

Algebra Lineare e Geometria

dispense del corso

Prof. Ernesto Dedò
Dipartimento di Matematica
Politecnico di Milano
`ernesto.dedo@polimi.it`

III edizione
febbraio 2012



Quest'opera è stata rilasciata con licenza Creative Commons Attribuzione - Non opere derivate 3.0 Unported. Per leggere una copia della licenza visita il sito web <http://creativecommons.org/licenses/by-nd/3.0/> o spediisci una lettera a Creative Commons, 171 Second Street, Suite 300, San Francisco, California, 94105, USA.

In particolare ne è vietato l'uso commerciale.

Queste dispense sono state composte con L^AT_EX

Indice

Indice	i
Elenco delle figure	v
Elenco delle tabelle	vii
Prefazione	ix
I Algebra lineare	1
1 Richiami di nozioni essenziali	3
1.1 Gli insiemi	3
1.2 Relazioni	4
1.3 Il simbolo di sommatoria	5
1.4 Il principio di induzione	6
1.5 Un po' di calcolo combinatorio	7
2 I sistemi lineari: teoria elementare	11
2.1 Concetti introduttivi	11
2.2 Risoluzione di un sistema	14
3 Matrici	15
3.1 Nomenclatura e prime operazioni	15
3.2 Operazioni sulle matrici	18
3.3 Polinomi di matrici	24
3.4 Matrici a blocchi	25
3.5 Applicazioni ai sistemi lineari	27
3.6 L'algoritmo di Gauss	30
4 Spazi vettoriali	35
4.1 Definizioni e prime proprietà	35
4.2 Sottospazi e basi	38
5 Determinante. Inversa	45

5.1	Definizioni di determinante	45
5.2	Proprietà del determinante	48
5.3	Un determinante particolare	50
5.4	Rango di una matrice	51
5.5	Calcolo del rango	53
5.6	Matrice inversa	54
6	Teoria dei sistemi lineari	57
6.1	Teoremi sui sistemi lineari	57
7	Applicazioni	63
7.1	Generalità	63
7.2	Applicazioni lineari, matrici, sistemi	67
7.3	Prodotto scalare, norma	68
7.4	Prodotto scalare	72
7.5	Generalizzazioni	75
8	Similitudine. Autovalori. Autovettori	77
8.1	Matrici simili	77
8.2	Autovalori ed autovettori di una matrice	78
9	Diagonalizzazione, matrici ortogonali	85
9.1	Diagonalizzazione di una matrice quadrata	85
9.2	Matrici ortogonali	88
9.3	Forme quadratiche	92
9.4	Matrici hermitiane e matrici unitarie	95
10	Polinomi di matrici	97
10.1	Teorema di Cayley Hamilton	97
10.2	Polinomio minimo	101
II	Geometria piana	103
11	La retta nel piano	105
11.1	Preliminari	105
11.2	Altri tipi di equazione della retta	107
11.3	Distanze	109
11.4	Fasci di rette	109
11.5	Coordinate omogenee	110
11.6	I sistemi di riferimento	112
11.7	Coordinate polari	114
12	La circonferenza nel piano	117
12.1	Generalità	117

12.2	Tangenti	119
12.3	Fasci di circonferenze	121
12.4	Circonferenza ed elementi impropri	123
13	Le coniche	125
13.1	Coniche in forma generale	130
13.2	Riconoscimento di una conica	131
13.3	Tangenti ad una conica in forma canonica	133
13.4	Conica per cinque punti	134
13.5	Le coniche in coordinate omogenee	135
13.6	Fasci di coniche	136
13.7	Fasci e punti impropri	139
III	Polarità piana	141
14	Proiettività ed involuzioni	143
14.1	Proiettività	143
14.2	Involuzioni	146
15	Polarità piana	149
15.1	Polare di un punto rispetto ad una conica irriducibile	149
15.2	Principali proprietà della polarità piana	152
15.3	Elementi coniugati rispetto ad una conica irriducibile	154
15.4	Triangoli autopolari	156
16	Centro ed assi	159
16.1	Centro e diametri	159
16.2	Assi di una conica	163
IV	Geometria dello spazio	165
17	Rette e piani nello spazio	167
17.1	Equazioni parametriche della retta nello spazio	167
17.2	Equazione di un piano nello spazio	168
17.3	Parallelismo e perpendicolarità nello spazio	170
17.4	La retta intersezione di due piani	171
17.5	Fasci di piani	173
17.6	Altri problemi su rette e piani	174
17.7	Simmetrie	180
17.8	Coordinate omogenee nello spazio	185
18	Sui sistemi di riferimento	189
18.1	Rototraslazioni	189

18.2	Coordinate polari e coordinate cilindriche	191
19	Linee e Superfici nello spazio	195
19.1	Superfici	195
19.2	Linee	198
20	Sfera e circonferenza nello spazio	201
20.1	La sfera	201
20.2	Piani tangenti ad una sfera	202
20.3	Circonferenze nello spazio	202
20.4	Fasci di sfere	204
21	Superfici rigate e di rotazione	207
21.1	Superfici rigate	207
21.2	Superfici di rotazione	208
22	Cilindri, coni e proiezioni	211
22.1	Coni	211
22.2	Cilindri	214
22.3	Proiezioni	216
22.4	Riconoscimento di una conica nello spazio	217
23	Superfici quadriche	219
23.1	Prime proprietà delle quadriche	219
23.2	Quadriche in forma canonica e loro classificazione	221
23.3	Natura dei punti e riconoscimento di una quadrica	226
	Indice analitico	233
	Indice analitico	233

Elenco delle figure

4.1	Matrici triangolari	44
7.1	Un sistema di riferimento cartesiano ortogonale nel piano	69
7.2	Un sistema di riferimento cartesiano ortogonale nello spazio	69
7.3	La regola del parallelogrammo per la somma di vettori	70
11.1	Distanza di due punti	109
11.2	Traslazione	113
11.3	Rotazione	114
11.4	Coordinate polari	115
12.1	La circonferenza in equazioni parametriche	119
12.2	Tangente in P ad una circonferenza	120
12.3	Tangenti da un punto ad una circonferenza	121
12.4	Fascio di circonferenze tangenti ad una retta	123
13.1	Le coniche	125
13.2	L'ellisse	126
13.3	L'iperbole	128
13.4	La parabola	129
13.5	Esempio 13.2	133
13.6	Esempio 13.3	136
13.7	Esempio 13.4	137
15.1	Tangenti da un punto ad una conica	153
15.2	Involuzione iperbolica, retta secante	155
15.3	Involuzione ellittica, retta esterna	155
15.4	Triangolo autopolare	157
16.1	Centro di simmetria	160
16.2	Centro graficamente	161
16.3	Diametro coniugato	162
16.4	Diametri e tangenti	163
17.1	Distanza di due punti	176

17.2	Distanza di un punto da un piano	177
17.3	Simmetrica di una retta rispetto ad un piano	182
17.4	Simmetrico di un piano rispetto ad un piano	183
17.5	Simmetrico di un piano rispetto ad un piano parallelo	184
17.6	Simmetrica di una retta rispetto ad un'altra retta	185
18.1	Coordinate polari	192
18.2	Coordinate cilindriche	193
20.1	Circonferenza nello spazio	203
21.1	Superficie di rotazione	209
22.1	Cono	211
22.2	Cilindro	215
22.3	Proiezione centrale	216
22.4	Proiezione parallela	218
23.1	Ellissoide	225
23.2	Gli iperboloidi	225
23.3	Paraboloide ellittico	225
23.4	Paraboloide iperbolico	226

Elenco delle tabelle

1	Lettere greche	xii
2	Simboli matematici usati in queste dispense	xiii
1.1	Relazioni di equivalenza	4
1.2	Relazioni d'ordine	5
3.1	Particolari matrici quadrate	18
3.2	Proprietà della somma di matrici	19
3.3	Proprietà del prodotto per uno scalare	20
3.4	Proprietà del prodotto di matrici	23
4.1	Proprietà degli spazi vettoriali	36
7.1	Proprietà della norma di un vettore	71
7.2	Proprietà della distanza di due punti	71
7.3	Proprietà del prodotto scalare	73
7.4	Proprietà delle forme bilineari	75
13.1	Le forme canoniche dell'equazione di una conica	132
23.1	Forma canonica delle quadriche specializzate	221
23.2	Forma canonica delle quadriche non specializzate	222
23.3	Riconoscimento di una quadrica non degenerare	227

Prefazione

MATHEMATICS

*is one of essential emanations of the human spirit – a thing to be valued in and for itself – like art or poetry*¹.

Oswald Veblen-1924²

Queste dispense contengono elementi di algebra lineare, con particolare riguardo all'algebra delle matrici ed agli spazi vettoriali e di geometria analitica nel piano e nello spazio in particolare lo studio delle coniche e delle quadriche, vi sono inoltre dei cenni di geometria proiettiva; esse rispecchiano fedelmente il corso dello stesso nome che da molti anni tengo presso la sede di Cremona del Politecnico di Milano. Sono corredate da numerosi esempi: alcuni di applicazione della teoria svolta, altri di approfondimento della stessa.

Nella seconda edizione, oltre a correggere numerosi errori, sono stati aggiunti i capitoli sulla polarità piana e sulle proprietà di centro ed assi di una conica; è stato inoltre aggiunto l'indice analitico.

Nella terza edizione sono stati eliminati altri refusi e sono state effettuate delle modifiche ad alcuni paragrafi con l'intento di renderli più chiari. È stato anche aumentato il numero degli esempi.

Consigli per affrontare meglio i corsi universitari di Matematica

COSE DA *non* FARE:

- i) *non* studiare su appunti presi da altri: generalmente ciascuno prende appunti a modo suo, mettendo in luce le cose che *per lui* sono più importanti o più difficili: di solito queste non coincidono con quelle che *a noi* sembrano più importanti o per noi sono più difficili;

¹LA MATEMATICA è una delle essenziali emanazioni dello spirito umano, una cosa che va valutata in sè e per sè, come l'arte o la poesia.

²Oswald Veblen- 24 June 1880 Decorah, Iowa, USA–10 Aug 1960 Brooklyn, Maine, USA.

- ii) *non* studiare sui testi di scuola superiore: per bene che vada sono impostati in maniera diversa e/o incompleti;
- iii) *non* imparare a memoria le formule o le dimostrazioni dei teoremi: la memoria tradisce molto più facilmente e più frequentemente del ragionamento: imparare a ricostruirle;
- iv) *non* aver paura di fare domande, ovviamente nei momenti opportuni: siete qui apposta per imparare;
- v) durante le lezioni *non* prendere freneticamente appunti: quello che spiega il docente di solito c'è sui libri o sulle dispense, mentre prendendo male gli appunti c'è un alto rischio di perdere il filo della lezione;
- vi) durante le esercitazioni *non* ricopiare pedissequamente la risoluzione degli esercizi;
- vii) *non* scoraggiarsi se, soprattutto all'inizio, sembra di non capire o sembra che il docente "vada troppo in fretta," seguendo i consigli si acquista presto il ritmo;
- viii) *non* perdere tempo: il fatto che all'università non ci siano compiti in classe e interrogazioni è una grande tentazione per rimandare il momento in cui "mettersi sotto" a studiare.

COSE DA FARE:

- i) Precedere *sempre* la lezione: cercare di volta in volta sui libri o sulle dispense gli argomenti che il docente tratterà nella prossima lezione e cominciare a leggerli per avere un'idea di che cosa si parlerà. In questo modo la lezione sarà più efficace e chiarirà molti dei dubbi e delle perplessità rimaste.
- ii) Iniziare a studiare dal primo giorno. La Matematica, soprattutto all'inizio, necessita di una lunga "digestione": di un ripensamento critico che non si può fare all'ultimo momento.
- iii) Studiare *sempre* con carta e penna a portata di mano, per poter rifare conti, dimostrazioni, figure ecc.
- iv) Se proprio si vuole farlo *imparare a prendere appunti* senza perdere il filo del discorso: appuntare solo i concetti base su cui poi riflettere e ricostruire da soli l'argomento; a esercitazioni appuntare *solo* il testo dell'esercizio ed il risultato finale e rifare per conto proprio l'esercizio.

- v) Imparare a *rifare* le dimostrazioni dei teoremi soffermandosi a riflettere sul ruolo che giocano le varie ipotesi del teorema, magari creando controesempi di situazioni in cui una o più delle ipotesi non valgono.
- vi) Ricordarsi che per imparare la matematica occorre far funzionare il cervello, e questo costa sempre un certo sforzo.

BUON LAVORO!

La tabella 1 a pagina xii fornisce un elenco di tutte le lettere greche, maiuscole e minuscole, con il loro nome in italiano, mentre la tabella 2 a pagina xiii elenca i simboli maggiormente usati nel testo.

Tabella 1 *Lettere greche*

<i>minuscole</i>	<i>maiuscole</i>	<i>nome</i>
α	A	alfa
β	B	beta
γ	Γ	gamma
δ	Δ	delta
ϵ o ε	E	epsilon
ζ	Z	zeta
η	H	eta
θ o ϑ	Θ	theta
ι	I	iota
κ	K	kappa
λ	Λ	lambda
μ	M	mi
ν	N	ni
ξ	Ξ	csi
o	O	omicron
π	Π	pi
ρ o ϱ	R	ro
σ o ς	Σ	sigma
τ	T	tau
υ	Y	ipsilon
ϕ o φ	Φ	fi
χ	X	chi
ψ	Ψ	psi
ω	Ω	omega

Tabella 2 *Simboli matematici usati in queste dispense*

$\mathbb{N}$	insieme dei numeri naturali
$\mathbb{Z}$	insieme dei numeri interi
$\mathbb{Q}$	insieme dei numeri razionali
$\mathbb{R}$	insieme dei numeri reali
$\mathbb{C}$	insieme dei numeri complessi
$\forall$	per ogni
$\exists$	esiste
$\exists!$	esiste un unico
$\in$	appartiene ad un insieme
$\cup$	unione di insiemi
$\cap$	intersezione di insiemi
Σ	somma
Π	prodotto
$\perp$	perpendicolare
$\langle \cdot, \cdot \rangle$	prodotto scalare
∞	infinito
$\wp$	insieme delle parti
$<$	minore
$>$	maggiore
$\leq$	minore o uguale
$\geq$	maggiore o uguale
$\subset$	sottoinsieme proprio
$\subseteq$	sottoinsieme
$\oplus$	somma diretta di insiemi
$\emptyset$	insieme vuoto

Parte I

Algebra lineare

Capitolo 1

Richiami di nozioni essenziali

In questo capitolo richiamiamo alcuni concetti e proprietà fondamentali, di cui faremo largo uso nel seguito, e che dovrebbero comunque essere noti ed acquisiti dagli studi precedenti, sia dalle Scuole superiori sia dal corso di Analisi Matematica.

Tuttavia *questo capitolo va letto con attenzione e non saltato a piè pari*. Se sussiste anche il minimo dubbio sui concetti qui esposti, lo studente deve correre rapidamente ai ripari, riprendendo in mano i testi ad esso relativi.

1.1 Gli insiemi

Il concetto di *insieme* è un concetto primitivo¹. Diciamo che l'elemento a appartiene all'insieme A e scriviamo $a \in A$ se a è un elemento di A . Indichiamo con il simbolo $\emptyset$ l'insieme *vuoto* cioè privo di elementi.² Due insiemi sono *uguali* se sono formati dagli stessi elementi, indipendentemente dall'ordine in cui i singoli elementi compaiono in ciascun insieme: per esempio gli insiemi $A = \{a, b, c, d\}$ e $B = \{b, a, d, c\}$ sono uguali.

Si dice che B è *sottoinsieme* di A e si scrive $B \subseteq A$ se tutti gli elementi di B sono anche elementi di A , in simboli

$$B \subseteq A \iff \forall a \in B: a \in B \Rightarrow a \in A \quad (1.1)$$

Se $B \neq A$ e vale la (1.1), B si chiama *sottoinsieme proprio* di A e si scrive, più propriamente, $B \subset A$.

¹Cioè non è un oggetto definito in base ad altri oggetti noti. Di esso si possono dare però proprietà e relazioni con altri oggetti noti.

²Attenzione, questo simbolo è quello dell'insieme vuoto e *non va usato* per indicare il numero 0 zero!

È noto che l'insieme vuoto si considera sottoinsieme di ogni insieme, cioè che

$$\emptyset \subseteq A \quad \forall A$$

Se A e B sono sottoinsiemi di uno stesso insieme U , si definisce

unione di A e B l'insieme $C = A \cup B$ tale che sia

$$C \equiv \{c \mid c \in A \text{ oppure } c \in B\}$$

cioè C è formato dagli elementi che appartengono ad A oppure a B .

intersezione di A e B l'insieme $C = A \cap B$ tale che

$$C \equiv \{c \mid \text{sia } c \in A, \text{ sia } c \in B\}$$

cioè C è l'insieme degli elementi che appartengono sia ad A sia a B .

differenza tra A e B è l'insieme $C = A \setminus B$ se $C = \{c \mid c \in A, c \notin B\}$ cioè l'insieme degli elementi di A che *non* appartengono a B .

Due insiemi tali che $A \cap B = \emptyset$ si chiamano *disgiunti*.

1.2 Relazioni

Ricordiamo che in un insieme A è definita una relazione di *equivalenza* $\mathcal{R}$ se, qualsiasi siano gli elementi $a, b, c \in A$, valgono le proprietà elencate nella tabella 1.1.

Per esempio è di equivalenza la relazione di parallelismo tra rette nel piano, pur di considerare parallele anche due rette coincidenti, mentre non lo è quella di perpendicolarità. (*perché?*)

Tabella 1.1 Relazioni di equivalenza

$a \mathcal{R} a$	<i>riflessiva</i>
$a \mathcal{R} b \iff b \mathcal{R} a$	<i>simmetrica</i>
$a \mathcal{R} b \text{ e } b \mathcal{R} c \Rightarrow a \mathcal{R} c$	<i>transitiva</i>

Se in un insieme A è definita una relazione di equivalenza $\mathcal{R}$ chiamiamo *classe di equivalenza* dell'elemento a l'insieme di tutti gli elementi di A che sono equivalenti ad a . L'insieme di tutte le classi di equivalenza di A si chiama insieme *quoziente di A rispetto a $\mathcal{R}$* .

Esempio 1.1. Per esempio, nell'insieme delle frazioni, la relazione

$$\frac{a}{b} \mathcal{R} \frac{ka}{kb} \quad (k \neq 0)$$

è una relazione di equivalenza. I numeri razionali si possono definire come l'insieme quoziente dell'insieme delle frazioni rispetto a $\mathbb{R}$ nel senso che il numero 0.5 rappresenta tutta la classe di frazioni equivalenti ad $\frac{1}{2}$.

Esempio 1.2. La relazione di parallelismo definisce una direzione: passando al quoziente, nell'insieme delle rette, rispetto alla relazione di parallelismo, si ottiene la direzione di una retta.

In un insieme A è definita una *relazione di ordine* $\preceq$ se per ogni $a, b, c \in A$ valgono le proprietà espresse nella tabella 1.2.

Tabella 1.2 Relazioni d'ordine

1.	$a \preceq b$ oppure $b \preceq a$	proprietà di dicotomia
2.	$a \preceq a$	proprietà riflessiva
3.	$a \preceq b$ e $b \preceq a \Rightarrow a = b$	proprietà antisimmetrica
4.	$a \preceq b$ e $b \preceq c \Rightarrow a \preceq c$	proprietà transitiva

Un insieme per tutti gli elementi del quale valgono tutte le proprietà elencate nella tabella 1.2 si chiama *totalmente ordinato*; in esso, data una qualunque coppia a e b di elementi, in virtù della proprietà di dicotomia, si può sempre dire se $a \preceq b$ oppure $b \preceq a$ cioè, come si suol dire, tutti gli elementi sono *confrontabili*. Esistono, però, insiemi in cui è definita una relazione d'ordine cosiddetta *parziale*, in cui non vale la proprietà 1 di dicotomia, cioè in cui non tutte le coppie di elementi sono confrontabili, come mostra il seguente

Esempio 1.3. Nell'insieme $\mathbb{N}$ dei numeri naturali, diciamo che $a \preceq b$ se a divide b . Questa è una relazione d'ordine parziale, infatti valgono le proprietà 2...4. –verificarlo per esercizio– ma non è totale perché esistono coppie di numeri non confrontabili, ad esempio 3 e 5 non sono confrontabili, perché nè 3 divide 5, nè, viceversa, 5 divide 3.

1.3 Il simbolo di sommatoria

Spesso in Matematica una somma di più addendi viene indicata con il simbolo

$$\Sigma$$

in cui gli addendi sono indicati da una lettera dotata di un indice numerico arbitrario, ad esempio

$$\sum_1^5 a_i = \sum_{j=1}^5 a_j.$$

Le seguenti scritture sono equivalenti:

$$\sum_{i=3}^6 x_i, \quad \sum_{k=3}^6 x_k, \quad \sum_{n=4}^7 x_{n-1},$$

e rappresentano la somma $x_3 + x_4 + x_5 + x_6$.

L'indice di sommatoria può essere sottinteso, quando è chiaro dal contesto. Dalle proprietà formali delle operazioni di somma e prodotto seguono subito queste uguaglianze:

$$\sum_{i=1}^n m x_i = m \sum_{i=1}^n x_i \quad \text{e} \quad \sum_{i=1}^n (x_i + y_i) = \sum_{i=1}^n x_i + \sum_{i=1}^n y_i.$$

Si considerano spesso anche somme su più indici, per esempio

$$\sum_{\substack{i=1\dots 2 \\ k=1\dots 3}} a_{ik} = a_{11} + a_{12} + a_{13} + a_{21} + a_{22} + a_{23}$$

e si verifica subito che vale la relazione

$$\sum_{\substack{i=1\dots 2 \\ k=1\dots 3}} a_{ik} = \sum_{i=1}^2 \left(\sum_{k=1}^3 a_{ik} \right) = \sum_{k=1}^3 \left(\sum_{i=1}^2 a_{ik} \right)$$

1.4 Il principio di induzione

Un procedimento di dimostrazione molto usato in Matematica, ed a cui ricorreremo varie volte, è quello basato sul *principio di induzione*. Esso è soprattutto adatto per dimostrare proprietà che dipendono dai numeri naturali.

DEFINIZIONE 1.1. Sia P_n una proposizione che dipende da un numero naturale n . Se

- i) P_1 è verificata
- ii) P_n si suppone vera
- iii) P_{n+1} è vera ogni volta che è vera P_n

allora P_n è vera per qualunque numero naturale n .

Vediamo qualche semplicissimo esempio.

Esempio 1.4. Consideriamo i primi n numeri dispari e verifichiamo che per ogni intero n si ha

$$1 + 3 + 5 + \cdots + (2n - 1) = n^2. \quad (1.2)$$

La (1.2) è ovviamente vera per $n = 1$, quindi P_1 è vera. Supponiamo ora che essa sia vera per un certo $n = k$, cioè che la somma dei primi k numeri dispari sia k^2 e dimostriamo che in questo caso è vera anche per $k + 1$. Per $n = k + 1$ la (1.2) diventa

$$1 + 3 + 5 + \cdots + (2(k + 1) - 1) = \\ \underbrace{(1 + 3 + 5 + \cdots + 2k - 1)}_{k^2} + 2k + 1 \stackrel{1.2}{=} k^2 + 2k + 1 = (k + 1)^2.$$

Dunque abbiamo fatto vedere che se la somma di k numeri dispari è k^2 ne segue che quella di $k + 1$ è $(k + 1)^2$. Essendo vera per $k = 1$ lo è per $k = 2$ quindi per $k = 3$ e così via per ogni n .

Esempio 1.5. Vogliamo far vedere che

$$1 + 2 + 3 + \cdots + n = \frac{n(n + 1)}{2} \quad (1.3)$$

Per $n = 2$ la (1.3) è ovviamente vera; supponiamo che sia vera per $n = k$ e dimostriamo che in tal caso è vera anche per $n = k + 1$. Infatti si ha:

$$1 + 2 + 3 + \cdots + k + (k + 1) = \frac{k(k + 1)}{2} + k + 1 = \\ = \frac{k(k + 1) + 2k + 2}{2} = \frac{k^2 + 3k + 2}{2} = \frac{(k + 1)(k + 2)}{2}.$$

1.5 Un po' di calcolo combinatorio

Come è noto, si chiama *fattoriale di n* ($n > 1$) e si scrive $n!$ il prodotto dei primi n interi

$$n! = n \cdot (n - 1) \cdots 3 \cdot 2$$

e si pone, per definizione, $1! = 1 = 0!$

Permutazioni

DEFINIZIONE 1.2. Si chiama *permutazione* di n oggetti un qualsiasi allineamento di questi oggetti.

Per esempio per contare quante sono le permutazioni delle tre lettere a, b, c possiamo provare ad elencarle:

$$\{abc\} \{acb\} \{bac\} \{bca\} \{cab\} \{cba\} \text{ quindi } P_3 = 6$$

Generalizzando, quante sono le permutazioni P_n di n oggetti? Osserviamo che il primo oggetto si può scegliere in n modi diversi, mentre per il secondo, una volta scelto il primo, posso operare $n - 1$ scelte, per il terzo le scelte sono $n - 2$ e così via, quindi

$$P_n = n(n - 1)(n - 2) \cdots 2 = n! \quad (1.4)$$

Consideriamo un'insieme $\mathcal{S}$ di n elementi e scegliamo una permutazione f degli elementi di $\mathcal{S}$ che considereremo *permutazione fondamentale*. Sia ora $p = \{x_1, x_2, \dots, x_n\}$ un'altra permutazione degli elementi dello stesso insieme $\mathcal{S}$. Se accade che due elementi si succedano in p in ordine inverso rispetto a come si succedono in f diciamo che essi formano una *inversione* o uno *scambio* rispetto ad f . Una permutazione p si dice di *classe pari* rispetto alla permutazione scelta come fondamentale f , se presenta un numero pari di scambi, di *classe dispari* in caso opposto.

Esempio 1.6. Sia $\{1, 2, 3, 4, 5\}$ la permutazione fondamentale dell'insieme formato dai primi cinque numeri interi. Consideriamo poi la permutazione $P\{3, 4, 2, 1, 5, \}$: in essa l'elemento 1 forma inversione con 2, 3 e 4 e l'elemento 2 con 3 e 4; quindi il numero totale di inversioni è 5. Dunque P è di classe dispari.

OSSERVAZIONE 1.1. Poiché scambiando due elementi si passa da una permutazione di classe pari ad una di classe dispari, il numero delle permutazioni di classe pari è uguale a quello delle permutazioni di classe dispari, cioè $\frac{n!}{2}$. Si può anche notare che se una permutazione p è di classe pari rispetto ad una permutazione fondamentale f , essa è di classe pari rispetto a qualunque permutazione q a sua volta di classe pari rispetto ad f .

Permutazioni con ripetizione

Consideriamo ora il caso in cui gli n oggetti non siano tutti distinti, per esempio consideriamo i 7 oggetti $a_1, a_2, b_1, b_2, b_3, c, d$ in cui

gli oggetti a sono indistinguibili tra loro e così pure gli oggetti b . Le permutazioni sono $7! = 5040$ ma siccome i tre oggetti b sono uguali avremo $3!$ permutazioni non distinguibili, quindi il numero diminuirà a $\frac{7!}{3!} = \frac{5040}{6} = 840$; siccome poi sono uguali anche i due elementi a abbiamo $\frac{7!}{3! \cdot 2!} = 420$.

Generalizzando, consideriamo l'insieme $X \equiv \{x_1, x_2, \dots, x_m\}$ se sappiamo che l'oggetto x_1 è ripetuto n_1 volte, l'oggetto x_2 è ripetuto n_2 volte e così via (con $\sum_{i=1}^m n_i = n$) si ha

$$P^* = \frac{n!}{n_1! n_2! \cdots n_m!} \quad (1.5)$$

Nel caso particolare in cui fra gli n oggetti ce ne siano k uguali tra loro e gli altri $n - k$ pure uguali tra loro (ma non ai precedenti) la formula (1.5) diventa

$$P^* = \frac{n!}{k!(n-k)!}$$

Combinazioni e disposizioni

Se abbiamo n oggetti e vogliamo suddividerli in gruppi di k oggetti ognuno di questi raggruppamenti si chiama *combinazione di n oggetti a k* se due di questi raggruppamenti sono diversi quando differiscono per almeno un oggetto. Si chiama invece *disposizione di n oggetti a k* ognuno di questi raggruppamenti quando consideriamo diversi due di essi non solo se differiscono per almeno un oggetto, ma anche se contengono gli stessi oggetti in ordine differente. Se indichiamo rispettivamente con $C_{n,k}$ e $D_{n,k}$ il numero delle combinazioni e quello delle disposizioni di n oggetti a k a k , dalla definizione segue subito che si ha

$$D_{n,k} = k! C_{n,k}$$

Ripetendo il ragionamento fatto per determinare il numero delle permutazioni, si ricava immediatamente che

$$D_{n,k} = n(n-1)(n-2) \cdots (n-k+1)$$

e di conseguenza

$$C_{n,k} = \frac{n(n-1)(n-2) \cdots (n-k+1)}{k!}$$

moltiplicando numeratore e denominatore per $(n - k)!$ si ottiene

$$C_{n,k} = \frac{n!}{k!(n - k)!} = \binom{n}{k}$$

che, come è noto, è il k -esimo coefficiente nello sviluppo della n -esima potenza di un binomio e prende perciò il nome di *coefficiente binomiale*.

Capitolo 2

I sistemi lineari: teoria elementare

2.1 Concetti introduttivi

In questo capitolo supporremo sempre, salvo esplicito avviso contrario, di trovarci nel campo reale $\mathbb{R}$, cioè che i coefficienti e le variabili delle equazioni che considereremo saranno numeri reali.

Un'equazione di primo grado è anche detta *lineare*; se è in una sola incognita essa si può sempre porre, con semplici passaggi algebrici, nella forma

$$ax + b = 0 \quad (2.1)$$

Per quanto riguarda la risoluzione dell'equazione (2.1) osserviamo che

Se $a \neq 0$ essa ammette sempre l'unica soluzione $x = -\frac{b}{a}$;
se $a = 0$ e $b \neq 0$ non ammette soluzioni;
se $a = 0$ e $b = 0$ essa è soddisfatta da *ogni* valore di x , quindi è più propriamente un'*identità*.

Se, invece, un'equazione lineare ha due variabili le soluzioni sono in generale infinite; per esempio consideriamo l'equazione

$$x + y = 1 \quad (2.2)$$

essa ammette come soluzione la coppia $x = 1, y = 0$, la coppia $x = 0, y = 1$, la coppia $x = \frac{1}{2}, y = \frac{1}{2} \dots$, e quindi, in generale, tutte le infinite coppie tali che

$$x = t, \quad y = 1 - t \quad \forall t \in \mathbb{R}.$$

Questa scrittura significa che per ogni valore del parametro t esiste una soluzione¹.

Sappiamo anche che una tale equazione rappresenta, nel piano riferito ad una coppia di assi cartesiani ortogonali, una retta, i cui infiniti punti hanno coordinate che sono appunto le soluzioni dell'equazione.

Se ora consideriamo, accanto alla (2.2) anche l'equazione $x - y = 1$ e ci proponiamo di trovare, se esistono, dei valori di x e y che soddisfanno *entrambe* le equazioni, abbiamo quello che si chiama un *sistema lineare* che si indica abitualmente con la scrittura

$$\begin{cases} x + y = 1 \\ x - y = 0 \end{cases} \quad (2.3)$$

Quindi un sistema di equazioni è, in generale, un insieme di equazioni di cui si cercano le eventuali soluzioni comuni; un sistema si chiama *lineare* se tutte le equazioni da cui è composto sono lineari.

Nel caso del sistema (2.3) si vede facilmente che esso è lineare e che la coppia $\left\{x = \frac{1}{2}, y = \frac{1}{2}\right\}$ soddisfa entrambe le equazioni, quindi è una soluzione del sistema. Non è difficile rendersi conto che essa è unica: dal punto di vista geometrico basta pensare che in un sistema di riferimento cartesiano ortogonale le equazioni lineari rappresentano rette, quindi esse, se si incontrano, hanno esattamente un punto in comune le cui coordinate sono proprio la soluzione del sistema stesso.

Possono però accadere altri casi: consideriamo per esempio il sistema

$$\begin{cases} x + y = 1 \\ x + y = 2 \end{cases} \quad (2.4)$$

È chiaro che le due equazioni si contraddicono l'una con l'altra, quindi il sistema *non ammette soluzioni* o è *impossibile*. Nell'interpretazione geometrica data precedentemente le due rette le cui equazioni formano il sistema (2.4) sono parallele, di conseguenza non hanno alcun punto in comune.

Un altro caso è rappresentato, per esempio, dal sistema

$$\begin{cases} x + y = 1 \\ 2x + 2y = 2 \end{cases} \quad (2.5)$$

In questo sistema appare chiaro che le due equazioni sono equivalenti nel senso che tutte le coppie soluzione della prima equazione lo

¹Attenzione, ogni soluzione è formata da una coppia ordinata di numeri reali.

sono anche della seconda, quindi questo sistema ammette *infinite soluzioni*. In questo caso è facile vedere che, dal punto di vista geometrico, le due rette coincidono.

Queste considerazioni si possono generalizzare al caso di sistemi con un numero qualunque (finito) di incognite. Parleremo allora, nel caso generale, di un sistema di m equazioni in n incognite².

$$\begin{cases} a_{11}x_1 + a_{12}x_2 + \cdots + a_{1n}x_n = b_1 \\ a_{21}x_1 + a_{22}x_2 + \cdots + a_{2n}x_n = b_2 \\ \dots\dots\dots \\ a_{m1}x_1 + a_{m2}x_2 + \cdots + a_{mn}x_n = b_m \end{cases} \quad (2.6)$$

La notazione che abbiamo usato per i coefficienti della (2.6) è la cosiddetta notazione *a doppio indice*, in cui il primo indice dei coefficienti rappresenta l'equazione, mentre il secondo la posizione dello stesso all'interno dell'equazione. Le incognite sono $x_1 \dots x_n$ i termini noti sono $b_1 \dots b_m$.

Dagli esempi visti si può concludere che un sistema lineare di m equazioni in n incognite può appartenere ad una ed una sola delle seguenti categorie:

sistema possibile questo a sua volta presenta due possibilità:

una sola soluzione costituita da una n -pla di valori.

infinite soluzioni dipendenti ciascuna da uno o più parametri, soluzioni che si possono determinare dando opportuni valori ai parametri.³

sistema impossibile non esistono soluzioni comuni alle varie equazioni.

Ci proponiamo, oltre che di imparare a risolvere un sistema lineare, anche di imparare a *discuterlo*, cioè a scoprire *a priori* (quindi senza doverlo risolvere), se è risolubile o no, e quante sono le sue soluzioni.

²L'interpretazione geometrica, nel caso di più di due incognite non è così immediata.

³Un tale sistema è a volte chiamato impropriamente *sistema indeterminato*; locuzione che può trarre in inganno in quanto le soluzioni, pur essendo infinite, possono però essere determinate.

2.2 Risoluzione di un sistema

Nelle scuole superiori avete imparato a risolvere semplici sistemi lineari con vari metodi, che si basano tutti sull'idea di fondo che è quella di trovare un sistema che abbia le stesse soluzioni di quello dato ma sia scritto in una forma più semplice.

DEFINIZIONE 2.1. Due sistemi lineari di m equazioni in n incognite si dicono *equivalenti* se hanno le stesse soluzioni, cioè se ogni n -pla che è soluzione dell'uno lo è anche dell'altro.

Ad esempio i sistemi

$$\begin{cases} x + 3y = 7 \\ 2x - y = 0 \end{cases} \quad \begin{cases} (1 + 6)x = 7 \\ y = 2x \end{cases} \quad \begin{cases} x = 1 \\ y = 2 \end{cases}$$

sono equivalenti, come si verifica facilmente.

Un'ottima tecnica per risolvere un sistema lineare è quella di trovare un sistema equivalente a quello dato ma con una struttura più semplice. Vedremo nei prossimi capitoli come si possa passare da un sistema lineare ad uno equivalente basandoci sull'osservazione che un sistema è definito quando sono dati i vari coefficienti *nelle rispettive posizioni*: per far questo nel prossimo capitolo introdurremo il concetto di *matrice*.

Capitolo 3

Matrici

Esistono molti metodi per applicare la strategia esposta alla fine del capitolo 2, per la maggior parte dei quali è comodo introdurre uno strumento matematico molto potente ed utilizzato nei più svariati campi della Matematica e di tutte le Scienze: il calcolo matriciale.

3.1 Nomenclatura e prime operazioni

Osserviamo che un sistema è completamente determinato quando siano dati i termini noti ed i coefficienti *nelle loro rispettive posizioni*. Ad esempio, riferendoci al sistema (2.6 a pagina 13), per tener conto dei coefficienti e delle loro posizioni possiamo scrivere la tabella:

$$A = \begin{bmatrix} a_{11} & a_{12} & \dots & a_{1n} \\ a_{21} & a_{22} & \dots & a_{2n} \\ \vdots & \vdots & \vdots & \vdots \\ a_{m1} & a_{m2} & \dots & a_{mn} \end{bmatrix} \quad (3.1)$$

che chiamiamo *matrice dei coefficienti* e i termini noti possiamo incolonnarli

$$B = \begin{bmatrix} b_1 \\ b_2 \\ \vdots \\ b_m \end{bmatrix}$$

ottenendo la *matrice (o vettore) dei termini noti*.

È anche importante, come vedremo, la matrice

$$C = \begin{bmatrix} a_{11} & a_{12} & \dots & a_{1n} & b_1 \\ a_{21} & a_{22} & \dots & a_{2n} & b_2 \\ \vdots & \vdots & \vdots & \vdots & \vdots \\ a_{m1} & a_{m2} & \dots & a_{mn} & b_m \end{bmatrix} \quad (3.2)$$

detta anche *matrice completa*, costruita a partire dalla A accostandole a destra la colonna B dei termini noti: possiamo anche scrivere $C = [A|B]$.

Più in generale chiamiamo *matrice di tipo* (m, n) una tabella di numeri, reali o complessi, organizzata in m righe e n colonne. Indicheremo sempre, da ora in poi, le matrici con lettere latine maiuscole e gli elementi con lettere minuscole dotate eventualmente di due indici che ne individuano la posizione nella tabella, ad esempio scriveremo $A = [a_{ik}]$ dove $1 \leq i \leq m$ e $1 \leq k \leq n$. L'elemento a_{ik} sarà allora l'elemento che appartiene alla i -esima riga e alla k -esima colonna.

Ad esempio nella matrice $A = \begin{bmatrix} 1 & 2 & 3 \\ 3 & 4 & 5 \\ 5 & 6 & 7 \end{bmatrix}$ si ha $a_{32} = 6$, in quanto il numero 6 è nella posizione $(3, 2)$ cioè appartiene alla terza riga ed alla seconda colonna, allo stesso modo si ha: $a_{23} = a_{31} = 5$.

Osserviamo che in alcuni testi le matrici sono indicate con $\begin{pmatrix} a & b \\ c & d \end{pmatrix}$ cioè con le parentesi tonde, anzichè con $\begin{bmatrix} a & b \\ c & d \end{bmatrix}$.

Due matrici sono uguali quando sono dello stesso tipo e sono formate dagli stessi elementi *nelle stesse posizioni*; formalmente scriviamo che se $A = [a_{ik}]$ e $B = [b_{ik}]$ sono due matrici, $A = B$ se e solo se sono dello stesso tipo e se $a_{ik} = b_{ik} \forall i, k$.

Come controesempio consideriamo le matrici $A = \begin{bmatrix} 1 & 2 \\ 3 & 4 \end{bmatrix}$ e $B = \begin{bmatrix} 2 & 1 \\ 3 & 4 \end{bmatrix}$; esse, pur essendo formate dagli stessi elementi, *non* sono uguali, infatti $a_{11} \neq b_{11}$ e $a_{12} \neq b_{12}$.

Se $A = [a_{ik}]$ è una matrice di tipo (m, n) chiamiamo *trasposta di* A la matrice A_T , di tipo (n, m) ottenuta scambiando ordinatamente le righe con le colonne¹, dunque B è la trasposta di A se $b_{ik} = a_{ki} \forall i, k$.

¹ Anche la notazione A_T per la trasposta di una matrice A non è univoca: a volte la trasposta

Ad esempio se $A = \begin{bmatrix} 1 & 2 \\ 3 & 4 \\ 5 & 6 \end{bmatrix}$ si avrà $A_T = \begin{bmatrix} 1 & 3 & 5 \\ 2 & 4 & 6 \end{bmatrix}$; ovviamente la trasposta di una matrice di tipo (m, n) è una matrice di tipo (n, m) e la trasposta di una matrice di tipo $n \times n$ è ancora dello stesso tipo.

Sussiste anche la proprietà $(A_T)_T = A$, cioè la trasposta della trasposta di una matrice A è ancora la matrice A . La semplice dimostrazione di questa proprietà è lasciata come esercizio al lettore.

Una matrice di tipo (m, n) in cui $m = n$, cioè in cui il numero delle righe è uguale a quello delle colonne, si chiama *matrice quadrata di ordine* n . Ci occuperemo ora di matrici quadrate.

Se $A = A_T$ (il che implica che A sia quadrata, dimostrarlo per esercizio), cioè se $a_{ik} = a_{ki} \forall i, k$ diciamo che A è *simmetrica*; se invece $A = -A_T$, cioè se $a_{ik} = -a_{ki} \forall i, k$ diciamo che A è *emisimmetrica*². Come esercizio dimostrare che una matrice emisimmetrica ha tutti gli elementi principali uguali a zero.

Esempio 3.1. *La matrice*

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

è *simmetrica*, mentre la matrice

$$\begin{bmatrix} 0 & 1 & -2 \\ -1 & 0 & 3 \\ 2 & -3 & 0 \end{bmatrix}$$

è *emisimmetrica*

Gli elementi a_{ik} con $i = k$ si chiamano *elementi principali* o *elementi appartenenti alla diagonale principale* e la loro somma si chiama *traccia* della matrice, si indica con $tr A$ e si ha quindi

$$tr A = \sum_{i=1}^n a_{ii}.$$

viene indicata con A_t o con ${}_t A$ o anche con A^t oppure $A^\top$ o in altri modi. Noi useremo *sempre* la notazione A_T .

²Qui, come faremo d'ora in poi, abbiamo indicato con $-A$ (leggere, ovviamente, *meno* A) la matrice che si ottiene da A cambiando segno a tutti i suoi elementi.

Per esempio se A è la matrice $\begin{bmatrix} 1 & 0 \\ 2 & -3 \end{bmatrix}$ la sua traccia è

$$\text{tr} A = 1 - 3 = -2$$

Sia ora A una matrice quadrata di ordine n ; nella tabella 3.1 definiamo alcune particolari matrici quadrate.

Tabella 3.1 Particolari matrici quadrate

<i>diagonale</i>	se $a_{ik} = 0$ per ogni $i \neq k$; cioè se gli elementi <i>non</i> principali sono tutti nulli;
<i>scalare</i>	se è diagonale e gli elementi principali (cioè gli elementi a_{ii}) sono uguali tra loro;
<i>unità</i>	se è scalare e $\forall i, a_{ii} = 1$; la indicheremo con I , sottintendendo l'ordine quando non c'è ambiguità;
<i>triangolare inferiore</i>	se $\forall i > k$ si ha $a_{ik} = 0$ cioè se tutti gli elementi al di sopra della diagonale principale sono nulli;
<i>triangolare superiore</i>	se $\forall i < k$ si ha $a_{ik} = 0$ cioè se tutti gli elementi al di sotto della diagonale principale sono nulli.

Osserviamo esplicitamente che non si fa nessuna ipotesi sugli elementi principali di una matrice diagonale o triangolare: essi potrebbero a loro volta essere tutti o in parte nulli.

Esempio 3.2. La matrice $A = \begin{bmatrix} 1 & 0 & 0 \\ 0 & 0 & 0 \\ 0 & 0 & 9 \end{bmatrix}$ è una matrice diagonale, mentre

la matrice $B = \begin{bmatrix} -2 & 0 & 0 \\ 0 & -2 & 0 \\ 0 & 0 & -2 \end{bmatrix}$ è una matrice scalare e $C = \begin{bmatrix} 3 & 0 & 0 \\ 1 & 2 & 0 \\ 0 & 2 & 0 \end{bmatrix}$ è una matrice triangolare.

Osserviamo anche che una matrice diagonale può essere considerata sia triangolare inferiore sia triangolare superiore.

3.2 Operazioni sulle matrici

Ci proponiamo, in questo paragrafo, di introdurre un' *Algebra delle matrici*, cioè di imparare a fare dei conti con le matrici; iniziamo con la somma di due matrici.

Somma di matrici

Se $A = [a_{ik}]$ e $B = [b_{ik}]$ sono due matrici dello stesso tipo, diciamo che $C = [c_{ik}]$ (dello stesso tipo di A e B) è la *somma* di A e B e scriviamo $C = A + B$ se ogni elemento di C è la somma degli elementi corrispondenti di A e B , cioè se si ha,

$$c_{ik} = a_{ik} + b_{ik}; \forall i, k.$$

La somma di matrici gode delle proprietà elencate nella tabella 3.2, che sono le proprietà di un *gruppo abeliano*. Nella tabella, A e B sono

Tabella 3.2 Proprietà della somma di matrici

i)	$A + (B + C) = (A + B) + C$	proprietà associativa
ii)	$A + B = B + A$	proprietà commutativa
iii)	$A + \mathbf{0} = A$	esistenza elemento neutro
iv)	$A + (-A) = \mathbf{0}$	esistenza opposto
v)	$(A + B)_T = A_T + B_T$	

matrici dello stesso tipo ed abbiamo indicato con $\mathbf{0}$ la matrice che ha tutti gli elementi nulli (che chiameremo *matrice nulla*). La verifica delle proprietà della somma di matrici è quasi immediata e costituisce un utile esercizio per il lettore volenteroso.

OSSERVAZIONE 3.1. *Attenzione!* non confondere il numero zero 0 con la matrice zero (o matrice nulla), che noi indicheremo sempre col simbolo $\mathbf{0}$ (in grassetto); osserviamo però che in alcuni testi ma soprattutto nella scrittura a mano, essa viene spesso indicata con $\underline{0}$. $\square$

A proposito di simbologia ricordiamo che il simbolo $\emptyset$ indica l'insieme vuoto — cioè l'insieme privo di elementi — ed unicamente a tale scopo va riservato e non il numero 0 , come invece molti hanno l'abitudine di fare. Questa confusione tra simboli diversi rappresenta una pessima abitudine da perdere nel minor tempo possibile.

OSSERVAZIONE 3.2. La somma di matrici si estende facilmente al caso di un numero qualsiasi di matrici.

Prodotto per uno scalare

Si introduce anche un *prodotto esterno* o *prodotto per*³ *uno scalare*, prodotto tra un numero ed una matrice: se α è un numero (o più

³Attenzione, esiste anche il *prodotto scalare* –v. § 7.4 a pagina 72– da non confondere con quello che introduciamo qui.

precisamente uno scalare⁴) e $A = [a_{ik}]$ è una matrice, si ha $B = \alpha A$ se

$$b_{ik} = \alpha a_{ik}, \quad \forall i, k.$$

Ovviamente, per il prodotto per uno scalare, valgono le proprietà elencate nella tabella 3.3, anch'esse di quasi immediata dimostrazione, che lasciamo come esercizio al lettore.

Tabella 3.3 Proprietà del prodotto per uno scalare

i)	$\alpha \mathbf{0} = \mathbf{0} = 0A;$
ii)	$1A = A;$
iii)	$(\alpha\beta)A = \alpha(\beta A) = \beta(\alpha A);$
iv)	$(\alpha + \beta)A = \alpha A + \beta A;$
v)	$\alpha(A + B) = \alpha A + \alpha B;$
vi)	$\alpha A = \mathbf{0} \implies \alpha = 0 \text{ oppure } A = \mathbf{0}.$

Combinazioni lineari di matrici

Le due operazioni di somma e di prodotto per uno scalare si uniscono nella definizione di *combinazione lineare* di matrici.

DEFINIZIONE 3.1. Siano date n matrici dello stesso tipo $A_1, A_2, \dots, A_n$ ed n scalari. $\alpha_1, \alpha_2, \dots, \alpha_n$. Diciamo che la matrice B è *combinazione lineare* delle A_i con *coefficienti* α_i se

$$B = \sum_{i=1}^n \alpha_i A_i = \alpha_1 A_1 + \alpha_2 A_2 + \dots + \alpha_n A_n. \quad (3.3)$$

Ovviamente la combinazione lineare (3.3) dà la matrice nulla se tutti gli α_i sono nulli.

DEFINIZIONE 3.2. Se una combinazione lineare di n matrici dà la matrice nulla anche se non tutti i coefficienti sono nulli, cioè se

$$\mathbf{0} = \sum_{i=1}^n \alpha_i A_i \text{ con qualche } \alpha_i \neq 0$$

diciamo che le matrici sono *linearmente dipendenti*. In caso contrario, cioè se posso ottenere la matrice nulla solo se tutti i coefficienti sono nulli, diciamo che le matrici sono *linearmente indipendenti*.

⁴Il termine "scalare" nasce dal fatto che, generalizzando, le matrici si possono definire su un campo qualsiasi $\mathbb{K}$, non necessariamente un campo numerico come $\mathbb{R}$ o $\mathbb{C}$ e da qui il nome di scalare, che indica un elemento del campo $\mathbb{K}$.

Sul concetto di dipendenza lineare sussiste il fondamentale

Teorema 3.1. *Siano $A_1, A_2, \dots, A_k$ k matrici dello stesso tipo, esse sono linearmente dipendenti se e solo se almeno una di esse si può scrivere come combinazione lineare delle altre.*

Dimostrazione. Se le matrici $A_1, \dots, A_k$ sono linearmente dipendenti, allora esiste una loro combinazione lineare che dà la matrice nulla senza che tutti i coefficienti siano nulli, cioè $\alpha_1 A_1 + \alpha_2 A_2 + \dots + \alpha_k A_k = \mathbf{0}$. In virtù della proprietà commutativa, possiamo supporre, senza ledere la generalità, che $\alpha_1 \neq 0$; allora possiamo scrivere $A_1 = -\frac{1}{\alpha_1}(\alpha_2 A_2 + \dots + \alpha_k A_k)$, quindi A_1 è combinazione lineare delle rimanenti, e si ha:

$$A_1 = -\sum_{i=2}^k \frac{\alpha_i}{\alpha_1} A_i.$$

Viceversa supponiamo che $A_1 = \sum_{i=2}^k \alpha_i A_i$, allora la combinazione lineare

$$A_1 - \alpha_2 A_2 - \dots - \alpha_k A_k = \mathbf{0}$$

dà la matrice nulla, ed almeno il primo coefficiente, quello di A_1 , è diverso da zero, dunque le matrici sono linearmente dipendenti. $\square$

Prodotto tra matrici

Il prodotto di matrici è un'operazione meno intuitiva delle due precedenti. Siano A e B , prese nell'ordine dato, due matrici di tipi rispettivamente (m, p) e (p, n) ; diciamo *prodotto righe per colonne* o semplicemente *prodotto* delle matrici A e B , nell'ordine dato, la matrice $C = AB = [c_{ik}]$ di tipo (m, n) se l'elemento generico c_{ik} è la somma dei prodotti degli elementi della i -esima riga di A per gli elementi della k -esima colonna di B cioè

$$c_{ik} = \sum_{j=1}^p a_{ij} b_{jk}$$

OSSERVAZIONE 3.3. La definizione di prodotto di una matrice A per una matrice B implica che il numero delle colonne della prima matrice deve coincidere con il numero delle righe della seconda; diciamo, in tal caso, che A è *conformabile* con B . $\square$

Esempio 3.3. Siano $A(2,2) = \begin{bmatrix} a & b \\ c & d \end{bmatrix}$ e $B(2,3) = \begin{bmatrix} x & y & z \\ u & v & w \end{bmatrix}$ allora A è conformabile con B e si ha

$$C(2,3) = AB = \begin{bmatrix} ax + bu & ay + bv & az + bw \\ cx + du & cy + dv & cz + dw \end{bmatrix}$$

OSSERVAZIONE 3.4. Si vede subito che, in generale, se A è di tipo (m, p) e B è di tipo (p, n) il prodotto AB è una matrice di tipo (m, n) , mentre il prodotto BA non ha senso, perché B non è conformabile con A . Se A è di tipo (n, p) e B di tipo (p, n) il prodotto si può fare nei due sensi, ma dà luogo a due matrici quadrate certamente diverse: infatti una è di ordine n ed una di ordine p . Infine, anche nel caso in cui A e B siano entrambe quadrate dello stesso ordine n il prodotto AB ed il prodotto BA , pur essendo entrambe matrici quadrate ancora di ordine n , non sono, in generale, uguali, come mostra il seguente

Esempio 3.4. Se $A = \begin{bmatrix} 0 & 0 \\ 0 & 1 \end{bmatrix}$ e $B = \begin{bmatrix} 0 & 0 \\ 1 & 0 \end{bmatrix}$ si verifica subito che è

$$AB = \begin{bmatrix} 0 & 0 \\ 1 & 0 \end{bmatrix} \quad \text{ma} \quad BA = \begin{bmatrix} 0 & 0 \\ 0 & 0 \end{bmatrix} = \mathbf{0} \quad (3.4)$$

OSSERVAZIONE 3.5. La (3.4) mette in luce, oltre alla non commutatività del prodotto di matrici, anche il fatto che per le matrici non vale la legge di annullamento del prodotto cioè che il prodotto di due matrici può dare la matrice nulla senza che nessuna delle due sia la matrice nulla.

Dunque per il prodotto di matrici non vale, in generale, la proprietà commutativa, tuttavia esistono coppie di matrici A e B tali che $AB = BA$: esse si chiamano *permutabili* o si dice anche che *commutano*. Ad esempio la matrice I commuta con tutte le matrici del suo stesso ordine.

Per il prodotto tra matrici valgono invece le proprietà elencate nella tabella 3.4 nella pagina successiva, dove, ovviamente, in ciascun caso si sottintende che le uguaglianze valgono solo se sono rispettate le conformabilità.

Anche queste proprietà sono di facile dimostrazione, ricordando le proprietà dei numeri reali (o complessi) e la definizione di prodotto righe per colonne.

OSSERVAZIONE 3.6. Osserviamo che il prodotto righe per colonne è in un certo senso “privilegiato” rispetto agli altri analoghi (di ovvio significato) *colonne per righe*, *colonne per colonne* o *righe per righe*, perché

Tabella 3.4 Proprietà del prodotto di matrici

i)	$A(BC) = (AB)C$	associativa
ii)	$(A + B)C = AC + BC \quad C(A + B) = CA + CB$	distributive
iii)	$\alpha(AB) = (\alpha A)B = A\alpha B = (AB)\alpha$	associatività mista
iv)	$AI = IA = A$	esistenza elemento neutro
v)	$A\mathbf{0} = \mathbf{0}A = \mathbf{0}$	
vi)	$(AB)_T = B_T A_T$	

è l'unico dei quattro che gode della proprietà associativa, come si verifica facilmente su esempi che il lettore è invitato a trovare come utile esercizio, per esempio scrivendo facili programmini per calcolatore.

Il prodotto elemento per elemento non è soddisfacente, per esempio perché considerando i determinanti (vedi definizione 5.2 a pagina 46) non sarebbe più vero il teorema 5.5 a pagina 50 che dice che il determinante di un prodotto di matrici è uguale al prodotto dei determinanti.

Applicando la definizione di prodotto tra matrici si vede anche che, per esempio, il sistema 2.6 a pagina 13 si può scrivere nella forma, più comoda

$$Ax = b$$

in cui A è la matrice, di tipo (m, n) , $A = \begin{bmatrix} a_{11} & a_{12} & \dots & a_{1n} \\ a_{21} & a_{22} & \dots & a_{2n} \\ \vdots & \vdots & \ddots & \vdots \\ a_{m1} & a_{m2} & \dots & a_{mn} \end{bmatrix}$, x è una matrice di tipo $(n, 1)$ (matrice colonna), $x = \begin{bmatrix} x_1 \\ x_2 \\ \vdots \\ x_n \end{bmatrix}$ e b una matrice di

tipo $(m, 1)$, $b = \begin{bmatrix} b_1 \\ b_2 \\ \vdots \\ b_m \end{bmatrix}$.

Potenza di una matrice quadrata

DEFINIZIONE 3.3. Si chiama *potenza ennesima* A^n ($n \in \mathbb{Z}$) di una matrice quadrata A la matrice

$$A^n = \begin{cases} \underbrace{A \cdot A \cdots A}_{n \text{ volte}} & \text{se } n \geq 2 \\ A & \text{se } n = 1 \\ I & \text{se } n = 0 \end{cases}$$

Per la potenza di matrici valgono le usuali proprietà delle potenze, con l'attenzione alla non commutatività del prodotto; in particolare, per esempio, si avrà

$$A^m \cdot A^n = A^{m+n} \quad (A^m)^n = A^{m \cdot n};$$

ma in generale la non commutatività implica che

$$(AB)^n \neq A^n \cdot B^n$$

così come per esempio $(A \pm B)^2$ sarà in generale diverso da $A^2 \pm 2AB + B^2$, a meno che, naturalmente, A e B non commutino; lo stesso si può dire per il prodotto $(A + B)(A - B)$ che, nella forma $A^2 - B^2$ dipende essenzialmente dalla commutatività.

Segnaliamo anche che esistono matrici non banali tali che $A^2 = A$ e $A^k = 0$; esse si chiamano, rispettivamente, matrici *idempotenti* e matrici *nilpotenti*. Per esempio la matrice $\begin{bmatrix} 1 & 1 \\ 0 & 0 \end{bmatrix}$ è idempotente, mentre la matrice $\begin{bmatrix} 0 & 0 \\ 3 & 0 \end{bmatrix}$ è nilpotente: verificarlo per esercizio.

3.3 Polinomi di matrici

Dato un polinomio di grado n

$$a_0x^n + a_1x^2 + \cdots + a_{n-1}x + a_n$$

possiamo ora formalmente sostituire alla variabile x una matrice A ottenendo il *polinomio matriciale*

$$a_0A^n + a_1A^2 + \cdots + a_{n-1}A + a_nI.$$

Osserviamo che il “termine noto” di un polinomio corrisponde al coefficiente del termine x^0 , da qui la sostituzione $A^0 = I$. I polinomi di matrici sono molto usati nella teoria delle matrici ed anche nelle altre Scienze.

3.4 Matrici a blocchi

Una matrice A può essere divisa, mediante linee orizzontali e/o verticali, in sottomatrici che prendono il nome di *blocchi*. Naturalmente ogni matrice può essere decomposta in blocchi in parecchi modi diversi, ad esempio la matrice $A = \begin{bmatrix} a & b & c \\ d & e & f \end{bmatrix}$ si può suddividere in blocchi, tra gli altri, nei seguenti modi:

$$\left[\begin{array}{c|cc} a & b & c \\ \hline d & e & f \end{array} \right] \quad \left[\begin{array}{ccc} a & b & c \\ \hline d & e & f \end{array} \right] \quad \left[\begin{array}{cc|c} a & b & c \\ \hline d & e & f \end{array} \right]$$

Se A e B sono due matrici dello stesso tipo posso sempre suddividere entrambe in blocchi dello stesso tipo, più precisamente possiamo supporre

$$A = \begin{bmatrix} A_{11} & A_{12} & \dots & A_{1n} \\ A_{21} & A_{22} & \dots & A_{2n} \\ \vdots & \vdots & \ddots & \vdots \\ A_{m1} & A_{m2} & \dots & A_{mn} \end{bmatrix} \quad \text{e} \quad B = \begin{bmatrix} B_{11} & B_{12} & \dots & B_{1n} \\ B_{21} & B_{22} & \dots & B_{2n} \\ \vdots & \vdots & \ddots & \vdots \\ B_{m1} & B_{m2} & \dots & B_{mn} \end{bmatrix}$$

dove le matrici A_{ik} sono rispettivamente dello stesso tipo delle matrici $B_{ik} \forall i, k$. A questo punto è ovvio, dalla definizione di somma tra matrici, che

$$A + B = \begin{bmatrix} A_{11} + B_{11} & A_{12} + B_{12} & \dots & A_{1n} + B_{1n} \\ \dots & \dots & \dots & \dots \\ A_{m1} + B_{m1} & A_{m2} + B_{m2} & \dots & A_{mn} + B_{mn} \end{bmatrix}$$

Allo stesso modo si può eseguire il prodotto a blocchi per uno scalare, moltiplicando ciascun blocco per lo scalare.

Per il prodotto a blocchi, invece, occorre che le matrici siano decomposte in blocchi a due a due conformabili, come illustra il seguente

Esempio 3.5. Siano A e B due matrici (A conformabile con B) divise a blocchi come segue

$$A = \left[\begin{array}{cc|c} a & b & c \\ d & e & f \\ \hline g & h & k \end{array} \right] \quad B = \left[\begin{array}{c|c} x & u \\ y & v \\ \hline z & w \end{array} \right]$$

e chiamiamo

$$\begin{aligned} A_1 &= \begin{bmatrix} a & b \\ c & d \end{bmatrix} & A_2 &= \begin{bmatrix} c \\ f \end{bmatrix} & A_3 &= [g \quad h] & A_4 &= [k] \\ B_1 &= \begin{bmatrix} x \\ y \end{bmatrix} & B_2 &= \begin{bmatrix} u \\ v \end{bmatrix} & B_3 &= [z] & B_4 &= [w] \end{aligned}$$

allora possiamo scrivere

$$A = \begin{bmatrix} A_1 & A_2 \\ A_3 & A_4 \end{bmatrix} \quad e \quad B = \begin{bmatrix} B_1 & B_2 \\ B_3 & B_4 \end{bmatrix}$$

ed il prodotto si può eseguire a blocchi in questo modo:

$$AB = \begin{bmatrix} A_1B_1 + A_2B_3 & A_1B_2 + A_2B_4 \\ A_3B_1 + A_4B_3 & A_3B_2 + A_4B_4 \end{bmatrix}$$

come se i singoli blocchi fossero elementi delle matrici.

