2 & \cos \theta \sin \varphi - v_3 & v_1 \\ \cos \theta - v_1 & \sin \theta \cos \varphi - v_2 & \sin \theta \sin \varphi - v_3 & v_2 \\ -v_1 & -\sin \varphi - v_2 & \cos \varphi - v_3 & v_3 \\ 0 & 0 & 0 & 1 \end{pmatrix}.$$

CAPITOLO 15

* Quaternioni e matrici di rotazione

In questo capitolo introduciamo un nuovo spazio vettoriale munito di una operazione di prodotto (associativa e distributiva rispetto alla somma, però non commutativa), che generalizza il campo dei numeri complessi, ed utilizziamo la sua operazione di prodotto per rappresentare in maniera computazionalmente efficiente le matrici di rotazione su $\mathbb{R}^3$. Per prima cosa, descriviamo in modo appropriato l'azione di una rotazione in $\mathbb{R}^3$ intorno al proprio asse.

15.1. Espressione delle rotazioni in forma assiale

Consideriamo un operatore di rotazione R su $\mathbb{R}^3$. Ogni tale operatore ha un asse di rotazione: chiamiamo $\mathbf{n}$ il versore di tale asse (è tradizione orientarlo in maniera tale che la rotazione avvenga in senso antiorario intorno a $\mathbf{n}$ quando visto guardando verso l'origine a partire dal punto terminale del versore, ma qui non è essenziale). Vediamo come R opera su un generico vettore $\mathbf{r}$ scomponendo $\mathbf{r}$ in una componente assiale $\mathbf{r}_{||}$ ed una trasversale $\mathbf{r}_{\perp}$: ossia,

$$\mathbf{r}_{||} = \langle \mathbf{r}, \mathbf{n} \rangle \mathbf{n}$$

$$\mathbf{r}_{\perp} = \mathbf{r} - \langle \mathbf{r}, \mathbf{n} \rangle \mathbf{n}$$

(qui e nel resto del Capitolo usiamo la notazione $\langle \mathbf{r}, \mathbf{n} \rangle \mathbf{n}$ invece che $\mathbf{r} \cdot \mathbf{n} \mathbf{n}$ per motivi di leggibilità). L'operatore R lascia invariata la componente assiale $\mathbf{r}_{||}$,

$$R\mathbf{r}_{||} = \mathbf{r}_{||},$$

e ruota la componente trasversale $\mathbf{r}_{\perp}$ nel piano ortogonale a $\mathbf{n}$. Esprimiamo dunque $R\mathbf{r}_{\perp}$ nella base di questo piano formata da $\mathbf{r}_{\perp}$ e da un vettore in questo piano ad esso ortogonale, ossia il prodotto vettore

$$\mathbf{v} = \mathbf{n} \times \mathbf{r} = \mathbf{n} \times \mathbf{r}_{\perp}.$$

Osserviamo che $\|\mathbf{v}\| = \|\mathbf{r}\| = \|R\mathbf{r}_\perp\|$. Pertanto, se denotiamo con θ l'angolo di rotazione, abbiamo

$$R\mathbf{r} = \cos \theta \mathbf{r}_\perp + \sin \theta \mathbf{v}.$$

Pertanto

$$\begin{aligned} R\mathbf{r} &= R\mathbf{r}_\parallel + R\mathbf{r}_\perp = \mathbf{r}_\parallel + \cos \theta \mathbf{r}_\perp + \sin \theta \mathbf{v} \\ &= \langle \mathbf{r}, \mathbf{n} \rangle \mathbf{n} + \cos \theta (\mathbf{r} - \langle \mathbf{r}, \mathbf{n} \rangle \mathbf{n}) + \sin \theta \mathbf{v} \\ &= \cos \theta \mathbf{r} + (1 - \cos \theta) \langle \mathbf{r}, \mathbf{n} \rangle \mathbf{n} + \sin \theta \mathbf{n} \times \mathbf{r}. \end{aligned} \quad (15.1.1)$$

15.2. Rotazioni in $\mathbb{R}^2$, numeri complessi ed estensione a tre dimensioni

Ora sviluppiamo in forma intuitiva l'idea di estensione dei numeri complessi ad uno spazio quadridimensionale dotato di una moltiplicazione ed utile a rappresentare le rotazioni in $\mathbb{R}^3$, come accennato all'inizio del capitolo. Riconsideriamo anzitutto le rotazioni bidimensionali in termini di numeri complessi. La rotazione di un angolo θ intorno all'origine di $\mathbb{R}^2$ si identifica con l'angolo θ , o equivalentemente con il punto $e^{i\theta}$ del cerchio unitario complesso. La moltiplicatività dell'esponenziale, $e^{i(\theta_1+\theta_2)} = e^{i\theta_1} e^{i\theta_2}$, assicura che il prodotto di due di queste matrici di rotazione (ossia la loro composizione) si associa al prodotto dei due corrispondenti numeri complessi (in altre parole, che la mappa dalle matrici di rotazione ai numeri complessi di modulo uno è un omomorfismo moltiplicativo, e quindi un isomorfismo - una mappa biunivoca che conserva le operazioni di prodotto - del gruppo delle matrici di rotazione su $\mathbb{R}^2$ nel gruppo dei numeri complessi di modulo 1). Vogliamo fare la stessa cosa per le matrici di rotazione su $\mathbb{R}^3$. Questo gruppo, però, non è commutativo, e quindi dobbiamo costruire una estensione del campo complesso la cui operazione di prodotto non sia commutativa. Invece che aggiungere una dimensione immaginaria a $\mathbb{R}$ per formare i numeri complessi $x + iy$, ora dobbiamo aggiungere tre dimensioni immaginarie, ossia una copia di $\mathbb{R}^3$, per formare una somma diretta $\mathbb{H} = \mathbb{R} \oplus \mathbb{R}^3$ di uno spazio reale unidimensionale ed uno spazio *immaginario* a dimensione tre sul campo $\mathbb{R}$: quindi abbiamo bisogno di tre unità immaginarie, che chiamiamo i , j e k . È opportuno considerare

queste unità immaginarie come tre vettori: li denotiamo con $\mathbf{i}$, $\mathbf{j}$, $\mathbf{k}$, e denotiamo il versore dell'asse reale con $\mathbf{1}$. (Anche se qui il prodotto scalare non viene impiegato, la visualizzazione geometrica viene facilitata dal pensare questi quattro vettori come versori ortogonali in $\mathbb{R}^4$.) Come nel caso dei numeri complessi, ora dovremmo scrivere espressioni come $x + iy + jz + kw$, oppure, in forma vettoriale, $x + y\mathbf{i} + z\mathbf{j} + w\mathbf{k}$. Però troveremo opportuno far agire sullo spazio $\mathbb{H}$ matrici, ovviamente a dimensione quattro, e saremo interessati in quelle matrici che preservano il sottospazio tridimensionale immaginario (ad esempio, visualizzeremo in questa maniera le matrici di rotazione su $\mathbb{R}^3$). È conveniente considerare queste sottomatrici tridimensionali alla stregua del blocco tre per tre che rappresenta la componente lineare di una trasformazione affine a dimensione 4. Per questo preferiamo riordinare i quattro versori nell'ordine $\mathbf{i}$, $\mathbf{j}$, $\mathbf{k}$, $\mathbf{1}$: quindi scriveremo $ix + jy + kz + w$ oppure $x\mathbf{i} + y\mathbf{j} + z\mathbf{k} + w$. Ovviamente stiamo per definire una operazione di somma con la proprietà commutativa, quindi in queste espressioni l'ordine è inessenziale, ma quando consideriamo azioni di matrici vogliamo che la quarta coordinata sia quella reale (che identifica i multipli di $\mathbf{1}$).

15.3. Quaternioni

DEFINIZIONE 15.3.1. (*Quaternioni.*) Un quaternione $\mathbf{q}$ è una espressione del tipo

$$\mathbf{q} = (\mathbf{q}_v, q_w) = (q_x, q_y, q_z, q_w) = iq_x + jq_y + kq_z + q_w, \quad (15.3.1)$$

dove q_x , q_y , q_z e q_w sono numeri reali, $\mathbf{q}_v \in \mathbb{R}^3$ ed i simboli i , j , k soddisfano le proprietà seguenti:

$$i^2 = j^2 = k^2 = -1 \quad (15.3.2)$$

$$jk = -kj = i$$

$$ki = -ik = j$$

$$ij = -ji = k,$$

Equivalentemente, lo spazio $\mathbb{H}$ dei quaternioni può essere definito come lo spazio delle combinazioni lineari di quattro vettori $\mathbf{i}$, $\mathbf{j}$, $\mathbf{k}$, $\mathbf{1}$ muniti della tabella di moltiplicazione (15.3.2).

Il vettore $\mathbf{q}_v = (q_x, q_y, q_z)$ è chiamato parte immaginaria del quaternion $\mathbf{q}$, mentre q_w è detto parte reale.

Se estendiamo per linearità l'operazione di somma dalle singole componenti, lo spazio dei quaternioni diventa uno spazio vettoriale a dimensione 4 sul campo $\mathbb{R}$, che denotiamo con $\mathbb{H}$. Estendendo anche la moltiplicazione in maniera che rispetti la proprietà associativa e distributiva rispetto alla somma, muniamo lo spazio $\mathbb{H}$ di una operazione di moltiplicazione.

In base a questa definizione, l'operazione di moltiplicazione tra due quaternioni $\mathbf{q}$ e $\mathbf{r}$ soddisfa la tabella moltiplicativa che presentiamo in (15.3.3). Dalla tabella si vede che la moltiplicazione tra quaternioni non gode della proprietà commutativa.

15.3.1. Proprietà e definizioni.

Definizioni.

- **Moltiplicazione:** Denotando con $\times$ il prodotto vettoriale in $\mathbb{R}^3$ e con $\cdot$ il prodotto scalare euclideo, abbiamo

$$\begin{aligned} \mathbf{qr} &= (iq_x + jq_y + kq_z + q_w)(ir_x + jr_y + kr_z + r_w) \\ &= i(q_y r_z - q_z r_y + r_w q_x + q_w r_x) \\ &\quad + j(q_z r_x - q_x r_z + r_w q_y + q_w r_y) \\ &\quad + k(q_x r_y - q_y r_x + r_w q_z + q_w r_z) \\ &\quad + q_w r_w - q_x r_x - q_y r_y - q_z r_z \\ &= (\mathbf{q}_v \times \mathbf{r}_v + r_w \mathbf{q}_v + q_w \mathbf{r}_v, q_w r_w - \mathbf{q}_v \cdot \mathbf{r}_v). \end{aligned} \quad (15.3.3)$$

Definiamo ora la somma tra quaternioni, il coniugato, la norma e l'elemento identità.

- **Somma:**

$$\mathbf{q} + \mathbf{r} = (\mathbf{q}_v, q_w) + (\mathbf{r}_v, r_w) = (\mathbf{q}_v + \mathbf{r}_v, q_w + r_w). \quad (15.3.4)$$

- **Coniugato:**

$$\mathbf{q}^* = (\mathbf{q}_v, q_w)^* = (-\mathbf{q}_v, q_w). \quad (15.3.5)$$

- **Norma:**

$$\begin{aligned}\|\mathbf{q}\|^2 &= \mathbf{q}\mathbf{q}^* = \mathbf{q}^*\mathbf{q} = \mathbf{q}_v \cdot \mathbf{q}_v + q_w^2 \\ &= q_x^2 + q_y^2 + q_z^2 + q_w^2.\end{aligned}\quad (15.3.6)$$

Si osservi che questa definizione fornisce la norma sui quaternioni, la quale peraltro coincide con la norma sullo spazio euclideo $\mathbb{R}^4$ a cui i quaternioni sono isomorfi come spazio vettoriale.

- **Elemento identità (per moltiplicazione):**

$$\mathbf{i} = (\mathbf{0}, 1) \quad (15.3.7)$$

Per ogni elemento di norma non nulla esiste il corrispettivo elemento inverso rispetto all'operazione di moltiplicazione, ovvero il reciproco:

- **Reciproco:**

$$\mathbf{q}^{-1} = \frac{1}{n(\mathbf{q})}\mathbf{q}^* \quad (15.3.8)$$

Nella formula precedente per il reciproco compare la moltiplicazione per uno scalare, che è definita nel modo ovvio seguente:

- **Moltiplicazione per scalare:**

$$\begin{aligned}s\mathbf{q} &= (\mathbf{0}, s)(\mathbf{q}_v, q_w) = (s\mathbf{q}_v, sq_w) \\ \mathbf{q}s &= (\mathbf{q}_v, q_w)(\mathbf{0}, s) = (s\mathbf{q}_v, sq_w).\end{aligned}\quad (15.3.9)$$

Vediamo quindi che la moltiplicazione con scalare è commutativa, quindi:

$$s\mathbf{q} = \mathbf{q}s = (s\mathbf{q}_v, sq_w).$$

Proprietà. Dalle precedenti definizioni seguono le proprietà che ora elenchiamo.

- **Proprietà del coniugato:**

$$\begin{aligned}(\mathbf{q}^*)^* &= \mathbf{q} \\ (\mathbf{q} + \mathbf{r})^* &= \mathbf{q}^* + \mathbf{r}^* \\ (\mathbf{qr})^* &= \mathbf{r}^* \mathbf{q}^*.\end{aligned}\quad (15.3.10)$$

- *Proprietà della norma:*

$$\|\mathbf{q}^*\| = \|\mathbf{q}\| \quad (15.3.11a)$$

$$\|\mathbf{qr}\| = \|\mathbf{q}\| \|\mathbf{r}\|. \quad (15.3.11b)$$

- *Proprietà della moltiplicazione:*

- Bilinearità

$$\begin{aligned} \mathbf{p}(\mathbf{sq} + \mathbf{tr}) &= \mathbf{spq} + \mathbf{tpr} \\ (\mathbf{sp} + \mathbf{tq})\mathbf{r} &= \mathbf{spr} + \mathbf{tqr}. \end{aligned} \quad (15.3.12)$$

- Associatività

$$\mathbf{p}(\mathbf{qr}) = (\mathbf{pq})\mathbf{r}. \quad (15.3.13)$$

- *Quaternioni unitari*

Consideriamo un quaternione $\mathbf{q} = (\mathbf{q}_v, q_w)$ *unitario*, ossia tale che $\|\mathbf{q}\| = 1$. Grazie a (15.3.6) un quaternione è unitario se e solo se può essere scritto come

$$\mathbf{q} = (\mathbf{u}_q \sin \omega, \cos \omega) \quad (15.3.14)$$

dove $\mathbf{u}_q$ è un versore in $\mathbb{R}^3$: $\|\mathbf{u}_q\|^2 = 1$.

Infatti, se $\mathbf{q}$ è del tipo in (15.3.14),

$$\begin{aligned} \|\mathbf{q}\|^2 &= \|(\mathbf{u}_q \sin \omega, \cos \omega)\|^2 = (\mathbf{u}_q \cdot \mathbf{u}_q) \sin^2 \omega + \cos^2 \omega \\ &= \sin^2 \omega + \cos^2 \omega = 1, \end{aligned}$$

e viceversa, se $\mathbf{u}_q$ è unitario, allora (15.3.14) segue dall'ultima espressione in (15.3.11) per la norma.

15.4. Rotazioni in $\mathbb{R}^3$ e coniugazione di quaternioni

Poiché la moltiplicazione di quaternioni è non commutativa, l'operazione di coniugazione,

$$\Omega_{\mathbf{q}}\mathbf{p} = \mathbf{qpq}^{-1},$$

è non banale. Inoltre, grazie alla moltiplicatività della norma (15.3.11b), se $\mathbf{q}$ è un quaternione unitario l'operatore lineare $\Omega_{\mathbf{q}}$ preserva la norma, e quindi è un operatore ortogonale sullo spazio dei quaternioni $\mathbb{H}$ pensato come spazio vettoriale reale a dimensione 4. Si noti che, se $\mathbf{q}$

è un quaternione unitario, $\mathbf{q}^{-1}$ coincide con il coniugato $\mathbf{q}$ in base a (15.3.8), e quindi $\Omega_{\mathbf{q}} = \mathbf{q}\mathbf{p}\mathbf{q}^*$.

TEOREMA 15.4.1. (Coniugazione di quaternioni e rotazioni in $\mathbb{R}^3$.) *Sia $\mathbf{p}$ un punto in uno spazio tridimensionale pensato come una classe di equivalenza per dilatazione in $\mathbb{R}^4$ grazie alle sue coordinate omogenee (p_x, p_y, p_z, p_w) . Identifichiamo $\mathbf{p}$ con il quaternione*

$$\mathbf{p} = ((p_x, p_y, p_z), p_w) = (\mathbf{p}_v, p_w).$$

Sia $\mathbf{q}$ un quaternione non nullo. Allora:

- (i) *L'operatore di coniugazione $\Omega_{\mathbf{q}}\mathbf{p} = \mathbf{q}\mathbf{p}\mathbf{q}^{-1}$ trasforma $\mathbf{p} = (\mathbf{p}_v, p_w)$ in un quaternione $\mathbf{p}' = (\mathbf{p}'_v, p_w) = (p'_x, p'_y, p'_z, p_w)$, tale che $\|\mathbf{p}_v\| = \|\mathbf{p}'_v\|$, e quindi, ristretto al sottospazio tridimensionale dei quaternioni immaginari puri $(\mathbf{p}_v, 0)$, è un operatore unitario (ossia di rotazione).*
- (ii) *Qualsiasi multiplo reale, diverso da zero, di $\mathbf{q}$ effettua la stessa trasformazione, e quindi la rotazione di punti tridimensionali, vista come coniugazione rispetto ad un opportuno quaternione, è compatibile con la rappresentazione di quei punti come classe di equivalenza di vettori in $\mathbb{R}^4$, ossia non dipende dalla scelta dei rappresentanti della classe di equivalenza.*
- (iii) *Se il quaternione $\mathbf{q}$ è unitario e come in (15.3.14) lo scriviamo $\mathbf{q} = (\mathbf{u}_{\mathbf{q}} \sin \omega, \cos \omega)$, allora $\Omega_{\mathbf{q}}$, ristretto al sottospazio tridimensionale dei quaternioni immaginari puri, è l'operatore di rotazione antioraria di angolo 2ω attorno all'asse $\mathbf{u}_{\mathbf{q}}$. Qui la rotazione si intende antioraria quando vista da un osservatore orientato come il vettore $\mathbf{u}_{\mathbf{q}}$ che guarda verso l'origine.*

DIMOSTRAZIONE. Osserviamo anzitutto che la parte (ii) è banale, in quanto l'inversa di $s\mathbf{q}$ è $\mathbf{q}^{-1}s^{-1}$, e la moltiplicazione con scalare gode della proprietà commutativa. Perciò $(s\mathbf{q})\mathbf{p}(s\mathbf{q})^{-1} = s\mathbf{q}\mathbf{p}\mathbf{q}^{-1}s^{-1} = \mathbf{q}\mathbf{p}\mathbf{q}^{-1}s s^{-1} = \mathbf{q}\mathbf{p}\mathbf{q}^{-1}$. Possiamo quindi assumere $\mathbf{q}$ come un quaternione unitario senza perdita di generalità. Per un quaternione unitario $\mathbf{q}$, $\mathbf{q}^{-1} = \mathbf{q}^*$; possiamo quindi scrivere $\mathbf{q}\mathbf{p}\mathbf{q}^{-1}$ come $\mathbf{q}\mathbf{p}\mathbf{q}^*$.

Dimostriamo la parte (i). La parte reale di un qualsiasi quaternione, $\text{Re } \mathbf{q}$, può essere estratta usando la formula $2 \text{Re } \mathbf{q} = \mathbf{q} + \mathbf{q}^*$. Consideriamo $2 \text{Re}(\mathbf{qpq}^*) = \mathbf{qpq}^* + (\mathbf{qpq}^*)^* = \mathbf{qpq}^* + \mathbf{qp}^* \mathbf{q}^*$. Visto che la moltiplicazione tra quaternioni è bilineare, possiamo scrivere l'ultimo membro come $\mathbf{q}(\mathbf{p} + \mathbf{q}^*)\mathbf{q}^* = 2\mathbf{q} \text{Re } \mathbf{pq}^* = 2 \text{Re } \mathbf{p}$. Quindi $\mathbf{q}$ coniuga $\mathbf{p} = (\mathbf{p}_v, p_w)$ in $\Omega_{\mathbf{q}}\mathbf{p} = \mathbf{p}' = (\mathbf{p}'_v, p_w)$, preservando la parte reale di $\mathbf{p}$. Inoltre l'operazione di moltiplicazione mantiene la norma perché la norma è moltiplicativa (si veda (15.3.11b)) e $\mathbf{q}$ è unitario: quindi $\|\mathbf{p}\| = \|\mathbf{p}'\|$. Infine, poiché p_w resta inalterata, $\|\mathbf{p}_v\| = \|\mathbf{p}'_v\|$.

Per ultimo dimostriamo (iii). Abbiamo visto che si può scegliere $\mathbf{q} = (\mathbf{q}_v, q_w)$ quaternione *unitario*. Sia $\mathbf{p}$ un quaternione *immaginario puro*. Dall'ultima identità della regola di moltiplicazione (15.3.3) si verifica facilmente che

$$\Omega_{\mathbf{q}}\mathbf{p} = ((q_w^2 - \mathbf{q}_v \cdot \mathbf{q}_v)\mathbf{p}_v + 2(\mathbf{q}_v \cdot \mathbf{p}_v)\mathbf{q}_v + 2q_w(\mathbf{q}_v \times \mathbf{p}_v), 0),$$

e poiché $\mathbf{q} = \mathbf{u}_{\mathbf{q}} \sin \omega, \cos \omega$ abbiamo

$$\begin{aligned} \Omega_{\mathbf{q}}\mathbf{p} &= ((\cos^2 \theta - \sin^2 \theta)\mathbf{p}_v + 2 \sin^2 \theta (\mathbf{p}_v \cdot \mathbf{u}_{\mathbf{q}})\mathbf{u}_{\mathbf{q}} + 2 \cos \theta \sin \theta \mathbf{u}_{\mathbf{q}} \times \mathbf{p}_v, 0) \\ &= (\cos 2\theta \mathbf{p}_v + (1 - \cos 2\theta)(\mathbf{p}_v \cdot \mathbf{u}_{\mathbf{q}})\mathbf{u}_{\mathbf{q}} + \sin 2\theta \mathbf{u}_{\mathbf{q}} \times \mathbf{p}_v, 0). \end{aligned}$$

□

ESEMPIO 15.4.2. Calcoliamo il quaternione unitario associato alla rotazione determinata di un angolo ω in senso antiorario rispetto all'asse x , o y , o z .

Dalla parte (iii) del Teorema 15.4.1 si trova che, per ciascun versore canonico di base $\mathbf{e}_1, \mathbf{e}_2, \mathbf{e}_3$, il quaternione è $\mathbf{q}_i = (\sin \omega \mathbf{e}_i, \cos \omega)$.

15.4.1. Composizione di rotazioni e prodotto di quaternioni. Vediamo come la composizione di rotazioni si associa al prodotto di quaternioni. Consideriamo due quaternioni unitari, $\mathbf{q}_1$ e $\mathbf{q}_2$, ed un vettore $\mathbf{p} \in \mathbb{R}^3$ trasformato in un quaternione mediante le sue coordinate omogenee, $\mathbf{p} = (p_x, p_y, p_z, p_w)$. In base al Teorema 15.4.1, l'operazione

$$\mathbf{q}_2(\mathbf{q}_1 \mathbf{p} \mathbf{q}_1^*) \mathbf{q}_2^*$$

ristretta al sottospazio tridimensionale dei quaternioni puramente immaginari coincide la composizione delle rotazioni associate a $\mathbf{q}_1$ ed

a $\mathbf{q}_2$, in questo ordine. Ponendo $\mathbf{q} = \mathbf{q}_2\mathbf{q}_1$ la formula precedente può essere scritta come

$$\mathbf{q}\mathbf{p}\mathbf{q}^*.$$

Riassumendo,

COROLLARIO 15.4.3. *La mappa Ω è un omomorfismo dal gruppo moltiplicativo $\mathbb{H}$ al gruppo delle matrici ortogonali reali su $\mathbb{R}^4$:*

15.4.2. Matrice di rotazione in termini di quaternioni. Ora troviamo la forma matriciale dell'operatore $\Omega_{\mathbf{q}}$ su $\mathbb{R}^4 \sim \mathbb{H}$ dato dalla coniugazione con il quaternione $\mathbf{q}$. Ritorniamo, per maggiore generalità, al caso di un quaternione $\mathbf{q}$ non necessariamente unitario.

Poiché la moltiplicazione fra quaternioni è bilineare, possiamo esprimere questa operazione sui quaternioni (pensati come vettori in $\mathbb{R}^4$) tramite matrici a dimensione 4, spezzandola nella moltiplicazione sulla sinistra, $\mathbf{q}\mathbf{p}$, e la moltiplicazione sulla destra, $\mathbf{p}\mathbf{q}^*$.

Scriviamo $\mathbf{L}^q\mathbf{p}$ la moltiplicazione sulla sinistra, $\mathbf{p} \mapsto \mathbf{q}\mathbf{p}$, con $\mathbf{q} = (q_x, q_y, q_z, q_w) = (\mathbf{q}_v, q_w)$. Segue immediatamente dalla regola di moltiplicazione (15.3.3) che la matrice associata all'operatore lineare $\mathbf{L}^q$ è

$$\mathbf{L}^q = \begin{pmatrix} q_w & -q_z & q_y & q_x \\ q_z & q_w & -q_x & q_y \\ -q_y & q_x & q_w & q_z \\ -q_x & -q_y & -q_z & q_w \end{pmatrix}$$

Scriviamo ora la moltiplicazione sulla destra, $\mathbf{p} \mapsto \mathbf{p}\mathbf{q}^*$ come $\mathbf{R}^{q^*}\mathbf{p}$, dove l'operatore $\mathbf{R}^{q^*}$ viene espresso, grazie alla regola di moltiplicazione, come

$$\mathbf{R}^{q^*} = \begin{pmatrix} q_w & -q_z & q_y & -q_x \\ q_z & q_w & -q_x & -q_y \\ -q_y & q_x & q_w & -q_z \\ q_x & q_y & q_z & q_w \end{pmatrix}$$

Ora un calcolo elementare ma tedioso mostra che la matrice $\mathbf{M}^{\mathbf{q}}$ associata all'operatore di coniugazione $\Omega_{\mathbf{q}}$ è

$$\begin{aligned} \mathbf{M}^{\mathbf{q}} &= \mathbf{L}^q \mathbf{R}^{q*} \\ &= \begin{pmatrix} q_w^2 + q_x^2 - q_y^2 - q_z^2 & 2(q_x q_y - q_w q_z) & 2(q_x q_z + q_w q_y) & 0 \\ 2(q_x q_y + q_w q_z) & q_w^2 - q_x^2 + q_y^2 - q_z^2 & 2(q_y q_z - q_w q_x) & 0 \\ 2(q_x q_z - q_w q_y) & 2(q_y q_z + q_w q_x) & q_w^2 - q_x^2 - q_y^2 + q_z^2 & 0 \\ 0 & 0 & 0 & q_w^2 + q_x^2 + q_y^2 + q_z^2 \end{pmatrix}. \end{aligned}$$

Semplificando otteniamo

$$\mathbf{M}^{\mathbf{q}} = \begin{pmatrix} \|\mathbf{q}\|^2 - 2(q_y^2 + q_z^2) & 2(q_x q_y - q_w q_z) & 2(q_x q_z + q_w q_y) & 0 \\ 2(q_x q_y + q_w q_z) & \|\mathbf{q}\|^2 - 2(q_x^2 + q_z^2) & 2(q_y q_z - q_w q_x) & 0 \\ 2(q_x q_z - q_w q_y) & 2(q_y q_z + q_w q_x) & \|\mathbf{q}\|^2 - 2(q_x^2 + q_y^2) & 0 \\ 0 & 0 & 0 & \|\mathbf{q}\|^2 \end{pmatrix}$$

Nel caso il quaternion $\mathbf{q}$ sia unitario si ha una ulteriore semplificazione:

$$\mathbf{M}^{\mathbf{q}} = \begin{pmatrix} 1 - 2(q_y^2 + q_z^2) & 2(q_x q_y - q_w q_z) & 2(q_x q_z + q_w q_y) & 0 \\ 2(q_x q_y + q_w q_z) & 1 - 2(q_x^2 + q_z^2) & 2(q_y q_z - q_w q_x) & 0 \\ 2(q_x q_z - q_w q_y) & 2(q_y q_z + q_w q_x) & 1 - 2(q_x^2 + q_y^2) & 0 \\ 0 & 0 & 0 & 1 \end{pmatrix} \quad (15.4.1)$$

Questo dimostra la prima parte del seguente risultato:

COROLLARIO 15.4.4. (Matrici di rotazione espresse in termini di quaternioni.) Consideriamo la matrice di rotazione (ossia ortogonale reale con determinante +1) $\mathbf{M}^{\mathbf{q}}$ dell'operatore di coniugazione $\Omega_{\mathbf{q}}$ dove $\mathbf{q}$ è un quaternion unitario: essa ha la forma espressa in (15.4.1). Viceversa, ogni matrice di rotazione su $\mathbb{R}^3$, espressa nel modo seguente in forma di matrice affine a dimensione quattro,

$$\mathbf{M}^{\mathbf{q}} = \begin{pmatrix} m_{00} & m_{01} & m_{02} & 0 \\ m_{10} & m_{11} & m_{12} & 0 \\ m_{20} & m_{21} & m_{22} & 0 \\ 0 & 0 & 0 & 1 \end{pmatrix},$$

è associata al quaternione unitario $\mathbf{q} = (q_x, q_y, q_z, q_w)$ dato da

$$\begin{aligned} q_w &= \pm \frac{1}{2} \sqrt{m_{00} + m_{11} + m_{22} + 1} = \pm \frac{1}{2} \sqrt{\text{tr}(\mathbf{M}^{\mathbf{q}})} \\ q_x &= \frac{m_{21} - m_{12}}{4q_w} \\ q_y &= \frac{m_{02} - m_{20}}{4q_w} \\ q_z &= \frac{m_{10} - m_{01}}{4q_w}. \end{aligned} \quad (15.4.2)$$

DIMOSTRAZIONE. Dobbiamo solo dimostrare la seconda parte dell'enunciato, ossia la ricostruzione del quaternione a partire dalla matrice. Questo significa invertire la prima parte dell'enunciato, ossia ricavare il quaternione dall'espressione (15.4.1). Da questa espressione si vede subito che la traccia (ossia la somma dei coefficienti diagonali) della matrice è $4(1 - q_x^2 - q_y^2 - q_z^2)$. Poiché $\mathbf{q}$ è un quaternione unitario, $q_x^2 + q_y^2 + q_z^2 + q_w^2 = 1$ e quindi

$$\text{tr}(\mathbf{M}^{\mathbf{q}}) = 4q_w^2.$$

Questo prova la prima identità in (15.4.2). Esaminando ancora (15.4.1) si vede che $m_{21} = 2(q_w q_x + q_y q_z)$ e $m_{12} = 2(q_y q_z - q_w q_x)$, da cui segue la seconda identità in (15.4.2). Le restanti due identità si provano allo stesso modo. $\square$

Esprimere le matrici di rotazione in termini di quaternioni è una notevole opportunità di ridurre la mole di calcoli in Computer Graphics, dove le rotazioni intervengono al cambiare della posizione dell'osservatore (come accade continuamente durante le animazioni). Eseguire il calcolo tramite moltiplicazione di quaternioni è numericamente vantaggioso, ma soprattutto questa procedura è facile da implementare in hardware, dal momento che la moltiplicazione di due quaternioni è lineare nelle coordinate di entrambi: l'implementazione in hardware permette di eseguirla in maniera velocissima. Pertanto può essere utile svolgere i seguenti esercizi.

ESERCIZIO 15.4.5. Rivedere il precedente esempio 15.4.2 e ritrovarne il risultato mediante il Corollario 15.4.4.

ESERCIZIO 15.4.6. Calcolare il quaternion unitario associato alla rotazione determinata dagli angoli di Eulero θ e ϕ introdotti nella Definizione 14.2.2.

ESERCIZIO 15.4.7. Calcolare il quaternion associato alla componente di rotazione della matrice affine di rototraslazione dello spostamento della macchina da ripresa, calcolata nel Corollario 14.2.3.

Parte 3

Matematica della prospettiva

CAPITOLO 16

* Trasformazioni prospettiche

Questo capitolo presenta vari tipi di trasformazioni prospettiche (dette anche *assonometriche*). Queste trasformazioni si suddividono in due categorie: *proiezione centrale*, se ci sono fasci di rette parallele che dopo la proiezione convergono verso opportuni punti di fuga, oppure *proiezioni parallele*, se questo non succede. Nel caso delle proiezioni parallele, i fasci di rette parallele rimangono paralleli, e la proiezione è determinata dalla scelta di un piano di proiezione e dalla direzione di proiezione. Se tale direzione è perpendicolare al piano di proiezione si dice che la proiezione è *ortogonale* (o *ortografica*), altrimenti *obliqua*. In generale il piano di proiezione non contiene l'origine, e quindi le proiezioni prospettiche non fissano l'origine e pertanto non sono operazioni lineari. Vedremo che la proiezione centrale si rappresenta con matrici proiettive, le altre con matrici affini. La proiezione centrale è quella più videorealistica, utilizzata in Computer Graphics, e la studiamo per prima (un caso particolarmente semplice nella Sezione 16.1, il caso generale nelle Sezioni 16.4 e 16.5). Rispetto alle proiezioni parallele, trattate sistematicamente nella Sezione 16.6, la proiezione ortogonale è anticipata alla Sezione 16.2, in modo da poter fornire una versione unificata della forma matriciale delle proiezioni centrale ed ortogonale nella Sezione 16.3. Sarebbe possibile estendere questa trattazione unificata anche alle altre proiezioni parallele, ma lasciamo questo tedioso compito al lettore.

Nel calcolare le matrici delle trasformazioni prospettiche dobbiamo mettere in guardia il lettore che la Computer Graphics ha una tradizione assai peculiare, quella di far agire le matrici sui vettori non da sinistra ma da destra. La ragione storica di ciò è che i primordi della Computer Graphics furono sviluppati non da matematici, bensì da studiosi a cui pareva strano che se due operatori A e B agiscono

su un vettore $\mathbf{p}$ *in questo ordine* allora si debba avere che la composizione dei due uno dopo l'altro si debba scrivere $B\mathbf{A}\mathbf{p} = B(\mathbf{A}\mathbf{p})$ invece che $AB\mathbf{p}$. Per rovesciare l'ordine con cui i due simboli si succedono sulla carta, questi studiosi preferirono scrivere l'azione da destra, in modo che la composizione diventasse $\mathbf{p}AB = (\mathbf{p}A)B$. Per nostra fortuna a questo punto il lettore avrà studiato la matematica e compreso la ragione dell'ordine naturale, e gli sarà facile capire gli articoli di Computer Graphics, nei quali purtroppo l'ordine è opposto.

16.1. Prospettiva centrale, proiezione standard

Riprendendo in esame l'Esempio 13.4.4, osserviamo anzitutto, a scanso di malintesi, che anche quando i rappresentanti delle classi in $\mathbb{P}^n$ si scelgono del tipo $(x_1, x_2, \dots, x_{n-1}, 1)$, la forma della matrice in $PGL_n(\mathbb{R})$ associata ad una trasformazione prospettica non ha necessariamente l'ultima riga $(0, \dots, 0, 1)$, come invece avviene per le matrici delle trasformazioni affini (Proposizione 14.1.4). Infatti questa sarebbe la forma giusta per l'azione sui punti al finito se la trasformazione mandasse l'iperpiano $\{x_n = 1\}$ in sé, come appunto avviene per le trasformazioni affini di $\mathbb{R}^{n-1}$ quando considerate come trasformazioni lineari in $\mathbb{R}^{n-1}$ che lasciano invariante tale iperpiano: ma ciò non avviene per le trasformazioni proiettive, per le quali l'azione sui vettori non è neppure definita, visto che i rappresentanti delle classi dell'equivalenza proiettiva possono essere dilatati (e quindi uscire dal suddetto piano) senza che l'azione ne risenta. Abbiamo visto un esempio concreto alla fine dell'Esempio 13.4.4, ed esattamente in (13.4.2), che riprendiamo in esame nella prossima Sezione.

Come è consuetudine in Computer Graphics, chiamiamo x e y le coordinate orizzontale e verticale, e z la profondità, orientata in modo da aumentare dall'origine verso l'osservatore. Il piano di visuale è parallelo agli assi x e y e quindi perpendicolare all'asse z , diciamo in posizione $z = d$ (qui stiamo supponendo che l'osservatore non si trovi verticalmente sopra l'origine, cioè sul piano x - y : in genere, in Computer Graphics, l'osservatore è poco al di sopra dell'asse z , nel senso che la sua posizione ha z positivo grande, y non molto grande e positivo, e $x = 0$). Si osservi che non stiamo richiedendo che il punto $\mathbf{p}$ abbia

valore positivo di z : anzi, se invece che modellare un osservatore che guarda una scena stessimo modellando una camera oscura come un cubo con un piccolo foro nell'origine disposto nel semispazio $z \leq 0$, il piano di visuale, cioè in questo caso il piano della pellicola, passerebbe attraverso l'interno del cubo, e quindi avrebbe d negativo: in tal caso l'immagine creata dalla proiezione sarebbe ribaltata rispetto alla scena reale, perché i raggi si incrociano nel passare tutti attraverso l'origine. Con questa scelta di coordinate, la proiezione sul piano di visuale porta gli assi x e y della scena tridimensionale su rispettivi assi orizzontale e verticale in tale piano: quindi i nomi delle coordinate sono quelli naturali per la compatibilità, perché se immaginiamo che questo sia il piano del monitor, è naturale chiamare questi assi del monitor x e y , rispettivamente.

Fissiamo ora in due modi diversi i parametri della proiezione prospettica della prospettiva centrale. Il primo modo è quello più naturale se si pensa di aver fissato una volta per tutte la posizione d del piano di visuale. In tal caso, collochiamo per ora l'osservatore nell'origine: cioè, l'origine è il *centro di proiezione*. Tratteremo nella prossima Sezione ?? il caso generale in cui il centro di proiezione è generico (non necessariamente l'origine).

La proiezione manda il generico punto $\mathbf{q} = (x, y, z)$ sul punto determinato sul piano di visuale dall'intersezione con la retta che passa per l'origine e per il punto $\mathbf{q}$. Riconosciamo immediatamente in questo modo di procedere l'implementazione del principio della proiezione stereografica nella geometria proiettiva (Sezione 13.2).

La proiezione avviene quindi tramite una similitudine, cioè una proporzione: il punto $\mathbf{q}$ viene mandato in

$$\mathbf{p} = (x_p, y_p, d) = \left(\frac{x}{z/d}, \frac{y}{z/d}, \frac{z}{z/d} \right) \equiv \left(\frac{x}{z/d}, \frac{y}{z/d}, d \right).$$

Possiamo riformulare questo risultato in termini di trasformazioni prospettiche: la trasformazione prospettica si scrive in termini di una (classe di equivalenza di) matrice $M_d^{(c)}$ in $PM_4(\mathbb{R})$ come in (13.4.1) dell'Esempio 13.4.4. Il punto $\mathbf{p}$ corrisponde ad un punto proiettivo al finito $[x, y, z, 1]$, e si ha

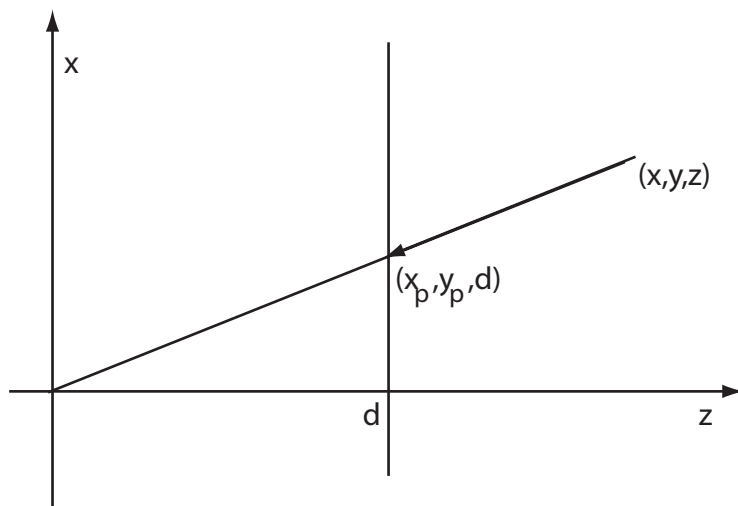


FIGURA 1. Proiezione della prospettiva centrale: piano di visuale in $z = d$

$$\begin{bmatrix} X \\ Y \\ Z \\ W \end{bmatrix} = M_d^{(c)} \begin{bmatrix} x \\ y \\ z \\ 1 \end{bmatrix} = \begin{pmatrix} 1 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 \\ 0 & 0 & 1 & 0 \\ 0 & 0 & 1/d & 0 \end{pmatrix} \begin{bmatrix} x \\ y \\ z \\ 1 \end{bmatrix}. \quad (16.1.1)$$

Quindi $[X, Y, Z, W] = [x, y, z, \frac{z}{d}]$, e la quarta coordinata omogenea vale $W = z/d$. Perciò come rappresentante della classe immagine

$$[X, Y, Z, W]$$

possiamo scegliere il consueto rappresentante stereografico

$$\left(\frac{X}{W}, \frac{Y}{W}, \frac{Z}{W}, 1 \right),$$

che corrisponde in $\mathbb{R}^3$ al punto $(x_p, y_p, z_p) = \left(\frac{x}{z/d}, \frac{y}{z/d}, d \right)$.

Ora veniamo al secondo modo utile di fissare i parametri prospettici. Poniamo il centro di prospettiva non più nell'origine, bensì in $(0, 0, -d)$, ed il piano di visuale in $z = 0$ (è consuetudine in Computer Graphics, per semplificare il processo di trasformazione da coordinate tridimensionali nel piano di visuale in $\mathbb{R}^3$ a coordinate bidimensionali del monitor, collocare il piano di visuale in $\{z = 0\}$). In tal modo, quando

facciamo crescere la distanza d), il piano di visuale non si sposta, e possiamo più agevolmente confrontare i risultati della trasformazione prospettica).

Lo stesso argomento di proporzionalità adottato prima ora porta alle seguenti equazioni:

$$\begin{aligned}\frac{x_p}{d} &= \frac{x}{z+d} \\ \frac{y_p}{d} &= \frac{y}{z+d}\end{aligned}\tag{16.1.2}$$

$$z_p = 0.\tag{16.1.3}$$

da cui si ricava

$$x_p = \frac{x}{1 + \frac{z}{d}}$$

$$y_p = \frac{y}{1 + \frac{z}{d}}$$

$$z_p = 0.$$

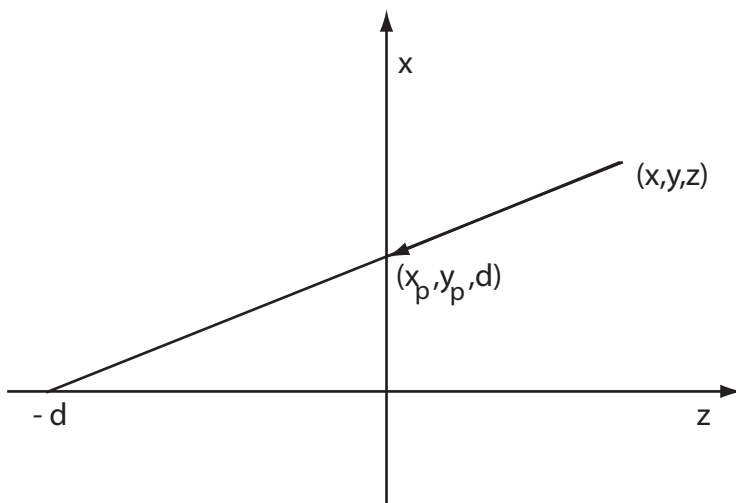


FIGURA 2. Proiezione della prospettiva centrale: piano di visuale in $z = 0$

Pertanto ora la “matrice” $M_d^{(c)}$ della trasformazione prospettica è la classe di equivalenza, cioè l’elemento in $PM_4(\mathbb{R})$, dato da

$$M_d^{(c)} = \begin{pmatrix} 1 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 \\ 0 & 0 & 0 & 0 \\ 0 & 0 & 1/d & 1 \end{pmatrix}. \quad (16.1.4)$$

Poiché gli elementi di $PM_4(\mathbb{R})$ sono classi di equivalenza per dilatazione, possiamo scegliere un altro rappresentante della stessa classe, dilatando quello appena scritto in modo da eliminare il denominatore. Così si ottiene la seguente forma della matrice della trasformazione prospettica:

$$M_d^{(c)} = \begin{pmatrix} d & 0 & 0 & 0 \\ 0 & d & 0 & 0 \\ 0 & 0 & 0 & 0 \\ 0 & 0 & 1 & d \end{pmatrix}. \quad (16.1.5)$$

Si noti che questa forma corrisponde nel modo più naturale alle equazioni (16.1.2) della trasformazione.

ESERCIZIO 16.1.1. Consideriamo la prospettiva centrale con centro di proiezione ubicato sull’asse z al punto $z = -d$ e piano di proiezione $\{z = 0\}$, la quale porta alla matrice di proiezione prospettica (16.1.4). Sia S una superficie emisferica di raggio 1 nel semispazio $\{x \geq 0\}$ con centro nel punto $(0, 0, 2)$, e C un cilindro di raggio $\frac{1}{2}$ avente per asse centrale la retta $\{y = 0, z = 2\}$. Sia $J = S \cap C$. Calcolare l’immagine prospettica di J sul piano di proiezione.

