

S. Benenti

Lezioni di
MECCANICA RAZIONALE

a.a. 1993-94

Parte prima

Istituto di Fisica Matematica

"J.-L. Lagrange"

Università di Torino

1994

CLU

INDICE

Capitolo I - Elementi di algebra lineare.

1. Richiami e notazioni	p. 1
2. Applicazioni lineari	5
3. Spazi duali	8
4. Endomorfismi lineari	11
4.1 Il commutatore di endomorfismi	12
4.2 Il determinante e la traccia di un endomorfismo	13
4.3 Il polinomio caratteristico di un endomorfismo	15
4.4 Autovalori ed autovettori di un endomorfismo	17
4.5 Il teorema di Hamilton-Cayley	20
5. Forme bilineari	22
5.1 Forme bilineari simmetriche e forme quadratiche	24
5.2 Forme bilineari antisimmetriche	29
6. Tensori	32
7. Algebra esterna	36
8. Spazi vettoriali euclidei	47
9. Endomorfismi in spazi euclidei	53
10. Endomorfismi ortogonali	62
11. Lo spazio vettoriale euclideo tridimensionale	68
11.1 La forma volume	69
11.2 L'aggiunzione	73
11.3 Il prodotto vettoriale	75
11.4 Le rotazioni	80
12. Spazi vettoriali iperbolici	85

Capitolo II - Calcolo vettoriale negli spazi affini.

1. Richiami sugli spazi affini	p. 95
2. Campi scalari e vettoriale	97
3. Coordinate non affini e riferimenti associati	102
4. Forme differenziali	110
5. Curve	122
6. Superfici	127
7. Sistemi dinamici	137
8. Sistemi dinamici sulla retta	142

9. Sistemi nel piano sulla retta	147
10. Integrali primi	152
11. Spazi affini euclidei	165
11.1 Il tensore metrico	166
11.2 Il gradiente	170
11.3 La forma volume	174
11.4 Il rotore	176
12. Curve e superfici nello spazio affine euclideo tridimensionale ..	178
13. Baricentro di un sistema di masse	188
14. Tensore d'inerzia di un sistema di masse	193
15. Sistemi di vettori applicati	201

Capitolo III - Cinematica del punto e del corpo rigido.

1. Cinematica del punto	p. 205
2. Moti centrali	212
3. Moti geodetici su di una superficie	216
4. Cinematica del corpo rigido	227
5. Cambiamenti di riferimento	235
6. Moti rigidi particolari	239
7. Moti rigidi composti	244

Hanno collaborato alla redazione di queste Lezioni i Dottori:

Manuelita Bonadies,
Paolo Cermelli,
Guido Magnano.

Torino, febbraio 1994

CAPITOLO I

ELEMENTI DI ALGEBRA LINEARE

L'algebra lineare fornisce strumenti di calcolo e strutture di base per la costruzione di modelli matematici della fisica e di altre scienze applicate. Ne esamineremo alcuni elementi, parte dei quali già trattati in altri corsi del primo biennio di matematica, il cui richiamo tuttavia ci consentirà di evidenziarne gli aspetti applicativi.

Questo capitolo può suddividersi in due parti. Nella prima parte, paragrafi da 1 a 7, sono trattati spazi vettoriali "puri", cioè senza alcuna struttura aggiuntiva. Nella seconda parte sono invece trattati gli spazi vettoriali dotati di un "prodotto scalare", cioè di una struttura euclidea.

1. Richiami e notazioni

Si dice che un insieme E è uno **spazio vettoriale** o **spazio lineare** sopra un corpo commutativo $\mathbb{K}$ se

(i) È definita un'operazione binaria interna su E , $+: E \times E \rightarrow E$, detta **somma**, soddisfacente agli assiomi di gruppo commutativo:

$$\left\{ \begin{array}{l} u + v = v + u, \quad \forall u, v \in E, \\ (u + v) + w = u + (v + w), \quad \forall u, v, w \in E, \\ \exists 0 \in E \mid u + 0 = u, \quad \forall u \in E, \\ \forall u \in E, \exists v \in E \mid u + v = 0. \end{array} \right.$$

(ii) È definita un'operazione $\mathbb{K} \times E \rightarrow E$, detta **prodotto per uno scalare**, tale da soddisfare alle seguenti condizioni:

$$\begin{cases} a(u+v) = au + av, & \forall a \in \mathbb{K}, u \in E, v \in E, \\ (a+b)u = au + bu, & \forall a, b \in \mathbb{K}, u \in E, \\ (ab)u = a(bu), & \forall a, b \in \mathbb{K}, u \in E, \\ 1u = u, & 1 = \text{unità in } \mathbb{K}, \forall u \in E. \end{cases}$$

Gli elementi di E si dicono **vettori**, gli elementi di $\mathbb{K}$ **scalari**. Denotiamo i vettori, salvo l'elemento neutro $0 \in E$ (**zero** o **vettore nullo**), con lettere grassetto (quasi sempre minuscole). Con lo stesso simbolo 0 denotiamo anche lo zero del campo $\mathbb{K}$. Per ogni vettore u si ha $0u = 0$. Si denota con $-u$ il **vettore opposto** a u , cioè il vettore per cui $u + (-u) = 0$. Scriviamo $v - u$ al posto di $v + (-u)$. Il vettore nullo e l'opposto di un qualunque vettore sono unici (perché unici sono, in un qualunque gruppo, l'elemento neutro e il reciproco di un qualunque elemento).

Tra gli spazi vettoriali sopra un assegnato campo $\mathbb{K}$ ritroviamo il campo stesso e tutte le sue potenze cartesiane $\mathbb{K}^n$, cioè l'insieme delle n -ple di elementi del campo.

Un sottoinsieme $S \subset E$ è un **sottospazio** se le proprietà (i) e (ii) valide in E valgono anche per tutti gli elementi di S . Se K e L sono due sottospazi, denotiamo con $K + L$ la loro somma, cioè l'insieme dei vettori di E costituito da tutte le possibili somme di elementi di K e di L . La somma $K + L$ e l'intersezione $K \cap L$ di due sottospazi vettoriali sono a loro volta dei sottospazi. Denotiamo con $E \oplus F$ la **somma diretta di due spazi vettoriali** (sul medesimo corpo K): è il prodotto cartesiano $E \times F$ dotato della naturale struttura di spazio vettoriale definita dalle operazioni $(u, w) + (u', w') = (u + u', w + w')$ e $a(u, w) = (au, aw)$. Denoteremo anche con $u \oplus w$ un generico elemento di $E \oplus F$. Si dice anche che uno spazio vettoriale E è la **somma diretta** di due suoi sottospazi K e L se $E = K + L$ e se $K \cap L = \{0\}$. In tal caso ogni vettore è la somma di due vettori univocamente determinati di K e L . E risulta pertanto identificabile con $K \oplus L$.

Una scrittura del tipo

$$a^1 v_1 + a^2 v_2 + \dots + a^k v_k = \sum_{i=1}^k a^i v_i$$

prende il nome di **combinazione lineare** dei vettori v_i con **coefficienti** $a^i \in K$. Due o più vettori si dicono **indipendenti** se la sola

combinazione lineare che fornisce il vettore nullo è quella a coefficienti tutti nulli:

$$\sum_{i=1}^k a^i v_i = 0 \implies a^i = 0.$$

Uno spazio vettoriale è detto a **dimensione finita** se ammette una **base** costituita da un numero finito n di vettori, cioè un insieme di n vettori indipendenti che generano tutto E tramite tutte le loro possibili combinazioni lineari. Il numero n coincide col massimo numero di vettori indipendenti in un qualunque insieme di vettori. Esso prende il nome di **dimensione** dello spazio. Scriveremo E_n oppure $\dim(E) = n$ per indicare che la dimensione dello spazio vettoriale E è n . Una generica base di E sarà denotata con $(e_1, e_2, \dots, e_n)$ o brevemente con (e_i) , dove l'indice i sarà inteso variare da 1 a n . La scrittura

$$v = v^1 e_1 + v^2 e_2 + \dots + v^n e_n = \sum_{i=1}^n v^i e_i$$

è la rappresentazione di un generico vettore $v \in E$ secondo la base (e_i) . I coefficienti di questa combinazione lineare sono le **componenti** del vettore v . Le componenti di un vettore, fissata la base, sono univocamente determinate. Si noti la diversa posizione degli indici riservata alle componenti (v^i), in alto anziché in basso. Si conviene tuttavia di omettere il simbolo di sommatoria per indici ripetuti in alto e in basso (*convenzione di Einstein*). Allora la rappresentazione precedente si scrive più semplicemente

$$v = v^i e_i.$$

Si noti che non ha importanza la scelta dell'indice, purché la lettera usata appartenga ad un insieme precedentemente convenuto (per esempio, lettere latine minuscole, lettere greche, lettere latine maiuscole dopo la I , e così via). Se per esempio conveniamo che le lettere latine minuscole a partire dalla h varino da 1 a n , la precedente scrittura è del tutto equivalente alle seguenti:

$$v = v^h e_h, \quad v = v^j e_j, \quad v = v^k e_k, \quad \dots$$

La diversa collocazione degli indici e la sommatoria sottintesa per indici ripetuti in alto e in basso sono convenzioni tipiche del **calcolo tensoriale**, i cui primi elementi saranno illustrati in questo capitolo. Il formalismo che ne deriva è agevole e sintetico.

Le formule riguardanti il calcolo vettoriale o tensoriale sono in genere di tre tipi. Una formula è di *tipo intrinseco* o *assoluto* se essa coinvolge vettori o tensori ed operazioni definite su questi senza l'intervento di basi. Ad una formula di tipo intrinseco corrisponde sempre una formula *con indici* o *in componenti*, conseguente alla scelta di una *base generica*. Vi sono infine formule, con indici, valide soltanto per una *base particolare* o per una *classe particolare di basi*. Per sottolineare questa circostanza useremo sovente un'uguaglianza con asterisco: $\doteq$.

Per quel che riguarda il concetto di dimensione è opportuno ricordare che per due sottospazi di uno spazio vettoriale E vale la **formula di Grassmann** (Hermann Günther Grassmann, 1809-1877)

$$\dim(K + L) + \dim(K \cap L) = \dim(K) + \dim(L).$$

Un'applicazione $A: E \rightarrow F$ da uno spazio vettoriale E ad uno spazio vettoriale F si dice **lineare** se

$$A(au + bv) = aA(u) + bA(v) \quad (\forall u, v \in E, \forall a, b \in K).$$

Si dice anche che A è un **omomorfismo lineare**. Denoteremo in genere gli omomorfismi lineari con lettere maiuscole in grassetto. Denoteremo con

$$Hom(E; F)$$

o semplicemente con

$$L(E; F)$$

l'insieme delle applicazioni lineari da E a F . Quest'insieme ha una naturale struttura di spazio vettoriale, con la somma $A + B$ e il prodotto aA per un elemento $a \in K$ definiti da

$$(A + B)(v) = A(v) + B(v), \quad (aA)(v) = aA(v).$$

Più in generale, se $(E_1, E_2, \dots, E_p, F)$ sono spazi vettoriali, un'applicazione

$$A: E_1 \times E_2 \times \dots \times E_p \rightarrow F$$

è detta **p-lineare** (genericamente **multilineare**) se è lineare su ogni argomento. Ad esempio l'applicazione $\mathbb{R}^2 \rightarrow \mathbb{R}: (x, y) \mapsto xy$ è *bilineare* (l'insieme dei numeri reali $\mathbb{R}$ ha un'ovvia struttura di spazio vettoriale su se stesso, di dimensione 1). Denoteremo con

$$L(E_1, E_2, \dots, E_p; F)$$

l'insieme delle applicazioni multilineari da $E_1 \times E_2 \times \dots \times E_p$ a F . Anch'esso è dotato di una struttura naturale di spazio vettoriale.

In queste lezioni ci limiteremo a considerare solo spazi vettoriali reali (cioè sul campo dei reali, $\mathbb{K} = \mathbb{R}$) e a dimensione finita.

Nei prossimi paragrafi esamineremo i seguenti spazi di applicazioni:

$$\begin{aligned} L(E; F) &= Hom(E, F) && \text{applicazioni lineari} \\ L(E; \mathbb{R}) &= E^* && \text{spazio duale di } E \\ L(E; E) &= End(E) && \text{endomorfismi lineari su } E \\ L(E, E; \mathbb{R}) &= T_2(E) && \text{forme bilineari su } E \\ L(\underbrace{E^*, \dots, E^*}_{p \text{ volte}}, \underbrace{E, \dots, E}_{q \text{ volte}}; \mathbb{R}) &= T_q^p(E) && \text{tensori di tipo } (p, q). \end{aligned}$$

2. Applicazioni lineari

Sia $A \in Hom(E; F)$. Denotiamo con $A(E)$ l'**immagine** di A , cioè il sottinsieme di F formato da quegli elementi che sono immagini secondo A di elementi di E , e con $Ker(A)$ o anche $A^{-1}(0)$ il **nucleo** di A , vale a dire il sottinsieme di E costituito dalle controimmagini dello zero di F :

$$\begin{aligned} A(E) &= \{v \in F \mid \exists u \in E \mid v = A(u)\}, \\ A^{-1}(0) &= Ker(A) = \{u \in E \mid A(u) = 0\}. \end{aligned}$$

Entrambi questi sottinsiemi sono dei sottospazi. Denotiamo con $A(S)$ l'immagine secondo A di un sottinsieme S di E . Se S è un sottospazio, anche $A(S)$ è un sottospazio. Vale l'uguaglianza

$$(1) \quad \dim(Ker(A)) + \dim(A(E)) = \dim(E),$$

che è conseguenza diretta del fatto che $A(E)$ è isomorfo allo spazio quoziente $E/Ker(A)$, costituito dalle classi della relazione di equivalenza su E definita da: $v \sim v' \iff v - v' \in Ker(A)$.

Un **isomorfismo** è un omomorfismo lineare biiettivo. Se $A: E \rightarrow F$ è un isomorfismo, i due spazi E ed F si dicono **isomorfi**. Essi hanno necessariamente la stessa dimensione. Un omomorfismo A è un isomorfismo se e solo se $Ker(A) = 0$.

Siano m ed n le dimensioni di E ed F rispettivamente. Scelta una base (e_i) di E ed una base (f_α) di F ($i = 1, \dots, m$; $\alpha = 1, \dots, n$), ad

ogni $\mathbf{A}$ risulta associata una matrice di numeri (A_i^α) definita dall'uguaglianza

$$(2) \quad \boxed{\mathbf{A}(\mathbf{e}_i) = A_i^\alpha \mathbf{f}_\alpha}$$

detta **matrice delle componenti** di $\mathbf{A}$. La (2) mostra che per ogni indice i i numeri A_i^α danno le componenti del vettore $\mathbf{A}(\mathbf{e}_i) \in F$ nella base $(\mathbf{f}_\alpha)$. La (2) è la scrittura sintetica del seguente sistema di uguaglianze:

$$(2') \quad \begin{cases} \mathbf{A}(\mathbf{e}_1) = A_1^1 \mathbf{f}_1 + A_1^2 \mathbf{f}_2 + \dots + A_1^n \mathbf{f}_n \\ \mathbf{A}(\mathbf{e}_2) = A_2^1 \mathbf{f}_1 + A_2^2 \mathbf{f}_2 + \dots + A_2^n \mathbf{f}_n \\ \dots \\ \mathbf{A}(\mathbf{e}_m) = A_m^1 \mathbf{f}_1 + A_m^2 \mathbf{f}_2 + \dots + A_m^n \mathbf{f}_n \end{cases}$$

La matrice delle componenti definisce completamente l'azione di $\mathbf{A}$: all'uguaglianza (assoluta) $\mathbf{v} = \mathbf{A}(\mathbf{u})$ corrisponde l'uguaglianza in componenti

$$(3) \quad \boxed{v^\alpha = A_i^\alpha u^i}$$

posto

$$\mathbf{v} = v^\alpha \mathbf{f}_\alpha, \quad \mathbf{u} = u^i \mathbf{e}_i.$$

Infatti per la linearità di $\mathbf{A}$ risulta:

$$\mathbf{v} = \mathbf{A}(u^i \mathbf{e}_i) = u^i \mathbf{A}(\mathbf{e}_i) = u^i A_i^\alpha \mathbf{f}_\alpha.$$

La (3) è l'espressione sintetica di un sistema di n equazioni lineari:

$$(3') \quad \begin{cases} v^1 = A_1^1 u^1 + A_2^1 u^2 + \dots + A_m^1 u^m \\ v^2 = A_1^2 u^1 + A_2^2 u^2 + \dots + A_m^2 u^m \\ \dots \\ v^n = A_1^n u^1 + A_2^n u^2 + \dots + A_m^n u^m \end{cases}$$

OSSERVAZIONE 1. Se si conviene che l'indice in basso i della matrice (A_i^α) sia *indice di riga* e l'indice in alto α sia *indice di colonna*, per cui

$$(A_i^\alpha) = \begin{pmatrix} A_1^1 & A_1^2 & \dots & A_1^n \\ A_2^1 & A_2^2 & \dots & A_2^n \\ \vdots & \vdots & \ddots & \vdots \\ A_m^1 & A_m^2 & \dots & A_m^n \end{pmatrix},$$

e quindi che le componenti dei vettori, (v^α) e (u^i) , formino dei **vettori riga**, allora la (3') si riscrive in termini di prodotto di matrici alla maniera seguente:

$$(3'') \quad (v^1, v^2, \dots, v^n) = (u^1, u^2, \dots, u^m) \begin{pmatrix} A_1^1 & A_1^2 & \dots & A_1^n \\ A_2^1 & A_2^2 & \dots & A_2^n \\ \vdots & \vdots & \ddots & \vdots \\ A_m^1 & A_m^2 & \dots & A_m^n \end{pmatrix}.$$

Si noti bene che a secondo membro il prodotto di matrici è *righe per colonne*.

Se invece si adotta la convenzione opposta, cioè che l'indice in alto α della matrice (A_i^α) sia indice di riga e l'indice in basso i sia indice di colonna, per cui

$$(A_i^\alpha) = \begin{pmatrix} A_1^1 & A_2^1 & \dots & A_m^1 \\ A_1^2 & A_2^2 & \dots & A_m^2 \\ \vdots & \vdots & \ddots & \vdots \\ A_1^n & A_2^n & \dots & A_m^n \end{pmatrix},$$

e quindi che le componenti dei vettori, (v^α) e (u^i) , formino dei **vettori colonna**, allora la (3') si riscrive ancora in termini di prodotto di matrici:

$$(3''') \quad \begin{pmatrix} v^1 \\ v^2 \\ \vdots \\ v^n \end{pmatrix} = \begin{pmatrix} A_1^1 & A_2^1 & \dots & A_m^1 \\ A_1^2 & A_2^2 & \dots & A_m^2 \\ \vdots & \vdots & \ddots & \vdots \\ A_1^n & A_2^n & \dots & A_m^n \end{pmatrix} \begin{pmatrix} u^1 \\ u^2 \\ \vdots \\ u^m \end{pmatrix}.$$

A secondo membro il prodotto di matrici è ancora *righe per colonne*.

È bene a questo proposito osservare che mentre la scrittura a secondo membro della (3), $A_i^\alpha u^i$, è del tutto equivalente a $u^i A_i^\alpha$, i prodotti matriciali nella (3'') e nella (3''') non sono certamente commutabili.

La corrispondenza tra endomorfismi e matrici definita dalla (2) dipende ovviamente dalla scelta delle due basi, in E e in F : se queste cambiano, cambiano anche le componenti degli endomorfismi. Una conseguenza immediata delle precedenti considerazioni è che

$$(4) \quad \dim(\text{Hom}(E, F)) = \dim(E) \cdot \dim(F).$$

OSSERVAZIONE 2. Determinare il nucleo di un endomorfismo A equivale a risolvere il sistema di equazioni lineari omogenee

$$(5) \quad A_i^\alpha x^i = 0,$$

composto da n equazioni, nelle m incognite (x^i) (che sono le componenti dei vettori di $\text{Ker}(A)$). Risolvere il sistema lineare non omogeneo

$$(6) \quad A_i^\alpha x^i = a^\alpha$$

significa invece determinare le controimmagini secondo A del vettore $a \in F$ di componenti (a^α) . Se a non appartiene all'immagine $A(E)$ questo sistema non ha soluzione. Se invece $a \in A(E)$ allora il sistema ha per soluzione tutte e sole quelle m -ple di numeri (x^i) che sono componenti dei vettori di una classe di equivalenza di $E/\text{Ker}(A)$. Quindi le soluzioni sono ∞^k , con $k = \dim(\text{Ker}(A))$. Se $k = 0$ (cioè se A è un isomorfismo) si ha una ed una sola soluzione per ogni scelta delle (a^α) . Si noti che per la (1), $k = m - r$, dove $m = \dim(E)$ e $r = \dim(A(E))$.

DEFINIZIONE 1. Dicesi **rango di un'applicazione lineare** A la dimensione della sua immagine. Lo denotiamo con $\text{rank}(A)$.

È notevole il fatto che il rango di un endomorfismo A coincide, qualunque siano le basi scelte, con il **rango della matrice** delle componenti (A_i^α) che è, per definizione, il *massimo ordine delle sue sottomatrici quadrate non singolari, cioè a determinante non nullo*. La dimostrazione è lasciata come esercizio.

3. Spazi duali

Un'applicazione lineare $\alpha: E \rightarrow \mathbb{R}$ da uno spazio vettoriale E allo spazio $\mathbb{R}$ dei numeri reali è detta **forma lineare** su E o **covettore**. Lo spazio $L(E; \mathbb{R})$ delle forme lineari è più semplicemente denotato con E^* ed è chiamato **spazio duale** di E .

Al posto di $\alpha(v)$ useremo il simbolo

$$\langle v, \alpha \rangle$$

per denotare il valore del covettore α sul vettore v . Diremo che $\langle v, \alpha \rangle$ è la **valutazione** di α su v o, viceversa, di v su α . Ponendo $F = \mathbb{R}$ in quanto si è detto nel paragrafo precedente, si vede in particolare che

$$\dim(E^*) = \dim(E).$$

Esiste un isomorfismo *canonico* $\iota: E \rightarrow E^{**}$ tra lo **spazio biduale** $E^{**} = L(E^*; \mathbb{R})$ e lo spazio E stesso, definito dall'uguaglianza

$$(2) \quad \langle \alpha, \iota(v) \rangle = \langle v, \alpha \rangle.$$

Gli elementi di E^{**} saranno di seguito sempre identificati con i vettori di E .

Ad ogni sottospazio $K \subset E$ associamo un sottospazio $K^\circ \subset E^*$, detto **polare** di K , costituito da tutti i covettori che annullano tutti gli elementi di K :

$$(3) \quad K^\circ = \{ \alpha \in E^* \mid \langle u, \alpha \rangle = 0, \forall u \in K \}.$$

Sussistono le seguenti proprietà, qualunque siano i sottospazi K e L di E :

$$(4) \quad \begin{cases} \dim(K) + \dim(K^\circ) = \dim(E), \\ K^{\circ\circ} = K, \\ K^\circ \subset L^\circ \iff L \subset K, \\ (K + L)^\circ = K^\circ \cap L^\circ, \\ K^\circ + L^\circ = (K \cap L)^\circ. \end{cases}$$

La loro dimostrazione è lasciata come esercizio. Si noti che la seconda di queste uguaglianze sottintende l'intervento dell'isomorfismo canonico tra lo spazio e il suo biduale e che inoltre, in virtù di questo isomorfismo, le ultime due uguaglianze (4) esprimono la medesima proprietà.

Ad ogni base (e_i) di E corrisponde un **base duale** in E^* , che denotiamo con (ϵ^i) , definita implicitamente dall'uguaglianza:

$$(5) \quad \boxed{\langle e_j, \epsilon^i \rangle = \delta_j^i}$$

dove δ_j^i è il **simbolo di Kronecker** (Leopold Kronecker, 1823-1851):

$$\boxed{\delta_j^i = \begin{cases} 1 & \text{se } i = j \\ 0 & \text{se } i \neq j \end{cases}}$$

Per ogni $v \in E$ e $\alpha \in E^*$ valgono le relazioni:

$$(6) \quad \boxed{v = v^i e_i} \iff \boxed{v^i = \langle v, \varepsilon^i \rangle}$$

$$(7) \quad \boxed{\alpha = \alpha_i \varepsilon^i} \iff \boxed{\alpha_i = \langle e_i, \alpha \rangle}$$

$$(8) \quad \boxed{\langle v, \alpha \rangle = \alpha_i v^i}$$

La loro dimostrazione è lasciata come esercizio. Si osservi la diversa posizione degli indici delle componenti di un covettore (in basso, anziché in alto come per le componenti dei vettori) e la simmetria formale di queste uguaglianze. L'equivalenza (6) mostra che le componenti di un vettore sono ottenute valutando sul vettore gli elementi della base duale. Analogamente la (7) mostra che le componenti di un covettore coincidono con le valutazioni di questo sugli elementi della base di E . Infine la (8) mostra che la valutazione tra un vettore ed un covettore è data dalla somma dei prodotti delle componenti omologhe.

OSSERVAZIONE 1. *Rappresentazione geometrica dei covettori.* Ad ogni covettore $\alpha \in E^*$ si possono far corrispondere due sottoinsiemi di E :

$$A_0 = \{x \in E \mid \langle x, \alpha \rangle = 0\}, \quad A_1 = \{x \in E \mid \langle x, \alpha \rangle = 1\}.$$

A_0 è un sottospazio di codimensione 1, A_1 è un suo laterale. La conoscenza di questi due sottoinsiemi (in effetti basterebbe la conoscenza del secondo) determina completamente α . Infatti, dato un vettore $v \in E$, o è $v \in A_0$, e allora poniamo $\langle v, \alpha \rangle = 0$, oppure è $v = \lambda x$ con $x \in A_1$, e allora poniamo $\langle v, \alpha \rangle = \lambda$. Si osservi che, denotate con (x^i) le componenti di un generico vettore x , i due sottoinsiemi A_0 e A_1 sono rispettivamente definiti dalle equazioni

$$\alpha_i x^i = 0, \quad \alpha_i x^i = 1.$$

Si tratta di una coppia di "iperpiani" paralleli, il primo passante per l'origine.

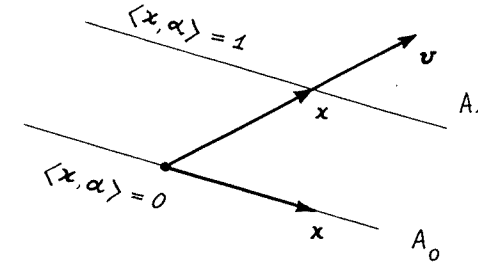


Fig. 3.1: rappresentazione geometrica dei covettori.

4. Endomorfismi lineari

Denotiamo con $End(E)$ lo spazio delle applicazioni lineari $L(E; E)$ di uno spazio vettoriale E in se stesso, cioè lo spazio degli **endomorfismi** di E . Denotiamo semplicemente con AB la composizione di due endomorfismi $A \circ B$. L'operazione di composizione degli endomorfismi conferisce allo spazio $End(E)$ la struttura di algebra associativa con unità; l'unità è l'applicazione identica di E in se stesso, che denotiamo con id_E o con 1_E o semplicemente con 1 . Denotiamo con $0: E \rightarrow E$ l'applicazione che manda tutti i vettori di E nello zero.

Gli endomorfismi invertibili sono detti **automorfismi** o **trasformazioni lineari** di E . Formano un gruppo che denotiamo con $Aut(E)$. Si denota con A^{-1} l'inverso di un automorfismo A .

Scelta una base (e_i) dello spazio E , denotata con (ε^i) la sua duale, le componenti di un endomorfismo A formano una matrice quadrata (A_i^j) di ordine n . Particolarizzando al caso presente quanto visto al §2, essa è definita implicitamente dall'uguaglianza

$$(1) \quad \boxed{A(e_i) = A_i^j e_j}$$

Applicando a quest'uguaglianza gli elementi della base duale, tenuto conto delle relazioni evidenziate nel §3, si trova per le componenti di un endomorfismo la definizione esplicita seguente:

$$(2) \quad \boxed{A_i^j = \langle A(e_i), \varepsilon^j \rangle}$$

L'uguaglianza $v = A(u)$ si traduce, in componenti, nell'uguaglianza

$$(3) \quad \boxed{v^i = A_j^i u^j}$$

OSSERVAZIONE 1. La matrice delle componenti dell'identità 1 è, qualunque sia la base scelta, la matrice unitaria, vale a dire:

$$(1)_j^i = \delta_j^i.$$

OSSERVAZIONE 2. La matrice delle componenti di un prodotto di endomorfismi è il prodotto delle due rispettive matrici:

$$(4) \quad \boxed{(AB)_i^j = A_k^j B_i^k}$$

Si ha infatti, applicando la (2) e la (1):

$$\begin{aligned} (AB)_i^j &= \langle A(B(e_i)), \epsilon^j \rangle = \langle A(B_i^k e_k), \epsilon^j \rangle \\ &= B_i^k \langle A(e_k), \epsilon^j \rangle = B_i^k A_k^j. \end{aligned}$$

Si noti che se si conviene che l'indice in basso sia indice di colonna, il prodotto matriciale nella (4) è *righe per colonne*: righe della prima matrice (B_i^k) per le colonne della seconda matrice (A_k^j) . Si noti bene però che la scrittura $A_k^j B_i^k$ a secondo membro della (4) è del tutto equivalente a $B_i^k A_k^j$.

OSSERVAZIONE 3. Gli automorfismi sono caratterizzati dalla risolubilità del sistema lineare (3) rispetto alle (u^j) , quindi dalla condizione

$$\det(A_j^i) \neq 0.$$

La matrice delle componenti di A^{-1} è l'inversa della matrice (A_j^i) .

4.1. Il commutatore di endomorfismi

Le due operazioni binarie interne fondamentali nello spazio degli endomorfismi $End(E)$, consentono di definire una terza operazione binaria interna, il **commutatore**. Il commutatore di due endomorfismi A e B è l'endomorfismo

$$(5) \quad [A, B] = AB - BA.$$

Si ha $[A, B] = 0$ se e solo se i due endomorfismi commutano. Il commutatore $[\cdot, \cdot]$ gode delle seguenti proprietà:

$$(6) \quad \begin{cases} [A, B] = -[B, A], \\ [aA + bB, C] = a[A, C] + b[B, C], \\ [A, [B, C]] + [B, [C, A]] + [C, [A, B]] = 0. \end{cases}$$

Le prime due, di verifica immediata, mostrano che il commutatore è antisimmetrico e bilineare. La terza, la cui verifica richiede un semplice calcolo, prende il nome di **identità di Jacobi** (Carl Gustav Jacob Jacobi, 1804-1851). Uno spazio vettoriale dotato di un'operazione binaria interna soddisfacente a queste tre proprietà prende il nome di **algebra di Lie** (Marius Sophus Lie, 1842-1899). Dunque $End(E)$ è un'algebra di Lie. Vedremo nel corso di queste lezioni altri esempi di algebre di Lie.

Nelle (6) la prima proprietà, cioè la proprietà anticommutativa, può essere sostituita da

$$(7) \quad [A, A] = 0.$$

È infatti ovvio che la $(6)_1$ implica la (7). Viceversa, se vale la (7) si ha successivamente:

$$\begin{aligned} 0 &= [A + B, A + B] = [A, A] + [B, B] + [A, B] + [B, A] \\ &= [A, B] + [B, A]. \end{aligned}$$

Di qui la $(6)_1$.

4.2. Il determinante e la traccia di un endomorfismo

Sebbene la matrice delle componenti di un endomorfismo cambi al cambiare della base (fanno eccezione l'endomorfismo nullo e quello identico), vi sono delle grandezze ad essa associate che invece non risentono di questo cambiamento. Queste grandezze si chiamano **invarianti** dell'endomorfismo: pur essendo calcolate attraverso le componenti, non dipendono dalla scelta della base e sono quindi intrinsecamente legate all'endomorfismo. Il seguente teorema mette in evidenza i due *invarianti fondamentali*.

PROPOSIZIONE 1. Il determinante $\det(A_j^i)$ e la traccia A_i^i (cioè la somma degli elementi della diagonale principale) della matrice delle componenti di un endomorfismo non dipendono dalla scelta della base.

Dimostrazione. Siano (e_i) e $(e_{i'})$ due basi di E . Sussistono delle relazioni del tipo

$$(8) \quad \boxed{e_i = E_i^{i'} e_{i'}} \iff \boxed{e_{i'} = E_i^i e_i}$$

dove $(E_i^{i'})$ e (E_i^i) sono matrici regolari una inversa dell'altra, cioè tali che

$$(9) \quad E_i^{i'} E_j^j = \delta_i^j, \quad E_i^i E_j^{j'} = \delta_i^{j'}.$$

Siano (ϵ^i) e $(\epsilon^{i'})$ le corrispondenti basi duali. Sussistono i legami

$$(10) \quad \boxed{\epsilon^{i'} = E_i^{i'} \epsilon^i} \iff \boxed{\epsilon^i = E_i^i \epsilon^{i'}}$$

Siano (A_i^j) e $(A_i^{j'})$ le matrici delle componenti di un endomorfismo A rispetto alle due basi date. Stante la definizione (2) e le relazioni (8) e (10), si ha successivamente:

$$\begin{aligned} A_{i'}^{j'} &= \langle A(e_{i'}), \epsilon^{j'} \rangle = E_i^i \langle A(e_i), \epsilon^{j'} \rangle \\ &= E_i^i E_j^{j'} \langle A(e_i), \epsilon^j \rangle = E_i^i E_j^{j'} A_i^j \end{aligned}$$

Si ottiene quindi l'uguaglianza

$$(11) \quad \boxed{A_{i'}^{j'} = E_i^i E_j^{j'} A_i^j}$$

che esprime la *legge di trasformazione delle componenti di un endomorfismo al variare della base*. Ciò posto, poiché il determinante del prodotto di matrici è il prodotto dei determinanti, dalla (11) segue

$$\det(A_{i'}^{j'}) = \det(A_i^j)$$

e quindi che il determinante è invariante. Inoltre: $A_{i'}^{i'} = E_i^i E_j^{j'} A_i^j = \delta_j^i A_i^i = A_i^i$. Dunque anche la somma degli elementi della diagonale principale è invariante. ■

È allora lecito introdurre la seguente definizione:

DEFINIZIONE 1. Il **determinante** e la **traccia** di un endomorfismo $A \in \text{End}(E)$ sono i numeri

$$\det(A) = \det(A_j^i), \quad \text{tr}(A) = A_i^i,$$

dove (A_j^i) è la matrice delle componenti di A rispetto ad una qualunque base di E .

OSSERVAZIONE 4. Per la traccia ed il determinante di un endomorfismo valgono le seguenti proprietà:

$$(12) \quad \begin{cases} \text{tr}(A + B) = \text{tr}(A) + \text{tr}(B), \\ \text{tr}(aA) = a \text{tr}(A), \\ \text{tr}(AB) = \text{tr}(BA), \\ \text{tr}(1) = n, \\ \det(AB) = \det(A) \det(B), \\ \det(aA) = a^n \det(A), \\ \det(1) = 1. \end{cases}$$

4.3. Il polinomio caratteristico di un endomorfismo

Introduciamo ora un altro concetto fondamentale nella teoria degli endomorfismi lineari, che come prima cosa ci consentirà di definire ulteriori invarianti.

DEFINIZIONE 2. Dicesi **polinomio caratteristico** di un endomorfismo $A \in \text{End}(E)$ il polinomio di grado $n = \dim(E)$ nell'indeterminata x definito da:

$$(13) \quad P(A, x) = (-1)^n \det(A - x 1)$$

Osserviamo che, essendo definito da un determinante, il polinomio caratteristico è un invariante, non dipende cioè dalla scelta della base. Si ha pertanto, comunque si scelga la base

$$(14) \quad \begin{aligned} P(A, x) &= (-1)^n \det(A_j^i - x \delta_j^i) \\ &= (-1)^n \det \begin{pmatrix} A_1^1 - x & A_1^2 & \dots & A_1^n \\ A_2^1 & A_2^2 - x & \dots & A_2^n \\ \vdots & \vdots & \ddots & \vdots \\ A_n^1 & A_n^2 & \dots & A_n^n - x \end{pmatrix} \end{aligned}$$

Il segnante $(-1)^n$ che interviene in questa definizione ha il compito di rendere uguale a 1 il coefficiente di x^n . È conveniente scrivere per esteso il polinomio caratteristico nella forma seguente:

$$(15) \quad P(\mathbf{A}, x) = \sum_{k=0}^n (-1)^k I_k(\mathbf{A}) x^{n-k} \\ = x^n - I_1(\mathbf{A}) x^{n-1} + I_2(\mathbf{A}) x^{n-2} - \dots + (-1)^n I_n(\mathbf{A}).$$

Siccome questo polinomio è invariante, sono invarianti tutti i suoi coefficienti $I_k(\mathbf{A})$ ($k = 1, \dots, n$) (come si è osservato si ha semplicemente $I_0(\mathbf{A}) = 1$). Essi prendono il nome di **invarianti principali** di $\mathbf{A}$.

La speciale scrittura a segni alterni (15) è motivata dal seguente notevole fatto: confrontando lo sviluppo del determinante (14) con la scrittura (15), si osserva che $I_k(\mathbf{A})$ risulta essere la somma dei minori principali di ordine k della matrice (A_i^j) . Ricordiamo che un **minore principale** è il determinante di una sottomatrice quadrata principale. Una **sottomatrice** quadrata si dice **principale** se la sua diagonale principale è formata da elementi della diagonale principale della matrice che la contiene. Quindi in particolare,

$$(16) \quad I_n(\mathbf{A}) = \det(\mathbf{A}), \quad I_1(\mathbf{A}) = \text{tr}(\mathbf{A}),$$

per cui tra gli invarianti principali ritroviamo il determinante e la traccia.

Si può constatare direttamente questo fatto per esempio nel caso $n = 3$. La (14) diventa

$$P(\mathbf{A}, x) = - \det \begin{pmatrix} A_1^1 - x & A_1^2 & A_1^3 \\ A_2^1 & A_2^2 - x & A_2^3 \\ A_3^1 & A_3^2 & A_3^3 - x \end{pmatrix}.$$

Ponendo $x = 0$ si vede subito che il termine noto del polinomio caratteristico è $-\det(A_i^j)$, per cui, essendo per la (15),

$$P(\mathbf{A}, x) = x^3 - I_1 x^2 + I_2 x - I_3$$

(scriviamo più semplicemente I_k al posto di $I_k(\mathbf{A})$), risulta $I_3 = \det(A_i^j)$. Per quel che riguarda il coefficiente I_2 di x , si osserva che esso si ottiene derivando il polinomio caratteristico rispetto a x e ponendo poi $x = 0$. Ricordando la regola di derivazione di un determinante (si derivano ad

una ad una le singole colonne, per esempio, e si sommano i determinanti così ottenuti), si trova che (usiamo il simbolo $| \cdot |$ per il determinante):

$$I_2 = - \begin{vmatrix} -1 & A_1^2 & A_1^3 \\ 0 & A_2^2 & A_2^3 \\ 0 & A_3^2 & A_3^3 \end{vmatrix} - \begin{vmatrix} A_1^1 & 0 & A_1^3 \\ A_2^1 & -1 & A_2^3 \\ A_3^1 & 0 & A_3^3 \end{vmatrix} - \begin{vmatrix} A_1^1 & A_1^2 & 0 \\ A_2^1 & A_2^2 & 0 \\ A_3^1 & A_3^2 & -1 \end{vmatrix} \\ = \begin{vmatrix} A_2^2 & A_2^3 \\ A_3^2 & A_3^3 \end{vmatrix} + \begin{vmatrix} A_1^1 & A_1^3 \\ A_3^1 & A_3^3 \end{vmatrix} + \begin{vmatrix} A_1^1 & A_1^2 \\ A_2^1 & A_2^2 \end{vmatrix},$$

come asserito. Infine, il coefficiente $-I_1$ di x^2 si ottiene soltanto, nello sviluppo del determinante, dal prodotto degli elementi della diagonale principale,

$$(A_1^1 - x)(A_2^2 - x)(A_3^3 - x) = (A_1^1 + A_2^2 + A_3^3)x^2 + \dots$$

per cui $I_1 = A_1^1 + A_2^2 + A_3^3$.

4.4. Autovalori e autovettori di un endomorfismo

Un **sottospazio** $S \subset E$ si dice **invariante** rispetto ad un endomorfismo $\mathbf{A}$ se $\mathbf{A}(S) \subset S$. Un vettore $\mathbf{v} \in E$ non nullo si dice **autovettore** di $\mathbf{A}$ se esiste un numero reale λ tale da aversi

$$(17) \quad \mathbf{A}(\mathbf{v}) = \lambda \mathbf{v}.$$

Un autovettore genera un sottospazio di autovettori (se si include lo zero), di dimensione 1 e invariante, che prende il nome di **autodirezione**.

L'equazione (17) può anche scriversi

$$(18) \quad (\mathbf{A} - \lambda \mathbf{1})(\mathbf{v}) = 0.$$

Scelta comunque una base di E , la (18) si traduce in un sistema di n equazioni lineari omogenee nelle incognite (v^i) , componenti dell'autovettore $\mathbf{v}$:

$$(19) \quad (A_i^j - \lambda \delta_i^j) v^i = 0.$$

Tale sistema ammette soluzioni non banali se e solo se λ annulla il determinante della matrice dei coefficienti, cioè, vista la definizione (13),

se e solo se λ è radice del polinomio caratteristico di A . Le radici del polinomio caratteristico, cioè le soluzioni dell'**equazione caratteristica**, algebrica di grado n ,

$$P(A, x) = 0,$$

prendono il nome di **autovalori** di A . Il loro insieme $\{\lambda_1, \lambda_2, \dots, \lambda_n\}$ prende il nome di **spettro** di A .

Gli autovalori sono invarianti, nel senso sopra detto. Vista la forma (14) con cui si scrive il polinomio caratteristico, dai noti teoremi sulle radici delle equazioni algebriche segue che l'invariante principale $I_k(A)$ è la somma di tutti i possibili prodotti di k autovalori. Risulta in particolare:

$$(20) \quad \begin{aligned} I_1(A) &= \text{tr}(A) = \lambda_1 + \lambda_2 + \dots + \lambda_n, \\ I_n(A) &= \det(A) = \lambda_1 \lambda_2 \dots \lambda_n. \end{aligned}$$

Chiamiamo **molteplicità algebrica** di un autovalore λ , e la denotiamo con $\text{ma}(\lambda)$, la sua molteplicità come radice del polinomio caratteristico. Gli autovalori possono essere reali o complessi. Ad ogni autovalore reale λ corrisponde un sottospazio $K_\lambda \subset E$ di autovettori ad esso associati, soluzioni cioè dell'equazione (17). Chiamiamo **molteplicità geometrica** di λ la dimensione del sottospazio invariante associato K_λ , e la denotiamo con $\text{mg}(\lambda)$. Si dimostra (la dimostrazione è in appendice al paragrafo) che è sempre

$$(21) \quad \text{ma}(\lambda) \geq \text{mg}(\lambda).$$

Tuttavia, per alcuni tipi di endomorfismi vale sempre l'uguaglianza, per esempio per gli endomorfismi simmetrici in spazi strettamente euclidei, che esamineremo più avanti. Osserviamo che sottospazi invarianti K_{λ_1} e K_{λ_2} corrispondenti ad autovalori distinti hanno intersezione nulla:

$$(22) \quad \lambda_1 \neq \lambda_2 \Rightarrow K_{\lambda_1} \cap K_{\lambda_2} = 0.$$

Se λ è un autovalore complesso allora gli autovettori associati sono complessi: il sistema (18) fornisce soluzioni complesse $v^i = u^i + i w^i$.

Si è così condotti a considerare la **complessificazione** dello spazio E , cioè lo spazio $\mathbb{C}E$ delle combinazioni $v = u + i w$, con $u, w \in E$, per le quali si definiscono le due operazioni di somma e di prodotto per un numero complesso in maniera formalmente analoga a quella della somma e del prodotto di numeri complessi.

Poiché l'equazione caratteristica ha coefficienti reali, ad ogni autovalore complesso $\lambda = a + ib$ corrisponde un autovalore coniugato $\bar{\lambda} = a - ib$ avente la stessa molteplicità. Si osserva inoltre che se $v = u + i w$ è autovettore associato a λ , allora il vettore coniugato $\bar{v} = u - i w$ è autovettore associato a $\bar{\lambda}$, e che ogni altro vettore complesso ottenuto moltiplicando v per un numero complesso qualunque è ancora autovettore associato a λ .

OSSERVAZIONE 5. Se $v = u + i w$ è un autovettore associato ad un autovalore *strettamente* complesso $\lambda = a + ib$ (cioè tale che $b \neq 0$), allora i vettori reali componenti (u, w) sono indipendenti. In caso contrario infatti, da una relazione del tipo $u = c w$ seguirebbe $v = (c + i) w$ e quindi dalla (17), dividendo per $c + i$, $A(w) = \lambda w$ con λ solo elemento complesso: assurdo. Decomposta allora la (17) nella sua parte reale e immaginaria, risulta

$$(23) \quad A(u) = a u - b w, \quad A(w) = a w + b u.$$

Di qui si osserva che il sottospazio generato dai vettori indipendenti (u, w) è invariante in A . Sia Π questo sottospazio (bidimensionale) e si Π' il sottospazio invariante corrispondente ad un secondo vettore complesso $v' = u' + i w'$ sempre associato a λ . Non può che essere $\Pi = \Pi'$ oppure $\Pi \cap \Pi' = 0$. Infatti, se un vettore non nullo x appartenesse all'intersezione, con opportuna moltiplicazione degli autovettori per un numero complesso, ci si può ricondurre al caso in cui $u = u' = x$. Dalla prima delle (23) seguirebbe

$$A(x) = a x - b w, \quad A(x) = a x + b w,$$

da cui $b(w - w') = 0$, quindi $w = w'$. Ciò significa che $v = v'$ e quindi che $\Pi = \Pi'$.

Da queste proprietà si deduce che ogni autovalore complesso λ individua un sottospazio invariante K_λ che è somma diretta di sottospazi invarianti bidimensionali. La sua dimensione è quindi pari e inoltre $K_\lambda = K_{\bar{\lambda}}$. In questo caso per molteplicità geometrica di λ s'intende la metà della dimensione di K_λ . Vale ancora la disuguaglianza (21), mentre la (22) va sostituita dalla più generale

$$(24) \quad \lambda_1 \neq \lambda_2, \quad \lambda_1 \neq \bar{\lambda}_2 \Rightarrow K_{\lambda_1} \cap K_{\lambda_2} = 0.$$

Dimostrazione della disuguaglianza (21). Sia λ un autovalore reale di un endomorfismo A . L'equazione caratteristica di A può ovviamente anche scriversi

$$\det((A - \lambda 1) - (x - \lambda) 1) = 0,$$

cioè

$$(\dagger) \quad (x - \lambda)^n - I_1(A - \lambda 1)(x - \lambda)^{n-1} + \dots + (-1)^n I_n(A - \lambda 1) = 0,$$

facendo intervenire gli invarianti principali dell'endomorfismo $A - \lambda 1$. Osserviamo ora che lo spazio invariante K_λ degli autovettori associati a λ è il nucleo di questo endomorfismo, e quindi

$$\text{mg}(\lambda) = n - \text{rank}(A - \lambda 1) = n - r.$$

Il rango r dell'endomorfismo $(A - \lambda 1)$ coincide con il rango della matrice $(A_i^j - \lambda \delta_i^j)$, comunque si scelga la base. Segue che gli invarianti

$$I_n(A - \lambda 1), \dots, I_{r+1}(A - \lambda 1),$$

sono tutti nulli perché somma di minori principali di ordine maggiore di r della matrice $(A_i^j - \lambda \delta_i^j)$. Pertanto l'equazione caratteristica $\dagger$ si riduce in effetti all'equazione

$$(x - \lambda)^n - I_1(A - \lambda 1)(x - \lambda)^{n-1} + \dots + (-1)^r I_r(A - \lambda 1)(x - \lambda)^{n-r} = 0.$$

Di qui si vede che λ è una radice di ordine $n - r$ almeno, perché, pur essendo per ipotesi i minori di ordine r non tutti nulli, non si può escludere che la somma di quelli principali, vale a dire l'invariante $I_r(A - \lambda 1)$, non si annulli. Dunque $\text{mg}(\lambda) \geq n - r$. Questa dimostrazione vale anche nel caso in cui λ è complesso, perché anche in questo caso $\text{mg}(\lambda) = n - r$ dove r è il rango della matrice $(A_i^j - \lambda \delta_i^j)$. ■

4.5. Il teorema di Hamilton-Cayley

Dato un endomorfismo A , se ne possono considerare le potenze:

$$A^0 = 1, \quad A^1 = A, \quad A^2 = AA, \quad \dots \quad A^k = \underbrace{AA \dots A}_{k \text{ volte}}, \quad \dots$$

Quindi ad ogni polinomio di grado k

$$P_k(x) = a_0 x^k + a_1 x^{k-1} + \dots + a_k$$

si può associare l'endomorfismo

$$P_k(A) = a_0 A^k + a_1 A^{k-1} + \dots + a_k 1.$$

Un polinomio $P_k(x)$ si dice **annichilante** di A se $P_k(A)$ si annulla identicamente.

Sussiste a questo proposito il notevole **teorema di Hamilton-Cayley** (William Rowan Hamilton, 1805-1865; Arthur Cayley, 1821-1895): *il polinomio caratteristico di un endomorfismo è annichilante:*

$$A^n - I_1(A) A^{n-1} + I_2(A) A^{n-2} - \dots + (-1)^n I_n(A) 1 = 0.$$

Una prima conseguenza di questo teorema è che tutte le potenze di un endomorfismo A formano un sottospazio di $\text{End}(E)$ al più di dimensione n , generato dall'identità 1 e dalle prime $n - 1$ potenze di A .

Dimostrazione del teorema di Hamilton-Cayley. Basiamo la dimostrazione sull'osservazione seguente. Se sussiste un'uguaglianza del tipo

$$(25) \quad C(x)(A - x 1) = P_k(x) 1$$

dove $P_k(x)$ è un polinomio di grado k e dove $C(x)$ è un polinomio di grado $k - 1$ in x a coefficienti in $\text{End}(E)$, cioè del tipo

$$C(x) = x^{k-1} C_1 + x^{k-2} C_2 + \dots + x C_{k-1} + C_k,$$

allora $P_k(A) = 0$ cioè $P_k(x)$ è annichilante di A . Infatti, esplicitata questa uguaglianza

$$\begin{aligned} & (x^{k-1} C_1 + x^{n-2} C_2 + \dots + x C_{n-1} + C_n)(A - x 1) \\ & = (a_0 x^k + a_1 x^{k-1} + \dots + a_k) 1 \end{aligned}$$

ed uguagliati i coefficienti delle varie potenze di x , si trova che

$$\begin{cases} -C_1 = a_0 1, \\ C_1 A - C_2 = a_1 1, \\ \dots \\ C_{k-1} A - C_k = a_{k-1} 1, \\ C_k A = a_k 1. \end{cases}$$

Moltiplicando a destra rispettivamente per $A^k, A^{k-1}, \dots, A$ le prime k di queste uguaglianze, e sommando poi membro a membro si trova $0 \doteq P_k(A)$, come asserito. Ciò posto, si tiene conto del fatto, dimostrato più avanti (§7, Oss. 5), che per ogni endomorfismo B esiste sempre un endomorfismo $\tilde{B}$, l'endomorfismo complementare, per cui

$$\tilde{B}B = (\det(B))1.$$

Se ora si prende l'endomorfismo $B = A - x1$, con indeterminata x , per la definizione di polinomio caratteristico l'uguaglianza precedente diventa

$$(-1)^n \tilde{B}(A - x1) = P(A, x)1.$$

Questa è del tipo (25) (con $k = n$ e $C(x) = (-1)^n \tilde{B}$) perché l'endomorfismo complementare $\tilde{B}$ ha come componenti dei determinanti di ordine $n-1$ della matrice $(B_j^i) = (A_j^i - x\delta_j^i)$, ed è quindi un polinomio di grado $n-1$ in x . Pertanto $P(A, A) = 0$. ■

5. Forme bilineari

Una **forma bilineare** su di uno spazio vettoriale E è un'applicazione bilineare di $E \times E$ in $\mathbb{R}$. Lo spazio delle forme bilineari, in accordo con quanto detto all'inizio di questo capitolo, è denotato con $L(E, E; \mathbb{R})$.

Un esempio elementare notevole di forma bilineare è dato dal **prodotto tensoriale** $\alpha \otimes \beta$ di due forme lineari $\alpha, \beta \in E^*$. Esso è definito dall'uguaglianza

$$(1) \quad (\alpha \otimes \beta)(u, v) = \langle u, \alpha \rangle \langle v, \beta \rangle$$

Se (ϵ^i) è la base duale di una base (e_i) di E allora gli n^2 prodotti $\epsilon^i \otimes \epsilon^j$ formano una base di $L(E, E; \mathbb{R})$. Data una forma bilineare φ , i numeri

$$(2) \quad \varphi_{ij} = \varphi(e_i, e_j)$$

sono le componenti della forma secondo la base $\epsilon^i \otimes \epsilon^j$:

$$(3) \quad \varphi = \varphi_{ij} \epsilon^i \otimes \epsilon^j$$

Con abuso di linguaggio essi si dicono anche componenti di φ secondo la base (e_i) . Per ogni coppia di vettori si ha:

$$(4) \quad \varphi(u, v) = \varphi_{ij} u^i v^j$$

La dimostrazione di quanto ora affermato è lasciata come esercizio.

Le componenti (φ_{ij}) formano una matrice quadrata $n \times n$. Conveniamo che il primo indice sia indice di riga, il secondo di colonna.

Al cambiare della base le componenti di una forma bilineare cambiano con la legge seguente,

$$(5) \quad \varphi_{i'j'} = E_{i'}^i E_{j'}^j \varphi_{ij},$$

deducibile dalla definizione (2) con un metodo analogo a quello visto per gli endomorfismi. Si vede dalla (5) che, a differenza di ciò che accade per gli endomorfismi, il determinante della matrice delle componenti di una forma bilineare non è un invariante. Si può però dimostrare che il rango della matrice è un invariante. Questo pertanto viene definito come rango della forma bilineare. In particolare, una forma bilineare si dice **regolare** o **non degenerare** se vale la condizione

$$(6) \quad \varphi(u, v) = 0, \forall v \in E \implies u = 0$$

oppure la condizione equivalente

$$(6') \quad \varphi(u, v) = 0, \forall u \in E \implies v = 0.$$

Entrambe sono equivalenti alla condizione

$$(6'') \quad \det(\varphi_{ij}) \neq 0.$$

In questo caso la forma bilineare ha rango massimo ($= n$). Che la (6'') sia equivalente alla (6) segue dal fatto che, in componenti, l'implicazione (6) si scrive:

$$\varphi_{ij} u^i v^j = 0, \forall (v^i) \in \mathbb{R}^n \implies u^i = 0.$$

Questa si semplifica in

$$\varphi_{ij} u^i = 0 \implies u^i = 0,$$

ed esprime la circostanza che il sistema lineare omogeneo (di n equazioni in n incognite) $\varphi_{ij} u^i = 0$ ammette solo la soluzione nulla.

DEFINIZIONE 1. La **trasposta** di una forma bilineare φ è la forma bilineare φ^T definita da

$$(7) \quad \boxed{\varphi^T(u, v) = \varphi(v, u)}$$

Si ha ovviamente $\varphi^{TT} = \varphi$.

DEFINIZIONE 2. Una forma bilineare si dice **simmetrica** (rispettivamente, **antisimmetrica**) se

$$(8) \quad \varphi^T = \varphi \quad (\text{resp. } \varphi^T = -\varphi),$$

cioè se

$$(9) \quad \varphi(u, v) = \varphi(v, u), \quad (\text{resp. } \varphi(u, v) = -\varphi(v, u)).$$

Questa condizione si esprime, qualunque sia la base scelta, nella simmetria (resp. nella antisimmetria) della matrice delle componenti (basta porre $(u, v) = (e_i, e_j)$ nella (7)):

$$(10) \quad \varphi_{ij} = \varphi_{ji}, \quad (\text{resp. } \varphi_{ij} = -\varphi_{ji}).$$

Una forma bilineare φ è decomponibile nella somma della sua **parte simmetrica** φ^S e della sua **parte antisimmetrica** φ^A . Vale cioè l'uguaglianza

$$(11) \quad \boxed{\varphi = \varphi^S + \varphi^A}$$

posto

$$\boxed{\varphi^S = \frac{1}{2}(\varphi + \varphi^T), \quad \varphi^A = \frac{1}{2}(\varphi - \varphi^T)}$$

5.1. Forme bilineari simmetriche e forme quadratiche

Esaminiamo ora alcune proprietà fondamentali delle forme bilineari simmetriche.

Sia φ una forma bilineare simmetrica. Due vettori si dicono **coniugati** (o anche **ortogonali**) rispetto a φ se $\varphi(u, v) = 0$. In particolare un vettore si dice **isotropo** se è coniugato con se stesso, $\varphi(u, u) = 0$, **unitario** se invece $\varphi(u, u) = \pm 1$. Una base (e_i) di E si dice **canonica** (sempre rispetto a φ) se tutti i suoi vettori sono unitari o isotropi e fra loro coniugati, cioè se, dopo un eventuale riordinamento della base, la matrice delle componenti assume la forma:

$$(12) \quad (\varphi_{ij}) = \begin{pmatrix} \mathbf{1}_p & 0 & 0 \\ 0 & -\mathbf{1}_q & 0 \\ 0 & 0 & 0 \end{pmatrix}$$

dove $\mathbf{1}_p$ e $\mathbf{1}_q$ denotano le matrici unitarie di ordine p e q rispettivamente. In una base canonica si ha quindi (si confronti con la (4)):

$$(13) \quad \varphi(u, v) = \sum_{a=1}^p u^a v^a - \sum_{b=p+1}^{p+q} u^b v^b.$$

Si dimostra (la dimostrazione è in appendice al paragrafo) che esistono basi canoniche e inoltre che in ogni base canonica la matrice delle componenti è sempre la (12), vale a dire che gli interi p e q sono invarianti rispetto alla scelta di basi canoniche. Questa notevole proprietà è nota come **teorema di Sylvester** o **legge di inerzia delle forme quadratiche** (James Joseph Sylvester, 1814-1897). Alla coppia di interi non negativi (p, q) si dà il nome di **segnatura** della forma bilineare simmetrica φ .

Si osserva dalla (12) che la somma $p + q$ è uguale al rango di φ .

Per calcolare la segnatura di una forma bilineare simmetrica φ basterebbe determinare una base canonica (o almeno una base di vettori fra loro ortogonali) (e_i) e valutare quindi i segni dei numeri $\varphi(e_i, e_i)$. La determinazione di una base canonica (si veda la dimostrazione della loro esistenza a fine paragrafo) può tuttavia, nella pratica, risultare laboriosa. Si può allora fare ricorso ad un metodo assai rapido, illustrato nel seguente enunciato^(†)

PROPOSIZIONE 1. Sia (φ_{ij}) la matrice delle componenti di una forma bilineare simmetrica φ relativa ad una base *qualsiasi*. Si consideri

^(†) Una dimostrazione si trova in F.G. TRICOMI, *Lezioni di Analisi Matematica I* (CEDAM, Padova).

una *qualunque* successione di sottomatrici quadrate principali, di ordine crescente da 1 a n , ciascuna propriamente contenuta nella successiva:

$$M_1, M_2, \dots, M_n = (\varphi_{ij}).$$

Allora l'intero q della segnatura (p, q) è uguale al numero delle variazioni di segno della successione di numeri

$$1, M_1, \det(M_2), \dots, \det(M_n) = \det(\varphi_{ij}),$$

dalla quale siano stati eliminati gli eventuali zeri.

Calcolato l'intero q , l'intero p viene poi calcolato conoscendo il rango r della matrice: $p = r - q$. Si noti bene che questo metodo per il calcolo della segnatura lascia piena libertà di scelta della base e delle sottomatrici della successione.

Una **forma quadratica** su di uno spazio vettoriale E è un'applicazione $\Phi: E \rightarrow \mathbb{R}$ tale che l'applicazione

$$\delta\Phi: E \times E \rightarrow \mathbb{R},$$

detta **polarizzazione** di Φ è definita da

$$(14) \quad \delta\Phi(u, v) = \frac{1}{2}(\Phi(u+v) - \Phi(u) - \Phi(v)),$$

è una forma bilineare (ovviamente simmetrica). Pertanto, per la sua stessa definizione, ad una forma quadratica corrisponde una forma bilineare simmetrica. Viceversa, se φ è una forma bilineare simmetrica allora l'applicazione $\Phi: E \rightarrow \mathbb{R}$ definita da

$$(15) \quad \boxed{\Phi(u) = \varphi(u, u) = \varphi_{ij} u^i u^j}$$

è una forma quadratica tale che $\delta\Phi = \varphi$. Infatti:

$$\begin{aligned} \delta\Phi(u, v) &= \frac{1}{2}(\Phi(u+v) - \Phi(u) - \Phi(v)) \\ &= \frac{1}{2}(\varphi(u+v, u+v) - \varphi(u, u) - \varphi(v, v)) \\ &= \frac{1}{2}(\varphi(u, u) + \varphi(v, v) + 2\varphi(u, v) - \varphi(u, u) - \varphi(v, v)) \\ &= \varphi(u, v). \end{aligned}$$

Vi è dunque una corrispondenza biunivoca fra forme quadratiche e forme bilineari simmetriche. Stante questa corrispondenza tutte le definizioni e le proprietà stabilite per le forme bilineari si trasferiscono alle forme quadratiche, e viceversa.

Una forma quadratica Φ si dice **semi-definita positiva** (rispettivamente, **semi-definita negativa**) se

$$\Phi(u) \geq 0 \quad (\text{resp. } \Phi(u) \leq 0) \quad \forall u \in E.$$

In particolare si dice **definita positiva** (risp. **definita negativa**) se essa si annulla solo sul vettore nullo, cioè se oltre alla condizione precedente vale l'implicazione

$$\Phi(u) = 0 \Rightarrow u = 0.$$

Si dice **non definita o indefinita** in tutti gli altri casi.

In una base canonica si ha (si veda la (13))

$$(16) \quad \Phi(u) = \sum_{a=1}^p (u^a)^2 - \sum_{b=p+1}^{p+q} (u^b)^2.$$

Di qui si osserva che: (a) una forma quadratica è semidefinita positiva se $q = 0$; (b) è definita positiva se $p = n$. Analogamente per il caso "negativo".

Per una forma bilineare simmetrica semi-definita positiva (o negativa) vale la **disuguaglianza di Schwarz** (Hermann Schwarz, 1843-1921)

$$(17) \quad (\varphi(u, v))^2 \leq \Phi(u)\Phi(v), \quad \forall u, v \in E.$$

Se in particolare φ è definita positiva (o negativa), allora si ha l'uguaglianza se e solo se u e v sono dipendenti.

Dimostrazione dell'esistenza di basi canoniche. Basta dimostrare l'esistenza di basi ortogonali, per cui cioè $\varphi(e_i, e_j) = 0$ per $i \neq j$. Se $\varphi = 0$ ogni base è ortogonale. Supponiamo $\varphi \neq 0$ e procediamo per induzione sull'intero $n = \dim(E)$. Per $n = 1$ ogni vettore non nullo $e \in E$ è una base ortogonale per qualunque forma bilineare simmetrica su E . Sia allora $n > 1$ e si supponga che ogni forma bilineare simmetrica

su di uno spazio di dimensione $< n$ ammetta una base ortogonale. Posto che $\varphi \neq 0$, esiste almeno un vettore $\mathbf{u} \in E$ per cui $\varphi(\mathbf{u}, \mathbf{u}) = \Phi(\mathbf{u}) \neq 0$. Si consideri allora il sottospazio $U = \{\mathbf{v} \in E \mid \varphi(\mathbf{u}, \mathbf{v}) = 0\}$ dei vettori ortogonali a $\mathbf{u}$. Siccome $\mathbf{u} \notin U$, abbiamo $m = \dim(U) \leq n - 1$. Per la restrizione a $U \times U$ di φ esiste allora una base ortogonale $(\mathbf{e}_1, \dots, \mathbf{e}_m)$. D'altra parte, posto per ogni $\mathbf{v} \in E$

$$v^n = \frac{\varphi(\mathbf{u}, \mathbf{v})}{\varphi(\mathbf{u}, \mathbf{u})},$$

si vede che $\varphi(\mathbf{u}, \mathbf{v} - v^n \mathbf{u}) = 0$. Dunque $\mathbf{v} - v^n \mathbf{u} \in U$ e quindi

$$\mathbf{v} = v^n \mathbf{u} + \sum_{i=1}^m v^i \mathbf{e}_i.$$

Ciò mostra che i vettori indipendenti $(\mathbf{e}_1, \dots, \mathbf{e}_m, \mathbf{u})$ formano una base ortogonale di E (e quindi anche che $m = n - 1$). ■

Dimostrazione del Teorema di Sylvester. Siano $(\mathbf{e}_i)$ e $(\mathbf{e}_{i'})$ due basi canoniche per la forma bilineare simmetrica φ . Per la (16) si ha, per ogni $\mathbf{u} \in E$, l'uguaglianza

$$(18) \quad \Phi(\mathbf{u}) = \sum_{a=1}^p (u^a)^2 - \sum_{b=p+1}^{p+q} (u^b)^2 = \sum_{a'=1}^{p'} (u^{a'})^2 - \sum_{b'=p'+1}^{p'+q'} (u^{b'})^2.$$

Si ha certamente $p' + q' = p + q = \text{rank}(\varphi)$. Occorre dimostrare che $p' = p$. Si supponga per assurdo $p' > p$ (e quindi $q' < q$). Considerata la relazione lineare $\mathbf{e}_i = E_i^{i'} \mathbf{e}_{i'}$ tra le due basi, si scriva il sistema lineare omogeneo

$$(19) \quad \sum_{b=p+1}^{p+q} E_b^{b'} x^b = 0, \quad b' = p' + 1, \dots, p' + q'.$$

Si tratta di q' equazioni nelle q incognite $(x^{p+1}, \dots, x^{p+q})$. Siccome $q' < q$, questo sistema ammette certamente una soluzione non nulla $(u^{p+1}, \dots, u^{p+q})$. Sia $\mathbf{u} \in E$ il vettore che ha tutte le altre componenti nulle: $u^1 = \dots = u^p = u^{p+q+1} = \dots = u^n = 0$. Per la prima parte della (18) si ha sicuramente $\Phi(\mathbf{u}) < 0$. Con una tale scelta si ha d'altra parte

$$u^{b'} = E_i^{b'} u^i = \sum_{b=p+1}^{p+q} E_b^{b'} u^b = 0$$

perché $(u^{p+1}, \dots, u^{p+q})$ è una soluzione del sistema (19). Dalla seconda della (18) segue allora $\Phi(\mathbf{u}) \geq 0$: assurdo. Per simmetria l'ipotesi $p' < p$ conduce pure ad un assurdo. Dunque $p' = p$. ■

Dimostrazione della disuguaglianza di Schwarz. Fissati i vettori $\mathbf{u}$ e $\mathbf{v}$, qualunque sia $x \in \mathbb{R}$, si ha $\Phi(\mathbf{u} + x\mathbf{v}) \geq 0$ perché Φ è semidefinita positiva. Questa disuguaglianza si traduce in

$$(20) \quad T(x) \equiv x^2 \Phi(\mathbf{v}) + 2x \varphi(\mathbf{u}, \mathbf{v}) + \Phi(\mathbf{u}) \geq 0,$$

dove $\Phi(\mathbf{v}) \geq 0$. Se $\Phi(\mathbf{v}) > 0$, il discriminante $\Delta = (\varphi(\mathbf{u}, \mathbf{v}))^2 - \Phi(\mathbf{u})\Phi(\mathbf{v})$ del trinomio $T(x)$ non può essere positivo, altrimenti $T(x)$ avrebbe due radici reali e distinte e la (20) sarebbe violata. Da $\Delta \leq 0$ segue la disuguaglianza (17). Se $\Phi(\mathbf{v}) = 0$, perché la (20) sia sempre soddisfatta deve essere necessariamente $\varphi(\mathbf{u}, \mathbf{v}) = 0$. La (17) è allora soddisfatta come uguaglianza. Se $\mathbf{u}$ è proporzionale a $\mathbf{v}$ (in particolare nullo), si verifica subito che vale l'uguaglianza nella (17). Viceversa, se vale l'uguaglianza risulta $\Delta = 0$ e pertanto il trinomio $T(x)$ ammette una radice doppia x_0 tale che $\Phi(\mathbf{u} + x_0 \mathbf{v}) = 0$. Se Φ è definita positiva deve necessariamente essere $\mathbf{u} + x_0 \mathbf{v} = 0$. ■

5.2. Forme bilineari antisimmetriche

Dati due covettori $\alpha, \beta \in E^*$ diciamo loro **prodotto esterno** la forma bilineare antisimmetrica $\alpha \wedge \beta$ definita da

$$(21) \quad (\alpha \wedge \beta)(\mathbf{u}, \mathbf{v}) = \langle \mathbf{u}, \alpha \rangle \langle \mathbf{v}, \beta \rangle - \langle \mathbf{u}, \beta \rangle \langle \mathbf{v}, \alpha \rangle$$

Richiamata la definizione (1) di prodotto tensoriale di due covettori, osserviamo che

$$(22) \quad \alpha \wedge \beta = \alpha \otimes \beta - \beta \otimes \alpha$$

Il prodotto esterno fra covettori è antisimmetrico. Per due covettori si ha

$$\alpha \wedge \beta = -\beta \wedge \alpha.$$

Inoltre

$$\alpha \wedge \alpha = 0.$$

Le forme bilineari antisimmetriche formano uno spazio vettoriale che denotiamo con $\Lambda^2(E)$. Se si considera una base duale (ϵ^i) , si possono di conseguenza considerare i prodotti esterni $\epsilon^i \wedge \epsilon^j$. Essi generano tutto lo spazio $\Lambda^2(E)$, valendo per ogni forma bilineare antisimmetrica φ la rappresentazione

$$(23) \quad \boxed{\varphi = \frac{1}{2} \varphi_{ij} \epsilon^i \wedge \epsilon^j}$$

posto, come nella (2), $\varphi_{ij} = \varphi(e_i, e_j)$. Infatti, valendo la (23) con (φ_{ij}) antisimmetriche, si ha successivamente, vista la definizione (21) di prodotto esterno:

$$\begin{aligned} \varphi(e_h, e_k) &= \frac{1}{2} \varphi_{ij} \epsilon^i \wedge \epsilon^j(e_h, e_k) \\ &= \frac{1}{2} \varphi_{ij} (\langle e_h, \epsilon^i \rangle \langle e_k, \epsilon^j \rangle - \langle e_h, \epsilon^j \rangle \langle e_k, \epsilon^i \rangle) \\ &= \frac{1}{2} \varphi_{ij} (\delta_h^i \delta_k^j - \delta_h^j \delta_k^i) \\ &= \frac{1}{2} (\varphi_{hk} - \varphi_{kh}) \\ &= \varphi_{hk}. \end{aligned}$$

Cosicché si ritrova la (2). Viceversa, se si considera la combinazione a secondo membro della (23) con i coefficienti dati dalla (2), si vede con calcolo analogo che tale combinazione coincide proprio con la φ :

$$\begin{aligned} \frac{1}{2} \varphi_{ij} \epsilon^i \wedge \epsilon^j(v_1, v_2) &= \frac{1}{2} \varphi(e_i, e_j) \epsilon^i \wedge \epsilon^j(v_1, v_2) \\ &= \frac{1}{2} \varphi(e_i, e_j) (\langle v_1, \epsilon^i \rangle \langle v_2, \epsilon^j \rangle - \langle v_1, \epsilon^j \rangle \langle v_2, \epsilon^i \rangle) \\ &= \frac{1}{2} \varphi(e_i, e_j) (v_1^i v_2^j - v_1^j v_2^i) \\ &= \frac{1}{2} \varphi(v_1, v_2) - \frac{1}{2} \varphi(v_2, v_1) \\ &= \varphi(v_1, v_2). \end{aligned}$$

I generatori $\epsilon^i \wedge \epsilon^j$ non sono tutti indipendenti, stante l'antisimmetria del prodotto esterno. Per ottenere una base, basta per esempio considerare tutti gli elementi $\epsilon^i \wedge \epsilon^j$ per cui $i < j$. Se ne deduce che

$$(24) \quad \dim(\Lambda^2(E)) = \frac{1}{2} n(n-1).$$

Si ritornerà su queste considerazioni, generalizzandole, al §7.