In generale se le matrici A e B sono decomposte in blocchi nel modo seguente

$$A = \begin{bmatrix} A_{11} & A_{12} & \dots & A_{1n} \\ A_{21} & A_{22} & \dots & A_{2p} \\ \vdots & \vdots & \ddots & \vdots \\ A_{q1} & A_{q2} & \dots & A_{qp} \end{bmatrix} \quad e \quad B = \begin{bmatrix} B_{11} & B_{12} & \dots & B_{1r} \\ B_{21} & B_{22} & \dots & B_{2r} \\ \vdots & \vdots & \ddots & \vdots \\ B_{p1} & B_{p2} & \dots & B_{pr} \end{bmatrix}$$

ed i blocchi A_{ij} sono conformabili con i blocchi $B_{jk} \forall i, j, k$ allora si può eseguire a blocchi il prodotto AB come se i singoli blocchi fossero elementi delle matrici.

OSSERVAZIONE 3.7. Il prodotto a blocchi è utile quando, per esempio, nelle matrici A e/o B si possono individuare blocchi che formano la matrice nulla oppure la matrice unità o anche nei seguenti casi:

- i) Se $A(m, p)$ e $B(p, m)$ sono due matrici tali che A sia conformabile con B e se indichiamo con $B_1, B_2, \dots, B_n$ le colonne di B , possiamo scrivere B a blocchi come $B = [B_1 \ B_2 \ \dots \ B_n]$ e quindi il prodotto AB si può scrivere

$$AB = [AB_1 \ AB_2 \ \dots \ AB_n]$$

- ii) Se ancora $B = [B_1 \ B_2 \ \dots \ B_n]$ e $D = \text{diag}(d_1, d_2, \dots, d_n)$ allora

$$BD = DB = [d_1B_1 \ d_2B_2 \ \dots \ d_nB_n]$$

- iii) Siano A e B quadrate e dello stesso ordine decomposte in blocchi nel modo seguente

$$A = \begin{bmatrix} A_1 & A_3 \\ \mathbf{0} & A_2 \end{bmatrix} \quad e \quad B = \begin{bmatrix} B_1 & B_3 \\ \mathbf{0} & B_2 \end{bmatrix}$$

dove A_1 e B_1 sono quadrate dello stesso ordine, allora si ha

$$AB = \begin{bmatrix} A_1B_1 & A_1B_3 + A_3B_1 \\ \mathbf{0} & A_2B_2 \end{bmatrix}$$

in particolare, se $A_3 = B_3 = \mathbf{0}$ si ha

$$\begin{bmatrix} A_1 & \mathbf{0} \\ \mathbf{0} & A_2 \end{bmatrix} \begin{bmatrix} B_1 & \mathbf{0} \\ \mathbf{0} & B_2 \end{bmatrix} = \begin{bmatrix} A_1B_1 & \mathbf{0} \\ \mathbf{0} & A_2B_2 \end{bmatrix}.$$

3.5 Applicazioni ai sistemi lineari

Equivalenza di matrici

Si chiamano *operazioni elementari* sulle righe (rispettivamente sulle colonne) di una matrice le seguenti operazioni

- i) scambio di due righe (colonne);
- ii) moltiplicazione di una riga (colonna) per una costante $k \neq 0$;
- iii) sostituzione di una riga (colonna) con la somma della riga (colonna) stessa con un'altra moltiplicata per una costante k .

DEFINIZIONE 3.4. Due matrici A e B si dicono *equivalenti per righe* (rispettivamente *per colonne*) se B si ottiene da A eseguendo un numero finito di operazioni elementari sulle righe (colonne).

Scriveremo che $A \sim B$. È facile dimostrare –farlo come esercizio– che se $A \sim B$ allora $B \sim A$ e che se $A \sim B$ e $B \sim C$ allora $A \sim C$, cioè che la relazione $\sim$ è una relazione di equivalenza.

Ed è importante il

Teorema 3.2. *Se le matrici complete di due sistemi lineari sono equivalenti per righe allora i sistemi sono equivalenti, cioè hanno le stesse soluzioni.*

Risoluzione di un sistema

Vediamo ora, su un esempio, come si possa utilizzare il Teorema 3.2 per la risoluzione di un sistema lineare.

Esempio 3.6. Consideriamo il sistema

$$\begin{cases} x + y + z - t = 1 \\ 2x + y + z + 3t = 2 \\ x - y - z = 0 \\ x + y + z - 3t = -1 \end{cases} \quad (3.5)$$

La matrice completa è la matrice

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

Se sommiamo alla II riga di A la prima moltiplicata per -2 , alla II la I moltiplicata per -1 ed alla IV la I moltiplicata per -1 otteniamo

$$\begin{bmatrix} 1 & 1 & 1 & -1 & 1 \\ 0 & -1 & -1 & 5 & 0 \\ 0 & -2 & -2 & 1 & -1 \\ 0 & 0 & 0 & -2 & -2 \end{bmatrix};$$

ora se in questa nuova matrice sommiamo alla III riga la seconda moltiplicata per -2 troviamo

$$\begin{bmatrix} 1 & 1 & 1 & -1 & 1 \\ 0 & -1 & -1 & 5 & 0 \\ 0 & 0 & 0 & -9 & -1 \\ 0 & 0 & 0 & -2 & -2 \end{bmatrix}.$$

Infine se sommiamo alla IV riga la II moltiplicata per $-\frac{2}{9}$ risulta

$$B = \begin{bmatrix} 1 & 1 & 1 & -1 & 1 \\ 0 & -1 & -1 & 5 & 0 \\ 0 & 0 & 0 & -9 & -1 \\ 0 & 0 & 0 & 0 & -\frac{16}{9} \end{bmatrix}.$$

La matrice B è ottenuta mediante operazioni elementari dalla A quindi è ad essa equivalente, allora, in virtù del Teorema 3.2 il sistema (3.5) ha le stesse soluzioni del sistema che ha la B come matrice completa; cioè equivale al

sistema

$$\begin{cases} x + y + z - t = 1 \\ -y - z + 5t = 0 \\ -9t = -1 \\ 0 = -\frac{16}{9} \end{cases}$$

che è palesemente impossibile.

Nella risoluzione data nell'esempio 3.6 abbiamo "ridotto" la matrice A ad una matrice equivalente "più semplice" nel senso precisato dalla seguente

DEFINIZIONE 3.5. Una matrice A di tipo (m, n) di elemento generico a_{ij} si dice *ridotta (per righe)* se in ogni riga che non contiene solo zeri esiste un elemento $a_{ij} \neq 0$ tale che per ogni k con $i < k \leq m$ si ha $a_{kj} = 0$.

Il procedimento illustrato è sempre possibile in quanto sussiste il

Teorema 3.3. Sia A una matrice qualsiasi, allora esiste una matrice B ridotta per righe equivalente alla A .

Possiamo ora introdurre ora il fondamentale concetto di *rango* di una matrice⁵.

DEFINIZIONE 3.6. Si dice *rango (o caratteristica)* di una matrice A il numero delle righe di una matrice ridotta A_1 equivalente alla A che non contengono solo zeri.

Scriveremo $r(A)$ per indicare il rango della matrice A ; dalla definizione 3.6 segue subito che $r(A)$ è un numero intero positivo (nullo se e solo se $A = 0$) cioè che $r(A) \in \mathbb{Z}^+$ ed è tale che, se A è di tipo (m, n) , vale la relazione $r(A) \leq \min(m, n)$.

Siccome la matrice ridotta per righe equivalente alla matrice A dipende dalle operazioni elementari effettuate, essa non è unica, dunque la definizione (3.6) è ben posta solo se il numero delle righe che non contengono solo zeri di una matrice ridotta A_i equivalente ad A non dipende dalle operazioni elementari effettuate sulla matrice A_i . Infatti si può dimostrare il

⁵Vedremo più avanti, a pagina 51, che il concetto di rango qui anticipato può essere definito in maniera formalmente molto diversa a partire dalla definizione di *determinante* di una matrice quadrata.

Teorema 3.4. *Siano A_1 e A_2 due matrici ridotte equivalenti alla stessa matrice A . Allora il numero delle righe che non contengono solo zeri è lo stesso in A_1 e A_2 , cioè $r(A) = r(A_1) = r(A_2)$.*

Con i concetti introdotti possiamo enunciare il fondamentale

Teorema 3.5 (di Rouché-Capelli). ⁶ *Sia A la matrice dei coefficienti di un sistema lineare e sia B la matrice completa, allora il sistema è possibile se e solo se*

$$r(A) = r(B).$$

Una dimostrazione del Teorema 3.5 verrà data più avanti, a pagina 58.

Una soluzione di un sistema lineare di m equazioni in n incognite è costituita, come abbiamo visto, da una n -pla ordinata di numeri. Abbiamo chiamato tale n -pla *vettore*; possiamo pensare ad un vettore come una matrice di tipo $(1, n)$ (vettore riga) o $(n, 1)$ (vettore colonna).

3.6 L'algoritmo di Gauss

Ricordiamo che un *sistema lineare* è costituito da un certo numero m di equazioni in un certo numero n di incognite.⁷ Una soluzione di un sistema lineare è quindi costituita da una n -pla ordinata di valori che sostituita alle incognite soddisfa *tutte le m equazioni*.

Come abbiamo visto, *un sistema lineare può ammettere una o infinite soluzioni oppure non essere risolubile*. Per risolvere un sistema lineare esistono vari metodi elementari, ben noti, ma che diventano scomodi quando il numero delle incognite è superiore a 2 o 3.

Un algoritmo spesso usato, che funziona sempre, è il cosiddetto *algoritmo di Gauss*⁸.

L'algoritmo si basa sul seguente

Teorema 3.6 (di Gauss). *Se un sistema lineare è ottenuto da un altro con una delle seguenti operazioni:*

- i) scambio di due equazioni*
- ii) moltiplicazione di ambo i membri di un'equazione per una costante non nulla*

⁶Eugene ROUCHE', 1832, Sommères, Gard (Francia) – 1910, Lunelle (Francia).
Alfredo CAPELLI, 1855, Milano – 1910, Napoli.

⁷Non si esclude, a priori, che possa essere $m = n$.

⁸Karl Friedrich Gauss 1777 (Brunswick) - 1855 (Göttingen) uno dei più grandi matematici di tutti i tempi.

iii) un'equazione è sostituita dalla somma di se stessa con un multiplo di un'altra.

allora i due sistemi hanno le stesse soluzioni.

e verrà illustrato mediante alcuni esempi.

Esempio 3.7. Sia da risolvere il sistema

$$\begin{cases} 3x_3 = 9 \\ x_1 + 5x_2 - 2x_3 = 2 \\ \frac{1}{3}x_1 + 2x_2 = 3 \end{cases}$$

Possiamo effettuare le seguenti operazioni:

i) scambiamo la prima riga con la terza:

$$\begin{cases} \frac{1}{3}x_1 + 2x_2 = 3 \\ x_1 + 5x_2 - 2x_3 = 2 \\ 3x_3 = 9 \end{cases};$$

ii) moltiplichiamo la prima riga per 3:

$$\begin{cases} x_1 + 6x_2 = 9 \\ x_1 + 5x_2 - 2x_3 = 2 \\ 3x_3 = 9 \end{cases};$$

iii) aggiungiamo la prima riga moltiplicata per -1 alla seconda:

$$\begin{cases} x_1 + 6x_2 = 9 \\ -x_2 - 2x_3 = -7 \\ 3x_3 = 9 \end{cases}.$$

Ormai il sistema è risolto: dalla terza equazione si ha $x_3 = 3$ da cui $x_2 = 1$ e $x_1 = 3$.

L'algoritmo di Gauss funziona anche quando il numero delle equazioni è diverso dal numero delle incognite cioè quando $m \neq n$

Esempio 3.8. Sia dato il sistema

$$\begin{cases} x + 3y = 1 \\ 2x + y = -3 \\ 2x + 2y = -2 \end{cases}$$

in cui $m > n$. Aggiungendo alla seconda ed alla terza riga la prima moltiplicata per -2 si ha il sistema equivalente

$$\begin{cases} x + 3y = 1 \\ -5y = -5 \\ -4y = -4 \end{cases}$$

A questo punto è già chiaro che l'unica soluzione è $x = -2$ e $y = 1$. In ogni caso, proseguendo l'algoritmo si può aggiungere alla terza riga la seconda moltiplicata per $-\frac{4}{5}$ ottenendo

$$\begin{cases} x + 3y = 1 \\ -5y = -5 \\ 0 = 0 \end{cases}$$

L'ultima uguaglianza è un'identità a riprova del fatto che le equazioni sono ridondanti.

Nel caso in cui il sistema fosse impossibile procedendo con l'algoritmo si arriva ad una contraddizione.

Esempio 3.9. Sia dato il sistema

$$\begin{cases} x + 3y = 1 \\ 2x + y = -3 \\ 2x + 2y = 0 \end{cases}$$

Sempre aggiungendo alla seconda ed alla terza riga la prima moltiplicata per -2 si ha il sistema equivalente

$$\begin{cases} x + 3y = 1 \\ -5y = -5 \\ -4y = -2 \end{cases}$$

a questo punto, però, aggiungendo alla terza riga la seconda moltiplicata per $-\frac{4}{5}$ si ottiene

$$\begin{cases} x + 3y = 1 \\ -5y = -5 \\ 0 = 2 \end{cases}$$

Quindi una palese contraddizione: possiamo concludere dunque che il sistema dato non ammette soluzioni, o, come più comunemente ma meno propriamente si dice, è impossibile.

Un sistema lineare può anche avere infinite soluzioni:

Esempio 3.10. Consideriamo il sistema $\begin{cases} x + y = 4 \\ 2x + 2y = 8 \end{cases}$ ed applicando l'algoritmo di Gauss (sommando alla seconda riga la prima moltiplicata per -2) si ottiene il sistema $\begin{cases} x + y = 4 \\ 0 = 0 \end{cases}$ equivalente a quello dato, da cui si vede che la seconda equazione è inutile, quindi la soluzione è data da tutte le infinite coppie di numeri che hanno come somma 4, che possiamo scrivere, ad esempio, come $x = t$ $y = 4 - t$.

Capitolo 4

Spazi vettoriali

Avete già sentito parlare di vettori, probabilmente avrete visto i vettori in Fisica, visualizzati come segmenti orientati. Completiamo e formalizziamo ora le nozioni viste.

4.1 Definizioni e prime proprietà

In generale diamo la seguente

DEFINIZIONE 4.1. Chiamiamo *spazio vettoriale su $\mathbb{R}$* e indichiamo con $\mathbb{R}^n$ l'insieme delle n -ple ordinate di numeri reali ed indichiamo con $\vec{v}$ ciascuna di queste n -ple, cioè $\vec{v} = [v_1, v_2, \dots, v_n]$, chiediamo inoltre che in $\mathbb{R}^n$ siano definite due operazioni: la somma componente per componente cioè

$$\vec{x} + \vec{y} = [x_1, x_2, \dots, x_n] + [y_1, y_2, \dots, y_n] = [x_1 + y_1, x_2 + y_2, \dots, x_n + y_n]$$

ed un prodotto esterno o prodotto per uno scalare

$$\alpha \vec{v} = \alpha [x_1, x_2, \dots, x_n] = [\alpha x_1, \alpha x_2, \dots, \alpha x_n]$$

dove α è un numero reale qualsiasi. Chiamiamo *vettori* le n -ple e *scalari* i numeri da cui essi sono formati.

Le proprietà di queste due operazioni sono elencate nella Tabella 4.1 nella pagina seguente, in cui α e β sono numeri reali e $\vec{v}$ e $\vec{w}$ sono vettori.

Si verifica immediatamente che lo spazio vettoriale $\mathbb{R}^n$ soddisfa alle proprietà indicate in Tabella 4.1 nella pagina successiva, in cui abbiamo indicato con $\mathbf{0}$ il vettore nullo, cioè quello che ha tutte le componenti uguali a 0 e con $-\vec{v}$ il vettore opposto di $\vec{v}$ cioè quello le cui componenti sono gli opposti delle componenti di $\vec{v}$.

Tabella 4.1 Proprietà degli spazi vettoriali

i)	$\vec{v} + (\vec{w} + \vec{u}) = (\vec{v} + \vec{w}) + \vec{u}$	associatività della somma
ii)	$\vec{v} + \vec{w} = \vec{w} + \vec{v}$	commutatività della somma
iii)	$\mathbf{0} + \vec{v} = \vec{v}$	esistenza del vettore nullo
iv)	$\vec{v} + (-\vec{v}) = \mathbf{0}$	esistenza del vettore opposto
v)	$\alpha(\vec{v} + \vec{w}) = \alpha\vec{v} + \alpha\vec{w}$	distributività rispetto alla somma di vettori
vi)	$(\alpha + \beta)\vec{v} = \alpha\vec{v} + \beta\vec{v}$	distributività rispetto alla somma di scalari
vii)	$\alpha(\beta\vec{v}) = (\alpha\beta)\vec{v}$	associatività mista
viii)	$1 \cdot \vec{v} = \vec{v}$	esistenza dell'unità
ix)	$0 \cdot \vec{v} = \mathbf{0}$	

Si può generalizzare il concetto di *spazio vettoriale* su un campo qualsiasi $\mathbb{K}$ dando la seguente

DEFINIZIONE 4.2. Uno spazio vettoriale V su $\mathbb{K}$ è costituito da un insieme V e da un campo $\mathbb{K}$; sugli elementi¹ di V sono definite due operazioni, una somma ed un prodotto esterno, che godono delle proprietà i) ... ix), elencate nella tabella 4.1 in cui $\vec{v}$ e $\vec{w}$ sono elementi di V e $\alpha, \beta \in \mathbb{K}$. Gli elementi di uno spazio vettoriale V si chiamano *vettori*, e *scalari* gli elementi del campo su cui esso è costruito.

Esempi di spazi vettoriali diversi da $\mathbb{R}^n$ sono (verificarlo per esercizio²):

- Le matrici di tipo (m, n) rispetto alle operazioni di somma e prodotto per uno scalare definite nel capitolo 3.
- I polinomi di grado n in una indeterminata sul campo reale, rispetto alle usuali somma e prodotto per uno scalare.

Riprendiamo ora un concetto fondamentale, già introdotto per le matrici nel § 3.2 a pagina 20

DEFINIZIONE 4.3. Come per le matrici, si dice che il vettore $\vec{w}$ è *combinazione lineare* dei k vettori $\vec{v}_1, \vec{v}_2, \dots, \vec{v}_k$ se esistono k scalari $\alpha_1, \alpha_2, \dots, \alpha_k$ tali che

$$\vec{w} = \sum_{i=1}^k \alpha_i \vec{v}_i = \mathbf{0}. \quad (4.1)$$

¹Sulla natura degli elementi di V non si fa nessuna ipotesi.

²Per verificare che un insieme V è uno spazio vettoriale bisogna verificare anzitutto che la somma di due elementi di V appartenga ancora a V e che il prodotto di un elemento di V per uno scalare sia ancora un elemento di V , poi che valgano le proprietà i)—ix) della tabella 4.1.

Fissiamo ora k vettori e consideriamo la loro combinazione lineare

$$\mathbf{0} = \sum_{i=1}^k \alpha_i \vec{v}_i. \quad (4.2)$$

È ovvio che la (4.2) è verificata quando *tutti* gli α_i sono nulli. Può però accadere che una combinazione lineare di vettori dia il vettore nullo senza che *tutti* i coefficienti siano nulli. In tal caso i vettori si chiamano *linearmente dipendenti* in caso contrario, cioè se la (4.2) vale *solo quando tutti* gli α_i sono nulli, si dice che i vettori sono *linearmente indipendenti*, quindi diamo la seguente

DEFINIZIONE 4.4. n vettori si dicono *linearmente dipendenti* se esiste una loro combinazione lineare

$$\sum_{i=1}^n \alpha_i \vec{v}_i$$

uguale al vettore nullo, senza che siano tutti nulli i coefficienti α_i

Ad esempio i vettori $\vec{v} = [1, 2, 1]$, $\vec{w} = [2, 0, -1]$ e $\vec{u} = [3, 2, 0]$ sono linearmente dipendenti, infatti si ha $\vec{v} + \vec{w} - \vec{u} = \mathbf{0}$. Invece i vettori $\vec{e}_1 = [1, 0, 0]$, $\vec{e}_2 = [0, 1, 0]$ e $\vec{e}_3 = [0, 0, 1]$, che chiameremo anche *vettori fondamentali*, sono linearmente indipendenti.

Sul concetto di dipendenza lineare sussiste il fondamentale

Teorema 4.1. Siano $\vec{v}_1, \vec{v}_2, \dots, \vec{v}_k$ k vettori, essi sono linearmente dipendenti se e solo se almeno uno di essi si può scrivere come combinazione lineare degli altri.

Esso è del tutto analogo al Teorema 3.1 a pagina 21 e si dimostra allo stesso modo.

Valgono, per la dipendenza ed indipendenza lineare, le seguenti proprietà:

- i) Se un insieme di vettori contiene il vettore nullo, esso è un insieme dipendente, infatti, ad esempio, $0 \cdot \vec{v} + 1 \cdot \mathbf{0}$ è una combinazione lineare non banale che genera il vettore nullo.
- ii) Aggiungendo un vettore qualsiasi ad un insieme linearmente dipendente, si ottiene ancora un insieme linearmente dipendente, infatti, se, per esempio, $\vec{v}_1, \vec{v}_2, \dots, \vec{v}_k$ sono linearmente dipendenti significa che $\sum_{i=1}^k \alpha_i \vec{v}_i = \mathbf{0}$ con qualche $\alpha_i \neq 0$ e quindi la combinazione lineare $\alpha_1 \vec{v}_1 + \alpha_2 \vec{v}_2 + \dots + \alpha_k \vec{v}_k + 0 \cdot \vec{v}_{k+1}$ dà il vettore nullo qualunque sia $\vec{v}_{k+1}$ sempre con gli stessi coefficienti non nulli.

- iii) Togliendo da un insieme indipendente un vettore qualsiasi si ottiene ancora un insieme indipendente, infatti se $\vec{v}_1, \dots, \vec{v}_n$ sono indipendenti, significa che $\sum_1^n \alpha_i \vec{v}_i = \mathbf{0}$ solo quando $\forall i, \alpha_i = 0$ quindi, a maggior ragione si ha $\sum_1^{n-1} \alpha_i \vec{v}_i = \mathbf{0}$ solo quando $\forall i, \alpha_i = 0$.

4.2 Sottospazi e basi

DEFINIZIONE 4.5. Sia V uno spazio vettoriale. Un sottoinsieme $W \subseteq V$ si chiama *sottospazio* di V se, rispetto alle stesse operazioni definite in V , è a sua volta uno spazio vettoriale.

In altre parole W è sottospazio di V se in esso continuano a valere le proprietà delle operazioni definite in V .

Ad esempio in $\mathbb{R}^3$ il sottoinsieme W formato dai vettori le cui componenti sono numeri dispari non è un sottospazio, perché la somma di due vettori cosiffatti è un vettore le cui componenti sono numeri pari e dunque non appartiene a W . Mentre invece lo è quello dei vettori che hanno, per esempio, la terza componente nulla.

In generale per verificare che un certo sottoinsieme W di uno spazio vettoriale V sia un sottospazio basta verificare che sia soddisfatta la seguente proprietà

- W è “chiuso” rispetto alle combinazioni lineari, cioè ogni combinazione lineare di vettori di W è ancora un vettore di W .

Infatti le proprietà degli spazi vettoriali (quelle elencate nella Tabella 4.1 a pagina 36) valgono in W in quanto valgono in V essendo $W \subseteq V$.

Si osserva subito che una immediata conseguenza di questa proprietà è che il vettore nullo $\mathbf{0}$ appartiene a qualunque sottospazio (infatti si ha $0 \cdot \vec{v} = \mathbf{0}$ qualunque sia $\vec{v}$).

Se ora consideriamo un certo numero k di vettori $\vec{v}_1, \vec{v}_2, \dots, \vec{v}_k$ di uno spazio vettoriale V e formiamo tutte le loro possibili combinazioni lineari otteniamo, come è facile verificare, uno spazio vettoriale W sottospazio di V , cioè tale che $W \subseteq V$. I vettori $\vec{v}_i$ si chiamano, in questo caso, *generatori* di W .

DEFINIZIONE 4.6. Un insieme di generatori linearmente indipendenti di uno spazio vettoriale V prende il nome di *base* di V .

Sulle basi è fondamentale il teorema

Teorema 4.2. *Se V è uno spazio vettoriale e $\mathcal{B}$ è una sua base, ogni vettore di V si può scrivere in maniera unica come combinazione lineare dei vettori di $\mathcal{B}$.*

Dimostrazione. Sia $\mathcal{B} \equiv \{\vec{e}_1, \dots, \vec{e}_n\}$ la base considerata, allora ogni vettore $\vec{v} \in V$ si scrive come combinazione lineare degli $\vec{e}_i$ in quanto questi ultimi sono generatori, inoltre, se, per assurdo, $\vec{v}$ si potesse scrivere in due modi diversi come combinazione lineare dei vettori di $\mathcal{B}$, si avrebbe sia $\vec{v} = \sum_1^n \alpha_i \vec{e}_i$ sia $\vec{v} = \sum_1^n \beta_i \vec{e}_i$ con qualche $\alpha_i \neq \beta_i$, in tal caso, sottraendo membro a membro, si avrebbe la combinazione lineare

$$\mathbf{0} = \sum (\alpha_i - \beta_i) \vec{e}_i \quad (4.3)$$

ma poichè, per ipotesi, gli $\vec{e}_i$ in quanto vettori di una base, sono linearmente indipendenti, si ha che tutti i coefficienti della combinazione (4.3) devono essere nulli, e quindi $\forall i, \alpha_i = \beta_i$. $\square$

Se in uno spazio vettoriale esiste una base $\mathcal{B}$ formata da n vettori, allora possiamo sostituire $p \leq n$ vettori di $\mathcal{B}$ con altrettanti vettori indipendenti ed ottenere ancora una base, come precisa il

Teorema 4.3. *Sia V uno spazio vettoriale e sia $\mathcal{B} = \{\vec{e}_1, \vec{e}_2, \dots, \vec{e}_n\}$ una sua base. Se $\vec{v}_1, \vec{v}_2, \dots, \vec{v}_p$ sono $p \leq n$ vettori indipendenti, allora si possono scegliere $n - p$ vettori di $\mathcal{B}$ in modo che l'insieme*

$$\{\vec{v}_1, \vec{v}_2, \dots, \vec{v}_p, \vec{e}_{p+1}, \dots, \vec{e}_n\}$$

sia ancora una base per V .

Dimostrazione. Dobbiamo far vedere che i vettori

$$\{\vec{v}_1, \vec{v}_2, \dots, \vec{v}_p, \vec{e}_{p+1}, \dots, \vec{e}_n\}$$

sono indipendenti e che generano tutto V . Poiché i $\vec{v}_i$ sono indipendenti, fra di essi non c'è il vettore nullo. Inoltre $\vec{v}_1$ è un vettore di V quindi

$$\vec{v}_1 = \sum_1^n \alpha_i \vec{e}_i \quad (4.4)$$

in quanto gli $\vec{e}_i$ formano una base. Mostriamo che $\{\vec{v}_1, \vec{e}_2, \dots, \vec{e}_n\}$ formano a loro volta una base: consideriamo una loro combinazione lineare che dia il vettore nullo, sia $\beta_1 \vec{v}_1 + \beta_2 \vec{e}_2 + \dots + \beta_n \vec{e}_n = \mathbf{0}$. Se $\beta_1 = 0$ allora si ha $\beta_2 \vec{e}_2 + \dots + \beta_n \vec{e}_n = \mathbf{0}$ e dall'indipendenza dei vettori $\vec{e}_i$ segue che $\forall i, \beta_i = 0$. Se fosse invece $\beta_1 \neq 0$ potrei

scrivere $\vec{v}_1 = -\frac{1}{\beta_1} \sum_{i=2}^n \beta_i \vec{e}_i$ ma ricordando che vale la (4.4) e per l'unicità della rappresentazione di un vettore mediante gli elementi di una base, concludiamo che i vettori sono indipendenti. Allo stesso modo si procede sostituendo di volta in volta un'altro vettore $\vec{v}_i$ ad uno della base. $\square$

Dalla definizione 4.6 segue inoltre che qualunque insieme di vettori indipendenti generatori V costituisce una base per V , quindi che uno stesso spazio vettoriale V ha infinite basi; di più sussiste anche il

Teorema 4.4. *Sia V uno spazio vettoriale, se una base di V è formata da n vettori, allora qualunque base è formata da n vettori.*

Dimostrazione. Infatti è facile vedere che qualunque $(n - k)$ -pla di vettori non può generare tutto V e viceversa che ogni insieme di $n + 1$ vettori è formato da vettori linearmente dipendenti in quanto ogni vettore di V si esprime, in virtù del Teorema 4.2, come combinazione lineare degli n vettori della base, quindi l' $n + 1$ -esimo non può essere indipendente dagli altri. $\square$

Da quanto detto fin qui si ricava che il numero n dei vettori di una base non dipende dalla scelta della base stessa, ma *caratterizza* lo spazio V e prende il nome di *dimensione* di V .

Lo spazio vettoriale costituito dal solo vettore nullo ha, per convenzione, dimensione 0. Si può anche osservare che se V è uno spazio vettoriale di dimensione n allora qualsiasi insieme di $k < n$ vettori indipendenti può essere completato ad una base di V aggiungendo altri $n - k$ vettori indipendenti.

OSSERVAZIONE 4.1. Da quanto detto si deduce anche che in uno spazio vettoriale V di dimensione n , qualunque n -pla di vettori indipendenti forma una base per V .

OSSERVAZIONE 4.2. Tutti gli spazi vettoriali fin qui considerati hanno dimensione n finita, ma è facile rendersi conto che esistono spazi vettoriali che non hanno dimensione finita, per esempio se $V = \mathcal{P}(x)$ è l'insieme dei polinomi in una indeterminata sul campo reale con le usuali operazioni di somma e di prodotto per uno scalare si vede subito che V è uno spazio vettoriale; tuttavia se si suppone che esistano p polinomi che generano V e se n è il grado massimo di questi polinomi, è ovvio che nessun polinomio di grado $m > n$ può essere generato da questi. Quindi $V = \mathcal{P}(x)$ rappresenta un'esempio di spazio vettoriale

Sorge il problema di come si possa, in uno stesso spazio vettoriale V , passare da una base ad un'altra, cioè quali siano le relazioni che legano i vettori di una base a quelli di un'altra. Siano $\mathcal{B}\{\vec{e}_1, \dots, \vec{e}_n\}$ e $\mathcal{B}'\{\vec{f}_1, \dots, \vec{f}_n\}$ due basi distinte di uno stesso spazio vettoriale V . I vettori $\vec{f}_i$, in quanto vettori di V si esprimono come combinazione lineare dei vettori $\vec{e}_i$ di $\mathcal{B}$, quindi si ha il sistema lineare:

[illegible]

$$A = \begin{bmatrix} a_{11} & a_{12} & \cdots & a_{1n} \\ a_{21} & a_{22} & \cdots & a_{2n} \\ \vdots & \vdots & \vdots & \vdots \\ a_{n1} & a_{n2} & \cdots & a_{nn} \end{bmatrix}$$

È facile verificare che l'intersezione insiemistica $V \cap W$ di sottospazi è a sua volta un sottospazio (dimostrarlo per esercizio); due spazi vettoriali V e W hanno sempre in comune almeno il vettore nullo, nel caso in cui questo sia l'unico vettore in comune si dice che sono *disgiunti* e si scrive $V \cap W = \{\mathbf{0}\}$: abbiamo già osservato che l'insieme formato dal solo vettore nullo è considerato uno spazio vettoriale.

Esempio 4.1. Consideriamo in $\mathbb{R}^3$ i due sottospazi $V = \{[x, y, z] \mid x = y\}$ formato dai vettori di $\mathbb{R}^3$ che hanno le prime due componenti uguali e $W = \{[x, y, z] \mid z = 0\}$, formato dai vettori di $\mathbb{R}^3$ che hanno la terza componente nulla, allora si ha

$$V \cup W = \{[x, y, z] \mid x = y \text{ oppure } z = 0\}$$

e se prendiamo $\vec{v} = [0, 0, 1] \in V$ e $\vec{w} = [1, 0, 0] \in W$ osserviamo che sia $\vec{v}$ sia $\vec{w}$ appartengono a $V \cup W$, ma il vettore $\vec{v} + \vec{w} = [1, 0, 1]$ non

ha uguali le prime due componenti nè ha la terza componente nulla, quindi $\vec{v} + \vec{w} \notin V \cup W$. Dunque l'insieme $V \cup W$ non è chiuso rispetto alla somma e di conseguenza non è un sottospazio di $\mathbb{R}^3$.

Sulla dimensione dei sottospazi di uno spazio vettoriale sussiste anche il

Teorema 4.5. *Sia V uno spazio vettoriale di dimensione n e sia W un suo sottospazio allora*

- i) W ha dimensione finita
- ii) $\dim(W) \leq \dim(V)$
- iii) se $\dim(W) = \dim(V)$ allora $W = V$

Dimostrazione. Se $W = \{0\}$ la i) e la ii) sono ovvie. Sia quindi $W \neq \{0\}$ e sia $\vec{w}_1$ un vettore di W , se $\vec{w}_1$ genera W allora $\dim W = 1$ e poiché in V vi è almeno un vettore indipendente segue che $\dim V \geq 1$ e quindi la i) e la ii) sono dimostrate. Se invece $\vec{w}_1$ non genera W allora esiste in W almeno un altro vettore $\vec{w}_2$ indipendente da $\vec{w}_1$. Se $\vec{w}_1$ e $\vec{w}_2$ generano W allora $\dim(W) = 2$ e, come nel caso precedente $\dim(V) \geq 2$ e così via. Il procedimento ha termine per un certo numero $m \leq n$ grazie al fatto che in V non possono esserci più di n vettori indipendenti. Se inoltre $\dim(W) = \dim(V)$ significa che W è generato da n vettori indipendenti, che per l'osservazione 4.1 generano anche V , quindi è dimostrata anche la iii). $\square$

Un'altra operazione tra sottospazi è la *somma* di sottospazi.

DEFINIZIONE 4.7. Se V e U sono due sottospazi di uno stesso spazio vettoriale diciamo che W è somma di U e V se i vettori di W sono tutti e soli quelli che si esprimono come somma di un vettore di U e di uno di V cioè

$$W = U + V = \{\vec{w} \mid \vec{w} = \vec{u} + \vec{v}, \text{ con } \vec{u} \in U, \vec{v} \in V\}$$

Si dimostra facilmente (farlo per esercizio) che la somma di sottospazi è a sua volta un sottospazio. Naturalmente non è detto che la scomposizione del vettore w sia unica, questo avviene, però, se e solo se i due spazi sono disgiunti; in questo caso diciamo che la somma è *diretta* e scriviamo

$$W = U \oplus V.$$

Sussiste infatti il

Teorema 4.6. *Siano U e W due sottospazi di V , e sia $\vec{v} \in V$; allora la scomposizione $\vec{v} = \vec{u} + \vec{w}$ con $\vec{u} \in U$ e $\vec{w} \in W$ è unica se e solo se $V = U \oplus W$.*

Dimostrazione. Sia $V = U \oplus W$; per definizione di somma di sottospazi esiste una coppia di vettori $\vec{u} \in U$ e $\vec{w} \in W$ tali che $\vec{v} = \vec{u} + \vec{w}$; dobbiamo dimostrare che tale coppia è unica, infatti se fosse anche $\vec{v} = \vec{u}' + \vec{w}'$ (con $\vec{u}' \in U$ e $\vec{w}' \in W$) si avrebbe $\vec{u} + \vec{w} = \vec{u}' + \vec{w}'$ e quindi $\vec{u} - \vec{u}' = \vec{w}' - \vec{w}$, da cui segue che il vettore $\vec{u} - \vec{u}'$ appartenendo sia a U che a W e di conseguenza alla loro intersezione è il vettore nullo; dunque dev'essere $\vec{u} = \vec{u}'$ e $\vec{w} = \vec{w}'$.

Viceversa supponiamo che sia $\vec{v} = \vec{u} + \vec{w}$ in un unico modo e facciamo vedere che $U \cap W = \{\mathbf{0}\}$. Supponiamo per assurdo che esista un vettore $\vec{z} \in U \cap W$ diverso dal vettore nullo, allora anche i vettori $\vec{u} + \vec{z}$ e $\vec{w} - \vec{z}$, appartenenti rispettivamente a U e W avrebbero come somma $\vec{v}$ contro l'ipotesi dell'unicità della decomposizione. $\square$

Sulla dimensione dello spazio somma di due sottospazi vale la relazione (di Grassmann³)

$$\dim(U + V) = \dim U + \dim V - \dim(U \cap V), \quad (4.6)$$

che diventa, se la somma è diretta,

$$\dim(U \oplus V) = \dim U + \dim V.$$

cioè la dimensione della somma diretta di due sottospazi è uguale alla somma delle loro dimensioni.

Esempio 4.2. *Sia T^n lo spazio vettoriale delle matrici triangolari alte e sia T_n quello delle triangolari basse. Vogliamo verificare la (4.6). Osserviamo che $T^n \cap T_n = D_n$ dove con D_n abbiamo indicato lo spazio vettoriale delle matrici diagonali e che $M_n = T^n + T_n$ lo spazio vettoriale delle matrici quadrate è somma (non diretta) di quello delle matrici triangolari basse e di quello delle matrici triangolari alte.*

Per calcolare la dimensione di T^n , ovviamente uguale a quella di T_n , procediamo così: dalla figura 4.1 nella pagina successiva si vede subito che gli elementi che non sono certamente nulli in T^n ed in T_n sono: 1 nella prima riga 2 nella seconda ecc. quindi in totale si ha $1 + 2 + \dots + n = \frac{n(n+1)}{2}$; dunque

³Herrmann Günther GRASSMANN, 1809, Stettino (Germania-odierna Polonia) – 1877, Stettino (Germania-odierna Polonia).

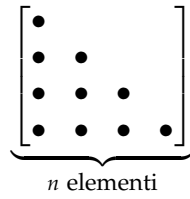


Figura 4.1 Matrici triangolari

possiamo dire che $\dim(T_n) = \dim(T^n) = \frac{n(n+1)}{2}$, inoltre $\dim(D_n) = n$ e $\dim(M_n) = n^2$ ed infatti, applicando la (4.6 nella pagina precedente), si ha

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

Capitolo 5

Determinante e rango di una matrice. Matrice inversa

Sia $\mathcal{M}_n$ l'insieme delle matrici quadrate di ordine n e sia $A \in \mathcal{M}_n$. Il determinante di A è il valore di una funzione che ha come dominio $\mathcal{M}_n$ e come codominio $\mathbb{R}$ o $\mathbb{C}$ (a seconda che gli elementi di A siano numeri reali o complessi¹), quindi $\det : \mathcal{M}_n \mapsto \mathbb{R}$ oppure $\det : \mathcal{M}_n \mapsto \mathbb{C}$.

5.1 Definizioni di determinante

Definiremo il determinante di una matrice quadrata prima in maniera ricorsiva poi in maniera classica.

Definizione ricorsiva

Definiamo il determinante in maniera ricorsiva, cominciando con il definire il determinante di una matrice di ordine 2 ed osservando come si può estendere questa definizione al caso di una matrice di ordine qualsiasi.

DEFINIZIONE 5.1. Sia $A = \begin{bmatrix} a & b \\ c & d \end{bmatrix}$ quadrata di ordine 2, allora il determinante di A è $\det(A) = ad - bc$.

Esempio 5.1. Il determinante della matrice $A = \begin{bmatrix} 1 & -2 \\ 3 & 1 \end{bmatrix}$ è $\det(A) = \begin{vmatrix} 1 & -2 \\ 3 & 1 \end{vmatrix} = 1 \cdot 1 - (-2) \cdot 3 = 7$.

¹Se le matrici sono definite su un campo qualsiasi $\mathbb{K}$ il codominio sarà $\mathbb{K}$.

Osserviamo che il determinante della matrice $A = \begin{bmatrix} x & y \\ z & t \end{bmatrix}$ si può indicare in uno qualsiasi dei seguenti modi: $\det A$, $\det(A)$, $|A|$, $\begin{vmatrix} x & y \\ z & t \end{vmatrix}$.

Diamo ora una definizione ricorsiva che permette di calcolare il determinante di una matrice di ordine qualsiasi quando si sappia calcolare il determinante di una matrice di ordine 2.

Per far questo introduciamo prima una notazione: se A è quadrata di ordine n chiameremo *minore complementare* dell'elemento a_{ik} e lo indicheremo con M_{ik} il determinante della sottomatrice che si ottiene da A cancellando la i -esima riga e la k -esima colonna. Chiamiamo poi *complemento algebrico* dell'elemento a_{ik} (e lo indicheremo con A_{ik}) il determinante della sottomatrice (di ordine $n - 1$) che si ottiene da A cancellando la i -esima riga e la k -esima colonna, con il proprio segno se $i + k$ è pari, col segno opposto se $i + k$ è dispari, cioè

$$A_{ik} = (-1)^{i+k} M_{ik}.$$

Esempio 5.2. Se $A = \begin{bmatrix} a & b & c \\ 1 & 2 & 3 \\ x & y & z \end{bmatrix}$ allora il complemento algebrico dell'elemento $a_{11} = a$ è $A_{11} = \begin{vmatrix} 2 & 3 \\ y & z \end{vmatrix} = 2z - 3y$ e quello dell'elemento $a_{23} = 3$ è $A_{23} = -\begin{vmatrix} a & b \\ x & y \end{vmatrix} = -(ay - bx) = bx - ay$.

DEFINIZIONE 5.2. Sia A una matrice quadrata di ordine n : si ha

$$\det(A) = \sum_{k=1}^n a_{ik} A_{ik} \quad (5.1)$$

cioè il determinante della matrice A è la somma dei prodotti degli elementi di una linea di A (riga o colonna) per i rispettivi complementi algebrici.

La definizione 5.2 ci dice, in sostanza, che per calcolare il determinante di una matrice quadrata di ordine n , possiamo calcolare un certo numero (al massimo n) di determinanti di matrici di ordine $n - 1$, a loro volta questi si determinano calcolando al più $n - 1$ determinanti di matrici di ordine $n - 2$ e così via, quindi, in pratica, basta saper calcolare il determinante di una matrice di ordine 2 con la definizione 5.1 nella pagina precedente ed applicare la ricorsione data nella definizione 5.2.

Esempio 5.3. Sia da calcolare il determinante di $A = \begin{bmatrix} 1 & 0 & -1 \\ 2 & 1 & 3 \\ 0 & -1 & 1 \end{bmatrix}$; prendiamo in considerazione la prima riga: abbiamo che il determinante di A è uguale a

$$1 \cdot \begin{vmatrix} 1 & 3 \\ -1 & 1 \end{vmatrix} - 0 \cdot \begin{vmatrix} 2 & 3 \\ 0 & 1 \end{vmatrix} + (-1) \cdot \begin{vmatrix} 2 & 1 \\ 0 & -1 \end{vmatrix} = 1 \cdot 4 - 0 \cdot 2 + (-1)(-2) = 6.$$

Avremmo potuto pervenire allo stesso risultato scegliendo un'altra qualsiasi riga o colonna (si consiglia di provare per esercizio).

Definizione classica

DEFINIZIONE 5.3. Se A è una matrice quadrata di ordine n , chiamiamo *prodotto associato* ad A il prodotto di n elementi di A presi in modo tale che in ciascuno di essi non ci siano due elementi appartenenti alla stessa riga od alla stessa colonna.

Esempio 5.4. Se $A = \begin{bmatrix} a & b & c \\ 1 & 2 & 3 \\ x & y & z \end{bmatrix}$ sono prodotti associati, per esempio i prodotti $a2z$ e $b1z$ ma non $a1y$ e neppure $c3x$

Se consideriamo la generica matrice di ordine n

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

ogni prodotto associato può essere indicato con $a_{1k_1}a_{2k_2} \cdots a_{nk_n}$ che abbiamo ordinato in ordine crescente rispetto ai primi indici e dove $\{k_1k_2 \dots k_n\}$ è una opportuna permutazione dei secondi indici.

Possiamo ora dare la definizione classica di determinante

DEFINIZIONE 5.4. Si chiama *determinante* della matrice quadrata A la somma di tutti i possibili prodotti associati ad A , presi ciascuno con il proprio segno o con il segno opposto a seconda che la permutazione dei secondi indici sia di classe pari o di classe dispari, cioè, formalmente

$$\det A = \sum (-1)^t \cdot a_{1k_1}a_{2k_2} \cdots a_{nk_n}$$

dove la somma è estesa a tutte le permutazioni dei secondi indici e t è il numero degli scambi che la permutazione dei secondi indici presenta rispetto alla prima.

Con questa impostazione la definizione 5.2 a pagina 46 diventa un teorema che prende il nome di

Teorema 5.1 (Primo teorema di Laplace²). *Sia A una matrice quadrata di ordine n : si ha*

$$\det(A) = \sum_{k=1}^n a_{ik} A_{ik} \quad (5.1)$$

cioè il determinante è la somma dei prodotti degli elementi di una linea (riga o colonna) per i rispettivi complementi algebrici.

Faremo uso anche del

Teorema 5.2 (Secondo teorema di Laplace). *La somma dei prodotti degli elementi di una linea (riga o colonna) per i complementi algebrici degli elementi di una linea parallela è nulla, cioè*

$$\sum_j a_{ij} A_{kj} = 0. \quad (5.2)$$

5.2 Proprietà del determinante

Dalla definizione 5.2, che, come abbiamo visto equivale al teorema 5.1, si ricava, come già detto, che calcolare il determinante di una matrice di ordine n equivale a calcolare al più n determinanti di ordine $n - 1$ e quindi al più $n(n - 1)$ di ordine $n - 2$ e così via al più $n!$ determinanti di ordine 2.

È possibile ridurre e semplificare di molto questi calcoli applicando alcune proprietà dei determinanti che si dimostrano facilmente e che sono elencate nel:

Teorema 5.3. *Sia A una matrice quadrata di ordine n , allora sussistono le seguenti proprietà:*

- i) $\det(A) = \det(A_T)$; *cioè il determinante di una matrice è uguale a quello della sua trasposta.*
- ii) *Scambiando tra loro due colonne³ di A si ottiene una matrice B tale che*
 $|A| = -|B|$.

²Pierre-Simon Laplace, 23 Marzo 1749 in Beaumont-en-Auge, Normandia, Francia-5 March 1827 in Parigi, Francia

³In forza della proprietà i), tutto quello che da qui in poi diciamo sulle colonne vale anche sulle righe.

- iii) Se una colonna di A è nulla, allora $\det(A) = 0$ e la matrice si chiama singolare.
- iv) Se ad una colonna di A si somma una combinazione lineare di altre colonne, si ottiene una matrice B tale che $\det(B) = \det(A)$.
- v) Se due colonne di A sono uguali allora A è singolare, cioè $\det(A) = 0$.
- vi) Moltiplicando per un numero α una colonna di A si ottiene una matrice B tale che $\det(B) = \alpha \det(A)$.
- vii) Se due colonne di A sono proporzionali allora A è singolare.
- viii) Se $B = \alpha A$ allora $\det(B) = \alpha^n \det(A)$.
- ix) Il determinante di una matrice triangolare (in particolare diagonale) è il prodotto dei suoi elementi principali.

Dimostrazione. Diamo un cenno della dimostrazione di alcune delle proprietà: i) è insita nella definizione 5.2; per quanto riguarda la ii) si osserva che scambiando tra loro due colonne cambia la parità di $i + k$ e quindi il segno del determinante; per la iii) basta pensare di sviluppare il determinante rispetto alla colonna nulla; la v) è conseguenza della ii) (perché?); per la vi) basta sviluppare il determinante rispetto alla colonna moltiplicata per α ; la vii) è conseguenza della vi) e della v); la viii) segue dalla vi) notando che ogni colonna è moltiplicata per α la ix) deriva dalla definizione classica e dal fatto che nelle matrici triangolari l'unico prodotto associato non certamente nullo è quello formato dagli elementi principali. $\square$

Osserviamo che applicazioni successive della proprietà iv) del Teorema 5.3 ci permettono di passare dalla matrice A ad un'altra matrice avente lo stesso determinante di A e che ha una linea formata da elementi tutti nulli tranne al più uno. Dunque per calcolare $\det(A)$ possiamo limitarci a calcolare un solo determinante di ordine $n - 1$; iterando il procedimento possiamo calcolare il determinante di A calcolando un solo determinante di ordine 2.

Esempio 5.5. Consideriamo la matrice $A = \begin{bmatrix} 1 & 1 & -1 \\ 1 & 2 & -1 \\ 1 & 1 & 1 \end{bmatrix}$; se alla seconda riga sottraiamo la prima otteniamo la matrice $B = \begin{bmatrix} 1 & 1 & -1 \\ 0 & 1 & 0 \\ 1 & 1 & 1 \end{bmatrix}$ che ha lo

stesso determinante di A e che possiamo sviluppare secondo gli elementi della seconda riga ottenendo $\det(A) = \det(B) = 1 \begin{vmatrix} 1 & -1 \\ 1 & 1 \end{vmatrix} = 2$

Dal Teorema 5.3 a pagina 48 si ricava anche l'importante risultato dato dal

Teorema 5.4. *Una matrice è singolare se e solo se le sue colonne (righe) formano un sistema di vettori linearmente dipendenti.*

Dimostrazione. Se i vettori colonna sono linearmente dipendenti allora uno di essi è combinazione lineare degli altri, quindi la matrice è singolare per la proprietà... , viceversa se $\det(A) = 0$ allora, per le proprietà... i vettori che formano le colonne di A ... $\square$

Per il prodotto di matrici vale il seguente teorema, di cui diamo solo l'enunciato

Teorema 5.5 (di Binet⁴). *Se A e B sono due matrici quadrate dello stesso ordine, allora*

$$\det(A \cdot B) = \det(A) \cdot \det(B). \quad (5.3)$$

Il teorema 5.5 si estende facilmente ad un numero qualsiasi di matrici quadrate dello stesso ordine.

ATTENZIONE l'analogo del Teorema di Binet per la somma di matrici in generale non vale cioè si ha che, in generale,

$$\det(A + B) \neq \det(A) + \det(B). \quad (5.4)$$

Come esercizio verificare la (5.4) su qualche esempio.

5.3 Un determinante particolare

Consideriamo ora il determinante

$$V(a_1, \dots, a_n) = \begin{vmatrix} 1 & a_1 & a_1^2 & \dots & a_1^{n-1} \\ 1 & a_2 & a_2^2 & \dots & a_2^{n-1} \\ \dots & \dots & \dots & \dots & \dots \\ 1 & a_n & a_n^2 & \dots & a_n^{n-1} \end{vmatrix}$$

in cui la prima colonna è formata da tutti 1 (cioè $a_i^0 \forall i$) e le altre colonne sono rispettivamente le successive potenze della seconda colonna; esso

⁴Jacques Philippe Marie BINET, 2 Feb 1786 in Rennes, Bretagne, France–12 May 1856 in Paris, France

prende il nome di *determinante di Vandermonde*⁵ o *determinante delle differenze*, infatti si verifica molto facilmente (farlo, come esercizio) che:

$$\begin{aligned} V(a_1, \dots, a_n) &= (a_2 - a_1)(a_3 - a_1) \dots \\ &\dots (a_3 - a_2) \dots \\ &\dots \dots \\ &\dots \dots (a_n - a_{n-1}). \end{aligned}$$

cioè il determinante è uguale al prodotto delle differenze dei suoi elementi a due a due in tutti i modi possibili. Da cui segue il

Teorema 5.6. *Il determinante di Vandermonde è uguale a zero se e solo se almeno due dei suoi argomenti sono uguali.*

5.4 Rango di una matrice

Siamo ora in grado di dare un'altra definizione di rango, formalmente diversa dalla definizione 3.6 a pagina 29, ma che si dimostra facilmente essere ad essa equivalente. Per far ciò dobbiamo premettere una definizione

DEFINIZIONE 5.5. Se A è di tipo (m, n) (quindi non necessariamente quadrata), si chiama *minore di ordine k* ($k \leq \min(m, n)$) *il determinante di una qualunque sottomatrice quadrata di A formata da k righe e k colonne di A .*

ATTENZIONE non è richiesto che le k righe e le k colonne siano adiacenti.

Esempio 5.6. Se $A = \begin{bmatrix} a & b & c \\ 1 & 2 & 3 \\ x & y & z \end{bmatrix}$ allora sono minori del secondo ordine i determinanti $\begin{vmatrix} 2 & 3 \\ y & z \end{vmatrix}$ oppure $\begin{vmatrix} a & c \\ x & z \end{vmatrix}$ invece il determinante $\begin{vmatrix} a & b \\ 1 & 3 \end{vmatrix}$ non lo è, infatti la seconda colonna non è formata da elementi tratti da una colonna di A .

DEFINIZIONE 5.6. [di rango] Si chiama *rango o caratteristica* di una matrice di tipo (m, n) l'ordine massimo dei minori non nulli che si possono estrarre da A .

⁵Alexandre-Théophile Vandermonde, 28 Feb 1735 in Paris, France – 1 Jan 1796 in Paris, France

Osserviamo esplicitamente che la definizione di rango si applica ad una matrice qualsiasi, non necessariamente quadrata, mentre si parla di determinante solo per le matrici quadrate.

La definizione 5.6 nella pagina precedente significa dunque che se dalla matrice A possiamo estrarre un minore di ordine k non nullo e tutti i minori di ordine più grande di k sono nulli, allora $r(A) = k$.

Dalla definizione segue subito che se A è quadrata di ordine n allora essa ha rango n se e solo se è non singolare, e che, se A è di tipo (m, n) , vale la relazione, già vista, $r(A) \leq \min(m, n)$.

Esempio 5.7. Sia $A = \begin{bmatrix} 1 & 2 & 3 & 4 \\ 0 & 1 & 0 & 3 \\ 3 & 1 & 4 & 2 \end{bmatrix}$; A è di tipo (3×4) quindi si osserva anzitutto che $r(A) \leq 3$ e che il minore $\begin{vmatrix} 1 & 2 \\ 0 & 1 \end{vmatrix}$ è non nullo, quindi è $2 \leq r(A) \leq 3$, inoltre si calcola rapidamente il determinante della sottomatrice formata dalle prime tre colonne che è diverso da 0, dunque $r(A) = 3$.

Osserviamo che per il rango di una matrice valgono, tra le altre, le proprietà espresse dal teorema 5.7, che si ricavano immediatamente dal Teorema 5.3 a pagina 48 e dal Teorema 5.4 a pagina 50.

Teorema 5.7. Due importanti proprietà del rango di una matrice sono:

- i) Due matrici ottenute una dall'altra mediante uno scambio di righe o colonne hanno lo stesso rango,
- ii) Il rango di una matrice è uguale al numero di righe o di colonne linearmente indipendenti presenti nella matrice stessa.

Sempre dalle proprietà dei determinanti e da quelle del rango (Teoremi 5.3 a pagina 48, 5.4 a pagina 50 e 5.7, rispettivamente) segue l'importante

Teorema 5.8. Se la matrice A ha rango r e b è un vettore colonna, allora la matrice $A|b$, ottenuta completando la A con la colonna b , ha rango r se e solo se b è combinazione lineare delle colonne di A .

Dimostrazione. L'aggiunta di una colonna non può far diminuire il rango e lo aumenta se e solo se la colonna è indipendente dalle altre. $\square$

5.5 Calcolo del rango

Per calcolare il rango di una matrice di tipo (m, n) dobbiamo esaminare tutti i minori di ordine $k = \min(m, n)$: se ce n'è uno non nullo il rango è k , se invece sono *tutti* nulli passeremo ad esaminare quelli di ordine $k - 1$ e così via; oppure, se “vediamo” un minore di ordine $p < k$ non nullo esamineremo *tutti* quelli di ordine $p + 1$: se sono tutti nulli il rango è p altrimenti sarà $r \geq p + 1$ e così via.

Questo procedimento può essere molto abbreviato applicando il Teorema di Kroneker 5.9, al quale dobbiamo però premettere la

DEFINIZIONE 5.7. Sia A una matrice qualsiasi e sia M una sua sottomatrice; *orlare* M significa completare la sottomatrice M con una riga ed una colonna di A non appartenenti a M .

Ovviamente la sottomatrice M non è detto sia formata da righe o colonne che in A sono adiacenti, pertanto la riga e la colonna che completano possono anche essere “interne” come nel seguente

Esempio 5.8. Se A è la matrice $\begin{bmatrix} 1 & 2 & 3 & 4 & 5 \\ a & b & c & d & e \\ x & y & z & k & t \end{bmatrix}$, una sua sottomatrice è

$M = \begin{bmatrix} 3 & 5 \\ z & t \end{bmatrix}$ che può essere orlata con la seconda riga e la quarta colonna di

A ottenendo la matrice $\begin{bmatrix} 3 & 4 & 5 \\ c & d & e \\ z & k & t \end{bmatrix}$ oppure con la seconda riga e, per esempio,

la prima colonna, ottenendo la matrice $\begin{bmatrix} 1 & 3 & 5 \\ a & c & e \\ x & z & t \end{bmatrix}$.

Siamo ora in grado di enunciare (senza dimostrazione) il seguente utilissimo

Teorema 5.9 (di Kroneker⁶). *Se in una matrice A esiste un minore M non nullo di ordine p e tutti i minori che orlano M sono nulli, allora il rango di A è p .*

Il teorema 5.9 permette quindi di limitare il controllo dei minori di ordine $p + 1$ a quelli che orlano il minore M .

⁶Leopold KRONEKER, 1813, Liegnitz, Prussia (oggi Polonia) – 1881, Berlino.

Ad esempio vogliamo il rango della matrice $A = \begin{bmatrix} 1 & 0 & 3 & 2 \\ 2 & 3 & 0 & 1 \\ 3 & 3 & 3 & 3 \end{bmatrix}$. Si osserva subito che il minore $M = \begin{bmatrix} 3 & 0 \\ 3 & 3 \end{bmatrix}$ è diverso da 0, quindi $2 \leq r(A) \leq 3$. In virtù del Teorema 5.9 nella pagina precedente possiamo limitarci a controllare se sono nulli i *due* minori del terz'ordine che orlano M (anziché controllare tutti e *quattro* i minori del terz'ordine di A); i minori che ci interessano sono quelli formati dalla I, II e III colonna e dalla II, III e IV, cioè: $M_1 = \begin{bmatrix} 1 & 0 & 3 \\ 2 & 3 & 0 \\ 3 & 3 & 3 \end{bmatrix}$ e $M_2 = \begin{bmatrix} 0 & 3 & 2 \\ 3 & 0 & 1 \\ 3 & 3 & 3 \end{bmatrix}$ entrambi palesemente nulli, in quanto in entrambi la terza riga è la somma delle prime due, dunque $r(A) = 2$.

Una proprietà del rango del prodotto di due matrici, spesso utile nelle applicazioni, è espressa dal seguente

Teorema 5.10. *Il rango del prodotto di due matrici A e B non supera il rango di ciascuna delle due, cioè $r(A \cdot B) \leq r(A)$ e $r(A \cdot B) \leq r(B)$, che equivale a scrivere*

$$r(AB) \leq \min(r(A), r(B)).$$

5.6 Matrice inversa

Sia A una matrice quadrata di ordine n ; una matrice B tale che

$$A \cdot B = B \cdot A = I$$

prende il nome di *inversa* di A . Una matrice che ammette inversa è detta *invertibile*.

Sorge il problema di stabilire quali siano le matrici invertibili e quante inverse abbia ciascuna di esse. Al primo quesito risponde il seguente

Teorema 5.11. *Una matrice A è invertibile se e solo se è non singolare, cioè se e solo se $\det(A) \neq 0$.*

Dimostrazione. Se A ammette come inversa B allora $AB = I$ e dal Teorema (di Binet) 5.5 a pagina 50 si ha

$$\det(A \cdot B) = \det(A) \cdot \det(B) = \det(I) = 1$$

dunque, per la legge di annullamento del prodotto, nè A nè B possono essere singolari.

Viceversa, supponiamo che $\det(A) \neq 0$, consideriamo la matrice A^* , detta *matrice dei complementi algebrici*, il cui elemento generico α_{ik} è il complemento algebrico dell'elemento a_{ki} , cioè $\alpha_{ik} = A_{ki}$. Allora se c_{ik} è il generico elemento della matrice AA^* si ha

$$c_{ik} = \sum_j a_{ij} \alpha_{jk} = \sum_j a_{ij} A_{kj} = \begin{cases} \det A & \text{se } i = k, \text{ per il I Teorema di Laplace 5.1} \\ 0 & \text{se } i \neq k, \text{ per il II Teorema di Laplace 5.2} \end{cases};$$

questo significa che

$$A \cdot A^* = \begin{bmatrix} \det(A) & 0 & \dots & 0 \\ 0 & \det(A) & \dots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \dots & \det(A) \end{bmatrix} = \det(A) \cdot I$$

e quindi che la matrice $\frac{A^*}{\det(A)}$ è un'inversa di A . $\square$

Al secondo quesito risponde il

Teorema 5.12. *Se l'inversa di A esiste, essa è unica.*

Dimostrazione. Dimostriamo il teorema per assurdo e supponiamo che B e C siano due inverse di A ; allora si ha: $AB = I = AC$ ma anche, moltiplicando a sinistra per B , $B(AB) = B(AC)$ e, per l'associatività del prodotto di matrici, $(BA)B = (BA)C$ da cui $B = C$. $\square$

L'unica inversa di una matrice invertibile A sarà, d'ora in poi, indicata con A^{-1} .