SVOLGIMENTO. Anzitutto determiniamo J . L’equazione di S è $x^2 + y^2 + (z - 2)^2 = 1, x \geq 0$. L’equazione di C è $y^2 + (z - 2)^2 = \frac{1}{4}$. Quindi i punti di J sono tutti e soli quelli che soddisfano le seguenti equazioni e disequazioni:

$$\begin{aligned} x &\geq 0 \\ x^2 + y^2 + (z - 2)^2 &= 1 \\ y^2 + (z - 2)^2 &= \frac{1}{4}. \end{aligned}$$

Se ne ricava

$$x^2 = \frac{3}{4}$$

$$x \geq 0$$

$$y^2 + (z - 2)^2 = \frac{1}{4}$$

ossia $x = \sqrt{3}/2$, $y^2 + (z - 2)^2 = \frac{1}{4}$. Si tratta, ovviamente, di una circonferenza sul piano $x = \sqrt{3}/2$, che parametrizziamo nel modo seguente:

$$x = \sqrt{3}/2 \quad (16.1.6)$$

$$y = \frac{1}{2} \cos t \quad (16.1.7)$$

$$z = 1 + \frac{\sin t}{2} \quad (16.1.8)$$

dove l'angolo t varia fra 0 e 2π . Scriviamo i punti di J in termini di coordinate omogenee (x, y, z, w) con quarta coordinata $w = 1$ (scelta del rappresentante stereografico per punti al finito), ed applichiamo la matrice di proiezione prospettica (16.1.4). Analogamente a (16.1.1), da (16.1.4) e 16.1.6 si ottiene

$$\begin{aligned} \begin{bmatrix} X \\ Y \\ Z \\ W \end{bmatrix} &= M_d^{(c)} \begin{bmatrix} x \\ y \\ z \\ 1 \end{bmatrix} = \begin{pmatrix} 1 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 \\ 0 & 0 & 0 & 0 \\ 0 & 0 & 1/d & 1 \end{pmatrix} \begin{bmatrix} \sqrt{3}/2 \\ \frac{1}{2} \cos t \\ 1 + \frac{1}{2} \sin t \\ 1 \end{bmatrix} \\ &= \begin{bmatrix} \sqrt{3}/2 \\ \frac{1}{2} \cos t \\ 0 \\ 1 + \frac{1 + \frac{1}{2} \sin t}{d} \end{bmatrix}. \quad (16.1.9) \end{aligned}$$

Ora riportiamo in coordinate omogenee con quarta componente 1 il punto proiettivo (X, Y, Z, W) così calcolato, dividendo per il valore di W : si ottiene così un altro rappresentante della classe di equivalenza

proiettiva $[X, Y, Z, W]$, e precisamente il punto $(x, y, z, 1)$ dove

$$\begin{aligned} x &= \frac{d}{d+1+\frac{1}{2}\sin t} \frac{\sqrt{3}}{2} \\ y &= \frac{d \cos t}{2(d+1+\frac{1}{2}\sin t)} \\ z &= 0. \end{aligned}$$

Quindi la curva immagine di J sul piano di proiezione $\{z=0\}$, munito delle coordinate x e y , ha la seguente equazione:

$$\begin{aligned} x &= \frac{\sqrt{3}d}{2d+2+\sin t} \\ y &= \frac{d \cos t}{2d+2+\sin t}. \end{aligned}$$

Si osservi che, se facciamo tendere d a $-\infty$ (ossia se passiamo alla proiezione ortogonale sul piano $\{z=0\}$, si ottiene, come previsto, $x = \sqrt{3}/2$, $y = \frac{1}{2}\cos t$: al variare di t questo punto immagine si muove sul segmento $x = \sqrt{3}/2$, $-1/2 \leq y \leq 1/2$, che è esattamente la proiezione ortogonale dell'intersezione fra sfera e cilindro, perché essa è una circonferenza nello spazio tridimensionale che giace sul piano $x = \sqrt{3}/2$, ha componente y del centro uguale a 0 e raggio $1/2$.

□

NOTA 16.1.2. (Punto di fuga della proiezione standard.) È interessante osservare che, in questo caso particolare della proiezione prospettica centrale, che si chiama la *proiezione standard*, il piano di visuale è ortogonale alla direzione dall'osservatore all'origine. Consideriamo un fascio di rette perpendicolare al piano di visuale: nella proiezione standard nella forma appena sviluppata, si tratta del fascio delle rette parallele all'asse z . Chiamiamo $\mathbf{r}_0 = (x_0, y_0, 0)$ il punto in cui una tale retta $\mathbf{r}$ interseca il piano $\{z=0\}$: allora la equazione parametrica di $\mathbf{r}$ è $\mathbf{r}(t) = \mathbf{r}_0 + t\mathbf{e}_3$. Poniamo $\mathbf{r}(t) = (x(t), y(t), z(t))$ e scriviamo i punti della retta $\mathbf{r}$ come i rappresentanti speciali delle loro classi di equivalenza proiettiva, ossia, in coordinate omogenee, $\mathbf{R}(t) = (x(t), y(t), z(t), 1)$; analogamente, scriviamo $\mathbf{S}$ per la corrispondente estensione quadridimensionale di ogni altra retta $\mathbf{s}$.

Quando applichiamo la trasformazione trovata in (16.1.5), la retta $\mathbf{R}$ viene trasformata nella retta $\mathbf{S} = T\mathbf{R}$ di equazioni parametriche

$$\mathbf{S}(t) = M_d^{(c)} \mathbf{R}(t) = \begin{pmatrix} d & 0 & 0 & 0 \\ 0 & d & 0 & 0 \\ 0 & 0 & 0 & 0 \\ 0 & 0 & 1 & d \end{pmatrix} \begin{pmatrix} a \\ b \\ t \\ 1 \end{pmatrix} = \begin{pmatrix} da \\ db \\ 0 \\ t + d \end{pmatrix}.$$

Riconduciamo il risultato alla consueta espressione stereografica delle classi di equivalenza proiettive, ossia con i rappresentanti speciali di ultima coordinata 1, rinormalizzando con la divisione per $t + d$. Si ottiene

$$\mathbf{S}(t) = \left(\frac{a}{t+d}, \frac{b}{t+d}, 0, 1 \right). \quad (16.1.10)$$

La retta immagine tridimensionale $\mathbf{s}$, che si ottiene considerando le prime tre componenti dopo questa normalizzazione, è

$$\mathbf{s}(t) = \left(\frac{a}{t+d}, \frac{b}{t+d}, 0 \right). \quad (16.1.11)$$

Osserviamo che essa giace, come deve essere, nel piano $\{z = 0\}$ (il piano di visuale), e quando $t \rightarrow \infty$ tende all'origine, quali che siano a e b . In altre parole, il fascio di rette parallele perpendicolare al piano di visuale viene trasformato nel fascio delle semirette radiali nel piano $\{z = 0\}$. L'origine è quindi il *punto di fuga* prospettico di questo fascio di rette.

Osserviamo che l'origine è esattamente il punto ottenuto sommando al centro di proiezione il versore direzionale del fascio di rette moltiplicato per la distanza d fra il centro di proiezione ed il piano di visuale. Se avessimo scelto un centro di proiezione diverso dall'origine, lo può ripetere lo stesso calcolo della trasformazione prospettica, che svolgeremo in dettaglio nella prossima Sezione 16.4. A partire da esso, riprenderemo in esame e generalizzeremo questo risultato nella Sezione 16.5. $\square$

ESEMPIO 16.1.3. (Trasformazione prospettica standard del cubo unitario.) Consideriamo il cubo Q i cui vertici sono $\{\pm \mathbf{e}_1 \pm \mathbf{e}_2 \pm \mathbf{e}_3\}$.

È improprio chiamare unitario questo cubo, perché ha lato 2, ma possiamo sempre dividere per due alla fine se proprio vogliamo, e quindi procediamo con questa scelta, che evita i denominatori. Consideriamo le classi di equivalenza proiettiva dei vettori $\mathbf{e}_i$, scrivendo ad esempio i loro rappresentanti standard come $\pm \mathbf{E}_1^\pm = (\pm 1, 0, 0, 1)$, $\mathbf{E}_2^\pm = (0, \pm 1, 0, 1)$, $\mathbf{E}_3^\pm = (0, 0, \pm 1, 1)$. Applichiamo agli $\mathbf{E}_i$ la matrice $M_d^{(c)}$ in (16.1.5) e scriviamo $\mathbf{V}_i^\pm = M_d^{(c)} \mathbf{E}_i^\pm$. Si ottiene $\mathbf{V}_i^\pm = dmbE_i^\pm$ per $i = 1$ e 2 , e $\mathbf{V}_3^\pm = (0, 0, 0, d \pm 1)$.

Rinormalizzando per ritornare ai rappresentanti proiettivi standard con 1 alla quarta componente, e considerando solo le prime tre componenti del risultato per ritrovare i vettori tridimensionali trasformati $\mathbf{v}_i$, e scrivendo M per la trasformazione prospettica tridimensionale in tal modo ottenuta da $M_d^{(c)}$, si verifica subito che si ha

$$\begin{aligned} T(1, 1, 1) &= \frac{(1, 1, 0)}{d+1} \\ T(1, 1, -1) &= \frac{(1, 1, 0)}{d-1} \\ T(1, -1, 1) &= \frac{(1, -1, 0)}{d+1} \\ T(1, -1, -1) &= \frac{(1, -1, 0)}{d-1} \\ T(-1, 1, 1) &= \frac{(-1, 1, 0)}{d+1} \\ T(-1, 1, -1) &= \frac{(-1, 1, 0)}{d-1} \\ T(-1, -1, 1) &= \frac{(-1, -1, 0)}{d+1} \\ T(-1, -1, -1) &= \frac{(-1, -1, 0)}{d-1}. \end{aligned}$$

Consideriamo allora i quattro vertici anteriori visti dall'osservatore (ovvero i punti $(\pm 1, \pm 1, -1)$: non dimentichiamo che l'osservatore è in $(0, 0, -d)$, sull'asse z *negativo*). Abbiamo appena visto che essi vengono compressi di un fattore $1/(d-1)$, mentre i quattro vertici posteriori di un fattore più grande, $1/(d+1)$. Questa è la compressione prospettica. Qui abbiamo supposto $d > 1$. Se $0 < d < 1$ allora $d-1 < 0$

ed i vertici anteriori sono in realtà alle spalle dell'osservatore, e vengono quindi scambiati di segno, ossia ribaltati, come succede ai raggi che passano per una lente quando la scena e l'osservatore sono ai lati opposti dell'obiettivo (centro di proiezione). Se poi $d = 1$, allora i vertici anteriori sono ai lati dell'osservatore, e la proiezione prospettica non è definita su di essi (le semirette da essi al centro di proiezione non passano per il piano di visuale $\{z = 0\}$, restano tutte nel piano $\{z = -1\}$). Osserviamo come è fatta la proiezione di una retta che attraversa il piano di visuale, riconsiderando la retta immagine $\mathbf{s}$ del fascio ortogonale al piano di visuale, $\mathbf{r}(t) = (a, b, t)$, che è stata calcolata in (16.1.11): $\mathbf{s}(t) = (\frac{a}{t+d}, \frac{b}{t+d}, 0)$. Facciamo variare il parametro t in modo che il punto $\mathbf{r}(t)$ passi dal semispazio anteriore all'osservatore posto nel centro di prospettiva $(0, 0, -d)$ a quello posteriore, ossia da $t > -d$ a $t < -d$. Quando t decresce verso $-d$ il punto $\mathbf{s}(t)$ si muove nel piano di visuale radialmente fuori dall'origine (ossia il punto di fuga) verso l'infinito: più precisamente, verso il punto all'infinito di questo piano le cui coordinate proiettive sono $(a, b, 0)$. Quando t diventa inferiore a $-d$ e continua a decrescere, il punto $\mathbf{s}(t)$ salta dalla parte opposta, sulla stessa retta radiale ma sulla semiretta opposta, e si avvicina dall'infinito al punto di fuga. Quindi una retta che attraversa il piano di visuale ha immagine che va all'infinito con un salto. Questo fatto crea una distorsione prospettica assai drastica: si provi ad immaginare l'immagine prospettica del cubo unitario quando l'osservatore si trova dentro il cubo! Per quato motivo, in Computer Graphics, il centro di prospettiva, ossia l'osservatore, si trova sempre da un lato del piano di visuale e la scena osservata dall'altro lato. $\square$

16.2. Proiezione prospettica ortogonale (o ortografica)

La proiezione ortogonale (detta anche ortografica) è la trasformazione prospettica introdotta nella Proposizione 6.6.5 (v), con la differenza che ora proiettiamo sul piano di visuale, che potrebbe essere un piano che non contiene l'origine, quindi non un sottospazio vettoriale ma un suo traslato (ed in tal caso l'origine non può essere preservata dalla

proiezione, che pertanto non è una applicazione lineare). Però manteniamo la scelta di localizzazione del piano di visuale fatta alla fine della precedente Sezione 16.4, e quindi il piano di visuale passa per l'origine: è il piano $\{z = 0\}$, quindi un sottospazio. La proiezione ortogonale all'asse z è quindi la trasformazione che manda il punto $\mathbf{p} = (x, y, z)$ nel punto di tale piano ottenuto ponendo uguale a zero la componente z , cioè $(x, y, 0)$. In coordinate omogenee stiamo ponendo

$$\begin{aligned}x_0 &= x \\y_0 &= y \\z_0 &= 0 \\w_0 &= w\end{aligned}$$

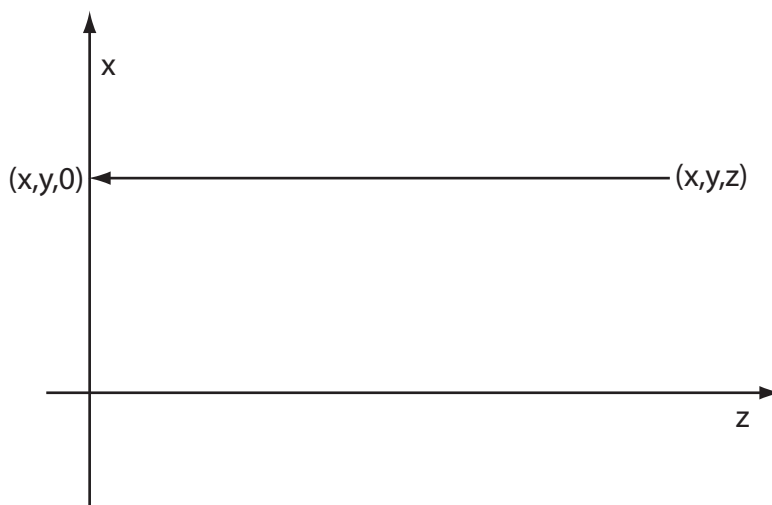


FIGURA 3. Proiezione della prospettiva ortogonale sul piano $z = 0$

Pertanto la “matrice” prospettica ora è

$$M^{(o)} = \begin{pmatrix} 1 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 \\ 0 & 0 & 0 & 1 \\ 0 & 0 & 0 & 1 \end{pmatrix}. \quad (16.2.1)$$

NOTA 16.2.1. Si osservi che la matrice $M^{(o)}$ si ottiene facendo tendere d ad infinito nell'espressione per $M_d^{(c)}$. In effetti questo fatto è naturale, perchè i raggi prospettici in questo tipo di prospettiva sono tutti paralleli all'asse di proiezione (nel nostro caso l'asse z , e quindi non convergono verso un centro di proiezione (punto di fuga) al finito, bensì verso un punto di fuga all'infinito (la direzione proiettiva dell'asse z , cioè, in coordinate omogenee, $[0, 0, 1, 0]$). Questo equivale a dire che la distanza d fra il centro di proiezione e l'origine tende ad infinito.

□

16.3. Un'unica matrice per prospettiva centrale e ortogonale

Nelle due trasformazioni prospettiche viste nelle Sezioni precedenti, la matrice della prospettiva centrale $M_d^{(c)}$ è quella che si applica quando il centro di proiezione è a distanza $d < \infty$ dal piano di visuale, mentre la matrice della proiezione ortogonale M^o si applica quando il centro di proiezione è all'infinito. È possibile unificare questi due casi nella seguente formulazione più generale di una trasformazione prospettica con un unico punto di fuga (al finito o all'infinito).

Sia $\mathbf{p} = (x, y, z)$ un punto generico da trasformare prospetticamente, sia $\mathbf{c}$ il centro di proiezione, ubicato in un punto arbitrario dello spazio, e collochiamo il piano di visuale in posizione $z = z_p$ (questa è la situazione abituale della Computer Graphics, nella quale il piano di visuale è ortogonale all'asse z). Consideriamo il punto $\mathbf{b}_p = (0, 0, z_p)$ dove questo piano interseca l'asse z , e sia $q = \|\mathbf{c} - \mathbf{b}_p\|$ e $\mathbf{d}$ il versore $\mathbf{d} = \frac{1}{q}(\mathbf{c} - \mathbf{b}_p)$. Come prima, denotiamo con $\mathbf{p}_p = (x_p, y_p, z_p)$ il punto sul piano di visuale ottenuto proiettando prospetticamente (con centro in $\mathbf{c}$) il generico punto $\mathbf{x} = (x, y, z)$. Allora $\mathbf{p}_p$ appartiene alla retta di equazione parametrica

$$\mathbf{r}(t) = (x'(t), y'(t), z'(t)) = \mathbf{c} + t(\mathbf{x} - \mathbf{c}) \quad (16.3.1)$$

per qualche $0 \leq t \leq 1$.

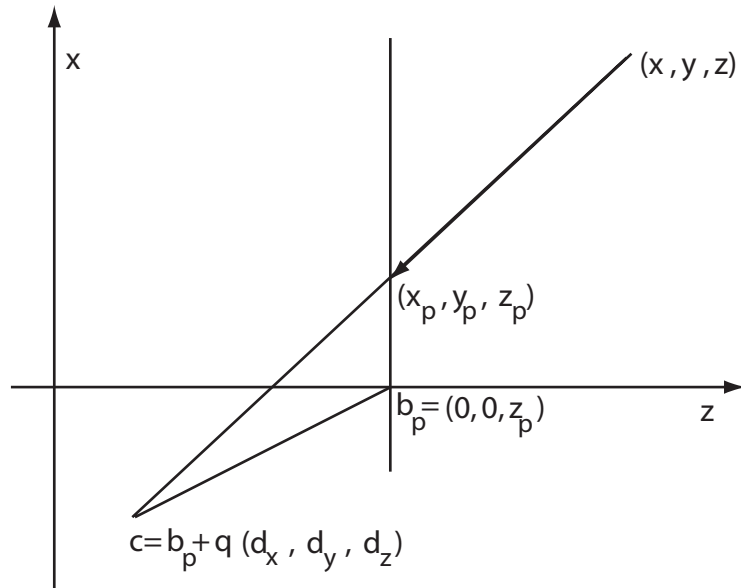


FIGURA 4. La proiezione prospettica: caso generale di punto di fuga al finito

Usando il fatto che $\mathbf{c} = \mathbf{b}_p + q\mathbf{d}$ ricaviamo da (16.3.1) che i punti del segmento parametrico verificano

$$x'(t) = qd_x + (x - qd_z)t$$

$$y'(t) = qd_y + (y - qd_z)t$$

$$z'(t) = z_p + qd_z + (z - (z_p + qd_z))t$$

Ora ricaviamo $\mathbf{p}_p$ ponendo $z' = z_p$ nell'ultima uguaglianza. In tal modo si ottiene il valore appropriato di t :

$$t = \frac{z_p - (z_p + qd_z)}{z - (z_p + qd_z)},$$

che sostituito nelle uguaglianze precedenti dà

$$\begin{aligned} x_p &= \frac{x - z \frac{d_x}{d_z} + z_p \frac{d_x}{d_z}}{1 + \frac{z_p - z}{q d_z}} \\ y_p &= \frac{y - z \frac{d_y}{d_z} + z_p \frac{d_y}{d_z}}{1 + \frac{z_p - z}{q d_z}} \\ z_p &= z_p \frac{1 + \frac{z_p - z}{q d_z}}{1 + \frac{z_p - z}{q d_z}} = \frac{-z \frac{z_p}{q d_z} + \frac{z_p(1 + q d_z)}{q d_z}}{1 + \frac{z_p - z}{q d_z}}. \end{aligned}$$

Queste uguaglianze equivalgono alla forma seguente della trasformazione prospettica $M \in PM_4(\mathbb{R})$:

$$\begin{bmatrix} x_p \\ y_p \\ z_p \\ 1 \end{bmatrix} = M \begin{bmatrix} x \\ y \\ z \\ 1 \end{bmatrix},$$

con

$$M = \begin{pmatrix} 1 & 0 & -\frac{d_x}{d_z} & z_p \frac{d_x}{d_z} \\ 0 & 1 & -\frac{d_y}{d_z} & z_p \frac{d_y}{d_z} \\ 0 & 0 & -\frac{z_p}{q d_z} & z_p + \frac{z_p^2}{q d_z} \\ 0 & 0 & -\frac{1}{q d_z} & 1 + \frac{z_p}{q d_z} \end{pmatrix}.$$

Si osservi che, scegliendo $\mathbf{d} = (0, 0, -1)$, otteniamo

- $M = M_d^{(c)}$ per $z_p = d$ e $q = d$ (Sezione 16.4),
- $M = M_d^{(c)}$ per $z_p = 0$ e $q = d$ (Sezione 16.4),
- $M = M^{(o)}$ per $z_p = 0$ e $q = \infty$ (Sezione 16.2),.

16.4. Forma matriciale generale della prospettiva centrale

Nella precedente Sezione 16.3 abbiamo visto la forma della matrice prospettica quando il centro di prospettiva è arbitrario ma, come di consuetudine in Computer Graphics, il piano di visuale è ortogonale all'asse z . Anche se questa è la consuetudine, in una animazione nella quale l'osservatore si sposta può essere necessario spostare anche il piano di visuale al fine di vedere parti interessanti della scena. Si pensi ad esempio ad un videogioco nel quale l'osservatore (ossia il centro di

proiezione) si muove in un labirinto, o anche in un appartamento dove attraversa una porta, ed ha bisogno di guardare a sinistra e a destra per vedere se ci sono nemici: mentre attraversa la porta, i suoi lati sinistro e destro giacciono nel piano di visuale o nelle sue immediate vicinanze, e quindi (Esempio 16.1.3) non sono visibili (sono mandati all'infinito dalla proiezione) oppure sono estremamente distorti. È come se l'osservatore guardasse con la coda dell'occhio: per vedere bene deve girare la testa, o, nel nostro caso, il piano di visuale. In questa Sezione deriviamo la forma generale della trasformazione prospettica centrale nella quale anche il piano di visuale è scelto arbitrariamente.

Cominciamo con un caso particolare che viola il nostro precedente proposito (legato al confronto con la prospettiva ortogonale) di collocare il centro di prospettiva in $(0, 0, -d)$, ed invece lo ricolloca nell'origine: ma siccome alla fine dovremo traslarlo ad un punto arbitrario, questa scelta ormai non comporta svantaggi, ed anzi rende più semplice effettuare in seguito la traslazione.

PROPOSIZIONE 16.4.1. (Prospettiva centrale con centro nell'origine e piano di proiezione arbitrario.) *La trasformazione prospettica con centro nell'origine e piano di visuale P a distanza d_0 dall'origine e con versore normale $\mathbf{n} = (n_x, n_y, n_z)$ è espressa, in coordinate proiettive (ovvero omogenee), dalla matrice*

$$M_0 = \begin{pmatrix} d_0 & 0 & 0 & 0 \\ 0 & d_0 & 0 & 0 \\ 0 & 0 & d_0 & 0 \\ n_x & n_y & n_z & 0 \end{pmatrix}.$$

DIMOSTRAZIONE. Poiché il centro di proiezione è l'origine, la trasformazione manda il punto generico $\mathbf{p} = (x, y, z)$ in

$$\mathbf{p}' = \alpha \mathbf{p} \in P. \quad (16.4.1)$$

Osserviamo che la distanza del piano P dall'origine è, a parte il segno, $d_0 = \mathbf{p}' \cdot \mathbf{n}$. Ora, il fattore di compressione α è il rapporto fra questa distanza e la distanza fra $\mathbf{p}$ e l'origine proiettata lungo la direzione normale al piano P : ossia, $\alpha = d_0 / \mathbf{p} \cdot \mathbf{n} = d_0 / (n_x x + n_y y + n_z z)$. La divisione per $n_x x + n_y y + n_z z$ non è altro che la divisione prospettica

per la distanza (in questo caso la distanza dall'origine proiettata lungo la direzione normale al piano P) se, in coordinate omogenee, si scrive la matrice M_0 come nell'enunciato. $\square$

Ora trattiamo il caso generale in cui anche il centro di prospettiva è arbitrario. A questo scopo è utile il seguente Lemma, che mette in rilievo la mancanza di linearità delle trasformazioni proiettive pensate come trasformazioni dello spazio tridimensionale visto in coordinate omogenee quadridimensionali.

NOTA 16.4.2. Siano $\mathbf{p} = (p_x, p_y, p_z) \in \mathbb{R}^3$, $\mathbf{q} = (q_x, q_y, q_z)$, $\mathbf{r} = (r_x, r_y, r_z)$ e si consideri la classe di equivalenza proiettiva $\bar{\mathbf{p}}$ di $(p_x, p_y, p_z, 1)$, che d'ora in poi chiameremo la *classe di equivalenza di $\mathbf{p}$* . Se due matrici M_1 e M_2 a quattro dimensioni mandano, rispettivamente $(p_x, p_y, p_z, 1)$ in (q_x, q_y, q_z, a) e (r_x, r_y, r_z, b) , con $a \neq -b$, allora la matrice $M_1 + M_2$ manda $(p_x, p_y, p_z, 1)$ in $(q_x + r_x, q_y + r_y, q_z + r_z, a + b)$, e quindi le corrispondenti trasformazioni proiettive mandano la classe $\bar{\mathbf{p}}$ nella classe di $(\mathbf{q} + \mathbf{r})/(a + b)$. In particolare, se le due matrici M_1 e M_2 hanno la stessa ultima riga, e quindi $a = b$, allora la matrice proiettiva la cui azione sulle classi di equivalenza manda la classe di $(p_x, p_y, p_z, 1)$ nella somma delle classi di $(q_x, q_y, q_z, 1)$ e $(r_x, r_y, r_z, 1)$ è la matrice con la stessa ultima riga di M_1 e M_2 e con le prime tre righe date dalla somma termine a termine di M_1 e M_2 . $\square$

NOTA 16.4.3. (*Il significato geometrico del fattore di distanza.*) Abbiamo ripetutamente incontrato, ad esempio nella Proposizione 16.4.1, il termine $d_0 = (\mathbf{p}', \mathbf{n})$, dove $\mathbf{p}'$ è un generico punto del piano di proiezione P . Chiaramente, $|d_0|$ è la distanza fra P e l'origine, ed il segno di d_0 è positivo se il versore normale $\mathbf{n}$ di P punta nel semispazio opposto all'origine. Nella suddetta Proposizione, l'origine è il centro di proiezione, e quindi la consuetudine della Computer Graphics è che la scena giaccia nel semispazio opposto, per evitare drastiche distorsioni prospettiche, come osservato alla fine dell'Esempio 16.1.3. Quindi in questo caso d_0 è positivo e misura la distanza dall'origine al piano. Nei prossimi enunciati sposteremo il centro di proiezione dall'origine ad un punto arbitrario $\mathbf{c}$ ed introdurremo la costante $d_1 = (\mathbf{c}, \mathbf{n})$.

Chiaramente, d_1 è la distanza fra i piani paralleli a P che passano per l'origine e per $\mathbf{c}$. Chiamiamo questi due piani paralleli P_0 e P_c , rispettivamente. Il segno di d_1 è positivo se il centro di proiezione $\mathbf{c}$ giace in quello dei due semispazi determinati da P_0 verso cui punta il versore normale $\mathbf{n}$ (qualche volta si parla del *semispazio positivo* determinato dalla scelta del versore normale).

Infine, introdurremo la quantità $d = d_0 - d_1$, che ha un ruolo significativo nella matrice della prospettiva centrale nel caso generale. Il significato geometrico di questa costante è comprensibile se consideriamo due casi. Il primo caso è quello in cui il piano P ed il centro $\mathbf{c}$ giacciono in semispazi opposti rispetto a P_0 . In tal caso i segni di d_0 e d_1 sono opposti, e quindi d , a parte il segno, misura la somma delle distanze fra P e P_0 e fra P_0 e P_c , ossia la distanza fra P ed il centro di proiezione $\mathbf{c}$ (esattamente come succedeva per d_0 nel caso della Proposizione 16.4.1. Nel secondo caso, i segni di d_0 e d_1 sono uguali, e quindi d misura la differenza delle suddette distanze. Però ora il piano P_c giace fra P e P_0 , e quindi d misura nuovamente, a parte il segno, la distanza fra P ed il centro di proiezione $\mathbf{c}$. Il segno di d è positivo o negativo a seconda della scelta di verso del versore $\mathbf{n}$, e più precisamente è positivo se il centro di proiezione (ossia l'osservatore) sta nel semispazio *negativo* determinato dal piano di visuale P . $\square$

TEOREMA 16.4.4. (Prospettiva centrale con centro e piano di proiezione arbitrari.) *Consideriamo la trasformazione prospettica con centro in un punto $\mathbf{c} = (c_x, c_y, c_z)$ e piano di visuale P con versore normale $\mathbf{n} = (n_x, n_y, n_z)$ a distanza con segno d_0 dall'origine (ossia $d_0 = \langle \mathbf{p}_0, \mathbf{n} \rangle$ per qualsiasi punto $\mathbf{p}_0 \in P$). Sia d_1 la distanza con segno fra i piani paralleli a P che passano per l'origine e per $\mathbf{c}$, rispettivamente (ossia $d_1 = \langle \mathbf{c}, \mathbf{n} \rangle$), e sia $d = d_0 - d_1$. Allora la trasformazione è espressa, in coordinate proiettive, dalla matrice $M_{\mathbf{c}}$, che in questo Teorema scriviamo brevemente M , data da*

$$M = M_{\mathbf{c}} = \begin{pmatrix} d + c_x n_x & c_x n_y & c_x n_z & -c_x d_0 \\ c_y n_x & d + c_y n_y & c_y n_z & -c_y d_0 \\ c_z n_x & c_z n_y & d + c_z n_z & -c_z d_0 \\ n_x & n_y & n_z & -d_x \end{pmatrix}.$$

DIMOSTRAZIONE. Potremmo svolgere la dimostrazione riportando il centro di proiezione all'origine mediante una traslazione, applicando allora la trasformazione prospettica la cui matrice è stata determinata nella Proposizione 16.4.1 ed infine ritornando indietro con la traslazione opposta. Questi calcoli li svolgeremo nel successivo Esercizio 16.4.5 qui invece preferiamo sviluppare il calcolo diretto in analogia alla dimostrazione della Proposizione 16.4.1.

Poiché ora la trasformazione prospettica proietta radialmente verso $\mathbf{c}$ il punto generico $\mathbf{p} = (x, y, z)$ in un punto $\mathbf{p}'$, l'identità (16.4.1) diventa

$$\mathbf{p}' - \mathbf{c} = \alpha(\mathbf{p} - \mathbf{c}). \quad (16.4.2)$$

Pertanto $\langle \mathbf{n}, \mathbf{p}' - \mathbf{c} \rangle = \alpha \langle \mathbf{n}, \mathbf{p} - \mathbf{c} \rangle$, da cui

$$\alpha = \frac{\langle \mathbf{n}, \mathbf{p}' \rangle - \langle \mathbf{n}, \mathbf{c} \rangle}{\langle \mathbf{n}, \mathbf{p} \rangle - \langle \mathbf{n}, \mathbf{c} \rangle} = \frac{d_0 - d_1}{\langle \mathbf{n}, \mathbf{p} \rangle - d_1} = \frac{d}{\langle \mathbf{n}, \mathbf{p} \rangle - d_1}, \quad (16.4.3)$$

perché, per definizione, $d_0 = \langle \mathbf{n}, \mathbf{p}' \rangle$, $d_1 = \langle \mathbf{n}, \mathbf{c} \rangle$ e $d = d_0 - d_1$.

Riscriviamo la trasformazione prospettica come $(p'_x, p'_y, p'_z) = T(x, y, z)$ con $p_x = \alpha x + (1 - \alpha)c_x$, e analogamente per p_y e p_z . Allora $T = T_1 + T_2$, dove $T_1(x, y, z) = \alpha(x, y, z)$ e $T_2(x, y, z) = (1 - \alpha)(c_x, c_y, c_z)$. Consideriamo le trasformazioni proiettive $\tilde{T}_1$ e $\tilde{T}_2$ indotte da T_1 e T_2 pensate in coordinate omogenee. Scriviamo le matrici (a dimensione 4) M_1 e M_2 associate alle trasformazioni proiettive $\tilde{T}_1$ e $\tilde{T}_2$. Chiaramente la prima è

$$M_1 = \begin{pmatrix} d & 0 & 0 & 0 \\ 0 & d & 0 & 0 \\ 0 & 0 & d & 0 \\ n_x & n_y & n_z & -d_x \end{pmatrix}.$$

Per scrivere M_2 osserviamo da (16.4.3) che

$$1 - \alpha = 1 - \frac{d}{\langle \mathbf{n}, \mathbf{p} \rangle - d_1} = \frac{\langle \mathbf{n}, \mathbf{p} \rangle - d_0}{\langle \mathbf{n}, \mathbf{p} \rangle - d_1}.$$

Pertanto $\tilde{T}_2$ diventa

$$\begin{aligned}\tilde{T}_2(x, y, z, 1) &= ((1 - \alpha)c_x, (1 - \alpha)c_y, (1 - \alpha)c_z, 1) \\ &\sim ((\langle \mathbf{n}, \mathbf{p} \rangle - d_0)c_x, (\langle \mathbf{n}, \mathbf{p} \rangle - d_0)c_y, (\langle \mathbf{n}, \mathbf{p} \rangle - d_0)c_z, \langle \mathbf{n}, \mathbf{p} \rangle - d_1) \\ &= ((n_x x + n_y y + n_z z)c_x - d_0 c_x, (n_x x + n_y y + n_z z)c_y - d_0 c_y, \\ &\quad (n_x x + n_y y + n_z z)c_z - d_0 c_z, n_x x + n_y y + n_z z - d_1),\end{aligned}$$

e quindi

$$M_2 = \begin{pmatrix} c_x n_x & c_x n_y & c_x n_z & -c_x d_0 \\ c_y n_x & c_y n_y & c_y n_z & -c_y d_0 \\ c_z n_x & c_z n_y & c_z n_z & -c_z d_0 \\ n_x & n_y & n_z & -d_1 \end{pmatrix}.$$

L'enunciato ora segue dalla Nota 16.4.2. $\square$

ESERCIZIO 16.4.5. Come annunciato all'inizio della dimostrazione del Teorema 16.4.4, ora diamo una formulazione alternativa di quella dimostrazione riportando all'origine il centro di proiezione $\mathbf{c}$ mediante la traslazione $\Lambda(x, y, z) = (x - c_x, y - c_y, z - c_z)$, poi applicando la trasformazione prospettica con centro l'origine della Proposizione 16.4.1 ed infine ritornando indietro con la traslazione opposta Λ^{-1} .

SVOLGIMENTO. La matrice della trasformazione affine associata alla traslazione Λ è

$$\Lambda = \begin{pmatrix} 1 & 0 & 0 & -c_x \\ 0 & 1 & 0 & -c_y \\ 0 & 0 & 1 & -c_z \\ 0 & 0 & 0 & 1 \end{pmatrix}$$

e la sua inversa ovviamente è la traslazione opposta,

$$\Lambda^{-1} = \begin{pmatrix} 1 & 0 & 0 & c_x \\ 0 & 1 & 0 & c_y \\ 0 & 0 & 1 & c_z \\ 0 & 0 & 0 & 1 \end{pmatrix}.$$

Per scrivere la matrice M_0 abbiamo bisogno della distanza con segno fra il piano di proiezione e l'origine *dopo aver applicato la traslazione* Λ . Nella Proposizione 16.4.1 questa distanza, chiaramente, era $d_0 =$

$\langle \mathbf{n}, \mathbf{p}' \rangle$, dove $\mathbf{p}'$ è un punto qualsiasi del piano di proiezione. La traslazione manda $\mathbf{p}'$ in $\mathbf{p}' - \mathbf{c}$, e quindi d_0 in

$$\langle \mathbf{n}, \mathbf{p}' - \mathbf{c} \rangle = \langle \mathbf{n}, \mathbf{p}' \rangle - \langle \mathbf{n}, \mathbf{c} \rangle = d_0 - d_1 = d$$

con la notazione del precedente Teorema 16.4.4. Pertanto, la matrice della Proposizione 16.4.1 ora diventa

$$M_0 = \begin{pmatrix} d & 0 & 0 & 0 \\ 0 & d & 0 & 0 \\ 0 & 0 & d & 0 \\ n_x & n_y & n_z & 0 \end{pmatrix}.$$

È ora immediato verificare che la composizione di matrici $\Lambda^{-1} M_0 \Lambda$ (si badi all'ordine!) è esattamente la matrice M_c del Teorema 16.4.4.

□

16.5. Punti di fuga della prospettiva centrale

Abbiamo introdotto il concetto di punto di fuga per la proiezione standard nella Nota 16.1.2. Rammentiamo che si tratta del punto in cui si incontrano le immagini sotto la trasformazione prospettica di un fascio di rette parallele, in quel caso ortogonali al piano di visuale, e quindi orientate secondo il versore normale dell'asse z . Ora, grazie al Teorema 16.4.4, possiamo estendere quel risultato al caso della prospettiva centrale generale, per un fascio di rette parallele diretto secondo un versore arbitrario $\mathbf{u} = (u_x, u_y, u_z)$, non necessariamente ortogonale al piano di visuale P . Come prima, sia $\mathbf{c} = (c_x, c_y, c_z)$ il centro di proiezione, e $\mathbf{n}$ il versore normale a P . Per qualsiasi punto $\mathbf{p}_0 = (x_0, y_0, z_0) \in P$, la distanza (con segno) di P dall'origine è $d_0 = \langle \mathbf{n}, \mathbf{p}_0 \rangle$, e la distanza (con segno) fra i piani paralleli a P che passano, rispettivamente, per l'origine e per $\mathbf{c}$ è $d_1 = \langle \mathbf{n}, \mathbf{c} \rangle$. Poniamo inoltre $d = d_0 - d_1$: quindi d è la distanza (con segno) fra il piano P ed il piano ad esso parallelo che passa per $\mathbf{c}$: il segno dipende dal verso scelto per il versore normale $\mathbf{n}$, e più precisamente si vede subito che è positivo se e solo se il versore $\mathbf{n}$ punta verso il semipiano determinato da P opposto a quello che contiene $\mathbf{c}$.

TEOREMA 16.5.1. *Adottiamo la notazione appena stabilita, e chiamiamo $\mathbf{r}$ le rette parallele dirette secondo il versore $\mathbf{u}$: $\mathbf{r}(t) = \mathbf{r}^0 + t\mathbf{u}$, con $\mathbf{r}^0$ arbitrario. Se il fascio non è diretto parallelamente al piano di visuale, ossia se $\langle \mathbf{n}, \mathbf{u} \rangle \neq 0$, allora la proiezione della prospettiva centrale generale del Teorema 16.4.4 manda questo fascio di rette in un insieme di rette $\mathbf{s}$ radiali rispetto al (ossia che si incontrano nel) punto di fuga*

$$\mathbf{f}_{\mathbf{u}} = \mathbf{c} + \frac{d}{\langle \mathbf{n}, \mathbf{u} \rangle} \mathbf{u}.$$

Altrimenti, se $\langle \mathbf{n}, \mathbf{u} \rangle = 0$, il fascio di rette rimane parallelo anche dopo la trasformazione prospettica, e non c'è alcun punto di incontro (per queste direzioni diciamo che il punto di fuga è all'infinito). In tal caso, tale punto all'infinito corrisponde al punto di coordinate omogenee $(u_x, u_y, u_z, 0)$ nello spazio proiettivo.

DIMOSTRAZIONE. Scriviamo $\mathbf{r}(t) = (x(t), y(t), z(t))$, e la retta immagine sotto la trasformazione prospettica come $\mathbf{s}(t) = (x'(t), y'(t), z'(t))$. Sia T la trasformazione della prospettiva centrale con piano di proiezione P e centro di proiezione $\mathbf{c}$: essa agisce in coordinate omogenee tramite la matrice $M_{\mathbf{c}}$ del Teorema 16.4.4. Scriviamo rispettivamente $\mathbf{R}(t)$, $\mathbf{S}(t)$ le classi di equivalenza proiettiva $[\mathbf{r}(t)]$, $[\mathbf{s}(t)]$, ossia in coordinate omogenee. Allora la trasformazione $T\mathbf{r} = \mathbf{s}$ diventa $M_{\mathbf{c}}\mathbf{R}(t) = \mathbf{S}(t)$. Scegliamo i rappresentanti speciali della classe di equivalenza omogenea $\mathbf{R}(t)$ dati da $(x(t), y(t), z(t), 1)$. L'azione di $M_{\mathbf{c}}$ manda questo vettore in un rappresentante di $\mathbf{S}$ che scriviamo (x'', y'', z'', g) . Dal Teorema 16.4.4 ora abbiamo

$$\begin{aligned} x''(t) &= (d + c_x n_x)(r_x^0 + tu_x) + c_x n_y(r_y^0 + tu_y) + c_x n_z(r_z^0 + tu_z) - c_x d_0 \\ y''(t) &= c_y n_x(r_x^0 + tu_x) + (d + c_y n_y)(r_y^0 + tu_y) + c_y n_z(r_z^0 + tu_z) - c_y d_0 \\ z''(t) &= c_z n_x(r_x^0 + tu_x) + c_z n_y(r_y^0 + tu_y) + (d + c_z n_z)(r_z^0 + tu_z) - c_z d_0 \\ g &= n_x(r_x^0 + tu_x) + n_y(r_y^0 + tu_y) + n_z(r_z^0 + tu_z) - d_1. \end{aligned}$$

Da qui passiamo ai valori tridimensionali delle componenti di $\mathbf{s}(t)$ dividendo x'' , y'' , z'' per g e poi facciamo tendere t ad infinito per trovare il punto di fuga $\mathbf{f}_{\mathbf{u}}$. Il limite esiste se il limite del denominatore non si

annulla. Osserviamo che, senza perdita di generalità, si può assumere che il punto $\mathbf{r}^0$ appartenga al piano P : infatti, se il fascio di rette non è parallelo a P , ossia se $\langle \mathbf{n}, \mathbf{u} \rangle \neq 0$, allora ogni retta del fascio ha un (unico) punto di intersezione con P e scegliamo $\mathbf{r}^0$ come questo punto di intersezione; se invece il fascio di rette è parallelo a P , possiamo scegliere per $\mathbf{r}(t)$ una retta che giace in P , ed allora ogni suo punto $\mathbf{r}^0$ sta in P . Con questa scelta di $\mathbf{r}^0$ sappiamo che $\langle \mathbf{n}, \mathbf{r}^0 \rangle$ è la distanza d_1 di P dall'origine. Ora vediamo dall'espressione di g che il limite per $t \rightarrow \infty$ è zero se e solo se $\langle \mathbf{n}, \mathbf{u} \rangle \neq 0$. In tal caso il limite è esattamente l'espressione dell'enunciato. Se invece $\langle \mathbf{n}, \mathbf{u} \rangle = 0$ il termine in t viene moltiplicato per zero e scompare, quindi le rette trasformate non si incontrano per alcun valore di t . In tal caso otteniamo $g = 0$, ossia un punto improprio con quarta coordinata omogenea 0, ed è chiaro (facendo tendere t ad infinito) che le altre coordinate omogenee tendono a (u_x, u_y, u_z) . $\square$

NOTA 16.5.2. La determinazione del punto di fuga di un fascio di rette in direzione $\mathbf{u}$, ricavata nel Teorema 16.5.1, si può ottenere senza calcoli con un argomento geometrico alternativo ed illuminante, che ora presentiamo.

Per prima cosa si osserva che il punto di fuga deve esistere, perchè la proiezione centrale diminuisce le distanze di punti che si muovono parallelamente, ossia a distanza costante prima della proiezione (questo fatto geometricamente evidente equivale alla divisione prospettica per la distanza).

In secondo luogo, L'immagine della proiezione prospettica centrale è il piano di visuale P . In particolare, il punto di fuga del fascio di rette parallele a $\mathbf{u}$ è un punto di P . In questo fascio basta allora considerare la retta che passa per il centro di proiezione $\mathbf{c}$, diciamo $\mathbf{r}(t) = \mathbf{c} + t\mathbf{u}$: il punto di fuga del fascio è l'intersezione di questa retta con P (esso esiste se e solo se $\mathbf{u}$ non è un vettore parallelo a P : se invece lo è, allora le rette del fascio sono tutte parallele al piano P , e poiché $\mathbf{c}$ si sceglie fuori di questo piano la retta $\mathbf{s}$ non lo interseca: in tal caso si dice che il punto di fuga è all'infinito, come vedremo meglio nella successiva Nota 16.5.3).