Sia data una forma bilineare antisimmetrica φ . Una base (e_i) di E si dice **canonica** rispetto a φ se la matrice delle componenti ha la forma seguente:

$$(25) \quad (\varphi_{ij}) = \begin{pmatrix} 0 & 1_p & 0 \\ -1_p & 0 & 0 \\ 0 & 0 & 0 \end{pmatrix}.$$

L'intero p è un invariante: infatti, come bene si vede dalla (25), il rango di (φ_{ij}) , che è invariante, è uguale a $2p$. In una base canonica si ha (si sostituisca la (25) nella (4)):

$$(26) \quad \varphi(u, v) = \sum_{a=1}^p (u^a v^{a+p} - u^{a+p} v^a),$$

vale a dire (si sostituisca la (25) nella (23)):

$$(27) \quad \varphi = \sum_{a=1}^p \epsilon^a \wedge \epsilon^{a+p}.$$

Dimostrazione dell'esistenza di basi canoniche per una forma bilineare antisimmetrica. Supposto $\varphi \neq 0$, altrimenti ogni base è canonica, si fissi una base *qualunque* $(\epsilon^{i'})$ di E^* e si consideri la rappresentazione

$$\varphi = \sum_{i' < j'} \varphi_{i'j'} \epsilon^{i'} \wedge \epsilon^{j'}.$$

A meno di un riordinamento possiamo sempre supporre $\varphi_{1'2'} \neq 0$ e quindi scivere

$$\varphi = \epsilon^{1'} \wedge \bar{\epsilon}^{1'} + \varphi^{1'}$$

posto

$$\begin{aligned} \epsilon^{1'} &= \epsilon^{1'} - \frac{1}{\varphi_{1'2'}} (\varphi_{2'3'} \epsilon^{3'} + \dots + \varphi_{2'n'} \epsilon^{n'}), \\ \bar{\epsilon}^{1'} &= \varphi_{1'2'} \epsilon^{2'} + \varphi_{1'3'} \epsilon^{3'} + \dots + \varphi_{1'n'} \epsilon^{n'}. \end{aligned}$$

Di conseguenza la forma φ^1 non coinvolge né $\varepsilon^{1'}$ né $\varepsilon^{2'}$, mentre i covettori $\varepsilon^1, \bar{\varepsilon}^1, \varepsilon^{3'}, \dots, \varepsilon^{n'}$ sono indipendenti. Se $\varphi^1 \neq 0$ si applica a questa lo stesso processo, e così via fino ad ottenere un'espressione del tipo

$$\varphi = \varepsilon^1 \wedge \bar{\varepsilon}^1 + \varepsilon^2 \wedge \bar{\varepsilon}^2 + \dots + \varepsilon^p \wedge \bar{\varepsilon}^p,$$

che è del tipo (27), posto $\bar{\varepsilon}^a = \varepsilon^{a+p}$. I covettori $(\varepsilon^1, \bar{\varepsilon}^1, \varepsilon^2, \bar{\varepsilon}^2, \dots, \varepsilon^p, \bar{\varepsilon}^p)$ formano con i rimanenti $\varepsilon^{b'}, b = 2p+1, \dots, n$, una base di E^* la cui duale è canonica per φ . ■

6. Tensori

Il concetto di tensore ha avuto origine dalla teoria dell'elasticità. Si è successivamente sviluppato nell'ambito dell'algebra lineare e della geometria differenziale, trovando notevoli applicazioni in vari rami della fisica-matematica.

DEFINIZIONE 1. Sia E uno spazio vettoriale. Dicesi **tensore di tipo (p, q)** sopra E un'applicazione multilineare del prodotto

$$(E^*)^p \times E^q = \underbrace{E^* \times E^* \times \dots \times E^*}_{p \text{ volte}} \times \underbrace{E \times E \times \dots \times E}_{q \text{ volte}}$$

in $\mathbb{R}$. Un tensore di tipo $(0, 0)$ è per definizione un numero reale. Denotiamo con $T_q^p(E)$ lo **spazio tensoriale di tipo (p, q)** su E , cioè lo spazio

$$T_q^p(E) = L(\underbrace{E^*, E^*, \dots, E^*}_{p \text{ volte}}, \underbrace{E, E, \dots, E}_{q \text{ volte}}; \mathbb{R}),$$

ponendo inoltre, per $(p, q) = (0, 0)$,

$$T_0^0(E) = \mathbb{R}.$$

Abbiamo già preso in esame alcuni tipi di spazi tensoriali. Per i primi valori della coppia di interi (p, q) abbiamo infatti:

$$(1) \quad \begin{cases} T_0^1(E) = L(E^*; \mathbb{R}) = E^{**} = E, \\ T_1^0(E) = L(E; \mathbb{R}) = E^*, \\ T_2^0(E) = L(E, E; \mathbb{R}) \quad (\text{forme bilineari}), \\ T_1^1(E) = L(E^*, E; \mathbb{R}) = \text{End}(E). \end{cases}$$

Quest'ultima uguaglianza deriva dall'esistenza di un isomorfismo canonico tra lo spazio $L(E^*, E; \mathbb{R})$ e lo spazio degli endomorfismi su E , definito dall'uguaglianza

$$(2) \quad A(\alpha, v) = \langle A(v), \alpha \rangle, \quad \forall \alpha \in E^*, v \in E.$$

A sinistra A è interpretato come elemento di $L(E^*, E; \mathbb{R})$, a destra come endomorfismo lineare.

Sia (e_i) una base di E . Sia (ε^i) la sua duale in E^* . Diciamo **componenti di un tensore $K \in T_q^p(E)$ secondo la base (e_i)** gli n^{p+q} numeri

$$(3) \quad K_{j_1 \dots j_q}^{i_1 \dots i_p} = K(\varepsilon^{i_1}, \dots, \varepsilon^{i_p}, e_{j_1}, \dots, e_{j_q})$$

Per esempio, le componenti di $K \in T_2^1(E)$ sono gli n^3 numeri

$$(4) \quad K_{jh}^i = K(\varepsilon^i, e_j, e_h).$$

Si osservi la corrispondenza nella posizione degli indici a primo e a secondo membro di queste uguaglianze. Le componenti servono, innanzitutto, a calcolare il valore del tensore sui vettori e covettori. Prendendo ad esempio il caso precedente, risulta:

$$(5) \quad K(\alpha, u, v) = K_{jh}^i \alpha_i u^j v^h.$$

Inoltre, al variare della base secondo le formule (8) e (10) del §4, le componenti di un tensore di tipo $(1, 2)$ variano secondo la legge

$$(6) \quad K_{j'h'}^{i'} = E_i^{i'} E_{j'}^j E_{h'}^h K_{jh}^i.$$

Questa formula si deduce direttamente dalla definizione (4) scritta per la base "accentata", tenendo conto quindi delle relazioni fra le basi, ed è facilmente estendibile al caso di un tensore di tipo qualunque, tenendo conto del formalismo degli indici ripetuti in alto e in basso. A questo proposito osserviamo che per passare dalla base (e_i) alla base "accentata" $(e_{i'})$ viene usata la matrice $(E_{i'}^i)$, con l'accento in basso. Gli indici in basso non accentati delle componenti vengono trasformati in indici accentati (sempre in basso) con la stessa matrice. Per questo gli indici in basso delle componenti vengono detti **indici covarianti**. Quelli in alto, invece, vengono detti **indici contravarianti**, perché nel trasformarsi fanno intervenire la matrice inversa. Pertanto, un tensore di tipo

(p, q) viene anche detto *tensore a p indici di contravarianza e q indici di covarianza*. Un tensore di tipo $(p, 0)$ viene detto **tensore contravariante di ordine p** , un tensore di tipo $(0, q)$ **tensore covariante di ordine q** . Denoteremo gli spazi tensoriali $T_0^p(E)$ e $T_q^0(E)$ semplicemente con $T^p(E)$ e $T_q(E)$.

Sui tensori possono definirsi varie operazioni che nel loro insieme costituiscono il cosiddetto **calcolo tensoriale**. Oltre alla somma di tensori dello stesso tipo si definisce il **prodotto tensoriale $K \otimes L$** (si legge **K tensore L**) fra due tensori di tipo qualsiasi K e L . Per esempio, se $K \in T_2^1(E)$ e $L \in T_1^1(E)$, $K \otimes L$ è il tensore di tipo $(2, 3)$ tale che

$$(7) \quad (K \otimes L)(\alpha, \beta, u, v, w) = K(\alpha, u, v) L(\beta, w),$$

per ogni $\alpha, \beta \in E^*$, $u, v, w \in E$. L'estensione di questa definizione al caso di tensori di altro tipo è ovvia. Si noti che $K \otimes L$ è in generale diverso da $L \otimes K$. Si ha $K \otimes L = L \otimes K$ se, per esempio, K è un tensore covariante e L un tensore contravariante.

Le componenti del prodotto tensoriale $K \otimes L$ sono il prodotto delle componenti di K e L rispettivamente. Vale cioè la formula (ci riferiamo ancora, per semplicità, al caso particolare considerato):

$$(8) \quad (K \otimes L)_{hkl}^{ij} = K_{hk}^i L_l^j.$$

Se uno dei due tensori è di tipo $(0, 0)$, cioè un numero reale, il prodotto tensoriale è l'ordinario prodotto per uno scalare:

$$(9) \quad a \otimes K = K \otimes a = aK \quad (a \in \mathbb{R}).$$

Il prodotto tensoriale è associativo:

$$K \otimes (L \otimes M) = (K \otimes L) \otimes M.$$

In molte circostanze torna utile considerare la somma diretta di tutti gli spazi tensoriali, di tutti i tipi, ordinati per esempio come segue:

$$T(E) = \mathbb{R} \oplus E \oplus E^* \oplus T^2(E) \oplus T_1^1(E) \oplus T_2(E) \oplus T^3(E) \oplus \dots$$

Un tensore K di tipo (p, q) può essere inteso come elemento di $T(E)$: lo si identifica con la successione costituita da tutti zeri e con K al posto (p, q) . Un siffatto elemento di $T(E)$ viene detto **omogeneo**. Pertanto ogni spazio tensoriale di tipo (p, q) può essere inteso come sottospazio di

$T(E)$. Su $T(E)$ si introduce l'operazione di **somma diretta $\oplus$** , definita sommando le singole componenti omogenee, nonché il prodotto tensoriale $\otimes$, per estensione lineare del prodotto di fattori omogenei. Con queste due operazioni lo spazio $T(E)$, detto **spazio tensoriale su E** , assume la struttura di algebra: è l'**algebra tensoriale su E** .

OSSERVAZIONE 1. È opportuno considerare una seconda definizione di tensore, suggerita dalle applicazioni alla fisica. Si supponga che ad un ente venga associato, per ogni scelta di una base (e_i) di uno spazio vettoriale E di dimensione n , un insieme di n^3 numeri contraddistinti da tre indici: $K_{h,i,j}$. Si supponga che la relazione tra due insiemi di questi numeri, corrispondenti a due basi distinte, risulti del tipo (6). Diciamo allora che l'ente è un tensore di tipo $(1, 2)$ e porremo gli indici in alto o in basso a secondo che essi siano contravarianti o covarianti. Più precisamente, si consideri l'insieme delle coppie

$$(e_i, K_{jh}^i)$$

dove (e_i) è una base di E e (K_{jh}^i) un insieme indicizzato di numeri reali. Diciamo *equivalenti* due coppie siffatte se, sussistendo tra le basi relazioni del tipo (8) del §4, sussistono tra i corrispondenti insiemi di numeri le relazioni (6). In tal modo si istituisce sull'insieme di queste coppie una relazione di equivalenza. Una classe di equivalenza è per definizione un tensore di tipo $(1, 2)$. In maniera analoga si dà la definizione di tensore di tipo (p, q) , con p e q qualsiasi, non entrambi nulli.

OSSERVAZIONE 2. I prodotti tensoriali

$$e_j \otimes \epsilon^i$$

formano una base di $T_1^1(E) = \text{End}(E)$, e le componenti di un endomorfismo secondo questa base sono proprio le componenti definite dalla (5) del n.4, vale a dire (si veda anche la (2)):

$$A = A_i^j e_j \otimes \epsilon^i, \quad A_i^j = A(\epsilon^j, e_i).$$

Allo stesso modo, i prodotti tensoriali

$$e_i \otimes \epsilon^j \otimes \epsilon^h$$

formano una base di $T_2^1(E)$, e le componenti di un tensore secondo questa base sono proprio le componenti definite nella (4), vale a dire:

$$K = K_{jh}^i e_i \otimes \epsilon^j \otimes \epsilon^h.$$

La dimostrazione è lasciata come esercizio. In generale dunque i prodotti tensoriali

$$e_{i_1} \otimes e_{i_2} \otimes \dots \otimes e_{i_p} \otimes \varepsilon^{j_1} \otimes \varepsilon^{j_2} \otimes \dots \otimes \varepsilon^{j_q},$$

costituiscono una base di $T_q^p(E)$ e quindi

$$\dim(T_q^p(E)) = n^{p+q}.$$

Inoltre per ogni $K \in T_q^p(E)$, definite le componenti come nella (3), si ha la rappresentazione

$$K = K_{j_1 j_2 \dots j_q}^{i_1 i_2 \dots i_p} e_{i_1} \otimes e_{i_2} \otimes \dots \otimes e_{i_p} \otimes \varepsilon^{j_1} \otimes \varepsilon^{j_2} \otimes \dots \otimes \varepsilon^{j_q}$$

OSSERVAZIONE 3. Un'altra operazione sui tensori è la **contrazione** (o **sturazione di indici**). Siano per esempio (L_{hkl}^{ij}) le componenti di un tensore di tipo (2,3). Poniamo il primo indice contravariante uguale al secondo indice covariante, eseguiamo cioè la somma rispetto a questi due indici, che si dicono *saturati*. Si ottiene un sistema a tre indici,

$$K_{hl}^j = L_{hil}^{ij},$$

per cui, come si verifica con breve calcolo, vale la legge di trasformazione (6). Questo sistema costituisce pertanto le componenti di un tensore K di tipo (1,2). Si dice che K è un tensore contratto di L . Altri esempi: la traccia A_i^i di un endomorfismo A è il tensore di tipo (0,0) contratto di A ; il vettore $u = A(v)$ può essere anche interpretato come contrazione del prodotto tensoriale $A \otimes v$, perché quest'ultimo ha componenti $A_i^j v^k$ e le componenti di u sono $u^j = A_i^j v^i$.

7. Algebra esterna

Per ogni intero $p > 1$ denotiamo con G_p il gruppo delle permutazioni di p elementi, detto **gruppo simmetrico** di ordine p . Denotiamo con ε_σ il segno di una permutazione $\sigma \in G_p$: uguale a +1 o -1 a seconda che σ sia pari o dispari, vale a dire composta da un numero pari o dispari di scambi. Ricordiamo che G_p è costituito da $p!$ elementi.

DEFINIZIONE 1. Un'applicazione p -lineare $\varphi: E^p \rightarrow \mathbb{R}$ (cioè un tensore covariante di ordine p su E) è detta **p -forma** (o anche **forma**

esterna di ordine o grado p) se è **antisimmetrica**, cioè se vale l'uguaglianza

$$(1) \quad \varphi(v_1, \dots, v_p) = \varepsilon_\sigma \varphi(\sigma(v_1, \dots, v_p)), \quad \forall v_1, \dots, v_p \in E,$$

ovvero

$$(2) \quad \varphi = \varepsilon_\sigma \varphi \circ \sigma,$$

per ogni $\sigma \in G_p$. Un covettore (caso $p = 1$) è per definizione una 1-forma, un numero reale una 0-forma.

Le p -forme costituiscono un sottospazio, che denotiamo con $\Lambda^p(E)$, dello spazio $T_p(E)$ dei tensori covarianti di ordine p . Si ha in particolare $\Lambda^1(E) = E^*$ e $\Lambda^0(E) = \mathbb{R}$.

Nel caso particolare $p = 2$ si ritrova la definizione data al §5 di forma bilineare antisimmetrica.

Le proprietà (a) e (b) seguenti sono caratteristiche delle p -forme e ne costituiscono quindi due definizioni alternative:

PROPOSIZIONE 1. Un'applicazione multilineare $\varphi: E^p \rightarrow \mathbb{R}$ è antisimmetrica se e solo se (a) cambia di segno quando si scambiano fra loro due argomenti, oppure se e solo se (b) si annulla quando due suoi argomenti coincidono.

Dimostrazione. L'equivalenza tra la (1) e la proprietà (a) segue direttamente dal fatto che ogni permutazione su p elementi è il prodotto di scambi, e che uno scambio è una permutazione dispari. Dalla (a) segue ovviamente la (b). Viceversa, per dimostrare che dalla (b) segue la (a) basta osservare che, stante la (b), si ha successivamente:

$$\begin{aligned} 0 &= \varphi(v_1 + v_2, v_1 + v_2, \dots, v_p) \\ &= \varphi(v_1, v_2, \dots, v_p) + \varphi(v_2, v_1, \dots, v_p), \end{aligned}$$

vale a dire

$$\varphi(v_1, v_2, \dots, v_p) = -\varphi(v_2, v_1, \dots, v_p).$$

Di qui l'antisimmetria sui due primi argomenti. Analogamente si dimostra l'antisimmetria sulle altre coppie di argomenti. ■

Di conseguenza:

PROPOSIZIONE 2. Una p -forma si annulla quando gli argomenti sono linearmente dipendenti.

PROPOSIZIONE 3. Ogni p -forma con $p > n = \dim(E)$ è nulla (per tanto ogni spazio $\Lambda^p(E)$ con $p > n$ si riduce ad un solo elemento, lo zero).

ESEMPIO 1. Sia E lo spazio vettoriale euclideo tridimensionale. Ponendo

$$\eta(u, v, w) = u \cdot v \times w \quad (\text{prodotto misto}).$$

si definisce un tensore covariante antisimmetrico η di ordine 3 detto **forma volume** (si veda il §11).

ESEMPIO 2. Sia $E = \mathbb{R}^n$: un vettore è una n -pla di numeri reali $v = (v^1, v^2, \dots, v^n)$. Una n -pla di vettori $(v_i) = (v_i^1, v_i^2, \dots, v_i^n)$ dà luogo ad una matrice (v_i^j) quadrata di ordine n . Ponendo

$$\varphi(v_1, v_2, \dots, v_n) = \det(v_i^j)$$

si definisce un tensore covariante antisimmetrico di ordine n sopra $\mathbb{R}^n$ (si veda l'Oss. 5 più avanti).

Per ogni $p > 1$ si consideri l'applicazione lineare

$$^A: T_p(E) \rightarrow \Lambda^p(E): \kappa \mapsto \kappa^A$$

definita da

$$(3) \quad \kappa^A = \frac{1}{p!} \sum_{\sigma \in G_p} \varepsilon_\sigma \kappa \circ \sigma$$

cioè da

$$(3') \quad \kappa^A(v_1, \dots, v_p) = \frac{1}{p!} \sum_{\sigma \in G_p} \varepsilon_\sigma \kappa(\sigma(v_1, \dots, v_p))$$

per ogni $v_1, \dots, v_p \in E$. Si tratta di un'applicazione, detta **operatore di antisimmetrizzazione**, che associa ad un qualunque tensore

covariante κ di ordine p una p -forma κ^A , cioè un tensore covariante antisimmetrico. Infatti, qualunque sia $\tau \in G_p$:

$$\begin{aligned} \varepsilon_\tau \kappa^A \circ \tau &= \frac{1}{p!} \sum_{\sigma \in G_p} \varepsilon_\tau \varepsilon_\sigma (\kappa \circ \sigma) \circ \tau \\ &= \frac{1}{p!} \sum_{\sigma \in G_p} \varepsilon_{\sigma \cdot \tau} \kappa \circ \sigma \circ \tau \\ &= \frac{1}{p!} \sum_{\sigma \cdot \tau \in G_p} \varepsilon_{\sigma \cdot \tau} \kappa \circ (\sigma \circ \tau) \\ &= \kappa^A \end{aligned}$$

tenuto conto che $\varepsilon_{\sigma \cdot \tau} = \varepsilon_\sigma \varepsilon_\tau$ e inoltre che, mentre σ varia in tutto G_p , anche $\sigma \circ \tau$ varia in tutto G_p .

Se $p = 1$ o $p = 0$ si pone per definizione $\kappa^A = \kappa$.

Si noti che nel caso particolare $p = 2$ si ritrova la definizione di parte antisimmetrica di una forma bilineare data al §5 e che per una p -forma φ , stante la (2), dalla (3) segue $\varphi^A = \varphi$.

L'operatore di antisimmetrizzazione consente di definire la seguente notevole operazione sulle forme esterne

DEFINIZIONE 2. Dicesi **prodotto esterno** di una p -forma φ per una q -forma ψ la $(p+q)$ -forma

$$(4) \quad \begin{aligned} \varphi \wedge \psi &= \frac{(p+q)!}{p! q!} (\varphi \otimes \psi)^A \\ &= \frac{1}{p! q!} \sum_{\sigma \in G_{p+q}} \varepsilon_\sigma (\varphi \otimes \psi) \circ \sigma \end{aligned}$$

ESEMPIO 3. Nel caso particolare di due covettori φ e ψ ($p = q = 1$) si ricade nella definizione di prodotto esterno data al §5:

$$\varphi \wedge \psi = \varphi \otimes \psi - \psi \otimes \varphi.$$

ESEMPIO 4. Nel caso in cui φ è un covettore e ψ una 2-forma si ha

$$(5) \quad \begin{aligned} (\varphi \wedge \psi)(u_1, u_2, u_3) &= \varphi(u_1) \psi(u_2, u_3) + \varphi(u_2) \psi(u_3, u_1) + \\ &+ \varphi(u_3) \psi(u_1, u_2). \end{aligned}$$

Si noti che a secondo membro compaiono le permutazioni cicliche dei tre vettori (u_1, u_2, u_3) . Infatti, per la definizione (4), si ha successivamente:

$$\begin{aligned} (\varphi \wedge \psi)(u_1, u_2, u_3) &= \frac{1}{1!2!} (\varphi(u_1) \psi(u_2, u_3) + \varphi(u_2) \psi(u_3, u_1) + \\ &\quad + \varphi(u_3) \psi(u_1, u_2) - \varphi(u_2) \psi(u_1, u_3) - \\ &\quad - \varphi(u_1) \psi(u_3, u_2) - \varphi(u_3) \psi(u_2, u_1)). \end{aligned}$$

Tenuto conto dell'antisimmetria di ψ , si trova la (5).

Il prodotto esterno gode delle seguenti proprietà:

$$(6) \quad \begin{aligned} \varphi \wedge \psi &= (-1)^{pq} \psi \wedge \varphi, \\ (a\varphi + b\pi) \wedge \psi &= a\varphi \wedge \psi + b\pi \wedge \psi, \\ (\varphi \wedge \psi) \wedge \pi &= \varphi \wedge (\psi \wedge \pi). \end{aligned}$$

La prima è una **proprietà commutativa graduata**. La seconda mostra che il prodotto esterno è bi-lineare (è lineare sul primo fattore, e quindi anche sul secondo, tenuto conto che vale la proprietà commutativa graduata). La terza mostra che il prodotto esterno è associativo. La dimostrazione è in appendice al paragrafo.

Insieme alla proprietà associativa (6)₃ si dimostra anche che

$$(7) \quad \begin{aligned} \varphi \wedge \psi \wedge \pi &= \frac{(p+q+r)!}{p!q!r!} (\varphi \otimes \psi \otimes \pi)^A \\ &= \frac{1}{p!q!r!} \sum_{\sigma \in G_{p+q+r}} \varepsilon_\sigma (\varphi \otimes \psi \otimes \pi) \circ \sigma \end{aligned}$$

essendo (p, q, r) i gradi della forme esterne (φ, ψ, π) rispettivamente. Questa formula si estende al caso del prodotto esterno di un numero qualsiasi di forme esterne. In particolare, per il prodotto esterno di due o più covettori $(\alpha^1, \alpha^2, \dots, \alpha^p)$ risulta semplicemente

$$(8) \quad \alpha^1 \wedge \alpha^2 \wedge \dots \wedge \alpha^p = p! (\alpha^1 \otimes \alpha^2 \otimes \dots \otimes \alpha^p)^A$$

vale a dire

$$(8') \quad \begin{aligned} (\alpha^1 \wedge \alpha^2 \wedge \dots \wedge \alpha^p)(v_1, v_2, \dots, v_p) &= \\ &= \sum_{\sigma \in G_p} \varepsilon_\sigma \langle v_{\sigma(1)}, \alpha^1 \rangle \langle v_{\sigma(2)}, \alpha^2 \rangle \dots \langle v_{\sigma(p)}, \alpha^p \rangle \end{aligned}$$

dove si è posto $(\sigma(1), \sigma(2), \dots, \sigma(p)) = \sigma(1, 2, \dots, p)$.

PROPOSIZIONE 4. Sia (e_i) una base di E . Sia (ε^i) la base duale in E^* . Una p -forma φ ammette una rappresentazione del tipo

$$(9) \quad \varphi = \frac{1}{p!} \varphi_{i_1 i_2 \dots i_p} \varepsilon^{i_1} \wedge \varepsilon^{i_2} \wedge \dots \wedge \varepsilon^{i_p}$$

con

$$(10) \quad \varphi_{i_1 i_2 \dots i_p} = \varphi(e_{i_1}, e_{i_2}, \dots, e_{i_p})$$

Dimostrazione. Come per ogni tensore (covariante) vale la rappresentazione

$$\varphi = \varphi_{i_1 i_2 \dots i_p} \varepsilon^{i_1} \otimes \varepsilon^{i_2} \otimes \dots \otimes \varepsilon^{i_p}$$

dove le componenti $\varphi_{i_1 i_2 \dots i_p}$ sono definite dalla (10). Applicando l'operatore di antisimmetrizzazione a entrambi i membri di questa uguaglianza, tenuto conto che $\varphi^A = \varphi$ e che per la (8)

$$(\varepsilon^{i_1} \otimes \varepsilon^{i_2} \otimes \dots \otimes \varepsilon^{i_p})^A = \frac{1}{p!} \varepsilon^{i_1} \wedge \varepsilon^{i_2} \wedge \dots \wedge \varepsilon^{i_p},$$

si trova la (9). ■

OSSERVAZIONE 1. La (9) mostra che i prodotti esterni $\varepsilon^{i_1} \wedge \varepsilon^{i_2} \wedge \dots \wedge \varepsilon^{i_p}$ generano tutto lo spazio $\Lambda^p(E)$ delle p -forme. Essi tuttavia non sono indipendenti, per la proprietà anticommutativa del prodotto esterno di covettori. Sono tuttavia indipendenti i prodotti esterni con indici tutti distinti, per esempio ordinati in successioni strettamente crescenti:

$$(11) \quad \varepsilon^{i_1} \wedge \varepsilon^{i_2} \wedge \dots \wedge \varepsilon^{i_p}, \quad i_1 < i_2 < \dots < i_p.$$

Questi costituiscono una base. Pertanto

$$(12) \quad \dim(\Lambda^p(E)) = \binom{n}{p} = \frac{n!}{p!(n-p)!}.$$

In particolare la dimensione dello spazio $\dim(\Lambda^n(E)) = 1$. Segue allora che ogni p -forma ammette una rappresentazione del tipo

$$(13) \quad \varphi = \sum_{i_1 < i_2 < \dots < i_p} \varphi_{i_1 i_2 \dots i_p} \varepsilon^{i_1} \wedge \varepsilon^{i_2} \wedge \dots \wedge \varepsilon^{i_p}$$

I numeri $\varphi_{i_1 i_2 \dots i_p}$ definiti dalla (10) prendono il nome di **componenti** della p -forma φ . Queste componenti sono antisimmetriche rispetto agli indici, cambiano cioè di segno quando si scambiano fra loro due indici (oppure si annullano quando due indici sono uguali). Le **componenti essenziali** sono dunque tante quante sono le successioni strettamente crescenti degli indici.

OSSERVAZIONE 2. Per poter interpretare il prodotto esterno come operazione binaria interna occorre estenderlo per linearità alla somma diretta di tutti gli spazi delle p -forme:

$$\begin{aligned}\Lambda(E) &= \bigoplus_{p=0}^{\infty} \Lambda^p(E) \\ &= \mathbb{R} \oplus E^* \oplus \Lambda^2(E) \oplus \dots \oplus \Lambda^n(E) \oplus 0 \oplus 0 \oplus \dots\end{aligned}$$

Si ottiene in tal modo un'algebra associativa, detta **algebra esterna** dello spazio E . Risulta

$$(14) \quad \dim(\Lambda(E)) = 2^n.$$

OSSERVAZIONE 3. Se $\kappa_{i_1 \dots i_p}$ sono le componenti di un tensore covariante $\kappa \in T_p(E)$ allora le componenti dell'antisimmetrizzato κ^A vengono denotate con $\kappa_{[i_1 \dots i_p]}$. Poiché le componenti sono definite da $\kappa_{i_1 \dots i_p} = \kappa(e_{i_1}, \dots, e_{i_p})$, dalla definizione (3') segue

$$(15) \quad \kappa_{[i_1 \dots i_p]} = \frac{1}{p!} \sum_{\sigma \in G_p} \varepsilon_{\sigma} \kappa_{\sigma(i_1, \dots, i_p)},$$

dove $\sigma(i_1, \dots, i_p)$ è la permutazione σ applicata alla successione dei p indici $(i_1, \dots, i_p)$. Nei casi $p = 2$ e $p = 3$ si ha per esempio:

$$\kappa_{[ij]} = \frac{1}{2}(\kappa_{ij} - \kappa_{ji}),$$

$$\kappa_{[hij]} = \frac{1}{6}(\kappa_{hij} + \kappa_{ijh} + \kappa_{jhi} - \kappa_{ihj} - \kappa_{hji} - \kappa_{jih}).$$

Considerazioni analoghe si ripetono per le componenti di un prodotto esterno. Per esempio, nel caso del prodotto di una 1-forma con una 2-forma (si veda l'Esempio 4, formula (5)):

$$(16) \quad (\varphi \wedge \psi)_{hij} = \varphi_h \psi_{ij} + \varphi_i \psi_{jh} + \varphi_j \psi_{hi}.$$

OSSERVAZIONE 4. Tutte le definizioni sopra date si trasportano immediatamente, per dualità, al caso di tensori completamente *contravarianti* e antisimmetrici. Per esempio, la definizione (8) di prodotto esterno di covettori si trasforma immediatamente, per dualità, nel prodotto esterno di p vettori. Pertanto con il simbolo $v_1 \wedge v_2 \wedge \dots \wedge v_p$ s'intende l'elemento di $T^p(E)$ tale che

$$\begin{aligned}(v_1 \wedge v_2 \wedge \dots \wedge v_p)(\alpha^1, \alpha^2, \dots, \alpha^p) &= \\ &= \sum_{\sigma \in G_p} \varepsilon_{\sigma} \langle v_{\sigma(1)}, \alpha^1 \rangle \langle v_{\sigma(2)}, \alpha^2 \rangle \dots \langle v_{\sigma(p)}, \alpha^p \rangle\end{aligned}$$

OSSERVAZIONE 5. Il determinante di un endomorfismo A di E può essere definito dall'uguaglianza

$$(17) \quad A(v_1) \wedge A(v_2) \wedge \dots \wedge A(v_n) = \det(A) v_1 \wedge v_2 \wedge \dots \wedge v_n.$$

Per riconoscerlo si consideri una base (e_i) di E . Per la proprietà anticommutativa del prodotto esterno di due vettori si ha

$$e_{i_1} \wedge e_{i_2} \wedge \dots \wedge e_{i_n} = \varepsilon_{i_1 i_2 \dots i_n} e_1 \wedge e_2 \wedge \dots \wedge e_n$$

dove il simbolo $\varepsilon_{i_1 i_2 \dots i_n}$, detto **simbolo o indicatore di Levi-Civita** (Tullio Levi-Civita, 1873-1941), vale +1 oppure -1 a seconda che la successione degli indici $(i_1, i_2, \dots, i_n)$ sia una permutazione pari o dispari della successione fondamentale $(1, 2, \dots, n)$, e vale zero in tutti gli altri casi (cioè quando almeno due indici coincidono). Il simbolo di Levi-Civita è antisimmetrico. Si ha allora:

$$\begin{aligned}A(e_1) \wedge A(e_2) \wedge \dots \wedge A(e_n) &= A_1^{i_1} A_2^{i_2} \dots A_n^{i_n} e_{i_1} \wedge e_{i_2} \wedge \dots \wedge e_{i_n} \\ &= A_1^{i_1} A_2^{i_2} \dots A_n^{i_n} \varepsilon_{i_1 i_2 \dots i_n} e_1 \wedge e_2 \wedge \dots \wedge e_n\end{aligned}$$

D'altra parte, utilizzando il simbolo di Levi-Civita, l'ordinaria definizione di determinante di una matrice può porsi nella forma sintetica

$$(18) \quad \boxed{\det(A_j^i) = \varepsilon_{i_1 i_2 \dots i_n} A_1^{i_1} A_2^{i_2} \dots A_n^{i_n}}$$

Si conclude quindi che la (17) vale per un qualunque sistema di vettori indipendenti. Essa vale ovviamente anche quando i vettori sono dipendenti, nel qual caso sono nulli ambo i membri.

Dalle (17) e (18) si possono dedurre agevolmente le proprietà fondamentali dei determinanti. Per esempio, dalla (17) si vede immediatamente che $\det(\mathbf{A} \mathbf{B}) = \det(\mathbf{A}) \det(\mathbf{B})$:

$$\begin{aligned} \det(\mathbf{A} \mathbf{B}) \mathbf{v}_1 \wedge \dots \wedge \mathbf{v}_n &= (\mathbf{A} \mathbf{B})(\mathbf{v}_1) \wedge \dots \wedge (\mathbf{A} \mathbf{B})(\mathbf{v}_n) \\ &= \mathbf{A}(\mathbf{B}(\mathbf{v}_1)) \wedge \dots \wedge \mathbf{A}(\mathbf{B}(\mathbf{v}_n)) \\ &= \det(\mathbf{A}) \mathbf{B}(\mathbf{v}_1) \wedge \dots \wedge \mathbf{B}(\mathbf{v}_n) \\ &= \det(\mathbf{A}) \det(\mathbf{B}) \mathbf{v}_1 \wedge \dots \wedge \mathbf{v}_n. \end{aligned}$$

Dalla (18) segue la regola dello sviluppo del determinante rispetto ad una linea. Posto infatti

$$(19) \quad \tilde{A}_{i_1}^1 = \varepsilon_{i_1 i_2 \dots i_n} A_2^{i_2} \dots A_n^{i_n},$$

la (18) diventa

$$(20) \quad \det(A_j^i) = A_1^{i_1} \tilde{A}_{i_1}^1.$$

Si riconosce allora nella (19), stante la definizione dell'indicatore di Levi-Civita e la sua antisimmetria, l'espressione del **complemento algebrico** del termine $A_1^{i_1}$ della prima riga della matrice (A_j^i) (*conveniamo che gli indici in basso siano indici di riga*) cioè del determinante della matrice quadrata di ordine $n-1$ ottenuta sopprimendo la riga 1 e la colonna i_1 , preceduto dal segno $(-1)^{1+i_1}$. La (20) rappresenta allora lo sviluppo del determinante rispetto alla prima riga. Analogamente si procede per lo sviluppo rispetto ad un'altra riga, o ad una colonna. Dalla (19), stante l'antisimmetria dell'indicatore di Levi-Civita, segue inoltre

$$A_2^{i_1} \tilde{A}_{i_1}^1 = \varepsilon_{i_1 i_2 \dots i_n} A_2^{i_1} A_2^{i_2} \dots A_n^{i_n} = 0.$$

Quindi: la somma dei prodotti degli elementi di una riga per i complementi algebrici di una riga diversa è nulla. Dalla (19) si osserva ancora che

$$\tilde{A}_{i_1}^1 = \frac{\partial}{\partial A_1^{i_1}} \det(A_j^i).$$

Da queste considerazioni, particolarizzate al caso di una riga (la prima), si trae la formula generale, valida per ogni endomorfismo lineare $\mathbf{A}$,

$$(21) \quad \boxed{\tilde{A}_i^j A_k^i = \det(\mathbf{A}) \delta_k^j}$$

dove

$$(22) \quad \boxed{\tilde{A}_i^j = \frac{\partial}{\partial A_j^i} \det(\mathbf{A})}$$

è il complemento algebrico dell'elemento A_j^i , vale a dire il determinante della matrice quadrata di ordine $n-1$ ottenuta cancellando la j -esima riga e la i -esima colonna della matrice delle componenti di $\mathbf{A}$, preceduta dal segno $(-1)^{j+i}$.

La matrice $(\tilde{A}_i^j)$ dei complementi algebrici definisce un endomorfismo $\tilde{\mathbf{A}}$ che chiamiamo **endomorfismo complementare** di $\mathbf{A}$. Cambiando base si ha infatti per la (22):

$$\begin{aligned} \tilde{A}_i^j &= \frac{\partial}{\partial A_j^i} \det(\mathbf{A}) \\ &= \frac{\partial}{\partial A_{j'}^{i'}} \det(\mathbf{A}) \frac{\partial A_{j'}^{i'}}{\partial A_j^i} \\ &= \tilde{A}_{i'}^{j'} \frac{\partial}{\partial A_j^i} (E_{j'}^h E_k^{i'} A_h^k) \\ &= \tilde{A}_{i'}^{j'} E_{j'}^j E_k^i. \end{aligned}$$

Dunque le $(\tilde{A}_i^j)$ cambiano secondo la legge di cambiamento delle componenti di un endomorfismo (formula (11), §4). La (21) rappresenta allora in componenti l'uguaglianza

$$(21') \quad \boxed{\tilde{\mathbf{A}} \mathbf{A} = (\det(\mathbf{A})) \mathbf{1}}$$

Si vede di qui che nel caso di un automorfismo, $\det(\mathbf{A}) \neq 0$,

$$(23) \quad \boxed{\mathbf{A}^{-1} = \frac{1}{\det(\mathbf{A})} \tilde{\mathbf{A}}}$$

Dimostrazione delle proprietà (6). La seguente notevole proprietà dell'operatore di antisimmetrizzazione implica la proprietà associativa del prodotto esterno (φ e κ sono tensori covarianti qualsiasi):

$$(24) \quad (\varphi^A \otimes \kappa)^A = (\varphi \otimes \kappa)^A = (\varphi \otimes \kappa^A)^A.$$

Infatti, tenuto conto della definizione (4) di prodotto esterno e della seconda uguaglianza (24), si ha successivamente:

$$\begin{aligned} (\varphi \wedge (\psi \wedge \pi)) &= \frac{(q+r)!}{q!r!} (\varphi \wedge (\psi \otimes \pi)^A) \\ &= \frac{(q+r)!}{q!r!} \left(\frac{(p+q+r)!}{p!(q+r)!} \varphi \otimes (\psi \otimes \pi)^A \right)^A \\ &= \frac{(p+q+r)!}{p!q!r!} (\varphi \otimes \psi \otimes \pi)^A \end{aligned}$$

Questo dimostra sia la proprietà associativa sia la formula (7). Non resta che dimostrare la prima delle uguaglianze (24) (la seconda si dimostra con procedimento analogo):

$$\begin{aligned} (\varphi^A \otimes \kappa)^A &= \left(\frac{1}{p!} \sum_{\sigma \in G_p} \varepsilon_\sigma \varphi \circ \sigma \otimes \kappa \right)^A \\ &= \frac{1}{p!} \sum_{\sigma \in G_p} \varepsilon_\sigma \left(\varphi \circ \sigma \otimes \kappa \right)^A \\ &= \left(\frac{1}{p!} \sum_{\sigma \in G_p} \varepsilon_\sigma^2 \right) (\varphi \otimes \kappa)^A \\ &= (\varphi \otimes \kappa)^A \end{aligned}$$

Per quel che riguarda infine la dimostrazione della proprietà commutativa graduata (6)₁, è sufficiente osservare che

$$\varphi \otimes \psi = \psi \otimes \varphi \circ \tau$$

dove τ è la permutazione che, operando su $p+q$ elementi, sposta gli ultimi q elementi ai primi q posti: $\tau(1, \dots, p, p+1, \dots, p+q) = (p+1, \dots, p+q, 1, \dots, p)$. Siccome τ è composta da pq scambi, $\varepsilon_\tau = (-1)^{pq}$. Allora, tenuto conto della definizione (4):

$$\begin{aligned} \varphi \wedge \psi &= (\varphi \otimes \psi)^A \\ &= (\psi \otimes \varphi \circ \tau)^A \\ &= \varepsilon_\tau (\psi \otimes \varphi)^A \\ &= (-1)^{pq} \psi \wedge \varphi. \end{aligned}$$

8. Spazi vettoriali euclidei

Uno **spazio vettoriale euclideo** è una coppia (E, g) costituita da uno spazio vettoriale E e da un **tensore metrico** g su E . Un tensore metrico (detto anche **metrica**) è per definizione una forma bilineare g su E *simmetrica e regolare*, cioè tale da soddisfare alle condizioni:

$$(1) \quad \begin{cases} g(u, v) = g(v, u), \\ g(u, v) = 0 \quad \forall v \in E \Rightarrow u = 0. \end{cases}$$

La forma bilineare simmetrica g è anche detta **prodotto scalare**. Per il prodotto scalare si usa comunemente la notazione $u \cdot v$, poniamo cioè ⁽¹⁾

$$(2) \quad u \cdot v = g(u, v).$$

Uno spazio vettoriale euclideo è quindi uno spazio vettoriale dotato di un prodotto scalare.

Chiamiamo **norma** e la denotiamo con $\|\cdot\|: E \rightarrow \mathbb{R}$ la forma quadratica corrispondente al tensore metrico

$$(3) \quad \|v\| = v \cdot v = g(v, v).$$

Chiamiamo invece **modulo** di un vettore il numero positivo ⁽²⁾

$$(4) \quad |v| = \sqrt{v \cdot v}.$$

Due vettori sono detti **ortogonali** se $u \cdot v = 0$. Diciamo che un vettore u è **unitario** se $\|u\| = \pm 1$. Diciamo che u è il **versore** del vettore v se u è unitario e se $v = k u$ con $k > 0$.

Seguendo una terminologia tipica della *teoria della Relatività Ristretta* si dice che un vettore v è del **genere spazio, luce, tempo** se la sua norma è rispettivamente positiva, nulla, negativa:

$$v \in E \quad \begin{cases} \text{genere spazio} & \|v\| > 0, \\ \text{genere luce} & \|v\| = 0, \\ \text{genere tempo} & \|v\| < 0. \end{cases}$$

⁽¹⁾ In molti testi il prodotto scalare è denotato con $u \times v$, simbolo che noi invece riserveremo al prodotto vettoriale nello spazio euclideo tridimensionale.

⁽²⁾ La norma $\|\cdot\|$ qui definita non è una *norma* nel senso usuale dell'Analisi. Lo è invece il modulo $|\cdot|$, che però viene in genere considerato solo nel caso in cui la forma quadratica g è definita positiva.

Un vettore del genere luce è anche detto **isotropo**. Un vettore isotropo è ortogonale a se stesso.

Una **base canonica** di uno spazio euclideo (E, g) è una base canonica rispetto alla forma bilineare simmetrica g , quindi una base costituita da vettori unitari fra loro ortogonali. Sappiamo dalla teoria delle forme quadratiche che due basi canoniche sono sempre formate da un egual numero p di vettori del genere spazio e quindi da un egual numero q di vettori del genere tempo. Siccome il tensore metrico è non degenere si ha sempre $p+q = n = \dim(E)$. La coppia di interi (p, q) è detta **segnatura** dello spazio euclideo. Si usa anche dire che lo spazio ha segnatura

$$\underbrace{+\dots+}_{p \text{ volte}} \underbrace{-\dots-}_{q \text{ volte}}.$$

Molti autori riservano il nome di *spazio euclideo* al caso in cui la metrica è definita positiva, cioè di segnatura $(n, 0)$. Per mettere in evidenza questa condizione noi useremo il termine **spazio strettamente euclideo**, riservando il termine di **spazio pseudo-euclideo** o **semi-euclideo** al caso di una metrica indefinita.

OSSERVAZIONE 1. In uno spazio strettamente euclideo la norma è sempre positiva (salvo che per il vettore nullo) e vale la **disuguaglianza di Schwarz** (§5)

$$(5) \quad (u \cdot v)^2 \leq \|u\| \|v\|,$$

dove l'uguaglianza sussiste se e solo se i vettori sono dipendenti. Utilizzando il modulo (4), che ora è semplicemente definito da

$$|v| = \sqrt{v \cdot v},$$

la disuguaglianza di Schwarz diventa

$$(6) \quad |u \cdot v| \leq |u| |v|.$$

Se nessuno dei due vettori è nullo si ha

$$\frac{|u \cdot v|}{|u| |v|} \leq 1.$$

Diciamo allora **angolo** dei due vettori il numero ϑ compreso tra 0 e π tale che

$$(7) \quad \cos \vartheta = \frac{u \cdot v}{|u| |v|}.$$

Vale infine la **disuguaglianza triangolare**

$$(8) \quad |u + v| \leq |u| + |v|,$$

conseguenza della disuguaglianza di Schwarz:

$$\begin{aligned} |u + v|^2 &= \|u + v\|^2 = \|u\|^2 + \|v\|^2 + 2u \cdot v \\ &\leq |u|^2 + |v|^2 + 2|u \cdot v| \\ &\leq |u|^2 + |v|^2 + 2|u| |v| \\ &= (|u| + |v|)^2. \end{aligned}$$

Considerata una base (e_i) dello spazio euclideo (E, g) , le componenti del tensore metrico

$$(9) \quad g_{ij} = g(e_i, e_j) = e_i \cdot e_j$$

formano una matrice, detta **matrice metrica**, simmetrica e regolare:

$$(10) \quad g_{ij} = g_{ji}, \quad \det(g_{ij}) \neq 0.$$

Queste condizioni sono in effetti equivalenti alle (1). Accanto alla matrice metrica (g_{ij}) torna utile considerare la matrice inversa, che denotiamo con (g^{ij}) . Essa è definita dall'uguaglianza

$$(11) \quad g_{ij} g^{jk} = \delta_i^k.$$

In una base canonica, detta anche **ortonormale**, la matrice metrica è (si veda la (12) del §5, tenendo conto che nel caso presente il rango della forma bilineare è massimo):

$$(12) \quad (g_{ij}) = \begin{pmatrix} 1_p & 0 \\ 0 & -1_q \end{pmatrix}$$

Questa matrice coincide con la sua inversa.

Mentre in generale si ha

$$(13) \quad u \cdot v = g_{ij} u^i v^j,$$

in una base canonica risulta (si veda la (13) del §5):

$$(14) \quad \mathbf{u} \cdot \mathbf{v} = \sum_{a=1}^p u^a v^a - \sum_{b=p+1}^n u^b v^b,$$

quindi

$$(15) \quad \|\mathbf{u}\|^2 = \sum_{a=1}^p (u^a)^2 - \sum_{b=p+1}^n (u^b)^2.$$

La presenza del tensore metrico consente di definire su di uno spazio euclideo, oltre al prodotto scalare, alcune altre fondamentali operazioni e precisamente:

- (I) un isomorfismo canonico tra E e il suo duale E^* ;
- (II) un **operatore ortogonale** sui sottospazi di E ;
- (III) un **operatore di trasposizione** sugli endomorfismi di E ;
- (IV) un isomorfismo canonico tra lo spazio degli endomorfismi su E e lo spazio delle forme bilineari su E .

Tratteremo i primi due argomenti in questo paragrafo, gli altri due nel prossimo.

L'applicazione lineare $\flat: E \rightarrow E^*: \mathbf{v} \mapsto \mathbf{v}^\flat$ definita dall'uguaglianza

$$(16) \quad \langle \mathbf{u}, \mathbf{v}^\flat \rangle = \mathbf{u} \cdot \mathbf{v}$$

è un isomorfismo, stante la regolarità del tensore metrico $\mathbf{g}$. Questo isomorfismo fa sì che uno spazio vettoriale euclideo possa essere identificato con il suo duale. Denotiamo con $\sharp: E^* \rightarrow E: \boldsymbol{\alpha} \mapsto \boldsymbol{\alpha}^\sharp$ l'inversa dell'applicazione $\flat$.

Introdotta una base di E , posto $\mathbf{v} = v^i \mathbf{e}_i$, si denotano con v_i le componenti del covettore $\mathbf{v}^\flat$, si pone cioè $\mathbf{v}^\flat = v_i \boldsymbol{\varepsilon}^i$. Allora la (16), scritta in componenti, diventa

$$(16') \quad u^i v_i = g_{ij} u^i v^j.$$

Valendo questa per ogni $(u^i) \in \mathbb{R}^n$ si trae

$$(17) \quad v_i = g_{ij} v^j$$

Questa è la definizione in componenti dell'applicazione $\flat$. Si osserva quindi che le componenti del covettore $\mathbf{v}^\flat$ si ottengono per **abbassamento degli indici** delle componenti (v^i) di $\mathbf{v}$ mediante la matrice metrica. L'applicazione inversa è definita dalle relazioni inverse delle (17):

$$(18) \quad v^i = g^{ij} v_j$$

Essa consiste quindi nell'**innalzamento degli indici** tramite l'inversa della matrice metrica.

OSSERVAZIONE 2. Le (g^{ij}) prendono il nome di **componenti contravarianti del tensore metrico**, mentre alle (g_{ij}) si dà il nome di **componenti covarianti**. Similmente, alle componenti (v_i) di $\mathbf{v}^\flat$ si dà il nome di **componenti covarianti del vettore $\mathbf{v}$** , riservando alle (v^i) il nome di **componenti contravarianti**. L'operazione di innalzamento e di abbassamento degli indici si può applicare anche ai tensori. Si possono così stabilire diversi isomorfismi fra spazi tensoriali aventi la stessa somma di indici covarianti e contravarianti, per esempio tra $T_1^1(E)$ e $T_2(E)$, come sarà visto nel prossimo paragrafo.

OSSERVAZIONE 3. Si è visto (§3) che un covettore $\boldsymbol{\alpha}$ è geometricamente rappresentato da una coppia di *iperpiani paralleli* (A_0, A_1) . Nell'identificazione ora vista di E con E^* , ad $\boldsymbol{\alpha}$, cioè a questa coppia di piani, corrisponde il vettore $\mathbf{v} = \boldsymbol{\alpha}^\sharp$ che è ortogonale al sottospazio A_0 e il cui prodotto scalare con i vettori del laterale A_1 vale 1. Se quindi denotiamo con $\mathbf{u}$ il vettore ortogonale ad A_0 appartenente ad A_1 , si ha

$$\mathbf{v} = \frac{\mathbf{u}}{\|\mathbf{u}\|}.$$

OSSERVAZIONE 4. Ad ogni base $(\mathbf{e}_i)$ di E si fa corrispondere un'altra base, detta **base reciproca** o **base coniugata** (e a volte, impropriamente, anche *base duale*), denotata con $(\mathbf{e}^i)$ e definita implicitamente dalle relazioni

$$(19) \quad \mathbf{e}^i \cdot \mathbf{e}_j = \delta_j^i.$$

Si verifica che

$$(20) \quad \mathbf{e}^i = g^{ij} \mathbf{e}_j, \quad \mathbf{e}_i = g_{ij} \mathbf{e}^j,$$

e che

$$(21) \quad e^i = (\varepsilon^i)^\sharp.$$

La Fig. 8.1 mostra per esempio il caso dello spazio euclideo bidimensionale. Assunta (c_1, c_2) come base canonica, le basi (e_1, e_2) e (e^1, e^2) sono basi una reciproca dell'altra, essendo $e_1 \cdot e^1 = e_2 \cdot e^2 = 1$, $e_1 \cdot e^2 = e_2 \cdot e^1 = 0$. In uno spazio strettamente euclideo, se (e_i) è una base canonica, allora $e_i = e^i$: la base coincide con la sua reciproca. Se lo spazio è pseudoeuclideo, di segnatura (p, q) , quest'uguaglianza vale solo per gli indici $i \leq p$. Per i rimanenti si ha $e_i = -e^i$.

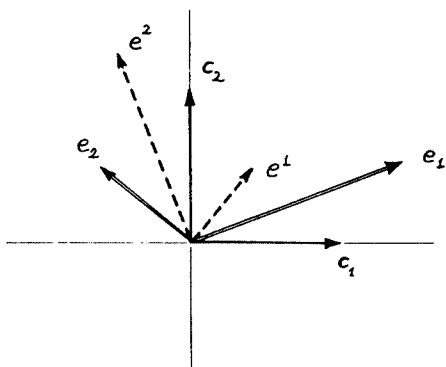


Fig. 8.1: basi reciproche.

Ad ogni sottospazio $K \subset E$ associamo il sottospazio

$$(22) \quad K^\perp = \{u \in E \mid u \cdot v = 0, \forall v \in K\},$$

detto **spazio ortogonale** a K . Confrontando questa definizione con quella del polare K° data al §3 e tenuto conto della definizione (16), si vede che

$$(23) \quad K^\perp = (K^\circ)^\sharp.$$

Pertanto dalle proprietà (4) del §3 seguono immediatamente analoghe proprietà per l'operatore ortogonale $\perp$:

$$(24) \quad \begin{cases} \dim(K) + \dim(K^\perp) = \dim(E), \\ K^{\perp\perp} = K, \\ K^\perp \subset L^\perp \iff L \subset K, \\ (K + L)^\perp = K^\perp \cap L^\perp, \\ K^\perp + L^\perp = (K \cap L)^\perp. \end{cases}$$

OSSERVAZIONE 5. Un **sottospazio** K si dice **isotropo** se $K \subset K^\perp$. Un sottospazio è isotropo se e solo se (a) tutti i suoi vettori sono isotropi, ovvero se e solo se (b) tutti i suoi vettori sono fra loro ortogonali (dimostrarlo).

OSSERVAZIONE 6. Si può dimostrare che la massima dimensione dei sottospazi isotropi è il minimo dei due numeri (p, q) della segnatura.

OSSERVAZIONE 7. Un **sottospazio** K si dice **euclideo** o **non degenere** se la restrizione del tensore metrico a K è ancora regolare (vale cioè la $(1)_2$ ristretta a K). Un sottospazio è non degenere se e solo se $K \cap K^\perp = 0$ (dimostrarlo).

OSSERVAZIONE 8. Il prodotto scalare si può estendere, in maniera naturale, ai vettori complessi. Il prodotto scalare dei vettori $v = u + iw$ e $v' = u' + iw'$ è il numero complesso $v \cdot v' = u \cdot u' - w \cdot w' + i(u \cdot w' + u' \cdot w)$. La norma di un vettore è il numero complesso $\|v\| = \|u\| - \|w\| + 2i u \cdot w$.

9. Endomorfismi in spazi euclidei

Sia (E, g) uno spazio euclideo (di segnatura qualsiasi). Ad ogni endomorfismo A su E si associa un altro endomorfismo $A^\top$ di E , detto **endomorfismo trasposto** di A , definito dall'uguaglianza

$$(1) \quad \boxed{A^\top(u) \cdot v = u \cdot A(v)}$$

Si verifica facilmente che valgono le seguenti proprietà dell'operatore di trasposizione T :

$$(2) \quad \begin{cases} 1^T = 1, \\ A^{TT} = A, \\ (A + B)^T = A^T + B^T, \\ (AB)^T = B^T A^T, \\ (A^k)^T = (A^T)^k \quad (k \in \mathbb{N}), \\ [A, B]^T = [B^T, A^T]. \end{cases}$$

Se in particolare A è un automorfismo:

$$(3) \quad (A^{-1})^T = (A^T)^{-1}.$$

Le componenti dell'endomorfismo trasposto sono date da

$$(4) \quad (A^T)_j^i = g_{jh} g^{ik} A_k^h$$

Occorre infatti ricordare la definizione di componenti di un endomorfismo e quanto detto al § precedente (in particolare la (21)), per dedurre successivamente:

$$\begin{aligned} (A^T)_j^i &= \langle A^T(e_j), \varepsilon^i \rangle \\ &= A^T(e_j) \cdot e^i \\ &= A(e^i) \cdot e_j \\ &= g^{ik} g_{jh} A(e_k) \cdot e^h \\ &= g^{ik} g_{jh} \langle A(e_k), \varepsilon^h \rangle \\ &= g^{ik} g_{jh} A_k^h. \end{aligned}$$

OSSERVAZIONE 1. Se lo spazio è strettamente euclideo (la metrica è definita positiva) e se la base scelta è canonica, allora $g_{ij} = \delta_{ij}$ e quindi

$$(5) \quad (A^T)_j^i = A_i^j.$$

Ciò significa che la matrice delle componenti dell'endomorfismo trasposto A^T coincide con la trasposta della matrice delle componenti di A (si scambiano cioè righe con colonne). Questa proprietà è valida solo in

questo caso. Se la metrica non è definita positiva e la base è canonica allora la (5) è verificata a meno del segno per certe coppie di indici (quali?).

Dalla (4) deduciamo che

$$(6) \quad \det(A^T) = \det(A), \quad \text{tr}(A^T) = \text{tr}(A).$$

Dalla prima di queste segue

$$\det(A - x1) = \det(A - x1)^T = \det(A^T - x1).$$

Dunque i polinomi caratteristici di A e A^T coincidono, e quindi coincidono anche tutti gli invarianti principali

$$(7) \quad I_k(A^T) = I_k(A),$$

nonché gli spettri.

Diciamo che un endomorfismo è **simmetrico** se $A^T = A$, che è **antisimmetrico** se $A^T = -A$. Ciò significa che si ha, rispettivamente:

$$A(u) \cdot v = A(v) \cdot u, \quad A(u) \cdot v = -A(v) \cdot u,$$

per ogni coppia di vettori (u, v) .

Ogni endomorfismo è decomponibile nella somma della sua **parte simmetrica** e della sua **parte antisimmetrica**:

$$(8) \quad \boxed{\begin{aligned} A^S &= \frac{1}{2}(A + A^T), \quad A^A = \frac{1}{2}(A - A^T) \\ A &= A^S + A^A \end{aligned}}$$

Segue che

$$(9) \quad A^T = A^S - A^A.$$

OSSERVAZIONE 2. (I) La potenza k -esima di un endomorfismo simmetrico è sempre simmetrica, qualunque sia k . Da $A^T = A$ segue infatti per la (2)₄ $(A^k)^T = (A^T)^k = A^k$. (II) La potenza k -esima di un endomorfismo antisimmetrico è simmetrica o antisimmetrica a seconda che k sia pari o dispari. Da $A^T = -A$ segue infatti $(A^k)^T = (A^T)^k =$

$(-A)^k = (-1)^k A^k$. (III) La traccia di un endomorfismo antisimmetrico è nulla. Infatti $\text{tr}(A) = \text{tr}(-A^T) = -\text{tr}(A^T) = -\text{tr}(A)$. (IV) Il determinante di un endomorfismo antisimmetrico è nullo se la dimensione n dello spazio è dispari. Infatti $\det(A) = \det(A^T) = \det(-A) = (-1)^n \det(A)$. (V) La traccia del prodotto di un endomorfismo simmetrico per uno antisimmetrico è nulla. Infatti se A è simmetrico e B antisimmetrico, $\text{tr}(AB) = \text{tr}(-AB) = -\text{tr}(AB)$. (VI) Il commutatore di due endomorfismi simmetrici, oppure di due endomorfismi antisimmetrici, è antisimmetrico. Segue dalla (2)₆ e dall'anticommutatività del commutatore.

OSSERVAZIONE 3. Da quest'ultima proprietà segue che gli endomorfismi antisimmetrici formano una sottoalgebra di Lie di $\text{End}(E)$.

OSSERVAZIONE 4. È importante sottolineare che le precedenti definizioni fanno intervenire in maniera essenziale la metrica g . Esse sono del tutto analoghe a quelle viste per le forme bilineari, dove però la nozione di trasposizione e quelle conseguenti di simmetria e di antisimmetria non richiedono la presenza sullo spazio vettoriale di alcuna struttura aggiuntiva. L'analogia risulta ulteriormente confermata nelle considerazioni seguenti.

Ad ogni endomorfismo A possiamo far corrispondere una forma bilineare, che denotiamo ancora con A , ponendo:

$$(10) \quad \boxed{A(u, v) = A(u) \cdot v}$$

Denotate con (A_{ij}) le componenti della forma bilineare A , risulta

$$(11) \quad \boxed{A_{ij} = g_{jk} A_i^k}$$

Infatti:

$$A_{ij} = A(e_i, e_j) = A(e_i) \cdot e_j = A_i^k e_k \cdot e_j = g_{jk} A_i^k.$$

Si osservi che, anche in questo caso, è presente un'operazione di abbassamento di indice. La relazione inversa delle (11) è:

$$(12) \quad A_i^k = g^{jk} A_{ij}.$$

L'operazione ora definita è un isomorfismo tra $\text{End}(E)$ e $T_2(E)$. Si osservi che accanto alla (10) sussiste una definizione alternativa analoga, che genera un altro isomorfismo: $A(u, v) = u \cdot A(v)$. Si osservi inoltre che A è un endomorfismo simmetrico (rispettivamente, antisimmetrico) se e solo se la corrispondente forma bilineare è simmetrica (risp. antisimmetrica).

OSSERVAZIONE 5. Il tensore metrico, come forma bilineare, corrisponde all'endomorfismo identico.

OSSERVAZIONE 6. Si consideri il prodotto tensoriale di due vettori $A = a \otimes b$. Si tratta di un tensore doppio contravariante. Le sue componenti in una base qualsiasi sono $A^{ij} = a^i b^j$. A questo tensore si fa corrispondere un endomorfismo, denotato ancora con A e detto **diade**, ponendo

$$(13) \quad A(u) = (a \cdot u) b,$$

e quindi ancora una forma bilineare ponendo

$$(14) \quad A(u, v) = a \cdot u b \cdot v.$$

Le componenti di quest'endomorfismo e di questa forma bilineare sono rispettivamente

$$(15) \quad A_i^j = a_i b^j, \quad A_{ij} = a_i b_j.$$

Più in generale ad ogni tensore doppio contravariante A , di componenti (A^{ij}) , per quanto visto al paragrafo precedente, corrisponde un endomorfismo ed una forma bilineare (che denotiamo sempre con A) le cui componenti si ottengono con l'abbassamento degli indici:

$$(16) \quad A_i^j = A^{kj} g_{ki}, \quad A_{ij} = A^{hk} g_{hi} g_{kj}.$$

Passiamo ora a considerare alcune fondamentali proprietà degli autovalori degli endomorfismi simmetrici e antisimmetrici, nell'ipotesi però che lo spazio E sia strettamente euclideo, cioè a metrica definita positiva (salvo per la Prop. 1, valida anche nel caso di uno spazio non strettamente euclideo). Cominciamo dagli endomorfismi simmetrici.

PROPOSIZIONE 1. Gli autovettori di un endomorfismo simmetrico corrispondenti ad autovalori distinti sono ortogonali.

Dimostrazione. Sia A un endomorfismo simmetrico. Sia $A(v_1) = \lambda_1 v_1$ e $A(v_2) = \lambda_2 v_2$. Moltiplicando scalarmente la prima uguaglianza per v_2 , la seconda per v_1 , sottraendo poi membro a membro e tenendo conto della simmetria di A si deduce $0 = (\lambda_1 - \lambda_2) v_1 \cdot v_2$. Da $\lambda_1 \neq \lambda_2$ segue $v_1 \cdot v_2 = 0$. ■

PROPOSIZIONE 2. Gli autovalori di un endomorfismo simmetrico in uno spazio strettamente euclideo sono reali.

Dimostrazione. Sia $A(v) = \lambda v$, con λ complesso e quindi v complesso: $v = u + iw$. Moltiplicando scalarmente membro a membro per $\bar{v} = u - iw$,

$$\bar{v} \cdot A(v) = \lambda (u - iw) \cdot (u + iw),$$

si trova, considerando A come forma bilineare simmetrica:

$$A(u, u) + A(w, w) = \lambda (\|u\| + \|w\|).$$

Essendo lo spazio strettamente euclideo è sicuramente $\|u\| + \|w\| > 0$ (a meno che non sia $v = 0$, caso a priori escluso). L'equazione precedente determina quindi l'autovalore λ attraverso un quoziente di numeri reali. ■

PROPOSIZIONE 3. Per un endomorfismo simmetrico in uno spazio strettamente euclideo la molteplicità geometrica degli autovalori coincide con la molteplicità algebrica.

Dimostrazione. Si consideri l'insieme degli autovalori distinti

$$(\lambda_1, \lambda_2, \dots, \lambda_l), \quad l \leq n,$$

nonché la successione delle rispettive molteplicità

$$(m_1, m_2, \dots, m_l)$$

e dei rispettivi spazi invarianti

$$(K_1, K_2, \dots, K_l)$$

(con riferimento alle notazioni del §4 abbiamo $m_i = m_{\lambda_i}$ e $K_i = K_{\lambda_i}$). Per la Prop. 1 questi sottospazi sono mutuamente ortogonali. Consideriamo il sottospazio $V \subset E$ dei vettori ortogonali a tutti i K_i . Questo è un sottospazio invariante di A ; infatti per ogni $v \in V$ e per ogni $v_i \in K_i$ si ha $A(v_i) \cdot v = 0$ perché $A(K_i) \subset K_i$. Per la simmetria di A di qui segue $A(v) \cdot v_i = 0$, il che significa appunto $A(v) \in V$. Allora la restrizione di A a V darebbe luogo ad ulteriori autovettori da aggiungersi alla lista già fatta: assurdo, a meno che non sia $V = \{0\}$. Abbiamo dunque che lo spazio E è somma diretta dei sottospazi mutuamente ortogonali associati agli autovalori:

$$(17) \quad E = K_1 \oplus K_2 \oplus \dots \oplus K_l.$$

Pertanto

$$\begin{aligned} \dim(E) = n &= \dim(K_1) + \dim(K_2) + \dots + \dim(K_l) \\ &= \text{mg}(\lambda_1) + \text{mg}(\lambda_2) + \dots + \text{mg}(\lambda_l). \end{aligned}$$

D'altra parte, per il teorema fondamentale dell'algebra:

$$\begin{aligned} n &= m_1 + m_2 + \dots + m_l \\ &= \text{ma}(\lambda_1) + \text{ma}(\lambda_2) + \dots + \text{ma}(\lambda_l). \end{aligned}$$

In generale, come si è detto al §4, vale la disuguaglianza $\text{ma}(\lambda_i) \geq \text{mg}(\lambda_i)$. Se vi fosse una sola disuguaglianza forte, stante le due uguaglianze sopra stabilite, si cadrebbe in un assurdo. ■

OSSERVAZIONE 7. Se gli autovalori di A sono tutti distinti, allora ognuno dei sottospazi invarianti K_i è unidimensionale e univocamente determinato.

PROPOSIZIONE 4. Ogni endomorfismo simmetrico in uno spazio strettamente euclideo ammette una rappresentazione del tipo

$$(18) \quad A = \sum_{\alpha=1}^n \lambda_{\alpha} c_{\alpha} \otimes c_{\alpha}$$

dove (c_{α}) è una base canonica di autovettori di A e dove (λ_{α}) sono i corrispondenti autovalori.

Per questa rappresentazione si deve tener conto dell'Oss. 6, cioè della corrispondenza tra tensori doppi contravarianti, endomorfismi e forme bilineari. Segue subito dalla scrittura (18) che i $(\mathbf{c}_\alpha)$ sono autovettori ed i (λ_α) sono autovalori. Infatti (si veda la (13)):

$$\mathbf{A}(\mathbf{c}_\beta) = \sum_{\alpha=1}^n \lambda_\alpha (\mathbf{c}_\alpha \cdot \mathbf{c}_\beta) \mathbf{c}_\alpha = \sum_{\alpha=1}^n \lambda_\alpha \delta_{\alpha\beta} \mathbf{c}_\alpha = \lambda_\beta \mathbf{c}_\beta.$$

Per ogni vettore $\mathbf{v}$ si ha quindi

$$(19) \quad \boxed{\mathbf{A}(\mathbf{v}) = \sum_{\alpha=1}^n \lambda_\alpha v_\alpha \mathbf{c}_\alpha}$$

essendo le (v_α) le sue componenti rispetto alla base canonica.

Dimostrazione. Si consideri la decomposizione (17) in sottospazi invarianti. Se $\mathbf{v}_1 \in K_1$ e $|\mathbf{v}_1| = 1$, allora

$$\mathbf{A}(\mathbf{v}_1, \mathbf{v}_1) = \mathbf{A}(\mathbf{v}_1) \cdot \mathbf{v}_1 = \lambda_1 \mathbf{v}_1 \cdot \mathbf{v}_1 = \lambda_1.$$

Analogamente, tenuto conto della Prop. 1, se $\mathbf{v}_2 \in K_2$ allora

$$\mathbf{A}(\mathbf{v}_1, \mathbf{v}_2) = \mathbf{A}(\mathbf{v}_1) \cdot \mathbf{v}_2 = \lambda_1 \mathbf{v}_1 \cdot \mathbf{v}_2 = 0.$$

Pertanto, scegliendo in ogni sottospazio $K_1, K_2, \dots, K_l$ una base canonica si determina una base canonica $(\mathbf{c}_\alpha)$ di tutto E per la quale vale la rappresentazione (18). ■

Si osservi che dalla (18) segue che nella base canonica $(\mathbf{c}_\alpha)$ le componenti $A_{\alpha\beta} = \mathbf{A}(\mathbf{c}_\alpha, \mathbf{c}_\beta)$ della forma bilineare $\mathbf{A}$ sono tutte nulle salvo quelle della diagonale principale, che coincidono con gli autovalori:

$$(20) \quad A_{\alpha\beta} = \delta_{\alpha\beta} \lambda_\alpha.$$

In forma matriciale:

$$(20') \quad (A_{\alpha\beta}) = \begin{pmatrix} \lambda_1 \mathbf{1}_{m_1} & 0 & \dots & 0 \\ 0 & \lambda_2 \mathbf{1}_{m_2} & \dots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \dots & \lambda_l \mathbf{1}_{m_l} \end{pmatrix}$$

Questa base canonica $(\mathbf{c}_\alpha)$ si trasforma in una base canonica $(\mathbf{c}'_\alpha)$ rispetto alla forma bilineare simmetrica $\mathbf{A}$ ponendo

$$\mathbf{c}'_\alpha = \frac{1}{\sqrt{|\lambda_\alpha|}} \mathbf{c}_\alpha$$

se $\lambda_\alpha \neq 0$, oppure $\mathbf{c}'_\alpha = \mathbf{c}_\alpha$ se $\lambda_\alpha = 0$. Si vede così che le componenti positive (risp. negative) della forma bilineare $\mathbf{A}$ in questa base sono tante quante gli autovalori positivi (risp. negativi) dell'endomorfismo $\mathbf{A}$. Si è pertanto dimostrato che

PROPOSIZIONE 5. Se un endomorfismo simmetrico in uno spazio strettamente euclideo ha, contati con la dovuta molteplicità, p autovalori positivi e q autovalori negativi, allora la corrispondente forma bilineare simmetrica ha segnatura (p, q) .

Tenuto conto che il tensore metrico è una forma bilineare simmetrica definita positiva, dalla dimostrazione di quest'ultima proprietà si deduce anche l'enunciato generale seguente:

PROPOSIZIONE 6. Due forme bilineari simmetriche, di cui una definita positiva, possono contemporaneamente porsi in forma diagonale, ammettono cioè basi ortogonali comuni.

Passiamo infine ad esaminare brevemente il caso antisimmetrico.

PROPOSIZIONE 7. Gli autovalori di un endomorfismo antisimmetrico in uno spazio strettamente euclideo sono nulli o immaginari puri.

Dimostrazione. La dimostrazione è analoga a quella della Prop. 2. Sia $\mathbf{A}(\mathbf{v}) = \lambda \mathbf{v}$, con λ complesso e quindi $\mathbf{v}$ complesso: $\mathbf{v} = \mathbf{u} + i\mathbf{w}$. Moltiplicando scalarmente membro a membro per $\bar{\mathbf{v}} = \mathbf{u} - i\mathbf{w}$, si trova, considerando $\mathbf{A}$ come forma bilineare antisimmetrica:

$$i 2 \mathbf{A}(\mathbf{w}, \mathbf{u}) = \lambda (|\mathbf{u}| + |\mathbf{w}|).$$

Essendo lo spazio strettamente euclideo quest'equazione determina l'autovalore λ come quoziente di un numero immaginario puro per un numero reale. ■

10. Endomorfismi ortogonali

Un **endomorfismo ortogonale** o **isometria** è un endomorfismo Q su di uno spazio euclideo o pseudoeuclideo (E, g) conservante il prodotto scalare, cioè tale che

$$(1) \quad Q(u) \cdot Q(v) = u \cdot v \quad (\forall u, v \in E)$$

Si osservi che la (1) equivale a

$$(2) \quad Q(u) \cdot Q(u) = u \cdot u \quad (\forall u \in E)$$

È infatti ovvio che la (1) implica la (2). Viceversa, se vale la (2) risulta

$$Q(u+v) \cdot Q(u+v) = (u+v) \cdot (u+v) \quad (\forall u, v \in E).$$

Sviluppando ambo i membri di quest'uguaglianza e riutilizzando la (2) si trova l'uguaglianza (1).

Dalla (1), per il fatto che la metrica non è degenera, segue che $\text{Ker}(Q) = 0$. Di conseguenza $Q(E) = E$. Ne deduciamo che un endomorfismo ortogonale è un isomorfismo.

Per definizione di endomorfismo trasposto la (1) equivale a

$$Q^T Q(u) \cdot v = u \cdot v,$$

e quindi, per l'arbitrarietà dei vettori, alla condizione

$$(3) \quad Q^T Q = 1$$

ovvero

$$(3') \quad Q^T = Q^{-1}$$

Dunque le formule (1), (2), (3) e (3') sono tutte definizioni equivalenti di endomorfismo ortogonale.