Le principali proprietà delle matrici invertibili sono date dal

Teorema 5.13. *Per le matrici invertibili valgono le seguenti proprietà:*

- i) *Se A è invertibile, allora $AB = AC$ implica $B = C$.*
- ii) *Se A è invertibile, allora $BA = CA$ implica $B = C$.*
- iii) *Se A e B sono due matrici quadrate tali che $AB = \mathbf{0}$ allora sussiste uno ed uno solo dei seguenti casi*
 - a) *nè A nè B sono invertibili.*
 - b) *A è invertibile, e allora B è la matrice nulla.*

c) B è invertibile, e allora A è la matrice nulla.

iv) Se A e B sono invertibili, allora il prodotto AB è invertibile e si ha $(AB)^{-1} = B^{-1}A^{-1}$;

Dimostrazione. i) Basta moltiplicare a sinistra ambo i membri per A^{-1} .

ii) Basta moltiplicare a destra ambo i membri per A^{-1} .

iii) Se $AB = 0$ almeno una delle due matrici è singolare per il Teorema di Binet (5.5 a pagina 50).

iv) Infatti, dal teorema di Binet si ricava che il prodotto di matrici non singolari è non singolare e si può scrivere

$$B^{-1}A^{-1}AB = B^{-1}(A^{-1}A)B = B^{-1}B = I$$

in cui abbiamo applicato l'associatività del prodotto.

□

Concludiamo il capitolo osservando che anche per le matrici non singolari si può parlare di potenza ad esponente negativo, definendo A^{-h} come l'inversa di A^h . Per esercizio dimostrare che A^{-h} pensata come inversa di A^h è anche la potenza h -esima di A^{-1} .

Capitolo 6

Teoria dei sistemi lineari

Abbiamo ora in mano tutti gli strumenti necessari per completare lo studio dei sistemi lineari; in particolare per decidere quando un sistema è possibile e quante soluzioni ammette, cioè per studiare la teoria dei sistemi lineari.

Ricordiamo che, per l'appunto, questa teoria si riferisce solo ai sistemi lineari, e non è applicabile a sistemi di grado superiore al primo.

6.1 Teoremi sui sistemi lineari

Cominciamo a considerare il caso particolare di un sistema in cui il numero delle equazioni (diciamo n) è uguale a quello delle incognite. In questo caso sussiste il

Teorema 6.1 (di Cramer¹). *Un sistema lineare di n equazioni in n incognite della forma $A\vec{x} = \vec{b}$ la cui matrice dei coefficienti sia quadrata e non singolare ammette una ed una sola soluzione costituita dalla n -pla*

$$x_1 = \frac{\det(A_1)}{\det(A)}, \quad x_2 = \frac{\det(A_2)}{\det(A)}, \quad \dots, \quad x_n = \frac{\det(A_n)}{\det(A)},$$

dove A_i è la matrice ottenuta dalla A sostituendo al posto della i -esima colonna la colonna dei termini noti.

Dimostrazione. Il sistema può essere scritto, in forma matriciale, come $A\vec{x} = \vec{b}$ e la matrice A è, per ipotesi, non singolare, dunque esiste A^{-1} . Allora si ha, moltiplicando a sinistra per A^{-1} entrambi i membri, $A^{-1}A\vec{x} = A^{-1}\vec{b}$ e quindi $\vec{x} = A^{-1}\vec{b}$. Ma ricordando che $A^{-1} =$

¹Gabriel CRAMER, 1704, Ginevra – 1752, Bagnols sur Céze (Francia).

$\frac{A^*}{\det(A)}$ si conclude che

$$\vec{x} = \frac{1}{\det(A)} A^* \vec{b};$$

osserviamo ora che $A^* \vec{b}$ è un vettore colonna, ciascuno dei componenti del quale è, ricordando la definizione 5.2 a pagina 46, il $\det(A_i)$. $\square$

Per un generico sistema lineare, in cui il numero delle equazioni non è necessariamente uguale a quello delle incognite, vale il Teorema (3.5 a pagina 30) di Rouché-Capelli, di cui qui diamo una dimostrazione basata sulla definizione di rango 5.6 a pagina 51 che abbiamo visto nel capitolo precedente.

Ricordiamo l'enunciato del Teorema 3.5:

Teorema (3.5 di Rouché–Capelli). *Sia $A\vec{x} = \vec{b}$ un sistema lineare di m equazioni in n incognite; esso ammette soluzioni se e solo se $r(A) = r(A|b)$ dove con $A|b$ abbiamo indicato la matrice ottenuta da A completandola con la colonna dei termini noti.*

Dimostrazione. Sia

$$A\vec{x} = \vec{b} \tag{6.1}$$

un sistema lineare di m equazioni in n incognite e immaginiamo di scrivere la matrice $A = [A_1 \ A_2 \ \dots \ A_n]$ scomposta in blocchi formati ciascuno da una delle sue colonne, che indicheremo con A_i . Allora la relazione (6.1) si può scrivere come

$$\vec{A}_1 x_1 + \vec{A}_2 x_2 + \dots + \vec{A}_n x_n = \vec{b}$$

e quindi ci dice che il sistema ammette soluzioni se e solo se $\vec{b}$ è combinazione lineare delle colonne di A . Ma allora, in virtù del Teorema 5.8 a pagina 52 questo accade se e solo se il rango di A è uguale al rango della matrice completa. $\square$

Segue anche, sia dal Teorema 3.5, sia dalla definizione di rango data nel precedente capitolo, che se un sistema possibile ha rango r , esistono esattamente r equazioni e r incognite indipendenti, dunque le altre $m - r$ equazioni sono combinazione lineare delle r indipendenti e non dicono nulla di nuovo, quindi si possono trascurare, e le altre $n - r$ incognite si possono considerare come parametri; dunque

Proposizione 6.2. *Un sistema lineare possibile di m equazioni in n incognite in cui il rango della matrice dei coefficienti sia $r < n$ ammette ∞^{n-r} soluzioni, cioè infinite soluzioni dipendenti da $n - r$ parametri. Se $r = n$ il sistema equivale ad un sistema di r equazioni in r incognite con matrice dei coefficienti non singolare, quindi ammette una ed una sola soluzione per il Teorema di Cramer (6.1 a pagina 57).*

Se il vettore $\vec{b} = \mathbf{0}$ è il vettore nullo, il sistema $A\vec{x} = \mathbf{0}$ si chiama *omogeneo*. Segue immediatamente dalla definizione (e dal teorema 3.5) che un sistema omogeneo è sempre possibile ed ammette sempre come soluzione banale il vettore nullo. Siamo quindi interessati ad eventuali soluzioni non banali (dette anche *autosoluzioni*). Una semplice conseguenza del Teorema di Cramer (6.1) è il

Corollario 6.3. *Un sistema omogeneo di n equazioni in n incognite $Ax = \mathbf{0}$ ammette soluzioni non banali se e solo se $\det(A) = 0$.*

Poichè aggiungendo ad una matrice una colonna nulla il rango non cambia, segue dal teorema di Rouché-Capelli (3.5 a pagina 30) e dalla Proposizione 6.2, il

Corollario 6.4. *Un sistema omogeneo di m equazioni in n incognite ammette autosoluzioni se e solo se $r(A) < n$.*

Se il rango della matrice dei coefficienti è r allora il sistema possiede $n - r$ soluzioni indipendenti, nel senso che tutte le altre, che, ricordiamo, sono infinite, sono combinazioni lineari delle precedenti.

Esempio 6.1. *Consideriamo il sistema*

$$\begin{cases} hx + z = h \\ (h+2)x + 3y = 3 \\ (h-2)y + z = h-1 \end{cases}.$$

Si tratta di un sistema di tre equazioni in tre incognite, applichiamo quindi il Teorema di Cramer. Il determinante dei coefficienti è: $\begin{vmatrix} h & 0 & 1 \\ h+2 & 3 & 0 \\ 0 & h-2 & 1 \end{vmatrix} = (h-1)(h+4)$.

Quindi per $h \neq 1$ e $h \neq -4$ il sistema ammette una ed una sola soluzione:

$$\begin{aligned} x &= \frac{\begin{vmatrix} h & 0 & 1 \\ 3 & 3 & 0 \\ h-1 & h-2 & 1 \end{vmatrix}}{(h-1)(h+4)}, \\ y &= \frac{\begin{vmatrix} h & h & 1 \\ h+2 & 3 & 0 \\ 0 & h-1 & 1 \end{vmatrix}}{(h-1)(h+4)}, \\ z &= \frac{\begin{vmatrix} h & 0 & h \\ h+2 & 3 & 3 \\ 0 & h-2 & h-1 \end{vmatrix}}{(h-1)(h+4)} \end{aligned}$$

Per $h = 1$ si ha

$$\begin{cases} x + z = 1 \\ x + y = 1 \\ y - z = 0 \end{cases}$$

che ammette le ∞^1 soluzioni $x = k$, $y = 1 - k$, $z = 1 - k$.

Per $h = -4$ il sistema diventa

$$\begin{cases} 4x - z = 4 \\ 2x - 3y = -3 \\ 6y - z = 5 \end{cases}$$

in cui la matrice dei coefficienti ha rango 2 mentre quella completa ha rango 3, pertanto il sistema è impossibile.

Esempio 6.2. *Discutiamo il sistema*

$$\begin{cases} hx + (h+2)y = 0 \\ (h+2)y = h \\ (h+1)x = 1 \end{cases}$$

di tre equazioni in due incognite. Se la matrice completa (di ordine 3) non è singolare, ha rango 3 e quindi il sistema è impossibile, perché la matrice dei coefficienti è di tipo $(3, 2)$, e di conseguenza ha rango al più uguale a 2; quindi i valori di h per cui il sistema può essere possibile sono da ricercare solo tra quelli che annullano il determinante della matrice completa, nel nostro caso

$h = 0$ e $h = -2$. Per tutti gli altri valori il sistema non ammette soluzioni. Per $h = 0$ il sistema diventa

$$\begin{cases} y = 0 \\ y = 0 \\ x = 1 \end{cases}$$

e quindi la soluzione è $x = 1, y = 0$; per $h = -2$ si ha

$$\begin{cases} x = 0 \\ 0 = -2 \\ -x = 1 \end{cases}$$

manifestamente impossibile (verificare che in questo caso i due ranghi sono diversi).

OSSERVAZIONE 6.1. Consideriamo un sistema lineare $\mathcal{S} : Ax = b$; il sistema lineare omogeneo $Ax = \mathbf{0}$ che ha la stessa matrice dei coefficienti si chiama sistema omogeneo associato a $\mathcal{S}$. Se si conoscono la soluzione generale x_0 del sistema omogeneo associato ed una soluzione particolare x_1 del sistema $\mathcal{S}$, la soluzione generale di quest'ultimo si può esprimere come

$$x = x_0 + x_1 \quad (6.2)$$

Infatti, ricordando che $Ax_0 = \mathbf{0}$, si ha $A(x_0 + x_1) = Ax_0 + Ax_1 = b$ e, viceversa se $Ax = b$ si ha $A(x - x_1) = Ax - Ax_1 = \mathbf{0}$, e quindi, posto $x_0 = x - x_1$ si ha $x = x_0 + x_1$.

Esempio 6.3. Si consideri il sistema

$$\begin{cases} x + 2y + 2z = 4 \\ 2x + y + z = 2 \end{cases}'$$

di due equazioni in tre incognite la cui matrice dei coefficienti ha rango 2, che ammette dunque ∞^1 soluzioni. Si vede subito che una soluzione particolare è data dalla terna $x = 0, y = 1, z = 1$. Per trovare la soluzione generale consideriamo il sistema omogeneo ad esso associato che è:

$$\begin{cases} x + 2y + 2z = 0 \\ 2x + y + z = 0 \end{cases}$$

la cui soluzione generale è $x = 0, y = t, z = -t$ dunque la soluzione generale del sistema dato sarà $x = 0, y = 1 + t, z = 1 - t$.

Capitolo 7

Applicazioni lineari, prodotto scalare

Il concetto di applicazione o funzione¹ è uno dei più importanti e dei più generali di tutta la Matematica. Abitualmente una funzione tra spazi vettoriali si chiama più propriamente *applicazione*.

7.1 Generalità

Ricordiamo che usualmente si scrive $f : A \mapsto B$ oppure $A \xrightarrow{f} B$ intendendo dire che stiamo considerando la funzione f dell'insieme A nell'insieme B .

L'elemento $y = f(x)$ appartiene a B e si chiama *immagine di x mediante f* , viceversa, l'elemento x di A di cui y è immagine si chiama *controimmagine di y* . L'insieme A si chiama *dominio* della funzione f , e si indica con $\text{dom } f$ e l'insieme B si chiama *codominio di f* . Dunque una funzione è data quando sono dati: la legge rappresentata dalla f , il dominio ed il codominio.

L'insieme di tutti gli elementi del codominio che sono immagini di qualche elemento del dominio A si chiama *immagine di f* e lo denoteremo con $\text{Im}_A(f)$, talvolta sottintendendo, nella notazione, il dominio di f .

Da quanto detto si ha subito che $\text{Im}(f) \subseteq B$.

Se accade che $\text{Im}(f) = B$ la funzione si chiama *suriettiva*; cioè una funzione è suriettiva se tutti gli elementi del codominio hanno una controimmagine.

¹Sono moltissimi i sinonimi del vocabolo funzione, alcuni usati più propriamente in contesti particolari, tra i tanti ricordiamo *applicazione*, *trasformazione*, *corrispondenza*, *mappa*, *operatore*, *morfismo*, ecc.

Se invece due elementi distinti dell'insieme di definizione hanno sempre immagini distinte, cioè se

$$\forall x \neq x' \in \text{def}(f) \implies f(x) \neq f(x')$$

allora la funzione si chiama *iniettiva*.

Una funzione che sia suriettiva ed iniettiva è detta *biiettiva* o *biunivoca*.

Osserviamo che la suriettività dipende dal codominio, mentre l'iniettività dal dominio della funzione.

Ad esempio la funzione $f : \mathbb{R} \mapsto \mathbb{R}$ data da $f(x) = x^2$ non è iniettiva, infatti, per esempio $f(-3) = f(3) = 9$ cioè esistono elementi distinti del dominio che hanno la stessa immagine e non è nemmeno suriettiva, perché, per esempio, -1 appartiene al codominio di f ma non ha controimmagine nel dominio; se ora invece cambiamo il codominio e consideriamo sempre la stessa funzione $f : \mathbb{R} \rightarrow \mathbb{R}^+$ con $f(x) = x^2$ allora la f non è iniettiva ma è suriettiva; viceversa se cambiamo il dominio, la funzione $f : \mathbb{R}^+ \rightarrow \mathbb{R}$ data da $f(x) = x^2$ è iniettiva ma non suriettiva.

Se dominio e codominio sono spazi vettoriali su uno stesso campo $\mathbb{K}$ (che indicheremo, rispettivamente, con V e W), si dice che $f : V \mapsto W$ è un'applicazione lineare o un omomorfismo tra i due spazi V e W se f conserva le combinazioni lineari, cioè se $\forall \vec{v}_1, \vec{v}_2 \in V$ e $\forall \alpha \in \mathbb{K}$ si ha:

$$\begin{aligned} \overrightarrow{f(\vec{v}_1 + \vec{v}_2)} &= \overrightarrow{f(\vec{v}_1)} + \overrightarrow{f(\vec{v}_2)}, \\ \overrightarrow{f(\alpha \vec{v})} &= \alpha \overrightarrow{f(\vec{v})} \end{aligned}$$

che si può scrivere anche come

$$\overrightarrow{f(\alpha \vec{v}_1 + \beta \vec{v}_2)} = \alpha \overrightarrow{f(\vec{v}_1)} + \beta \overrightarrow{f(\vec{v}_2)},$$

Esempio 7.1. L'applicazione $\mathbb{R}^2 \rightarrow \mathbb{R}^2$ che associa al vettore $\vec{v} = [x, y]$ il vettore $\overrightarrow{f(v)} = [y, x^2]$ non è lineare, come è facile verificare. Infatti se $\vec{v} = [a, b]$ e $\vec{w} = [c, d]$ si avrà $\overrightarrow{f(\vec{v})} = [b, a^2]$ e $\overrightarrow{f(\vec{w})} = [d, c^2]$ e quindi $\overrightarrow{f(\vec{v})} + \overrightarrow{f(\vec{w})} = [b + d, a^2 + c^2]$ che è diverso da $\overrightarrow{f(\vec{v} + \vec{w})} = [b + d, (a + c)^2]$.

Invece l'applicazione $f : \mathbb{R} \rightarrow \mathbb{R}$ tale che $[x, y] \mapsto [x + y, x - y]$ è lineare ma questa verifica la lasciamo come esercizio.

Abbiamo già parlato dell'immagine di f , che è un sottoinsieme del codominio; è anche importante il sottoinsieme del dominio dato dalla

DEFINIZIONE 7.1. Si chiama *nucleo*² dell'applicazione $V \xrightarrow{f} W$ e si indica con $\text{Ker } f$ il sottoinsieme di V formato dagli elementi che hanno come immagine lo zero di W :

$$\text{Ker } f = \{ \vec{v} \in V \mid \overrightarrow{f(\vec{v})} = \mathbf{0}_W \}$$

È molto facile dimostrare il

Teorema 7.1. Il nucleo di un'applicazione lineare $V \xrightarrow{f} W$ tra due spazi vettoriali V e W è un sottospazio di V .

Dimostrazione. Basta far vedere che il nucleo è chiuso rispetto alle combinazioni lineari: infatti se $\vec{v}_1$ e $\vec{v}_2$ sono vettori del nucleo, si ha, sfruttando la linearità di f ,

$$\overrightarrow{f(\alpha \vec{v}_1 + \beta \vec{v}_2)} = \alpha \overrightarrow{f(\vec{v}_1)} + \beta \overrightarrow{f(\vec{v}_2)} = \alpha \mathbf{0}_W + \beta \mathbf{0}_W = \mathbf{0}_W$$

quindi anche $\alpha \vec{v}_1 + \beta \vec{v}_2$ è un vettore del nucleo. $\square$

Da questo teorema segue anche, ovviamente, che il vettore nullo appartiene al nucleo.

Altrettanto facile è dimostrare il

Teorema 7.2. L'immagine dell'applicazione lineare $V \xrightarrow{f} W$ è un sottospazio di W .

Dimostrazione. La dimostrazione è lasciata come esercizio ... basta far vedere che l'immagine è... $\square$

Il nucleo e l'immagine di un'applicazione lineare sono legati all'iniettività ed alla suriettività dal

Teorema 7.3. Un'applicazione lineare $V \xrightarrow{f} W$ è:

- i) suriettiva se e solo se $\text{Im } f = W$,
- ii) iniettiva se e solo se $\text{Ker } f = \{\mathbf{0}_V\}$ cioè il nucleo consiste solo nel vettore nullo.

²*kern* viene dalla parola inglese *kernel* che significa, appunto, nucleo.

Dimostrazione. la *i*) segue dalla definizione di suriettività. Dimostriamo quindi la *ii*). Sia f iniettiva, allora, poiché il vettore nullo appartiene al nucleo si ha $\overrightarrow{f(0_V)} = 0_W$ e, per l'iniettività, non può esistere un altro vettore $\vec{v}$ tale che $\overrightarrow{f(\vec{v})} = 0_W$ con $\vec{v} \neq 0_V$.

Viceversa sia $\text{Ker } f = \{0_V\}$ e sia $\overrightarrow{f(\vec{v}_1)} = \overrightarrow{f(\vec{v}_2)}$ allora si ha, per la linearità di f , $\overrightarrow{f(\vec{v}_1)} - \overrightarrow{f(\vec{v}_2)} = \overrightarrow{f(\vec{v}_1 - \vec{v}_2)} = 0_W$ e quindi $\vec{v}_1 - \vec{v}_2 \in \text{Ker } f$, dunque, per l'ipotesi, $\vec{v}_1 - \vec{v}_2 = 0_V$ da cui $\vec{v}_1 = \vec{v}_2$, quindi l'applicazione è iniettiva. $\square$

Naturalmente non tutte le applicazioni lineari sono iniettive o suriettive, per esempio l'applicazione $\mathbb{R}^2 \rightarrow \mathbb{R}^2$ che manda il vettore $[x, y]$ nel vettore $[x + y, 0]$ non è iniettiva, perché il nucleo è formato da tutti i vettori del tipo $[a, -a]$, nè suriettiva, perché il vettore $[a, b]$ con $b \neq 0$ non è immagine di alcun vettore del dominio.

Esistono anche applicazioni lineari che sono simultaneamente suriettive ed iniettive, tali applicazioni, come abbiamo detto prendono il nome di *applicazioni biunivoche* o *isomorfismi*; ad esempio l'applicazione $\mathcal{M}_2 \mapsto \mathbb{R}^4$ che associa alla matrice quadrata $\begin{bmatrix} a & b \\ c & d \end{bmatrix}$ il vettore $[a, b, c, d]$ è un isomorfismo (dimostrarlo per esercizio).

Se $f : V \rightarrow W$ e $\overrightarrow{f(\vec{v})} = \vec{w}$ può esistere una applicazione $g : W \rightarrow V$ tale che $\overrightarrow{g(\vec{w})} = \vec{v}$ in tal caso si dice che g è l'*applicazione inversa* di f e si indica con f^{-1} .

È molto utile nelle applicazioni il

Teorema 7.4. *Se $f : V \rightarrow W$ vale la relazione*

$$\dim(V) = \dim(\text{ker } f) + \dim(\text{Im } f)$$

cioè la somma tra dimensione del nucleo di una applicazione lineare e la dimensione dell'immagine uguaglia la dimensione di V .

Dimostrazione. Sia U uno spazio supplementare di $\text{Ker } f$ cioè sia U tale che $V = \text{Ker } f \oplus U$, e sia $\dim(U) = s$. Sia inoltre $\mathcal{B}' = \{\vec{e}'_1, \dots, \vec{e}'_q\}$ una base di V tale che i suoi primi s vettori costituiscano una base di U ed i successivi $q - s$ vettori siano una base per $\text{ker } f$: dimostriamo che i vettori $\overrightarrow{f(\vec{e}'_1)}, \dots, \overrightarrow{f(\vec{e}'_s)}$ sono linearmente indipendenti. Da $a_1 f(\vec{e}'_1) + \dots + a_s f(\vec{e}'_s) = 0_W$ segue, per la linearità di f , che è $f(a_1 \vec{e}'_1 + \dots + a_s \vec{e}'_s) = 0_W$, dunque $a_1 \vec{e}'_1 + \dots + a_s \vec{e}'_s \in \text{ker}(f) \cap U$ e quindi, tenendo

5

È facile rendersi conto che un'applicazione lineare f tra due spazi vettoriali V e W è nota quando si sa come si trasformano i vettori di una base di V ; in altre parole quando si conoscono i trasformati mediante f dei vettori di una base di V . Consideriamo una base $\mathcal{B} = \{\vec{e}_1, \vec{e}_2, \dots, \vec{e}_n\}$ di V ed una base $\mathcal{B}' = \{\vec{e}'_1, \vec{e}'_2, \dots, \vec{e}'_m\}$ di W , se $f(\vec{e}_1), \dots, f(\vec{e}_n)$ sono i trasformati mediante la f dei vettori di $\mathcal{B}$ è chiaro che ciascuno di essi, appartenendo a W , si scrive come combinazione lineare dei vettori di $\mathcal{B}'$, dunque si ha il sistema

[illegible]

$$\Gamma = \begin{bmatrix} a_{11} & a_{12} & \dots & a_{1m} \\ a_{21} & a_{22} & \dots & a_{2m} \\ \vdots & \vdots & \dots & \vdots \\ a_{n1} & a_{n2} & \dots & a_{nm} \end{bmatrix},$$

Esempio 7.2. Consideriamo l'applicazione $f : \mathbb{R}^2 \mapsto \mathbb{R}^2$ che manda il vettore $\vec{v} = [x, y]$ nel vettore $\overrightarrow{f(\vec{v})} = [x + y, 0]$ essa è lineare (verificarlo per esercizio), inoltre si ha

$$\begin{aligned}\vec{e}_1 &= [1, 0] \longrightarrow [1, 0] = \vec{e}'_1 \\ \vec{e}_2 &= [0, 1] \longrightarrow [1, 0] = \vec{e}'_1\end{aligned}$$

e quindi la matrice associata a f rispetto alle basi canoniche è $\Gamma = \begin{bmatrix} 1 & 0 \\ 1 & 0 \end{bmatrix}$.

Esempio 7.3. Sia ora $f : \mathbb{R}^3 \rightarrow \mathbb{R}^2$ tale che $[x, y, z] \rightarrow [x + z, y + z]$. Si ha

$$\begin{aligned}\vec{e}_1 = [1, 0, 0] &\longrightarrow [1, 0] = \vec{e}'_1 \\ \vec{e}_2 = [0, 1, 0] &\longrightarrow [0, 1] = \vec{e}'_2 \\ \vec{e}_3 = [0, 0, 1] &\longrightarrow [1, 1] = \vec{e}'_1 + \vec{e}'_2\end{aligned}$$

dunque la matrice associata a questa applicazione, rispetto alle basi canoniche,

è la matrice $\begin{bmatrix} 1 & 0 \\ 0 & 1 \\ 1 & 1 \end{bmatrix}$.

Scegliendo basi diverse la matrice associata ad f ovviamente cambia, però, ricordando che anche il cambiamento di base si rappresenta mediante una matrice, si capisce che la matrice Γ' associata ancora ad f rispetto a due nuove basi è legata alla Γ dalla relazione $\Gamma' = M\Gamma N$ dove M e N sono le matrici del cambiamento di base in V ed in W .

Le matrici associate ad una applicazione lineare, rispetto a qualunque scelta di basi, hanno tutte lo stesso rango, di più, nel caso particolare in cui $V \equiv W$ le matrici associate ad una stessa applicazione lineare sono tutte *simili*³.

Si può anche dimostrare che se A è la matrice associata all'applicazione $f : V \mapsto W$ la dimensione dell'immagine di f è il rango di A cioè

$$r(A) = \dim(\text{Im} f)$$

rispetto ad una qualunque coppia di basi.

7.3 Prodotto scalare, norma

Norma di un vettore in $\mathbb{R}^2$ o $\mathbb{R}^3$

Riprendiamo ora lo studio dei vettori da un punto di vista più geometrico. È noto, per esempio dalla Fisica, che spesso è comodo visualizzare un vettore del piano o dello spazio come una “frecciolina”, caratterizzata da una lunghezza, una direzione ed un verso (o orientamento); si vede subito che se il vettore v è spiccato dall'origine, esso è completamente individuato dal suo secondo estremo.

Abbiamo già visto che i punti di una retta orientata possono essere messi in corrispondenza biunivoca con i numeri reali; se ora consideriamo, invece, due rette, per comodità perpendicolari, e fissiamo su ciascuna di esse un'unità di misura (in genere la stessa), ad ogni

³Il concetto di matrici simili verrà introdotto nel paragrafo 8.2 a pagina 78.

punto del piano corrisponde una coppia ordinata di numeri reali, come illustrato nella Figura 7.1. Così facendo abbiamo fissato quello che si chiama un *sistema di riferimento cartesiano*⁴ *ortogonale*. Nella figura 7.1 è raffigurato un sistema di riferimento cartesiano ortogonale nel piano.

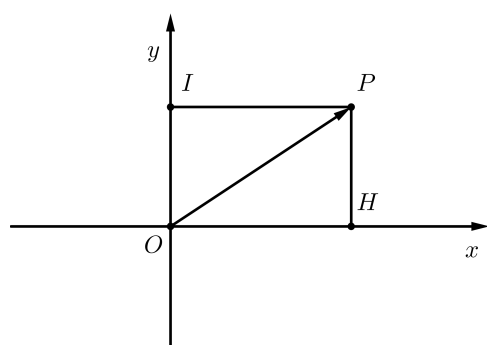


Figura 7.1 Un sistema di riferimento cartesiano ortogonale nel piano

Nello spazio la generalizzazione non è immediata: bisogna considerare come assi tre rette orientate concorrenti⁵ e a due a due perpendicolari e chiamare *coordinate cartesiane* del punto P le tre distanze di P dai tre piani che queste rette a due a due formano; se le tre rette si chiamano rispettivamente x , y e z si ha, per esempio: $x_P = \text{distanza}(P, yz)$, dove, con yz abbiamo indicato il piano individuato dai due assi y e z .

Il punto P può anche essere visto come il secondo estremo di un vettore spiccato dall'origine, che, come già detto, è completamente individuato dalle sue coordinate; in $\mathbb{R}^2$ scriviamo quindi $\overrightarrow{OP} = \vec{v} = [x, y]$, in cui x, y sono le coordinate cartesiane del punto $P(x, y)$ nel sistema di riferimento scelto, mettendo così in luce che si tratta di un vettore del piano, cioè di $\mathbb{R}^2$.

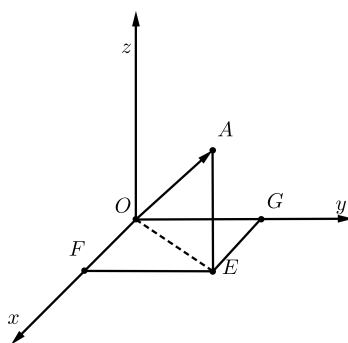


Figura 7.2 Un sistema di riferimento cartesiano ortogonale nello spazio

In modo analogo parliamo di vettori nello spazio come di vettori di $\mathbb{R}^3$: $\overrightarrow{OP} = \vec{v} = [x, y, z]$ in cui le coordinate cartesiane del punto

⁴da Cartesio: René DESCARTES, 1569, La Haye (Francia) – 1650, Stoccolma (Svezia).

⁵cioè passanti tutte e tre per un medesimo punto.

$P(x, y, z)$ sono x, y, z . Vedi la figura 7.2 nella pagina precedente dove è rappresentato un sistema di riferimento cartesiano nello spazio.

Le operazioni tra vettori che abbiamo imparato a conoscere nei paragrafi precedenti si visualizzano tra i vettori geometrici. La somma di due vettori si definisce con la *regola del parallelogrammo*: vedi la figura 7.3 in cui è $\vec{w} = \vec{v} + \vec{u}$.

Per il prodotto di un vettore per uno scalare λ prendiamo in considerazione i tre casi seguenti:

$\lambda > 0$ allora $\lambda \vec{OP}$ è il vettore $\vec{OP}$ con la lunghezza moltiplicata per λ ;

$\lambda = 0$ allora $\lambda \vec{OP}$ è il vettore nullo;

$\lambda < 0$ allora $\lambda \vec{OP}$ è il vettore $\vec{OP}$ con la lunghezza moltiplicata per $|\lambda|$ e di verso opposto.

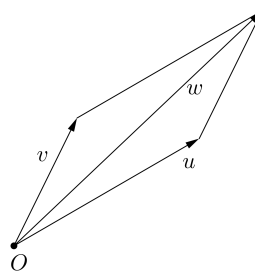


Figura 7.3 La regola del parallelogrammo per la somma di vettori

Queste operazioni corrispondono alle operazioni già viste

$$\lambda[x, y] = [\lambda x, \lambda y]$$

e

$$[x, y] + [x', y'] = [x + x', y + y']$$

in $\mathbb{R}^2$ e

$$\lambda[x, y, z] = [\lambda x, \lambda y, \lambda z]$$

e

$$[x, y, z] + [x', y', z'] = [x + x', y + y', z + z']$$

in $\mathbb{R}^3$.

Anche i concetti di dipendenza ed indipendenza lineare hanno una facile interpretazione geometrica, infatti si vede subito che *in $\mathbb{R}^2$ due vettori sono linearmente dipendenti se e solo se stanno su una stessa retta ed in $\mathbb{R}^3$ tre vettori sono dipendenti se e solo se sono complanari.*

Se identifichiamo gli elementi dello spazio vettoriale $\mathbb{R}^2$ con i segmenti orientati spiccati dall'origine, nel piano riferito ad un sistema di

coordinate cartesiane ortogonali, si nota che ad ogni vettore di $\mathbb{R}^2$ risulta associato un numero reale non negativo: la lunghezza del segmento OP . Definiamo allora la *norma* di un vettore $\vec{v} = [x, y] \in \mathbb{R}^2$ come

$$\|\vec{v}\| = \|[x, y]\| = \sqrt{x^2 + y^2} \quad (7.1)$$

e in $\mathbb{R}^3$

$$\|\vec{v}\| = \|[x, y, z]\| = \sqrt{x^2 + y^2 + z^2}. \quad (7.2)$$

La norma verifica le proprietà elencate nella tabella 7.1, dove, per noi,

Tabella 7.1 *Proprietà della norma di un vettore*

i)	$\ \vec{v}\ \geq 0$	$\forall \vec{v} \in V;$
ii)	$\ \vec{v}\ = 0 \iff \vec{v} = \mathbf{0}$	$\forall \vec{v} \in V;$
iii)	$\ \lambda \vec{v}\ = \lambda \cdot \ \vec{v}\ $	$\forall \vec{v} \in V \text{ e } \forall \lambda \in \mathbb{R};$
iv)	$\ \vec{u} + \vec{v}\ \leq \ \vec{u}\ + \ \vec{v}\ ,$	$\forall \vec{u}, \vec{v} \in V.$

$V \equiv \mathbb{R}^2$ oppure $V \equiv \mathbb{R}^3$ (in realtà si può dare una definizione di norma e di prodotto scalare in uno spazio vettoriale qualsiasi, e non solo in $\mathbb{R}^n$ come vedremo più avanti).

Le prime tre proprietà sono banali; la quarta è la disuguaglianza triangolare che abbiamo già visto.

Dalle proprietà della norma appare chiaro come si può definire una distanza in $\mathbb{R}^2$ o in $\mathbb{R}^3$. Infatti se poniamo

$$d(u, v) = \|u - v\| \quad (7.3)$$

si verifica immediatamente che valgono per ogni u, v, w le proprietà elencate nella tabella 7.2 (caratterizzanti una distanza).

Tabella 7.2 *Proprietà della distanza di due punti*

i)	$d(u, v) \geq 0;$	$\forall u, v$
ii)	$d(u, v) = 0 \iff v = u;$	$\forall v = u$
iii)	$d(u, v) = d(v, u);$	$\forall u, v$
iv)	$d(u, v) \leq d(u, w) + d(w, v).$	$\forall v, u, w$

Osserviamo esplicitamente che se si identificano $\mathbb{R}^3$ e $\mathbb{R}^2$ con lo spazio ed il piano riferiti a coordinate cartesiane ortogonali, la distanza definita dalla (7.3) coincide con quella usualmente definita.

7.4 Prodotto scalare

Definiamo *prodotto scalare*⁶ di due vettori $\vec{v} = [x_1, y_1]$ e $\vec{w} = [x_2, y_2]$ di $\mathbb{R}^2$ il numero reale

$$\langle \vec{v}, \vec{w} \rangle = \langle [x_1, y_1], [x_2, y_2] \rangle = x_1x_2 + y_1y_2 = \begin{bmatrix} x_1 & y_1 \end{bmatrix} \begin{bmatrix} x_2 \\ y_2 \end{bmatrix} = \vec{v} \vec{w}_T \quad (7.4)$$

e in $\mathbb{R}^3$ il numero reale

$$\begin{aligned} \langle \vec{v}, \vec{w} \rangle &= \langle [x_1, y_1, z_1], [x_2, y_2, z_2] \rangle = \\ &= x_1x_2 + y_1y_2 + z_1z_2 = \begin{bmatrix} x_1 & y_1 & z_1 \end{bmatrix} \begin{bmatrix} x_2 \\ y_2 \\ z_2 \end{bmatrix} = \\ &= \vec{v} \vec{w}_T \quad (7.5) \end{aligned}$$

Occorre precisare che, formalmente, il prodotto scalare come è stato definito dalle (7.4) e (7.5) equivale al prodotto $\vec{v} \vec{w}_T$, pensando $\vec{v}$ come matrice costituita da una sola riga e $\vec{w}_T$ come matrice di una sola colonna. In realtà questa equivalenza è solo formale, in quanto, in quest'ultimo caso, otteniamo una matrice composta da un solo elemento: uno scalare⁷ (v. nota 4 a pagina 20), che, in virtù dell'isomorfismo tra $\mathcal{M}_1$ ed $\mathbb{R}$, possiamo ritenere equivalenti.

Per esempio

$$\langle [4, 3, -1], [0, -3, 4] \rangle = 4 \cdot 0 + 3(-3) + (-1)4 = -13$$

oppure

$$\left\langle \begin{bmatrix} 2 \\ 3 \end{bmatrix}, \begin{bmatrix} -3 \\ 1 \end{bmatrix} \right\rangle = 2(-3) + 3 \cdot 1 = -3$$

(nel secondo esempio abbiamo usato vettori colonna, per sottolineare l'assoluta intercambiabilità, in questo contesto, delle due notazioni).

Il prodotto scalare è legato alla norma dalla relazione

$$\sqrt{\langle \vec{u}, \vec{u} \rangle} = \|\vec{u}\|. \quad (7.6)$$

Inoltre valgono le proprietà elencate nella tabella 7.3 a fronte, semplicissime

⁶Da non confondere con il prodotto *per uno scalare*: si noti che il prodotto scalare $\langle \cdot, \cdot \rangle$ associa ad una coppia ordinata di vettori un numero reale, mentre il prodotto per uno scalare associa ad una coppia scalare-vettore un vettore.

⁷nel nostro caso un numero reale.

Tabella 7.3 Proprietà del prodotto scalare

i)	$\langle \vec{u}, \vec{u} \rangle \geq 0$	$\forall \vec{u} \in V$
ii)	$\langle \vec{u}, \vec{u} \rangle = 0 \iff \vec{u} = \mathbf{0}$	$\forall \vec{u} \in V$
iii)	$\langle \vec{u}, \vec{v} \rangle = \langle \vec{v}, \vec{u} \rangle$	$\forall \vec{u}, \vec{v} \in V$
iv)	$\langle \lambda \vec{v}, \vec{u} \rangle = \lambda \langle \vec{v}, \vec{u} \rangle = \langle \vec{u}, \lambda \vec{v} \rangle$	$\forall \vec{u}, \vec{v} \in V \text{ e } \forall \lambda \in \mathbb{R}$
v)	$\langle \vec{u}, \vec{v} + \vec{w} \rangle = \langle \vec{u}, \vec{v} \rangle + \langle \vec{u}, \vec{w} \rangle$	$\forall \vec{u}, \vec{v}, \vec{w} \in V$

da dimostrare tenendo conto della definizione, e la cui dimostrazione è proposta come esercizio.

Sussiste il seguente

Teorema 7.5. Siano $\vec{u}$ e $\vec{v}$ due vettori di $\mathbb{R}^2$ o di $\mathbb{R}^3$ e si considerino i segmenti orientati ad essi associati nel piano o nello spazio riferiti a sistemi di coordinate ortogonali. Allora

$$\langle \vec{u}, \vec{v} \rangle = \|\vec{u}\| \cdot \|\vec{v}\| \cos \varphi \quad (7.7)$$

dove $\varphi \in [0, \pi]$ è l'ampiezza dell'angolo fra i due segmenti.

Dimostrazione. La dimostrazione, che invitiamo il lettore a sviluppare per esteso, è un'immediata conseguenza del Teorema del coseno e delle proprietà del prodotto scalare. $\square$

Esempio 7.4. Siano dati in $\mathbb{R}^2$ i due vettori $\vec{u} = [1, -3]$ e $\vec{v} = [2, 5]$, vogliamo conoscere l'angolo φ che formano i segmenti orientati ad essi associati; il loro prodotto scalare è $\langle \vec{u}, \vec{v} \rangle = 1 \cdot 2 + (-3)5 = -13$. Inoltre $\|\vec{u}\| = \sqrt{10}$ e $\|\vec{v}\| = \sqrt{29}$ dunque l'angolo fra i due segmenti orientati sarà tale che $\cos \varphi = \frac{\langle \vec{u}, \vec{v} \rangle}{\|\vec{u}\| \cdot \|\vec{v}\|} = -\frac{13}{\sqrt{290}}$; osservando che $-\frac{\sqrt{3}}{2} < -\frac{13}{\sqrt{290}} < -\frac{\sqrt{2}}{2}$ si deduce che $\frac{3\pi}{4} < \varphi < \frac{5\pi}{6}$.

Un vettore $\vec{u} = [a, b, c]$ si dice *unitario* o *versore* se ha norma 1, cioè se $\|\vec{u}\| = 1$, quindi se $a^2 + b^2 + c^2 = 1$.

Dal Teorema 7.5 si ricava che le componenti a , b e c di u sono i coseni degli angoli che u forma con i versori fondamentali $e_1 = [1, 0, 0]$, $e_2 = [0, 1, 0]$ ed $e_3 = [0, 0, 1]$ e prendono il nome di *coseni direttori del vettore* $\vec{u}$.

Ogni vettore non nullo può essere normalizzato dividendolo per la propria norma, infatti è facile verificare che $\left\| \frac{\vec{v}}{\|\vec{v}\|} \right\| = \frac{\|\vec{v}\|}{\|\vec{v}\|} = 1$.

Due vettori diversi dal vettore nullo si dicono *perpendicolari* o *ortogonali* se l'ampiezza dell'angolo tra u e v è $\varphi = \frac{\pi}{2}$.

Segue immediatamente dal Teorema 7.5 che

$$\langle \vec{u}, \vec{v} \rangle = 0 \iff \varphi = \frac{\pi}{2} \text{ con } \vec{u} \neq \mathbf{0}, \vec{v} \neq \mathbf{0} \quad (7.8)$$

scriveremo dunque $\vec{v} \perp \vec{u} \iff \langle \vec{u}, \vec{v} \rangle = 0$ cioè due vettori sono ortogonali se e solo se il loro prodotto scalare è nullo.

Il discorso si generalizza:

DEFINIZIONE 7.2. Si dice che n vettori $\vec{v}_1, \vec{v}_2, \dots, \vec{v}_n$ sono *mutuamente ortogonali* se

$$\langle \vec{v}_i, \vec{v}_j \rangle = 0 \quad (7.9)$$

per ogni i, j con $i, j = 1 \dots n$.

Se i $\vec{v}_i$ sono anche normalizzati (cioè sono dei versori) diciamo che sono *ortonormali*.

Una *base ortogonale* è una base costituita da vettori mutuamente ortogonali; se i vettori sono anche normalizzati, allora abbiamo a che fare con una *base ortonormale*.

Osserviamo che *vettori ortogonali sono sempre indipendenti*, mentre non vale in generale il viceversa. (Verificarlo per esercizio)

Esempio 7.5. In $\mathbb{R}^3$ la base $\mathcal{B} = \{[1, 3, 1], [-1, 0, 1], [6, -4, 6]\}$ è una base ortogonale ma non ortonormale (verificarlo per esercizio).

Sappiamo che in uno spazio vettoriale ogni vettore si può esprimere come combinazione lineare dei vettori di una base; se la base è ortogonale o ortonormale, si possono determinare in maniera semplice i coefficienti della combinazione lineare:

Teorema 7.6. Sia $\{\vec{e}_1, \vec{e}_2, \vec{e}_3\}$ una base ortogonale, allora $\vec{v} = \lambda_1 \vec{e}_1 + \lambda_2 \vec{e}_2 + \lambda_3 \vec{e}_3$ in cui

$$\lambda_i = \frac{\langle \vec{v}, \vec{e}_i \rangle}{\langle \vec{e}_i, \vec{e}_i \rangle}; \quad (7.10)$$

se invece la base è ortonormale

$$\lambda_i = \langle \vec{v}, \vec{e}_i \rangle.$$

Esempio 7.6. Sia $\mathcal{B}$ la base ortogonale

$$\mathcal{B} = \{\vec{e}_1 = [1, 3, 1], \vec{e}_2 = [-1, 0, 1], \vec{e}_3 = [6, -4, 6]\}$$

dell'esempio 7.5: abbiamo $\langle \vec{e}_1, \vec{e}_1 \rangle = 11$, $\langle \vec{e}_2, \vec{e}_2 \rangle = 2$, $\langle \vec{e}_3, \vec{e}_3 \rangle = 88$. Se $\vec{v} = [3, 2, 5]$ usando i risultati trovati e la (7.10) si ha

$$\begin{aligned} \vec{v} &= \frac{\langle \vec{v}, \vec{e}_1 \rangle}{11} \vec{e}_1 + \frac{\langle \vec{v}, \vec{e}_2 \rangle}{2} \vec{e}_2 + \frac{\langle \vec{v}, \vec{e}_3 \rangle}{88} \vec{e}_3 = \\ &= \frac{14}{11} \vec{e}_1 + \vec{e}_2 + \frac{5}{11} \vec{e}_3. \end{aligned}$$

Dal punto di vista geometrico la (7.10) significa che $\langle \vec{v}, \vec{e}_i \rangle \vec{e}_i$ è la componente del vettore $\vec{v}$ nella direzione di $\vec{e}_i$ o anche che è la proiezione ortogonale di $\vec{v}$ sulla retta su cui giace il vettore $\vec{e}_i$. Come è noto la proiezione ortogonale ha lunghezza $\|\vec{v}\| \cdot |\cos \varphi_i|$; lo stesso risultato si trova applicando il teorema 7.5:

$$\|\langle \vec{v}, \vec{e}_i \rangle \vec{e}_i\| = |\langle \vec{v}, \vec{e}_i \rangle| \cdot \|\vec{e}_i\| = \|\vec{v}\| \cdot \|\vec{e}_i\|^2 \cdot |\cos \varphi_i| = \|\vec{v}\| \cdot |\cos \varphi_i|.$$

7.5 Generalizzazioni

Il concetto di prodotto scalare è molto più generale di quello qui definito (che è il prodotto scalare standard in uno spazio vettoriale isomorfo a $\mathbb{R}^2$ od a $\mathbb{R}^3$).

Ricordiamo che si chiama *prodotto cartesiano* di due insiemi V e W e si indica con $V \times W$ l'insieme delle coppie ordinate (v, w) essendo $v \in V$, e $w \in W$.

In generale dati due spazi vettoriali sul medesimo campo $\mathbb{K}$ un'applicazione g da $V \times W$ a $\mathbb{K}$ per cui valgano le proprietà elencate nella tabella 7.4 si chiama applicazione o *forma bilineare* (nel senso che è lineare

Tabella 7.4 Proprietà delle forme bilineari

i)	$g(v_1 + v_2, w) = g(v_1, w) + g(v_2, w)$	$\forall v_1, v_2 \in V, w \in W$
ii)	$g(v, w_1 + w_2) = g(v, w_1) + g(v, w_2)$	$\forall v \in V, w_1, w_2 \in W$
iii)	$g(\alpha v, w) = g(v, \alpha w) = \alpha g(v, w)$	$\forall v \in V, w \in W, \alpha \in \mathbb{K}$

re rispetto a tutt'e due le variabili; nello stesso senso si parla talvolta anche di forma *multilineare*).

Un'applicazione bilineare tale che si abbia

$$g(v, w) = g(w, v) \quad \forall v \in V, w \in W$$

si chiama *simmetrica*. Un' applicazione bilineare simmetrica $g : V \times V \rightarrow \mathbb{R}$ per cui sia

$$i) \quad g(v, w) \geq 0 \quad \forall v, w \in V$$

$$ii) \quad g(v, v) = 0 \iff v = \mathbf{0}$$

si chiama *prodotto scalare* e si preferisce indicare $g(v, w)$ con $\langle \vec{v}, \vec{w} \rangle$. Uno spazio vettoriale in cui sia stato definito un prodotto scalare si chiama *euclideo*.

Come utile esercizio, il lettore verifichi che il prodotto scalare standard definito nel paragrafo precedente è una forma bilineare simmetrica che gode delle proprietà *i)* e *ii)*

Come ulteriore esempio si verifichi che in $\mathbb{R}^2$ è un prodotto scalare

$$\langle \vec{x}, \vec{y} \rangle = [x_1, x_2] \cdot \begin{bmatrix} 2 & 1 \\ 1 & 5 \end{bmatrix} \cdot \begin{bmatrix} y_1 \\ y_2 \end{bmatrix}.$$

Ovviamente due vettori ortogonali rispetto ad un prodotto scalare possono non esserlo rispetto ad un altro. Quando parleremo di vettori ortogonali senza precisare rispetto a quale prodotto scalare ci riferiremo sempre al prodotto scalare standard.

Per i prodotti scalari vale il

Teorema 7.7. *Sia V uno spazio vettoriale euclideo, cioè uno spazio vettoriale dotato di un prodotto scalare, allora si ha:*

$$\langle \vec{u}, \vec{v} \rangle = \frac{1}{2} (\langle \vec{u} + \vec{v}, \vec{u} + \vec{v} \rangle - \langle \vec{u}, \vec{u} \rangle - \langle \vec{v}, \vec{v} \rangle)$$

Dimostrazione. la dimostrazione, che il lettore è invitato a scrivere in maniera esplicita, è un semplice calcolo basato sulla bilinearità e sulla simmetria del prodotto scalare. $\square$

Sia V uno spazio vettoriale euclideo e sia U un suo sottospazio. Indichiamo con $U^\perp$ l'insieme di tutti i vettori di V che sono ortogonali a vettori di U (rispetto ad un fissato prodotto scalare) e lo chiamiamo *complemento ortogonale* di U (rispetto a quel certo prodotto scalare).

Capitolo 8

Matrici simili. Autovalori ed autovettori di una matrice quadrata

8.1 Matrici simili

Siano A e B due matrici quadrate di ordine n

DEFINIZIONE 8.1. Diciamo che la matrice A è *simile* alla matrice B se esiste una matrice di passaggio P , non singolare, tale che

$$P^{-1}AP = B. \quad (8.1)$$

Si vede subito che

Teorema 8.1. *La similitudine di matrici è una relazione di equivalenza.*

Dimostrazione. Infatti ogni matrice è simile a se stessa (basta prendere, nella (8.1) come matrice di passaggio la matrice I); ricordando che $P = (P^{-1})^{-1}$ si ricava subito che se A è simile a B con matrice di passaggio P allora B sarà simile ad A con matrice di passaggio P^{-1} ; infine da $P^{-1}AP = B$ e $Q^{-1}BQ = C$ si ha $Q^{-1}P^{-1}APQ = C$ dunque, ricordando che $(PQ)^{-1} = Q^{-1}P^{-1}$ si conclude che A è simile a C con matrice di passaggio PQ . $\square$

È immediato dimostrare il

Teorema 8.2. *Due matrici simili hanno lo stesso determinante.*

Dimostrazione. Siano A e B le due matrici. Dalla formula 8.1 e dal teorema di Binet 5.5 a pagina 50 si ricava che

$$\begin{aligned}\det(P^{-1}AP) &= \det B \\ \det(P^{-1}) \cdot \det A \cdot \det P &= \det B\end{aligned}$$

e la tesi segue immediatamente dal fatto che $\det(P^{-1}) \det P = 1$. $\square$

OSSERVAZIONE 8.1. Attenzione! L'inverso del Teorema 8.2 non vale: cioè esistono matrici che hanno lo stesso determinante e *non* sono simili. È un utile esercizio trovare degli esempi di questo fatto.

8.2 Autovalori ed autovettori di una matrice

Sia ora A una matrice quadrata di ordine n , per esempio associata ad un endomorfismo¹, e consideriamo la relazione

$$A\vec{x} = \lambda\vec{x} \quad (8.2)$$

con $\vec{x} \neq \mathbf{0}$. La (8.2) in sostanza dice che applicando al vettore $\vec{x}$ la matrice A si ottiene un vettore proporzionale ad $\vec{x}$, cioè che $\vec{x}$ *non cambia direzione*.

Fissata la A ci chiediamo se esistono degli scalari λ e dei vettori $\vec{x}$ per cui valga la (8.2), cioè, fissata la trasformazione, ci chiediamo se ci sono vettori che, trasformati, non cambiano direzione.

Gli scalari λ che compaiono nella (8.2) si chiamano *autovalori* o *valori propri* ed i vettori $\vec{x}$ si chiamano *autovettori* o *vettori propri* della matrice A .

È lecito ora porsi la domanda: fissata una matrice A esistono autovalori? quanti? ed autovettori?

Per rispondere a queste domande riscriviamo la (8.2) nella forma

$$\lambda\vec{x} - A\vec{x} = \mathbf{0}$$

equivalente a

$$(\lambda I - A)\vec{x} = \mathbf{0} \quad (8.3)$$

che possiamo pensare come un *sistema lineare omogeneo* di n equazioni in n incognite la cui matrice dei coefficienti è la matrice $\lambda I - A$. Sappiamo che un tale sistema ammette soluzioni non banali se e solo se il determinante della matrice dei coefficienti è nullo.

¹Ricordiamo che si chiama endomorfismo un'applicazione lineare $f : V \rightarrow V$ di uno spazio vettoriale su se stesso.

Il determinante della matrice dei coefficienti del sistema (8.2) è ovviamente funzione di λ

$$\det(\lambda I - A) = \begin{vmatrix} \lambda - a_{11} & -a_{12} & -a_{13} & \dots & -a_{1n} \\ -a_{21} & \lambda - a_{22} & -a_{23} & \dots & -a_{2n} \\ -a_{31} & -a_{32} & \lambda - a_{33} & \dots & -a_{3n} \\ \vdots & \vdots & \vdots & \ddots & \vdots \\ -a_{n1} & -a_{n2} & -a_{n3} & \dots & \lambda - a_{nn} \end{vmatrix}$$

e si può scrivere come polinomio in λ

$$\varphi(\lambda) = \det(\lambda I - A) = \lambda^n + c_1 \lambda^{n-1} + c_2 \lambda^{n-2} + \dots + c_{n-1} \lambda + c_n \quad (8.4)$$

che è un polinomio di grado n nella variabile λ con coefficiente direttore² uguale a 1, che prende il nome di *polinomio caratteristico* della matrice A ed i cui zeri sono tutti e soli gli autovalori di A .

Il coefficiente di λ^n proviene solo dal prodotto degli elementi principali ed è quindi uguale a 1; i vari successivi coefficienti $c_1 \dots c_n$ di $\varphi(\lambda)$ si possono trovare sviluppando normalmente il determinante di $\lambda I - A$ oppure tenendo conto del legame tra di essi ed i minori estratti dalla matrice A (Teorema 8.7 a pagina 81).

Per il Teorema fondamentale dell'Algebra sappiamo che ogni polinomio di grado n ha, nel campo complesso, esattamente n radici contate ciascuna con la propria molteplicità, dunque sussiste il

Teorema 8.3. *Una matrice quadrata di ordine n ha, nel campo complesso $\mathbb{C}$, esattamente n autovalori, ciascuno contato con la propria molteplicità.*

Una banale ed immediata conseguenza è che se anche la matrice A ha tutti gli elementi reali, i suoi autovalori possono essere in tutto o in parte numeri complessi.

Se $\lambda_1, \lambda_2, \dots, \lambda_s$ ($s \leq n$) sono gli autovalori distinti di A e $k_1, k_2, \dots, k_s$ rispettivamente le loro molteplicità algebriche (quindi $\sum_{j=1}^s k_j = n$), chiamiamo l'espressione

$$\begin{pmatrix} \lambda_1 & \lambda_2 & \dots & \lambda_s \\ k_1 & k_2 & \dots & k_s \end{pmatrix}$$

lo *spettro* della matrice A . In questa scrittura sotto ad ogni autovalore è indicata la sua molteplicità come radice del polinomio caratteristico detta anche *molteplicità algebrica*.

Vale il

²cioè coefficiente del termine di grado massimo.

Teorema 8.4. *Due matrici simili A e B hanno lo stesso polinomio caratteristico.*

Dimostrazione. Sia $B = P^{-1}AP$: si ha successivamente

$$\begin{aligned} |\lambda I - B| &= |\lambda I - P^{-1}AP| = |P^{-1}\lambda IP - P^{-1}AP| = \\ &= |P^{-1}(\lambda I - A)P| = |P^{-1}||\lambda I - A||P| \\ &= |\lambda I - A| \end{aligned}$$

e quindi i polinomi caratteristici sono uguali. $\square$

OSSERVAZIONE 8.2. Dal Teorema 8.4 segue che due matrici simili hanno gli stessi autovalori con le stesse molteplicità, quindi lo stesso spettro.

OSSERVAZIONE 8.3. ATTENZIONE Il Teorema 8.4 non è invertibile!, cioè due matrici che hanno lo stesso polinomio caratteristico possono anche non essere simili, per esempio le matrici $A = \begin{bmatrix} 0 & 0 \\ 0 & 0 \end{bmatrix}$ e $B = \begin{bmatrix} 0 & 1 \\ 0 & 0 \end{bmatrix}$ hanno entrambe polinomio caratteristico $\varphi(\lambda) = \lambda^2$ ma la A , essendo la matrice nulla, è simile solo a se stessa.

Sugli autovalori valgono le proprietà espresse dai seguenti teoremi

Teorema 8.5. *Se $\lambda_1, \lambda_2, \dots, \lambda_n$ sono gli n autovalori di A e se c_n è il termine noto del polinomio caratteristico di A sussiste la relazione*

$$\det(A) = (-1)^n c_n = \lambda_1 \cdot \lambda_2 \cdots \lambda_n.$$

Dimostrazione. Infatti si ha, tenendo conto della (8.4):

$$\begin{aligned} \varphi(\lambda) = \det(\lambda I - A) &= \lambda^n + c_1 \lambda^{n-1} + \cdots + c_n = \\ &= (\lambda - \lambda_1)(\lambda - \lambda_2) \cdots (\lambda - \lambda_n) \end{aligned}$$

relazione che vale $\forall \lambda$ e quindi, in particolare, anche per $\lambda = 0$. Per questo valore si ha

$$\varphi(0) = \det(-A) = c_n = (-\lambda_1)(-\lambda_2) \cdots (-\lambda_n)$$

cioè $(-1)^n \det(A) = c_n = (-1)^n \lambda_1 \lambda_2 \cdots \lambda_n$ che è la tesi. $\square$

Ne segue il

Corollario 8.6. *Una matrice è singolare se e solo se ha almeno un autovalore nullo.*

OSSERVAZIONE 8.4. Da quanto detto si ricava anche che in una matrice triangolare (in particolare diagonale) gli autovalori coincidono con gli elementi della diagonale principale.

DEFINIZIONE 8.2. Sia A una matrice quadrata; chiamiamo *minore principale di ordine k* e lo indichiamo con M_k il determinante di una sottomatrice quadrata di ordine k i cui elementi principali sono solo elementi principali di A .

Esempio 8.1. I minori principali di ordine 2 della matrice $\begin{bmatrix} a & b & c \\ 1 & 2 & 3 \\ d & e & f \end{bmatrix}$ sono

$\begin{vmatrix} a & c \\ d & f \end{vmatrix}, \begin{vmatrix} a & b \\ 1 & 2 \end{vmatrix}, \begin{vmatrix} 2 & 3 \\ e & f \end{vmatrix}$ ma non, ad esempio $\begin{vmatrix} 1 & 2 \\ d & e \end{vmatrix}$ perché i suoi elementi principali non sono tutti elementi principali di A

Si dimostra allora che

Teorema 8.7. Se

$$\varphi(\lambda) = \lambda^n + c_1\lambda^{n-1} + \cdots + c_{n-1}\lambda + c_n$$

è il polinomio caratteristico di una matrice A , allora

$$c_k = (-1)^k \sum M_k \quad (8.5)$$

dove la somma è estesa a tutti i minori principali di ordine k estratti da A .

Quindi c_i è –a meno del segno– la somma dei minori principali di ordine i , da cui, per esempio, $c_1 = -\text{tr}(A)$; c_2 è la somma di tutti i minori principali di ordine 2 e così via.

Inoltre sussiste il

Teorema 8.8. Se λ è un autovalore di molteplicità k allora

$$r(\lambda I - A) \geq n - k. \quad (8.6)$$

Il numero $n - r(\lambda I - A)$ è il numero delle soluzioni indipendenti del sistema (8.2), cioè degli autovettori indipendenti associati all'autovalore λ , che prende il nome di *molteplicità geometrica* dell'autovalore. Un autovalore per cui la molteplicità algebrica sia uguale a quella geometrica, cioè per il quale vale il segno = nella (8.6), si chiama *regolare*. Segue subito da questa definizione e dal teorema 8.8 il

Teorema 8.9. Ogni autovalore semplice è regolare.

Dimostrazione. Infatti $r(\lambda I - A) < n$ in quanto $\det(\lambda I - A) = 0$ e dal Teorema 8.8 segue che è $r \geq n - 1$ dunque $n - 1 \leq r < n$ da cui $r = n - 1$. $\square$

Siano ora $\vec{x}$ e $\vec{y}$ due autovettori della matrice A associati entrambi all'autovalore λ ; sussiste il

Teorema 8.10. *Ogni combinazione lineare di autovettori di A associati a λ è un autovettore di A associato a λ .*

Dimostrazione. Siano $A\vec{x} = \lambda\vec{x}$ e $A\vec{y} = \lambda\vec{y}$, consideriamo il vettore $\alpha\vec{x} + \beta\vec{y}$ e vogliamo dimostrare che anch'esso è autovettore associato a λ . Si ha, infatti

$$\begin{aligned} A(\alpha\vec{x} + \beta\vec{y}) &= A\alpha\vec{x} + A\beta\vec{y} = \\ \alpha A\vec{x} + \beta A\vec{y} &= \alpha\lambda\vec{x} + \beta\lambda\vec{y} = \lambda(\alpha\vec{x} + \beta\vec{y}). \end{aligned}$$

$\square$

Quindi l'insieme degli autovettori associati ad un autovalore, con l'aggiunta del vettore nullo, costituisce uno spazio vettoriale, che prende il nome di *autospazio* associato a λ e la molteplicità geometrica dell'autovalore, che corrisponde al numero degli autovettori indipendenti, è la dimensione di questo autospazio.

Per il Teorema 8.8 si ha subito che *la molteplicità geometrica di un autovalore λ non supera quella algebrica e la uguaglia se e solo se λ è regolare.*

Vale il

Teorema 8.11. *Siano $\lambda_1, \lambda_2, \dots, \lambda_s$ s autovalori distinti di A ($s \leq n$) e siano $\vec{x}_1, \vec{x}_2, \dots, \vec{x}_s$ s autovettori associati ordinatamente ai λ_i . Allora i vettori $\vec{x}_i$ sono linearmente indipendenti.*

Di conseguenza gli autospazi associati ad ogni autovalore sono disgiunti, e la loro somma è un sottospazio V' dello spazio vettoriale V in cui è definito l'endomorfismo rappresentato dalla matrice A .

Da quanto detto segue facilmente che gli autovalori di A sono tutti regolari se e solo se $V' = V$.

