

ANALISI MATEMATICA

Ottavio Caligaris - Pietro Oliva



CAPITOLO 1

SPAZI EUCLIDEI n -DIMENSIONALI.

Per lo studio delle funzioni di più variabili reali occorre aver presenti alcune proprietà degli spazi euclidei ad n dimensioni.

È inoltre indispensabile conoscere qualche proprietà delle applicazioni lineari in $\mathbb{R}^n$ e delle forme bilineari e quadratiche.

1. Norma e Prodotto scalare

DEFINIZIONE 1.1. *Indichiamo con $\mathbb{R}^n$ lo spazio vettoriale costituito dalle $n - ple$ ordinate di numeri reali; in altre parole*

$$x \in \mathbb{R}^n \Leftrightarrow x = (x_1, x_2, \dots, x_n) \text{ con } x_k \in \mathbb{R}.$$

In $\mathbb{R}^n$ si definiscono le operazioni di somma e di prodotto per uno scalare mediante le

$$x + y = (x_1 + y_1, x_2 + y_2, \dots, x_n + y_n) \text{ , } x, y \in \mathbb{R}^n$$

e

$$\alpha x = (\alpha x_1, \alpha x_2, \dots, \alpha x_n) \text{ , } \alpha \in \mathbb{R} \text{ , } x \in \mathbb{R}^n.$$

L'insieme dei vettori

$$e_1 = (1, 0, 0, \dots, 0)$$

$$e_2 = (0, 1, 0, \dots, 0)$$

.....

$$e_n = (0, 0, 0, \dots, 1)$$

costituisce una base di $\mathbb{R}^n$; si avrà pertanto che, se $x \in \mathbb{R}^n$

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

DEFINIZIONE 1.2. *Si definisce norma in $\mathbb{R}^n$ una funzione che si indica con*

$$\| \cdot \| : \mathbb{R}^n \rightarrow \mathbb{R}$$

che verifica le seguenti proprietà:

- (1) $\|x\| \geq 0 \quad \forall x \in \mathbb{R}^n$
- (2) $\|x\| = 0 \Leftrightarrow x = 0$
- (3) $\|\alpha x\| = |\alpha| \|x\| \quad \forall \alpha \in \mathbb{R}, \forall x \in \mathbb{R}^n$
- (4) $\|x + y\| \leq \|x\| + \|y\| \quad \forall x, y \in \mathbb{R}^n.$

DEFINIZIONE 1.3. Si definisce *prodotto scalare* in $\mathbb{R}^n$ una funzione

$$\langle \cdot, \cdot \rangle : \mathbb{R}^n \times \mathbb{R}^n \longrightarrow \mathbb{R}$$

tale che

- (1) $\langle x, x \rangle \geq 0 \quad \forall x \in \mathbb{R}^n$
- (2) $\langle x, y \rangle = \langle y, x \rangle \quad \forall x, y \in \mathbb{R}^n$
- (3) $\langle x, x \rangle = 0 \Leftrightarrow x = 0$
- (4) $\langle \alpha x + \beta y, z \rangle = \alpha \langle x, z \rangle + \beta \langle y, z \rangle \quad \forall x, y, z \in \mathbb{R}^n, \forall \alpha, \beta \in \mathbb{R}.$

Come nel caso del valore assoluto per i numeri reali, si ha

$$||x| - |y|| \leq |x - y|$$

LEMMA 1.1. Siano $x, y \in \mathbb{R}^n$, allora

- **-Disuguaglianza di Schwarz-Holder-** se $1/p + 1/q = 1$, $p, q \geq 1$,

$$\left| \sum x_i y_i \right| \leq \sum |x_i| |y_i| \leq \left(\sum |x_i|^p \right)^{1/p} \left(\sum |y_i|^q \right)^{1/q}$$

- **-Disuguaglianza di Minkowski-**

$$\left(\sum |x_i + y_i|^p \right)^{1/p} \leq \left(\sum |x_i|^p \right)^{1/p} + \left(\sum |y_i|^p \right)^{1/p}$$

La disuguaglianza di Holder si riduce alla più nota disuguaglianza di Schwarz per $p = q = 2$.