Poiché $\det(Q^T) = \det(Q)$, la (3) implica $(\det(Q))^2 = 1$ cioè

$$(4) \quad \det(Q) = \pm 1.$$

Gli endomorfismi ortogonali il cui determinante vale $+1$ sono chiamati **rotazioni**; quelle con determinante -1 **rotazioni improprie**. L'identità 1 è ovviamente una rotazione. Il prodotto di due (o più) rotazioni è una rotazione. Il prodotto di due rotazioni improprie è una rotazione. Il prodotto di una rotazione per una rotazione impropria è una rotazione impropria.

L'insieme di tutti gli endomorfismi ortogonali sopra uno spazio (E, g) è un gruppo rispetto all'ordinaria composizione degli endomorfismi, cioè un sottogruppo di $\text{Aut}(E)$, detto **gruppo ortogonale** di (E, g) . Lo denotiamo con $O(E, g)$. Le rotazioni formano un sottogruppo, detto **gruppo ortogonale speciale**, che denotiamo con $SO(E, g)$. Le rotazioni improprie non formano ovviamente sottogruppo. Se due spazi vettoriali hanno uguale dimensione e segnatura, allora i rispettivi gruppi ortogonali sono isomorfi.

Passando alle componenti, osserviamo che un endomorfismo Q è ortogonale se e solo se le sue componenti (Q_i^j) rispetto ad una base generica verificano le uguaglianze

$$(5) \quad g_{ij} Q_h^i Q_k^j = g_{hk}$$

che sono in tutto $\frac{1}{2}n(n+1)$. Queste si ottengono dalla (1) scritta per una generica coppia (e_h, e_k) di vettori della base.

Se lo spazio è strettamente euclideo e se la base scelta è canonica, allora le (5) diventano

$$(6) \quad \sum_{i=1}^n Q_h^i Q_k^i = \delta_{hk}$$

Ciò significa che il prodotto della matrice (Q_j^i) per la sua trasposta è la matrice unitaria. Le matrici soddisfacenti a questa proprietà sono dette **matrici ortogonali**.

Se la base è canonica ma lo spazio non è strettamente euclideo la formula (6) non è più valida. In ogni caso, le matrici delle componenti in una base canonica degli endomorfismi ortogonali (o delle rotazioni) di uno spazio di segnatura (p, q) formano un gruppo denotato con $O(p, q)$ (oppure $SO(p, q)$). Si denotano in particolare con $O(n)$ e $SO(n)$ i gruppi delle matrici $n \times n$ ortogonali, cioè soddisfacenti alla (6), e di quelle ortogonali a determinante unitario. La scelta di una base canonica stabilisce un isomorfismo tra $O(E, g)$ e $O(p, q)$.

OSSERVAZIONE 1. Gli autovalori, complessi o reali, di un endomorfismo ortogonale in uno spazio strettamente euclideo sono unitari: $|\lambda| = 1$. Infatti, dall'equazione $Q(v) = \lambda v$ segue, applicando la coniugazione, l'equazione $Q(\bar{v}) = \bar{\lambda}\bar{v}$. Moltiplicando scalarmente membro a membro queste due uguaglianze si ottiene $Q(v) \cdot Q(\bar{v}) = \lambda\bar{\lambda}v \cdot \bar{v}$. Poiché Q è ortogonale si ha anche $Q(v) \cdot Q(\bar{v}) = v \cdot \bar{v}$. Di qui segue $\lambda\bar{\lambda} = 1$, come asserito.

* * *

Vediamo ora alcuni esempi di rappresentazione delle rotazioni, che suggeriscono alcuni esercizi e ulteriori approfondimenti.

ESEMPIO 1. **Rappresentazione mediante simmetrie.** Sia $a \in E$ un vettore non isotropo. Si consideri l'endomorfismo S_a definito da

$$(7) \quad S_a(v) = v - 2 \frac{a \cdot v}{a \cdot a} a.$$

Esso è una **simmetria** rispetto al piano ortogonale ad a , nel senso che, come si verifica immediatamente,

$$S_a(a) = -a, \quad S_a(v) = v$$

per ogni vettore v ortogonale ad a . Di conseguenza, S_a è un'isometria. Più precisamente è una rotazione impropria. Infatti a è un autovettore di autovalore -1 mentre ogni vettore ortogonale ad a è autovettore di autovalore $+1$, sicché lo spettro è $(-1, 1, \dots, 1)$ e quindi $\det(S_a) = -1$. Si può dimostrare, ma la dimostrazione non è semplice, che *ogni isometria è il prodotto di simmetrie*. È invece più facile dimostrare (e la dimostrazione è lasciata come esercizio) che due simmetrie S_a e S_b commutano se e solo se i corrispondenti vettori a e b sono dipendenti (allora le simmetrie coincidono) oppure sono ortogonali.

ESEMPIO 2. **Decomposizione simmetrica delle isometrie.** Decomponiamo un endomorfismo Q nelle sue parti simmetrica e antisimmetrica:

$$(8) \quad Q = Q^S + Q^A, \quad Q^S = \frac{1}{2}(Q + Q^T), \quad Q^A = \frac{1}{2}(Q - Q^T).$$

Segue che

$$QQ^T = (Q^S + Q^A)(Q^S - Q^A) = (Q^S)^2 - (Q^A)^2 + Q^A Q^S - Q^S Q^A,$$

$$Q^T Q = (Q^S - Q^A)(Q^S + Q^A) = (Q^S)^2 - (Q^A)^2 - Q^A Q^S + Q^S Q^A.$$

Se Q è un'isometria, allora entrambe queste espressioni devono valere 1, sicché sottraendo membro a membro segue subito

$$(9) \quad Q^A Q^S - Q^S Q^A = [Q^A, Q^S] = 0$$

e quindi

$$(10) \quad (Q^S)^2 - (Q^A)^2 = 1.$$

Viceversa, se valgono entrambe queste condizioni allora le predette espressioni valgono 1. È così dimostrato che le condizioni (9) e (10) sono necessarie e sufficienti affinché un endomorfismo Q sia ortogonale. Dalla (10) si deduce in particolare che un endomorfismo ortogonale Q è simmetrico se e solo se esso è involutorio, cioè se $Q^2 = 1$. Lo si riconosce comunque anche dal fatto che, valendo la (3'), la condizione $QQ = 1$ è equivalente a $Q = Q^T$. Osserviamo infine che se l'endomorfismo ortogonale Q è simmetrico allora, qualunque sia il vettore v , il vettore $v \pm Q(v)$ è autovettore di Q con autovalore ± 1 . Infatti:

$$Q(v \pm Q(v)) = Q(v) \pm Q^2(v) = Q(v) \pm v = \pm(v \pm Q(v)).$$

ESEMPIO 3. **Rappresentazione esponenziale.** Dato un endomorfismo A , consideriamone tutte le sue potenze $\{A^k; k \in \mathbb{N}\}$ e quindi la serie

$$(11) \quad e^A = \sum_{k=0}^{\infty} \frac{1}{k!} A^k,$$

detta **esponenziale** di A . Questa è convergente, nel senso che qualunque sia A e comunque si fissi una base di E , le n^2 serie di numeri reali date dalle componenti di e^A sono tutte convergenti (anzi, assolutamente convergenti) o, se si vuole, che tale è la serie vettoriale

$$\sum_{k=0}^{\infty} \frac{1}{k!} A^k(v),$$

qualunque sia il vettore v . Omettiamo la discussione sulla convergenza di questa serie e su serie analoghe costruibili a partire da funzioni analitiche. Ci limitiamo qui ad osservare che se $(\lambda_1, \dots, \lambda_n)$ è lo spettro di

A allora lo spettro di e^A è $(e^{\lambda_1}, \dots, e^{\lambda_n})$. Infatti da $A(v) = \lambda v$ segue $A^k(v) = \lambda^k v$ e quindi

$$e^A(v) = \sum_{k=0}^{\infty} \frac{1}{k!} A^k(v) = \sum_{k=0}^{\infty} \frac{1}{k!} \lambda^k v = e^{\lambda} v.$$

L'esponenziale di un endomorfismo non gode di tutte le proprietà della corrispondente funzione analitica. Per esempio l'uguaglianza

$$e^{A+B} = e^A e^B$$

non è più vera in generale. Lo è se i due endomorfismi A e B commutano (col che commutano anche gli esponenziali). Vale comunque la proprietà

$$(e^A)^T = e^{A^T}.$$

Di qui segue che: se A è antisimmetrico allora e^A è una rotazione. Si ha infatti:

$$e^A (e^A)^T = e^A e^{A^T} = e^A e^{-A} = e^{(A-A)} = 1.$$

Inoltre la condizione $\det(e^A) = 1$ segue dal fatto che nello spettro di un endomorfismo antisimmetrico gli autovalori non nulli si distribuiscono in coppie di segno opposto, sicché la loro somma è uguale a zero. Allora, per quanto sopra detto,

$$\det(e^A) = \prod_{i=1}^n e^{\lambda_i} = e^{\sum_{i=1}^n \lambda_i} = e^0 = 1.$$

ESEMPIO 4. Rappresentazione di Cayley. Se A è un endomorfismo antisimmetrico tale che $A - 1$ è invertibile allora l'endomorfismo

$$(12) \quad Q = (1 + A)(1 - A)^{-1}$$

è ortogonale. Si verifica infatti facilmente, utilizzando le ben note proprietà della trasposizione e il fatto che i due endomorfismi $1 \pm A$ commutano, che $Q Q^T = 1$. Un po' meno immediata è la verifica del fatto che $\det(Q) = 1$ (essa è lasciata come esercizio). Si ottiene quindi un'ulteriore

rappresentazione delle rotazioni in termini di endomorfismi antisimmetrici. Si osservi che la condizione $\det(1 - A) \neq 0$, richiesta per la validità della scrittura (12), è sempre soddisfatta in uno spazio strettamente euclideo. Infatti, questa equivale alla condizione che 1 sia un autovalore di A , cosa impossibile perché in tali spazi gli autovalori non nulli di un endomorfismo antisimmetrico sono immaginari puri.

ESEMPIO 5. Rappresentazione quaternionale. Veniamo infine ad una interessante rappresentazione delle rotazioni nello spazio euclideo tridimensionale $\mathbb{E}_3$. Si consideri nella somma diretta di spazi vettoriali $\mathbb{R} \oplus \mathbb{E}_3$ il prodotto (interno) definito da

$$(13) \quad (a, u)(b, v) = (ab - u \cdot v, av + bu + u \times v)$$

nonché il prodotto scalare definito da

$$(14) \quad (a, u) \cdot (b, v) = ab + u \cdot v.$$

L'elemento inverso è definito da

$$(15) \quad (a, u)^{-1} = \frac{(a, -u)}{\|(a, u)\|}$$

dove si è posto

$$\|(a, u)\| = (a, u) \cdot (a, u) = a^2 + u^2.$$

Quindi tutti gli elementi sono invertibili, fuorché quello nullo $(0, 0)$. Si ottiene un'algebra associativa, non commutativa, con unità $(1, 0)$. Si pone per comodità $(a, 0) = a$ e $(0, v) = v$. Il prodotto scalare è conservato dal prodotto interno, nel senso che

$$(16) \quad \|(a, u)(b, v)\| = \|(a, u)\| \|(b, v)\|.$$

Da questa proprietà segue per esempio che gli elementi unitari, cioè quelli per cui $\|(a, u)\| = 1$, formano un sottogruppo del gruppo degli elementi non nulli.

Un modo conveniente per rappresentare quest'algebra, e quindi riconoscerne facilmente le proprietà fondamentali, consiste nella scelta di una base canonica (i, j, k) dello spazio vettoriale e nella rappresentazione di un generico elemento $(a, u = bi + cj + dk)$ nella somma formale

$$(17) \quad a + bi + cj + dk,$$

detta **quaternione**. Il prodotto di due quaternioni, in conformità alla definizione (13), è l'estensione lineare dei seguenti prodotti fondamentali:

$$(18) \quad i^2 = j^2 = k^2 = -1, \quad ij = k, \quad jk = i, \quad ki = j.$$

Definita quest'algebra, ad ogni suo elemento non nullo (a, u) si associa un'applicazione $Q: \mathbb{E}_3 \rightarrow \mathbb{E}_3$ definita da

$$(19) \quad Q(v) = (a, u)v(a, u)^{-1}.$$

Lasciamo come esercizio di dimostrare che: (i) il secondo membro definisce effettivamente un vettore; (ii) che l'applicazione così definita è un endomorfismo ortogonale, anzi una rotazione; (iii) che l'applicazione

$$f: (\mathbb{R} \times \mathbb{E}_3) \setminus \{0\} \rightarrow SO(\mathbb{E}_3)$$

definita dalla (19) è un omomorfismo di gruppi.

In particolare possiamo restringere quest'applicazione agli elementi unitari di $\mathbb{R} \times \mathbb{E}_3$. Osserviamo allora che, scegliendo una base canonica nello spazio euclideo, questi elementi sono caratterizzati dall'equazione (si veda la (17))

$$a^2 + b^2 + c^2 + d^2 = 1,$$

e quindi descrivono tutta la sfera unitaria $\mathbb{S}_3 \subset \mathbb{R}^4$. Risulta pertanto che: (i) la sfera $\mathbb{S}^3$ ha una struttura di gruppo; (ii) esiste un omomorfismo di gruppi

$$\varphi: \mathbb{S}_3 \rightarrow SO(3)$$

che costituisce un *ricoprimento universale* del gruppo delle rotazioni dello spazio euclideo tridimensionale (questo termine è proprio della teoria dei *gruppi di Lie*, argomento di corsi superiori). Si osservi che ogni elemento di $SO(3)$ ha come controimmagine due punti opposti di $\mathbb{S}_3$.

Sulla rappresentazione delle rotazioni nello spazio tridimensionale euclideo si ritornerà più avanti (§11.4).

11. Lo spazio vettoriale euclideo tridimensionale

In questo paragrafo esamineremo le proprietà fondamentali di uno spazio vettoriale sul campo reale, tridimensionale, dotato di un tensore metrico, cioè di un prodotto scalare, definito positivo. Lo denotiamo

semplicemente con $\mathbb{E}_3$. Adotteremo tutte le definizioni e notazioni già introdotte al §8. Siccome dovremo distinguere tra basi generiche e basi canoniche, una base canonica sarà denotata con (c_α) ($\alpha = 1, 2, 3$), una base generica con (e_i) ($i = 1, 2, 3$). Le componenti relative a una base canonica saranno quindi contrassegnate con indici greci, quelle relative ad una base generica con indici latini. Quando non occorrono notazioni con indici, una base canonica sarà denotata con (i, j, k) . Una base canonica è anche detta **ortonormale**.

11.1. La forma volume

Due **basi ordinate** di uno spazio vettoriale (ci riferiamo ad uno spazio vettoriale di dimensione qualsiasi n) si dicono **concordi** o **e-quiorientate** se la matrice della trasformazione da una base all'altra ha determinante positivo. Si stabilisce in questo modo una relazione di equivalenza nell'insieme delle basi ordinate. Le classi di equivalenza sono due e prendono il nome di **orientamenti**. Quando si sceglie una base e la si ordina, come si fa usualmente, corredandone gli elementi con un indice che va da 1 a n , si sceglie automaticamente un orientamento dello spazio. Se si scambiano fra di loro due vettori della base, oppure se si cambia verso ad un numero dispari di vettori della base, si cambia orientamento.

Consideriamo ora lo spazio vettoriale euclideo tridimensionale $\mathbb{E}_3$, pur avvertendo che le considerazioni svolte in questo sottoparagrafo si estendono a spazi euclidei di dimensione superiore. Fissata una base canonica $(c_\alpha) = (c_1, c_2, c_3)$, col che risulta anche fissato un orientamento, definiamo una 3-forma $\eta: \mathbb{E}_3 \times \mathbb{E}_3 \times \mathbb{E}_3 \rightarrow \mathbb{R}$ ponendo

$$(1) \quad \eta(u, v, w) = \begin{vmatrix} u^1 & u^2 & u^3 \\ v^1 & v^2 & v^3 \\ w^1 & w^2 & w^3 \end{vmatrix}$$

Che l'applicazione η così definita sia multilineare e antisimmetrica segue immediatamente dalle proprietà del determinante.

È importante la circostanza che la definizione (1) di η non dipende dalla scelta della base canonica, purché questa appartenga all'orientamento prefissato. Per riconoscerlo basta osservare che, ricordata la definizione di determinante, la definizione (1) si può porre in forma sintetica

utilizzando il **simbolo o indicatore di Levi-Civita**:

$$(2) \quad \boxed{\eta(u, v, w) = \varepsilon_{\alpha\beta\gamma} u^\alpha v^\beta w^\gamma}$$

Il simbolo $\varepsilon_{\alpha\beta\gamma}$ vale +1 se la successione di indici (α, β, γ) è una permutazione pari della disposizione fondamentale $(1, 2, 3)$, -1 se è invece una permutazione dispari, 0 in tutti gli altri casi, cioè quando almeno due degli indici sono uguali. La (2) mostra subito che le componenti del tensore η nella base canonica (c_α) coincidono proprio con il simbolo di Levi-Civita:

$$(3) \quad \eta_{\alpha\beta\gamma} = \varepsilon_{\alpha\beta\gamma}.$$

Quest'ultima proprietà è manifestamente indipendente dalla base canonica di partenza: se si ripetesse la definizione (1) scegliendo un'altra base canonica si giungerebbe al medesimo risultato (3). Pertanto la definizione di η data dalla (1) è invariante rispetto alla scelta della base canonica.

Il significato di η è notevole: il determinante nella definizione (1) rappresenta il volume del parallelepipedo individuato dai vettori (u, v, w) , più precisamente il volume con segno: positivo se questi vettori formano, nell'ordine, una base concorde con quella canonica prefissata, negativo in caso contrario; nullo se i vettori sono dipendenti.

Per riconoscerlo si può scegliere una base canonica (c_i) tale che c_1 sia parallelo ad u e il sottospazio individuato da (c_1, c_2) contenga v . Allora:

$$u = u^1 c_1, \quad v = v^1 c_1 + v^2 c_2, \quad w = w^1 c_1 + w^2 c_2 + w^3 c_3.$$

Per l'antisimmetria di η segue

$$\eta(u, v, w) = u^1 v^2 w^3.$$

Si riconosce che quest'ultimo prodotto è proprio il volume del parallelepipedo individuato dai vettori (u, v, w) (Fig. 11.1).

Per quanto ora visto il tensore antisimmetrico covariante η prende il nome di **forma volume**. La sua definizione si estende facilmente a spazi euclidei di dimensione qualsiasi. Si osservi che di forme volume ne esistono due, una opposta all'altra, una per ogni orientamento. L'orientamento che di solito si assume nello spazio euclideo tridimensionale è quello della **mano destra**: i vettori di una base canonica ordinata

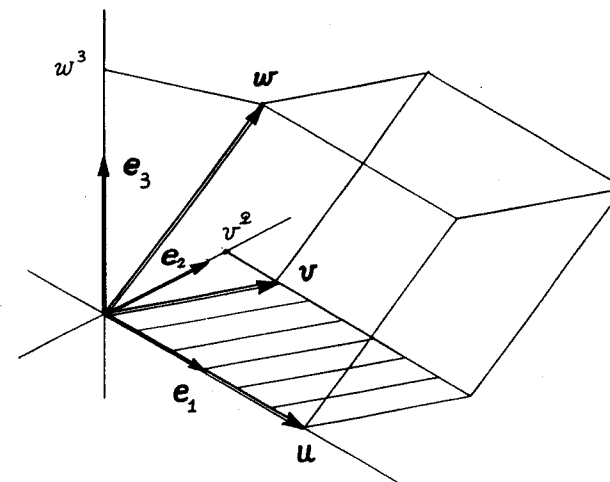


Fig. 11.1: la forma volume.

$(c_1, c_2, c_3) = (i, j, k)$ sono rispettivamente orientati come indice-medio-pollice (o pollice-indice-medio) della mano destra.

La forma volume è un ente fondamentale nello spazio euclideo. Per suo tramite, come si vedrà nei prossimi paragrafi, si possono definire due operazioni fondamentali: l'**aggiunzione** e il **prodotto vettoriale**.

OSSERVAZIONE 1. Mostriamo che le componenti di η rispetto ad una qualunque base (e_i) sono date da

$$(4) \quad \boxed{\eta_{hij} = J \varepsilon_{hij}, \quad J = \det(E_i^\alpha)}$$

dove ε_{hij} è ancora il simbolo di Levi-Civita, corredato però di indici latini, oppure da

$$(5) \quad \boxed{\eta_{hij} = \pm \sqrt{g} \varepsilon_{hij}, \quad g = \det(g_{ij})}$$

col segno + se la base (e_i) è concorde con la base fondamentale (c_α) . Infatti, posto che le due basi sono legate da relazioni del tipo

$$e_i = E_i^\alpha c_\alpha,$$

e posto che la base (c_α) è canonica e quindi che le componenti del tensore metrico rispetto a questa coincidono col simbolo di Kronecker, risulta

$$g_{ij} = E_i^\alpha E_j^\beta \delta_{\alpha\beta},$$

per cui

$$(6) \quad \det(g_{ij}) = (\det(E_i^\alpha))^2.$$

D'altra parte, per la legge di trasformazione delle componenti di un tensore, si ha anche, stante la (3),

$$\eta_{hij} = E_h^\alpha E_i^\beta E_j^\gamma \varepsilon_{\alpha\beta\gamma}.$$

Scelto per esempio $(h, i, j) = (1, 2, 3)$ e ricordata la definizione di determinante, si vede che quest'ultima formula implica

$$\eta_{123} = E_1^\alpha E_2^\beta E_3^\gamma \varepsilon_{\alpha\beta\gamma} = \det(E_i^\alpha).$$

Se invece $(h, i, j) = (2, 1, 3)$, allora

$$\eta_{213} = -\det(E_i^\alpha).$$

Si riconosce così che vale la formula generale (4). Dalla (4) segue poi la (5) in virtù della (6).

OSSERVAZIONE 2. Data una base canonica (c_α) e denotata con (γ^α) la sua duale, la forma volume risulta definita dal doppio prodotto esterno (si veda il §7, formula (8)):

$$(7) \quad \boxed{\eta = \gamma^1 \wedge \gamma^2 \wedge \gamma^3}$$

Infatti, utilizzando il simbolo di Levi-Civita, si può scrivere

$$\gamma^1 \wedge \gamma^2 \wedge \gamma^3 = 3! (\gamma^1 \otimes \gamma^2 \otimes \gamma^3)^A = \varepsilon_{\alpha\beta\gamma} \gamma^\alpha \otimes \gamma^\beta \otimes \gamma^\gamma.$$

Per cui

$$\begin{aligned} (\gamma^1 \wedge \gamma^2 \wedge \gamma^3)(u, v, w) &= \varepsilon_{\alpha\beta\gamma} (\gamma^\alpha \otimes \gamma^\beta \otimes \gamma^\gamma)(u, v, w) \\ &= \varepsilon_{\alpha\beta\gamma} u^\alpha v^\beta w^\gamma, \end{aligned}$$

come nella (2). Dunque la forma volume ammette la definizione equivalente (7). Tuttavia, grazie all'identificazione fra vettori e covettori (siamo in uno spazio euclideo) e al fatto che si tratta di basi canoniche si può anche scrivere (si tenga presente l'Oss. 4, §7):

$$(7') \quad \boxed{\eta = c_1 \wedge c_2 \wedge c_3}$$

11.2. L'aggiunzione

Introdotta la forma volume, ad ogni vettore $u \in \mathbb{E}_3$ si fa corrispondere una forma bilineare antisimmetrica, sempre su $\mathbb{E}_3$, chiamata **aggiunta** di u , denotata con $*u$ e definita dall'uguaglianza

$$(1) \quad \boxed{(*u)(v, w) = \eta(u, v, w)}$$

Posto

$$\varphi = *u,$$

si ha successivamente:

$$\begin{aligned} \varphi_{ij} &= \varphi(e_i, e_j) = (*u)(e_i, e_j) \\ &= \eta(u, e_i, e_j) = u^h \eta(e_h, e_i, e_j) = u^h \eta_{hij}. \end{aligned}$$

Quindi le componenti φ_{ij} della forma aggiunta, in una base generica dell'orientamento scelto, sono date da

$$(2) \quad \boxed{\varphi_{ij} = \eta_{hij} u^h}$$

Se in particolare la base è canonica

$$(3) \quad \boxed{\varphi_{\alpha\beta} \stackrel{*}{=} \varepsilon_{\gamma\alpha\beta} u^\gamma}$$

vale a dire:

$$(4) \quad \boxed{(\varphi_{\alpha\beta}) = \begin{pmatrix} 0 & u^3 & -u^2 \\ -u^3 & 0 & u^1 \\ u^2 & -u^1 & 0 \end{pmatrix}}$$

La corrispondenza

$$*: \mathbb{E}_3 \rightarrow \Lambda^2(\mathbb{E}_3): u \mapsto *u$$

che così si genera è un isomorfismo. Con lo stesso simbolo denotiamo la corrispondenza inversa, denotiamo cioè con $*\varphi$ il **vettore aggiunto** della 2-forma φ . In questo modo si ha $**u = u$ e $**\varphi = \varphi$. Dalla (4) si vede che

$$(5) \quad \varphi_{23} = u^1, \quad \varphi_{31} = u^2, \quad \varphi_{12} = u^3,$$

cioè che la relazione inversa della (3) può essere scritta, considerando il simbolo di Levi-Civita con gli indici in alto,

$$(6) \quad u^\alpha = \frac{1}{2} \varepsilon^{\alpha\beta\gamma} \varphi_{\beta\gamma}$$

In una base qualsiasi, sempre nell'orientamento scelto, risulta

$$(7) \quad u^i = \frac{1}{2} \eta^{ijh} \varphi_{jh}$$

dove le (η^{ijh}) sono le componenti contravarianti della forma volume:

$$(8) \quad \eta^{ijh} = g^{ik} g^{jl} g^{hm} \eta_{klm}.$$

Queste si possono a loro volta esprimere mediante il simbolo di Levi-Civita:

$$(9) \quad \eta^{ijh} = \frac{1}{\sqrt{g}} \varepsilon^{ijh}.$$

Infatti dalla (8) si trae:

$$\eta^{ijh} = \sqrt{g} g^{ik} g^{jl} g^{hm} \varepsilon_{klm} = \sqrt{g} \varepsilon^{ijh} \det(g^{rs}) = \frac{1}{\sqrt{g}} \varepsilon^{ijh},$$

poiché $\det(g^{rs}) = g^{-1}$.

La forma volume permette anche di stabilire un isomorfismo fra $\mathbb{R} = \Lambda^0(\mathbb{E}_3)$ e $\Lambda^3(\mathbb{E}_3)$, che chiamiamo ancora **aggiunzione** e che denotiamo ancora con $*$. Ad un numero reale a corrisponde la 3-forma

$$(10) \quad *a = a \eta$$

e, viceversa, ad una 3-forma φ corrisponde il numero reale

$$(11) \quad *\varphi = \frac{1}{3!} \eta^{hij} \varphi_{hij} = \frac{1}{3!} \varepsilon^{\alpha\beta\gamma} \varphi_{\alpha\beta\gamma}$$

Per riconoscere la coerenza tra la (10) e la (11), cioè che $**a = a$, occorre osservare che per il simbolo di Levi-Civita vale l'uguaglianza

$$(12) \quad \varepsilon^{\alpha\beta\gamma} \varepsilon_{\alpha\beta\gamma} = 3!$$

Accanto a questa valgono anche le seguenti, la cui verifica è lasciata come esercizio:

$$(13) \quad \varepsilon^{\alpha\beta\gamma} \varepsilon_{\mu\beta\gamma} = 2 \delta_\mu^\alpha, \quad \varepsilon^{\alpha\beta\gamma} \varepsilon_{\mu\nu\gamma} = \delta_\mu^\alpha \delta_\nu^\beta - \delta_\nu^\alpha \delta_\mu^\beta.$$

11.3. Il prodotto vettoriale

Per mezzo della forma volume definiamo un'operazione binaria interna su $\mathbb{E}_3$, detta **prodotto vettoriale** e denotata con $\times$, ponendo

$$(1) \quad u \times v \cdot w = \eta(u, v, w)$$

Si tenga presente che il prodotto vettoriale è in alcuni testi denotato con $\wedge$ ed è anche chiamato *prodotto esterno*. Tuttavia il simbolo $\wedge$ è oggi universalmente usato per il prodotto esterno di forme antisimmetriche (§7).

Segue subito dalla definizione (1) che il prodotto vettoriale $u \times v$ è ortogonale ai due fattori:

$$(2) \quad u \times v \cdot u = 0, \quad u \times v \cdot v = 0.$$

Segue inoltre che, se si prendono i tre vettori di una base canonica appartenente all'orientamento scelto, si ha

$$(3) \quad c_1 \times c_2 = c_3, \quad c_2 \times c_3 = c_1, \quad c_3 \times c_1 = c_2.$$

Il prodotto vettoriale è anticommutativo, bilineare e soddisfa all'identità di Jacobi:

$$(4) \quad \begin{cases} u \times v = -v \times u, \\ (a u + b v) \times w = a u \times w + b v \times w, \\ u \times (v \times w) + v \times (w \times u) + w \times (u \times v) = 0. \end{cases}$$

Pertanto il prodotto vettoriale istituisce sullo spazio euclideo tridimensionale una struttura di *algebra di Lie*. Le prime due proprietà sono di immediata verifica. La terza si può far seguire, con qualche calcolo, dalla formula del **doppio prodotto vettoriale**

$$(5) \quad \boxed{\mathbf{u} \times (\mathbf{v} \times \mathbf{w}) = \mathbf{u} \cdot \mathbf{w} \mathbf{v} - \mathbf{u} \cdot \mathbf{v} \mathbf{w}}$$

che è equivalente a

$$(5') \quad \boxed{(\mathbf{u} \times \mathbf{v}) \times \mathbf{w} = \mathbf{u} \cdot \mathbf{w} \mathbf{v} - \mathbf{v} \cdot \mathbf{w} \mathbf{u}}$$

Per dimostrare la (5) si osserva che il primo membro è un vettore ortogonale a $\mathbf{v} \times \mathbf{w}$ e quindi, essendo questo prodotto ortogonale a $\mathbf{v}$ e $\mathbf{w}$, appartenente allo spazio generato dai vettori $(\mathbf{v}, \mathbf{w})$. Si può allora scrivere

$$\mathbf{u} \times (\mathbf{v} \times \mathbf{w}) = a \mathbf{v} + b \mathbf{w}.$$

Moltiplicando scalarmente per $\mathbf{u}$ si trova

$$0 = a \mathbf{v} \cdot \mathbf{u} + b \mathbf{w} \cdot \mathbf{u}.$$

Possiamo pertanto scrivere

$$a = \rho \mathbf{w} \cdot \mathbf{u}, \quad b = -\rho \mathbf{v} \cdot \mathbf{u},$$

con ρ costante indeterminata. Segue allora:

$$\mathbf{u} \times (\mathbf{v} \times \mathbf{w}) = \rho ((\mathbf{w} \cdot \mathbf{u}) \mathbf{v} - (\mathbf{v} \cdot \mathbf{u}) \mathbf{w}).$$

Per come è definito il prodotto vettoriale il primo membro è lineare rispetto ad ogni fattore. Così dicasi della quantità tra parentesi a secondo membro. Allora il fattore ρ deve essere un numero indipendente dai vettori $(\mathbf{u}, \mathbf{v}, \mathbf{w})$ che compaiono nell'uguaglianza. Se si considera per esempio una base canonica appartenente all'orientamento scelto e si prende $(\mathbf{u}, \mathbf{v}, \mathbf{w}) = (\mathbf{c}_1, \mathbf{c}_1, \mathbf{c}_2)$, posto che (si veda la (3)) $\mathbf{c}_1 \times \mathbf{c}_2 = \mathbf{c}_3$ e $\mathbf{c}_1 \times \mathbf{c}_3 = -\mathbf{c}_2$, risulta l'uguaglianza $-\mathbf{c}_3 = \rho(-\mathbf{c}_3)$. Deve quindi essere $\rho = 1$ e la (5) è dimostrata.

Ponendo $\mathbf{w} = \mathbf{e}_i$ nella (1) si trovano le componenti *covarianti* del prodotto vettoriale in una base generica:

$$(6) \quad (\mathbf{u} \times \mathbf{v})_i = \eta_{ihj} u^h v^j.$$

Le componenti contravarianti sono di conseguenza

$$(7) \quad (\mathbf{u} \times \mathbf{v})^i = \eta^{ihj} u_h v_j.$$

In una base canonica $(\mathbf{c}_\alpha)$, tenendo presente che non vi è più differenza numerica tra componenti covarianti e contravarianti, si ha

$$(8) \quad \boxed{(\mathbf{u} \times \mathbf{v})_\alpha = \varepsilon_{\alpha\beta\gamma} u^\beta v^\gamma}$$

ovvero anche

$$(8') \quad (\mathbf{u} \times \mathbf{v})^\alpha = \varepsilon^{\alpha\beta\gamma} u_\beta v_\gamma.$$

Di qui la possibilità di esprimere il prodotto vettoriale mediante un determinante formale:

$$(9) \quad \boxed{\mathbf{u} \times \mathbf{v} = \begin{vmatrix} \mathbf{c}_1 & \mathbf{c}_2 & \mathbf{c}_3 \\ u^1 & u^2 & u^3 \\ v^1 & v^2 & v^3 \end{vmatrix}}$$

Tenuto conto delle (8) e (8'), nonché della (13)₂ del §11.1, si può ridimostrare la formula (5) del doppio prodotto vettoriale alla maniera seguente:

$$\begin{aligned} (\mathbf{u} \times (\mathbf{v} \times \mathbf{w}))_\alpha &= \varepsilon_{\alpha\beta\gamma} u^\beta (\mathbf{v} \times \mathbf{w})^\gamma \\ &= \varepsilon_{\alpha\beta\gamma} u^\beta \varepsilon^{\gamma\mu\nu} v_\mu w_\nu \\ &= \varepsilon_{\alpha\beta\gamma} \varepsilon^{\mu\nu\gamma} u^\beta v_\mu w_\nu \\ &= u^\beta w_\beta v_\alpha - u^\beta v_\beta w_\alpha. \end{aligned}$$

OSSERVAZIONE 1. Partendo dalla definizione (1) si dimostra che se ϑ è l'angolo compreso tra i vettori non nulli $\mathbf{u}$ e $\mathbf{v}$ (valutato tra 0 e π) allora

$$(10) \quad \boxed{|\mathbf{u} \times \mathbf{v}| = |\mathbf{u}| |\mathbf{v}| \sin \vartheta}$$

Infatti, per le proprietà già viste:

$$\begin{aligned}
 |\mathbf{u} \times \mathbf{v}|^2 &= (\mathbf{u} \times \mathbf{v}) \cdot (\mathbf{u} \times \mathbf{v}) \\
 &= (\mathbf{u} \times \mathbf{v}) \times \mathbf{u} \cdot \mathbf{v} \\
 &= (\mathbf{u} \cdot \mathbf{u} \mathbf{v} - \mathbf{u} \cdot \mathbf{v} \mathbf{u}) \cdot \mathbf{v} \\
 &= |\mathbf{u}|^2 |\mathbf{v}|^2 - (\mathbf{u} \cdot \mathbf{v})^2 \\
 &= |\mathbf{u}|^2 |\mathbf{v}|^2 \left(1 - \frac{(\mathbf{u} \cdot \mathbf{v})^2}{|\mathbf{u}|^2 |\mathbf{v}|^2}\right) \\
 &= |\mathbf{u}|^2 |\mathbf{v}|^2 (1 - \cos^2 \vartheta).
 \end{aligned}$$

OSSERVAZIONE 2. L'operazione combinata di prodotto vettoriale e prodotto scalare, $\mathbf{u} \times \mathbf{v} \cdot \mathbf{w}$, è chiamata **prodotto misto**. Dalla definizione (1) seguono subito, stante l'antisimmetria di η , la **proprietà anticommutativa**, la **proprietà ciclica** e la **proprietà di scambio** del prodotto misto:

$$(11) \quad \begin{cases} \mathbf{u} \times \mathbf{v} \cdot \mathbf{w} = -\mathbf{v} \times \mathbf{u} \cdot \mathbf{w}, & \mathbf{u} \times \mathbf{v} \cdot \mathbf{w} = -\mathbf{u} \times \mathbf{w} \cdot \mathbf{v}, & \dots \\ \mathbf{u} \times \mathbf{v} \cdot \mathbf{w} = \mathbf{v} \times \mathbf{w} \cdot \mathbf{u} = \mathbf{w} \times \mathbf{u} \cdot \mathbf{v}, \\ \mathbf{u} \times \mathbf{v} \cdot \mathbf{w} = \mathbf{u} \cdot \mathbf{v} \times \mathbf{w}. \end{cases}$$

OSSERVAZIONE 3. *Endomorfismi assiali*. Ad ogni vettore $\mathbf{a} \in \mathbb{E}_3$ corrisponde per aggiunta una forma bilineare antisimmetrica $*\mathbf{a}$. Interpretata questa come endomorfismo (antisimmetrico), vale la formula

$$(12) \quad \boxed{*\mathbf{a}(\mathbf{v}) = \mathbf{a} \times \mathbf{v}}$$

Infatti, in virtù delle definizioni sopra date, si ha successivamente

$$*\mathbf{a}(\mathbf{v}) \cdot \mathbf{u} = *\mathbf{a}(\mathbf{v}, \mathbf{u}) = \eta(\mathbf{a}, \mathbf{v}, \mathbf{u}) = \mathbf{a} \times \mathbf{v} \cdot \mathbf{u},$$

per cui, stante l'arbitrarietà del vettore $\mathbf{u}$, segue la (12). Siccome l'aggiunzione è una corrispondenza biunivoca fra vettori e forme bilineari antisimmetriche, e quindi endomorfismi antisimmetrici, concludiamo che nello spazio euclideo tridimensionale ogni endomorfismo antisimmetrico è del tipo (12), cioè del tipo $\mathbf{a} \times$, con $\mathbf{a} \in \mathbb{E}_3$. Si noti che il vettore $\mathbf{a}$ è un autovettore di $*\mathbf{a}$, corrispondente all'autovalore nullo; infatti

$*\mathbf{a}(\mathbf{a}) = \mathbf{a} \times \mathbf{a} = 0$. Per questo motivo gli endomorfismi antisimmetrici sono anche chiamati **endomorfismi assiali**.

OSSERVAZIONE 4. La corrispondenza biunivoca così stabilita tra i vettori e gli endomorfismi antisimmetrici è un *isomorfismo d'algebra di Lie* (si ricordi che, Oss. 3 del §9, in ogni spazio euclideo gli endomorfismi antisimmetrici formano, con il commutatore, un'algebra di Lie). Sussiste infatti l'uguaglianza

$$(13) \quad \boxed{*[A, B] = (*A) \times (*B)}$$

ovvero l'uguaglianza equivalente

$$(13') \quad \boxed{[*a, *b] = *(a \times b)}$$

Infatti, posto $*\mathbf{a} = A$ e $*\mathbf{b} = B$, si ha

$$AB(\mathbf{u}) = \mathbf{a} \times (\mathbf{b} \times \mathbf{u})$$

e quindi, utilizzando l'identità di Jacobi per il doppio prodotto vettoriale,

$$\begin{aligned}
 [A, B](\mathbf{u}) &= \mathbf{a} \times (\mathbf{b} \times \mathbf{u}) - \mathbf{b} \times (\mathbf{a} \times \mathbf{u}) \\
 &= \mathbf{a} \times (\mathbf{b} \times \mathbf{u}) + \mathbf{b} \times (\mathbf{u} \times \mathbf{a}) \\
 &= -\mathbf{u} \times (\mathbf{a} \times \mathbf{b}) = (\mathbf{a} \times \mathbf{b}) \times \mathbf{u} = *(a \times b)(\mathbf{u}).
 \end{aligned}$$

OSSERVAZIONE 5. Sappiamo che in $\mathbb{E}_3$, come in ogni spazio euclideo, esiste una corrispondenza biunivoca fra vettori e covettori (§8). Al §5 abbiamo definito il *prodotto esterno di due covettori*. Il legame tra questo prodotto esterno e il prodotto vettoriale è espresso dall'uguaglianza

$$(14) \quad \boxed{*(\mathbf{u} \times \mathbf{v}) = \mathbf{u}^b \wedge \mathbf{v}^b}$$

Infatti, tenuto conto della formula (5') sul doppio prodotto vettoriale, si ha successivamente:

$$\begin{aligned}
 *(a \times b)(\mathbf{u}, \mathbf{v}) &= \eta(a \times b, \mathbf{u}, \mathbf{v}) \\
 &= (a \times b) \times \mathbf{u} \cdot \mathbf{v} \\
 &= a \cdot \mathbf{u} b \cdot \mathbf{v} - a \cdot \mathbf{v} b \cdot \mathbf{u} \\
 &= \langle \mathbf{u}, \mathbf{a}^b \rangle \langle \mathbf{v}, \mathbf{b}^b \rangle - \langle \mathbf{v}, \mathbf{a}^b \rangle \langle \mathbf{u}, \mathbf{b}^b \rangle \\
 &= (\mathbf{a}^b \wedge \mathbf{b}^b)(\mathbf{u}, \mathbf{v}).
 \end{aligned}$$

11.4. Le rotazioni

Consideriamo le rotazioni dello spazio $\mathbb{E}_3$. Ricordiamo, dalla teoria generale svolta al §10, che una *rotazione* è una isometria con determinante uguale a +1 e che i suoi autovalori, se lo spazio è strettamente euclideo, sono unitari, quindi del tipo $e^{i\vartheta} = \cos \vartheta + i \sin \vartheta$. Per $n = 3$ gli autovalori sono radici di un polinomio di terzo grado a coefficienti reali. Quindi almeno uno di essi è reale, gli altri due complessi coniugati o reali. Imponendo la condizione che il prodotto di tutti e tre gli autovalori sia uguale a 1 = $\det(Q)$ si trova che l'unico spettro possibile di una rotazione è del tipo

$$(1), e^{i\vartheta}, e^{-i\vartheta}).$$

Ciò premesso, se escludiamo il caso $Q = 1$, il cui spettro è (1,1,1), l'autovalore 1 determina un sottospazio unidimensionale di autovettori $K_1 \subset \mathbb{E}_3$ che chiamiamo **asse della rotazione**. Il numero ϑ che compare nella (1), determinato a meno del segno e di multipli di 2π , prende il nome di **angolo della rotazione**.

Si può rappresentare la rotazione Q mediante una coppia (u, ϑ) dove u è un versore dell'asse e ϑ è l'angolo della rotazione, orientato in maniera tale che per ogni vettore c ortogonale a u il passaggio da c a $Q(c)$ secondo il verso di ϑ sia tale da produrre un "avvitamento" nel verso di u (è sottintesa la scelta del solito orientamento dello spazio, secondo la regola della mano destra). Si dice allora che (u, ϑ) sono **versore e angolo della rotazione** Q o anche che u è il **versore dell'angolo** ϑ (Fig. 11.2). Si osservi che le coppie (u, ϑ) e $(-u, -\vartheta)$ danno luogo alla stessa rotazione come anche, in particolare, le coppie (u, π) e $(u, -\pi)$.

Se si sceglie una base canonica (c_α) tale che $c_3 = u$, allora la matrice delle componenti della rotazione Q generata dalla coppia (u, ϑ) , con u versore di ϑ , ha la forma seguente:

$$(2) \quad (Q_\alpha^\beta) = \begin{pmatrix} \cos \vartheta & \sin \vartheta & 0 \\ -\sin \vartheta & \cos \vartheta & 0 \\ 0 & 0 & 1 \end{pmatrix}.$$

Infatti, tenuto presente che in una base canonica

$$Q_\alpha^\beta = Q_{\alpha\beta} = Q(c_\alpha) \cdot c_\beta,$$

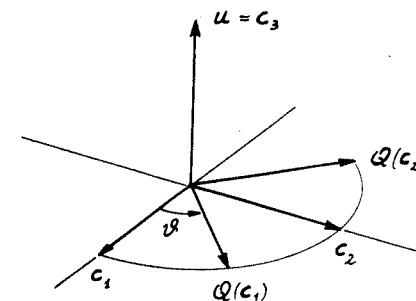


Fig. 11.2: versore-angolo di una rotazione.

si vede con l'ausilio della Fig. 11.2 che $Q_{11} = Q(c_1) \cdot c_1 = \cos \vartheta$, $Q_{12} = Q(c_1) \cdot c_2 = \sin \vartheta$, e così via.

Oltre alla coppia versore-angolo vi sono varie altre rappresentazioni della rotazioni, per esempio quella esponenziale e quella quaternionale, già esaminate al §10. Oltre a queste è da prendersi in considerazione quella fornita dagli **angoli di Eulero** (Leonhard Euler, 1707-1783), basata sulla relazione fra una terna ortonormale di vettori, detta **terna fissa**, e la **terna ruotata**: la prima è una base canonica $(c_\alpha) = (i, j, k)$; la seconda è la terna $(Q(c_\alpha)) = (i', j', k')$. Gli angoli di Euler (θ, ϕ, ψ) , detti rispettivamente **angolo di nutazione**, **angolo di rotazione propria**, **angolo di precessione** e soddisfacenti alle limitazioni

$$(3) \quad \begin{cases} 0 < \theta < \pi, \\ 0 \leq \phi < 2\pi, \\ 0 \leq \psi < 2\pi, \end{cases}$$

sono definiti, nell'ipotesi che $k' = Q(c_3)$ sia distinto da $k = c_3$, come segue (Fig. 11.3). (I) L'angolo di nutazione θ è l'angolo compreso tra k e k' . (II) Si consideri l'intersezione fra il piano (i', j') e il piano (i, j) : è una "retta" detta **linea dei nodi**. Questa è individuata dal versore del prodotto vettoriale $k \times k'$ che denotiamo con N (infatti la retta in questione è ortogonale sia a k che a k'). (III) L'angolo di rotazione propria ϕ è l'angolo formato da (N, i') di versore k' e l'angolo di precessione ψ è l'angolo formato da (i, N) di versore k . Da questa

definizione discende che assegnati comunque tre angoli soddisfacenti alle limitazioni (3) risulta univocamente definita la terna ruotata (si veda anche l'Esercizio 2 seguente).

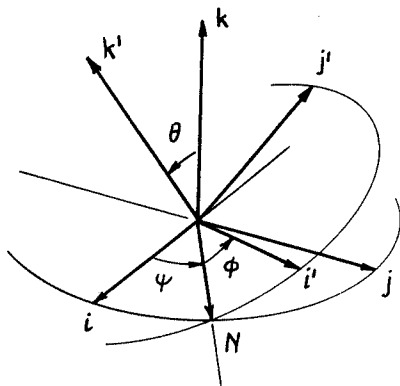


Fig. 11.3: gli angoli di Euler.

* * *

ESERCIZIO 1. Poniamoci il problema di (I) come determinare la coppia $(\mathbf{u}, \vartheta)$ a partire dalla rotazione $\mathbf{Q}$ e, viceversa, di (II) come determinare la rotazione $\mathbf{Q}$ a partire dalla coppia $(\mathbf{u}, \vartheta)$. Osserviamo innanzitutto che, data $\mathbf{Q}$, il versore $\mathbf{u}$ è un autovettore associato all'autovalore 1, quindi determinabile risolvendo un sistema di equazioni lineari. Tuttavia, è più semplice la sua determinazione attraverso la parte antisimmetrica di $\mathbf{Q}$, almeno quando questa non è nulla, perché si ha

$$(4) \quad \sin \vartheta \mathbf{u} = *\mathbf{Q}^A.$$

Se infatti si guarda alla matrice (2), che è ottenuta scegliendo opportunamente una base canonica, e si ricorda la relazione tra le componenti

di una forma bilineare antisimmetrica e quelle del vettore aggiunto, data per esempio dalla (4) del §11.2, si trova l'uguaglianza

$$(A_{\alpha}^{\beta}) = \begin{pmatrix} 0 & a^3 & -a^2 \\ -a^3 & 0 & a^1 \\ a^2 & -a^1 & 0 \end{pmatrix} = \begin{pmatrix} 0 & \sin \vartheta & 0 \\ -\sin \vartheta & 0 & 0 \\ 0 & 0 & 0 \end{pmatrix},$$

posto $\mathbf{A} = \mathbf{Q}^A$ e $\mathbf{a} = *\mathbf{A}$. Di qui la (4). La formula (4), fissata una delle due possibili scelte del versore $\mathbf{u}$, determina solamente $\sin \vartheta$ e non ϑ . Occorre quindi, per eliminare questa indeterminazione, estrarre qualche altra informazione da $\mathbf{Q}$. Si può per esempio calcolare $\cos \vartheta$ attraverso la formula (che lasciamo da dimostrare)

$$(5) \quad 1 + 2 \cos \vartheta = I_2(\mathbf{Q}) = \text{tr}(\mathbf{Q}).$$

Questa ci consente fra l'altro di osservare che l'invariante primo (la traccia) e l'invariante secondo di una rotazione coincidono. Il problema inverso (II) può essere risolto con l'uso della rappresentazione esponenziale. Infatti, data una qualunque coppia $(\mathbf{u}, \vartheta)$, e introdotto l'endomorfismo aggiunto $\mathbf{U} = *\mathbf{u}$, la rotazione

$$(6) \quad \mathbf{Q} = e^{\vartheta \mathbf{U}}$$

è proprio la rotazione di versore-angolo $(\mathbf{u}, \vartheta)$. Per riconoscerlo cominciamo con l'osservare che:

$$\begin{aligned} \mathbf{U}(\mathbf{v}) &= (*\mathbf{U}) \times \mathbf{v} = \mathbf{u} \times \mathbf{v}, \\ \mathbf{U}^2(\mathbf{v}) &= \mathbf{u} \times \mathbf{U}(\mathbf{v}) = \mathbf{u} \times (\mathbf{u} \times \mathbf{v}) = \mathbf{u} \cdot \mathbf{v} \mathbf{u} - \mathbf{v}, \\ \mathbf{U}^3(\mathbf{v}) &= \mathbf{u} \times \mathbf{U}^2(\mathbf{v}) = -\mathbf{u} \times \mathbf{v} = -\mathbf{U}(\mathbf{v}), \end{aligned}$$

quindi che

$$\mathbf{U}^{2n+2} = (-1)^n \mathbf{U}^2, \quad \mathbf{U}^{2n+1} = (-1)^n \mathbf{U}.$$

Si ha allora successivamente:

$$\begin{aligned} e^{\vartheta \mathbf{U}} &= \sum_{n=0}^{\infty} \frac{1}{2n!} \vartheta^{2n} \mathbf{U}^{2n} + \sum_{n=0}^{\infty} \frac{1}{(2n+1)!} \vartheta^{2n+1} \mathbf{U}^{2n+1} \\ &= 1 + \left(\sum_{n=1}^{\infty} \frac{(-1)^{n-1}}{2n!} \vartheta^{2n} \right) \mathbf{U}^2 + \left(\sum_{n=0}^{\infty} \frac{(-1)^n}{(2n+1)!} \vartheta^{2n+1} \right) \mathbf{U}, \end{aligned}$$

vale a dire

$$(7) \quad Q = e^{\vartheta U} = 1 + \sin \vartheta U + (1 - \cos \vartheta) U^2.$$

Questa formula è in accordo con le precedenti (4) e (5) perché implica

$$*Q^A = \sin \vartheta *U = \sin \vartheta u,$$

$$\text{tr}(Q) = 3 + (1 - \cos \vartheta) \text{tr}(U^2) = 2 \cos \vartheta + 1,$$

valendo la formula generale, la cui verifica è immediata,

$$(8) \quad \text{tr}(U^2) = -2 \|u\|,$$

che quindi in questo caso produce $\text{tr}(U^2) = -2$. La formula (7), tenuto conto delle espressioni $U(v)$ e $U^2(v)$ sopra scritte, consente anche di esprimere il generico vettore $Q(v)$ mediante la coppia (u, ϑ) :

$$(9) \quad Q(v) = \cos \vartheta v + (1 - \cos \vartheta) u \cdot v u + \sin \vartheta u \times v.$$

ESERCIZIO 2. Dimostrare che la matrice delle componenti della rotazione Q determinata dagli angoli di Euler (θ, ϕ, ψ) è:

$$(10) \quad (Q_{\alpha}^{\beta}) = \begin{pmatrix} \cos \phi \cos \psi & \cos \phi \sin \psi & \sin \phi \sin \theta \\ -\sin \phi \sin \psi \cos \theta & +\sin \phi \cos \psi \cos \theta & \sin \phi \sin \theta \\ -\sin \phi \cos \psi & -\sin \phi \sin \psi & \cos \phi \sin \theta \\ -\cos \phi \sin \psi \cos \theta & +\cos \phi \cos \psi \cos \theta & \cos \phi \sin \theta \\ \sin \psi \sin \theta & -\cos \psi \sin \theta & \cos \theta \end{pmatrix}$$

Per giungere alla (10) si può osservare che la rotazione Q è composta da tre rotazioni successive Φ, Θ, Ψ , caratterizzate dalle seguenti coppie versore-angolo:

$$\Phi = (k, \phi), \quad \Theta = (i, \theta), \quad \Psi = (k, \psi).$$

Le matrici delle componenti sono rispettivamente:

$$(\Phi_{\alpha}^{\beta}) = \begin{pmatrix} \cos \phi & \sin \phi & 0 \\ -\sin \phi & \cos \phi & 0 \\ 0 & 0 & 1 \end{pmatrix}, \quad (\Theta_{\alpha}^{\beta}) = \begin{pmatrix} 1 & 0 & 0 \\ 0 & \cos \theta & \sin \theta \\ 0 & -\sin \theta & \cos \theta \end{pmatrix},$$

$$(\Psi_{\alpha}^{\beta}) = \begin{pmatrix} \cos \psi & \sin \psi & 0 \\ -\sin \psi & \cos \psi & 0 \\ 0 & 0 & 1 \end{pmatrix}.$$

Si deve quindi eseguire il prodotto matriciale

$$Q_{\alpha}^{\beta} = \Phi_{\alpha}^{\rho} \Theta_{\rho}^{\sigma} \Psi_{\sigma}^{\beta}.$$

Eseguendo il prodotto (righe per colonne, secondo le convenzioni adottate) delle prime due matrici si ottiene

$$(\Phi_{\alpha}^{\rho} \Theta_{\rho}^{\sigma}) = \begin{pmatrix} \cos \phi & \sin \phi \cos \theta & \sin \phi \sin \theta \\ -\sin \phi & \cos \phi \cos \theta & \cos \phi \sin \theta \\ 0 & -\sin \theta & \cos \theta \end{pmatrix}$$

Moltiplicando questa per la terza si ottiene la (10).

12. Spazi vettoriali iperbolici

Uno **spazio vettoriale iperbolico** è uno spazio vettoriale pseudoeuclideo (H_n, g) di segnatura $(n-1, 1)$. Ciò significa (si veda il §8) che in ogni base canonica (c_i) un vettore ha norma negativa (si dice che è del *genere tempo*) e gli altri $n-1$ hanno norma positiva (si dicono del *genere spazio*). Conveniamo che gli indici latini varino da 0 a $n-1$, che in una base canonica (c_i) il vettore del genere tempo sia c_0 , che i rimanenti vettori del genere spazio siano denotati con (c_{α}) , e che gli indici greci varino quindi da 1 a $n-1$. La matrice metrica in una base canonica avrà pertanto la forma

$$(1) \quad (g_{ij}) = \begin{pmatrix} -1 & 0 & \dots & 0 \\ 0 & 1 & \dots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \dots & 1 \end{pmatrix}.$$

Lo spazio iperbolico di dimensione 4 prende il nome di **spazio vettoriale di Minkowski** (Hermann Minkowski, 1864-1909). Questo spazio è di fondamentale importanza per la teoria della Relatività Ristretta.

Ricordiamo ancora che un vettore ortogonale a se stesso è detto *isotropo* o *del genere luce*.

Ci limitiamo in questo paragrafo ad esaminare alcune proprietà fondamentali degli spazi iperbolici in generale e, in particolare, dello spazio iperbolico bidimensionale. Si tenga presente che per *spazio iperbolico* alcuni autori intendono uno spazio di dimensione pari $2n$ e segnatura (n, n) .

PROPOSIZIONE 1. Un vettore (non nullo) ortogonale ad un vettore del genere tempo è un vettore del genere spazio. Il sottospazio ortogonale ad un vettore del genere tempo è uno spazio strettamente euclideo (di dimensione $n - 1$).

Dimostrazione. Sia v un vettore del genere tempo. Si consideri il sottospazio unidimensionale K generato da v e il suo ortogonale $K^\perp$. Non può essere $v \in K^\perp$ perché v sarebbe anche ortogonale a se stesso, quindi isotropo. Dunque $K \cap K^\perp = 0$ e il sottospazio $K^\perp$ è a metrica non degenera (si veda l'Oss. 6 del §8). Si può quindi determinare su $K^\perp$, che è di dimensione $n - 1$, una base canonica (c_α) . Se consideriamo anche il versore c_0 di v , otteniamo una base canonica. Siccome c_0 è del genere tempo tutti gli altri vettori della base sono del genere spazio, altrimenti si violerebbe il teorema di Sylvester, posto che la segnatura è $(n - 1, 1)$. Dunque tutti i vettori di $K^\perp$ sono del genere spazio e questo è strettamente euclideo. ■

In maniera analoga si dimostra che

PROPOSIZIONE 2. Il sottospazio ortogonale ad un vettore del genere spazio è uno spazio iperbolico se $n > 2$ o del genere tempo se $n = 2$.

Dalla Proposizione 1 segue che i sottospazi del genere tempo hanno dimensione 1. Segue anche che, fissato un sottospazio $K \subset H$ del genere tempo, risulta definita una **decomposizione spazio-temporale** dello spazio iperbolico H nella somma diretta

$$H = K \oplus S,$$

dove S è il sottospazio strettamente euclideo ortogonale a K ($S = K^\perp$).

PROPOSIZIONE 3. Il sottospazio ortogonale ad un vettore (non nullo) del genere luce è a metrica degenera contenente il vettore stesso.

Tutti i vettori del sottospazio indipendenti da questo sono del genere spazio.

Dimostrazione. Sia v un vettore isotropo, K il sottospazio da esso generato. Siccome v è ortogonale a se stesso per definizione, risulta $K \subset K^\perp$. Quindi $K \cap K^\perp \neq 0$ e il sottospazio $K^\perp$ è a metrica degenera. Se in esso vi fosse un vettore del genere tempo, si contrasterebbe con la Proposizione 1, che implica v del genere spazio. Se vi fossero due vettori indipendenti del genere luce, si contrasterebbe con la Proposizione 4 seguente. ■

PROPOSIZIONE 4. Due vettori del genere luce e ortogonali sono dipendenti (cioè *paralleli*).

Dimostrazione. Un generico vettore isotropo $v \neq 0$ può sempre mettersi nella forma $v = a(u + c_0)$, dove u è un vettore unitario, del genere spazio, ortogonale a c_0 , generico vettore unitario del genere tempo; basta infatti porre $u = \frac{1}{a}v - c_0$ con $a = -v \cdot c_0$ (si osservi che la condizione $v \cdot c_0 = 0$ violerebbe la Proposizione 1). Se $v' = a'(u' + c_0)$ è un altro vettore isotropo decomposto allo stesso modo, dalla condizione $v \cdot v' = 0$ segue $u \cdot u' - 1 = 0$. Ma due vettori unitari del genere spazio hanno prodotto scalare uguale a 1 se e solo se sono uguali. Da $u = u'$ segue che v e v' sono dipendenti. ■

Dalla proposizione precedente segue che i sottospazi isotropi hanno dimensione 1, a conferma della proprietà generale messa in evidenza nell'Oss. 5 del n.8.

ESERCIZIO 1. Si dimostri che se $I \subset H$ è un sottospazio isotropo, il sottospazio quoziente H/I ha una naturale struttura di spazio vettoriale strettamente euclideo.

DEFINIZIONE 1. Denotiamo con T il sottoinsieme dei vettori del genere tempo. Diciamo che due vettori $u, v \in T$ sono **ortocroni** se $u \cdot v < 0$

PROPOSIZIONE 5. La relazione di ortocronismo è una relazione di equivalenza che divide T in due classi.

Dimostrazione. La relazione di ortocronismo è *riflessiva* perché si ha $v \cdot v = \|v\| < 0$ per ogni vettore del genere tempo (per definizione).

È ovviamente *simmetrica*, perché è simmetrico il prodotto scalare. Per dimostrare la proprietà *transitiva* dobbiamo dimostrare l'implicazione

$$(2) \quad \mathbf{u} \cdot \mathbf{v} < 0, \quad \mathbf{v} \cdot \mathbf{w} < 0 \implies \mathbf{u} \cdot \mathbf{w} < 0.$$

Per dimostrare che le classi di equivalenza sono due dobbiamo dimostrare l'implicazione

$$(3) \quad \mathbf{u} \cdot \mathbf{v} > 0, \quad \mathbf{v} \cdot \mathbf{w} > 0 \implies \mathbf{u} \cdot \mathbf{w} < 0.$$

Per questo scopo, supposto senza ledere la generalità che il vettore $\mathbf{v}$ sia unitario, si scelga una base canonica tale da aversi $\mathbf{c}_0 = \mathbf{v}$. Si ha allora:

$$\mathbf{u} \cdot \mathbf{w} = \sum_{\alpha=1}^{n-1} u^\alpha w^\alpha - u^0 w^0, \quad \mathbf{u} \cdot \mathbf{v} = g_{00} u^0 = -u^0, \quad \mathbf{w} \cdot \mathbf{v} = -w^0.$$

Le condizioni $\mathbf{u} \in T$ e $\mathbf{w} \in T$ sono equivalenti alle disuguaglianze

$$\sum_{\alpha=1}^{n-1} (u^\alpha)^2 - (u^0)^2 < 0, \quad \sum_{\alpha=1}^{n-1} (w^\alpha)^2 - (w^0)^2 < 0.$$

Per la disuguaglianza di Schwarz, valida nello spazio strettamente euclideo ortogonale a $\mathbf{c}_0$ e generato dai vettori $(\mathbf{c}_\alpha)$, si ha anche

$$\left(\sum_{\alpha=1}^{n-1} u^\alpha w^\alpha \right)^2 \leq \sum_{\alpha=1}^{n-1} (u^\alpha)^2 \sum_{\alpha=1}^{n-1} (w^\alpha)^2 < (u^0 w^0)^2.$$

Abbiamo allora le seguenti implicazioni: $(\mathbf{u} \cdot \mathbf{v} < 0, \mathbf{v} \cdot \mathbf{w} < 0) \iff (u^0 > 0, w^0 > 0) \implies u^0 w^0 > 0 \implies \mathbf{u} \cdot \mathbf{w} < 0$. Quindi l'implicazione (2) è dimostrata. Analogamente si ha: $(\mathbf{u} \cdot \mathbf{v} > 0, \mathbf{v} \cdot \mathbf{w} > 0) \iff (u^0 < 0, w^0 < 0) \implies u^0 w^0 > 0 \implies \mathbf{u} \cdot \mathbf{w} < 0$. Quindi anche l'implicazione (3) è dimostrata. ■

Denotiamo con T^+ e T^- le due classi di equivalenza in cui è diviso l'insieme T dei vettori del genere tempo; essi si chiamano rispettivamente **futuro** e **passato**. Denotiamo con L l'insieme dei vettori del genere luce. Essi formano il cosiddetto **cono di luce**. Se infatti denotiamo con $\mathbf{x} = x^i \mathbf{c}_i$ il generico vettore dello spazio e se imponiamo che esso sia isotropo, cioè che $\mathbf{x} \cdot \mathbf{x} = 0$, troviamo la condizione

$$(4) \quad \sum_{\alpha=1}^{n-1} (x^\alpha)^2 - (x^0)^2 = 0.$$

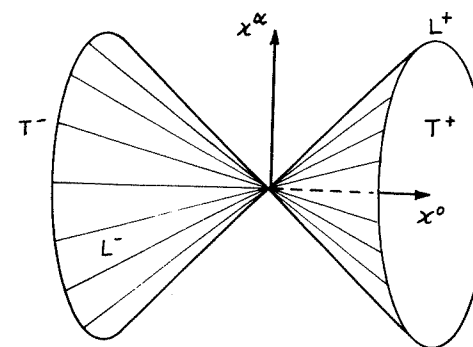


Fig. 12.1: il cono di luce.

Questa è l'equazione di un cono di centro l'origine (vettore nullo). Il cono di luce è diviso in due falde, L^+ e L^- , ciascuna delle quali separa dai vettori del genere spazio i vettori delle classi T^+ e T^- rispettivamente, i quali stanno all'interno di ciascuna falda (Fig. 12.1).

I vettori unitari del genere spazio sono caratterizzati dall'equazione

$$(5) \quad \sum_{\alpha=1}^{n-1} (x^\alpha)^2 - (x^0)^2 = 1.$$

Si tratta di un iperboloide a una falda. I vettori unitari del genere tempo sono invece caratterizzati dall'equazione

$$(6) \quad \sum_{\alpha=1}^{n-1} (x^\alpha)^2 - (x^0)^2 = -1.$$

Si tratta di un iperboloide a due falde. Osserviamo inoltre che

PROPOSIZIONE 6. La somma di due o più vettori di T^+ (risp. di T^-) è ancora un vettore di T^+ (risp. di T^-).

Dimostrazione. Siano $u, v \in T^+$. Ciò significa che $\|u\| < 0$, $\|v\| < 0$, $u \cdot v < 0$. Di qui segue $\|u + v\| = \|u\| + \|v\| + 2u \cdot v < 0$; quindi la somma $u + v$ è del genere tempo. Inoltre: $u \cdot (u + v) = u \cdot u + u \cdot v < 0$. Dunque u e la somma $u + v$ sono ortocroni. ■

ESERCIZIO 2. (I) Si verifichi che nello spazio $\text{End}(E)$ degli endomorfismi sopra uno spazio vettoriale E si definisce un prodotto scalare ponendo

$$(7) \quad A \cdot B = \text{tr}(AB).$$

(II) Si dimostri che, nel caso $n = 2$, la segnatura del tensore metrico così definito è $(3,1)$. Lo spazio $H_4 = \text{End}(E_2)$ degli endomorfismi su di uno spazio bidimensionale è quindi uno spazio di Minkowski. (III) Si dimostri che l'assegnazione di un tensore metrico definito positivo su E_2 implica la definizione di una decomposizione spazio-temporale di H_4 . (IV) Si estendano queste considerazioni al caso dello spazio $H_{n^2} = \text{End}(E_n)$, dimostrando che la segnatura del prodotto scalare definito dalla (7) è

$$\left(\frac{n(n+1)}{2}, \frac{n(n-1)}{2} \right).$$

Esaminiamo ora, per meglio comprendere la struttura di uno spazio iperbolico, il caso più semplice dello spazio iperbolico bidimensionale H_2 .

Fissata una base canonica (c_0, c_1) , con c_0 del genere tempo, osserviamo innanzitutto che il cono di luce, insieme dei vettori isotropi, ha ora equazione

$$(x^1)^2 - (x^0)^2 = 0,$$

e si riduce pertanto alle *rette isotrope* di equazione $x^1 - x^0 = 0$ e $x^1 + x^0 = 0$, che sono le bisettrici dei quadranti individuati dai vettori della base. Volendo raffigurare su di un piano i vettori (c_0, c_1) , unitari e ortogonali fra loro, dobbiamo tener presente che il foglio su cui riportiamo la figura (Fig. 12.2), fissato un punto corrispondente al vettore nullo, è la rappresentazione intuitiva dello spazio bidimensionale strettamente euclideo, non di quello iperbolico. Per esempio, contrariamente al nostro senso euclideo, i vettori che si trovano sulle bisettrici hanno *lunghezza nulla*. Così, i vettori unitari del genere tempo sono caratterizzati dall'equazione

$$(x^0)^2 - (x^1)^2 = 1,$$

e descrivono quindi un'iperbole di vertici i punti $(\pm 1, 0)$ e asintoti le rette isotrope, mentre i vettori unitari del genere spazio sono caratterizzati dall'equazione

$$(x^1)^2 - (x^0)^2 = 1,$$

e descrivono pertanto l'iperbole di vertici i punti $(0, \pm 1)$ con gli stessi asintoti. Se ora prendiamo un vettore u unitario del genere tempo, la retta descritta dai vettori ortogonali a questo ha equazione

$$u^1 x^1 - u^0 x^0 = 0.$$

Si consideri quindi uno dei due vettori v unitari e ortogonali a u ; sappiamo che questi sono del genere spazio. Allora

$$\|u - v\| = \|u\| + \|v\| + 2u \cdot v = -1 + 1 + 0 = 0.$$

Ciò significa che il vettore differenza $u - v$ è isotropo, quindi giacente su una delle due bisettrici (asintoti). Concludiamo che, nella rappresentazione euclidea, due vettori del piano iperbolico non isotropi, unitari e ortogonali, stanno in posizione simmetrica rispetto ad una delle due bisettrici (Fig. 12.2).

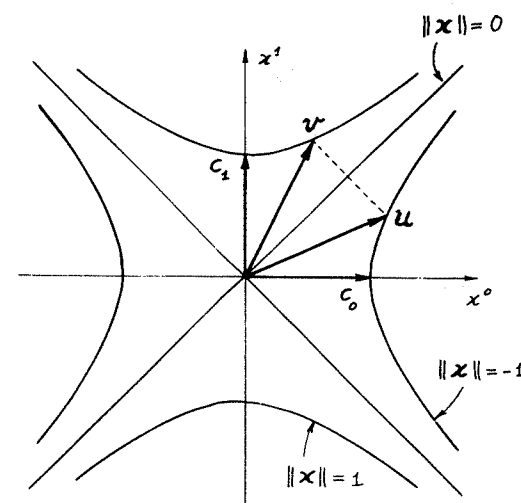


Fig. 12.2: ortogonalità dei vettori nel piano iperbolico.

Si osservi ancora, guardando la figura, che per un "osservatore iperbolico" la coppia $(\mathbf{u}, \mathbf{v})$ è una base canonica, quindi del tutto equivalente alla coppia $(\mathbf{c}_0, \mathbf{c}_1)$, ottenibile da questa mediante una *rotazione*, cioè mediante un elemento di $SO(1,1)$. Siamo quindi condotti ad esaminare il gruppo delle isometrie sullo spazio iperbolico bidimensionale. Ci limitiamo a quelle di determinante +1, cioè alle rotazioni. Ricordiamo innanzitutto che le componenti (Q_i^j) di un qualunque endomorfismo Q sono legate alle componenti (Q_{ij}) della forma bilineare corrispondente dalle equazioni $Q_{ij} = g_{jk} Q_i^k$. Pertanto, essendo nel caso presente

$$g_{00} = -1, \quad g_{11} = 1, \quad g_{01} = g_{10} = 0,$$

risulta:

$$Q_{00} = -Q_0^0, \quad Q_{01} = Q_0^1, \quad Q_{10} = -Q_1^0, \quad Q_{11} = Q_1^1.$$

Posto inoltre

$$Q(\mathbf{c}_0) = a \mathbf{c}_0 + b \mathbf{c}_1, \quad Q(\mathbf{c}_1) = c \mathbf{c}_0 + d \mathbf{c}_1,$$

cioè

$$(Q_i^j) = \begin{pmatrix} a & b \\ c & d \end{pmatrix},$$

dalle condizioni $Q(\mathbf{c}_i) \cdot Q(\mathbf{c}_j) = \mathbf{c}_i \cdot \mathbf{c}_j$ e $\det(Q_i^j) = 1$, caratteristiche di una rotazione, seguono le uguaglianze:

$$a^2 - b^2 = 1, \quad d^2 - c^2 = 1, \quad bd - ac = 0, \quad ad - bc = 1.$$

Dalle ultime due si trae $a = d$ e $b = c$. Dalla prima segue $a^2 \geq 1$. Possiamo allora porre

$$(8) \quad a = \cosh \chi, \quad b = \sinh \chi,$$

oppure

$$(9) \quad a = -\cosh \chi, \quad b = \sinh \chi.$$

Poiché $a = -Q(\mathbf{c}_0) \cdot \mathbf{c}_0$, nel primo caso abbiamo $a > 0$, quindi una **rotazione ortocrona**, conservante cioè la classe di ogni vettore del

genere tempo. La matrice delle componenti di una rotazione ortocrona ha quindi la forma:

$$(10) \quad (Q_i^j) = \begin{pmatrix} \cosh \chi & \sinh \chi \\ \sinh \chi & \cosh \chi \end{pmatrix}.$$

Nel secondo caso si ha una rotazione non ortocrona e la matrice corrispondente è:

$$(11) \quad (Q_i^j) = \begin{pmatrix} -\cosh \chi & \sinh \chi \\ \sinh \chi & -\cosh \chi \end{pmatrix}.$$

Il parametro χ che interviene in queste due rappresentazioni prende il nome di **pseudoangolo della rotazione**. Esso può variare da $-\infty$ a $+\infty$, ricoprendo tutte le rotazioni ortocrone. Per $\chi = 0$ si ha ovviamente $Q = 1$. La composizione di due rotazioni di pseudoangolo χ_1 e χ_2 è la rotazione di pseudoangolo $\chi_1 + \chi_2$: lo si desume dalla (10) utilizzando le proprietà elementari delle funzioni trigonometriche iperboliche. Sempre in base a tali proprietà, si osserva che si può scegliere come parametro rappresentativo la tangente iperbolica dello pseudoangolo,

$$\beta = \tanh \chi = \frac{\sinh \chi}{\cosh \chi},$$

variabile nell'intervallo aperto $(-1, 1)$. La (10) diventa allora

$$(12) \quad (Q_i^j) = \frac{1}{\sqrt{1 - \beta^2}} \begin{pmatrix} 1 & \beta \\ \beta & 1 \end{pmatrix}.$$

CAPITOLO II

CALCOLO VETTORIALE
NEGLI SPAZI AFFINI

Gli spazi affini costituiscono il supporto di alcuni fondamentali modelli matematici della meccanica. Primo fra tutti lo spazio affine tridimensionale euclideo, che fornisce il modello dello spazio fisico osservato da un riferimento (almeno nelle teorie "classiche" o "euclidee"). Sono pure spazi affini, a quattro dimensioni, lo spazio-tempo della Meccanica Newtoniana e lo spazio-tempo della Relatività Ristretta.