Sussiste anche il

Teorema 8.12. *Se $\vec{x}$ è autovettore di A associato all'autovalore λ allora esso è anche autovettore di A^k associato all'autovalore $\lambda^k \forall k > 0$.*

Dimostrazione. Da $A\vec{x} = \lambda\vec{x}$ si ricava, moltiplicando a sinistra per A ,

$$A^2\vec{x} = A\lambda\vec{x} = \lambda A\vec{x} = \lambda\lambda\vec{x} = \lambda^2\vec{x}$$

e quindi x è autovettore di A associato all'autovalore λ^2 . Iterando il procedimento si perviene alla tesi. $\square$

Capitolo 9

Diagonalizzazione, matrici ortogonali

9.1 Diagonalizzazione di una matrice quadrata

Una matrice quadrata si dice *diagonalizzabile* se è simile ad una matrice diagonale, cioè se esiste una matrice non singolare P tale che

$$P^{-1}AP = \Delta$$

con Δ matrice diagonale.

Ci proponiamo ora di stabilire dei criteri di diagonalizzabilità, il che equivale a dare delle condizioni *necessarie e sufficienti* per stabilire quali sono tutte e sole le matrici diagonalizzabili.

Sussiste a questo proposito il

Teorema 9.1. *Una matrice quadrata A di ordine n è diagonalizzabile se e solo se ammette n autovettori indipendenti.*

Dimostrazione. Siano $\vec{X}_1, \vec{X}_2, \dots, \vec{X}_n$ n autovettori indipendenti di A e siano associati rispettivamente agli autovalori $\lambda_1, \lambda_2, \dots, \lambda_n$. Indichiamo con $P = [\vec{X}_1 \ \vec{X}_2 \ \dots \ \vec{X}_n]$ la matrice che ha come colonne gli $\vec{X}_i$; allora, ricordando i punti *i*) e *ii*) dell'osservazione 3.7 a pagina 26, sussistono le uguaglianze

$$\begin{aligned} AP &= \begin{bmatrix} A\vec{X}_1 & A\vec{X}_2 & \dots & A\vec{X}_n \end{bmatrix} \\ P \cdot \text{diag}(\lambda_1, \lambda_2, \dots, \lambda_n) &= \begin{bmatrix} \lambda_1 \vec{X}_1 & \lambda_2 \vec{X}_2 & \dots & \lambda_n \vec{X}_n \end{bmatrix} \end{aligned} \quad (9.1)$$

poiché gli $\vec{X}_i$ sono autovettori di A associati ordinatamente agli autovalori λ_i , per ogni i si ha $A\vec{X}_i = \lambda_i \vec{X}_i$, dalla (9.1) segue che

$$AP = P \cdot \text{diag}(\lambda_1, \lambda_2, \dots, \lambda_n)$$

ed essendo P non singolare, in quanto formata da vettori indipendenti segue che

$$P^{-1}AP = \text{diag}(\lambda_1, \lambda_2, \dots, \lambda_n) \quad (9.2)$$

quindi se A ammette n autovettori indipendenti, essa è diagonalizzabile.

Viceversa se supponiamo che A sia diagonalizzabile allora esistono una matrice invertibile P ed una matrice diagonale $\Delta = \text{diag}(\lambda_1, \lambda_2, \dots, \lambda_n)$ tali che

$$AP = P\Delta$$

ma se chiamiamo $\vec{X}_1, \vec{X}_2, \dots, \vec{X}_n$ le colonne di P , dalla (9.2) e ricordando ancora l'osservazione 3.7 a pagina 26, otteniamo

$$A \begin{bmatrix} \vec{X}_1 & \vec{X}_2 & \dots & \vec{X}_n \end{bmatrix} = \begin{bmatrix} \vec{X}_1 & \vec{X}_2 & \dots & \vec{X}_n \end{bmatrix} \cdot \text{diag}(\lambda_1, \lambda_2, \dots, \lambda_n),$$

da cui

$$\begin{bmatrix} A\vec{X}_1 & A\vec{X}_2 & \dots & A\vec{X}_n \end{bmatrix} = \begin{bmatrix} \lambda_1\vec{X}_1 & \lambda_2\vec{X}_2 & \dots & \lambda_n\vec{X}_n \end{bmatrix}$$

e quindi, per ogni i , si ha $A\vec{X}_i = \lambda_i\vec{X}_i$, dunque gli $\vec{X}_i$ sono autovettori di A , linearmente indipendenti in quanto P è non singolare. $\square$

OSSERVAZIONE 9.1. Il Teorema 9.1 significa, in sostanza, che una matrice che rappresenta un endomorfismo di uno spazio vettoriale V è diagonalizzabile se e solo se esiste una base formata da autovettori di V . Cioè se $V = V_1 \oplus V_2 \oplus \dots \oplus V_k$ dove i V_i sono gli autospazi associati, ordinatamente, agli autovalori λ_i .

OSSERVAZIONE 9.2. Le matrici P che trasformano la A in una matrice diagonale sono tutte (e sole) quelle formate da n autovettori indipendenti di A , quindi sono *infinite*.

OSSERVAZIONE 9.3. Se $\text{diag}(\lambda_1, \lambda_2, \dots, \lambda_n)$ è una qualunque matrice diagonale simile ad A , i suoi elementi principali sono tutti (e soli) gli autovalori di A .

OSSERVAZIONE 9.4. Una matrice diagonalizzabile è univocamente determinata dai suoi autovettori e dai suoi autovalori.

Abbiamo detto che l'avere lo stesso polinomio caratteristico non basta affinché due matrici siano simili, tuttavia

Teorema 9.2. Se A e B sono entrambe diagonalizzabili, esse sono simili se e solo se hanno lo stesso polinomio caratteristico.

Dimostrazione. Se A e B sono simili, hanno lo stesso polinomio caratteristico per il Teorema 8.4, viceversa se A e B hanno lo stesso polinomio caratteristico esse hanno gli stessi autovalori con le stesse molteplicità; dunque sono entrambe simili alla medesima matrice diagonale, e quindi simili tra loro. $\square$

Per verificare se due matrici A e B sono simili è necessario, anzitutto, controllare che abbiano lo stesso polinomio caratteristico; a questo punto possono accadere tre casi:

- o sono entrambe diagonalizzabili, e allora per il Teorema 9.2 sono simili tra loro,
- oppure una è diagonalizzabile e l'altra no, e allora non sono simili,
- oppure ancora nessuna delle due è diagonalizzabile, e allora occorre ricorrere a metodi più sofisticati.

Un'altra condizione necessaria e sufficiente per la diagonalizzabilità, meno semplice da dimostrare ma più comoda da usare, è quella espressa dal

Teorema 9.3. *Una matrice è diagonalizzabile se e solo se ha tutti gli autovalori regolari.*

Poiché un autovalore semplice è regolare (vedi Teorema 8.9 a pagina 81) dal Teorema 9.3 segue immediatamente il

Corollario 9.4. *Una matrice è diagonalizzabile se tutti i suoi autovalori sono distinti.*

OSSERVAZIONE 9.5 (ATTENZIONE!). L'inverso del corollario 9.4 in generale non vale, cioè una matrice può essere diagonalizzabile anche se i suoi autovalori non sono tutti distinti, basta infatti che siano tutti regolari (vedi il teorema 9.3).

Esempio 9.1. Sia $A = \begin{bmatrix} a & a+1 \\ a+3 & a+2 \end{bmatrix}$. Vogliamo vedere per quali valori di a essa è diagonalizzabile. Il polinomio caratteristico di A è

$$\varphi(\lambda) = \lambda^2 - (2a+2)\lambda + a(a+2) - (a+3)(a+1).$$

Le radici di $\varphi(\lambda)$ sono -1 e $2a+3$, che coincidono se e solo se $a = -2$, quindi per $a \neq -2$ la matrice è diagonalizzabile perché i due autovalori sono distinti,

quindi semplici entrambi; per $a = -2$ la matrice diventa $\begin{bmatrix} -2 & -1 \\ 1 & 0 \end{bmatrix}$ il cui polinomio caratteristico è $(\lambda + 1)^2$ cioè ammette l'autovalore -1 doppio. Esso è regolare se $r(-I - A) = 0$ il che palesemente non è. Concludiamo dunque che la matrice A è diagonalizzabile per ogni $a \neq -2$.

Esempio 9.2. Vogliamo determinare per quali valori dei parametri è diagonalizzabile la matrice

$$\begin{bmatrix} 1 & 0 & 0 \\ a & 0 & 0 \\ b & a & 1 \end{bmatrix}.$$

Si vede subito (A è triangolare) che gli autovalori sono 0 semplice e 1 doppio. L'autovalore 0 è regolare in quanto semplice. Esaminiamo la regolarità di

1. Esso è regolare quando $r(I - A) = 3 - 2 = 1$; la matrice $1I - A$ è

$$\begin{bmatrix} 0 & 0 & 0 \\ -a & 1 & 0 \\ -b & -a & 0 \end{bmatrix} \text{ il cui unico minore del second'ordine non certamente nullo è}$$

$\begin{bmatrix} -a & 1 \\ -b & -a \end{bmatrix}$ esso però si annulla per $a^2 + b = 0$, dunque per questi valori, e solo per questi, A è diagonalizzabile.

9.2 Matrici ortogonali

DEFINIZIONE 9.1. Diciamo che una matrice U è ortogonale se è reale e se

$$UU_T = U_T U = I. \quad (9.3)$$

Sulle matrici ortogonali sussiste il

Teorema 9.5. Una matrice U è ortogonale se e solo se le sue colonne formano un sistema ortonormale di vettori.

Dimostrazione. Infatti la relazione (9.3) implica che se $U = [u_{ik}]$ e $U_T = [u_{ki}]$ si ha $\delta_{ik} = \sum_j u_{ji} u_{jk}$ dove il generico elemento del prodotto è

l'elemento generico della matrice I cioè $\delta_{ik} = \begin{cases} 1 & \text{se } i = k \\ 0 & \text{se } i \neq k \end{cases}$ ovvero le colonne di U formano un sistema ortonormale. $\square$

Sulle matrici ortogonali vale il

Teorema 9.6. Sia U una matrice ortogonale. Allora:

- i) $\det(U) = \pm 1$;
- ii) U è invertibile e $U^{-1} = U_T$;
- iii) U_T è ortogonale.

Dimostrazione. Essendo $U_T U = I$, dal Teorema (di Binet) 5.5 a pagina 50 si ha $\det(U) \det(U_T) = 1$ ma poiché $\det(U) = \det(U_T)$ si ha $\det^2(U) = 1$ da cui $\det(U) = \pm 1$. Il punto ii) segue dal precedente e dall'unicità della matrice inversa. Il punto iii) dal fatto che $(U_T)_T = U$ $\square$

Se U e V sono due matrici ortogonali, allora la matrice $W = UV$ è ortogonale, infatti $WW_T = UV(UV)_T = UVV_T U_T = I$.

Se A è una matrice diagonalizzabile e tra le matrici che la diagonalizzano esiste una matrice ortogonale, diciamo che A è *ortogonalmente diagonalizzabile*.

Dimostriamo ora il

Teorema 9.7. *Gli autovalori di una matrice reale simmetrica sono reali*

Dimostrazione. Sia A simmetrica: si ha $A = \overline{A}_T$. Se λ è un autovalore di A ed $\vec{x}$ un autovettore ad esso associato si ha:

$$\lambda \vec{x} = A \vec{x} \quad (9.4)$$

da cui segue subito che è $\lambda \vec{x}_T = (\overline{A \vec{x}})_T = \overline{\vec{x}_T A} = \overline{\vec{x}_T} A$. Allora, utilizzando ancora la (9.4),

$$\overline{\lambda \vec{x}_T} \vec{x} = \overline{\vec{x}_T} A \vec{x} = \overline{\vec{x}_T} \lambda \vec{x} = \lambda \overline{\vec{x}_T} \vec{x};$$

poiché $\overline{\vec{x}_T} \vec{x} \neq 0$ in quanto $\vec{x}$ autovettore, concludiamo che $\overline{\lambda} = \lambda$ e quindi che λ è reale. $\square$

Possiamo ora enunciare il

Teorema 9.8. *Una matrice reale simmetrica è sempre ortogonalmente diagonalizzabile.*

Questo significa non solo che una matrice A simmetrica e reale è sempre diagonalizzabile ma anche che tra le varie matrici che la diagonalizzano se ne può sempre trovare almeno una ortogonale.

OSSERVAZIONE 9.6. Vogliamo esplicitamente notare che l'ipotesi del Teorema 9.8 che la matrice sia reale è essenziale, infatti, per esempio, la matrice simmetrica $A = \begin{bmatrix} 2i & 1 \\ 1 & 0 \end{bmatrix}$ non è neanche diagonalizzabile (verificarlo per esercizio).

Anche l'inverso del Teorema 9.8 è vero, cioè

Teorema 9.9. *Una matrice ortogonalmente simile ad una matrice diagonale reale è simmetrica.*

Dimostrazione. Se $U_T A U = \Delta$ trasponendo si ha $(U_T A U)_T = \Delta_T$ ma Δ è simmetrica, quindi $U_T A_T U = \Delta = U_T A U$ da cui, moltiplicando a destra per U_T e a sinistra per U segue che $A = A_T$. $\square$

Un altro teorema utile soprattutto nelle applicazioni è il

Teorema 9.10. *Se $\vec{x}$ e $\vec{y}$ sono due autovettori della matrice reale simmetrica A associati rispettivamente agli autovalori λ e μ con $\lambda \neq \mu$, allora $\vec{x}$ e $\vec{y}$ sono ortogonali.*

Dimostrazione. Dalle ipotesi abbiamo che $\lambda \vec{x} = A \vec{x}$ e $\mu \vec{y} = A \vec{y}$. Siccome A è simmetrica e quindi λ reale (per il Teorema 9.7 nella pagina precedente), dalla prima uguaglianza segue che $\lambda \vec{x}_T = \vec{x}_T A$ e quindi anche $\lambda \vec{x}_T \vec{y} = \vec{x}_T A \vec{y}$, ma tenendo conto della seconda, si ha $\lambda \vec{x}_T \vec{y} = \mu \vec{x}_T \vec{y}$ quindi $(\lambda - \mu) \vec{x}_T \vec{y} = 0$. Poiché $\lambda \neq \mu$ segue che $\vec{x}_T \vec{y} = 0$. $\square$

Vediamo ora, su alcuni esempi, come si può costruire una matrice ortogonale che diagonalizza una data matrice simmetrica.

Esempio 9.3. Sia $A = \begin{bmatrix} 0 & 1 & 1 \\ 1 & 0 & 1 \\ 1 & 1 & 0 \end{bmatrix}$. Il polinomio caratteristico di A è $\lambda^3 -$

$3\lambda - 2$; quindi gli autovalori sono $\lambda_1 = \lambda_2 = -1$ e $\lambda_3 = 2$.

Il sistema
$$\begin{cases} \lambda x - y - z = 0 \\ -x + \lambda y - z = 0 \\ -x - y + \lambda z = 0 \end{cases}$$
 fornisce gli autovettori di A .

Per $\lambda = -1$ abbiamo la famiglia di autovettori $\begin{bmatrix} \alpha \\ \beta \\ -\alpha - \beta \end{bmatrix}$ con α e β non entrambi nulli, e per $\lambda = 2$ l'autovettore $\begin{bmatrix} \gamma \\ \gamma \\ \gamma \end{bmatrix}$ con $\gamma \neq 0$. Dovremo scegliere due autovettori associati a λ_1 ed uno associato a λ_3 , quindi possiamo costruire la matrice

$$\begin{bmatrix} \alpha & \alpha' & \gamma \\ \beta & \beta' & \gamma \\ -\alpha - \beta & -\alpha' - \beta' & \gamma \end{bmatrix}$$

che dobbiamo rendere ortogonale scegliendo opportuni valori per i parametri α , α' , β , β' e γ . Per il Teorema 9.10 l'ultima colonna è ortogonale a ciascuna delle altre due, quindi basta imporre la condizione $\alpha\alpha' + \beta\beta' + (\alpha + \beta)(\alpha' + \beta') = 0$ che equivale a

$$2\alpha\alpha' + 2\beta\beta' + \alpha\beta' + \beta\alpha' = 0.$$

Poniamo $\alpha = 0$ allora $\beta(2\beta' + \alpha') = 0$ e poiché, in questo caso dev'essere $\beta \neq 0$ bisognerà prendere $\alpha' = -2\beta'$. Se scegliamo allora $\beta = 1$, $\beta' = 1$ e $\gamma = 1$ otteniamo la matrice

$$\begin{bmatrix} 0 & -2 & 1 \\ 1 & 1 & 1 \\ -1 & 1 & 1 \end{bmatrix}$$

le cui colonne sono a due a due ortogonali. Normalizzando le colonne otteniamo la matrice

$$\begin{bmatrix} 0 & \frac{-2}{\sqrt{6}} & \frac{1}{\sqrt{3}} \\ \frac{1}{\sqrt{2}} & \frac{1}{\sqrt{6}} & \frac{1}{\sqrt{3}} \\ \frac{-1}{\sqrt{2}} & \frac{1}{\sqrt{6}} & \frac{1}{\sqrt{3}} \end{bmatrix}$$

che è una matrice ortogonale che diagonalizza la A .

Esempio 9.4. Consideriamo ora la matrice $A = \begin{bmatrix} 1 & 0 & 0 \\ 0 & 2 & 3 \\ 0 & 3 & 2 \end{bmatrix}$. È facile verificare che i suoi autovalori sono $\lambda_1 = 1$, $\lambda_2 = -1$ e $\lambda_3 = 5$ a cui sono associati, rispettivamente, gli autovettori

$$\begin{bmatrix} h \\ 0 \\ 0 \end{bmatrix} \quad \begin{bmatrix} 0 \\ j \\ -j \end{bmatrix} \quad \begin{bmatrix} 0 \\ k \\ k \end{bmatrix}$$

che sono ortogonali, in quanto associati ad autovalori distinti (ancora il Teorema 9.10 a pagina 90). Se scegliamo $h = k = j = 1$ otteniamo la matrice

$$P = \begin{bmatrix} 1 & 0 & 0 \\ 0 & 1 & 1 \\ 0 & -1 & 1 \end{bmatrix}$$

che diagonalizza la A ma non è ortogonale, in quanto le sue colonne non sono normalizzate. Normalizzando si ottiene

$$\begin{bmatrix} 1 & 0 & 0 \\ 0 & \frac{1}{\sqrt{2}} & \frac{1}{\sqrt{2}} \\ 0 & \frac{-1}{\sqrt{2}} & \frac{1}{\sqrt{2}} \end{bmatrix}$$

che è la matrice cercata.

OSSERVAZIONE 9.7. Dai due esempi precedenti, ed in particolare dall'esempio 9.3 a pagina 90 si nota che se A è reale e simmetrica vi sono in generale più matrici ortogonali che la diagonalizzano.

9.3 Forme quadratiche

Un polinomio omogeneo di grado m nelle variabili $x_1, x_2, \dots, x_n$ si chiama *forma di grado m* . In particolare se il polinomio è di secondo grado si parla di *forma quadratica*.

Esempio 9.5. Si consideri il polinomio

$$\Phi(x, y, z) = 2x^2 + 10xy + 9y^2 + 6xz$$

esso è una forma quadratica; osserviamo che si può anche scrivere come

$$\Phi(x, y, z) = 2x^2 + 5xy + 5yx + 9y^2 + 3xz + 3zx \quad (9.5)$$

spezzando i termini rettangolari.

In generale un polinomio omogeneo di secondo grado in n variabili si può scrivere nella forma

$$\Phi(x_1, x_2, \dots, x_n) = \sum_{i,k}^{1\dots n} a_{ik} x_i x_k. \quad (9.6)$$

con $a_{ik} = a_{ki}$.

In tal modo possiamo associare ad ogni forma quadratica in n variabili una matrice simmetrica $A = [a_{ik}]$ di ordine n e possiamo scrivere

$$\Phi = \vec{X}_T A \vec{X}$$

dove $\vec{X}_T = [x_1, x_2, \dots, x_n]$.

Quindi, per esempio, la matrice associata alla forma (9.5) dell'esempio 9.5 è la $A = \begin{bmatrix} 2 & 5 & 3 \\ 5 & 9 & 0 \\ 3 & 0 & 0 \end{bmatrix}$.

Esempio 9.6. La matrice associata alla forma

$$\Phi \equiv x^2 + 2xy - 4yz + 3z^2$$

sarà la matrice simmetrica

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

Si chiama *rango* della forma quadratica Φ il rango della matrice ad essa associata.

Se si opera sulle variabili una trasformazione lineare di matrice B si ottiene una nuova forma quadratica Ψ la cui matrice associata sarà $B_T A B$ da cui risulta ovvio che le due forme hanno lo stesso rango. Si chiama *forma canonica* ogni forma quadratica la cui matrice associata è diagonale, quindi una forma canonica sarà:

$$\Phi = a_{11}x_1^2 + a_{22}x_2^2 + \dots + a_{nn}x_n^2$$

cioè una somma di quadrati.

Ci poniamo il problema: *È possibile, mediante una trasformazione lineare invertibile, ridurre una forma quadratica qualsiasi a forma canonica?*

Consideriamo la forma quadratica $\Phi(x, y) = 5x^2 + 4xy + 2y^2$.

La generica trasformazione lineare

$$\begin{cases} x = au + bv \\ y = cu + dv \end{cases}$$

la trasforma in

$$5(au + bv)^2 + 4(au + bv)(cu + dv) + 2(cu + dv)^2 \quad (9.7)$$

che diventa

$$\begin{aligned} & (5a^2 + 4ac + 2c^2)u^2 + \\ & (10ab + 4ad + 4bc + 4cd)uv + \\ & (5b^2 + 4bd + 2d^2)v^2 \end{aligned} \quad (9.8)$$

Si vede subito che la riducono a forma canonica tutte quelle trasformazioni con $ad \neq bc$ per cui è nullo il coefficiente del termine rettangolare della (9.8), cioè per cui è

$$10ab + 4ad + 4bc + 4cd = 0$$

Questo si può ottenere in infiniti modi, per esempio ponendo $a = 1$, $b = 2$, $c = -2$, $d = 1$ otteniamo la forma canonica

$$\Phi = 5u^2 + 30v^2$$

mentre se prendiamo $a = 0$, $b = 1$, $c = 2$, $d = -1$ abbiamo la forma canonica

$$\Phi = 8u^2 + 3v^2.$$

Sulla riduzione a forma canonica di una forma quadratica sussiste il teorema

Teorema 9.11 (di Lagrange). *Ogni forma quadratica a coefficienti complessi (reali) di rango $r > 0$ si può ridurre, mediante una trasformazione lineare invertibile a coefficienti complessi (reali) alla forma canonica*

$$c_1x_1^2 + c_2x_2^2 + \cdots + c_nx_n^2$$

dove i c_i sono numeri complessi (reali) non tutti nulli.

Se la forma quadratica è in particolare *reale* il teorema di Lagrange si precisa meglio:

Teorema 9.12. *Ogni forma quadratica reale $\Phi = X_TAX$ si può ridurre, mediante una trasformazione ortogonale¹, alla forma canonica*

$$\Psi = \lambda_1x_1^2 + \lambda_2x_2^2 + \cdots + \lambda_nx_n^2$$

dove $\lambda_1, \lambda_2, \dots, \lambda_n$ sono gli autovalori di A .

¹Cioè la cui matrice è ortogonale.

Dimostrazione. Infatti essendo A reale simmetrica, esiste una matrice ortogonale U tale che $U_T A U = \text{diag}(\lambda_1, \lambda_2, \dots, \lambda_n)$ $\square$

OSSERVAZIONE 9.8. Se la forma quadratica ha rango r allora gli autovalori non nulli di A sono esattamente r .

OSSERVAZIONE 9.9. Se la riduzione a forma canonica viene effettuata mediante una generica trasformazione lineare invertibile, non necessariamente ortogonale, non si può garantire che i coefficienti siano gli autovalori di A , ad esempio la forma (9.5) $\Phi = 5x^2 + 4xy + 2y^2$ può essere ridotta a forma canonica in $\Psi = 8u^2 + 3v^2$ ma gli autovalori di A sono 1 e 6.

Si chiama *indice* di una forma quadratica $\Phi = X_T A X$ il numero $p \geq 0$ degli autovalori positivi di A .

Vale il

Teorema 9.13. Ogni forma canonica Ψ ottenuta da una forma quadratica reale Φ mediante una trasformazione lineare invertibile ha il numero dei coefficienti positivi uguale all'indice p di Φ .

Una forma quadratica $\Phi(x_1, x_2, \dots, x_n)$ si dice *definita positiva* (rispettivamente *semidefinita positiva*) se per ogni scelta delle variabili, non tutte nulle, si ha $\Phi(x_1, x_2, \dots, x_n) > 0$ (rispettivamente $\Phi(x_1, x_2, \dots, x_n) \geq 0$).

Vale il

Teorema 9.14. Una forma quadratica $\Phi = X_T A X$ è definita positiva se e solo se tutti gli autovalori di A sono positivi.

Analogamente si parla di forme quadratiche *definite* (semidefinite) *negative*

Vale anche l'analogo del teorema 9.14:

Teorema 9.15. Una forma quadratica $\Phi = X_T A X$ è definita negativa se e solo se tutti gli autovalori di A sono negativi.

Naturalmente esistono forme quadratiche la cui matrice ha sia autovalori positivi, sia autovalori negativi cioè nè definite positive nè definite negative. Qualcuno chiama queste forme *non definite*.

9.4 Matrici hermitiane e matrici unitarie

DEFINIZIONE 9.2. Una matrice A si chiama *hermitiana* (o *autoaggiunta*) se $A = \overline{A}_T$.

DEFINIZIONE 9.3. Una matrice U si chiama *unitaria* se $\bar{U}_T U = I$.

DEFINIZIONE 9.4. Una matrice si chiama *normale* se $A\bar{A}_T = \bar{A}_T A$ cioè se commuta con la sua coniugata trasposta (detta anche *aggiunta*).

OSSERVAZIONE 9.10. È ovvio dalle definizioni che le matrici reali simmetriche sono matrici hermitiane reali e che le matrici ortogonali sono matrici unitarie reali.

OSSERVAZIONE 9.11. Le matrici hermitiane e quelle unitarie sono particolari matrici normali.

Teorema 9.16. Una matrice è unitariamente simile ad una matrice diagonale se e solo se è normale

OSSERVAZIONE 9.12. Ne scende che ogni matrice reale, permutabile con la sua trasposta, è diagonalizzabile in particolare che ogni matrice ortogonale è diagonalizzabile

Con lo stesso procedimento usato per le matrici simmetriche reali nel Teorema 9.10 a pagina 90 si dimostra il

Teorema 9.17. Gli autovalori di una matrice hermitiana sono reali.

Segue anche che

Teorema 9.18. Ogni matrice hermitiana è unitariamente simile ad una matrice diagonale reale.

Sussiste anche il

Teorema 9.19. Una matrice A è normale se e solo se esiste un polinomio $f(\lambda)$ tale che $\bar{A}_T = f(A)$.

Segue il

Corollario 9.20. Una matrice reale A è permutabile con la sua trasposta se e solo se quest'ultima si può esprimere come polinomio in A .

Capitolo 10

Polinomi di matrici

Come abbiamo visto nel paragrafo 3.3 a pagina 24 si definisce un *polinomio di matrici* (o polinomio matriciale) come segue: se

$$p(\lambda) = c_1\lambda^n + c_2\lambda^{n-1} + \cdots + c_n\lambda + c_{n+1}$$

è un polinomio di grado n nella variabile λ , possiamo formalmente sostituire a λ la matrice quadrata A ed ottenere il polinomio matriciale

$$p(A) = c_1A^n + c_2A^{n-1} + \cdots + c_nA + c_{n+1}I. \quad (10.1)$$

OSSERVAZIONE 10.1. Ovviamente $p(A)$ è a sua volta una matrice quadrata dello stesso ordine di A e che si ottiene sviluppando i conti nella (10.1); notiamo anche che il coefficiente c_{n+1} è in realtà coefficiente di λ^0 e quindi, nella sostituzione formale che operiamo, diventa coefficiente di $A^0 = I$.

10.1 Teorema di Cayley Hamilton

Introduciamo ora un Teorema, che va sotto il nome di Teorema di Cayley¹ - Hamilton² e che ha parecchie applicazioni, anche al di fuori dell'ambito dell'Algebra Lineare.

Abbiamo già accennato che una matrice quadrata A si può sempre ridurre a forma triangolare eseguendo su di essa operazioni elementari sulle righe o sulle colonne, quindi possiamo aggiungere che ogni matrice quadrata è "triangolarizzabile" cioè è simile ad una matrice triangolare che ha come elementi principali gli autovalori di A . Dunque esiste una matrice non singolare P , tale che

$$P^{-1}AP = T$$

¹Arthur CAYLEY, 1821, Richmond, Inghilterra - 1895, Cambridge, Inghilterra.

²William HAMILTON, 1788, Glasgow, Scozia - 1856, Edimburgo, Scozia.

con T matrice triangolare.

Esempio 10.1. Consideriamo ora tre matrici triangolari, per esempio alte, per semplicità di ordine tre, B_1 , B_2 e B_3 tali che $\forall i$ in B_i sia nullo l' i -esimo elemento principale. Per esempio siano

$$B_1 = \begin{bmatrix} 0 & a_1 & b_1 \\ 0 & c_1 & d_1 \\ 0 & 0 & e_1 \end{bmatrix} \quad B_2 = \begin{bmatrix} a_2 & b_2 & c_2 \\ 0 & 0 & d_2 \\ 0 & 0 & e_2 \end{bmatrix} \quad B_3 = \begin{bmatrix} a_3 & b_3 & c_3 \\ 0 & d_3 & e_3 \\ 0 & 0 & 0 \end{bmatrix}$$

Un facile calcolo mostra che

$$B_1 B_2 B_3 = \mathbf{0}$$

L'esempio 10.1 è solo un caso particolare di quanto percisa il seguente

Lemma 10.1. Siano $B_1, B_2, \dots, B_n$ n matrici triangolari alte di ordine n tali che, per ogni $i = 1, \dots, n$, l'elemento b_{ii} della i -esima matrice B_i sia nullo. Allora il prodotto delle matrici B_i dà la matrice nulla, cioè

$$B_1 B_2 \cdots B_n = \mathbf{0}$$

Dimostrazione. Decomponiamo in blocchi ciascuna matrice B_i in modo che sia

$$B_i = \begin{bmatrix} H_i & k_i \\ L_i & M_i \end{bmatrix}.$$

Se prendiamo H_1 di tipo $(1, 1)$, H_2 di tipo $(1, 2)$, H_3 di tipo $(2, 3) \dots H_{n-1}$ di tipo $(n-2, n-1)$ ed H_n di tipo $(n-1, n)$ si può moltiplicare a blocchi ciascuna matrice per la successiva. Osservando inoltre che le matrici L_i sono tutte nulle in forza del fatto che le B_i sono triangolari alte, si constata facilmente che il prodotto $B_2 B_3 \cdots B_{n-1}$ è una matrice ripartita a blocchi della forma $\begin{bmatrix} P & Q \\ \mathbf{0} & R \end{bmatrix}$; inoltre, per come sono state costruite, si

$$\text{ha } B_1 = \begin{bmatrix} 0 & H_1 \\ \mathbf{0} & K_1 \end{bmatrix} \text{ e } B_n = \begin{bmatrix} H_n & K_n \\ \mathbf{0} & 0 \end{bmatrix}.$$

Si ha quindi

$$B_1 B_2 \cdots B_n = \begin{bmatrix} 0 & H_1 \\ \mathbf{0} & K_1 \end{bmatrix} \begin{bmatrix} P & Q \\ \mathbf{0} & R \end{bmatrix} \begin{bmatrix} H_n & K_n \\ \mathbf{0} & 0 \end{bmatrix}$$

che dà, evidentemente, la matrice nulla. □

Siamo ora in grado di dimostrare il

Teorema 10.2 (di Cayley-Hamilton). *Ogni matrice quadrata è radice del proprio polinomio caratteristico. Cioè se $\varphi_A(\lambda) = \lambda^n + c_1\lambda^{n-1} + \dots + c_{n-1}\lambda + c_n$ è il polinomio caratteristico della matrice quadrata A , allora vale la relazione matriciale*

$$A^n + c_1A^{n-1} + \dots + c_{n-1}A + c_nI = \mathbf{0} \quad (10.2)$$

Dimostrazione. Sia A quadrata di ordine n e siano $\lambda_1, \lambda_2, \dots, \lambda_n$ i suoi autovalori. Poiché A è triangolarizzabile, esiste una matrice P tale che

$$P^{-1}AP = B = \begin{bmatrix} \lambda_1 & b_{12} & \dots & b_{1n} \\ 0 & \lambda_2 & \dots & b_{2n} \\ \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \dots & \lambda_n \end{bmatrix}$$

Se ora poniamo, $\forall i = 1 \dots n$ $B_i = B - \lambda_i I$ si riconosce subito che le matrici B_i hanno le caratteristiche richieste dal lemma 10.1 nella pagina precedente, inoltre si ha:

$$\begin{aligned} & (A - \lambda_1 I)(A - \lambda_2 I) \dots (A - \lambda_n I) = \\ & = PP^{-1}(A - \lambda_1 I)PP^{-1}(A - \lambda_2 I)PP^{-1} \dots PP^{-1}(A - \lambda_n I)PP^{-1} = \\ & = P(B - \lambda_1 I)(B - \lambda_2 I) \dots (A - \lambda_n I) = \\ & = PB_1B_2 \dots B_nP^{-1} \end{aligned}$$

quindi

$$(A - \lambda_1 I)(A - \lambda_2 I) \dots (A - \lambda_n I) = \mathbf{0} \quad (10.3)$$

e siccome il polinomio caratteristico di A si può scrivere come

$$\varphi(\lambda) = (\lambda - \lambda_1)(\lambda - \lambda_2) \dots (\lambda - \lambda_n) \quad (10.4)$$

sostituendo formalmente la matrice A nella (10.4) in virtù della (10.3) si perviene alla tesi. $\square$

OSSERVAZIONE 10.2 (ATTENZIONE). Il Teorema 10.2 non è invertibile questo significa che se una matrice quadrata di ordine n è radice di un certo polinomio $p(x)$ cioè se si ha che $p(A) = \mathbf{0}$ non è detto che $p(\lambda)$ sia il polinomio caratteristico di A .

Esempio 10.2. Come esempio di quanto affermato nell'osservazione 10.2, consideriamo la matrice $A = I_2$ essa è radice del polinomio matriciale $A^2 - 3A + 2I$ ma il suo polinomio caratteristico è, come si vede immediatamente, $\varphi(\lambda) = \lambda^2 - 2\lambda + 1$

Applicazioni del Teorema di Cayley–Hamilton

Se $\varphi_A(\lambda) = \lambda^n + c_1\lambda^{n-1} + \cdots + c_{n-1}\lambda + c_n$ è il polinomio caratteristico della matrice A , quadrata, di ordine n il Teorema 10.2 nella pagina precedente ci assicura che

$$A^n + c_1A^{n-1} + \cdots + c_{n-1}A + c_nI = \mathbf{0} \quad (10.5)$$

cioè che le successive potenze di A dalla 0 alla n sono *linearmente dipendenti* e quindi, per esempio, che

$$A^n = -(c_1A^{n-1} + \cdots + c_{n-1}A + c_nI)$$

cioè che la potenza n -esima di A si può scrivere come combinazione lineare delle potenze di grado inferiore e che *i coefficienti di questa combinazione lineare sono proprio i coefficienti del polinomio caratteristico*. Inoltre sappiamo che se A è invertibile $c_n \neq 0$ e quindi si può scrivere

$$I = -\frac{1}{c_n}(A^n + c_1A^{n-1} + \cdots + c_{n-1}A)$$

da cui, moltiplicando ambo i membri per A^{-1} si ottiene

$$A^{-1} = -\frac{1}{c_n}(A^{n-1} + c_1A^{n-2} + \cdots + c_{n-1}I) \quad (10.6)$$

La (10.6) ci dice che *la matrice inversa di una matrice invertibile A è combinazione lineare delle potenze di A* e che i coefficienti della combinazione lineare si ricavano facilmente da quelli del polinomio caratteristico di A .

Esempio 10.3. Vogliamo calcolare l'inversa della matrice $A = \begin{bmatrix} 1 & 0 & -1 \\ 1 & 1 & 1 \\ 2 & 0 & 1 \end{bmatrix}$

usando il Teorema di Cayley–Hamilton, cioè usando la (10.6). Il polinomio caratteristico di A sarà $\varphi(\lambda) = \lambda^3 + a\lambda^2 + b\lambda + c$. Essendo $a = -\text{tr}(A) = -3$, $b = \begin{vmatrix} 1 & 0 \\ 1 & 1 \end{vmatrix} + \begin{vmatrix} 1 & 2 \\ 0 & 1 \end{vmatrix} + \begin{vmatrix} 1 & -1 \\ 2 & 1 \end{vmatrix} = 5$ e $c = -\det A = -3$ esso diventa

$$\varphi(\lambda) = \lambda^3 - 3\lambda^2 + 5\lambda - 3,$$

quindi $A^3 - 3A^2 + 5A - 3I = \mathbf{0}$ da cui $A^{-1} = \frac{1}{3}A^2 - A + \frac{5}{3}I$. Allo-

ra si ricava facilmente che $A^2 = \begin{bmatrix} -1 & 0 & -2 \\ 4 & 1 & 1 \\ 4 & 0 & 1 \end{bmatrix}$ e quindi che $A^{-1} =$

$$\frac{1}{3} \begin{bmatrix} 1 & 0 & 1 \\ 1 & 3 & -2 \\ -2 & 0 & 1 \end{bmatrix}.$$

10.2 Polinomio minimo

Nell'Osservazione 10.2 a pagina 99 e nell'esempio successivo abbiamo visto che una matrice può essere anche radice di un polinomio diverso dal suo polinomio caratteristico.

DEFINIZIONE 10.1. Si chiama *polinomio minimo* della matrice quadrata A e si indica con $\mu(\lambda)$ il polinomio di grado minimo e di coefficiente direttore³ uguale a 1 che ammette A come radice.

È facile dimostrare che il polinomio minimo esiste ed è unico: l'esistenza è garantita dal Teorema 10.2 a pagina 99 (infatti esiste per lo meno il polinomio caratteristico) e l'unicità si dimostra per assurdo – farlo per esercizio.

Enunciamo e dimostriamo ora alcune proprietà del polinomio minimo.

Teorema 10.3. *Il polinomio minimo $\mu(\lambda)$ di una matrice quadrata A divide tutti i polinomi che ammettono A come radice, quindi, in particolare, divide il polinomio caratteristico di A .*

Dimostrazione. Sia $p(\lambda)$ un polinomio che ammette A come radice; cioè tale che sia $p(A) = \mathbf{0}$. Se dividiamo $p(\lambda)$ per $\mu(\lambda)$ otteniamo un quoziente ed un resto, cioè

$$p(\lambda) = q(\lambda)\mu(\lambda) + r(\lambda) \quad (10.7)$$

se $r(\lambda) = 0$ abbiamo la tesi, se invece fosse $r(\lambda) \neq 0$, la (10.7) porterebbe alla identità matriciale

$$p(A) = q(A)\mu(A) + r(A)$$

dalla quale, tenendo conto che p e μ ammettono A come radice, segue che $r(A) = \mathbf{0}$, ma siccome r è un polinomio di grado inferiore a μ e μ è quello di grado minimo che ammette A come radice, segue che r è identicamente nullo. $\square$

Dal Teorema 10.3 segue subito che il polinomio minimo divide anche il polinomio caratteristico, e che, quindi, le sue radici sono anche radici di $\varphi(\lambda)$, cioè sono autovalori di A , come precisato dal

Teorema 10.4. *Le radici del polinomio minimo $\mu(\lambda)$ di una matrice A sono tutti e soli gli autovalori di A con molteplicità non maggiori di quelle che hanno come radici del polinomio caratteristico $\varphi(\lambda)$.*

³Ricordiamo la nota 2 a pagina 79.

Questo significa che se, per esempio, una matrice ammette come polinomio caratteristico $\varphi(\lambda) = (\lambda - 1)^2(\lambda - 2)$ il suo polinomio minimo può essere solo $\varphi(\lambda)$ stesso oppure $(\lambda - 1)(\lambda - 2)$.

Si può anche dimostrare che

Teorema 10.5. *Due matrici simili hanno lo stesso polinomio minimo.*

Naturalmente esistono matrici che, pur avendo lo stesso polinomio minimo, non sono simili, anzi, che non hanno nemmeno lo stesso polinomio caratteristico, come si vede nel seguente

Esempio 10.4. Siano $A = \begin{bmatrix} 0 & 0 & 0 \\ 0 & 1 & 1 \\ 0 & 0 & 1 \end{bmatrix}$ e $B = \begin{bmatrix} 0 & 0 & 0 \\ 0 & 0 & 0 \\ 0 & 0 & 1 \end{bmatrix}$. Si vede subito che entrambe hanno come polinomio minimo $\mu(\lambda) = \lambda^2 - \lambda$, infatti $A^2 = A$ e $B^2 = B$, ma $\varphi_A(\lambda) = \lambda(\lambda - 1)^2$ mentre $\varphi_B(\lambda) = \lambda^2(\lambda - 1)$.

Anche l'avere lo stesso polinomio caratteristico non garantisce affatto che i polinomi minimi siano uguali come si vede nel seguente

Esempio 10.5. Consideriamo le matrici $A = \begin{bmatrix} 0 & 0 & 0 \\ 1 & 0 & 0 \\ 2 & 3 & 0 \end{bmatrix}$ e $B = \begin{bmatrix} 0 & 0 & 0 \\ 1 & 0 & 0 \\ 2 & 0 & 0 \end{bmatrix}$. Tutte e due hanno come polinomio caratteristico $\varphi(\lambda) = \lambda^3$, ma siccome si vede subito che $B^2 = \mathbf{0}$ si ha $\mu_B(\lambda) = \lambda^2$ mentre essendo $A^3 = \mathbf{0} \neq A^2$ si ha $\mu_A(\lambda) = \varphi(\lambda) = \lambda^3$.

OSSERVAZIONE 10.3. Con un facile calcolo, che si propone come esercizio, si verifica che il polinomio minimo di una matrice diagonale D_n avente $t \leq n$ autovalori distinti è $(\lambda - \lambda_1)(\lambda - \lambda_2) \cdots (\lambda - \lambda_t)$.

Riferendoci a questa osservazione possiamo dimostrare il

Teorema 10.6. *Una matrice A è diagonalizzabile se e solo se il suo polinomio minimo ammette solo radici semplici.*

Dimostrazione. Dimostriamo, per semplicità, solo la parte “solo se”. Sia A diagonalizzabile, allora essa è simile ad una matrice diagonale il cui polinomio minimo, per l'osservazione precedente, è privo di radici multiple, quindi la tesi segue dal teorema 10.5. $\square$

Il Teorema 10.6 fornisce un altro potente criterio per stabilire se sono diagonalizzabili matrici di cui è facile determinare il polinomio minimo. Ad esempio si ricava da esso che ogni matrice idempotente, cioè per cui sia $A^2 = A$ che ha quindi come polinomio minimo $\lambda^2 - \lambda$ è diagonalizzabile.

Parte II

Geometria piana

Capitolo 11

La retta nel piano

11.1 Preliminari

In un sistema di riferimento cartesiano ortogonale è noto dagli studi precedenti che una retta si rappresenta con un'equazione lineare

$$ax + by + c = 0, \quad (11.1)$$

in cui a e b non siano entrambi nulli. È anche noto che, se la retta non è parallela all'asse y , si può scrivere anche nella forma, cosiddetta *canonica*

$$y = mx + q. \quad (11.2)$$

È altresì noto che la forma canonica (11.2) dell'equazione di una retta mette in luce la "pendenza" della retta: il coefficiente m , che si chiama *coefficiente angolare*, rappresenta la tangente goniometrica dell'angolo che la retta forma con l'asse x : dunque se indichiamo con φ l'ampiezza di quest'angolo, si ha $m = \tan \varphi$.

Da un altro punto di vista una retta r è determinata da un vettore $\vec{v}$ che ne fissa la direzione e da un punto $P_0 \in r$: un punto P appartiene alla retta r se e solo se il segmento PP_0 ha la stessa direzione di $\vec{v}$.

Siano ora (x_0, y_0) le coordinate di P_0 e (x, y) le coordinate di P , il punto P appartiene a r se e solo se il vettore $[x - x_0, y - y_0]$ ed il vettore $\vec{v} = [-b, a]$ sono linearmente dipendenti, cioè se e solo se $a(x - x_0) = -b(y - y_0)$ cioè $ax + by = ax_0 + by_0$ da cui si ottiene la (11.1) ponendo $c = -ax_0 - by_0$.

Dunque la retta di equazione $ax + by = 0$ ha la direzione del vettore $\vec{v} = [-b, a]$.

Poiché $\langle [-b, a], [a, b] \rangle = -ab + ba = 0$, possiamo anche affermare che la retta $ax + by + c = 0$ è ortogonale al vettore $\vec{w} = [a, b]$; i numeri $-b$ e a sono noti come *parametri direttori* della retta, essi sono definiti

a meno di un fattore di proporzionalità non nullo. Normalizzando il vettore $\vec{v}$ si ottiene il vettore $\vec{v}' = [-b', a']$ in cui $-b'$ e a' sono proprio i coseni degli angoli che la retta (orientata) forma con la direzione positiva degli assi coordinati e si chiamano *coseni direttori* della retta.

Da quanto detto segue il

Teorema 11.1. *Sia r la retta di equazione $ax + by + c = 0$ allora:*

- i) *r è parallela all'asse x se e solo se $a = 0$,*
- ii) *r è parallela all'asse y se e solo se $b = 0$,*
- iii) *se r_1 è la retta di equazione $a_1x + b_1y + c_1 = 0$ allora r è parallela a r_1 se e solo se*

$$ab_1 = ba_1,$$

- iv) *r e r_1 sono perpendicolari se e solo se*

$$aa_1 + bb_1 = 0.$$

OSSERVAZIONE 11.1. L'equazione della retta è, ovviamente, definita a meno di un fattore di proporzionalità non nullo, nel senso che le equazioni $ax + by + c = 0$ e $kax + kby + kc = 0$ rappresentano la stessa retta $\forall k \neq 0 \in \mathbb{R}$. D'ora in avanti parleremo comunque, come si fa abitualmente, dell'equazione di una retta (usando l'articolo determinativo), sottintendendo che ci riferiamo ad una qualsiasi delle possibili equazioni della retta, di solito quella la cui scrittura è più semplice, tuttavia questa proprietà non va dimenticata, perché il non tenerne conto può portare a gravi errori, soprattutto nelle applicazioni.

Esempio 11.1. *Scriviamo l'equazione della retta perpendicolare al vettore $[5, 1]$ e passante per il punto $(-1, 2)$. Si vede subito che la retta ha un'equazione della forma $5x + y + c = 0$ e che passa per il punto $(-1, 2)$ se e solo se $5 \cdot (-1) + 2 + c = 0$, da cui $c = 3$ e quindi la retta cercata ha equazione $5x + y + 3 = 0$.*

Esempio 11.2. *Vogliamo l'equazione della retta che passa per i punti $A(4, 1)$ e $B(3, 2)$. Sia essa di equazione $ax + by + c = 0$; si ha, imponendo il passaggio per il primo punto $a(x - 4) + b(y - 1) = 0$ e per il secondo $a(3 - 4) + b(2 - 1) = 0 \iff -a + b = 0 \iff a = b$, dunque, scegliendo $a = 1$ si ha $x + y = 5$, risultato a cui si poteva pervenire direttamente, osservando che 5 è proprio la somma dell'ascissa e dell'ordinata sia di A che di B .*

Si può dimostrare anche che se $a_1x + b_1y + c_1 = 0$ e $a_2x + b_2y + c_2 = 0$ sono due rette che formano un angolo φ si ha

$$\cos \varphi = \frac{|\langle [a_1, b_1], [a_2, b_2] \rangle|}{\|[a_1, b_1]\| \cdot \|[a_2, b_2]\|} = \frac{|a_1a_2 + b_1b_2|}{\sqrt{a_1^2 + b_1^2} \cdot \sqrt{a_2^2 + b_2^2}}$$

11.2 Altri tipi di equazione della retta

L'equazione

$$\frac{x}{p} + \frac{y}{q} = 1 \quad (11.3)$$

è un'equazione lineare e rappresenta quindi una retta; è facile vedere che in questa forma l'equazione della retta (che è detta *equazione segmentaria*) mette in risalto le intercette sugli assi, precisamente la retta di equazione (11.3) interseca gli assi coordinati nei punti $P(p, 0)$ e $Q(0, q)$. Viceversa la retta che passa, per esempio, per i punti $(2, 0)$ e $(0, -3)$ ha equazione $\frac{x}{2} + \frac{y}{-3} = 1$ cioè $3x - 2y - 6 = 0$.

Viceversa tenendo presente l'equazione segmentaria (11.3) possiamo trovare in maniera semplice le intercette sugli assi di una retta scritta in forma generale dividendo ambo i membri per un multiplo comune dei coefficienti delle incognite. Per esempio la retta $3x - 4y = 12$ si scrive anche, dividendo entrambi i membri per 12, nella forma $\frac{x}{4} + \frac{y}{-3} = 1$ mettendo in evidenza che passa per i punti $P(4, 0)$ e $Q(0, -3)$.

Abbiamo visto che una retta r del piano che passa per il punto $P_0(x_0, y_0)$ ed ha la direzione del vettore $\vec{u} = [p, q]$ è l'insieme dei punti (x, y) tali che il vettore $\vec{v} = [x - x_0, y - y_0]$ sia linearmente dipendente dal vettore $\vec{u}$; ma sappiamo anche che due vettori sono linearmente dipendenti se e solo se sono proporzionali, quindi dev'essere $\vec{v} = \lambda \cdot \vec{u}$, da cui

$$\begin{cases} x = x_0 + \lambda p \\ y = y_0 + \lambda q \end{cases} \quad (11.4)$$

le (11.4) si chiamano *equazioni parametriche* della retta; ogni punto della retta ha coordinate espresse dalle (11.4), e per ogni valore del parametro nelle (11.4) si ottiene un punto della retta. Le componenti p e q del vettore $\vec{u}$ sono una coppia¹ di parametri direttori della retta.

¹Abbiamo già osservato che i parametri direttori sono definiti a meno di una costante moltiplicativa non nulla.

Ad esempio se vogliamo le equazioni parametriche della retta

$$r \equiv 3x - 2y = 2$$

possiamo cominciare scegliendo un punto di r , per esempio $P(0, -1)$. La retta data ha direzione del vettore $[2, 3]$ e quindi abbiamo le equazioni parametriche

$$\begin{cases} x = 2t \\ y = 3t - 1 \end{cases} . \quad (11.5)$$

OSSERVAZIONE 11.2. Nel sistema (11.5) abbiamo chiamato t il parametro: il nome è ovviamente arbitrario, ed è consuetudine indicarlo con la lettera t , ma lo studente dovrebbe abituarsi a lavorare con una rappresentazione parametrica della retta in cui il parametro può essere indicato da una lettera qualsiasi.

OSSERVAZIONE 11.3. Mentre l'equazione cartesiana di una retta r è unica a meno di un fattore di proporzionalità non nullo, come notato nell'osservazione 11.1 a pagina 106, le equazioni parametriche di una stessa retta possono assumere aspetti molto diversi.

Per esempio la retta di equazioni parametriche (11.5) si può anche scrivere con le equazioni parametriche

$$\begin{cases} x = \frac{2+t}{3} \\ y = \frac{t}{2} \end{cases} . \quad (11.6)$$

Viceversa per passare dalle equazioni parametriche ad un'equazione cartesiana si può procedere in vari modi:

- dando due valori al parametro si ottengono due punti della retta, poi si scrive l'equazione della retta per i due punti trovati.
- oppure si può eliminare il parametro dalle equazioni, trovando un'equazione cartesiana della retta data.

Per esempio se r ha equazioni parametriche

$$\begin{cases} x = 2t \\ y = 2t - 1 \end{cases}$$

posiamo eliminare il parametro sottraendo membro a membro le due equazioni si ottiene così l'equazione cartesiana della r che risulta quindi $x - y = 1$.

11.3 Distanze

La distanza di due punti nel piano segue da una immediata applicazione del Teorema di Pitagora, infatti dalla figura 11.1 si nota subito che la distanza d tra i punti $A(x_1, y_1)$ e $B(x_2, y_2)$ è l'ipotenusa di un triangolo rettangolo i cui cateti sono $AC = x_2 - x_1$ e $BC = y_2 - y_1$ da cui

$$d = \sqrt{(x_2 - x_1)^2 + (y_2 - y_1)^2}$$

Ricordiamo inoltre che la distanza del punto P di coordinate (x_0, y_0) dalla retta di equazione $ax + by + c = 0$ è

$$d = \frac{|ax_0 + by_0 + c|}{\sqrt{a^2 + b^2}}.$$

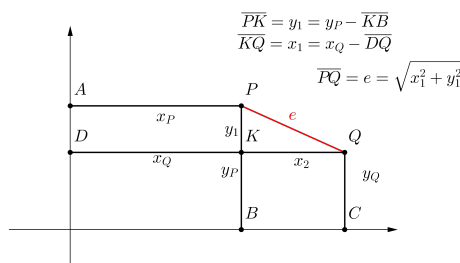


Figura 11.1 Distanza di due punti

11.4 Fasci di rette

L'insieme di tutte le rette del piano che passano per un punto P prende il nome di *fascio*² di rette che ha come *centro* o *sostegno* il punto P . L'equazione globale di tutte le rette del fascio $\mathcal{F}$ che ha per sostegno P si può ottenere facilmente come combinazione lineare non banale delle equazioni di due qualsiasi rette di $\mathcal{F}$. Questo significa che se abbiamo due rette distinte r e r' entrambe passanti per P e rispettivamente di equazioni $ax + by + c = 0$ e $a'x + b'y + c' = 0$ l'equazione

$$\lambda(ax + by + c) + \mu(a'x + b'y + c') = 0 \quad (11.7)$$

rappresenta *tutte e sole* le rette che passano per P se λ e μ non sono entrambi nulli. Infatti per ogni coppia di valori (non entrambi nulli) di λ e μ , la (11.7) rappresenta una retta passante per P . Viceversa una qualunque retta per P è rappresentata da un'equazione del tipo (11.7), cioè qualunque retta per P è individuata da una particolare coppia di valori di λ e μ .

OSSERVAZIONE 11.4. L'equazione (11.7) è omogenea (rispetto ai parametri λ e μ) cioè è definita a meno di un fattore di proporzionalità non

²Il concetto di fascio –di rette, di piani, di circonferenze...– è molto generale e lo incontreremo ancora.

nullo. Può essere comodo scriverla, invece, usando un solo parametro, per esempio nella forma

$$k(ax + by + c) + (a'x + b'y + c') = 0 \quad (11.8)$$

in cui abbiamo posto $k = \frac{\lambda}{\mu}$. In questo caso, però, quando $\mu = 0$ il parametro k perde di significato, e quindi l'equazione del fascio, nella forma (11.8) non rappresenta tutte le rette del fascio, mancando quella per cui $\mu = 0$, cioè la retta $ax + by + c = 0$. Se si considera, però, che $\lim_{\mu \rightarrow 0} k = \infty$, si può accettare il valore infinito per il parametro k , convenendo che in questo caso la (11.8) rappresenta la retta $ax + by + c = 0$.

I fasci sono uno strumento molto potente e comodo: vediamo un primo esempio.

Esempio 11.3. Vogliamo scrivere l'equazione della la retta che passa per i punti $A(1, 0)$ e $B(2, -3)$. Essa appartiene al fascio che ha per sostegno, per esempio, il punto A e quindi di equazione

$$x - 1 + ky = 0 \quad (11.9)$$

ottenuta come combinando linearmente le equazioni $x = 1$ e $y = 0$ delle rette parallele agli assi passanti per A . La retta cercata si otterrà imponendo il passaggio della generica retta del fascio per il punto P , quindi sostituendo le coordinate di P nella (11.9), cioè scrivendo $2 - 1 - 3k = 0$ da cui si ottiene $k = \frac{1}{3}$ ed ottenendo quindi l'equazione $3x + y - 3 = 0$.

Esempio 11.4. Vogliamo l'equazione della retta passante per $P(1, 0)$ e perpendicolare alla retta $3x - 2y + 1 = 0$. Scriviamo l'equazione del fascio di rette per P ; essa sarà:

$$\lambda x + \mu(y - 1) = 0$$

per la condizione di perpendicolarità si avrà $3\lambda - 2\mu = 0$ da cui, per esempio, $\lambda = 2$ e $\mu = 3$ a cui corrisponde la retta $2x + 3y - 3 = 0$

11.5 Coordinate omogenee

Siccome due rette non parallele hanno in comune un punto e due rette parallele una direzione, fa comodo, in certi contesti, assimilare una direzione ad un punto "improprio" o punto "all'infinito". Con questa convenzione due rette distinte nel piano hanno sempre un punto

(proprio o improprio) in comune, quindi due rette sono parallele se (e solo se) hanno in comune un punto improprio. I punti impropri verranno denotati con l'apice ∞ , per esempio P_∞ .

Sorge il problema di come "coordinatizzare" i punti impropri. Nel piano, come sappiamo, un punto al finito può essere rappresentato da una coppia ordinata di numeri reali, per poter rappresentare anche i punti impropri conviene utilizzare una terna di coordinate, cosiddette *omogenee*, precisamente, se $P(X, Y)$ attribuiamo a P le tre coordinate, non tutte nulle, x, y e u legate alle precedenti dalla relazione

$$X = \frac{x}{u} \quad \text{e} \quad Y = \frac{y}{u} \quad (11.10)$$

Possiamo allora dire che P ha coordinate *omogenee* x, y, u e scriviamo $P(x : y : u)$. Se $u \neq 0$ possiamo scrivere che $P(X, Y)$ ha coordinate omogenee $P(X : Y : 1)$, per esempio possiamo attribuire all'origine le coordinate omogenee $O = (0 : 0 : 1)$; i punti impropri saranno allora tutti e soli quelli la cui terza coordinata omogenea è nulla.

OSSERVAZIONE 11.5. Le coordinate omogenee di un punto sono definite a meno di un fattore di proporzionalità, questo significa che il punto di coordinate omogenee $P(3 : 2 : 1)$ coincide con il punto di coordinate omogenee $P(6 : 4 : 2)$, e quindi che un punto a coordinate razionali si può sempre considerare come un punto a coordinate omogenee intere: per esempio il punto proprio $P = \left(\frac{1}{5}, \frac{3}{5}\right)$ ha coordinate omogenee $P(1 : 3 : 5)$

Consideriamo ora l'equazione generale della retta: in coordinate non omogenee essa sarà $aX + bY + c = 0$: applicando le (11.10) diventerà $ax + by + cu = 0$ ed avrà come punto improprio il punto per cui $u = 0$ che sarà dunque $P_\infty(-b : a : 0)$.

In questo contesto, l'equazione $u = 0$ rappresenta tutti e soli i punti la cui terza coordinata omogenea è nulla, quindi tutti (e soli) i punti impropri, essendo un'equazione lineare possiamo dire che essa rappresenta una retta, precisamente quella che chiamiamo *retta impropria*, cioè il luogo dei punti impropri del piano.

Concludiamo il paragrafo con la seguente

OSSERVAZIONE 11.6. Un sistema lineare omogeneo di tre equazioni *in tre incognite* può sempre essere interpretato geometricamente, in coordinate omogenee, come la ricerca dell'eventuale punto comune di tre rette.

Il seguente esempio chiarisce la situazione.

Esempio 11.5. *Il sistema*

$$\begin{cases} hx - y + hu = 0 \\ hx - y - u = 0 \\ x - hy + u = 0 \end{cases}$$

è lineare omogeneo, quindi ammette la soluzione banale $(0 : 0 : 0)$ che in coordinate omogenee non rappresenta alcun punto. La matrice dei coefficienti

è: $\begin{bmatrix} h & -1 & h \\ h & -1 & -1 \\ 1 & -h & 1 \end{bmatrix}$; essa ha rango $r = 3$ per $h \neq \pm 1$, ha $r = 2$ per $h = 1$ e

$r = 1$ per $h = -1$ quindi per $h \neq \pm 1$ le rette non hanno in comune alcun punto nè proprio nè improprio, per $h = 1$ ci sono ∞^1 soluzioni, che sono le coordinate omogenee di uno ed un solo punto (proprio o improprio) e per $h = -1$ ci sono ∞^2 soluzioni che rappresentano le coordinate omogenee dei punti di una retta; quindi le tre rette coincidono.

11.6 I sistemi di riferimento**Cambiamento del sistema di riferimento**

Abbiamo visto che un sistema di riferimento cartesiano ortogonale è determinato da un'origine O , che corrisponde al punto di coordinate $(0,0)$ e dai due versori fondamentali degli assi u_1 e u_2 e che il punto $P(x, y,)$ è rappresentato dal vettore $OP = xu_1 + yu_2$, cioè è una combinazione lineare dei versori fondamentali.

Se $P(x, y)$ è il generico punto del piano e se si effettua una traslazione di assi che porta l'origine nel nuovo punto $O'(\alpha, \beta)$, le coordinate (x', y') di P rispetto ai nuovi assi saranno:

$$\begin{cases} x' = x + \alpha \\ y' = y + \beta \end{cases}$$

come si vede chiaramente nella figura 11.2 nella pagina successiva.

Ci si rende conto molto facilmente che le traslazioni sono trasformazioni lineari di $\mathbb{R}^2$ in sè che conservano le distanze –verificarlo per esercizio–.

Ci si può chiedere che cosa succede delle coordinate di $P(x, y)$ quando si cambia il sistema di riferimento, prendendo altri due vettori indipendenti come base. Questo significa effettuare una rotazione del sistema di riferimento. Al punto P saranno associati altri due numeri (x', y') e ci si chiede qual è il legame tra queste coppie di numeri.