Sia come sempre d la distanza fra $\mathbf{c}$ e P . Poiché $\mathbf{u}$ ha norma 1, il valore di t per il quale la retta $\mathbf{s}(t) = \mathbf{c} + t\mathbf{u}$ interseca P è la lunghezza $d_{\mathbf{u}}$ del segmento da $\mathbf{c}$ a P nella direzione di $\mathbf{u}$. La lunghezza della proiezione di questo segmento sulla direzione normale a P (quella del versore $\mathbf{n}$) è quindi d : pertanto $d_{\mathbf{u}} = d_1 / \langle \mathbf{n}, \mathbf{u} \rangle = d_1 / \cos \theta$, dove θ è l'angolo sotteso da $\mathbf{u}$ e $\mathbf{n}$. In altre parole, il punto di fuga è precisamente $\mathbf{f}_{\mathbf{u}} = \mathbf{c} + \mathbf{u}d / \cos \theta$, come trovato con i calcoli della dimostrazione del Teorema 16.5.1. $\square$

NOTA 16.5.3. È interessante vedere come cambia la posizione del punto di fuga al variare del versore $\mathbf{u}$ del fascio nella sfera unitaria. Se $\mathbf{u}$ varia nella sfera mantenendo costante la propria deviazione angolare θ da $\mathbf{n}$ (ossia se si muove sul parallelo C_{θ} rispetto al polo $\mathbf{n}$), allora il punto di fuga descrive anch'esso una circonferenza nello spazio, ed esattamente il parallelo C_{θ} della sfera con centro $\mathbf{c}$ e raggio $d / \cos \theta$, che tende ad infinito quando θ tende a $\pi/2$. Questo è vero fin tanto che $\theta \neq \pi/2$, ossia purché $\mathbf{u}$ non sia perpendicolare a $\mathbf{n}$, ovvero sull'equatore. Il caso di $\mathbf{u}$ sull'equatore è il caso limite in cui il raggio di queste circonferenze diverge, e diciamo allora che i punti di fuga sono all'infinito: in tal caso il fascio di rette parallele rimane tale anche dopo la trasformazione prospettica.

Quindi punti di fuga a distanza piccola dal centro di proiezione si ottengono quando il fascio si avvicina ad essere perpendicolare al piano di proiezione (ossia l'angolo θ vicino a zero, $\mathbf{u}$ vicino ad essere allineato con $\mathbf{n}$), ed in aggiunta quando d è piccolo. Per quanto visto nella Nota 16.4.3, l'ultima condizione significa che la distanza fra il centro di proiezione (ovvero l'osservatore) ed il piano di proiezione P è piccola, ossia la prospettiva è accentuata. In altre parole, punti di fuga ravvicinati corrispondono a un osservatore vicino alla scena, con l'accentuazione (o distorsione) prospettica tipica di ciò che si ottiene quando si fotografa una scena da vicino. In questo caso, perché si riesca a riprendere la scena nella sua interezza, bisogna usare un grandangolo, e questa distorsione prospettica si chiama appunto *distorsione grandangolare*. $\square$

Dal Teorema 16.5.1 otteniamo immediatamente il risultato seguente. Si noti che la terza identità è coerente con quella trovata nella Nota 16.1.2.

COROLLARIO 16.5.4. *Scriviamo $\mathbf{f}_i$ invece che $\mathbf{f}_{\mathbf{e}_i}$ ($i=1,2,3$) per i punti di fuga principali della prospettiva centrale, ossia i punti di fuga dei fasci di rette paralleli, rispettivamente, ai tre versori canonici di base. Essi sono, rispettivamente,*

$$\mathbf{f}_1 = (c_x + \frac{d}{n_x}, c_y, c_z) \quad (16.5.1a)$$

$$\mathbf{f}_2 = (c_x, c_y + \frac{d}{n_y}, c_z) \quad (16.5.1b)$$

$$\mathbf{f}_3 = (c_x, c_y, c_z + \frac{d}{n_z}). \quad (16.5.1c)$$

In queste identità, si noti che, per ogni coppia di versori canonici, i corrispondenti punti di fuga principali hanno componente lungo il restante versore uguale a quella del centro di proiezione: ad esempio, i punti di fuga principali nelle direzioni degli assi x e y hanno la stessa altezza sul piano $\{z = 0\}$ del centro di proiezione, e così via ruotando gli indici.

ESERCIZIO 16.5.5. (i) Un pittore sceglie la prospettiva per un quadro nel quale, oltre ad altri dettagli, dipinge due edifici a forma di parallelepipedo. Il primo edificio è disposto in modo frontale, diciamo come il cubo di lato 1 nel primo ottante che contiene $(0, 0, 3)$ ed ha la faccia più vicina al piano di proiezione parallela ad esso e disposta sul piano $\{z = 3\}$, e le altre due facce che contengono $(0, 0, 3)$ disposte rispettivamente sui piani $\{x = 0\}$ e $\{y = 0\}$. Per dipingere questo edificio il pittore sceglie la distanza d del proprio punto di osservazione dal piano della tela (supponiamo che il piano sia $\{z = 2\}$ e sia $d = 1$), ed un punto di fuga per i lati del cubo diretti parallelamente all'asse z (ossia ortogonalmente alla tela). Supponiamo che come punto di fuga scelga $\mathbf{f}_3 = (1, 1, 2)$.

L'altro edificio è anch'esso a forma cubica, con angolo di $\pi/4$ (45 gradi) rispetto alla direzione di osservazione: ad esempio il cubo di altezza $\sqrt{2}$ sopra il quadrato sul piano $\{y = 0\}$ di

vertici $(x = 4, z = 3)$, $(5, 4)$, $(4, 5)$, $(3, 4)$. Il pittore come disegna in maniera prospetticamente corretta i lati di questo secondo edificio?

- (ii) La stessa scena dipinta dal pittore viene fotografata con una macchina fotografica disposta dove era il pittore, con la stessa angolazione (quindi frontalmente rispetto al primo edificio, ossia con le stesse angolazioni orizzontale e verticale). Quella che era la distanza del pittore dalla tela ora diventa la distanza fra il punto di messa a fuoco ed il piano di proiezione, ossia il piano della pellicola (o del sensore, in fotografia digitale). Poiché il piano del sensore è abbastanza vicino al piano centrale dell'obiettivo, con buona approssimazione possiamo rimpiazzare questa distanza d con la lunghezza focale dell'obiettivo, che, quando guardiamo la fotografia, non ci è nota. Come possiamo ricostruire dalla fotografia l'angolazione nello spazio tridimensionale del secondo edificio?

SVOLGIMENTO. (i). Sappiamo che $\mathbf{n} = (0, 0, 1)$, quindi $n_z = 1$. Dall'equazione (16.5.1c) del punto di fuga principale $\mathbf{f}_3$ ricaviamo $c_z + d = 2$, quindi $c_z = 1$, ed anche le altre coordinate del centro di proiezione valgono 1. In altre parole $\mathbf{c} = (1, 1, 1)$.

Le due facciate visibili del secondo edificio sono dirette come $\mathbf{u}^- = (-1/\sqrt{2}, 0, -1/\sqrt{2})$ e $\mathbf{u}^+ = (1/\sqrt{2}, 0, -1/\sqrt{2})$. I punti di fuga per questi versori si ottengono dal Teorema 16.5.1:

$$\mathbf{f}_- = (1, 1, 1) - \sqrt{2}(-1/\sqrt{2}, 0, -1/\sqrt{2}) = (2, 1, 2)$$

$$\mathbf{f}_+ = (1, 1, 1) - \sqrt{2}(1/\sqrt{2}, 0, -1/\sqrt{2}) = (0, 1, 2).$$

Il pittore disegna quindi i bordi del secondo edificio facendo convergere le linee parallele ai rispettivi punti di fuga appena trovati.

(ii). Dalla fotografia misuriamo le coordinate del punto di fuga dei lati del primo edificio ortogonali al piano della pellicola, che abbiamo chiamato $\mathbf{f}_3$. Grazie all'equazione (16.5.1c) questa misura ci dà il vettore $\mathbf{c} + d\mathbf{e}_3$, ossia le prime due coordinate di $\mathbf{c}$ ed il valore della quantità $c_z + d$. Ora consideriamo i punti di fuga associati agli spigoli non verticali del secondo edificio (quelli verticali rimangono paralleli al piano della

pellicola, perché l'edificio e questo piano sono entrambi verticali, quindi paralleli: il terzo punto di fuga è all'infinito). Vogliamo determinare i versori $\mathbf{u}_{\pm}$ delle due famiglie ortogonali di tali spigoli. Osserviamo che $\mathbf{n} = \mathbf{e}_3$, e pertanto $\langle \mathbf{u}^{\pm}, \mathbf{n} \rangle = u_z^{\pm}$. Ora dalle equazioni (16.5.1a) e (16.5.1c) ricaviamo i vettori $\mathbf{c} + \frac{d}{u_z^{\pm}} \mathbf{u}^{\pm}$. Poiché le prime due componenti di $\mathbf{c}$ sono già note, questo ci restituisce i valori di $du_x^{\pm}/u_z^{\pm}$ e $du_y^{\pm}/u_z^{\pm}$ (ma attenzione: sappiamo già che $u_y^{\pm} = 0$, perché l'edificio è verticale, non inclinato, quindi l'ultima informazione non ci serve). Per le terze componenti, troviamo in entrambi i casi il valore di $c_z + d$, ma sapevamo già che questo valore è 1, quindi anche questa informazione non ci dà nulla di nuovo. Infine sappiamo che $\mathbf{u}^+$ e $\mathbf{u}^-$ sono ortogonali e di norma 1. Ricapitolando tutte queste informazioni nell'ordine dato, abbiamo:

$$\begin{aligned} c_z + d &= 1 \\ d \frac{u_x^+}{u_z^+} &= \text{costante nota} \\ d \frac{u_x^-}{u_z^-} &= \text{costante nota} \\ u_x^+ u_x^- + u_z^+ u_z^- &= 0 \\ (u_x^+)^2 + (u_z^+)^2 &= 1 \\ (u_x^-)^2 + (u_z^-)^2 &= 1. \end{aligned}$$

Si tratta di sei equazioni quadratiche nelle sei incognite $d, c_z, u_x^{\pm}, u_z^{\pm}$ (si rammenti che $u_y^{\pm} = 0$). Se il sistema è risolvibile, in tal modo si determina d ed i due versori $\mathbf{u}^{\pm}$ richiesti. Lasciamo al lettore considerare le condizioni di risolubilità e sviluppare i calcoli numerici dopo aver fissato a proprio piacere le posizioni dei tre punti di fuga utilizzati (ma ci si rammenti che questi tre punti debbono essere scelti tutti alla stessa altezza y , perché giacciono sulla retta che rappresenta l'orizzonte sul piano della pellicola, e questa retta è orizzontale perché la macchina fotografica non è inclinata di lato).

Si noti anche che a questo punto conosciamo anche $\mathbf{c}$ e d , e quindi l'intera matrice della trasformazione prospettica. Ritorneremo su

questa idea e svilupperemo i dettagli in una situazione analoga nel successivo Esempio 16.5.8. $\square$

ESERCIZIO 16.5.6. Per accentuare la drammaticità della fotografia, un fotografo vuol riprendere un grattacielo in modo che i due spigoli adiacenti del tetto visibili dalla sua posizione si dispongano parallelamente a due lati adiacenti del fotogramma (si veda la Figura ??). È possibile? Come si deve orientare la macchina fotografica?

SVOLGIMENTO. Senza perdita di generalità poniamo nell'origine il punto di messa a fuoco dentro la macchina fotografica: in tal modo la trasformazione prospettica corrisponde alla matrice M_0 particolarmente semplice del Teorema 16.4.4, dove d è la distanza fra il centro di proiezione (l'origine) ed il piano di proiezione (il piano del sensore o della pellicola). Dobbiamo trovare se esiste un versore $\mathbf{n}$ verso cui orientare la macchina fotografica onde ottenere l'inquadratura voluta. Osserviamo che, se tale versore esiste, allora possiamo ruotare l'immagine nel piano di proiezione semplicemente ruotando assialmente (intorno a $\mathbf{n}$) la macchina fotografica in verso opposto: pertanto non importa se i due lati della sommità del grattacielo si proiettano parallelamente ai lati del fotogramma, quello che conta è che si proiettino in maniera perpendicolare. Chiamiamo $\mathbf{p}$ il punto di congiunzione dei due lati del tetto che il fotografo vede: i due lati sono disposti lungo due rette uscenti da $\mathbf{p}$ e dirette come due vettori canonici di base, diciamo $\mathbf{e}_1$ e $\mathbf{e}_3$ qui adottiamo la consueta terminologia della Computer Graphics e disponiamo il lato della profondità del grattacielo lungo l'asse z : l'asse x rappresenta la direzione, diciamo, quasi frontale rispetto al fotografo, e l'asse y è l'altezza). Una volta proiettate prospetticamente, queste due rette passano per il punto $\mathbf{p}'$ immagine di $\mathbf{p}$ sotto la trasformazione prospettica, e sono dirette verso i punti di fuga principali corrispondenti ai versori $\mathbf{e}_1$ e $\mathbf{e}_3$. Diciamo che, a meno di cambiamento di scala, l'edificio abbia larghezza 1 (ovvero estensione 1 lungo l'asse x), profondità w (lungo l'asse z) ed altezza h (lungo l'asse y). Per comodità, ma senza perdita di generalità, disponiamo la base dell'edificio sul rettangolo $(1, 0, 1)$, $(1, 0, w)$, $(2, 0, w)$ e $(2, 0, 1)$: pertanto il punto $\mathbf{p}$ è $(1, h, 1)$, ed i due lati visibili del tetto hanno come rispettivi punti terminali

$\mathbf{p}_1 := \mathbf{p} + \mathbf{e}_1 = (2, h, 1)$ e $\mathbf{p}_3 := \mathbf{p} + w\mathbf{e}_3 = (1, h, w)$. In realtà non abbiamo bisogno di utilizzare la individuazione dei punti di fuga principali dato dal Corollario 16.5.4: basta effettuare la trasformazione prospettica (data nel nostro caso dalla matrice particolarmente semplice M_0) ai punti $\mathbf{p}_1$ e $\mathbf{p}_3$ ed imporre che i loro trasformati $\mathbf{p}'_1$ e $\mathbf{p}'_3$ verifichino la relazione di ortogonalità

$$\mathbf{p}'_1 - \mathbf{p}' \perp \mathbf{p}'_3 - \mathbf{p}'.$$

Ma è altrettanto semplice, e a questo punto più elegante, utilizzare le formule per i punti di fuga principali del Corollario 16.5.4: la precedente relazione di ortogonalità allora diventa

$$\mathbf{p}'_1 - \mathbf{p}' \perp \mathbf{p}'_3 - \mathbf{p}' \quad (16.5.2)$$

(osserviamo che non occorre moltiplicare per w il vettore $\mathbf{f}_{\mathbf{e}_3}$ del punto di fuga (16.5.1c) per w onde ottenere proprio $\mathbf{p}'_3$: l'ortogonalità non dipende dalla lunghezza di questo vettore).

Vista la forma della matrice M_0 del Teorema 16.4.4 ed il fatto che $\mathbf{c} = 0$, troviamo $\mathbf{p}' = (d/\langle \mathbf{n}, \mathbf{p} \rangle)\mathbf{p}$. Allora, in base alle identità (16.5.1a) e (16.5.1c), la relazione di ortogonalità (16.5.2) diventa

$$(d/n_x)\mathbf{e}_1 - (d/\langle \mathbf{n}, \mathbf{p} \rangle)\mathbf{p} \perp (d/n_z)\mathbf{e}_3 - (d/\langle \mathbf{n}, \mathbf{p} \rangle)\mathbf{p}.$$

Vista la scelta di $\mathbf{p} = (1, h, 1)$, un facile calcolo diretto permette di riscrivere questa condizione come

$$n_x^2 + n_z^2 - n_x n_z h^2 + (n_x + n_z)n_y h = 0. \quad (16.5.3)$$

Rammentando anche la condizione di normalizzazione $n_x^2 + n_y^2 + n_z^2 = 1$, abbiamo due equazioni quadratiche, che corrispondono a superficie quadratiche che si intersecano in una curva. Tutti i versori in questa curva risolvono il problema. Ad esempio, una soluzione banale è $\mathbf{n} = \pm \mathbf{e}_2 = (0, \pm 1, 0)$: in tal caso la macchina fotografica è orientata verticalmente, e se il grattacielo fosse trasparente in effetti si vedrebbero i lati del tetto rimanere paralleli dopo la proiezione, perché sono paralleli al piano di proiezione (ovviamente il grattacielo non è trasparente, ma questa soluzione rimane valida se invece che mettere la macchina fotografica nell'origine la mettiamo alta sopra l'edificio in un punto dell'asse y , e la orientiamo in basso, lungo il versore $(0, -1, 0)$. Oltre a

questa soluzione banale c'è una famiglia ad un parametro di soluzioni non banali.

La curva delle soluzioni non è degenera, ossia non consiste dei soli punti $(0, \pm 0)$, almeno se h è abbastanza grande. Infatti, poniamo $\mathbf{n} = (t, \sqrt{1-2t^2}, t)$, con $0 < t < 1/\sqrt{2}$: allora la condizione (16.5.3) diventa

$$0 = 2t^2 - 4t^2h^2 + 2t\sqrt{1-2t^2}h = (1-2h^2)t^2 + t\sqrt{1-2t^2}h.$$

C'è una soluzione per $t = 0$; per piccoli $t > 0$, il primo termine all'ultimo membro cresce quadraticamente ed è negativo se $h > 1/\sqrt{2}$, mentre il secondo cresce linearmente ed è positivo perché $h > 0$, quindi la somma è positiva. Quando t si avvicina a $1/\sqrt{2}$ la somma si avvicina a $(1-2h^2)/2$, che è negativo: quindi c'è una soluzione non banale nell'intervallo $0 < t < 1/\sqrt{2}$. $\square$

ESEMPIO 16.5.7. La prospettiva centrale con un solo punto di fuga è quella in cui il piano di proiezione è parallelo a due dei versori canonici, ossia perpendicolare al terzo. Ad esempio, se come d'abitudine in Computer Graphics la direzione perpendicolare è quella dell'asse z , si ottiene la prospettiva del tipo della proiezione standard della Sezione 16.1. Abbiamo già determinato nell'Esempio 16.1.3 come appare il cubo unitario in questa prospettiva: lo disegniamo in Figura 5.

Questa è la prospettiva più frequente nella pittura (ed infatti fu, almeno in parte, inventata o reinventata da pittori e artisti), ed è spesso usata in modo da comunicare sensazioni di maestosità o di tensione. Nelle prossime Figure vediamo vari esempi, cominciando da un esempio dove la prospettiva è completamente sbagliata (particolarmente sulle piastrelle del pavimento), ma la bravura dell'illustratore (Jean Fouquet, un illustratore di manoscritti miniati del primo Rinascimento francese) dissimula la distorsione prospettica. Continuiamo poi con esempi assai più precisi del Rinascimento italiano... per finire con qualche imprecisione moderna.

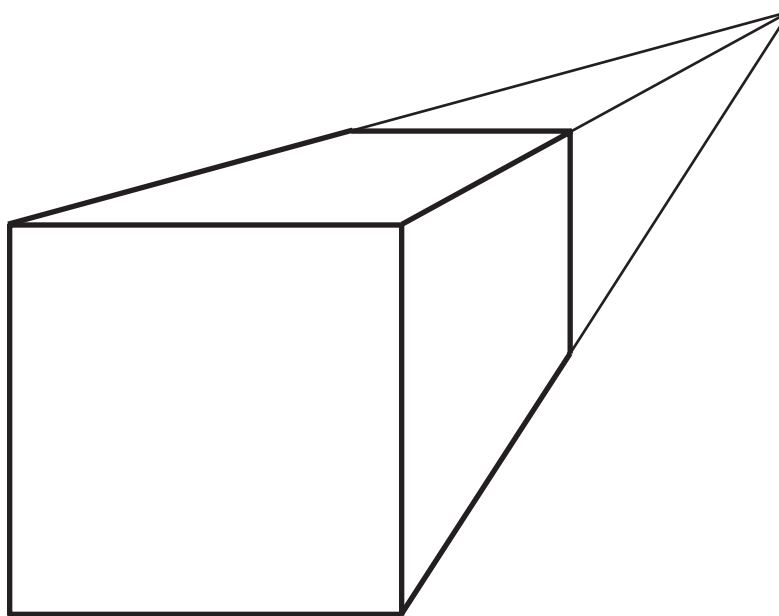


FIGURA 5. Prospettiva centrale con un solo punto di punto di fuga al finito

Il caso di due punti di fuga è illustrato nella Figura 11. La retta a cui appartengono i due punti di fuga si chiama l'*orizzonte*. Si osservi che, in pittura, è il pittore a decidere a che altezza collocare la retta dell'orizzonte, esattamente come, nel caso di un solo punto di fuga, l'altezza di tale punto nel disegno può essere scelta arbitrariamente. In effetti, la precedente Figura 5 illustra il caso di orizzonte a media altezza (e centro di prospettiva alto, ovvero punto di osservazione dall'alto: la faccia superiore del cubo è visibile, il centro di prospettiva (ossia l'osservatore) è più alto del cubo). In Figura 11 disegniamo il cubo con orizzonte e centro alti, in Figura 12 lo ridisegniamo con orizzonte e centro più bassi (la proiezione sul piano di visuale dell'orizzonte è la linea tratteggiata). Se abbassassimo ancora di più il centro di prospettiva, vedremmo anche la faccia inferiore del cubo, dal di sotto.



FIGURA 6. Prospettiva ad un solo punto di fuga, ma completamente sbagliata, nella illustrazione “Ciro II il Grande e gli ebrei”, di Jean Fouquet (circa 1420–1477), dal volume “Antichità giudaiche”.

Vediamo qualche esempio di prospettiva a due punti di fuga in pittura nelle Figure 16.5.7 e 16.5.7.

□

ESEMPIO 16.5.8. (Costruzione della prospettiva a due punti di fuga in pittura.) Ecco il modo in cui un pittore costruisce la prospettiva centrale con due punti di fuga. In questa prospettiva il piano di visuale deve intersecare due dei tre assi coordinati, diciamo x e y . Il pittore sceglie la distanza dall’origine del piano di visuale P (il piano su cui sta la tela) e la sua angolazione; sul piano P traccia poi la linea

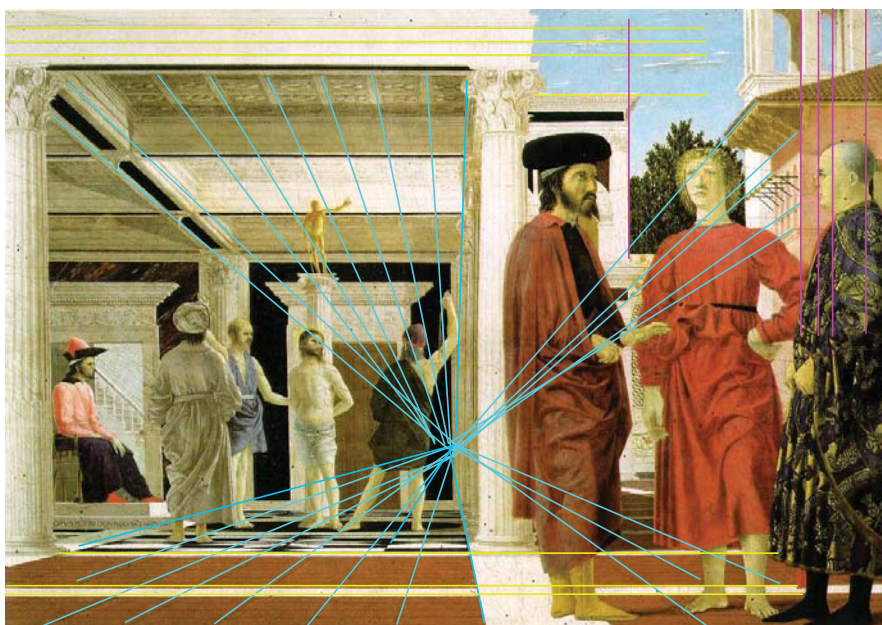


FIGURA 7. Un solo punto di fuga nella Flagellazione di Cristo di Piero della Francesca. Sono evidenziate con colori diversi le rette ortogonali al piano del dipinto, che convergono al punto di fuga, e quelle dirette orizzontalmente e verticalmente. La prospettiva è scelta in base a finalità artistiche: il centro di proiezione (ossia l'osservatore) ad altezza medio-bassa ingigantisce le figure vicine al piano di proiezione (ossia il dipinto), conferendogli solennità. Piero della Francesca fu uno degli inventori della prospettiva, eppure in una scena così ricca di linee parallele si può notare qualche lieve imprecisione nel punto di convergenza. Però, però... questo dipinto non è su tela ma su una tavola di legno, leggermente curva... è possibile che la curvatura spieghi i piccoli errori nel punto di incontro del fascio di rette? Se un dipinto piano sul quale immaginiamo un fascio di rette radiali che si incontrano, diciamo al suo centro, venisse incollato su una superficie curva e poi fotografato (come accade in questa figura), nell'immagine piana così ottenuta le rette si incurverebbero, e le loro tangenti sulle regioni laterali della fotografia non si incontrerebbero tutte esattamente nel punto di centro.

dell'orizzonte, e su di essa sceglie i due punti di fuga $\mathbf{f}_1$ e $\mathbf{f}_2$. Mostriamo che con queste scelte la prospettiva è data dalla seguente matrice



FIGURA 8. Un solo punto di fuga nell'Annunciazione di Carlo Crivelli. Le rette ortogonali al piano del dipinto sono evidenziate. La proiezione prospettica del punto di fuga si proietta sul dipinto nella teca al centro, conferendogli un senso di misteriosa aspettativa e di importanza. La prospettiva è molto accentuata: la distanza del centro di proiezione dal piano di visuale quindi è assai ridotta (si veda la Nota 16.5.3). La cura della convergenza della prospettiva qui è magistrale.



FIGURA 9. Un solo punto di fuga nell'Altare Paumgartner di Albrecht Dürer. Si noti l'errore prospettico dato dalla linea gialla della tettoia a sinistra: o forse la tettoia non è perfettamente parallela al muro?

prospettica M , espressa solo in termini di h , α_x , d_0 , t_1 e t_2 :

$$d_0 \begin{pmatrix} 1 - t_1 & \tan \alpha_x (1 - t_2) & 0 & -d_0 (1 - t_2) / \cos \alpha_x \\ (\cot \alpha_x) t_1 & t_2 & 0 & -d_0 t_1 / \sin \alpha_x \\ 0 & 0 & t_2 - t_1 & 0 \\ \cos \alpha_x & \sin \alpha_x & 0 & -(\cos \alpha_x (1 - t_2) + \sin \alpha_x t_1) / d_0 \end{pmatrix}.$$

Cominciamo con l'osservare che scegliere l'angolazione di P significa

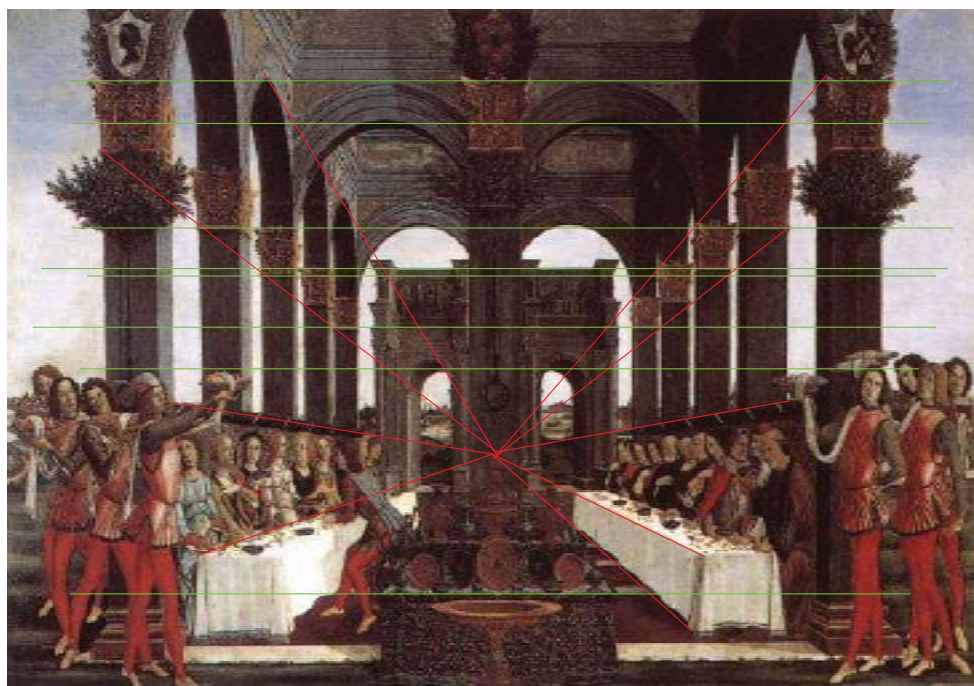


FIGURA 10. Un solo punto di fuga nel Banchetto Nuziale di Sandro Botticelli. Il centro di proiezione medio-basso è abbastanza alto da permetterci di vedere leggermente dall'alto i commensali seduti, come se chi guarda fosse in piedi, ma sufficientemente basso da ingigantire le arcate della loggia in primissimo piano, conferendogli maestosità e solennità. La prospettiva è mediamente accentuata: la distanza del centro di proiezione dal piano del dipinto non è elevata.

scegliere i suoi coseni direttori $n_x = \cos \alpha_x$, $n_y = \cos \alpha_y$, $n_z = \cos \alpha_z$, in altre parole il suo versore normale $\mathbf{n} = (n_x, n_y, n_z)$. In questo caso, però, ci sono solo due punti di fuga, e quindi intersezioni di P solo con i due assi coordinati x e y (chiamiamoli E_1 ed E_2 rispettivamente), ma non con il terzo: quindi P è parallelo all'asse z e l'angolo α_z è nullo, mentre α_x e α_y sono complementari: $\alpha_x + \alpha_y = \pi/2$. Allora si ha $\mathbf{n} = (n_x, n_y, 0)$ con $n_x^2 + n_y^2 = 1$. La retta s che passa per E_1 e E_2 giace in P e quindi è perpendicolare a $\mathbf{n}$: pertanto essa taglia gli assi x e y con angoli α_y e α_x , rispettivamente (Figura 15).

Sia, come sempre, d_0 la distanza dall'origine che il pittore ha scelto

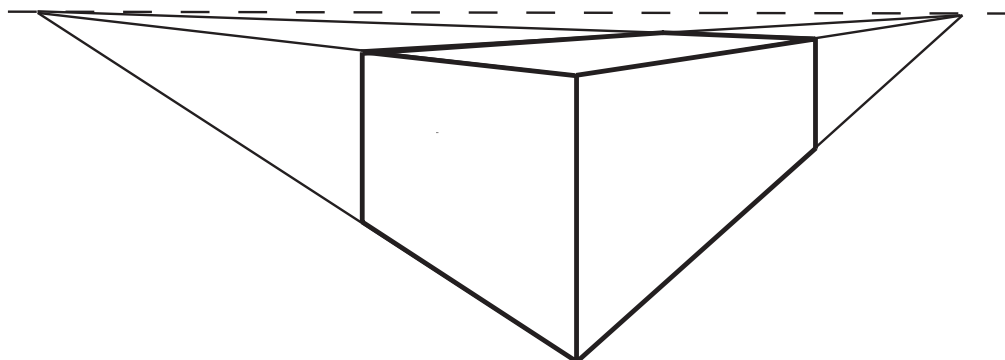


FIGURA 11. Prospettiva centrale con due punti di fuga, centro di proiezione alto, e quindi orizzonte alto.

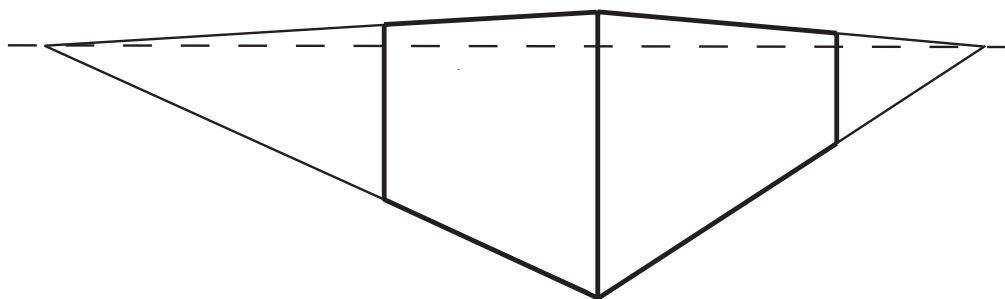


FIGURA 12. Prospettiva centrale con due punti di fuga e centro di proiezione ad altezza media, e quindi orizzonte ad altezza media.

per P . Allora

$$E_1 = (d_0 / \cos \alpha_x, 0, 0) = d_0 (1/n_x, 0, 0) \quad (16.5.4)$$

e

$$E_2 = d_0 (0, 1/n_y, 0). \quad (16.5.5)$$

In altre parole, $n_x = d_0 / \|\mathbf{E}_1\|$ e $n_y = d_0 / \|\mathbf{E}_2\|$; come già visto, $n_z = 0$. La retta $\mathbf{s}$, essendo ortogonale a $\mathbf{n}$, ha equazione $xn_x + yn_y = d_0$, ovvero $x/\|\mathbf{E}_1\| + y/\|\mathbf{E}_2\| = 1$ (questa formulazione è anche ovvia



FIGURA 13. Prospettiva centrale con due punti di fuga nella Festa Nuziale di Pieter Breughel. Centro di proiezione ed orizzonte a media altezza. In questo dipinto le linee parallele sono poche (la scena è meno geometrica delle altre) ed i punti di fuga molto lontani (la prospettiva non è accentuata, il centro di proiezione è lontano dal dipinto), ma sembra di poter notare qualche incertezza nella prospettiva. Si osservi il bordo della panca trasportata da due uomini in primo piano: la sua linea (gialla) non converge neppure approssimativamente al punto di fuga, quindi nella scena la panca non è parallela alla tavola imbandita.

geometricamente), od alternativamente anche, in forma parametrica,

$$\mathbf{s}(t) = \mathbf{E}_1 + (\mathbf{E}_2 - \mathbf{E}_1)t$$

(il segmento che congiunge E_1 ed E_2 corrisponde ai valori di t in $[0, 1]$).

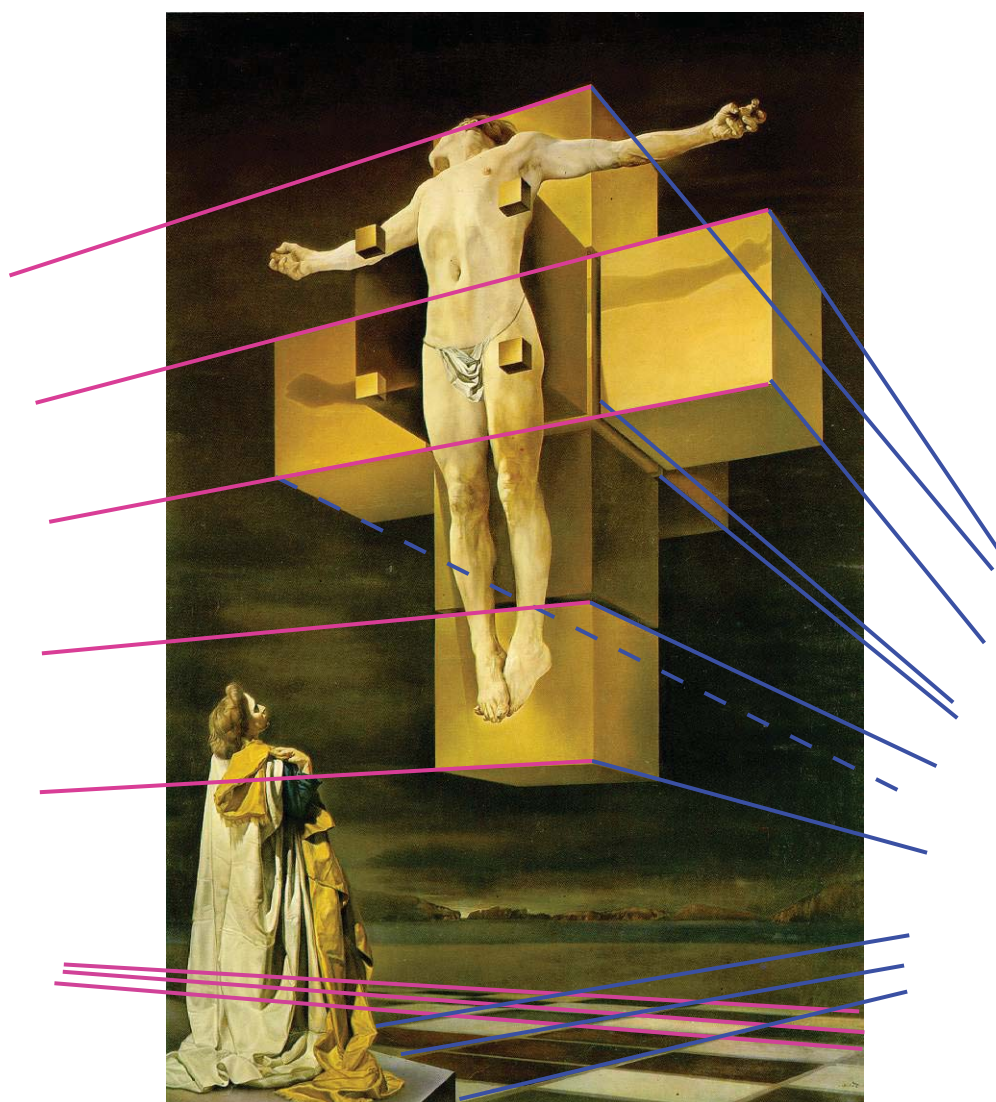


FIGURA 14. Prospettiva centrale con due punti di fuga nella Crocifissione di Salvador Dalí. Centro di proiezione (ovvero osservatore) ed orizzonte bassi per dare maestosità al crocifisso. Si noti la linea tratteggiata, che non converge al punto di fuga del suo fascio: un errore prospettico o un crocifisso a lati non paralleli?

Riscriviamo le equazioni parametriche componente per componente:

$$r_x(t) = \frac{d_0}{n_x} (1 - t)$$

$$r_y(t) = \frac{d_0}{n_y} t$$

$$r_z(t) = 0.$$

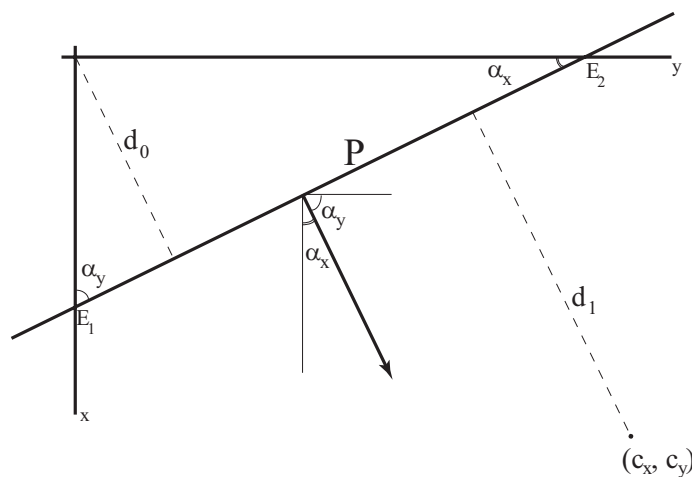


FIGURA 15. Prospettiva centrale con due punti di fuga: angolazione e coseni direttori del piano di visuale rispetto agli assi coordinati.

Sul piano di visuale P il pittore fissa la linea dell'orizzonte $\mathbf{h}$ ad altezza h rispetto al suolo, ossia parallela alla retta $\mathbf{s}$ e a distanza h sopra di essa. A partire dalle equazioni parametriche di $\mathbf{r}$ che abbiamo appena calcolato ora troviamo

$$h_x(t) = \frac{d_0}{n_x} (1 - t)$$

$$h_y(t) = \frac{d_0}{n_y} t$$

$$h_z(t) = h.$$

Sulla retta di orizzonte $\mathbf{h}$ il pittore fissa i due punti di fuga $\mathbf{f}_1$ e $\mathbf{f}_2$, scegliendo due valori t_1 e t_2 di t . Quindi si ha

$$\mathbf{f}_1 = \left(\frac{d_0}{n_x} (1 - t_1), \frac{d_0}{n_y} t_1, h \right)$$

$$\mathbf{f}_2 = \left(\frac{d_0}{n_x} (1 - t_2), \frac{d_0}{n_y} t_2, h \right).$$

Ora calcoliamo le coordinate del centro di proiezione $\mathbf{c}$ e la sua distanza d da P a partire da quelle dei punti di fuga, grazie al Corollario 16.5.4. È già stato esplicitamente notato in quel Corollario che l'altezza rispetto al suolo (ossia la terza componente) di $\mathbf{c}$ coincide con quella di $\mathbf{f}_1$ e $\mathbf{f}_2$,

ossia vale h . Per le altre componenti, usando $\mathbf{f}_1$ troviamo

$$\begin{aligned} c_x + \frac{d}{n_x} &= \frac{d_0}{n_x} (1 - t_1) \\ c_y &= \frac{d_0}{n_y} t_1 \\ c_z &= h, \end{aligned} \tag{16.5.6}$$

ed usando invece $\mathbf{f}_2$,

$$\begin{aligned} c_x &= \frac{d_0}{n_x} (1 - t_2) \\ c_y + \frac{d}{n_y} &= \frac{d_0}{n_y} t_2 \\ c_z &= h. \end{aligned} \tag{16.5.7}$$

Osserviamo anzitutto, per uso futuro, che da (16.5.7) e (16.5.6) segue

$$d_1 = \langle \mathbf{n}, \mathbf{c} \rangle = \cos \alpha_x (1 - t_2) + \cos \alpha_y t_1 = \cos \alpha_x (1 - t_2) + \sin \alpha_x t_1 \tag{16.5.8}$$

(l'ultima identità segue ovviamente dal fatto che gli angoli α_x e α_y sono complementari).

In secondo luogo, sostituendo le seconde equazioni nelle prime, troviamo

$$\frac{d}{n_x} = \frac{d_0}{n_x} (t_2 - t_1) \tag{16.5.9}$$

$$\frac{d}{n_y} = \frac{d_0}{n_y} (t_2 - t_1). \tag{16.5.10}$$

Pertanto si ha

$$d = d_0 (t_2 - t_1).$$

Ora consideriamo il cubo con un vertice all'origine e le facce sui piani coordinati la cui proiezione sul piano di base $\{y = 0\}$ è il quadrato Q che tocca P nel vertice opposto all'origine, e sia q la lunghezza del suo lato. La parte di E_1 che non sta in Q è un segmento sull'asse x che col piano P ed il lato di Q forma un triangolo il cui angolo opposto è α_x . Pertanto, come illustrato anche in Figura 16, si ha $E_1 = q(1 + \tan \alpha_x)$. Nello stesso modo si vede che $E_2 = q(1 + \tan \alpha_y) = q(1 + \cot \alpha_x)$ dal momento che $\alpha_y = \frac{\pi}{2} - \alpha_x$.

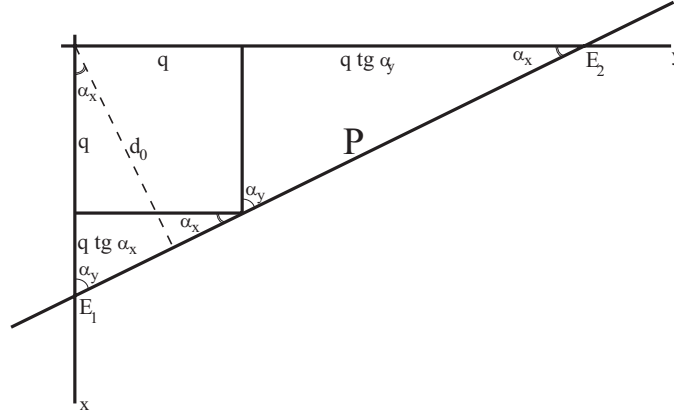


FIGURA 16. Intersezioni del piano di visuale con gli assi coordinati in funzione dei suoi coseni direttori e della distanza dall'origine.

Da queste identità e (16.5.4) abbiamo

$$q = \frac{d_0}{(1 + \tan \alpha_x) \cos \alpha_x} = \frac{d_0}{\cos \alpha_x + \sin \alpha_x} = \frac{d_0}{n_x + n_y} \quad (16.5.11)$$

di nuovo perché $\cos \alpha_y = \sin \alpha_x$.

Lo stesso calcolo con cui abbiamo dimostrato (16.5.9) ora ci dà

$$\frac{d}{n_x} = q (1 + \tan \alpha_x) (t_2 - t_1) \quad (16.5.12)$$

$$\frac{d}{n_y} = q (1 + \cot \alpha_x) (t_2 - t_1). \quad (16.5.13)$$

Scriviamo di nuovo $n_x = \cos \alpha_x$, $n_y = \cos \alpha_y = \sin \alpha_x = \sqrt{1 - \cos^2 \alpha_x}$.