Per $p = q = 2$ la disuguaglianza di Schwarz può essere riscritta come

$$|\langle x, y \rangle| \leq \|x\| \|y\|$$

e può essere dedotta osservando che, $\forall t \in \mathbb{R}$

$$0 \leq \|x + ty\|^2 = \langle x + ty, x + ty \rangle = t^2\|y\|^2 + 2t\langle x, y \rangle + \|x\|^2$$

Ciò implica infatti

$$\langle x, y \rangle^2 - \|x\|^2\|y\|^2 \leq 0$$

La corrispondente disuguaglianza triangolare segue da

$$\|x + y\|^2 = \|x\|^2 + \|y\|^2 + 2\langle x, y \rangle \leq \|x\|^2 + \|y\|^2 + 2\|x\|\|y\|$$

Osserviamo che

$$|\langle x, y \rangle| = \|x\|\|y\|$$

se e solo se esiste $t \in \mathbb{R}$ tale che $x + ty = 0$, ovvero x e y sono paralleli.

Pertanto

$$\|x\| = \sup\{\langle x, y \rangle : \|y\| \leq 1\} = \max\{|\langle x, y \rangle| : \|y\| \leq 1\}.$$

Sono esempi di norme in $\mathbb{R}^n$ le seguenti

$$\|x\|_p = \left(\sum_{k=1}^n |x_k|^p \right)^{1/p} \quad p \geq 1$$

$$\|x\|_{\infty} = \max\{|x_i| : i = 1, \dots, n\}.$$

Mentre un esempio di prodotto scalare è dato da

$$\langle x, y \rangle = \left(\sum_{k=1}^n x_k y_k \right)$$

Ovviamente si ha $\langle x, x \rangle = \|x\|_2^2$.

In $\mathbb{R}^n$ useremo abitualmente la $\|\cdot\|_2$ che è detta norma euclidea in quanto $\|x\|_2$ coincide con la distanza euclidea del vettore x dall'origine.

Nel seguito faremo riferimento, a meno di espliciti avvisi contrari, a tale norma e scriveremo $\|\cdot\|$ in luogo di $\|\cdot\|_2$.

Osserviamo altresì che il prodotto scalare sopra definito, pur non essendo l'unico possibile, sarà l'unico da noi considerato.

Si vede subito che

$$\langle P_1, P_2 \rangle = |P_1| |P_2| \cos(\theta_2 - \theta_1)$$

Infatti, facendo riferimento ad $\mathbb{R}^2$ e alla figura 4.5, si ha

$$\begin{aligned}
 \langle P_1, P_2 \rangle &= \\
 &= x_1 x_2 + y_1 y_2 = |P_1| |P_2| \cos(\theta_2) \cos(\theta_1) + |P_1| |P_2| \sin(\theta_2) \sin(\theta_1) = \\
 &= |P_1| |P_2| \cos(\theta_2 - \theta_1)
 \end{aligned}$$

L'osservazione appena fatta giustifica il fatto che

Diciamo che due vettori $x, y \in \mathbb{R}^n$ sono ortogonali se $\langle x, y \rangle = 0$.

Diciamo che sono paralleli se esiste $\lambda \in \mathbb{R}$ tale che $x = \lambda y$. Se x ed y sono paralleli $\langle x, y \rangle = \|x\| \|y\|$

Se $\|\cdot\|_a$ e $\|\cdot\|_b$ sono due norme in $\mathbb{R}^n$ si dice che sono equivalenti se esistono due costanti reali H e K tali che

$$H \|x\|_b \leq \|x\|_a \leq K \|x\|_b.$$

Si può dimostrare che

In $\mathbb{R}^n$ tutte le norme sono equivalenti.

È anche interessante osservare che

La funzione $p \longrightarrow \|x\|_p$ è decrescente $\forall x \in \mathbb{R}^n$ e si ha

$$\|x\|_\infty = \lim_p \|x\|_p = \inf\{\|x\|_p : p \geq 1\}$$

inoltre

$$\|x\|_\infty \leq \|x\|_p \leq \|x\|_1 \leq n\|x\|_\infty \leq n\|x\|_q \leq n\|x\|_1 \quad \forall p, q \geq 1$$

e pertanto le norme $\|\cdot\|_p$ sono tutte equivalenti.

Le notazioni vettoriali introdotte consentono di esprimere facilmente condizioni che individuano rette piani e sfere.

Possiamo individuare i punti di una retta che passa per il punto (x_0, y_0, z_0) ed è parallela alla direzione (a, b, c) semplicemente sommando (x, y, z) con il vettore $t(a, b, c)$ al variare di $t \in \mathbb{R}$.