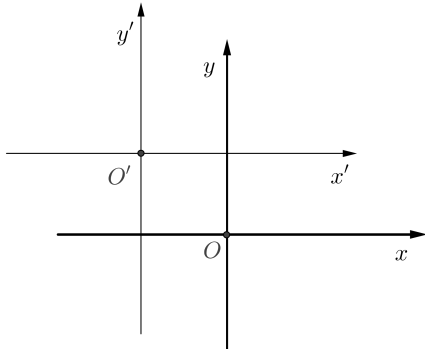


Figura 11.2 Traslazione

Nella figura 11.3 nella pagina seguente OQ ed OR sono rispettivamente l'ascissa e l'ordinata del punto P rispetto al sistema di riferimento non ruotato e OQ' ed OR' quelle rispetto al sistema ruotato.

Si effettua in questo modo un cambiamento di base nello spazio vettoriale $\mathbb{R}^2$ che, come sappiamo, è rappresentato da una matrice quadrata di ordine 2.

Se consideriamo i cambiamenti di sistema di riferimento che lasciano ferma l'origine (escludiamo il caso noto delle traslazioni di assi), osserviamo che se P ha coordinate (x, y) in un sistema di riferimento e (x', y') nell'altro, le equazioni che legano le coordinate di P nei due sistemi di riferimento sono date dal sistema

$$\begin{cases} x = ax' + by' \\ y = cx' + dy' \end{cases} \quad (11.11)$$

o, in forma più compatta da

$$\vec{x} = A \vec{x}'$$

dove $\vec{x} = \begin{bmatrix} x \\ y \end{bmatrix}$ e $\vec{x}' = \begin{bmatrix} x' \\ y' \end{bmatrix}$ ed $A = \begin{bmatrix} a & b \\ c & d \end{bmatrix}$,

Dimostriamo ora il

Teorema 11.2. *Si consideri un cambiamento di sistema di riferimento che lascia ferma l'origine, quindi rappresentato dalle equazioni (11.11) e che conservi la distanza di due punti³. Allora la matrice A è ortogonale e siamo in presenza di una rotazione di assi vedi Figura 11.3 nella pagina successiva.*

Dimostrazione. Sia $P(x_0, y_0)$ un punto. Senza ledere la generalità, possiamo supporre che il secondo punto sia l'origine $O(0, 0)$ (se non lo fosse possiamo effettuare prima una opportuna traslazione di assi). La distanza d di P da O è $\sqrt{x_0^2 + y_0^2}$ quindi deve essere $d^2 = x_0^2 + y_0^2 = x_0'^2 + y_0'^2$.

³cioè se la distanza di due punti P e Q è d nel sistema non accentato, essa resta d in quello accentato

Ma si ha

$$\begin{aligned} x_0^2 + y_0^2 &= (ax'_0 + by'_0)^2 + (cx'_0 + dy'_0)^2 = \\ &= a^2 x_0'^2 + b^2 x_0'^2 + 2abx'_0 y'_0 + c^2 x_0'^2 + d^2 x_0'^2 + 2cdx'_0 y'_0 = \\ &= (a^2 + c^2)x_0'^2 + (b^2 + d^2)y_0'^2 + 2(ab + cd)x'_0 y'_0. \end{aligned}$$

che è uguale a $x_0'^2 + y_0'^2$ se e solo se

$$\begin{cases} a^2 + c^2 = 1 \\ b^2 + d^2 = 1 \\ ab + cd = 0 \end{cases}$$

che sono le condizioni per cui A è una matrice ortogonale. □

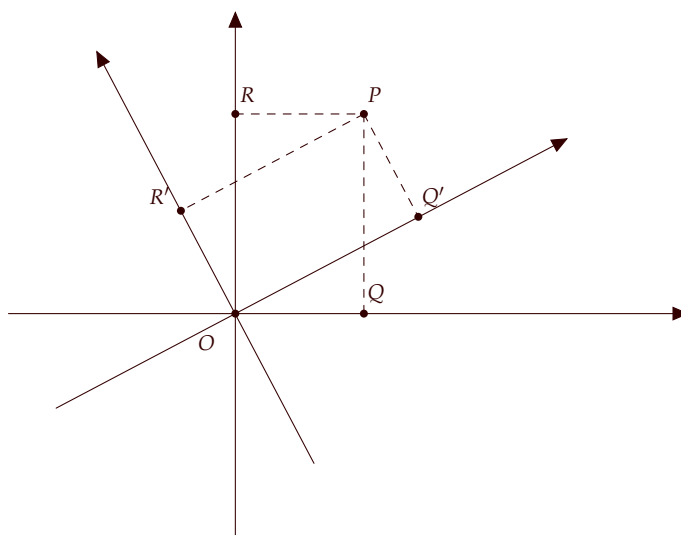


Figura 11.3 Rotazione

11.7 Coordinate polari

Presentiamo ora un altro sistema di riferimento che può essere comodo in varie circostanze.

Parliamo delle *coordinate polari nel piano*. Sia P un punto distinto dall'origine e siano (x, y) le sue coordinate in un sistema di riferimento cartesiano. Se $\rho = \sqrt{x^2 + y^2}$ è la distanza di P dall'origine, e ϑ l'angolo

che il vettore $\overrightarrow{OP}$ forma con l'asse, si vede subito dalla figura 11.4 che si ha

$$\begin{cases} x = \rho \cos \vartheta \\ y = \rho \sin \vartheta \end{cases}$$

con $\rho > 0$, $0 \leq \vartheta < 2\pi$.

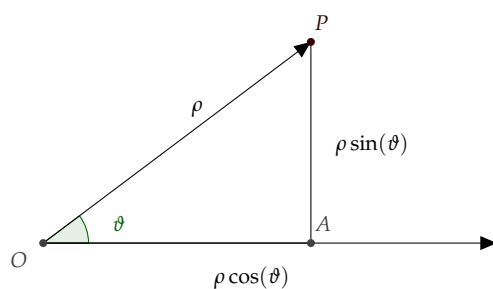


Figura 11.4 *Coordinate polari*

Possiamo allora dire che il punto P ha *coordinate polari* (ρ, ϑ) . Se P coincide con l'origine, ovviamente ϑ non è definito, e possiamo dire che esso è individuato dalla sola coordinata $\rho = 0$.

OSSERVAZIONE 11.7. Mentre le coordinate polari sono rappresentate da una coppia ordinata di numeri reali che rappresentano due lunghezze, i due numeri reali che rappresentano le coordinate polari, invece, sono rispettivamente una lunghezza e l'ampiezza di un angolo (misurata in radianti), quindi $\rho \in \mathbb{R}^+$ e $0 \leq \vartheta < 2\pi$.

Capitolo 12

La circonferenza nel piano

12.1 Generalità

Consideriamo un riferimento cartesiano ortogonale. La distanza tra i punti $P(x_1, y_1)$ e $Q(x_2, y_2)$ è, come è noto,

$$d = \sqrt{(x_1 - x_2)^2 + (y_1 - y_2)^2}.$$

Sappiamo che una circonferenza Γ di centro $C(x_0, y_0)$ e raggio r è l'insieme dei punti del piano che distano r da C . Per trovare l'equazione di tale circonferenza basta tradurre in equazione la definizione, cioè, se $P(x, y)$ è un punto qualsiasi del piano, esso sta sulla circonferenza se e solo se la sua distanza da C è r ; cioè

$$r = \sqrt{(x - x_0)^2 + (y - y_0)^2}$$

da cui

$$(x - x_0)^2 + (y - y_0)^2 = r^2 \quad (12.1)$$

che assume la più usuale forma

$$x^2 + y^2 + ax + by + c = 0 \quad (12.2)$$

pur di porre $a = -2x_0$, $b = -2y_0$ e $c = x_0^2 + y_0^2 - r^2$.

Ma si può anche far vedere che ogni equazione del tipo (12.2) rappresenta una circonferenza. Infatti, confrontando la (12.1) con la (12.2) si vede subito che essa rappresenta la circonferenza di centro $C\left(-\frac{a}{2}, -\frac{b}{2}\right)$

e di raggio $r = \frac{\sqrt{a^2 + b^2 - 4c}}{2}$ purché si abbia

$$a^2 + b^2 - 4c > 0. \quad (12.3)$$

Esempio 12.1. La circonferenza che ha centro nel punto $C(1, 0)$ e raggio $r = \sqrt{2}$ ha equazione $(x - 1)^2 + y^2 = 2$ che, sviluppata, diventa $x^2 + y^2 - 2x - 1 = 0$.

Esempio 12.2. Sia γ la circonferenza

$$2x^2 + 2y^2 - 4x - 8y + 1 = 0. \quad (12.4)$$

Vogliamo trovarne centro e raggio. L'equazione (12.4) si può scrivere anche nella forma (canonica) $x^2 + y^2 - 2x - 4y + \frac{1}{2} = 0$ da cui si ha subito che il centro ha coordinate $C(1, 2)$ ed il raggio è $r = \frac{\sqrt{4+16-2}}{2} = \frac{3\sqrt{2}}{2}$.

OSSERVAZIONE 12.1. È comodo rimuovere l'eccezione (12.3) in modo da poter affermare che *tutte* le equazioni del tipo (12.2) cioè *tutte* le equazioni di secondo grado in cui manca il termine rettangolare ed in cui i coefficienti di x^2 ed y^2 sono uguali rappresentano una circonferenza. Per far ciò bisogna ampliare il piano cartesiano con i punti a coordinate complesse e quindi ammettere che sia una circonferenza anche la curva rappresentata dall'equazione

$$x^2 + y^2 = 0$$

che ha un solo punto reale, e rappresenta la circonferenza con centro nell'origine e raggio nullo o peggio, l'equazione

$$x^2 + y^2 + 1 = 0$$

che rappresenta la circonferenza, completamente immaginaria, di centro l'origine e raggio immaginario $r = i$.

La circonferenza nel piano ha anche un'interessante rappresentazione parametrica: come si vede dalla figura 12.1 nella pagina successiva, se il centro è il punto $C(x_0, y_0)$ ed il raggio è r si ha che $\cos \varphi = \frac{x - x_0}{r}$ e $\sin \varphi = \frac{y - y_0}{r}$ dove φ è l'angolo che il raggio CP forma con la direzione dell'asse x . Dunque la circonferenza che ha centro in $C(x_0, y_0)$ e raggio r ha equazioni parametriche

$$\begin{cases} x = x_0 + r \cos \varphi \\ y = y_0 + r \sin \varphi \end{cases}. \quad (12.5)$$

Due circonferenze possono intersecarsi, essere esterne una all'altra, essere interne una all'altra o tangenti, internamente od esternamente.

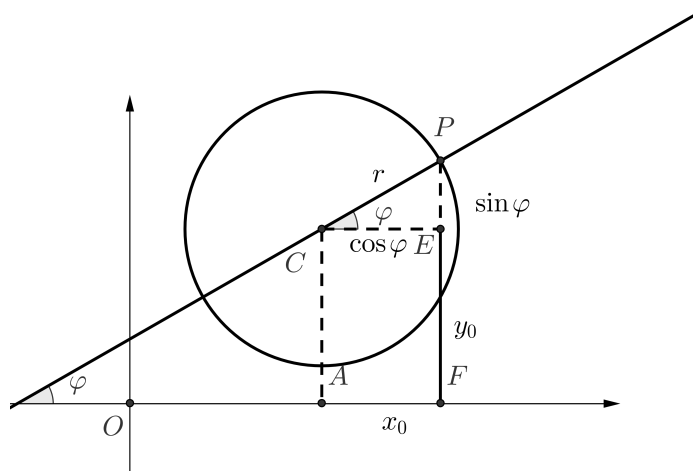


Figura 12.1 La circonferenza in equazioni parametriche

Esempio 12.3. Siano $\gamma_1 \equiv x^2 + y^2 - 2x = 0$ e $\gamma_2 \equiv 2x^2 + 2y^2 - y = 0$. Vogliamo trovare le loro intersezioni. Si osserva subito che passano entrambe per l'origine, quindi o sono ivi tangenti o sono secanti. Si potrebbe stabilirlo esaminando la relazione che c'è tra la distanza dei centri e la somma o la differenza dei raggi, ma qui ci interessano le intersezioni. Dobbiamo studiare quindi il sistema

$$\begin{cases} x^2 + y^2 - 2x = 0 \\ 2x^2 + 2y^2 - y = 0 \end{cases} \text{ equivalente a } \begin{cases} x^2 + y^2 - 2x = 0 \\ x^2 + y^2 - \frac{1}{2} = 0 \end{cases}$$

e anche, sottraendo membro a membro, al sistema

$$\begin{cases} x^2 + y^2 - 2x = 0 \\ -2x + \frac{1}{2}y = 0 \end{cases}$$

che fornisce le soluzioni $x_1 = 0$, $y_1 = 0$ e $x_2 = \frac{2}{17}$, $y_2 = \frac{8}{17}$.

Una retta ed una circonferenza hanno in comune due punti: se i due punti sono reali e distinti la retta è *secante*, se sono reali e coincidenti la retta è *tangente* se sono immaginarie la retta è *esterna*

12.2 Tangenti

Sussiste il ben noto

Teorema 12.1. Siano P un punto e γ una circonferenza. Se P è esterno a γ da P escono due e due sole tangenti alla circonferenza, invece se $P \in \gamma$ la tangente è una sola.

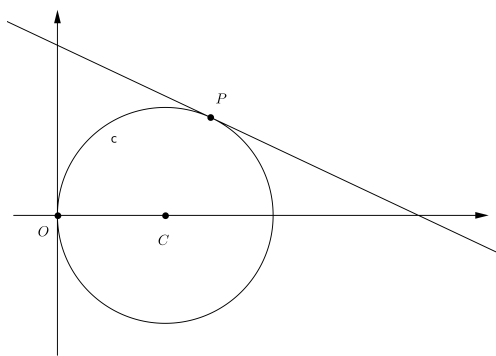


Figura 12.2 Tangente in P ad una circonferenza

Per determinare l'equazione di una tangente può essere utile ricordare una proprietà elementare illustrata nella figura 12.2: la tangente ad una circonferenza in un suo punto è perpendicolare al raggio passante per quel punto.

Si può dimostrare il

Teorema 12.2. Sia $\gamma : (x - x_0)^2 + (y - y_0)^2 = r^2$ l'equazione di una circonferenza e sia $T(x_1, y_1) \in \Gamma$. Allora la retta tangente a γ in T ha equazione

$$(x - x_1)(x_1 - x_0) + (y - y_1)(y_1 - y_0) = 0 \quad (12.6)$$

Dimostrazione. Poiché la tangente passa per P , essa ha equazione $a(x - x_1) + b(y - y_1) = 0$. Questa retta deve essere perpendicolare alla retta PC cioè al vettore $[x_1 - x_0, y_1 - y_0]$ e quindi possiamo porre $a = x_1 - x_0$ e $b = y_1 - y_0$. $\square$

Nel caso di P esterno per trovare le tangenti uscenti da P si può procedere come nel seguente

Esempio 12.4. Vogliamo determinare le tangenti alla circonferenza

$$\gamma \equiv x^2 + y^2 - 2x = 0$$

passanti per $P(0, 2)$. (vedi la figura 12.3 nella pagina successiva). La γ ha centro nel punto $C(1, 0)$ e raggio $r = 1$. Si nota subito da questo fatto che una delle tangenti è l'asse y . Per trovare l'altra tangente si può procedere in vari modi

- i) Nel fascio di rette che ha per sostegno P si scelgono quelle a distanza 1 da C .

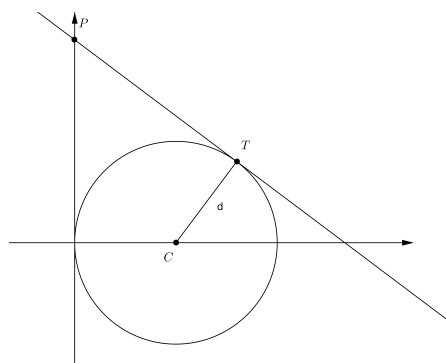


Figura 12.3 Tangenti da un punto ad una circonferenza

ii) L'equazione canonica della generica retta per P è $y = mx + 2$. Quindi, intersecando con la γ si ha il sistema:

$$\begin{cases} y = mx + 2 \\ x^2 + y^2 - 2x = 0 \end{cases}$$

che ammette come equazione risolvente $x^2 + (mx + 2)^2 - 2x = 0$ che diventa, con semplici passaggi,

$$(m^2 + 1)x^2 + 2(2m - 1)x + 4 = 0. \quad (12.7)$$

La (12.7) ammette due soluzioni coincidenti quando

$$\frac{\Delta}{4} = (2m - 1)^2 - 4(m^2 + 1) = 0$$

cioè quando $-4m - 3 = 0$ da cui si ha la retta

$$y = -\frac{3}{4}x + 2 \quad \text{cioè} \quad 3x + 4y - 8 = 0.$$

OSSERVAZIONE 12.2. Nell'esempio 12.4 a fronte abbiamo trovato una sola retta: questo è dovuto al fatto che abbiamo usato l'equazione canonica che, come è noto, non individua rette parallele all'asse y .

12.3 Fasci di circonferenze

L'insieme di tutte le circonferenze che passano per due punti fissi A e B si chiama *fascio di circonferenze* ed i due punti prendono il nome di *punti base* del fascio; se γ_1 e γ_2 sono due circonferenze di equazioni rispettive

$$x^2 + y^2 + a_1x + b_1y + c_1 = 0 \quad (12.8)$$

e

$$x^2 + y^2 + a_2x + b_2y + c_2 = 0 \quad (12.9)$$

si vede subito che l'equazione

$$\lambda(x^2 + y^2 + a_1x + b_1y + c_1) + \mu(x^2 + y^2 + a_2x + b_2y + c_2) = 0 \quad (12.10)$$

rappresenta, per $\lambda \neq -\mu$ ancora una circonferenza che passa per A e B . Per $\lambda = -\mu$ la (12.10) rappresenta una retta passante per A e B che è detta *asse radicale* del fascio e che può esser considerata la circonferenza di raggio massimo (infinito) del fascio. Quindi

Teorema 12.3. *L'equazione del fascio che ha per sostegno i punti A e B cioè di tutte e sole le circonferenze che passano per questi punti si ottiene come combinazione lineare non banale di quelle di due qualsiasi circonferenze passanti per A e B .*

Come abbiamo osservato, l'asse radicale, che appartiene al fascio perché si ottiene ponendo $\lambda = -\mu$ nella (12.10), viene considerato come una circonferenza di raggio infinito, e come tale viene spesso usato per scrivere l'equazione del fascio stesso.

Facendo riferimento all'osservazione 11.4 a pagina 109 notiamo che anche i fasci di circonferenze si possono descrivere con un solo parametro non omogeneo, pur di tener conto delle condizioni elencate appunto nell'osservazione 11.4.

Esempio 12.5. *Se vogliamo scrivere l'equazione della circonferenza che passa per i tre punti $A(1,0)$, $B(3,0)$ e $C(2,3)$ possiamo scrivere anzitutto l'equazione del fascio che ha come punti base A e B combinando linearmente due qualsiasi circonferenze per A e B : le più comode sono quella di raggio massimo (cioè l'asse radicale, di equazione $y = 0$) e quella di raggio minimo, cioè quella che ha per diametro AB , quindi centro nel punto $O(2,0)$ e raggio 1 cioè di equazione $(x-2)^2 + y^2 = 1$, da cui l'equazione del fascio*

$$x^2 + y^2 - 4x + \lambda y + 3 = 0$$

dove, in accordo con quanto detto nell'osservazione 11.4 a pagina 109, abbiamo usato un solo parametro, ovviamente posto nella posizione più comoda. A questo punto la circonferenza che cerchiamo sarà quella del fascio che passa per il punto C , cioè quella per cui

$$2^2 + 3^2 - 4 \cdot 2 + 3\lambda + 3 = 0.$$

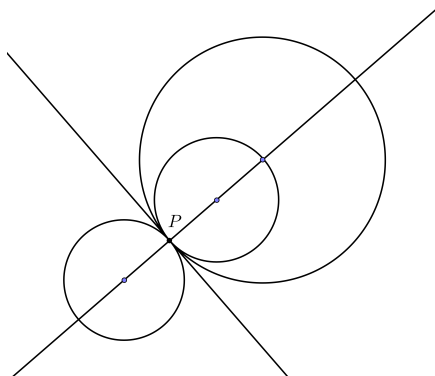


Figura 12.4 Fascio di circonferenze tangenti ad una retta

I due punti base del fascio possono anche essere coincidenti: è il caso di un fascio di circonferenze tangenti in un punto ad una retta (Fig. 12.4). In questo caso la circonferenza di raggio massimo è sempre l'asse radicale, cioè la retta tangente, e quella di raggio minimo si riduce alla circonferenza (immaginaria) che ha centro nel punto e raggio nullo.

Esempio 12.6. Vogliamo l'equazione del fascio di circonferenze tangenti nel punto $P(1,2)$ alla retta di

equazione $y = 2x$.

Possiamo combinare linearmente l'equazione della retta con quella della circonferenza che ha centro in P e raggio 0, che è: $(x - 1)^2 + (y - 2)^2 = 0$ ottenendo l'equazione del fascio

$$(x - 1)^2 + (y - 2)^2 + \lambda(2x - y) = 0$$

Si parla di fascio anche quando le due circonferenze hanno in comune solo punti immaginari, questo è il caso, per esempio, di fasci di circonferenze concentriche.

Esempio 12.7. Per esempio il fascio di circonferenze che hanno centro nel punto $C(1,0)$ può essere rappresentato dall'equazione

$$(x - 1)^2 + y^2 + k[(x - 1)^2 + y^2 - 1]$$

ottenuto combinando linearmente la circonferenza di raggio nullo e quella di raggio 1, entrambe con centro in C .

12.4 Circonferenza ed elementi impropri

Il problema delle intersezioni di due circonferenze dà luogo, come abbiamo visto, ad un sistema di quarto grado (due equazioni di secondo), tuttavia le soluzioni, reali o complesse che siano, sono al più due. La ragione di questo strano fatto risiede nella considerazione dell'esistenza di due punti impropri che appartengono a tutte le circonferenze del piano; infatti, in coordinate omogenee l'equazione della generica circonferenza è

$$x^2 + y^2 + axu + byu + cu^2 = 0.$$

I punti di intersezione di questa circonferenza con la retta impropria sono le soluzioni del sistema omogeneo

$$\begin{cases} x^2 + y^2 = 0 \\ u = 0 \end{cases}$$

e cioè i punti di coordinate omogenee $(1 : \pm i : 0)$ che, come si verifica immediatamente, soddisfano l'equazione di una *qualsiasi* circonferenza; essi prendono il nome di *punti ciclici* del piano.

Esempio 12.8. *Nel caso di un fascio di circonferenze concentriche, ad esempio (scrivendo le equazioni in coordinate omogenee)*

$$\gamma_1 \equiv x^2 + y^2 - 2xu + 2yu - 3u^2 = 0 \text{ e } \gamma_2 \equiv x^2 + y^2 - 2xu + 2yu = 0$$

se cerchiamo l'asse radicale otteniamo l'equazione $3u^2 = 0$, cioè la retta impropria contata due volte. Questo significa che tutte le circonferenze di questo fascio sono tangenti alla retta impropria nei punti ciclici.

Capitolo 13

Le coniche

In questo paragrafo esamineremo le principali proprietà di alcune curve piane note con il nome di *coniche*¹ e cioè delle curve che vanno sotto il nome di *ellisse*, *parabola* ed *iperbole*.

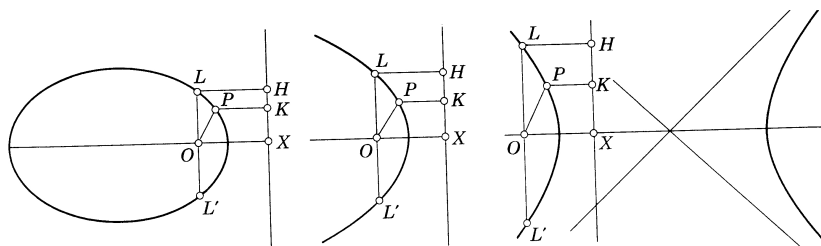


Figura 13.1 *Le coniche*

In generale se O è un punto fissato del piano ed r una retta non passante per O chiamiamo *conica* il luogo dei punti P tali che la distanza PO sia $\varepsilon > 0$ volte la distanza tra P ed r . Il numero ε che si chiama *eccentricità della conica* viene spesso indicato anche con la lettera e^2 . Una conica si chiama *ellisse* se $\varepsilon < 1$, *parabola* se $\varepsilon = 1$ e *iperbole* se $\varepsilon > 1$. (Fig.13.1)

Nello specifico possiamo dire che

¹Il nome deriva dal fatto che l'ellisse, la parabola e l'iperbole sono sezioni piane di un cono circolare.

²Da non confondere con il numero $e = 2,71 \dots$ di Nepero, base dei logaritmi naturali

DEFINIZIONE 13.1. Dati due punti F_1 e F_2 del piano, si chiama *ellisse* (Fig. 13.2) di fuochi F_1 e F_2 l'insieme dei punti P del piano tali che sia costante la somma delle distanze di P da F_1 e F_2

$$d(PF_1) + d(PF_2) = 2a$$

dove a è una costante tale che $2a > d(F_1F_2)$.

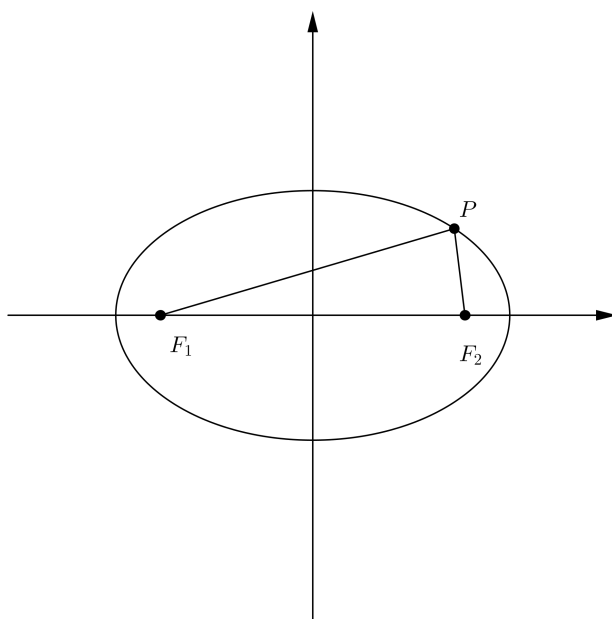


Figura 13.2 L'ellisse

Determiniamo l'equazione dell'ellisse scegliendo il sistema di riferimento in modo che i fuochi abbiano coordinate $F_1(-c, 0)$ e $F_2(c, 0)$, tenendo conto che è $0 < c < a$. Precisamente dimostriamo che l'ellisse che ha fuochi F_1 e F_2 ha equazione

$$\frac{x^2}{a^2} + \frac{y^2}{b^2} = 1. \quad (13.1)$$

dove $b = \sqrt{a^2 - c^2}$.

Sia $P(x, y)$; allora deve essere

$$2a = d(PF_1) + d(PF_2) = \sqrt{(x+c)^2 + y^2} + \sqrt{(x-c)^2 + y^2}.$$

Da cui

$$\begin{aligned}(x+c)^2 + y^2 &= \left(2a - \sqrt{(x-c)^2 + y^2}\right)^2 \\ &= 4a^2 + (x-c)^2 + y^2 - 4a\sqrt{(x-c)^2 + y^2}\end{aligned}$$

e quindi

$$4a\sqrt{(x-c)^2 + y^2} = 4(a^2 - cx)$$

da cui

$$\begin{aligned}a^2(x^2 - 2xc + c^2 + y^2) &= (a^2 - cx)^2 = a^4 - 2a^2xc + c^2x^2 \\ (a^2 - c^2)x^2 + a^2y^2 &= a^2(a^2 - c^2) \iff \frac{x^2}{a^2} + \frac{y^2}{a^2 - c^2} = 1\end{aligned}$$

e la (13.1) segue ponendo $b^2 = a^2 - c^2$.

OSSERVAZIONE 13.1. Per come è stato definito segue naturalmente che $b < a$; se $b = a$ si ottiene una circonferenza e se $b > a$ si scambiano il ruolo della x e y e dunque si ottiene l'ellisse di fuochi $(0, \pm c)$.

Per esempio vogliamo determinare i fuochi dell'ellisse di equazione

$$4x^2 + y^2 = 1. \quad (13.2)$$

Possiamo scrivere la (13.2) nella forma

$$\frac{x^2}{\left(\frac{1}{2}\right)^2} + y^2 = 1$$

da cui si ricava immediatamente che $a = \frac{1}{2}$ e $b = 1$; essendo $b > a$ si ottiene $c = \sqrt{b^2 - a^2}$ e dunque $c = \sqrt{1 - \frac{1}{4}} = \frac{\sqrt{3}}{2}$ e quindi $F_1 \left(0, -\frac{\sqrt{3}}{2}\right)$

e $F_2 \left(0, \frac{\sqrt{3}}{2}\right)$.

Tenendo presente l'equazione (13.1) si vede subito che si può porre $\frac{x}{a} = \cos \varphi$ e $\frac{y}{b} = \sin \varphi$ da cui si ottengono le comode *equazioni parametriche* dell'ellisse riferita ai propri assi di simmetria

$$\begin{cases} x = a \cos \varphi \\ y = b \sin \varphi \end{cases} \quad \varphi_0 \leq \varphi \leq \varphi_0 + 2\pi. \quad (13.3)$$

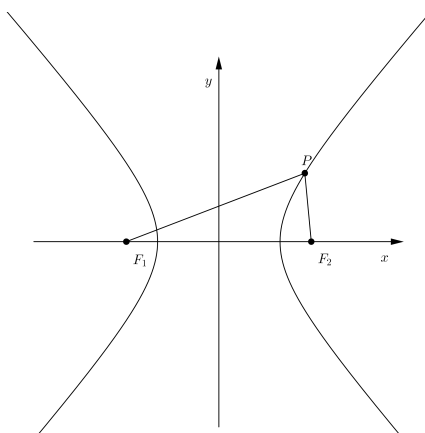


Figura 13.3 L'iperbole

DEFINIZIONE 13.2. Dati nel piano due punti F_1 e F_2 si chiama *iperbole* di fuochi F_1 e F_2 l'insieme dei punti P del piano tali che è costante il valore assoluto della differenza delle distanze di P da F_1 e F_2

$$|d(PF_1) - d(PF_2)| = 2a$$

dove a è una costante che soddisfa la relazione $2a < d(F_1F_2)$.

Se $F_1(-c, 0)$ e $F_2(c, 0)$ si ottiene, con calcoli del tutto analoghi a quelli svolti precedentemente – e che lasciamo per esercizio al lettore – che l'iperbole avente fuochi in F_1 e F_2 ha equazione

$$\frac{x^2}{a^2} - \frac{y^2}{b^2} = 1 \quad (13.4)$$

dove $c = \sqrt{a^2 + b^2}$.

Si osserva anche che per grandi valori di $|x|$, l'iperbole "si avvicina" alle due rette di equazioni $ay = \pm bx$, infatti dalla (13.4) si ha

$$ay = \pm \sqrt{-a^2b^2 + b^2x^2},$$

da cui, se $y > 0$

$$\begin{aligned} ay - b|x| &= \sqrt{-a^2b^2 + b^2x^2} - b|x| \\ &= \frac{(\sqrt{-a^2b^2 + b^2x^2} + b|x|)(\sqrt{-a^2b^2 + b^2x^2} - b|x|)}{\sqrt{-a^2b^2 + b^2x^2} + b|x|} \\ &= \frac{-a^2b^2}{b|x| + \sqrt{-a^2b^2 + b^2x^2}}. \end{aligned}$$

questa è una quantità che diventa sempre più piccola al crescere di $|x|$. Le rette di equazioni $y = \pm \frac{b}{a}x$ si chiamano *asintoti* dell'iperbole.³ L'iperbole si chiama *equilatera* se gli asintoti sono perpendicolari.

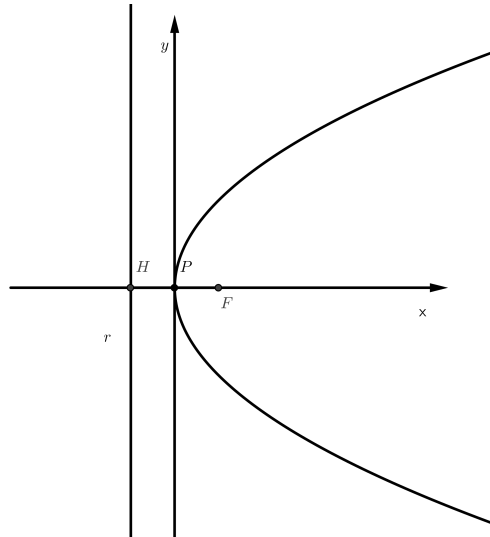


Figura 13.4 La parabola

Consideriamo ora la *parabola*.

DEFINIZIONE 13.3. Dati nel piano un punto F ed una retta r tali che $F \notin r$, una *parabola* di fuoco F e direttrice r è l'insieme dei punti P del piano equidistanti da F e da r cioè:

$$d(FP) = d(Pr).$$

Per trovarne l'equazione canonica (v. Fig. 13.4) sia $p > 0$, se $F\left(\frac{p}{2}, 0\right)$ allora r ha equazione $x = -\frac{p}{2}$ e l'equazione della parabola è:

$$y^2 = 2px \quad (13.5)$$

infatti

$$d(PF) = \sqrt{\left(x - \frac{p}{2}\right)^2 + y^2}$$

e

$$d(P, r) = \left|x + \frac{p}{2}\right|$$

³Il concetto generale di asintoto di una curva è stato chiarito nei corsi di Analisi.

da cui

$$x^2 + px + \frac{p^2}{4} = x^2 - px + \frac{p^2}{4} + y^2$$

che, semplificata, è la (13.5).

Dalle equazioni che abbiamo trovato notiamo che l'ellisse e l'iperbole sono curve simmetriche: esse posseggono due assi di simmetria tra loro ortogonali che nel nostro caso coincidono con gli assi del sistema di riferimento, quindi hanno anche un centro di simmetria, che coincide con il punto di incontro degli assi e che, in forma canonica, è l'origine del sistema di riferimento; la parabola, invece, ha un solo asse di simmetria che, in forma canonica, coincide con l'asse x e quindi non ha un centro di simmetria.

Quelle che abbiamo esaminato sono le cosiddette *equazioni canoniche* delle coniche, cioè quelle in cui appunto gli assi di simmetria delle coniche coincidono con gli assi coordinati, per le coniche a centro e con l'asse x per la parabola. Se ciò non accade la forma dell'equazione può essere molto diversa.

13.1 Coniche in forma generale

In generale l'equazione di una conica è una generica equazione di secondo grado, quindi ha la forma

$$ax^2 + bxy + cy^2 + dx + ey + f = 0 \quad (13.6)$$

con a, b, c non tutti nulli; oppure, in forma matriciale,

$$\vec{x}A\vec{x}_T = 0$$

dove $\vec{x}$ è il vettore $\begin{bmatrix} x & y & 1 \end{bmatrix}$ ed A è la matrice simmetrica

$$A = \begin{bmatrix} a & \frac{b}{2} & \frac{d}{2} \\ \frac{b}{2} & c & \frac{e}{2} \\ \frac{d}{2} & \frac{e}{2} & f \end{bmatrix}$$

OSSERVAZIONE 13.2. Tra le equazioni della forma (13.6) dobbiamo accettare anche equazioni del tipo $x^2 + y^2 = 0$ (circonferenza che ha un solo punto reale, già vista) o $x^2 + 2y^2 + 1 = 0$ (ellisse completamente immaginaria) oppure $x^2 - 2y^2 = 0$ spezzata nelle due rette reali $x + \sqrt{2}y = 0$ e $x - \sqrt{2}y = 0$ ed altre "stranezze" del genere. Quindi, per completezza, dobbiamo chiamare coniche anche curve a punti di coordinate complesse o curve *spezzate* in coppie di rette, queste ultime prendono anche il nome di *coniche degeneri*.

13.2 Riconoscimento di una conica

Sorge allora il problema di “riconoscere” la conica, cioè di sapere se l’equazione 13.6 a fronte rappresenti un’ellisse, un’iperbole o una parabola, degenerare o no.

Il problema del riconoscimento di una conica si può affrontare in vari modi; per i nostri scopi possiamo notare subito che la (13.6) rappresenta una circonferenza se e solo se $b = 0$ e $a = c$. Inoltre, nel caso generale, si dimostra che mediante un opportuno cambiamento di sistema di riferimento l’equazione 13.6 nella pagina precedente) si può portare in una delle tre forme canoniche (13.1 a pagina 126), (13.4 a pagina 128) e (13.5 a pagina 129) che non contengono il termine “rettangolare”⁴.

Se il polinomio a primo membro della (13.6) si scompone in fattori lineari, la conica è detta *degenere* e spezzata in due rette (reali o immaginarie, coincidenti o no). Si verifica facilmente, con passaggi elementari ma un po’ laboriosi, che una conica degenerare rimane tale in qualunque sistema di riferimento cartesiano ortogonale; quindi l’essere degenerare è un carattere *invariante* rispetto ad una qualsiasi rototraslazione di assi. Un’altra caratteristica invariante di una conica è la sua natura, cioè il fatto di essere un’ellisse piuttosto che una parabola od un’iperbole, equilatera o no.

Questi caratteri invarianti si traducono in termini algebrici esaminando la matrice A vista nel paragrafo precedente, e la sua sottomatrice formata dalle prime due righe e dalle prime due colonne: $B = \begin{bmatrix} a & \frac{b}{2} \\ \frac{b}{2} & c \end{bmatrix}$. Si può infatti dimostrare che

Teorema 13.1. *Una conica è degenerare se e solo se $I_3 = \det A = 0$; la conica è una parabola se $I_2 = \det B = 0$; è un’ellisse se $I_2 > 0$ ed è un’iperbole se $I_2 < 0$. In particolare se la traccia di B cioè $I_1 = a + c = 0$ è nulla si tratta di una iperbole equilatera.*

OSSERVAZIONE 13.3. Dal teorema 13.1 si vede dunque che la natura di una conica è completamente determinata solo dai coefficienti dei termini di secondo grado della sua equazione ed in particolare che la conica è una parabola se e solo se il complesso dei termini di secondo grado è il quadrato di un opportuno binomio.

⁴In realtà le equazioni canoniche dell’ellisse e dell’iperbole non contengono nemmeno termini lineari, ma questi ultimi si possono facilmente eliminare con una traslazione degli assi.

OSSERVAZIONE 13.4. Osserviamo anche che la parabola degenera in due rette reali parallele (eventualmente sovrapposte⁵: per esempio quelle rappresentate dall'equazione $x^2 = 0$); l'ellisse degenera in due rette incidenti entrambe prive di punti reali (tranne il loro punto di intersezione) infine l'iperbole degenera è costituita da due rette reali incidenti in un punto, che sono perpendicolari se e solo se l'iperbole è equilatera.

Si può anche dimostrare che con un'opportuna rototraslazione di assi l'equazione generica di una conica (13.6 a pagina 130) si può sempre portare in una ed una sola delle forme canoniche elencate nella tabella 13.1.

Tabella 13.1 Le forme canoniche dell'equazione di una conica

$\frac{x^2}{a^2} + \frac{y^2}{a^2} = 1$	(ellisse reale)
$\frac{x^2}{a^2} + \frac{y^2}{a^2} = -1$	(ellisse immaginaria)
$\frac{x^2}{a^2} + \frac{y^2}{a^2} = 0$	(ellisse degenera)
$\frac{x^2}{a^2} - \frac{y^2}{a^2} = 1$	(iperbole non degenera)
$\frac{x^2}{a^2} - \frac{y^2}{a^2} = 0$	(iperbole degenera)
$y^2 = 2px$	(parabola non degenera)
$y^2 = 0$	(parabola degenera)

In riferimento alla tabella 13.1, notiamo che:

- i) Nei due casi dell'ellisse, se $a = b$ si ha una circonferenza, rispettivamente reale o immaginaria.
- ii) Nell'ellisse immaginaria se $a = b$ si ha la circonferenza di raggio nullo, degenera in due rette immaginarie: di equazioni $x \pm iy = 0$ che si chiamano *rette isotrope*.
- iii) Nell'equazione dell'iperbole se $a = b$ si ha l'iperbole equilatera.

Osserviamo anche che operare la rotazione che riduce a forma canonica l'equazione di una conica equivale a diagonalizzare ortogonalmente la matrice A , il che è sempre possibile, essendo A simmetrica (vedi Teorema 9.8 a pagina 89).

⁵in questo caso si parla spesso di *una retta contata due volte*.

Esempio 13.1. Sia da riconoscere la conica

$$x^2 - 4xy + 4y^2 + x - 2y = 0. \quad (13.7)$$

Si vede subito che la (13.7) si può scrivere come $(x - 2y)^2 + x - 2y = 0$, quindi, poiché il complesso dei termini di secondo grado è un quadrato, si tratta di una parabola; inoltre, raccogliendo opportunamente la (13.7) si scrive anche $(x - 2y)(x - 2y + 1) = 0$ dunque è degenera. del resto si vede anche

subito che la matrice $A = \begin{bmatrix} 1 & -2 & \frac{1}{2} \\ -2 & 4 & -1 \\ \frac{1}{2} & -1 & 0 \end{bmatrix}$ è singolare.

Esempio 13.2. Vogliamo riconoscere la conica di equazione

$$x^2 - 2xy + 6y^2 - 2x = 0.$$

La matrice dei coefficienti è $A = \begin{bmatrix} 1 & -1 & -1 \\ -1 & 6 & 0 \\ -1 & 0 & 0 \end{bmatrix}$ che non è singolare. Si

vede subito che la sottomatrice $B = \begin{bmatrix} 1 & -1 \\ -1 & 6 \end{bmatrix}$ ha determinante uguale a 5 quindi positivo: si tratta dunque di un'ellisse non degenera (vedi Figura 13.5).

13.3 Tangenti ad una conica in forma canonica

Sia γ una conica. Una retta r ha in comune con γ al massimo due punti, infatti il sistema formato dalle equazioni della retta e della conica ammette al più due soluzioni, in quanto la risolvente del sistema è al più di secondo grado. Questi due punti, però, possono essere distinti, reali o complessi oppure coincidenti. In questo caso diciamo che la retta r è tangente alla conica γ .

Nel caso delle equazioni in forma canonica si verifica agevolmente che se $P(x_0, y_0)$ è un punto

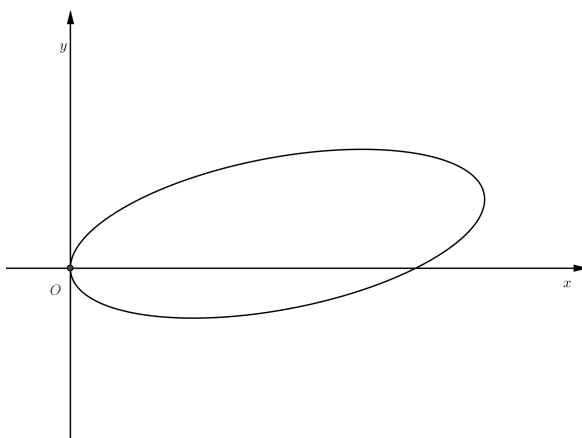


Figura 13.5 Esempio 13.2

appartenente alla conica l'equazione della tangente in P alla conica è, per le coniche reali a centro,

$$\frac{xx_0}{a^2} \pm \frac{yy_0}{b^2} = 1 \quad (13.8)$$

ovviamente con il segno $+$ se si tratta di un'ellisse e con il segno $-$ se si tratta di un'iperbole. ed invece per la parabola *in forma canonica* l'equazione della tangente in P è:

$$yy_0 = p(x + x_0). \quad (13.9)$$

ATTENZIONE Le (13.8) e (13.9) rappresentano le tangenti *solo se* P appartiene alla conica e quest'ultima è scritta in forma canonica.

13.4 Conica per cinque punti

L'equazione generica della conica (13.6 a pagina 130) è un'equazione che dipende da sei coefficienti omogenei⁶; quindi imporre il passaggio per un punto dà luogo ad una equazione in sei variabili. Dunque il passaggio per cinque punti distinti $P_i(x_i, y_i)$, $i = 1 \dots 5$ dà luogo al sistema

$$\begin{cases} ax_1^2 + bx_1y_1 + cy_1^2 + dx_1 + ey_1 + f_1 = 0 \\ ax_2^2 + bx_2y_2 + cy_2^2 + dx_2 + ey_2 + f_2 = 0 \\ ax_3^2 + bx_3y_3 + cy_3^2 + dx_3 + ey_3 + f_3 = 0 \\ ax_4^2 + bx_4y_4 + cy_4^2 + dx_4 + ey_4 + f_4 = 0 \\ ax_5^2 + bx_5y_5 + cy_5^2 + dx_5 + ey_5 + f_5 = 0 \end{cases}$$

che è lineare omogeneo di 5 equazioni nelle 6 incognite a, b, c, d, e ed f . Allora se il rango della matrice dei coefficienti è massimo, cioè, detta A la matrice dei coefficienti, se $r(A) = 5$ il sistema ammette ∞^1 soluzioni, cioè infinite soluzioni che differiscono solo di un fattore di proporzionalità. Quanto qui esposto si traduce nel

Teorema 13.2. *Per cinque punti, a tre a tre non allineati, passa una ed una sola conica non degenera.*

Dimostrazione. Se i cinque punti sono a tre a tre non allineati le equazioni sono indipendenti quindi il sistema ammette ∞^1 soluzioni, infatti, se così non fosse, cioè se una delle equazioni fosse combinazione lineare delle altre quattro, si avrebbe che tutte le coniche che passano per quattro punti $A B C D$ passerebbero anche per il quinto E , il che è assurdo,

⁶cioè, ricordiamo, definiti a meno di un fattore di proporzionalità non nullo.

perché la conica che si spezza nelle rette AB e CD dovrebbe passare per E che per ipotesi non può appartenere nè alla retta AB nè alla CD . L'unica conica che passa per i cinque punti è anche irriducibile, perché se così non fosse almeno tre dei cinque punti sarebbero allineati. $\square$

Le condizioni poste non vietano che due dei cinque punti coincidano. In tal caso la conica è tangente ad una retta passante per i due punti coincidenti (e per nessuno dei rimanenti). Di più le coppie di punti coincidenti possono essere due, in tal caso la conica sarà tangente a due rette che passano ciascuna per una delle coppie di punti coincidenti.

13.5 Le coniche in coordinate omogenee

In coordinate omogenee, cioè lavorando nel piano ampliato con gli elementi impropri, l'equazione (13.6 a pagina 130) si scrive

$$ax^2 + bxy + cy^2 + dxu + eyu + fu^2 = 0 \quad (13.10)$$

Per riconoscere la natura di una conica è più elegante studiarne il comportamento all'infinito, intersecandola con la retta impropria. Infatti segue dalle considerazioni svolte sin qui, che una retta ed una conica hanno sempre due punti in comune, distinti o coincidenti, reali o meno, propri o impropri.

Una rototraslazione di assi, che, come abbiamo visto, non altera la natura di una conica, manda punti propri in punti propri e punti impropri in punti impropri, quindi la natura di una conica equivale al suo comportamento all'infinito, che possiamo esaminare sulle equazioni canoniche.

Per $\frac{x^2}{a^2} \pm \frac{y^2}{b^2} = u^2$ i punti impropri sono quelli delle due rette $\frac{x^2}{a^2} \pm \frac{y^2}{b^2} = 0$ e cioè, rispettivamente $\left(1 : \pm i \frac{b}{a} : 0\right)$ per l'ellisse e $\left(1 : \pm \frac{b}{a} : 0\right)$ per l'iperbole, dunque *l'iperbole ha due punti impropri reali e distinti: quelli dei suoi asintoti mentre i punti impropri dell'ellisse sono immaginari.*

Per quanto riguarda l'equazione $y^2 = 2pxu$ della parabola si ha $y^2 = 0$ che equivale all'asse x contato due volte, quindi la parabola è tangente alla retta impropria nel punto $X_\infty(1 : 0 : 0)$. Nel caso generale l'intersezione tra una parabola e la retta impropria produce il sistema

$$\begin{cases} ax^2 + bxu + cy^2 = 0 \\ u = 0 \end{cases}$$

in cui la prima equazione è il quadrato di un binomio ($\alpha x + \beta y$) e quindi la parabola avrà il punto improprio ($\beta : -\alpha : 0$).

Dunque possiamo concludere che:

L'ellisse non ha punti impropri reali.

La parabola ha due punti impropri reali coincidenti, cioè è tangente alla retta impropria nel suo punto improprio.

L'iperbole ha due punti impropri reali e distinti, in direzione ortogonale se e solo se essa è equilatera.

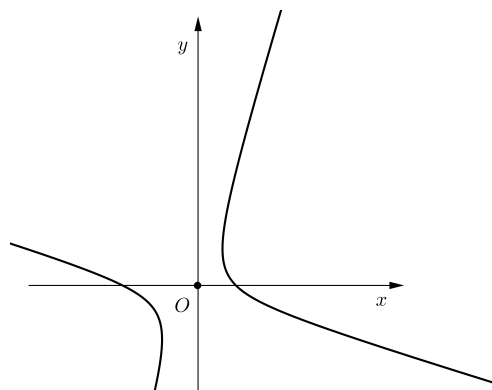


Figura 13.6 Esempio 13.3

Esempio 13.3. Vogliamo riconoscere la natura della conica γ che in coordinate non omogenee ha equazione $x^2 + 3xy - y^2 + x - 2 = 0$. Se passiamo a coordinate omogenee e intersechiamo la γ con la retta impropria $u = 0$ otteniamo il sistema

$$\begin{cases} x^2 + 3xy - y^2 + xu - 2u^2 = 0 \\ u = 0 \end{cases} \text{ equivalente a } \begin{cases} x^2 + 3xy - y^2 = 0 \\ u = 0 \end{cases}$$

Dividendo la prima equazione per x^2 otteniamo i coefficienti angolari delle rette in cui è spezzata la conica del secondo sistema: $m^2 - 3m - 1 = 0$ che ammette le due radici reali e distinte $m_{12} = \frac{3 \pm \sqrt{13}}{2}$, quindi la conica, avendo i due punti impropri reali e distinti $P_\infty(2 : 3 + \sqrt{13} : 0)$ e $Q_\infty(2 : 3 - \sqrt{13} : 0)$ è un'iperbole; di più poiché i due punti impropri sono in direzioni ortogonali, si tratta di un'iperbole equilatera. (Vedi figura 13.6)

Esempio 13.4. Riconosciamo la conica $x^2 + xy + 2y^2 + x - 1 = 0$. Passando a coordinate omogenee ed intersecando con la retta impropria, abbiamo $x^2 + xy + 2y^2 = 0$ da cui $1 + m + 2m^2 = 0$ che ha radici complesse, quindi la conica è un'ellisse, come si vede dalla figura 13.7 a fronte)

13.6 Fasci di coniche

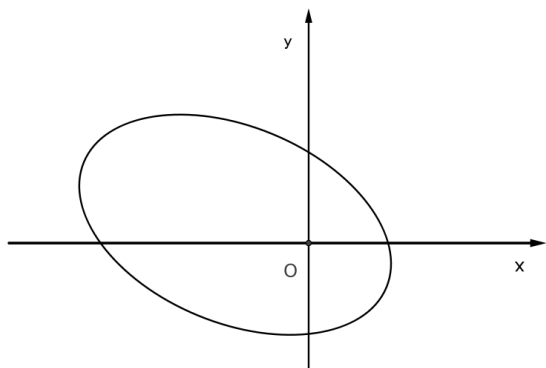


Figura 13.7 Esempio 13.4

In questo paragrafo estenderemo al caso delle coniche il concetto di fascio già visto per le rette e le circonferenze; in generale, infatti, un fascio è un insieme di enti geometrici le cui equazioni dipendono linearmente da una coppia di parametri omogenei o da un solo parametro non omogeneo.

Generalità sui fasci di coniche

DEFINIZIONE 13.4. Chiamiamo *fascio di coniche* generato da due coniche γ_1 e γ_2 , che si incontrano in quattro punti (reali o no, propri o impropri, a tre a tre non allineati ma eventualmente a due a due coincidenti) l'insieme $\mathcal{F}$ di tutte le coniche la cui equazione si ottiene come combinazione lineare non banale delle equazioni di γ_1 e γ_2 . In tale situazione chiamiamo *punti base del fascio* i quattro punti comuni a γ_1 e γ_2 .

Si verifica facilmente che *tutte e sole* le coniche di $\mathcal{F}$ passano per tutti e quattro i punti base.

Per quattro punti a tre a tre non allineati passano sei rette, che, opportunamente considerate a due a due, formano le (uniche) tre coniche degeneri di $\mathcal{F}$. Osserviamo inoltre che imporre ad una conica di passare per quattro punti non allineati vuol dire fornire quattro condizioni lineari indipendenti.

Per scrivere l'equazione del fascio di coniche che ha come punti base i punti A, B, C e D basta combinare linearmente le equazioni di due *qualsiasi* di queste coniche; in generale le più comode sono proprio quelle degeneri; esse sono tre: quella spezzata nella coppia di rette AB e CD che indicheremo con $AB \cdot CD$, la $AC \cdot BD$ e la $AD \cdot BC$; possiamo usare due di queste per scrivere l'equazione del fascio.

Se due dei quattro punti base coincidono si ha una condizione di tangenza in un punto: è il caso del fascio di coniche tangenti in un punto ad una data retta e passanti per altri due punti.

Ci possono essere anche due coppie di punti coincidenti: in tal caso si parla di fascio di coniche bitangenti, cioè di coniche tangenti in un punto P ad una retta r ed in un punto Q ad una seconda retta s .

Se le due coniche da cui partiamo per costruire il fascio hanno rispettivamente equazioni:

$$\begin{aligned}\gamma_1 &\equiv a_1x^2 + b_1xy + c_1y^2 + d_1x + e_1y + f_1 = 0 \\ \gamma_2 &\equiv a_2x^2 + b_2xy + c_2y^2 + d_2x + e_2y + f_2 = 0\end{aligned}$$

L'equazione del fascio sarà:

$$\lambda(a_1x^2 + b_1xy + c_1y^2 + d_1x + e_1y + f_1) + \mu(a_2x^2 + b_2xy + c_2y^2 + d_2x + e_2y + f_2) = 0 \quad (13.11)$$

L'equazione (13.11) può assumere anche la forma

$$\begin{aligned}(\lambda a_1 + \mu a_2)x^2 + (\lambda b_1 + \mu b_2)xy + (\lambda c_1 + \mu c_2)y^2 + \\ + (\lambda d_1 + \mu d_2)x + (\lambda e_1 + \mu e_2)y + \lambda f_1 + \mu f_2 = 0\end{aligned} \quad (13.12)$$

Ci chiediamo quando l'equazione (13.11) rappresenta un'iperbole equilatera. Ricordiamo che una conica $ax^2 + bxy + cy^2 + dx + ey + f = 0$ è un'iperbole equilatera se e solo se $a + c = 0$ quindi nella (13.12) se e solo se $\lambda a_1 + \mu a_2 + \lambda c_1 + \mu c_2 = 0$. Questa è un'equazione lineare che può ammettere una sola soluzione, oppure essere identicamente verificata (se $a_1 + c_1 = a_2 + c_2 = 0$) o non ammettere soluzioni (se, ad esempio, $a_1 + c_1 \neq 0$, $a_2 + c_2 = 0$); quindi in un fascio di coniche possono esserci esattamente una iperbole equilatera, solo iperboli equilateri oppure nessuna iperbole equilatera.

Analogamente possiamo dire che la (13.12) rappresenta una parabola se e solo se

$$(\lambda b_1 + \mu b_2)^2 + 4(\lambda a_1 + \mu a_2)(\lambda c_1 + \mu c_2) = 0 \quad (13.13)$$

L'equazione (13.13) può ammettere esattamente due soluzioni reali, distinte o coincidenti, e allora nel fascio ci sono due parabole distinte o coincidenti⁷ oppure essere identicamente soddisfatta, e allora si tratta di un fascio di sole parabole o impossibile, e allora si tratta di un fascio di coniche che non contiene parabole.

Esempio 13.5. Cerchiamo se esistono parabole passanti per i punti comuni alle coniche $\gamma_1 \equiv x^2 + y^2 = 4$ e $\gamma_2 \equiv x^2 - 2xy - y^2 + 2x = 0$ (circonferenza

⁷eventualmente degeneri in due rette parallele o sovrapposte.

e iperbole equilatera non degenera). Il fascio individuato dalle due coniche ha equazione $\lambda(x^2 + y^2 - 4) + \mu(x^2 - 2xy - y^2 + 2x) = 0$ che si può anche scrivere come $x^2 - 2xy - y^2 + 2x + k(x^2 + y^2 - 4) = 0$ o anche come

$$(k+1)x^2 - 2xy + (k-1)y^2 + 2x - 4k = 0$$

Per avere una parabola dovrà essere $(k+1)(k-1) - 1 = 0$ da cui $k = \pm\sqrt{2}$. Quindi si hanno due parabole, di equazioni rispettive

$$(1 + \sqrt{2})x^2 - 2xy - (1 + \sqrt{2})y^2 + 2x - 4\sqrt{2} = 0$$

$$(1 - \sqrt{2})x^2 - 2xy - (1 - \sqrt{2})y^2 + 2x + 4\sqrt{2} = 0$$

13.7 Fasci e punti impropri

Le questioni riguardanti le coniche per quattro o cinque punti hanno senso anche quando uno o due dei punti sono impropri; ad esempio se si dice che una conica ammette un asintoto parallelo ad una retta r data, si intende che la conica passa per il punto improprio R_∞ della retta r ; se si conosce proprio l'equazione dell'asintoto, allora le condizioni sono due: il passaggio per il punto improprio e la tangenza ivi alla retta asintoto.

La circonferenza è una particolare conica che viene individuata dando solo tre condizioni lineari indipendenti; questo fatto si spiega considerando che, come abbiamo detto alla fine del capitolo 12 tutte le circonferenze del piano passano per due punti: i punti ciclici. Un fascio di circonferenze secanti è dunque un caso particolare di quello delle coniche per quattro punti distinti.

Per quanto riguarda l'asse radicale di un fascio di circonferenze, esaminiamo che cosa accade in coordinate omogenee intersecando due circonferenze: si ha il sistema

$$\begin{cases} x^2 + y^2 + axu + byu + cu^2 = 0 \\ x^2 + y^2 + \alpha xu + \beta yu + \gamma u^2 = 0 \end{cases}$$

Sottraendo membro a membro e raccogliendo opportunamente si ha l'equazione

$$u[(a - \alpha)x + (b - \beta)y + c - \gamma] = 0$$

che è l'equazione di una conica spezzata nella retta impropria e in un'altra retta, quella che abbiamo chiamato l'asse radicale del fascio.

Parte III

Polarità piana

Capitolo 14

Proiettività ed involuzioni

Chiamiamo *forma proiettiva di prima specie* qualunque insieme di enti geometrici la cui equazione è definita al variare di due parametri lineari omogenei od uno non omogeneo che possa però assumere anche il valore infinito: le rette, viste come insieme di punti (*punteggiate*) sono forme di prima specie, così come lo sono i fasci di rette o di coniche.

In questo capitolo ci occuperemo di particolari trasformazioni tra forme di prima specie dette proiettività e di un particolare ed importante tipo di tali trasformazioni dette involuzioni.

14.1 Proiettività

Si considerino due rette proiettive (cioè completate con il loro punto improprio e pensate come insieme di punti) r e r' . Su ciascuna di esse si può fissare un sistema di ascisse in modo che $(x : u)$ siano le coordinate omogenee del generico punto $P \in r$ e $(x' : u')$ quelle del generico punto $P' \in r'$. Chiamiamo *proiettività* tra r ed r' la corrispondenza biunivoca π che associa al punto $P(x : u)$ di r il punto P' di r' di coordinate

$$\begin{cases} \rho x' = ax + bu \\ \rho u' = cx + du \end{cases} \quad (14.1)$$

con $ac - bd \neq 0$ e $\rho \in \mathbb{R}, \rho \neq 0$. Facendo nel sistema (14.1) il rapporto membro a membro e passando dalle coordinate omogenee a quelle ordinarie si ottiene

$$x' = \frac{ax + b}{cx + d} \quad (14.2)$$

che è la forma più usata per descrivere una proiettività. Esse valgono soltanto per i punti propri, infatti il corrispondente del punto $P_\infty(1 :$

0) dalla (14.1) ha coordinate tali che sia $\begin{cases} \rho x' = a \\ \rho y' = c \end{cases}$ e se $c \neq 0$ si ha $P'(a : c)$ cioè $P'(\frac{a}{c})$ mentre se è $c = 0$ si ottiene il punto di coordinate omogenee $P'(a : 0)$ cioè il punto improprio della retta r' . Esso può anche essere determinato, usando la (14.2), come $x' = \lim_{x \rightarrow 0} \frac{ax + b}{cx + d}$. In modo analogo si osserva che il punto di r che ha come corrispondente il punto improprio di r' è $P(-d : c)$, dunque per $c \neq 0$ si ha $P(-\frac{d}{c})$ e per $c = 0$ si ha il punto improprio di r .

Un altro modo per rappresentare una proiettività si può ricavare eliminando il denominatore nella (14.2) ottenendo una relazione del tipo

$$\alpha x x' + \beta x + \gamma x' + \delta = 0 \quad (14.3)$$

con la condizione $\alpha\delta - \beta\gamma = 0$.

Se nella (14.1 nella pagina precedente) lasciamo cadere la condizione $ad - bc = 0$ si perde la biunivocità, come mostra il seguente

Esempio 14.1. Si consideri la corrispondenza $x' = \frac{2x-4}{-x+2}$ in cui, appunto, è $ad - bc = 2 \cdot 2 - (-4) \cdot (-1) = 0$. Possiamo scrivere $x' = \frac{2(x-2)}{-(x-2)}$ e quindi $x' = -2$ per ogni $x \neq 2$ mentre per $x = 2$ si ha la forma di indecisione $\frac{0}{0}$ che corrisponde alla coppia di coordinate omogenee nulle ma non rappresenta alcun punto.