La condizione di normalizzazione $n_x^2 + n_y^2 = 1$ e le identità (16.5.11) e (16.5.12) portano a

$$\begin{aligned} 1 &= d^2 d_0^2 \frac{1}{(t_2 - t_1)^2 (\cos \alpha_x + \sin \alpha_x)^2} \left(\frac{1}{(1 + \tan \alpha_x)^2} + \frac{1}{(1 + \cot \alpha_x)^2} \right) \\ &= d^2 d_0^2 \frac{1}{(t_2 - t_1)^2 \left(\frac{1}{\cos^2 \alpha_x} + \frac{1}{\sin^2 \alpha_x} \right)} = d^2 d_0^2 \frac{\cos^2 \alpha_x \sin^2 \alpha_x}{(t_2 - t_1)^2} \end{aligned}$$

da cui

$$d = \frac{t_2 - t_1}{d_0 \cos \alpha_x \sin \alpha_x}.$$

Abbiamo pertanto espresso d , d_1 , $\mathbf{c}$ e $\mathbf{n}$ in termini di h , α_x , d_0 , t_1 e t_2 . Se si mettono insieme le varie formule dimostrate sopra, calcoli elementari mostrano che la matrice della trasformazione prospettica del Teorema 16.4.4 in termini di questi parametri ora diventa la matrice M dell'enunciato. $\square$

NOTA 16.5.9. Nel caso in cui la scena da dipingere consista di oggetti con lati paralleli, come edifici a forma di cubo o parallelepipedo, abbiamo mostrato nella parte (i) dell'Esercizio 16.5.5 che il pittore non ha bisogno di calcolare la trasformazione prospettica per dipingere correttamente gli edifici in prospettiva. Invece, se la scena si compone di oggetti non a facce parallele, come persone od alberi, in teoria il pittore dovrebbe seguire la procedura illustrata nel precedente Esempio 16.5.8 per ricavare la trasformazione prospettica ed applicarla per disporre questi oggetti o persone sulla tela. Naturalmente, nessun pittore si cercherà mai di derivare la trasformazione prospettica. Il pittore invece dipinge immaginando di inscrivere gli oggetti entro opportuni cubi, ed in tal modo ritorna al caso precedente: fa slittare la proiezione dei cubi sulla tela verso i rispettivi punti di fuga e li utilizza come guide per aiutarsi nel dipingere gli oggetti nella prospettiva giusta. Questa tecnica è illustrata in Figura 17.

Però un elaboratore elettronico, invece, potrebbe utilizzare la trasformazione prospettica appena ricavata, non per la pittura, ma per il seguente altro scopo. Supponiamo che l'elaboratore riceva una immagine di un dipinto, o ancora meglio una fotografia di una scena vera (quindi non viziata da errori prospettici del pittore), e sia programmato per ricostruire, quanto più accuratamente possibile, la modellazione tridimensionale della scena reale. Questo significa invertire la matrice prospettica ed applicarla all'immagine. Purtroppo, però, la trasformazione prospettica manda lo spazio nel piano di visuale (quello della pellicola o del sensore fotografico, nel nostro caso), e quindi non è invertibile, in particolare non invertibile. Punti che giacciono sulla stessa semiretta uscente dal centro di proiezione vengono identificati dalla trasformazione prospettica: quindi per ogni punto del piano di visuale

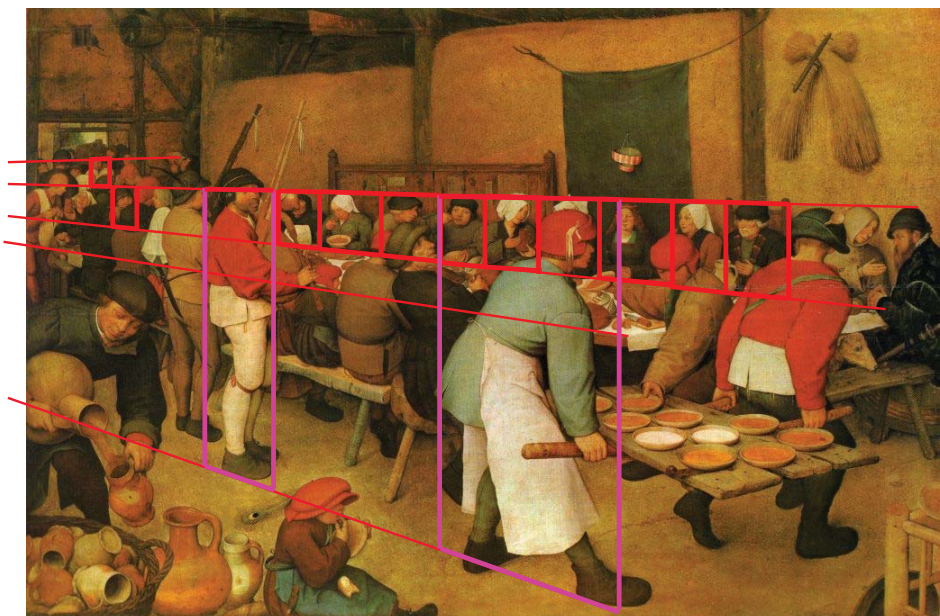


FIGURA 17. Nel caso si debba dipingere un gruppo di persone invece che oggetti a forma di cubo o parallelepipedo (come ad esempio edifici), una approssimazione consiste nell'inscatolare le persone allineate in scatole cubiche o parallelepipedi, immaginare di disegnare le scatole in prospettiva ed usarle come guide.

l'incertezza della ricostruzione ha un solo grado di libertà, dato dalla distanza in $\mathbb{R}^3$ dal centro di proiezione. Per ciascun oggetto, questa incertezza deve essere risolta da una scelta esplicita, che l'elaboratore può compiere solo tramite indizi basati su statistiche di scene reali, ma che una versione interattiva dell'applicativo può lasciar compiere all'operatore sulla base del buon senso e dell'immaginazione tridimensionale.

□

16.6. Prospettive parallele

Come sappiamo, le trasformazioni prospettiche sono applicazioni da $\mathbb{R}^3$ ad un piano di visuale $V \sim \mathbb{R}^2$. V è un piano in $\mathbb{R}^3$ che non passa necessariamente per l'origine, e quindi non costituisce in generale un sottospazio vettoriale.

Si chiamano *parallele* le trasformazioni prospettiche senza punti di fuga, ossia le proiezioni che conservano il parallelismo delle rette. A

differenza della prospettiva centrale, le prospettive parallele non comportano una divisione prospettica per la distanza, e quindi si possono descrivere in coordinate usuali tridimensionali, senza bisogno di coordinate omogenee quadridimensionali. Poiché V in generale non contiene l'origine, anche le prospettive parallele, come già quella centrale, non sono applicazioni lineari (non conservano l'origine!), però almeno sono trasformazioni affini, a differenza della prospettiva centrale. Quindi in generale esse sono rappresentate da matrici affini, ossia a dimensione quattro, tranne che nel caso in cui il piano V passi per l'origine, nel qual caso basta una matrice a dimensione tre (la corrispondente matrice affine avrebbe l'ultima colonna e l'ultima riga uguali a $(0, 0, 0, 1)$, e quindi componente di traslazione nulla: essa si identifica in modo naturale con il suo minore a dimensione tre relativo alle prime tre componenti).

Le proiezioni parallele si suddividono in ortogonali (o ortografiche) quando la direzione di proiezione coincide con il versore normale al piano di proiezione, ed oblique altrimenti. Il caso delle proiezioni ortogonali è assai semplice ed è stato già trattato nelle Sezioni 16.2 e 16.3.

16.6.1. Proiezione parallela. In questa e nelle prossime sottosezioni studiamo la *prospettiva parallela* in varie sue presentazioni. La presente sottosezione calcola la matrice di questa trasformazione prospettica in termini del vettore della direzione di proiezione: il risultato si chiama la matrice della *proiezione parallela*. Fissato un vettore di direzione $\mathbf{q}$, questa trasformazione prospettica manda il generico punto $\mathbf{p} = (x, y, z)$ nel punto $\mathbf{p}' = (x', y', z') \in V$ che appartiene alla retta che passa per $\mathbf{p}$ ed ha direzione $\mathbf{q}$. In altre parole,

$$\mathbf{p}' - \mathbf{p} = t\mathbf{q} \quad (16.6.1)$$

dove $t \in \mathbb{R}$ varia al variare di $\mathbf{p}$ (si noti che il suo segno è diverso nei due semispazi determinati dal piano V , e se $\mathbf{q}$ è normalizzato, come d'abitudine, allora $|t|$ è la distanza da $\mathbf{p}$ al punto proiettato $\mathbf{p}'$). Si noti anche che $|t|$ non è la distanza da $\mathbf{p}$ al piano V a meno che il vettore $\mathbf{q}$ sia normale a V : in tal caso si dice che la proiezione è *ortogonale*, ovvero

ortografica, e questa proiezione l'abbiamo già studiata nella Sezione [16.2](#).

I punti del piano di visuale V ovviamente rimangono fissi: per loro $t = 0$. Vediamo prima un esempio semplice.

ESEMPIO 16.6.1. (**Proiezione parallela sul piano di base.**) Scegliamo $V = \{(x, y, 0)\}$. Allora [\(16.6.1\)](#) diventa

$$x' - x = t q_x \quad y' - y = t q_y \quad z' - z = t q_z$$

ma poiché $z' = 0$ (dal momento che $\mathbf{p}'$ sta nel piano di base V) si ottiene

$$t = -\frac{z}{q_z} \quad x' = x - \frac{q_x}{q_z} z \quad y' = y - \frac{q_y}{q_z} z$$

e quindi la matrice prospettica è

$$M_{\mathbf{q}}^{\text{par}} = \begin{pmatrix} 1 & 0 & -\frac{q_x}{q_z} \\ 0 & 1 & -\frac{q_y}{q_z} \\ 0 & 0 & 0 \end{pmatrix}.$$

Il fatto che la terza riga sia tutta nulla è ovvio perché l'immagine della proiezione consiste di vettori con terza componente $z = 0$. In maniera naturale, questa proiezione si potrebbe quindi rappresentare come una matrice di due righe e tre colonne, da $\mathbb{R}^3$ al *sottospazio* $V \sim \mathbb{R}^2$ (in questo esempio V è un sottospazio vettoriale). $\square$

Ora studiamo il caso generale, in cui il piano V è generico. Sia $\mathbf{n} = (n_x, n_y, n_z)$ il suo versore normale (ovvero, sia $n_x x + n_y y + n_z z = d_0$ l'equazione di V : qui, come al solito, il significato geometrico di d_0 è la distanza di V dall'origine). Ora [\(16.6.1\)](#) diventa il sistema lineare

$$x' - x = t q_x$$

$$y' - y = t q_y$$

$$z' - z = t q_z$$

$$n_x x + n_y y + n_z z = d_0.$$

Calcoli elementari portano alla soluzione seguente:

PROPOSIZIONE 16.6.2. (**Proiezione parallela.**) *La proiezione parallela in direzione del vettore $\mathbf{q}$ sul piano V con versore normale $\mathbf{n} = (n_x, n_y, n_z)$ e distanza dall'origine d_0 è rappresentata dalla matrice*

$$M_{\mathbf{q},V}^{\text{par}} = \begin{pmatrix} d_1 - q_x n_x & -q_x n_y & -q_x n_z & q_x d_0 \\ -q_y n_x & d_1 - q_y n_y & -q_y n_z & q_y d_0 \\ -q_z n_x & -q_z n_y & d_1 - q_z n_z & q_z d_0 \\ 0 & 0 & 0 & d_1 \end{pmatrix},$$

dove $d_1 = \mathbf{q} \cdot \mathbf{n}$ è la componente del vettore direzionale della proiezione perpendicolare a V (se $\mathbf{q}$ è un versore, d_1 è il coseno dell'angolo formato da questo versore e la normale a V).

ESERCIZIO 16.6.3. Si ritrovi la matrice della Proposizione 16.6.2 riconducendo il calcolo all'Esempio 16.6.1.

SUGGERIMENTO. Si fissi un punto $\mathbf{r} = (x_0, y_0, z_0) \in V$, lo si trasli all'origine tramite la matrice affine di traslazione

$$T_{-\mathbf{r}} = \begin{pmatrix} 0 & 0 & 0 & -x_0 \\ 0 & 0 & 0 & -y_0 \\ 0 & 0 & 0 & -z_0 \\ 0 & 0 & 0 & 1 \end{pmatrix},$$

poi si applichi una matrice di rotazione $R = R_{\mathbf{n},\mathbf{e}_3}$ che porta il versore $\mathbf{n}$ in $(0, 0, 1)$ ossia che porta il piano traslato di V dopo la fase precedente di traslazione (che è il piano parallelo a V che passa per l'origine) nel piano $\{z = 0\}$, poi si applichi la matrice $M_{\mathbf{q}'}^{\text{par}}$ dell'Esempio 16.6.1, dove però ora si deve usare il vettore direzionale ruotato $\mathbf{q}' = R\mathbf{q}$, e poi si ritorni indietro: la matrice coniugata così ottenuta è la matrice della proiezione centrale richiesta:

$$T_{-\mathbf{r}}^{-1} R_{\mathbf{n},\mathbf{e}_3}^{-1} M_{\mathbf{q}'}^{\text{par}} R_{\mathbf{n},\mathbf{e}_3} T_{-\mathbf{r}} = M_{\mathbf{q},V}^{\text{par}}.$$

L'unica difficoltà consiste nel ricavare la matrice di rotazione R . Come matrice affine, R è in realtà lineare, e quindi la sua ultima riga e colonna coincidono con $(0, 0, 0, 1)$. L'inverso R^{-1} di R è uguale al suo trasposto R^* , perché R è una matrice ortonormale. D'altra parte $R^* \mathbf{e}_3 = R^{-1} \mathbf{e}_3 = \mathbf{n}$, e quindi la terza riga di R ha come prime tre componenti quelle di $\mathbf{n}$. Abbiamo un grado di libertà nella scelta di R , perché

se prima di far agire R ruotiamo tutto lo spazio intorno all'asse di rotazione $\mathbf{n}$ la trasformazione R richiesta non ne risente. Quindi possiamo scegliere la prima componente della seconda riga di R uguale a 0, ed allora, per l'ortogonalità con la terza riga, si ottiene che le prime tre componenti della seconda riga di R devono essere $(0, n_z, -n_y)$ (a parte il segno, che è irrilevante, ma che comunque sceglieremo, nel prossimo passaggio, in modo che R abbia determinante $+1$ invece che -1 : ma si troverebbe la soluzione anche con la scelta opposta). Occorre però normalizzare la seconda riga della matrice, perché le righe di una matrice ortonormale sono una base ortonormale. A questo punto, per l'ortonormalità, la prima riga è il prodotto vettoriale fra la seconda e la terza. Si verifichi che in tal modo si trova la seguente forma matriciale di R :

$$R = \begin{pmatrix} 1 & \frac{-n_x n_y}{\sqrt{1+n_x^2}\sqrt{n_y^2+n_z^2}} & \frac{-n_x n_z}{\sqrt{1+n_x^2}\sqrt{n_y^2+n_z^2}} & 0 \\ 0 & \frac{n_z}{\sqrt{n_y^2+n_z^2}} & \frac{-n_y}{\sqrt{n_y^2+n_z^2}} & 0 \\ n_x & n_y & n_z & 0 \\ 0 & 0 & 0 & 1 \end{pmatrix}.$$

Si noti che abbiamo scritto R come una matrice a dimensione quattro, ma in realtà la rotazione è in tre dimensioni. Però R deve essere moltiplicata per matrici di trasformazioni affini che mandano il piano di base in un piano che non passa per l'origine, e quindi non è un sottospazio bidimensionale di $\mathbb{R}^3$. Pertanto queste trasformazioni affini hanno matrici a dimensione quattro, e per questo dobbiamo immergere la matrice di R in una matrice a dimensione quattro, prendendo la sua somma diretta con la matrice identità a dimensione 1. $\square$

16.6.2. Proiezione obliqua. La matrice della *proiezione obliqua* è un'altra presentazione della trasformazione prospettica parallela, questa volta espressa non in termini del vettore direzionale bensì di un *fattore di accorciamento* e di un *angolo di azimuth*.

Supponiamo che il piano di proiezione sia il piano $\{z = 0\}$, e che il punto $\mathbf{e}_3 = (0, 0, 1)$ sia proiettato sul punto $\mathbf{p} = (x', y', 0)$. Se $f = \sqrt{x'^2 + y'^2}$ e θ è l'angolo fra il semiasse x positivo ed il vettore $\mathbf{p}$, allora $x' = f \cos \theta$ e $y' = f \sin \theta$. Il fattore f misura l'accorciamento

prospettico dell'asse della profondità (l'asse perpendicolare al piano di proiezione).

NOTA 16.6.4. *Sebbene di solito il fattore di accorciamento f sia minore di 1, e quindi rappresenti un vero accorciamento dell'asse della profondità, possiamo avere prospettive con $f > 1$, in cui quindi abbiamo un allungamento. Ne vedremo un esempio in seguito nella Sottosezione 16.7.2. Naturalmente, l'allungamento dell'asse delle profondità provoca una distorsione prospettica, particolarmente evidente se f è grande.*

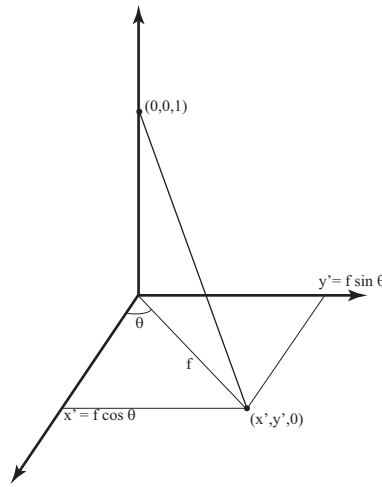


FIGURA 18. Proiezione obliqua sul piano $\{z = 0\}$: fattore di accorciamento f e azimuth θ .

Pertanto il vettore della direzione di proiezione è

$$\mathbf{q} = \mathbf{p} - \mathbf{e}_3 = (f \cos \theta, f \sin \theta, -1),$$

e dall'Esempio 16.6.1 troviamo la forma seguente della matrice di proiezione:

$$M_{f,\theta}^{\text{par}} = M_{\mathbf{q}}^{\text{par}} = \begin{pmatrix} 1 & 0 & f \cos \theta \\ 0 & 1 & f \sin \theta \\ 0 & 0 & 0 \end{pmatrix}. \quad (16.6.2)$$

Se si preferisce, si può immergere questa matrice in una a quattro dimensioni, come abbiamo fatto prima nell'Esercizio 16.6.3:

$$M_{f,\theta}^{\text{par}} = \begin{pmatrix} 1 & 0 & f \cos \theta & 0 \\ 0 & 1 & f \sin \theta & 0 \\ 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 1 \end{pmatrix}.$$

La trasformazione diventa

$$\begin{aligned} x' &= x + zf \cos \theta \\ y' &= y + zf \sin \theta. \end{aligned} \tag{16.6.3}$$

Da questo caso particolare si passa al caso generale di piano di proiezione arbitrario tramite coniugazione con la matrice della trasformazione affine che mappa tale piano sul piano $\{z = 0\}$, come accennato nell'Esercizio 16.6.3. Come visto in quell'Esercizio, quando il piano di proiezione non passa per l'origine, e quindi non è un sottospazio vettoriale, questa matrice affine è una vera matrice a dimensione quattro (non è l'immersione quadridimensionale di una matrice a dimensione tre), e per questo è opportuno immergere R in quattro dimensioni.

In realtà, procedendo in base a semplici illustrazioni geometriche, possiamo talvolta evitare il procedimento di moltiplicazione matriciale, algebricamente accurato ma geometricamente piuttosto oscuro. Osserviamo anzitutto che la trasformazione prospettica data dalla matrice 16.6.2 cambia aspetto a seconda del piano di proiezione. Anche nel caso in cui questo piano passi per l'origine, ovvero sia un sottospazio vettoriale, la forma della trasformazione dipende dalla scelta di una base in questo sottospazio. In generale, se il piano di proiezione è un sottospazio affine (ossia non passa per l'origine di $\mathbb{R}^3$), la forma della trasformazione, rispetto ad una sua base affine, dipende dalla scelta di tale base.

Ad esempio, consideriamo il caso del piano di proiezione $\{z = -1\}$. Rispetto a prima, dobbiamo traslare da $z = 0$ a $z = -1$. Ma consideriamo adesso la trasformazione prospettica come mappa dalle coordinate canoniche di $\mathbb{R}^3$ ad una scelta di coordinate intrinseche sul piano di

proiezione: in tal modo è inutile considerare la traslazione nella variabile z , la quale individua solo la traslazione dell'origine di $\mathbb{R}^3$ alla nuova origine $(0, 0, -1)$ del piano affine $\{z = -1\}$. Pertanto, in tal modo ci riconduciamo sempre alla situazione in cui il piano di proiezione è $\{z = 0\}$: assumiamo dunque d'ora in poi che questo sia il piano di proiezione. (Se vogliamo ritornare al caso della proiezione su un piano V in $\mathbb{R}^3$ con le coordinate di V indotte dalle coordinate canoniche di $\mathbb{R}^3$, basta comporre con una trasformazione affine da V al piano di base $B = \{z = 0\}$, ovvero con l'inverso di una trasformazione affine da B a V : ad esempio, nel caso di cui sopra, $V = \{z = -1\}$, e una trasformazione affine ovvia è $z \mapsto z - 1$, il cui inverso è $z \mapsto z + 1$: quindi, in tutte le prossime formule basta sostituire z con $z + 1$. Osserviamo che la scelta di questa trasformazione affine ammette un solo grado di libertà: data una trasformazione affine da V a B , il piano di arrivo B resta lo stesso solo se dopo la trasformazione affine operiamo una rotazione intorno al versore normale di B , ossia l'asse z . Quindi possiamo supporre, senza perdita di generalità, che la proiezione del punto $\mathbf{e}_3 = (0, 0, 1)$ giaccia sull'asse x , ossia si abbia $\theta = 0$ nella matrice (16.6.2).

Scriviamo quindi la trasformazione relativamente alla scelta di base, in questo piano, concorde con la restrizione della base canonica di $\mathbb{R}^3$, ossia formata dai versori $\mathbf{e}_1$ dell'asse x e $\mathbf{e}_2$ dell'asse y : in generale indichiamo con $\mathbf{i}$ e $\mathbf{j}$ i due versori di base nel sottospazio di proiezione (o i vettori della base affine se, come nel caso presente, il piano di proiezione non passa per l'origine), e chiamiamo x' e y' i corrispondenti assi coordinati in tale piano. Con la scelta naturale di base nel piano di proiezione data come sopra dalla restrizione, nelle nuove coordinate x' e y' la trasformazione ritorna ad essere quella data dalle identità (16.6.3).

Per comodità, scambiamo fra loro gli assi y e z in $\mathbb{R}^3$: l'asse delle profondità è ora l'asse y , il piano di proiezione è $\{y = 0\}$ e la trasformazione diventa

$$x' = x + yf \cos \theta \quad (16.6.4)$$

$$y' = yf \sin \theta + z. \quad (16.6.5)$$

Ora operiamo un'altra scelta di base, in maniera tale che gli assi x e y siano visti, in prospettiva, come nella Figura 19 il semiasse y' positivo coincide nella visualizzazione con il semiasse z positivo, il semiasse y positivo forma un angolo α (con $0 < \alpha < \frac{\pi}{2}$) con il semiasse x' positivo ed è orientato nel verso positivo di x' , ed infine il semiasse x positivo forma un angolo β (con $0 < \beta < \frac{\pi}{2}$) con il semiasse x' negativo (e quindi è orientato nel verso *negativo* di x'). Questa è la trasformazione obliqua più naturale per rappresentare il grafico di una funzione definita su un quadrato centrato nell'origine in maniera che la vista sia non frontale bensì ad angolo. Adesso, per tener conto del fatto che il piano di proiezione può essere traslato, e quindi allontanato dall'origine in $\mathbb{R}^3$, è opportuno utilizzare non più un solo fattore di accorciamento f , ma tre, per descrivere gli accorciamenti dei versori di tutti e tre gli assi: indichiamo questi fattori con f_x , f_y e f_z (l'ultimo è quello che prima abbiamo chiamato f): i loro valori sono le lunghezze dei trasformati prospettici dei versori della base canonica $\mathbf{e}_1$, $\mathbf{e}_2$, $\mathbf{e}_3$, rispettivamente. Ovviamente, per ragioni di simmetria, sarebbe ragionevole porre $f_x = f_y$ (la scelta di f_z può essere diversa per via della proiezione obliqua, che privilegia l'asse z). In effetti, in seguito useremo questa prospettiva per la rimozione di linee nascoste in grafici prospettici tridimensionali, e nella versione più realistica porremo $f_x = f_y$ (Nota 16.7.1 e Sottosezione 16.7.3), ma avremo bisogno di disegnare il grafico con fattori di scala diversi sui due assi x e y ; se i fattori di scala sono scelti uguali, allora trarremo vantaggio della possibilità di scegliere $f_x \neq f_y$ (si veda l'algoritmo che conduce all'identità (16.7.6)). In seguito alla rotazione e traslazione, e quindi alla scelta dei due angoli di inclinazione prospettica ed ai tre fattori di accorciamento, la trasformazione ora diventa:

$$x' = -xf_x \cos \beta + yf_y \cos \alpha \quad (16.6.6)$$

$$y' = -xf_x \sin \beta - yf_y \sin \alpha + zf_z \quad (16.6.7)$$

(si osservi che i segni meno corrispondono al fatto che, al crescere di x , x' e y' decrescono, ed al crescere di y si ha che x' decresce ma y'

cresce, per come la scelta dei versi di visualizzazione degli assi (si veda la Figura 19).

NOTA 16.6.5. *La scelta di usare f_x e f_y diversi può essere molto conveniente in Computer Graphics, dove spesso occorre adattare le dimensioni orizzontali e verticali sulla finestra di visuale per adattare l'output alle proporzioni dei monitor, che non sono di solito quadrate, bensì in rapporto 4:3 o 16:9. In tal caso, si può scrivere il codice grafico in maniera che, alla scelta delle proporzioni del monitor, corrisponda una scelta analoga del rapporto $f_x : f_y$: in tal modo l'output si riscalda nel passaggio da un monitor ad un altro (ad esempio, può riempire esattamente entrambi i monitor, senza tagli).*

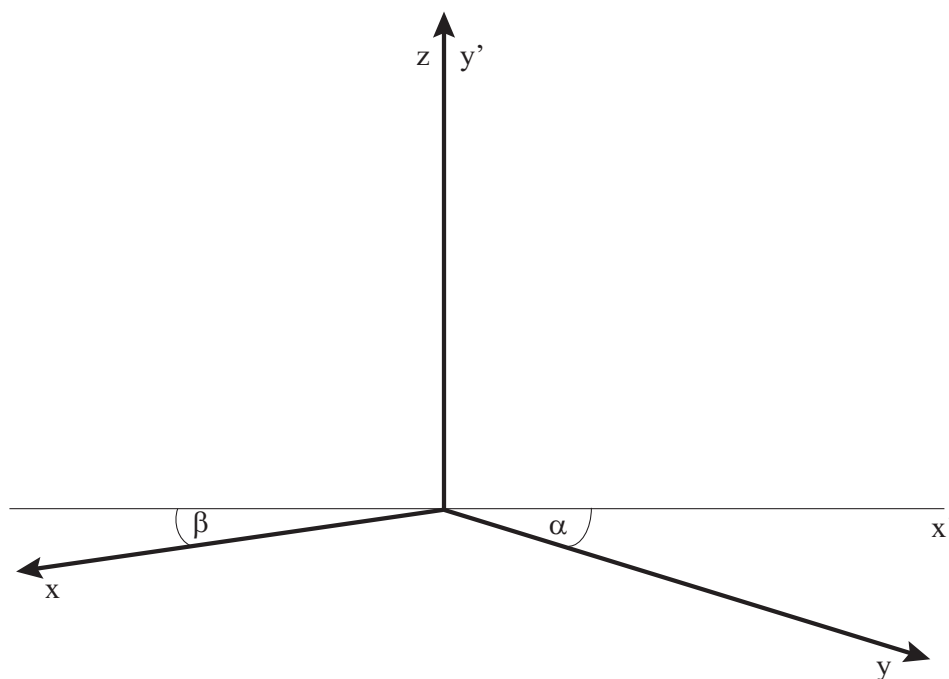


FIGURA 19. Visualizzazione in proiezione obliqua con vista non frontale bensì angolare.

Proiezioni cavaliera e cabinet. Il caso della proiezione obliqua senza accorciamento, ossia con fattore di accorciamento $f = 1$, si chiama

proiezione cavaliera. Sappiamo da (16.6.2) che la forma tridimensionale della sua matrice è

$$M_{1,\theta}^{\text{par}} = \begin{pmatrix} 1 & 0 & \cos \theta \\ 0 & 1 & \sin \theta \\ 0 & 0 & 0 \end{pmatrix}. \quad (16.6.8)$$

Questa proiezione è spesso usata nel disegno tecnico per scopi militari, poiché la sua immagine fornisce la mappa prospettica di una zona vista dall'alto senza distorsioni di accorciamento. Se invece proiettiamo sul piano $\{x = 0\}$ oppure $\{y = 0\}$ otteniamo invece le proiezioni laterali o frontali senza accorciamenti, anche esse utili nel campo militare o architettonico e nel disegno tecnico, le cui matrici si ottengono dalla precedente tramite coniugazione per la matrice ortogonale del cambiamento di base che manda $\mathbf{e}_3$ in $\mathbf{e}_1$, $\mathbf{e}_1$ in $\mathbf{e}_2$ e $\mathbf{e}_2$ in $\mathbf{e}_3$, ossia

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

oppure, rispettivamente, la sua trasposta: in tal modo si vede che la proiezione cavaliera laterale (sul piano $\{x = 0\}$) ha per matrice

$$\begin{pmatrix} 0 & 0 & 0 \\ \cos \theta & 1 & 0 \\ \sin \theta & 0 & 1 \end{pmatrix}$$

e quella frontale (sul piano $\{y = 0\}$) ha per matrice

$$\begin{pmatrix} 1 & \cos \theta & 0 \\ 0 & 0 & 0 \\ 0 & \sin \theta & 1 \end{pmatrix}. \quad (16.6.9)$$

Un'altra scelta tipica di fattore di accorciamento, usata spesso nel disegno tecnico, è $f = \frac{1}{2}$. La proiezione corrispondente si chiama *proiezione cabinet*. Anche in tal caso possiamo trascrivere la matrice corrispondente, ad esempio sul piano di base $\{z = 0\}$:

$$M_{\frac{1}{2},\theta}^{\text{par}} = \begin{pmatrix} 1 & 0 & \frac{1}{2} \cos \theta \\ 0 & 1 & \frac{1}{2} \sin \theta \\ 0 & 0 & 0 \end{pmatrix}.$$

Le scelte tipiche di azimuth per un disegno realistico sono $\theta = \frac{\pi}{4}$ e $\theta = \frac{\pi}{6}$. La prima porta alla matrice di proiezione

$$M_{\frac{1}{2}, \frac{\pi}{4}}^{\text{par}} = \begin{pmatrix} 1 & 0 & \frac{\sqrt{2}}{2} \\ 0 & 1 & \frac{\sqrt{2}}{2} \\ 0 & 0 & 0 \end{pmatrix}, \quad (16.6.10)$$

la seconda alla matrice

$$M_{\frac{1}{2}, \frac{\pi}{4}}^{\text{par}} = \begin{pmatrix} 1 & 0 & \frac{\sqrt{3}}{4} \\ 0 & 1 & \frac{1}{4} \\ 0 & 0 & 0 \end{pmatrix}. \quad (16.6.11)$$

16.6.3. Vari tipi di proiezioni ortogonali. Rammentiamo che le proiezioni ortogonali sono le proiezioni parallele la cui direzione di proiezione è perpendicolare al piano di proiezione. Nella Sezione 16.2 abbiamo trattato il caso in cui la direzione di proiezione sia parallela all'asse z , o ad uno degli assi coordinati. In generale, si dice che la proiezione è

- *isometrica* se la direzione di proiezione determina angoli uguali con tutti e tre gli assi coordinati;
- *dimetrica* se la direzione di proiezione forma angoli uguali con due assi coordinati;
- *trimetrica* se i tre angoli formati dalla direzione di proiezione con gli assi coordinati sono tutti diversi.

Nel caso isometrico, il vettore normale al piano di proiezione è proporzionale a $(\pm 1, \pm 1, \pm 1)$: ci sono otto direzioni possibili, quelle delle bisettrici degli otto ottanti (se non si ha interesse ad eliminare linee nascoste, conta solo la direzione di proiezione e non il verso, e quindi le direzioni sono a due a due equivalenti: rimangono solo quattro casi).

16.7. Applicazione: rimozione di linee nascoste in grafici 3D

Prendiamo in esame un problema tipico della Computer Graphics, quello del disegno del grafico di una funzione F di due variabili, che per il momento, seguendo la notazione usuale in matematica, indichiamo con x e y . Vogliamo proiettare i punti $(x, y, z = F(x, y))$

sul piano frontale $\{y = 0\}$. Supponiamo per semplicità che la funzione sia stata calcolata sui punti (x_j, y_m) di una griglia di passo τ , con $j, m = 1, \dots, N$. Per semplicità poniamo $x_j = j\tau$, $y_m = m\tau$. Osserviamo che in questo modo i valori y_m di profondità della griglia sono tutti positivi, e quindi dallo stesso lato del piano di proiezione $\{y = 0\}$: è opportuno che sia così affinché il disegno sia realistico. Poniamo $z_{jm} = (x_j, y_m, F(x_j, y_m))$. Vogliamo proiettare prospetticamente i punti $\mathbf{p}_{jm} = (x_j, y_m, z_{jm})$ sul piano $\{y = 0\}$ e poi unire i punti proiettati con segmenti, in modo che la proiezione prospettica di ciascun punto $\mathbf{p}_{jm}$ sia unita a quella del suo predecessore e del suo successore sulla stessa linea $x = \text{costante}$ (ossia $\mathbf{p}_{j\pm 1, m}$) e $y = \text{costante}$ (ossia $\mathbf{p}_{j, m\pm 1}$). In questo modo otteniamo una famiglia a due parametri di spezzate nel piano di proiezione, che chiamiamo le *poligonal* $x = \text{costante}$ e $y = \text{costante}$. Esse producono un disegno a maglie del grafico, fatto con una rete poligonale, la quale però, se i punti della griglia sono vicini rispetto al tasso di variazione di F (ossia se la griglia è fitta: τ piccolo), approssima bene un disegno a curve lisce. Nella Figura 20 illustriamo il grafico a maglie della funzione $\sin(x^2 + y^2)$ su un quadrato centrato nell'origine.

Il problema però è che questo disegno a maglie è trasparente: ci si vede in mezzo, le parti posteriori del grafico risultano visibili attraverso le maglie anteriori. Dovremmo quindi rimuovere le zone posteriori nascoste. Se lo facciamo, si ottiene il grafico in Figura 21. Ma come si deve procedere per eliminare le linee nascoste?

16.7.1. Rimozione di linee prospetticamente nascoste. Proviamo a trovare un algoritmo per rimuovere le parti nascoste. Per semplicità, restringiamo per ora l'attenzione alle poligonal $x = \text{costante}$, ossia quelle con punti nodali $\mathbf{p}_{jm}$, $j = 0, \dots, N$, dove $0 \leq m \leq N$ è fissato. Osserviamo che la loro profondità in tre dimensioni varia proporzionalmente a m : infatti la m -sima poligonale $x = \text{costante}$ giace in un piano, questi piani sono paralleli ed all'aumentare di m si allontanano dal piano di proiezione (nel caso non si sia scelto di far aumentare la profondità all'aumentare di m , occorre ora cominciare il

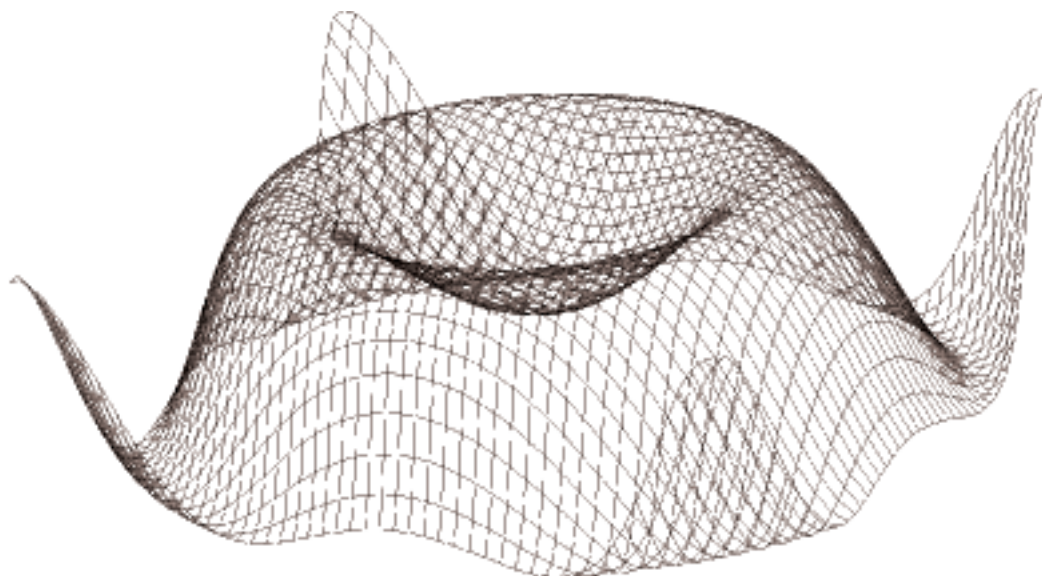


FIGURA 20. Grafico a maglie di $\sin(x^2 + y^2)$ senza rimozione delle linee nascoste.

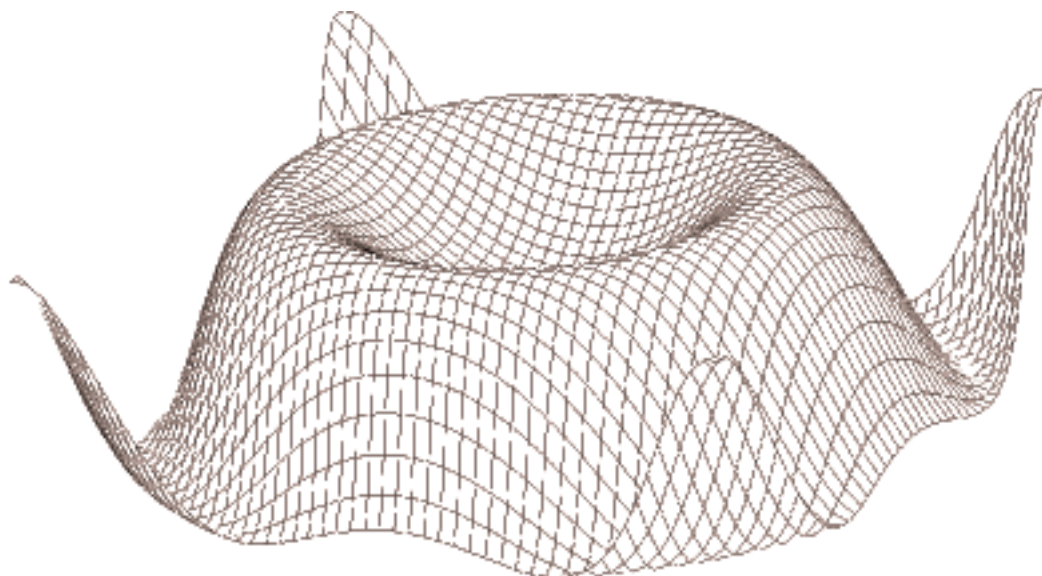


FIGURA 21. Grafico a maglie di $\sin(x^2 + y^2)$ con rimozione delle linee nascoste.

disegno a partire da $m = 0$ oppure da $m = N$ a seconda di quale delle due poligonali, ossia di quale dei quattro vertici del rettangolo della griglia, sia più vicino al punto di visuale).

La prima di queste poligonali, per $m = 0$, viene proiettata sul piano di proiezione ed interamente disegnata: è la più frontale, nulla può nascerla. Anche la seconda viene interamente disegnata: non può infatti essere nascosta dalla prima, può certo accavallarsi e passarle da sopra a sotto o viceversa, ma in entrambi i casi, o da sopra o da sotto, è visibile. Si ottiene così lo schizzo della Figura 22.

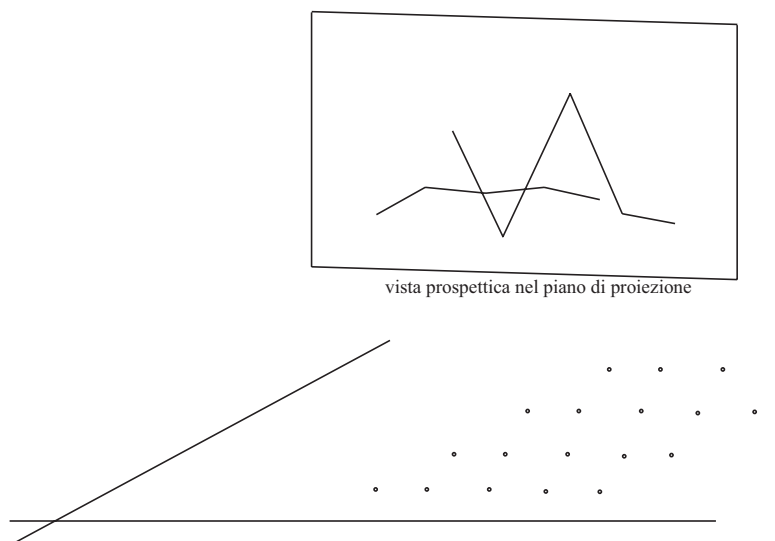


FIGURA 22. Le prime due spezzate $x = \text{costante}$ in un grafico 3D prospettico.

Ma a partire dalla terza poligonale le cose cambiano. Infatti, le prime due ora racchiudono un'area opaca nel piano di proiezione, che nasconde la zona retrostante, e quindi attraverso la quale non si deve vedere dietro. L'area può avere l'aspetto di una striscia se le prime due poligonali non si accavallano, ed altrimenti ha una o più strozzature. I segmenti proiettati della terza poligonale vengono rimossi per la parte che penetra dentro tale area. La terza poligonale si aggiunge alle prime due ed aumenta l'area di opacità, ossia di invisibilità. In effetti, dopo aver tracciato m poligonali, tale area è delimitata dai *profili massimo e minimo* che il disegno forma fino a quel momento sul piano di proiezione. In Figura 23 schizziamo il risultato del tracciamento della terza e quarta poligonale.

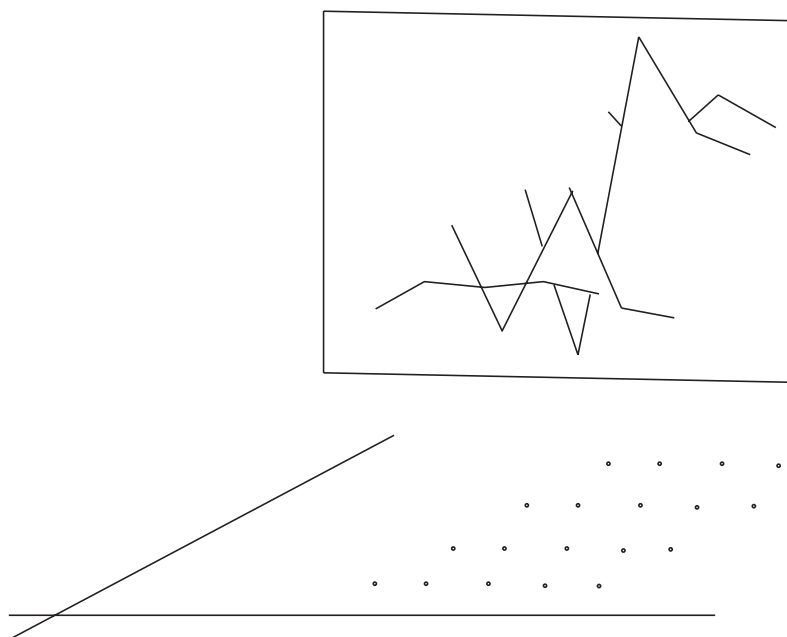


FIGURA 23. Le prime quattro spezzate $x = \text{costante}$: a partire dalla terza, esse non sono visibili nell'area fra i profili massimo e minimo precedenti.

Nel caso generale nel quale si intenda tracciare il grafico a maglia completo, quindi sia le poligonali $x = \text{costante}$ sia quelle $y = \text{costante}$, per tener conto del fatto che la visibilità di segmenti delle poligonali del primo tipo potrebbe essere affetta dall'andamento di quelle del secondo tipo e viceversa, modifichiamo l'ordine di tracciamento in modo che le due famiglie di poligonali siano simultaneamente disegnate all'aumentare della profondità. Ovviamente, al fine di garantire il corretto aggiornamento dei profili all'aumentare della profondità, è opportuno, in teoria, cominciare dalla poligonale più frontale: prima poligonale $x = \text{costante}$ oppure $y = \text{costante}$ a seconda di quale fra i lati esterni x e y del rettangolo della griglia sia più vicino al punto di osservazione. In realtà, però, questa cautela è inessenziale nel nostro caso: pur di procedere incrementando la distanza dall'osservatore, qualunque dei due ordini nella scelta fra $x = \text{costante}$ e $y = \text{costante}$ produce lo stesso risultato. Infatti le prime due poligonali $x = \text{costante}$ e $y = \text{costante}$ si posizionano su aree del piano di visuale non sovrapposte verticalmente,

in quanto questo è ciò che accade per la griglia di base (si veda la Figura 19) e la posizione dei nodi delle poligonali differisce da quello dei nodi della griglia solo per spostamenti verticali, in base alla forma della trasformazione prospettica (16.6.6). Da qui segue che, nel disegnare ciascuna maglia rettangolare (o meglio a quadrilatero) del disegno, non può succedere che, una volta tracciati i segmenti corrispondenti al lato di griglia più frontale diciamo $x = \text{costante}$ e quello più frontale di tipo $y = \text{costante}$, il segmento relativo al lato posteriore $y = \text{costante}$ venga indebitamente nascosto dai profili massimo e minimo stabiliti nel tracciamento dei segmenti precedenti; analoga impossibilità si ha se si traccia prima il segmento $y = \text{costante}$ e poi i due segmenti limitrofi $x = \text{costante}$. Lasciamo al lettore il tracciare opportuni disegni per visualizzare questo fatto.