Otterremo in tal caso che

$$(x, y, z) = (x_0, y_0, z_0) + t(a, b, c)$$

che scritta componente per componente

$$\begin{cases} x = x_0 + ta \\ y = y_0 + tb \\ z = z_0 + tc \end{cases}$$

fornisce le equazioni parametriche della retta.

Se $t \in \mathbb{R}_+$ avremo una delle due semirette in cui (x_0, y_0, z_0) divide la retta intera, mentre se $t \in [a, b]$ ci limitiamo ad un segmento della retta stessa.

Un piano passante per l'origine può essere individuato dai vettori perpendicolari ad un vettore assegnato; l'equazione del piano si potrà quindi scrivere come

$$\langle (x, y, z), (a, b, c) \rangle = 0$$

mentre il piano parallelo che passa per (x_0, y_0, z_0) è dato da

$$\langle (x - x_0, y - y_0, z - z_0), (a, b, c) \rangle = 0$$

come abbiamo già visto una sfera può essere individuata come l'insieme dei punti che hanno distanza dal centro (x_0, y_0, z_0) minore del raggio R ;

Una sfera sarà pertanto individuata dalla condizione

$$\|(x - x_0, y - y_0, z - z_0)\| \leq R$$

2. Applicazioni Lineari

DEFINIZIONE 1.4. Si chiama *applicazione lineare* una funzione

$$f : \mathbb{R}^n \rightarrow \mathbb{R}^m$$

tale che

$$f(\alpha x + \beta y) = \alpha f(x) + \beta f(y) \quad \forall x, y \in \mathbb{R}^n, \quad \forall \alpha, \beta \in \mathbb{R}$$

$\mathcal{L}(\mathbb{R}^n, \mathbb{R}^m)$ è l'insieme delle applicazioni lineari su $\mathbb{R}^n$ a valori in $\mathbb{R}^m$. $\mathcal{L}(\mathbb{R}^n, \mathbb{R})$ si chiama anche *spazio duale di $\mathbb{R}^n$* .

Gli elementi di $\mathcal{L}(\mathbb{R}^n, \mathbb{R}^m)$ possono essere messi in corrispondenza biunivoca con le matrici aventi m righe ed n colonne, $\mathcal{M}^{m \times n}$.

Più precisamente $\mathcal{L}(\mathbb{R}^n, \mathbb{R}^m)$ ed $\mathcal{M}^{m \times n}$ sono isomorfi in quanto ogni applicazione lineare f può essere scritta nella forma

$$f(x) = Ax \quad \text{con } A \in \mathcal{M}^{m \times n}.$$

e d'altro canto $f(x) = Ax$ è lineare.

In particolare

Le applicazioni lineari da $\mathbb{R}^n$ in $\mathbb{R}$ sono tutte e sole quelle della forma

$$f(x) = \langle x^*, x \rangle \quad \text{con } x^* \in \mathbb{R}^n.$$

DEFINIZIONE 1.5. Se $f \in \mathcal{L}(\mathbb{R}^n, \mathbb{R}^m)$ definiamo norma di f ,

$$\|f\|_0 = \sup\{\|f(x)\| : \|x\| \leq 1\}.$$

Possiamo identificare f con la matrice A per la quale risulta $f(x) = Ax$, per cui possiamo anche definire

$$\|A\|_0 = \sup\{\|Ax\| : \|x\| \leq 1\}.$$

D'altro canto si può definire

$$\|A\|_p = \left(\sum_{ij} |a_{ij}|^p \right)^{1/p} \quad p \geq 1$$

e

$$\|A\|_\infty = \max\{|a_{ij}| : i = 1, \dots, m, j = 1, \dots, n\}$$

esattamente come negli spazi euclidei e si può osservare che

$$\|A\|_0 \leq \|A\|_2 = \|A\|$$

la disuguaglianza essendo stretta ad esempio se $A = I$.

Possiamo anche provare che

Si ha che

$$\|A\|_0 = \sup\{|\langle Ax, y \rangle| : \|x\| \leq 1, \|y\| \leq 1\}$$

ed inoltre, se $A \in \mathcal{M}^{n \times n} = \mathcal{M}^n$ è simmetrica

$$\|A\|_0 = \sup\{|\langle Ax, x \rangle| : \|x\| \leq 1\} = \Lambda$$

dove

$$\Lambda = \max\{\lambda_i : i = 1, 2, \dots, n\} = \lambda_{i_0}$$

3. Forme Bilineari e Quadratiche

DEFINIZIONE 1.6. Si chiama *forma bilineare in $\mathbb{R}^n$* una funzione

$$f : \mathbb{R}^n \times \mathbb{R}^n \rightarrow \mathbb{R}$$

tale che $f(\cdot, y)$ e $f(x, \cdot)$ siano funzioni lineari su $\mathbb{R}^n$.