Esempio 14.2. Consideriamo la proiettività $x' = \frac{x-1}{x+1}$ e calcoliamo il corrispondente di qualche punto. Ad $A(1)$ corrisponde il punto A' di ascissa $x' = \frac{1-1}{1+1} = 0$, al punto $B(-1)$ il punto di ascissa $x' = \frac{-1-1}{-1+1} = \infty$ cioè il punto improprio della retta r' ed al punto P_∞ improprio della r il punto di ascissa $x' = \lim_{x \rightarrow \infty} \frac{x-1}{x+1} = 1$

Nel caso in cui le due rette (o più in generale le due forme) coincidano, cioè, come si suol dire, la proiettività sia tra forme sovrapposte si chiamano *uniti* o *fissi* i punti che sono corrispondenti di se stessi.

Esempio 14.3. Vogliamo trovare i punti uniti della proiettività di equazione $xx' - x - 2x' = 0$. Un punto è unito se e solo se è corrispondente di se stesso, cioè se $x = x'$. Tenendo conto di questa condizione si ottiene l'equazione $x^2 - 3x = 0$ da cui i due punti uniti $O(0)$ ed $A(3)$.

Dall'esempio precedente si vede che, se nella (14.3) è $\alpha \neq 0$ la ricerca dei punti uniti si riconduce alla ricerca delle radici di un'equazione di secondo grado, quindi una proiettività può avere:

- due punti uniti distinti (*proiettività iperbolica*)
- un punto unito doppio (*proiettività parabolica*)
- nessun punto unito (*proiettività ellittica*)

Nel caso in cui sia $\alpha = 0$ si verifica facilmente che un punto unito è il punto improprio: infatti da $\beta x + \gamma x' + \delta = 0$ si ha $x' = \frac{-\beta x - \delta}{\gamma}$ (qui γ è sicuramente non nullo perché se no sarebbe $\alpha\beta - \gamma\delta = 0$ e non avremmo una proiettività) e quindi $\lim_{x \rightarrow \infty} \frac{-\beta x - \delta}{\gamma} = \infty$. L'altro punto unito sarà, se $\beta \neq -\gamma$ $x = -\frac{\delta}{\beta + \gamma}$ e se $\beta = -\gamma$ il punto improprio è un punto unito doppio.

Ogni proiettività è rappresentata da un'equazione omogenea che dipende da 4 coefficienti, quindi è perfettamente determinata quando sono date tre coppie di punti corrispondenti (uniti o meno):

Esempio 14.4. Vogliamo l'equazione della proiettività che ammette come uniti i punti $A(0)$, $B(1)$ e fa corrispondere al punto $C(2)$ il punto $C'(-1)$. Le con-

dizioni poste, sostituite nella (14.3), danno luogo al sistema
$$\begin{cases} \delta = 0 \\ \alpha + \beta + \gamma + \delta = 0 \\ -2\alpha + 2\beta - \gamma + \delta = 0 \end{cases}$$
 da cui si ricava $\delta = 0$, $\alpha = 3\beta$ e $\gamma = -4\beta$ e quindi l'equazione della proiettività è $3xx' + x - 4x' = 0$.

Finora abbiamo considerato proiettività tra rette (punteggiate), ma come abbiamo detto, si possono considerare allo stesso modo proiettività tra fasci di rette o di coniche, in questo caso le coordinate (omogenee) del singolo elemento del fascio sono i parametri che lo individuano nel fascio stesso.

Esempio 14.5. Sia $\mathcal{F}$ il fascio $\lambda x + \mu y = 0$ di rette per l'origine e $\mathcal{F}'$ $\lambda'(x-1) + \mu'y = 0$ quello delle rette per $P(1,0)$, e sia π la corrispondenza che associa ad ogni retta r di $\mathcal{F}$ la retta r' di $\mathcal{F}'$ ad essa perpendicolare. Per la condizione di perpendicolarità si deve avere $\lambda\mu' + \lambda'\mu = 0$ cioè
$$\begin{cases} \rho\lambda' = -\mu \\ \rho\mu' = \lambda \end{cases}$$
 che diventa, utilizzando nei due fasci un solo parametro non omogeneo e

scrivendo : $y = mx$ e $\mathcal{F}' : y = m'(x - 1)$ $m' = -\frac{1}{m}$ cioè $mm' - 1 = 0$. Si hanno quindi delle relazioni analoghe alle (14.1), (14.2) ed (14.3); dunque la relazione di ortogonalità definisce una proiettività tra i due fasci considerati.

Le rette, pensate come punteggiate ed i fasci (di rette o di coniche) vengono globalmente indicate, in questo contesto, come abbiamo visto, come forme di prima specie. D'ora in avanti, salvo avviso contrario, sottintendiamo di estendere a tutte le forme di prima specie le considerazioni che facciamo per le punteggiate.

14.2 Involuzioni

Se in una proiettività π tra forme di prima specie sovrapposte si ha una coppia di punti A e B tale che $\pi(A) = B \implies \pi(B) = A$ si dice che la coppia è *involutoria* o che i punti *si corrispondono in doppio modo*.

DEFINIZIONE 14.1. Una proiettività tra forme di prima specie sovrapposte in cui ogni coppia di elementi è involutoria si chiama *involuzione*.

L'equazione di una involuzione si può scrivere, in analogia con la (14.3) come

$$\alpha xx' + \beta(x + x') + \delta = 0 \quad (14.4)$$

che è un'equazione simmetrica, cioè non cambia scambiando x con x' .

Lemma 14.1. Se in una proiettività di equazione (14.3) esiste una coppia involutoria, allora si ha $\beta = \gamma$ e quindi la proiettività è in particolare una involuzione.

Dimostrazione. Infatti se a x_1 , ascissa di un punto non unito corrisponde $x_2 \neq x_1$ si ha $\alpha x_1 x_2 + \beta x_1 + \gamma x_2 + \delta = 0$, ma poiché anche x_1 corrisponde ad x_2 si avrà anche $\alpha x_2 x_1 + \beta x_2 + \gamma x_1 + \delta = 0$; sottraendo membro a membro otteniamo $\beta(x_1 - x_2) + \gamma(x_2 - x_1) = 0$ da cui $(x_1 - x_2)(\beta - \gamma) = 0$ e poiché per ipotesi $x_2 \neq x_1$, deve essere $\beta - \gamma = 0$, quindi la proiettività è un'involuzione. $\square$

Per le involuzioni sussiste il

Teorema 14.2. La proiettività di equazione (14.3) è un'involuzione se e soltanto se $\beta = \gamma$.

Dimostrazione. Se $\beta = \gamma$, risolvendo la (14.3) sia rispetto ad x che ad x' si ha

$$x' = \frac{\beta x + \delta}{\alpha x + \beta} \quad \text{e} \quad x = \frac{\beta x' + \delta}{\alpha x' + \beta}$$

quindi ogni coppia di punti è involutoria; viceversa se la proiettività è un'involuzione, allora ogni coppia di punti è involutoria e quindi, per il lemma 14.1 si ha $\beta = \gamma$. $\square$

A questo punto è chiaro che

Corollario 14.3. *Una proiettività π tra forme di prima specie sovrapposte che ammetta una coppia involutoria è una involuzione, cioè tutte le coppie di elementi che si corrispondono in π sono involutorie*

Dimostrazione. Segue immediatamente dal Lemma 14.1 e dal Teorema 14.2. $\square$

L'equazione (14.4) dipende da tre coefficienti omogenei, quindi essa è univocamente determinata da due coppie di punti corrispondenti¹, in particolare un'involuzione è univocamente determinata dai suoi punti uniti.

Anche le involuzioni si possono classificare in base ai loro punti uniti: precisamente un'involuzione è

iperbolica se ammette due punti uniti reali e distinti;

ellittica se non ha punti uniti reali.

Non esistono involuzioni paraboliche, infatti ponendo nella 14.4 $x = x'$ si ottiene l'equazione di secondo grado $\alpha x^2 + 2\beta x + \delta = 0$ il cui discriminante $\Delta = 4\beta^2 - 4\alpha\gamma = 4(B^2 - \alpha\gamma)$ non si annulla mai perché dev'essere $\alpha\delta - \beta\gamma \neq 0$ e quindi, siccome si tratta di una involuzione $\alpha\delta - \beta^2 \neq 0$

Vediamo ora qualche esempio

Esempio 14.6. *Calcoliamo i punti uniti dell'involuzione $xx' + 1 = 0$. Si perviene all'equazione $x^2 + 1 = 0$ che non ha radici (in $\mathbb{R}$). Quindi l'involuzione non ha punti uniti reali ed è dunque ellittica.*

Esempio 14.7. *Troviamo l'equazione dell'involuzione che ha come punti uniti $A(0)$ e $B(1)$. Le coordinate dei punti uniti sono soluzioni dell'equazione $x^2 - x = 0$ da cui si ottiene $xx' - \frac{1}{2}(x + x') = 0$.*

Esempio 14.8. *Sia data, in un fascio di rette, l'involuzione che fa corrispondere alla retta di coefficiente angolare $m = 3$ quella di coefficiente angolare $m' = \infty$ e che ammette come unita la retta per cui $m = 1$. Vogliamo trovare*

¹Due punti che si corrispondono in una involuzione vengono anche spesso detti *punti coniugati*.

l'altra retta unita. Dalla relazione $m' = -\frac{\beta m + \delta}{\alpha m + \beta}$, tenendo conto delle condizioni date si hanno le relazioni $3\alpha + \beta = 0$ e $1 = -\frac{\beta + \delta}{\alpha + \beta}$ che danno luogo al sistema $\begin{cases} \alpha + 2\beta = -\delta \\ 3\alpha + \beta = 0 \end{cases}$ che ha come soluzione $\begin{cases} \delta = 5\alpha \\ \beta = -3\alpha \end{cases}$. L'involuzione cercata ha dunque equazione $mm' - 3(m + m') + 5 = 0$. Ponendo $m = m'$ si ha l'equazione ai punti uniti $m^2 - 6m + 5 = 0$, le cui soluzioni sono $m = 1$, che sapevamo, ed $m = 5$, coefficiente angolare dell'altra retta unita.

Capitolo 15

Polarità piana

In questo capitolo esamineremo e studieremo una particolare corrispondenza biunivoca tra punti e rette del piano indotta da una conica.

15.1 Polare di un punto rispetto ad una conica irriducibile

Sia data una conica irriducibile¹ γ che, come già visto, può essere rappresentata, in coordinate omogenee, dall'equazione

$$f(x, y, u) = a_{11}x^2 + 2a_{12}xy + a_{22}y^2 + 2a_{13}xu + 2a_{23}yu + a_{33}u^2 = 0 \quad (15.1)$$

o anche, in notazione matriciale

$$\vec{x}A\vec{x}_T = 0 \quad (15.2)$$

dove $\vec{x} = [x, y, u]$ ed $A = [a_{ik}]$ è la matrice simmetrica dei coefficienti della (15.1) e $\det A \neq 0$ in quanto consideriamo solo coniche irriducibili.

DEFINIZIONE 15.1. Chiamiamo *polare* del punto $P(x_0 : t_0 : u_0)$ rispetto alla conica di equazione (15.1) la retta che può scriversi indifferente-

¹Da qui in poi, e per tutto il capitolo, anche se non espressamente detto, le coniche prese in esame dovranno essere considerate, salvo esplicito avviso contrario, irriducibili.

mente² in ciascuna delle due seguenti forme

$$x \left(\frac{\partial f}{\partial x} \right)_P + y \left(\frac{\partial f}{\partial y} \right)_P + u \left(\frac{\partial f}{\partial u} \right)_P = 0 \quad (15.3)$$

$$x_0 \left(\frac{\partial f}{\partial x} \right) + y_0 \left(\frac{\partial f}{\partial y} \right) + u_0 \left(\frac{\partial f}{\partial u} \right) = 0 \quad (15.4)$$

L'equazione (15.3) è quella di una retta i cui coefficienti sono le tre derivate parziali calcolate nel punto P ; si vede subito che anche la (15.4) rappresenta una retta in quanto è sicuramente un'equazione lineare e che le tre derivate parziali non si annullano contemporaneamente in alcun punto del piano. Inoltre si può facilmente dimostrare che la corrispondenza che associa ad ogni punto del piano la sua polare rispetto ad una conica γ è biunivoca.

Se la retta r è la polare del punto P il punto P si chiama *polo* della retta r rispetto alla conica γ .

Se la conica è data con l'equazione (15.2), l'equazione della polare è data da

$$\vec{x} A \vec{x}_0^T = 0 \quad \text{oppure} \quad \vec{x}_0 A \vec{x}_T \quad (15.5)$$

dove $\vec{x}_0$ è il vettore $[x_0, y_0, u_0]$; l'equivalenza delle due forme nella (15.5) deriva dal fatto che A è simmetrica, si ha infatti $\vec{x} A \vec{x}_0^T = \vec{x}_0^T A^T \vec{x}_T = \vec{x}_0 A \vec{x}_T$.

L'equivalenza delle (15.5) con le (15.3) ed (15.4) può essere ricavata con semplici calcoli, che costituiscono un utile esercizio, osservando che le righe della matrice A (e quindi anche le colonne, essendo A simmetrica) sono i coefficienti delle tre derivate parziali del polinomio a primo membro della (15.1) rispetto ad x , y e u rispettivamente.

Esempio 15.1. Vogliamo calcolare la polare del punto $P(1 : 0 : 1)$ rispetto all'iperbole di equazione $x^2 - xy + xu - u^2 = 0$. Applicando la (15.4) si ottiene $1 \cdot (2x - y + u) + 0 \cdot (-x) + 1 \cdot (x - 2u) = 0$ e quindi l'equazione della polare è $3x - 2y - 2u = 0$

Come si vede dall'esempio 15.1 è spesso più comodo usare la forma (15.4) piuttosto che la (15.3), che torna invece utile in casi come quello del seguente

²L'equivalenza delle due forme dell'equazione della polare si può facilmente verificare sostituendo semplicemente nella (15.4) le espressioni delle tre derivate parziali del polinomio di cui alla (15.1) calcolate in un punto generico.

Esempio 15.2. Vogliamo determinare il punto proprio, di ascissa unitaria, che rispetto all'ellisse irriducibile γ di equazione $f(x, y, u) = x^2 + 4y^2 - u^2 = 0$ ha la polare perpendicolare rispetto a quella del punto $P(-1, 1)$. La polare di P rispetto a γ ha equazione

$$[x, y, u] \begin{bmatrix} 1 & 0 & 0 \\ 0 & 4 & 0 \\ 0 & 0 & -1 \end{bmatrix} \begin{bmatrix} -1 \\ 1 \\ 1 \end{bmatrix} = 0$$

cioè $[x, y, u] \begin{bmatrix} -1 \\ 4 \\ -1 \end{bmatrix} = 0$ da cui $x - 4y + u = 0$; sia ora $A(1, t)$ il generico punto proprio del piano di ascissa unitaria: i coefficienti di x e y nell'equazione della sua polare sono $a = \left(\frac{\partial f}{\partial x}\right)_A$ e $b = \left(\frac{\partial f}{\partial y}\right)_A$; le due polari sono perpendicolari se e soltanto se è $1 \cdot \left(\frac{\partial f}{\partial x}\right)_A - 4 \cdot \left(\frac{\partial f}{\partial y}\right)_A = 0$ cioè $(2x)_A - 4(8y)_A = 0$ da cui $2 - 32t = 0$ e quindi $t = \frac{1}{16}$. Il punto cercato è allora $A\left(1, \frac{1}{16}\right)$.

Dare una coppia polo-polare, cioè un punto e la sua polare rispetto ad una conica, significa imporre sui coefficienti dell'equazione della conica due condizioni lineari indipendenti.

Esempio 15.3. Consideriamo le coniche che ammettono come polare del punto $X_\infty(1 : 0 : 0)$ la retta di equazione $y + 1 = 0$. Pensiamo alla conica nella forma 15.1 a pagina 149: l'equazione $\left(\frac{\partial f}{\partial x}\right) \cdot 1 + \left(\frac{\partial f}{\partial y}\right) \cdot 0 + \left(\frac{\partial f}{\partial u}\right) \cdot 0 = 0$ che diventa $a_{11}x + a_{12}y + a_{13}u = 0$ dovrà essere quella della retta data, quindi si deve avere $a_{11} = 0 = a_{12} - a_{13}$ dunque proprio due condizioni lineari indipendenti.

Si può dimostrare che se P appartiene alla conica (e soltanto in questo caso) la sua polare rispetto alla conica è la tangente in P alla conica: diventa allora molto semplice trattare molte questioni di tangenza mediante lo strumento della polarità. Ad esempio la tangente nell'origine ad una conica γ si può pensare come la polare dell'origine $(0 : 0 : 1)$ cioè il complesso dei termini lineari del polinomio che rappresenta la γ . Inoltre gli asintoti di un'iperbole sono le tangenti all'iperbole stessa nei suoi punti impropri, quindi possono essere determinati come le polari di questi ultimi.

Esempio 15.4. Vogliamo gli asintoti dell'iperbole $6x^2 - 5xy + y^2 - 2xu = 0$. I suoi punti impropri sono $P_\infty(1 : 2 : 0)$ e $Q_\infty(1 : 3 : 0)$. Gli asintoti, che

sono le polari di tali punti, avranno rispettivamente equazioni $1 \cdot (12x - 5y - 2u) + 2 \cdot (-5x + 2y) = 0$ e $1 \cdot (12x - 5y - 2u) + 3 \cdot (-5x + 2y) = 0$, cioè, in coordinate non omogenee $2x - y - 2 = 0$ e $3x - y + 2 = 0$

15.2 Principali proprietà della polarità piana

Sussiste il seguente fondamentale

Teorema 15.1 (Legge di reciprocità o di Plücker). *Se la polare del punto P rispetto ad una conica irriducibile $\gamma : f(x, y, u) = 0$ passa per il punto Q allora la polare di Q rispetto alla medesima conica passa per il punto P .*

Dimostrazione. Scriviamo l'equazione della polare di P rispetto a γ nella forma (15.3):

$$x \left(\frac{\partial f}{\partial x} \right)_P + y \left(\frac{\partial f}{\partial y} \right)_P + u \left(\frac{\partial f}{\partial u} \right)_P = 0$$

per ipotesi essa passa per $Q(x_1, y_1, u_1)$ quindi la

$$x_1 \left(\frac{\partial f}{\partial x} \right)_P + y_1 \left(\frac{\partial f}{\partial y} \right)_P + u_1 \left(\frac{\partial f}{\partial u} \right)_P = 0 \quad (15.6)$$

è un'identità. Scriviamo ora l'equazione della polare di Q nella forma (15.4):

$$x_1 \left(\frac{\partial f}{\partial x} \right) + y_1 \left(\frac{\partial f}{\partial y} \right) + u_1 \left(\frac{\partial f}{\partial u} \right) = 0$$

che, grazie alla (15.6), è soddisfatta dalle coordinate di P . $\square$

I seguenti corollari sono immediate conseguenze della legge di reciprocità (Teorema 15.1) (le semplicissime dimostrazioni sono proposte come esercizio)

Corollario 15.2. *Se il polo Q della retta q appartiene alla retta p allora il polo P di p appartiene alla retta q .*

Corollario 15.3. *Se il punto P si muove su una retta r , allora la polare di P descrive il fascio di rette che ha per sostegno il polo R di r .*

Corollario 15.4. *Se una retta r descrive un fascio che ha per sostegno il punto P allora il suo polo descrive una retta che è la polare di P .*

La legge di reciprocità ci fornisce un metodo comodo per determinare le coordinate del polo di una retta:

Esempio 15.5. Date la retta $r : x + y - 3u = 0$ e la parabola $x^2 - 2xy + y^2 - 2xu = 0$ vogliamo determinare il polo R di r . Per la legge di reciprocità le polari di due qualsiasi punti di r passano per il polo di r , quindi possiamo scegliere su r due punti comodi, per esempio il punto improprio $P_\infty(1 : -1 : 0)$ ed il punto (proprio) $Q(3 : 0 : 1)$. La polare di P_∞ ha equazione $2x - 2y - u = 0$ mentre la polare di Q ha equazione $2x - 3y - 3u = 0$. Il punto cercato, in quanto intersezione delle due polari, avrà le coordinate che sono soluzione del sistema $\begin{cases} 2x - 2y - u = 0 \\ 2x - 3y - 3u = 0 \end{cases}$ e cioè $R(-\frac{3}{2} : -2 : 1)$ che si può anche scrivere come $R(-3 : -4 : 2)$.

Dalle proprietà della polarità piana segue anche il

Teorema 15.5. Per un punto P non appartenente ad una conica γ passano esattamente due tangenti alla γ , reali o meno.

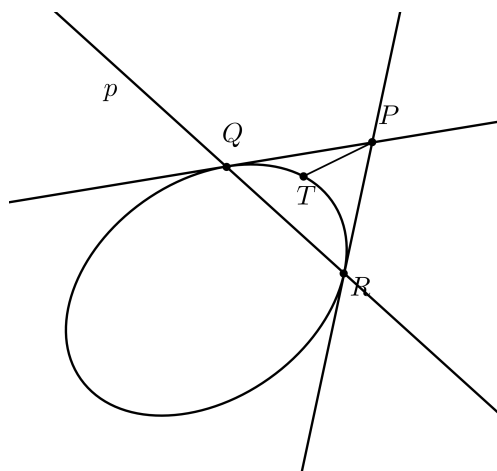


Figura 15.1 Tangenti da un punto ad una conica

Dimostrazione. Riferiamoci alla Figura 15.1. Siano Q e R le intersezioni della polare di P con la conica³, allora la polare di Q è tangente alla conica, perché Q sta sulla conica e passa per P per la legge di reciprocità. In modo analogo si può dire che la polare di R , a sua volta tangente, passa anch'essa per P . Quindi per P passano almeno due tangenti alla conica. Dimostriamo che sono solo due. Supponiamo, per assurdo che anche la retta PT sia tangente alla conica; sempre per la legge di reciprocità, il punto T deve stare sulla polare di P e sulla conica, quindi,

³Nel piano ampliato con elementi impropri ed elementi immaginari una retta ed una conica hanno sempre due punti in comune, a meno che il punto non appartenga alla conica.

ricordando che una conica ed una retta hanno in comune al massimo due punti, esso deve coincidere con Q o con R . $\square$

Dal teorema 15.5 segue ovviamente che

Teorema 15.6. *La polare rispetto ad una conica irriducibile γ di un punto P non appartenente ad essa congiunge i punti di tangenza delle tangenti passanti per P .*

Questo risultato si può sfruttare per scrivere in maniera rapida ed elegante le equazioni delle tangenti condotte da un punto ad una conica.

Esempio 15.6. *Sia γ la conica di equazione $x^2 - 2xy + x - 1 = 0$; si vogliano le tangenti alla γ uscenti dall'origine. La polare dell'origine ha, in coordinate omogenee, equazione $x - 2u = 0$; essa ha in comune con la conica il punto improprio ed il punto di coordinate $(8 : 5 : 4)$. Le tangenti cercate, che congiungono tali punti con O hanno rispettivamente equazioni $x = 0$ e $5x - 8y = 0$.*

15.3 Elementi coniugati rispetto ad una conica irriducibile

Diciamo che il punto A è *coniugato* o *reciproco* del punto B rispetto ad una conica irriducibile γ se la polare di A passa per B . Poiché, in questo caso, in virtù della legge di reciprocità (Teorema 15.1 a pagina 152) anche la polare di B passa per A possiamo dire che i due punti sono (mutuamente) reciproci rispetto a γ . In modo analogo diciamo che due rette a e b sono reciproche se il polo dell'una sta sull'altra. Dalle considerazioni precedenti risulta evidente che il luogo dei punti reciproci, rispetto ad una conica irriducibile γ , di un punto P è la polare di P rispetto a γ . Analogamente le rette reciproche di una retta a sono quelle del fascio che ha per sostegno il polo di a .

Esempio 15.7. *Vogliamo il punto reciproco dell'origine rispetto alla conica $\gamma : x^2 - y^2 - u^2 = 0$ che appartiene all'asse x . Poiché la polare dell'origine rispetto a γ è la retta impropria, il punto cercato sarà il punto improprio dell'asse x , cioè il punto $X_\infty(1 : 0 : 0)$.*

Da quanto detto segue anche che fissata una conica irriducibile γ ed una retta r non tangente a γ ,

Teorema 15.7. *I punti di r reciproci rispetto a γ si corrispondono in una involuzione, detta dei punti reciproci o dei punti coniugati i cui punti uniti sono gli eventuali punti di intersezione reali tra la retta e la conica.*

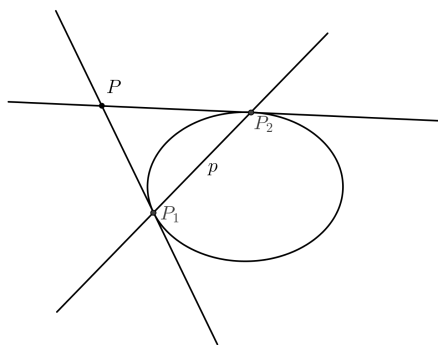


Figura 15.2 *Involuzione iperbolica, retta secante*

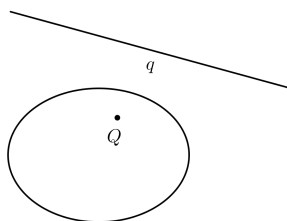


Figura 15.3 *Involuzione ellittica, retta esterna*

La natura di questa involuzione permette di distinguere le rette non tangenti alla conica in: secanti, se l'involuzione è iperbolica, ed esterne se essa è ellittica. È escluso il caso in cui la retta sia tangente in T alla conica, perché in tal caso non si avrebbe un'involuzione, in quanto ogni punto di r sarebbe reciproco di T . La situazione è illustrata nelle figure 15.2 e 15.3.

Vale anche, dualmente, il

Teorema 15.8. *Le rette di un fascio $\mathcal{F}$ avente per sostegno il punto P non appartenente ad una conica γ che siano reciproche rispetto alla stessa γ si corrispondono in una involuzione le cui rette unite sono le eventuali rette reali passanti per P e tangenti alla γ .*

Esempio 15.8. *Se per una conica γ non degenera l'involuzione delle rette reciproche nel fascio che ha per sostegno l'origine $O(0,0)$ ha equazione $mm' + (m + m') = 0$, allora le tangenti condotte da O alla conica, che sono le rette unite dell'involuzione, si possono determinare come le rette i cui coefficienti angolari sono soluzioni dell'equazione $m^2 - 2m = 0$, cioè $m = 0$ ed $m = 2$ e sono quindi le rette $y = 0$ e $y = 2x$.*

Dimostriamo ora il

Teorema 15.9. *Una conica γ è il luogo dei punti che appartengono alla propria polare, cioè sono autoconiugati rispetto alla γ*

Dimostrazione. Se $P \in \gamma$, la sua polare rispetto alla conica è la tangente in P alla conica, quindi passa per P .

Viceversa supponiamo che $P(x_0 : y_0 : u_0)$ appartenga alla propria polare rispetto alla conica $\gamma : f(x, y, u) = 0$. Sostituendo le coordinate di P nell'equazione della polare, si deve avere

$$x_0 \left(\frac{\partial f}{\partial x} \right)_P + y_0 \left(\frac{\partial f}{\partial y} \right)_P + u_0 \left(\frac{\partial f}{\partial u} \right)_P = 0$$

Ma essendo $f(x, y, u)$ omogenea, dal Teorema di Eulero⁴ sulle funzioni omogenee segue che $x \left(\frac{\partial f}{\partial x} \right) + y \left(\frac{\partial f}{\partial y} \right) + u \left(\frac{\partial f}{\partial u} \right) = 2f(x, y, u)$ e quindi $f(x_0, y_0, u_0) = 0$. $\square$

15.4 Triangoli autopolari

DEFINIZIONE 15.2. Si chiama *autopolare* per una conica irriducibile γ un triangolo⁵ tale che ogni vertice sia il polo del lato opposto.

Nella figura 15.4 a fronte il triangolo PQR è autopolare per l'ellisse.

Tenendo conto del fatto che la polare di un punto esterno ad una conica taglia la conica stessa in due punti reali e distinti, mentre quella di un punto interno non ha punti reali in comune con la conica, è facile verificare che un triangolo autopolare ha sempre esattamente un vertice interno alla conica e due esterni ad essa, inoltre può accadere che uno o due dei vertici di un triangolo autopolare siano punti impropri. Nel primo caso, il lato opposto al vertice improprio è un diametro della conica, nel secondo il vertice proprio è il centro della stessa, come preciseremo meglio nel prossimo capitolo.

L'insieme di tutte le coniche che ammettono un certo triangolo come autopolare è costituito da ∞^2 coniche e come tale prende il nome di *rete* di coniche, questo equivale a dire che dare un triangolo autopolare per una conica equivale a dare tre condizioni lineari indipendenti, infatti, pur essendo date tre coppie polo–polare (quindi sei condizioni lineari) per la legge di reciprocità (Teorema 15.1 a pagina 152) le condizioni *indipendenti* sono solo tre.

⁴Noto teorema di Analisi sulle funzioni omogenee.

⁵In Geometria proiettiva un triangolo è una figura piana costituita da tre *rette* non concorrenti, (cioè non passanti tutte tre per un medesimo punto) e *non da tre segmenti* come in Geometria elementare. Le tre rette si dicono *lati* del triangolo ed i tre punti (*propri* od *impropri che siano*) in cui le rette si intersecano a due a due sono chiamati *vertici* del triangolo.

Si può dimostrare che l'equazione complessiva della rete di coniche che ammetta come autopolare il triangolo di vertici A , B e C si ottiene come combinazione lineare dei quadrati delle equazioni dei tre lati, cioè, in forma simbolica,

$$\lambda AB^2 + \mu AC^2 + \nu BC^2 = 0 \quad (15.7)$$

Anche in questo caso, invece di usare tre parametri omogenei, se ne possono usare due non omogenei, ad esempio $h = \frac{\lambda}{\nu}$ e $k = \frac{\mu}{\nu}$, perdendo così le coniche che si ottengono dalla (15.7) per $\nu = 0$: in questo caso, tuttavia, non è necessario pensare di riammetterle per il valore $\nu = \infty$ in quanto le questioni relative alla polarità riguardano solo coniche irriducibili e dunque nella (15.7) non può annullarsi nessuno dei tre coefficienti. Si dimostra che in realtà vale anche il viceversa, cioè nessuna delle coniche della rete considerata individuata dai valori non nulli dei tre parametri è degenera.

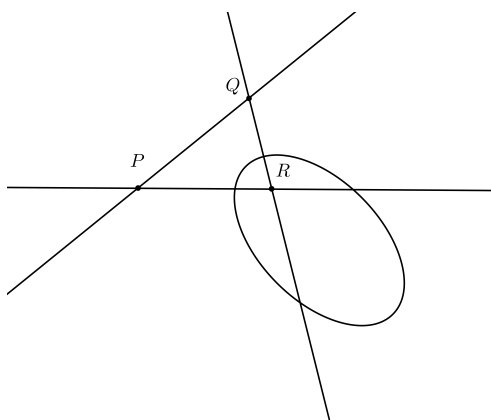


Figura 15.4 Triangolo autopolare

Esempio 15.9. Vogliamo determinare l'equazione dell'iperbole che ammette come autopolare il triangolo formato dagli assi coordinati e dalla retta di equazione $x + 2y - 1 = 0$ ed ammette come asintoto la retta $x + 2y + 2 = 0$.

Abbiamo tre condizioni fornite dal triangolo autopolare e due dall'asintoto (tangente nel punto improprio, quindi coppia polo-polare). La rete di coniche che ammette il triangolo dato come autopolare ha equazione $\lambda x^2 + \mu y^2 + \nu(x + 2y - 1)^2 = 0$; la polare di $P_\infty(2 : -1 : 0)$, punto improprio dell'asintoto, rispetto alla generica conica della rete ha equazione $(\lambda + 2\nu)x + (\mu + 4\nu)y - 4\nu = 0$, che coincide con la retta data quando $\lambda + 2\mu = \frac{\mu + 4\nu}{2} = \frac{-4\nu}{2}$ cioè quando $\lambda = -4\nu$ e $\mu = -8\nu$. La conica ha dunque equazione $-4x^2 - 8y^2 + (x + 2y - 1)^2 = 0$ che diventa facilmente $3x^2 - 4xy + 4y^2 + 2x + 4y - 1 = 0$.

Nella discussione del precedente esempio abbiamo usato tre parametri omogenei invece di due non omogenei per ricordare, ancora una volta, come due rette coincidano quando le due equazioni siano individuate da due polinomi non necessariamente uguali ma anche solo proporzionali.

Capitolo 16

Centro, diametri ed assi di una conica irriducibile

In questo capitolo vediamo come le proprietà della polarità piana viste nel capitolo precedente possano essere utilmente utilizzate per determinare ulteriori elementi di una conica, notevoli da un punto di vista geometrico e già accennati in modo elementare nei precedenti capitoli sulle coniche.

16.1 Centro e diametri di una conica

DEFINIZIONE 16.1. Se γ è una conica irriducibile chiamiamo *centro* di γ il polo della retta impropria rispetto a γ . Chiamiamo *diametro* di γ una qualunque retta che passi per il centro, cioè, per la legge di reciprocità la polare di un qualunque punto improprio.

La parabola, in quanto tangente alla retta impropria nel suo punto improprio, ha centro improprio: il punto improprio del suo asse. Ne segue che la parabola ha tutti i diametri paralleli, in quanto passanti tutti per il medesimo punto improprio.

L'ellisse e l'iperbole, invece, hanno centro proprio (e sono perciò dette anche *coniche a centro*), infatti se esso fosse improprio, sarebbe autoconiugato, visto che apparterebbe alla propria polare e dunque per il Teorema 15.9 a pagina 155 apparterebbe alla conica e la sua polare, la retta impropria, sarebbe tangente alla conica, contro l'ipotesi che sia un'ellisse od un'iperbole.

Si mostra facilmente che il polo della retta impropria è centro di simmetria per una conica a centro, cioè ogni corda AA' che passa per C è tale che C sia il punto medio di AA' . (vedi Figura 16.1 nella pagina successiva). Si può facilmente verificare questo fatto considerando

la conica in forma canonica, cioè operando una rototraslazione del sistema di riferimento, trasformazione che non altera le proprietà della polarità piana. In questo caso la conica assume equazione $a_{11}x^2 + a_{22}y^2 + a_{33} = 0$. Una qualunque retta per il centro, che qui è l'origine, ha equazione $y = mx$ da cui l'equazione $(a_{11} - a_{22}m^2)x^2 + a_{33} = 0$ che è un'equazione di secondo grado tale che la semisomma delle radici è nulla, dunque esse sono simmetriche.

Esempio 16.1. Vogliamo trovare il centro della conica $x^2 - xy + x - 3 = 0$. Basta intersecare le polari di due punti impropri qualsiasi, per esempio i più comodi sono i punti impropri degli assi coordinati $X_\infty(1 : 0 : 0)$ e $Y_\infty(0 : 1 : 0)$. Le polari sono rispettivamente le rette $2x - y + 1 = 0$ e $x = 0$; il punto comune è soluzione del sistema $\begin{cases} 2x - y + 1 = 0 \\ x = 0 \end{cases}$ e cioè il punto $C(0, 1)$.

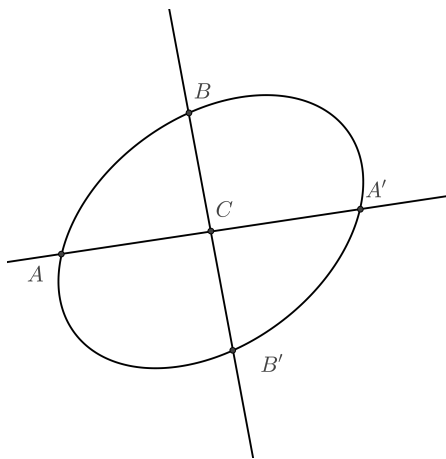


Figura 16.1 Polo della retta impropria, centro di simmetria

I diametri di una conica irriducibile formano un fascio, proprio se la conica è a centro ed improprio se la conica è una parabola. Riveste una particolare importanza, in questo fascio, l'involuzione che ad ogni diametro fa corrispondere il suo coniugato essa prende il nome di *involuzione dei diametri coniugati*: due diametri sono coniugati se e soltanto se il polo dell'uno è il punto improprio dell'altro.

Gli asintoti di una iperbole sono gli elementi uniti di tale involuzione, essendo diametri ed autoconiugati (tangenti alla conica, quindi contenenti il loro polo).

Più in generale possiamo dire che l'*involuzione dei diametri coniugati* di una conica γ è ellittica se e solo se la γ è un'ellisse ed è iperbolica se e solo se essa è un'iperbole. Quindi l'esame dell'involuzione dei diametri coniugati fornisce un ulteriore strumento per il riconoscimento di una conica. Si può inoltre dimostrare che ogni diametro di una conica a centro dimezza le corde in direzione coniugata cioè ogni diametro è asse di simmetria obliqua.

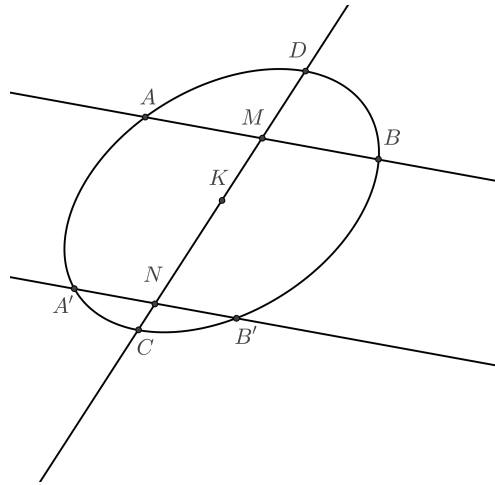


Figura 16.2 Determinazione grafica del centro di una conica

Questo risultato permette di determinare graficamente il centro di una conica ed il diametro coniugato ad un diametro dato, come mostrano i seguenti esempi e le seguenti figure.

Esempio 16.2. Vogliamo determinare graficamente il centro della conica γ di Figura 16.2. Tracciamo due rette a e b parallele che intersecano la conica, e siano A e B e, rispettivamente A' e B' i punti di intersezione delle rette con la conica. Indicati con M ed N i punti medi delle corde AB e $A'B'$ rispettivamente; tracciamo la retta MN . Essa è un diametro (coniugato alla direzione delle rette a e b); se indichiamo con C e D i punti di intersezione di questo diametro con la conica, il centro è il punto medio K della corda CD .

Esempio 16.3. Vogliamo determinare il diametro coniugato al diametro d nella conica γ di Figura 16.3 nella pagina successiva. Tracciamo una retta parallela alla d che intersechi la conica. Siano A e B le due intersezioni della retta con la conica. e sia M il punto medio di AB allora la retta CM è il diametro coniugato cercato.

Dimostriamo ora l'importante

Teorema 16.1. Sia

$$a_{11}x^2 + 2a_{12}xy + a_{22}y^2 + 2a_{13}xu + 2a_{23}yu + a_{33} = 0$$

l'equazione di una conica irriducibile γ , allora l'equazione dell'involuzione dei diametri coniugati di γ è

$$a_{22}mm' + a_{12}(m + m') + a_{11} = 0 \quad (16.1)$$

Dimostrazione. Consideriamo due qualsiasi diametri della conica e supponiamo che abbiano coefficienti angolari m e m' rispettivamente e pertanto passino per i punti impropri $P_\infty(1 : m : 0)$ e $Q_\infty(1 : m' : 0)$, essi sono coniugati se e soltanto se la polare p di $P_\infty(1 : m : 0)$ passa per $Q_\infty(1 : m' : 0)$. La p ha equazione $2a_{11}x + 2a_{12}y + 2a_{13} + m(2a_{12}x + 2a_{22}y + 2a_{23}) = 0$ che diventa $(a_{11} + ma_{12})x + (ma_{22} + a_{12})y + a_{13} + ma_{23} = 0$; essa passerà per $Q_\infty(1 : m' : 0)$ se e solo se è $a_{11} + ma_{12} + m'(ma_{22} + a_{12}) = 0$ che con facili passaggi si riconduce alla (16.1). $\square$

Esempio 16.4. La conica di equazione $3x^2 - 2xy + y^2 - 3x + y - 7 = 0$ ammette come involuzione dei diametri coniugati l'equazione $mm' - (m + m') + 3 = 0$. Le rette unite hanno coefficienti angolari che sono soluzioni dell'equazione $m^2 - 2m + 3 = 0$ che non sono reali. L'involuzione è perciò ellittica e quindi la conica è un'ellisse.

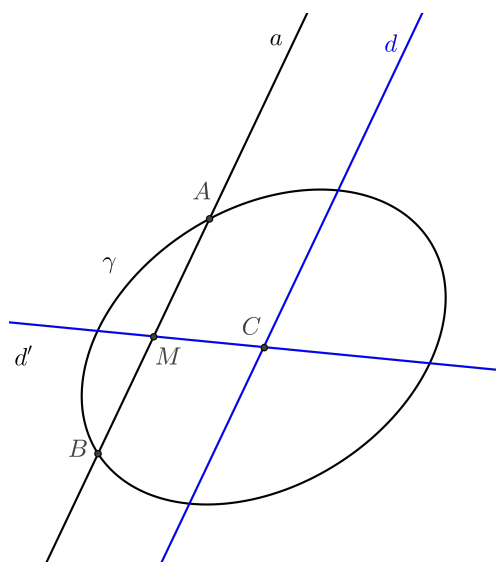


Figura 16.3 Determinazione grafica del diametro coniugato

È facile verificare che le tangenti agli estremi A e B di un diametro d sono parallele al diametro coniugato (v. fig. 16.4 a fronte) infatti d' è la polare del punto improprio P_∞ di d' e quindi le tangenti alla conica in A e B passano per P_∞ .

Osserviamo anche che ogni coppia di diametri coniugati forma, con la retta impropria un triangolo autopolare, infatti la polare dell'intersezione di due diametri, che è il centro della conica, è la retta impropria e la polare del punto improprio di ciascun diametro è quello ad esso coniugato. Se chiamiamo r e s i due diametri coniugati, la rete di coniche

Se la conica è una circonferenza l'equazione dell'involuzione dei diametri coniugati si riduce a

$$mm' + 1 = 0 \quad (16.2)$$

relazione che si chiama anche *involuzione circolare*, viceversa ogni conica che ammetta la (16.2) come involuzione dei diametri coniugati è una circonferenza. Osserviamo anche che la (16.2) è la relazione che lega due rette ortogonali quindi in una circonferenza i diametri coniugati sono ortogonali e viceversa ogni conica per cui tutti i diametri ortogonali sono coniugati è una circonferenza.

che ammette queste due rette come diametri coniugati ha equazione $\lambda r^2 + \mu s^2 + \nu u^2 = 0$.

Esempio 16.5. La conica che ammette le rette r ed s rispettivamente di equazioni $y = x$ e $y = -2x + 1$ come diametri coniugati e passa per $A(1, 2)$ e $B(0, 2)$ può essere cercata nella rete di coniche che ha equazione $\alpha(x - y)^2 + \beta(2x + y - 1)^2 + 1 = 0$; imponendo il passaggio per i due punti dati si ottiene il sistema $\begin{cases} \alpha + 9\beta + 1 = 0 \\ 4\alpha + \beta + 1 = 0 \end{cases}$ che ha come soluzione $\alpha = -\frac{8}{35}$ e $\beta = -\frac{3}{35}$ da cui l'equazione della conica $20x^2 - 4xy + 11y^2 - 12x - 6y - 32 = 0$ che è un'ellisse.

16.2 Assi di una conica

Si dice *asse* di una conica un diametro proprio che sia coniugato alla direzione ad esso ortogonale; per quanto detto prima un asse dimezza le corde in direzione ortogonale, quindi è *asse di simmetria ortogonale*.

Per una parabola $\mathcal{P}$ tutti i diametri sono paralleli e passano per il punto improprio di $\mathcal{P}$, di conseguenza l'asse è la polare del punto improprio in direzione ortogonale a quello di $\mathcal{P}$; quindi la parabola ha un solo asse.

Consideriamo ora una conica a centro γ e siano a un suo asse e a' il diametro coniugato ad a . Per definizione di asse il polo di a è il punto improprio di a' , quindi, per la legge di reciprocità il polo di a' è il punto improprio di a dunque anche A' è un asse di γ . Abbiamo così dimostrato il

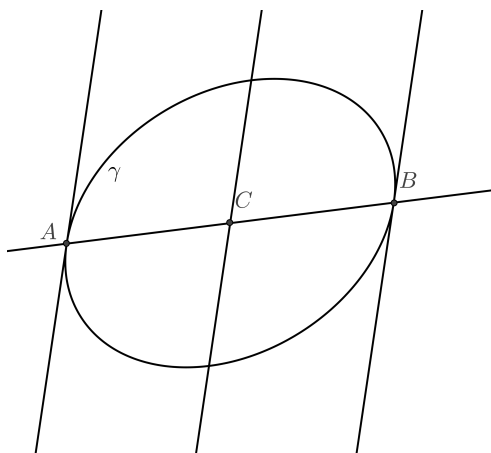


Figura 16.4 Tangenti agli estremi di un diametro

Teorema 16.2. Se una conica a centro possiede un asse, ne possiede almeno un altro.

Di più, gli assi di una conica a centro diversa da una circonferenza sono esattamente due. Infatti, se nella (16.1 a pagina 161) poniamo

$m' = -\frac{1}{m}$ otteniamo

$$-a_{22} + a_{12} \left(m - \frac{1}{m} \right) + a_{11} = 0$$

che diventa

$$a_{12}m^2 + (a_{11} - a_{22})m - a_{12} = 0 \quad (16.3)$$

Per $a_{12} \neq 0$ la (16.3) ammette due radici reali e distinte, quindi la conica ammette esattamente una coppia di assi. Se invece è $a_{12} = 0$ e $a_{11} \neq a_{22}$ l'equazione (16.3) si abbassa di grado ed ammette una radice nulla: si hanno quindi anche in questo caso esattamente due assi, di coefficienti angolari $m = 0$ e $m = \infty$.

Infine se è $a_{12} \neq 0$ e $a_{11} = a_{22}$ cioè se la conica è una circonferenza, la (16.3) diventa un'identità in accordo col fatto che l'involuzione circolare è l'involuzione dei diametri coniugati della circonferenza, per la quale, dunque, tutti i diametri sono assi.

Esempio 16.6. Vogliamo gli assi della conica $\gamma : x^2 + xy - 2 = 0$. In questo caso la (16.3) diventa $m^2 + m - 1 = 0$ che ha come radici $m = \frac{-1 \pm \sqrt{5}}{2}$: gli assi della γ possono essere determinati o come polari dei due punti impropri $\left(1 : \frac{-1 \pm \sqrt{5}}{2} : 0\right)$ oppure osservando che il centro di γ è l'origine in quanto la sua equazione manca dei termini lineari. Dunque gli assi sono le rette di equazioni $y = \frac{-1 \pm \sqrt{5}}{2}x$.

Parte IV

Geometria dello spazio

Capitolo 17

Rette e piani nello spazio

17.1 Equazioni parametriche della retta nello spazio

Quanto detto nel paragrafo 11.2 a pagina 107 sulle equazioni parametriche della retta nel piano si generalizza facilmente al caso di rette nello spazio, precisamente diciamo che, nello spazio, una retta che passa per il punto $P(x_0, y_0, z_0)$ ed ha la direzione del vettore $\vec{u} = [a, b, c]$ è l'insieme dei punti $P(x, y, z)$ tali che il vettore $[x - x_0, y - y_0, z - z_0]$ sia proporzionale al vettore $\vec{u} = [a, b, c]$, ottenendo così le *equazioni parametriche di una retta nello spazio*.

$$\begin{cases} x = x_0 + at \\ y = y_0 + bt \\ z = z_0 + ct \end{cases} \quad (17.1)$$

Osserviamo che in questo caso l'eliminazione del parametro porta a *due* equazioni cartesiane, quindi *non è possibile descrivere una retta¹ nello spazio mediante una sola equazione cartesiana*. La direzione della retta è data da quella del vettore $[a, b, c]$ le cui componenti si chiamano *parametri direttori* di r .

Esempio 17.1. Vogliamo le equazioni parametriche della retta che passa per i punti $P(3, 1, 0)$ e $Q(0, -1, 1)$. Essa ha la direzione del vettore $[3, 1, 0] - [0, -1, 1] = [3, 2, -1]$ e quindi le sue equazioni parametriche sono

$$\begin{cases} x = 3 + 3t \\ y = 1 + 2t \\ z = -t \end{cases}$$

Infatti per $t = 0$ otteniamo il punto P e per $t = -1$ il punto Q .

¹In generale una linea.

OSSERVAZIONE 17.1. Da quanto detto discende che i parametri direttori di una retta sono definiti a meno di un fattore di proporzionalità non nullo, in accordo col fatto che il vettore $\vec{v} = [a, b, c]$ ha la stessa direzione del vettore $[ka, kb, kc]$ se $k \neq 0$. Se invece di $\vec{v}$ consideriamo il versore $\vec{v}' = \frac{\vec{v}}{\|\vec{v}\|}$ le componenti di $\vec{v}'$ sono i cosiddetti *coseni direttori* della retta, infatti sono proprio i coseni degli angoli² che la retta forma con la direzione positiva degli assi coordinati.

Esempio 17.2. Vogliamo trovare i coseni direttori della retta

$$r : \begin{cases} x = 1 + t \\ y = -2t \\ z = -t \end{cases}.$$

Una terna di parametri direttori è $1, -2, -1$ e quindi, normalizzando i parametri direttori otteniamo i coseni direttori che sono, ordinatamente, $\frac{1}{\sqrt{1+4+1}}, \frac{-2}{\sqrt{6}}, \frac{-1}{\sqrt{6}}$.

Esempio 17.3. Vogliamo scrivere le equazioni parametriche delle rette che passano per il punto $P(1, 2, 3)$ e non tagliano il piano xy . Ogni retta che non taglia il piano xy ha la direzione di un qualunque vettore del piano xy diverso dal vettore nullo, cioè $[a, b, 0]$ con a e b non contemporaneamente nulli; allora le equazioni parametriche delle rette cercate saranno

$$\begin{cases} x = 1 + at \\ y = 2 + bt \\ z = 3 \end{cases}.$$

Come nel piano i parametri direttori sono utili per controllare perpendicolarità e parallelismo di rette nello spazio, come vedremo nel § 17.3 a pagina 170.

17.2 Equazione di un piano nello spazio

Un piano π è orientato quando in esso è dato il verso delle rotazioni positive. Per consuetudine una retta n ortogonale ad un piano orientato,

²Dobbiamo precisare che nello spazio si parla di angolo φ tra due rette orientate r ed s anche se esse non si incontrano e non sono parallele – cioè sono, come si suol dire, sghembe – quando formano l'angolo φ le rette r' e s' , rispettivamente parallele ed equiverse ad r ed s e passanti per un medesimo punto, per esempio l'origine del sistema di riferimento.

detta anche *normale*, è orientata in modo che un osservatore in piedi sul piano π dalla parte positiva di n veda le rotazioni positive sul piano operare in senso antiorario.

Sia ora dato un punto $P(x_0, y_0, z_0) \in \pi$ e la sua normale n , di parametri direttori a, b, c e passante per P . Osserviamo che il generico punto $Q(x, y, z)$ appartiene a π se e solo se appartiene a rette passanti per P ortogonali ad n . Una terna di parametri direttori della generica retta per P è $x - x_0, y - y_0$ e $z - z_0$ quindi deve essere

$$a(x - x_0) + b(y - y_0) + c(z - z_0) = 0$$

Che diventa

$$ax + by + cz + d = 0 \quad (17.2)$$

pur di porre $d = -ax_0 - by_0 - cz_0$. Abbiamo così dimostrato il

Teorema 17.1. *Nello spazio tutte e sole le equazioni della forma 17.2 rappresentano un piano in cui a, b e c sono numeri non tutti e tre nulli³.*

Si suol dire che a, b e c sono i parametri direttori del piano, cioè i parametri direttori di un piano coincidono con quelli di una sua qualsiasi normale.

OSSERVAZIONE 17.2. Abbiamo visto che i numeri a, b e c non devono essere tutti e tre nulli. Tuttavia da quanto visto si può dedurre che:

- Se uno dei tre è nullo l'equazione (17.2) rappresenta un piano parallelo all'asse avente lo stesso nome della variabile che manca.
- Se due dei tre sono nulli, il piano rappresentato dalla (17.2) è parallelo al piano individuato dai due assi delle variabili che mancano.

Esempio 17.4. *Il piano $2x - y + 3 = 0$ è parallelo all'asse z ed interseca il piano xy lungo la retta di equazioni*

$$\begin{cases} 2x - y + 3 = 0 \\ z = 0 \end{cases}$$

mentre il piano $x = 5$ è parallelo al piano yz .

³in quanto parametri direttori di una retta.

17.3 Parallelismo e perpendicolarità nello spazio

Nel paragrafo 11.1 a pagina 105 abbiamo visto dei criteri per stabilire quando nel piano due rette sono parallele o perpendicolari; discuteremo ora l'analogo problema nello spazio.

Due rette si dicono parallele se hanno la stessa direzione⁴. Due piani si dicono paralleli se non hanno punti in comune o se coincidono; un piano ed una retta si dicono paralleli se non hanno punti in comune o se la retta giace sul piano. Vale il

Teorema 17.2. *I piani π_1 e π_2 di equazioni rispettive $a_1x + b_1y + c_1z + d_1 = 0$ e $a_2x + b_2y + c_2z + d_2 = 0$ sono paralleli se e solo se i vettori $\vec{u} = [a_1, b_1, c_1]$ e $\vec{v} = [a_2, b_2, c_2]$ sono proporzionali, cioè se $\vec{u} = k\vec{v}$ con $k \neq 0$, quindi se*

$$a_1 = ka_2; \quad b_1 = kb_2; \quad c_1 = kc_2 \quad k \neq 0. \quad (17.3)$$

Se inoltre $d_1 = kd_2$ i piani sono coincidenti; inoltre π_1 e π_2 sono ortogonali se e solo se

$$\langle [a_1, b_1, c_1], [a_2, b_2, c_2] \rangle = 0. \quad (17.4)$$

cioè se e solo se

$$a_1a_2 + b_1b_2 + c_1c_2 = 0 \quad (17.5)$$

Analogo teorema vale per le rette nello spazio:

Teorema 17.3. *Le rette*

$$r_1 : \begin{cases} x = x_1 + a_1t \\ y = y_1 + b_1t \\ z = z_1 + c_1t \end{cases} \quad e \quad r_2 : \begin{cases} x = x_2 + a_2t \\ y = y_2 + b_2t \\ z = z_2 + c_2t \end{cases}$$

sono parallele se e solo se i vettori $\vec{u} = [a_1, b_1, c_1]$ e $\vec{v} = [a_2, b_2, c_2]$ sono proporzionali, cioè se e solo se $\vec{u} = k\vec{v}$ con $k \neq 0$, quindi se vale la (17.3). Inoltre r_1 e r_2 sono ortogonali se e solo se vale la (17.5)

OSSERVAZIONE 17.3. Ribadiamo (vedi nota 2 a pagina 168) che nello spazio due rette perpendicolari possono anche non intersecarsi, cioè essere sghembe!

Per le relazioni di perpendicolarità e parallelismo tra una retta ed un piano vale il teorema seguente

⁴Abbiamo già visto che è comodo considerare parallele anche due rette coincidenti, in modo che la relazione di parallelismo sia una relazione di equivalenza.

Teorema 17.4. *Sia r la retta di equazioni*

$$\begin{cases} x = x_0 + a_1 t \\ y = y_0 + b_1 t \\ z = z_0 + c_1 t \end{cases}$$

e π il piano di equazione $a_2 x + b_2 y + c_2 z + d_2 = 0$ allora r e π sono paralleli se e solo se $\langle [a_1, b_1, c_1], [a_2, b_2, c_2] \rangle = 0$ e sono perpendicolari se e solo se i vettori $\vec{u} = [a_1, b_1, c_1]$ e $\vec{v} = [a_2, b_2, c_2]$ sono proporzionali, cioè se $\vec{u} = k\vec{v}$ con $k \neq 0$, quindi se vale la (17.3).

Dunque le relazioni di parallelismo e perpendicolarità tra due rette o tra due piani nello spazio per così dire si “scambiano” nel caso di parallelismo e perpendicolarità tra una retta ed un piano; ciò avviene perché per un piano π di equazione

$$ax + by + cz + d = 0$$

le componenti del vettore $[a, b, c]$ formano, come abbiamo visto, una terna di parametri direttori di una retta ortogonale a π .

17.4 La retta intersezione di due piani

Siano dati i due piani

$$a_1 x + b_1 y + c_1 z + d_1 = 0 \quad (\pi_1)$$

$$a_2 x + b_2 y + c_2 z + d_2 = 0 \quad (\pi_2)$$

Per quanto detto finora essi saranno non paralleli se i vettori $\vec{v}_1 = [a_1, b_1, c_1]$ e $\vec{v}_2 = [a_2, b_2, c_2]$ sono linearmente indipendenti. Dal punto di vista geometrico, due piani non paralleli hanno in comune una retta, in accordo col fatto che $P(x, y, z) \in \pi_1 \cap \pi_2$ se e solo se (x, y, z) è soluzione del sistema

$$\begin{cases} a_1 x + b_1 y + c_1 z + d_1 = 0 \\ a_2 x + b_2 y + c_2 z + d_2 = 0 \end{cases} \quad (17.6)$$

Dal teorema 3.5 a pagina 30 segue che, nell'ipotesi che i piani non siano paralleli, questo sistema è possibile, dato che, se chiamiamo A la matrice dei coefficienti e B quella completa, si ha $r(A) = r(B) = 2$, dunque il sistema ammette ∞^1 soluzioni che sono effettivamente i punti di una retta nello spazio. Dunque una retta può essere individuata come intersezione di due piani.

Esempio 17.5. Sia data la retta

$$r : \begin{cases} x + y + z = 1 \\ x - z = 0 \end{cases}; \quad (17.7)$$

i vettori $[1, 1, 1]$ e $[1, 0, -1]$ sono linearmente indipendenti dunque i piani non sono paralleli ed il sistema rappresenta una retta. Per trovare le equazioni parametriche di questa retta possiamo osservare che il sistema equivale a

$$\begin{cases} y + 2z = 1 \\ x = z \end{cases}$$

e quindi una possibile soluzione è data da

$$\begin{cases} x = t \\ y = 1 - 2t \\ z = t \end{cases}$$

che sono le coordinate del generico punto di r e quindi rappresentano una terna di equazioni parametriche della retta.

OSSERVAZIONE 17.4. Le soluzioni del sistema (17.7) si possono anche scrivere come

$$\begin{cases} x = \frac{1 - \tau}{2} \\ y = \tau \\ z = \frac{1 - \tau}{2} \end{cases}$$

ed in altri infiniti modi, significativamente diversi, ciascuno dei quali rappresenta una terna di equazioni parametriche della r .

Se accade che $r(A) = 1$, cioè che i vettori $\vec{v}_1 = [a_1, b_1, c_1]$ e $\vec{v}_2 = [a_2, b_2, c_2]$ sono linearmente dipendenti, i due piani sono paralleli se $r(B) = 2$ e coincidenti se $r(B) = 1$; nel primo caso il sistema risulta impossibile, infatti i piani non si incontrano, nel secondo il sistema ammette ∞^2 soluzioni, cioè i due piani π_1 e π_2 coincidono e qualunque punto di essi ha coordinate che sono soluzioni del sistema.

OSSERVAZIONE 17.5. Sia data la retta

$$r : \begin{cases} ax + by + cz + d = 0 \\ a'x + b'y + c'z + d' = 0 \end{cases}$$

la retta r' parallela ad r (quindi avente gli stessi parametri direttori) e passante per l'origine è:

$$r' : \begin{cases} ax + by + cz = 0 \\ a'x + b'y + c'z = 0 \end{cases} \quad (17.8)$$

Poiché r' passa per l'origine, una terna di parametri direttori di r' è data dalle coordinate di un suo punto qualsiasi, che sono quindi le soluzioni del sistema lineare omogeneo (17.8) e quindi proporzionali ai tre minori $\begin{vmatrix} a & b \\ a' & b' \end{vmatrix}, -\begin{vmatrix} a & c \\ a' & c' \end{vmatrix}$ e $\begin{vmatrix} b & c \\ b' & c' \end{vmatrix}$ estratti dalla matrice dei coefficienti.

Quanto detto nell'Osservazione 17.5 rappresenta un modo pratico e veloce per trovare una terna di parametri direttori di una retta scritta come intersezione di due piani, senza passare dalle sue equazioni parametriche.

17.5 Fasci di piani

L'insieme di tutti i piani che passano per una stessa retta si chiama *fascio di piani*

Teorema 17.5. Se $\pi_1 : a_1x + b_1y + c_1z + d_1 = 0$ e $\pi_2 : a_2x + b_2y + c_2z + d_2 = 0$ sono due piani non paralleli che definiscono la retta r , l'equazione del generico piano passante per r è

$$\lambda(a_1x + b_1y + c_1z + d_1) + \mu(a_2x + b_2y + c_2z + d_2) = 0 \quad (17.9)$$

cioè l'equazione del fascio di piani che ha per sostegno la retta r si ottiene come combinazione lineare delle equazioni di due piani qualsiasi passanti per r .

Una conseguenza del Teorema 17.5 è che qualunque coppia di piani appartenenti al fascio⁵ di equazione (17.9) rappresenta la retta r le cui equazioni cartesiane possono essere dunque molto differenti.

OSSERVAZIONE 17.6. Come già più volte osservato, nella pratica può essere comodo usare un solo parametro non omogeneo invece di due omogenei, con le solite avvertenze sul valore infinito del parametro.

OSSERVAZIONE 17.7. Ha senso anche parlare di fasci di *piani paralleli*: combinando linearmente le equazioni di due piani paralleli si ottengono piani le cui equazioni differiscono solo per il termine noto. Essi definiscono una *retta impropria* dello spazio.

⁵Osserviamo che ogni piano del fascio è individuato da una coppia di coefficienti λ e μ della (17.9), anch'essi definiti a meno di un fattore di proporzionalità non nullo.