Questo capitolo, dedicato al complesso degli enti che si possono definire sugli spazi affini e al calcolo che ne consegue (che chiamiamo brevemente "calcolo vettoriale"), si sviluppa in paragrafi strettamente correlati tra loro, raggruppabili in due parti. La prima parte (paragrafi da 1 a 10) è dedicata alle strutture fondamentali definibili in uno spazio affine (campi scalari, campi vettoriali, forme differenziali, curve, superfici) e ai sistemi dinamici, argomento centrale intorno al quale si sviluppa non solo questo capitolo ma l'intero corso. Come sarà visto in un capitolo successivo, tutte queste nozioni sono estendibili a spazi di natura più generale: le varietà differenziabili. La seconda parte è invece dedicata ad una miscellanea di argomenti e concetti per la maggior parte conseguenti all'istituzione sugli spazi affini (in particolare su quello tridimensionale) di una struttura euclidea. Essi si possono raccogliere in tre gruppi: (i) operatori differenziali su campi scalari e vettoriali (§11), (ii) curve e superfici nello spazio affine euclideo tridimensionale (§12), (iii) geometria dei sistemi di masse e di vettori (§13, §14, §15).

1. Richiami sugli spazi affini

Uno **spazio affine** è una terna $(\mathcal{A}, E, \delta)$ costituita da un insieme $\mathcal{A}$, i cui elementi sono detti **punti**, da uno spazio vettoriale E (reale e di dimensione finita) detto **soggiacente** o **associato** ad $\mathcal{A}$, e da un'applicazione $\delta: \mathcal{A} \times \mathcal{A} \rightarrow E$, soddisfacente alle due seguenti condizioni: (i) per ogni coppia $(P, v) \in \mathcal{A} \times E$ esiste un unico punto $Q \in \mathcal{A}$ tale che $\delta(P, Q) = v$; (ii) per ogni terna di punti (P, Q, R) vale l'uguaglianza

$$(1) \quad \delta(P, Q) + \delta(Q, R) = \delta(P, R).$$

Uno spazio affine $(\mathcal{A}, E, \delta)$ è indicato brevemente con $\mathcal{A}$, mentre il vettore $\delta(P, Q)$ è denotato semplicemente con PQ , o anche con $Q - P$, per cui l'uguaglianza (1) si scrive più semplicemente

$$PQ + QR = PR,$$

o anche

$$(Q - P) + (R - Q) = R - P,$$

assumendo in questo secondo caso l'aspetto di una identità algebrica. La **dimensione** di uno spazio affine è per definizione la dimensione dello spazio vettoriale associato.

Un vettore $v \in E$ è anche detto **vettore libero**. Una coppia $(P, v) \in \mathcal{A} \times E$ è detta **vettore applicato**. Un vettore applicato (P, v) si identifica con la coppia ordinata di punti (P, Q) dove Q è tale che $PQ = v$.

Un sottoinsieme $\mathcal{B} \subset \mathcal{A}$ è un **sottospazio affine** di dimensione m se l'immagine di $\mathcal{B} \times \mathcal{B}$ secondo δ è un sottospazio di E di dimensione m . In particolare, un sottospazio affine di dimensione 1 è detto **retta**.

Se si fissa un punto $O \in \mathcal{A}$, detto **origine** o **centro**, allora lo spazio $\mathcal{A}$ si identifica con lo spazio vettoriale associato E : infatti, in base agli assiomi (i) e (ii), la corrispondenza che ad ogni vettore $x \in E$ associa il punto $P \in \mathcal{A}$ tale che $x = OP$ è biunivoca. L'origine viene identificata col vettore nullo.

Viceversa, uno spazio vettoriale E può assumere la struttura di spazio affine; se infatti si pone $\delta(u, v) = v - u$, si verifica facilmente che la terna (E, E, δ) soddisfa alle condizioni richieste.

Se oltre ad un'origine O si fissa anche una base $(c_\alpha) = (c_1, \dots, c_n)$ di E , se cioè si fissa un **riferimento cartesiano** (O, c_α) , (gli indici greci $\alpha, \beta, \dots$ si intendono variabili da 1 a n , dove n è la dimensione

dello spazio affine) risulta definita una corrispondenza biunivoca fra lo spazio affine $\mathcal{A}$ e lo spazio $\mathbb{R}^n$: ad ogni punto $P \in \mathcal{A}$ si fa corrispondere l'ennupla reale (x^α) costituita dalle componenti secondo la base $(\mathbf{c}_\alpha)$ del vettore OP ,

$$(2) \quad OP = \mathbf{x} = x^\alpha \mathbf{c}_\alpha.$$

Risultano al contempo definite delle applicazioni $x^\alpha: \mathcal{A} \rightarrow \mathbb{R}$, dette **coordinate cartesiane** o **affini**, tali che $x^\alpha(P)\mathbf{c}_\alpha = OP$. Si noti che l'identificazione di $\mathcal{A}$ con $\mathbb{R}^n$ implica, fra l'altro, l'istituzione sullo spazio affine della topologia di $\mathbb{R}^n$.

Se si considerano due riferimenti cartesiani $(O, \mathbf{c}_\alpha)$ e $(O', \mathbf{c}_{\alpha'})$ (conviene, come vedremo, porre l'apice sull'indice) allora i due sistemi di coordinate (x^α) e $(x^{\alpha'})$ sono legati da **trasformazioni affini** cioè da relazioni del tipo

$$(3) \quad x^\alpha = a_{\alpha'}^\alpha x^{\alpha'} + b^\alpha, \quad x^{\alpha'} = a_\alpha^{\alpha'} x^\alpha + b^{\alpha'}.$$

Le matrici $(a_{\alpha'}^\alpha)$ e $(a_\alpha^{\alpha'})$, una inversa dell'altra, sono le matrici di trasformazione delle basi:

$$(4) \quad \mathbf{c}_{\alpha'} = a_{\alpha'}^\alpha \mathbf{c}_\alpha, \quad \mathbf{c}_\alpha = a_\alpha^{\alpha'} \mathbf{c}_{\alpha'},$$

mentre le (b^α) e $(b^{\alpha'})$ sono le componenti del vettore OO' e del vettore opposto $O'O$ nelle due basi:

$$(5) \quad OO' = b^\alpha \mathbf{c}_\alpha, \quad O'O = b^{\alpha'} \mathbf{c}_{\alpha'}.$$

Le (3) seguono infatti dalle uguaglianze $OP = OO' + O'P$ e $O'P = O'O + OP$, tenuto conto della (2), insieme all'analoga $O'P = x^{\alpha'} \mathbf{c}_{\alpha'}$, e delle (4) e (5).

Se lo spazio vettoriale E soggiacente ad uno spazio affine è dotato di un tensore metrico $\mathbf{g}$, vale a dire di un prodotto scalare, si dice che lo spazio affine è **euclideo**. Uno **spazio affine euclideo** è pertanto una quaterna $(\mathcal{A}, E, \delta, \mathbf{g})$ dove $(\mathcal{A}, E, \delta)$ è uno spazio affine e $\mathbf{g}$ è un tensore metrico sopra lo spazio vettoriale E . Si dice che uno spazio affine ha segnatura (p, q) se questa è la segnatura della forma bilineare simmetrica $\mathbf{g}$; in particolare lo spazio affine si dice **strettamente euclideo** se la segnatura è $(n, 0)$, cioè se il tensore metrico è definito positivo. Dato un riferimento affine $(O, \mathbf{c}_\alpha)$ dove la base è canonica, le corrispondenti coordinate (x^α) si dicono **coordinate canoniche** o **coordinate ortonormali**.

La presenza del tensore metrico consente di definire su di uno spazio affine le operazioni tipiche della **geometria euclidea**, quali la misura di distanze, di angoli, di lunghezze di curve, di aree di superfici, di volumi, e così via, nonché speciali operazioni su campi scalari, vettoriali e tensoriali. Alcune di queste operazioni traggono origine da analoghe operazioni definite sugli spazi vettoriali euclidei. Prima di esaminarle (§11) occorre però trattare concetti ed operazioni che prescindono dalla struttura euclidea.

NOTAZIONI. Per spazi affini di dimensione 2 e 3 denoteremo generici riferimenti cartesiani con $(O, \mathbf{i}, \mathbf{j})$ e $(O, \mathbf{i}, \mathbf{j}, \mathbf{k})$ e le coordinate associate con (x, y) e (x, y, z) . In spazi affini euclidei le basi $(\mathbf{i}, \mathbf{j})$ e $(\mathbf{i}, \mathbf{j}, \mathbf{k})$ si intenderanno, salvo esplicito avviso contrario, canoniche (ortonormali): i vettori sono di modulo unitario e mutuamente ortogonali. Precisata la base, per i vettori sarà talvolta conveniente l'uso della notazione matriciale; a seconda della convenzione adottata potranno essere rappresentati da vettori colonna o vettori riga. Le scritture

$$\mathbf{v} = v^\alpha \mathbf{c}_\alpha, \quad \mathbf{v} = \begin{pmatrix} v^1 \\ v^2 \\ \vdots \\ v^n \end{pmatrix}, \quad \mathbf{v} = (v^1, v^2, \dots, v^n)$$

saranno quindi equivalenti. Nel caso $n = 3$ si avranno in particolare scritture del tipo

$$\mathbf{v} = a\mathbf{i} + b\mathbf{j} + c\mathbf{k}, \quad \mathbf{v} = \begin{pmatrix} a \\ b \\ c \end{pmatrix}, \quad \mathbf{v} = (a, b, c).$$

2. Campi scalari e vettoriali

Una funzione reale sopra uno spazio affine $f: \mathcal{A} \rightarrow \mathbb{R}$ è anche detta **campo scalare**. Assegnate delle coordinate cartesiane, ogni campo scalare $f: \mathcal{A} \rightarrow \mathbb{R}$ è rappresentato da una funzione reale a n variabili reali denotata con $f(x^1, x^2, \dots, x^n)$ o brevemente con $f(x^\alpha)$. Un campo scalare f è detto di **classe** C^k se la sua funzione rappresentativa $f(x^\alpha)$ ammette derivate parziali continue fino all'ordine k . Diciamo che f è

di classe C^∞ o differenziabile se la $f(x^\alpha)$ ammette derivate parziali continue di qualunque ordine.

Sull'insieme di tutti i campi scalari sopra uno spazio affine si introducono le seguenti tre operazioni fondamentali: la *somma*, definita da $(f + g)(P) = f(P) + g(P)$, il *prodotto per un numero reale* $a \in \mathbb{R}$, definito da $(af)(P) = a(f(P))$, il *prodotto numerico*, definito da $(fg)(P) = f(P)g(P)$. Con queste operazioni l'insieme assume la struttura di anello commutativo e di algebra associativa e commutativa. Sovente restringeremo la nostra attenzione, per semplicità, all'anello dei campi scalari di classe C^∞ , che denoteremo con $C^\infty(\mathcal{A}, \mathbb{R})$, o brevemente con $\mathcal{F}(\mathcal{A})$. Denoteremo in particolare con $C^\infty(U, \mathbb{R})$, o brevemente con $\mathcal{F}(U)$, l'anello dei campi scalari di classe C^∞ sopra un aperto $U \subset \mathcal{A}$.

Un **campo vettoriale** è un'applicazione $\mathbf{X}: \mathcal{A} \rightarrow \mathcal{A} \times E$ che associa ad ogni punto $P \in \mathcal{A}$ un vettore $\mathbf{X}(P) \in E$ applicato in P . Fissato un riferimento cartesiano, ogni campo vettoriale $\mathbf{X}$ risulta rappresentato da un insieme di n funzioni reali $X^\alpha: \mathcal{A} \rightarrow \mathbb{R}$, dette **componenti cartesiane**, tali da aversi per ogni punto $P \in \mathcal{A}$

$$\mathbf{X}(P) = X^\alpha(P) \mathbf{c}_\alpha.$$

Scriveremo più semplicemente

$$\mathbf{X} = X^\alpha \mathbf{c}_\alpha$$

Le componenti X^α del campo, essendo dei campi scalari, sono a loro volta rappresentabili da funzioni $X^\alpha(x^\beta)$ nelle n coordinate cartesiane (x^β) . Un campo vettoriale si dice di classe C^k (rispettivamente, di classe C^∞ o differenziabile) se tutte le sue componenti sono di classe C^k (rispettivamente, di classe C^∞).

OSSERVAZIONE 1. Nel piano affine riferito a coordinate cartesiane $(x^\alpha) = (x, y)$ ogni campo vettoriale $\mathbf{X}$ è definito da due funzioni reali $X(x, y)$ e $Y(x, y)$ nelle due variabili (x, y) . Un campo vettoriale ammetterà dunque una rappresentazione del tipo (usiamo qui la rappresentazione *colonna*)

$$\mathbf{X} = X(x, y) \mathbf{i} + Y(x, y) \mathbf{j} = \begin{pmatrix} X(x, y) \\ Y(x, y) \end{pmatrix},$$

dove $(\mathbf{i}, \mathbf{j})$ sono i vettori del riferimento. Analogamente, in uno spazio affine tridimensionale,

$$\mathbf{X} = X(x, y, z) \mathbf{i} + Y(x, y, z) \mathbf{j} + Z(x, y, z) \mathbf{k} = \begin{pmatrix} X(x, y, z) \\ Y(x, y, z) \\ Z(x, y, z) \end{pmatrix}.$$

Sull'insieme di tutti i campi vettoriali sopra uno spazio affine si introducono le seguenti due operazioni fondamentali: la *somma*, definita da $(\mathbf{X} + \mathbf{Y})(P) = \mathbf{X}(P) + \mathbf{Y}(P)$, il *prodotto per un numero reale* $a \in \mathbb{R}$, definito da $(a\mathbf{X})(P) = a(\mathbf{X}(P))$. Con queste operazioni l'insieme assume la struttura di modulo su $\mathcal{F}(\mathcal{A})$ e di spazio vettoriale su $\mathbb{R}$ (di dimensione infinita). Denoteremo con $\mathcal{X}^k(\mathcal{A})$ lo spazio dei campi vettoriali di classe C^k su tutto lo spazio affine ed in particolare con $\mathcal{X}^k(U)$ lo spazio dei campi vettoriali di classe C^k sopra un aperto $U \subset \mathcal{A}$. Denoteremo semplicemente con $\mathcal{X}(\mathcal{A})$ o con $\mathcal{X}(U)$ gli spazi dei campi vettoriali di classe C^∞ .

OSSERVAZIONE 2. Oltre a campi scalari e vettoriali, sopra uno spazio affine si possono definire **campi tensoriali**. Un campo tensoriale di tipo (p, q) è un'applicazione $\mathbf{K}: \mathcal{A} \rightarrow \mathcal{A} \times T_q^p(E)$ che associa ad ogni punto $P \in \mathcal{A}$ un tensore $\mathbf{K}(P) \in T_q^p(E)$ applicato in P . Le componenti cartesiane di un campo tensoriale si definiscono in maniera del tutto analoga a quelle dei campi vettoriali. Campi tensoriali sono per esempio le *forme differenziali*, che esamineremo più avanti.

Esaminiamo una prima fondamentale operazione coinvolgente campi scalari e vettoriali.

DEFINIZIONE 1. Diciamo **derivata** di un campo scalare f rispetto ad un campo vettoriale $\mathbf{X}$ il campo scalare $\mathbf{X}f$ definito da

$$(1) \quad \mathbf{X}f = X^\alpha \frac{\partial f}{\partial x^\alpha}$$

dove con

$$\frac{\partial f}{\partial x^\alpha}$$

s'intende la derivata parziale rispetto alla x^α della funzione rappresentativa $f(x^1, \dots, x^n)$ in un qualunque sistema di coordinate cartesiane.

OSSERVAZIONE 3. Questa definizione ha significato intrinseco (o *assoluto*), non dipende cioè dalla scelta delle coordinate cartesiane. Si osserva infatti che, rappresentato il campo secondo due basi $(\mathbf{c}_\alpha)$ e $(\mathbf{c}_{\alpha'})$,

$$\mathbf{X} = X^\alpha \mathbf{c}_\alpha = X^{\alpha'} \mathbf{c}_{\alpha'},$$

e tenuto conto delle relazioni del tipo (4) del §1 tra le due basi, tra le componenti del campo sussiste il legame

$$(2) \quad X^\alpha = a_{\alpha'}^\alpha X^{\alpha'}.$$

D'altra parte, interpretata la f come funzione delle (x^α) per il tramite delle $(x^{\alpha'})$, dalle relazioni (3) del §1 segue che

$$\frac{\partial f}{\partial x^\alpha} = \frac{\partial f}{\partial x^{\alpha'}} \frac{\partial x^{\alpha'}}{\partial x^\alpha} = \frac{\partial f}{\partial x^{\alpha'}} a_{\alpha'}^\alpha.$$

Si ha quindi:

$$X^\alpha \frac{\partial f}{\partial x^\alpha} = X^\alpha \frac{\partial f}{\partial x^{\alpha'}} a_{\alpha'}^\alpha = X^{\alpha'} \frac{\partial f}{\partial x^{\alpha'}}.$$

Ciò mostra l'indipendenza della definizione (1) dalla scelta delle coordinate cartesiane.

ESERCIZIO 1. Per ben comprendere il significato dell'osservazione precedente si provi a considerare, per esempio, l'operazione di "derivazione seconda" definita da

$$X^\alpha \frac{\partial^2 f}{\partial x^{\alpha^2}}$$

e si verifichi che questa dipende dalla scelta delle coordinate cartesiane (x^α) .

OSSERVAZIONE 4. Con l'operazione di derivazione così definita ogni campo vettoriale determina un'applicazione $X: \mathcal{F}(\mathcal{A}) \rightarrow \mathcal{F}(\mathcal{A})$ (si usa ancora lo stesso simbolo X) che è *lineare*

$$(3) \quad X(af + bg) = aXf + bXg \quad (a, b \in \mathbb{R}),$$

e soddisfa alla **regola di Leibniz**:

$$(4) \quad X(fg) = gXf + fXg.$$

Si può dimostrare che, viceversa, ogni operatore sulle funzioni soddisfacente a queste due condizioni si identifica con un campo vettoriale. Come sarà giustificato più avanti, useremo anche la notazione

$$\langle X, df \rangle$$

al posto di Xf .

OSSERVAZIONE 5. Si può anche definire la derivata di un campo scalare f rispetto ad un vettore applicato (P, v) . È il numero reale

$$(5) \quad (P, v)f = v^\alpha \left(\frac{\partial f}{\partial x^\alpha} \right)_P,$$

dove le (v^α) sono le componenti del vettore v e la derivata parziale è calcolata nel punto P .

OSSERVAZIONE 6. L'interpretazione dei campi vettoriali come operatori di derivazione consente di definire il **commutatore** di due campi vettoriali X e Y . È il campo vettoriale $[X, Y]$ definito da

$$(6) \quad [X, Y]f = X(Yf) - Y(Xf), \quad f \in \mathcal{F}(\mathcal{A})$$

Si istituisce così sullo spazio $\mathcal{X}(\mathcal{A})$ dei campi vettoriali un'operazione binaria interna $[\cdot, \cdot]$ detta anche **parentesi di Lie**, *anticommutativa*, *bilineare* e per la quale vale l'**identità di Jacobi**:

$$(7) \quad \begin{cases} [X, Y] = -[Y, X], \\ [aX + bY, Z] = a[X, Z] + b[Y, Z], \\ [X, [Y, Z]] + [Y, [Z, X]] + [Z, [X, Y]] = 0. \end{cases}$$

Con tale operazione lo spazio vettoriale $\mathcal{X}(\mathcal{A})$ assume pertanto la struttura di **algebra di Lie**.

ESERCIZIO 2. Si verifichi che le componenti cartesiane del commutatore sono date da:

$$(8) \quad [X, Y]^\alpha = X^\beta \partial_\beta Y^\alpha - Y^\beta \partial_\beta X^\alpha, \quad \left(\partial_\beta = \frac{\partial}{\partial x^\beta} \right).$$

Di qui si osservi che la (6) definisce effettivamente un campo vettoriale, cioè che per il commutatore vale la linearità e la regola di Leibniz. Si verifichi inoltre che vale l'identità di Jacobi (7)₃.

OSSERVAZIONE 7. Ad ogni campo vettoriale X si può associare un campo scalare detto la **divergenza** di X , denotato con $\text{div}(X)$ e definito da:

$$(9) \quad \text{div}(X) = \frac{\partial X^1}{\partial x^1} + \dots + \frac{\partial X^n}{\partial x^n} = \partial_\alpha X^\alpha$$

L'applicazione

$$\operatorname{div}: \mathcal{X}(\mathcal{A}) \rightarrow \mathcal{F}(\mathcal{A}): \mathbf{X} \mapsto \operatorname{div}(\mathbf{X})$$

così definita gode delle seguenti *proprietà caratteristiche*:

$$(10) \quad \begin{cases} \operatorname{div}(\mathbf{X} + \mathbf{Y}) = \operatorname{div}(\mathbf{X}) + \operatorname{div}(\mathbf{Y}), \\ \operatorname{div}(f\mathbf{X}) = f \operatorname{div}(\mathbf{X}) + \mathbf{X}f \quad (f \in \mathcal{F}(\mathcal{A})), \\ \mathbf{X} = \text{cost.} \implies \operatorname{div}(\mathbf{X}) = 0. \end{cases}$$

ESERCIZIO 3. Dimostrare che la definizione (9) non dipende dalla scelta delle coordinate cartesiane (si veda l'Oss. 3 e l'Eserc. 1). Dimostrare che se un operatore da campi vettoriali a campi scalari soddisfa alle (10) allora esso è necessariamente la divergenza.

3. Coordinate non affini e riferimenti associati

Sebbene le coordinate cartesiane abbiano carattere "privilegiato" perchè direttamente connesse con la struttura affine, è talvolta conveniente, in particolari circostanze, far uso di altri tipi di coordinate. La nozione generale di **sistema di coordinate** è strettamente connessa con la nozione di **carta** sopra un insieme, come mostrato dalla seguente definizione.

DEFINIZIONE 1. Sia $\mathcal{A}$ un insieme (quindi, non necessariamente uno spazio affine). Diciamo **carta** di dimensione n sull'insieme $\mathcal{A}$ un'applicazione biettiva $\varphi: U \rightarrow \mathbb{R}^n$ di un sottoinsieme $U \subset \mathcal{A}$, detto **dominio** della carta, in un aperto $\varphi(U)$ di $\mathbb{R}^n$. Si chiamano **coordinate** associate alla carta φ le n funzioni $q^i: U \rightarrow \mathbb{R}$ definite da $q^i = pr_i \circ \varphi$, dove $pr_i: \mathbb{R}^n \rightarrow \mathbb{R}$ è la proiezione i -esima, l'applicazione cioè che ad ogni n -pla reale $(r^1, r^2, \dots, r^n)$ associa l'elemento i -esimo r^i .

Se $\mathcal{A}$ è in particolare uno spazio affine, fissate delle coordinate cartesiane (x^α) su $\mathcal{A}$ e assegnata una carta $\varphi: U \rightarrow \mathbb{R}^n$ ⁽¹⁾, le coordinate (q^i) a questa associate, essendo delle funzioni reali sul dominio U , sono

⁽¹⁾ Senza entrare nei dettagli di una discussione più approfondita, accettiamo il fatto che la dimensione della carta abbia la stessa dimensione n dello spazio affine.

rappresentabili, come si è osservato al paragrafo precedente, da funzioni nelle n variabili reali (x^α) :

$$(1) \quad q^i = q^i(x^\alpha).$$

Siccome l'applicazione φ è biettiva, questo sistema di n funzioni è invertibile; le (x^α) possono cioè esprimersi in funzione delle (q^i) :

$$(2) \quad x^\alpha = x^\alpha(q^i).$$

Le relazioni (1) o (2) vanno sotto il nome di **cambiamenti o trasformazioni di coordinate**. Diciamo che le coordinate (q^i) sono di classe C^∞ se tali sono tutte le funzioni (1) e (2).

Alle relazioni invertibili (1) e (2) sono associate due matrici $n \times n$ di funzioni, date dalle derivate parziali

$$(3) \quad E_\alpha^i = \frac{\partial q^i}{\partial x^\alpha}, \quad E_i^\alpha = \frac{\partial x^\alpha}{\partial q^i},$$

dette **matrici jacobiane** della trasformazione. L'invertibilità del sistema di funzioni (1) implica, per un fondamentale teorema di Analisi (il **teorema della funzione inversa**), che queste matrici sono regolari e l'una inversa dell'altra.

ESEMPIO 1. *Coordinate polari del piano*. Nello spazio affine bi-dimensionale (piano affine), riferito a coordinate cartesiane $(x^1, x^2) = (x, y)$, le equazioni

$$(5) \quad \begin{cases} x = r \cos \vartheta, \\ y = r \sin \vartheta, \end{cases}$$

definiscono nuove coordinate $(q^1, q^2) = (r, \vartheta)$, dette **coordinate polari**. Inteso $r > 0$, si deve scegliere il campo di variabilità dell'angolo ϑ . Se si pone per esempio $-\pi < \vartheta < \pi$ allora queste coordinate descrivono il dominio aperto U del piano ottenuto asportando l'origine e tutto il semiasse negativo delle x . Le (5), che corrispondono alle equazioni generali (2), sono invertibili (ad ogni punto del dominio U corrisponde infatti un ben determinato valore delle coordinate (r, ϑ)). Le relazioni inverse (corrispondenti alle (1)) possono scriversi in vari modi, utilizzando le

funzioni trigonometriche inverse. Per esempio:

$$(4) \quad \begin{cases} r = \sqrt{x^2 + y^2}, \\ \vartheta = \begin{cases} \arcsin \frac{y}{\sqrt{x^2 + y^2}}, & x \geq 0, \\ \pi - \arcsin \frac{y}{\sqrt{x^2 + y^2}}, & x < 0, y > 0, \\ -\pi - \arcsin \frac{y}{\sqrt{x^2 + y^2}}, & x < 0, y < 0. \end{cases} \end{cases}$$

intendendo la funzione arcsin a valori nell'intervallo chiuso $[-\frac{\pi}{2}, \frac{\pi}{2}]$.

ESEMPIO 2. Coordinate cilindriche dello spazio. Nello spazio affine tridimensionale, riferito a coordinate cartesiane $(x^1, x^2, x^3) = (x, y, z)$, le stesse equazioni (4), insieme alla $q^3 = z$, definiscono nuove coordinate $(q^1, q^2, q^3) = (r, \vartheta, z)$, dette **coordinate cilindriche**.

ESEMPIO 3. Coordinate polari dello spazio o sferiche. Nello spazio affine tridimensionale, riferito a coordinate cartesiane (x, y, z) , le equazioni

$$(6) \quad \begin{cases} x = r \sin \varphi \cos \vartheta, \\ y = r \sin \varphi \sin \vartheta, \\ z = r \cos \varphi. \end{cases}$$

definiscono implicitamente nuove coordinate $(q^1, q^2, q^3) = (r, \varphi, \vartheta)$ sopra il dominio aperto U ottenuto asportando dallo spazio il semiasse positivo delle x e tutto l'asse z ed i cui valori coprono l'aperto $\{(r, \varphi, \vartheta) \in \mathbb{R}; r > 0, 0 < \varphi < \pi, 0 < \vartheta < 2\pi\}$. Esse prendono rispettivamente il nome di **raggio**, **colatitudine** e **longitudine**, e nell'insieme di **coordinate polari sferiche**.

Le trasformazioni del tipo (2) possono sintetizzarsi nella scrittura vettoriale

$$(7) \quad \boxed{OP = \mathbf{x}(q^i)}$$

dove $\mathbf{x}$ è il vettore posizione rispetto all'origine O del generico punto P , le cui componenti rispetto alla base $(\mathbf{c}_\alpha)$ sono appunto le (x^α) . Possiamo allora considerare le derivate parziali della funzione vettoriale (7),

$$(8) \quad \boxed{\mathbf{E}_i = \frac{\partial \mathbf{x}}{\partial q^i}}$$

Si tratta di n campi vettoriali sul dominio U delle coordinate (q^i) . Le componenti di questi campi secondo la base fondamentale $(\mathbf{c}_\alpha)$ formano la seconda delle matrici jacobiane (3):

$$(9) \quad \mathbf{E}_i = E_i^\alpha \mathbf{c}_\alpha.$$

La (9) mostra che i campi $(\mathbf{E}_i)$ sono ottenuti dalla base $(\mathbf{c}_\alpha)$ mediante una trasformazione lineare coinvolgente la matrice jacobiana (E_i^α) che è ovunque regolare, cioè invertibile. I vettori $(\mathbf{E}_i)$ sono pertanto indipendenti in ogni punto del dominio U . Si dice che essi costituiscono il **riferimento naturale** associato alle coordinate (q^i) .

Questo riferimento può essere utilizzato per rappresentare campi vettoriali. Se infatti $\mathbf{X}$ è un campo vettoriale definito in un insieme contenente il dominio U della carta, in U può essere rappresentato da una combinazione lineare del tipo

$$(10) \quad \boxed{\mathbf{X} = X^i \mathbf{E}_i}$$

Le funzioni (X^i) prendono il nome di **componenti secondo le coordinate** (q^i) del campo $\mathbf{X}$. Il legame tra queste e quelle cartesiane si ottiene combinando la (9) con la (10):

$$(11) \quad X^\alpha = E_i^\alpha X^i \iff X^i = E_\alpha^i X^\alpha.$$

Se si interpretano i campi vettoriali $(\mathbf{E}_i)$ come derivazioni, allora

$$(12) \quad \boxed{\mathbf{E}_i f = \frac{\partial f}{\partial q^i}}$$

intendendo con

$$\frac{\partial f}{\partial q^i}$$

la derivata parziale della funzione rappresentativa $f(q^1, \dots, q^n)$ rispetto alla coordinata q^i . Infatti, tenuto conto della (9) e della (3), si ha successivamente:

$$\mathbf{E}_i f = E_i^\alpha \frac{\partial f}{\partial x^\alpha} = \frac{\partial x^\alpha}{\partial q^i} \frac{\partial f}{\partial x^\alpha} = \frac{\partial f}{\partial q^i}.$$

Una conseguenza immediata di questa proprietà è che, stante la (10), la derivata di un campo scalare rispetto ad un campo vettoriale risulta espressa, qualunque sia il sistema di coordinate scelto, dalla formula

$$(13) \quad \boxed{Xf = X^i \frac{\partial f}{\partial q^i}}$$

dove le (X^i) sono le componenti del campo vettoriale in quelle coordinate. È bene osservare che tali componenti coincidono con le derivate rispetto al campo X delle coordinate (q^i) , interpretate come funzioni reali sul dominio U . Vale cioè la formula

$$(14) \quad \boxed{X^i = Xq^i}$$

che si ottiene dalla stessa (13) ponendo $f = q^i$.

OSSERVAZIONE 1. I campi vettoriali E_i sono tangenti alle rispettive **curve coordinate**. Queste curve sono il luogo dei punti caratterizzati da valori costanti per tutte le coordinate meno una. Per esempio, se si fissano i valori delle coordinate $(q^2, \dots, q^n)$ si ottiene una curva parametrizzata dalla coordinata q^1 ; il campo E_1 è tangente a tutte le curve ottenute in questo modo. Questi concetti saranno ripresi più avanti.

OSSERVAZIONE 2. Il riferimento naturale associato a coordinate cartesiane (x^α) è proprio il sistema dei campi vettoriali costanti (c_α) .

OSSERVAZIONE 3. Siccome i campi vettoriali (E_i) possono rappresentarsi come funzioni delle coordinate (q^i) , ha senso considerarne le derivate parziali

$$\frac{\partial E_i}{\partial q^j} = \frac{\partial^2 x}{\partial q^j \partial q^i}.$$

Essendo queste derivate parziali ancora dei campi vettoriali, è naturale considerarne la rappresentazione secondo il riferimento (E_i) . Denotiamo genericamente con Γ_{ji}^k le loro componenti; poniamo cioè

$$(15) \quad \boxed{\partial_j E_i = \partial_j \partial_i x = \Gamma_{ji}^k E_k}$$

adottando, per comodità di scrittura, la notazione abbreviata

$$\partial_i = \frac{\partial}{\partial q^i}.$$

Le Γ_{ji}^k sono in effetti delle funzioni sopra il dominio U della carta. Esse prendono il nome di **simboli di Christoffel** (Elwin Bruno Christoffel, 1829-1900) associati alle coordinate (q^i) . Questi simboli, stante la presenza nella (15) di una derivata parziale seconda, sono simmetrici rispetto agli indici in basso:

$$(16) \quad \Gamma_{ji}^k = \Gamma_{ij}^k.$$

OSSERVAZIONE 4. I simboli di Christoffel sono identicamente nulli se e solo se le coordinate sono cartesiane. Dalla definizione (15) si vede infatti che è identicamente $\Gamma_{ji}^k = 0$ se e solo se i campi E_i sono costanti. Ciò accade se e solo se le coordinate sono cartesiane (Oss. 3).

ESERCIZIO 1. Si dimostri che le componenti del commutatore di due campi vettoriali sono, qualunque siano le coordinate scelte,

$$(17) \quad \boxed{[X, Y]^i = X^j \partial_j Y^i - Y^j \partial_j X^i}$$

ESERCIZIO 2. Utilizzando le proprietà caratteristiche della divergenza (Oss. 7, §2) si dimostri che, qualunque siano le coordinate, essa è espressa da

$$(18) \quad \boxed{\operatorname{div}(X) = \partial_i X^i + \Gamma_{ji}^j X^i}$$

ESERCIZIO 3. Calcolo dei vettori (E_i) e dei simboli di Christoffel per le coordinate polari del piano. Essendo $(q^1, q^2) = (r, \vartheta)$, è più comodo usare la notazione (E_r, E_ϑ) al posto di (E_1, E_2) . Per le (5) è

$$x = r(\cos \vartheta i + \sin \vartheta j).$$

Applicando la (8) si trova che

$$(19) \quad \begin{cases} E_r = \cos \vartheta i + \sin \vartheta j = u, \\ E_\vartheta = r(-\sin \vartheta i + \cos \vartheta j) = r\tau, \end{cases}$$

dove u è il **versore radiale**, cioè il versore di OP , e τ il **versore trasverso**, ortogonale ad u nel verso di ϑ crescente. È sottintesa, per questa

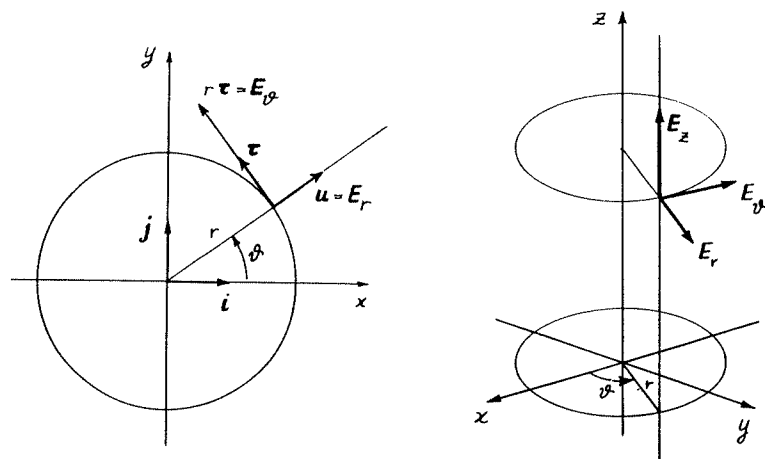


Fig. 3.1: coordinate polari e cilindriche.

definizione, la presenza di una struttura euclidea tale che i versori (i, j) formino una base canonica. Secondo una tale struttura la coordinata r rappresenta la distanza dall'origine O (Fig. 3.1). Derivando ancora rispetto alle coordinate, si trova

$$\begin{cases} \partial_r \mathbf{E}_r = 0, \\ \partial_r \mathbf{E}_\vartheta = \partial_\vartheta \mathbf{E}_r = -\sin \vartheta \mathbf{i} + \cos \vartheta \mathbf{j} = \boldsymbol{\tau}, \\ \partial_\vartheta \mathbf{E}_\vartheta = -r(\cos \vartheta \mathbf{i} + \sin \vartheta \mathbf{j}) = -r \mathbf{u}. \end{cases}$$

Dal confronto con le (19) segue ancora

$$\partial_r \mathbf{E}_\vartheta = \boldsymbol{\tau} = \frac{1}{r} \mathbf{E}_\vartheta, \quad \partial_\vartheta \mathbf{E}_\vartheta = -r \mathbf{u} = -r \mathbf{E}_r.$$

Quindi dal confronto con la definizione (15) segue che i simboli di Christoffel non nulli sono soltanto

$$(20) \quad \Gamma_{22}^1 = -r, \quad \Gamma_{12}^2 = \Gamma_{21}^2 = \frac{1}{r}.$$

ESERCIZIO 4. Calcolo dei vettori $(\mathbf{E}_i)$ e dei simboli di Christoffel per le coordinate cilindriche. Si ripete quanto visto per le coordinate polari, osservando che in più si ha $\mathbf{E}_z = \mathbf{k}$.

ESERCIZIO 5. Calcolo dei vettori $(\mathbf{E}_i)$ e dei simboli di Christoffel per le coordinate polari sferiche. Per le (6) è

$$\mathbf{x} = r(\sin \varphi \cos \vartheta \mathbf{i} + \sin \varphi \sin \vartheta \mathbf{j} + \cos \varphi \mathbf{k}),$$

e quindi (Fig. 3.2):

$$(21) \quad \begin{cases} \mathbf{E}_r = \sin \varphi \cos \vartheta \mathbf{i} + \sin \varphi \sin \vartheta \mathbf{j} + \cos \varphi \mathbf{k} = \mathbf{u}, \\ \mathbf{E}_\varphi = r(\cos \varphi \cos \vartheta \mathbf{i} + \cos \varphi \sin \vartheta \mathbf{j} - \sin \varphi \mathbf{k}), \\ \mathbf{E}_\vartheta = r(-\sin \varphi \sin \vartheta \mathbf{i} + \sin \varphi \cos \vartheta \mathbf{j}). \end{cases}$$

Derivando ancora una volta rispetto alle coordinate, tenuto conto delle (21) e della (15), si trova che i simboli di Christoffel non nulli sono:

$$(22) \quad \begin{cases} \Gamma_{22}^1 = -r, & \Gamma_{12}^2 = \Gamma_{21}^2 = \Gamma_{31}^3 = \Gamma_{13}^3 = \frac{1}{r}, \\ \Gamma_{33}^1 = -r \sin^2 \varphi, & \Gamma_{33}^2 = -\sin \varphi \cos \varphi, & \Gamma_{23}^3 = \Gamma_{32}^3 = \cot \varphi. \end{cases}$$

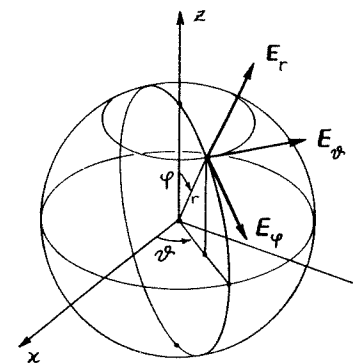


Fig. 3.2: coordinate polari sferiche.

4. Forme differenziali

DEFINIZIONE 1. Sia $\mathcal{A}$ uno spazio affine, $\mathcal{F}(\mathcal{A})$ l'anello dei campi scalari (funzioni reali) di classe C^∞ su $\mathcal{A}$ e $\mathcal{X}(\mathcal{A})$ lo spazio dei campi vettoriali C^∞ su $\mathcal{A}$. Una **forma lineare** o **1-forma** sopra uno spazio affine $\mathcal{A}$ è un'applicazione $\mathcal{F}(\mathcal{A})$ -lineare $\varphi: \mathcal{X}(\mathcal{A}) \rightarrow \mathcal{F}(\mathcal{A})$, cioè tale che

$$(1) \quad \varphi(fX + gY) = f\varphi(X) + g\varphi(Y)$$

per ogni scelta delle funzioni $f, g \in \mathcal{F}(\mathcal{A})$ e dei campi vettoriali $X, Y \in \mathcal{X}(\mathcal{A})$.

L'insieme delle 1-forme sopra uno spazio affine $\mathcal{A}$, che denotiamo con $\Phi^1(\mathcal{A})$, è un modulo sull'anello $\mathcal{F}(\mathcal{A})$ e uno spazio vettoriale su $\mathbb{R}$. La somma di due forme lineari e il prodotto per una funzione sono definiti da:

$$(\varphi + \psi)(X) = \varphi(X) + \psi(X), \quad (f\varphi)(X) = f \cdot \varphi(X).$$

Denotiamo con

$$\langle X, \varphi \rangle,$$

anziché con $\varphi(X)$ il valore della forma lineare φ sul campo vettoriale X . Si definisce in tal modo un'applicazione bilineare

$$\langle \cdot, \cdot \rangle: \mathcal{X}(\mathcal{A}) \times \Phi^1(\mathcal{A}) \rightarrow \mathcal{F}(\mathcal{A}),$$

che prende il nome di **valutazione** tra una forma lineare e un campo vettoriale.

OSSERVAZIONE 1. Una forma lineare può anche essere interpretata come **campo di covettori** cioè come applicazione

$$\varphi: \mathcal{A} \rightarrow \mathcal{A} \times E^*$$

che associa ad ogni punto $P \in \mathcal{A}$ un covettore applicato in P . Il collegamento tra questa e la precedente definizione è dato dalla formula

$$\langle X, \varphi \rangle(P) = \langle X(P), \varphi(P) \rangle.$$

Ha allora senso valutare una 1-forma φ su di un vettore applicato (P, v) . Il risultato $\langle v, \varphi(P) \rangle$ è un numero reale.

Un esempio fondamentale di forma lineare è il **differenziale** df di un **campo scalare** f . È definito dall'uguaglianza

$$(2) \quad \langle X, df \rangle = Xf$$

La linearità dell'applicazione $df: \mathcal{X}(\mathcal{A}) \rightarrow \mathcal{F}(\mathcal{A}): X \mapsto \langle X, df \rangle$ segue dal fatto che la derivata Xf , tenuto fisso il campo scalare f , è lineare rispetto al campo vettoriale X (si veda la definizione (1) del §2). Dalla regola di Leibniz per la derivata rispetto ad un campo vettoriale (formula (4) del §2) segue la regola di Leibniz per il differenziale:

$$(3) \quad d(fg) = gdf + f dg$$

Siano (q^i) generiche coordinate sopra un aperto U . Siccome si tratta di funzioni reali sopra U ha senso considerarne i differenziali (dq^i) . Per quanto visto al paragrafo precedente, essi fanno corrispondere ad un campo X le sue componenti (X^i) :

$$(4) \quad \langle X, dq^i \rangle = X^i$$

e quindi in particolare

$$(5) \quad \langle E_k, dq^i \rangle = \delta_k^i$$

Nel dominio U delle coordinate ogni forma lineare è combinazione lineare dei differenziali delle coordinate, ammette cioè una **rappresentazione locale** del tipo

$$(6) \quad \varphi = \varphi_i dq^i$$

dove le (φ_i) sono funzioni reali su U , dette **componenti** di φ nelle coordinate (q^i) , definite da

$$(7) \quad \varphi_i = \langle E_i, \varphi \rangle$$

Infatti, valendo la (6), si ha successivamente:

$$\langle E_k, \varphi \rangle = \varphi_i \langle E_k, dq^i \rangle = \varphi_i \delta_k^i = \varphi_k,$$

da cui segue la (7). Viceversa, se si considerano le funzioni definite dalla (7) e quindi la combinazione lineare (6), si ha per ogni campo $\mathbf{X}$:

$$\langle \mathbf{X}, \varphi_i dq^i \rangle = \varphi_i \langle \mathbf{X}, dq^i \rangle = \langle \mathbf{E}_i, \varphi \rangle X^i = \langle X^i \mathbf{E}_i, \varphi \rangle = \langle \mathbf{X}, \varphi \rangle.$$

Dunque la combinazione lineare $\varphi_i dq^i$ coincide proprio con φ e la (6) è dimostrata.

Dalle formule precedenti segue che la valutazione di una forma lineare sopra un campo vettoriale è data, qualunque siano le coordinate scelte, dalla somma dei prodotti delle componenti omologhe:

$$(8) \quad \boxed{\langle \mathbf{X}, \varphi \rangle = X^i \varphi_i}$$

OSSERVAZIONE 2. Le formule precedenti sono del tutto analoghe alle formule concernenti le forme lineari su di uno spazio vettoriale (Cap. I, §3): in luogo di una base ($\mathbf{e}_i$) abbiamo i campi vettoriali ($\mathbf{E}_i$) del riferimento associato alle coordinate ed in luogo della base duale (ϵ^i) abbiamo i differenziali (dq^i) delle coordinate; qui le componenti (φ_i) sono, anziché dei numeri, delle funzioni.

OSSERVAZIONE 3. Particolarizzando la (6) al caso del differenziale di una funzione si ottiene la formula

$$(9) \quad \boxed{df = \frac{\partial f}{\partial q^i} dq^i}$$

la quale mostra che le componenti del differenziale di una funzione sono le derivate parziali rispetto alle coordinate. Si ha infatti, per ogni campo vettoriale,

$$\langle \mathbf{X}, \varphi \rangle = X^i \frac{\partial f}{\partial q^i} = \langle \mathbf{X}, dq^i \rangle \frac{\partial f}{\partial q^i} = \langle \mathbf{X}, \frac{\partial f}{\partial q^i} dq^i \rangle.$$

Stante la rappresentazione (6), le forme lineari sono anche chiamate **forme differenziali lineari**. Si definiscono forme differenziali di grado superiore al primo, e sull'insieme di queste forme si istituisce un calcolo di grande utilità: il **calcolo differenziale esterno**.

DEFINIZIONE 2. Una **forma differenziale** o **forma esterna di grado p** o **p -forma** sopra uno spazio affine $\mathcal{A}$ è un'applicazione p -lineare antisimmetrica dello spazio $\mathcal{X}(\mathcal{A})^p$ nello spazio $\mathcal{F}(\mathcal{A})$:

$$\varphi: \underbrace{\mathcal{X}(\mathcal{A}) \times \mathcal{X}(\mathcal{A}) \times \dots \times \mathcal{X}(\mathcal{A})}_{p \text{ volte}} \rightarrow \mathcal{F}(\mathcal{A}).$$

Ricordiamo che l'antisimmetria equivale alla proprietà seguente: la funzione $\varphi(\mathbf{X}_1, \mathbf{X}_2, \dots, \mathbf{X}_p)$ cambia di segno se si scambiano tra loro due qualsiasi argomenti. Denotiamo con $\Phi^p(\mathcal{A})$ lo spazio delle p -forme sopra $\mathcal{A}$. Si tratta di uno spazio vettoriale reale a dimensione infinita. Si pone per definizione $\Phi^0(\mathcal{A}) = \mathcal{F}(\mathcal{A})$. Si ha $\Phi^p(\mathcal{A}) = 0$ quando $p > n$: ogni p -forma è la forma nulla se $p > n = \dim(\mathcal{A})$.

Alle forme differenziali si possono adattare le definizioni e le operazioni algebriche viste per le forme multilineari antisimmetriche sugli spazi vettoriali, in particolare il **prodotto esterno**. A questo proposito è importante osservare (si veda l'Oss. 1 per il caso delle 1-forme) che una forma differenziale può essere intesa come *campo tensoriale* (Oss. 2, §2) cioè come applicazione

$$\varphi: \mathcal{A} \rightarrow \mathcal{A} \times \Lambda^p(E)$$

che assegna ad ogni punto $P \in \mathcal{A}$ un tensore covariante antisimmetrico di ordine p .

Per esempio, il prodotto esterno di due 1-forme ξ e η è la 2-forma

$$(10) \quad \boxed{(\xi \wedge \eta) = \xi \otimes \eta - \eta \otimes \xi}$$

che opera quindi sui campi vettoriali alla maniera seguente:

$$(11) \quad (\xi \wedge \eta)(\mathbf{X}, \mathbf{Y}) = \langle \mathbf{X}, \xi \rangle \langle \mathbf{Y}, \eta \rangle - \langle \mathbf{X}, \eta \rangle \langle \mathbf{Y}, \xi \rangle$$

Il prodotto esterno fra due 1-forme è quindi bilineare e anticommutativo,

$$(12) \quad \xi \wedge \eta = -\eta \wedge \xi,$$

quindi in particolare tale che

$$(13) \quad \xi \wedge \xi = 0.$$

Più in generale, il prodotto esterno di p 1-forme $(\xi^1, \xi^2, \dots, \xi^p)$ è la p -forma $\xi^1 \wedge \xi^2 \wedge \dots \wedge \xi^p$ definita da

$$(14) \quad \xi^1 \wedge \xi^2 \wedge \dots \wedge \xi^p = \sum_{\sigma \in G_p} \varepsilon(\sigma) \xi^{\sigma(1)} \otimes \xi^{\sigma(2)} \otimes \dots \otimes \xi^{\sigma(p)}$$

dove G_p è il gruppo simmetrico d'ordine p , $(\sigma(1), \sigma(2), \dots, \sigma(p))$ è la permutazione σ degli indici $(1, 2, \dots, p)$ e $\varepsilon(\sigma) = \pm 1$ a seconda che σ sia pari o dispari.

Sempre in virtù dell'analogia con il calcolo esterno sugli spazi vettoriali, si osserva che, assegnato un generico sistema di coordinate (q^i) , ogni p -forma φ ammette una rappresentazione del tipo

$$(15) \quad \varphi = \frac{1}{p!} \varphi_{i_1 i_2 \dots i_p} dq^{i_1} \wedge dq^{i_2} \wedge \dots \wedge dq^{i_p}$$

dove le funzioni $\varphi_{i_1 i_2 \dots i_p}$, dette **componenti** di φ , sono i valori della forma φ sui campi vettoriali (E_i) :

$$(16) \quad \varphi_{i_1 i_2 \dots i_p} = \varphi(E_{i_1}, E_{i_2}, \dots, E_{i_p})$$

Queste componenti sono antisimmetriche negli indici. Nel caso $p = 2$ si ha per esempio:

$$(17) \quad \varphi = \frac{1}{2} \varphi_{ij} dq^i \wedge dq^j, \quad \varphi_{ij} = \varphi(E_i, E_j), \quad \varphi_{ij} = -\varphi_{ji}.$$

Oltre che sull'operazione di prodotto esterno, il calcolo differenziale esterno si basa sull'operazione di differenziale d estesa alle p -forme. L'estensione è definita alla maniera seguente.

Una **p -forma elementare** o **forma differenziale elementare di grado p** (con p intero positivo) è il prodotto di una funzione per il prodotto esterno di p differenziali di funzioni, cioè una scrittura del tipo

$$(18) \quad \varphi = f dg^1 \wedge dg^2 \wedge \dots \wedge dg^p,$$

dove $(f, g^1, g^2, \dots, g^p)$ sono funzioni differenziabili. Come si vede dalla (15), ogni p -forma si può sempre pensare come una somma di p -forme elementari (almeno sul dominio di un sistema di coordinate).

Il **differenziale** $d\varphi$ della forma elementare (18) è per definizione la $p+1$ -forma

$$(19) \quad d\varphi = df \wedge dg^1 \wedge dg^2 \wedge \dots \wedge dg^p$$

Si noti che questa è ancora una forma elementare, dove però la funzione è la funzione costante 1. Se quindi si applica alla stessa $d\varphi$ la definizione di differenziale ora data si trova

$$(20) \quad dd\varphi = 0$$

La definizione di differenziale si estende per linearità alle somme di forme elementari e quindi, per quanto si è detto, a p -forme qualsiasi. Pertanto, tenuto conto della rappresentazione (15), il differenziale di una p -forma qualsiasi φ è la $(p+1)$ -forma $d\varphi$ definita da

$$(21) \quad \begin{aligned} d\varphi &= \frac{1}{p!} d\varphi_{i_1 i_2 \dots i_p} \wedge dq^{i_1} \wedge dq^{i_2} \wedge \dots \wedge dq^{i_p} \\ &= \frac{1}{p!} \partial_{i_0} \varphi_{i_1 i_2 \dots i_p} dq^{i_0} \wedge dq^{i_1} \wedge dq^{i_2} \wedge \dots \wedge dq^{i_p} \end{aligned}$$

Si determina quindi, per ogni intero non negativo p , un'applicazione $\mathbb{R}$ -lineare

$$d: \Phi^p(\mathcal{A}) \rightarrow \Phi^{p+1}(\mathcal{A})$$

che gode della proprietà (20), cioè tale che

$$(20') \quad d^2 = 0.$$

OSSERVAZIONE 4. Esaminiamo in particolare il differenziale di 1-forme. Data una 1-forma φ , se ne consideri la rappresentazione locale $\varphi = \varphi_i dq^i$ in un generico sistema di coordinate (q^i) . Il differenziale $d\varphi$ è la 2-forma

$$(22) \quad d\varphi = d\varphi_i \wedge dq^i$$

Dimostriamo che questa definizione non dipende dalla scelta delle coordinate. Per far questo si consideri un secondo sistema di coordinate $(q^{i'})$. Considerata la matrice jacobiana

$$E_i^{i'} = \frac{\partial q^{i'}}{\partial q^i}$$

della trasformazione di coordinate, dall'uguaglianza

$$dq^{i'} = \frac{\partial q^{i'}}{\partial q^i} dq^i$$

segue il legame tra le componenti della 1-forma:

$$\varphi_i = E_i^{i'} \varphi_{i'}.$$

Si ha allora:

$$d\varphi_i \wedge dq^i = d(E_i^{i'} \varphi_{i'}) \wedge dq^i = d\varphi_{i'} \wedge E_i^{i'} dq^i + \varphi_{i'} \partial_j E_i^{i'} dq^j \wedge dq^i.$$

Tuttavia, osservato che $E_i^{i'} dq^i = dq^{i'}$ e che $\partial_j E_i^{i'}$ è simmetrico rispetto agli indici (i, j) (perché è la derivata seconda rispetto a q^i e q^j di $q^{i'}$) mentre il termine $dq^i \wedge dq^j$ è antisimmetrico, il secondo termine si annulla e si ha semplicemente

$$d\varphi_i \wedge dq^i = d\varphi_{i'} \wedge dq^{i'}.$$

Pertanto la definizione di $d\varphi$ è formalmente la stessa qualunque sia il sistema di coordinate scelto. Ciò stabilito, va osservato che il termine $d\varphi_i$ è il differenziale di una funzione ed è quindi esprimibile attraverso le sue derivate parziali:

$$d\varphi_i = \partial_j \varphi_i dq^j.$$

Pertanto la (22) si sviluppa nella formula

$$(23) \quad d\varphi = \partial_j \varphi_i dq^j \wedge dq^i$$

Di qui segue che le componenti di $d\varphi$ sono

$$(24) \quad (d\varphi)_{ij} = d\varphi(E_i, E_j) = \partial_i \varphi_j - \partial_j \varphi_i.$$

Se in particolare φ è a sua volta il differenziale di una funzione, cioè se $\varphi = df$, poiché $\varphi_i = \partial_i f$, risulta $d\varphi = 0$, quindi: $ddf = 0$.

Tutte le precedenti considerazioni si estendono immediatamente al caso in cui si consideri non tutto lo spazio affine $\mathcal{A}$ (che è identificabile con $\mathbb{R}^n$) ma un suo sottinsieme aperto U . Denotiamo allora con $\Phi^p(U)$ lo spazio delle p -forme su U .

Le definizioni seguenti sono fondamentali per il calcolo differenziale esterno.

DEFINIZIONE 3. Una p -forma φ si dice **chiusa** se $d\varphi = 0$. Una p -forma φ si dice **esatta** se è il differenziale di una $(p-1)$ -forma ψ , chiamata **potenziale** di φ : $\varphi = d\psi$.

Dalla (20) segue ovviamente che

PROPOSIZIONE 1. Ogni forma esatta è necessariamente chiusa.

Non così immediata è invece l'implicazione inversa, che anzi non è valida se non sotto opportune condizioni coinvolgenti il grado della forma ed la topologia del dominio di definizione. Si può tuttavia affermare che:

PROPOSIZIONE 2. Se il dominio D di definizione di una p -forma chiusa φ è omeomorfo a $\mathbb{R}^n$ allora tale forma è esatta, esiste cioè in D una $(p-1)$ -forma ψ tale che $\varphi = d\psi$.

Quest'enunciato, noto come **lemma di Poincaré-Volterra** (Jules Henri Poincaré, 1854-1912; Vito Volterra, 1860-1940), è della massima importanza. Ne omettiamo la dimostrazione, osservando però che, almeno nel caso in cui φ è una 1-forma, esso è trattato nei corsi di Analisi, nel capitolo dedicato all'integrazione delle forme differenziali lineari⁽¹⁾.

Stante il fatto che ogni punto di un qualunque dominio di definizione D ammette un intorno omeomorfo a $\mathbb{R}^n$ contenuto in D (per esempio una ipersfera aperta di centro il punto e raggio opportuno) e che quindi in questo intorno si può applicare il lemma, segue immediatamente che:

PROPOSIZIONE 3. Ogni forma chiusa è localmente esatta.

⁽¹⁾ Per le 1-forme e per $n = 2$ o $n = 3$ si dimostra più precisamente che se il dominio D è *semplicemente connesso*, cioè tale che al suo interno ogni curva chiusa è deformabile con continuità in un punto, allora la chiusura implica l'esattezza. Ciò mostra che l'ipotesi che il dominio sia omeomorfo ad $\mathbb{R}^n$ può, per certe dimensioni e per certi valori del grado delle forme, essere sostituita da ipotesi diverse o più deboli, quali la semplice connessione. Queste considerazioni sono oggetto della Geometria Differenziale e della Topologia Algebrica.

Ciò vuol dire che essa ammette nell'intorno di ogni punto del suo dominio di definizione dei potenziali, detti **potenziali locali**. Giustapposendo e prolungando potenziali locali si può giungere alla costruzione di un **potenziale globale**, definito cioè su tutto il dominio D , ma solo nel caso in cui la forma è esatta.

Possiamo schematizzare la precedente discussione nel quadro seguente:

forma esatta	$\Rightarrow$	forma chiusa
forma chiusa	$\Rightarrow$	forma localmente esatta
forma chiusa su dominio $\simeq \mathbb{R}^n$	$\Rightarrow$	forma esatta

OSSERVAZIONE 5. Se una forma è esatta allora il suo potenziale è definito a meno di una forma chiusa. Vale a dire: se $\varphi = d\psi$, allora ogni forma del tipo $\psi' = \psi + \eta$ con $d\eta = 0$ è ancora un potenziale di φ , perché $d\psi' = d\psi$. In particolare, se φ è una 1-forma, i suoi potenziali sono delle funzioni e differiscono per delle costanti (sulle componenti connesse del dominio di definizione). Infatti le 0-forme chiuse sono funzioni tali che $df = 0$ e quindi localmente costanti.

Vediamo ora alcuni esempi.

ESEMPIO 1. Per ogni 1-forma φ ed ogni funzione f vale la formula

$$(25) \quad d(f\varphi) = df \wedge \varphi + f d\varphi$$

Infatti, in virtù delle formule sopra stabilite:

$$\begin{aligned} d(f\varphi) &= d(f\varphi_i dq^i) \\ &= d(f\varphi_i) \wedge dq^i \\ &= \varphi_i df \wedge dq^i + f d\varphi_i \wedge dq^i \\ &= df \wedge \varphi + f d\varphi. \end{aligned}$$

Una formula analoga vale anche se φ è una p -forma: dimostrarlo.

ESEMPIO 2. Si consideri in $\mathbb{R}^2$ (coordinate (x, y)) la generica 1-forma

$$(26) \quad \varphi = A(x, y) dx + B(x, y) dy.$$

Calcoliamone il differenziale. Tenuto conto che il prodotto esterno di due differenziali di funzioni è anticommutativo, si ha:

$$\begin{aligned} d\varphi &= dA \wedge dx + dB \wedge dy \\ &= \left(\frac{\partial A}{\partial x} dx + \frac{\partial A}{\partial y} dy \right) \wedge dx + \left(\frac{\partial B}{\partial x} dx + \frac{\partial B}{\partial y} dy \right) \wedge dy \\ &= \left(\frac{\partial B}{\partial x} - \frac{\partial A}{\partial y} \right) dx \wedge dy. \end{aligned}$$

La condizione di chiusura $d\varphi = 0$ equivale pertanto a:

$$(27) \quad \frac{\partial A}{\partial y} = \frac{\partial B}{\partial x}.$$

L'analogo calcolo fatto per una 1-forma in $\mathbb{R}^3$,

$$(28) \quad \varphi = A(x, y, z) dx + B(x, y, z) dy + C(x, y, z) dz,$$

mostra che le condizioni di chiusura sono

$$(29) \quad \frac{\partial A}{\partial y} = \frac{\partial B}{\partial x}, \quad \frac{\partial B}{\partial z} = \frac{\partial C}{\partial y}, \quad \frac{\partial C}{\partial x} = \frac{\partial A}{\partial z}.$$

In generale, la chiusura $d\varphi = 0$ di una 1-forma è espressa, qualunque siano le coordinate scelte, dalle equazioni (si veda la (24))

$$(30) \quad \partial_i \varphi_j = \partial_j \varphi_i$$

ESEMPIO 3. La chiusura di una 2-forma φ è espressa dalle equazioni

$$(31) \quad \partial_h \varphi_{ij} + \partial_i \varphi_{jh} + \partial_j \varphi_{hi} = 0$$

dove, si osservi, gli indici sono permutati ciclicamente. Si ha infatti, particolarizzando la (21) al caso $p = 2$,

$$d\varphi = \frac{1}{2} \partial_h \varphi_{ij} dq^h \wedge dq^i \wedge dq^j.$$

Tenuto conto delle (16) e (14), le componenti di $d\varphi$ si ottengono sommando $\partial_h \varphi_{ij}$ agli analoghi termini ottenuti da tutte le permutazioni dei tre indici, con il segno della parità. Quindi:

$$(d\varphi)_{hij} = \frac{1}{2} (\partial_h \varphi_{ij} + \partial_i \varphi_{jh} + \partial_j \varphi_{hi} - \partial_h \varphi_{ji} - \partial_j \varphi_{ih} - \partial_i \varphi_{hj}).$$

Quest'espressione si semplifica, tenuto conto che le φ_{ij} sono antisimmetriche, in

$$(32) \quad (d\varphi)_{hij} = \partial_h \varphi_{ij} + \partial_i \varphi_{jh} + \partial_j \varphi_{hi}$$

ESEMPIO 4. Si consideri il piano affine riferito a coordinate cartesiane (x, y) e su di esso la forma differenziale lineare

$$\xi = a dx + b dy$$

con a e b costanti. La forma ξ è esatta. Un suo potenziale è la funzione $f(x, y) = ax + by$.

ESEMPIO 5. Si consideri il piano affine riferito a coordinate cartesiane (x, y) e su di esso, con eventuale esclusione dell'origine, la forma differenziale lineare

$$\xi = h(r^2)(x dx + y dy),$$

dove $h(\cdot)$ è una funzione differenziabile sull'asse reale positivo ed $r^2 = x^2 + y^2$. Se ξ è definita anche nell'origine, allora in quel punto si annulla. Verifichiamo che è chiusa:

$$\begin{aligned} d\xi &= dh \wedge (x dx + y dy) + h d(x dx + y dy) \\ &= \left(\frac{\partial h}{\partial x} dx + \frac{\partial h}{\partial y} dy \right) \wedge (x dx + y dy) + 0 \\ &= y \frac{\partial h}{\partial x} dx \wedge dy + x \frac{\partial h}{\partial y} dy \wedge dx \\ &= \left(y \frac{\partial h}{\partial x} - x \frac{\partial h}{\partial y} \right) dx \wedge dy \\ &= h'(2yx - 2xy) dx \wedge dy = 0. \end{aligned}$$

Lo si riconosce più facilmente passando a coordinate polari (r, ϑ) . Da $r^2 = x^2 + y^2$ segue, differenziando ambo i membri e semplificando,

$$dr = \frac{1}{r}(x dx + y dy).$$

Posto allora $\phi(r) = r h(r^2)$, segue che ξ è semplicemente data da

$$\xi = \phi(r) dr.$$

Di qui segue immediatamente non solo che ξ è chiusa, ma che è esatta. Un potenziale di ξ è infatti dato da una qualunque primitiva $f(r)$ di $\phi(r)$.

ESEMPIO 6. Ancora sul piano affine riferito a coordinate cartesiane (x, y) consideriamo la forma differenziale lineare

$$\xi = -y dx + x dy.$$

Non è chiusa:

$$d\xi = -dy \wedge dx + dx \wedge dy = 2 dx \wedge dy.$$

Nel piano esclusa l'origine si consideri la forma differenziale

$$\eta = \frac{1}{r^k} \xi, \quad r = \sqrt{x^2 + y^2}.$$

Il suo differenziale è dato da

$$d\eta = \frac{2-k}{r^k} dx \wedge dy.$$

Per $k = 2$ si ha $d\eta = 0$. In questo caso la forma η , pur essendo chiusa, non ammette un potenziale globale (si osservi che il dominio di definizione di η non è omeomorfo a $\mathbb{R}^2$). Lo si riconosce facilmente passando a coordinate polari. In queste coordinate risulta infatti:

$$\xi = -y dx + x dy = -r \sin \vartheta d(r \cos \vartheta) + r \cos \vartheta d(r \sin \vartheta),$$

e, a conti fatti:

$$\xi = r^2 d\vartheta.$$

Allora (sempre per $k = 2$) $\eta = d\vartheta$. L'angolo ϑ è dunque un potenziale di η . Ma questa funzione è continua e differenziabile solo se si esclude una semiretta uscente dall'origine. Pertanto, se si persiste nel considerare il piano affine esclusa l'origine come dominio di definizione della forma η , in tale dominio essa è chiusa ma non esatta.

5. Curve

Chiamiamo **curva parametrizzata** (brevemente: **curva**) in uno spazio affine $\mathcal{A}$ un'applicazione $\gamma: I \rightarrow \mathcal{A}$ di un intervallo reale aperto $I \subset \mathbb{R}$ nello spazio affine $\mathcal{A}$. Ad ogni valore reale $t \in I$ corrisponde un punto $P = \gamma(t)$ dello spazio affine. Considerata un'origine $O \in \mathcal{A}$, per la curva γ sussiste la **rappresentazione vettoriale**

$$(1) \quad \mathbf{x} = \gamma(t)$$

dove $\mathbf{x} = 0P$ è il vettore posizione del punto P rispetto al punto O . Se si considerano coordinate cartesiane (x^α) aventi origine in O , la (1) si traduce nelle **equazioni parametriche**

$$(2) \quad x^\alpha = \gamma^\alpha(t) \quad (\alpha = 1, \dots, n).$$

Si dice che la curva è di classe C^k se le funzioni $\gamma^\alpha(t)$ sono tutte di classe C^k .

Una curva può essere interpretata come **moto** di un punto P nello spazio affine, se il parametro t è interpretato come **tempo**. Quest'**interpretazione cinematica** di una curva parametrizzata è di particolare interesse nel caso in cui lo spazio affine ha dimensione 3 oppure 4 e sarà nel seguito ripresa e ulteriormente sviluppata. Nel caso in cui la curva rappresenti il moto di un punto nello spazio affine tridimensionale euclideo il generico vettore $OP = \mathbf{x}$ verrà anche denotato con $\mathbf{r}$ e chiamato **vettore posizione** o **raggio vettore**.

L'immagine della curva, cioè l'insieme $\gamma(I) = \{P \in \mathcal{A} \mid \exists t \in I \mid \gamma(t) = P\}$, è detta **traiettoria** o **orbita**. Va osservato che in geometria una curva è intesa di solito come **curva non parametrizzata**, cioè come **luogo di punti** definito da $n-1$ equazioni coinvolgenti coordinate cartesiane (x^α) .

Diciamo **vettore tangente** alla curva γ nel punto $\gamma(t)$ il vettore denotato con $\dot{\gamma}(t)$ e definito dal limite

$$(3) \quad \dot{\gamma}(t) = \lim_{h \rightarrow 0} \frac{1}{h} (\gamma(t+h) - \gamma(t)).$$

Nel contesto cinematico esso prende il nome di **velocità (istantanea)** e lo si denota con $\mathbf{v}(t)$. Nella rappresentazione vettoriale del moto abbiamo

$$(4) \quad \mathbf{v}(t) = \lim_{h \rightarrow 0} \frac{1}{h} (\mathbf{x}(t+h) - \mathbf{x}(t))$$

Scriveremo allora più semplicemente:

$$(5) \quad \mathbf{v} = \frac{d\mathbf{x}}{dt}$$

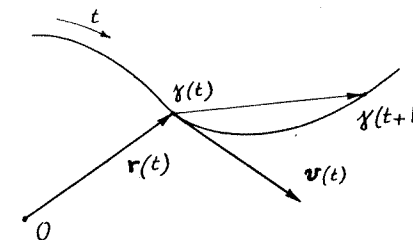


Fig. 5.1: rappresentazione vettoriale di una curva e vettore tangente (velocità).

OSSERVAZIONE 1. Si conviene di interpretare il vettore tangente $\dot{\gamma}(t) = \mathbf{v}(t)$ come vettore applicato nel punto $\gamma(t)$. L'applicazione

$$\dot{\gamma}: I \rightarrow \mathcal{A} \times E: t \mapsto (\gamma(t), \dot{\gamma}(t))$$

viene detta **curva tangente** della curva $\gamma: I \rightarrow \mathcal{A}$.

Oltre al vettore tangente, cioè alla velocità, si considera anche il vettore **accelerazione (istantanea)** definito dal limite

$$(6) \quad \mathbf{a}(t) = \lim_{h \rightarrow 0} \frac{1}{h} (\mathbf{v}(t+h) - \mathbf{v}(t))$$

È la derivata del vettore velocità:

$$(7) \quad \mathbf{a} = \frac{d\mathbf{v}}{dt} = \frac{d^2\mathbf{x}}{dt^2}$$

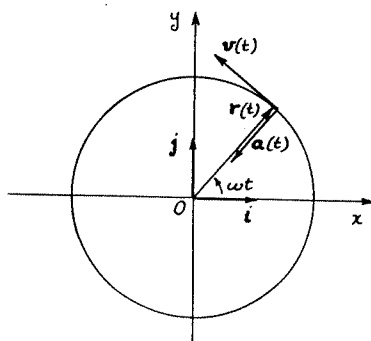


Fig. 5.2: moto circolare uniforme.

Nella rappresentazione parametrica cartesiana della curva le componenti dei vettori velocità ed accelerazione sono date rispettivamente dalle derivate prime e seconde delle funzioni γ^α :

$$(8) \quad v^\alpha = \frac{d\gamma^\alpha}{dt}, \quad a^\alpha = \frac{dv^\alpha}{dt} = \frac{d^2\gamma^\alpha}{dt^2}.$$

ESEMPIO 1. Nel piano affine bidimensionale si consideri la curva di equazioni parametriche

$$\begin{cases} x = r \cos \omega t, \\ y = r \sin \omega t, \end{cases}$$

con r e ω numeri reali positivi. Poiché da queste equazioni segue $x^2 + y^2 = r^2$, la traiettoria è la circonferenza di raggio r e centro l'origine (per comodità di interpretazione conviene pensare il piano dotato della struttura euclidea rispetto alla quale le coordinate cartesiane scelte (x, y) sono ortonormali). La rappresentazione vettoriale del moto è

$$\mathbf{x}(t) = r(\cos \omega t \mathbf{i} + \sin \omega t \mathbf{j}).$$

Di conseguenza la velocità e l'accelerazione sono

$$\begin{cases} \mathbf{v}(t) = r\omega(-\sin \omega t \mathbf{i} + \cos \omega t \mathbf{j}), \\ \mathbf{a}(t) = -r\omega^2(\cos \omega t \mathbf{i} + \sin \omega t \mathbf{j}) = -\omega^2 \mathbf{x}(t). \end{cases}$$

Questa curva rappresenta un **moto circolare uniforme** di **velocità angolare** ω (Fig. 5.2).