Le funzioni bilineari su $\mathbb{R}^n$ sono tutte e sole quelle definite da

$$f(x, y) = \langle x, Ay \rangle = \langle Bx, y \rangle \quad \text{con } A, B \in \mathcal{M}^n$$

dove B è la matrice trasposta di A . (si ottiene da A scambiando le righe con le colonne. Solitamente si denota $B = A^*$).

DEFINIZIONE 1.7. Se f è una forma bilineare in $\mathbb{R}^n$; la funzione

$$g : \mathbb{R}^n \longrightarrow \mathbb{R}$$

definita da

$$g(x) = f(x, x)$$

si chiama forma quadratica in $\mathbb{R}^n$.

Si può sempre trovare una matrice $A \in \mathcal{M}^n$, non unica, tale che

$$g(x) = \langle x, Ax \rangle$$

possiamo inoltre sempre scegliere A in modo che sia una matrice simmetrica; in tal caso A si dice matrice associata alla forma quadratica e risulta univocamente determinata.

Se g è una forma quadratica; g si dice semidefinita positiva (negativa) se

$$g(x) \geq 0 \quad (\leq 0) \quad \forall x \in \mathbb{R}^n$$

g si dice definita positiva (negativa) se

$$g(x) > 0 \quad (< 0) \quad \forall x \in \mathbb{R}^n \setminus \{0\}.$$

Si possono provare i seguenti:

TEOREMA 1.1. *Sia g una forma quadratica e sia A la matrice ad essa associata; allora*

- *g è definita positiva se e solo se, detto A_k il minore principale di ordine k di A , si ha*

$$\det A_k > 0 \quad \forall k ;$$

- *g è definita negativa se solo se*

$$(-1)^k \det A_k > 0 \quad \forall k.$$

TEOREMA 1.2. *Sia g una forma quadratica e sia A la matrice ad essa associata; allora*

- *g è definita positiva (negativa) se e solo se, detti λ_k i suoi autovalori, si ha*

$$\lambda_k > 0 \quad (< 0) \quad \forall k ;$$

- *g è semidefinita positiva (negativa) se e solo se*

$$\lambda_k \geq 0 \quad (\leq 0) \quad \forall k.$$

Il prodotto scalare in $\mathbb{R}^n$

$$f(x, y) = \langle x, y \rangle$$

è il più semplice esempio di funzione bilineare; la forma quadratica

$$g(x) = f(x, x) = \langle x, x \rangle = \|x\|^2$$

si riduce alla norma euclidea in $\mathbb{R}^n$. La matrice di rappresentazione della forma bilineare associata al prodotto scalare è la matrice identica.

4. Proprietà Topologiche

Una successione in $\mathbb{R}^n$ è una applicazione

$$\begin{aligned} x : \mathbb{N} &\longrightarrow \mathbb{R}^n \\ \mathbb{N} \ni k &\mapsto x_k \in \mathbb{R}^n \end{aligned}$$

Le proprietà di una successione in $\mathbb{R}^n$ possono essere studiate componente per componente.

DEFINIZIONE 1.8. Chiamiamo sfera aperta di centro x_0 e raggio r , l'insieme

$$S(x_0, r) = \{x \in \mathbb{R}^n : \|x - x_0\| < r\}$$

Se $A \subset \mathbb{R}^n$

- $x_0 \in A$ si dice interno ad A se esiste $r > 0$ tale che $S(x_0, r) \subset A$.
- x_0 è punto di frontiera per $A \subset \mathbb{R}^n$ se

$$\forall \delta > 0 \quad S(x_0, \delta) \cap A \neq \emptyset \quad e \quad S(x_0, \delta) \cap A^c \neq \emptyset$$

L'insieme dei punti di A che sono interni si indica con $\text{int } A$.

$A \subset \mathbb{R}^n$ si dice aperto se tutti i suoi punti sono interni, cioè se $A = \text{int } A$.

$A \subset \mathbb{R}^n$ si dice chiuso se il suo complementare è aperto.

∂A è l'insieme dei punti di frontiera di A .