17.6 Altri problemi su rette e piani

Esaminiamo in questo paragrafo, prevalentemente su esempi, alcuni altri problemi sulle rette e sui piani nello spazio.

Intersezione tra retta e piano

Siano dati il piano di equazione $ax + by + cz + d = 0$ e la retta di equazioni

$$\begin{cases} x = x_0 + \alpha t \\ y = y_0 + \beta t \\ z = z_0 + \gamma t \end{cases}.$$

Piano e retta hanno un solo punto comune se e solo se retta e piano non sono paralleli, cioè se e solo se $a\alpha + b\beta + c\gamma \neq 0$: in questo caso, sostituendo rispettivamente x, y, z nell'equazione del piano si ottiene un valore di t che, sostituito a sua volta nelle equazioni della retta dà le coordinate del punto di intersezione.

Esempio 17.6. Si considerino il piano $\pi : 3x - y + z - 1 = 0$ e la retta

$$r : \begin{cases} x = t \\ y = t - 1 \\ z = 2t \end{cases}$$

sostituendo le coordinate del generico punto della retta nell'equazione del piano, si ottiene $3t - t + 1 + 2t - 1 = 0$ da cui $t = 0$ e dunque l'intersezione è $\pi \cap r = P(0, -1, 0)$.

Se la retta è data come intersezione di due piani gli eventuali punti comuni tra piano e retta sono le soluzioni del sistema formato dalle tre equazioni; piano e retta hanno un solo punto in comune se e solo se tale sistema ammette una ed una sola soluzione, cioè se $r(A) = 3$, il che equivale a dire che $\det A \neq 0$, se con A abbiamo indicato la matrice dei coefficienti del sistema.

Esempio 17.7. La retta

$$r) \equiv \begin{cases} x = y \\ y = z \end{cases}$$

ed il piano $\alpha \equiv x - 2y + 3z - 1 = 0$ hanno in comune il punto le cui coordinate sono la soluzione del sistema

$$\begin{cases} x - y = 0 \\ y - z = 0 \\ x - 2y + 3z = 1 \end{cases}$$

cioè $P = \left(\frac{1}{2}, \frac{1}{2}, \frac{1}{2}\right)$.

Se $r(A) = 2$ ed il sistema è possibile, allora $r \in \pi$, altrimenti se il sistema è impossibile r e π non hanno punti in comune, quindi sono paralleli. Osserviamo che è sempre $r(A) \neq 1$ altrimenti i tre piani coinciderebbero e non sarebbe individuata alcuna retta.

Rette sghembe

Due rette che non hanno punti in comune e non sono parallele, cioè due rette non complanari, si dicono *sghembe*. Se entrambe le rette sono date come intersezione di due piani, esse sono sghembe se non sono parallele e se il sistema formato dalle quattro equazioni è impossibile.

Esempio 17.8. *Le rette*

$$r : \begin{cases} x + z = 0 \\ y - z = 1 \end{cases} \quad e \quad s : \begin{cases} 2x + y = 1 \\ x + 2z = 0 \end{cases}$$

non sono parallele; consideriamo allora il sistema formato dalle quattro equazioni:

$$\begin{cases} x + z = 0 \\ y - z = 1 \\ 2x + y = 1 \\ x + 2z = 0 \end{cases}$$

la matrice dei coefficienti sarà $A = \begin{bmatrix} 1 & 0 & 1 \\ 0 & 1 & -1 \\ 2 & 1 & 0 \\ 1 & 0 & 2 \end{bmatrix}$ che ha rango 3 e quella

completa sarà $B = \begin{bmatrix} 1 & 0 & 1 & 0 \\ 0 & 1 & -1 & 1 \\ 2 & 1 & 0 & 1 \\ 1 & 0 & 2 & 0 \end{bmatrix}$, anch'essa di rango 3, dunque il sistema

ammette una soluzione, quindi le rette si incontrano, dunque sono complanari

e quindi non sghembe. L'equazione del piano che le contiene entrambe si può determinare in vari modi, per esempio trovando il fascio $\mathcal{F}$ di piani per una delle due e, scelto un punto comodo P sull'altra (ovviamente diverso dal punto di intersezione), scrivere l'equazione del piano di $\mathcal{F}$ passante per P ; oppure, determinato il punto Q comune alle due rette, scegliere un punto comodo $R \in r$ ($R \neq Q$) ed uno comodo $S \in s$ ($S \neq Q$) e poi determinando il piano per i tre punti Q, R ed S .

Se una delle rette è data come intersezione di due piani e l'altra con le sue equazioni parametriche basta sostituire l'espressione del parametro nelle equazioni dei due piani.

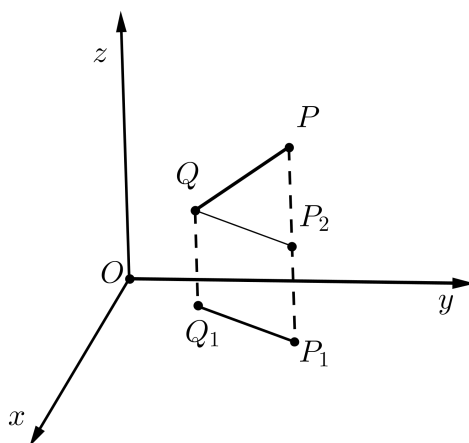


Figura 17.1 Distanza di due punti

Esempio 17.9. Siano date le rette

$$r : \begin{cases} x = t \\ y = 2t \\ z = 1 - 3t \end{cases} \quad \text{ed} \quad s : \begin{cases} x + y - z = 0 \\ x + z = 1 \end{cases}$$

si ha, sostituendo le coordinate del generico punto di r nelle equazioni due piani che formano s , il sistema

$$\begin{cases} t + 2t - 1 + 3t = 0 \\ t + 1 - 3t = 1 \end{cases} \quad \text{che diventa} \quad \begin{cases} 6t = 1 \\ 2t = 0 \end{cases}$$

formato da due equazioni palesemente in contraddizione, quindi il sistema è impossibile e concludiamo che le rette sono sghembe.

Distanze

- *Distanza di due punti:* dalla figura 17.1 nella pagina precedente si vede subito che se sono dati i punti $P(x_1, y_1, z_1)$ e $Q(x_2, y_2, z_2)$, la distanza PQ è l'ipotenusa del triangolo rettangolo che ha come cateti la differenza delle quote $z_1 - z_2$ e la distanza delle proiezioni ortogonali P_1 e Q_1 rispettivamente di P e Q sul piano xy . Sul piano xy dal Teorema di Pitagora si ricava immediatamente che la distanza $P_1(x_1, y_1) - Q_1(x_2, y_2)$ è

$$d(P_1Q_1) = \sqrt{(x_1 - x_2)^2 + (y_1 - y_2)^2}.$$

e quindi un'ulteriore applicazione del Teorema di Pitagora fornisce la distanza di due punti nello spazio:

$$d(PQ) = \sqrt{(x_1 - x_2)^2 + (y_1 - y_2)^2 + (z_1 - z_2)^2}.$$

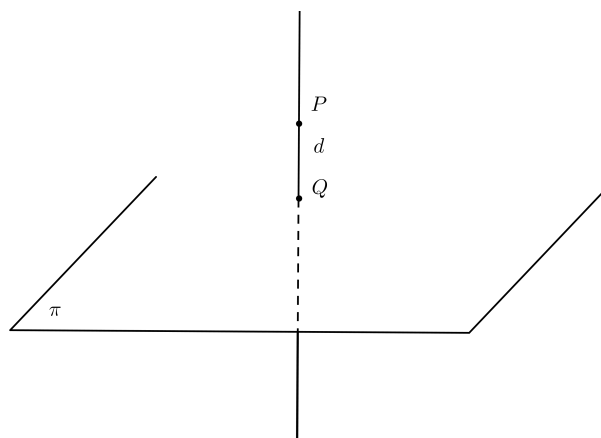


Figura 17.2 Distanza di un punto da un piano

- *Distanza di un punto da un piano:* (fig. 17.2) se $\pi : ax + by + cz + d = 0$ è l'equazione del piano allora la distanza di $P(x_0, y_0, z_0)$ da π è:

$$\frac{|ax_0 + by_0 + cz_0 + d|}{\sqrt{a^2 + b^2 + c^2}} \quad (17.10)$$

infatti distanza si può calcolare considerando la perpendicolare al piano per P e, chiamato Q l'intersezione di questa retta con il piano, calcolando la distanza PQ .

La retta perpendicolare a π passante per il punto P ha equazioni

$$\begin{cases} x = x_0 + at \\ y = y_0 + bt, \\ z = z_0 + ct \end{cases}$$

dunque $Q(x_0 + at, y_0 + bt, z_0 + ct)$ è il punto di intersezione di tale retta con π , e la sua distanza da P è

$$|t|\sqrt{a^2 + b^2 + c^2}, \quad (17.11)$$

se appartiene a π , cioè se

$$\begin{aligned} & a(x_0 + at) + b(y_0 + bt) + c(z_0 + ct) + d = 0 \\ \iff & t(a^2 + b^2 + c^2) = -(ax_0 + by_0 + cz_0 + d) \\ \iff & t = -\frac{ax_0 + by_0 + cz_0 + d}{a^2 + b^2 + c^2}. \end{aligned}$$

Sostituendo questo valore nella formula 17.11 che esprime la distanza di P da Q si ha la 17.10 nella pagina precedente.

- *Distanza di un punto da una retta*: siano P un punto e r una retta tali che $P \notin r$, il piano che passa per P ed è perpendicolare a r incontra la retta r in un punto R (che è la *proiezione ortogonale di P su r*) la distanza del punto dalla retta sarà allora la distanza PR .
- *Distanza di piani paralleli*: se $\pi_1 : ax + by + cz + d_1 = 0$ e $\pi_2 : ax + by + cz + d_2 = 0$ sono due piani paralleli, allora la loro distanza è:

$$\frac{|d_1 - d_2|}{\sqrt{a^2 + b^2 + c^2}} \quad (17.12)$$

Infatti se $P(x_0, y_0, z_0) \in \pi_1$, è evidente che la distanza cercata è

$$d(P, \pi_2) \text{ cioè } \frac{|ax_0 + by_0 + cz_0 + d_2|}{\sqrt{a^2 + b^2 + c^2}} = \frac{|d_2 - d_1|}{\sqrt{a^2 + b^2 + c^2}}$$

- *Distanza di due rette sghembe*: se r_1 e r_2 sono due rette sghembe, si vede facilmente che esistono due punti $P \in r_1$ e $Q \in r_2$ tali che la retta PQ è perpendicolare sia ad r_1 che ad r_2 . La distanza PQ è la distanza delle due rette. Dal punto di vista geometrico, però è più comodo considerare il piano σ passante per r_1 e parallelo a r_2 ; la distanza cercata sarà quella di un qualsiasi punto $P \in r_2$ da σ .

Esempio 17.10. Vogliamo calcolare la distanza delle rette sghembe dell'esempio 17.9 a pagina 176. Calcoliamo l'equazione del piano per s parallelo ad r di cui una terna di parametri direttori è $(1, 2, -3)$. Il fascio di piani che ha per sostegno s ha equazione $x + y - z + k(x + z - 1) = 0$. Dovrà essere $k + 1 + 2 - 3(k - 1) = 0$ da cui $k = 3$, quindi il piano cercato ha equazione $4x + y + 2z - 3 = 0$. La distanza richiesta sarà la distanza di questo piano da un punto qualsiasi della r , per esempio il punto $P(0, 0, 1)$ (corrispondente al valore $t = 0$) che è $\frac{1}{\sqrt{21}}$.

Angoli tra rette, tra piani, tra rette e piani

Siano $\pi_1 : a_1x + b_1y + c_1z + d_1 = 0$ e $\pi_2 : a_2x + b_2y + c_2z + d_2 = 0$ due piani. L'ampiezza $\varphi \in [0, \frac{\pi}{2}]$ dell'angolo tra π_1 e π_2 è, ponendo $\vec{v} = [a_1, b_1, c_1]$ e $\vec{w} = [a_2, b_2, c_2]$, tale che:

$$\cos \varphi = \frac{\langle \vec{v}, \vec{w} \rangle}{\|\vec{v}\| \cdot \|\vec{w}\|}.$$

Essa deriva dalla formula che dà l'angolo tra due vettori e dal fatto che se α è l'angolo tra $\vec{v}$ e $\vec{w}$, essendo rispettivamente $\vec{v}$ e $\vec{w}$ ortogonali ai due piani, si ha $\varphi = \alpha$ se $\alpha \leq \frac{\pi}{2}$ e $\varphi = \pi - \alpha$ se $\alpha > \frac{\pi}{2}$ e quindi $\cos \varphi = |\cos \alpha|$.

Siano ora

$$r : \begin{cases} x = x_1 + a_1t \\ y = y_1 + b_1t \\ z = z_1 + c_1t \end{cases} \quad \text{e} \quad s : \begin{cases} x = x_2 + a_2s \\ y = y_2 + b_2s \\ z = z_2 + c_2s \end{cases}$$

due rette nello spazio; e sia φ l'angolo da esse formato⁶, con $\varphi \in [0, \frac{\pi}{2}]$; se esse sono parallele diremo che $\varphi = 0$, in generale indicando con $\vec{v} = [a_1, b_1, c_1]$ e $\vec{w} = [a_2, b_2, c_2]$ si ha

$$\cos \varphi = \frac{\langle \vec{v}, \vec{w} \rangle}{\|\vec{v}\| \cdot \|\vec{w}\|}.$$

Concludiamo considerando una retta

$$r : \begin{cases} x = x_1 + at \\ y = y_1 + bt \\ z = z_1 + ct \end{cases}$$

⁶Ricordiamo, ancora una volta, che non è necessario che r e s si incontrino.

ed un piano $\pi : \alpha x + \beta y + \gamma z + \delta = 0$, se $v = [a, b, c]$ e $w = [\alpha, \beta, \gamma]$ e $\varphi \in \left[0, \frac{\pi}{2}\right]$ è l'angolo tra la retta ed il piano, si ha

$$\sin \varphi = \frac{\langle \vec{v}, \vec{w} \rangle}{\|\vec{v}\| \cdot \|\vec{w}\|} \quad (17.13)$$

Infatti, detto $\psi \in \left[0, \frac{\pi}{2}\right]$ l'angolo tra r e la normale al piano, si ha

$$\cos \psi = \frac{\langle \vec{v}, \vec{w} \rangle}{\|\vec{v}\| \cdot \|\vec{w}\|}$$

e poiché $\varphi = \frac{\pi}{2} - \psi$ si ha $\cos \psi = \sin \varphi$ e quindi la (17.13).

17.7 Simmetrie

In questo paragrafo presentiamo, per lo più mediante esempi, alcuni problemi nello spazio riguardanti le simmetrie rispetto a punti a rette ed a piani.

Simmetrie rispetto ad un punto

Siano P e Q due punti distinti dello spazio; il simmetrico P' di P rispetto a Q è l'unico punto appartenente alla retta PQ e tale che Q sia punto medio del segmento PP' .

Sia ora P un punto dello spazio e sia f una qualunque figura, la figura f' simmetrica di f rispetto a P è il luogo dei punti simmetrici rispetto a P dei punti di f .

Esempio 17.11. Vogliamo il simmetrico A' del punto $A(1, 0, 0)$ rispetto al punto $M(-1, 2, 1)$. Dalla definizione si ha

$$\begin{cases} x_M = \frac{x_A + x_{A'}}{2} \\ y_M = \frac{y_A + y_{A'}}{2} \\ z_M = \frac{z_A + z_{A'}}{2} \end{cases} \quad \text{da cui} \quad \begin{cases} x_{A'} = 2x_M - x_A \\ y_{A'} = 2y_M - y_A \\ z_{A'} = 2z_M - z_A \end{cases}$$

quindi, nel caso in esame,

$$\begin{cases} x_{A'} = 2(-1) - 1 = -3 \\ y_{A'} = 2 \cdot 2 - 0 = 4 \\ z_{A'} = 2 \cdot 1 - 0 = 2 \end{cases}$$

da cui $A'(-3, 4, 2)$.

Esempio 17.12. Sia data la retta

$$r : \begin{cases} x = t + 1 \\ y = 2t \\ z = t - 1 \end{cases}.$$

Vogliamo le equazioni della retta r' simmetrica di r rispetto al punto $P(1, -1, 2)$; essa è costituita dal luogo dei punti simmetrici di quelli di r rispetto a P , quindi sarà, come è facile dimostrare, una retta parallela ad r . Scegliamo due punti "comodi" su r , per esempio quello per cui $t = 0$ cioè $A(1, 0, 1)$ e quello per $t = 1$, cioè $B(2, 2, 0)$. La retta cercata è quella che passa per A' e B' , simmetrici di A e B rispettivamente. Procedendo come nell'esercizio 17.11 a fronte troviamo $A'(1, -2, 5)$ e $B'(0, -4, 4)$ e dunque sarà

$$r' : \begin{cases} x = -t \\ y = -4 - 2t \\ z = 4 - t \end{cases}.$$

Esempio 17.13. Per determinare l'equazione del piano π' simmetrico di $\pi : x - y + z = 0$ rispetto a $T(2, -1, 0)$ consideriamo il generico punto $P(x, y, z)$, scriviamo le coordinate del suo simmetrico P' rispetto a T procedendo come nell'esercizio 17.11 nella pagina precedente ed otteniamo $x + x' = 2x_T$, $y + y' = 2y_T$ e $z + z' = 2z_T$ da cui

$$\begin{cases} x = 4 - x' \\ y = -2 - y' \\ z = -z' \end{cases}.$$

Poiché P sta su π se e solo se $x - y + z = 0$ si ha che le coordinate di P' devono soddisfare l'equazione $4 - x' + 2 + y' - z' = 0$ che si può scrivere come $x' - y' + z' - 6 = 0$ e che rappresenta un piano parallelo a π .

Simmetria rispetto ad un piano

Siano P un punto e π un piano tali che $P \notin \pi$; chiamiamo proiezione ortogonale di P su π il punto H intersezione tra π stesso e la retta per P perpendicolare a π . Il simmetrico di un punto P rispetto ad un piano π è il simmetrico di P rispetto ad H . Come nel caso della simmetria rispetto ad un punto⁷ anche in questo caso si dimostra che la simmetrica di una retta rispetto ad un piano è una retta ed il simmetrico di un piano rispetto ad un piano è ancora un piano. Il tutto è illustrato nella figura 17.3

⁷ detta anche *simmetria centrale*.

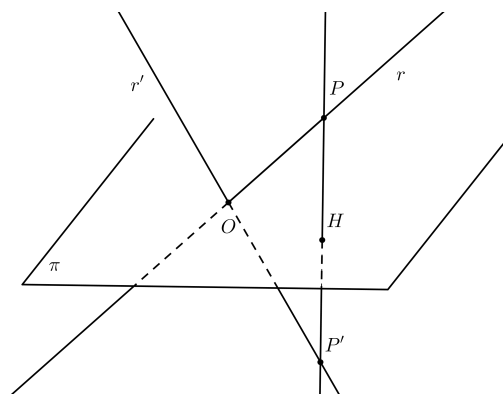


Figura 17.3 Simmetrica di una retta rispetto ad un piano

Esempio 17.14. Vogliamo le equazioni della retta simmetrica di

$$r : \begin{cases} x = t \\ y = t \\ z = t \end{cases}$$

rispetto al piano $\pi : x - y + z = 0$. La retta taglia il piano nell'origine. La retta cercata sarà dunque (vedi figura 17.3) la congiungente dell'origine con il punto P' simmetrico rispetto al piano di un qualsiasi punto $P \in r$ diverso dal punto comune (nel nostro caso l'origine $O(0,0,0)$). Scegliamo, per esempio $P(1,1,1)$. La retta per P ortogonale a π ha equazioni⁸

$$\begin{cases} x = 1 + \tau \\ y = 1 - \tau \\ z = 1 + \tau \end{cases}$$

da cui si ha $1 + \tau - (1 - \tau) + 1 + \tau = 0$ e quindi $H\left(\frac{2}{3}, \frac{4}{3}, \frac{2}{3}\right)$ dunque le coordinate di P' sono

$$\begin{cases} x_{P'} = 2 \cdot \frac{2}{3} - 1 = \frac{1}{3} \\ y_{P'} = 2 \cdot \frac{4}{3} - 1 = \frac{5}{3} \\ z_{P'} = 2 \cdot \frac{2}{3} - 1 = \frac{1}{3} \end{cases}$$

⁸Per evitare equivoci usiamo un parametro diverso, il parametro τ . È ovvio che il nome del parametro è del tutto arbitrario.

allora la retta r' , che congiunge O con H ha equazioni

$$\begin{cases} x = \frac{1}{3}k \\ y = \frac{5}{3}k \\ z = \frac{1}{3}k \end{cases} \quad \text{che possiamo anche scrivere (vedi la nota 8)} \quad \begin{cases} x = \alpha \\ y = 5\alpha \\ z = \alpha \end{cases}$$

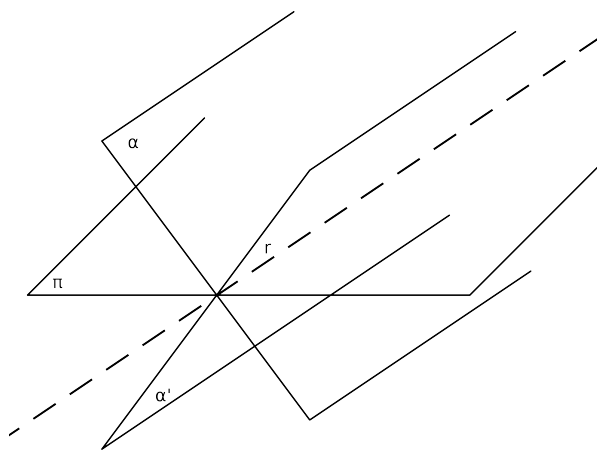


Figura 17.4 Simmetrico di un piano rispetto ad un piano

Esempio 17.15. Per determinare l'equazione del piano α' simmetrico di $\alpha : x = 1$ rispetto a $\pi : x - y - z = 0$ conviene utilizzare la teoria dei fasci di piani, infatti il piano cercato passerà per la retta r intersezione tra α e π (v. Fig. 17.4). Dunque l'equazione di α' sarà del tipo $x - y - z + k(x - 1) = 0$. Il valore di k può essere determinato procedendo come nel caso precedente: se consideriamo su α un punto, per esempio $P(1, 0, 0)$ la sua proiezione H su π stando sulla retta

$$\begin{cases} x = 1 + t \\ y = -t \\ z = -t \end{cases}$$

perpendicolare a π e passante per P risulta individuata da $t = -\frac{1}{3}$ quindi è $H = (\frac{2}{3}, \frac{1}{3}, \frac{1}{3})$. Dunque il simmetrico P' di P sarà $P'(\frac{1}{3}, \frac{2}{3}, \frac{2}{3})$ e sostituendo le sue coordinate nell'equazione del fascio si ottiene, con semplici calcoli, l'equazione del piano α' che è $x + 2y + 2z - 3 = 0$.

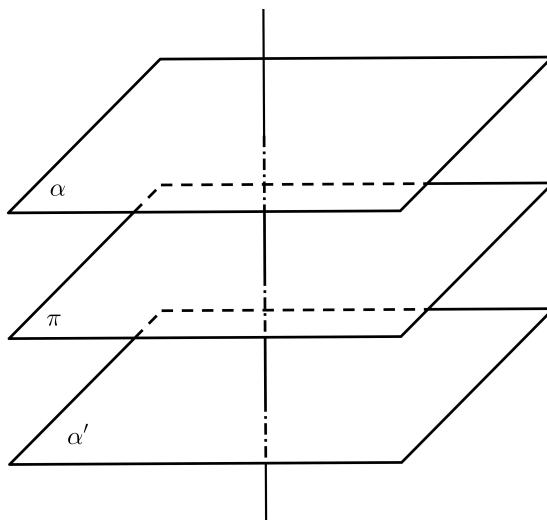


Figura 17.5 Simmetrico di un piano rispetto ad un piano parallelo

Nel caso in cui i due piani siano paralleli, vedi fig. 17.5 si può anche, per esempio, considerare la loro distanza: se essa è d basta scrivere l'equazione del piano a distanza d da α .

Simmetrie rispetto ad una retta

Se P è un punto dello spazio ed r una retta che non lo contiene, chiamando H la proiezione ortogonale di P su r , il simmetrico P' di P rispetto ad r è il punto dello spazio simmetrico di P rispetto ad H .

La retta r' simmetrica di r rispetto ad s si determina (v. fig. 17.6 nella pagina successiva) considerando i simmetrici di due punti di r

Esempio 17.16. Vogliamo le equazioni della retta r' simmetrica di

$$r : \begin{cases} x = 1 \\ y = t \\ z = -t \end{cases} \quad \text{rispetto ad } s : \begin{cases} x = z \\ y = 0 \end{cases}.$$

Si verifica subito che le due rette date sono sghembe; per determinare r' occorre e basta determinare i simmetrici di due punti comodi di r e scrivere le equazioni della retta che li congiunge: per esempio siano $P(1, 0, 0)$ e $Q(1, 1, -1)$ (corrispondenti, rispettivamente, ai valori $t = 0$ e $t = 1$) i due punti: i due simmetrici sono, come si ricava facilmente $P'(0, 0, -\frac{1}{2})$ e $Q'(-1, -1, 1)$ dunque una terna di parametri direttori di s' è $1, 1, -\frac{3}{2}$ quindi la retta s' ha

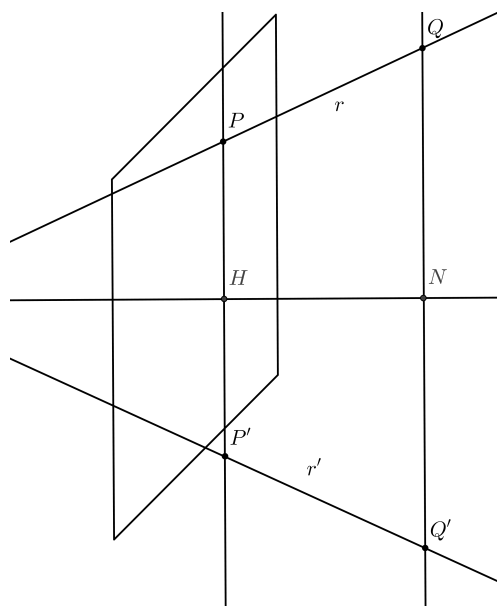


Figura 17.6 Simmetrica di una retta rispetto ad un'altra retta

equazioni

$$\begin{cases} x = 2t - 1 \\ y = 2t + -1 . \\ z = -3t + 1 \end{cases}$$

17.8 Coordinate omogenee nello spazio

In analogia con quanto visto nel piano, si può introdurre una sistema di coordinate omogenee anche nello spazio: i punti saranno definiti da quaterne di coordinate omogenee. Quando la quarta coordinata omogenea (anche qui indicata con la lettera u) sarà non nulla si avrà a che fare con punti dello spazio ordinario, mentre, anche qui, i punti la cui quarta coordinata omogenea è nulla verranno detti *impropri* o *all'infinito* e corrisponderanno, nello spazio ordinario, alle varie direzioni di rette parallele. Saremo dunque in presenza di uno spazio ampliato con i punti impropri, cioè i punti per cui $u = 0$ costituenti il cosiddetto *piano improprio*. Su ogni piano di equazione $ax + by + cz + du = 0$ esisterà una retta impropria, che sarà l'intersezione del piano stesso con il piano improprio ed avrà quindi equazioni

$$\begin{cases} ax + by + cz + du = 0 \\ u = 0 \end{cases} \quad \text{o anche} \quad \begin{cases} ax + by + cz = 0 \\ u = 0 \end{cases}$$

Su ogni retta r ci sarà un punto improprio, intersezione tra le rette improprie dei piani passanti per r come mostra il seguente

Esempio 17.17. *Determiniamo le coordinate del punto improprio della retta*

$$r : \begin{cases} x = 1 + t \\ y = 2t \\ z = -t \end{cases}.$$

In coordinate cartesiane avremo

$$r : \begin{cases} y = 2x - 2 \\ z = -x + 1 \end{cases} \quad \text{che in coordinate omogenee diventa } r : \begin{cases} 2x - y - 2u = 0 \\ x + z - u = 0 \end{cases}.$$

Intersecando con il piano improprio si ottiene

$$\begin{cases} y = 2x \\ z = -x \\ u = 0 \end{cases}$$

da cui le coordinate del punto improprio $P_{\infty}(1 : 2 : -1 : 0)$ che è il punto comune alle rette improprie dei due piani $y = 2x$ e $z = -x$ le cui equazioni sono

$$\begin{cases} 2x - y = 0 \\ u = 0 \end{cases} \quad e \quad \begin{cases} x + z = 0 \\ u = 0 \end{cases}.$$

Nello spazio ampliato con gli elementi impropri il termine “complanarità” è sinonimo di “incidenza”: incidenza in un punto proprio per le rette incidenti dello spazio ordinario, ed in un punto improprio per rette parallele.

In questo contesto risulta semplice l’interpretazione geometrica di sistemi lineari in quattro incognite come insiemi di piani

Esempio 17.18. *Vogliamo vedere se le rette*

$$r : \begin{cases} x + ky - z = -2 \\ 2x + kz = -1 \end{cases} \quad \text{ed} \quad s : \begin{cases} x + 3y + (k-1)z = -1 \\ (k+1)y + z = k \end{cases}$$

sono sghembe. Dopo averle riscritte in coordinate omogenee, avremo a che fare con un sistema lineare omogeneo di quattro equazioni in quattro incognite che indicherà la complanarità delle rette, quando ammette autosoluzioni, ed il loro essere sghembe, quando ammetterà solo la soluzione banale⁹.

⁹Ricordando che la quaterna di coordinate omogenee $(0 : 0 : 0 : 0)$ non rappresenta alcun punto.

Quindi, nel nostro caso, avremo il sistema

$$\begin{cases} x + ky - z = -2 \\ 2x + kz = -1 \\ x + 3y + (k-1)z = -1 \\ (k+1)y + z = k \end{cases}$$

di matrice dei coefficienti

$$A : \begin{bmatrix} 1 & k & -1 & 2 \\ 0 & 2 & k & -1 \\ 1 & 3 & k-1 & 1 \\ 0 & k+1 & 1 & -k \end{bmatrix}$$

il cui rango è $r = 4$ per $k \neq \pm 1$, $r = 3$ per $k = -1$ ed $r = 2$ per $k = 1$. Questo significa che per $k \neq \pm 1$ il sistema non ha autosoluzioni, quindi i piani non hanno, nello spazio ampliato con i punti impropri, alcun punto in comune: le due rette sono dunque sghembe; per $k = -1$ ci sono ∞^1 soluzioni, dunque esattamente¹⁰ un punto di intersezione, proprio o improprio (basta risolvere il sistema per scoprirlo, cioè per scoprire se le rette sono effettivamente incidenti o parallele); per $k = 1$ ci sono ∞^2 soluzioni, dunque le due rette coincidono e sono, ovviamente complanari.

Concludiamo il paragrafo notando che l'introduzione dei punti impropri dello spazio permette di rimuovere, tra le altre, la dissimmetria nella trattazione dei fasci di piani passanti per una retta o paralleli, dal momento che, in questo contesto, i piani paralleli passano tutti per una medesima retta impropria.

¹⁰Ricordiamo che anche le coordinate omogenee dei punti dello spazio sono definite a meno di un fattore di proporzionalità.

Capitolo 18

Sui sistemi di riferimento

In questo capitolo ci occuperemo di alcune particolari trasformazioni del sistema di riferimento cartesiano ortogonale in sè e forniremo poi un breve cenno ad altri possibili sistemi di riferimento diversi da quello cartesiano, ma usati spesso in vari settori della Matematica, nelle altre Scienze e nella Tecnologia.

18.1 Rototraslazioni

Risulta quasi immediato verificare che una trasformazione di assi, cioè il passaggio da un sistema di riferimento monometrico ad un altro con gli assi paralleli ed equiversi a quelli del precedente e con la stessa unità di misura¹ è descritto dalle equazioni²

$$\begin{cases} x' = x - \alpha \\ y' = y - \beta \\ z' = z - \gamma \end{cases} \quad (18.1)$$

dove (α, β, γ) sono le coordinate della nuova origine nel sistema di riferimento iniziale; x, y, z sono le coordinate del generico punto in questo sistema e le coordinate accentate sono le coordinate nel nuovo sistema di riferimento.

Le traslazioni conservano le distanze: se due punti hanno distanza d rispetto ad un sistema di riferimento, essi hanno distanza d rispetto

¹Noi considereremo sempre le trasformazioni geometriche nel modo cosiddetto passivo, nel senso, cioè, che saranno sempre gli assi del sistema di riferimento e non i punti a muoversi: risulterebbe interessante trattare l'argomento anche in modo contrario, ma ciò comporterebbe un livello di astrazione superiore che esula dagli scopi di queste dispense.

²cioè tali equazioni legano le coordinate di un punto generico dello spazio nel vecchio sistema di riferimento a quelle del nuovo.

a qualsiasi altro sistema traslato rispetto al primo quindi non vi è cambiamento di unità di misura.

Per quanto riguarda le rotazioni vale un teorema, analogo al Teorema 11.2 a pagina 113, e che si dimostra con la stessa tecnica, secondo cui la matrice di una rotazione è una matrice ortogonale. Più precisamente, con considerazioni puramente geometriche, si fa vedere che le equazioni di questa trasformazione sono:

$$\begin{cases} x' = x \cos \alpha_1 + y \cos \beta_1 + z \cos \gamma_1 \\ y' = x \cos \alpha_2 + y \cos \beta_2 + z \cos \gamma_2 \\ z' = x \cos \alpha_3 + y \cos \beta_3 + z \cos \gamma_3 \end{cases} \quad (18.2)$$

dove $\cos \alpha_1, \cos \beta_1, \cos \gamma_1; \cos \alpha_2, \cos \beta_2, \cos \gamma_2; \cos \alpha_3, \cos \beta_3, \cos \gamma_3$ sono, rispettivamente, i coseni direttori dei nuovi assi rispetto ai vecchi. Poiché la matrice dei coefficienti è la matrice

$$\Gamma = \begin{bmatrix} \cos \alpha_1 & \cos \beta_1 & \cos \gamma_1 \\ \cos \alpha_2 & \cos \beta_2 & \cos \gamma_2 \\ \cos \alpha_3 & \cos \beta_3 & \cos \gamma_3 \end{bmatrix}$$

che è ortogonale, la matrice inversa è quindi la trasposta

$$\Gamma^{-1} = \Gamma_T = \begin{bmatrix} \cos \alpha_1 & \cos \alpha_2 & \cos \alpha_3 \\ \cos \beta_1 & \cos \beta_2 & \cos \beta_3 \\ \cos \gamma_1 & \cos \gamma_2 & \cos \gamma_3 \end{bmatrix}$$

da cui si ricavano facilmente le equazioni della trasformazione inversa.

La composizione di una rotazione e di una traslazione è ancora una isometria³ che prende il nome di *rototraslazione*, le cui equazioni si ottengono, mettendo insieme la (18.1) e la (18.2), con il sistema

$$\begin{cases} x' = x \cos \alpha_1 + y \cos \beta_1 + z \cos \gamma_1 + a \\ y' = x \cos \alpha_2 + y \cos \beta_2 + z \cos \gamma_2 + b \\ z' = x \cos \alpha_3 + y \cos \beta_3 + z \cos \gamma_3 + c \end{cases} \quad (18.3)$$

OSSERVAZIONE 18.1. La composizione di una rotazione e di una traslazione non è, in generale, commutativa; cioè applicando prima la rotazione e poi la traslazione si ottiene in generale un risultato diverso da quello che si ottiene applicando prima la traslazione e poi la rotazione.

³cioè una trasformazione dello spazio in sè che conserva le distanze.

Esempio 18.1. Consideriamo la trasformazione

$$\tau : \begin{cases} x' = x + 2 \\ y' = y - 1 \\ z' = z - 2 \end{cases} \text{ e la rotazione } \varrho : \begin{cases} x'' = \frac{1}{\sqrt{2}}x' - \frac{1}{\sqrt{2}}y' \\ y'' = \frac{1}{\sqrt{2}}x' + \frac{1}{\sqrt{2}}y' \\ z'' = z' \end{cases}$$

sostituendo si ha :

$$\begin{cases} x'' = \frac{1}{\sqrt{2}}(x + 2) - \frac{1}{\sqrt{2}}(y - 1) = \frac{1}{\sqrt{2}}x - \frac{1}{\sqrt{2}}y + \frac{3}{\sqrt{2}} \\ y'' = \frac{1}{\sqrt{2}}(x + 2) + \frac{1}{\sqrt{2}}(y - 1) = \frac{1}{\sqrt{2}}x + \frac{1}{\sqrt{2}}y + \frac{3}{\sqrt{2}} \\ z'' = z - 2 \end{cases}$$

Facendo agire, invece le due trasformazioni in ordine inverso, cioè eseguendo prima la rotazione e poi la traslazione, si ha

$$\varrho : \begin{cases} x' = \frac{1}{\sqrt{2}}x - \frac{1}{\sqrt{2}}y \\ y' = \frac{1}{\sqrt{2}}x + \frac{1}{\sqrt{2}}y \\ z' = z \end{cases} \quad e \quad \tau : \begin{cases} x'' = x' + 2 \\ y'' = y' - 1 \\ z'' = z' - 2 \end{cases}$$

si ottiene facilmente

$$\begin{cases} x'' = \frac{1}{\sqrt{2}}x - \frac{1}{\sqrt{2}}y + 2 \\ y'' = \frac{1}{\sqrt{2}}x + \frac{1}{\sqrt{2}}y - 1 \\ z'' = z - 2 \end{cases}$$

che differisce dalla precedente.

18.2 Coordinate polari e coordinate cilindriche

Anche per i punti dello spazio si parla, come nel piano, di coordinate polari. Vediamo come si possono introdurre e come si passa da un sistema polare ad uno cartesiano e viceversa.

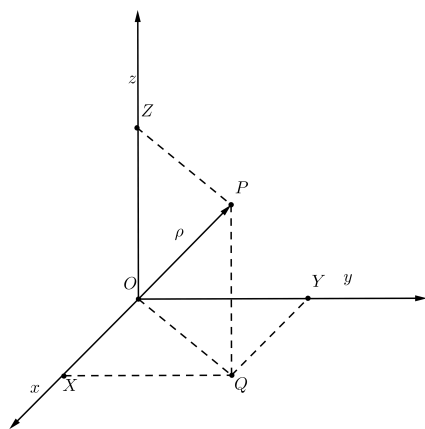


Figura 18.1 Coordinate polari

Fissato nello spazio un sistema di coordinate cartesiane ortogonali⁴ (v. figura 18.1), consideriamo come *asse polare* l'asse z e la semiretta ortogonale all'asse polare coincidente con la direzione positiva dell'asse x ; otteniamo un semipiano ω giacente sul piano xy che prende il nome di *semipiano polare*.

Ciò posto un qualunque punto P dello spazio individua un segmen-

to $\overline{OP} = \rho$ detto *raggio vettore* di P , inoltre il vettore $\overrightarrow{OP}$ (che ha quindi norma ρ) forma, con l'asse polare, un angolo ϑ detto *colatitudine* o *angolo* (o *distanza*) *zenitale* di P (ovviamente si ha $0 \leq \vartheta \leq \pi$), infine il semipiano individuato da P e dall'asse polare forma, con il semipiano ω un angolo φ detto *azimut* o *longitudine* di P (si ha $0 \leq \varphi \leq 2\pi$). A volte, invece della colatitudine ϑ si considera la latitudine $\psi = \frac{\pi}{2} - \vartheta$ che è l'angolo formato dal raggio vettore con il *piano equatoriale* ε passante per O e perpendicolare all'asse polare.

In tal modo ogni punto P viene individuato dalle sue *coordinate polari*⁵ $(\rho, \vartheta, \varphi)$ o (ρ, ψ, φ) molto usate, specialmente in Geografia ed in Astronomia.

Osservando ancora la figura 18.1 si possono facilmente scrivere le formule di passaggio dalle coordinate polari alle cartesiane e viceversa. Infatti se con Q e Z si indicano le proiezioni ortogonali di P sul piano xy e sull'asse z rispettivamente e con X e Y le proiezioni ortogonali di P sugli assi x e y , si hanno le formule $OZ = z = \rho \cos \vartheta$, $OQ = \rho \sin \vartheta$, $= X = x = OQ \cos \varphi = \rho \sin \vartheta \cos \varphi$ e $XQ = OY = y = OQ \sin \varphi = \rho \sin \vartheta \sin \varphi$, quindi

⁴Nello spazio, come, del resto anche nel piano, si può, ovviamente, introdurre un riferimento polare anche in assenza di uno cartesiano preesistente; qui preferiamo invece partire da un sistema cartesiano, in quanto siamo interessati alla determinazione del legame tra le coordinate di un punto in sistemi di riferimento diversi in qualche modo però legati tra loro.

⁵In questo modo, infatti, ogni punto dello spazio viene individuato da una ed una sola terna di numeri; fa eccezione l'origine del riferimento, per cui si ha $\rho = 0$ e φ e ϑ indeterminati

$$\begin{cases} x = \rho \sin \vartheta \cos \varphi \\ y = \rho \sin \vartheta \sin \varphi \\ z = \rho \cos \vartheta \end{cases}$$

e, all'inverso, come è facile verificare,

$$\begin{cases} \rho = \sqrt{x^2 + y^2 + z^2} \\ \vartheta = \arccos \frac{z}{\sqrt{x^2 + y^2 + z^2}} \end{cases}$$

mentre, noti ρ e ϑ si ottiene φ , con le relazioni

$$\cos \varphi = \frac{x}{\rho \sin \vartheta} \quad \text{e} \quad \sin \varphi = \frac{y}{\rho \sin \vartheta}$$

Un altro tipo di coordinate per individuare un punto P dello spazio, che, in un certo senso, è una via di mezzo tra quelle cartesiane e quelle polari è costituito da quelle che si chiamano *coordinate cilindriche*: un punto P è rappresentato dalle coordinate polari della sua proiezione Q sul piano xy (vedi figura 18.2) e dalla quota $z_0 = z$ di P .

Dunque si ha $\rho_0 = OQ$ e $\vartheta = \widehat{XOQ}$ detti rispettivamente *distanza orizzontale* e *azimut*. Le coordinate cilindriche di P sono allora (ρ_0, ϑ, z_0) e le formule di passaggio si ricavano immediatamente dalla definizione:

$$\begin{cases} x = \rho_0 \cos \vartheta \\ y = \rho_0 \sin \vartheta \\ z = z_0 \end{cases}$$

e, all'inverso, $\rho_0 = \sqrt{x^2 + y^2}$ e, noto ρ_0 si può calcolare ϑ dalle relazioni $\cos \vartheta = \frac{x}{\rho_0}$ e $\sin \vartheta = \frac{y}{\rho_0}$.

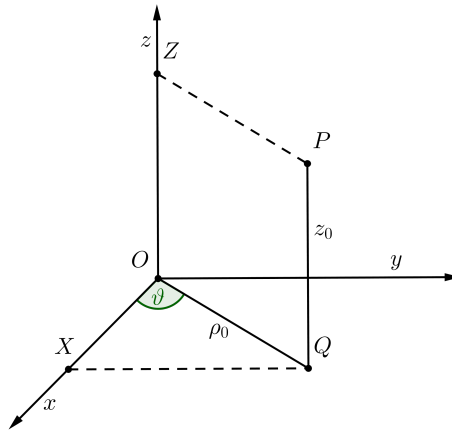


Figura 18.2 Coordinate cilindriche

Capitolo 19

Linee e Superfici nello spazio

In questo capitolo vedremo come si rappresentano in generale linee e superfici nello spazio.

19.1 Superfici

DEFINIZIONE 19.1. Si chiama *superficie* il luogo dei punti dello spazio, che riferito ad un sistema di coordinate cartesiane ortogonali, soddisfano un'equazione cartesiana

$$f(x, y, z) = 0 \quad (19.1)$$

od una terna di equazioni parametriche del tipo

$$\begin{cases} x = g(u, v) \\ y = h(u, v) \\ z = i(u, v) \end{cases} \quad (19.2)$$

dove nella (19.1) la f è funzione reale di al più tre variabili, in genere reali e nella (19.2) g , h ed i sono tre funzioni reali di al più due variabili.

Nella definizione 19.1 si parla di funzioni qualsiasi, che quindi possono dar luogo a superfici anche molto lontane dal concetto intuitivo che abbiamo di superficie.

La superficie più semplice è il piano che, come abbiamo visto, è rappresentato da un'equazione cartesiana lineare, $ax + by + cz + d = 0$ purché i coefficienti a , b e c non siano tutti e tre nulli, ma anche da una terna di equazioni parametriche lineari

$$\begin{cases} x = x_0 + \alpha u + \beta v \\ y = y_0 + \gamma u + \delta v \\ z = z_0 + \epsilon u + \zeta v \end{cases} \quad (19.3)$$

con α, γ, ϵ non tutti e tre nulli ed ugualmente non contemporaneamente nulla la terna (β, δ, ζ) . Le equazioni (19.2) rappresentano effettivamente un piano, perché eliminando da esse i due parametri u e v si ottiene una equazione (ed una sola) lineare.

OSSERVAZIONE 19.1. Mentre l'equazione cartesiana di un piano è sempre lineare e, viceversa in un sistema di riferimento cartesiano ogni equazione lineare rappresenta un piano, le equazioni parametriche di un piano possono anche non essere lineari, addirittura non algebriche.

Esempio 19.1. Mostriamo infatti che il sistema

$$\begin{cases} x = u^3 + \log v \\ y = 2u^3 + 3 \log v \\ z = 3u^3 + 7 \log v \end{cases}$$

rappresenta un piano. Per far ciò basta eliminare i due parametri u e v ricavando u^3 e $\log v$ dalle prime due equazioni: si ottiene $u^3 = 3x - y$ e, similmente, $\log v = y - 2x$; sostituendo poi i valori trovati nella terza equazione, si perviene all'equazione cartesiana $5x - 4y + z = 0$ che rappresenta appunto un piano.

Per la precisione in questo caso si dovrebbe parlare di un piano privato di una semiretta, infatti, poiché v è argomento di logaritmo, dev'essere $v > 0$ e cioè $y - 2x > 0$ quindi bisognerebbe scrivere $\begin{cases} 5x - 4y + z = 0 \\ y - 2x > 0 \end{cases}$.

Tra le superfici di equazione (19.1 nella pagina precedente) esistono, ad esempio, anche quelle rappresentate dall'equazione

$$x^2 + y^2 + z^2 = k \quad (19.4)$$

che per $k > 0$ rappresentano una superficie reale, quella formata da tutti (e soli) i punti che distano $\sqrt{k}$ dall'origine: si tratta di una sfera, superficie di cui parleremo diffusamente nel prossimo capitolo. Ma se $k = 0$ si ottiene una superficie che ha un solo punto reale, l'origine $O(0, 0, 0)$ e se $k < 0$ nessun punto della superficie, che per comodità chiamiamo ancora sfera, è reale.

Quindi anche nel caso dello spazio occorre accettare superfici formate solo da punti a coordinate complesse, ossevando, di passata, che anche le superfici reali contengono, in generale punti a coordinate complesse.

L'equazione cartesiana di una superficie, equazione 19.1 a pagina 195, si può, sotto certe ipotesi¹ esplicitare, cioè scrivere nella forma

$$z = \varphi(x, y) \quad (19.5)$$

che in molte occasioni può essere più comoda.

DEFINIZIONE 19.2. Una superficie si dice *algebraica di ordine n* se può essere rappresentata da un'equazione di tipo (19.1) in cui il primo membro è un polinomio di grado n . Le superfici che non possono essere rappresentate in questo modo si chiamano *trascendenti*.

Quindi, per esempio, esaminando ancora un caso patologico, e seguendo la definizione 19.2 la superficie di equazione $x = 0$ deve essere considerata diversa dalla superficie $x^2 = 0$ pur essendo entrambe formate dagli stessi punti.

Sia ora $f(x, y, z) = 0$ una superficie algebrica di ordine n e sia r una generica retta di equazioni parametriche

$$\begin{cases} x = x_0 + at \\ y = y_0 + bt \\ z = z_0 + ct \end{cases}.$$

Le intersezioni di r con la superficie sono ovviamente le soluzioni dell'equazione

$$f(x_0 + at, y_0 + bt, z_0 + ct) = 0 \quad (19.6)$$

che è un polinomio uguagliato a zero nella variabile t e di grado non maggiore di n , il quale, per il Teorema Fondamentale dell'Algebra, ammette $m \leq n$ radici, reali o complesse contate con le opportune molteplicità. Quindi possiamo concludere che una retta ed una superficie algebrica di ordine n hanno al più n punti in comune.

Una retta si dice *tangente* alla superficie se l'equazione (19.6) risolvente l'intersezione ammette almeno una radice multipla. Questa radice è costituita dalle coordinate del punto di tangenza.

Una retta tangente può ovviamente intersecare la superficie anche in altri punti diversi da quello di tangenza o addirittura essere tangente in più punti.

DEFINIZIONE 19.3. Un punto P di una superficie algebrica Σ si dice *multiplo* per la superficie Σ se ogni retta passante per P è tangente a Σ .

¹che qui non importa precisare: si tratta di un noto teorema di Analisi sulle funzioni implicite di più variabili.

19.2 Linee

Una linea (o curva) $\mathcal{L}$ nello spazio può essere rappresentata da equazioni parametriche del tipo

$$\mathcal{L} \equiv \begin{cases} x = f(t) \\ y = g(t) \\ z = h(t) \end{cases} \quad (19.7)$$

che individuano le coordinate del generico punto di $\mathcal{L}$ al variare del parametro t . Se proviamo ad eliminare il parametro dalle equazioni (19.7) otteniamo sempre *due* equazioni cartesiane, quindi una curva può anche essere rappresentata come intersezione di due superfici.

Se tutti i punti di una curva appartengono ad un medesimo piano la curva si dice *piana*, nel caso opposto si parla di curva *sgheмба* o *gobba*. Un esempio di curva gobba è fornito dal filetto di una comune vite.

Per stabilire se una linea è piana oppure gobba, si può procedere in vari modi, a seconda di come è rappresentata la linea:

- Se la linea è individuata dall'intersezione di due superfici, cioè da un sistema di due equazioni cartesiane, allora si può cercare un sistema equivalente che contenga un'equazione lineare. Se lo si trova allora si può affermare che la curva è piana e l'equazione lineare trovata è quella del piano che la contiene. Ovviamente il fatto di non riuscire a trovare il sistema più semplice non sempre garantisce che non esista e quindi non garantisce che la curva sia gobba.
- Se la curva, invece, è data mediante le sue equazioni parametriche (caso peraltro a cui ci si può sempre ricondurre, ed in generale con relativa facilità) basta sostituire le coordinate del punto generico della curva (19.7) nell'equazione del generico piano $ax + by + cz + d = 0$ ottenendo

$$a \cdot f(t) + b \cdot g(t) + c \cdot h(t) + d = 0 \quad (19.8)$$

La curva è piana se e solo se la (19.8) è identicamente verificata, cioè è verificata da *ogni valore* del parametro t .

Nel caso in cui il primo membro dell'equazione (19.8) sia un polinomio di grado n la condizione è verificata se e solo se ogni coefficiente del polinomio (compreso quindi anche il termine noto) è nullo. Ponendo uguali a zero tutti i coefficienti si ottiene un

sistema lineare omogeneo di $k \leq n + 1$ equazioni nelle quattro incognite a, b, c e d . La curva è piana se e solo se un tale sistema ammette autosoluzioni, cioè se e solo se la matrice dei suoi coefficienti ha rango < 4 .

Esempio 19.2. Sia data la curva di equazioni

$$\begin{cases} x^2 + y^2 + z^2 - 3x + 2y + z - 1 = 0 \\ x^2 + y^2 + z^2 - x + 1 = 0 \end{cases}$$

sottraendo membro a membro si ottiene il sistema equivalente

$$\begin{cases} x^2 + y^2 + z^2 - 3x + 2y + z - 1 = 0 \\ -2x + 2y + z + 2 = 0 \end{cases}$$

quindi la curva è piana ed appartiene al piano $2x - 2y - z - 2 = 0$.

Esempio 19.3. Sia data la curva di equazioni parametriche

$$\begin{cases} x = t^2 + t \\ y = t - 1 \\ z = t^2 - 2t \end{cases}.$$

Vogliamo vedere se è piana.

In questo caso la (19.8) diventa $a(t^2 + t) + b(t - 1) + c(t^2 - 2t) + d = 0$ che, ordinando il polinomio, diventa $(a + c)t^2 + (a + b - 2c)t + d - b = 0$, identicamente verificata se e solo se il sistema

$$\begin{cases} a + c = 0 \\ a + b - 2c = 0 \\ d - b = 0 \end{cases}$$

ammette autosoluzioni, il che accade, in quanto il rango della matrice dei coefficienti è palesemente non maggiore di 3. L'equazione del piano si ottiene facilmente: i coefficienti sono, ordinatamente, una di queste autosoluzioni.

Esempio 19.4. Sia $\mathcal{L}$ la linea

$$\begin{cases} x = t^3 \\ y = t^3 + t^2 \\ z = t \end{cases}$$

si ha $at^3 + b(t^3 + t^2) + ct + d = 0$ cioè $(a + b)t^3 + bt^2 + ct + d = 0$ da cui il sistema $\begin{cases} a + b = 0 \\ b = 0 \\ c = 0 \\ d = 0 \end{cases}$ che non ammette autosoluzioni, quindi possiamo concludere che la curva non è piana.

Capitolo 20

Sfera e circonferenza nello spazio

20.1 La sfera

Come è noto si chiama *sfera* il luogo dei punti dello spazio equidistanti da un punto fissato. Se tale punto è $C(\alpha, \beta, \gamma)$ e la distanza è $R \geq 0$ si vede subito che, indicando con $P(x, y, z)$ un generico punto dello spazio, l'equazione della sfera è

$$(x - \alpha)^2 + (y - \beta)^2 + (z - \gamma)^2 = R^2 \quad (20.1)$$

che diventa

$$x^2 + y^2 + z^2 - 2\alpha x - 2\beta y - 2\gamma z + \alpha^2 + \beta^2 + \gamma^2 - R^2 = 0$$

che si può scrivere

$$x^2 + y^2 + z^2 + ax + by + cz + d = 0$$

pur di porre $a = -2\alpha$, $b = -2\beta$, $c = -2\gamma$ e $d = \alpha^2 + \beta^2 + \gamma^2 - R^2$.

Viceversa l'equazione

$$x^2 + y^2 + z^2 + ax + by + cz + d = 0 \quad (20.2)$$

caratterizzata dal fatto di avere i coefficienti dei termini quadratici uguali e non contenere alcun termine rettangolare rappresenta la sfera con

centro nel punto $C = \left(-\frac{a}{2}, -\frac{b}{2}, -\frac{c}{2}\right)$ e raggio $R = \frac{1}{2}\sqrt{a^2 + b^2 + c^2 - 4d}$

come si può facilmente mostrare con ragionamenti analoghi a quelli effettuati per la circonferenza nel piano. Ampliando lo spazio con i punti a coordinate complesse, come abbiamo fatto per il piano, possiamo eliminare nella (20.1) l'eccezione $R \geq 0$ ed affermare che *ogni equazione della forma (20.2) rappresenta una sfera nello spazio.*

20.2 Piani tangenti ad una sfera

Esaminiamo ora, su esempi, come si possa determinare l'equazione del piano tangente¹ ad una sfera in un suo punto P .

Esempio 20.1. Sia data la sfera $\Sigma \equiv x^2 + y^2 + z^2 - 2x + 2y = 0$ che ha centro nel punto $C(1, -1, 0)$ e raggio $R = \sqrt{2}$. Vogliamo l'equazione del piano π tangente nell'origine $O(0, 0, 0)$ alla Σ . Il piano cercato è perpendicolare alla retta OC che ha parametri direttori $1, -1, 0$, quindi possiamo scegliere questi parametri direttori per π ; poiché, infine π passa per l'origine, la sua equazione è: $x - y = 0$.

Per inciso, si può mostrare che il piano tangente nell'origine ad una sfera² è rappresentato dall'equazione che si ottiene eguagliando a zero il complesso dei termini lineari dell'equazione che la rappresenta.

Esempio 20.2. Siano Σ la sfera di equazione $x^2 + y^2 + z^2 = 1$ ed r la retta di equazioni $\begin{cases} x = 3 \\ z = 0 \end{cases}$. Vogliamo gli eventuali piani tangenti a Σ e passanti per r .

Si vede subito che la sfera data ha centro nell'origine e raggio $R = 1$ i piani cercati appartengono allora al fascio che ha per sostegno la retta r , cioè al fascio $x + \lambda z - 3 = 0$ ed hanno distanza 1 dall'origine, quindi dev'essere $\frac{3}{\sqrt{1+\lambda^2}} = 1$ da cui si ricava $\lambda = \pm 2\sqrt{2}$; in accordo con il fatto che i piani tangenti ad una sfera e passanti per una retta sono al più due.

20.3 Circonferenze nello spazio

L'intersezione di una sfera Σ con un piano π rappresenta una circonferenza, reale, se il piano taglia la sfera in punti reali, cioè se la sua distanza dal centro è minore del raggio R , ridotta ad un solo punto se il piano è tangente, cioè se la distanza è R e immaginaria se il piano non taglia la sfera, cioè se la distanza è maggiore di R , e viceversa ogni circonferenza dello spazio può essere vista come l'intersezione di un piano con una sfera. Quindi, in generale, una circonferenza si può rappresentare nello spazio con le equazioni

$$\gamma \equiv \begin{cases} x^2 + y^2 + z^2 + ax + by + cz + d = 0 \\ \alpha x + \beta y + \gamma z + \delta = 0 \end{cases} \quad (20.3)$$

¹Il piano tangente ad una superficie in un suo punto P può essere inteso come il luogo delle rette tangenti in P alla superficie.

²anzi, ad una qualsiasi superficie algebrica.

Ci proponiamo di calcolare le coordinate del centro e la lunghezza del raggio di γ . Se indichiamo con C e R rispettivamente il centro ed il raggio della sfera Σ , è facile rendersi conto (vedi Figura 20.1) che il centro C' della γ è la proiezione ortogonale di C sul piano π , quindi si può determinare come intersezione tra il piano π

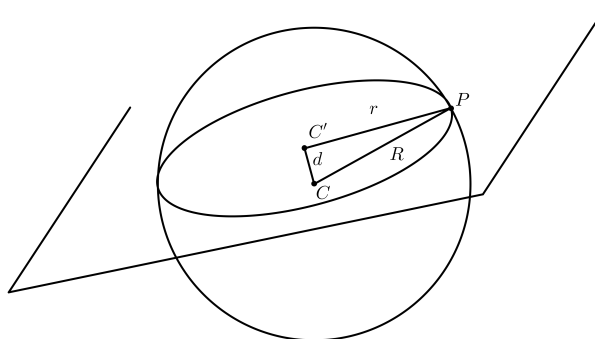


Figura 20.1 Circonferenza nello spazio

e la retta ad esso perpendicolare passante per C . Per quanto riguarda il raggio r della circonferenza, osservando sempre la figura 20.1, il Teorema di Pitagora ci permette di scrivere che $R^2 = r^2 + d^2$, dove con d abbiamo indicato la distanza di C dal piano π da cui immediatamente $r = \sqrt{R^2 - d^2}$.

Esempio 20.3. Siano date le equazioni

$$\begin{cases} x^2 + y^2 + z^2 - 2x + 2y - z = 0 \\ x + y - z = 0 \end{cases}$$

che rappresentano l'intersezione della sfera di centro $C \left(1, -1, \frac{1}{2}\right)$ e raggio

$R = \frac{3}{2}$ con il piano $x + y - z = 0$; esse rappresentano una circonferenza reale, infatti il piano taglia la sfera in punti reali dato che la distanza d di π da C è $\frac{|1 - 1 - \frac{1}{2}|}{\sqrt{3}}$, quindi $d = \frac{1}{2\sqrt{3}} < R$. Dunque $r = \sqrt{\frac{9}{4} - \frac{1}{12}} = \frac{1}{2}\sqrt{\frac{26}{3}}$. Per trovare le coordinate del centro di γ osserviamo che una terna di parametri direttori del piano è $1, 1, -1$, quindi la retta perpendicolare al piano passante

per C avrà equazioni $\begin{cases} x = 1 + t \\ y = -1 + t \\ z = \frac{1}{2} - t \end{cases}$ intersecandola con il piano si ottiene $1 +$

$t - 1 + t - \frac{1}{2} + t = 0$ da cui $t = \frac{1}{6}$ cui corrisponde il punto $C' \left(\frac{7}{6}, -\frac{5}{6}, \frac{1}{3}\right)$ centro della circonferenza data.

20.4 Fasci di sfere

L'insieme $\mathcal{F}$ di tutte e sole le sfere che passano per una data circonferenza γ (reale o completamente immaginaria, degenerare in un punto o meno) e dal piano che la contiene costituisce quello che si chiama un *fascio di sfere*³ il piano che ne fa parte si chiama *piano radicale del fascio* e la circonferenza data si chiama *sostegno* del fascio. Il luogo dei centri delle sfere di $\mathcal{F}$ risulta ovviamente costituito dalla retta che passa per il centro di γ ed è ortogonale al piano su cui γ giace.

In modo del tutto analogo ai fasci di piani si dimostra che *l'equazione di tutti e soli gli elementi di un fascio di sfere si scrive come combinazione lineare non banale di quelle di due qualsiasi di esse*.