A parte la determinazione di quale sia il lato più vicino, è invece essenziale tracciare il grafico a partire dal suo vertice più vicino al punto di visuale (poiché sia i vertici della griglia sia il punto di visuale sono fissati prima dell'inizio della procedura di disegno, queste scelte sono semplici da effettuare; in questa presentazione abbiamo scelto gli indici della griglia in maniera da cominciare con il vertice di indici massimi).

Il procedimento è il seguente. Invece di disegnare consecutivamente i segmenti corrispondenti ai tratti relativi ai punti nodali $\mathbf{p}_{jm}$, $j = 0, \dots, N$, con $0 \leq m \leq N$ è fissato (ossia la m -sima poligonale $x = \text{costante}$) e poi quelli del tipo $y = \text{costante}$, tracciamo invece il primo segmento di ciascuna poligonale $x = \text{costante}$, ossia il trasformato prospettico del segmento da $\mathbf{p}_{0m}$ a $\mathbf{p}_{1m}$, poi i trasformati dei due segmenti delle poligonali $y = \text{costante}$ che vi si appoggiano, ossia i segmenti da $\mathbf{p}_{0m}$ a $\mathbf{p}_{0,m+1}$ e da $\mathbf{p}_{1m}$ a $\mathbf{p}_{1,m+1}$, e così via a maglie adiacenti. Durante il disegno, il gruppo di segmenti j -simo è da $\mathbf{p}_{jm}$ a $\mathbf{p}_{j+1,m}$ seguito dai due segmenti trasversali $\mathbf{p}_{jm}$ a $\mathbf{p}_{j,m+1}$ e da $\mathbf{p}_{j+1,m}$ a $\mathbf{p}_{j+1,m+1}$: di ciascuno di questi segmenti tridimensionali disegniamo, in questo ordine, i trasformati prospettici. Ma ovviamente non occorre tracciare due volte lo stesso segmento, quindi per $j > 0$ basta tracciare

i segmenti nell'ordine

$$\mathbf{P}_{jm} \rightarrow \mathbf{P}_{j+1,m} \quad (16.7.1)$$

$$\mathbf{P}_{j,m} \rightarrow \mathbf{P}_{j,m+1} . \quad (16.7.2)$$

In questo modo i profili minimo e massimo vengono aggiornati correttamente al crescere della profondità.

Il problema ora consiste nel rimuovere dal disegno, ossia evitare di tracciare, quelle parti dei segmenti delle poligonaliche penetrano nella striscia di opacità, ossia le cui altezze sul piano di proiezione sono comprese fra i profili minimo e massimo. In linea di principio, il procedimento è ovvio: basta tracciare i segmenti prospettici sul piano di proiezione un pixel per volta e, prima di tracciare, verificare se il pixel verifica la condizione di invisibilità (altezza compresa fra i profili minimo e massimo): se è così si omette il disegno, altrimenti si disegna. Naturalmente occorre ricavare, per ogni colonna di pixel, le coordinate del pixel di intersezione fra il segmento e quella colonna; la discretizzazione del segmento ai singoli pixel può provocare errori di arrotondamento, ma limitatamente ad uno scarto di un pixel.

Purtroppo in una tipica immagine ci sono almeno un migliaio di pixel in ascissa, ossia un migliaio di colonne. Anche considerando grafici a maglie di bassa risoluzione (una quarantina di linee $x = \text{costante}$), abbiamo bisogno di invocare la routine grafica di tracciamento delle linee circa 40.000 volte per un pixel alla volta, anche se stiamo solo cercando di disegnare 1600 segmenti. Questo accade per ciascun singolo disegno, con un aggravio piuttosto oneroso di elaborazione grafica.

È opportuno procedere diversamente. Osserviamo che il tracciamento di una linea sul viewport discretizzato (composto di singoli pixel) viene effettuato tramite un algoritmo incrementale (algoritmo di Bresenham) che a ciascun passo decide se accendere il pixel successivo orizzontalmente o verticalmente a seconda della pendenza della linea e di quanto fatto al passo precedente: si veda un libro di testo di Computer Graphics, ad esempio [10] (E' quindi sufficiente modificare l'algoritmo di Bresenham in modo che, oltre alla scelta se accendere il pixel in direzione orizzontale o verticale, si faccia anche la scelta se accenderlo o no a seconda del fatto che la condizione di visibilità sia soddisfatta.

Questo procedimento funziona quale che sia la prospettiva utilizzata: però occorre riscrivere la routine della libreria grafica responsabile del tracciamento delle linee, e quindi ricompilarla. Se si preferisce evitare questo fastidio, si può ricorrere ad un diverso algoritmo, che presentiamo qui di seguito, basato su un allineamento intelligente dei punti della griglia di base e l'impiego di una prospettiva appropriata.

16.7.2. Algoritmo veloce di rimozione di linee nascoste in grafici 3D: allineamento della griglia. Per un algoritmo veloce di rimozione delle linee nascoste impieghiamo la proiezione obliqua introdotta nella Sottosezione [16.6.2](#).

Rammentiamo che abbiamo denotato con $z_{jm} = (x_j, y_m)$ i nodi della griglia di base e con $\mathbf{p}_{jm} = (x_j, y_m, z_{jm})$ i punti nodali del grafico a maglie della funzione F su questa griglia. Chiamiamo $\mathbf{r}_{jm}$ la proiezione prospettica di $\mathbf{p}_{jm}$ sul piano $\{y = 0\}$ e $\rho_{(ij),(i+1,j)}$ (rispettivamente, $\rho_{(ij),(i,j+1)}$) il segmento su questo piano che congiunge $\mathbf{r}_{ij}$ con $\mathbf{r}_{i+1,j}$ (rispettivamente, con $\mathbf{r}_{i,j+1}$).

Il problema, come abbiamo visto, consiste nel fatto che i segmenti $\rho_{(ij),(i+1,j)}$ e $\rho_{(ij),(i,j+1)}$ devono essere disegnati solo per le parti che non penetrano nell'area del piano di visuale compresa fra i profili minimo e massimo. Se questo controllo viene eseguito prima di disegnare ciascun pixel il procedimento risulta accurato ma lento. L'algoritmo che proponiamo si basa sull'impiego della prospettiva obliqua e sullo scegliere i suoi parametri in modo che i punti iniziali e finali di questi segmenti siano allineati verticalmente. Questo è possibile perché, in base alle regole di trasformazione ([16.6.4](#)) e ([16.6.6](#)), il valore della coordinata z , ossia il valore della funzione di cui tracciamo il grafico, non influenza il valore dell'ascissa x' sul piano di visuale, e quindi i punti estremi dei segmenti $\rho_{(ij),(i+1,j)}$ e $\rho_{(ij),(i,j+1)}$ sono allineati verticalmente se e solo se lo sono i trasformati prospettici dei punti nodali della griglia di base.

In tal modo, occorre solo verificare la condizione di non appartenenza all'area di invisibilità all'inizio ed alla fine del segmento che stiamo disegnando: se entrambi i punti estremi sono visibili tracciamo l'intero segmento, altrimenti la linearità ci permette di calcolare immediatamente quale parte tracciare. È quindi necessario memorizzare soltanto

le altezze dei profili minimo e massimo in corrispondenza delle colonne di pixel che corrispondono ai punti estremi dei segmenti $\rho_{(ij),(i+1,j)}$ e $\rho_{(ij),(i,j+1)}$, ossia, come appena osservato, i trasformati dei punti nodali della griglia di base.

Qui però dobbiamo affrontare un problema. Nel caso un segmento debba essere tracciato solo in parte, diciamo nella parte iniziale, il punto terminale di questa parte tracciata diventa un nuovo punto per il quale occorre memorizzare i valori dei profili minimo e massimo, in aggiunta a quelli provenienti dalla griglia di base. Ma in questo modo si perderebbe l'allineamento verticale, e la dimensione degli array in cui memorizzare i valori di tali profili potrebbe raddoppiarsi ad ogni passo: questa crescita esponenziale distruggerebbe i vantaggi di velocità di calcolo dell'algoritmo. Conviene pertanto omettere la memorizzazione di questi valori intermedi dei profili minimo e massimo. In tal caso, all'ascissa del punto intermedio il profilo massimo viene sovrastimato e quello minimo sottostimato. Come conseguenza, il resto del disegno può omettere qualche piccolo tratto, una imprecisione tollerabile se la risoluzione è elevata, ossia se il passo di incremento di griglia τ è piccolo.

Naturalmente, all'inizio del procedimento bisogna determinare se cominciare dalla prima poligonale $x = \text{costante}$ oppure $y = \text{costante}$. Rammentiamo che si deve cominciare da quella delle due più vicina al punto di osservazione: pertanto, dall'asse che ha angolo di inclinazione minore. Questa semplice scelta è necessaria solo se il punto di vista è angolato, e quindi ci sono due angoli di inclinazione non nulli: nel caso di punto di vista frontale, l'angolo dell'asse frontale è zero e si comincia con questo.

Aggiorniamo la notazione. Abbiamo chiamato τ l'incremento delle coordinate x e y fra due consecutive rette $x = \text{costante}$ e $y = \text{costante}$ della griglia di base. Denotiamo con Δx l'incremento delle ascisse sul piano di visuale (ossia della variabile x') ottenuto dopo la trasformazione prospettica quando si passa da una retta $x = \text{costante}$ della griglia di base alla successiva (ossia da (x_j, y_m) a (x_{j+1}, y_m)) e con Δy l'incremento *ancora delle ascisse x'* al progredire delle rette $y = \text{costante}$ (ossia

da (x_j, y_m) a (x_j, y_{m+1})). La condizione di allineamento verticale è

$$\Delta x = \Delta y. \quad (16.7.3)$$

Per ottenere l'allineamento voluto, consideriamo dapprima il caso di vista frontale dato dalla trasformazione di coordinate (16.6.4), dalla quale si ottiene

$$\Delta x = \tau \quad (16.7.4)$$

$$\Delta y = \tau f \cos \theta. \quad (16.7.5)$$

Pertanto la condizione di allineamento è soddisfatta se e solo se $f \cos \theta = 1$. Poiché ovviamente $\cos \theta \leq 1$, si deve avere $f \geq 1$. L'allineamento dopo la trasformazione prospettica è illustrato in Figura 24. Naturalmente, se θ non è piccolo (ovvero se è più vicino a $\frac{\pi}{2}$ che a 0), allora f è grande e si ha una forte distorsione prospettica, come osservato nella Nota 16.6.4.

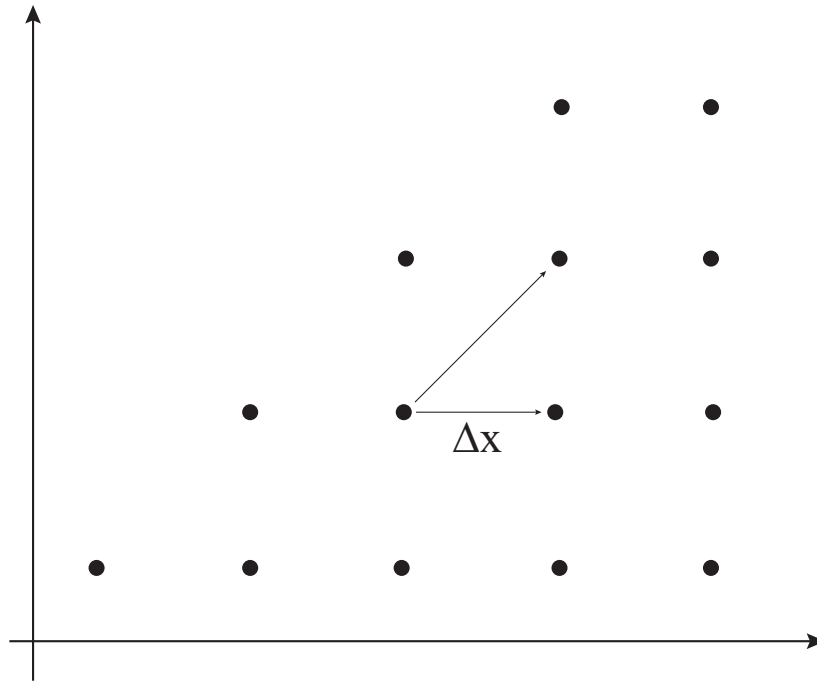


FIGURA 24. Allineamento prospettico verticale della griglia di base in proiezione obliqua frontale.

Per evitare la distorsione prospettica, possiamo adottare una variante più complessa dell'algoritmo, nella quale, invece di allineare verticalmente ogni riga $y = \text{costante}$ di punti della griglia, si allinea verticalmente ogni seconda riga successiva, mentre i punti nodali della righe dispari vengono disposti verticalmente a metà strada fra le colonne di allineamento delle righe pari (si veda la Fig. 24). In questo modo la condizione da soddisfare diventa $f \cos \theta = \frac{1}{2}$, e quindi possiamo scegliere f più piccolo di prima, ma si raddoppia la lunghezza degli array dei profili minimo e massimo.

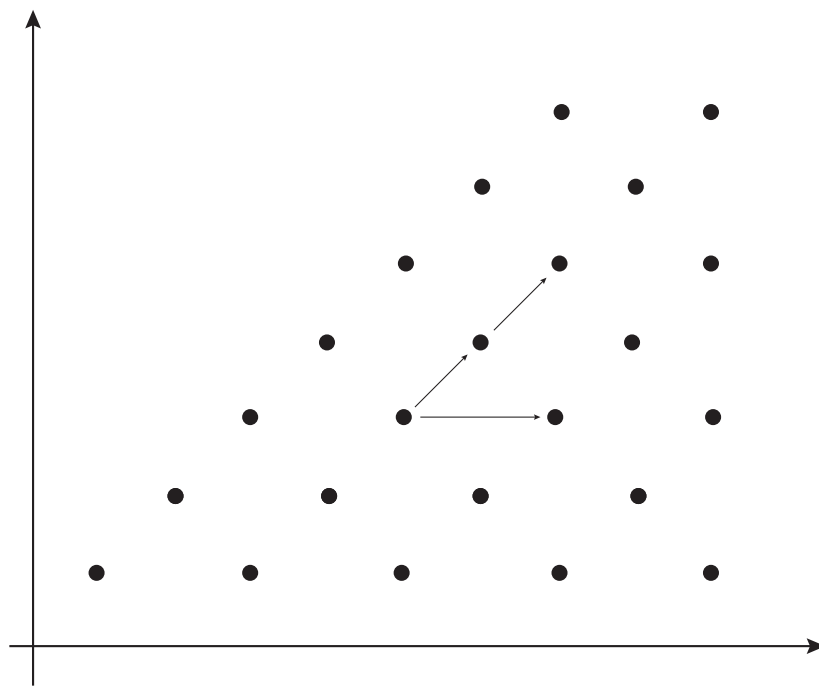


FIGURA 25. Allineamento prospettico a due passi della griglia, in proiezione obliqua frontale.

Veniamo ora al caso, più verisimile, della prospettiva obliqua con vista angolata, data dalla trasformazione (16.6.6). In base a questa trasformazione, invece delle identità (16.7.2), ora abbiamo

$$\Delta x = -\tau f_x \cos \beta$$

$$\Delta y = \tau f_y \cos \alpha,$$

e quindi la condizione di allineamento (16.7.3) diventa

$$f_x \cos \beta = f_y \cos \alpha. \quad (16.7.6)$$

Una volta scelti i fattori di accorciamento, possiamo adattare gli angoli α e β di inclinazione prospettica degli assi in modo da soddisfare questa relazione, o viceversa. Questo completa l'algoritmo.

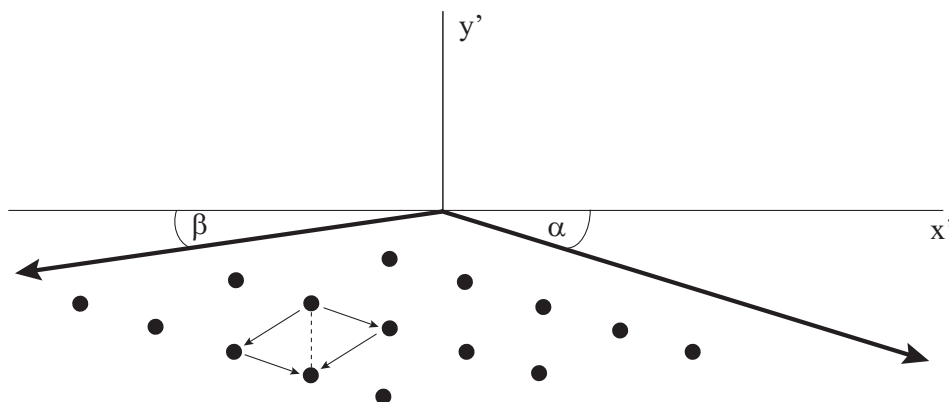


FIGURA 26. Allineamento prospettico verticale della griglia di base in proiezione obliqua angolata.

NOTA 16.7.1. Talvolta, per visualizzare un grafico in due variabili in un rettangolo, fissiamo il centro (X_0, Y_0) di un rettangolo sul piano xy con i lati paralleli agli assi, e le lunghezze ΔX e ΔY dei due lati, e scegliamo una griglia equiripartita suddividendo in $N - 1$ parti i due lati. In tal modo ciascuna cella della griglia ha lati non necessariamente quadrati, di lunghezze rispettive $\Delta X/(N - 1)$ e $\Delta Y/(N - 1)$ invece che τ . Allora la condizione di allineamento (16.7.6) diventa

$$\Delta X = \Delta Y \frac{f_x \cos \alpha}{f_y \cos \beta}.$$

In tal caso, la versatilità aggiuntiva data dai parametri di accorciamento f_x e f_y non serve più, perché viene rimpiazzata dalla libertà di scegliere a piacimento, e non necessariamente uguali, le lunghezze ΔX e ΔY dei lati (tranne che per le esigenze di invarianza dell'output al cambiare del

monitor spiegate nella Nota [16.6.5](#)). Pertanto di solito si pone $f_x = f_y$ e la regola di allineamento diventa

$$\Delta X = \Delta Y \frac{\cos \alpha}{\cos \beta}.$$

Per ogni scelta dei rapporti dei lati, si fissa un angolo di inclinazione α e si determina l'angolo β in modo che soddisfi questa regola (badando a non compiere scelte di angoli troppo estreme per non distorcere la prospettiva).

16.7.3. Appendice: pseudocodice per la rimozione di linee nascoste. Per concludere l'analisi dell'algoritmo, ora scegliamo la versione più tipica e verisimile, quella illustrata nella Nota [16.7.1](#), e ne presentiamo lo pseudocodice, omettendo le parti di scelta dei parametri grafici (numero di punti della griglia, angoli di inclinazione, scelta dell'asse da cui cominciare il disegno, rinormalizzazioni delle scale x e y per adattare l'output alle proporzioni del monitor, calcolo dell'array dei valori della funzione, inizializzazioni degli array dei profili). Per motivi di chiarezza didattica, lo pseudocodice è scritto in Pascal (le istruzioni grafiche di tracciamento delle linee e spostamento del cursore usano la sintassi di TurboPascal), ma da qui è facile trasportare il codice ad altri linguaggi, ad esempio C o Fortran. Occorre comunque assicurarsi che la routine di tracciamento di un segmento $(x1, y1) - (x2, y2)$ sulla periferica di output abbia la sintassi `line(x1,y1,x2,y2)`, o altrimenti modificare la sintassi appropriatamente. Questo codice è una nuova redazione, dove sono stati corretti alcuni errori e reso notevolmente più chiaro lo sviluppo logico, di versioni precedenti nei linguaggi Basic, Fortran e Pascal pubblicate in alcune tesi di laurea fra il 1982 ed il 1986 [\[4\]](#), [\[11\]](#), [\[8\]](#).

PROCEDURE DRAW

```

type punto, puntoprec: array[0..100,0..1] of real;
type fun array[0..100,0..100] of real;
type profmax, profmin: array[0..100] of real;

var alpha,beta,tau1,tau2,tau3,tau4,tau5: real;
    axmax,axmin,aymax,aymin,x1,x2,y1,y2: real;
    scalex,scaley,offsetx,offsety,xmax,xmin,ymax,ymin: real;
    x1,y1,x2,y2,xprof1,yprof1,xprof2,yprof2: real;
    num: integer;
    fullyvisible,invisible,drawsegment1,drawsegment2: boolean;
```

16.7. APPLICAZIONE: RIMOZIONE DI LINEE NASCOSTE IN GRAFICI 3D345

```
PROCEDURE INIT
{inizializza i valori da disegnare e gli angoli prospettici}
BEGIN {of procedure init}
    {... codice di inizializzazione, e acquisizione dei valori
della funzione (array fun) ...}
    tau1=cos(beta);
    tau2=cos(alpha);
    tau3=sin(beta);
    tau4=sin(alpha);
    tau5=f;

END; {of procedure init}


PROCEDURE WINDOW (var wxmin,wymin,wxmax,wymax: real);
{Calcola i coefficienti per la trasformazione
dalle coordinate del piano di proiezione a coordinate del viewport
a partire dagli estremi del viewport passati come parametri.
xs, ys: coordinate del piano di proiezione; xmonitor, ymonitor:
coordinate del viewport; wxmax, ymax, wxmin, wymin: estremi
della finestra di visuale in coordinate del piano di proiezione;
xmax, ymax: coordinate del vertice in basso a destra del monitor
(ad esempio 1024 e 768 per risoluzione VGA), assumendo 0,0 le
coordinate del vertice in alto a sinistra.
La regola di trasformazione e':
xmonitor = xmax*(xs-wxmin)/(wxmax-wxmin)  $\equiv$  scalex*xs-offsetx
ymonitor = ymax*(ys-wymin)/(wymax-wymin)  $\equiv$  scaley*ys-offsety}

    var wx, wy, scalex, scaley, offsetx, offsety: real;

BEGIN {of procedure window}
    if wxmin>wxmax then {scambiamo wxmin e wxmax}
    begin
        wx:=wxmax;
```

```

    wxmax:=wxmin;
    wxmin:=wx
end;
if wymin>wymax then  {scambiamo wymin e wymax}
begin
    wy:=wymax;
    wymax:=wymin;
    wymin:=wy
end;

{determiniamo i parametri della trasformazione dal piano di
proiezione al monitor}
    scalex:=xmax/(wxmax-wxmin);
    offsetx:=xmax*wxmin/(wxmax-wxmin);
    scaley:=ymax/(wymax-wymin);
    offsety:=ymax*wymin/(wymax-wymin);

END; {of procedure window}

PROCEDURE TRANSFORM (var x,y: real);
{Esegue la trasformazione dalle coordinate del piano di proiezione
ad opportune coordinate del viewport.
xs, ys: coordinate del piano di proiezione; xmonitor, ymonitor:
coordinate del viewport.
La regola di trasformazione e':
xmonitor = xmax*(xs-wxmin)/(wxmax-wxmin)  $\equiv$  scalex*xs-offsetx
ymonitor = ymax*(ys-wymin)/(wymax-wymin)  $\equiv$  scaley*ys-offsety
Qui scriviamo x, y per le coordinate xs, ys prima della trasformazione
ed anche per le coordinate xmonitor, ymonitor dopo la trasformazione.}

BEGIN {of procedure transform}

    x:=scalex*x-offsetx;
    y:=scaley*y-offsety;

```

16.7. APPLICAZIONE: RIMOZIONE DI LINEE NASCOSTE IN GRAFICI 3D347

{Queste coordinate sono float, ma prima di disegnare dobbiamo arrotondarle all'intero più vicino se la routine di tracciamento delle linee richiede la posizione dei pixel del monitor con indici interi: questa è l'ipotesi che faremo nella procedure DRAW responsabile del disegno.

In caso contrario basta omettere le prossime due istruzioni.}

x:=round(x);

y:=round(y);

END; {of procedure transform}

PROCEDURE INTERSECT (var x1in,y1in,x1fin,y1fin,x2in,y2in,x2fin,y2fin: real,

x1,y1,x2,y2,yprofmax1, yprofmax2, yprofmin1, yprofmin2: real);

{Calcola i punti di intersezione fra tre segmenti non verticali, il secondo ed il terzo dei quali si possono pensare come i due profili, e la porzione del primo segmento da omettere dal disegno perché nascosta.

Il parametro interno intersmax è true se il profilo massimo viene intersecato in un punto interno, e analogamente per intersmin per quanto concerne il profilo minimo.

Il parametro invisible è true se il segmento è completamente nascosto.

Pertanto questa procedura restituisce 5 possibili casi determinati da variabili globali boolean

che INTERSECT modifica, come segue:

fullyvisible=true: si deve disegnare l'intero segmento (x1,y1)-(x2,y2);

invisible=true: non si deve disegnare niente (questa variabile non occorre quindi restituirla al chiamante;

drawsegment1=true: si deve disegnare la parte iniziale del segmento, che viene restituita da questa procedura col nome (x1in,y1in)-(x1fin,y1fin)

(in questo caso x1in=x1 e y1in=y1, ma preferiamo cambiargli

nome per gestire la routine di disegno in maniera più simmetrica);
 drawsegment2=true: si deve disegnare la parte finale del segmento, chiamata (x2in,y2in)-(x2fin,y2fin); anche qui, x2fin=x2 e y2fin=y2.

Se drawsegment1=false o drawsegment2=false non si deve disegnare la parte corrispondente del segmento.}

{In questa procedura possiamo supporre che i segmenti abbiano estremi allineati verticalmente, quindi comincino tutti all'ascissa x1 e finiscano a x2.}

BEGIN {of procedure intersect}

{Inizializzazione:}

var xintmax, xintmin, yintmax, yintmin: real;
 pendprofmax, pendprofmin, pendenza, pendenzaprof, yprof1,
 yprof2, temp: real;
 var intersmax, intersmin: boolean; fullyvisible:=false;
 invisible:=true;
 intersmax:=true;
 intersmin:=true;
 drawsegment1:=false;
 drawsegment2:=false;

{Vogliamo calcolare il punto di intersezione delle rette che contengono i segmenti (x1,y1)-(x2,y2) e (x1,yprof1)-(x2,yprof2)}

{Riordiniamo le ascisse in ordine crescente se necessario:}

if x2>x1 then begin
 temp:=x1; x1:=x2; x2:=temp
end;
if x2>x1 then begin
temp:=x1; x1:=x2; x2:=temp
end;

16.7. APPLICAZIONE: RIMOZIONE DI LINEE NASCOSTE IN GRAFICI 3D349

{Per i nostri scopi possiamo assumere che le due rette non siano verticali.

Calcoliamo prima la pendenza delle due rette:}

```
pendenza:=(y2-y1)/(x2-x1);  
pendprofmax:=(yprofmax2-yprofmax1)/(x2-x1);  
pendprofmin:=(yprofmin2-yprofmin1)/(x2-x1);
```

{Trattiamo prima il caso di totale visibilità o di invisibilità.}

```
if ((y1>=yprofmax1) and (y2>=yprofmax2)) then
```

```
begin
```

```
    fullyvisible:=true; intersmax:=false
```

```
end;
```

```
if ((y1<=yprofmin1) and (y2<=yprofmin2)) then
```

```
begin
```

```
    fullyvisible:=true; intersmin:=false
```

```
end;
```

{In questi casi il segmento giace sopra il profilo massimo o sotto quello minimo, quindi è tutto visibile. Se coincide con uno dei profili allora non ci sono punti di intersezione interni.}

```
if (y1<yprofmax1) and (y1>yprofmin1) and (y2<yprofmax2) and  
(y2<yprofmin2) then
```

```
begin
```

```
    {In questo caso il segmento è invisibile.}
```

```
    invisible:=true;
```

```
    return
```

```
end;
```

{Osserviamo che il codice precedente ha già trattato il caso di segmento parallelo ad uno dei profili, e più in generale di segmenti completamente visibili o invisibili.}

{Calcoliamo ora l'intersezione col profilo massimo:}

```
yprof1:=yprofmax1;
```

```
yprof2:=yprofmax2;
```

```
pendenzaprof:=pendprofmax;
```

```

{Ora possiamo assumere che il segmento non sia parallelo
al profilo e che sia parzialmente visibile.
  Trattiamo allora il caso di visibilità parziale:
risolviamo il sistema lineare delle due rette.
Il punto di intersezione verifica le equazioni
yinters =x1 + pendenza * (xinters - x1)
yinters =x1 + pendenzaprof * (xinters - x1)
quindi la soluzione xinters si ottiene da
x1+pendenza*(xinters-x1)=x1+pendenzaprof*(xinters-x1)
ossia
      xinters =(x1*(1-pendenzaprof)
              - x1*(1-pendenza))/(pendenza-pendenzaprof) .
Scriviamo xintmax o xintmin invece di xinters
a seconda che il profilo sia quello massimo
o quello minimo.}
xintmax:=(x1*(1-pendenzaprof)
-x1*(1-pendenza))/(pendenza-pendenzaprof);
yintmax:=x1+pendenza*(xintmax-x1);
if (x1<=xintmax) and (xintmax<=x2) then xintersmax:=true;

{Calcoliamo infine l'intersezione col profilo minimo:}
yprof1:=yprofmin1;
yprof2:=yprofmin2;
pendenzaprof:=pendprofmin;
xintmin:=(x1*(1-pendenzaprof)
-x1*(1-pendenza))/(pendenza-pendenzaprof);
yintmin:=x1+pendenza*(xintmin-x1);
if (x1<=xintmin) and (xintmin<=x2) then xintersmin:=true;

{A questo punto almeno uno dei profili viene intersecato. Calcoliamo
allora l'estremo sinistro (x1in,y1in) e destro (x1fin,y1fin)
della prima parte di segmento da disegnare e l'estremo sinistro
(x2in,y2in) e destro (x2fin,y2fin) della seconda parte.}
if (intersmax=true) then begin {Intersezione con prof max.}
  if (y1>=yprofmax1) then begin

```

16.7. APPLICAZIONE: RIMOZIONE DI LINEE NASCOSTE IN GRAFICI 3D351

```
        {bisogna disegnare almeno la prima parte del segmento:}
        x1in:=x1;
        y1in:=y1;
        x1fin:=xintmax;
        y1fin:=yintmax;
        drawsegment1:=true
    end
else begin
    {bisogna disegnare almeno la seconda parte del segmento:}
    x2in:=xintmax;
    y2in:=yintmax;
    x2fin:=x2;
    y2fin:=y2;
    drawsegment2:=true
end
end;
end;
if (intersmin=true) then begin    {Intersezione con prof min.}

    if (y1<=yprofmin1) then begin
        x1in:=x1;
        y1in:=y1;
        x1fin:=xintmin;
        y1fin:=yintmin;
        drawsegment1:=true
    end
    else begin
        x2in:=xintmin;
        y2in:=yintmin;
        x2fin:=x2;
        y2fin:=y2;
        drawsegment2:=true
    end
    end;    end;
{Caso di una sola intersezione:}
```

```

    if (intersmax=false) then begin    {Niente intersezione con
prof max, solo con prof min}
    if (y1<yprofmin1) then begin
        {bisogna disegnare solo la prima parte del segmento:}
        x1in:=x1;
        y1in:=y1;
        x1fin:=xintmin;
        y1fin:=yintmin;
        drawsegment1:=true;
        drawsegment2:=false
    end
    else begin
        {bisogna disegnare solo la seconda parte del segmento:}
        x2in:=xintmin;
        y2in:=yintmin;
        x2fin:=x2;
        y2fin:=y2;
        drawsegment1:=false;
        drawsegment2:=true
    end
    end
end;

END; {of procedure intersect}

PROCEDURE PERSPECTIVE (var xout,yout: real; x,y,z: real);
{Esegue la trasformazione prospettica in prospettiva obliqua
dalle coordinate 3D x,y,z a coordinate xout, yout del viewplane.}

BEGIN {of procedure perspective}
    xout = -x*tau1 + y*tau2;
    yout = -x*tau3 - y*tau4 + z*tau5;

END; {of procedure perspective}

```

```
BEGIN {of procedure draw}
```

{num è il numero di quadratini in cui è suddiviso ogni lato della griglia di base.

Gli array x e y contengono le ascisse e le ordinate dei nodi della griglia di base; l'array bidimensionale fun contiene i valori della funzione da disegnare sui nodi di questa griglia. L'array bidimensionale punto contiene i punti bidimensionali (valori x' e y') sul piano di proiezione dei trasformati prospettici di una riga dei nodi del wireframe 3D. Rammentiamo che, prima di disegnarli, questi valori devono essere adattati alla scelta di coordinate nel viewport, tramite la procedura TRANSFORM, che utilizza gli estremi di viewport calcolati dalla procedura WINDOW.

L'array puntoprec ha le stesse dimensioni di punto, e serve per memorizzare la riga precedente. Questa memorizzazione è necessaria quando si devono tracciare segmenti trasversali alla poligonale corrente nel corso del disegno a maglie spiegato in (16.7.1): i punti iniziali di questi segmenti trasversali sono infatti stati già aggiornati nel corso del ciclo di aggiornamento dell'array unidimensionale punto.

Gli array profmax, profmin sono unidimensionali e memorizzano i valori minimo e massimo dei profili della zona invisibile. I corrispondenti array tempmax, tempmin memorizzano la versione di profmax e profmin che si riferisce alle prime due poligonali x=costante: in seguito vengono travasati in profmax e profmin rispettando gli indici in maniera che la capienza di profmax e profmin si raddoppi e tempmax e tempmin occupino la seconda metà di profmax e profmin (che corrisponde alla metà di sinistra se si va all'indietro in profondità, come illustrato nella Figura 19). Da quel momento in poi si ha bisogno solo degli array profmax e profmin per memorizzare i profili, ma bisogna traslare gli indici di un passo ogni volta che si passa alla poligonale un

passo più indietro.

axmin, aymin, axmax, aymax sono le coordinate sul piano $z=0$ dei vertici estremi della griglia di base.}

{Scelta degli estremi del viewport, in conformità alla Sezione **16.7.2.**

Come si vede in Figura 19, le ascisse e le ordinate del viewplane diminuiscono all'aumentare di x e non dipendono da z , mentre le ordinate del viewplane diminuiscono all'aumentare di x ma aumentano all'aumentare di y e z . Nel viewplane l'asse delle ordinate è orientato verso il basso; quindi l'ascissa del punto in alto a sinistra del viewplane corrisponde ad ascissa massima della griglia axmax e ordinata minima axmin, mentre la sua ordinata corrisponde a axmax, aymax e al valore minimo della funzione zmin. Per il punto in basso a destra naturalmente succede l'opposto.}

```
x1 = -axmax*tau1 + aymin*tau2;
y1 = -axmax*tau3 - aymax*tau4 + zmin*tau5;
x2 = -axmin*tau1 + aymax*tau2;
y2 = -axmin*tau3 - aymin*tau4 + zmax*tau5;
```

```
window (x1, y1, x2, y2);
```

{Tracciamo la prima curva $x=\text{costante}$, a partire dalla curva più vicina al punto di osservazione (ossia $x=\text{num}$). Per le due prime curve $x=\text{costante}$ e $y=\text{costante}$ non si devono effettuare rimozioni di linee nascoste, e quindi l'ordine di tracciamento dei segmenti e di aggiornamento dei profili massimo e minimo è inessenziale: comunque seguiamo l'ordine di tracciamento dal davanti all'indietro, ossia dall'indice num all'indice 0, come osservato più sopra.}

```
for iy:=num downto 0 do
begin
  xx:=x[num]; yy:=y[iy]; zz:=fun[num,iy];
```

16.7. APPLICAZIONE: RIMOZIONE DI LINEE NASCOSTE IN GRAFICI 3D355

```

perspective(xout, yout, xx, yy, zz);
punto[iy,0]:=xout ;
punto[iy,1]:=yout;
{Inizializzazione dei profili: gli array tempmax e tempmin
memorizzano i profili massimo e minimo delle prime due
poligonali x=costante, ossia x=num e x=num-1.}
tempmax[iy,0]:=punto[iy,0];
tempmax[iy,1]:=punto[iy,1];
tempmin[iy,0]:=punto[iy,0];
tempmin[iy,1]:=punto[iy,1]
end;
for iy:=num downto 1 do
begin
  {Per disegnare i segmenti dobbiamo adattare le coordinate
  del piano di proiezione a coordinate di viewport,
  ossia di monitor, chiamando la procedura TRANSFORM.}
  x1:=punto[iy,0];
  y1:=punto[iy,1];
  x2:=punto[iy-1,0];
  y2:=punto[iy-1,1];
  transform(x1,y1);
  transform(x2,y2);
  line(x1,y1,x2,y2);
end;

{Tracciamo la seconda curva x=costante(x=num-1):}
for iy:=num downto 0 do
begin
  {Memorizziamo in puntoprec i valori correnti dell'array punto
  prima di aggiornarlo: pertanto puntoprec contiene i valori
  di punto al passo precedente di aggiornamento. Questa
  memorizzazione è necessaria per tracciare i segmenti
  trasversali alla poligonale corrente: i punti iniziali di
  questi segmenti trasversali sono stati infatti modificati
  nel corso del ciclo di aggiornamento dell'array punto.}

```

```

    puntoprec[iy,0]:=punto[iy,0];
    puntoprec[iy,1]:=punto[iy,1];
    xx:=x[num-1]; yy:=y[iy]; zz:=fun[num-1,iy];
    perspective(xout, yout, xx, yy, zz);
    punto[iy,0]:=xout ;
    punto[iy,1]:=yout
end; for iy:=0 to num-1 do
begin
{Le posizioni dei due array si sfalsano di un passo ogni volta
che l'indice della x aumenta di una unità, e quindi nel passaggio
dalla poligonale x=num a quella x=num-1 occorre far slittare
di un passo gli array perché gli indici siano consistenti.}
    tempmax[iy,0]:=tempmax[iy+1,0];
    tempmax[iy,1]:=tempmax[iy+1,1];
    tempmin[iy,0]:=tempmin[iy+1,0];
    tempmin[iy,1]:=tempmin[iy+1,1]
end;
for iy:=num downto 1 do
begin {Disegno a maglie, come spiegato in (16.7.1):}
    {tratto x=costante}
    x1:=punto[iy,0];
    y1:=punto[iy,1];
    y2:=punto[iy-1,1];
    x2:=punto[iy-1,0];
    transform(x1,y1);
    transform(x2,y2);
    line(x1,y1,x2,y2);
    {tratto y=costante: questo è il segmento trasversale}
    x1:=punto[iy-1,0];
    y1:=punto[iy-1,1];
    {è qui che abbiamo bisogno del valore di punto alla riga precedente
    all'ultimo aggiornamento:}
    x2:=puntoprec[iy-1,0];    y2:=puntoprec[iy-1,1];
    transform(x1,y1);
    transform(x2,y2);

```

```

    line(x1,y1,x2,y2);
end;

{Aggiorniamo i profili massimo e minimo:}

for iy:=num-1 downto 0 do
begin
    if punto[iy,1]>tempmax[iy,1] then
    begin
        tempmax[iy,0]:=punto[iy,0];
        tempmax[iy,1]:=punto[iy,1]
    end;
    if punto[iy,1]<tempmin[iy,1] then
    begin
        tempmin[iy,0]:=punto[iy,0];
        tempmin[iy,1]:=punto[iy,1]
    end
end;
end;

{Tracciamo la prima curva y=costante(y=num):}

for ix:=num downto 0 do
begin
    x:=x[ix]; yy:=y[num]; zz:=fun[ix,num];
    perspective(xout, yout, xx, yy, zz);
    punto[ix,0]:=xout ;
    punto[ix,1]:=yout;
    profmax[ix,0]:=punto[ix,0];
    profmax[ix,1]:=punto[ix,1];
    profmin[ix,0]:=punto[ix,0];
    profmin[ix,1]:=punto[ix,1]
end;
for ix:=num downto 1 do
begin
    x1:=punto[ix,0];

```

```

    y1:=punto[ix,1];
    x2:=punto[ix-1,0];
    y2:=punto[ix-1,1];
    transform(x1,y1);
    transform(x2,y2);
    line(x1,y1,x2,y2)
end;

{Tracciamo la seconda curva y=costante(y=num-1):}
for ix:=num downto 0 do
begin
    puntoprec[ix,0]:=punto[ix,0];
    puntoprec[ix,1]:=punto[ix,1];
    x:=x[ix]; yy:=y[num-1]; zz:=fun[ix,num-1];
    perspective(xout, yout, xx, yy, zz);
    punto[ix,0]:=xout;
    punto[ix,1]:=yout
end;
for ix:=0 to num-1 do
begin
    {Le posizioni dei due array si sfalsano di un passo ogni volta
    che l'indice della y aumenta di una unità, e quindi nel passaggio
    dalla poligonale y=num a quella y=num-1 occorre far slittare
    di un passo gli array perché gli indici siano consistenti.}
    profmax[ix,0]:=profmax[ix+1,0];
    profmax[ix,1]:=profmax[ix+1,1];
    profmin[ix,0]:=profmin[ix+1,0];
    profmin[ix,1]:=profmin[ix+1,1]
end;
for ix:=num-1 downto 1 do
begin
    {tratto y=costante}
    x1:=punto[ix,0];
    y1:=punto[ix,1];
    x2:=punto[ix-1,0];

```

16.7. APPLICAZIONE: RIMOZIONE DI LINEE NASCOSTE IN GRAFICI 3D359

```

y2:=punto[ix-1,1];
transform(x1,y1);
transform(x2,y2);
line(x1,y1,x2,y2);
{tratto x=costante}
x1:=punto[ix-1,0];
y1:=punto[ix-1,1];
x2:=puntoprec[ix-1,0];
y2:=puntoprec[ix-1,1];
transform(x1,y1);
transform(x2,y2);
line(x1,y1,x2,y2);
end;

{Aggiorniamo i profili massimo e minimo:}
for ix:=num-1 downto 0 do
begin
  if punto[ix,1]>profmax[ix,1] then
  begin
    profmax[ix,0]:=punto[ix,0];
    profmax[ix,1]:=punto[ix,1]
  end;
  if punto[ix,1]<profmin[ix,1] then
  begin
    profmin[ix,0]:=punto[ix,0];
    profmin[ix,1]:=punto[ix,1]
  end
end;
end;
for ix:=1 to num-1 do
begin
  {Gli array tempmax, tempmin memorizzavano i profili massimo
  e minimo delle prime due poligonali x=costante: ora li travasiamo
  in profmax e profmin in maniera da raddoppiare la capienza di
  profmax e profmin e disporre i dati di tempmax e tempmin nella
  seconda metà di profmax e profmin (che corrisponde alla metà
```

di sinistra se si va all'indietro in profondità, come illustrato nella Figura 19).}

```
    profmax[num+(num-ix)-1,0]:=tempmax[ix-1,0];
    profmax[num+(num-ix)-1,1]:=tempmax[ix-1,1];
    profmin[num+(num-ix)-1,0]:=tempmin[ix-1,0];
    profmin[num+(num-ix)-1,1]:=tempmin[ix-1,1]
end;
```

{Tracciamo alternativamente un segmento di curva $y=\text{cost}$ e $x=\text{cost}$ a partire dalla seconda poligonale in poi, come spiegato in (16.7.1):}

```
for iy:=num-2 downto 0 do
begin
    for ix:=num downto 0 do
    begin
        puntoprec[ix,0]:=punto[ix,0];
        puntoprec[ix,1]:=punto[ix,1];
        x:=x[ix]; yy:=y[iy]; zz:=fun[ix,iy];
        perspective(xout, yout, xx, yy, zz);
        punto[ix,0]:=xout;
        punto[ix,1]:=yout
    end;
end;
```

{Poichè facciamo avanzare le y, i num+1 indici dei profili massimo e minimo sono posizioni consecutive in una sequenza di $2*\text{num}+1$ posizioni che corrispondono alle verticali dei nodi della griglia vista in prospettiva, metà per i nodi del lato più vicino al variare di x e l'altra metà al variare di y. Queste num posizioni si sfalsano di un passo ogni volta che l'indice della y aumenta di una unità, e quindi ogni volta occorre far slittare di un passo l'array perché gli indici siano consistenti. Nella scansione di una poligonale si usano solo num+1 di questi indici.}

```
for ix:=0 to (2*num-1) do
begin
    profmax[ix,0]:=profmax[ix+1,0];
```

```

    profmax[ix,1]:=profmax[ix+1,1];
    profmin[ix,0]:=profmin[ix+1,0];
    profmin[ix,1]:=profmin[ix+1,1];
end;
for ix:=num-1 downto 1 do
begin
{tratto y=costante}
    x1:=punto[ix,0];
    y1:=punto[ix,1];
    x2:=punto[ix-1,0];
    y2:=punto[ix-1,1];
    yprofmax1:=profmax[ix,1];
    yprofmax2:=profmax[ix-1,1];
    yprofmin1:=profmin[ix,1];
    yprofmin2:=profmin[ix-1,1];
    intersect (x1in,y1in,x1fin,y1fin,x2in,y2in,x2fin,y2fin,
    x1,y1,x2,y2,yprofmax1,yprofmax2,yprofmin1,yprofmin2);
    if fullyvisible then
    begin
        transform(x1,y1);
        transform(x2,y2);
        line(x1,y1,x2,y2)
    end
    else
    begin
        if drawsegment1 then
        begin
            transform(x1in,y1in);
            transform(x2in,y2in);
            line(x1in,y1in,x2in,y2in)
        end;
        if drawsegment2 then
        begin
            transform(x1fin,y1fin);
            transform(x2fin,y2fin);

```

```
        line(x1fin,y1fin,x2fin,y2fin)
    end
end;

{Aggiorniamo i profili massimo e minimo in corrispondenza al
punto finale a cui siamo arrivati (quello con indice ix-1).}

    if punto[ix-1,1]>=profmax[ix-1,1] then
    begin
        profmax[ix-1,0]:=punto[ix-1,0];
        profmax[ix-1,1]:=punto[ix-1,1]
    end;
    if punto[ix-1,1]<=profmin[ix-1,1] then
    begin
        profmin[ix-1,0]:=punto[ix-1,0];
        profmin[ix-1,1]:=punto[ix-1,1]
    end
end

END; {of procedure draw}
```

Parte 4

Appendice: norme, prodotti scalari, forme bilineari

* Appendice: norme, prodotti scalari e forme bilineari

Questa Appendice espande ed approfondisce il contenuto delle Sezioni [6.3](#), [6.5](#), [6.6](#) e [6.7](#).