La chiusura di A è l'insieme

$$\text{cl } A = \{x \in \mathbb{R}^n : \exists x_k \in A, x_k \rightarrow x\}.$$

Si può verificare che

Se $A \subset \mathbb{R}^n$

- $\text{cl } A = \partial A \cup A$
- **A è aperto se e solo se $\forall x \in A, \forall x_k, x_k \rightarrow x$, si ha $x_k \in A$ definitivamente;**
- **A è chiuso se e solo se $\forall x_k \in A, x_k \rightarrow x$, si ha $x \in A$.**

DEFINIZIONE 1.9. Un insieme $A \subset \mathbb{R}^n$ si dice

- *limitato*, se esiste $r > 0$ tale che $A \subset S(0, r)$;
- *convesso*, se $\forall x, y \in A, \forall \lambda \in [0, 1], \lambda x + (1 - \lambda)y \in A$;
- *compatto*, se $\forall x_k \in A$, esiste un'estratta $x_{k_h} \rightarrow x \in A$;
- *connesso*, se non esistono due insiemi aperti, A_1, A_2 tali che $A_1 \cap A \neq \emptyset, A_2 \cap A \neq \emptyset, A \cap (A_1 \cap A_2) = \emptyset, A \subset A_1 \cup A_2$.

Si può dimostrare che

- A è compatto se e solo se A è chiuso e limitato.
- In $\mathbb{R}$ gli insiemi connessi sono tutti e soli gli intervalli.
- Gli insiemi connessi ed aperti di $\mathbb{R}^n$ possono essere caratterizzati dalla seguente condizione
 $\forall x, y \in A$ esiste una funzione

$$\phi : [a, b] \longrightarrow A$$

continua, lineare a tratti (*il cui grafico è costituito da segmenti paralleli agli assi*) **tale che**
 $\phi(a) = x$ e $\phi(b) = y$.

- Se A è convesso, allora A è connesso.

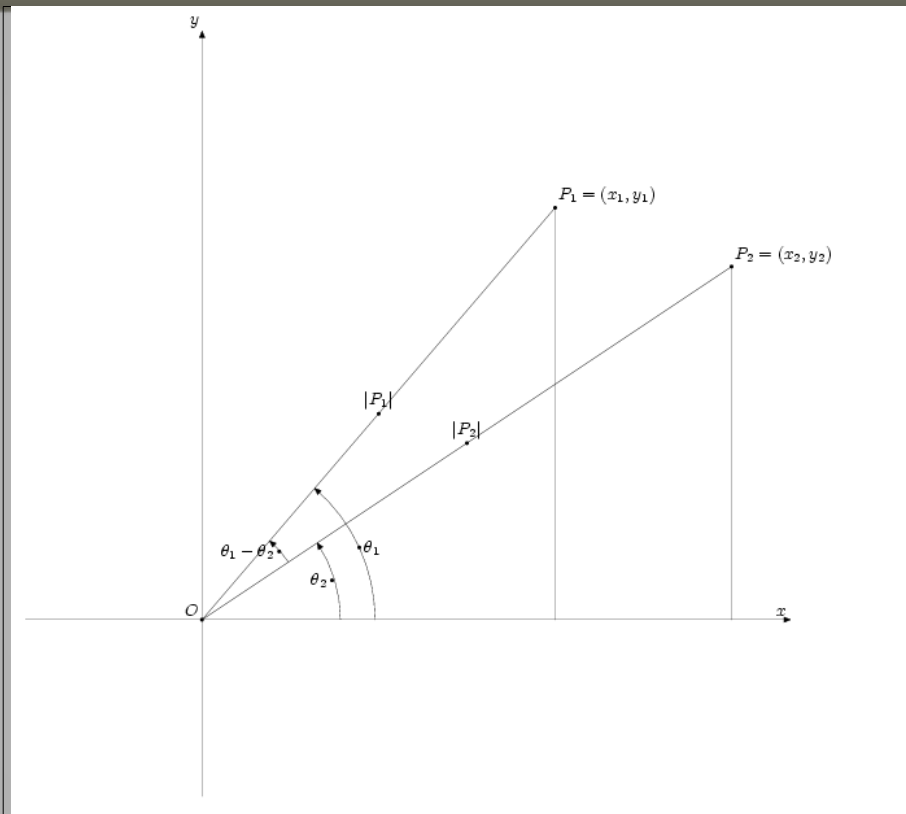


FIGURA 1.1.