Esempio 20.4. Sia data la circonferenza γ intersezione delle due sfere

$$\begin{cases} x^2 + y^2 + z^2 - 2x - 2y = 0 \\ x^2 + y^2 + z^2 + x - 2y + z - 1 = 0 \end{cases}$$

Il fascio che ha per sostegno la γ ha equazione

$$\lambda(x^2 + y^2 + z^2 - 2x - 2y) + \mu(x^2 + y^2 + z^2 + x - 2y + z - 1) = 0$$

che, come al solito e con le solite avvertenze sull'eventuale valore infinito del parametro, si può scrivere, con un parametro solo, nella forma

$$x^2 + y^2 + z^2 - 2x - 2y + k(x^2 + y^2 + z^2 + x - 2y + z - 1) = 0 \quad (20.4)$$

che diventa

$$(k+1)(x^2 + y^2 + z^2) + (k-2)x + -2(k+1)y + kz - k = 0.$$

Per $k = -1$ la (20.4) diventa $3x + z - 1 = 0$ che è l'equazione del piano radicale del fascio. Essa si ottiene comunque sottraendo membro a membro le equazioni della γ .

Esempio 20.5. Vogliamo l'equazione di una sfera che ha raggio $R = \frac{\sqrt{6}}{2}$ e passa per la circonferenza di equazioni

$$\gamma : \begin{cases} x^2 + y^2 + z^2 - x = 1 \\ x + y - z + 1 = 0 \end{cases}.$$

³ Anche qui, come nel caso delle circonferenze del piano, si considerano anche altri tipi di fasci: i fasci costituiti dalle sfere tangenti un piano dato in un punto dato (e qui la circonferenza sostegno ha raggio nullo) o le sfere che hanno un dato centro.

Le sfere che soddisfano tali condizioni son al massimo due ed appartengono al fascio $\mathcal{F}$ che ha per sostegno la γ , la cui equazione è

$$x^2 + y^2 + z^2 - x - 1 + k(x + y - z + 1) = 0$$

L'equazione della generica sfera di $\mathcal{F}$ si può scrivere come

$$x^2 + y^2 + z^2 + (k - 1)x + ky - kz + k - 1 = 0$$

e quindi il suo raggio è

$$R = \frac{1}{2} \sqrt{(k - 1)^2 + k^2 + k^2 - 4(k - 1)}$$

Uguagliando R al valore dato si ottiene l'equazione

$$\frac{6}{4} = \frac{1}{4}(k^2 - 2k + 1 + 2k^2 - 4k + 4)$$

che, con semplici passaggi, diventa $3k^2 - 6k - 1 = 0$ da cui si ottengono i due valori di k .

Capitolo 21

Superfici rigate e di rotazione

21.1 Superfici rigate

Si dice *rigata* una superficie Σ tale che per ogni suo punto passi almeno una retta reale tutta contenuta in Σ . Le rette che formano la Σ si chiamano *generatrici* ed ogni linea che appartiene a Σ ed incontra ogni generatrice si chiama (*curva*) *direttrice*. Per esempio è rigata la superficie Σ di equazione

$$xy - z^2 + 1 = 0 \quad (21.1)$$

infatti l'equazione (21.1) si può scrivere

$$x \cdot y = (z + 1)(z - 1)$$

che ammette le stesse soluzioni dell'insieme costituito dai due sistemi di rette

$$\begin{cases} x = k(z - 1) \\ ky = z - 1 \end{cases} \quad \begin{cases} x = h(z + 1) \\ ky = z + 1 \end{cases};$$

al variare dei parametri non nulli h e k queste sono rette che giacciono per intero sulla superficie e si verifica che per ogni punto della (21.1) passa almeno una di queste rette, quindi la superficie è rigata.

È facile scrivere le equazioni parametriche di una superficie rigata quando siano date una direttrice $\mathcal{L}$ ed un sistema di generatrici: se le equazioni della direttrice sono

$$\mathcal{L} = \begin{cases} x = f(t) \\ y = g(t) , \\ z = h(t) \end{cases}$$

per ogni punto $P_t(f(t), g(t), h(t))$ di essa passa una generatrice i cui parametri direttori $a(t), b(t), c(t)$ dipendono da P_t . La generatrice

passante per P_t ha quindi equazioni

$$\begin{cases} x = f(t) + a(t)\tau \\ y = g(t) + b(t)\tau \\ z = h(t) + c(t)\tau \end{cases} \quad (21.2)$$

Il sistema (21.2), pensato come sistema nella sola incognita τ rappresenta la generica generatrice, mentre pensato come sistema nelle due incognite t e τ rappresenta la superficie rigata cercata.

Naturalmente per avere l'equazione cartesiana della superficie bisogna eliminare i due parametri t e τ , cosa che può comportare conti a volte un po' laboriosi.

In particolare se f, g ed h sono funzioni lineari di t ed a, b e c funzioni costanti, si ha a che fare con le equazioni parametriche di un piano: lasciamo come esercizio al lettore la costruzione di un esempio in cui si possano chiaramente individuare la direttrice e le generatrici di una superficie rigata costituita da un piano.

Esempio 21.1. Cerchiamo l'equazione della superficie rigata che ha come direttrice la "cubica gobba" di equazioni¹

$$\begin{cases} x = t \\ y = t^2 \\ z = t^3 \end{cases}$$

e come generatrici rette di parametri direttori 1, 1, 1.

Scriveremo allora il sistema

$$\begin{cases} x = t + \tau \\ y = t^2 + \tau \\ z = t^3 + \tau \end{cases}$$

che costituiscono una terna di equazioni parametriche della superficie cercata.

Eliminando i parametri per passare alla forma cartesiana della superficie otteniamo, dopo qualche calcolo, lungo ma non difficile, l'equazione

$$(x - y)^3 - (y - z)(x - 2y + z) = 0.$$

21.2 Superfici di rotazione

Si chiama *di rotazione* o *rotonda* una superficie Σ (fig. 21.1 nella pagina successiva) ottenuta facendo ruotare una linea $\mathcal{L}$ attorno ad una retta r , detta di solito *asse di rotazione*.

¹Questo esempio mostra che la direttrice può anche non essere una curva piana.

Per scrivere l'equazione di una tale superficie basta considerare il generico punto $P \in \mathcal{L}$ ed imporre che sul piano π passante per P ed ortogonale a r esso descriva una circonferenza con centro l'intersezione tra r e π . In questo modo si ottengono due equazioni (di una sfera e del piano π); ambedue queste equazioni dipendono dal parametro che individua la posizione di P sulla linea $\mathcal{L}$; eliminando questo parametro si ottiene l'equazione della superficie, come mostra il seguente esempio.

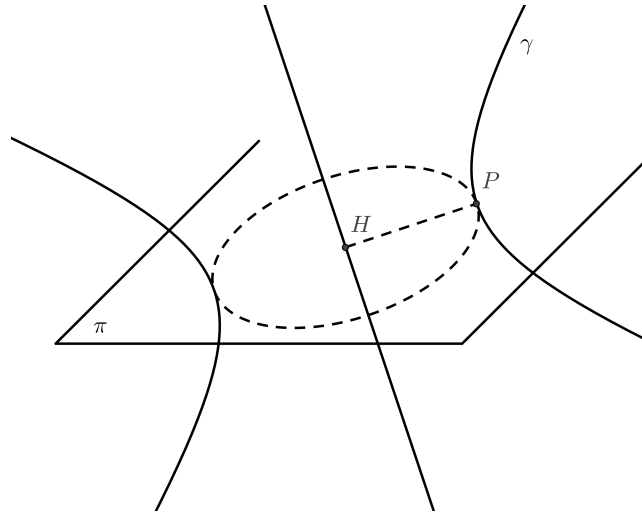


Figura 21.1 Superficie di rotazione

Esempio 21.2. Vogliamo l'equazione cartesiana della superficie che si ottiene facendo ruotare la linea $\mathcal{L}$ di equazioni parametriche

$$\begin{cases} x = (t - 1)^2 \\ y = 0 \\ z = t \end{cases} \quad (21.3)$$

attorno alla retta r di equazioni

$$\begin{cases} x = 1 \\ y = 1 \end{cases}.$$

Un punto generico dello spazio P appartiene alla $\mathcal{L}$ se e solo se le sue coordinate soddisfano la (21.3) e quindi sono $P((t - 1)^2, 0, t)$; il piano π passante per P e ortogonale ad r ha equazione $z = t$. La circonferenza è individuata tagliando il piano P con una qualunque sfera avente centro C su r e raggio CP . Prendendo $C(1, 1, 0)$ si ha l'equazione della sfera

$$(x - 1)^2 + (y - 1)^2 + z^2 = [(t - 1)^2 - 1]^2 + 1 + t^2$$

ottenendo così la circonferenza di equazioni

$$\begin{cases} (x - 1)^2 + (y - 1)^2 + z^2 = [(t - 1)^2 - 1]^2 + 1 + t^2 \\ z = t \end{cases}$$

In questo caso è particolarmente semplice l'eliminazione del parametro, dopo la quale, con opportune semplificazioni, si ottiene l'equazione

$$(x - 1)^2 + (y - 1)^2 + z^2(z - 2)^2 - 1 = 0$$

che rappresenta la superficie di rotazione cercata.

Capitolo 22

Cilindri, coni e proiezioni

22.1 Coni

DEFINIZIONE 22.1. Fissati nello spazio un punto V ed una curva algebrica¹ $\mathcal{L}$, si chiama *cono*² di vertice V e direttrice $\mathcal{L}$ l'insieme di tutte e sole le rette passanti per V e secanti la $\mathcal{L}$ (Figura 22.1), quindi il cono è una superficie rigata.

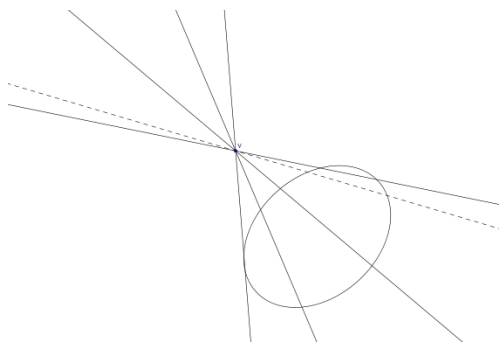


Figura 22.1 Cono

Ricordiamo che un'equazione in tre variabili $f(x, y, z) = 0$ si dice *omogenea di grado k* se $\forall t \in \mathbb{R}$ e $\forall (x, y, z) \in \mathbb{R}^3$ accade che

$$f(tx, ty, tz) = t^k f(x, y, z)$$

il che equivale a dire, se essa è rappresentata da un polinomio eguagliato a zero, che il polinomio $f(x, y, z)$ è costituito da monomi tutti dello stesso grado. Le equazioni omogenee in tre variabili

hanno una particolare importanza, come appare dal seguente

Teorema 22.1. *Le equazioni omogenee intere nelle incognite x, y, z rappresentano tutti e soli i coni con il vertice nell'origine.*

Dimostrazione. Infatti, supponiamo che l'equazione $f(x, y, z) = 0$ sia omogenea e che il punto $P(x_0, y_0, z_0)$ appartenga alla superficie da essa

¹cioè rappresentabile mediante l'intersezione di due superfici algebriche.

²La nozione di cono si potrebbe estendere anche al caso di direttrici non algebriche, ma in questo caso occorrerebbe fare alcune precisazioni che appesantirebbero la trattazione.

rappresentata, cioè sia tale che $f(x_0, y_0, z_0) = 0$; allora, poiché la f è omogenea, è anche $f(tx_0, ty_0, tz_0) = 0 \forall t \in \mathbb{R}$, quindi ogni punto della retta

$$\begin{cases} x = x_0 t \\ y = y_0 t, \\ z = z_0 t \end{cases}$$

che congiunge P con l'origine, appartiene alla superficie, la quale risulta quindi un cono con vertice nell'origine.

Viceversa se $f(x, y, z) = 0$ rappresenta un cono con vertice nell'origine, vuol dire che quando $P(x_0, y_0, z_0)$ appartiene al cono, ad esso appartiene l'intera retta che congiunge P con O e cioè la retta di equazioni

$$\begin{cases} x = x_0 t \\ y = y_0 t, \\ z = z_0 t \end{cases}$$

il che significa che $f(x_0, y_0, z_0) = 0 = f(tx_0, ty_0, tz_0) = 0 \forall t \in \mathbb{R}$ e $\forall (x, y, z) \in R$, quindi la f è omogenea. $\square$

Segue immediatamente il più generale

Corollario 22.2. *Le equazioni omogenee nelle tre incognite $x - \alpha$, $y - \beta$, $z - \gamma$ rappresentano tutti e soli i coni con vertice nel punto $V(\alpha, \beta, \gamma)$.*

Dimostrazione. Basta operare la traslazione d'assi che porta l'origine nel punto V . $\square$

Abbiamo visto che un cono è dato quando siano dati il vertice ed una direttrice, vediamo come sfruttare questo fatto per scriverne l'equazione.

Sia $V(a, b, c)$ il vertice e $\mathcal{L}$ la direttrice di equazioni

$$\begin{cases} f(x, y, z) = 0 \\ g(x, y, z) = 0 \end{cases}.$$

Se $P(x_0, y_0, z_0)$ è un punto del cono esso dovrà stare su una retta che passa per V e che taglia la direttrice. La generica retta per V e per P ha equazioni parametriche

$$\begin{cases} x = a - (x_0 - a)t \\ y = b - (y_0 - b)t \\ z = c - (z_0 - c)t \end{cases} \quad (22.1)$$

Per imporre il fatto che la retta (22.1) tagli la direttrice, basta sostituire nelle equazioni della $\mathcal{L}$ i valori dati dalle (22.1) ed eliminare il parametro dalle due equazioni.

Esempio 22.1. Vogliamo l'equazione del cono che ha vertice in $V(0, 1, 0)$ e come direttrice la circonferenza di equazioni

$$\begin{cases} x^2 + y^2 + z^2 - 2x = 0 \\ x - 2y + 3z = 0 \end{cases}.$$

Se $P(x_0, y_0, z_0)$ è un punto generico, esso sta sul cono anzitutto se appartiene alla retta PV che ha equazioni

$$\begin{cases} x = x_0 t \\ y = 1 + (y_0 - 1)t ; \\ z = z_0 t \end{cases}$$

poi la retta PV deve tagliare la direttrice, quindi si ha il sistema:

$$\begin{cases} (x_0 t)^2 + (1 + (y_0 - 1)t)^2 + (z_0 t)^2 - 2(x_0 t) = 0 \\ x_0 t - 2(1 + (y_0 - 1)t) + 3z_0 t = 0 \end{cases}.$$

Eliminando il parametro t dalle due equazioni si ottiene la relazione che deve intercorrere tra le coordinate di P affinché quest'ultimo stia su rette passanti per V che intersecano la direttrice, cioè affinché P stia sul cono cercato. Nel nostro caso, con facili conti si trova l'equazione $x^2 - 13z^2 + 8x(y - 1) - 6xz = 0$ che è effettivamente omogenea nelle differenze tra le coordinate x, y, z e le coordinate del vertice.

Quando il vertice è nell'origine è semplice verificare l'omogeneità, quando invece il vertice è un altro punto occorre spesso raccogliere opportunamente i termini.

Alcuni coni sono di rotazione, perché ottenuti dalla rotazione di una retta attorno ad un'altra ad essa incidente

Esempio 22.2. Siano r e s rispettivamente le rette di equazioni

$$r : \begin{cases} x = t \\ y = t \\ z = t \end{cases} \quad e \quad s : \begin{cases} x = -t \\ y = t \\ z = 3t \end{cases} ;$$

vogliamo l'equazione del cono che si ottiene facendo ruotare la retta r attorno alla s . Si può osservare che tale cono ha vertice $V \equiv O(0, 0, 0)$, che è il

punto comune a r e s ; per trovare una direttrice possiamo considerare il generico piano perpendicolare ad s e non passante per V , che ha equazione $x - y - 3z + k = 0$ con $k \neq 0$; esso interseca la s nel punto $S \left(-\frac{1}{5}k, \frac{1}{5}k, \frac{3}{5}k \right)$ e la r nel punto $R \left(\frac{1}{3}k, \frac{1}{3}k, \frac{1}{3}k \right)$ e tagliarlo con la sfera che ha centro in S e raggio uguale a $\overline{RS}$, determinabile facilmente con il teorema di Pitagora (vedi figura ??). Prendendo un valore comodo di k (per esempio qui se si prende $k = 15$ si ottengono coordinate intere) si ricavano facilmente le equazioni della direttrice. Oppure tale cono si può scrivere come superficie di rotazione, facendo appunto ruotare la retta r attorno alla retta s .

OSSERVAZIONE 22.1. Facendo ruotare una retta attorno ad un'altra retta, otteniamo un cono se e solo se le rette sono incidenti; se sono parallele otteniamo, come vedremo tra poco, un cilindro, se sono sghembe una superficie del second'ordine detta iperboloide ad una falda di cui parleremo diffusamente nel capitolo 23 sulle quadriche.

22.2 Cilindri

DEFINIZIONE 22.2. Fissata nello spazio una curva algebrica γ ed una retta r che interseca γ si chiama *cilindro* di direttrice γ e generatrice r la superficie (rigata) formata da tutte e sole le rette parallele ad r che intersecano la γ . (Figura 22.2 a fronte).

Quindi il cilindro è ben di più del cilindro circolare che siamo abituati a conoscere.

OSSERVAZIONE 22.2. Si vede subito che ogni equazione del tipo

$$f(x, y) = 0 \quad (22.2)$$

rappresenta, nello spazio, un cilindro con le generatrici parallele all'asse z ed avente come generatrice sul piano xy la curva di equazioni

$$\begin{cases} f(x, y) = 0 \\ z = 0 \end{cases};$$

infatti si verifica subito che, se $P(x_0, y_0)$ è un punto le cui coordinate soddisfano l'equazione (22.2), allora (e solo allora) qualunque punto della retta

$$\begin{cases} x = x_0 \\ y = y_0 \end{cases}$$

(che è la generica retta parallela all'asse z passante per P) soddisfa l'equazione (22.2).

Quindi, generalizzando l'osservazione 22.2, possiamo dire che *un'equazione in due variabili rappresenta sempre, nello spazio, un cilindro con le generatrici parallele all'asse che ha lo stesso nome della variabile che manca.*

Per scrivere l'equazione di un cilindro con le generatrici in direzione generica, cioè non necessariamente parallele ad uno degli assi coordinati, si può sfruttare la definizione, vale a dire che si può pensare al fatto che un generico punto $P(x_0, y_0, z_0)$ dello spazio appartiene al cilindro se e solo se sta su una retta parallela alla generatrice che interseca la curva direttrice.

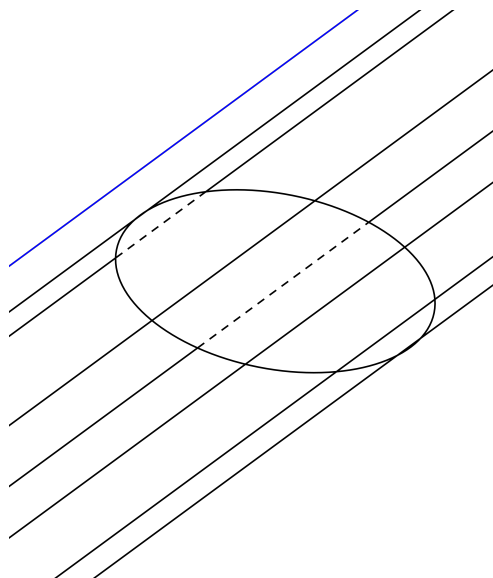


Figura 22.2 Cilindro

Esempio 22.3. Vogliamo l'equazione del cilindro con le generatrici parallele alla retta r di equazioni $x = y = z$ che taglia il piano xy sulla parabola di equazione $y^2 = x$ avente quindi come direttrice la curva

$$\begin{cases} y^2 = x \\ z = 0 \end{cases}.$$

Sia $P(x_0, y_0, z_0)$ un generico punto dello spazio; la retta passante per P e parallela ad r ha equazioni

$$\begin{cases} x = x_0 + t \\ y = y_0 + t \\ z = z_0 + t \end{cases}.$$

Il punto P appartiene al cilindro se e soltanto se quest'ultima retta taglia la parabola. Quindi deve essere verificato il sistema

$$\begin{cases} (y_0 + t)^2 = x_0 + t \\ z_0 + t = 0 \end{cases}.$$

Eliminando, con semplici calcoli, il parametro t dalle due equazioni del sistema si ottiene la relazione che deve intercorrere tra le coordinate di P perché questi appartenga al cilindro, che è quindi

$$(y - z)^2 = x - z.$$

OSSERVAZIONE 22.3. Esistono, ovviamente, cilindri rotondi, cioè che hanno una direttrice formata da una circonferenza. La loro equazione si può anche scrivere come quella di una superficie di rotazione.

OSSERVAZIONE 22.4. Non sempre i calcoli per l'eliminazione del parametro sono così immediati come nell'esempio 22.3 nella pagina precedente: a volte possono essere piuttosto lunghi e laboriosi.

22.3 Proiezioni

DEFINIZIONE 22.3. Si chiama *proiezione di una curva γ da un punto V sul piano α* l'intersezione tra il piano α stesso ed il cono avente vertice in V e come direttrice γ . (v. figura 22.3)

Si chiama *proiezione di una curva γ dalla direzione della retta r sul piano α* (v. fig. 22.4 a pagina 218) l'intersezione tra il piano α stesso ed il cilindro avente γ come generatrice e generatrici parallele a r .

Dalla definizione 22.3 escludiamo, per evitare casi patologici, che $V \in \alpha$ oppure r sia parallela ad α .

Esempio 22.4. Cerchiamo le equazioni della proiezione della curva

$$\gamma \equiv \begin{cases} x = y^2 + 1 \\ z = x + 1 \end{cases}$$

dall'origine $O(0,0,0)$ sul piano $x = 1$.

Il cono che ha vertice nell'origine e come direttrice la γ ha equazione (verificarlo per esercizio)

$$y^2 + x(x - z) + (x - z)^2 = 0$$

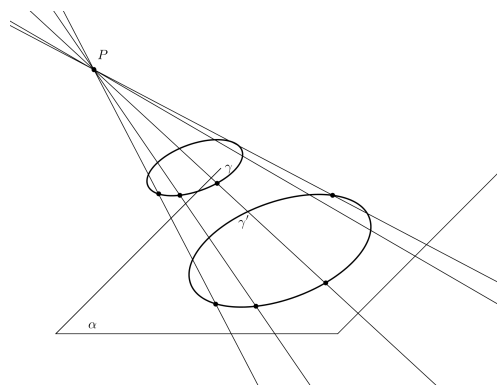


Figura 22.3 Proiezione centrale

evidentemente omogenea nelle tre variabili; quindi la curva proiezione può essere individuata dal sistema

$$\begin{cases} y^2 + x(x - z) + (x - z)^2 = 0 \\ x = 1 \end{cases}.$$

Se sostituiamo nella prima equazione il valore di x dato dalla seconda, otteniamo la stessa curva, rappresentata però come intersezione di un cilindro con le generatrici perpendicolari al piano:

$$\begin{cases} y^2 + z^2 - 3z + 2 = 0 \\ x = 1 \end{cases}.$$

22.4 Riconoscimento di una conica nello spazio

Si può mostrare che nello spazio ogni conica irriducibile può essere vista come sezione piana di un cono circolare. Più precisamente se tagliamo un cono rotondo tale che la generatrice formi un angolo α con l'asse di rotazione con un piano non passante per il vertice e che formi un angolo β con lo stesso asse di rotazione otteniamo rispettivamente

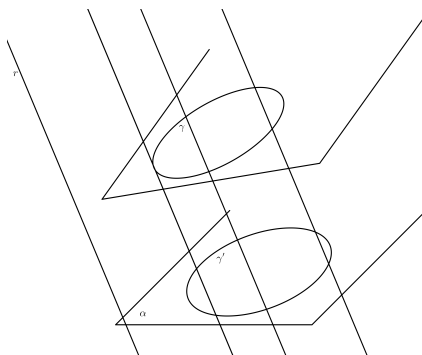
- una parabola se $\alpha = \beta$
- un'ellisse se $\alpha < \beta$
- un'iperbole se $\alpha > \beta$

Si può inoltre mostrare che qualunque sezione piana di una superficie del secondo ordine è una conica, eventualmente degenera. Per riconoscerla, si può usare il seguente teorema, che enunciamo senza dimostrazione.

Teorema 22.3. *Ogni proiezione parallela di una conica non ne altera la natura, se effettuata da una direzione non parallela al piano su cui giace la conica stessa.*

Quindi per studiare la natura di una conica nello spazio basta proiettarla, se possibile, su uno dei piani coordinati³.

³ovviamente se la conica giace già su un piano parallelo ad uno dei piani coordinati questo non è possibile, ma in tal caso il riconoscimento procede come già abbiamo visto nel piano.

**Figura 22.4** Proiezione parallela

Esempio 22.5. Vogliamo riconoscere la conica di equazioni:

$$\begin{cases} x^2 + xy + z^2 - x + 1 = 0 \\ x - y - z = 0 \end{cases}$$

Eliminando dal sistema una delle incognite, per esempio la z , si ottiene l'equazione del cilindro che proietta la conica ortogonalmente al piano xy , che ha quindi equazione $x^2 + xy + (x - y)^2 - x + 1 = 0$, la proiezione sarà dunque.

$$\begin{cases} x^2 + xy + (x - y)^2 - x + 1 = 0 \\ z = 0 \end{cases}$$

che sul piano xy rappresenta un'ellisse; dunque la conica data è un'ellisse.

Capitolo 23

Superfici quadriche

In questo capitolo parleremo delle superfici algebriche del secondo ordine, cioè rappresentabili con equazioni polinomiali di secondo grado, classificandole, indicando le loro forme canoniche, ed alcuni procedimenti atti ad ottenere il loro riconoscimento.

23.1 Prime proprietà delle quadriche

DEFINIZIONE 23.1. Si chiama *superficie quadrica* o più brevemente *quadrica* ogni superficie rappresentata, nello spazio riferito ad un sistema di coordinate cartesiane ortogonali, da un'equazione di secondo grado nelle variabili x, y e z .

Per esempio abbiamo già visto che sono quadriche le sfere ed alcune superfici, per esempio quelle che si ottengono dalla rotazione di una retta intorno ad un'altra retta.

La più generale equazione di secondo grado in x, y e z ha la forma

$$ax^2 + by^2 + cz^2 + dxy + exz + fyz + gx + hy + iz + l = 0. \quad (23.1)$$

che dipende da 10 coefficienti omogenei (cioè definiti a meno di un fattore di proporzionalità non nullo) e quindi da 9 coefficienti non omogenei. Partendo dal fatto che per determinare questi nove coefficienti occorre e basta imporre 9 condizioni lineari indipendenti, si può dimostrare in modo analogo a come si è fatto per le coniche, che

Teorema 23.1. *Per 9 punti, a 4 a 4 non complanari, passa una ed una sola quadrica irriducibile.*

Come abbiamo fatto per le coniche nel piano, possiamo associare ad una quadrica una matrice simmetrica costruita a partire dai suoi

coefficienti:

$$A = \begin{bmatrix} a & \frac{d}{2} & \frac{e}{2} & \frac{g}{2} \\ \frac{d}{2} & b & \frac{f}{2} & \frac{h}{2} \\ \frac{e}{2} & \frac{f}{2} & c & \frac{i}{2} \\ \frac{g}{2} & \frac{h}{2} & \frac{i}{2} & l \end{bmatrix}$$

In tal modo possiamo riscrivere l'equazione di una generica quadrica (23.1 nella pagina precedente) nel seguente modo

$$\vec{x}A\vec{x}_T = 0 \quad (23.2)$$

essendo $\vec{x}$ il vettore $[x, y, z, 1]$.

Il rango della matrice A , invariante per rototraslazioni, ci fornisce molte informazioni sulla quadrica. Se il determinante di A è uguale a zero diciamo che la quadrica è *specializzata* con indice di specializzazione uguale a $4 - r(A)$: si dimostra che una quadrica non specializzata non ha punti multipli e che una quadrica specializzata con indice di specializzazione uguale a 1 (o, come si dice *specializzata una volta*, cioè per cui è $r(A) = 3$ è un cono o un cilindro quadrico. Se è un cono, ammette come unico punto multiplo il suo vertice, mentre se è un cilindro non possiede punti multipli al finito.

Si può inoltre provare che se l'indice di specializzazione è maggiore di 1, cioè se $r(A) \leq 2$ la quadrica è riducibile: più precisamente se $r(A) = 2$ la quadrica è spezzata in una coppia di piani distinti, e, se tali piani non sono paralleli, la superficie ammette come punti multipli tutti e soli i punti della retta comune ai due piani, mentre se $r(A) = 1$ si ha una coppia di piani coincidenti.

Esempio 23.1. *La quadrica di equazione*

$$x^2 + y^2 + z^2 + 2xy - 2xz - 2yz - 5x - 5y + 5z + 6 = 0$$

è spezzata. Infatti è facile vedere che la matrice

$$A = \begin{bmatrix} 1 & 1 & -1 & -\frac{5}{2} \\ 1 & 1 & -1 & -\frac{5}{2} \\ -1 & -1 & 1 & \frac{5}{2} \\ -\frac{5}{2} & -\frac{5}{2} & \frac{5}{2} & 6 \end{bmatrix}$$

ha rango 1. Del resto la sua equazione diventa facilmente

$$(x + y - z)^2 - 5(x + y - z) + 6 = 0$$

che si scompone in

$$[(x + y - z) - 2] [(x + y - z) - 3] = 0$$

che mette in luce i due piani paralleli in cui è spezzata.

23.2 Quadriche in forma canonica e loro classificazione

L'equazione di una qualunque quadrica si può ricondurre, con una opportuna rototraslazione del sistema di riferimento, ad una ed una sola delle forme canoniche elencate nella tabella 23.1 che riguarda le quadriche specializzate e 23.2 nella pagina seguente che riguarda quelle non specializzate.

Tabella 23.1 Forma canonica delle quadriche specializzate

irriducibili	
$\frac{x^2}{a^2} + \frac{y^2}{b^2} + \frac{z^2}{c^2} = 0$	cono quadrico immaginario
$\frac{x^2}{a^2} + \frac{y^2}{b^2} + \frac{z^2}{c^2} = 0$	cono quadrico reale
$\frac{x^2}{a^2} + \frac{y^2}{b^2} = 1$	cilindro quadrico a sezione ellittica
$\frac{x^2}{a^2} - \frac{y^2}{b^2} = 1$	cilindro quadrico a sezione iperbolica
$\frac{x^2}{a^2} + \frac{y^2}{b^2} = -1$	cilindro quadrico completamente immaginario
$y^2 = 2px, p = 0$	cilindro quadrico a sezione parabolica
riducibili	
$\frac{x^2}{a^2} - \frac{y^2}{b^2} = 0$	piani reali incidenti
$\frac{x^2}{a^2} + \frac{y^2}{b^2} = 0$	piani immaginari incidenti
$y^2 = a^2, a \neq 0$	piani reali paralleli
$y^2 = -a^2, a \neq 0$	piani immaginari paralleli
$y^2 = 0$	piani reali coincidenti

OSSERVAZIONE 23.1. In riferimento alle tabelle citate osserviamo che le quadriche dette *a centro* sono quelle che ammettono un centro di simmetria che, nella forma canonica, è l'origine del riferimento, mentre il paraboloide *a sella* è così chiamato a causa della sua caratteristica forma. A pagina 225 trovate le figure delle quadriche. La denominazione *completamente immaginaria* è riservata a superfici su cui non vi

Tabella 23.2 Forma canonica delle quadriche non specializzate

a centro	
$\frac{x^2}{a^2} + \frac{y^2}{b^2} + \frac{z^2}{c^2} = 1$	ellissoide reale
$\frac{x^2}{a^2} + \frac{y^2}{b^2} + \frac{z^2}{c^2} = -1$	ellissoide completamente immaginario
$\frac{x^2}{a^2} + \frac{y^2}{b^2} - \frac{z^2}{c^2} = 1$	iperboloide iperbolico (o ad una falda)
$\frac{x^2}{a^2} - \frac{y^2}{b^2} - \frac{z^2}{c^2} = 1$	iperboloide ellittico (a due falde)
non a centro	
$\frac{x^2}{a^2} + \frac{y^2}{b^2} = 2pz, p \neq 0$	paraboloide ellittico
$\frac{x^2}{a^2} - \frac{y^2}{b^2} = 2pz, p \neq 0$	paraboloide iperbolico (o a sella)

sono punti a coordinate reali, mentre la dizione *immaginaria* si riferisce a quelle superfici per cui i punti a coordinate reali sono al massimo quelli di una retta: nel caso del cono immaginario il vertice ha coordinate reali, mentre la quadrica di equazione $\frac{x^2}{a^2} + \frac{y^2}{b^2} = 0$ si spezza nei due piani di equazioni $\frac{x}{a} \pm i\frac{y}{b} = 0$ che hanno in comune una retta reale, l'asse z . Aggiungiamo che la denominazione *ellittica* o *iperbolica* dipende dalla natura dei loro punti (vedi § 23.3 a pagina 226); invece i cilindri vengono distinti a seconda della natura delle coniche che si ottengono intersecandoli con un piano perpendicolare alle generatrici.

È semplice verificare che, tra le quadriche non specializzate, vi sono superfici rigate: esse sono l'iperboloide ed il paraboloide iperbolici.

Esempio 23.2. La forma canonica del paraboloide a sella, che è

$$\frac{x^2}{a^2} - \frac{y^2}{b^2} = 2pz$$

si può anche scrivere nella forma

$$\left(\frac{x}{a} - \frac{y}{b}\right) \left(\frac{x}{a} + \frac{y}{b}\right) = 2p \cdot z$$

e quindi si vede subito che su di esso ci sono due sistemi di rette, definiti, al variare dei parametri k e h , dai sistemi

$$\left\{ \begin{array}{l} \frac{x}{a} - \frac{y}{b} = kz \\ k \left(\frac{x}{a} + \frac{y}{b} \right) = 2p \end{array} \right. \quad e \quad \left\{ \begin{array}{l} \frac{x}{a} + \frac{y}{b} = 2hp \\ h \left(\frac{x}{a} - \frac{y}{b} \right) = z \end{array} \right.$$

Riferendoci ancora alla forma canonica delle quadriche, possiamo notare che tutte le volte che vi figurano coefficienti uguali per due delle tre incognite si ha a che fare con una quadrica ottenuta dalla rotazione di una opportuna curva attorno all'asse con il nome della terza incognita.

Possiamo infine dare equazioni parametriche¹ delle quadriche non specializzate (tabella 23.2 a fronte):

- il sistema

$$\left\{ \begin{array}{l} x = a \cos \vartheta \cos \varphi \\ y = b \cos \vartheta \sin \varphi \\ z = c \sin \vartheta \end{array} \right.$$

(con $0 \leq \varphi < 2\pi$ e $0 \leq \vartheta < 2\pi$) rappresenta un ellissoide reale: se $a = b = c = R$ si ha la sfera di centro nell'origine e raggio R .

- il sistema

$$\left\{ \begin{array}{l} x = a \cosh \varphi \cos \vartheta \\ y = b \cosh \varphi \sin \vartheta \\ z = c \sinh \varphi \end{array} \right.$$

(con $0 \leq \vartheta < 2\pi$ e $\varphi \in \mathbb{R}$) rappresenta un iperboloide ad una falda.

- il sistema

$$\left\{ \begin{array}{l} x = a \cos \vartheta \\ y = b \sin \vartheta \cosh \varphi \\ z = c \sin \vartheta \sinh \varphi \end{array} \right.$$

(con $0 \leq \vartheta < 2\pi$ e $\varphi \in \mathbb{R}$) rappresenta un iperboloide a due falde.

¹A questo punto dev'essere chiaro che queste non sono le *uniche possibili* equazioni parametriche delle quadriche, bensì quelle maggiormente usate.

- il sistema

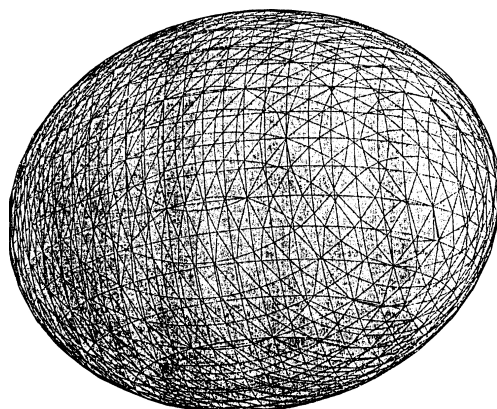
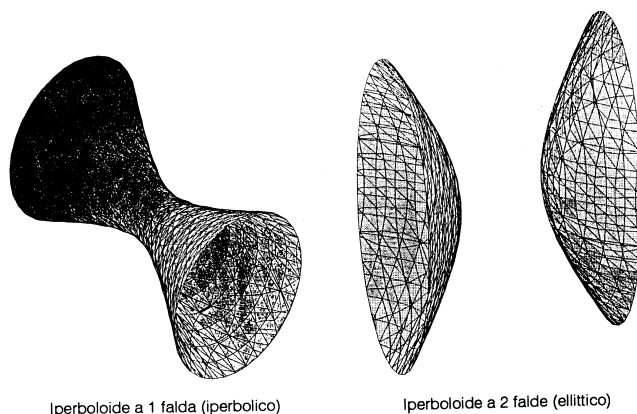
$$\begin{cases} x = at \\ y = bs \\ z = \frac{t^2 + s^2}{2p} \end{cases}$$

(con $s, t \in \mathbb{R}$) rappresenta un paraboloide ellittico.

- il sistema

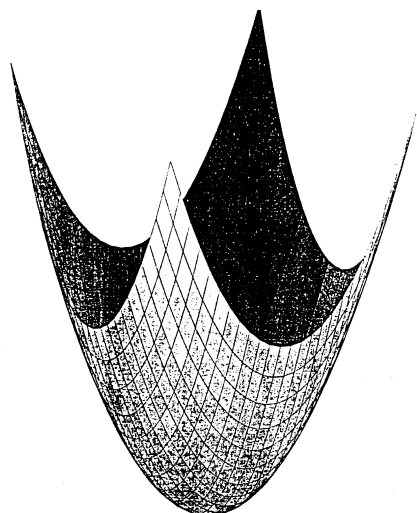
$$\begin{cases} x = at \\ y = bs \\ z = \frac{t^2 - s^2}{2p} \end{cases}$$

(con $s, t \in \mathbb{R}$) rappresenta un paraboloide iperbolico.

**Figura 23.1** *Ellissoide*

Iperboloide a 1 falda (iperbolico)

Iperboloide a 2 falde (ellittico)

Figura 23.2 *Gli iperboloidi***Figura 23.3** *Paraboloide ellittico*

23.3 Natura dei punti e riconoscimento di una quadrica

Poiché le quadriche sono superfici del second'ordine, le loro intersezioni con un piano sono delle coniche. Se $P(x_0, y_0, z_0)$ è un punto semplice di una quadrica irriducibile Σ e se consideriamo i vettori $\vec{p} = [x_0, y_0, z_0, 1]$ e $\vec{x} = [x, y, z, 1]$ allora l'equazione $\vec{p}A\vec{x}_T = 0$, dove A è la matrice dei coefficienti della quadrica, rappresenta il piano tangente in $P(x_0, y_0, z_0)$ alla quadrica.

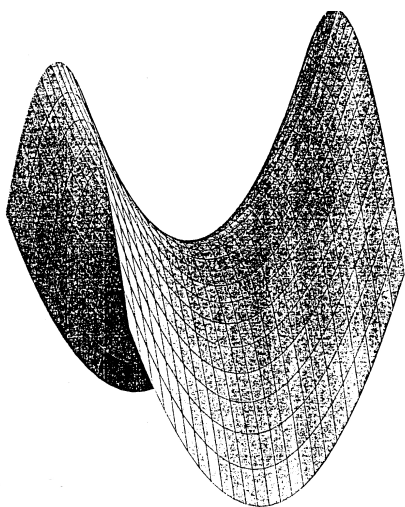


Figura 23.4 Paraboloide iperbolico

Si può anche provare che l'intersezione tra una quadrica ed il piano tangente in un suo punto semplice P è sempre una conica degenera in una coppia di rette; a questo proposito sussiste la

DEFINIZIONE 23.2. Un punto P di una quadrica Σ si dice *parabolico*, *ellittico* o *iperbolico* a seconda che il piano tangente in P tagli la quadrica secondo rispettivamente due rette coincidenti, due rette immaginarie o due rette distinte.

Si dimostra anche che tutti i punti di una quadrica hanno la stessa natura, cioè sono tutti di uno ed uno solo dei tre tipi considerati e che i diversi

tipi sono discriminati dal segno del determinante della matrice A dei coefficienti della quadrica: precisamente *i punti sono iperbolici se $\det A > 0$, ellittici se $\det A < 0$ e parabolici se $\det A = 0$.*

Uno dei modi per riconoscere una quadrica passa attraverso lo studio della natura dei suoi punti:

- Una quadrica a punti iperbolici è ovviamente rigata e non specializzata, quindi può essere solo un paraboloide a sella oppure un iperboloide ad una falda: i due casi si distinguono osservando se la sottomatrice B di A formata dai termini di secondo grado (cioè dalle prime tre righe e dalle prime tre colonne di A) ammette o no un autovalore nullo, cioè se $\det B = 0$.
- Se i punti sono ellittici si ha a che fare con un paraboloide ellittico, un ellissoide od un iperboloide a due falde, anche qui i tre casi si discriminano osservando il determinante di B , precisamente se

$|B| = 0$ si tratta del paraboloide, se $\det B > 0$ la quadrica è un ellissoide, se invece $\det B < 0$ si tratta dell'iperboloide.

- una quadrica a punti parabolici è sempre specializzata.

I casi esaminati sono riassunti nella tabella 23.3.

Tabella 23.3 Riconoscimento di una quadrica non degenera

Punti iperbolici $\det A > 0$	$ B = 0$	Paraboloide iperbolico
	$ B < 0$	Iperboloide ad una falda
	$ B > 0$	Ellissoide completamente immaginario
Punti ellittici $\det A < 0$	$ B = 0$	Paraboloide ellittico
	$ B < 0$	Iperboloide a due falde
	$ B > 0$	Ellissoide reale
Punti parabolici $\det A = 0$	$ B = 0$	Cilindro quadrico
	$ B < 0$	Cono quadrico reale
	$ B > 0$	Cono quadrico immaginario

Concludiamo il paragrafo dicendo che per distinguere i vari cilindri quadrici (a sezione ellittica, iperbolica o parabolica), si possono considerare i tre autovalori della matrice B introdotta prima: se uno di essi è nullo e gli altri sono discordi, la quadrica è un cilindro a sezione iperbolica, mentre se sono concordi il cilindro è a sezione ellittica, se invece gli autovalori nulli sono due, la quadrica è un cilindro a sezione parabolica².

Esempio 23.3. Vogliamo riconoscere la quadrica di equazione $xy - z = 0$. La sua matrice dei coefficienti è

$$A = \begin{bmatrix} 1 & \frac{1}{2} & 0 & 0 \\ \frac{1}{2} & 0 & 0 & 0 \\ 0 & 0 & 0 & -\frac{1}{2} \\ 0 & 0 & -\frac{1}{2} & 0 \end{bmatrix}$$

che, come si verifica subito, ha $\det A = \frac{1}{16} > 0$, quindi è una quadrica irriducibile non specializzata ed a punti iperbolici, fatto che si può anche constatare tenendo conto che il piano tangente nell'origine ad una qualsiasi superficie algebrica ha come equazione il complesso dei termini di primo grado;

²Non è difficile dimostrare che gli autovalori di B non possono essere tutti e tre nulli.

intersecando questo piano con la quadrica, si vede subito che l'intersezione è spezzata nelle due rette

$$\begin{cases} x = 0 \\ z = 0 \end{cases} \quad e \quad \begin{cases} y = 0 \\ z = 0 \end{cases}$$

reali e distinte, quindi l'origine è un punto iperbolico, di conseguenza la quadrica è a punti iperbolici. La matrice $B = \begin{bmatrix} 0 & \frac{1}{2} & 0 \\ \frac{1}{2} & 0 & 0 \\ 0 & 0 & 0 \end{bmatrix}$ è singolare e di rango 2, quindi la quadrica è un paraboloide iperbolico.

Esempio 23.4. Consideriamo la quadrica $x^2 + 2xy - 2xz - 2 = 0$. La matrice dei coefficienti è:

$$\begin{bmatrix} 1 & 1 & -1 & 0 \\ 1 & 0 & 0 & 0 \\ -1 & 0 & 0 & 0 \\ 0 & 0 & 0 & -2 \end{bmatrix}$$

singolare e di rango 3: si tratta di una quadrica irriducibile specializzata una

sola volta; la matrice $B = \begin{bmatrix} 1 & 1 & -1 \\ 1 & 0 & 0 \\ -1 & 0 & 0 \end{bmatrix}$ è singolare e di rango 2, quindi

abbiamo a che fare con un cilindro a sezione ellittica o iperbolica: siccome gli autovalori non nulli di B sono soluzioni dell'equazione $\lambda^2 - \lambda - 2 = 0$ essi sono discordi, quindi la quadrica è un cilindro a sezione iperbolica.

Riduzione a forma canonica

Per ridurre a forma canonica l'equazione di una quadrica osserviamo prima di tutto che la matrice B relativa ai termini di secondo grado è reale e simmetrica, quindi, in virtù del teorema 9.8 a pagina 89, è diagonalizzabile; la matrice che la diagonalizza si costruisce, come abbiamo visto, a partire da una base ortonormale di B .

È facile dimostrare che ciascuno dei tre autovettori che formano la base ammette come componenti i coseni direttori degli assi della quadrica: ciò significa che esiste sempre una rototraslazione che porta ad un sistema di riferimento i cui assi coincidono con gli assi di simmetria della quadrica, e quindi l'equazione assume una delle forme elencate nelle tabelle 23.1 a pagina 221 e 23.2 a pagina 222.

Esempio 23.5. Supponiamo di aver già operato la la rotazione del riferimento che porta la forma quadratica dell'equazione di una quadrica Σ a forma canonica, quindi di aver a che fare, ad esempio con la quadrica rappresentata, in un opportuno sistema di riferimento, dall'equazione

$$x^2 - 2y^2 + z^2 - 2x + y + 3 = 0$$

(per i passaggi necessari basta ricordare come ridurre a forma canonica una forma quadratica: vedi § 9.3 a pagina 92). I tre autovalori della forma quadratica sono non nulli, il che equivale a dire che il suo determinante $|B|$ è diverso da zero. Σ è dunque una quadrica a centro. Cerchiamo allora la traslazione del riferimento

$$\begin{cases} x = X + x_0 \\ y = Y + y_0 \\ z = Z + z_0 \end{cases}$$

che porta la sua equazione ad una delle forme canoniche: avremo così, sostituendo

$$(X + x_0)^2 - 2(Y + y_0)^2 + (Z + z_0)^2 - 2(X + x_0) + (Y + y_0) + 3 = 0$$

cioè, semplificando

$$X^2 - 2Y^2 + Z^2 + 2(x_0 - 1)X - (4y_0 - 1)Y + 2z_0Z + x_0^2 - 2y_0^2 + z_0^2 - 2x_0 + y_0 + 3 = 0.$$

Per eliminare i termini lineari occorre e basta avere

$$\begin{cases} x_0 - 1 = 0 \\ 4y_0 - 1 = 0 \\ 2z_0 = 0 \end{cases}$$

e quindi $X^2 - 2Y^2 + Z^2 + \frac{17}{8} = 0$ da cui la forma

$$-\frac{X^2}{\frac{17}{8}} + \frac{Y^2}{\frac{17}{16}} - \frac{Z^2}{\frac{17}{8}} = 1$$

ne segue anche che il centro ha coordinate $C(1, \frac{1}{4}, 0)$.

Operiamo ora la trasformazione

$$\begin{cases} x' = Z \\ y' = X \\ z' = Y \end{cases}$$

che, come si verifica subito è una rotazione in quanto la sua matrice è ortogonale a determinante positivo e che ha lo scopo di cambiare il nome agli assi da cui:

$$\frac{x'^2}{\frac{17}{16}} - \frac{y'^2}{\frac{17}{8}} - \frac{z'^2}{\frac{17}{8}} = 1$$

si tratta quindi di un iperboloide ellittico o a due falde.

I punti impropri delle quadriche

Se consideriamo, come abbiamo fatto nel piano, lo spazio ampliato con i suoi punti impropri, sorgono alcune questioni interessanti che qui brevemente accenniamo. Ad esempio in questa ambientazione coni e cilindri non sono più distinguibili, nel senso che i cilindri sono particolari coni il cui vertice è il punto improprio della retta generatrice. Nella stessa ottica, la quadrica di equazione $x^2 - 1 = 0$, che si spezza nei due piani paralleli $\pi_1 : x + 1 = 0$ e $\pi_2 : x - 1 = 0$ ammette come punti multipli tutti e soli quelli della retta impropria $r_\infty : \pi_1 \cap \pi_2$. Se $P(x_0 : y_0 : z_0 : u_0)$ è un punto semplice (e $\vec{\rho} = [x_0, y_0, z_0, u_0]$ il vettore delle sue coordinate) della quadrica irriducibile Σ di equazione

$$\begin{aligned} f(x, y, z, u) = & a_{11}x^2 + \cdots + a_{33}z^2 + \\ & + 2a_{12}xy + \cdots + a_{13}xz + \\ & + 2a_{14}xu + \cdots + a_{34}zu + a_{44}u^2 = 0 \end{aligned} \quad (23.3)$$

che si scrive anche

$$\vec{x}A\vec{x}_T = 0 \quad (23.4)$$

allora si dimostra che il piano tangente in P ha come equazione una delle quattro seguenti forme che si possono dimostrare equivalenti:

$$x \left(\frac{\partial f}{\partial x} \right)_P + y \left(\frac{\partial f}{\partial y} \right)_P + z \left(\frac{\partial f}{\partial z} \right)_P + u \left(\frac{\partial f}{\partial u} \right)_P = 0 \quad (23.5)$$

$$x_0 \frac{\partial f}{\partial x} + y_0 \frac{\partial f}{\partial y} + z_0 \frac{\partial f}{\partial z} + u_0 \frac{\partial f}{\partial u} = 0 \quad (23.6)$$

$$\vec{\rho}A\vec{x}_T \quad (23.7)$$

$$\vec{x}A\vec{\rho}_T \quad (23.8)$$

Un semplicissimo calcolo (che proponiamo come esercizio al lettore) mostra che le (23.5...23.8) sono equazioni lineari che rappresentano il medesimo piano: per l'appunto il piano tangente in P alla Σ .

Per riconoscere una quadrica irriducibile Σ si può allora studiare l'intersezione della Σ con il piano improprio $\pi_\infty : u = 0$: si tratta ovviamente di una conica, che prende il nome di *conica all'infinito* della Σ e che è rappresentata dalle equazioni

$$\begin{cases} a_{11}x^2 + a_{12}y^2 + a_{33}z^2 + 2a_{12}xy + a_{13}xz + a_{23}yz = 0 \\ u = 0 \end{cases} \quad (23.9)$$

o, similmente il cono che la proietta dall'origine, detto *cono asintotico* per la quadrica, la cui equazione è la prima delle due del sistema (23.9).

La conica all'infinito è:

<i>reale e non degenera</i>	per gli iperboloidi ed il cono quadratico reale;
<i>immaginaria non degenera</i>	per gli ellissoidi ed il cono quadratico immaginario;
<i>degenera in due rette reali ed incidenti</i>	per il paraboloide iperbolico ed il cilindro quadratico a sezione iperbolica;
<i>degenera in due rette immaginarie</i>	per il paraboloide ellittico e l'ellissoide, ma anche per il cilindro quadratico a sezione ellittica completamente immaginario.
<i>degenera in due rette reali coincidenti</i>	per il cilindro quadratico a sezione parabolica.

Esempio 23.6. Consideriamo la quadrica di equazione $xy = z$. In coordinate omogenee, la sua intersezione con il piano improprio ha equazioni

$$\begin{cases} xy - zu = 0 \\ u = 0 \end{cases} \quad \text{equivalente a} \quad \begin{cases} xy = 0 \\ u = 0 \end{cases}$$

cioè le due rette improprie reali e incidenti

$$\begin{cases} x = 0 \\ u = 0 \end{cases} \quad e \quad \begin{cases} y = 0 \\ u = 0 \end{cases}$$

quindi si tratta di un paraboloide iperbolico o di un cilindro a sezione iperbolica.

Se esaminiamo il rango della matrice $A = \frac{1}{2} \begin{bmatrix} 0 & 1 & 0 & 0 \\ 1 & 0 & 0 & 0 \\ 0 & 0 & 0 & 1 \\ 0 & 0 & 1 & 0 \end{bmatrix}$ si vede immedia-

tamente che è $\det A \neq 0$ quindi la quadrica non è specializzata dunque si tratta di un paraboloide iperbolico.

Vale la pena di notare in conclusione che, come già accadeva nel piano, l'ampliamento dello spazio con i punti impropri permette di eliminare fastidiose dissimmetrie.

Indice analitico

A

Angolo di due rette
 nel piano, 107
 nello spazio, 168

Applicazione, 63
 biiettiva, 64
 biunivoca, 64
 iniettiva, 64
 inversa, 66
 lineare, 64
 suriettiva, 63

Asintoti di un'iperbole, 129

Asse

 di una conica, 163

Autosoluzioni

 di un sistema omogeneo, 59

Autospazio, 82

Autovalori, 78

 regolari, 81

Autovettori, 78

B

Base, 38

 ortogonale, 74

 ortonormale, 74

Binet

 Teorema di, 50

C

Caratteristica

vedi Rango 51

Cayley-Hamilton

 Teorema di, 99

Centro

 di una conica, 159

Cilindri, 214

Circonferenza

 nello spazio, 202

Codominio, 63

Coefficiente angolare, 105

Coefficiente binomiale, 10

Combinazione lineare

 di matrici, 20

 di vettori, 36

Combinazioni, 9

Complemento algebrico, 46

Concetto primitivo, 3

Coni, 211

Conica all'infinito, 231

Conica per cinque punti, 134

Coniche, 125

Coniche a centro, 159

Coniche degeneri, 130

Cono asintotico, 231

Controimmagine, 63

Coordinate

 cilindriche, 191

 polari, 114

 nello spazio, 191

Coordinate omogenee

 nello spazio, 185

Coordinate polari

 nel piano, 114

Coppia involutoria, 146

- Coseni direttori, 168
 - di un vettore, 73
- Cramer
 - Teorema di, 57
- Curva
 - direttrice, 207
 - gobba, 198
 - sghemba, 198
- D
- Determinante, 45
 - definizione classica, 47
 - Definizione ricorsiva, 45
 - proprietà, 48
- Diametro
 - di una conica, 159
- Dimensione
 - di uno spazio vettoriale, 40
- Disposizioni, 9
- Distanza
 - di due vettori, 71
 - proprietà, 71
- Distanze nello spazio, 177
 - di due punti, 177
 - di due rette sghembe, 178
 - di piani paralleli, 178
 - di un punto da un piano, 177
 - di un punto da una retta, 178
- Dominio, 63
- E
- Eccentricità di una conica, 125
- Ellisse, 126
- Ellissoide, 222
- Endomorfismo, 78
- Equazione
 - del piano, 168
 - della circonferenza
 - nel piano, 117
 - della retta
 - canonica, 105
 - normale, 105
 - segmentaria, 107
 - lineare, 11
 - omogenea, 211
- Equazioni canoniche
 - delle coniche, 130, 132
- Equazioni parametriche
 - dell'ellisse, 127
 - della circonferenza, 118
 - di una retta
 - nel piano, 107
 - nello spazio, 167
- F
- Fasce
 - di sfere, 204
- Fascio
 - di circonferenze, 121
 - di coniche, 137
 - di piani, 173
 - di rette, 109
- Forma bilineare, 75
 - proprietà, 75
- Forma canonica
 - delle quadriche non specializzate, 222
 - delle quadriche specializzate, 221
- Forma proiettiva di prima specie, 143
- Forma quadratica
 - definita negativa, 95
 - definita positiva, 95
 - semidefinita negativa, 95
 - semidefinita positiva, 95
- Forme
 - quadratiche, 92
- Funzione, 63

G

Gauss

- algoritmo di , 30

Generatori

- di uno spazio vettoriale, 38

Generatrici, 207

Grassmann

- formula di, 43

I

Immagine, 63

Insieme, 3

- differenza, 4
- intersezione, 4
- parzialmente ordinato, 5
- sottoinsieme, 3
- totalmente ordinato, 5
- unione, 4
- vuoto, 3

- Intersezione tra retta e piano, 174

Invarianti

- di una conica, 131

Inversa

- di un'applicazione lineare, 66

Involuzione

- circolare, 162
- dei diametri coniugati, 160
- dei punti coniugati, 155
- dei punti reciproci, 155

Involuzioni, 146

- ellittiche, 147
- iperboliche, 147

Iperbole, 128

Iperboloide

- ellittico, 222
- iperbolico, 222

Isomorfismo, 66

K

Kroneker

- Teorema di, 53

L

Lagrange

- Teorema di, 94

Laplace

- primo teorema di, 48
- secondo teorema di, 48

Linee nello spazio, 198

M

Matrice, 16

- aggiunta, 96
- associata ad una forma quadratica, 93
- autoaggiunta, 95
- dei complementi algebrici, 55
- del cambiamento di base, 41
- diagonale, 18
- diagonale principale, 17
- diagonalizzabile, 85
- elementi principali, 17
- elevazione a potenza, 24
- emisimmetrica, 17
- hermitiana, 95
- inversa, 54
- invertibile, 54
- normale, 96
- nulla, 19
- operazioni elementari
 - sulle colonne, 27
 - sulle righe, 27
- orlare, 53
- ortogonale, 88
- ortogonalmente diagonalizzabile, 89
- quadrata, 17
- ridotta, 29

- scalare, 18
 - simmetrica, 17
 - singolare, 49
 - traccia di una, 17
 - trasposta, 16
 - triangolare, 18
 - unità, 18
 - unitaria, 96
- Matrice associata
 - ad un'applicazione lineare, 67
- Matrici
 - a blocchi, 25
 - conformabili, 21
 - equivalenti, 27
 - idempotenti, 24
 - invertibili, proprietà, 55
 - linearmente dipendenti, 20
 - linearmente indipendenti, 20
 - nilpotenti, 24
 - polinomi di, 24
 - prodotto, 21
 - prodotto per uno scalare, 19
 - simili, 77
 - somma, 18
 - somma di, 19
 - uguali, 16
- Minore, 51
 - complementare, 46
 - principale, 81
- Molteplicità
 - algebrica, 79
 - geometrica, 81
- N
- Natura dei punti di una quadrica, 226
- Norma di un vettore, 71
 - proprietà, 71
- Nucleo
 - di un'applicazione lineare, 65
- O
- Omomorfismo, 64
- P
- Paraboloide
 - ellittico, 222
 - iperbolico, 222
- Parallelismo e perpendicolarità
 - tra due piani nello spazio, 170
 - tra due rette nello spazio, 170
 - tra un piano ed una retta nello spazio, 171
- Parallelismo e perpendicolarità nello spazio, 170
- Parametri direttori
 - di un piano, 169
 - di una retta, 167
 - di una retta nel piano, 105
- Permutazione, 8
 - con ripetizione, 8
 - scambio, 8
- Piano tangente
 - ad una sfera, 202
- Plücker
 - legge di, 152
- Polare di un punto, 149
- Polinomio caratteristico, 79
- Polinomio minimo, 101
- Polo, 150
- Principio di induzione, 6
- Prodotto associato, 47
- Prodotto cartesiano
 - di due insiemi, 75
- Prodotto scalare
 - in $\mathbb{R}^n$, 72
 - in generale, 76

proprietà, 72
Proiettività, 143
 ellittica, 145
 iperbolica, 145
 parabolica, 145
Proiezione ortogonale, 181
Proiezioni, 216
 centrali, 216
 parallele, 216
Pundi base di un fascio
 di coniche, 137
Punti
 coniugati, 154
 ellittici, 226
 iperbolici, 226
 multipli, 197
 parabolici, 226
 reciproci, 154
Punti base di un fascio
 di circonferenze, 121
Punti ciclici, 124
Punti impropri
 di una quadrica, 230
Punti uniti, 144
Punto improprio, 110

Q
Quadriche, 219
Quadriche in forma canonica
 loro classificazione, 221

R
Rango
 calcolo del, 53
 di una forma quadratica, 93
 di una matrice, 29, 51
 proprietà, 52
Reciprocità
 legge di, 152
Regola del parallelogrammo, 70
Relazioni, 4

 d'ordine, 5
 di equivalenza, 4
 insieme quoziente, 4
 di equivalenza, 4
 di ordine, 5
Retta impropria, 111
Rette
 punteggiate, 143
Rette isotrope, 132
Rette sghembe, 175
Riconoscimento di una conica
 nello spazio, 217
Riconoscimento di una quadrica
 non degenerare, 227
Riduzione a forma canonica
 dell'equazione di una quadrica, 228
Rototraslazioni, 189
Rouché-Capelli
 Teorema di, 30, 58

S
Sfera, 201
Simmetrie, 180
 rispetto ad un piano, 181
 rispetto ad un punto, 180
 rispetto ad una retta, 184
Sistema di equazioni, 12
Sistema di riferimento
 nel piano, 69
 nello spazio, 69
Sistema lineare, 12
 omogeneo, 59
 risoluzione elementare, 14
 soluzioni, 13
Sistemi lineari
 teoria dei, 57
 teoria elementare, 11
Somma
 di sottospazi, 42
Somma diretta

- di sottospazi, 42
- Σ sommatoria, 5
- Sostegno
 - di un fascio, 109
- Sottospazio, 38
- Spazi vettoriali
 - proprietà, 36
- Spazio euclideo, 76
- Spazio vettoriale, 36
 - su $\mathbb{R}$, 35
- Spettro
 - di una matrice, 79
- Superfici, 195
 - di rotazione, 208
 - rigate, 207
- Superficie
 - algebrica, 197
- T
- Tangente
 - ad una superficie, 197
- Tangenti
 - ad una circonferenza, 119
- Tangenti ad una conica
 - in forma canonica, 133
- Triangoli autopolari, 156
- V
- Vandermonde
 - determinante di, 50
- Versore, 73
- Vettore
 - unitario, 73
- Vettori
 - disgiunti, 41
 - fondamentali, 37
 - linearmente dipendenti, 37
 - linearmente indipendenti, 37
 - ortogonali, 74
 - ortonormali, 74