Se anziché coordinate cartesiane si considerano generiche coordinate (q^i) , la curva (1) è rappresentata da equazioni parametriche del tipo

$$q^i = \gamma^i(t).$$

Nella rappresentazione parametrica della curva in coordinate generiche le componenti della velocità, rispetto al riferimento associato $(\mathbf{E}_i)$, sono ancora le derivate prime delle funzioni rappresentatrici γ^i . Infatti, facendo ricorso alla rappresentazione vettoriale, osserviamo che il vettore posizione $\mathbf{x}$ può pensarsi dipendente dal parametro t attraverso le coordinate (q^i) . Risulta pertanto

$$\mathbf{v} = \frac{d\mathbf{x}}{dt} = \frac{\partial \mathbf{x}}{\partial q^i} \frac{dq^i}{dt},$$

dove alle (q^i) vanno sostituite le funzioni $\gamma^i(t)$. Ricordando la definizione dei vettori $(\mathbf{E}_i)$, questa uguaglianza si traduce in

(9)

$$\mathbf{v} = v^i \mathbf{E}_i$$

posto che si abbia

(10)

$$v^i = \frac{dq^i}{dt}$$

col che è dimostrato quanto asserito. Le componenti dell'accelerazione invece non sono più semplicemente le derivate prime di queste componenti, cioè le derivate seconde delle funzioni (γ^i) . Derivando la (9) e tenendo conto che i vettori $(\mathbf{E}_i)$ dipendono dal parametro t attraverso le coordinate (q^i) , si ha infatti successivamente:

$$\begin{aligned} \frac{d\mathbf{v}}{dt} &= \frac{dv^i}{dt} \mathbf{E}_i + v^i \frac{d\mathbf{E}_i}{dt} \\ &= \frac{dv^i}{dt} \mathbf{E}_i + v^i \partial_j \mathbf{E}_i \frac{dq^j}{dt} \\ &= \frac{dv^i}{dt} \mathbf{E}_i + v^i v^j \Gamma_{ji}^k \mathbf{E}_k \\ &= \left(\frac{dv^i}{dt} + v^j v^h \Gamma_{jh}^i \right) \mathbf{E}_i. \end{aligned}$$

Risulta quindi

$$(11) \quad \boxed{\mathbf{a} = a^i \mathbf{E}_i}$$

con

$$(12) \quad \boxed{a^i = \frac{dv^i}{dt} + \Gamma_{jh}^i v^j v^h}$$

Come si vede, per ottenere le componenti dell'accelerazione, alle derivate prime delle componenti delle velocità vanno aggiunti dei termini quadratici in tali componenti, i cui coefficienti sono i simboli di Christoffel. Formule di questo tipo saranno ritrovate nel seguito.

OSSERVAZIONE 2. Si consideri un'equazione parametrica del tipo

$$q^1 = t, \quad q^2 = c^2, \quad \dots, \quad q^n = c^n,$$

con $(c^2, \dots, c^n)$ costanti. Si tratta di una curva coordinata (Oss. 1, §3). Dalla (10) segue che le componenti del vettore tangente sono

$$v^1 = 1, \quad v^2 = \dots = v^n = 0.$$

Quindi per la (9), $\mathbf{v} = \mathbf{E}_1$. Il campo vettoriale $\mathbf{E}_1$ è dunque tangente a questa curva. È così verificato che i campi vettoriali $\mathbf{E}_i$ sono tangenti alle rispettive curve coordinate.

Nel seguito sarà utilizzata la proprietà seguente.

PROPOSIZIONE 1. Siano $\gamma: I \rightarrow \mathcal{A}$ una curva ed $F: \mathcal{A} \rightarrow \mathbb{R}$ un campo scalare (entrambi di classe C^1). Per ogni $t \in I$ vale la formula

$$(13) \quad \boxed{D(F \circ \gamma)(t) = \langle \mathbf{v}(t), dF \rangle}$$

dove D è il simbolo di derivata di funzione reale a variabile reale.

Dimostrazione. Si noti che la composizione $F \circ \gamma$ è una funzione reale sopra l'intervallo I . Introdotto un qualunque sistema di coordinate cartesiane e ricordata la regola della derivazione di funzioni composte, si ha successivamente:

$$\begin{aligned} D(F \circ \gamma)(t) &= \frac{\partial F}{\partial x^\alpha} D\gamma^\alpha(t) \\ &= \frac{\partial F}{\partial x^\alpha} v^\alpha(t) \\ &= \langle \mathbf{v}(t), dF \rangle. \quad \blacksquare \end{aligned}$$

6. Superfici

Una **superficie** in uno spazio affine $\mathcal{A}$ è un insieme di punti (o **luogo di punti**) soddisfacenti ad una ben definita proprietà. Vi sono essenzialmente due modi per definire una superficie: o attraverso un'equazione (o più di una) o per via parametrica.

DEFINIZIONE 1. Assegnata una funzione reale $F: \mathcal{A} \rightarrow \mathbb{R}$ sullo spazio affine (la supponiamo di classe sufficientemente alta, C^∞ per semplicità), l'insieme

$$(1) \quad Q = \{P \in \mathcal{A} \mid F(P) = 0\}$$

dei punti di $\mathcal{A}$ che annullano F , cioè l'insieme degli **zeri** della funzione F , supposto che non sia vuoto, è una **superficie** in $\mathcal{A}$. Un punto $P \in \mathcal{A}$ si dice **punto singolare** o **punto critico** della funzione F se in P il differenziale di F si annulla: $dF(P) = 0$. Un punto della superficie Q definita dalla (1) si dice **punto regolare** se esso non è punto singolare di F . Una superficie Q si dice **superficie regolare** (di dimensione $n-1$ o **codimensione** 1) se ammette una rappresentazione del tipo (1) senza punti singolari.

Scelte sullo spazio affine delle coordinate (x^α) (anche non cartesiane) ($\alpha = 1, \dots, n$), la (1) equivale all'equazione

$$(2) \quad F(x^\alpha) = 0,$$

dove $F(x^\alpha)$ è la funzione rappresentativa di F . La (2) è detta **equazione della superficie** nelle coordinate (x^α) . Un punto è singolare se in corrispondenza ai valori delle sue coordinate tutte le derivate parziali di F , che sono le componenti del differenziale dF , si annullano.

ESEMPIO 1. Consideriamo, in questo e negli esempi seguenti, lo spazio affine tridimensionale. L'equazione (1) in coordinate cartesiane (x, y, z) diventa

$$(3) \quad F(x, y, z) = 0.$$

Se si considera per esempio la funzione $F(x, y, z) = x^2 + y^2 + z^2 - 1$, la corrispondente equazione

$$(4) \quad x^2 + y^2 + z^2 - 1 = 0$$

definisce una superficie regolare Q (la sfera di raggio unitario con centro l'origine) perché la funzione F ha differenziale non nullo in tutti i punti dello spazio ad eccezione dell'origine O , che non appartiene a Q . Le sue derivate parziali sono infatti

$$(2x, 2y, 2z)$$

e non si annullano contemporaneamente se non nell'origine. In coordinate polari sferiche (r, φ, ϑ) la stessa superficie ha equazione

$$r = 1.$$

Posto $G = F^2$, la stessa superficie è definita dall'equazione $G = 0$. In questo caso però ogni punto di Q è punto critico di G , perché $dG = 2F dF$ ed $F = 0$ su Q .

ESEMPIO 2. L'equazione

$$x^2 + y^2 - z^2 = 0.$$

Definisce una superficie (un cono) con un punto singolare nell'origine (il vertice). Infatti nell'origine si annullano tutte e tre le derivate parziali della funzione a primo membro, cioè il vettore

$$(2x, 2y, -2z).$$

ESEMPIO 3. Sia $f(x, y)$ una funzione in due variabili. Il suo grafico, di equazione

$$(5) \quad z = f(x, y),$$

definisce una superficie regolare. La (5) si può infatti mettere nella forma (3) ponendo $F(x, y, z) = z - f(x, y)$, e si ha sempre $dF \neq 0$ perché

$$\frac{\partial F}{\partial z} = 1.$$

Viceversa, un'equazione del tipo (3) può porsi nella forma (5) se e solo se è risolubile rispetto a z . Questo è possibile se e solo se nei punti soddisfacenti alla (3) si ha

$$(6) \quad \frac{\partial F}{\partial z} \neq 0.$$

È da notarsi che questa circostanza può o non verificarsi per tutti i punti o verificarsi in maniera non univoca. Per esempio, l'equazione (4) della sfera è risolubile rispetto a z in due modi,

$$(7) \quad z = \pm \sqrt{1 - x^2 - y^2}.$$

Si tratta delle due semisfere con bordo costituito dai punti della circonferenza $x^2 + y^2 = 1$ del piano $z = 0$, punti nei quali la condizione (6) non è verificata (Fig. 6.1).

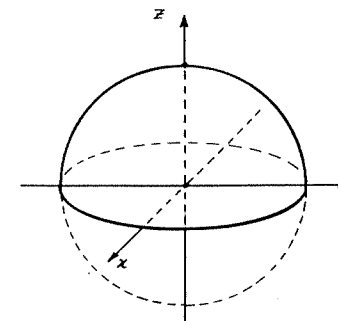


Fig. 6.1

ESEMPIO 4. La scelta di opportune coordinate può facilitare la rappresentazione di una superficie (vedi l'esempio della sfera in coordinate polari sferiche). Per esempio un'equazione del tipo

$$F(r, z) = 0$$

con $r^2 = x^2 + y^2$ definisce una *superficie di rotazione* intorno all'asse z .

OSSERVAZIONE 1. Per ogni $c \in \mathbb{R}$, l'insieme

$$(8) \quad Q_c = \{P \in \mathcal{A} \mid F(P) = c\}$$

è, se non vuoto, una superficie (basta sostituire $F - c$ a F nelle precedenti considerazioni). L'equazione (2), che ora diventa

$$(9) \quad F(x^\alpha) = c,$$

definisce al variare di c un insieme di superfici, dette **superfici** o **insiemi di livello** della funzione F . Due insiemi di livello corrispondenti a due distinti valori della costante c sono ovviamente disgiunti (hanno intersezione vuota). Gli insiemi di livello formano quindi una partizione di M detta **fogliettamento**.

ESEMPIO 5. Le equazioni

$$x^2 + y^2 + z^2 = c$$

con $c > 0$ rappresentano un fogliettamento di sfere su tutto lo spazio esclusa l'origine. Per $c = 0$ la superficie si riduce ad un punto: l'origine. Per $c < 0$ l'insieme di livello è vuoto.

OSSERVAZIONE 2. Dobbiamo anche considerare il caso in cui una superficie (od un fogliettamento) è definito da più di un'equazione. Per far questo occorre premettere una definizione. Date $k \leq n$ funzioni reali (sempre di classe C^∞) $(F_a) = (F_1, \dots, F_k)$, sopra un aperto $M \subset \mathcal{A}$, queste si dicono **funzioni indipendenti** in un punto $P \in M$ se in quel punto i loro differenziali (dF_a) sono linearmente indipendenti. Questa condizione si traduce, qualunque siano le coordinate (x^α) (anche non cartesiane), nella massimalità del rango della matrice $k \times n$ delle derivate parziali:

$$(10) \quad \left(\frac{\partial F_a}{\partial x^\alpha} \right).$$

Infatti l'indipendenza dei differenziali equivale, per definizione, all'implicazione

$$\kappa^a dF_a = 0 \implies \kappa^a = 0,$$

quindi all'implicazione

$$\kappa^a \frac{\partial F_a}{\partial x^\alpha} = 0 \implies \kappa^a = 0.$$

Il primo membro è un sistema lineare omogeneo di n equazioni nelle k incognite (κ^a) . Affinché dia come unica soluzione $\kappa^a = 0$ occorre e basta che la matrice dei coefficienti (10) abbia rango massimo.

OSSERVAZIONE 3. Per ogni scelta delle costanti $\mathbf{c} = (c_a)$ l'insieme

$$(11) \quad Q_{\mathbf{c}} = \{P \in M \mid F_a(P) = c_a\}$$

o è vuoto o è una superficie (in senso lato). Se le funzioni (F_a) sono indipendenti in un suo punto, questo si dice **punto regolare** della superficie. Se tutti i punti sono regolari, allora si dice che $Q_{\mathbf{c}}$ è una **superficie regolare di codimensione k** (ovvero di **dimensione $n - k$**). Al variare di $\mathbf{c}$ le $Q_{\mathbf{c}}$ definiscono un fogliettamento. La definizione (11) si traduce nelle k equazioni

$$(12) \quad \boxed{F_a(x_\alpha) = c_a \quad (a = 1, \dots, k)}$$

Vediamo ora la definizione parametrica delle superfici. Si consideri un'equazione vettoriale del tipo

$$(13) \quad \boxed{OP = \mathbf{x}(q^i) \quad (i = 1, \dots, m)}$$

dove a secondo membro la funzione vettoriale $\mathbf{x}(q^i)$ nelle $m < n$ variabili (q^i) è supposta sufficientemente regolare (per semplicità di classe C^∞) in un dominio aperto $D \subset \mathbb{R}^m$. Al variare di questi parametri in D , il punto P descrive un insieme $Q \subset \mathcal{A}$. In altri termini, l'equazione (13) trasforma il dominio $D \subset \mathbb{R}^m$ in un sottoinsieme $Q \subset \mathcal{A}$. In ogni punto dell'insieme Q sono definiti gli m vettori

$$(14) \quad \boxed{\mathbf{E}_i = \frac{\partial \mathbf{x}}{\partial q^i}}$$

DEFINIZIONE 2. Un punto $P \in Q$ si dice **punto regolare** o **non singolare** se i vettori $(\mathbf{E}_i)$ sono in quel punto indipendenti. Se tutti i punti sono regolari si dice che l'insieme Q è una **superficie immersa** di dimensione m (Fig. 6.2) e la (13) prende il nome di **rappresentazione vettoriale parametrica** della superficie; si dice anche che l'equazione (13) è un'**immersione**.

OSSERVAZIONE 4. In coordinate cartesiane (x^α) di origine O le (13) si traducono in n equazioni scalari

$$(15) \quad x^\alpha = x^\alpha(q^i)$$

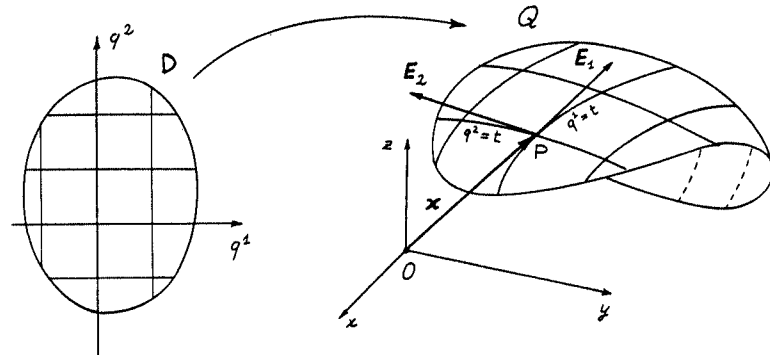


Fig. 6.2: rappresentazione parametrica di una superficie.

dette **equazioni parametriche** della superficie. L'indipendenza dei vettori (E_i) equivale alla massimalità del rango della matrice $m \times n$

$$(16) \quad \left(\frac{\partial x^\alpha}{\partial q^i} \right).$$

Infatti l'indipendenza equivale, per definizione, all'implicazione

$$a^i E_i = 0 \implies a^i = 0,$$

quindi all'implicazione

$$a^i \frac{\partial x^\alpha}{\partial q^i} = 0 \implies a^i = 0.$$

Il primo membro è un sistema lineare omogeneo di n equazioni nelle m incognite (a^i). Affinché dia come unica soluzione $a^i = 0$ occorre e basta che la matrice dei coefficienti (16) abbia rango massimo.

OSSERVAZIONE 5. Se l'immersione (13) è iniettiva (a m -ple distinte di valori dei parametri corrispondono punti distinti di Q) i parametri

(q^i) vengono detti **coordinate superficiali**. Essi infatti possono interpretarsi come *funzioni* sulla superficie, $q^i: Q \rightarrow \mathbb{R}$, ponendo, per ogni punto $P \in Q$, $q^i(P)$ uguale al valore del parametro q^i che lo determina.

OSSERVAZIONE 6. I campi vettoriali E_i sono tangenti alle corrispondenti **curve coordinate**, ottenute dalla rappresentazione (13) facendo variare la sola coordinata q^i e assegnando alle rimanenti dei valori costanti (Fig. 6.2). Per esempio le equazioni parametriche

$$q^1 = t, \quad q^2 = c^2, \quad \dots, \quad q^m = c^m,$$

con $(c^2, \dots, c^m)$ costanti, definiscono una famiglia di curve coordinate a cui il campo E_1 è tangente (si possono ripetere qui le considerazioni del § precedente). Si noti che le curve coordinate sono le immagini secondo l'applicazione (13) delle rette coordinate di $\mathbb{R}^m$. Segue che i campi vettoriali E_i sono in ogni punto tangenti alla superficie e generano quindi **piano tangente**. Si dice anche che essi formano il **riferimento naturale tangente** associato alle coordinate (q^i) (Fig. 6.2).

Il legame tra i due tipi di rappresentazione di una superficie, quella mediante equazioni e quella parametrica, è descritto dal seguente enunciato

PROPOSIZIONE 1. Sia Q una superficie regolare di equazioni $F_a(x^\alpha) = 0$ ($a = 1, \dots, k$). Comunque si fissi un punto $P \in Q$ esiste un'immersione iniettiva $OP = x(q^i)$ ($i = 1, \dots, m$; $m = n - k$) che manda un aperto $D \subset \mathbb{R}^m$ in un intorno $U \subset Q$ di P .

Dimostrazione. La massimalità del rango della matrice (10) a k righe ed $n > k$ colonne in un punto P (dovuta all'ipotesi di regolarità) implica che una sua sottomatrice quadrata di ordine massimo k è regolare. Supposto che questa sia costituita dalle ultime k colonne (se così non è basta cambiare l'ordine delle coordinate (x^α)), è possibile risolvere le k equazioni $F_a(x^\alpha) = 0$ rispetto alle ultime k coordinate ($x^{n-k+1}, \dots, x^n$) e quindi esprimere queste in funzione delle prime $n - k$. Posto pertanto

$$(17) \quad x^1 = q^1, \quad \dots, \quad x^m = q^m \quad (m = n - k),$$

è possibile esprimere *tutte* le n coordinate (x^α) in funzione dei parametri (q^i) = ($q^1, \dots, q^m$), ottenendo così una rappresentazione del tipo (15) iniettiva. La massimalità del rango della matrice del tipo (16) è infine

assicurata dal fatto che, stante le (17), le sue prime m colonne formano la matrice identica di ordine m . ■

OSSERVAZIONE 7. Questa proposizione mostra che una superficie regolare è in genere *solo localmente* rappresentabile da equazioni parametriche, ovvero che non è possibile in genere descriverla interamente mediante una sola rappresentazione parametrica. Una rappresentazione parametrica del tipo considerato (immersione inettiva) prende anche il nome di **carta** di Q . Questo concetto, che sarà precisato in un successivo capitolo del corso, è alla base della nozione di **varietà differenziabile**, nozione che generalizza quella di superficie.

ESEMPIO 6. Consideriamo la sfera nello spazio affine tridimensionale (Esempi 1 e 3). Risolvendo l'equazione rispetto a z (equazioni (7)) si ottengono due rappresentazioni parametriche della sfera,

$$(18) \quad \begin{cases} x = u, \\ y = v, \\ z = \pm\sqrt{1 - u^2 - v^2}, \end{cases}$$

nei parametri $q^1 = u, q^2 = v$, variabili nel disco unitario aperto $D \subset \mathbb{R}^2$ di equazione $u^2 + v^2 < 1$ (Fig. 6.1). Ciascuna di queste non rappresenta globalmente la sfera, ma solo una semisfera aperta. Le equazioni

$$(19) \quad \begin{cases} x = \sin \varphi \cos \vartheta, \\ y = \sin \varphi \sin \vartheta, \\ z = \cos \varphi, \end{cases}$$

forniscono un'ulteriore rappresentazione parametrica della sfera nelle coordinate $(q^1, q^2) = (\vartheta, \varphi)$ variabili nel dominio aperto $D \subset \mathbb{R}^2$ definito da $0 < \varphi < \pi$ e $0 < \vartheta < 2\pi$. Sono quindi esclusi dalla rappresentazione i punti della sfera appartenenti al meridiano che interseca il semiasse positivo delle x , nonché i poli (Fig. 6.3).

ESERCIZIO 1. Nello spazio affine euclideo tridimensionale riferito a coordinate cartesiane ortonormali (x, y, z) si scriva l'equazione del **toro** e se ne dia anche una rappresentazione parametrica. Questa superficie è ottenuta facendo ruotare una circonferenza intorno ad un asse (per es. l'asse z) esterno ad essa e giacente nel suo piano (per es. il piano (x, z)).

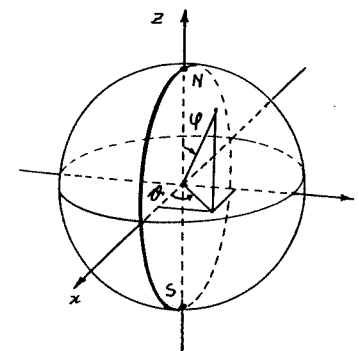


Fig. 6.3

Nel seguito sarà utilizzato il seguente criterio di tangenza di un vettore ad una superficie.

PROPOSIZIONE 2. Un vettore $\mathbf{v}$ è tangente ad una superficie definita da equazioni del tipo (12) in un suo punto regolare se e solo se valgono le condizioni

$$(20) \quad \langle \mathbf{v}, dF_a \rangle = 0 \quad (a = 1, \dots, k)$$

Dimostrazione. Combinando una qualunque rappresentazione parametrica della superficie con le sue equazioni (12) si ottengono le identità

$$F_a(x^\alpha(q^i)) = c_a.$$

Derivandole parzialmente rispetto alle coordinate (q^i) si ottengono le identità

$$\frac{\partial F_a}{\partial x^\alpha} \frac{\partial x^\alpha}{\partial q^i} = 0.$$

Queste possono scriversi, più sinteticamente

$$\langle \mathbf{E}_i, dF_a \rangle = 0.$$

Ciò significa che sulla superficie Q si annullano le derivate rispetto ai campi vettoriali tangenti $\mathbf{E}_i$ delle funzioni F_a che la definiscono. Siccome, per l'ipotesi di regolarità della rappresentazione, i campi $\mathbf{E}_i$ sono

ovunque indipendenti e generano i piani tangenti alla superficie, concludiamo che valgono le (20) per ogni vettore $\mathbf{v}$ tangente alla superficie. Viceversa, stante l'indipendenza dei k differenziali dF_a nel punto regolare considerato, i vettori soddisfacenti alle equazioni (20) formano un sottospazio di dimensione $n - k$, quindi coincidente necessariamente con lo spazio tangente alla superficie. ■

OSSERVAZIONE 8. Una seconda dimostrazione di questa proposizione può essere sviluppata rendendo rigoroso nei dettagli il seguente ragionamento intuitivo. Un vettore $\mathbf{v}$ è tangente ad una superficie Q se e solo se esiste una curva $\gamma(t)$ la cui immagine è tutta contenuta in Q e tale che $\mathbf{v} = \dot{\gamma}(0)$, cioè tale che $\mathbf{v}$ è il vettore tangente alla curva per $t = 0$. Ciò posto, se F è una funzione reale sullo spazio affine costante su Q , si ha $F \circ \gamma(t) = \text{cost.}$ perché γ giace sulla superficie, e quindi (si veda il § precedente)

$$\langle \mathbf{v}, dF \rangle = D(F \circ \gamma)(0) = 0.$$

Ciò vale in particolare per le funzioni F_a che definiscono la superficie. Di qui le (20).

OSSERVAZIONE 9. Il criterio di tangenza per un vettore si estende immediatamente ad un campo vettoriale: condizione necessaria e sufficiente affinché un campo vettoriale $\mathbf{X}$ su $\mathcal{A}$ sia tangente ad una superficie di equazioni $F_a = \text{cost.}$ è che nei punti regolari sia

$$(21) \quad \langle \mathbf{X}, dF_a \rangle = 0, \quad (a = 1, \dots, k)$$

In coordinate cartesiane la (20) diventa

$$(20') \quad v^\alpha \frac{\partial F_a}{\partial x^\alpha} = 0,$$

e la (21)

$$(21') \quad X^\alpha \frac{\partial F_a}{\partial x^\alpha} = 0.$$

Ad esempio: nello spazio affine tridimensionale un campo vettoriale $\mathbf{X}$ di componenti cartesiane (X, Y, Z) risulta tangente ad una superficie di equazione $F = 0$ se e solo se nei punti di questa si ha

$$(22) \quad X \frac{\partial F}{\partial x} + Y \frac{\partial F}{\partial y} + Z \frac{\partial F}{\partial z} = 0.$$

7. Sistemi dinamici

Si consideri uno spazio affine $\mathcal{A}$ (per esempio il piano $\mathbb{R}^2$) invaso da un fluido in **moto stazionario**, cioè tale che in ogni prefissato punto di $\mathcal{A}$ il vettore velocità delle particelle del fluido che vi transitano è costante nel tempo. In queste condizioni i vettori velocità danno luogo ad un campo vettoriale $\mathbf{X}$ indipendente dal tempo. Viceversa, se si assegna un campo vettoriale $\mathbf{X}$ su $\mathcal{A}$, allora questo può essere interpretato come **campo di velocità** delle particelle di un fluido in moto stazionario (Fig. 7.1). In questo caso si pone il problema della determinazione dei moti delle particelle, cioè delle *curve integrali* del campo.

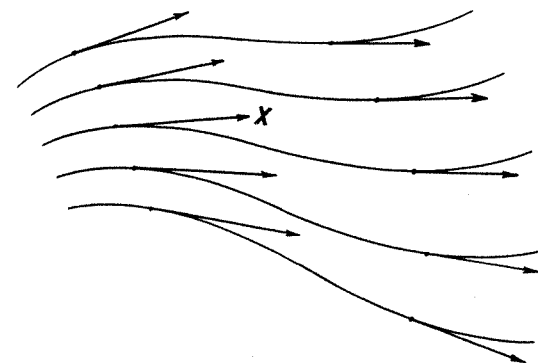


Fig. 7.1: campo di velocità di un fluido in moto stazionario.

DEFINIZIONE 1. Dicesi **curva integrale** di un campo vettoriale $\mathbf{X}$ una curva $\gamma: I \rightarrow \mathcal{A}$, con l'intervallo di definizione $I \subset \mathbb{R}$ contenente lo zero, tale che in ogni suo punto $\gamma(t)$ il vettore tangente $\dot{\gamma}(t)$ coincide con il valore del campo $\mathbf{X}$ in quel punto: $\dot{\gamma}(t) = \mathbf{X}(\gamma(t))$, cioè tale che

$$(1) \quad \dot{\gamma} = \mathbf{X} \circ \gamma$$

con a secondo membro la composizione delle applicazioni $\gamma: I \rightarrow \mathcal{A}$ e $\mathbf{X}: \mathcal{A} \rightarrow \mathcal{A} \times E$. Diciamo inoltre che la curva integrale è **basata nel punto** P_0 se $\gamma(0) = P_0$, cioè se per $t = 0$ la curva "passa" per il punto P_0 .

Le curve integrali rappresentano i moti delle particelle del fluido secondo l'interpretazione del campo $\mathbf{X}$ come campo di velocità. Un campo vettoriale, interpretato come campo di velocità, viene detto **sistema dinamico**.

Nella rappresentazione vettoriale, una curva $\mathbf{x} = \gamma(t)$ è curva integrale del campo $\mathbf{X}$ se soddisfa identicamente all'equazione differenziale vettoriale

$$(2) \quad \boxed{\frac{d\mathbf{x}}{dt} = \mathbf{X}(\mathbf{x})}$$

che in coordinate cartesiane si spezza in un sistema di n equazioni

$$(3) \quad \boxed{\frac{dx^\alpha}{dt} = X^\alpha(x^1, \dots, x^n) \quad (\alpha = 1, \dots, n)}$$

dove a secondo membro compaiono le componenti del campo (funzioni delle coordinate). Le rappresentazioni parametriche delle curve integrali sono pertanto tutte e sole le soluzioni di questo sistema. Se si considerano coordinate non cartesiane (q^i) la (2) si traduce ancora in un sistema di n equazioni dello stesso tipo:

$$(3') \quad \boxed{\frac{dq^i}{dt} = X^i(q^1, \dots, q^n) \quad (i = 1, \dots, n)}$$

In entrambi i casi si tratta di un **sistema di n equazioni differenziali ordinarie del primo ordine**, perché coinvolge solo le derivate prime delle funzioni incognite, in forma **normale**, perché le derivate prime sono esplicitate a primo membro, e **autonomo**, perché la variabile indipendente t non compare esplicitamente nei secondi membri.

Integrare il sistema (3) (o il sistema (3')) significa determinare *tutte* le sue possibili soluzioni, quindi tutte le possibili curve integrali del campo, il cui insieme forma il cosiddetto **spazio delle soluzioni** o **spazio dei moti**. Per distinguere una soluzione dall'altra risulta conveniente fare riferimento alle **condizioni iniziali** (o **dati iniziali**). Le condizioni iniziali impongono la posizione del punto mobile, cioè gli n valori delle sue coordinate, per $t = 0$. Nei corsi di Analisi Matematica si studiano le condizioni necessarie e sufficienti per l'esistenza e l'unicità delle soluzioni del sistema (3) con assegnate condizioni iniziali: **teorema di Cauchy** (Augustin Louis Cauchy, 1789-1857). Si può dimostrare che:

TEOREMA 1. Sia $\mathbf{X}$ un campo vettoriale di classe C^k ($k \geq 1$) su di un dominio M . Comunque si fissi il punto $P_0 \in M$ esiste una ed una sola **curva integrale massimale** basata in P_0 .

Va osservato che, subordinatamente alle ipotesi di regolarità del campo, di curve integrali basate in P_0 ne esistono infinite. Tutte quante però coincidono nelle intersezioni dei loro rispettivi intervalli di definizione (che contengono tutti lo zero). Vale a dire: se $\gamma_1: I_1 \rightarrow \mathcal{A}$ e $\gamma_2: I_2 \rightarrow \mathcal{A}$ sono due curve integrali del campo basate nello stesso punto P_0 , allora $\gamma_1|_{I_1 \cap I_2} = \gamma_2|_{I_1 \cap I_2}$. La curva integrale massimale basata in P_0 , che denotiamo con

$$\gamma_{P_0}: I_{P_0} \rightarrow M$$

è quella curva integrale tale che

$$I \subset I_{P_0}, \quad \gamma_{P_0}|_I = \gamma$$

per ogni altra curva integrale $\gamma: I \rightarrow \mathcal{A}$ basata in P_0 .

In genere l'intervallo I_{P_0} di definizione della curva integrale massimale dipende dal punto base P_0 e può anche non coincidere con tutto l'asse reale. Se accade però che tutte le curve integrali massimali sono definite su tutto l'asse reale, cioè che $I_{P_0} = \mathbb{R}$ per ogni P_0 , allora si dice che il **campo vettoriale è completo**.

In ogni caso, poiché per ogni condizione iniziale P_0 risulta determinata una ed una sola curva integrale massimale, possiamo descrivere lo spazio delle soluzioni, cioè l'insieme di *tutte* le curve integrali del campo $\mathbf{X}$, mediante una sola applicazione

$$\varphi: D \subset \mathbb{R} \times M \rightarrow M$$

ponendo

$$(4) \quad \varphi(t, P_0) = \gamma_{P_0}(t)$$

Ad ogni coppia (t, P_0) tale che $t \in I_{P_0}$ l'applicazione φ fa corrispondere il valore per t della curva integrale massimale γ_{P_0} basata in P_0 . L'applicazione φ prende il nome di **flusso** del campo $\mathbf{X}$. Il teorema di Cauchy si completa allora nel seguente enunciato, che mette in evidenza le proprietà del flusso (cioè la dipendenza delle soluzioni dai dati iniziali).

TEOREMA 2. Se X è un campo vettoriale di classe C^k ($k \geq 1$) su di un dominio M , allora:

- (i) Il dominio D del flusso corrispondente è un aperto di $\mathbb{R} \times M$ e φ è di classe C^k su D ⁽¹⁾.
- (ii) Se U è un aperto di M e $\delta > 0$ un numero tale che $(-\delta, \delta) \times U \subset D$, allora per ogni $t \in (-\delta, \delta)$ l'applicazione

$$\varphi_t: U \rightarrow U_t: P \mapsto \varphi(t, P)$$

è un omeomorfismo di classe C^k ⁽¹⁾ di U sopra un aperto $U_t \subset M$ e $\varphi_{-t}: P \mapsto \varphi(-t, P)$ è l'omeomorfismo inverso.

- (iii) Sussiste l'uguaglianza

$$(5) \quad \varphi(t, \varphi(s, P)) = \varphi(t+s, P),$$

per ogni terna (t, s, P) per cui i due membri hanno significato.

OSSERVAZIONE 1. Nel caso di un campo completo risulta $D = \mathbb{R} \times M$ e per ogni $t \in \mathbb{R}$ l'applicazione

$$(6) \quad \varphi_t: M \rightarrow M: P \mapsto \varphi(t, P) = \gamma_P(t)$$

è una trasformazione (di classe C^k) del dominio M del campo. La (5) si traduce nella seguente **proprietà di gruppo o proprietà di evoluzione**,

$$(7) \quad \boxed{\varphi_t \circ \varphi_s = \varphi_{t+s}}$$

dalla quale segue

$$(8) \quad \boxed{\varphi_t \circ \varphi_s = \varphi_s \circ \varphi_t, \quad \varphi_0 = \text{id}_M, \quad (\varphi_t)^{-1} = \varphi_{-t}}$$

L'insieme delle trasformazioni $\{\varphi_t; t \in \mathbb{R}\}$ forma dunque un gruppo abeliano, detto **gruppo di trasformazioni ad un parametro**.

OSSERVAZIONE 2. La proprietà (iii) del Teorema 2, cioè la proprietà di gruppo (7)₁, è una conseguenza diretta, oltre che dell'unicità della soluzione determinata da un dato iniziale, della seguente notevole

⁽¹⁾ Si veda l'Oss. 4, più avanti.

proprietà del sistema differenziale (2): lo spazio delle soluzioni è invariante rispetto alle *traslazioni temporali*, vale a dire, se $\gamma(t)$ è una curva integrale (definita sull'intervallo I) allora per ogni fissato numero reale $s \in I$ la curva $\gamma_s(t) = \gamma(t+s)$ è ancora una curva integrale (definita sull'intervallo $I_s = I - s$). Si vede infatti che se $x^\alpha = \gamma^\alpha(t)$ sono delle funzioni che risolvono il sistema (3) e se a queste funzioni si sostituiscono le funzioni $x^\alpha = \gamma_s^\alpha(t) = \gamma^\alpha(t+s)$ il sistema risulta ancora soddisfatto perché, posto $u = t+s$,

$$\frac{d\gamma_s^\alpha(t)}{dt} = \frac{d\gamma^\alpha(u)}{du} \frac{du}{dt} = \frac{d\gamma^\alpha(u)}{du} = X^\alpha(\gamma^\beta(u)) = X^\alpha(\gamma_s^\beta(t)).$$

Ciò posto, per come è stato definito il flusso si può affermare da un lato che l'applicazione $t \mapsto \varphi(t, \varphi(s, P))$ è la curva integrale massimale basata in $\varphi(s, P)$ e dall'altro, per la proprietà di invarianza, che la curva $t \mapsto \varphi(t+s, P)$ è ancora una curva integrale. Quest'ultima è basata in $\varphi(0+s, P) = \varphi(s, P)$, dunque nello stesso punto della curva precedente. Per l'unicità le due curve coincidono: di qui l'uguaglianza $\varphi(t, \varphi(s, P)) = \varphi(t+s, P)$.

OSSERVAZIONE 3. Si può dimostrare che le proprietà descritte nell'Oss. 1 non sono solo necessarie ma anche sufficienti affinché un insieme di curve $\varphi(t, P)$ sia lo spazio delle curve integrali di un certo campo vettoriale.

OSSERVAZIONE 4. Il flusso φ associato ad un campo vettoriale X ammette una rappresentazione vettoriale del tipo

$$(9) \quad x = \varphi(t, x_0)$$

che in coordinate cartesiane si spezza in un sistema di n funzioni

$$(10) \quad x^\alpha = \varphi^\alpha(t, x_0^\beta).$$

Per il Teorema 2, punto (i), queste sono di classe C^k . La (9) rappresenta una famiglia di curve dipendenti da $x_0 = (x_0^\beta)$. Per ogni fissato valore di t le relazioni (10) sono invertibili (punto (ii) del Teorema 2), e quindi la matrice jacobiana

$$\left(\frac{\partial \varphi^\alpha}{\partial x_0^\beta} \right)$$

è ovunque regolare.

OSSERVAZIONE 5. Le immagini delle curve integrali di un campo vettoriale sono dette **orbite del campo**. Il loro insieme costituisce il cosiddetto **ritratto di fase** del campo. Si noti che un'orbita è l'immagine di tutte le curve integrali basate nei suoi punti. Nei testi di Fisica le orbite di un campo vettoriale vengono dette **linee di flusso**.

OSSERVAZIONE 6. La completezza o meno di un campo vettoriale può essere determinata da considerazioni di ordine topologico. Si può per esempio dimostrare che se un campo vettoriale ha supporto compatto allora è completo. Il supporto è la chiusura dell'insieme dei punti non singolari del campo (si veda l'Oss. 7 seguente).

OSSERVAZIONE 7. Un punto P si dice **punto singolare** di un campo vettoriale $\mathbf{X}$ se $\mathbf{X}(P) = 0$. Una curva integrale massimale basata in un punto singolare è costante: assume cioè sempre il valore P per ogni $t \in \mathbb{R}$. Il flusso del campo lascia quindi unito ogni suo punto singolare. L'analisi del comportamento di un campo nell'intorno dei suoi punti singolari riveste particolare interesse. Si osservi che in corrispondenza ai punti singolari si annullano tutti i secondi membri del sistema differenziale (3).

OSSERVAZIONE 8. La nozione di campo vettoriale e quindi di sistema dinamico, qui introdotta per gli spazi affini, può in realtà prescindere dalla struttura affine ed estendersi non solo alle superfici ma anche a spazi di natura più generale: le *varietà differenziabili*.

8. Sistemi dinamici sulla retta

Adattiamo i concetti ed i teoremi enunciati nel § precedente al caso $n = 1$, cioè al caso in cui lo spazio affine è la retta reale $\mathbb{R}$. La retta reale ha una naturale struttura di spazio affine: lo spazio vettoriale soggiacente è ancora $\mathbb{R}$ e si pone $\delta(x, y) = y - x$. Il più semplice ed elementare modello matematico della meccanica è il **moto di un punto su di una retta**. Assegnata sulla retta una coordinata affine x , avente origine in un punto O , il moto di un punto su di questa è rappresentato da una funzione $x = \gamma(t)$ nella variabile reale t (il tempo). Questa funzione determina la **posizione** del punto sulla retta in corrispondenza

ad ogni istante t appartenente al campo di definizione della funzione (che supporremo essere un intervallo, limitato o no, estremi inclusi o no). La derivata di questa funzione, ammessa la sua esistenza, rappresenta la **velocità istantanea** del punto. La denotiamo v o anche con $\dot{x}$. La derivata seconda, cioè la derivata prima della velocità, rappresenta a sua volta l'**accelerazione istantanea** del punto. La denotiamo con a o anche con $\ddot{x}$.

Un modo per *generare* dei moti su di una retta è quello di fissare a *priori* un legame tra la posizione del punto e la sua velocità, legame che può essere imposto da un'equazione del tipo

$$(1) \quad F(x, \dot{x}) = 0,$$

con F prefissata funzione reale nelle due variabili reali $(x, \dot{x})$, oppure, più in particolare, da un'equazione del tipo

$$(2) \quad \dot{x} = X(x),$$

dove X è una funzione reale nella variabile reale x . In questo secondo caso dunque la velocità del punto è univocamente determinata dalla sua posizione. Si osservi che un'equazione del tipo (2) può sempre porsi nella forma (1): $X(x) - \dot{x} = 0$. Viceversa un'equazione del tipo (1) può porsi nella forma (2) solo quando è possibile *risolvere* la (1) rispetto alla variabile $\dot{x}$. È conveniente, per rendersi conto della questione, interpretare l'equazione (1) come equazione di una curva sul piano $\mathbb{R}^2$ di coordinate $(x, \dot{x})$. Un'equazione del tipo (1) prende il nome di **equazione differenziale del primo ordine autonoma in forma implicita**, un'equazione del tipo (2) di **equazione differenziale del primo ordine autonoma in forma normale**.

Assegnata un'equazione del tipo (1) (ovvero (2)), ci si pone il problema di determinare le sue soluzioni, cioè tutte le funzioni (o moti sulla retta) $x = x(t)$ che la soddisfano identicamente, posto $\dot{x} = \frac{dx}{dt}$. Il caso dell'equazione normale (2) è il caso particolare ($n = 1$) della teoria generale vista al § precedente. La funzione $X(x)$ rappresenta la componente di un campo vettoriale sulla retta.

Per una migliore comprensione degli esempi che seguiranno conviene precisare meglio le condizioni di esistenza ed unicità delle soluzioni di un'equazione del tipo (2). Si può dimostrare che:

TEOREMA. (i) Se la funzione $X(x)$ è continua in un punto x_0 , allora esiste una soluzione tale che $x(0) = x_0$. (ii) Se la funzione $X(x)$

è lipschitziana in x_0 , allora esiste un'unica soluzione $x(t)$ definita in un intervallo massimale I_{x_0} tale che $x(0) = x_0$. (iii) Se la funzione $X(x)$ è uniformemente lipschitziana, allora per ogni x_0 si ha $I_{x_0} = \mathbb{R}$ (il campo è completo).

Si ricordi che una funzione $X(x)$ è detta **lipschitziana** (Rudolf Lipschitz, 1832-1903) in un punto x_0 se in un intorno di quel punto essa ha rapporto incrementale limitato, vale a dire se esiste un intorno U di x_0 ed un numero reale positivo A tale che

$$|X(x_1) - X(x_2)| < A |x_1 - x_2|, \quad \forall x_1, x_2 \in U.$$

La lipschitzianità implica la continuità ed è implicata dalla derivabilità; è cioè una condizione *intermedia* tra la continuità e la derivabilità. Una funzione $X(x)$ è detta **uniformemente lipschitziana** se, comunque si fissi un $\varepsilon > 0$ esiste un $A > 0$ tale che

$$|X(x_1) - X(x_2)| < A |x_1 - x_2|, \quad \forall x_1, x_2 : |x_1 - x_2| < \varepsilon.$$

ESEMPIO 1. Si consideri la semplice equazione differenziale

$$(3) \quad \boxed{\dot{x} = x}$$

La funzione $X(x) = x$ a secondo membro è di classe C^∞ su tutto l'asse reale. Vale pertanto il teorema di Cauchy di esistenza ed unicità delle soluzioni. Il suo integrale generale è $x = c e^t$ con c costante arbitraria. Siccome per $t = 0$ si ha $x(0) = c$ questa costante coincide con il dato iniziale x_0 . Ne segue che ogni curva integrale è definita su tutto l'asse reale (il campo X è completo) ed il flusso, cioè l'insieme di tutte le sue curve integrali, è dato dalla funzione

$$\varphi(t, x_0) = x_0 e^t.$$

La proprietà di gruppo (5), §3 è chiaramente soddisfatta:

$$\varphi(s, \varphi(t, x_0)) = \varphi(t, x_0) e^s = x_0 e^t e^s = x_0 e^{t+s} = \varphi(t+s, x_0).$$

ESEMPIO 2. Nell'equazione

$$(4) \quad \boxed{\dot{x} = x^2}$$

il secondo membro è di classe C^∞ su tutto $\mathbb{R}$, pertanto esiste ed è unica la soluzione per ogni dato iniziale. Quest'equazione, come la precedente e come tutte le equazioni del tipo

$$(5) \quad \boxed{\dot{x} = f(x)}$$

si integra per separazione delle variabili: la si riscrive sotto la forma

$$\frac{dx}{f(x)} = dt$$

e si integrano ambo i membri. Nel nostro caso dalla scrittura

$$\frac{dx}{x^2} = dt$$

segue $t + c = -x^{-1}$ e quindi $x = -(t + c)^{-1}$. Ponendo $t = 0$ si vede che $x_0 = -\frac{1}{c}$, per cui in conclusione il flusso risulta dato da

$$\varphi(t, x_0) = \frac{x_0}{1 - x_0 t}.$$

Come si vede le curve integrali non sono definite su tutto l'asse reale: il punto $t = \frac{1}{x_0}$ è una singolarità. Si noti a questo proposito che la funzione $X(x) = x^2$ non è uniformemente lipschitziana (si veda il Teorema). L'intervallo di definizione delle soluzioni è

$$I_{x_0} = \begin{cases} (-\infty, \frac{1}{x_0}) & \text{per } x_0 > 0, \\ \mathbb{R} & \text{per } x_0 = 0, \\ (\frac{1}{x_0}, +\infty) & \text{per } x_0 < 0, \end{cases}$$

La proprietà di gruppo si verifica immediatamente:

$$\begin{aligned} \varphi(s, \varphi(t, x_0)) &= \frac{\varphi(t, x_0)}{1 - \varphi(t, x_0) s} = \frac{x_0}{1 - x_0 t} \frac{1}{1 - \frac{s x_0}{1 - x_0 t}} \\ &= \frac{x_0}{1 - x_0 (t + s)} = \varphi(t + s, x_0) \end{aligned}$$

ESEMPIO 3. Nell'equazione

$$(6) \quad \boxed{\dot{x} = |x|}$$

il secondo membro è una funzione definita su tutto $\mathbb{R}$, non differenziabile nell'origine ma uniformemente lipschitziana: per ogni coppia di numeri (x_1, x_2) si ha infatti sempre $|X(x_1) - X(x_2)| = ||x_1| - |x_2|| \leq |x_1 - x_2|$. Dobbiamo dunque attenderci l'esistenza e l'unicità delle soluzioni per ogni dato iniziale e la loro estendibilità a tutti i valori $t \in \mathbb{R}$. In effetti l'equazione si spezza in due equazioni

$$\dot{x} = \begin{cases} x, & \text{per } x \geq 0, \\ -x, & \text{per } x < 0, \end{cases}$$

che si integrano in

$$\varphi(t, x_0) = \begin{cases} x_0 e^t, & \text{per } x_0 \geq 0, \\ x_0 e^{-t}, & \text{per } x_0 < 0. \end{cases}$$

ESEMPIO 4. Si consideri l'equazione differenziale implicita

$$(7) \quad \boxed{\dot{x}^2 - x = 0}$$

Questa si spezza in due equazioni differenziali distinte in forma normale sulla semiretta reale $x \geq 0$:

$$(8) \quad \dot{x} = \begin{cases} +\sqrt{x}, \\ -\sqrt{x}. \end{cases}$$

Per entrambe la condizione di Lipschitz non è soddisfatta nell'origine (il rapporto incrementale non è limitato). Integrandole separatamente (per separazione delle variabili) si trovano rispettivamente le soluzioni

$$x(t) = \left(\sqrt{x_0} + \frac{t}{2}\right)^2, \quad x(t) = \left(\sqrt{x_0} - \frac{t}{2}\right)^2.$$

Per $x_0 = 0$ le due soluzioni coincidono, ma si osserva che anche $x(t) = 0$ è una soluzione. Dunque per ciascuna delle equazioni normali (8) non vi è unicità in corrispondenza al dato iniziale $x_0 = 0$. Le soluzioni trovate confermano la seguente proprietà riscontrabile direttamente dall'equazione (7): l'insieme delle soluzioni dell'equazione (7) è invariante rispetto all'inversione temporale $t \rightarrow -t$, vale a dire se $x(t)$ è una soluzione, $x(-t)$ è ancora una soluzione. Per ognuno dei flussi generati dalle

equazioni (8) continua a valere la proprietà di gruppo. Posto infatti che nel primo caso

$$\varphi(t, x_0) = \left(\sqrt{x_0} + \frac{t}{2}\right)^2,$$

si ha successivamente:

$$\varphi(s, \varphi(t, x_0)) = \left(\sqrt{\varphi(t, x_0)} + \frac{s}{2}\right)^2 = \left(\sqrt{x_0} + \frac{t}{2} + \frac{s}{2}\right)^2 = \varphi(t+s, x_0).$$

ESEMPIO 5. Si consideri infine l'equazione differenziale implicita

$$(9) \quad \boxed{(\dot{x} + 1)^2 - x = 0}$$

Non esiste alcuna soluzione per il dato iniziale $x_0 = 0$. In corrispondenza a tale posizione iniziale si avrebbe infatti la velocità iniziale $\dot{x}_0 = -1 < 0$ che tenderebbe a spostare il punto sul semiasse $x < 0$ dove l'equazione non è definita.

9. Sistemi dinamici nel piano

Nel piano affine riferito a coordinate cartesiane (x, y) un campo vettoriale ammette una rappresentazione del tipo

$$(1) \quad \mathbf{X} = X(x, y) \mathbf{i} + Y(x, y) \mathbf{j}.$$

Il corrispondente sistema di equazioni del primo ordine è di conseguenza

$$(2) \quad \begin{cases} \dot{x} = X(x, y), \\ \dot{y} = Y(x, y). \end{cases}$$

Consideriamo in questo paragrafo alcuni esempi elementari. Altri esempi sono visti al §10.

ESEMPIO 1. Al campo vettoriale definito da $\mathbf{X}(\mathbf{x}) = \mathbf{x} = x \mathbf{i} + y \mathbf{j}$ corrisponde il sistema di equazioni

$$(3) \quad \begin{cases} \dot{x} = x, \\ \dot{y} = y. \end{cases}$$

Si tratta di equazioni differenziali *separate* cioè coinvolgenti ciascuna una singola coordinata. La curva integrale basata in $P_0 = (x_0, y_0)$ ha equazioni

$$x = x_0 e^t, \quad y = y_0 e^t,$$

ed è definita su tutto l'asse reale. Il campo è pertanto completo. Il flusso è dato da

$$\varphi(t; x_0) = e^t x_0,$$

posto $x_0 = x_0 i + y_0 j$. Le orbite delle curve integrali non basate nell'origine sono le semirette aperte uscenti dall'origine, che è punto singolare. Per $t \rightarrow -\infty$ le curve integrali tendono all'origine (Fig. 9.1). La trasformazione φ_t è l'omotetia di centro l'origine e coefficiente e^t . Il campo è invariante rispetto a rotazioni intorno all'origine.

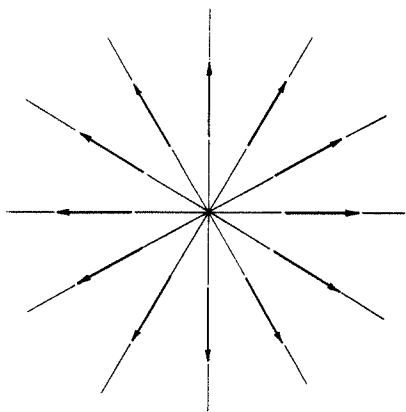


Fig. 9.1

ESEMPIO 2. Si consideri la famiglia di curve del piano rappresentate dalle equazioni parametriche

$$(4) \quad \begin{cases} x = x_0 + y_0 t + \frac{1}{2} g t^2, \\ y = y_0 + g t. \end{cases}$$

Posto (utilizziamo la notazione matriciale per i vettori)

$$\begin{pmatrix} x \\ y \end{pmatrix} = \varphi_t \begin{pmatrix} x_0 \\ y_0 \end{pmatrix},$$

le applicazioni $\varphi_t: \mathbb{R}^2 \rightarrow \mathbb{R}^2$ soddisfano alle proprietà di gruppo. Infatti:

$$\begin{aligned} \varphi_s \left(\varphi_t \begin{pmatrix} x_0 \\ y_0 \end{pmatrix} \right) &= \varphi_s \begin{pmatrix} x_0 + y_0 t + \frac{1}{2} g t^2 \\ y_0 + g t \end{pmatrix} \\ &= \begin{pmatrix} x_0 + y_0 t + \frac{1}{2} g t^2 + (y_0 + g t) s + \frac{1}{2} g s^2 \\ y_0 + g t + g s \end{pmatrix} \\ &= \begin{pmatrix} x_0 + y_0 (t+s) + \frac{1}{2} g (t+s)^2 \\ y_0 + g (t+s) \end{pmatrix} = \varphi_{t+s} \begin{pmatrix} x_0 \\ y_0 \end{pmatrix}. \end{aligned}$$

Le curve considerate forniscono dunque lo spazio delle soluzioni di un sistema del tipo (2). Derivando si ottiene

$$\dot{x} = y_0 + g t = y, \quad \dot{y} = g.$$

Il campo vettoriale che genera il flusso considerato è pertanto

$$(5) \quad X = y i + g j.$$

ESEMPIO 3. *Equazioni differenziali ordinarie come sistemi dinamici.* Un'equazione differenziale del primo ordine in forma normale, cioè del tipo

$$(6) \quad y' = f(x, y),$$

può essere interpretata come sistema dinamico. Sia $D \subset \mathbb{R}^2$ il dominio (aperto) di definizione della $f(x, y)$. Si consideri su D il campo vettoriale

$$(7) \quad X = i + f(x, y)j.$$

Il sistema differenziale corrispondente è

$$(8) \quad \begin{cases} \dot{x} = 1, \\ \dot{y} = f(x, y). \end{cases}$$

Poiché la prima equazione dà $x = t + c$ con c costante arbitraria, la variabile x può essere identificata con t e quindi le soluzioni $y = y(t)$ sono

soluzioni dell'equazione differenziale (6). Anche un'equazione normale del secondo ordine

$$(9) \quad y'' = f(x, y, y'),$$

può essere interpretata come sistema dinamico. A questa si associa infatti il campo vettoriale

$$(10) \quad \mathbf{X} = \mathbf{i} + z\mathbf{j} + f(x, y, z)\mathbf{k},$$

definito nel dominio (aperto) $D \subset \mathbb{R}^3$ della funzione f , il cui sistema differenziale è

$$(11) \quad \begin{cases} \dot{x} = 1, \\ \dot{y} = z, \\ \dot{z} = f(x, y, z). \end{cases}$$

Le soluzioni $y = y(t)$ di questo sistema, identificato t con x grazie alla prima equazione, forniscono le soluzioni dell'equazione (9). Questo procedimento si estende alle equazioni differenziali di qualunque ordine, purché in forma normale, cioè con la derivata di ordine massimo esplicitata a primo membro.

ESEMPIO 4. Equazione del primo ordine non autonoma. Si supponga che il moto di un punto su di una retta sia governato da un'equazione del tipo

$$(12) \quad \dot{x} = F(x, t).$$

Questa impone che la velocità dipenda con legge ben definita non solo dalla posizione ma anche dal tempo t . Quest'equazione può trasformarsi in un sistema dinamico in $\mathbb{R}^2$:

$$(13) \quad \begin{cases} \dot{x} = F(x, y), \\ \dot{y} = 1. \end{cases}$$

La seconda equazione significa infatti $y = t + c$ con $c \in \mathbb{R}$ arbitrario. Le soluzioni della (12) sono dunque quelle del sistema (13) con dato iniziale $y_0 = 0$ (per il quale y si identifica con t).

ESEMPIO 5. Equazioni del secondo ordine sulla retta. Si supponga che il moto di un punto su di una retta sia dovuto ad un'equazione del tipo

$$(14) \quad \ddot{x} = F(x, \dot{x}).$$

Questa impone che la sua accelerazione sia una funzione ben determinata della sua posizione e della sua velocità. Quest'equazione equivale ad un sistema differenziale del primo ordine su $\mathbb{R}^2$:

$$(15) \quad \begin{cases} \dot{x} = y, \\ \dot{y} = F(x, y). \end{cases}$$

La coordinata y rappresenta dunque la velocità del punto. Pertanto, subordinatamente alle solite condizioni di regolarità della funzione F , l'equazione (14) ha una ed una sola soluzione per ogni coppia di dati iniziali (x_0, y_0) , posizione e velocità iniziali. Lo spazio delle soluzioni o dei moti dipende dunque da due parametri. Per esempio la semplice equazione del secondo ordine che regge il moto di un grave lungo una retta verticale (accelerazione costante)

$$(16) \quad \ddot{x} = g \quad (g \in \mathbb{R})$$

equivale al sistema del primo ordine

$$(17) \quad \begin{cases} \dot{x} = y, \\ \dot{y} = g, \end{cases}$$

cioè al campo vettoriale (5) già considerato nell'Esempio 2. Altri due semplici esempi tratti dalla meccanica di un punto su di una retta sono le equazioni

$$(18) \quad \ddot{x} = -\omega^2 x, \quad \ddot{x} = \omega^2 x.$$

La prima, che si scrive anche

$$(19) \quad \ddot{x} + \omega^2 x = 0,$$

è chiamata **equazione dell'oscillatore armonico a una dimensione** o **equazione del moto armonico** e regola il moto di un punto su di una retta soggetto ad una forza lineare attrattiva (forza elastica) centrata nell'origine. La seconda è l'**equazione del moto centrifugo**, dove il punto è invece soggetto ad una forza lineare repulsiva. Le equazioni (18) equivalgono rispettivamente ai sistemi

$$(20) \quad \begin{cases} \dot{x} = y, \\ \dot{y} = -\omega^2 x, \end{cases} \quad \begin{cases} \dot{x} = y, \\ \dot{y} = \omega^2 x. \end{cases}$$

Sul calcolo dei loro flussi si ritornerà al § seguente.

10. Integrali primi

La determinazione delle curve integrali e quindi del flusso di un dato campo vettoriale è un problema che può presentare notevoli difficoltà. Non esiste infatti alcun metodo *generale* per integrare un sistema di equazioni differenziali in forma normale. L'Analisi Matematica ha tuttavia sviluppato sia dei *metodi qualitativi* atti a stabilire il comportamento delle curve integrali (come la loro *stabilità*, la *periodicità*, etc.), sia dei *metodi numerici* che, con l'ausilio di elaboratori elettronici, ne consentono la tabulazione con l'approssimazione desiderata.

Nell'ambito di questi metodi, notevoli semplificazioni dei calcoli e utili informazioni sul comportamento delle curve integrali possono scaturire dalla conoscenza di certe funzioni, dette *integrali primi*, le quali godono della proprietà di mantenersi costanti lungo le curve integrali. Anzi, se si conoscono $n - 1$ integrali primi indipendenti (ma questo numero può scendere in particolari situazioni) si è in grado di determinare le curve integrali con una integrazione (o, come si usa dire, con una *quadratura*).

Gli integrali primi possono determinarsi risolvendo una particolare equazione alle derivate parziali. Si pone così un problema alternativo a quello dell'integrazione del sistema differenziale associato al campo, ma che può a sua volta presentare notevoli difficoltà. Sovente però la ricerca degli integrali primi ha successo, o perché si è in grado di integrare la corrispondente equazione differenziale, o perché il campo vettoriale gode di proprietà di *simmetria* o di *invarianza*, tali da produrre, in base a opportuni teoremi, degli integrali primi (esempi notevoli sono, nel corso di Meccanica Razionale, l'*integrale primo dell'energia*, l'*integrale primo della quantità di moto*, l'*integrale primo delle aree*).

DEFINIZIONE 1. Dicesi **funzione integrale** o **integrale primo** di un campo vettoriale $\mathbf{X}$ su di uno spazio affine $\mathcal{A}$ una funzione reale F sopra $\mathcal{A}$ che si mantiene costante lungo ogni curva integrale di $\mathbf{X}$, cioè tale che per ogni curva integrale $\gamma: I \rightarrow \mathcal{A}$

$$F \circ \gamma(t_1) = F \circ \gamma(t_2), \quad \forall t_1, t_2 \in I,$$

ovvero tale che,

$$(1) \quad \boxed{D(F \circ \gamma)(t) = 0, \quad \forall t \in I}$$

dove D è il simbolo di derivata di funzione a una variabile reale.

Perché la (1) abbia senso occorre ovviamente che la funzione $F \circ \gamma: I \rightarrow \mathbb{R}$ sia derivabile, condizione senz'altro soddisfatta nel caso in cui sia il campo vettoriale sia la funzione F sono di classe C^1 almeno. La condizione (1) si può anche scrivere, più sinteticamente,

$$(1') \quad \dot{F} = 0,$$

oppure

$$(1'') \quad \frac{dF}{dt} = 0.$$

OSSERVAZIONE 1. Si può più in generale intendere come *integrale primo* una funzione a valori in uno spazio vettoriale (per esempio lo spazio vettoriale E associato allo spazio affine), costante lungo le curve integrali. Incontreremo nel seguito degli esempi.

Per verificare che una funzione è un integrale primo occorre, secondo la Def. 1, conoscere le curve integrali del campo, le quali sono di solito incognite. Esiste tuttavia una proprietà caratteristica degli integrali primi che prescinde da tale conoscenza. Essa è illustrata nell'enunciato seguente.

PROPOSIZIONE 1. Condizione necessaria e sufficiente affinché una funzione F sia integrale primo di un campo vettoriale $\mathbf{X}$ è che la sua derivata rispetto a $\mathbf{X}$ sia nulla:

$$(2) \quad \boxed{\langle \mathbf{X}, dF \rangle = 0}$$

Dimostrazione. Per ogni curva integrale γ del campo $\mathbf{X}$ vale la formula

$$(3) \quad \boxed{D(F \circ \gamma) = \langle \mathbf{X}, dF \rangle \circ \gamma}$$

immediata conseguenza della Prop. 1 (formula (13)) del §5. Stante l'arbitrarietà di γ segue che la condizione (2) è equivalente alla (1). ■

In un qualunque sistema di coordinate (q^i) la (2) si traduce nell'equazione

$$(4) \quad \boxed{X^i \frac{\partial F}{\partial q^i} = 0}$$

Si tratta di un'equazione differenziale alle derivate parziali, del primo ordine, lineare, omogenea, nella funzione incognita F . La ricerca degli integrali primi è quindi ricondotta all'integrazione di quest'equazione.

OSSERVAZIONE 2. Dalla definizione di integrale primo segue che se $(F_1, F_2, \dots, F_m)$ sono m integrali primi, allora una qualunque loro composizione $f \circ (F_1, F_2, \dots, F_m)$ tramite una funzione $f: \mathbb{R}^m \rightarrow \mathbb{R}$ è ancora un integrale primo.

Tra le soluzioni dell'equazione (2) vi sono le funzioni costanti. Esse sono ovviamente prive di interesse (si dicono **integrali primi banali**). Si tratta quindi di stabilire se un campo vettoriale X ammette integrali primi non banali. In genere un campo vettoriale ammette **integrali primi locali**, vale a dire definiti nell'intorno di punti del suo dominio di definizione, ma può non ammettere **integrali primi globali**, cioè definiti su tutto il suo dominio di definizione. Integrali primi globali si possono costruire, in linea di principio, prolungando o giustapponendo integrali primi locali, tenendo presente l'Oss. 2. Tale operazione può tuttavia non dare esito positivo, come mostra l'esempio seguente.

ESEMPIO 1. Si consideri il campo vettoriale dell'Esempio 1 del §9: $X = xi + yj$. Se escludiamo l'origine e consideriamo per esempio la circonferenza S_1 di raggio unitario e centro l'origine, una qualunque funzione $H: S_1 \rightarrow \mathbb{R}$ sopra di questa si può estendere al piano $\mathbb{R}^2$, esclusa l'origine, per valori costanti lungo le semirette uscenti dall'origine. Siccome le orbite del campo sono tutte queste semirette (aperte) più l'origine, le funzioni così costruite, a partire da funzioni H di classe C^∞ , sono gli integrali primi di X . Si vede chiaramente che questi integrali primi non sono estendibili per continuità all'origine se non nel caso in cui la funzione H sulla circonferenza è costante: ma in questo caso l'integrale primo è anch'esso una funzione costante, quindi banale. Concludiamo allora che il campo X su $\mathbb{R}^2$ non ammette integrali primi globali non banali, pur ammettendo nell'intorno di ogni punto diverso dall'origine infiniti integrali primi.

L'importanza della conoscenza di integrali primi è messa in evidenza dalle considerazioni seguenti.

Sia F una funzione reale sopra lo spazio affine $\mathcal{A}$. Come si è visto al §6, gli insiemi $Q_c \subset \mathcal{A}$ definiti dall'equazione $F = c$ con $c \in \mathbb{R}$ prendono

il nome di **insiemi di livello** o **superfici di livello**. Alcuni di questi possono essere vuoti. In ogni caso essi formano una partizione in sottoinsiemi disgiunti di tutto il campo di definizione di F . Più in generale si può considerare un insieme di k campi scalari $(F_a) = (F_1, \dots, F_k)$ e per ogni $c = (c_1, \dots, c_k)$ considerare l'insieme di livello Q_c definito dalle equazioni

$$F_1 = c_1, \quad F_2 = c_2, \quad \dots \quad F_k = c_k.$$

Si ottiene così ancora una partizione (fogliettamento) in sottoinsiemi disgiunti del campo di definizione delle funzioni (F_a) .

PROPOSIZIONE 2. Se una curva integrale di X ha un punto di intersezione con un insieme di livello generato da uno o più integrali primi, allora è tutta contenuta in questo.

Si vuol dire che se $\gamma: I \rightarrow M$, con M dominio di definizione di X , è una curva integrale di X e se per un qualche $t_0 \in I$ si ha $\gamma(t_0) \in Q_c$ allora $\gamma(t) \in Q_c$ per ogni $t \in I$, vale a dire $\gamma(I) \subset Q_c$.

Dimostrazione. La condizione $P_0 = \gamma(t_0) \in Q_c$ implica $F(P_0) = F(\gamma(t_0)) = c$, quindi $F(\gamma(t)) = c$ per ogni t , e quindi $\gamma(t) \in Q_c$ per ogni t . ■

Dalla proprietà ora vista segue in particolare che se il numero degli integrali primi (F_a) è tale che gli insiemi di livello sono o punti o curve (non parametrizzate) allora essi si identificano con le orbite del campo (vale a dire con le immagini delle curve integrali del campo).

Consideriamo a questo proposito tre esempi, elementari ma significativi. Il primo (Esempio 2) mostra come si possa giungere alla completa determinazione delle curve integrali attraverso l'integrazione dell'equazione degli integrali primi. Il secondo (Esempio 3) è un esempio di sistema dinamico su di una superficie (il cilindro), le cui orbite si determinano attraverso un integrale primo. Il terzo (Esempio 4) è un esempio di sistema dinamico dove la periodicità delle curve integrali è riconosciuta per mezzo di un integrale primo.

ESEMPIO 2. **L'oscillatore armonico.** Sul piano affine riferito a coordinate cartesiane (x, y) si consideri il campo (si veda la fine del §9)

$$X = yi - \omega^2 xj$$

il cui sistema differenziale del primo ordine è

$$(5) \quad \begin{cases} \dot{x} = y, \\ \dot{y} = -\omega^2 x. \end{cases}$$

Posto che

$$\dot{F} = \frac{\partial F}{\partial x} \dot{x} + \frac{\partial F}{\partial y} \dot{y} = \frac{\partial F}{\partial x} y - \frac{\partial F}{\partial y} \omega^2 x,$$

l'equazione (4) degli integrali primi diventa:

$$(6) \quad y \frac{\partial F}{\partial x} - \omega^2 x \frac{\partial F}{\partial y} = 0.$$

Questa si può integrare per *separazione delle variabili*. Si cerca cioè una sua soluzione del tipo $F(x, y) = A(x) + B(y)$. Con quest'ipotesi l'equazione (6) diventa (' è il simbolo di derivata prima)

$$y A'(x) - \omega^2 x B'(y) = 0,$$

e quindi, subordinatamente alla condizione che i denominatori non si annullino,

$$\frac{y}{B'(y)} = \omega^2 \frac{x}{A'(x)}.$$

Poiché il primo membro è funzione solo della x e il secondo solo della y , entrambi devono essere costanti. Posta per esempio questa costante uguale a 1, quest'ultima equazione si spezza in due equazioni differenziali ordinarie separate, coinvolgenti cioè ciascuna una singola variabile:

$$\omega^2 \frac{x}{A'(x)} = 1, \quad \frac{y}{B'(y)} = 1,$$

ovvero

$$A'(x) = \omega^2 x, \quad B'(y) = y.$$

Si ha quindi, a meno di inessenziali costanti additive,

$$A(x) = \frac{1}{2} \omega^2 x^2, \quad B(y) = \frac{1}{2} y^2.$$

Il procedimento della separazione delle variabili ha quindi successo. Omettendo l'inessenziale fattore $\frac{1}{2}$, abbiamo infatti trovato il seguente integrale primo, soluzione della (6):

$$F(x, y) = \omega^2 x^2 + y^2.$$

L'insieme Q_c definito dall'equazione $F = c$ è vuoto (nel campo reale) per $c < 0$; si riduce all'origine per $c = 0$ ed è un'ellisse per $c > 0$. Infatti in quest'ultimo caso

$$(7) \quad \omega^2 x^2 + y^2 = b^2$$

posto $c = b^2$. Si tratta di un'ellisse di semiassi $a = \frac{b}{\omega}$ e b . Al variare di c si ottiene una famiglia di ellissi, che sono orbite del campo $\mathbf{X}$ (percorse in senso orario, come si riconosce facilmente valutando il campo in qualche punto del piano, per esempio sull'asse x , vedi Fig. 10.1). Concludiamo pertanto che le curve integrali sono *chiuse* e quindi *periodiche*. Dalla (7) si trae

$$y = \pm b \sqrt{1 - \frac{x^2}{a^2}} \quad \left(a = \frac{b}{\omega} \right).$$

Scegliendo il segno + e sostituendo nella prima delle equazioni (5), si ottiene un'equazione differenziale del primo ordine nella sola x , a variabili separabili:

$$\frac{dx}{dt} = b \sqrt{1 - \frac{x^2}{a^2}}.$$

Il suo integrale generale risulta essere

$$(8) \quad x = a \sin(\omega t + \phi),$$

con ϕ costante arbitraria. Sostituendo questa funzione nella seconda delle (6), si trova

$$\frac{dy}{dt} = -\omega^2 a \sin(\omega t + \phi),$$

quindi:

$$(9) \quad y = b \cos(\omega t + \phi).$$

La curva integrale così determinata, di equazioni parametriche (8) e (9), è basata nel punto $(x_0, y_0) = (a \sin \phi, b \cos \phi)$. Per ottenere il flusso del campo occorre sostituire le costanti (a, b) con i dati iniziali (x_0, y_0) . Sviluppando la (8) e la (9) si trova:

$$(10) \quad \begin{cases} x = x_0 \cos \omega t + \frac{1}{\omega} y_0 \sin \omega t, \\ y = -\omega x_0 \sin \omega t + y_0 \cos \omega t. \end{cases}$$

È interessante osservare che in questo caso il legame tra il vettore $\mathbf{x}_0 = (x_0, y_0)$ e il vettore $\mathbf{x} = (x, y)$ è lineare, sicché il legame $\mathbf{x} = \varphi(t, \mathbf{x}_0)$ può porsi in forma matriciale:

$$(10') \quad \begin{pmatrix} x \\ y \end{pmatrix} = \begin{pmatrix} \cos \omega t & \frac{1}{\omega} \sin \omega t \\ -\omega \sin \omega t & \cos \omega t \end{pmatrix} \begin{pmatrix} x_0 \\ y_0 \end{pmatrix}.$$

In altre parole: il gruppo ad un parametro $\{\varphi_t; t \in \mathbb{R}\}$ generato dal campo $\mathbf{X}$ è un gruppo di trasformazioni lineari rappresentato dalle matrici

$$\varphi_t = \begin{pmatrix} \cos \omega t & \frac{1}{\omega} \sin \omega t \\ -\omega \sin \omega t & \cos \omega t \end{pmatrix}$$

È un utile esercizio verificare che queste matrici godono della proprietà di gruppo $\varphi_t \circ \varphi_s = \varphi_{t+s}$.

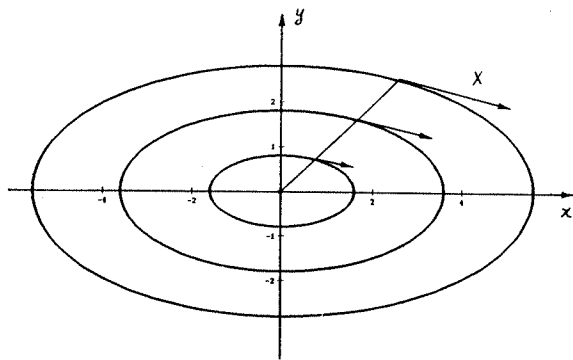


Fig. 10.1: le orbite dell'oscillatore armonico.