17.1. Completezza e compattezza negli spazi metrici

Il contenuto di questa Sezione è contenuto in molti libri di testo di Analisi Matematica; si veda anche la [Sezione 1.1](#) di [?].

DEFINIZIONE 17.1.1. Uno *spazio metrico* è un insieme $\mathbf{X}$ munito con una *funzione distanza* (o *metrica*) $d : \mathbf{X} \times \mathbf{X} \longrightarrow [0, +\infty)$, cioè una funzione che assegna ad ogni due elementi $\mathbf{x}, \mathbf{y} \in \mathbf{X}$ un numero $d(\mathbf{x}, \mathbf{y}) \geq 0$ (la loro distanza) in tal modo che siano soddisfatte le condizioni

- (i) $d(\mathbf{x}, \mathbf{y}) = 0$ se e solo se $\mathbf{x} = \mathbf{y}$,
- (ii) (*simmetria*) $d(\mathbf{x}, \mathbf{y}) = d(\mathbf{y}, \mathbf{x})$ per tutti i $\mathbf{x}, \mathbf{y} \in \mathbf{X}$,
- (iii) (*disuguaglianza triangolare*) $d(\mathbf{x}, \mathbf{z}) \leq d(\mathbf{x}, \mathbf{y}) + d(\mathbf{y}, \mathbf{z})$ per tutti i $\mathbf{x}, \mathbf{y}, \mathbf{z} \in \mathbf{X}$.

ESEMPIO 17.1.2.

- *Metrica euclidea in* $\mathbb{R}$: $\mathbf{X} = \mathbb{R}$, $d(t, s) = |t - s|$ per ogni $t, s \in \mathbb{R}$.
- *Metrica* l^∞ *in* $\mathbb{R}^n$: $\mathbf{X} = \mathbb{R}^n$, $d(\mathbf{x}, \mathbf{y}) = \max_{1 \leq j \leq n} |x_j - y_j|$

$$\text{dove } \mathbf{x} = \begin{pmatrix} x_1 \\ \vdots \\ x_n \end{pmatrix}, \mathbf{y} = \begin{pmatrix} y_1 \\ \vdots \\ y_n \end{pmatrix}.$$

□

NOTA 17.1.3. Un'altra forma della disuguaglianza triangolare è

$$|d(\mathbf{x}, \mathbf{z}) - d(\mathbf{y}, \mathbf{z})| \leq d(\mathbf{x}, \mathbf{y}), \quad \mathbf{x}, \mathbf{y}, \mathbf{z} \in \mathbf{X}. \quad (17.1.1)$$

Infatti, la disuguaglianza triangolare implica

$$d(\mathbf{x}, \mathbf{z}) \leq d(\mathbf{x}, \mathbf{y}) + d(\mathbf{y}, \mathbf{z}),$$

cioè

$$d(\mathbf{x}, \mathbf{z}) - d(\mathbf{y}, \mathbf{z}) \leq d(\mathbf{x}, \mathbf{y}),$$

e da qui, scambiando $\mathbf{x}$ con $\mathbf{y}$ ed usando la simmetria di d , otteniamo

$$d(\mathbf{y}, \mathbf{z}) - d(\mathbf{x}, \mathbf{z}) \leq d(\mathbf{y}, \mathbf{x}) = d(\mathbf{x}, \mathbf{y}),$$

cioè

$$-d(\mathbf{x}, \mathbf{y}) \leq d(\mathbf{x}, \mathbf{z}) - d(\mathbf{y}, \mathbf{z}).$$

Viceversa, da (17.1.1) si ha $d(\mathbf{x}, \mathbf{z}) - d(\mathbf{y}, \mathbf{z}) \leq |d(\mathbf{x}, \mathbf{z}) - d(\mathbf{y}, \mathbf{z})| \leq d(\mathbf{x}, \mathbf{y})$, che è la disuguaglianza triangolare. $\square$

In uno spazio metrico $\mathbf{X}$ possiamo definire una nozione di convergenza.

DEFINIZIONE 17.1.4. Si dice che una successione $\{\mathbf{x}_k\}_{k \geq 1}$ in $\mathbf{X}$ converge ad $\mathbf{x} \in \mathbf{X}$ se per ogni $\varepsilon > 0$ esiste un indice k_ε tale che $d(\mathbf{x}, \mathbf{x}_k) \leq \varepsilon$ per qualsiasi $k \geq k_\varepsilon$.

NOTA 17.1.5. La successione $\{\mathbf{x}_k\}_{k \geq 1}$ non può convergere a due elementi diversi. Infatti, se essa converge sia ad $\mathbf{x}$ che ad $\mathbf{x}'$, allora per ogni $\varepsilon > 0$ ed ogni k abbastanza grande,

$$d(\mathbf{x}, \mathbf{x}') \leq d(\mathbf{x}, \mathbf{x}_k) + d(\mathbf{x}_k, \mathbf{x}') \leq \varepsilon + \varepsilon = 2\varepsilon,$$

perciò $d(\mathbf{x}, \mathbf{x}') = 0$. $\square$

NOTAZIONE 17.1.6. Se la successione $\{\mathbf{x}_k\}_{k \geq 1}$ converge, allora l'unico elemento di $\mathbf{X}$ a quale la successione converge si chiama il limite della successione e si indica con

$$\lim_{k \rightarrow \infty} \mathbf{x}_k, \quad \text{ovvero} \quad \lim_k \mathbf{x}_k.$$

La convergenza di $\{\mathbf{x}_k\}_{k \geq 1}$ ad $\mathbf{x}$ si esprime scrivendo $\mathbf{x}_k \longrightarrow \mathbf{x}$.

Sia $\{\mathbf{x}_k\}_{k \geq 1}$ una successione in uno spazio metrico $\mathbf{X}$. Ricordiamo che, per ogni successione strettamente crescente $1 \leq k_1 < k_2 < \dots$ di numeri naturali, la successione $\{\mathbf{x}_{k_j}\}_{j \geq 1}$ si chiama *sottosuccessione* di $\{\mathbf{x}_k\}_{k \geq 1}$. È facile vedere che, se $\{\mathbf{x}_k\}_{k \geq 1}$ converge ad $\mathbf{x}$, allora tutte le sue sottosuccessioni convergono ad $\mathbf{x}$.

LEMMA 17.1.7. *Ogni successione convergente $\{\mathbf{x}_k\}_{k \geq 1}$ in uno spazio metrico $\mathbf{X}$ è una successione di Cauchy, cioè gode della proprietà:*

$$\forall \varepsilon > 0 \quad \exists k_\varepsilon \geq 1 \quad \text{tale che } d(\mathbf{x}_k, \mathbf{x}_l) \leq \varepsilon \quad \forall k, l \geq k_\varepsilon. \quad (17.1.2)$$

DIMOSTRAZIONE. Se $\mathbf{x}$ è il limite della successione e k_ε è tale che $d(\mathbf{x}, \mathbf{x}_k) \leq \varepsilon/2$ per $k \geq k_\varepsilon$, allora per $k, l \geq k_\varepsilon$ vale

$$d(\mathbf{x}_k, \mathbf{x}_l) \leq d(\mathbf{x}_k, \mathbf{x}) + d(\mathbf{x}, \mathbf{x}_l) \leq \frac{\varepsilon}{2} + \frac{\varepsilon}{2} = \varepsilon.$$

NOTAZIONE 17.1.8. *Uno spazio metrico, nel quale, viceversa, ogni successione di Cauchy converge, si chiama completo.*

LEMMA 17.1.9. *Ogni successione di Cauchy $\{\mathbf{x}_k\}_{k \geq 1}$ in uno spazio metrico $\mathbf{X}$ è limitata, cioè esistono $\mathbf{a} \in \mathbf{X}$ e $r \geq 0$ tali che $d(\mathbf{a}, \mathbf{x}_k) \leq r$ per ogni $k \geq 1$.*

DIMOSTRAZIONE. Per (17.1.2) esiste $k_1 \geq 1$ tale che $d(\mathbf{x}_k, \mathbf{x}_l) \leq 1$ per qualsiasi $k, l \geq k_1$. Quindi, ponendo

$$\mathbf{a} := \mathbf{x}_1, \quad r := 1 + \max_{1 \leq k \leq k_1} d(\mathbf{x}_1, \mathbf{x}_k),$$

abbiamo $d(\mathbf{a}, \mathbf{x}_k) \leq r$ per ogni $k \geq 1$.

Nella verifica della completezza può essere utile il fatto seguente:

LEMMA 17.1.10. *Una successione di Cauchy $\{\mathbf{x}_k\}_{k \geq 1}$ in uno spazio metrico $\mathbf{X}$ converge se e solo se esiste una sua sottosuccessione convergente.*

DIMOSTRAZIONE. La parte “solo se” essendo ovvia, supponiamo che esista una sottosuccessione $\{\mathbf{x}_{k_j}\}_{j \geq 1}$ convergente ad $\mathbf{x}$. Ora sia $\varepsilon > 0$

arbitrario. Allora esistono $k_0 \geq 1$ e $j_0 \geq 1$ tali che

$$d(\mathbf{x}_k, \mathbf{x}_l) \leq \frac{\varepsilon}{2} \quad \text{per } k, l \geq k_0, \quad d(\mathbf{x}, \mathbf{x}_{k_j}) \leq \frac{\varepsilon}{2} \quad \text{per } j \geq j_0.$$

Scegliamo un $j_1 \geq j_0$ tale che $k_{j_1} \geq k_0$. Allora per ogni $k \geq k_0$ abbiamo:

$$d(\mathbf{x}, \mathbf{x}_k) \leq d(\mathbf{x}, \mathbf{x}_{k_{j_1}}) + d(\mathbf{x}_{k_{j_1}}, \mathbf{x}_k) \leq \frac{\varepsilon}{2} + \frac{\varepsilon}{2} = \varepsilon.$$

TEOREMA 17.1.11. (Completezza di $\mathbb{R}$.) *Munito della metrica euclidea, $\mathbb{R}$ è uno spazio metrico completo.*

Per la dimostrazione di questo Teorema dobbiamo rammentare i seguenti fatti (si veda un libro di testo di Analisi Matematica).

LEMMA 17.1.12. *Ogni sottoinsieme limitato non vuoto $S \subset \mathbb{R}$ ha estremo inferiore m ed estremo superiore M , cioè esistono numeri reali m ed M tali che*

$$m \leq t \quad \forall t \in S, \quad \forall \lambda > m \quad \exists t_\lambda \in S \text{ tale che } t_\lambda < \lambda, \quad (17.1.3)$$

$$t \leq M \quad \forall t \in S, \quad \forall \mu < M \quad \exists t_\mu \in S \text{ tale che } \mu < t_\mu. \quad (17.1.4)$$

Poniamo

$$m =: \inf S \quad \text{o} \quad \inf_{t \in S} t, \quad M =: \sup S \quad \text{o} \quad \sup_{t \in S} t.$$

LEMMA 17.1.13. *Ogni successione non decrescente e superiormente limitata $\{t_k\}_{k \geq 1}$ di numeri reali converge a $\lim_{k \rightarrow \infty} t_k = \sup_{k \geq 1} t_k$.*

DIMOSTRAZIONE. Poniamo $M = \sup_{k \geq 1} t_k$: dal lemma precedente sappiamo che $M < \infty$. Inoltre, per ogni $\varepsilon > 0$, esiste un k_ε con $M - \varepsilon < t_{k_\varepsilon}$. Allora, per la crescita della successione $\{t_k\}_{k \geq 1}$,

$$k \geq k_\varepsilon \implies t_k \geq t_{k_\varepsilon} > M - \varepsilon \implies 0 \leq M - t_k < \varepsilon \implies |M - t_k| < \varepsilon.$$

Analogamente,

COROLLARIO 17.1.14. *Ogni successione non crescente ed inferiormente limitata $\{s_k\}_{k \geq 1}$ di numeri reali converge a $\lim_{k \rightarrow \infty} s_k = \inf_{k \geq 1} s_k$.*

Il seguente Lemma permette di ridurre la dimostrazione di molte affermazioni riguardanti successioni generali di numeri reali al caso di successioni monotone (cioè non decrescenti o non crescenti).

LEMMA 17.1.15. *Ogni successione di numeri reali ammette una sottosuccessione monotona.*

DIMOSTRAZIONE. Sia $\{t_k\}_{k \geq 1}$ una successione di numeri reali. Indichiamo con S l'insieme di tutti gli indici k tale che t_k maggiore qualsiasi elemento successivo della successione considerata, cioè

$$S := \{k \geq 1; t_l \leq t_k \quad \forall l > k\}.$$

Se S è infinito, allora $S = \{k_1, k_2, \dots\}$ dove $k_1 < k_2 < \dots$ e chiaramente la sottosuccessione infinita $\{t_{k_j}\}_{j \geq 1}$ è non crescente.

Se invece S è finito, allora esiste un k_1' tale che nessun k con $k \geq k_1'$ appartiene ad S . In particolare $k_1' \notin S$, perciò esiste (almeno) un $k_2' > k_1'$ tale che $t_{k_2'} > t_{k_1'}$. Poiché neanche k_2' appartiene ad S , esiste $k_3' > k_2'$ con $t_{k_3'} > t_{k_2'}$, e così via: otteniamo numeri naturali $k_1' < k_2' < \dots$ tale che $t_{k_1'} < t_{k_2'} < \dots$ e concludiamo che $\{t_k\}_{k \geq 1}$ ammette una sottosuccessione strettamente crescente.

DIMOSTRAZIONE DEL TEOREMA 17.1.11. Se $\{t_k\}_{k \geq 1}$ è una successione di Cauchy di numeri reali, allora per il Lemma 17.1.15 esiste una sua sottosuccessione monotona $\{t_{k_j}\}_{j \geq 1}$. Poiché le successioni di Cauchy sono limitate (Lemma 17.1.9), anche la sottosuccessione $\{t_{k_j}\}_{j \geq 1}$ è limitata, quindi, essendo monotona, convergente (Lemma 17.1.13 e Corollario 17.1.14). Ne segue che la successione di Cauchy $\{t_k\}_{k \geq 1}$ ammette una sottosuccessione convergente e perciò converge.

□

Analogamente:

TEOREMA 17.1.16. (**Completezza di $\mathbb{R}^n$.**) *Munito della metrica l^∞ , $\mathbb{R}^n$ è uno spazio metrico completo per ogni n .*

DIMOSTRAZIONE. Sia $\mathbf{x}^{(k)} = \begin{pmatrix} x_1^{(k)} \\ \vdots \\ x_n^{(k)} \end{pmatrix}$, $k \geq 1$, una successione di

Cauchy in $\mathbb{R}^n$. Allora per ogni $\varepsilon > 0$ esiste un $k_\varepsilon \geq 1$ tale che

$$\max_{1 \leq j \leq n} |x_j^{(k)} - x_j^{(l)}| \leq \varepsilon \text{ per ogni } k, l \geq k_\varepsilon :$$

quindi tutte le n successioni di numeri reali $\{x_j^{(k)}\}_{k \geq 1}$, $1 \leq j \leq n$, sono di Cauchy, perciò convergono. Poniamo

$$x_j := \lim_{k \rightarrow \infty} x_j^{(k)} \text{ e } \mathbf{x} = \begin{pmatrix} x_1 \\ \vdots \\ x_n \end{pmatrix} :$$

si verifica facilmente che $\mathbf{x}^{(k)} \rightarrow \mathbf{x}$.

DEFINIZIONE 17.1.17. (i) Un sottoinsieme F di uno spazio metrico $\mathbf{X}$ si dice *chiuso* se appartiene ad F il limite di ogni successione convergente contenuta in F . In altre parole, $F \subset \mathbf{X}$ è chiuso se

$$F \ni \mathbf{x}_k \rightarrow \mathbf{x} \implies \mathbf{x} \in F.$$

(ii) Un sottoinsieme K di $\mathbf{X}$ si chiama *compatto* se ogni successione appartenente a K ammette una sottosuccessione convergente ad un elemento di K .

LEMMA 17.1.18. Ogni sottoinsieme compatto di $\mathbf{X}$ è chiuso.

DIMOSTRAZIONE. Se $K \subset \mathbf{X}$ è compatto e $K \ni \mathbf{x}_k \rightarrow \mathbf{x}$, allora esiste una sottosuccessione $\{\mathbf{x}_{k_j}\}_{j \geq 1}$ convergente ad un $\mathbf{y} \in K$ e, poiché $\{\mathbf{x}_{k_j}\}_{j \geq 1}$ resta convergente ad $\mathbf{x}$, dall'unicità del limite risulta che $\mathbf{x} = \mathbf{y} \in K$.

LEMMA 17.1.19. Ogni sottoinsieme compatto di $\mathbf{X}$ è limitato.

DIMOSTRAZIONE. Sia $S \subset \mathbf{X}$ un insieme non limitato. Allora S non è vuoto, perciò esiste un $\mathbf{x}_1 \in S$. Poiché S non è limitato, non possiamo avere $d(\mathbf{x}_1, \mathbf{x}) \leq 1$ per tutti i $\mathbf{x} \in S$, perciò esiste un $\mathbf{x}_2 \in S$ tale che

$$d(\mathbf{x}_1, \mathbf{x}_2) > 1.$$

Iteriamo lo stesso argomento: se avessimo $d(\mathbf{x}_1, \mathbf{x}) \leq 1$ o $d(\mathbf{x}_2, \mathbf{x}) \leq 1$ per ogni $\mathbf{x} \in S$, allora risulterebbe $d(\mathbf{x}_1, \mathbf{x}) \leq 1 + d(\mathbf{x}_1, \mathbf{x}_2)$ per ogni $\mathbf{x} \in S$ e quindi S sarebbe limitato. Pertanto esiste $\mathbf{x}_3 \in S$ tale che

$$d(\mathbf{x}_1, \mathbf{x}_3) > 1, \quad d(\mathbf{x}_2, \mathbf{x}_3) > 1.$$

Continuando così si ottiene una successione $\{\mathbf{x}_k\}_{k \geq 1}$ in S tale che $d(\mathbf{x}_k, \mathbf{x}_l) > 1$ per qualsiasi $k \neq l$. Ma allora nessuna sottosuccessione di questa successione è di Cauchy ed a maggior ragione nessuna sottosuccessione può convergere (Lemma 17.1.7). Di conseguenza S non è compatto.

TEOREMA 17.1.20. *un sottoinsieme di $\mathbb{R}^n$ è compatto se e solo se è chiuso e limitato.*

DIMOSTRAZIONE. Poiché gli insiemi compatti sono chiusi e limitati in tutti gli spazi metrici (Lemmi 17.1.18 e 17.1.19), dobbiamo solo dimostrare che ogni sottoinsieme chiuso e limitato K di $\mathbb{R}^n$ è compatto.

Rammentiamo quanto osservato nella dimostrazione del Teorema 17.1.11 sulla completezza di $\mathbb{R}$: *ogni successione limitata di numeri reali ammette una sottosuccessione convergente.*

Sia allora $\mathbf{x}^{(k)} = \begin{pmatrix} x_1^{(k)} \\ \vdots \\ x_n^{(k)} \end{pmatrix}$, $k \geq 1$, una successione in K . Poiché tutte

le n successioni di numeri reali $\{x_j^{(k)}\}_{k \geq 1}$, $1 \leq j \leq n$, sono limitati, applicando iterativamente l'osservazione precedente otteniamo

- una sottosuccessione $\{\mathbf{x}^{(k)}\}_{k \in \mathcal{K}_1}$ di $\{\mathbf{x}^{(k)}\}_{k \geq 1}$, ove $\mathcal{K}_1$ è un sottoinsieme infinito di $\mathbb{N}$, cioè una successione strettamente crescente di numeri naturali tale che la sottosuccessione $\{x_1^{(k)}\}_{k \in \mathcal{K}_1}$ di $\{x_1^{(k)}\}_{k \geq 1}$ converga,
- poi una sottosuccessione $\{\mathbf{x}^{(k)}\}_{k \in \mathcal{K}_2}$, ove $\mathcal{K}_2$ è un sottoinsieme infinito di $\mathcal{K}_1$, tale che la sottosuccessione $\{x_2^{(k)}\}_{k \in \mathcal{K}_2}$ converga,
-
- ed infine una sottosuccessione $\{\mathbf{x}^{(k)}\}_{k \in \mathcal{K}_n}$, ove $\mathcal{K}_n$ è un sottoinsieme infinito di $\mathcal{K}_{n-1}$, tale che anche $\{x_n^{(k)}\}_{k \in \mathcal{K}_n}$ converga.

Quindi convergono tutte le n sottosuccessioni reali $\{x_j^{(k)}\}_{k \in \mathcal{K}_n}$, $1 \leq j \leq n$, il che significa che la sottosuccessione $\{\mathbf{x}^{(k)}\}_{k \in \mathcal{K}_n}$ converge ad un vettore $\mathbf{x} \in \mathbb{R}^n$. Poiché K è chiuso, $K \ni \mathbf{x}^{(k)} \longrightarrow \mathbf{x}$ implica che $\mathbf{x} \in K$.

DEFINIZIONE 17.1.21. Siano $\mathbf{X}$ e $\mathbf{Y}$ spazi metrici. Diciamo che una applicazione $f : \mathbf{X} \longrightarrow \mathbf{Y}$ è *continua* se trasforma successioni convergenti in successioni convergenti.

Se f è continua abbiamo

$$\mathbf{x}_k \longrightarrow \mathbf{x} \in \mathbf{X} \implies f(\mathbf{x}_k) \longrightarrow f(\mathbf{x}) \in \mathbf{Y}.$$

Infatti, la convergenza di $\{f(\mathbf{x}_k)\}_{k \geq 1}$ ad $f(\mathbf{x})$ risulta dalla convergenza della successione

$$f(\mathbf{x}_1), f(\mathbf{x}), f(\mathbf{x}_2), f(\mathbf{x}), f(\mathbf{x}_3), f(\mathbf{x}), \dots$$

che è l'immagine tramite f della successione convergente

$$\mathbf{x}_1, \mathbf{x}, \mathbf{x}_2, \mathbf{x}, \mathbf{x}_3, \mathbf{x}, \dots$$

NOTA 17.1.22. Segue immediatamente dalle definizioni che, se l'applicazione $f : \mathbf{X} \longrightarrow \mathbf{Y}$ è continua e K è un sottoinsieme compatto di $\mathbf{X}$, allora $f(K) = \{f(\mathbf{x}); \mathbf{x} \in \mathbf{X}\}$ è un sottoinsieme compatto di $\mathbf{Y}$. $\square$

La conservazione della compattezza sotto passaggio ad una immagine continua porta all'importante risultato seguente.

TEOREMA 17.1.23. (**Weierstrass.**) Se $\mathbf{X}$ è uno spazio metrico, $f : \mathbf{X} \rightarrow \mathbb{R}$ è una funzione continua e $K \subset \mathbf{X}$ è un insieme compatto, allora f è limitata su K ed esistono un punto di minimo $\mathbf{x}_{\min}$ ed uno di massimo $\mathbf{x}_{\max}$ per f su K , cioè due elementi di K tali che

$$f(\mathbf{x}_{\min}) \leq f(\mathbf{x}) \leq f(\mathbf{x}_{\max}), \quad \forall \mathbf{x} \in K.$$

DIMOSTRAZIONE. La limitatezza di f su K segue dalla compattezza, e quindi limitatezza dell'immagine continua $f(K)$ dell'insieme compatto K .

Siano $m := \inf f(K)$ e $M := \sup f(K)$. In base alla definizione dell'estremo inferiore m (17.1.4), per ogni numero intero $k \geq 1$ esiste un $t_k \in f(K)$ tale che $m \leq t_k < m + \frac{1}{k}$. Allora $f(K) \ni t_k \rightarrow m$ ed il fatto che l'insieme compatto $f(K)$ sia chiuso (Lemma 17.1.18) implica $m \in f(K)$, cioè che esiste un $\mathbf{x}_{\min} \in K$ tale che $m = f(\mathbf{x}_{\min})$.

In maniera analoga si verifica l'esistenza di un $\mathbf{x}_{\max} \in K$ per quale $M = f(\mathbf{x}_{\max})$.

17.2. Norme su spazi vettoriali reali

DEFINIZIONE 17.2.1. Sia $\mathbf{X}$ uno spazio vettoriale reale. Una *norma* su $\mathbf{X}$ è una funzione

$$\mathbf{X} \ni \mathbf{x} \mapsto \|\mathbf{x}\| \in [0, +\infty)$$

tale che

- (i) $\|\mathbf{x}\| = 0$ se e solo se $\mathbf{x} = \mathbf{0}_{\mathbf{X}}$,
- (ii) $\|\lambda \mathbf{x}\| = |\lambda| \|\mathbf{x}\|$ per ogni scalare $\lambda \in \mathbb{R}$ e vettore $\mathbf{x}$,
- (iii) (*disuguaglianza triangolare*) $\|\mathbf{x} + \mathbf{y}\| \leq \|\mathbf{x}\| + \|\mathbf{y}\|$ per tutti i vettori $\mathbf{x}$ e $\mathbf{y}$.

Uno spazio vettoriale reale munito di una norma si chiama *spazio vettoriale normato reale*.

ESEMPIO 17.2.2. Se $\mathbf{X}$ è di dimensione finita n e

$$\mathcal{B} = \{\mathbf{b}_1, \dots, \mathbf{b}_n\}$$

è una base di $\mathbf{X}$, allora ognuna delle formule

$$\left\| \sum_{j=1}^n \alpha_j \mathbf{b}_j \right\|_1^{(\mathcal{B})} := \sum_{j=1}^n |\alpha_j|, \quad \left\| \sum_{j=1}^n \alpha_j \mathbf{b}_j \right\|_{\infty}^{(\mathcal{B})} := \max_{1 \leq j \leq n} |\alpha_j|$$

definisce una norma su $\mathbf{X}$. Nel caso di $\mathbf{X} = \mathbb{R}^n$ e $\mathcal{B}$ base naturale di $\mathbb{R}^n$ useremo le notazioni $\|\cdot\|_1$ e $\|\cdot\|_{\infty}$, senza specificare la base. $\square$

NOTA 17.2.3. Un'altra forma della disuguaglianza triangolare è

$$|\|\mathbf{x}\| - \|\mathbf{y}\|| \leq \|\mathbf{x} - \mathbf{y}\|, \quad \mathbf{x}, \mathbf{y} \in \mathbf{X}. \quad (17.2.1)$$

Infatti, la disuguaglianza triangolare implica

$$\|\mathbf{x}\| = \|\mathbf{y} + (\mathbf{x} - \mathbf{y})\| \leq \|\mathbf{y}\| + \|\mathbf{x} - \mathbf{y}\|,$$

cioè

$$\|\mathbf{x}\| - \|\mathbf{y}\| \leq \|\mathbf{x} - \mathbf{y}\|,$$

e da qui, scambiando nella disuguaglianza precedente $\mathbf{x}$ con $\mathbf{y}$, otteniamo

$$\|\mathbf{y}\| - \|\mathbf{x}\| \leq \|\mathbf{y} - \mathbf{x}\| = \|\mathbf{x} - \mathbf{y}\|,$$

cioè

$$-\|\mathbf{x} - \mathbf{y}\| \leq \|\mathbf{x}\| - \|\mathbf{y}\|.$$

Viceversa, da (17.2.1) segue $\|\mathbf{x} + \mathbf{y}\| - \|\mathbf{y}\| \leq \|\mathbf{x} + \mathbf{y}\| - \|\mathbf{y}\| \leq \|\mathbf{x}\|$, quindi la disuguaglianza triangolare. $\square$

Ogni norma $\|\cdot\|$ su $\mathbf{X}$ definisce una *funzione distanza*

$$d = d_{\|\cdot\|} : \mathbf{X} \times \mathbf{X} \ni (\mathbf{x}, \mathbf{y}) \mapsto \|\mathbf{x} - \mathbf{y}\| \in [0, +\infty),$$

con la quale $\mathbf{X}$ diventa uno spazio metrico. Perciò in uno spazio vettoriale normato possiamo parlare di convergenza, compattezza e continuità.

DEFINIZIONE 17.2.4. Due norme $\|\cdot\|_1$ e $\|\cdot\|_2$ su $\mathbf{X}$ si dicono *equivalenti* se esistono costanti $0 < c_1 \leq c_2$ tale che

$$c_1 \|\mathbf{x}\|_1 \leq \|\mathbf{x}\|_2 \leq c_2 \|\mathbf{x}\|_1 \text{ per tutti i vettori } \mathbf{x}. \quad (17.2.2)$$

Chiaramente, l'equivalenza di norme è una relazione di equivalenza, nel senso della Definizione 1.3.5. La proprietà transitiva non è altro che l'ovvia osservazione seguente: se $\|\cdot\|_1, \|\cdot\|_2, \|\cdot\|_3$ sono norme su $\mathbf{X}$, allora

$$\left. \begin{array}{l} \|\cdot\|_1 \text{ equivalente a } \|\cdot\|_2 \\ \|\cdot\|_2 \text{ equivalente a } \|\cdot\|_3 \end{array} \right\} \implies \|\cdot\|_1 \text{ equivalente a } \|\cdot\|_3.$$

Nel precedente esempio 17.2.2 le due norme $\|\cdot\|_1^{(\mathcal{B})}$ e $\|\cdot\|_\infty^{(\mathcal{B})}$ sono equivalenti:

$$\left\| \sum_{j=1}^n \alpha_j \mathbf{b}_j \right\|_\infty^{(\mathcal{B})} \leq \left\| \sum_{j=1}^n \alpha_j \mathbf{b}_j \right\|_1^{(\mathcal{B})} \leq n \left\| \sum_{j=1}^n \alpha_j \mathbf{b}_j \right\|_\infty^{(\mathcal{B})}.$$

Vale di più:

TEOREMA 17.2.5. (Equivalenza delle norme a dimensione finita.) *Ogni due norme su uno spazio vettoriale di dimensione finita sono equivalenti.*

DIMOSTRAZIONE. Sia $\mathbf{X}$ uno spazio vettoriale di dimensione n e

$$\mathcal{B} = \{\mathbf{b}_1, \dots, \mathbf{b}_n\}$$

una base di $\mathbf{X}$. Basta provare che qualsiasi norma $\|\cdot\|$ su $\mathbf{X}$ è equivalente a $\|\cdot\|_\infty^{(\mathcal{B})}$.

Una delle due disuguaglianze è immediata:

$$\begin{aligned} \left\| \sum_{j=1}^n \alpha_j \mathbf{b}_j \right\| &\leq \sum_{j=1}^n |\alpha_j| \|\mathbf{b}_j\| \leq \max_{1 \leq j \leq n} |\alpha_j| \cdot \sum_{j=1}^n \|\mathbf{b}_j\| \\ &= \sum_{j=1}^n \|\mathbf{b}_j\| \cdot \left\| \sum_{j=1}^n \alpha_j \mathbf{b}_j \right\|_\infty^{(\mathcal{B})}. \end{aligned}$$

Per dimostrare l'altra disuguaglianza, consideriamo la funzione

$$f : \mathbb{R}^n \ni \begin{pmatrix} \alpha_1 \\ \vdots \\ \alpha_n \end{pmatrix} \mapsto \left\| \sum_{j=1}^n \alpha_j \mathbf{b}_j \right\| \in [0, +\infty).$$

La funzione f è continua, perché

$$\begin{aligned} &\left| \left\| \sum_{j=1}^n \alpha_j \mathbf{b}_j \right\| - \left\| \sum_{j=1}^n \beta_j \mathbf{b}_j \right\| \right| \\ (17.2.1) \quad &\leq \left\| \sum_{j=1}^n \alpha_j \mathbf{b}_j - \sum_{j=1}^n \beta_j \mathbf{b}_j \right\| = \left\| \sum_{j=1}^n (\alpha_j - \beta_j) \mathbf{b}_j \right\| \\ &\leq \sum_{j=1}^n |\alpha_j - \beta_j| \|\mathbf{b}_j\| \leq \max_{1 \leq j \leq n} |\alpha_j - \beta_j| \cdot \sum_{j=1}^n \|\mathbf{b}_j\| \\ &= \sum_{j=1}^n \|\mathbf{b}_j\| \cdot \left\| \begin{pmatrix} \alpha_1 \\ \vdots \\ \alpha_n \end{pmatrix} - \begin{pmatrix} \beta_1 \\ \vdots \\ \beta_n \end{pmatrix} \right\|_\infty. \end{aligned}$$

D'altra parte, il sottoinsieme

$$\begin{aligned}
 S_\infty &:= \left\{ \begin{pmatrix} \alpha_1 \\ \vdots \\ \alpha_n \end{pmatrix} \in \mathbb{R}^n; \left\| \begin{pmatrix} \alpha_1 \\ \vdots \\ \alpha_n \end{pmatrix} \right\|_\infty = 1 \right\} \\
 &= \left\{ \begin{pmatrix} \alpha_1 \\ \vdots \\ \alpha_n \end{pmatrix} \in \mathbb{R}^n; |\alpha_j| = 1 \text{ per almeno un } 1 \leq j \leq n \right\} \\
 &= \bigcup_{j=1}^n \left\{ \begin{pmatrix} \alpha_1 \\ \vdots \\ \alpha_n \end{pmatrix} \in \mathbb{R}^n; |\alpha_j| = 1 \right\}
 \end{aligned}$$

di $\mathbb{R}^n$ (la *sfera unitaria* (nel senso di superficie sferica di raggio 1) in $\mathbb{R}^n$ rispetto alla norma $\|\cdot\|_\infty$) è chiaramente chiuso e limitato, quindi compatto. Ora il Teorema di Weierstrass [17.1.23](#) implica l'esistenza di

un punto di minimo $\boldsymbol{\alpha}^{(o)} = \begin{pmatrix} \alpha_1^{(o)} \\ \vdots \\ \alpha_n^{(o)} \end{pmatrix}$ di f su S_∞ . Poiché non tutti gli

$\alpha_j^{(o)}$ possono essere zero e $\mathbf{b}_1, \dots, \mathbf{b}_n$ sono linearmente indipendenti, si ha $c := f(\boldsymbol{\alpha}^{(o)}) > 0$.

Ora, per qualsiasi $\mathbf{0}_n \neq \boldsymbol{\alpha} = \begin{pmatrix} \alpha_1 \\ \vdots \\ \alpha_n \end{pmatrix} \in \mathbb{R}^n$ abbiamo $\|\boldsymbol{\alpha}\|_\infty^{-1} \boldsymbol{\alpha} \in S_\infty$.

Perciò si ha

$$\|\boldsymbol{\alpha}\|_\infty^{-1} \left\| \sum_{j=1}^n \alpha_j \mathbf{b}_j \right\| = f\left(\|\boldsymbol{\alpha}\|_\infty^{-1} \boldsymbol{\alpha}\right) \geq f(\boldsymbol{\alpha}^{(o)}) = c,$$

e quindi

$$\left\| \sum_{j=1}^n \alpha_j \mathbf{b}_j \right\| \geq c \|\boldsymbol{\alpha}\|_\infty = c \left\| \sum_{j=1}^n \alpha_j \mathbf{b}_j \right\|_\infty^{(\mathcal{B})}.$$

Il Teorema di equivalenza [17.2.5](#) implica:

TEOREMA 17.2.6. (Caratterizzazione della compattezza). *Ogni spazio vettoriale normato di dimensione finita è completo. Un sottoinsieme di uno spazio vettoriale normato di dimensione finita è compatto se e solo se è chiuso e limitato.*

DIMOSTRAZIONE. Sia $\mathbf{X}$ uno spazio vettoriale normato di dimensione finita n . Scegliamo una base

$$\mathcal{B} = \{\mathbf{b}_1, \dots, \mathbf{b}_n\}$$

per $\mathbf{X}$ e consideriamo l'applicazione lineare invertibile

$$\mathbb{R}^n \ni \begin{pmatrix} \alpha_1 \\ \vdots \\ \alpha_n \end{pmatrix} \mapsto \sum_{j=1}^n \alpha_j \mathbf{b}_j \in \mathbf{X}.$$

Per il precedente Teorema di equivalenza 17.2.5 le due norme $\alpha \mapsto \|T(\alpha)\|$ e $\|\cdot\|_\infty$ su $\mathbb{R}^n$ sono equivalenti, ossia esistono $0 < c_1 \leq c_2$ tali che

$$c_1 \|\alpha\|_\infty \leq \|T(\alpha)\| \leq c_2 \|\alpha\|_\infty, \quad \alpha \in \mathbb{R}^n.$$

Quindi per una successione $\{\mathbf{x}_k\}_{k \geq 1}$ in $\mathbf{X}$

$$\{\mathbf{x}_k\}_{k \geq 1} \text{ converge ad } \mathbf{x} \iff \{T^{-1}(\mathbf{x}_k)\}_{k \geq 1} \text{ converge a } T^{-1}(\mathbf{x})$$

$$\{\mathbf{x}_k\}_{k \geq 1} \text{ è di Cauchy} \iff \{T^{-1}(\mathbf{x}_k)\}_{k \geq 1} \text{ è di Cauchy.}$$

Di conseguenza la completezza di $\mathbb{R}^n$ implica la completezza di $\mathbf{X}$.

Pertanto, se per $K \subset \mathbf{X}$ poniamo

$$T^{-1}(K) := \{T^{-1}(\mathbf{x}) ; \mathbf{x} \in K\} = \{\alpha \in \mathbb{R}^n ; T(\alpha) \in K\},$$

abbiamo

$$K \text{ è chiuso} \iff T^{-1}(K) \text{ è chiuso,}$$

$$K \text{ è limitato} \iff T^{-1}(K) \text{ è limitato,}$$

$$K \text{ è compatto} \iff T^{-1}(K) \text{ è compatto.}$$

Poiché gli insiemi compatti in $\mathbb{R}^n$ sono esattamente gli insiemi chiusi e limitati (Teorema 17.1.20), per le equivalenze di cui sopra la stessa cosa vale in $\mathbf{X}$.

In particolare, per ogni spazio vettoriale normato $\mathbf{X}$ di dimensione finita la *palla unità chiusa* $\{\mathbf{x} \in \mathbf{X}; \|\mathbf{x}\| \leq 1\}$ e la *sfera unità* $\{\mathbf{x} \in \mathbf{X}; \|\mathbf{x}\| = 1\}$ sono compatte.

è naturale porsi la seguente domanda. Dato uno spazio vettoriale normato $\mathbf{X}$ di dimensione finita n , possiamo trovare una base $\mathcal{B}$ di $\mathbf{X}$ tale che nelle disuguaglianze (17.2.2), che esprimono l'equivalenza di $\|\cdot\|_1^{(\mathcal{B})}$ o $\|\cdot\|_\infty^{(\mathcal{B})}$ con la norma data di $\mathbf{X}$, le costanti siano quanto possibile di vicine ad 1 e la loro dipendenza da n sia *esplicita*? Una risposta è data dal

TEOREMA 17.2.7. (Auerbach.) *Sia $\mathbf{X}$ uno spazio vettoriale normato di dimensione finita n . Allora esiste una base*

$$\mathcal{B} = \{\mathbf{b}_1, \dots, \mathbf{b}_n\}$$

di $\mathbf{X}$ tale che

$$\|\mathbf{x}\|_\infty^{(\mathcal{B})} \leq \|\mathbf{x}\| \leq \|\mathbf{x}\|_1^{(\mathcal{B})}, \quad \mathbf{x} \in \mathbf{X},$$

ovvero

$$\|\mathbf{b}_k\| = 1, \quad 1 \leq k \leq n,$$

$$|\lambda_k| \leq \left\| \sum_{j=1}^n \lambda_j \mathbf{b}_j \right\|, \quad 1 \leq k \leq n, \quad \lambda_1, \dots, \lambda_n \in \mathbb{R}.$$

DIMOSTRAZIONE. Fissiamo anzitutto una base

$$\mathcal{E} = \{\mathbf{e}_1, \dots, \mathbf{e}_n\}$$

di $\mathbf{X}$, e consideriamo sullo spazio vettoriale prodotto

$$\mathbf{X}^n = \underbrace{\mathbf{X} \times \dots \times \mathbf{X}}_{n \text{ volte}} = \underbrace{\mathbf{X} \oplus \dots \oplus \mathbf{X}}_{n \text{ volte}}$$

la norma $\|(\mathbf{x}_1, \dots, \mathbf{x}_n)\| := \max_{1 \leq k \leq n} \|\mathbf{x}_k\|$. È immediato verificare che la palla unità chiusa di $\mathbf{X}^n$ è $(\mathbf{X}_1)^n = \underbrace{\mathbf{X}_1 \times \dots \times \mathbf{X}_1}_{n \text{ volte}}$, ove $\mathbf{X}_1$ indica la palla unità chiusa di $\mathbf{X}$. Per il Teorema 17.2.6 sulla caratterizzazione della compattezza, $(\mathbf{X}_1)^n$ è compatto.

Consideriamo la funzione

$$f : \mathbf{X}^n \ni \left(\sum_{j=1}^n \alpha_{j1} \mathbf{e}_j, \dots, \sum_{j=1}^n \alpha_{jn} \mathbf{e}_j \right) \mapsto \left| \det \begin{pmatrix} \alpha_{11} & \dots & \alpha_{1n} \\ \vdots & \ddots & \vdots \\ \alpha_{n1} & \dots & \alpha_{nn} \end{pmatrix} \right|.$$

Notiamo che $f(\mathbf{x}_1, \dots, \mathbf{x}_n)$ è il valore assoluto del determinante della matrice che ha per colonna k -esima i coefficienti dello sviluppo di $\mathbf{x}_k$ rispetto alla base $\mathcal{E}$. Tenendo conto che la norma di $\mathbf{X}$ è equivalente a $\|\cdot\|_{\infty}^{(\mathcal{B})}$, si vede subito che la funzione f è continua. Perciò il Teorema di Weierstrass 17.1.23 implica che esiste un punto di massimo $(\mathbf{b}_1, \dots, \mathbf{b}_n)$ per f sull'insieme compatto $(\mathbf{X}_1)^n$.

Poiché

$$\begin{aligned} f(\mathbf{b}_1, \dots, \mathbf{b}_n) &\geq f(\|\mathbf{e}_1\|^{-1} \mathbf{e}_1, \dots, \|\mathbf{e}_n\|^{-1} \mathbf{e}_n) \\ &= \left| \det \begin{pmatrix} \|\mathbf{e}_1\|^{-1} & \dots & 0 \\ \vdots & \ddots & \vdots \\ 0 & \dots & \|\mathbf{e}_n\|^{-1} \end{pmatrix} \right| \\ &= (\|\mathbf{e}_1\| \dots \|\mathbf{e}_n\|)^{-1} > 0, \end{aligned}$$

si ha che, se si pone $\mathbf{b}_k = \sum_{j=1}^n \beta_{jk} \mathbf{e}_j$ per $1 \leq k \leq n$, la matrice $B = (\beta_{jk})_{1 \leq j, k \leq n}$ soddisfa $\det B \neq 0$. Perciò

$$\sum_{j=1}^n \lambda_j \mathbf{b}_j = \mathbf{0}_{\mathbf{X}} \implies B \cdot \begin{pmatrix} \lambda_1 \\ \vdots \\ \lambda_n \end{pmatrix} = \begin{pmatrix} 0 \\ \vdots \\ 0 \end{pmatrix} \implies \begin{pmatrix} \lambda_1 \\ \vdots \\ \lambda_n \end{pmatrix} = \begin{pmatrix} 0 \\ \vdots \\ 0 \end{pmatrix},$$

ossia i vettori $\mathbf{b}_1, \dots, \mathbf{b}_n$ sono linearmente indipendenti. Perciò

$$\mathcal{B} = \{\mathbf{b}_1, \dots, \mathbf{b}_n\}$$

è una base di $\mathbf{X}$.

Per ogni $1 \leq k \leq n$ abbiamo $\|\mathbf{b}_k\| = 1$. Infatti, se fosse $\|\mathbf{b}_k\| < 1$, allora avremmo

$$f(\mathbf{b}_1, \dots, \|\mathbf{b}_k\|^{-1} \mathbf{b}_k, \dots, \mathbf{b}_n) = \|\mathbf{b}_k\|^{-1} f(\mathbf{b}_1, \dots, \mathbf{b}_n) > f(\mathbf{b}_1, \dots, \mathbf{b}_n),$$

in contraddizione col fatto che $(\mathbf{b}_1, \dots, \mathbf{b}_n)$ è un punto di massimo per f su $(\mathbf{X}_1)^n$.

Ora dimostriamo che

$$|\lambda_{k_0}| \leq \left\| \sum_{k=1}^n \lambda_k \mathbf{b}_k \right\|, \quad 1 \leq k_0 \leq n, \lambda_1, \dots, \lambda_n \in \mathbb{R}. \quad (17.2.3)$$

Per assurdo, supponiamo il contrario, cioè che esista un indice $1 \leq k_0 \leq n$ ed una combinazione lineare $\mathbf{x} = \sum_{k=1}^n \lambda_k \mathbf{b}_k$ tale che $|\lambda_{k_0}| > \|\mathbf{x}\|$.