CAPITOLO 2

LE FUNZIONI DI PIÙ VARIABILI

Questo capitolo è dedicato allo studio delle proprietà di continuità e differenziabilità delle funzioni

$$f : \mathbb{R}^n \rightarrow \mathbb{R}^m$$

con $n, m \geq 1$,

DEFINIZIONE 2.1. *È data una funzione*

$$f : \mathbb{R}^n \longrightarrow \mathbb{R}^m$$

se sono assegnati

- *un insieme $A \subset \mathbb{R}^n$*
- *una corrispondenza*

$$x \in A \mapsto f(x) \in \mathbb{R}^m$$

che ad ogni $x \in A$ associa uno ed un solo vettore $f(x) \in \mathbb{R}^m$.

Si dice che A è il dominio di f e si scrive $D(f) = A$ e, nel caso che tale dominio non sia esplicitamente indicato, si suppone la corrispondenza f definita per tutti gli $x \in \mathbb{R}^n$ per cui è possibile considerare $f(x)$.

Si definisce rango di f

$$R(f) = \{y \in \mathbb{R}^m : \exists x \in A, y = f(x)\}$$

e grafico di f .

$$G(f) = \{(x, y) \in \mathbb{R}^n \times \mathbb{R}^m : y = f(x)\}$$

Restrizione e composizione di funzioni sono definite come nel caso reale e parimenti simile è la definizione di iniettività, surgettività, bigettività.

1. Limiti

Lo studio dei limiti di una funzione di più può essere condotto semplicemente ripercorrendo i risultati ottenuti nel caso di una funzione reale di una variabile reale, avendo cura di puntualizzare solo qualche particolare sui valori infiniti.

Quando $n = 1$, si estende $\mathbb{R}$ mediante due punti all'infinito che vengono denominati $+\infty$ e $-\infty$, in quanto è ben chiaro che due punti di $\mathbb{R}$ possono sempre essere confrontati nella relazione d'ordine ($\mathbb{R}$ è totalmente ordinato). Se $n > 1$ cade la possibilità di ordinare totalmente $\mathbb{R}^n$ e pertanto si preferisce estendere $\mathbb{R}^n$, $n > 1$, con un solo punto all'infinito che viene denominato semplicemente ∞ . Ricordiamo che, anche se meno utile, questa possibilità esiste anche in $\mathbb{R}$

DEFINIZIONE 2.2. Se $n = 1$ e $x_0 \in \mathbb{R} \cup \{\pm\infty\}$, definiamo, per $\rho > 0$

$$I(x_0, \rho) = \begin{cases} (\rho, +\infty) & \text{se } x = +\infty \\ (x_0 - \rho, x_0 + \rho) & \text{se } x \in \mathbb{R} \\ (-\infty, -\rho) & \text{se } x = -\infty \end{cases}$$

Definiamo inoltre $I^0(x_0, \rho) = I(x_0, \rho) \setminus \{x_0\}$.

Se $n > 1$, $x_0 \in \mathbb{R}^n \cup \{\infty\}$, definiamo per $\rho > 0$

$$I(x_0, \rho) = \begin{cases} \{x_0 \in \mathbb{R}^n : \|x_0\| < \rho\} = S(x_0, \rho) & x_0 \in \mathbb{R}^n \\ \{x_0 \in \mathbb{R}^n : \|x_0\| > \rho\} & x_0 = \infty \end{cases}$$

anche qui poniamo $I^0(x_0, \rho) = I(x_0, \rho) \setminus \{x_0\}$

DEFINIZIONE 2.3. Sia $A \subset \mathbb{R}^n$, si dice che $x_0 \in \mathbb{R}^n \cup \{\infty\}$ è un punto di accumulazione per A se $\forall r > 0 \ I^0(x_0, r) \cap A \neq \emptyset$.

Indichiamo con $\mathcal{D}(A)$ l'insieme dei punti di accumulazione di A .

Sia $A \subset \mathbb{R}^n$, $x_0 \in \mathbb{R}^n \cup \{\infty\}$; $x_0 \in \mathcal{D}(A)$ **se e solo se** $\exists x_k \in A, x_k \neq x_0, x_k \rightarrow x_0$.