OSSERVAZIONE 3. Il sistema dinamico ora considerato rientra in una particolare classe di sistemi dinamici, i **sistemi dinamici lineari**, cioè del tipo

$$\dot{\mathbf{x}} = \mathbf{A} \mathbf{x}$$

dove $\mathbf{A}$ è una matrice costante. Il flusso generato dal sistema (10) è ancora lineare e dato dalla matrice

$$\varphi_t = \exp(\mathbf{A} t) = \sum_{k=0}^{+\infty} \frac{t^k}{k!} \mathbf{A}^k.$$

Nel caso precedente è

$$\mathbf{A} = \begin{pmatrix} 0 & 1 \\ -\omega^2 & 0 \end{pmatrix}$$

e quindi, essendo

$$\mathbf{A}^2 = -\omega^2 \mathbf{1},$$

risulta

$$\varphi_t = \cos \omega t \mathbf{1} + \frac{1}{\omega} \sin \omega t \mathbf{A},$$

a conferma della (10).

ESEMPIO 3. Il pendolo semplice. Consideriamo lo spazio affine tridimensionale riferito sia a coordinate cartesiane (x, y, z) , sia a coordinate cilindriche (r, ϑ, z) . Consideriamo la superficie regolare Q definita dall'equazione $r = R$, con R costante positiva. Si tratta di un cilindro di raggio R e asse l'asse delle z . Intendendo l'angolo ϑ variabile nell'intervallo aperto $(-\pi, \pi)$, le coordinate (ϑ, z) si possono interpretare come coordinate superficiali, che mappano il cilindro Q , tolta una sua direttrice, nella striscia aperta di $\mathbb{R}^2$ definita da $\{(\vartheta, z) \in \mathbb{R}^2 \mid -\pi < \vartheta < \pi\}$. Introdotti i vettori $(\mathbf{E}_i)$ associati a queste, cioè i vettori $(\mathbf{E}_\vartheta = R\boldsymbol{\tau}, \mathbf{E}_z = \mathbf{k})$, si consideri il campo vettoriale

$$(11) \quad \mathbf{X} = z \mathbf{E}_\vartheta - \omega^2 \sin \vartheta \mathbf{E}_z,$$

prolungato per continuità sulla direttrice omessa. Il sistema differenziale del primo ordine corrispondente è:

$$(12) \quad \begin{cases} \frac{d\vartheta}{dt} = z, \\ \frac{dz}{dt} = -\omega^2 \sin \vartheta. \end{cases}$$

La combinazione di queste due equazioni fornisce l'equazione differenziale del secondo ordine

$$(13) \quad \frac{d^2 \vartheta}{dt^2} + \omega^2 \sin \vartheta = 0,$$

detta **equazione del pendolo semplice**. L'equazione (4) degli integrali primi diventa:

$$(14) \quad z \frac{\partial F}{\partial \vartheta} - \omega^2 \sin \vartheta \frac{\partial F}{\partial z} = 0.$$

Anche in questo caso, come nel precedente, il procedimento di integrazione per separazione delle variabili ha successo. Si trova, a conti fatti, un integrale primo del tipo

$$(15) \quad F(\vartheta, z) = \frac{1}{2}z^2 - \omega^2 \cos \vartheta.$$

Il campo ha due punti singolari: $P_0 = (\vartheta = 0, z = 0)$ e il punto opposto $P_1 = (\pm\pi, 0)$. Le orbite di $\mathbf{X}$ sono le curve Q_c sul cilindro di equazione

$$(16) \quad \frac{1}{2}z^2 - \omega^2 \cos \vartheta = c,$$

che possiamo anche scrivere

$$(17) \quad z^2 = 2\omega^2(\cos \vartheta + e), \quad e = \frac{c}{\omega^2}.$$

Per la realtà di z deve essere $c + \omega^2 \cos \vartheta \geq 0$, quindi $-\omega^2 \leq c < +\infty$, vale a dire $-1 \leq e < +\infty$. Si hanno allora quattro possibilità. (I) Per $c = -\omega^2$ (cioè $e = -1$) dalla (16) si vede che non può che essere $\vartheta = z = 0$. In questo caso l'orbita Q_c si riduce al punto singolare P_0 . (II) Per $c \in (-\omega^2, \omega^2)$ (cioè $e \in (-1, 1)$) Q_c è una curva chiusa simmetrica rispetto a P_0 , all'asse z e all'asse ϑ , di equazione

$$(18) \quad z^2 = 2\omega^2(\cos \vartheta + \cos \vartheta_0) \quad (e = \cos \vartheta_0).$$

(III) Per $c = \omega^2$ (cioè $e = 1$) l'orbita Q_c ha equazione

$$z^2 = 2\omega^2(\cos \vartheta + 1) = 4\omega^2 \cos^2 \left(\frac{\vartheta}{2} \right),$$

e si spezza pertanto nelle due curve

$$z = \pm 2\omega \cos \left(\frac{\vartheta}{2} \right)$$

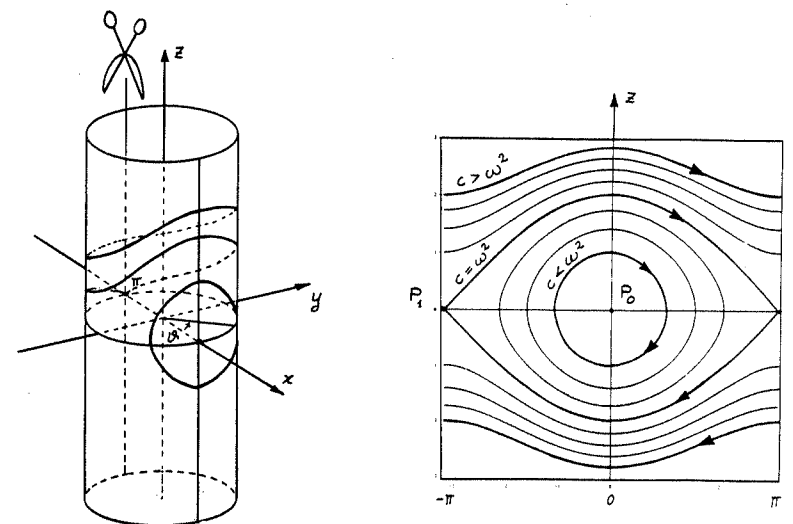


Fig. 10.2: le orbite del pendolo semplice.

che si chiudono nel secondo punto singolare P_1 e che chiamiamo **curve limite**. (IV) Per $c > \omega^2$ (cioè $e > 1$) Q_c si sdoppia in due curve chiuse simmetriche esterne alle curve limite (Fig. 10.2).

ESEMPIO 4. Il sistema dinamico di Lotka-Volterra (Alfred James Lotka, 1880-1849). Nello spazio affine bidimensionale riferito a coordinate cartesiane (x, y) si consideri sul dominio aperto $D = \{(x, y) \in \mathbb{R}^2 \mid x > 0, y > 0\}$ costituito dai punti aventi entrambe le coordinate positive (primo quadrante) il campo vettoriale

$$(19) \quad \mathbf{X} = (a - by)x \mathbf{i} + (cx - d)y \mathbf{j},$$

con (a, b, c, d) numeri reali positivi. Il corrispondente sistema del primo ordine è:

$$(20) \quad \begin{cases} \dot{x} = (a - by)x, \\ \dot{y} = (cx - d)y. \end{cases}$$

Questo sistema costituisce il celebre modello di Lotka-Volterra per la dinamica di due popolazioni, il cui numero di individui è rappresentato dalle funzioni $y(t)$ e $x(t)$, la prima *predatrice* della seconda, viventi in un ambiente isolato ideale (si veda in appendice al paragrafo il processo costruttivo di questo modello).

L'equazione (4) degli integrali primi diventa:

$$(a - by)x \frac{\partial F}{\partial x} + (cx - d)y \frac{\partial F}{\partial y} = 0.$$

Ancora una volta questa si integra per separazione delle variabili, ponendo cioè $F(x, y) = A(x) + B(y)$. Si trova che una soluzione è

$$F(x, y) = cx - d \log x + by - a \log y.$$

L'analisi della funzione $z = F(x, y)$ mostra che la superficie rappresentativa in $\mathbb{R}^3$ è convessa verso il basso con un minimo nel punto P_0 di coordinate

$$x_0 = \frac{d}{c}, \quad y_0 = \frac{a}{b},$$

che è punto singolare del campo $\mathbf{X}$ (Fig. 10.3). Le orbite Q_c , di equazione $F = c$, si ottengono allora intersecando questa superficie con il piano $z = c$. Stante la convessità della superficie, le orbite risultano delle curve chiuse, che circondano il punto singolare P_0 (Fig. 10.3). Di qui si deduce che le curve integrali sono periodiche e che il campo è completo (ciò significa che le due popolazioni hanno un andamento periodico a meno che le condizioni iniziali non corrispondano al punto singolare, nel qual caso il numero degli individui, per entrambe le specie, risulta costante nel tempo). Se si valuta il campo $\mathbf{X}$ in qualche punto di D , per esempio sulle rette $x = x_0$ o $y = y_0$ passanti per il punto singolare, si riconosce che queste vengono percorse in senso antiorario.

La costruzione del sistema dinamico di Lotka-Volterra. Consideriamo due popolazioni viventi in uno stesso ambiente *chiuso* (senza influenze dall'esterno). Denotiamo con $x(t)$ e $y(t)$ il numero degli individui di ciascuna popolazione, funzioni del tempo t . La popolazione x è costituita da *prede* (p.es. topi) della popolazione y , costituita da *predatori* (gatti). Le condizioni di vita sono le seguenti: (i) L'ambiente sarebbe "ideale" per i topi, se non vi fossero i gatti; vale a dire: i topi dispongono

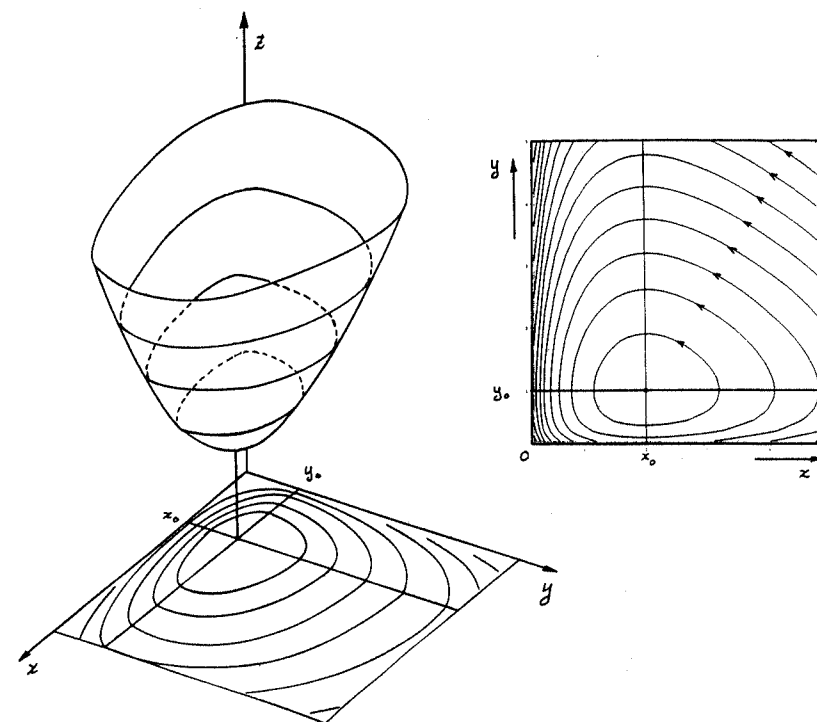


Fig. 10.3: le orbite del sistema di Lotka-Volterra.

di risorse alimentari illimitate e la loro mortalità è dovuta solo a cause naturali non eccezionali e alla presenza dei gatti. (ii) I gatti si nutrono solo di topi e muoiono solo di cause naturali.

Cerchiamo un modello per la dinamica delle due popolazioni. Analizziamo il tipo di crescita dei predatori e quello delle prede separatamente.

(a) Crescita delle prede. In un ambiente ideale dove le risorse nutritive sono illimitate (ciò vale per ogni tipo di popolazione) l'incremento Δx della popolazione x in un intervallo Δt di tempo è proporzionale a Δt (più il tempo passa, più gli individui si accoppiano) e al numero x degli individui (più ci sono individui, più aumenta la possibilità di

incontri e quindi di accoppiamenti). Scriviamo pertanto

$$(21) \quad \Delta x = a x \Delta t,$$

dove a è un coefficiente di proporzionalità, costante e positivo. Siccome per Δt che tende a zero il rapporto $\frac{\Delta x}{\Delta t}$ tende alla velocità di crescita $\dot{x}$, dalla relazione precedente segue l'equazione differenziale della *crescita ideale o malthusiana di una popolazione*

$$(22) \quad \boxed{\dot{x} = a x} \quad \left(\dot{x} = \frac{dx}{dt} \right).$$

Le sue soluzioni sono le funzioni esponenziali (si dice che la crescita ideale di una popolazione è "esponenziale"): si tratta di funzioni a crescita molto rapida (la cui velocità di crescita aumenta con la crescita, come afferma la (22)). La costante a che compare nella (22) si può determinare sperimentalmente ed è una costante caratteristica della "salute" della popolazione.

Consideriamo ora la presenza dei predatori. Le prede divorate in un intervallo di tempo Δt sono un numero proporzionale a Δt con un coefficiente di proporzionalità $f(x, y)$ che dipende dal numero degli individui di entrambe le specie, vale a dire del tipo: $f(x, y)\Delta t$. Dal canto suo il coefficiente di proporzionalità $f(x, y)$ è: (i) proporzionale a y (se 1 gatto mangia 1 topo, 2 gatti mangiano 2 topi, etc.); (ii) proporzionale a x (più sono i topi più cresce la possibilità che un gatto catturi un topo). Quindi possiamo supporre che il numero delle prede catturate nell'intervallo Δt è del tipo $b x y \Delta t$ con b costante positiva. Questa quantità va sottratta alla (21), per ottenere l'incremento complessivo dei topi in presenza dei gatti nell'intervallo di tempo Δt . Così facendo al posto della legge (22) di crescita ideale si trova una legge del tipo

$$(23) \quad \boxed{\dot{x} = x(a - b y)}$$

Questa non consente di calcolare subito, come per la (22), l'andamento della popolazione x , perché per far questo occorre conoscere anche l'andamento $y(t)$ della popolazione y .

(b) Crescita dei predatori. In abbondanza di prede vale quanto detto per la crescita ideale. Si può tuttavia osservare che deve esistere un numero m di quantità di prede necessario per il sostentamento dei predatori. Se $x > m$ la crescita è positiva (esubero di prede), se $x < m$

la crescita è negativa (carestia), se $x = m$ la crescita è nulla (equilibrio). In questo caso l'incremento dei gatti Δy in Δt , che in condizioni ideali sarebbe del tipo (21), $\Delta y = c y \Delta t$, con c costante positiva, si modifica in

$$\Delta y = c(x - m)y \Delta t.$$

Di qui, dividendo per Δt e facendolo tendere a zero, si trova la relazione

$$(24) \quad \boxed{\dot{y} = y(c x - d)}$$

dove $d = cm$ è una costante positiva.

Conclusione: le relazioni (23) e (24) governano completamente la dinamica delle due popolazioni. Le costanti (a, b, c, d) , che caratterizzano lo stato delle due popolazioni ed i loro rapporti, devono determinarsi sperimentalmente.

11. Spazi affini euclidei

Come già si è detto al §1, uno **spazio affine euclideo** è una quaterna $(\mathcal{A}, E, \delta, \mathbf{g})$ dove $(\mathcal{A}, E, \delta)$ è uno spazio affine e $\mathbf{g}$ è un tensore metrico sopra lo spazio vettoriale E (che pertanto è uno spazio vettoriale euclideo). Si dice che uno spazio affine ha segnatura (p, q) se questa è la segnatura della forma bilineare simmetrica $\mathbf{g}$; in particolare lo spazio affine si dice **strettamente euclideo** se la segnatura è $(n, 0)$, cioè se il tensore metrico è definito positivo. Dato un riferimento affine $(O, \mathbf{c}_\alpha)$ dove la base è canonica, le corrispondenti coordinate (x^α) si dicono **coordinate canoniche** o **coordinate ortonormali**.

Negli spazi affini a 2 o 3 dimensioni strettamente euclidei denotiamo con $(\mathbf{i}, \mathbf{j})$ o $(\mathbf{i}, \mathbf{j}, \mathbf{k})$ una generica base canonica e con (x, y) o (x, y, z) le corrispondenti coordinate cartesiane ortonormali. In dimensione 3 si conviene tacitamente che l'orientamento sia quello della mano destra.

Questo paragrafo è dedicato alle operazioni fondamentali che la presenza del tensore metrico consente di definire su di uno spazio affine euclideo.

11.1. Il tensore metrico

Cominciamo con l'esaminare il **prodotto scalare di due campi vettoriali**: ad ogni coppia di campi vettoriali (X, Y) sopra lo spazio affine $\mathcal{A}$ corrisponde il campo scalare

$$(1) \quad \boxed{X \cdot Y = g(X, Y)}$$

definito da

$$(X \cdot Y)(P) = g(X(P), Y(P)) \quad (P \in \mathcal{A}).$$

Ciò significa che in ogni punto P il valore della funzione $X \cdot Y$ è il prodotto scalare, in E , dei due vettori $X(P)$ e $Y(P)$. Si definisce in tal modo una forma bilineare simmetrica sui campi vettoriali a valori nei campi scalari, che continuiamo a denotare con g ed a chiamare **tensore metrico** o semplicemente **metrica**:

$$(2) \quad g: \mathcal{X}(\mathcal{A}) \times \mathcal{X}(\mathcal{A}) \rightarrow \mathcal{F}(\mathcal{A}): (X, Y) \mapsto X \cdot Y.$$

Le sue componenti rispetto a generiche coordinate (q^i) sono definite da

$$(3) \quad g_{ij} = E_i \cdot E_j.$$

Sono funzioni sopra il dominio di definizione delle coordinate e formano una matrice simmetrica, ovunque regolare, detta **matrice metrica**. Se in particolare le coordinate sono cartesiane le componenti della metrica sono costanti. Data la matrice metrica, è possibile calcolare il prodotto scalare di due qualunque campi vettoriali attraverso le loro componenti in quelle coordinate:

$$(4) \quad X \cdot Y = g_{ij} X^i Y^j.$$

Infatti: $X \cdot Y = X^i E_i \cdot Y^j E_j = X^i Y^j E_i \cdot E_j = X^i Y^j g_{ij}$. Ricordato inoltre che i differenziali dq^i delle coordinate associano ad un campo vettoriale le sue componenti, $\langle X, dq^i \rangle = X^i$, dalle (3) e (4) segue che il tensore metrico, inteso come campo tensoriale, può esprimersi come combinazione dei prodotti tensoriali di questi differenziali. Si può cioè scrivere:

$$(5) \quad \boxed{g = g_{ij} dq^i \otimes dq^j}$$

Le coordinate (q^i) sono dette **ortogonali** se la corrispondente matrice metrica (g_{ij}) è diagonale:

$$g_{ij} = 0, \quad \text{per } i \neq j.$$

Ciò significa che in ogni punto del dominio di definizione delle coordinate i corrispondenti vettori (E_i) sono mutuamente ortogonali.

OSSERVAZIONE 1. In molti testi il tensore metrico è sostituito dal simbolo ds^2 e ds è chiamato *elemento d'arco euclideo*. L'espressione (5) è sostituita dalla scrittura

$$(6) \quad \boxed{ds^2 = g_{ij} dq^i dq^j}$$

Questa notazione rende in effetti più automatico il calcolo delle componenti della metrica. Se infatti la metrica è strettamente euclidea (definita positiva) si scrive

$$(7) \quad ds^2 = \sum_{\alpha=1}^n (dx^\alpha)^2,$$

dove le (x^α) sono delle coordinate cartesiane *ortonormali*. Se si considerano altre coordinate (q^i) , una volta date le equazioni di trasformazione $x^\alpha = x^\alpha(q^i)$, è sufficiente calcolare i differenziali

$$dx^\alpha = \frac{\partial x^\alpha}{\partial q^i} dq^i$$

e sostituirli nella sommatoria (7). Sviluppatisi i calcoli, si trova un polinomio di secondo grado omogeneo nelle (dq^i) , cioè un'espressione del tipo (6), i cui coefficienti sono

$$g_{ij} = \sum_{\alpha} \frac{\partial x^\alpha}{\partial q^i} \frac{\partial x^\alpha}{\partial q^j} = \sum_{\alpha} E_i^\alpha E_j^\alpha = E_i \cdot E_j,$$

a conferma della (3). Se la metrica è di segnatura (p, q) , in qualunque sistema di coordinate cartesiane ortonormali per il ds^2 sussiste l'espressione

$$(8) \quad ds^2 = \sum_{\alpha=1}^p (dx^\alpha)^2 - \sum_{\beta=p+1}^n (dx^\beta)^2,$$

dalla quale si parte, procedendo allo stesso modo, per calcolare la matrice metrica in un qualunque altro sistema di coordinate.

ESERCIZIO 1. Si calcolino le componenti del tensore metrico corrispondenti a: (I) coordinate polari del piano euclideo, (II) coordinate cilindriche dello spazio tridimensionale euclideo, (III) coordinate polari (o sferiche) dello spazio tridimensionale euclideo (si veda il §2). Si può procedere tramite la (3), richiamando le espressioni degli E_i viste al §2 per tutte queste coordinate. Si trova che le matrici metriche sono rispettivamente date da

$$\begin{pmatrix} 1 & 0 \\ 0 & r^2 \end{pmatrix}, \quad \begin{pmatrix} 1 & 0 & 0 \\ 0 & r^2 & 0 \\ 0 & 0 & 1 \end{pmatrix}, \quad \begin{pmatrix} 1 & 0 & 0 \\ 0 & r^2 & 0 \\ 0 & 0 & r^2 \sin^2 \varphi \end{pmatrix}.$$

Si osservi che tutte queste coordinate sono ortogonali. In questi tre casi l'espressione (5) diventa, rispettivamente:

$$(9) \quad \begin{cases} g = dr \otimes dr + r^2 d\vartheta \otimes d\vartheta, \\ g = dr \otimes dr + r^2 d\vartheta \otimes d\vartheta + dz \otimes dz, \\ g = dr \otimes dr + r^2 (d\varphi \otimes d\varphi + \sin^2 \varphi d\vartheta \otimes d\vartheta). \end{cases}$$

Si può anche procedere come nell'Oss. 1, partendo dalla scrittura

$$ds^2 = dx^2 + dy^2 + dz^2$$

(che si riduce a $ds^2 = dx^2 + dy^2$ nel caso del piano), sostituendovi le espressioni delle coordinate cartesiane in funzione delle coordinate considerate. La (5) diventa rispettivamente nei tre casi, a conferma delle (9),

$$(10) \quad \begin{cases} ds^2 = dr^2 + r^2 d\vartheta^2, \\ ds^2 = dr^2 + r^2 d\vartheta^2 + dz^2, \\ ds^2 = dr^2 + r^2 (d\varphi^2 + \sin^2 \varphi d\vartheta^2), \end{cases}$$

In analogia con quanto accade negli spazi vettoriali euclidei, la presenza del tensore metrico stabilisce una corrispondenza biunivoca fra

campi vettoriali a 1-forme. Ponendo, per ogni coppia di campi vettoriali,

$$(11) \quad \langle Y, X^b \rangle = Y \cdot X$$

si definisce un'applicazione lineare invertibile

$$b: \mathcal{X}(\mathcal{A}) \rightarrow \Phi^1(\mathcal{A}): X \mapsto X^b.$$

L'applicazione inversa è denotata con $\sharp$. In componenti l'applicazione b si traduce nell'abbassamento degli indici delle componenti dei campi vettoriali per opera della matrice metrica (g_{ij}) ; l'applicazione inversa $\sharp$ nell'innalzamento degli indici delle componenti delle 1-forme per opera della **matrice metrica inversa** (g^{ij}) . Valgono cioè le relazioni:

$$(12) \quad \begin{aligned} \boxed{\xi = X^b} &\iff \boxed{\xi_i = g_{ij} X^j} \\ \boxed{X = \xi^\sharp} &\iff \boxed{X^i = g^{ij} \xi_j} \end{aligned}$$

OSSERVAZIONE 2. I simboli di Christoffel (Γ_{ij}^k) relativi ad un sistema di coordinate (q^i) sono definiti dall'equazione (formula (15), §3)

$$(13) \quad \partial_i E_j = \Gamma_{ij}^k E_k.$$

Questi possono anche calcolarsi con la formula

$$(14) \quad \boxed{\Gamma_{ij}^k = g^{kh} \Gamma_{ijh}}$$

dove le funzioni (Γ_{ijh}) , dette **simboli di Christoffel di prima specie**, sono calcolabili attraverso le derivate delle componenti del tensore metrico secondo la formula

$$(15) \quad \boxed{\Gamma_{ijh} = \frac{1}{2} (\partial_i g_{jh} + \partial_j g_{hi} - \partial_h g_{ij})}$$

Infatti, tenuto conto che le (14) si invertono nelle relazioni

$$(16) \quad \boxed{\Gamma_{ijh} = g_{kh} \Gamma_{ij}^k}$$

moltiplicata la (13) scalarmente per $\mathbf{E}_h$, risulta

$$(17) \quad \Gamma_{ijh} = \partial_i \mathbf{E}_j \cdot \mathbf{E}_h.$$

Ma il primo membro di queste uguaglianze è anche uguale a

$$\partial_i \mathbf{E}_j \cdot \mathbf{E}_h = \partial_i (\mathbf{E}_j \cdot \mathbf{E}_h) - \partial_i \mathbf{E}_h \cdot \mathbf{E}_j = \partial_i g_{jh} - \partial_i \mathbf{E}_h \cdot \mathbf{E}_j,$$

per cui, combinando i due risultati si ottiene l'uguaglianza

$$(18) \quad \Gamma_{ijh} + \Gamma_{ihj} = \partial_i g_{jh}.$$

Con la permutazione degli indici si ottengono altre due uguaglianze simili:

$$\begin{cases} \Gamma_{jhi} + \Gamma_{jih} = \partial_j g_{hi}, \\ \Gamma_{hij} + \Gamma_{hji} = \partial_h g_{ij}. \end{cases}$$

Sommando membro a membro le prime due e sottraendo la terza, tenuto conto della simmetria rispetto ai primi due indici dei simboli di Christoffel di prima specie, si ricava proprio l'uguaglianza (15).

11.2. Il gradiente

Su di uno spazio affine ad ogni campo scalare U viene associata una 1-forma: il suo differenziale dU . Se sullo spazio affine è presente un tensore metrico, allora al differenziale dU si può associare il campo vettoriale $dU^\sharp$. Esso prende il nome di **gradiente** del campo scalare U e viene denotato con $\text{grad}(U)$. Si pone cioè per definizione:

$$(1) \quad \boxed{\text{grad}(U) = dU^\sharp}$$

Tenuto conto della definizione (11) del §11.1, si osserva che il gradiente è caratterizzato dalla seguente proprietà, valida qualunque siano il campo scalare U e il campo vettoriale $\mathbf{Y}$:

$$(2) \quad \boxed{\mathbf{Y} \cdot \text{grad}(U) = \langle \mathbf{Y}, dU \rangle}$$

Dalle (12) di §11.1 segue che le sue componenti in generiche coordinate sono

$$(3) \quad \boxed{(\text{grad}(U))^i = g^{ij} \partial_j U}$$

Se in particolare lo spazio affine è strettamente euclideo e se le coordinate scelte (x^α) sono cartesiane ortonormali, poiché la matrice metrica è la matrice unitaria, le componenti del gradiente risultano essere le derivate parziali della funzione:

$$(4) \quad \boxed{(\text{grad}(U))^\alpha = \frac{\partial U}{\partial x^\alpha}}$$

per cui

$$(4') \quad \text{grad}(U) = \frac{\partial U}{\partial x^\alpha} \mathbf{c}_\alpha = \frac{\partial U}{\partial x^1} \mathbf{c}_1 + \frac{\partial U}{\partial x^2} \mathbf{c}_2 + \dots + \frac{\partial U}{\partial x^n} \mathbf{c}_n.$$

Nello spazio affine tridimensionale euclideo si ha quindi in particolare:

$$(5) \quad \boxed{\text{grad}(U) = \frac{\partial U}{\partial x} \mathbf{i} + \frac{\partial U}{\partial y} \mathbf{j} + \frac{\partial U}{\partial z} \mathbf{k}}$$

Il gradiente gode delle seguenti proprietà, conseguenze immediate di analoghe proprietà del differenziale:

$$(6) \quad \begin{cases} \text{grad}(U + V) = \text{grad}(U) + \text{grad}(V), \\ \text{grad}(UV) = V \text{grad}(U) + U \text{grad}(V), \\ \text{grad}(U) = 0 \iff U = \text{cost.}, \\ \langle \text{grad}(U), dV \rangle = \text{grad}(U) \cdot \text{grad}(V). \end{cases}$$

L'implicazione $\implies$ nella (6)₃ è valida sulle componenti connesse del dominio di definizione del campo scalare F .

Un campo vettoriale $\mathbf{F}$ si dice **conservativo** se è il gradiente di un campo scalare U , che prende il nome di **potenziale**:

$$(7) \quad \mathbf{F} = \text{grad}(U).$$

Ciò significa che la corrispondente 1-forma $\varphi = \mathbf{F}^\flat$ è esatta:

$$(8) \quad \varphi = dU.$$

Condizione necessaria affinché un campo sia conservativo è che la 1-forma corrispondente sia chiusa: $d\varphi = 0$. Questa condizione è anche sufficiente per l'esistenza di potenziali locali (o di un potenziale globale, se il dominio di definizione del campo è omeomorfo a $\mathbb{R}^n$ oppure, per

$n = 2, 3$, se è semplicemente connesso). Essa si traduce in componenti nelle equazioni (si veda la (30) del §4)

$$(9) \quad \partial_i \varphi_j = \partial_j \varphi_i.$$

Poiché $\varphi_i = g_{ij} F^j$, si hanno di conseguenza per le componenti del campo vettoriale le condizioni

$$(10) \quad \partial_i (g_{jh} F^h) = \partial_j (g_{ih} F^h).$$

Se lo spazio è strettamente euclideo e le coordinate sono cartesiane ortonormali queste si semplificano in

$$(11) \quad \partial_i F^j = \partial_j F^i.$$

OSSERVAZIONE 1. Le definizioni ora date hanno un particolare interesse in Fisica, dove i **campi di forza** sono rappresentati da campi vettoriali nello spazio affine euclideo tridimensionale. Ne consideriamo alcuni esempi fondamentali, richiamando anche gli esempi sulle forme differenziali visti al §4. (I) **Campi costanti o uniformi**. In ogni punto dello spazio hanno lo stesso valore. Sono campi conservativi: con una opportuna scelta del riferimento cartesiano possiamo porre $\mathbf{F} = a \mathbf{i}$, ed un potenziale è la funzione $U(x) = ax$. (II) **Campi centrali simmetrici**. Un campo vettoriale si dice **campo centrale** se esiste un punto O , detto **centro del campo**, tale che per ogni punto P , con l'esclusione al più del punto O stesso, il vettore $\mathbf{F}(P)$ è parallelo a OP . Inoltre il campo si dice **simmetrico**, o più precisamente a **simmetria sferica**, se la sua intensità dipende solo dalla distanza dal centro, cioè se esiste una funzione reale ϕ tale che, escluso il centro O , si ha $\mathbf{F}(P) = \phi(r) \mathbf{u}$ con $r = |OP|$ e $\mathbf{u}$ versore di OP . Se si considerano coordinate polari sferiche di centro O , si osserva che in questo caso è $\mathbf{F} = \phi(r) \mathbf{E}_r$ e quindi che la 1-forma associata è $\mathbf{F}^b = \phi(r) dr$. Di qui segue, per quanto visto al §3, che il campo $\mathbf{F}$ è conservativo e che un suo potenziale è una qualunque funzione $U(r)$ primitiva della $\phi(r)$. (III) **Campi coulombiani**. Sono campi centrali (non definiti nel centro) del tipo

$$(12) \quad \mathbf{F} = \frac{k}{r^2} \mathbf{u} = \frac{k}{r^3} \mathbf{r} \quad (k \in \mathbb{R}).$$

Un potenziale è la funzione

$$(13) \quad U(r) = -\frac{k}{r}.$$

Se $k > 0$ il campo è di tipo *repulsivo*, se $k < 0$ è invece *attrattivo*. In questo secondo caso il campo $\mathbf{F}$ è detto **campo newtoniano**. (IV) **Campo elastico e campo centrifugo**. Sono i campi, definiti su tutto lo spazio affine, del tipo

$$(14) \quad \mathbf{F} = k r \mathbf{u} = k \mathbf{r} \quad (k \in \mathbb{R}),$$

con $k < 0$ e $k > 0$ rispettivamente. Un potenziale è la funzione

$$(15) \quad U(r) = \frac{k}{2} r^2.$$

Il gradiente gode di una notevole proprietà geometrica nota come **teorema del gradiente**.

TEOREMA. Il gradiente di una funzione U è ortogonale alle superfici di livello $U = \text{cost.}$ (dette anche **superfici equipotenziali**) ed è orientato nel verso di U crescente.

Dimostrazione. Escludiamo i punti singolari ($dU = 0$) della funzione U , dove il gradiente è nullo, vale a dire i punti singolari delle superfici di livello. Si è visto (§6, Prop. 1) che un vettore $\mathbf{v}$ è tangente ad una superficie $U = \text{cost.}$ se e solo se $\langle \mathbf{v}, dU \rangle = 0$. Questa condizione equivale a $\mathbf{v} \cdot \text{grad}(U) = 0$ (si veda la (2)); pertanto $\text{grad}(U)$ è ortogonale a tutti i vettori tangenti alle superfici equipotenziali. D'altra parte la derivata del campo scalare U rispetto al campo vettoriale $\text{grad}(U)$ è positiva:

$$\langle \text{grad}(U), dU \rangle = \text{grad}(U) \cdot \text{grad}(U) > 0.$$

Ciò dimostra la seconda parte dell'enunciato. ■

OSSERVAZIONE 2. La combinazione della divergenza e del gradiente fornisce il **laplaciano** di un campo scalare U , denotato in genere con ΔU :

$$(16) \quad \boxed{\Delta U = \text{div}(\text{grad}(U))}$$

Combinando la definizione di divergenza (9) del §2 con la (4') si vede che, in uno spazio strettamente euclideo riferito a coordinate cartesiane ortonormali,

$$(17) \quad \Delta U = \sum_{\alpha=1}^n \frac{\partial^2 U}{\partial x^{\alpha 2}}.$$

11.3. La forma volume

In un qualunque spazio affine esistono due orientamenti, corrispondenti ai due orientamenti dello spazio vettoriale sottostante. Se lo spazio affine è euclideo, ad ogni orientamento corrisponde una **forma volume** definita, in analogia con quanto visto per gli spazi vettoriali euclidei tridimensionali (§11.1, Cap. I) da

$$(1) \quad \eta = dx^1 \wedge dx^2 \wedge \dots \wedge dx^n,$$

essendo le (x^α) coordinate cartesiane ortonormali. Nel caso strettamente euclideo e tridimensionale si ha per definizione

$$(2) \quad \boxed{\eta = dx \wedge dy \wedge dz}$$

OSSERVAZIONE 1. Se si considera un qualunque altro sistema di coordinate (q^i) , con calcolo del tutto analogo a quello svolto per gli spazi vettoriali, si trova che

$$(3) \quad \eta = J dq^1 \wedge dq^2 \wedge dq^3 = \pm \sqrt{g} dq^1 \wedge dq^2 \wedge dq^3.$$

dove J è lo **jacobiano** corrispondente alle coordinate (q^i) ,

$$(4) \quad J = \det(E_i^\alpha), \quad E_i^\alpha = \frac{\partial x^\alpha}{\partial q^i},$$

e dove, nella seconda uguaglianza, posto

$$(5) \quad g = \det(g_{ij}),$$

va scelto il segno $+$ se le coordinate (q^i) sono concordi all'orientamento scelto sullo spazio affine. In coordinate generiche (q^i) la forma volume ammette una rappresentazione del tipo

$$(6) \quad \eta = \frac{1}{3!} \eta_{hij} dq^h \wedge dq^i \wedge dq^j.$$

dove, in conformità con le (3), le componenti η_{hij} sono date da

$$(7) \quad \eta_{hij} = J \varepsilon_{hij} = \pm \sqrt{g} \varepsilon_{hij}.$$

OSSERVAZIONE 2. Nel caso tridimensionale la forma volume fa corrispondere ad ogni terna di campi vettoriali (X, Y, Z) il campo scalare definito, qualunque siano le coordinate scelte, da

$$(8) \quad \eta(X, Y, Z) = \eta_{hij} X^h Y^i Z^j.$$

Si possono di conseguenza estendere ai campi vettoriali le operazioni viste per i vettori nello spazio euclideo tridimensionale: in primo luogo il prodotto vettoriale $X \times Y$ di due campi vettoriali ponendo

$$(9) \quad \eta(X, Y, Z) = X \times Y \cdot Z,$$

e quindi l'operazione di aggiunzione $*$, che mette in corrispondenza biunivoca campi vettoriali con 2-forme e campi scalari con 3-forme ponendo

$$(10) \quad *X(Y, Z) = X \cdot Y \times Z, \quad *F = F \eta.$$

Valgono per queste operazioni proprietà analoghe a quelle viste per gli spazi vettoriali.

OSSERVAZIONE 3. Su di uno spazio affine qualsiasi si definisce la divergenza di un campo vettoriale (Oss. 7, §2). Se lo spazio è strettamente euclideo e tridimensionale, la divergenza può anche essere definita alla maniera seguente:

$$(11) \quad \boxed{\operatorname{div}(X) = *d*X}$$

Si osservi intanto che questa formula ha senso perché: $*X$ è una 2-forma, $d*X$ è una 3-forma, $*d*X$ è un campo scalare. La dimostrazione dell'equivalenza della (11) con la definizione (9) del §2 è lasciata come esercizio. Un campo a divergenza nulla

$$\operatorname{div}(X) = 0$$

è detto **campo solenoidale** o anche **campo incompressibile**. Il primo termine trae origine dalla teoria dell'elettromagnetismo, il secondo dalla meccanica dei fluidi: le trasformazioni φ_t associate al campo X conservano il volume delle porzioni di spazio a cui sono applicate.

OSSERVAZIONE 4. L'ultima affermazione dell'osservazione precedente, la cui dimostrazione esula dal programma di questo corso, richiama la definizione di **volume di un dominio** T , sottinsieme dello spazio

affine. Osserviamo soltanto che il volume di un dominio T nello spazio tridimensionale è l'integrale esteso a T della forma volume η , cioè il numero (si vedano la (2) e la (3)₁)

$$(12) \quad \int_T \eta = \int_T dx \wedge dy \wedge dz = \int_D J dq^1 \wedge dq^2 \wedge dq^3,$$

dove $D \subset \mathbb{R}^3$ è l'immagine in $\mathbb{R}^3$ del dominio T secondo la carta corrispondente alle coordinate (q^i) e l'integrale a secondo membro va inteso come integrale triplo in queste variabili. Per cui, eliminando i simboli di prodotto esterno in accordo con le notazioni correnti, risulta

$$(13) \quad \int_T \eta = \iiint_T dx dy dz = \iiint_D J dq^1 dq^2 dq^3.$$

In quanto precede è implicita l'ipotesi che il dominio delle coordinate (q^i) ricopra tutto T .

11.4. Il rotore

In uno spazio strettamente euclideo tridimensionale si definisce il **rotore** di un campo vettoriale; è il campo vettoriale

$$(1) \quad \boxed{\text{rot}(\mathbf{F}) = *d\mathbf{F}^b}$$

Si osservi che questa definizione è ben data: $\mathbf{F}^b$ è una 1-forma, $d\mathbf{F}^b$ è una 2-forma, $*d\mathbf{F}^b$ è un campo vettoriale. In base alle definizioni precedenti, segue che se si prende un riferimento cartesiano ortonormale $(O, \mathbf{c}_\alpha) = (O, \mathbf{i}, \mathbf{j}, \mathbf{k})$ di coordinate $(x^\alpha) = (x, y, z)$, allora

$$(2) \quad \text{rot}(\mathbf{F}) = \varepsilon_{\alpha\beta\gamma} \frac{\partial F^\gamma}{\partial x^\beta} \mathbf{c}_\alpha,$$

ovvero

$$(3) \quad \text{rot}(\mathbf{F}) = \begin{vmatrix} \mathbf{i} & \mathbf{j} & \mathbf{k} \\ \frac{\partial}{\partial x} & \frac{\partial}{\partial y} & \frac{\partial}{\partial z} \\ X & Y & Z \end{vmatrix},$$

essendo $(F^\alpha) = (X, Y, Z)$ le componenti cartesiane del campo. Si noti che, posto $\varphi = \mathbf{F}^b$, la definizione (1) equivale a

$$(4) \quad * \text{rot}(\mathbf{F}) = d\varphi.$$

Dalle proprietà dei differenziali di 1-forme seguono dunque delle proprietà del rotore. Le principali sono le seguenti (la loro dimostrazione è lasciata come esercizio):

$$(5) \quad \begin{cases} \text{rot}(\mathbf{X} + \mathbf{Y}) = \text{rot}(\mathbf{X}) + \text{rot}(\mathbf{Y}), \\ \text{rot}(a \mathbf{X}) = a \text{rot}(\mathbf{X}) \quad (a \in \mathbb{R}), \\ \text{rot}(U \mathbf{X}) = U \text{rot}(\mathbf{X}) + \text{grad}(U) \times \mathbf{X}, \\ \text{rot}(\text{grad}(U)) = 0, \\ \text{div}(\text{rot}(\mathbf{X})) = 0, \\ \text{div}(\mathbf{X} \times \mathbf{Y}) = \text{rot}(\mathbf{X}) \cdot \mathbf{Y} - \text{rot}(\mathbf{Y}) \cdot \mathbf{X}. \end{cases}$$

Un campo vettoriale $\mathbf{F}$ si dice **irrotazionale** se

$$\text{rot}(\mathbf{X}) = 0.$$

La (5)₄ mostra che un campo conservativo è irrotazionale (una forma esatta è chiusa). Viceversa, per il lemma di Poincaré, un campo è irrotazionale ammette potenziali locali; se il suo dominio è semplicemente connesso allora ammette un potenziale globale.

OSSERVAZIONE 1. In molti testi il gradiente di un campo scalare, la divergenza e il rotore di un campo vettoriale sono rispettivamente denotati con:

$$(6) \quad \begin{cases} \text{grad}(U) = \nabla U, \\ \text{div}(\mathbf{X}) = \nabla \cdot \mathbf{X} \quad (\text{se } \cdot \text{ è il prodotto scalare}), \\ \text{rot}(\mathbf{X}) = \nabla \times \mathbf{X} \quad (\text{se } \times \text{ è il prodotto vettoriale}). \end{cases}$$

Va sottolineato il fatto che, mentre gli operatori di gradiente e divergenza sono definibili in spazi di qualunque dimensione n , il rotore è definito solo per $n = 3$.

12. Curve e superfici nello spazio euclideo tridimensionale

Sia $\gamma: I \rightarrow \mathcal{A}$ una curva nello spazio affine (strettamente) euclideo tridimensionale di equazioni parametriche, rispetto a generiche coordinate (q^i) ,

$$(1) \quad q^i = \gamma^i(t) \quad (i = 1, 2, 3; t \in I \subset \mathbb{R}).$$

Si supponga, per semplicità, che essa sia di classe C^∞ , e che l'intervallo I di definizione contenga lo zero. Sia $v(t) = \dot{\gamma}(t)$ il vettore tangente (ovvero la velocità istantanea, interpretata la curva come moto), le cui componenti, ricordiamolo, sono date dalle derivate delle funzioni parametriche (1):

$$v^i = \frac{dq^i}{dt}.$$

Per ogni $t \in I$ risulta definito l'integrale

$$(2) \quad s(t) = \int_0^t |v(x)| dx = \int_0^t \sqrt{g_{ij} \frac{dq^i}{dx} \frac{dq^j}{dx}} dx,$$

ed è di conseguenza definita una funzione $s: I \rightarrow \mathbb{R}: t \mapsto s(t)$, monotona crescente (l'integrando è infatti sempre non negativo), tale che

$$(3) \quad \frac{ds}{dt} = |v|.$$

Questa funzione consente innanzitutto di calcolare la **lunghezza della curva** nell'intervallo $[t_1, t_2] \subset I$: è il numero positivo o nullo

$$s(t_2) - s(t_1).$$

Se inoltre è sempre $v(t) \neq 0$ (cioè se, dal punto di vista cinematico, non vi sono *istanti di arresto*), allora anche la funzione monotona inversa $t = t(s)$ è differenziabile e può sostituire t nella parametrizzazione della curva. Il parametro s prende allora il nome di **parametro euclideo** o di **ascissa euclidea**.

Consideriamo una rappresentazione vettoriale della curva:

$$OP = r(t).$$

Il vettore tangente alla curva nella parametrizzazione in t è il vettore

$$v = \frac{dr}{dt}.$$

Il vettore tangente alla curva nella parametrizzazione in s è invece il vettore

$$(4) \quad t = \frac{dr}{ds} = \frac{dt}{ds} v$$

È un vettore unitario detto **versore tangente** alla curva. Il fatto che sia unitario segue dalla (3):

$$|t| = \frac{dt}{ds} |v| = 1.$$

Consideriamo quindi la derivata del versore tangente t rispetto al parametro euclideo s e poniamo

$$(5) \quad \frac{dt}{ds} = c n$$

assumendo $c \geq 0$ e $|n| = 1$. Definiamo in tal modo due grandezze: uno scalare non negativo $c = c(s)$ detto **curvatura** della curva e un versore $n = n(s)$ detto **versore normale principale** della curva. Il versore n è in ogni punto ortogonale alla curva, cioè al versore tangente t , per effetto della seguente generale proprietà, valida qualunque sia la segnatura dello spazio:

PROPOSIZIONE 1. Sia $u(x)$ una funzione vettoriale derivabile nella variabile reale x . Se la norma di u è costante la sua derivata è un vettore ortogonale:

$$\|u(x)\| = \text{cost.} \implies u \cdot \frac{du}{dx} = 0.$$

Per dimostrarlo è sufficiente ricorrere alla seguente proprietà, di dimostrazione immediata:

PROPOSIZIONE 2. Per la derivata di un prodotto scalare vale la regola di Leibniz:

$$\frac{d}{dx} (u \cdot v) = \frac{du}{dx} \cdot v + u \cdot \frac{dv}{dx}.$$

Infatti da $\mathbf{u} \cdot \mathbf{u} = \text{cost.}$ segue:

$$0 = \frac{d}{dx}(\mathbf{u} \cdot \mathbf{u}) = 2 \frac{d\mathbf{u}}{dx} \cdot \mathbf{u}. \quad \blacksquare$$

Pure di immediata dimostrazione è la proprietà seguente, valida nel caso tridimensionale:

PROPOSIZIONE 3. Per la derivata di un prodotto vettoriale vale la regola di Leibniz,

$$\frac{d}{dx}(\mathbf{u} \times \mathbf{v}) = \frac{d\mathbf{u}}{dx} \times \mathbf{v} + \mathbf{u} \times \frac{d\mathbf{v}}{dx},$$

Dalla (5) si nota che lo scalare c misura la *rapidità* con cui cambia la tangente alla curva al variare del parametro euclideo s . Se $c \neq 0$ il suo inverso

$$(6) \quad \varrho = \frac{1}{c}$$

prende il nome di **raggio di curvatura** della curva. Se $\mathbf{t}$ è costante, la curvatura è nulla e il versore normale principale è indeterminato. È questo il caso in cui la traccia della curva γ è una retta. Infatti, se $\mathbf{t} = \text{cost.}$ dalla (4) segue $\mathbf{r}(s) = s\mathbf{t} + \mathbf{r}_0$, con $\mathbf{r}_0 = \mathbf{r}(0)$.

In ogni punto P della curva il piano individuato dalla terna $(P, \mathbf{t}, \mathbf{n})$ prende il nome di **piano osculatore** alla curva nel punto P , e il punto C tale che $PC = \varrho \mathbf{n}$ è detto **centro di curvatura** della curva nel punto P .

Il vettore unitario

$$(7) \quad \mathbf{b} = \mathbf{t} \times \mathbf{n}$$

prende il nome di **versore binormale**. In ogni punto della curva risulta così definita una terna di versori indipendenti $(\mathbf{t}, \mathbf{n}, \mathbf{b})$ che prende il nome di **terna fondamentale** o **triedro fondamentale**. Siccome per la derivazione di un prodotto vettoriale vale la regola di Leibniz, vista la (5), si ha

$$\frac{d\mathbf{b}}{ds} = \frac{d\mathbf{t}}{ds} \times \mathbf{n} + \mathbf{t} \times \frac{d\mathbf{n}}{ds} = \mathbf{t} \times \frac{d\mathbf{n}}{ds}.$$

Quindi la derivata di $\mathbf{b}$ è ortogonale a $\mathbf{t}$. Ma essa è anche ortogonale a $\mathbf{b}$, per la Prop. 1. È dunque necessariamente parallela a $\mathbf{n}$. Si pone allora:

$$(8) \quad \boxed{\frac{d\mathbf{b}}{ds} = \tau \mathbf{n}}$$

Lo scalare $\tau = \tau(s)$ così definito prende il nome di **torsione** della curva. La torsione dà la misura di quanto la curva si scosti da una curva piana (una curva piana è una curva la cui immagine appartiene ad un piano immerso nello spazio affine). Si veda a questo proposito l'Esercizio 3, più avanti.

Le formule precedenti (5) e (8) esprimono la derivata rispetto al parametro euclideo s dei versori $\mathbf{t}$ e $\mathbf{b}$. La derivata di $\mathbf{n}$ è invece data dalla formula seguente, la quale, si noti, non fa comparire nuove entità scalari:

$$(9) \quad \boxed{\frac{d\mathbf{n}}{ds} = -c \mathbf{t} - \tau \mathbf{b}}$$

Infatti da $\mathbf{n} = \mathbf{b} \times \mathbf{t}$ segue:

$$\frac{d\mathbf{n}}{ds} = \frac{d\mathbf{b}}{ds} \times \mathbf{t} + \mathbf{b} \times \frac{d\mathbf{t}}{ds} = \tau \mathbf{n} \times \mathbf{t} + c \mathbf{b} \times \mathbf{n}.$$

Le formule (5), (8) e (9) sono nel loro complesso chiamate **formule di Frenet** (Jean Frédéric Frenet, 1816-1900).

OSSERVAZIONE 1. Le grandezze scalari e vettoriali che abbiamo qui sopra introdotto sono intrinsecamente legate alla curva γ , non dipendono cioè dalla sua rappresentazione parametrica (salvo l'orientamento, cioè il verso con cui si fa crescere l'ascissa euclidea) e riflettono quindi delle sue proprietà puramente geometriche.

ESERCIZIO 1. Dimostrare che una circonferenza di raggio R ha curvatura costante pari a $\frac{1}{R}$.

ESERCIZIO 2. Calcolare la curvatura e la torsione di un'elica di raggio R e passo p . Ricordiamo che l'elica è una curva ottenuta dalla composizione di un moto circolare uniforme e da un moto rettilineo uniforme in direzione ortogonale al piano del moto circolare; il passo è

la distanza tra due punti che si corrispondono dopo un giro completo nel moto circolare.

ESERCIZIO 3. Dimostrare che una curva è piana se e solo se la sua torsione è ovunque nulla. Suggerimento: la condizione $\tau = 0$ equivale (si veda la (8)) a $\mathbf{b} = \text{cost.}$, e questa implica $\mathbf{b} \cdot \mathbf{r} = \text{cost.}$; di qui far seguire che il moto è piano. Viceversa, se la curva sta su di un piano, sia il versore tangente che il versore normale principale sono tangenti al piano, e quindi $\mathbf{b}$ è ortogonale a questo, dunque è costante; per cui $\tau = 0$.

Consideriamo ora, nello spazio affine euclideo tridimensionale, una superficie regolare Q di rappresentazione parametrica vettoriale

$$(10) \quad OP = \mathbf{r}(q^1, q^2).$$

Denotiamo al solito con

$$\mathbf{E}_i = \frac{\partial \mathbf{r}}{\partial q^i} \quad (i = 1, 2)$$

i vettori tangenti alla superficie corrispondenti a questa parametrizzazione (d'ora in poi gli indici latini assumono i valori (1,2)). Ricordiamo che essi sono indipendenti in ogni punto della superficie e generano pertanto i piani tangenti.

Sia $T_P Q$ lo spazio vettoriale dei vettori tangenti alla superficie Q in un suo punto P e TQ l'insieme di tutti i vettori tangenti alla superficie Q . Denotiamo con $TQ \times_Q TQ$ l'insieme delle coppie di vettori tangenti nel medesimo punto di Q . La restrizione del tensore metrico $\mathbf{g}$ alle coppie di vettori di quest'insieme, definisce una forma bilineare simmetrica $\mathbf{A}$ sui vettori tangenti alla superficie:

$$(11) \quad \mathbf{A}: TQ \times_Q TQ \rightarrow \mathbb{R}: (\mathbf{u}, \mathbf{v}) \mapsto \mathbf{u} \cdot \mathbf{v}.$$

Essa è detta **prima forma fondamentale** della superficie. Poiché per questi vettori si ha $\mathbf{u} \cdot \mathbf{v} = u^i v^j \mathbf{E}_i \cdot \mathbf{E}_j$, quindi

$$(11) \quad \mathbf{u} \cdot \mathbf{v} = A_{ij} u^i v^j,$$

posto

$$(12) \quad A_{ij} = \mathbf{E}_i \cdot \mathbf{E}_j,$$

si può scrivere

$$(13) \quad \boxed{\mathbf{A} = A_{ij} dq^i \otimes dq^j}$$

Le componenti (A_{ij}) della prima forma fondamentale formano una matrice simmetrica 2×2 detta **matrice metrica superficiale**. Nei testi classici esse sono denotate con (E, F, G) . Si pone cioè:

$$A_{11} = E, \quad A_{12} = A_{21} = F, \quad A_{22} = G.$$

In ogni punto di una superficie regolare Q è possibile definire due versori ortogonali a Q , uno opposto all'altro. Se è possibile estendere questi versori a tutta la superficie in modo da ottenere due campi vettoriali (differenziabili) distinti, si dice che la superficie è **orientabile**. Se la superficie Q è definita da un'equazione $F(x, y, z) = 0$ con F campo scalare differenziabile senza punti critici su Q (cioè $dF \neq 0$ nei punti di Q), allora un campo di versori normali è per il teorema del gradiente dato da

$$(14) \quad \mathbf{N} = \frac{\text{grad}(F)}{|\text{grad}(F)|}.$$

Se invece la superficie è data attraverso una rappresentazione parametrica del tipo (10) possiamo considerare il campo dei versori ortogonali definito da

$$(15) \quad \mathbf{N} = \frac{\mathbf{E}_1 \times \mathbf{E}_2}{|\mathbf{E}_1 \times \mathbf{E}_2|}.$$

I campi vettoriali tangenti $(\mathbf{E}_i)$ sono funzioni delle coordinate superficiali (q^i) . Consideriamone le derivate parziali e decomponiamole secondo il piano tangente e la direzione ortogonale alla superficie:

$$(16) \quad \boxed{\partial_i \mathbf{E}_j = \Gamma_{ij}^h \mathbf{E}_h + B_{ij} \mathbf{N}}$$

Risultano definiti due sistemi di funzioni, (Γ_{ij}^h) e (B_{ij}) , entrambi simmetrici negli indici in basso,

$$(17) \quad \Gamma_{ij}^h = \Gamma_{ji}^h, \quad B_{ij} = B_{ji},$$

perché $\partial_i \mathbf{E}_j = \partial_i \partial_j \mathbf{r} = \partial_j \partial_i \mathbf{r} = \partial_j \mathbf{E}_i$. Le funzioni (Γ_{ij}^h) prendono il nome di **simboli di Christoffel di seconda specie** associati alle coordinate (q^i) . Le funzioni (B_{ij}) possono interpretarsi come componenti di una forma bilineare simmetrica sui vettori tangenti alla superficie

$$(18) \quad B = B_{ij} dq^i \otimes dq^j$$

detta **seconda forma fondamentale** o **tensore di curvatura** della superficie. Nei testi classici esse sono denotate con (L, M, N) . Si pone cioè

$$B_{11} = L, \quad B_{12} = B_{21} = M, \quad B_{22} = N.$$

OSSERVAZIONE 2. I simboli di Christoffel possono calcolarsi direttamente dai coefficienti (A_{ij}) della prima forma fondamentale attraverso la formula

$$(19) \quad \Gamma_{ij}^h = A^{hk} \Gamma_{ijk}$$

dove (A^{hk}) è la matrice inversa della matrice metrica (A_{ij}) e dove le funzioni (Γ_{ijk}) , dette **simboli di Christoffel di prima specie**, sono definite da

$$(20) \quad \Gamma_{ijk} = \frac{1}{2} (\partial_i A_{jk} + \partial_j A_{ki} - \partial_k A_{ij})$$

Per riconoscerlo basta moltiplicare (16) scalarmente per $\mathbf{E}_k$, tenuto conto che $\mathbf{E}_k \cdot \mathbf{N} = 0$, e procedere formalmente come nell'Oss. 2 del §11.1 (ora il ruolo di tensore metrico è svolto dalle (A_{ij})).

OSSERVAZIONE 3. Se si moltiplica la (16) scalarmente per il versore $\mathbf{N}$ si ottiene una definizione esplicita delle componenti della seconda forma fondamentale:

$$(21) \quad B_{ij} = \partial_i \mathbf{E}_j \cdot \mathbf{N} = -\partial_i \mathbf{N} \cdot \mathbf{E}_j$$

Come si vede da quest'ultima espressione, le (B_{ij}) forniscono una misura del variare del versore normale alla superficie e quindi la *curvatura* della superficie stessa. Per rendere più precisa quest'idea si consideri una curva tracciata sulla superficie, con parametro euclideo s . Il versore

$\mathbf{t}$ tangente alla curva è ovviamente tangente alla superficie, quindi esprimibile come combinazione lineare dei vettori $(\mathbf{E}_i)$: $\mathbf{t} = t^i \mathbf{E}_i$. Lungo la curva tutte queste grandezze possono pensarsi funzioni di s attraverso i parametri superficiali (q^i) , per cui si ha successivamente:

$$\frac{d\mathbf{t}}{ds} = \frac{dt^i}{ds} \mathbf{E}_i + t^i \frac{d\mathbf{E}_i}{ds} = \frac{dt^i}{ds} \mathbf{E}_i + t^i \frac{dq^j}{ds} \partial_j \mathbf{E}_i.$$

Sicché, moltiplicando scalarmente per $\mathbf{N}$, si trova:

$$\frac{d\mathbf{t}}{ds} \cdot \mathbf{N} = t^i t^j \partial_i \mathbf{E}_j \cdot \mathbf{N},$$

vale a dire, tenuto conto della (5) e della (21),

$$(22) \quad B(\mathbf{t}, \mathbf{t}) = c \mathbf{n} \cdot \mathbf{N}$$

Questa formula fornisce una definizione intrinseca della seconda forma fondamentale.

OSSERVAZIONE 4. Poiché la prima forma fondamentale $\mathbf{A}$ svolge il ruolo di tensore metrico sulla superficie ed è definita positiva, si è condotti a considerare gli autovalori (λ_1, λ_2) della seconda forma $\mathbf{B}$, che sono ovunque reali, per quanto visto al Capitolo I, §9. Essi sono le radici dell'equazione caratteristica

$$\det(B_{ij} - x A_{ij}) = 0$$

e prendono il nome di **curvature principali**. Le corrispondenti auto-direzioni (cioè i sottospazi tangenti di dimensione 1 generati dagli autovettori) sono dette **direzioni principali di curvatura**. Dove $\lambda_1 \neq \lambda_2$ esse sono univocamente determinate e ortogonali fra loro. Nei punti in cui $\lambda_1 = \lambda_2$ esse sono indeterminate. Questi punti si dicono **punti ombelicali** della superficie. Una curva sulla superficie, tangente in ogni punto ad una direzione principale di curvatura, prende il nome di **linea di curvatura**. I punti di una superficie in cui entrambe le curvature principali sono positive o negative si dicono **punti ellittici**, quelli in cui hanno segno diverso **punti iperbolici**, quelli in cui una soltanto è nulla **punti parabolici**. Le quantità

$$(23) \quad H = \frac{1}{2}(\lambda_1 + \lambda_2), \quad K = \lambda_1 \lambda_2,$$

sono dette rispettivamente **curvatura media** e **curvatura totale** (o **curvatura gaussiana**). Si noti che esse sono degli invarianti associati alla seconda forma fondamentale:

$$(24) \quad H = \frac{1}{2} I_1(\mathbf{B}), \quad K = I_2(\mathbf{B}).$$

Lo studio della prima e della seconda forma fondamentale di una superficie è argomento dei corsi di Geometria Differenziale. Per sottolinearne l'importanza, osserviamo quanto segue. Immaginiamo una superficie nello spazio affine tridimensionale euclideo materializzata in una membrana sottile flessibile ma inestendibile (realizzata per esempio con carta). Le proprietà metriche delle figure che possono tracciarsi su tale foglio e che prescindono da relazioni con lo spazio ambiente esterno fanno parte della cosiddetta **geometria interna** della superficie. Tali proprietà sono tutte determinate dalla prima forma quadratica fondamentale $\mathbf{A}$ dunque, assegnato un qualunque sistema di coordinate sulla superficie, dalle sue componenti (A_{ij}) e da tutte le sue derivate parziali rispetto a queste coordinate (quindi dai simboli di Christoffel, etc.). D'altra parte, la **geometria esterna** della superficie, che tiene conto del modo con cui la membrana è immersa nello spazio ambiente, dipende anche dalla seconda forma fondamentale $\mathbf{B}$. È tuttavia sorprendente il fatto che la curvatura gaussiana, pur essendo definita come invariante della seconda forma fondamentale, dipende solo dalla geometria interna della superficie cioè dalla prima forma fondamentale. Questa proprietà, scoperta da Gauss (**theorema egregium**, Karl Friedrich Gauss, 1777-1855), mostra che due superfici applicabili una sull'altra (con sole flessioni e senza "stiramenti") hanno nei punti corrispondenti la stessa curvatura totale. In particolare, posto che, come subito si verifica, un piano immerso nello spazio affine ha la seconda forma fondamentale quindi anche la curvatura gaussiana identicamente nulla, si trova che: una superficie è localmente sviluppabile in un piano se e solo se $K = 0$.

ESERCIZIO 4. Mostrare, calcolandola, che un cilindro e un cono hanno curvatura gaussiana nulla.

ESERCIZIO 5. Mostrare che per una sfera di raggio R si ha $\mathbf{B} = \frac{1}{R} \mathbf{A}$ (quindi le due curvature principali coincidono entrambe con $\frac{1}{R}$: tutti i punti sono ombelicali).

OSSERVAZIONE 5. La forma aggiunta $\sigma = * \mathbf{N}$ del versore normale fornisce l'**elemento d'area** della superficie. Infatti per ogni coppia

$(\mathbf{u}, \mathbf{v})$ di vettori tangenti alla superficie il numero

$$(25) \quad \sigma(\mathbf{u}, \mathbf{v}) = \eta(\mathbf{N}, \mathbf{u}, \mathbf{v}) = \mathbf{N} \cdot \mathbf{u} \times \mathbf{v}$$

fornisce il volume del parallelepipedo individuato dai vettori $(\mathbf{N}, \mathbf{u}, \mathbf{v})$ e quindi, essendo $\mathbf{N}$ unitario e ortogonale ai vettori $(\mathbf{u}, \mathbf{v})$, l'area del parallelogramma individuato da questi. Le componenti dell'elemento d'area σ rispetto a coordinate generiche (q^i) si possono calcolare ponendo nella (25) $(\mathbf{u}, \mathbf{v}) = (\mathbf{E}_i, \mathbf{E}_j)$ e tenendo conto di quanto si è visto al §11.1 del Cap. I. Si trova che

$$(26) \quad \sigma_{ij} = \eta(\mathbf{N}, \mathbf{E}_i, \mathbf{E}_j) = \pm \sqrt{A} \varepsilon_{ij}$$

dove

$$(27) \quad A = \det(A_{ij})$$

e ε_{ij} è il simbolo di Levi-Civita a due indici:

$$\varepsilon_{12} = 1, \quad \varepsilon_{21} = -1, \quad \varepsilon_{11} = 0, \quad \varepsilon_{22} = 0,$$

con il segno $+$ se la scelta dell'ordine delle coordinate e del versore ortogonale alla superficie $\mathbf{N}$ sono tali da rendere la base $(\mathbf{N}, \mathbf{E}_1, \mathbf{E}_2)$ concorde all'orientamento scelto nello spazio. In tal caso risulta:

$$(28) \quad \sigma = \frac{1}{2} \sigma_{ij} dq^i \wedge dq^j = \sqrt{A} dq^1 \wedge dq^2.$$

Si può di conseguenza definire come **area di un dominio** $T \subset Q$ il numero

$$(29) \quad \int_T \sigma = \iint_D \sqrt{A} dq^1 dq^2,$$

dove $D \subset \mathbb{R}^2$ è il dominio immagine di T secondo le coordinate (q^i) .

Per dimostrare direttamente la (26) si osserva innanzitutto che l'area del parallelogramma $(\mathbf{E}_1, \mathbf{E}_2)$ è, in valore assoluto, il modulo del vettore $\mathbf{E}_1 \times \mathbf{E}_2$. D'altra parte è

$$|\mathbf{E}_1 \times \mathbf{E}_2| = \sqrt{A}.$$

Infatti:

$$\begin{aligned}
 \|\mathbf{E}_1 \times \mathbf{E}_2\| &= (\mathbf{E}_1 \times \mathbf{E}_2) \cdot (\mathbf{E}_1 \times \mathbf{E}_2) \\
 &= (\mathbf{E}_1 \times \mathbf{E}_2) \times \mathbf{E}_1 \cdot \mathbf{E}_2 \\
 &= (\mathbf{E}_1 \cdot \mathbf{E}_1 \mathbf{E}_2 - \mathbf{E}_1 \cdot \mathbf{E}_2 \mathbf{E}_1) \cdot \mathbf{E}_2 \\
 &= A_{11} A_{22} - (A_{12})^2 \\
 &= \det(A_{ij}).
 \end{aligned}$$

Pertanto, assunta la (15) come definizione del versore $\mathbf{N}$, risulta

$$\sigma(\mathbf{E}_1, \mathbf{E}_2) = \mathbf{N} \cdot \mathbf{E}_1 \times \mathbf{E}_2 = \sqrt{A}.$$

13. Baricentro di un sistema di masse

I concetti trattati in questo paragrafo sono puramente affini, sussistono cioè in uno spazio affine qualsiasi, senza struttura euclidea. La teoria del baricentro trova comunque notevole applicazione nel caso dello spazio affine euclideo tridimensionale, a cui quindi ci riferiamo.

DEFINIZIONE 1. Un **sistema di masse** o **di punti materiali** è un insieme finito

$$\mathcal{M} = \{(P_\nu, m_\nu); \nu = 1, \dots, N\}$$

costituito da N coppie (P_ν, m_ν) dove P_ν è un punto dello spazio affine $\mathcal{A}$ ed m_ν è un numero *positivo* detto **massa** del punto P_ν . Ogni coppia (P_ν, m_ν) si dice **punto materiale**.

DEFINIZIONE 2. Dicesi **baricentro** (o anche **centro di massa**) di un sistema di masse $\mathcal{M}$ il punto $G \in \mathcal{A}$ definito dall'equazione

$$(1) \quad \boxed{\sum_{\nu} m_{\nu} G P_{\nu} = 0} \quad \left(\sum_{\nu} = \sum_{\nu=1}^N \right),$$

ovvero da

$$(2) \quad \boxed{OG = \frac{1}{m} \sum_{\nu} m_{\nu} O P_{\nu}, \quad m = \sum_{\nu} m_{\nu}}$$

dove O è un qualsiasi punto di riferimento ed m è la **massa totale** del sistema.

Osserviamo innanzitutto che la definizione (2) non dipende dalla scelta del punto O . Infatti se P è un altro punto:

$$\begin{aligned}
 PG &= PO + OG = PO + \frac{1}{m} \sum_{\nu} m_{\nu} O P_{\nu} \\
 &= \frac{1}{m} \sum_{\nu} m_{\nu} (PO + O P_{\nu}) = \frac{1}{m} \sum_{\nu} m_{\nu} P P_{\nu}.
 \end{aligned}$$

La (2) definisce G in maniera univoca. Ponendo $O = G$ nella (2) si trova la (1). Viceversa, con la decomposizione $GP_{\nu} = GO + OP_{\nu} = OP_{\nu} - OG$ dalla (1) si ricava la (2). Le definizioni (1) e (2) sono dunque equivalenti.

Sia (x^{α}) un qualunque sistema di coordinate cartesiane. Se denotiamo con (x_{ν}^{α}) le coordinate del punto P_{ν} e con (x_G^{α}) le coordinate del baricentro, dalla (2) segue che

$$(3) \quad x_G^{\alpha} = \frac{1}{m} \sum_{\nu} m_{\nu} x_{\nu}^{\alpha}.$$

Il baricentro di un sistema di masse gode delle seguenti proprietà.

PROPOSIZIONE 1. Proprietà di appartenenza: se tutti i punti materiali stanno in un semispazio delimitato da un piano allora il baricentro giace in quel semispazio.

Dimostrazione. Basta considerare un riferimento cartesiano per il quale il piano ha per esempio equazione $x^1 = 0$ ed i punti sono tutti situati nel semispazio $x^1 \geq 0$. Dalla (3) segue necessariamente $x_G^1 \geq 0$ e il baricentro sta nello stesso semispazio. ■

Poiché un piano può essere inteso come intersezione di due semispazi si ha subito che:

COROLLARIO 1. Se i punti materiali appartengono tutti ad un piano (o ad una retta) anche il baricentro appartiene a quel piano (a quella retta).

Un altro immediato corollario è il seguente:

COROLLARIO 2. Se i punti sono contenuti in un dominio chiuso convesso allora anche il baricentro è contenuto in quel dominio.

Infatti un tale dominio è l'intersezione di semispazi.

PROPOSIZIONE 2. Proprietà di partizione o distributiva: se il sistema di masse $\mathcal{M}$ è l'unione di due sistemi di masse disgiunti $\mathcal{M}'$ e $\mathcal{M}''$ allora il baricentro G di $\mathcal{M}$ coincide col baricentro del sistema $\{(G', m'), (G'', m'')\}$ costituito dai baricentri di $\mathcal{M}'$ e $\mathcal{M}''$, dotati delle rispettive masse totali.

Dimostrazione. Si consideri il sistema suddiviso in due sottosistemi disgiunti $\mathcal{M}' = \{(P_{\nu'}, m_{\nu'}); \nu' = 1, \dots, N'\}$ e $\mathcal{M}'' = \{(P_{\nu''}, m_{\nu''}); \nu'' = N' + 1, \dots, N\}$. Dalla definizione (2) segue

$$m' OG' = \sum_{\nu'} m_{\nu'} OP_{\nu'}, \quad m'' OG'' = \sum_{\nu''} m_{\nu''} OP_{\nu''}.$$

Dunque:

$$m' OG' + m'' OG'' = \sum_{\nu} m_{\nu} OP_{\nu} = m OG. \quad \blacksquare$$

Si osservi che la proprietà distributiva vale più in generale per una qualunque partizione del sistema, in più di due sottosistemi disgiunti.

DEFINIZIONE 3. Un piano $\Pi \subset \mathcal{A}$ si dice **piano di simmetria materiale coniugato alla direzione u** (u è un vettore non nullo, eventualmente unitario) se per ogni coppia (P_{ν}, m_{ν}) del sistema di masse esiste una coppia (P_{ρ}, m_{ρ}) tale che: (i) $m_{\nu} = m_{\rho}$, (ii) $P_{\nu}P_{\rho}$ è parallelo a u , (iii) il punto medio di $P_{\nu}P_{\rho}$ appartiene a Π . Non si esclude che sia $P_{\nu} = P_{\rho} \in \Pi$.

PROPOSIZIONE 3. Proprietà di simmetria: se esiste un piano di simmetria materiale allora il baricentro appartiene a questo piano.