Poiché λ_{k_0} non si annulla, abbiamo $\mathbf{x} \neq \mathbf{0}_{\mathbf{X}}$. Perciò possiamo dividere tutti i λ_k per $\|\mathbf{x}\|$, cioè assumere che $|\lambda_{k_0}| > \|\mathbf{x}\| = 1$. Ora, usando la linearità dei determinanti rispetto ad ogni colonna ed il fatto che si annullano nel caso di due colonne uguali, otteniamo

$$\begin{aligned} & f(\mathbf{b}_1, \dots, \mathbf{b}_{k_0-1}, \mathbf{x}, \mathbf{b}_{k_0+1}, \dots, \mathbf{b}_n) \\ &= \left| \det \begin{pmatrix} \beta_{11} & \dots & \beta_{1,k_0-1} & \sum_{k=1}^n \lambda_k \beta_{1k} & \beta_{1,k_0+1} & \dots & \beta_{1n} \\ \vdots & \ddots & \vdots & \vdots & \vdots & \ddots & \vdots \\ \beta_{n1} & \dots & \beta_{n,k_0-1} & \sum_{k=1}^n \lambda_k \beta_{nk} & \beta_{n,k_0+1} & \dots & \beta_{nn} \end{pmatrix} \right| \\ &= \left| \det \begin{pmatrix} \beta_{11} & \dots & \beta_{1,k_0-1} & \lambda_{k_0} \beta_{1k_0} & \beta_{1,k_0+1} & \dots & \beta_{1n} \\ \vdots & \ddots & \vdots & \vdots & \vdots & \ddots & \vdots \\ \beta_{n1} & \dots & \beta_{n,k_0-1} & \lambda_{k_0} \beta_{nk_0} & \beta_{n,k_0+1} & \dots & \beta_{nn} \end{pmatrix} \right| \\ &= |\lambda_{k_0}| \left| \det \begin{pmatrix} \beta_{11} & \dots & \beta_{1k_0} & \dots & \beta_{1n} \\ \vdots & \ddots & \vdots & \ddots & \vdots \\ \beta_{n1} & \dots & \beta_{nk_0} & \dots & \beta_{nn} \end{pmatrix} \right| \\ &= \underbrace{|\lambda_{k_0}|}_{>1} f(\mathbf{b}_1, \dots, \mathbf{b}_{k_0}, \dots, \mathbf{b}_n) > f(\mathbf{b}_1, \dots, \mathbf{b}_n), \end{aligned}$$

in contraddizione con il fatto che $(\mathbf{b}_1, \dots, \mathbf{b}_n)$ è un punto di massimo di f su $(\mathbf{X}_1)^n$.

Per completare la dimostrazione basta mostrare che le uguaglianze (17.2.3) e $\|\mathbf{b}_k\| = 1$ per $1 \leq k \leq n$ sono equivalenti a

$$\|\mathbf{x}\|_{\infty}^{(\mathcal{B})} \leq \|\mathbf{x}\| \leq \|\mathbf{x}\|_1^{(\mathcal{B})}, \quad \mathbf{x} \in \mathbf{X}. \quad (17.2.4)$$

In effetti le uguaglianze $\|\mathbf{b}_k\| = 1$ per $1 \leq k \leq n$ implicano

$$\left\| \sum_{j=1}^n \lambda_j \mathbf{b}_j \right\| \leq \sum_{j=1}^n |\lambda_j| = \left\| \sum_{j=1}^n \lambda_j \mathbf{b}_j \right\|_1^{(\mathcal{B})},$$

e da (17.2.3) segue che

$$\left\| \sum_{j=1}^n \lambda_j \mathbf{b}_j \right\| \geq \max_{1 \leq j \leq n} |\lambda_j| = \left\| \sum_{j=1}^n \lambda_j \mathbf{b}_j \right\|_{\infty}^{(\mathcal{B})}.$$

Viceversa, da una parte (17.2.4) implica che $1 = \|\mathbf{b}_j\|_{\infty}^{(\mathcal{B})} \leq \|\mathbf{b}_j\| \leq \|\mathbf{b}_j\|_1^{(\mathcal{B})} = 1$, cioè $\|\mathbf{b}_j\| = 1$ per ogni $1 \leq j \leq n$. D'altra parte, la seconda disuguaglianza in (17.2.4) implica (17.2.3).

17.3. Prodotti scalari su spazi vettoriali reali

Sia $\mathbf{X}$ uno spazio vettoriale reale. Una funzione $\theta : \mathbf{X} \times \mathbf{X} \longrightarrow \mathbb{R}$ si dice *forma biadditiva* su $\mathbf{X}$ se è additiva rispetto ad entrambe le variabili :

$$\begin{aligned} \theta(\mathbf{x}_1, \mathbf{y}) + \theta(\mathbf{x}_2, \mathbf{y}) &= \theta(\mathbf{x}_1 + \mathbf{x}_2, \mathbf{y}), & \mathbf{x}_1, \mathbf{x}_2, \mathbf{y} \in \mathbf{X}, \\ \theta(\mathbf{x}, \mathbf{y}_1) + \theta(\mathbf{x}, \mathbf{y}_2) &= \theta(\mathbf{x}, \mathbf{y}_1 + \mathbf{y}_2), & \mathbf{x}, \mathbf{y}_1, \mathbf{y}_2 \in \mathbf{X}, \end{aligned}$$

e si dice *biomogenea* se è omogenea rispetto ad entrambe le variabili :

$$\theta(\lambda \mathbf{x}, \mathbf{y}) = \lambda \theta(\mathbf{x}, \mathbf{y}) = \theta(\mathbf{x}, \lambda \mathbf{y}), \quad \mathbf{x}, \mathbf{y} \in \mathbf{X}, \lambda \in \mathbb{R}$$

Una forma biadditiva e biomogenea su $\mathbf{X}$ si chiama *forma bilineare* su $\mathbf{X}$.

D'altra parte, una funzione $\theta : \mathbf{X} \times \mathbf{X} \longrightarrow \mathbb{R}$ si dice *simmetrica* se

$$\theta(\mathbf{y}, \mathbf{x}) = \theta(\mathbf{x}, \mathbf{y}), \quad \mathbf{x}, \mathbf{y} \in \mathbf{X}.$$

Esempi:

1) Su $\mathbb{R}^n$ possiamo definire una forma bilineare simmetrica θ tramite la

$$\text{formula } \theta(\mathbf{x}, \mathbf{y}) = \sum_{j=1}^n x_j y_j, \text{ ove } \mathbf{x} = \begin{pmatrix} x_1 \\ \vdots \\ x_n \end{pmatrix} \text{ e } \mathbf{y} = \begin{pmatrix} y_1 \\ \vdots \\ y_n \end{pmatrix}.$$

2) La forma bilineare

$$\mathbb{R}^2 \times \mathbb{R}^2 \ni \left(\begin{pmatrix} x_1 \\ x_2 \end{pmatrix}, \begin{pmatrix} y_1 \\ y_2 \end{pmatrix} \right) \longmapsto x_1 y_1 + x_2 y_2 + x_1 y_2$$

non è simmetrica.

Rimarchiamo che la biaddittività basta per la biomogeneità rispetto alla moltiplicazione con i scalari razionali:

Lemma 1. *Siano $\mathbf{X}$ uno spazio vettoriale reale e θ una forma biadditiva su $\mathbf{X}$. Allora*

$$\theta(\lambda \mathbf{x}, \mathbf{y}) = \lambda \theta(\mathbf{x}, \mathbf{y}) = \theta(\mathbf{x}, \lambda \mathbf{y}), \quad \mathbf{x}, \mathbf{y} \in \mathbf{X}, \lambda \in \mathbb{Q}.$$

Dimostrazione. Usando l'additività di θ rispetto alla prima variabile, per ogni $\mathbf{x}, \mathbf{y} \in \mathbf{X}$ ed ogni intero $m \geq 1$ si deduce successivamente

$$\theta(m \mathbf{x}, \mathbf{y}) = \theta(\underbrace{\mathbf{x} + \dots + \mathbf{x}}_{m \text{ volte}}, \mathbf{y}) = \underbrace{\theta(\mathbf{x}, \mathbf{y}) + \dots + \theta(\mathbf{x}, \mathbf{y})}_{m \text{ volte}} = m \theta(\mathbf{x}, \mathbf{y}),$$

$$m \theta\left(\frac{1}{m} \mathbf{x}, \mathbf{y}\right) = \theta\left(m \frac{1}{m} \mathbf{x}, \mathbf{y}\right) = \theta(\mathbf{x}, \mathbf{y}), \text{ quindi}$$

$$\theta\left(\frac{1}{m} \mathbf{x}, \mathbf{y}\right) = \frac{1}{m} \theta(\mathbf{x}, \mathbf{y}).$$

Cosicch 

$$\theta\left(\frac{n}{m} \mathbf{x}, \mathbf{y}\right) = n \theta\left(\frac{1}{m} \mathbf{x}, \mathbf{y}\right) = n \frac{1}{m} \theta(\mathbf{x}, \mathbf{y}) = \frac{n}{m} \theta(\mathbf{x}, \mathbf{y})$$

per ogni $\mathbf{x} \in \mathbf{X}$, $n \in \mathbb{Z}$ e $m \geq 1$ intero.

Similmente si dimostra anche l'identit  $\theta\left(\mathbf{x}, \frac{n}{m} \mathbf{y}\right) = \frac{n}{m} \theta(\mathbf{x}, \mathbf{y})$. □

Per ogni forma biadditiva θ su $\mathbf{X}$ vale l'identit  del parallelogramma :

$$\theta(\mathbf{x} + \mathbf{y}, \mathbf{x} + \mathbf{y}) + \theta(\mathbf{x} - \mathbf{y}, \mathbf{x} - \mathbf{y}) = 2 \theta(\mathbf{x}, \mathbf{x}) + 2 \theta(\mathbf{y}, \mathbf{y}), \quad \mathbf{x}, \mathbf{y} \in \mathbf{X}.$$

Per la dimostrazione basta sommare le uguaglianze

$$\begin{aligned} \theta(\mathbf{x} + \mathbf{y}, \mathbf{x} + \mathbf{y}) &= \theta(\mathbf{x}, \mathbf{x}) + \theta(\mathbf{x}, \mathbf{y}) + \theta(\mathbf{y}, \mathbf{x}) + \theta(\mathbf{y}, \mathbf{y}), \\ \theta(\mathbf{x} - \mathbf{y}, \mathbf{x} - \mathbf{y}) &= \theta(\mathbf{x}, \mathbf{x}) - \theta(\mathbf{x}, \mathbf{y}) - \theta(\mathbf{y}, \mathbf{x}) + \theta(\mathbf{y}, \mathbf{y}). \end{aligned}$$

Perci  la funzione

$$q_\theta : \mathbf{X} \ni \mathbf{x} \longmapsto \theta(\mathbf{x}, \mathbf{x}) \in \mathbb{R},$$

verifica l'identit  del parallelogramma sotto la forma

$$q_\theta(\mathbf{x} + \mathbf{y}) + q_\theta(\mathbf{x} - \mathbf{y}) = 2 q_\theta(\mathbf{x}) + 2 q_\theta(\mathbf{y}), \quad \mathbf{x}, \mathbf{y} \in \mathbf{X}.$$

Se θ   biomogenea, allora q_θ soddisfa la condizione di omogeneit  quadratica ovvia

$$q_\theta(\lambda \mathbf{x}) = \lambda^2 q_\theta(\mathbf{x}), \quad \mathbf{x} \in \mathbf{X}, \lambda \in \mathbb{R}.$$

Chiamiamo *forma quadratica* su $\mathbf{X}$ ogni funzione $q : \mathbf{X} \rightarrow \mathbb{R}$ che soddisfa l'identità del parallelogramma e la condizione di omogeneità quadratica :

$$\begin{aligned} q(\mathbf{x} + \mathbf{y}) + q(\mathbf{x} - \mathbf{y}) &= 2q(\mathbf{x}) + 2q(\mathbf{y}), & \mathbf{x}, \mathbf{y} \in X, \\ q(\lambda \mathbf{x}) &= \lambda^2 q(\mathbf{x}), & \mathbf{x} \in \mathbf{X}, \lambda \in \mathbb{R}. \end{aligned}$$

Chiaramente, ogni forma quadratica è pari e si annulla in $\mathbf{0}_X$:

$$\begin{aligned} q(-\mathbf{x}) &= q((-1)\mathbf{x}) = (-1)^2 q(\mathbf{x}) = q(\mathbf{x}), \\ q(\mathbf{0}_X) &= q(0 \cdot \mathbf{0}_X) = 0^2 q(\mathbf{0}_X) = 0. \end{aligned}$$

Se θ è una forma biadditiva simmetrica su $\mathbf{X}$, allora vale la *formula di polarizzazione* :

$$\theta(\mathbf{x}, \mathbf{y}) = \frac{1}{4} \left(\theta(\mathbf{x} + \mathbf{y}, \mathbf{x} + \mathbf{y}) - \theta(\mathbf{x} - \mathbf{y}, \mathbf{x} - \mathbf{y}) \right), \quad \mathbf{x}, \mathbf{y} \in X.$$

Per la dimostrazione basta considerare la differenza delle identità

$$\begin{aligned} \theta(\mathbf{x} + \mathbf{y}, \mathbf{x} + \mathbf{y}) &= \theta(\mathbf{x}, \mathbf{x}) + 2\theta(\mathbf{x}, \mathbf{y}) + \theta(\mathbf{y}, \mathbf{y}), \\ \theta(\mathbf{x} - \mathbf{y}, \mathbf{x} - \mathbf{y}) &= \theta(\mathbf{x}, \mathbf{x}) - 2\theta(\mathbf{x}, \mathbf{y}) + \theta(\mathbf{y}, \mathbf{y}). \end{aligned}$$

Perciò q_θ determina θ unicamente :

$$\theta(\mathbf{x}, \mathbf{y}) = \frac{1}{4} \left(q_\theta(\mathbf{x} + \mathbf{y}) - q_\theta(\mathbf{x} - \mathbf{y}) \right), \quad \mathbf{x}, \mathbf{y} \in X.$$

Diciamo che una forma biadditiva θ su $\mathbf{X}$ è

semidefinita positiva se è simmetrica e

$$\theta(\mathbf{x}, \mathbf{x}) \geq 0 \text{ per ogni } \mathbf{x} \in X,$$

definita positiva se è simmetrica e

$$\theta(\mathbf{x}, \mathbf{x}) > 0 \text{ per ogni } \mathbf{0}_X \neq \mathbf{x} \in X.$$

Notiamo che la proprietà $\theta(\mathbf{x}, \mathbf{x}) > 0$, $\mathbf{0}_X \neq \mathbf{x} \in X$, non implica la simmetria di θ , nemmeno quando θ è bilineare: vedi la forma bilineare dell'Esempio 2) di cui sopra.

Possiamo chiederci : è vero che per qualsiasi forma quadratica q la formula di polarizzazione

$$\theta(\mathbf{x}, \mathbf{y}) = \frac{1}{4} \left(q(\mathbf{x} + \mathbf{y}) - q(\mathbf{x} - \mathbf{y}) \right), \quad \mathbf{x}, \mathbf{y} \in X$$

definisce una forma bilineare θ su $\mathbf{X}$? Secondo il prossimo lemma manca poco che la risposta sia completamente affermativa :

Lemma 2. *Siano $\mathbf{X}$ uno spazio vettoriale reale e $q : \mathbf{X} \longrightarrow \mathbb{R}$ una forma quadratica. Allora la formula di polarizzazione*

$$\theta(\mathbf{x}, \mathbf{y}) = \frac{1}{4} \left(q(\mathbf{x} + \mathbf{y}) - q(\mathbf{x} - \mathbf{y}) \right), \quad \mathbf{x}, \mathbf{y} \in \mathbf{X}.$$

definisce una forma $\theta : \mathbf{X} \times \mathbf{X} \longrightarrow \mathbb{R}$ biadditiva e simmetrica :

$$\begin{aligned} \theta(\mathbf{x}_1, \mathbf{y}) + \theta(\mathbf{x}_2, \mathbf{y}) &= \theta(\mathbf{x}_1 + \mathbf{x}_2, \mathbf{y}), & \mathbf{x}_1, \mathbf{x}_2, \mathbf{y} \in \mathbf{X}, \\ \theta(\mathbf{x}, \mathbf{y}_1) + \theta(\mathbf{x}, \mathbf{y}_2) &= \theta(\mathbf{x}, \mathbf{y}_1 + \mathbf{y}_2), & \mathbf{x}, \mathbf{y}_1, \mathbf{y}_2 \in \mathbf{X}, \\ \theta(\mathbf{y}, \mathbf{x}) &= \theta(\mathbf{x}, \mathbf{y}), & \mathbf{x}, \mathbf{y} \in \mathbf{X}, \end{aligned}$$

che ridefinisce q :

$$q(\mathbf{x}) = \theta(\mathbf{x}, \mathbf{x}), \quad \mathbf{x} \in \mathbf{X}.$$

Dimostrazione. Usando che q è pari e si annulla in $\mathbf{0}_{\mathbf{X}}$, deduciamo

$$\theta(\mathbf{y}, \mathbf{x}) = \frac{1}{4} \left(q(\mathbf{y} + \mathbf{x}) - q(\mathbf{y} - \mathbf{x}) \right) = \frac{1}{4} \left(q(\mathbf{x} + \mathbf{y}) - q(\mathbf{x} - \mathbf{y}) \right) = \theta(\mathbf{x}, \mathbf{y}),$$

$$\theta(\mathbf{0}_{\mathbf{X}}, \mathbf{x}) = \frac{1}{4} \left(q(\mathbf{x}) - q(-\mathbf{x}) \right) = 0,$$

$$\theta(\mathbf{x}, \mathbf{x}) = \frac{1}{4} \left(q(\mathbf{x} + \mathbf{x}) - q(\mathbf{x} - \mathbf{x}) \right) = \frac{1}{4} q(2\mathbf{x}) = q(\mathbf{x}).$$

Per ogni $\mathbf{x}, \mathbf{y}, \mathbf{z} \in \mathbf{X}$, usando l'identità del parallelogramma si ottiene

$$\begin{aligned} 2\theta(\mathbf{x} + \mathbf{y}, 2\mathbf{z}) &= q(\mathbf{x} + \mathbf{y} + 2\mathbf{z}) - q(\mathbf{x} + \mathbf{y} - 2\mathbf{z}) \\ &= q((\mathbf{x} + \mathbf{z}) + (\mathbf{y} + \mathbf{z})) + q((\mathbf{x} + \mathbf{z}) - (\mathbf{y} + \mathbf{z})) \\ &\quad - q((\mathbf{x} - \mathbf{z}) + (\mathbf{y} - \mathbf{z})) - q((\mathbf{x} - \mathbf{z}) - (\mathbf{y} - \mathbf{z})) \\ &= \left(2q(\mathbf{x} + \mathbf{z}) + 2q(\mathbf{y} + \mathbf{z}) \right) - \left(2q(\mathbf{x} - \mathbf{z}) + 2q(\mathbf{y} - \mathbf{z}) \right) \\ &= 4\theta(\mathbf{x}, \mathbf{z}) + 4\theta(\mathbf{y}, \mathbf{z}), \end{aligned}$$

quindi

$$\theta(\mathbf{x} + \mathbf{y}, 2\mathbf{z}) = 2\theta(\mathbf{x}, \mathbf{z}) + 2\theta(\mathbf{y}, \mathbf{z}). \quad (17.3.1)$$

Ponendo in questa uguaglianza $\mathbf{y} = \mathbf{0}_{\mathbf{X}}$, risulta

$$\theta(\mathbf{x}, 2\mathbf{z}) = 2\theta(\mathbf{x}, \mathbf{z}), \quad \mathbf{x}, \mathbf{z} \in \mathbf{X}$$

e così possiamo cambiare nell'uguaglianza (17.3.1) $\theta(\mathbf{x} + \mathbf{y}, 2\mathbf{z})$ con $2\theta(\mathbf{x} + \mathbf{y}, \mathbf{z})$, ottenendo

$$\theta(\mathbf{x} + \mathbf{y}, \mathbf{z}) = \theta(\mathbf{x}, \mathbf{z}) + \theta(\mathbf{y}, \mathbf{z}), \quad \mathbf{x}, \mathbf{y}, \mathbf{z} \in \mathbf{X}.$$

Così abbiamo dimostrato l'additività di θ rispetto alla prima variabile e per la simmetria di θ risulta anche la sua additività rispetto alla seconda variabile. $\square$

Sorprendentemente, la forma θ ottenuta nel lemma precedente non è in generale bilineare. Un controesempio (non elementare) nel caso $\mathbf{X} = \mathbb{R}^2$ si trova in [5].

Per questa ragione spesso si chiamano forme quadratiche solo le funzioni del tipo q_θ con θ una forma bilineare. In questo caso $q_\theta = q_{\theta_{\text{simm}}}$ per l'unica forma bilineare simmetrica θ_{simm} data dalla formula

$$\theta_{\text{simm}}(\mathbf{x}, \mathbf{y}) = \frac{1}{2} (\theta(\mathbf{x}, \mathbf{y}) + \theta(\mathbf{y}, \mathbf{x})).$$

Ciò nonostante vedremo che se i valori presi dalla forma quadratica q sono contenuti in $[0, +\infty)$, allora la formula di polarizzazione definisce una forma bilineare θ (che sarà automaticamente semidefinita positiva).

Come preparazione, proviamo due disuguaglianze per forme biaddittive positivamente semidefinite :

La disuguaglianza di Schwarz. *Siano $\mathbf{X}$ uno spazio vettoriale reale e $\theta : \mathbf{X} \times \mathbf{X} \rightarrow \mathbb{R}$ una forma biaddittiva semidefinita positiva. Allora*

$$\theta(\mathbf{x}, \mathbf{y})^2 \leq \theta(\mathbf{x}, \mathbf{x}) \theta(\mathbf{y}, \mathbf{y}), \quad \mathbf{x}, \mathbf{y} \in \mathbf{X}.$$

Dimostrazione. Siano $\mathbf{x}, \mathbf{y} \in \mathbf{X}$ vettori arbitrari. Poiché nel caso $\theta(\mathbf{x}, \mathbf{y}) = 0$ la disuguaglianza è ovvia, resta a trattare il caso $\theta(\mathbf{x}, \mathbf{y}) \neq 0$.

Usando la biaddittività di θ e Lemma 1, deduciamo per ogni $\alpha, \beta \in \mathbb{Q}$:

$$0 \leq \theta(\alpha \mathbf{x} + \beta \mathbf{y}, \alpha \mathbf{x} + \beta \mathbf{y}) = \alpha^2 \theta(\mathbf{x}, \mathbf{x}) + \beta^2 \theta(\mathbf{y}, \mathbf{y}) + 2\alpha\beta \theta(\mathbf{x}, \mathbf{y}).$$

Poiché ogni numero reale è limite di una successione di numeri razionali, per passaggio a limite otteniamo che

$$\alpha^2 \theta(\mathbf{x}, \mathbf{x}) + \beta^2 \theta(\mathbf{y}, \mathbf{y}) + 2\alpha\beta \theta(\mathbf{x}, \mathbf{y}) \geq 0, \quad \alpha, \beta \in \mathbb{R}.$$

Ponendo ora $\alpha = -\frac{|\theta(\mathbf{x}, \mathbf{y})|}{\theta(\mathbf{x}, \mathbf{x})} t$, $\beta = s$, ove $t, s \in \mathbb{R}$ sono arbitrari, risulta che

$$t^2 \theta(\mathbf{x}, \mathbf{x}) + s^2 \theta(\mathbf{y}, \mathbf{y}) - 2ts |\theta(\mathbf{x}, \mathbf{y})| \geq 0$$

per ogni $t, s \in \mathbb{R}$.

Se fosse $\theta(\mathbf{x}, \mathbf{x}) = 0$, allora con $s = 1$ risulterebbe $\theta(\mathbf{y}, \mathbf{y}) - 2t |\theta(\mathbf{x}, \mathbf{y})| \geq 0$ per ogni $t \in \mathbb{R}$, il che non è vero per $t > 0$ abbastanza grande. Cosicché $\theta(\mathbf{x}, \mathbf{x}) > 0$. Similmente, $\theta(\mathbf{y}, \mathbf{y}) > 0$.

Ora, con $t = \theta(\mathbf{y}, \mathbf{y})^{1/2}$ ed $s = \theta(\mathbf{x}, \mathbf{x})^{1/2}$ risulta

$$\theta(\mathbf{y}, \mathbf{y}) \theta(\mathbf{x}, \mathbf{x}) + \theta(\mathbf{x}, \mathbf{x}) \theta(\mathbf{y}, \mathbf{y}) - 2 \theta(\mathbf{y}, \mathbf{y})^{1/2} \theta(\mathbf{x}, \mathbf{x})^{1/2} |\theta(\mathbf{x}, \mathbf{y})| \geq 0$$

e semplificando con $2 \theta(\mathbf{y}, \mathbf{y})^{1/2} \theta(\mathbf{x}, \mathbf{x})^{1/2} > 0$ concludiamo :

$$\theta(\mathbf{x}, \mathbf{x})^{1/2} \theta(\mathbf{y}, \mathbf{y})^{1/2} - |\theta(\mathbf{x}, \mathbf{y})| \geq 0. \quad \square$$

La disuguaglianza di Schwarz nel caso della forma bilineare definita positiva dell'Esempio 1) è la classica disuguaglianza di Cauchy-Buniakovski :

$$\left(\sum_{j=1}^n \alpha_j \beta_j \right)^2 \leq \sum_{j=1}^n \alpha_j^2 \sum_{j=1}^n \beta_j^2, \quad \alpha_1, \dots, \alpha_n, \beta_1, \dots, \beta_n \in \mathbb{R}. \quad (17.3.2)$$

La disuguaglianza di Minkowski. Siano $\mathbf{X}$ uno spazio vettoriale reale e $\theta : \mathbf{X} \times \mathbf{X} \rightarrow \mathbb{R}$ una forma biadditiva semidefinita positiva. Allora

$$\theta(\mathbf{x} + \mathbf{y}, \mathbf{x} + \mathbf{y})^{1/2} \leq \theta(\mathbf{x}, \mathbf{x})^{1/2} + \theta(\mathbf{y}, \mathbf{y})^{1/2}, \quad \mathbf{x}, \mathbf{y} \in \mathbf{X}.$$

Dimostrazione. Usando la disuguaglianza di Schwarz, deduciamo per ogni $\mathbf{x}, \mathbf{y} \in \mathbf{X}$

$$\begin{aligned} \theta(\mathbf{x} + \mathbf{y}, \mathbf{x} + \mathbf{y}) &= \theta(\mathbf{x}, \mathbf{x}) + \theta(\mathbf{y}, \mathbf{y}) + \theta(\mathbf{x}, \mathbf{y}) + \theta(\mathbf{y}, \mathbf{x}) \\ &\leq \theta(\mathbf{x}, \mathbf{x}) + \theta(\mathbf{y}, \mathbf{y}) + 2 \theta(\mathbf{x}, \mathbf{x})^{1/2} \theta(\mathbf{y}, \mathbf{y})^{1/2} \\ &= \left(\theta(\mathbf{x}, \mathbf{x})^{1/2} + \theta(\mathbf{y}, \mathbf{y})^{1/2} \right)^2. \quad \square \end{aligned}$$

Nel caso della forma bilineare definita positiva dell'Esempio 1) la disuguaglianza di Minkowski prende la forma

$$\left(\sum_{j=1}^n (\alpha_j + \beta_j)^2\right)^{1/2} \leq \left(\sum_{j=1}^n \alpha_j^2\right)^{1/2} + \left(\sum_{j=1}^n \beta_j^2\right)^{1/2}, \quad (17.3.3)$$

$$\alpha_1, \dots, \alpha_n, \beta_1, \dots, \beta_n \in \mathbb{R}.$$

Teorema di Jordan - von Neumann. *Per ogni spazio vettoriale reale $\mathbf{X}$ e forma quadratica $q : \mathbf{X} \rightarrow [0, +\infty)$ la formula di polarizzazione*

$$\theta(\mathbf{x}, \mathbf{y}) = \frac{1}{4} \left(q(\mathbf{x} + \mathbf{y}) - q(\mathbf{x} - \mathbf{y}) \right), \quad \mathbf{x}, \mathbf{y} \in X.$$

definisce una forma bilineare semidefinita positiva $\theta : \mathbf{X} \times \mathbf{X} \rightarrow \mathbb{R}$ tale che

$$q(\mathbf{x}) = \theta(\mathbf{x}, \mathbf{x}), \quad \mathbf{x} \in \mathbf{X}.$$

Dimostrazione. Per Lemma 2 θ è una forma diadditiva semidefinita positiva tale che $q(\mathbf{x}) = \theta(\mathbf{x}, \mathbf{x})$, $\mathbf{x} \in \mathbf{X}$. Dobbiamo provare ancora che θ è biomogenea.

Siano $\mathbf{x}, \mathbf{y} \in X$ e $\lambda \in \mathbb{R}$ arbitrari. Esiste una successione $\{r_k\}_{k \geq 1}$ di numeri razionali tale che $r_k \rightarrow \lambda$. Per la disuguaglianza di Schwarz abbiamo per ogni $k \geq 1$

$$\begin{aligned} |\theta(\lambda \mathbf{x}, \mathbf{y}) - \theta(r_k \mathbf{x}, \mathbf{y})| &= |\theta((\lambda - r_k) \mathbf{x}, \mathbf{y})| \\ &\leq \theta((\lambda - r_k) \mathbf{x}, (\lambda - r_k) \mathbf{x})^{1/2} \theta(\mathbf{y}, \mathbf{y})^{1/2} \\ &= |\lambda - r_k| \theta(\mathbf{x}, \mathbf{x})^{1/2} \theta(\mathbf{y}, \mathbf{y})^{1/2}. \end{aligned}$$

Perciò $\theta(r_k \mathbf{x}, \mathbf{y}) \rightarrow \theta(\lambda \mathbf{x}, \mathbf{y})$. Poiché, per Lemma 1,

$$\theta(r_k \mathbf{x}, \mathbf{y}) = r_k \theta(\mathbf{x}, \mathbf{y}) \text{ per ogni } k \geq 1,$$

passando al limite per $k \rightarrow \infty$, risulta $\theta(\lambda \mathbf{x}, \mathbf{y}) = \lambda \theta(\mathbf{x}, \mathbf{y})$.

Finalmente, per la simmetria di θ , l'omogeneità di θ rispetto alla prima variabile implica la sua omogeneità anche rispetto alla seconda variabile. $\square$

Sia $\mathbf{X}$ uno spazio vettoriale reale. Per il teorema di Jordan - von Neumann e per la disuguaglianza di Minkowski esiste una corrispondenza biunivoca fra

- forme bilineari $\theta : \mathbf{X} \times \mathbf{X} \longrightarrow \mathbb{R}$ positivamente definite,
- forme quadratiche $q : \mathbf{X} \longrightarrow [0, +\infty)$ che si annullano solo in $\mathbf{0}_{\mathbf{X}}$,
- norme $\mathbf{X} \ni \mathbf{x} \longmapsto \|\mathbf{x}\|$ che soddisfano l'identità del parallelogramma

$$\|\mathbf{x} + \mathbf{y}\|^2 + \|\mathbf{x} - \mathbf{y}\|^2 = 2\|\mathbf{x}\|^2 + 2\|\mathbf{y}\|^2, \quad \mathbf{x}, \mathbf{y} \in \mathbf{X},$$

tale che

$$q(\mathbf{x}) = \theta(\mathbf{x}, \mathbf{x}), \quad \theta(\mathbf{x}, \mathbf{y}) = \frac{1}{4} \left(q(\mathbf{x} + \mathbf{y}) - q(\mathbf{x} - \mathbf{y}) \right), \quad \|\mathbf{x}\| = \sqrt{q(\mathbf{x})}.$$

Notiamo che non tutte le norme soddisfano l'identità del parallelogramma. Per esempio, per la norma $\|\cdot\|_{\infty}$ su $\mathbb{R}^2$ abbiamo

$$\left\| \begin{pmatrix} 1 \\ 0 \end{pmatrix} + \begin{pmatrix} 0 \\ 1 \end{pmatrix} \right\|_{\infty}^2 + \left\| \begin{pmatrix} 1 \\ 0 \end{pmatrix} - \begin{pmatrix} 0 \\ 1 \end{pmatrix} \right\|_{\infty}^2 = 2,$$

mentre

$$2 \left\| \begin{pmatrix} 1 \\ 0 \end{pmatrix} \right\|_{\infty}^2 + 2 \left\| \begin{pmatrix} 0 \\ 1 \end{pmatrix} \right\|_{\infty}^2 = 4.$$

Una forma bilineare definita positiva si chiama *prodotto scalare*. La forma bilineare dell'Esempio 1) è un prodotto scalare che si chiama *il prodotto scalare naturale di $\mathbb{R}^n$* . Per il prodotto scalare naturale di $\mathbb{R}^n$ si usa di solito la notazione $\mathbf{x} \cdot \mathbf{y}$.

Più generale, se $\mathbf{X}$ è uno spazio vettoriale reale di dimensione finita n e

$$\mathcal{B} = \{\mathbf{b}_1, \dots, \mathbf{b}_n\}$$

è una base di $\mathbf{X}$, allora la formula

$$\left\langle \sum_{j=1}^n \alpha_j \mathbf{b}_j, \sum_{j=1}^n \beta_j \mathbf{b}_j \right\rangle^{(\mathcal{B})} := \sum_{j=1}^n \alpha_j \beta_j$$

definisce un prodotto scalare su $\mathbf{X}$ e la norma associata a questo è data dalla formula

$$\left\| \sum_{j=1}^n \alpha_j \mathbf{b}_j \right\|_2^{(\mathcal{B})} = \sqrt{\left\langle \sum_{j=1}^n \alpha_j \mathbf{b}_j, \sum_{j=1}^n \alpha_j \mathbf{b}_j \right\rangle^{(\mathcal{B})}} = \sqrt{\sum_{j=1}^n \alpha_j^2}.$$

Per il teorema sull'equivalenza delle norme, ogni norma su $\mathbf{X}$ è equivalente alla norma $\|\cdot\|_2^{(\mathcal{B})}$, quindi ad una norma che proviene da un prodotto scalare. Per di più, secondo il teorema di Auerbach, data una norma arbitraria $\|\cdot\|$ su $\mathbf{X}$, la base $\mathcal{B}$ può essere scelta in tal modo che

$$\frac{1}{\sqrt{n}} \|\mathbf{x}\|_2^{(\mathcal{B})} \leq \|\mathbf{x}\| \leq \sqrt{n} \|\mathbf{x}\|_2^{(\mathcal{B})}, \quad \mathbf{x} \in \mathbf{X}.$$

Infatti, le disuguaglianze precedenti risultano dalla scelta fatta nel teorema di Auerbach, tenendo conto che $\|\mathbf{x}\|_1^{(\mathcal{B})} \leq \sqrt{n} \|\mathbf{x}\|_2^{(\mathcal{B})}$ e $\|\mathbf{x}\|_2^{(\mathcal{B})} \leq \sqrt{n} \|\mathbf{x}\|_\infty^{(\mathcal{B})}$:

$$\begin{aligned} \sum_{j=1}^n |\alpha_j| &\stackrel{(17.3.2)}{\leq} \left(\sum_{j=1}^n 1^2 \right)^{1/2} \left(\sum_{j=1}^n \alpha_j^2 \right)^{1/2} = \sqrt{n} \left(\sum_{j=1}^n \alpha_j^2 \right)^{1/2}, \\ \left(\sum_{j=1}^n \alpha_j^2 \right)^{1/2} &\leq \max_{1 \leq j \leq n} |\alpha_j| \left(\sum_{j=1}^n 1 \right)^{1/2} = \sqrt{n} \max_{1 \leq j \leq n} |\alpha_j|. \end{aligned}$$

17.4. Norme e prodotti scalari su spazi vettoriali complessi

Sia $\mathbf{X}$ uno spazio vettoriale complesso. Come nel caso reale, una *norma* su $\mathbf{X}$ è una funzione

$$\mathbf{X} \ni \mathbf{x} \mapsto \|\mathbf{x}\| \in [0, +\infty)$$

tale che

- (i) $\|\mathbf{x}\| = 0$ se e solo se $\mathbf{x} = \mathbf{0}_{\mathbf{X}}$,
- (ii) $\|\lambda \mathbf{x}\| = |\lambda| \|\mathbf{x}\|$ per ogni scalare $\lambda \in \mathbb{C}$ e vettore $\mathbf{x}$,
- (iii) (*disuguaglianza triangolare*) $\|\mathbf{x} + \mathbf{y}\| \leq \|\mathbf{x}\| + \|\mathbf{y}\|$ per tutti i vettori $\mathbf{x}$ e $\mathbf{y}$.

$\mathbf{X}$ munito di una norma si chiama *spazio vettoriale normato complesso*.

Ricordiamo che qualsiasi spazio vettoriale complesso $\mathbf{X}$ diventa in modo naturale uno spazio vettoriale reale se restringiamo la moltiplicazione

$$\mathbb{C} \times \mathbf{X} \ni (\lambda, \mathbf{x}) \mapsto \lambda \mathbf{x}$$

dei vettori con i scalari solo ai scalari reali:

$$\mathbb{R} \times \mathbf{X} \ni (\lambda, \mathbf{x}) \mapsto \lambda \mathbf{x}.$$

Indichiamo con $\mathbf{X}_{\mathbb{R}}$ lo spazio vettoriale reale così ottenuto. Si vede facilmente che se $\mathcal{B} = \{\mathbf{b}_1, \dots, \mathbf{b}_n\}$ è un insieme di vettori linearmente

indipendenti in $\mathbf{X}$, allora $\mathcal{B} \cup i\mathcal{B} = \{\mathbf{b}_1, i\mathbf{b}_1, \dots, \mathbf{b}_n, i\mathbf{b}_n\}$ è un insieme di vettori linearmente indipendenti in $\mathbf{X}_{\mathbb{R}}$. Perciò se $\mathcal{B}$ è una base dello spazio vettoriale complesso $\mathbf{X}$, allora $\mathcal{B} \cup i\mathcal{B}$ è una base dello spazio vettoriale reale $\mathbf{X}_{\mathbb{R}}$. In particolare, la dimensione reale di $\mathbf{X}$ è il doppio della sua dimensione complessa.

Tutti i teoremi che sono stati dimostrati per norme su spazi vettoriali reali di dimensione finita, come

- il teorema sull’equivalenza delle norme,
- il teorema sulla caratterizzazione della compattezza ed
- il teorema di Auerbach,

valgono per le norme sui spazi vettoriali complessi. Le dimostrazioni per i spazi vettoriali complessi sono identiche alle dimostrazioni nell’ambito reale.

La teoria dei prodotti scalari si cambia invece sostanzialmente nel caso complesso. Invece di usare forme bilineari complesse, si lavora con le forme “sesquilineare”: se $\mathbf{X}$ è uno spazio vettoriale complesso, allora una *forma sesquilineare* su $\mathbf{X}$ è una funzione $\theta : \mathbf{X} \times \mathbf{X} \longrightarrow \mathbb{C}$ che è lineare rispetto alla prima variabile ed *antilineare* rispetto alla seconda. In altre parole, θ è sesquilineare se è biadditiva :

$$\begin{aligned}\theta(\mathbf{x}_1, \mathbf{y}) + \theta(\mathbf{x}_2, \mathbf{y}) &= \theta(\mathbf{x}_1 + \mathbf{x}_2, \mathbf{y}), & \mathbf{x}_1, \mathbf{x}_2, \mathbf{y} \in \mathbf{X}, \\ \theta(\mathbf{x}, \mathbf{y}_1) + \theta(\mathbf{x}, \mathbf{y}_2) &= \theta(\mathbf{x}, \mathbf{y}_1 + \mathbf{y}_2), & \mathbf{x}, \mathbf{y}_1, \mathbf{y}_2 \in \mathbf{X},\end{aligned}$$

omogenea rispetto alla prima variabile ed *antiomogenea* rispetto alla seconda variabile :

$$\theta(\lambda \mathbf{x}_1, \mathbf{y}) = \lambda \theta(\mathbf{x}_1, \mathbf{y}), \quad \theta(\mathbf{x}_1, \lambda \mathbf{y}) = \bar{\lambda} \theta(\mathbf{x}_1, \mathbf{y}), \quad \mathbf{x}, \mathbf{y} \in \mathbf{X}, \lambda \in \mathbb{C}.$$

Poi, una funzione $\theta : \mathbf{X} \times \mathbf{X} \longrightarrow \mathbb{C}$ si dice *hermitiana* se

$$\theta(\mathbf{y}, \mathbf{x}) = \overline{\theta(\mathbf{x}, \mathbf{y})}, \quad \mathbf{x}, \mathbf{y} \in \mathbf{X}.$$

Esempio: Su $\mathbb{C}^n$ possiamo definire una forma bilineare hermitiana

$$\theta \text{ tramite la formula } \theta(\mathbf{x}, \mathbf{y}) = \sum_{j=1}^n x_j \overline{y_j}, \text{ ove } \mathbf{x} = \begin{pmatrix} x_1 \\ \vdots \\ x_n \end{pmatrix} \text{ e } \mathbf{y} = \begin{pmatrix} y_1 \\ \vdots \\ y_n \end{pmatrix}.$$

Similmente come nel caso reale, ogni forma biadditiva $\theta : \mathbf{X} \times \mathbf{X} \longrightarrow \mathbb{C}$ soddisfa l'identità del parallelogramma

$$\theta(\mathbf{x} + \mathbf{y}, \mathbf{x} + \mathbf{y}) + \theta(\mathbf{x} - \mathbf{y}, \mathbf{x} - \mathbf{y}) = 2\theta(\mathbf{x}, \mathbf{x}) + 2\theta(\mathbf{y}, \mathbf{y}), \quad \mathbf{x}, \mathbf{y} \in X.$$

Perciò la funzione

$$q_\theta : \mathbf{X} \ni \mathbf{x} \longmapsto \theta(\mathbf{x}, \mathbf{x}) \in \mathbb{C},$$

verifica l'identità del parallelogramma sotto la forma

$$q_\theta(\mathbf{x} + \mathbf{y}) + q_\theta(\mathbf{x} - \mathbf{y}) = 2q_\theta(\mathbf{x}) + 2q_\theta(\mathbf{y}), \quad \mathbf{x}, \mathbf{y} \in X.$$

Se θ è sesquilineare, allora q_θ soddisfa anche la condizione di omogeneità quadratica assoluta ovvia

$$q_\theta(\lambda \mathbf{x}) = |\lambda|^2 q_\theta(\mathbf{x}), \quad \mathbf{x} \in \mathbf{X}, \lambda \in \mathbb{C}.$$

Chiamiamo *forma quadratica (complessa)* su $\mathbf{X}$ ogni funzione $q : \mathbf{X} \longrightarrow \mathbb{C}$ che soddisfa l'identità del parallelogramma e la condizione di omogeneità quadratica assoluta :

$$\begin{aligned} q(\mathbf{x} + \mathbf{y}) + q(\mathbf{x} - \mathbf{y}) &= 2q(\mathbf{x}) + 2q(\mathbf{y}), & \mathbf{x}, \mathbf{y} \in X, \\ q(\lambda \mathbf{x}) &= |\lambda|^2 q(\mathbf{x}), & \mathbf{x} \in \mathbf{X}, \lambda \in \mathbb{C}. \end{aligned}$$

Contrario al caso reale, per ogni forma sesquilineare θ (non solo per quelle hermitiane !) su $\mathbf{X}$ vale la *formula di polarizzazione (complessa)* :

$$\theta(\mathbf{x}, \mathbf{y}) = \frac{1}{4} \sum_{k=0}^3 i^k \theta(\mathbf{x} + i^k \mathbf{y}, \mathbf{x} + i^k \mathbf{y}), \quad \mathbf{x}, \mathbf{y} \in X.$$

Infatti, le identità

$$\begin{aligned} \theta(\mathbf{x} \pm \mathbf{y}, \mathbf{x} \pm \mathbf{y}) &= \theta(\mathbf{x}, \mathbf{x}) \pm \theta(\mathbf{x}, \mathbf{y}) \pm \theta(\mathbf{y}, \mathbf{x}) + \theta(\mathbf{y}, \mathbf{y}), \\ \theta(\mathbf{x} \pm i\mathbf{y}, \mathbf{x} \pm i\mathbf{y}) &= \theta(\mathbf{x}, \mathbf{x}) \mp i\theta(\mathbf{x}, \mathbf{y}) \pm i\theta(\mathbf{y}, \mathbf{x}) + \theta(\mathbf{y}, \mathbf{y}) \end{aligned}$$

implicano

$$\begin{aligned} &\theta(\mathbf{x} + \mathbf{y}, \mathbf{x} + \mathbf{y}) - \theta(\mathbf{x} - \mathbf{y}, \mathbf{x} - \mathbf{y}) + i\theta(\mathbf{x} + i\mathbf{y}, \mathbf{x} + i\mathbf{y}) - i\theta(\mathbf{x} - i\mathbf{y}, \mathbf{x} - i\mathbf{y}) \\ &= \theta(\mathbf{x}, \mathbf{x}) + \theta(\mathbf{x}, \mathbf{y}) + \theta(\mathbf{y}, \mathbf{x}) + \theta(\mathbf{y}, \mathbf{y}) \\ &\quad - \theta(\mathbf{x}, \mathbf{x}) + \theta(\mathbf{x}, \mathbf{y}) + \theta(\mathbf{y}, \mathbf{x}) - \theta(\mathbf{y}, \mathbf{y}) \\ &\quad + i\theta(\mathbf{x}, \mathbf{x}) + \theta(\mathbf{x}, \mathbf{y}) - \theta(\mathbf{y}, \mathbf{x}) + i\theta(\mathbf{y}, \mathbf{y}) \end{aligned}$$

$$\begin{aligned}
& -i\theta(\mathbf{x}, \mathbf{x}) + \theta(\mathbf{x}, \mathbf{y}) - \theta(\mathbf{y}, \mathbf{x}) - i\theta(\mathbf{y}, \mathbf{y}) \\
& = 4\theta(\mathbf{x}, \mathbf{y}).
\end{aligned}$$

Una prima implicazione della formula di polarizzazione è che se $\mathbf{X}$ è uno spazio vettoriale su $\mathbb{C}$ e θ è una forma sesquilineare su $\mathbf{X}$, allora

$$\theta \text{ è hermitiana} \iff \theta(\mathbf{x}, \mathbf{x}) \in \mathbb{R}, \mathbf{x} \in \mathbf{X}. \quad (17.4.1)$$

Infatti, se θ è hermitiana, allora abbiamo $\theta(\mathbf{x}, \mathbf{x}) = \overline{\theta(\mathbf{x}, \mathbf{x})}$, perciò $\theta(\mathbf{x}, \mathbf{x})$ è reale. Viceversa, se $\theta(\mathbf{x}, \mathbf{x}) \in \mathbb{R}$ per ogni $\mathbf{x} \in \mathbf{X}$, allora otteniamo, usando due volte la formula di polarizzazione :

$$\begin{aligned}
\overline{\theta(x, y)} &= \overline{\frac{1}{4} \sum_{k=0}^3 i^k \theta(x + i^k y, x + i^k y)} = \frac{1}{4} \sum_{k=0}^3 (-i)^k \theta(x + i^k y, x + i^k y) \\
&= \frac{1}{4} \sum_{k=0}^3 (-i)^k \theta\left(y + (-i)^k x, y + (-i)^k x\right) = \theta(y, x).
\end{aligned}$$

Se $\mathbf{X}$ è uno spazio vettoriale su $\mathbb{C}$, allora una forma sesquilineare θ su $\mathbf{X}$ si chiama

semidefinita positiva se $\theta(\mathbf{x}, \mathbf{x}) \geq 0$ per ogni $\mathbf{x} \in \mathbf{X}$,
definita positiva se $\theta(\mathbf{x}, \mathbf{x}) > 0$ per ogni $0 \neq \mathbf{x} \in \mathbf{X}$.