DEFINIZIONE 2.4. Sia $f : A \rightarrow \mathbb{R}^n$, $A \subset \mathbb{R}^n$ e sia $x_0 \in \mathcal{D}(A)$; diciamo che

$$\lim_{x \rightarrow x_0} f(x) = \ell$$

se

$\forall \varepsilon > 0 \ \exists \delta(\varepsilon) > 0$ tale che se $x \in I^0(x_0, \delta(\varepsilon)) \cap A$ si ha $f(x) \in I(\ell, \varepsilon)$

Si può facilmente provare che:

- **ogni funzione che ammette limite finito è localmente limitata;**
- **il limite di una funzione, se esiste, è unico;**
- **se $m = 1$ vale il teorema della permanenza del segno;**
- **il limite di una somma è uguale alla somma dei limiti, ove questi esistono finiti;**
- **il limite del prodotto di una funzione a valori reali per una funzione a valori vettoriali è uguale al prodotto dei limiti, ove questi esistono finiti;**
- **se $m = 1$ il limite del reciproco di una funzione è uguale al reciproco del limite della funzione stessa, ammesso che non sia nullo.**
- **se $m = 1$ valgono i risultati sul confronto dei limiti del tipo considerato per le funzioni reali di una variabile;**
- **il limite di una funzione può essere caratterizzato per successioni come per le funzioni di una variabile.**

Ricordiamo anche l'enunciato che permette di calcolare il limite di una funzione composta

Sia $f : A \longrightarrow \mathbb{R}^m$, $A \subset \mathbb{R}^n$, $x_0 \in \mathcal{D}(A)$ e sia $g : B \longrightarrow A$, $B \subset \mathbb{R}^p$, $y_0 \in \mathcal{D}(B)$, $g(B) \subset A$; supponiamo che

$$\lim_{x \rightarrow x_0} f(x) = \ell \quad \text{e} \quad \lim_{y \rightarrow y_0} g(y) = x_0$$

Allora, se una delle due seguenti condizioni è verificata

- $x_0 \notin \text{dom } f$
- $f(x_0) = \ell$

si ha

$$\lim_{y \rightarrow y_0} f(g(y)) = \ell.$$

2. Continuità

DEFINIZIONE 2.5. Sia $f : A \longrightarrow \mathbb{R}^m$, $x_0 \in A \subset \mathbb{R}^n$, diciamo che f è una funzione continua in x_0 se $\forall \varepsilon > 0 \exists \delta(\varepsilon) > 0$ tale che

se $x \in A$, $\|x - x_0\| < \delta(\varepsilon)$ si ha $\|f(x) - f(x_0)\| < \varepsilon$.

Nel caso in cui $x_0 \in A \cap \mathcal{D}(A)$, la condizione sopra espressa è equivalente alla

$$\lim_{x \rightarrow x_0} f(x) = f(x_0)$$

f si dice continua in A se è continua in ogni punto di A .

Come nel caso delle funzioni reali di una variabile reale si prova che:

- **la somma di funzioni continue è continua;**
- **il prodotto di una funzione a valori vettoriali per una funzione a valori scalari, entrambe continue, è continua;**
- **se $m = 1$, il reciproco di una funzione continua è continuo dove ha senso definirlo;**
- **vale la caratterizzazione della continuità per successioni data, nel caso reale;**
- **la composta di funzioni continue è una funzione continua.**
- **se $\lim_{x \rightarrow x_0} f(x) = \lambda$ e $\lim_{x \rightarrow x_0} g(x) = \mu$, con $\lambda, \mu \in \mathbb{R}^m$ si ha**

$$\lim_{x \rightarrow x_0} \langle f(x), g(x) \rangle = \langle \lambda, \mu \rangle$$

- **In particolare la funzione $\langle \cdot, \cdot \rangle : \mathbb{R}^m \times \mathbb{R}^m \rightarrow \mathbb{R}$ è continua**

Valgono per le funzioni continue i soliti teoremi

TEOREMA 2.1. -degli zeri - Sia $f : A \longrightarrow \mathbb{R}$, $A \subset \mathbb{R}^n$, A aperto e connesso e supponiamo che f sia una funzione continua; allora se esistono $x_1, x_2 \in A$ tali che $f(x_1)f(x_2) < 0$ esiste anche $x_0 \in A$ tale che $f(x_0) = 0$.

DIMOSTRAZIONE. Poichè A è connesso è possibile congiungere x_1 ed x_2 con una linea spezzata costituita di segmenti paralleli agli assi coordinati.

Siano x_j gli estremi di ciascuno dei segmenti, nel caso in cui $f(x_j) = 0$ per qualche j , il teorema è dimostrato, in caso contrario esisteranno x_k, x_{k+1} tali che $f(x_k)f(x_{k+1}) < 0$.