Dimostrazione. Basta applicare la proprietà distributiva e poi il Corollario 1 al sistema, considerato come unione del sistema dei punti che stanno sul piano Π di simmetria e delle coppie dei punti materiali simmetrici, il cui baricentro giace su Π . ■

La nozione di baricentro, insieme a tutte le proprietà viste per il caso di un sistema finito di punti materiali, si estende al caso di un **sistema materiale continuo**. Un tale sistema è definito da una coppia (M, μ) dove M è un **dominio** di dimensione k e μ una k -forma su M , detta **densità di massa**, tale che abbia senso considerare l'integrale

$$(4) \quad m = \int_M \mu,$$

detto **massa totale** del sistema, nonché l'integrale

$$(5) \quad OG = \frac{1}{m} \int_M OP \mu,$$

che è la scrittura sintetica delle equazioni

$$(6) \quad x_G^{\alpha} = \frac{1}{m} \int_M x^{\alpha} \mu.$$

In particolare, nel caso dello spazio tridimensionale euclideo, se $k = 3$ si può porre

$$\mu = \mu \eta,$$

dove η è la forma volume (pensata ristretta al dominio tridimensionale M); siamo nel caso di un **solido materiale**. Se $k = 2$ ed M è una superficie regolare orientabile, si può porre

$$\mu = \mu \sigma,$$

dove σ è l'elemento d'area; abbiamo allora una **superficie materiale**. Infine, nel caso di una **curva materiale**, $k = 1$, si può porre

$$\mu = \mu ds,$$

dove s è un'ascissa euclidea. In tutti questi casi viene evidenziata una funzione reale $\mu: M \rightarrow \mathbb{R}$, che chiamano **densità scalare di massa**, e che, nella maggior parte dei casi, si richiede essere non negativa. Se questa è costante si dice che il **sistema materiale** è **omogeneo**. Tutti questi concetti possono essere rigorosamente definiti nell'ambito della Teoria della misura e della Teoria dell'integrazione delle forme. Per l'uso che ne faremo nel nostro corso, basta osservare che nel caso di un dominio tridimensionale, l'integrale (4) diventa l'integrale triplo

$$m = \iiint_D J \mu dq^1 dq^2 dq^3,$$

dove $D \subset \mathbb{R}^3$ è il dominio di integrazione relativo alle coordinate (q^i) e J è lo jacobiano corrispondente a queste; il prodotto $J\mu$ è rappresentato da una funzione nelle stesse coordinate. Se le coordinate scelte sono cartesiane ortonormali, allora si ha più semplicemente

$$m = \iiint_D \mu \, dx \, dy \, dz.$$

Gli integrali (6) diventano a loro volta:

$$x_G^\alpha = \frac{1}{m} \iiint_D J \mu x^\alpha \, dq^1 \, dq^2 \, dq^3.$$

Scritture analoghe valgono per le superfici materiali e le curve materiali. Per i sistemi omogenei, dove la densità di massa è costante, le formule ottenute continuano a valere sostituendovi $\mu = 1$ e intendendo la massa totale m uguale al volume, o all'area, o alla lunghezza del dominio. Sussistono a questo proposito i seguenti classici **teoremi di Guldino** (Paul Guldin, 1577-1643):

PROPOSIZIONE 4. Il volume del solido generato da un dominio piano ruotato intorno ad una retta del suo piano che non lo interseca, è uguale al prodotto dell'area del dominio per la lunghezza della circonferenza descritta dal suo baricentro.

PROPOSIZIONE 5. L'area della superficie generata da un arco di curva piana ruotata intorno ad una retta del suo piano, è uguale al prodotto della lunghezza dell'arco per quella della circonferenza descritta dal suo baricentro.

Dimostrazione. Si consideri un sistema di coordinate cartesiane ortonormali tale che l'asse z sia asse di rotazione e il dominio piano stia sul piano (x, z) . Allora la coordinata x del suo baricentro è

$$x_G = \frac{1}{A} \iint_D x \, dx \, dz,$$

essendo A la sua area. D'altra parte il volume V del solido di rotazione è la somma, cioè l'integrale, dei volumi elementari degli anelli ottenuti facendo ruotare elementi del dominio, vale a dire:

$$V = \iint_D 2\pi x \, dx \, dz.$$

Quindi: $V = 2\pi x_G A$. In maniera analoga si dimostra il secondo teorema di Guldino. ■

14. Tensore d'inerzia di un sistema di masse

Nello spazio affine euclideo tridimensionale $(\mathcal{A}, E, \delta, g)$ sia dato un sistema di masse $\mathcal{M} = \{(P_\nu, m_\nu); \nu = 1, \dots, N\}$.

DEFINIZIONE 1. Chiamiamo **tensore d'inerzia** del sistema $\mathcal{M}$ nel punto $O \in \mathcal{A}$ l'endomorfismo $I_O \in \text{End}(E)$ definito da

$$(1) \quad I_O(v) = \sum_{\nu} m_{\nu} (OP_{\nu} \times v) \times OP_{\nu}$$

Per la formula del doppio prodotto vettoriale la definizione (1) equivale a

$$(1') \quad I_O(v) = \sum_{\nu} m_{\nu} (|OP_{\nu}|^2 v - OP_{\nu} \cdot v OP_{\nu}).$$

Facendo variare il punto O nello spazio affine si ottiene un campo tensoriale di tipo $(1, 1)$

$$I : \mathcal{A} \rightarrow \mathcal{A} \times T_1^1(E) : O \mapsto I_O.$$

al cui studio è dedicato questo paragrafo. Com'è noto dalla Fisica e come si vedrà più avanti, esso trova notevoli applicazioni nella meccanica dei sistemi materiali.

Dalla definizione (1) segue, per la proprietà di scambio del prodotto misto,

$$I_O(v) \cdot u = \sum_{\nu} m_{\nu} (OP_{\nu} \times v) \cdot (OP_{\nu} \times u),$$

quindi che il tensore d'inerzia I è simmetrico:

$$I_O(v) \cdot u = I_O(u) \cdot v.$$

Denotiamo ancora con lo stesso simbolo, com'è consuetudine, la forma bilineare simmetrica associata. Poniamo cioè

$$(3) \quad I_O(v, u) = I_O(v) \cdot u.$$

Denotiamo invece con I_O la forma quadratica associata a $\mathbf{I}_O$, detta **forma d'inerzia** nel punto O , definita da

$$(4) \quad I_O(\mathbf{v}) = \mathbf{I}_O(\mathbf{v}) \cdot \mathbf{v} = \sum_{\nu} m_{\nu} |\mathbf{OP}_{\nu} \times \mathbf{v}|^2$$

OSSERVAZIONE 1. Segue immediatamente dalla definizione (1) che il tensore d'inerzia $\mathbf{I}_O$ (e quindi anche la forma d'inerzia I_O) sono "funzioni additive" di sistemi materiali, vale a dire: se $\mathcal{M}'$ e $\mathcal{M}''$ sono due sistemi materiali disgiunti e $\mathbf{I}'_O$ e $\mathbf{I}''_O$ sono i rispettivi tensori d'inerzia, allora $\mathbf{I}_O = \mathbf{I}'_O + \mathbf{I}''_O$ è il tensore d'inerzia del sistema $\mathcal{M} = \mathcal{M}' \cup \mathcal{M}''$.

DEFINIZIONE 2. Chiamiamo **momento d'inerzia** rispetto ad una retta $a \subset \mathcal{A}$ il numero I^a , positivo o nullo, dato dalla somma dei prodotti delle masse m_{ν} per le distanze al quadrato dalla retta a dei rispettivi punti P_{ν} :

$$(5) \quad I^a = \sum_{\nu} m_{\nu} (d(a, P_{\nu}))^2$$

Questa definizione, anch'essa assai importante per le applicazioni, è strettamente connessa alle precedenti. Infatti, per le proprietà del prodotto vettoriale,

$$(d(a, P_{\nu}))^2 = |\mathbf{OP}_{\nu} \times \mathbf{u}|^2,$$

dove $\mathbf{u}$ è un versore della retta a ed O un suo qualunque punto. Quindi dalla (4) segue

$$(6) \quad I^a = I_O(\mathbf{u}), \quad \forall O \in a, \quad \mathbf{u} \text{ versore di } a$$

Vediamo ora come varia il tensore d'inerzia $\mathbf{I}_O$, quindi anche la forma d'inerzia I_O , al variare del punto O . Le formule seguenti, che chiamiamo **formule di trasposizione**, consentono di calcolare $\mathbf{I}_O$ e I_O in un qualunque punto $O \in \mathcal{A}$ quando siano noti nel baricentro G del sistema materiale:

$$(7) \quad \begin{aligned} \mathbf{I}_O(\mathbf{v}) &= \mathbf{I}_G(\mathbf{v}) + m(\mathbf{OG} \times \mathbf{v}) \times \mathbf{OG}, \\ I_O(\mathbf{v}) &= I_G(\mathbf{v}) + m|\mathbf{OG} \times \mathbf{v}|^2 \end{aligned}$$

essendo m la massa totale. Si ha infatti successivamente:

$$\begin{aligned} \mathbf{I}_O(\mathbf{v}) &= \sum_{\nu} m_{\nu} (\mathbf{OP}_{\nu} \times \mathbf{v}) \times \mathbf{OP}_{\nu} \\ &= \sum_{\nu} m_{\nu} ((\mathbf{OG} + \mathbf{GP}_{\nu}) \times \mathbf{v}) \times (\mathbf{OG} + \mathbf{GP}_{\nu}) \\ &= m(\mathbf{OG} \times \mathbf{v}) \times \mathbf{OG} + (\mathbf{OG} \times \mathbf{v}) \times \sum_{\nu} m_{\nu} \mathbf{GP}_{\nu} + \\ &\quad + \left(\left(\sum_{\nu} m_{\nu} \mathbf{GP}_{\nu} \right) \times \mathbf{v} \right) \times \mathbf{OG} + \sum_{\nu} m_{\nu} (\mathbf{GP}_{\nu} \times \mathbf{v}) \times \mathbf{GP}_{\nu} \\ &= m(\mathbf{OG} \times \mathbf{v}) \times \mathbf{OG} + \mathbf{I}_G(\mathbf{v}), \end{aligned}$$

ricordato che, per definizione di baricentro, è $\sum_{\nu} m_{\nu} \mathbf{GP}_{\nu} = 0$. La seconda della (7) segue immediatamente dalla prima. Immediata conseguenza delle formule di trasposizione è il seguente **teorema di Huygens** (Christiaan Hygens, 1629-1695) sui momenti d'inerzia:

PROPOSIZIONE 1. Sia a una retta e sia g la retta parallela ad a passante per il baricentro G di un sistema di masse $\mathcal{M}$, allora

$$(8) \quad I^a = I^g + m(d(a, g))^2$$

dove m è la massa totale e $d(a, g)$ è la distanza delle due rette.

Dimostrazione. Si consideri la definizione (6) e si applichi la seconda delle (7):

$$I^a = I_O(\mathbf{u}) = I_G(\mathbf{u}) + m|\mathbf{OG} \times \mathbf{u}|^2 = I^g + m(d(a, g))^2. \quad \blacksquare$$

COROLLARIO 1. Al variare della retta a in un fascio di rette parallele, il momento d'inerzia I^a è minimo quando a passa per il baricentro.

Studiamo il comportamento del tensore e della forma d'inerzia in un punto O al variare del vettore $\mathbf{v}$. Osserviamo innanzitutto dalla (4) che $I_O(\mathbf{v}) \geq 0$. Si ha $I_O(\mathbf{v}) = 0$ se $\mathbf{v} = 0$ oppure se $\mathbf{OP}_{\nu}$ è parallelo a $\mathbf{v}$ per ogni punto P_{ν} . In quest'ultimo caso tutti i punti del sistema materiale si trovano su di una retta passante per O e parallela a $\mathbf{v}$. Chiamiamo questo caso **caso degenere**. Se il sistema si riduce al punto O il tensore

d'inerzia in O è identicamente nullo. In conclusione: *la forma d'inerzia I_O è definita positiva, al di fuori del caso degenerare in cui è semi-definita positiva.*

Di conseguenza, richiamate le proprietà degli endomorfismi simmetrici in spazi strettamente euclidei, possiamo affermare che: *il tensore d'inerzia I_O ammette tre autovalori reali (I_1, I_2, I_3) detti **momenti principali d'inerzia** relativi al punto O . Nel caso non degenerare sono tutti e tre positivi, nel caso degenerare uno di essi è nullo e gli altri due sono uguali fra loro (per ragioni di simmetria, come sarà chiarito fra poco). Le autodirezioni di I_O , tra loro ortogonali, determinano tre rette passanti per O , (a_1, a_2, a_3) dette **assi principali d'inerzia** relativi al punto O . Questi assi sono univocamente determinati se e solo se i momenti principali d'inerzia sono tutti distinti. I momenti principali d'inerzia e gli assi principali d'inerzia relativi al baricentro si chiamano rispettivamente **momenti centrali d'inerzia** e **assi centrali d'inerzia**.*

OSSERVAZIONE 2. I momenti principali d'inerzia (I_1, I_2, I_3) in un punto O sono i momenti d'inerzia relativi ai rispettivi assi principali (a_1, a_2, a_3). Infatti, se per esempio u_1 è un autovettore unitario associato all'autovalore I_1 , da $I_O(u_1) = I_1 u_1$ segue $I_O(u_1) = I_O(u_1) \cdot u_1 = I_1 u_1 \cdot u_1 = I_1$.

Come vedremo tra poco, la conoscenza degli assi principali d'inerzia è importante perché consente notevoli semplificazioni nel calcolo dei momenti d'inerzia. La loro determinazione è facilitata dalla seguente importante **proprietà di simmetria**:

PROPOSIZIONE 2. Se Π è un piano di **simmetria materiale ortogonale** del sistema materiale $\mathcal{M}$ (cioè se è un piano di simmetria coniugato ad un versore ortogonale u), allora qualunque sia $O \in \Pi$ la retta per O ortogonale a Π è asse principale d'inerzia in O .

Dimostrazione. Si prenda un versore u ortogonale al piano Π e si considerino due punti materiali simmetrici: (P_1, m_1) e (P_2, m_2) . Per definizione di simmetria materiale ortogonale si ha $m_1 = m_2$, $P_1 P_2$ parallelo ad u , e il punto medio di $P_1 P_2$ appartiene a Π : allora, qualunque sia $O \in \Pi$, si ha $|OP_1| = |OP_2|$ e $OP_1 \cdot u = -OP_2 \cdot u$. Per il sistema

materiale costituito da questi due soli punti si ha di conseguenza:

$$\begin{aligned} I_O(u) &= m_1 (|OP_1|^2 u - OP_1 \cdot u OP_1) + m_2 (|OP_2|^2 u - OP_2 \cdot u OP_2) \\ &= m_1 (2|OP_1|^2 u + OP_1 \cdot u P_1 P_2) \\ &= m_1 (2|OP_1|^2 \pm OP_1 \cdot u |P_1 P_2|) u \\ &= I u. \end{aligned}$$

Ciò mostra che u è autovettore di I_O . Invece per un sistema costituito da un solo punto materiale (P, m) posto su Π , si ha semplicemente $I_O(u) = m|OP|^2 u$, perché $OP \cdot u = 0$. Anche in questo caso u è autovettore. Ciò basta a dimostrare il teorema in virtù della proprietà additiva del tensore d'inerzia (Oss. 1). ■

Poiché gli assi principali d'inerzia sono fra loro ortogonali, l'enunciato del teorema può continuare con: ... e gli altri due assi principali d'inerzia giacciono su Π .

La dimostrazione del seguente corollario, assai utile nelle applicazioni, è lasciata come esercizio.

COROLLARIO 2. Se un sistema materiale $\mathcal{M}$ ammette due piani di simmetria materiale ortogonale fra loro ortogonali, allora per ogni punto della retta intersezione gli assi principali d'inerzia sono la retta intersezione stessa e le due rette ortogonali giacenti sui due piani. Se i due piani di simmetria ortogonale non sono fra loro ortogonali, allora la retta intersezione è asse principale d'inerzia per ogni suo punto, gli altri due assi principali sono a questa ortogonali ma indeterminati ed i due momenti d'inerzia corrispondenti sono uguali.

Veniamo infine al calcolo delle componenti del tensore d'inerzia. Si considerino un riferimento ortonormale di origine il punto O , $(O, c_\alpha) = (O, i, j, k)$ e le corrispondenti coordinate $(x^\alpha) = (x, y, z)$. Denotiamo con $(x_\nu^\alpha) = (x_\nu, y_\nu, z_\nu)$ le coordinate del punto P_ν . Tenuto conto della definizione (1') e del fatto che $OP_\nu = x_\nu^\alpha c_\alpha$ e $|OP_\nu|^2 = \sum_{\alpha=1}^3 (x_\nu^\alpha)^2$, posto

$$I_{\alpha\beta} = I_O(c_\alpha) \cdot c_\beta,$$

si ottiene:

$$(9) \quad \boxed{\begin{aligned} I_{\alpha\alpha} &= \sum_{\nu} m_{\nu} ((x_{\nu}^{\alpha+1})^2 + (x_{\nu}^{\alpha+2})^2) \\ I_{\alpha\beta} &= - \sum_{\nu} m_{\nu} x_{\nu}^{\alpha} x_{\nu}^{\beta} \quad (\alpha \neq \beta) \end{aligned}}$$

Nei testi classici la matrice simmetrica $(I_{\alpha\beta})$ delle componenti della forma d'inerzia, detta **matrice d'inerzia**, è denotata con

$$(I_{\alpha\beta}) = \begin{pmatrix} A & -C' & -B' \\ -C' & B & -A' \\ -B' & -A' & C \end{pmatrix},$$

cosicché le (9) diventano:

$$(10) \quad \begin{aligned} A &= \sum_{\nu} m_{\nu} (y_{\nu}^2 + z_{\nu}^2), & A' &= \sum_{\nu} m_{\nu} y_{\nu} z_{\nu}, \\ B &= \sum_{\nu} m_{\nu} (z_{\nu}^2 + x_{\nu}^2), & B' &= \sum_{\nu} m_{\nu} z_{\nu} x_{\nu}, \\ C &= \sum_{\nu} m_{\nu} (x_{\nu}^2 + y_{\nu}^2), & C' &= \sum_{\nu} m_{\nu} x_{\nu} y_{\nu}. \end{aligned}$$

OSSERVAZIONE 3. Se gli assi coordinati sono assi principali d'inerzia per il punto O , allora la matrice metrica si deve diagonalizzare, quindi necessariamente

$$A' = B' = C' = 0$$

e (A, B, C) sono i momenti principali d'inerzia.

OSSERVAZIONE 4. Dalle prime tre equazioni (10) si traggono le disuguaglianze

$$A + B \geq C, \quad B + C \geq A, \quad C + A \geq B.$$

La matrice d'inerzia consente il calcolo del momento d'inerzia rispetto ad una qualunque retta a passante per O . Infatti dalla (6) e (4) si trae

$$(11) \quad I^a = I_{\alpha\beta} u^{\alpha} u^{\beta},$$

dove si sono denotati con (u^{α}) le componenti di un versore u della retta a . Queste componenti sono anche dette **coseni direttori** della retta a e sono classicamente denotati con (α, β, γ) . Cosicché in notazioni classiche la (11) diventa:

$$(12) \quad I^a = A\alpha^2 + B\beta^2 + C\gamma^2 - 2(A'\beta\gamma + B'\gamma\alpha + C'\alpha\beta).$$

Se gli assi sono assi principali d'inerzia allora si ha più semplicemente:

$$(13) \quad \boxed{I^a = A\alpha^2 + B\beta^2 + C\gamma^2}$$

OSSERVAZIONE 5. *L'ellissoide d'inerzia.* Si consideri la quadrica Q di centro O definita dalle equazioni

$$(14) \quad I_{\alpha\beta} x^{\alpha} x^{\beta} = 1,$$

ovvero dalle equazioni (se gli assi sono principali d'inerzia):

$$(15) \quad Ax^2 + By^2 + Cz^2 = 1.$$

Siccome la forma d'inerzia è definita positiva, almeno nel caso non degenero, questa quadrica è un ellissoide, detto **ellissoide d'inerzia** nel punto O , ed i suoi assi sono gli assi principali d'inerzia. Nel caso degenero, dove tutti i punti sono distribuiti su di una retta passante per O , la quadrica Q definita dalla (15) è un cilindro retto circolare avente come asse la retta; infatti uno dei due momenti principali, quello relativo alla retta materiale, si annulla e gli altri due sono uguali (si veda il Corollario 2 oppure si guardino le (10)). In ogni caso, dal confronto della (14) con la (11) si osserva che se P^a è il punto d'intersezione della quadrica Q con la retta a passante per O , essendo $x^{\alpha} = |OP^a| u^{\alpha}$, risulta:

$$(16) \quad |OP^a| = \frac{1}{\sqrt{I^a}}.$$

OSSERVAZIONE 6. *Sistemi materiali piani.* Si consideri un **sistema materiale piano**: tutti i punti materiali giacciono su di un piano Π . Questo piano è ovviamente piano di simmetria materiale ortogonale e quindi vale l'enunciato del Teorema 2. Se C è il momento principale d'inerzia rispetto all'asse ortogonale al piano passante per un suo qualunque punto O (che diventa asse z) allora

$$(17) \quad C = A + B$$

Per riconoscerlo basta utilizzare le formule (10), tenendo conto che $z_\nu = 0$.

OSSERVAZIONE 7. *Sistemi materiali continui.* Le definizioni di tensore d'inerzia rispetto ad un punto e di momento d'inerzia rispetto ad una retta sono estendibili a un sistema continuo di masse, secondo quanto già osservato per il baricentro. Valgono le stesse proprietà viste per il caso di un sistema finito di punti. Per esempio, la definizione (1) per un sistema continuo tridimensionale $(M, \mu\eta)$ diventa

$$(18) \quad I_O(v) = \int_M \mu (\mathbf{r} \times \mathbf{v}) \times \mathbf{r} \eta,$$

dove $\mathbf{r}$ è il vettore posizione rispetto ad O del generico punto P del sistema continuo, η è la forma volume, μ è la densità di massa. Così le (10) si trasformano in integrali tripli:

$$(19) \quad \begin{aligned} A &= \iiint_D \mu J (y^2 + z^2) dq^1 dq^2 dq^3, & A' &= \iiint_D \mu J y z dq^1 dq^2 dq^3, \\ B &= \iiint_D \mu J (z^2 + x^2) dq^1 dq^2 dq^3, & B' &= \iiint_D \mu J z x dq^1 dq^2 dq^3, \\ C &= \iiint_D \mu J (x^2 + y^2) dq^1 dq^2 dq^3, & C' &= \iiint_D \mu J x y dq^1 dq^2 dq^3, \end{aligned}$$

dove $D \subset \mathbb{R}^3$ è il dominio di integrazione nelle coordinate (q^i) , e dove la densità μ , lo jacobiano J e le coordinate cartesiane (x, y, z) devono essere espresse in funzione delle variabili (q^i) . Formule analoghe valgono per sistemi bidimensionali e unidimensionali.

15. Sistemi di vettori applicati

DEFINIZIONE 1. Sia $S = \{(P_\nu, \mathbf{v}_\nu); \nu = 1, \dots, N\}$ un **sistema di vettori applicati** nello spazio affine euclideo tridimensionale. Diciamo **risultante** del sistema S il vettore

$$(1) \quad \boxed{\mathbf{R} = \sum_{\nu} \mathbf{v}_\nu}$$

momento risultante rispetto al punto o **polo** O il vettore

$$(2) \quad \boxed{\mathbf{M}_O = \sum_{\nu} OP_\nu \times \mathbf{v}_\nu}$$

Il momento $OP \times \mathbf{v}$ di un singolo vettore applicato $(P, \mathbf{v})$ è, per le proprietà del prodotto vettoriale, un vettore ortogonale al piano individuato dai vettori OP e $\mathbf{v}$, di modulo uguale al prodotto del modulo di $\mathbf{v}$ per la distanza del polo O dalla **retta di applicazione** di $(P, \mathbf{v})$, cioè della retta passante per P e parallela a $\mathbf{v}$.

Facendo variare il polo nello spazio affine si ottiene un campo vettoriale

$$\mathbf{M}: A \rightarrow A \times E: P \mapsto \mathbf{M}_P.$$

La seguente **formula di trasposizione dei momenti** consente di calcolare il momento risultante nel polo P , noto il suo valore in un polo O e il risultante $\mathbf{R}$:

$$(3) \quad \boxed{\mathbf{M}_P = \mathbf{M}_O + \mathbf{R} \times OP}$$

Per dimostrarla basta considerare la decomposizione $OP_\nu = OP + PP_\nu$ nella definizione (2). Dalla (3) si deducono le proprietà seguenti:

$$(4) \quad \mathbf{M}_O = \mathbf{M}_P \iff \mathbf{R} \text{ parallelo a } OP,$$

$$(5) \quad \mathbf{M}_O = \mathbf{M}_P, \forall O, P \in A \iff \mathbf{R} = 0,$$

$$(6) \quad \mathbf{M}_O \cdot \mathbf{R} = \mathbf{M}_P \cdot \mathbf{R}.$$

La proprietà (4) mostra che il momento risultante non varia spostando il polo lungo una retta parallela al risultante $\mathbf{R}$. La (5) mostra che il momento risultante non dipende dalla scelta del polo se e solo se il risultante è nullo. La (6) mostra che la componente del momento risultante rispetto al risultante è invariante rispetto alla scelta del polo. L'andamento del campo vettoriale $\mathbf{M}$, nel caso in cui $\mathbf{R} \neq 0$ è ulteriormente chiarito dalla seguente

PROPOSIZIONE 1. Dato un sistema di vettori applicati con $\mathbf{R} \neq 0$, esiste un luogo di punti $a \subset \mathcal{A}$ tale che per ogni $A \in a$ il momento risultante $\mathbf{M}_A$ è parallelo ad $\mathbf{R}$. Questo luogo è una retta parallela ad $\mathbf{R}$, detta **asse centrale** del sistema.

Dimostrazione. Il luogo di punti a cercato è caratterizzato dall'equazione

$$\mathbf{M}_A \times \mathbf{R} = 0.$$

Per la formula di trasposizione $\mathbf{M}_A = \mathbf{M}_O + \mathbf{R} \times OA$, da questa si trae l'equazione

$$(7) \quad \mathbf{M}_O \times \mathbf{R} + (\mathbf{R} \times OA) \times \mathbf{R} = 0$$

nel vettore OA , dove il punto O ed i vettori $\mathbf{R}$ e $\mathbf{M}_O$ sono noti. Se quest'equazione è soddisfatta da un punto A essa è soddisfatta anche per tutti i punti $A' = A + k\mathbf{R}$ che stanno sulla retta per A parallela ad $\mathbf{R}$:

$$\mathbf{M}_O \times \mathbf{R} + (\mathbf{R} \times OA') \times \mathbf{R} = \mathbf{M}_O \times \mathbf{R} + (\mathbf{R} \times OA) \times \mathbf{R} = 0.$$

Dunque il luogo cercato è costituito da rette parallele a $\mathbf{R}$. Per dimostrare che esso si riduce ad una sola retta, si consideri il piano per O contenente A e perpendicolare a $\mathbf{R}$. Qui sopra l'equazione (7), sviluppato il doppio prodotto vettoriale ed essendo $OA \cdot \mathbf{R} = 0$, si riduce a

$$\mathbf{R} \times \mathbf{M}_O - |\mathbf{R}|^2 OA = 0.$$

Quest'equazione definisce un solo punto A , intersezione del luogo cercato col piano. ■

Riscriviamo la formula di trasposizione dei momenti prendendo come punto di riferimento un punto A dell'asse centrale:

$$(8) \quad \mathbf{M}_P = \mathbf{M}_A + \mathbf{R} \times AP.$$

Si osserva allora che il generico momento risultante $\mathbf{M}_P$ è somma di un vettore $\mathbf{M}_A$ parallelo ad $\mathbf{R}$, cioè all'asse centrale, e di un vettore $\mathbf{R} \times AP$ ortogonale ad $\mathbf{R}$ il cui modulo cresce in maniera proporzionale alla distanza del polo P dall'asse centrale. La Fig. 15.1 rappresenta il variare di $\mathbf{M}_P$ al variare di P su di una semiretta ortogonale all'asse centrale a . Per rappresentare l'andamento del campo $\mathbf{M}$ in tutto lo spazio occorre ruotare tale figura intorno all'asse centrale e traslarla lungo questo.

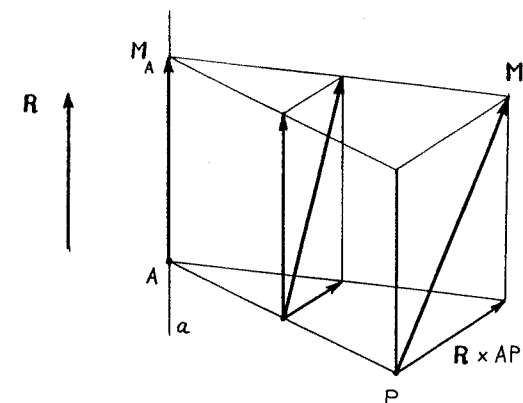


Fig. 15.1

Poiché i vettori $\mathbf{M}_A$ ed $\mathbf{R}$ sono paralleli, dalla (8) si trae

$$M_P^2 = M_A^2 + (\mathbf{R} \times AP)^2.$$

Quest'uguaglianza dimostra la seguente proprietà caratteristica dell'asse centrale (che potrebbe essere assunta come definizione):

PROPOSIZIONE 2. L'asse centrale è il luogo dei punti in cui il momento risultante ha minima intensità.

OSSERVAZIONE 1. Due sistemi di vettori applicati $\mathcal{S}$ e $\mathcal{S}'$ si dicono **equivalenti** se i corrispondenti momenti risultanti coincidono in ogni punto: $\mathbf{M} = \mathbf{M}'$. Per questo occorre e basta che i due sistemi abbiano lo stesso risultante e lo stesso momento risultante rispetto ad un polo prefissato:

$$\mathcal{S} \sim \mathcal{S}' \quad \Longleftrightarrow \quad \mathbf{R} = \mathbf{R}', \quad \mathbf{M}_O = \mathbf{M}'_O.$$

CAPITOLO III

CINEMATICA DEL PUNTO
E DEL CORPO RIGIDO

La cinematica studia il moto dei corpi indipendentemente dalle cause che lo provocano e dalle leggi fisiche che lo governano. È quindi un ramo della meccanica avente carattere puramente descrittivo, che utilizza un linguaggio e dei modelli matematici atti a rappresentare lo spazio fisico ed i corpi che in esso si muovono. La scelta di questi modelli non è unica.

Il modello classico dello spazio fisico osservato da un riferimento è lo spazio affine tridimensionale euclideo. È un modello immediatamente accettabile dal nostro intuito. Quando osserviamo lo spazio che ci circonda, anche a estreme distanze, presupponiamo sempre la scelta di un riferimento, per esempio la nostra aula di lezione, il treno su cui viaggiamo, la Terra, un sistema di tre assi uscenti dal Sole e diretti verso tre stelle fisse, etc. . In ogni caso identifichiamo un riferimento con un "corpo rigido" ideale, invadente tutto l'universo, o almeno quella parte di universo in cui si evolve il sistema il cui moto vogliamo descrivere e studiare. Ebbene, questo "corpo rigido" è modellato nello spazio affine euclideo.

Noi concepiamo l'idea dell'esistenza di più riferimenti, ma la scelta ricade su di uno solamente. La scelta di un riferimento è, a livello puramente cinematico, una questione di pura convenienza descrittiva. Noi

concepiamo anche l'idea che due riferimenti debbano essere considerati distinti se i corrispondenti "corpi rigidi" sono in moto l'uno rispetto all'altro. Così, negli esempi sopra citati, la nostra aula e la Terra costituiscono un medesimo riferimento, anche se la loro estensione, e quindi la scala dei fenomeni a cui li rapportiamo, è ben diversa. Si pone comunque il problema di stabilire i legami che intercorrono tra gli enti cinematici associati ad un medesimo moto ma relativi a due diversi riferimenti. A questo problema è dedicato un paragrafo di questo capitolo, ma su di esso si ritornerà in capitoli successivi.

Il concetto di moto è strettamente connesso non solo con quello di spazio ma anche con quello di tempo. Dobbiamo quindi anche scegliere un modello per il tempo. In questo capitolo il ruolo del tempo è ridotto a quello di parametro o di variabile indipendente. Si vedrà nei capitoli successivi come invece il concetto di tempo possa assumere un ruolo diverso, fondendosi con quello di spazio per costituire un'unica struttura assoluta: lo "spazio-tempo".

Occorre infine scegliere un modello per il corpo o i corpi il cui moto si vuole studiare. Il modello fondamentale, ed anche il più semplice, è quello di "punto". Il moto di un punto è rappresentato da una curva nello spazio affine tridimensionale euclideo. Questo modello è accettabile solo se l'estensione del corpo è del tutto trascurabile rispetto all'estensione del suo moto e alle altre grandezze significative che intervengono nel fenomeno studiato. Un secondo modello fondamentale è quello di "corpo rigido": esso è adottato per quei corpi estesi le cui particelle mantengono sensibilmente invariate, durante il moto, le mutue distanze.

1 . Cinematica del punto

Il moto di un punto è rappresentato da una curva $\gamma: I \rightarrow \mathcal{A}: t \mapsto \gamma(t)$ nello spazio affine tridimensionale euclideo $\mathcal{A}$. Il parametro t è il tempo, l'intervallo reale di definizione I rappresenta l'estensione temporale del moto. Ad ogni istante $t \in I$ corrisponde un punto $\gamma(t) \in \mathcal{A}$ che rappresenta la posizione occupata dal corpo in quell'istante, punto che denotiamo anche con $P(t)$. Come si è già visto al §5 del Cap. II, scelto un punto O di riferimento, il moto di un punto è descritto da una **rappresentazione vettoriale**

$$(1) \quad OP = r(t),$$

la quale, scelto un riferimento affine (O, c_α) , si traduce in **equazioni parametriche** cartesiane

$$(2) \quad x^\alpha = \gamma^\alpha(t) \quad (\alpha = 1, 2, 3).$$

Noi supporremo che queste funzioni siano di classe C^2 almeno. Da questa rappresentazione, con due derivate temporali successive, si ricavano la **velocità** e l'**accelerazione** del punto

$$(3) \quad v = \frac{dr}{dt}, \quad a = \frac{dv}{dt},$$

che sono a loro volta funzioni vettoriali di t . Le tre funzioni vettoriali del tempo $(r(t), v(t), a(t))$ sono gli **enti cinematici fondamentali** del punto. Conveniamo di considerare i vettori $v(t)$ e $a(t)$ applicati nel punto $P(t)$.

L'immagine $\gamma(I)$ della curva γ rappresentante il moto prende il nome di **orbita** o **traiettoria**. Nei casi più comuni l'orbita è un sottinsieme (connesso) di una curva regolare dello spazio affine. Se su questa curva si considera un'ascissa euclidea s avente origine in un punto P_0 (per esempio quello corrispondente a $t = 0$), allora il moto è completamente determinato da una funzione reale a variabile reale,

$$(4) \quad s = s(t)$$

che prende il nome di **legge oraria** del moto. Orbita e legge oraria forniscono dunque una descrizione completa del moto, scindendolo nella sua parte puramente geometrica e nella sua parte temporale. In base a questa distinzione possiamo dare della velocità e dell'accelerazione le seguenti **rappresentazioni intrinseche**:

$$(5) \quad v = \dot{s} t, \quad a = \ddot{s} t + c \dot{s}^2 n$$

dove t e n sono il versore tangente e il versore normale principale alla traiettoria, c è la sua curvatura, e dove $\dot{s}$ e $\ddot{s}$ denotano le derivate prime e seconde della legge oraria. Infatti (si veda il §12 del Cap. II), il versore tangente è definito da

$$t = \frac{dr}{ds} = \frac{dt}{ds} v,$$

e quindi

$$v = \frac{ds}{dt} t = \dot{s} t.$$

Ricordato poi che

$$\frac{dt}{ds} = c n,$$

si ha successivamente:

$$a = \frac{dv}{dt} = \frac{d\dot{s}}{dt} t + \dot{s} \frac{dt}{dt} = \ddot{s} t + \dot{s} \frac{dt}{ds} \frac{ds}{dt} = \ddot{s} t + c \dot{s}^2 n.$$

Da questa rappresentazione si osserva che l'accelerazione è, in ogni istante ovvero in ogni punto dell'orbita, la somma di due vettori diretti secondo la tangente all'orbita e alla normale principale,

$$(6) \quad a_t = \ddot{s} t, \quad a_n = c \dot{s}^2 n,$$

che chiamiamo rispettivamente **accelerazione tangente** e **accelerazione normale**. Di conseguenza, in ogni punto dell'orbita, l'accelerazione appartiene al piano osculatore. Le funzioni reali $\dot{s}$ e $\ddot{s}$ sono dette rispettivamente **velocità scalare** e **accelerazione scalare**.

Alla rappresentazione vettoriale del moto (1) possiamo affiancare la **rappresentazione radiale di centro O** , costituita da una coppia di funzioni del tempo $(r(t), u(t))$, una reale l'altra vettoriale, tali che:

$$(7) \quad r = r u, \quad |u| = 1, \quad r > 0$$

Pertanto, in ogni istante t , u è il versore del vettore posizione del punto rispetto al centro O mentre r è la sua distanza da questo. Una tale rappresentazione è valida per moti non passanti per il centro O . Se ad un certo istante t_0 il punto mobile passa per O , nel quale ovviamente u risulta indeterminato, può essere possibile e conveniente una rappresentazione radiale $(r(t), u(t))$ estesa per continuità a t_0 , con la funzione r non ristretta a valori positivi (si pensi, per esempio, al semplice caso del moto di un punto su di una retta per O di versore, costante, u : nell'attraversare O può restare valida la rappresentazione (7) pur di lasciar assumere ad r valori positivi, negativi e nulli).

Un ulteriore importante ente cinematico è la **velocità areale** rispetto ad un punto O :

$$(8) \quad v_{ar} = \frac{1}{2} r \times v$$

È, a meno del fattore $\frac{1}{2}$, il momento rispetto al punto fisso O del vettore velocità v pensato applicato in P .

Vediamo ora le definizioni di alcuni tipi particolari ma fondamentali di moti.

DEFINIZIONE 1. Moto uniforme: è un moto a velocità scalare costante, quindi caratterizzato dalla condizione $|\mathbf{v}| = \text{cost.}$ ovvero $\dot{s} = \text{cost.}$

Non esistono istanti di arresto, a meno che non sia proprio $\mathbf{v} = 0$, nel qual caso il punto è immobile e la sua orbita si riduce ad un punto. Al di fuori di questo caso l'orbita ammette sempre il versore tangente $\mathbf{t}$. Dalla (5)₂ si vede che un moto è uniforme se e solo se la sua accelerazione è sempre normale (cioè ortogonale alla traiettoria). Un esempio è il moto circolare uniforme, già considerato al §5 del Cap. II.

DEFINIZIONE 2. Moto rettilineo uniforme o moto inerziale: è un moto con velocità vettoriale costante, $\mathbf{v} = \text{cost.}$, quindi caratterizzato dalla condizione $\mathbf{a} = 0$.

La rappresentazione vettoriale del moto è di conseguenza

$$(9) \quad \mathbf{r} = \mathbf{v} t + \mathbf{r}_0,$$

dove $\mathbf{r}_0$ è la posizione del punto per $t = 0$. L'orbita è una retta, percorsa con velocità costante.

DEFINIZIONE 3. Moto piano: è un moto la cui orbita giace su di un piano.

Esaminiamo la rappresentazione radiale $(r(t), \mathbf{u}(t))$ nel caso di un moto piano. Si considerino nel piano del moto coordinate polari (r, ϑ) centrate in un punto O . L'angolo ϑ è detto **anomalia** ed r **raggio**. Si consideri inoltre nello spazio un riferimento canonico $(O, \mathbf{i}, \mathbf{j}, \mathbf{k})$ tale che il piano del moto coincida col piano $(O, \mathbf{i}, \mathbf{j})$, e che $\mathbf{k}$ sia il versore di ϑ (Fig. 1.1). Ciò posto, la derivata rispetto al tempo del versore $\mathbf{u}$ di OP risulta essere

$$(10) \quad \dot{\mathbf{u}} = \dot{\vartheta} \mathbf{k} \times \mathbf{u}$$

dove $\dot{\vartheta}$ è la derivata di ϑ rispetto al tempo, che chiamiamo **velocità angolare**. Per dimostrare la (10) si può innanzitutto osservare che

$$\mathbf{u} = \cos \vartheta \mathbf{i} + \sin \vartheta \mathbf{j}.$$

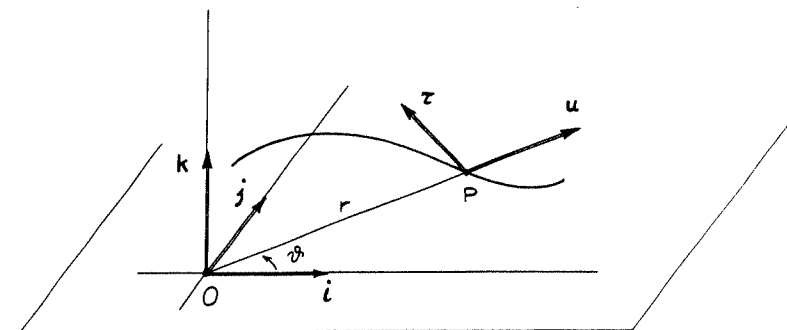


Fig. 1.1

Derivando quest'uguaglianza rispetto al tempo si ha successivamente:

$$\dot{\mathbf{u}} = \frac{d\mathbf{u}}{d\vartheta} \frac{d\vartheta}{dt} = (-\sin \vartheta \mathbf{i} + \cos \vartheta \mathbf{j}) \dot{\vartheta} = \dot{\vartheta} \mathbf{k} \times \mathbf{u}.$$

Convienne introdurre il **versore trasverso**

$$(11) \quad \boldsymbol{\tau} = \mathbf{k} \times \mathbf{u}.$$

Allora la (10) diventa

$$(10') \quad \dot{\mathbf{u}} = \dot{\vartheta} \boldsymbol{\tau}.$$

Poiché $\boldsymbol{\tau}$ è ortogonale a $\mathbf{u}$ e orientato secondo ϑ crescente, risulta

$$(12) \quad \dot{\boldsymbol{\tau}} = -\dot{\vartheta} \mathbf{u}.$$

Derivando due volte rispetto al tempo la rappresentazione radiale (7) $\mathbf{r} = r \mathbf{u}$, si ottengono le espressioni della velocità e dell'accelerazione in coordinate polari:

$$(13) \quad \begin{aligned} \mathbf{v} &= \dot{r} \mathbf{u} + r \dot{\vartheta} \boldsymbol{\tau}, \\ \mathbf{a} &= (\ddot{r} - r \dot{\vartheta}^2) \mathbf{u} + (2 \dot{r} \dot{\vartheta} + r \ddot{\vartheta}) \boldsymbol{\tau}. \end{aligned}$$

La velocità e l'accelerazione risultano decomposte nella somma di due vettori, uno parallelo al versore radiale $\mathbf{u}$, la **velocità radiale** e l'**accelerazione radiale** rispettivamente, e uno ortogonale a questo, la **velocità trasversa** e l'**accelerazione trasversa**:

$$(14) \quad \begin{aligned} \mathbf{v}_{rad} &= \dot{r} \mathbf{u}, & \mathbf{a}_{rad} &= (\ddot{r} - r \dot{\vartheta}^2) \mathbf{u}, \\ \mathbf{v}_{trasv} &= r \dot{\vartheta} \boldsymbol{\tau}, & \mathbf{a}_{trasv} &= (2 \dot{r} \dot{\vartheta} + r \ddot{\vartheta}) \boldsymbol{\tau}. \end{aligned}$$

OSSERVAZIONE 1. Ricordiamo (§3, Cap. II) che i vettori del riferimento $(\mathbf{E}_r, \mathbf{E}_\vartheta)$ associato alle coordinate polari sono dati da

$$\mathbf{E}_r = \mathbf{u}, \quad \mathbf{E}_\vartheta = r \boldsymbol{\tau}.$$

Le (13) possono quindi anche scriversi

$$\begin{cases} \mathbf{v} = \dot{r} \mathbf{E}_r + \dot{\vartheta} \mathbf{E}_\vartheta, \\ \mathbf{a} = (\ddot{r} - r \dot{\vartheta}^2) \mathbf{E}_r + \left(\ddot{\vartheta} + \frac{2}{r} \dot{r} \dot{\vartheta} \right) \mathbf{E}_\vartheta. \end{cases}$$

Quindi, denotate con $(v^i) = (v^r, v^\vartheta)$ e $(a^i) = (a^r, a^\vartheta)$ le componenti della velocità e dell'accelerazione in coordinate polari, risulta

$$v^r = \dot{r}, \quad v^\vartheta = \dot{\vartheta},$$

a conferma della formula generale (formula (10), §5, Cap. II)

$$v^i = \frac{dq^i}{dt} = \dot{q}^i,$$

e, per l'accelerazione,

$$a^r = \ddot{r} - r \dot{\vartheta}^2, \quad a^\vartheta = \ddot{\vartheta} + \frac{2}{r} \dot{r} \dot{\vartheta},$$

a conferma della formula generale (formula (12), §5, Cap. II)

$$a^i = \frac{dv^i}{dt} + \Gamma_{jh}^i v^j v^h = \ddot{q}^i + \Gamma_{jh}^i \dot{q}^j \dot{q}^h,$$

essendo i simboli di Christoffel non nulli dati da (formula (20), §3, Cap. II)

$$\Gamma_{22}^1 = \Gamma_{\vartheta\vartheta}^r = -r, \quad \Gamma_{12}^2 = \Gamma_{r\vartheta}^\vartheta = \frac{1}{r}.$$

Dalla (7) e dalla prima delle (13) segue per la velocità areale di un moto piano l'espressione

$$(15) \quad \boxed{v_{ar} = \frac{1}{2} r^2 \dot{\vartheta} \mathbf{k}}$$

Essa è dunque un vettore ortogonale al piano del moto il cui modulo è pari alla derivata rispetto al tempo della funzione $A(t)$ data dall'area descritta dal vettore $\mathbf{r}(t)$ a partire da un istante prefissato qualsiasi. Infatti, considerato un qualunque istante t e un incremento positivo Δt del tempo, denotati con $\Delta\vartheta$ e ΔA i corrispondenti incrementi dell'anomalia e dell'area, risulta, nell'ipotesi che la funzione r sia crescente nell'intervallo $(t, t + \Delta t)$,

$$\frac{1}{2} r^2(t) \Delta\vartheta \leq \Delta A \leq \frac{1}{2} r^2(t + \Delta t) \Delta\vartheta,$$

quindi, dividendo per $\Delta t > 0$,

$$\frac{1}{2} r^2(t) \frac{\Delta\vartheta}{\Delta t} \leq \frac{\Delta A}{\Delta t} \leq \frac{1}{2} r^2(t + \Delta t) \frac{\Delta\vartheta}{\Delta t}.$$

Per $\Delta t \rightarrow 0$ si trova

$$(16) \quad \dot{A} = \frac{1}{2} r^2 \dot{\vartheta}.$$

Si osservi dalla (14)₄ che l'accelerazione trasversa è, a meno di un fattore $\frac{2}{r}$, la derivata della velocità areale scalare:

$$(17) \quad \mathbf{a}_{trasv} = \frac{1}{r} \frac{d}{dt} (r^2 \dot{\vartheta}) \boldsymbol{\tau} = \frac{2}{r} \ddot{A} \boldsymbol{\tau}.$$

DEFINIZIONE 4. **Moto uniformemente accelerato**: è un moto ad accelerazione costante, $\mathbf{a} = \text{cost.}$

Con due integrazioni successive si ottiene:

$$(18) \quad \begin{aligned} \mathbf{v} &= \mathbf{a} t + \mathbf{v}_0, \\ \mathbf{r} &= \frac{1}{2} \mathbf{a} t^2 + \mathbf{v}_0 t + \mathbf{r}_0, \end{aligned}$$

con $\mathbf{v}_0$ ed $\mathbf{r}_0$ velocità e posizione all'istante t_0 . Se $\mathbf{a} = 0$ si ritrova il moto rettilineo uniforme. Se $\mathbf{a} \neq 0$ e $\mathbf{v}_0$ è parallelo ad $\mathbf{a}$ (in particolare $\mathbf{v}_0 = 0$) il moto è rettilineo. Se $\mathbf{v}_0$ non è parallelo ad $\mathbf{a}$, il moto è piano: il piano del moto è individuato dalla terna $(P_0, \mathbf{v}_0, \mathbf{a})$, con P_0 punto posizione iniziale. La traiettoria è una parabola. Esempio fondamentale di moto uniformemente accelerato è il **moto di un grave**.

DEFINIZIONE 5. Moto periodico: un moto è detto periodico se è definito su tutto l'asse temporale e se esiste un numero $T > 0$ tale che per ogni $t \in \mathbb{R}$

$$(19) \quad \mathbf{r}(t) = \mathbf{r}(t + T).$$

Se un tal numero esiste, non è unico (si osservi infatti che per ogni intero k , il numero kT sostituito a T rende ancora valida la (19)). Il più piccolo dei numeri positivi T per cui vale la (19) prende il nome di **periodo** del moto, il numero $\nu = \frac{1}{T}$ di **frequenza**, il numero $\omega = 2\pi\nu = \frac{2\pi}{T}$ di **pulsazione**.

2. Moti centrali

DEFINIZIONE 1. Un moto si dice **centrale** se esiste un punto fisso $O \in \mathcal{A}$, detto **centro del moto**, tale che l'accelerazione $\mathbf{a}(t)$ è in ogni istante t parallela al vettore $OP = \mathbf{r}(t)$.

Questa condizione di parallelismo si traduce nell'equazione

$$(1) \quad \boxed{\mathbf{a} \times \mathbf{r} = 0}$$

che è pertanto l'equazione caratteristica dei moti centrali. Un'altra caratterizzazione dei moti centrali è la seguente.

PROPOSIZIONE 1. Un moto è centrale se e solo se la velocità areale (vettoriale) è costante:

$$(2) \quad \boxed{v_{ar} = \text{cost.}}$$

Dimostrazione. Per la derivazione di un prodotto vettoriale vale la regola di Leibniz, quindi:

$$\frac{d}{dt}(\mathbf{r} \times \mathbf{v}) = \frac{d\mathbf{r}}{dt} \times \mathbf{v} + \mathbf{r} \times \frac{d\mathbf{v}}{dt} = \mathbf{v} \times \mathbf{v} + \mathbf{r} \times \mathbf{a} = \mathbf{r} \times \mathbf{a},$$

per cui si ha $2\mathbf{v}_{ar} = \mathbf{r} \times \mathbf{v} = \text{cost.}$ se e solo se $\mathbf{r} \times \mathbf{a} = 0$. ■

È notevole il fatto che

PROPOSIZIONE 2. Un moto centrale è un moto piano.

Dimostrazione. Da $OP \times \mathbf{v} = 2\mathbf{v}_{ar} = \text{cost.}$ segue che il punto $P(t)$ si trova sempre sul piano per O ortogonale al vettore costante $\mathbf{v}_{ar}$. ■

In particolare si ha un moto rettilineo se e solo se la velocità areale è nulla. Infatti la condizione $OP \times \mathbf{v} = 0$ equivale al parallelismo tra il vettore posizione OP e la velocità $\mathbf{v}$, e il punto si muove su di una retta passante per O . Chiamiamo un moto di questo tipo **moto centrale degenero**.

Nello studio di un moto centrale è conveniente utilizzare coordinate polari (r, ϑ) sul piano del moto, aventi polo coincidente con il centro del moto O . Si denota con $\mathbf{k}$ il versore dell'angolo ϑ , ortogonale a tale piano. Dalla formula (15) del §1 e dal Teorema 1, ovvero anche dalla (17) del §1 (osservato che i moti centrali sono caratterizzati dalla condizione $\mathbf{a}_{trasv} = 0$), segue che

COROLLARIO 1. In un moto centrale la quantità

$$(3) \quad \boxed{c = r^2 \dot{\vartheta}}$$

è costante (è detta **costante delle aree**).

OSSERVAZIONE 1. Si vede dalla (3) che in un moto centrale con $c \neq 0$ le funzioni $r(t)$ e $\dot{\vartheta}(t)$ non si annullano mai. Dunque il punto non passa mai per il centro del moto e per un osservatore posto nel centro O il punto P ruota sempre nella stessa direzione, con velocità angolare tanto più piccola quanto più il punto P è distante da O .

OSSERVAZIONE 2. Se di un moto centrale si conoscono la costante delle aree $c \neq 0$, l'equazione polare $r = r(\vartheta)$ della traiettoria e l'anomalia

ϑ_0 all'istante $t = 0$, allora il moto è completamente determinato. Basta infatti calcolare l'integrale (si veda la (3))

$$(4) \quad t(\vartheta) = \frac{1}{c} \int_{\vartheta_0}^{\vartheta} r^2(x) dx.$$

Esso fornisce una funzione monotona, quindi invertibile in una funzione $\vartheta = \vartheta(t)$ che insieme all'equazione della traiettoria $r = r(\vartheta)$ definisce completamente il moto.

PROPOSIZIONE 3. In un moto centrale non degenere, nota la costante delle aree $c \neq 0$ e l'equazione della traiettoria $r = r(\vartheta)$, risultano determinate in ogni punto di questa la velocità e l'accelerazione. Queste sono date dalle formule

$$(5) \quad \boxed{v = c \left(-\frac{d}{d\vartheta} \frac{1}{r} \mathbf{u} + \frac{1}{r} \boldsymbol{\tau} \right), \quad \mathbf{a} = -\frac{c^2}{r^2} \left(\frac{1}{r} + \frac{d^2}{d\vartheta^2} \frac{1}{r} \right) \mathbf{u}}$$

Questo teorema mostra che in un moto centrale la velocità e l'accelerazione del punto possono dedursi da dati puramente geometrici (equazione dell'orbita) e dalla costante delle aree. La seconda delle (5) è chiamata **formula di Binet** (Jacques Binet, 1786-1856), per quanto fosse già nota a Isaac Newton (1642-1727); come subito si vedrà, per suo tramite è possibile dedurre la legge di gravitazione universale.

Dimostrazione. Essendo

$$\dot{\vartheta} = \frac{c}{r^2},$$

risulta:

$$\begin{aligned} \dot{r} &= \frac{dr}{d\vartheta} \dot{\vartheta} = c \frac{1}{r^2} \frac{dr}{d\vartheta} = -c \frac{d}{d\vartheta} \frac{1}{r}, \\ \ddot{r} &= \frac{d\dot{r}}{dt} = \frac{d\dot{r}}{d\vartheta} \dot{\vartheta} = -c \dot{\vartheta} \frac{d^2}{d\vartheta^2} \frac{1}{r} = -\frac{c^2}{r^2} \frac{d^2}{d\vartheta^2} \frac{1}{r}. \end{aligned}$$

Si sostituiscono allora queste espressioni di $\dot{\vartheta}$, $\dot{r}$ e $\ddot{r}$ nelle rappresentazioni (13) del §1 della velocità e dell'accelerazione in coordinate polari, tenendo conto, che per definizione di moto centrale, l'accelerazione trasversa si annulla identicamente. ■

Una delle più notevoli applicazioni della teoria dei moti centrali si ha nella deduzione, operata da Newton, della legge di gravitazione universale a partire dalle **leggi di Keplero** (Johann Kepler, 1571-1630), di solito enunciate come segue:

- I. Le orbite dei pianeti sono ellittiche e il Sole occupa uno dei fuochi.
- II. Le aree descritte dal raggio vettore che va dal Sole ad un pianeta, sono proporzionali ai tempi impiegati a descriverle.
- III. I quadrati dei tempi impiegati dai vari pianeti a percorrere le loro orbite sono proporzionali ai cubi dei semiassi maggiori.

Queste leggi furono dedotte da una enorme massa di dati astronomici raccolti dal maestro di Kepler, Tycho Brahe (1546-1601). I pianeti vengono rappresentati da punti mobili nel riferimento centrato nel Sole, che è un punto fisso, con orientamento invariabile rispetto alle stelle fisse.

La prima legge mostra innanzitutto che il moto di un pianeta è piano. La seconda legge afferma che la velocità areale rispetto al Sole è costante. Di qui segue che il moto di ogni pianeta è centrale di centro il Sole.

La prima legge afferma che l'orbita di ogni pianeta è una ellisse. Sappiamo che l'equazione dell'ellisse in coordinate polari centrate in uno dei fuochi è

$$(6) \quad r = \frac{p}{1 + e \cos \vartheta},$$

dove, detti (a, b) i semiassi maggiore e minore rispettivamente,

$$(7) \quad p = \frac{b^2}{a}$$

è il cosiddetto **parametro** e

$$(8) \quad e = \sqrt{1 - \frac{b^2}{a^2}} \leq 1$$

è l'**eccentricità**. Essendo il moto centrale, per il Teorema 3 sappiamo che è possibile, applicando la formula di Binet, determinare l'accelerazione a partire dall'orbita. Sostituendo allora l'espressione

$$\frac{1}{r} = \frac{1}{p} (1 + e \cos \vartheta)$$

nella seconda delle (5) si ottiene semplicemente

$$(9) \quad \mathbf{a} = -\gamma \frac{1}{r^2} \mathbf{u},$$

posto

$$(10) \quad \gamma = \frac{c^2}{p}.$$

È così dimostrato che l'accelerazione di ogni pianeta è diretta e orientata verso il Sole ed è inversamente proporzionale al quadrato della distanza dal Sole. Tutto questo segue dalle prime due leggi di Keplero.

La terza legge, denotato con T il periodo di rivoluzione del pianeta, afferma che la quantità

$$\frac{a^3}{T^2}$$

non dipende dal pianeta. È quindi un invariante del sistema solare. Da questa invarianza segue l'universalità della legge (9), cioè che la costante γ che vi compare non dipende dal pianeta. Infatti, per il significato di velocità areale, l'area $A = \pi ab$ dell'orbita ellittica di un pianeta è

$$A = \frac{1}{2} \int_0^T r^2 \dot{\vartheta} dt = \frac{c}{2} \int_0^T dt = \frac{1}{2} c T.$$

Quindi dall'uguaglianza

$$\pi ab = \frac{1}{2} c T,$$

per il coefficiente γ si trae l'espressione

$$\gamma = \frac{c^2}{p} = \frac{4\pi^2 a^2 b^2}{T^2} \frac{a}{b^2} = 4\pi^2 \frac{a^3}{T^2},$$

che per la terza legge è un numero indipendente dal pianeta.

3. Moti geodetici

Sia data una superficie regolare Q nello spazio euclideo tridimensionale, rappresentata localmente dall'equazione vettoriale $OP = \mathbf{r}(q^1, q^2)$.

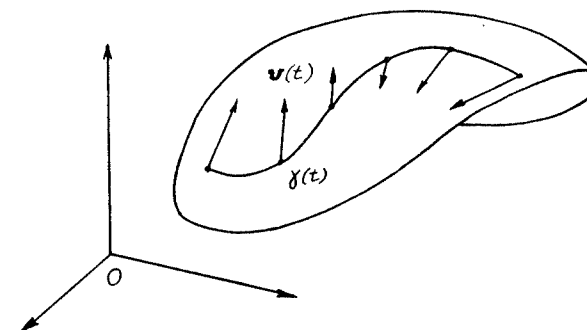


Fig. 3.1: campo di vettori su di una curva, tangenti ad una superficie.

Su questa sia data una curva $\gamma(t)$ di equazioni parametriche $q^i = q^i(t)$, ($i = 1, 2$). Sia inoltre assegnata una funzione vettoriale $\mathbf{v}(t)$ che ad ogni valore del parametro t associa un vettore tangente alla superficie nel punto $\gamma(t)$ (Fig. 3.1).

Il vettore derivato $D\mathbf{v}(t)$, che denotiamo anche con

$$\frac{d\mathbf{v}}{dt},$$

non è più in generale un vettore tangente alla superficie. Rappresentato infatti il vettore $\mathbf{v}$ secondo i vettori $(\mathbf{E}_i)$ associati alle coordinate (q^i) , posto cioè

$$\mathbf{v} = v^i \mathbf{E}_i,$$

tenuto conto delle formule del §12, Cap. II, e del fatto che i vettori $(\mathbf{E}_i)$ dipendono dal parametro t tramite le coordinate superficiali (q^i) , si ha successivamente:

$$(1) \quad \begin{aligned} \frac{d\mathbf{v}}{dt} &= \frac{dv^i}{dt} \mathbf{E}_i + v^i \frac{d\mathbf{E}_i}{dt} \\ &= \frac{dv^i}{dt} \mathbf{E}_i + v^i \partial_j \mathbf{E}_i \frac{dq^j}{dt} \\ &= \left(\frac{dv^k}{dt} + \Gamma_{ji}^k \frac{dq^j}{dt} v^i \right) \mathbf{E}_k + B_{ji} \frac{dq^j}{dt} v^i \mathbf{N}. \end{aligned}$$

La derivata di $\mathbf{v}(t)$ risulta pertanto decomposta nella somma di un vettore ortogonale alla superficie e di un vettore tangente,

$$(2) \quad D_{(i)}\mathbf{v} = \left(\frac{dv^k}{dt} + \Gamma_{ji}^k \frac{dq^j}{dt} v^i \right) \mathbf{E}_k$$

a cui diamo il nome di **derivata intrinseca** o **derivata interna** del vettore tangente alla superficie $\mathbf{v}(t)$. In quest'espressione intervengono infatti solo elementi della geometria interna della superficie (le componenti v^i di $\mathbf{v}$, le componenti $\frac{dq^i}{dt}$ del vettore tangente alla curva, i simboli di Christoffel, che dipendono solo dalla prima forma fondamentale). Ciò significa che le componenti della derivata intrinseca, cioè le funzioni

$$(3) \quad \frac{dv^k}{dt} + \Gamma_{ji}^k \frac{dq^j}{dt} v^i,$$

sono invarianti rispetto a trasformazioni della superficie che conservano la prima forma fondamentale, cioè a "flessioni" della superficie, pensata inestendibile (come fatta di uno strato sottile di carta).

È spontaneo considerare il caso particolare, ma fondamentale, in cui il vettore $\mathbf{v}(t)$ coincide proprio con il vettore $\dot{\gamma}(t)$ tangente alla curva (Fig. 3.2). In tal caso la derivata del vettore $\mathbf{v}$ coincide con l'accelerazione $\mathbf{a}$ del moto rappresentato da γ ed essendo

$$v^i = \frac{dq^i}{dt},$$

la (1) diventa

$$(4) \quad \mathbf{a} = \frac{d\mathbf{v}}{dt} = \left(\frac{d^2 q^k}{dt^2} + \Gamma_{ij}^k \frac{dq^i}{dt} \frac{dq^j}{dt} \right) \mathbf{E}_k + B_{ij} \frac{dq^i}{dt} \frac{dq^j}{dt} \mathbf{N}.$$

Chiamiamo allora **accelerazione intrinseca** o **interna** del moto $\gamma(t)$ la sola parte tangente di questo vettore, cioè la derivata intrinseca del vettore velocità:

$$(5) \quad \mathbf{a}_{(i)} = D_{(i)}\mathbf{v}$$

Le sue componenti sono date da

$$(6) \quad a_{(i)}^k = \frac{d^2 q^k}{dt^2} + \Gamma_{ij}^k \frac{dq^i}{dt} \frac{dq^j}{dt}$$

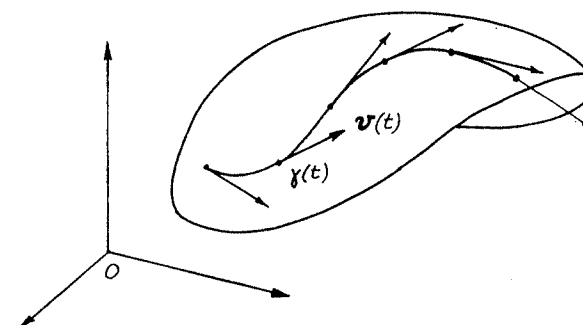


Fig. 3.2: campo delle velocità.

È importante osservare l'analogia tra questa formula e quella delle componenti dell'accelerazione di un punto in uno spazio affine riferito a coordinate generiche. Incontreremo ancora formule di questo tipo, nell'analisi delle equazioni dinamiche di un punto su di una superficie liscia e, più in generale, delle equazioni dinamiche di Lagrange.

La (4) mette anche in evidenza che la componente normale alla superficie dell'accelerazione fa intervenire solo la seconda forma fondamentale $\mathbf{B}$ della superficie, valutata sul vettore velocità. Posto allora

$$(7) \quad \mathbf{a}_{(N)} = \mathbf{B}(\mathbf{v}, \mathbf{v})\mathbf{N},$$

la (4) diventa

$$(8) \quad \mathbf{a} = \mathbf{a}_{(i)} + \mathbf{a}_{(N)}$$

Vediamo ora un'importante sviluppo dei concetti sopra introdotti.

Nello spazio affine euclideo tridimensionale un *moto rettilineo uniforme* o *inerziale* è per definizione un moto la cui velocità vettoriale istantanea $\mathbf{v}$ è costante, ovvero un moto ad accelerazione nulla: $\mathbf{a} = 0$ (Def. 2, §1). La traiettoria è una retta. Il concetto di moto inerziale si può estendere al caso di un punto mobile sopra una superficie alla maniera seguente:

DEFINIZIONE 1. Il moto di un punto su di una superficie si dice **moto geodetico** o **moto inerziale** se l'accelerazione intrinseca è identicamente nulla, $\mathbf{a}_{(i)} = 0$, ovvero (vista la (8)) se l'accelerazione è sempre ortogonale alla superficie, $\mathbf{a} = \mathbf{a}_{(N)}$. Chiamiamo **geodetiche della superficie** le orbite dei moti geodetici.

Vista la (6), i moti geodetici sono caratterizzati dalle equazioni

$$(9) \quad \boxed{\frac{d^2 q^k}{dt^2} + \Gamma_{ij}^k \frac{dq^i}{dt} \frac{dq^j}{dt} = 0}$$

dette **equazioni delle geodetiche**.

È importante osservare che queste equazioni equivalgono al sistema di quattro equazioni differenziali

$$(10) \quad \begin{cases} \frac{dq^i}{dt} = v^i, \\ \frac{dv^k}{dt} = -\Gamma_{ij}^k v^i v^j, \end{cases}$$

nelle quattro funzioni incognite $(q^i(t), v^k(t))$. Ne consegue che i moti geodetici sono le curve integrali di un sistema dinamico $\mathbf{X}$ sopra lo spazio TQ dei vettori tangenti alla superficie Q . Questo spazio può essere interpretato come superficie di dimensione 4 immersa nello spazio TA a 6 dimensioni dei vettori applicati di tutto lo spazio affine euclideo tridimensionale $\mathcal{A}$, oppure come varietà differenziabile di dimensione 4 (come meglio si vedrà in un successivo capitolo), rappresentata dalle coordinate (q^i, v^j) ; le prime due sono coordinate sulla superficie Q , le seconde sono le componenti dei vettori tangenti secondo queste coordinate. Al campo $\mathbf{X}$ si dà il nome (con abuso di linguaggio) di **flusso geodetico** della superficie. A questo campo vettoriale possiamo applicare tutte le considerazioni fatte per i sistemi dinamici in generale. In primo luogo, in base al teorema di Cauchy, possiamo affermare che:

PROPOSIZIONE 1. Assegnato un punto P_0 della superficie Q e un vettore $\mathbf{v}_0$ tangente a questa in P_0 , esiste ed è unica una curva geodetica massimale (cioè una soluzione del sistema (10) definita su di un intervallo temporale massimale) basata in P_0 e avente come vettore tangente per $t = 0$ il vettore $\mathbf{v}_0$.

In secondo luogo possiamo considerare **integrali primi delle geodetiche**, cioè integrali primi del sistema dinamico (10). Questi sono, per

definizione, funzioni sullo spazio TQ a valori reali (o anche funzioni a valori vettoriali), costanti lungo le soluzioni del sistema (10), cioè lungo le geodetiche. Si tratta, nel nostro caso, di funzioni del tipo $F(\mathbf{r}, \mathbf{v})$, cioè dipendenti dal vettore posizione $\mathbf{r}$ e dalla velocità $\mathbf{v}$, le quali, subordinatamente alla condizione che il moto rappresentato da una funzione $\mathbf{r}(t)$ stia sulla superficie e che la corrispondente accelerazione intrinseca sia nulla, soddisfano alla condizione

$$(11) \quad \boxed{\frac{d}{dt} F(\mathbf{r}(t), \mathbf{v}(t)) = 0}$$

Ad esempio (esempio notevole) abbiamo

$$\frac{d}{dt} v^2 = 2 \mathbf{v} \cdot \mathbf{a} = 0$$

perché $\mathbf{v}$ è tangente alla superficie mentre $\mathbf{a}$, riducendosi al solo termine $\mathbf{a}_{(N)}$ (vedi la (8)), è ortogonale a questa. Pertanto la semplice funzione v^2 , ovvero $|\mathbf{v}|$, è un integrale primo. Di qui segue che

PROPOSIZIONE 2. I moti geodetici sono moti uniformi.

Possiamo riesaminare questa circostanza da un altro punto di vista, considerando la rappresentazione intrinseca dell'accelerazione:

$$\mathbf{a} = \ddot{s} \mathbf{t} + c \dot{s}^2 \mathbf{n}.$$

Siccome il versore $\mathbf{t}$ tangente alla curva è anche tangente alla superficie, la condizione $\mathbf{a} = \mathbf{a}_{(N)}$, caratteristica dei moti geodetici, equivale alle due condizioni

$$(12) \quad \ddot{s} = 0, \quad c \dot{s}^2 \mathbf{n} = \mathbf{B}(\mathbf{v}, \mathbf{v}) \mathbf{N}.$$

La prima di queste afferma che il moto è uniforme ($|\mathbf{v}| = \dot{s} = \text{cost.}$). La seconda mostra in particolare che il versore normale principale alla curva è parallelo (cioè uguale o opposto) al versore normale alla superficie. Osservato che l'uniformità del moto è un fatto puramente cinematico, mentre la condizione $\mathbf{n} = \pm \mathbf{N}$ è puramente geometrica, si deduce la seguente proprietà caratteristica delle geodetiche, proprietà che può essere assunta come definizione:

PROPOSIZIONE 3. Una geodetica è una curva tale che in ogni suo punto il versore normale principale è ortogonale alla superficie.

OSSERVAZIONE 1. *Interpretazione variazionale delle geodetiche.* Le geodetiche di una superficie sono caratterizzate da una notevole proprietà geometrica la cui trattazione rientra nell'ambito del *calcolo delle variazioni*. Non disponendo per ora di questa teoria, possiamo ricorrere ad una affermazione di carattere intuitivo: *le geodetiche sono le curve a lunghezza stazionaria*. Si vuol dire che se si fissano due punti qualunque della superficie, tra tutte le curve tracciate su questa e aventi i due punti come estremi, le geodetiche sono quelle di lunghezza stazionaria (in particolare, minima) rispetto a curve *prossime*. Gli esempi seguenti chiariranno questo concetto (si osserverà anche che di geodetiche congiungenti due punti prefissati ne possono esistere più di una, anche infinite). Una semplice interpretazione fisica di questa definizione è la seguente: se un filo flessibile e inestendibile (o anche elastico) è teso su di una superficie liscia, esso traccia su di questa una geodetica.

OSSERVAZIONE 2. *Interpretazione dinamica delle geodetiche.* Anticipiamo alcune considerazioni di carattere dinamico che saranno riprese nel capitolo seguente. Si consideri un punto mobile sopra una superficie fissa. L'equazione dinamica del moto è $m\mathbf{a} = \mathbf{F}_a + \mathbf{F}_r$, dove m è la massa, $\mathbf{F}_a$ è il vettore rappresentante la *forza attiva* agente sul punto ed $\mathbf{F}_r$ è la *forza reattiva* o *reazione vincolare* esercitata dalla superficie sul punto. Si dice che la superficie è *liscia* se la reazione vincolare è sempre ortogonale a questa. Si dice inoltre che il punto si muove di **moto spontaneo** se la forza attiva è nulla, cioè se il punto è soggetto alla sola reazione vincolare. In questo caso l'equazione del moto diventa semplicemente $m\mathbf{a} = \mathbf{F}_r$. Se la superficie è liscia, segue da questa equazione che l'accelerazione è sempre ortogonale alla superficie, cioè che la sua accelerazione intrinseca è sempre nulla. Pertanto: *il moto spontaneo di un punto su di una superficie liscia è un moto geodetico*.

ESEMPIO 1. *Superfici sviluppabili.* Sono superfici la cui curvatura totale è nulla. La loro geometria interna è, localmente, quella del piano euclideo. Consideriamo, per esempio, un *cilindro* (nel senso ordinario del termine, cioè un cilindro, retto a sezione circolare, indefinito). Se lo si taglia lungo una sua direttrice, esso risulta sviluppabile in una striscia del piano, compresa tra due rette parallele. Un punto di una delle due rette viene identificato con il punto ortogonalmente opposto sull'altra retta.

Poiché le geodetiche del piano sono le rette, le geodetiche del cilindro sono rappresentate, nello sviluppo, o da rette parallele alle rette limite della striscia, o da segmenti di rette trasversi alla striscia fra loro paralleli ed equidistanti. Pertanto, sul cilindro, esse danno luogo a delle eliche (Fig. 3.3), eventualmente degeneri, cioè direttrici o circonferenze. Due punti distinti, non giacenti su di una circonferenza, sono congiungibili da infinite geodetiche.

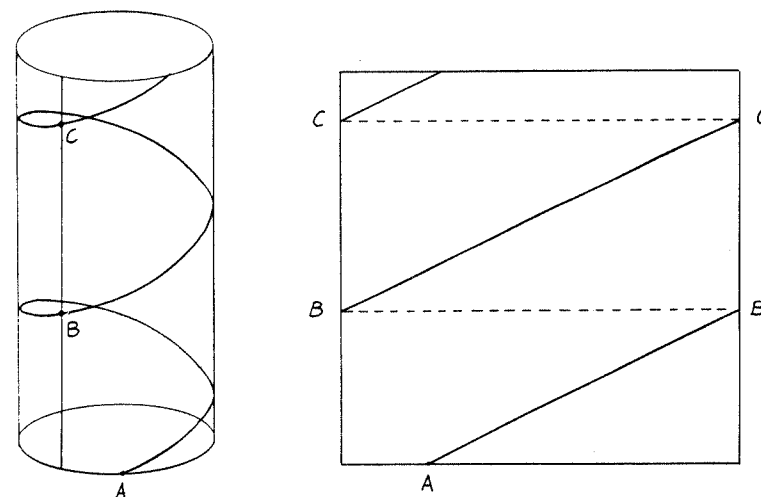


Fig. 3.3: le geodetiche del cilindro.