(17.4.1) implica che ogni forma sesquilineare complessa semidefinita positiva è automaticamente hermitiana. Per questa ragione, diversamente dal caso reale, non abbiamo richiesto esplicitamente l'hermitianità nella definizione di una forma sesquilineare complessa semidefinita positiva.

Diversamente del caso reale, secondo un risultato relativamente recente di Svetozar Kurepa, ogni forma quadratica complessa proviene da una unica forma sesquilineare, che poi risulta essere hermitiana se la forma quadratica prendeva solo valori reali, e semidefinita positiva se la forma quadratica prendeva valori in $[0, +\infty)$. La dimostrazione originale si trova in [6], mentre dimostrazioni semplificati sono state date in [12], [7].

Qui noi ci limitiamo a dimostrare la versione complessa del

Teorema di Jordan - von Neumann. *Se $\mathbf{X}$ è uno spazio vettoriale complesso e $q : \mathbf{X} \rightarrow [0, +\infty)$ è una forma quadratica, allora la formula di polarizzazione*

$$\theta(\mathbf{x}, \mathbf{y}) = \frac{1}{4} \sum_{k=0}^3 i^k \theta(\mathbf{x} + i^k \mathbf{y}, \mathbf{x} + i^k \mathbf{y}), \quad \mathbf{x}, \mathbf{y} \in \mathbf{X}.$$

definisce una forma sesquilineare semidefinita positiva $\theta : \mathbf{X} \times \mathbf{X} \rightarrow \mathbb{C}$ tale che

$$q(\mathbf{x}) = \theta(\mathbf{x}, \mathbf{x}), \quad \mathbf{x} \in \mathbf{X}.$$

Dimostrazione. q è una forma quadratica sullo spazio vettoriale reale $\mathbf{X}_{\mathbb{R}}$, quindi per la versione reale del teorema di Jordan - von Neumann esiste una forma bilineare semidefinita positiva $\theta_0 : \mathbf{X}_{\mathbb{R}} \times \mathbf{X}_{\mathbb{R}} \rightarrow \mathbb{R}$ tale che

$$q(\mathbf{x}) = \theta_0(\mathbf{x}, \mathbf{x}), \quad \mathbf{x} \in \mathbf{X}.$$

Per di più, θ_0 è data dalla formula di polarizzazione (reale)

$$\theta_0(\mathbf{x}, \mathbf{y}) = \frac{1}{4} \left(q(\mathbf{x} + \mathbf{y}) - q(\mathbf{x} - \mathbf{y}) \right), \quad \mathbf{x}, \mathbf{y} \in \mathbf{X}.$$

Per ogni $\mathbf{x}, \mathbf{y} \in \mathbf{X}$, applicando la formula di polarizzazione precedente, otteniamo che

$$\begin{aligned} & \theta_0(i\mathbf{x}, \mathbf{y}) + \theta_0(\mathbf{x}, i\mathbf{y}) \\ &= \frac{1}{4} \left(\theta_0(i\mathbf{x} + \mathbf{y}, i\mathbf{x} + \mathbf{y}) - \theta_0(i\mathbf{x} - \mathbf{y}, i\mathbf{x} - \mathbf{y}) \right. \\ & \quad \left. + \theta_0(\mathbf{x} + i\mathbf{y}, \mathbf{x} + i\mathbf{y}) - \theta_0(\mathbf{x} - i\mathbf{y}, \mathbf{x} - i\mathbf{y}) \right) \\ &= \frac{1}{4} \left(\underbrace{q(i\mathbf{x} + \mathbf{y}) - q(i\mathbf{x} - \mathbf{y})}_{=0} + \underbrace{q(\mathbf{x} + i\mathbf{y}) - q(\mathbf{x} - i\mathbf{y})}_{=0} \right) \\ &= 0, \end{aligned}$$

ove abbiamo usato

$$q(i\mathbf{x} \pm \mathbf{y}) = q(i(\mathbf{x} \mp i\mathbf{y})) = |i|^2 q(\mathbf{x} \mp i\mathbf{y}) = q(\mathbf{x} \mp i\mathbf{y}).$$

Cosicché

$$\theta_0(\mathbf{x}, i\mathbf{y}) = -\theta_0(i\mathbf{x}, \mathbf{y}), \quad \mathbf{x}, \mathbf{y} \in \mathbf{X}. \quad (17.4.2)$$

D'altra parte, per ogni $\mathbf{x} \in \mathbf{X}$,

$$\theta_0(\mathbf{x} + i\mathbf{x}, \mathbf{x} + i\mathbf{x}) = q((1+i)\mathbf{x}) = |1+i|^2 q(\mathbf{x}) = 2q(\mathbf{x})$$

e

$$\begin{aligned}
 \theta_0(\mathbf{x} + i\mathbf{x}, \mathbf{x} + i\mathbf{x}) &= \theta_0(\mathbf{x}, \mathbf{x}) + \theta_0(i\mathbf{x}, \mathbf{x}) + \theta_0(\mathbf{x}, i\mathbf{x}) + \theta_0(i\mathbf{x}, i\mathbf{x}) \\
 &= q(\mathbf{x}) + 2\theta_0(\mathbf{x}, i\mathbf{x}) + q(i\mathbf{x}) \\
 &= 2q(\mathbf{x}) + 2\theta_0(\mathbf{x}, i\mathbf{x})
 \end{aligned}$$

implicano che

$$\theta_0(\mathbf{x}, i\mathbf{x}) = 0, \quad \mathbf{x} \in \mathbf{X}. \quad (17.4.3)$$

Definiamo ora $\theta : \mathbf{X} \times \mathbf{X} \longrightarrow \mathbb{C}$ ponendo

$$\begin{aligned}
 \theta(\mathbf{x}, \mathbf{y}) &:= \theta_0(\mathbf{x}, \mathbf{y}) + i\theta_0(\mathbf{x}, i\mathbf{y}) \\
 &= \frac{1}{4} \left(q(\mathbf{x} + \mathbf{y}) - q(\mathbf{x} - \mathbf{y}) + i q(\mathbf{x} + i\mathbf{y}) - i q(\mathbf{x} - i\mathbf{y}) \right).
 \end{aligned}$$

Chiaramente, θ è biadditiva ed abbiamo $\theta(\lambda\mathbf{x}, \mathbf{y}) = \lambda\theta(\mathbf{x}, \mathbf{y}) = \theta(\mathbf{x}, \lambda\mathbf{y})$ per ogni $\mathbf{x}, \mathbf{y} \in \mathbf{X}$ e $\lambda \in \mathbb{R}$. Per la linearità complessa di θ rispetto alla prima variabile basta provare solo $\theta(i\mathbf{x}, \mathbf{y}) = i\theta(\mathbf{x}, \mathbf{y})$, mentre per l'antilinearità di θ rispetto alla seconda variabile, l'identità $\theta(\mathbf{x}, i\mathbf{y}) = -i\theta(\mathbf{x}, \mathbf{y})$:

$$\begin{aligned}
 \theta(i\mathbf{x}, \mathbf{y}) &= \theta_0(i\mathbf{x}, \mathbf{y}) + i\theta_0(i\mathbf{x}, i\mathbf{y}) \stackrel{(17.4.2)}{=} -\theta_0(\mathbf{x}, i\mathbf{y}) + i\theta(\mathbf{x}, \mathbf{y}) \\
 &= i(\theta_0(\mathbf{x}, \mathbf{y}) + i\theta_0(\mathbf{x}, i\mathbf{y})) = i\theta(\mathbf{x}, \mathbf{y}).
 \end{aligned}$$

$$\begin{aligned}
 \theta(\mathbf{x}, i\mathbf{y}) &= \theta_0(\mathbf{x}, i\mathbf{y}) + i\theta_0(\mathbf{x}, -\mathbf{y}) = \theta_0(\mathbf{x}, i\mathbf{y}) - i\theta_0(\mathbf{x}, \mathbf{y}) \\
 &= -i(\theta_0(\mathbf{x}, \mathbf{y}) + i\theta_0(\mathbf{x}, i\mathbf{y})) = -i\theta(\mathbf{x}, \mathbf{y}).
 \end{aligned}$$

Finalmente, $\theta(\mathbf{x}, \mathbf{x}) = \theta_0(\mathbf{x}, \mathbf{x}) + i\theta_0(\mathbf{x}, i\mathbf{x}) \stackrel{(17.4.3)}{=} \theta_0(\mathbf{x}, \mathbf{x}) = q(\mathbf{x})$, $\mathbf{x} \in \mathbf{X}$. In particolare, poiché i valori di q sono contenuti in $[0, +\infty)$, θ è semidefinita positiva. $\square$

Anche nell'ambito complesso vale

La disuguaglianza di Schwarz. *Se $\mathbf{X}$ è uno spazio vettoriale complesso e $\theta : \mathbf{X} \times \mathbf{X} \longrightarrow \mathbb{C}$ è una forma sesquilineare semidefinita positiva, allora*

$$|\theta(\mathbf{x}, \mathbf{y})|^2 \leq \theta(\mathbf{x}, \mathbf{x}) \theta(\mathbf{y}, \mathbf{y}), \quad \mathbf{x}, \mathbf{y} \in \mathbf{X}.$$

Dimostrazione. Poiché $\mathbf{X} \times \mathbf{X} \ni (\mathbf{x}, \mathbf{y}) \mapsto \Re \theta(\mathbf{x}, \mathbf{y})$ è una forma $\mathbb{R}$ -bilinare semidefinita positiva, secondo la disuguaglianza di Schwarz nell'ambito reale abbiamo

$$(\Re \theta(\mathbf{x}, \mathbf{y}))^2 \leq \theta(\mathbf{x}, \mathbf{x}) \theta(\mathbf{y}, \mathbf{y}), \quad \mathbf{x}, \mathbf{y} \in \mathbf{X}. \quad (17.4.4)$$

Siano adesso $\mathbf{x}, \mathbf{y} \in \mathbf{X}$ arbitrari. Se $\theta(\mathbf{x}, \mathbf{y}) = 0$, allora vale trivialmente $|\theta(\mathbf{x}, \mathbf{y})|^2 \leq \theta(\mathbf{x}, \mathbf{x}) \theta(\mathbf{y}, \mathbf{y})$. Supponiamo quindi che $\theta(\mathbf{x}, \mathbf{y}) \neq 0$. Ponendo

$$\lambda := \frac{|\theta(\mathbf{x}, \mathbf{y})|}{\theta(\mathbf{x}, \mathbf{y})},$$

abbiamo

$$\begin{aligned} \Re \theta(\lambda \mathbf{x}, \mathbf{y}) &= \Re (\lambda \theta(\mathbf{x}, \mathbf{y})) = |\theta(\mathbf{x}, \mathbf{y})|, \\ \theta(\lambda \mathbf{x}, \lambda \mathbf{x}) &= \lambda \bar{\lambda} \theta(\mathbf{x}, \mathbf{x}) = \theta(\mathbf{x}, \mathbf{x}). \end{aligned}$$

Perciò (17.4.4) implica

$$|\theta(\mathbf{x}, \mathbf{y})|^2 = (\Re \theta(\lambda \mathbf{x}, \mathbf{y}))^2 \leq \theta(\lambda \mathbf{x}, \lambda \mathbf{x}) \theta(\mathbf{y}, \mathbf{y}) = \theta(\mathbf{x}, \mathbf{x}) \theta(\mathbf{y}, \mathbf{y}).$$

□

La disuguaglianza di Schwarz implica, esattamente come nel caso reale,

La disuguaglianza di Minkowski. *Siano $\mathbf{X}$ uno spazio vettoriale complesso e $\theta : \mathbf{X} \times \mathbf{X} \rightarrow \mathbb{R}$ una forma sesquilineare semidefinita positiva. Allora*

$$\theta(\mathbf{x} + \mathbf{y}, \mathbf{x} + \mathbf{y})^{1/2} \leq \theta(\mathbf{x}, \mathbf{x})^{1/2} + \theta(\mathbf{y}, \mathbf{y})^{1/2}, \quad \mathbf{x}, \mathbf{y} \in \mathbf{X}.$$

□

Come nel caso reale, anche nel caso di uno spazio vettoriale complesso $\mathbf{X}$, per il teorema di Jordan - von Neumann e per la disuguaglianza di Minkowski esiste una corrispondenza biunivoca fra

- forme sesquilineari $\theta : \mathbf{X} \times \mathbf{X} \rightarrow \mathbb{C}$ positivamente definite,
- forme quadratiche $q : \mathbf{X} \rightarrow [0, +\infty)$ che si annullano solo in $\mathbf{0}_{\mathbf{X}}$,
- norme $\mathbf{X} \ni \mathbf{x} \mapsto \|\mathbf{x}\|$ che soddisfano l'identità del parallelogramma

$$\|\mathbf{x} + \mathbf{y}\|^2 + \|\mathbf{x} - \mathbf{y}\|^2 = 2\|\mathbf{x}\|^2 + 2\|\mathbf{y}\|^2, \quad \mathbf{x}, \mathbf{y} \in \mathbf{X},$$

tale che

$$q(\mathbf{x}) = \theta(\mathbf{x}, \mathbf{x}), \quad \theta(\mathbf{x}, \mathbf{y}) = \frac{1}{4} \sum_{k=0}^3 i^k \theta(\mathbf{x} + i^k \mathbf{y}, \mathbf{x} + i^k \mathbf{y}), \quad \|\mathbf{x}\| = \sqrt{q(\mathbf{x})}.$$

Una forma sesquilineare definita positiva si chiama *prodotto scalare (complesso)*. La forma sesquilineare dell'Esempio è un prodotto scalare che si chiama *il prodotto scalare naturale di $\mathbb{C}^n$* . Per il prodotto scalare naturale di $\mathbb{C}^n$ si usa di solito la notazione $\mathbf{x} \cdot \mathbf{y}$.

Più generale, se $\mathbf{X}$ è uno spazio vettoriale complesso di dimensione finita n e

$$\mathcal{B} = \{\mathbf{b}_1, \dots, \mathbf{b}_n\}$$

è una base di $\mathbf{X}$, allora la formula

$$\left\langle \sum_{j=1}^n \alpha_j \mathbf{b}_j, \sum_{j=1}^n \beta_j \mathbf{b}_j \right\rangle^{(\mathcal{B})} := \sum_{j=1}^n \alpha_j \overline{\beta_j}$$

definisce un prodotto scalare su $\mathbf{X}$ e la norma associata a questo è data dalla formula

$$\left\| \sum_{j=1}^n \alpha_j \mathbf{b}_j \right\|_2^{(\mathcal{B})} = \sqrt{\left\langle \sum_{j=1}^n \alpha_j \mathbf{b}_j, \sum_{j=1}^n \alpha_j \mathbf{b}_j \right\rangle^{(\mathcal{B})}} = \sqrt{\sum_{j=1}^n |\alpha_j|^2}.$$

Per il teorema sull'equivalenza delle norme, ogni norma su $\mathbf{X}$ è equivalente alla norma $\|\cdot\|_2^{(\mathcal{B})}$, quindi ad una norma che proviene da un prodotto scalare. Per di più, secondo il teorema di Auerbach, data una norma arbitraria $\|\cdot\|$ su $\mathbf{X}$, la base $\mathcal{B}$ può essere scelta in tal modo che

$$\frac{1}{\sqrt{n}} \|\mathbf{x}\|_2^{(\mathcal{B})} \leq \|\mathbf{x}\| \leq \sqrt{n} \|\mathbf{x}\|_2^{(\mathcal{B})}, \quad \mathbf{x} \in \mathbf{X}.$$

Infatti, le disuguaglianze precedenti risultano dalla scelta fatta nel teorema di Auerbach, tenendo conto che, similmente al caso reale, $\|\mathbf{x}\|_1^{(\mathcal{B})} \leq \sqrt{n} \|\mathbf{x}\|_2^{(\mathcal{B})}$ e $\|\mathbf{x}\|_2^{(\mathcal{B})} \leq \sqrt{n} \|\mathbf{x}\|_\infty^{(\mathcal{B})}$.

17.5. Basi ortogonali e proiezione ortogonale

Per evitare ripetizioni, trattiamo qui il caso di un prodotto scalare reale e quello di un prodotto scalare complesso simultaneamente. Perciò $\mathbb{K}$ indicherà un campo che può essere uguale sia ad $\mathbb{R}$ che a $\mathbb{C}$.

Sia $\mathbf{X}$ uno spazio vettoriale su $\mathbb{K}$, munito di un prodotto scalare

$$\mathbf{X} \times \mathbf{X} \ni (\mathbf{x}, \mathbf{y}) \longmapsto \langle \mathbf{x}, \mathbf{y} \rangle \in \mathbb{K},$$

e $\|\cdot\|$ la norma definita da questa tramite la formula $\|\mathbf{x}\| := \sqrt{\langle \mathbf{x}, \mathbf{x} \rangle}$. L'uguaglianza di due vettori $\mathbf{x}_1, \mathbf{x}_2 \in \mathbf{X}$ può essere caratterizzata tramite i loro prodotti scalari con i vettori in $\mathbf{X}$:

$$\mathbf{x}_1 = \mathbf{x}_2 \iff \langle \mathbf{x}_1, \mathbf{y} \rangle = \langle \mathbf{x}_2, \mathbf{y} \rangle, \quad \mathbf{y} \in \mathbf{X}.$$

Infatti, $\langle \mathbf{x}_1, \mathbf{y} \rangle = \langle \mathbf{x}_2, \mathbf{y} \rangle$ è equivalente a $\langle \mathbf{x}_1 - \mathbf{x}_2, \mathbf{y} \rangle = 0$, che per $\mathbf{y} = \mathbf{x}_1 - \mathbf{x}_2$ significa $\|\mathbf{x}_1 - \mathbf{x}_2\|^2 = 0$. Poi, la norma di ogni $\mathbf{x} \in \mathbf{X}$ si esprime mediante i suoi prodotti scalari con i vettori in $\mathbf{X}$, senza ricorrere a radici quadrate:

$$\|\mathbf{x}\| = \sup \{ |\langle \mathbf{x}, \mathbf{y} \rangle|; \mathbf{y} \in \mathbf{X}, \|\mathbf{y}\| \leq 1 \}.$$

Infatti, la disuguaglianza di Schwarz implica $\geq$ e nel caso $\mathbf{x} = \mathbf{0}_{\mathbf{X}}$ abbiamo addirittura uguaglianza. Se invece $\mathbf{x} \neq \mathbf{0}_{\mathbf{X}}$, allora

$$\|\mathbf{x}\| = \left\langle \mathbf{x}, \frac{1}{\|\mathbf{x}\|} \mathbf{x} \right\rangle \leq \sup \{ |\langle \mathbf{x}, \mathbf{y} \rangle|; \mathbf{y} \in \mathbf{X} \}.$$

Due vettori $\mathbf{x}, \mathbf{y} \in \mathbf{X}$ si dicono *ortogonali* o *perpendicolari* e si scrive $\mathbf{x} \perp \mathbf{y}$ se $\langle \mathbf{x}, \mathbf{y} \rangle = 0$. Se $S \subset \mathbf{X}$ e $\mathbf{x} \in \mathbf{X}$, allora $\mathbf{x} \perp S$ significa che $\mathbf{x}$ è ortogonale a tutti gli elementi di S . L'*ortogonale* di un sottoinsieme $S \subset \mathbf{X}$ è il sottospazio lineare

$$S^\perp := \{ \mathbf{x} \in \mathbf{X}; \mathbf{x} \perp S \} = \{ \mathbf{x} \in \mathbf{X}; \langle \mathbf{x}, \mathbf{y} \rangle = 0 \text{ per ogni } \mathbf{y} \in S \}$$

di $\mathbf{X}$. Chiaramente, se $\mathbf{x}$ è ortogonale a $\{\mathbf{x}_1, \dots, \mathbf{x}_n\}$, allora è ortogonale ad ogni combinazione lineare di questi vettori, quindi al sottospazio lineare $\text{lin}(\{\mathbf{x}_1, \dots, \mathbf{x}_n\})$ generato da loro:

$$\{\mathbf{x}_1, \dots, \mathbf{x}_n\}^\perp = \left(\text{lin}(\{\mathbf{x}_1, \dots, \mathbf{x}_n\}) \right)^\perp.$$

Uno dei vantaggi del lavoro con prodotti scalari ed ortogonalità consiste nel fatto che i coefficienti di una combinazione lineare di vettori

non zero e a due a due ortogonali sono determinati dalla combinazione lineare tramite una formula semplice :

PROPOSIZIONE 17.5.1. Coefficienti delle combinazioni lineari di vettori ortogonali.) *Sia $\mathbf{X}$ uno spazio vettoriale su $\mathbb{K}$, munito di un prodotto scalare $\langle \cdot, \cdot \rangle$. Se*

$$\mathbf{0}_{\mathbf{X}} \neq \mathbf{b}_1, \dots, \mathbf{b}_n \in \mathbf{X}$$

sono a due a due ortogonali, allora per ogni $\alpha_1, \dots, \alpha_n \in \mathbb{K}$ abbiamo

$$\alpha_k = \frac{1}{\langle \mathbf{b}_k, \mathbf{b}_k \rangle} \left\langle \sum_{j=1}^n \alpha_j \mathbf{b}_j, \mathbf{b}_k \right\rangle, \quad 1 \leq k \leq n. \quad (17.5.1)$$

In particolare, un insieme finito di vettori non nulli ed a due a due ortogonali in $\mathbf{X}$ è linearmente indipendente.

DIMOSTRAZIONE. Basta osservare che

$$\left\langle \sum_{j=1}^n \alpha_j \mathbf{b}_j, \mathbf{b}_k \right\rangle = \alpha_k \langle \mathbf{b}_k, \mathbf{b}_k \rangle + \sum_{j \neq k} \alpha_j \underbrace{\langle \mathbf{b}_j, \mathbf{b}_k \rangle}_{=0} = \alpha_k \underbrace{\langle \mathbf{b}_k, \mathbf{b}_k \rangle}_{>0}.$$

Se $\mathbf{X}$ è uno spazio vettoriale su $\mathbb{K}$ munito di un prodotto scalare $\langle \cdot, \cdot \rangle$, allora un sistema di generatori di $\mathbf{X}$, consistente da vettori $\mathbf{b}_1, \dots, \mathbf{b}_n \neq \mathbf{0}_{\mathbf{X}}$ a due a due ortogonali, si chiama una *base ortogonale* di $\mathbf{X}$. In questo caso lo sviluppo di $\mathbf{x} \in \mathbf{X}$ rispetto a $\mathbf{b}_1, \dots, \mathbf{b}_n$ è dato dalla formula (17.5.1) :

$$\mathbf{x} = \sum_{k=1}^n \frac{\langle \mathbf{x}, \mathbf{b}_k \rangle}{\langle \mathbf{b}_k, \mathbf{b}_k \rangle} \mathbf{b}_k.$$

Una *base ortonormale* di $\mathbf{X}$ è un sistema di generatori di $\mathbf{X}$, consistente da vettori $\mathbf{b}_1, \dots, \mathbf{b}_n$ di lunghezza una, cioè con $\|\mathbf{b}_j\|^2 = \langle \mathbf{b}_j, \mathbf{b}_j \rangle = 1$ per ogni $1 \leq j \leq n$, a due a due ortogonali. In questo caso lo sviluppo di un vettore $\mathbf{x} \in \mathbf{X}$ rispetto a $\mathbf{b}_1, \dots, \mathbf{b}_n$ è ancora più semplice :

$$\mathbf{x} = \sum_{k=1}^n \langle \mathbf{x}, \mathbf{b}_k \rangle \mathbf{b}_k.$$

Notiamo che ad ogni base ortogonale $\mathbf{b}_1, \dots, \mathbf{b}_n$ di $\mathbf{X}$ possiamo associare la base ortonormale

$$\frac{1}{\|\mathbf{b}_1\|} \mathbf{b}_1, \dots, \frac{1}{\|\mathbf{b}_n\|} \mathbf{b}_n,$$

perciò non c'è una differenza essenziale fra questi due tipi di base.

La base naturale di $\mathbb{K}^n$ è chiaramente una base ortonormale rispetto al prodotto scalare naturale.

ESEMPIO 17.5.2. Sia $\mathbf{b}_1, \dots, \mathbf{b}_n$ una base di $\mathbb{K}^n$. Allora, per trovare lo sviluppo $\mathbf{x} = \sum_{k=1}^n \alpha_k \mathbf{b}_k$ di un vettore $\mathbf{x} \in \mathbb{K}^n$, dobbiamo risolvere il sistema di equazioni lineari

$$\begin{pmatrix} | & & | \\ \mathbf{b}_1 & \dots & \mathbf{b}_n \\ | & & | \end{pmatrix} \cdot \begin{pmatrix} \alpha_1 \\ \vdots \\ \alpha_n \end{pmatrix} = \mathbf{x}.$$

Perciò se la base $\mathbf{b}_1, \dots, \mathbf{b}_n$ è ortogonale, allora la formula dello sviluppo rispetto ad essa esprime esattamente la soluzione del sistema lineare di cui sopra. $\square$

Il prossimo teorema ci fa capire meglio la nozione di “proiezione ortogonale” su un sottospazio lineare, introdotta nella Proposizione 6.6.5 (v) e che sarà approfondita successivamente nel Teorema 17.5.4:

TEOREMA 17.5.3. (Proprietà di minimo della proiezione ortogonale.) *Siano $\mathbf{X}$ uno spazio vettoriale su $\mathbb{K}$, dotato con un prodotto scalare $\langle \cdot, \cdot \rangle$, ed $\mathbf{Y}$ un sottospazio lineare di $\mathbf{X}$. Allora per $\mathbf{x}_0 \in \mathbf{X}$ e $\mathbf{y}_0 \in \mathbf{Y}$ sono equivalenti:*

- (i) $\|\mathbf{x}_0 - \mathbf{y}_0\| \leq \|\mathbf{x}_0 - \mathbf{y}\|, \quad \mathbf{y} \in \mathbf{Y};$
- (ii) $\langle \mathbf{x}_0 - \mathbf{y}_0, \mathbf{y} \rangle = 0, \quad \mathbf{y} \in \mathbf{Y}.$

Inoltre, dato $\mathbf{x}_0 \in \mathbf{X}$, può esistere al più un elemento $\mathbf{y}_0 \in \mathbf{Y}$ tale che queste due condizioni equivalenti siano soddisfatte.

DIMOSTRAZIONE. Per ogni $\mathbf{y} \in \mathbf{X}$ abbiamo

$$\|\mathbf{x}_0 - \mathbf{y}_0\|^2 - \|\mathbf{x}_0 - \mathbf{y}\|^2 = 2 \Re \langle \mathbf{x}_0 - \mathbf{y}_0, \mathbf{y} - \mathbf{y}_0 \rangle - \|\mathbf{y} - \mathbf{y}_0\|^2. \quad (17.5.2)$$

Infatti,

$$\begin{aligned}
\|\mathbf{x}_0 - \mathbf{y}\|^2 &= \langle (\mathbf{x}_0 - \mathbf{y}_0) - (\mathbf{y} - \mathbf{y}_0), (\mathbf{x}_0 - \mathbf{y}_0) - (\mathbf{y} - \mathbf{y}_0) \rangle \\
&= \langle \mathbf{x}_0 - \mathbf{y}_0, \mathbf{x}_0 - \mathbf{y}_0 \rangle - \langle \mathbf{x}_0 - \mathbf{y}_0, \mathbf{y} - \mathbf{y}_0 \rangle \\
&\quad - \langle \mathbf{y} - \mathbf{y}_0, \mathbf{x}_0 - \mathbf{y}_0 \rangle + \langle \mathbf{y} - \mathbf{y}_0, \mathbf{y} - \mathbf{y}_0 \rangle \\
&= \|\mathbf{x}_0 - \mathbf{y}_0\|^2 - 2 \Re \langle \mathbf{x}_0 - \mathbf{y}_0, \mathbf{y} - \mathbf{y}_0 \rangle + \|\mathbf{y} - \mathbf{y}_0\|^2.
\end{aligned}$$

Per (i) $\Rightarrow$ (ii) : (i) e (17.5.2) implicano che

$$2 \Re \langle \mathbf{x}_0 - \mathbf{y}_0, \mathbf{y} - \mathbf{y}_0 \rangle \leq \|\mathbf{y} - \mathbf{y}_0\|^2, \quad \mathbf{y} \in \mathbf{Y},$$

cioè

$$2 \Re \langle \mathbf{x}_0 - \mathbf{y}_0, \mathbf{y} \rangle \leq \|\mathbf{y}\|^2, \quad \mathbf{y} \in \mathbf{Y}.$$

Risulta che per ogni $\mathbf{y} \in \mathbf{Y}$ e $\lambda > 0$ abbiamo $2 \Re \langle \mathbf{x}_0 - \mathbf{y}_0, \lambda \mathbf{y} \rangle \leq \|\lambda \mathbf{y}\|^2$, cioè $2 \Re \langle \mathbf{x}_0 - \mathbf{y}_0, \mathbf{y} \rangle \leq \lambda \|\mathbf{y}\|^2$, e passando al limite per $\lambda \rightarrow 0$ otteniamo che $\Re \langle \mathbf{x}_0 - \mathbf{y}_0, \mathbf{y} \rangle \leq 0$. Se cambiamo in questa disuguaglianza $\mathbf{y}$ con $-\mathbf{y}$, risulta anche la disuguaglianza inversa, quindi $\Re \langle \mathbf{x}_0 - \mathbf{y}_0, \mathbf{y} \rangle = 0$.

Nel caso $\mathbb{K} = \mathbb{R}$ questo significa (ii). Nel caso $\mathbb{K} = \mathbb{C}$ invece risulta, per ogni $\mathbf{y} \in \mathbf{Y}$, che

$$|\langle \mathbf{x}_0 - \mathbf{y}_0, \mathbf{y} \rangle|^2 = \Re \langle \mathbf{x}_0 - \mathbf{y}_0, \overline{\langle \mathbf{x}_0 - \mathbf{y}_0, \mathbf{y} \rangle} \mathbf{y} \rangle = 0.$$

Per (ii) $\Rightarrow$ (i) : (i) risulta subito dalla condizione (ii) grazie a (17.5.2).

Per l'unicità della proiezione ortogonale basta notare che, se ogni uno di $\mathbf{y}_0, \mathbf{y}'_0 \in \mathbf{Y}$ soddisfa (ii), allora $\mathbf{y}'_0 - \mathbf{y}_0 = (\mathbf{x}_0 - \mathbf{y}_0) - (\mathbf{x}_0 - \mathbf{y}'_0)$ è ortogonale ad ogni vettore in $\mathbf{Y}$, quindi anche a se stesso. Perciò $\mathbf{y}'_0 - \mathbf{y}_0 = \mathbf{0}_X$.

Ora siano $\mathbf{X}$ uno spazio vettoriale su $\mathbb{K}$ e $\mathbf{Y}$ un suo sottospazio lineare. La *proiezione ortogonale* di un $\mathbf{x}_0 \in \mathbf{X}$ su $\mathbf{Y}$ è, se esiste, l'unico $\mathbf{y}_0 \in \mathbf{Y}$ che soddisfa le condizioni equivalenti del teorema precedente. Indichiamo $P_{\mathbf{Y}}(\mathbf{x}_0) := \mathbf{y}_0$. Riteniamo che $P_{\mathbf{Y}}(\mathbf{x}_0)$ è l'unico elemento di $\mathbf{Y}$ che minimizza la distanza di $\mathbf{x}_0$ a $\mathbf{Y}$, cioè

$$\|\mathbf{x}_0 - P_{\mathbf{Y}}(\mathbf{x}_0)\| = \inf \{ \|\mathbf{x}_0 - \mathbf{y}\| ; \mathbf{y} \in \mathbf{Y} \},$$

ovvero soddisfa la condizione d'ortogonalità equivalente

$$\mathbf{x}_0 - P_{\mathbf{Y}}(\mathbf{x}_0) \perp \mathbf{Y}.$$

Se assumiamo l'esistenza di una base ortogonale di $\mathbf{Y}$, allora l'esistenza della proiezione ortogonale su $\mathbf{Y}$ è garantita e, per di più, abbiamo una formula esplicita per essa :

TEOREMA 17.5.4. (Esistenza e forma della proiezione ortogonale.) *Siano $\mathbf{X}$ uno spazio vettoriale su $\mathbb{K}$, munito di un prodotto scalare $\langle \cdot, \cdot \rangle$, e $\mathbf{Y}$ un sottospazio lineare di $\mathbf{X}$, con una base ortogonale $\mathbf{b}_1, \dots, \mathbf{b}_m$. Allora per ogni $\mathbf{x} \in \mathbf{X}$ esiste la proiezione ortogonale $P_{\mathbf{Y}}(\mathbf{x})$ di $\mathbf{x}$ su $\mathbf{Y}$ e si ha*

$$P_{\mathbf{Y}}(\mathbf{x}) = \sum_{k=1}^m \frac{\langle \mathbf{x}, \mathbf{b}_k \rangle}{\langle \mathbf{b}_k, \mathbf{b}_k \rangle} \mathbf{b}_k.$$

DIMOSTRAZIONE. Stiamo cercando una combinazione lineare

$$P_{\mathbf{Y}}(\mathbf{x}) = \sum_{k=1}^m \alpha_k \mathbf{b}_k$$

tale che

$$\mathbf{x} - \sum_{k=1}^m \alpha_k \mathbf{b}_k \in \mathbf{Y}^\perp = \{\mathbf{b}_1, \dots, \mathbf{b}_m\}^\perp.$$

Le condizioni di ortogonalità richieste sono

$$\begin{aligned} 0 &= \left\langle \mathbf{x} - \sum_{k=1}^m \alpha_k \mathbf{b}_k, \mathbf{b}_j \right\rangle = \langle \mathbf{x}, \mathbf{b}_j \rangle - \sum_{k=1}^m \alpha_k \langle \mathbf{b}_k, \mathbf{b}_j \rangle \\ &= \langle \mathbf{x}, \mathbf{b}_j \rangle - \alpha_j \langle \mathbf{b}_j, \mathbf{b}_j \rangle, \end{aligned}$$

cioè

$$\alpha_j = \frac{\langle \mathbf{x}, \mathbf{b}_j \rangle}{\langle \mathbf{b}_j, \mathbf{b}_j \rangle},$$

ove $1 \leq j \leq m$. Perciò

$$\sum_{k=1}^m \alpha_k \mathbf{b}_k = \sum_{k=1}^m \frac{\langle \mathbf{x}, \mathbf{b}_k \rangle}{\langle \mathbf{b}_k, \mathbf{b}_k \rangle} \mathbf{b}_k$$

soddisfa le condizioni che caratterizzano la proiezione ortogonale di $\mathbf{x}$ su $\mathbf{Y}$.

Notiamo che nella situazione precedente abbiamo $(\mathbf{Y}^\perp)^\perp = \mathbf{Y}$. Infatti, l'inclusione $\mathbf{Y} \subset (\mathbf{Y}^\perp)^\perp$ è ovvia. Per l'inclusione inversa, sia

$\mathbf{x} \in (\mathbf{Y}^\perp)^\perp$ arbitrario. Allora $\mathbf{x} - P_{\mathbf{Y}}(\mathbf{x})$ si trova nello stesso tempo in $\mathbf{Y}^\perp$ ed in $(\mathbf{Y}^\perp)^\perp$, perciò è ortogonale a se stesso e risulta che deve annullarsi. Di conseguenza $\mathbf{x} = P_{\mathbf{Y}}(\mathbf{x}) \in \mathbf{Y}$.

Usando il rimarco precedente, risulta subito che $\mathbf{x} - P_{\mathbf{Y}}(\mathbf{x})$ è la proiezione ortogonale di $\mathbf{x}$ su $\mathbf{Y}^\perp$. Questa osservazione può essere utile nel calcolo di una proiezione ortogonale.

Un esempio numerico : Si calcoli la proiezione ortogonale del vettore

$$\mathbf{x} = \begin{pmatrix} 3 \\ -1 \\ 7 \end{pmatrix}$$

sul sottospazio lineare

$$\mathbf{Y} = \left\{ \begin{pmatrix} x_1 \\ x_2 \\ x_3 \end{pmatrix} \in \mathbb{R}^3; 2x_1 - x_2 + 3x_3 = 0 \right\}$$

di $\mathbb{R}^3$ (rispetto al prodotto scalare naturale).

Osservando che $\mathbf{Y} = \mathbf{Z}^\perp$, ove $\mathbf{Z}$ è il sottospazio lineare di $\mathbb{R}^3$ generato dal vettore

$$\begin{pmatrix} 2 \\ -1 \\ 3 \end{pmatrix},$$

che costituisce nello stesso tempo una base ortogonale per $\mathbf{Z}$, ci conviene calcolare prima $P_{\mathbf{Z}}(\mathbf{x})$ ed ottenere $P_{\mathbf{Y}}(\mathbf{x})$ quale $\mathbf{x} - P_{\mathbf{Z}}(\mathbf{x})$:

$$\begin{aligned} P_{\mathbf{Y}}(\mathbf{x}) &= \begin{pmatrix} 3 \\ -1 \\ 7 \end{pmatrix} - \frac{\begin{pmatrix} 3 \\ -1 \\ 7 \end{pmatrix} \cdot \begin{pmatrix} 2 \\ -1 \\ 3 \end{pmatrix}}{\begin{pmatrix} 2 \\ -1 \\ 3 \end{pmatrix} \cdot \begin{pmatrix} 2 \\ -1 \\ 3 \end{pmatrix}} \begin{pmatrix} 2 \\ -1 \\ 3 \end{pmatrix} \\ &= \begin{pmatrix} 3 \\ -1 \\ 7 \end{pmatrix} - \frac{28}{14} \begin{pmatrix} 2 \\ -1 \\ 3 \end{pmatrix} = \begin{pmatrix} -1 \\ 1 \\ 1 \end{pmatrix}. \quad \square \end{aligned}$$

A questo punto, il procedimento di ortogonalizzazione di Gram-Schmidt, introdotto nella Sezione 6.7) e che per completezza richiama qui di seguito, implica che ogni spazio vettoriale di dimensione finita, munito di un prodotto scalare, ammette una base ortogonale.

PROPOSIZIONE 17.5.5. (Procedimento di ortogonalizzazione di Gram-Schmidt.) *Sia $\mathbf{X}$ uno spazio vettoriale su $\mathbb{K}$, munito di un prodotto scalare $\langle \cdot, \cdot \rangle$, e sia $\mathbf{a}_1, \dots, \mathbf{a}_n$ una base di $\mathbf{X}$. Allora possiamo costruire una base ortogonale $\mathbf{b}_1, \dots, \mathbf{b}_n$ di $\mathbf{X}$ ponendo successivamente*

$$\begin{aligned}\mathbf{b}_1 &:= \mathbf{a}_1, \\ \mathbf{b}_k &:= \mathbf{a}_k - \sum_{j=1}^{k-1} \frac{\langle \mathbf{a}_k, \mathbf{b}_j \rangle}{\langle \mathbf{b}_j, \mathbf{b}_j \rangle} \mathbf{b}_j \quad \text{per } 2 \leq k \leq n.\end{aligned}$$

DIMOSTRAZIONE. Si verifica facilmente che, per ogni $1 \leq k \leq n$, i vettori $\mathbf{a}_1, \dots, \mathbf{a}_k$ generano lo stesso sottospazio lineare $\mathbf{X}_k$ di $\mathbf{X}$ di $\mathbf{b}_1, \dots, \mathbf{b}_k$. Con questa notazione abbiamo, grazie al Teorema 17.5.4 sulla proiezione ortogonale,

$$\mathbf{b}_k = \mathbf{a}_k - P_{\mathbf{X}_{k-1}}(\mathbf{a}_k), \quad 2 \leq k \leq n.$$

Quindi $\mathbf{b}_1, \dots, \mathbf{b}_n$ sono a due a due ortogonali.

D'altro canto, l'indipendenza lineare dei vettori $\mathbf{a}_1, \dots, \mathbf{a}_n$ implica che $\mathbf{X}_k \neq \mathbf{X}_{k-1}$ per ogni $2 \leq k \leq n$. Risulta che $\mathbf{b}_k \neq \mathbf{0}_{\mathbf{X}}$ per ogni $2 \leq k \leq n$. Ricordiamo che ovviamente vale anche $\mathbf{b}_1 = \mathbf{a}_1 \neq \mathbf{0}_{\mathbf{X}}$.

Infine, $\mathbf{b}_1, \dots, \mathbf{b}_n$ generano lo stesso sottospazio lineare di $\mathbf{X}$ che $\mathbf{a}_1, \dots, \mathbf{a}_n$, cioè l'intero $\mathbf{X}$. Cosicché $\mathbf{b}_1, \dots, \mathbf{b}_n$ è una base ortogonale di $\mathbf{X}$.

Bibliografia

- [1] M. Abate, *Algebra Lineare*, McGraw-Hill Italia, 2000. **1**
- [2] H. Anton, Ch. Rorres, *Elementary Linear Algebra with Applications*, J. Wiley & Sons, New York, 1987. **7.7**
- [3] C. Ciliberto, F. Tovenà, *Appunti di Metodi Numerici per la Grafica mod. 1*, <http://www.mat.uniroma2.it/~tovenà/media2.pdf>, Università di Roma “Tor Vergata”, 2005.
- [4] E. Franti, *Algoritmi per grafica tridimensionale prospettica*, Tesi di Laurea (relatore A. M. Picardello), Università di Roma “La Sapienza”, anno acc. 1982-83. **16.7.3**
- [5] A. M. Gleason, *The definition of a quadratic form*, Amer. Math. Monthly **73** (1966), 1049–1056. **17.3**
- [6] S. Kurepa, *Quadratic and sesquilinear functionals*, Glasnik Mat. Fiz. Astr. **20** (1965), 79–92. **17.4**
- [7] S. Kurepa, *On the definition of a quadratic form*, Publications de l’Institut Mathématique **42 (56)** (1987), 35–41. **17.4**
- [8] M. P. Merola, *Integrazione dell’equazione delle onde mediante trasformata rapida di Fourier e visualizzazione grafica della loro soluzione*, Tesi di Laurea (relatore A. M. Picardello), Università di Roma “La Sapienza”, anno acc. 1986-87. **16.7.3**
- [9] M. Picardello, *Analisi di Fourier e trattamento numerico dei segnali*, http://www.mat.uniroma2.it/~picard/SMC/didattica/materiali_did/Anal.Armon./LIBRO.pdf, Università di Roma “Tor Vergata”, 2009. **6.6.6, 10.2, 10.2**
- [10] M. Picardello, *Algoritmi e metodi analitici, numerici e statistici in Computer Graphics*, http://www.mat.uniroma2.it/~picard/SMC/didattica/materiali_did/Comp.Graph./Note_di_Computer_Graphics.pdf, Università di Roma “Tor Vergata”, 2009. **16.7.1**
- [11] M. Sabbatucci, *Integrazione numerica di equazioni alle derivate parziali ed algoritmi di rappresentazione prospettica delle loro soluzioni*, Tesi di Laurea (relatore A. M. Picardello), Università di Roma “La Sapienza”, anno acc. 1984-85. **16.7.3**
- [12] P. Vrbová, *Quadratic functionals and bilinear forms*, Časopis Pest. Mat. **98** (1973), 159–161.