Un cono è invece sviluppabile in una parte di piano compresa fra due semirette, a e a' , uscenti da un punto V (vertice). Un punto $P \in a$ è identificato col punto $P' \in a'$ equidistante da V . Un caso particolare di geodetica è fornito dalle semirette uscenti da V : sono le direttrici del cono. Se vogliamo invece altre geodetiche, cominciamo col tracciare una semiretta a partire da un punto $A \in a$, interna allo sviluppo del cono. Preso il punto $A' \in a'$ corrispondente all'estremo $A \in a$, proseguiamo la geodetica con una semiretta che parte da A' con lo stesso angolo di incidenza che la semiretta precedente ha in A , ma in senso opposto rispetto al vertice V , come indicato in Fig. 3.4. Analogamente si procede sulle eventuali altre intersezioni di queste semirette con le semirette limite.

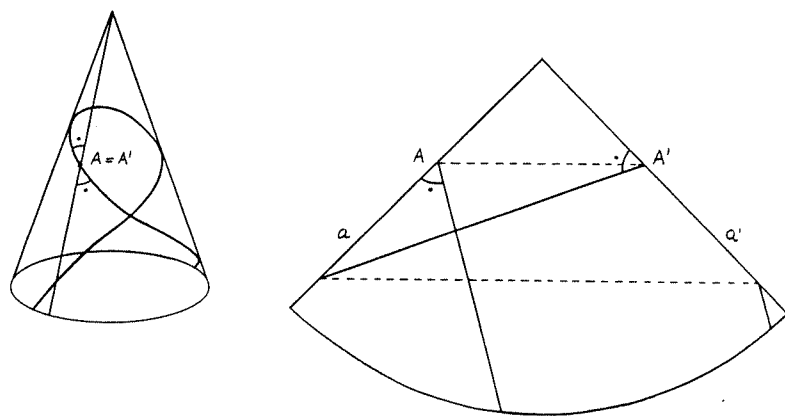


Fig. 3.4: le geodetiche del cono.

ESERCIZIO 1. Quante intersezioni avrà una geodetica del cono (non direttrice) con se stessa? Si studi come varia questo numero al variare dell'angolo di apertura del cono. Si applichi il risultato ottenuto per risolvere il seguente problema di statica: ad un anello di filo flessibile e sottile viene attaccato un peso puntiforme, l'anello viene quindi posto su di un cono circolare retto liscio, ad asse verticale. Per quali aperture del cono l'anello si sfilà?

ESEMPIO 2. *La sfera.* Sia O il centro di una sfera e sia $\mathbf{r} = OP$ il vettore posizione di un generico punto P di questa. Osserviamo subito che la funzione (vettoriale) $\mathbf{r} \times \mathbf{v}$ è un integrale primo delle geodetiche. Infatti, derivando questo prodotto e imponendo la condizione $\mathbf{a} = \mathbf{a}_{(N)}$, caratteristica delle geodetiche, si trova

$$\frac{d}{dt}(\mathbf{r} \times \mathbf{v}) = \mathbf{v} \times \mathbf{v} + \mathbf{r} \times \mathbf{a} = \mathbf{r} \times \mathbf{a}_{(N)} = 0,$$

perché $\mathbf{a}_{(N)}$ è parallelo ad N ed i vettori $(\mathbf{r}, N)$ sono paralleli. Allora, lungo una geodetica della sfera, il vettore $\mathbf{K} = \mathbf{r} \times \mathbf{v}$ è costante (si noti che tale vettore non può essere nullo). Questo implica che il vettore posizione $\mathbf{r}$ è sempre ortogonale ad un vettore costante $\mathbf{K}$. Dunque il

vettore $\mathbf{r}$ descrive necessariamente una circonferenza di raggio massimo, intersezione della sfera col piano passante per il suo centro O e ortogonale a $\mathbf{K}$. La stessa proprietà può dimostrarsi, in via più breve, utilizzando quanto già noto per i moti centrali. Infatti la condizione $\mathbf{a} = \mathbf{a}_{(N)}$ sulla sfera implica che l'accelerazione è sempre diretta verso il suo centro: i moti geodetici della sfera sono dunque moti centrali di centro il centro della sfera. Siccome i moti centrali sono piani e contengono il centro del moto, si conclude che le orbite sono circonferenze massime.

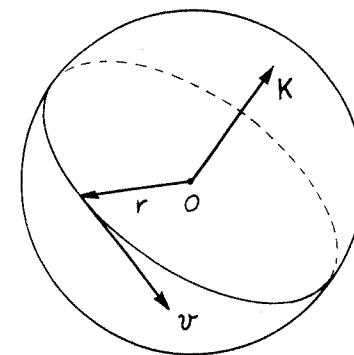


Fig. 3.5: le geodetiche della sfera.

ESEMPIO 3. *Superfici di rotazione.* Si consideri una superficie regolare di rotazione intorno ad un asse di versore $\mathbf{u}$. Sia O un punto qualunque di questo asse e sia $\mathbf{r} = OP$ il vettore di posizione di un generico punto P della superficie. Ebbene: la funzione $\mathbf{u} \times \mathbf{r} \cdot \mathbf{v}$ è un integrale primo delle geodetiche. Infatti, derivando questo prodotto misto, e imponendo la condizione $\mathbf{a} = \mathbf{a}_{(N)}$, caratteristica delle geodetiche, si ha successivamente:

$$\frac{d}{dt}(\mathbf{u} \times \mathbf{r} \cdot \mathbf{v}) = \mathbf{u} \times \mathbf{v} \cdot \mathbf{v} + \mathbf{u} \times \mathbf{r} \cdot \mathbf{a} = \mathbf{u} \times \mathbf{r} \cdot \mathbf{a}_{(N)} = 0,$$

perché $\mathbf{a}_{(N)}$ è parallelo a N ed i vettori $(\mathbf{u}, \mathbf{r}, N)$ sono complanari. Osserviamo ora che il vettore $\mathbf{u} \times \mathbf{r}$ è tangente al parallelo passante per P e

ha modulo pari alla distanza ρ del punto P dall'asse di rotazione. Sicché, detto ϑ l'angolo formato tra questo e il vettore velocità $\mathbf{v}$, ed osservato che quest'ultimo per i moti geodetici ha modulo costante, dall'integrale primo delle geodetiche $\mathbf{u} \times \mathbf{r} \cdot \mathbf{v} = \text{cost.}$ deduciamo che: *per ogni curva geodetica su di una superficie di rotazione si ha*

$$\rho \cos \vartheta = \text{cost.},$$

dove ρ è la distanza dall'asse di rotazione e ϑ è l'angolo formato con il parallelo. Questo è il **teorema di Clairaut** (Alexis Claude Clairaut, 1713-1765).

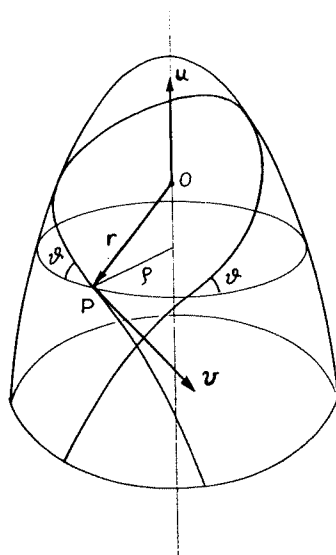


Fig. 3.6: le geodetiche di una superficie di rotazione.

ESERCIZIO 2. Si consideri il teorema di Clairaut nei casi sopra considerati (sono tutte superfici di rotazione). Nel caso del cilindro si riconosce immediatamente che le geodetiche sono delle eliche, perché $\vartheta = \text{cost.}$ Nel caso del cono si vede che una geodetica, allontanandosi dal vertice, tende asintoticamente ad essere perpendicolare ai paralleli, perché crescendo ρ , $\cos \vartheta$ deve tendere a zero.

4. Cinematica del corpo rigido

Un **corpo rigido** è un insieme di punti mobili che durante il moto mantengono inalterate le mutue distanze. Nell'analisi cinematica di un corpo rigido si prescinde inizialmente dalla sua *forma*, in particolare dal fatto che esso sia costituito da un numero finito o infinito di punti. Il moto di un corpo rigido prende il nome di **moto rigido**. Un corpo rigido può essere rappresentato, in maniera essenziale, da una quaterna di punti (Q, P_1, P_2, P_3) , vertici di un tetraedro non degenere, le cui distanze sono costanti nel tempo. Noto il moto di questo tetraedro, è ben determinato il moto di ogni altro punto ad esso rigidamente collegato (cioè avente da questo distanza costante), detto **punto solidale**. La scelta del tetraedro rappresentante il corpo rigido, per quanto in principio del tutto arbitraria, può essere in alcuni casi dettata da ragioni di opportunità (per esempio, nel caso in cui un punto solidale al corpo sia fisso, conviene sceglierlo come uno dei vertici del tetraedro). In genere si sceglie una quaterna di punti tali che i vettori $\mathbf{u}_\alpha = QP_\alpha$, con $\alpha = 1, 2, 3$, siano unitari e fra loro mutuamente ortogonali. Pensati applicati nel punto Q , essi costituiscono un **riferimento canonico solidale**, o, come si usa dire brevemente, una **terna solidale**. Noto il moto del punto Q e la dipendenza dal tempo dei versori $\mathbf{u}_\alpha$, noto cioè il moto di una terna solidale, risulta completamente determinato il moto di un qualunque punto solidale P , avendosi

$$(1) \quad P(t) = Q(t) + x^\alpha \mathbf{u}_\alpha(t),$$

dove le (x^α) sono le coordinate, costanti, del punto P rispetto alla terna solidale. La dipendenza dal tempo dei versori $(\mathbf{u}_\alpha)$ deve essere compatibile con il vincolo di ortonormalità, devono cioè ad ogni istante essere verificate le uguaglianze

$$(2) \quad \mathbf{u}_\alpha \cdot \mathbf{u}_\beta = \delta_{\alpha\beta} \quad (\alpha, \beta = 1, 2, 3).$$

Subordinatamente al vincolo (2), la (1) fornisce la rappresentazione del più generale moto rigido.

È però utile considerare una rappresentazione alternativa che ha il pregio di prescindere dalla scelta di una terna solidale. Essa è basata sulle seguenti definizioni.

DEFINIZIONE 1. Un **endomorfismo affine** su di uno spazio affine $(\mathcal{A}, E, \delta)$ (nelle considerazioni che ora faremo non è necessario restringersi al caso tridimensionale) è un'applicazione $\varphi: \mathcal{A} \rightarrow \mathcal{A}$ tale che esiste

un endomorfismo lineare $Q: E \rightarrow E$, detto **soggiacente** o **associato** a φ , per cui

$$\delta(\varphi(A), \varphi(B)) = Q(\delta(A, B)), \quad \forall A, B \in \mathcal{A},$$

ovvero, con le usuali notazioni,

$$(3) \quad \varphi(B) - \varphi(A) = Q(B - A).$$

Si ha in particolare un **automorfismo affine**, detto anche **trasformazione affine**, se φ è biettiva (per la qual cosa occorre e basta che l'endomorfismo lineare associato sia un automorfismo).

DEFINIZIONE 2. Un **isometria affine** è una trasformazione affine su di uno spazio affine euclideo $(\mathcal{A}, E, \delta, g)$ conservante la distanza dei punti,

$$\|\varphi(A) - \varphi(B)\| = \|A - B\|,$$

quindi tale che l'endomorfismo associato Q è un'isometria, cioè un elemento del gruppo ortogonale su (E, g) . Diciamo che un **isometria affine** è **propria** se Q è una rotazione, cioè se $\det(Q) = 1$.

DEFINIZIONE 3. Un **moto rigido** è una isometria affine propria dipendente da un parametro reale t (il tempo) variabile in un intervallo reale I ,

$$\varphi_t: \mathcal{A} \rightarrow \mathcal{A},$$

a cui corrisponde una rotazione dipendente dal tempo

$$Q_t: E \rightarrow E.$$

La dipendenza dal tempo è supposta di classe C^2 almeno, nel senso che sarà precisato nel seguito.

Ad ogni prefissato punto $P_0 \in \mathcal{A}$ corrisponde il punto mobile

$$(4) \quad P(t) = \varphi_t(P_0).$$

Due punti siffatti mantengono invariata la loro distanza. Infatti:

$$\|P(t) - Q(t)\| = \|\varphi_t(P_0) - \varphi_t(Q_0)\| = \|Q_t(P_0 - Q_0)\| = \|P_0 - Q_0\|.$$

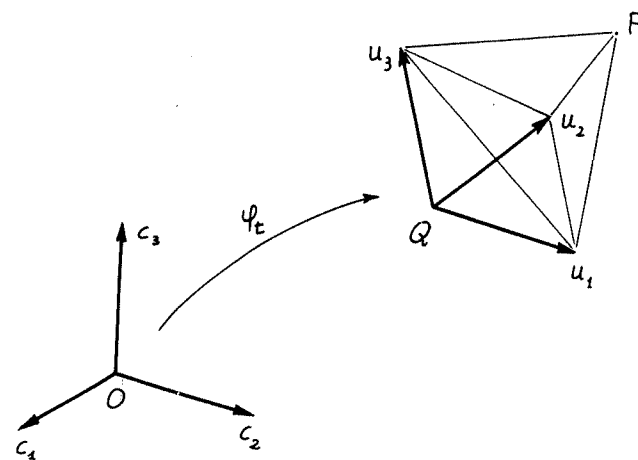


Fig. 4.1: rappresentazione di un moto rigido.

Se in particolare si considera in $\mathcal{A}$ un riferimento cartesiano ortonormale (O, c_α) , cioè una **terna fissa**, risulta di conseguenza definito il moto di una terna solidale (Q, u_α) ponendo

$$(5) \quad Q(t) = \varphi_t(O), \quad u_\alpha(t) = Q_t(c_\alpha).$$

e si ricade così nella rappresentazione (1) (Fig. 4.1).

Con questa generale interpretazione di moto rigido siamo condotti a considerare endomorfismi lineari, in particolare rotazioni, dipendenti da un parametro. Ci limitiamo alle osservazioni seguenti, sufficienti ai nostri scopi.

OSSERVAZIONE 1. *Endomorfismi dipendenti da un parametro.* Su di uno spazio vettoriale E si consideri un endomorfismo lineare dipendente da un parametro reale t variabile in un intervallo reale I ,

$$A_t: E \rightarrow E, \quad t \in I.$$

Denotiamo con

$$\dot{A}_t: E \rightarrow E$$

la derivata rispetto a t di questo endomorfismo. È ancora un endomorfismo definito da

$$(6) \quad \dot{A}_t(v) = \lim_{h \rightarrow 0} \frac{1}{h} (A_{t+h}(v) - A_t(v)).$$

Quest'operazione di derivazione degli endomorfismi lineari dipendenti da un parametro è lineare e per la derivata della composizione di due endomorfismi vale la regola di Leibniz, vale a dire (omettiamo il suffisso t quando questo non è strettamente necessario):

$$(7) \quad (aA + bB)' = a\dot{A} + b\dot{B}, \quad (AB)' = \dot{A}B + A\dot{B}.$$

La dimostrazione della regola di Leibniz è formalmente analoga a quella relativa al prodotto di funzioni. Applicandola per esempio al prodotto $AA^{-1} = 1$ si ottiene l'uguaglianza

$$(8) \quad \dot{A}A^{-1} = -A(A^{-1})'$$

valida per ogni isomorfismo A . Per gli endomorfismi in uno spazio euclideo vale l'uguaglianza

$$(9) \quad (A^T)' = \dot{A}^T.$$

OSSERVAZIONE 2. *Rotazioni dipendenti da un parametro.* Sia (E, g) uno spazio euclideo. Una rotazione $Q_t \in SO(E, g)$ funzione di un parametro t (di classe C^1 almeno) definisce un *moto rigido* nello spazio vettoriale euclideo. Applicata la rotazione Q_t ad un qualunque prefissato vettore $u_0 \in E$, si ottiene un vettore dipendente dal parametro t , $u_t = Q_t(u_0)$. La derivata rispetto a t di questo vettore è data da

$$\dot{u} = \dot{Q}(u_0),$$

quindi, essendo $u_0 = Q^{-1}(u)$, da

$$(10) \quad \boxed{\dot{u} = \Omega(u)}$$

posto

$$(11) \quad \boxed{\Omega = \dot{Q}Q^{-1}}$$

La formula (10) consente di calcolare la derivata del vettore u_t tramite l'endomorfismo dipendente dal tempo Ω_t , chiamato **velocità angolare**. È notevole il fatto che Ω è antisimmetrico. Se lo si interpreta come forma bilineare, si ha infatti successivamente

$$\Omega(u, u) = \Omega(u) \cdot u = \dot{u} \cdot u = 0,$$

perché u ha modulo costante. Un'altra dimostrazione dell'antisimmetria tiene conto delle proprietà (8) e (9) e del fatto che $Q^{-1} = Q^T$:

$$(\dot{Q}Q^{-1})^T = Q\dot{Q}^T = Q(Q^T)' = Q(Q^{-1})' = -\dot{Q}Q^{-1}.$$

Nel caso di uno spazio vettoriale tridimensionale, al tensore antisimmetrico Ω_t corrisponde per aggiunta un vettore

$$(12) \quad \boxed{\omega = *\Omega}$$

detto **vettore velocità angolare**, per cui la (10) diventa

$$(13) \quad \boxed{\dot{u} = \omega \times u}$$

OSSERVAZIONE 3. Nel caso tridimensionale del vettore velocità angolare si può dare un'altra definizione. Si considerino tre versori mutuamente ortogonali $(u_\alpha) = (u_1, u_2, u_3)$ trasportati da Q_t e le loro derivate rispetto a t . Si ponga

$$(14) \quad \boxed{\omega = \frac{1}{2} \sum_{\alpha=1}^3 u_\alpha \times \dot{u}_\alpha}$$

Si ha successivamente:

$$\begin{aligned} \omega \times u_\beta &= \frac{1}{2} \sum_{\alpha=1}^3 (u_\alpha \times \dot{u}_\alpha) \times u_\beta \\ &= \frac{1}{2} \sum_{\alpha=1}^3 (\delta_{\alpha\beta} \dot{u}_\alpha - \dot{u}_\alpha \cdot u_\beta u_\alpha) \\ &= \frac{1}{2} \left(\dot{u}_\beta + \sum_{\alpha=1}^3 \dot{u}_\beta \cdot u_\alpha u_\alpha \right) \\ &= \frac{1}{2} (\dot{u}_\beta + \dot{u}_\beta) = \dot{u}_\beta, \end{aligned}$$

osservato che dalla condizione (2) segue per derivazione

$$\dot{\mathbf{u}}_\beta \cdot \mathbf{u}_\alpha + \mathbf{u}_\beta \cdot \dot{\mathbf{u}}_\alpha = (\mathbf{u}_\beta \cdot \mathbf{u}_\alpha)' = 0,$$

quindi

$$\dot{\mathbf{u}}_\beta \cdot \mathbf{u}_\alpha = -\mathbf{u}_\beta \cdot \dot{\mathbf{u}}_\alpha,$$

e che inoltre, essendo $(\mathbf{u}_\alpha)$ una terna ortonormale, per qualunque vettore $\mathbf{v}$ vale l'uguaglianza

$$\mathbf{v} = \sum_{\alpha=1}^3 \mathbf{v} \cdot \mathbf{u}_\alpha \mathbf{u}_\alpha.$$

Pertanto, definito ω con la (14), risulta

$$(15) \quad \boxed{\dot{\mathbf{u}}_\beta = \omega \times \mathbf{u}_\beta \quad (\beta = 1, 2, 3)}$$

a conferma della (13). Le formule (15) sono anche dette **formule di Poisson** (Siméon Denis Poisson, 1781-1840).

In queste premesse abbiamo raccolto tutti gli elementi per enunciare e dimostrare i seguenti due teoremi fondamentali della cinematica del corpo rigido:

TEOREMA 1. In un moto rigido le velocità $\mathbf{v}_P$ e $\mathbf{v}_Q$ di due qualsiasi punti P e Q solidali sono tali che ad ogni istante

$$(16) \quad \boxed{\mathbf{v}_P \cdot \mathbf{QP} = \mathbf{v}_Q \cdot \mathbf{QP}}$$

Dimostrazione. Per due punti mobili vale l'uguaglianza

$$(\mathbf{QP})' = (\mathbf{P} - \mathbf{Q})' = \mathbf{v}_P - \mathbf{v}_Q.$$

Se si deriva rispetto al tempo la condizione di rigidità $\mathbf{QP} \cdot \mathbf{QP} = \text{cost.}$ si trova la (16). ■

La (16) prende il nome di **prima formula fondamentale** della cinematica del corpo rigido (Fig. 4.2).

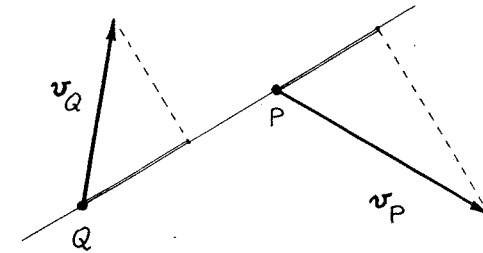


Fig. 4.2: significato geometrico della prima formula fondamentale.

TEOREMA 2. In un moto rigido, ad ogni istante t , esiste ed è unico un vettore $\omega(t)$, detto **velocità angolare**, tale che qualunque siano i punti P e Q solidali al corpo le loro velocità $\mathbf{v}_P$ e $\mathbf{v}_Q$ in quell'istante sono legate dalla relazione

$$(17) \quad \boxed{\mathbf{v}_P = \mathbf{v}_Q + \omega \times \mathbf{QP}}$$

Dimostrazione. Diamo di questo teorema due dimostrazioni. La prima fa capo alla Definizione 3 di moto rigido, la seconda alla rappresentazione (1) e alle formule di Poisson. (I) Il moto dei punti solidali è definito da (si veda la (4))

$$\mathbf{P}(t) = \varphi_t(\mathbf{P}_0), \quad \mathbf{Q}(t) = \varphi_t(\mathbf{Q}_0).$$

Derivando rispetto al tempo otteniamo:

$$\begin{aligned} \mathbf{v}_P - \mathbf{v}_Q &= (\varphi_t(\mathbf{P}_0) - \varphi_t(\mathbf{Q}_0))' \\ &= \dot{\mathbf{Q}}(\mathbf{Q}_0\mathbf{P}_0) \\ &= \dot{\mathbf{Q}}\mathbf{Q}^{-1}(\mathbf{QP}) \\ &= \boldsymbol{\Omega}(\mathbf{QP}) \\ &= \omega \times \mathbf{QP}. \end{aligned}$$

(II) Derivando la (1) rispetto al tempo otteniamo:

$$\begin{aligned} \mathbf{v}_P - \mathbf{v}_Q &= x^\alpha \dot{\mathbf{u}}_\alpha \\ &= x^\alpha \boldsymbol{\omega} \times \mathbf{u}_\alpha \\ &= \boldsymbol{\omega} \times (x^\alpha \mathbf{u}_\alpha) \\ &= \boldsymbol{\omega} \times \mathbf{QP}. \end{aligned}$$

Il fatto che $\boldsymbol{\omega}$ sia unico è conseguenza ovvia della stessa (17): se essa sussistesse per un altro vettore $\boldsymbol{\omega}'$, sottraendo membro a membro le due equazioni si otterrebbe $(\boldsymbol{\omega} - \boldsymbol{\omega}') \times \mathbf{QP} = 0$ per qualunque coppia di punti solidali; quindi, necessariamente, $\boldsymbol{\omega} - \boldsymbol{\omega}' = 0$. ■

La (17) prende il nome di **seconda formula fondamentale** della cinematica del corpo rigido. Si noti che, moltiplicata scalarmente per $\mathbf{QP}$, la seconda formula fondamentale implica la prima.

La seconda formula fondamentale consente di calcolare, in ogni prefissato istante t , la velocità di un qualunque punto solidale P , nota la velocità angolare e la velocità di un solo punto solidale Q in quell'istante. Essa consente quindi di descrivere completamente l'**atto di moto rigido** cioè il campo vettoriale delle velocità dei punti solidali in un istante t . Si osserva che essa è formalmente analoga alla formula di trasposizione dei momenti per un sistema di vettori applicati (§15, Cap. II),

$$\mathbf{M}_P = \mathbf{M}_Q + \mathbf{R} \times \mathbf{QP},$$

dove il risultante $\mathbf{R}$ è l'analogo della velocità angolare $\boldsymbol{\omega}$. Si possono allora adattare all'atto di moto rigido tutte le osservazioni e le proprietà del momento risultante, in particolare quelle riguardanti l'asse centrale. Ne risulta il seguente fondamentale **teorema di Mozzi** (Giulio Mozzi, 1730-1813):

TEOREMA 3. In ogni istante t in cui $\boldsymbol{\omega} \neq 0$ esiste ed è unica una retta a i cui punti hanno velocità parallela a $\boldsymbol{\omega}$, ovvero velocità minima. La retta a , detta **asse di Mozzi** o **asse di moto**, è parallela a $\boldsymbol{\omega}$.

OSSERVAZIONE 3. Le velocità dei punti di una retta parallela a $\boldsymbol{\omega}$, quindi all'asse di Mozzi, sono fra loro uguali. Infatti la (17) implica $\mathbf{v}_P = \mathbf{v}_Q$ quando $\mathbf{QP}$ è parallelo a $\boldsymbol{\omega}$.

OSSERVAZIONE 4. Atti di moto particolari. La (17) mostra che il generico **atto di moto** è **elicoidale** (si veda la Fig. 15.1 del §15,

Cap. II, relativa all'andamento del momento risultante di un sistema di vettori applicati). Se $\boldsymbol{\omega} = 0$ dalla (17) si vede che tutti i punti hanno, nell'istante considerato, la stessa velocità. Si è in presenza di un **atto di moto traslatorio**. Se la velocità dei punti dell'asse di Mozzi è nulla si dice che l'**atto di moto** è **rotatorio**.

OSSERVAZIONE 5. Distribuzione delle accelerazioni. Derivando la (17) rispetto al tempo si trova la formula seguente, che fornisce l'accelerazione istantanea $\mathbf{a}_P$ di un qualunque punto solidale P , nota l'accelerazione $\mathbf{a}_Q$ di un punto prefissato Q , la velocità angolare $\boldsymbol{\omega}$ e la sua derivata $\dot{\boldsymbol{\omega}}$:

$$(18) \quad \begin{aligned} \mathbf{a}_P &= \mathbf{a}_Q + \dot{\boldsymbol{\omega}} \times \mathbf{QP} + \boldsymbol{\omega} \times (\boldsymbol{\omega} \times \mathbf{QP}) \\ &= \mathbf{a}_Q + \dot{\boldsymbol{\omega}} \times \mathbf{QP} + \boldsymbol{\omega} \cdot \mathbf{QP} \boldsymbol{\omega} - |\boldsymbol{\omega}|^2 \mathbf{QP} \end{aligned}$$

5. Cambiamenti di riferimento

I concetti di moto, di velocità, di accelerazione sono **concetti relativi**: la loro definizione presuppone la scelta di un riferimento. Come si è già detto, scegliere un riferimento significa assumere come modello dello spazio fisico, nell'ambito delle teorie classiche, lo spazio affine tridimensionale euclideo. I punti di questo spazio affine si identificano con delle particelle ideali costituenti un corpo rigido invadente tutto l'universo. A ciascuna di queste particelle, immobili, quindi con mutue distanze costanti, attribuiamo un "nome convenzionale", ricorrendo a sistemi di coordinate. Possiamo quindi descrivere il moto di un punto esprimendo in funzione del tempo le coordinate dei punti del riferimento via via occupati dal punto mobile. In genere, per rappresentare i punti di un riferimento, si usano coordinate cartesiane ortonormali, associate ad un riferimento cartesiano $(O, \mathbf{c}_\alpha)$.

Siano dunque: $\mathcal{R}$ un riferimento, (A, E, δ, g) lo spazio affine euclideo tridimensionale corrispondente, $(O, \mathbf{c}_\alpha)$ un riferimento cartesiano ortonormale rappresentante $\mathcal{R}$, che chiamiamo brevemente **terna fissa** (i versori ortogonali $(\mathbf{c}_\alpha)$ si pensano applicati nel punto O). Si consideri un corpo rigido in moto rispetto ad $\mathcal{R}$. Questo corpo rigido costituisce un secondo riferimento $\mathcal{R}'$. Sia $(O', \mathbf{c}_{\alpha'})$ una **terna solidale** a $\mathcal{R}'$. Si

consideri inoltre un vettore $\mathbf{u}(t)$ dipendente dal tempo e per questo la doppia rappresentazione

$$(1) \quad \mathbf{u} = u^\alpha \mathbf{c}_\alpha = u^{\alpha'} \mathbf{c}_{\alpha'}.$$

Denotiamo con $D\mathbf{u}$ la derivata rispetto al tempo del vettore $\mathbf{u}$ relativa al riferimento originario $\mathcal{R}$. Essa è data da

$$(2) \quad D\mathbf{u} = \dot{u}^\alpha \mathbf{c}_\alpha,$$

dove $\dot{u}^\alpha$ denota la derivata della funzione reale $u^\alpha(t)$. Stante la seconda delle (1) si ha anche, tenuto conto delle formule di Poisson,

$$(3) \quad D\mathbf{u} = \dot{u}^{\alpha'} \mathbf{c}_{\alpha'} + u^{\alpha'} \boldsymbol{\omega} \times \mathbf{c}_{\alpha'},$$

dove $\boldsymbol{\omega}$ è la velocità angolare di $\mathcal{R}'$ rispetto a $\mathcal{R}$. D'altra parte, rispetto ad un osservatore posto in $\mathcal{R}'$ i versori ($\mathbf{c}_{\alpha'}$) sono costanti, quindi se denotiamo con $D'\mathbf{u}$ la derivata rispetto al tempo relativa al riferimento $\mathcal{R}'$, si ha

$$(4) \quad D'\mathbf{u} = \dot{u}^{\alpha'} \mathbf{c}_{\alpha'}.$$

Pertanto dal confronto delle (2), (3) e (4) segue l'uguaglianza

$$(5) \quad \boxed{D\mathbf{u} = D'\mathbf{u} + \boldsymbol{\omega} \times \mathbf{u}}$$

che esprime il legame tra le derivate temporali di un vettore $\mathbf{u}$ rispetto ai due riferimenti. Si noti che per vettori solidali a $\mathcal{R}'$, per i quali $D'\mathbf{u} = 0$, si ritrovano le formule di Poisson.

OSSERVAZIONE 1. Dalla (5) segue in particolare

$$(6) \quad \boxed{D\boldsymbol{\omega} = D'\boldsymbol{\omega}}$$

Il vettore velocità angolare ha la stessa derivata in entrambi i riferimenti.

Per stabilire anche la relazione tra le derivate seconde, applichiamo la (5) a se stessa. Otteniamo successivamente:

$$\begin{aligned} D^2\mathbf{u} &= D(D'\mathbf{u} + \boldsymbol{\omega} \times \mathbf{u}) \\ &= D'(D'\mathbf{u} + \boldsymbol{\omega} \times \mathbf{u}) + \boldsymbol{\omega} \times (D'\mathbf{u} + \boldsymbol{\omega} \times \mathbf{u}) \\ &= D'^2\mathbf{u} + D\boldsymbol{\omega} \times \mathbf{u} + 2\boldsymbol{\omega} \times D'\mathbf{u} + \boldsymbol{\omega} \times (\boldsymbol{\omega} \times \mathbf{u}). \end{aligned}$$

Stante la (6) possiamo porre per comodità $\dot{\boldsymbol{\omega}} = D\boldsymbol{\omega} = D'\boldsymbol{\omega}$. Abbiamo allora ottenuto la formula

$$(7) \quad \boxed{D^2\mathbf{u} = D'^2\mathbf{u} + 2\boldsymbol{\omega} \times D'\mathbf{u} + \dot{\boldsymbol{\omega}} \times \mathbf{u} + \boldsymbol{\omega} \times (\boldsymbol{\omega} \times \mathbf{u})}$$

Si consideri il moto di un punto $P(t)$. Esso può essere osservato dal riferimento $\mathcal{R}$ oppure dal riferimento $\mathcal{R}'$ (Fig. 5.1). Tra il vettore posizione $\mathbf{r} = OP$ nel riferimento $\mathcal{R}$ e il vettore posizione $\mathbf{r}' = O'P$ nel riferimento $\mathcal{R}'$ sussiste la relazione

$$(8) \quad \boxed{\mathbf{r} = OO' + \mathbf{r}'}$$

Denotiamo con

$$\mathbf{v} = D\mathbf{r}, \quad \mathbf{v}' = D'\mathbf{r}'$$

le velocità relative ai due riferimenti. Applicando l'operatore D alla (8), vista la (5), si ottiene l'uguaglianza

$$(9) \quad \mathbf{v} = \mathbf{v}_{O'} + \mathbf{v}' + \boldsymbol{\omega} \times \mathbf{r}',$$

dove

$$\mathbf{v}_{O'} = D(OO')$$

è la velocità del punto O' rispetto ad $\mathcal{R}$. In virtù della seconda formula fondamentale della cinematica del corpo rigido il vettore

$$(10) \quad \mathbf{v}_{tr} = \mathbf{v}_{O'} + \boldsymbol{\omega} \times O'P,$$

che compare a secondo membro dell'uguaglianza (9), è la velocità del punto P pensato, nell'istante considerato, solidale al riferimento, ovvero *la velocità del punto solidale a $\mathcal{R}'$ che nell'istante considerato è occupato dal punto mobile $P(t)$* . A questa velocità si dà il nome di **velocità di trascinamento**. Possiamo allora enunciare il seguente

TEOREMA 1. (**teorema dei moti relativi**). La velocità di un punto relativa ad un riferimento $\mathcal{R}$ è la somma della velocità relativa ad un altro riferimento $\mathcal{R}'$ e della velocità di trascinamento, cioè della velocità del punto solidale a $\mathcal{R}'$ istantaneamente occupato dal punto mobile:

$$(11) \quad \boxed{\mathbf{v} = \mathbf{v}' + \mathbf{v}_{tr}}$$

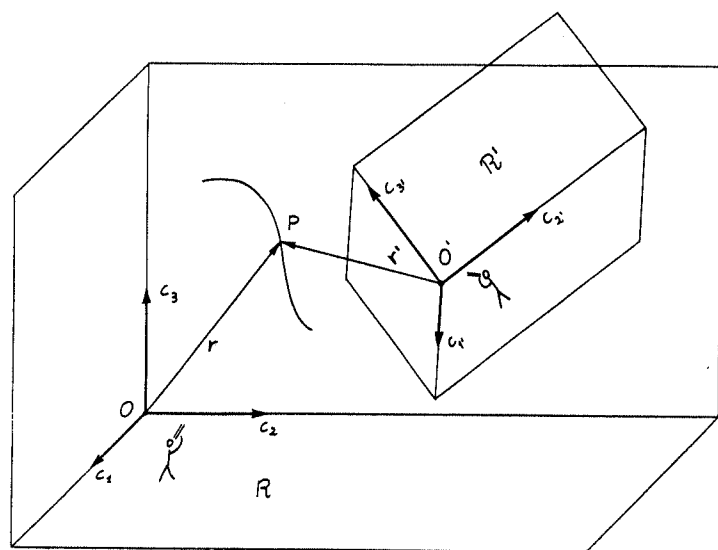


Fig. 5.1: un moto osservato da due riferimenti.

Si consideri ora la derivata seconda dell'uguaglianza (8),

$$D^2 \mathbf{r} = D^2(OO') + D^2 \mathbf{r}'.$$

Tenuto conto della (7) applicata ad $\mathbf{r}'$ si ottiene subito

$$(12) \quad \mathbf{a} = \mathbf{a}_{O'} + \mathbf{a}' + 2\boldsymbol{\omega} \times \mathbf{v}' + \dot{\boldsymbol{\omega}} \times \mathbf{r}' + \boldsymbol{\omega} \times (\boldsymbol{\omega} \times \mathbf{r}'),$$

dove

$$\mathbf{a}_{O'} = D\mathbf{v}_{O'} = D^2(OO')$$

è l'accelerazione del punto O' rispetto a $\mathcal{R}$. Per la formula (18) del § precedente il vettore

$$(13) \quad \mathbf{a}_{tr} = \mathbf{a}_{O'} + \dot{\boldsymbol{\omega}} \times \mathbf{r}' + \boldsymbol{\omega} \times (\boldsymbol{\omega} \times \mathbf{r}')$$

che compare a secondo membro della (12) è l'accelerazione del punto solidale a $\mathcal{R}'$ che nell'istante considerato è occupato dal punto mobile $P(t)$. Ad esso si dà il nome di **accelerazione di trascinamento**. Al vettore

$$(14) \quad \boxed{\mathbf{a}_c = 2\boldsymbol{\omega} \times \mathbf{v}'}$$

che pure compare nella (12), si dà il nome di **accelerazione complementare** o **accelerazione di Coriolis** (Gustave Gaspard Coriolis, 1792-1843). Possiamo allora enunciare il seguente

TEOREMA 2. (teorema di Coriolis). L'accelerazione di un punto relativa ad un riferimento $\mathcal{R}$ è la somma dell'accelerazione relativa ad un altro riferimento $\mathcal{R}'$, dell'accelerazione di trascinamento, cioè dell'accelerazione del punto solidale a $\mathcal{R}'$ istantaneamente occupato dal punto mobile, e dell'accelerazione complementare:

$$(15) \quad \boxed{\mathbf{a} = \mathbf{a}' + \mathbf{a}_{tr} + \mathbf{a}_c}$$

Si noti bene che l'accelerazione complementare $\mathbf{a}_c$ si annulla identicamente quando è sempre nulla la velocità angolare: in tal caso, come vedremo al prossimo paragrafo, il moto di $\mathcal{R}'$ rispetto a $\mathcal{R}$ è *traslatorio*. Se in più è anche $\mathbf{a}_{O'} = 0$, cioè se il moto di $\mathcal{R}'$ è *traslatorio rettilineo uniforme* (si veda il prossimo paragrafo), si annulla anche l'accelerazione di trascinamento e quindi

$$(16) \quad \boxed{\mathbf{a} = \mathbf{a}'}$$

Le accelerazioni nei due riferimenti sono dunque rappresentate dal medesimo vettore. Il riferimento $\mathcal{R}'$ si dice allora **equivalente** a $\mathcal{R}$.

6. Moti rigidi particolari

Esaminiamo in questo paragrafo alcuni particolari tipi di moti rigidi.

DEFINIZIONE 1. Dicesi **moto traslatorio** un moto rigido con velocità angolare identicamente nulla: $\boldsymbol{\omega} = 0$ ad ogni istante t .

Un moto è traslatorio se e solo se ogni vettore solidale al corpo è costante; lo si deduce dalla formula (5) del §5: se $D'u = 0$ si ha l'equivalenza $Du = 0 \iff \omega = 0$.

Un moto è traslatorio se e solo se ad ogni istante le velocità di tutti i punti solidali coincidono (cioè se e solo se in ogni istante l'atto di moto è traslatorio); lo si deduce dalla seconda formula fondamentale: $v_P = v_Q, \forall P, Q$ solidali $\iff \omega = 0$.

Tra i moti traslatori troviamo in particolare il **moto traslatorio rettilineo**, in cui ogni punto solidale si muove di moto rettilineo, e il **moto traslatorio rettilineo uniforme**, in cui ogni punto solidale si muove di moto rettilineo uniforme.

DEFINIZIONE 2. Dicesi **moto rigido con punto fisso** un moto rigido in cui uno dei punti solidali ha sempre velocità nulla.

Se O è punto fisso, quindi $v_O = 0$, la seconda formula fondamentale diventa

$$(1) \quad v_P = \omega \times OP.$$

Ogni punto solidale si muove su di una sfera di centro O .

L'asse di Mozzi passa in ogni istante t per il punto fisso O , perché questo ha sempre velocità minima. Pertanto l'asse di Mozzi è la retta parallela a ω passante per O . I suoi punti hanno tutti velocità istantanea nulla.

L'asse di Mozzi varia in genere da istante ad istante, e descrive quindi una superficie rigata, un cono $\mathcal{C}$ di vertice O , detta **cono fisso**. Nel corpo rigido, inteso come riferimento $\mathcal{R}'$, esso descrive un secondo cono $\mathcal{C}'$, detto **cono solidale** o **cono mobile**. Entrambi i coni prendono il nome di **coni di Poincot** (Louis Poincot, 1777-1859).

TEOREMA 1. In un moto rigido con punto fisso il cono solidale rotola senza strisciare sul cono fisso.

Quest'enunciato, che fornisce una notevole interpretazione geometrica dei moti rigidi con punto fisso, richiede l'intervento della definizione seguente.

DEFINIZIONE 3. Si dice che una superficie regolare rigida $\mathcal{C}'$ **rotola senza strisciare** su di una superficie regolare rigida fissa $\mathcal{C}$ (si dice anche che il moto di $\mathcal{C}'$ su $\mathcal{C}$ è di **puro rotolamento**), se in ogni istante

nei punti di contatto coincidono i piani tangenti alle due superfici e i punti solidali a $\mathcal{C}'$ a contatto con $\mathcal{C}$ hanno velocità nulla.

Dimostrazione. In ogni istante i due coni di Poincot si toccano su una direttrice (l'asse di Mozzi a quell'istante). Consideriamo una curva non parametrizzata sul cono solidale $\mathcal{C}'$ trasversale alle direttrici. Per ogni istante t risulta definito un punto di contatto $P(t)$ che sta sulla curva. Il punto P si muove rispetto al riferimento $\mathcal{R}'$ solidale a $\mathcal{C}'$ con una velocità v' tangente a $\mathcal{C}'$ e non parallela all'asse di Mozzi. Siccome il punto si muove anche sul cono fisso $\mathcal{C}$, la sua velocità v relativa al riferimento fisso, rispetto al quale studiamo il moto del corpo rigido, è tangente a $\mathcal{C}$. Per il teorema dei moti relativi $v = v' + v_{tr}$. Ma in questo caso la velocità di trascinamento è la velocità di un punto solidale che sta sull'asse di Mozzi, quindi nulla. Dunque $v = v'$. Ma entrambi questi vettori sono tangenti ai rispettivi coni e non paralleli all'asse di Mozzi. Dunque i piani tangenti ai due coni coincidono. Inoltre, come già si è osservato, i punti solidali a $\mathcal{C}'$ che sono a contatto con $\mathcal{C}$ si trovano sull'asse di Mozzi e quindi hanno velocità nulla e le condizioni di puro rotolamento richieste dalla Def. 3 sono soddisfatte. ■

DEFINIZIONE 4. Dicesi **moto rigido con asse fisso** un moto rigido in cui esiste una retta solidale di punti fissi.

Tale retta è evidentemente l'asse di Mozzi. Ogni punto del corpo rigido si muove di moto circolare intorno all'asse. Se questo è individuato dal versore k di un riferimento ortonormale (O, i, j, k) , allora la velocità angolare è data da

$$(2) \quad \omega = \dot{\vartheta} k$$

dove ϑ è l'angolo di rotazione, funzione del tempo, di versore k . Per dimostrarlo possiamo per esempio ricorrere alla definizione (14) del §4, considerata una terna solidale (u_α) con $k = u_3$. Per quanto si è già visto in cinematica del punto (formula (10), §1), si ha $\dot{u}_1 = \dot{\vartheta} u_2$ e $\dot{u}_2 = -\dot{\vartheta} u_1$. Per cui, essendo $\dot{u}_3 = 0$, si ha

$$\omega = \frac{1}{2}(u_1 \times \dot{u}_1 + u_2 \times \dot{u}_2) = \frac{1}{2} \dot{\vartheta} (u_1 \times u_2 - u_2 \times u_1) = \dot{\vartheta} u_3 = \dot{\vartheta} k.$$

DEFINIZIONE 5. Dicesi **moto rigido piano** un moto rigido in cui un piano di punti solidali scorre cioè su di un piano fisso.

Un tale moto è dunque, essenzialmente, il moto di un piano rigido Π' sopra un piano fisso Π . La velocità angolare, quindi l'asse di Mozzi, è sempre ortogonale al piano. Infatti se $\mathbf{k}$ è un versore ortogonale a Π' , da $\dot{\mathbf{k}} = \boldsymbol{\omega} \times \mathbf{k} = 0$ segue che $\boldsymbol{\omega}$ è sempre parallelo a $\mathbf{k}$.

Per la velocità angolare di un moto rigido piano vale ancora la formula (2), dove ϑ è l'angolo compreso tra una qualunque retta solidale a Π' e una qualunque retta fissa in Π . La dimostrazione è formalmente la stessa, scelte le terne in maniera tale che $\mathbf{k} = \mathbf{u}_3$ siano ortogonali ai piani.

Il punto intersezione $C(t)$ dell'asse di Mozzi $a(t)$ con il piano fisso Π prende il nome di **centro d'istantanea rotazione**. Il punto solidale a Π' che nell'istante considerato è sovrapposto a C ha velocità nulla. Il punto C descrive su Π una curva $\mathcal{P}$ detta **polare fissa**, e su Π' una curva $\mathcal{P}'$ solidale a questo, detta **polare mobile**. In maniera del tutta analoga al caso del moto rigido con un punto fisso (nel caso presente i coni di Poincot degenerano in cilindri e le loro intersezioni coi piani danno luogo alle polari), si prova che:

TEOREMA 2. In un moto rigido piano la polare mobile ruota senza strisciare sulla polare fissa.

Quindi, note queste due curve, il moto rigido piano è geometricamente determinato. Per determinarlo completamente occorre per esempio conoscere la legge oraria del moto di C .

Per determinare il centro d'istantanea rotazione, e quindi le due polari, si può in molti casi utilizzare il **teorema di Chasles** (Michel Chasles, 1793-1880):

TEOREMA 3. In un moto rigido piano, ad ogni istante t per cui $\boldsymbol{\omega}(t) \neq 0$, il centro d'istantanea rotazione è l'intersezione di due rette condotte per due punti solidali e ortogonali alle rispettive velocità.

Dimostrazione. Siano P e Q due punti solidali al corpo. Se $\boldsymbol{\omega}(t) \neq 0$, l'atto di moto è rotatorio di centro $C(t)$, quindi $\mathbf{v}_P = \boldsymbol{\omega} \times \mathbf{CP}$ e $\mathbf{v}_Q = \boldsymbol{\omega} \times \mathbf{CQ}$, e le rette $\mathbf{CP}$ e $\mathbf{CQ}$ sono rispettivamente ortogonali a $\mathbf{v}_P$ e $\mathbf{v}_Q$. Se $\mathbf{PC}$ è parallela a $\mathbf{QC}$, l'intersezione di cui all'enunciato non esiste ("va all'infinito"). Se così fosse per tutti le coppie di punti solidali, l'atto di moto sarebbe traslatorio, contro l'ipotesi $\boldsymbol{\omega}(t) \neq 0$. ■

ESEMPIO 1. Si consideri il moto di un'asta rigida, cioè di un segmento rigido, di lunghezza l , i cui estremi A e B sono vincolati a scorrere

su due assi incidenti e ortogonali (assi x e y rispettivamente) (Fig. 6.2). Si pensi al moto del piano rigidamente collegato all'asta sul piano fisso (O, x, y) . Per determinare le polari si può applicare il teorema di Chasles. Le velocità dei due estremi sono necessariamente parallele ai rispettivi assi di scorrimento. Quindi il centro di istantanea rotazione C si trova come intersezione delle rette per A e B e ortogonali agli assi. Per un osservatore solidale col piano fisso questo punto ha sempre distanza pari a l dall'origine O : la polare fissa $\mathcal{P}$ è quindi la circonferenza di raggio l e centro l'origine. Per un osservatore solidale all'asta invece, il punto C sottende sempre gli estremi A e B sotto un angolo retto, e quindi la polare mobile è la circonferenza di raggio $\frac{l}{2}$ e centro il punto medio dell'asta. Il moto dell'asta è dunque geometricamente equivalente al moto di rotolamento di una circonferenza di raggio $\frac{l}{2}$ all'interno di una circonferenza fissa di raggio l .

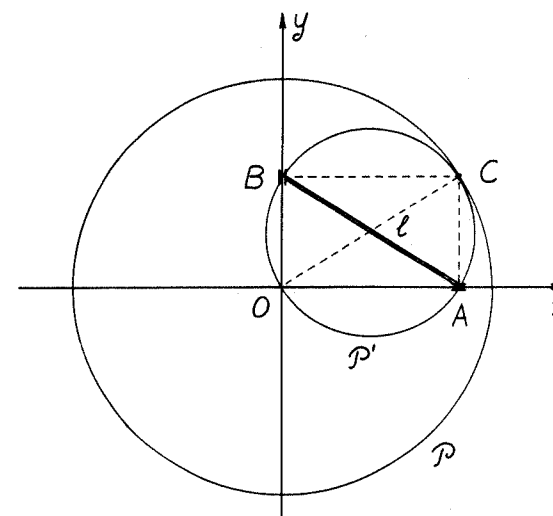


Fig. 6.2

7. Moti rigidi composti

Si consideri ora la composizione di due moti rigidi

$$\varphi_t = \varphi_{2t} \circ \varphi_{1t},$$

e la parte lineare ad essa associata:

$$Q_t = Q_{2t} Q_{1t}.$$

Il corrispondente tensore velocità angolare è

$$(1) \quad \Omega = \Omega_2 + Q_2 \Omega_1 Q_2^{-1}$$

dove Ω_1 e Ω_2 sono le velocità angolari dei moti componenti. Infatti:

$$\begin{aligned} \Omega &= \dot{Q} Q^{-1} \\ &= (\dot{Q}_2 Q_1 + Q_2 \dot{Q}_1) (Q_1^{-1} Q_2^{-1}) \\ &= \dot{Q}_2 Q_2^{-1} + Q_2 \dot{Q}_1 Q_1^{-1} Q_2^{-1} \\ &= \Omega_2 + Q_2 \Omega_1 Q_2^{-1}, \end{aligned}$$

Passando ai vettori aggiunti, risulta la formula

$$(2) \quad \omega = \omega_2 + Q_2(\omega_1)$$

Si osservi infatti che vale la formula generale

$$(3) \quad *Q\Omega Q^{-1} = Q(\omega), \quad \omega = *\Omega,$$

dove Ω è antisimmetrico e Q è una rotazione. Questo perché una rotazione, conservando tutte le grandezze definite tramite il tensore metrico, conserva anche il prodotto vettoriale:

$$(4) \quad Q(u \times v) = Q(u) \times Q(v).$$

Quindi:

$$Q\Omega Q^{-1}(v) = Q(\omega \times Q^{-1}(v)) = Q(\omega) \times v,$$

da cui la (3).

La formula (2) fornisce la *legge di composizione delle velocità angolari* di due moti rigidi.

OSSERVAZIONE 1. La formula (2) può essere reinterpretata come legame tra le velocità angolari di un corpo rigido in due riferimenti diversi. Sia ω la velocità angolare di un corpo rigido $\mathcal{R}''$ relativa ad un riferimento $\mathcal{R}$. Sia invece ω' la sua velocità angolare osservata da un altro riferimento $\mathcal{R}'$. Sia infine ω_{tr} la velocità angolare di $\mathcal{R}'$ rispetto a $\mathcal{R}$, che possiamo chiamare **velocità angolare di trascinamento**. Allora:

$$(5) \quad \omega = \omega' + \omega_{tr}$$

Si ha così il **teorema della somma delle velocità angolari**, analogo al teorema dei moti relativi. Una dimostrazione diretta elementare di questo teorema è la seguente. Per la seconda formula fondamentale di cinematica rigida si ha, per qualunque coppia (P, Q) di punti solidali al corpo $\mathcal{R}''$,

$$\begin{cases} v_P - v_Q = \omega \times QP, \\ v'_P - v'_Q = \omega' \times QP. \end{cases}$$

D'altra parte, per il teorema dei moti relativi,

$$\begin{cases} v_P = v'_P + v_{O'} + \omega_{tr} \times O'P, \\ v_Q = v'_Q + v_{O'} + \omega_{tr} \times O'Q, \end{cases}$$

essendo O' un prefissato punto solidale a $\mathcal{R}'$. Sottraendo membro a membro e tenendo conto delle uguaglianze precedenti si trova

$$\omega \times QP = \omega' \times QP + \omega_{tr} \times QP.$$

Di qui la (5) per l'arbitrarietà di QP .

ESERCIZIO 1. Applicando le considerazioni precedenti, si dimostri la formula

$$(6) \quad \omega = \dot{\psi} \mathbf{k} + \dot{\phi} \mathbf{k}' + \dot{\theta} \mathbf{N}$$

che esprime la velocità angolare tramite le derivate degli angoli di Eulero. Si può applicare la formula (2) tenendo presente quanto detto nell'Esercizio 2, §11.4, Cap. I. Si può anche applicare la formula (5) considerando però, oltre al riferimento fisso $\mathcal{R}$ rispetto al quale si calcolano gli angoli di Eulero, due riferimenti intermedi: il riferimento $\mathcal{R}'$ solidale alla linea dei nodi e al versore $\mathbf{k}$ e il riferimento $\mathcal{R}''$ solidale alla linea dei nodi e al versore $\mathbf{k}'$.

OSSERVAZIONE 2. Un esempio notevole di composizione di moti rigidi è il **moto di precessione regolare**: si tratta di un moto rigido composto di due moti rigidi con punto fisso O , il primo di velocità angolare costante ω_r , detta **velocità di rotazione propria**, il secondo di velocità angolare ω_p , pure costante, detta **velocità di precessione**. Lo si realizza facendo ruotare un corpo rigido intorno ad un suo asse solidale f , detto **asse di figura**, con velocità angolare ω_r , ovviamente parallela a f , e facendo quindi ruotare l'asse f intorno ad un asse fisso p passante per un suo punto O , detto **asse di precessione**, con velocità angolare costante ω_p . La velocità angolare del corpo rigido risulta essere la somma

$$(7) \quad \omega = \omega_p + \omega_r$$

dove, in base alla (2), il vettore ω_r è inteso solidale al corpo. Si dice allora, brevemente, che in un moto di precessione la velocità angolare è la somma di un vettore costante nello spazio e di un vettore costante nel corpo.

In un moto di precessione i coni di Poinot sono rotondi. L'asse di Mozzi a forma infatti angolo costante sia con l'asse di precessione che con l'asse di figura (Fig. 7.1). Dunque ogni moto di precessione è realizzabile facendo rotolare uniformemente un cono circolare retto C' su di un cono fisso C , pure circolare retto. Una precessione si dice **progressiva** o **retrograda** a seconda che i due vettori ω_p e ω_r formino angolo acuto o ottuso. Nella Fig. 7.1 la precessione è progressiva. Nella Fig. 7.2 sono raffigurati i coni di Poinot in due moti di precessione retrograda.

Un esempio notevole di moto di precessione è quello della Terra intorno al suo centro, rispetto alle stelle fisse. Sappiamo che la Terra ruota intorno al suo asse polare f , con velocità costante. Ma l'asse polare f non conserva rispetto alle stelle fisse direzione invariabile. Esso infatti ruota uniformemente intorno ad un asse fisso p passante per il suo centro e ortogonale al piano dell'eclittica (cioè della sua orbita intorno al Sole), compiendo però un solo giro in circa 26 mila anni! (questo periodo prende il nome di **anno platonico**). Si tratta di un moto di precessione retrograda. I coni di Poinot si configurano come nella prima della due Fig. 7.2. Soltanto che, mentre il cono fisso ha un'apertura lievemente superiore a $23^\circ 30'$, il cono solidale che gli rotola all'interno ha un'apertura di soli 8,67 millesimi di secondo circa, dovendo compiere un giro completo intorno all'asse p dopo aver compiuto su se stesso un numero di giri pari al numero dei giorni in un anno platonico (circa 9.400.000 giorni).

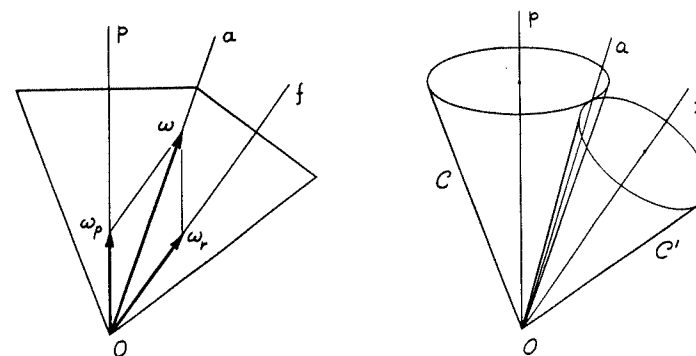


Fig. 7.1: coni di Poinot in un moto di precessione progressiva

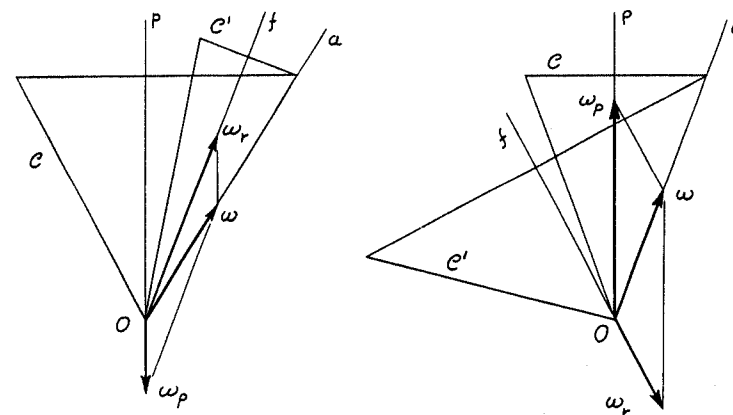


Fig. 7.2: coni di Poinot in due moti di precessione retrograda

INDICE ANALITICO

A

Abbassamento degli indici	51
Accelerazione	123-206
Accelerazione complementare	239
Accelerazione di Coriolis	239
Accelerazione di trascinamento	239
Accelerazione interna	218
Accelerazione intrinseca	218
Accelerazione normale	207
Accelerazione radiale	210
Accelerazione scalare	207
Accelerazione tangente	207
Accelerazione trasversa	210
Aggiunzione	73
Algebra di Lie	13-101
Algebra esterna	42
Algebra tensoriale	35
Angoli di Euler	81
Angolo di una rotazione	80
Angolo di nutazione	81
Angolo di precessione	81
Angolo di rotazione propria	81
Angolo fra vettori	48
Anno platonico	246
Anomalia	208
Antisimmetrizzazione	38
Applicazione lineare	4-5
Applicazione multilineare	4-5
Area di un dominio	187
Ascissa euclidea	178
Asse di una rotazione	80
Asse centrale	202
Asse di figura	246
Asse di moto	234
Asse di Mozzi	234

Asse di precessione	246
Assi centrali d'inerzia	196
Assi principali d'inerzia	196
Atto di moto elicoidale	234
Atto di moto rigido	234
Atto di moto rotatorio	235
Atto di moto traslatorio	235
Autodirezione	17
Automorfismo affine	228
Automorfismo lineare	11
Autovalori	18
Autovettori	17

B

Baricentro	188
Base canonica di una forma bilineare antisimmetrifca	31
Base canonica di una f. bilin. simm. .	25
Base canonica di uno sp. euclideo .	48
Base coniugata	51
Base duale	9
Base ordinata	69
Base ortonormale	49-69
Base reciproca	51
Basi concordi	69
Basi equiorientate	69

C

Calcolo differenziale esterno	112
Cambiamenti di coordinate	103
Campo centrale simmetrico	172
Campo centrifugo	173
Campo costante o uniforme	172
Campo coulombiano	172

Campo di forza	172
Campo di velocità	137
Campo elastico	173
Campo incompressibile	175
Campo irrotazionale	177
Campo newtoniano	173
Campo scalare	97
Campo solenoidale	176
Campo tensoriale	99
Campo vettoriale	98
Campo vettoriale completo	139
Campo vettoriale conservativo	171
Carta	102
Centro d'istantanea rotazione	242
Centro del moto	212
Centro di curvatura	180
Centro di massa	188
Classe di un campo scalare	97
Classe di un'applicazione	141
Classe C^∞	98
Codimensione	127-131
Colatitudine	104
Commutatore di campi vett.	101
Commutatore di endomorf. lineari ..	12
Complemento algebrico	44
Complessificazione di uno sp. vett. ..	18
Componenti cartesiane	98
Componenti contravarianti	51
Componenti covarianti	51
Componenti di un tensore	33
Componenti essenziali	42
Condizioni iniziali	138
Coni di Poincot	240
Cono di luce	88
Cono fisso	240
Cono solidale	240
Conservativo (campo vett.)	171
Contrazione	36
Coordinate canoniche	96
Coordinate cartesiane o affini	96
Coordinate cilindriche	104
Coordinate di una carta	102
Coordinate ortonormali	96
Coordinate polari	103

Coordinate polari sferiche	104
Coordinate superficiali	133
Corpo rigido	227
Coseni direttori	199
Costante delle aree	213
Curva integrale	137
Curva integrale massimale	139
Curva materiale	191
Curva non parametrizzata	122
Curva parametrizzata	122
Curva tangente	122
Curvatura di una curva	179
Curvatura gaussiana	186
Curvatura media	186
Curvatura totale	186
Curvature principali	185
Curve coordinate	106-133

D

Decomposizione spazio-temporale ...	86
Densità di massa	191
Densità scalare di massa	191
Derivata di un campo scalare	99
Derivata interna	218
Derivata intrinseca	218
Determinante di un endomorfismo ...	15
Differenziabile	98
Differenziale di un campo scalare ...	111
Differenziale di una p-forma	115
Direzioni principali di curvatura ...	185
Disuguaglianza di Schwarz	27-48
Disuguaglianza triangolare	49
Divergenza	101
Doppio prodotto vettoriale	76

E

Eccentricità	215
Elemento d'area	186
Ellissoide d'inerzia	186
Endomorfismo complementare ...	22-45
Endomorfismo lineare	11
Endomorfismo affine	227

Endomorfismo antisimmetrico	55
Endomorfismo assiale	79
Endomorfismo ortogonale	62
Endomorfismo simmetrico	55
Endomorfismo trasposto	53
Equazione caratteristica	18
Equazione del moto armonico	151
Equazione del moto centrifugo	151
Equazione del pendolo semplice	160
Equazione dell'oscillatore armonico ..	151
Equazione di una superficie	127
Equazioni delle geodetiche	220
Equazioni parametriche di una curva ..	122
Equazioni parametriche di una sup. ...	132
Equazioni parametriche di un moto ...	206
Esponenziale di un endomorfismo ...	65

F

Flusso	139
Flusso geodetico	220
Fogliettamento	130
Forma aggiunta	73
Forma bilineare	22
Forma bilineare antisimmetrica	24
Forma bilineare non degenerare	23
Forma bilineare regolare	23
Forma bilineare simmetrica	24
Forma bilineare trasposta	24
Forma chiusa	117
Forma differenziale	113
Forma differenziale lineare	112
Forma differenziale elementare	114
Forma esatta	117
Forma esterna	37-113
Forma lineare	8-110
Forma quadratica	26
Forma quadratica definita	27
Forma quadratica indefinita	27
Forma quadratica semi-definita	27
Forma volume	70-174
Formula di Binet	214
Formula di Grassmann	4
Formula di trasposizione dei momenti ..	201

Formule di Frenet	181
Formule di Poisson	232
Formule di trasposizione	194
Frequenza	212
Funzione integrale	152
Funzioni indipendenti	130
Futuro	88

G

Genere luce (vettore del)	47
Genere spazio (vettore del)	47
Genere tempo (vettore del)	47
Geodetiche	220
Geometria esterna	186
Geometria euclidea	97
Geometria interna	186
Gradiente	170
Gruppo di trasformazioni ad un parametro	140
Gruppo ortogonale	63
Gruppo ortogonale speciale	63

I

Identità di Jacobi	13-101
Immagine di un'appl. lineare	5
Immersione	131
Indicatore di Levi-Civita	43-70
Indici contravarianti	33
Indici covarianti	33
Innalzamento degli indici	51
Insiemi di livello	130-155
Integrale primo	152
Integrali primi banali	154
Integrali primi delle geodetiche	220
Integrali primi globali	154
Integrali primi locali	154
Invarianti principali	16
Isometria	62
Isometria affine	228
Isomorfismo lineare	5
Isotropo (vettore)	25-48

J			
Jacobiano	174		
L			
Laplaciano	173		
Legge d'inerzia	25		
Legge oraria	206		
Leggi di Keplero	215		
Lemma di Poincaré-Volterra	117		
Linee di flusso	142		
Lipschitziana (funzione)	144		
Longitudine	104		
Lunghezza di una curva	178		
M			
Massa totale	189		
Matrice d'inerzia	198		
Matrice delle componenti di un'applic. lineare	6		
Matrice jacobiana	103		
Matrice metrica	49-166		
Matrice metrica superficiale	183		
Matrice ortogonale	63		
Metrica	47-166		
Minore principale	16		
Modulo	47		
Molteplicità algebrica	18		
Molteplicità geometrica	18		
Momenti centrali d'inerzia	196		
Momenti principali d'inerzia	196		
Momento d'inerzia	194		
Momento risultante	201		
Moto centrale	212		
Moto centrale degenerare	213		
Moto circolare uniforme	124		
Moto di precessione regolare	246		
Moto di un grave	212		
Moto di un punto	122-205		
Moto geodetico	220		
Moto inerziale	208-220		
Moto periodico	212		
Moto piano	208		
Moto rettilineo uniforme	208		
Moto rigido	227-228		
Moto rigido con asse fisso	241		
Moto rigido con punto fisso	240		
Moto rigido piano	241		
Moto spontaneo	222		
Moto stazionario	137		
Moto traslatorio	239		
Moto traslatorio rettilineo	240		
Moto traslatorio rettilineo uniforme	240		
Moto uniforme	208		
Moto uniformemente accelerato	211		
N			
Norma	47		
Nucleo di un'applic. lineare	5		
O			
Omomorfismo lineare	5		
Operatore di antisimmetrizzazione	38		
Operatore di trasposizione	50-54		
Operatore ortogonale	50		
Orbita	122-206		
Orbite di un campo vett.	142		
Orientamenti	69		
Ortocroni (vettori)	87		
Ortogonale (operatore)	50		
Oscillatore armonico	155		
P			
Parametro	183		
Parametro di una conica	215		
Parametro euclideo	178		
Parentesi di Lie	101		
Parte antisimmetrica di un endom.	55		
Parte antisimmetrica di una f. bilin.	24		
Parte simmetrica di un endom.	55		
Parte simmetrica di una f. bilin.	24		
Passato	88		
Pendolo semplice	159		
Periodo	212		
p-forma	113		
Piano di simmetria materiale	190		
Piano di simmetria mat. ortogonale	196		
Piano osculatore	180		
Piano tangente	133		
Polare di un sottospazio	9		
Polare fissa	242		
Polare mobile	242		
Polarizzazione	26		
Polinomio annichilante	21		
Polinomio caratteristico	15		
Polo	201		
Potenziale	117-171		
Potenziale globale	118		
Potenziali locali	118		
Precessione regolare	246		
Precessione progressiva	246		
Precessione retrograda	246		
Prima forma fondamentale	182		
Prima formula fondamentale	232		
Prodotto esterno	29-39-116		
Prodotto misto	78		
Prodotto scalare	47		
Prodotto scalare di campi vett.	166		
Prodotto tensoriale	22-34		
Prodotto vettoriale	75		
Proprietà ciclica	78		
Proprietà commutativa graduata	40		
Proprietà di appartenenza	189		
Proprietà di evoluzione	140		
Proprietà di gruppo	140		
Proprietà di partizione	190		
Proprietà di scambio	78		
Proprietà di simmetria	190-196		
Proprietà distributiva	190		
Pseudoangolo di rotazione	93		
Pulsazione	212		
Punto critico di una funzione	127		
Punto ellittico	185		
Punto iperbolico	185		
Punto materiale	188		
Punto ombelicale	185		
Punto parabolico	185		
Punto regolare di una superficie	127-131		
Punto singolare di un campo vett.	142		
Punto singolare di una funzione	127		
Punto solidale	227		
Q			
Quaternione	67		
R			
Raggio	208		
Raggio di curvatura	180		
Raggio vettore	122		
Rango di una matrice	8		
Rango di un'applic. lineare	8		
Rappresentazione di Cayley	66		
Rappresentazione esponenziale	65		
Rappresentazione intrinseca	206		
Rappresentazione quaternionale	67		
Rappresentazione radiale	207		
Rappresentazione vettoriale di una curva - di un moto	122-205		
Regola di Leibniz	100		
Retta di applicazione	201		
Riferimento cartesiano	95		
Riferimento equivalente	239		
Riferimento naturale	105		
Riferimento naturale tangente	133		
Riferimento solidale	227		
Risultante	201		
Ritratto di fase	142		
Rotazione	63		
Rotazione impropria	63		
Rotazione ortocrona	92		
Rotore	176		
S			
Saturazione	36		
Seconda forma fondamentale	184		
Seconda formula fondamentale	234		
Segnatura	25		
Segnatura di uno spazio euclideo	48		
Simboli di Christoffel	107		

Simboli di Christoffel di 1.a specie ..	169
Simboli di Christoffel di 2.a specie ..	184
Simbolo di Kronecker	9
Simbolo di Levi-Civita	43-70
Simmetria	64
Simmetria materiale ortogonale	196
Sistema autonomo	138
Sistema del primo ordine	138
Sistema di coordinate	102
Sistema di equazioni differenziali ...	138
Sistema di masse	188
Sistema dinamico	110
Sistema dinamico di Lotka-Volterra ..	161
Sistema di vettori applicati	201
Sistema materiale continuo	191
Sistema materiale omogeneo	191
Sistema materiale piano	200
Sistema normale	138
Sistemi di vettori equivalenti	203
Solido materiale	191
Somma diretta di spazi	2
Sottomatrice principale	16
Sottospazio affine	55
Sottospazio euclideo	53
Sottospazio invariante	16
Sottospazio isotropo	53
Sottospazio non-degenere	53
Spazio affine	95
Spazio affine euclideo	96-165
Spazio biduale	9
Spazio ortogonale	52
Spazio pseudo-euclideo	48
Spazio semi-euclideo	48
Spazio strettamente euclideo	48
Spazio tensoriale	35
Spazio vettoriale	1
Spazio vettoriale di Minkowski	85
Spazio vettoriale euclideo	47
Spazio vettoriale iperbolico	85
Spazio vettoriale soggiacente	95
Spettro	18
Superfici di livello	130-155
Superfici equipotenziali	173
Superficie	127

Superficie immersa	131
Superficie materiale	191
Superficie orientabile	183
Superficie regolare	127

T

Tensore contravariante	34
Tensore covariante	34
Tensore d'inerzia	193
Tensore di curvatura	184
Tensore di tipo (p, q)	32
Tensore metrico	47-166
Tensore omogeneo	34
Teorema dei moti relativi	237
Teorema del gradiente	173
Teorema della somma delle velocità angolari	245
Teorema di Cauchy	138
Teorema di Chasles	242
Teorema di Clairaut	226
Teorema di Coriolis	239
Teorema di Hamilton-Cayley	21
Teorema di Huygens	195
Teorema di Mozzi	234
Teorema di Sylvester	25
Teoremi di Guldino	192
Terna fissa	229-235
Terna fondamentale	180
Terna solidale	227-235
Theorema egregium	186
Toro	134
Torsione	181
Traccia di un endomorfismo	15
Traiettoria	122-206
Trasformazione affine	96-228
Trasformazione lineare	11
Trasformazioni di coordinate	103
Trasposta di una forma bilin.	24
Trasposto endomorfismo	53
Triedro fondamentale	180

V

Valutazione tra forme e vettori ...	9-110
-------------------------------------	-------

Velocità	122-206
Velocità angolare	124-208-231-233
Velocità angolare di trascinamento ..	245
Velocità areale	207
Velocità di precessione	246
Velocità di rotazione propria	246
Velocità di trascinamento	237
Velocità radiale	210
Velocità scalare	207
Velocità trasversa	210
Versore	47
Versore binormale	180
Versore di un angolo	80
Versore di una rotazione	80
Versore normale principale	179
Versore radiale	107
Versore tangente	179
Versore trasverso	107-209

Vettore aggiunto	74
Vettore applicato	95
Vettore colonna	7
Vettore del genere luce	47
Vettore del genere spazio	47
Vettore del genere tempo	47
Vettore isotropo	25-48
Vettore libero	95
Vettore posizione	122
Vettore riga	7
Vettore tangente ad una curva	122
Vettore unitario	25-47
Vettori coniugati	25
Vettori ortocroni	87
Vettori ortogonali	25-47
Volume di un dominio	175

Finito di stampare nel mese di marzo 1994
nella M.S./Litografia s.r.l. di Torino
Via Mazzini, 24