Allora la funzione $[0, 1] \ni t \mapsto \varphi(t) = f(x_k + t(x_{k+1} - x_k)) \in \mathbb{R}$ è continua e si può applicare a φ il teorema degli zeri. $\square$

TEOREMA 2.2. - Weierstraß- Sia $f : A \longrightarrow \mathbb{R}$ una funzione continua e supponiamo che A sia un insieme compatto; allora esistono $x_1, x_2 \in A$ tali che

$$f(x_1) = \min\{f(x) : x \in A\}$$

$$f(x_2) = \max\{f(x) : x \in A\}$$

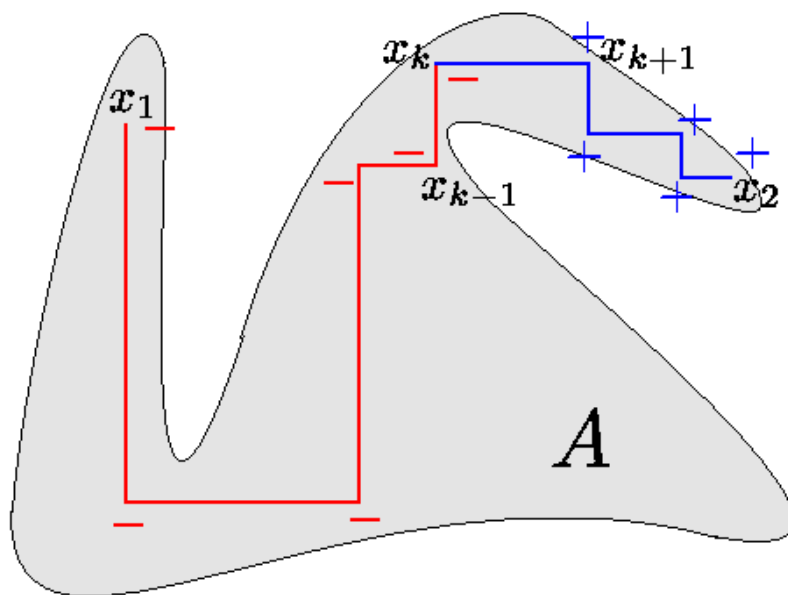


FIGURA 2.1. Il teorema degli zeri

DIMOSTRAZIONE. Sia, ad esempio $\lambda = \inf\{f(x) : x \in A\}$; allora esiste una successione $x_k \in A$ tale che $f(x_k) \rightarrow \lambda$ e, dal momento che A è compatto, è possibile trovare una successione $x_{k_h} \rightarrow x_1 \in A$. Si ha pertanto $f(x_{k_h}) \rightarrow \lambda$ e, per la continuità di f , $f(x_{k_h}) \rightarrow f(x_1)$. Ne segue che $\lambda = f(x_1)$ e la tesi. $\square$

TEOREMA 2.3. - Weierstraß generalizzato - Sia $f : \mathbb{R}^n \longrightarrow \mathbb{R}$ una funzione continua e supponiamo che esista $\hat{x} \in \mathbb{R}^n$ tale che

$$\lim_{x \rightarrow \infty} f(x) > f(\hat{x})$$

allora esiste $x_1 \in \mathbb{R}^n$ tale che

$$f(x_1) = \min\{f(x) : x \in \mathbb{R}^n\}$$

DIMOSTRAZIONE. Sia

$$\lambda = \inf\{f(x) : x \in \mathbb{R}^n\}$$

si ha $f(\hat{x}) \geq \lambda$.

Sia poi $\delta > 0$ tale che se $\|x\| > \delta$ si abbia $f(x) \leq f(\hat{x}) + \varepsilon$, con $\varepsilon > 0$.

Sia ancora $x_k \in \mathbb{R}^n$ tale che $f(x_k) \rightarrow \lambda$.

Allora, per k abbastanza grande, si ha $f(x_k) < f(\hat{x}) + \varepsilon$ e quindi $\|x_k\| \leq \delta$.

Si può pertanto estrarre da x_k una successione x_{k_h} tale che $x_{k_h} \rightarrow x_1$ e si può concludere utilizzando le stesse argomentazioni del teorema precedente. $\square$

DEFINIZIONE 2.6. Sia $f : A \longrightarrow \mathbb{R}^m$, $A \subset \mathbb{R}^n$; f si dice uniformemente continua in A se $\forall \varepsilon > 0 \exists \delta(\varepsilon) > 0$ tale che se $x, y \in A$ e $\|x - y\| < \delta(\varepsilon)$, si ha

$$\|f(x) - f(y)\| < \varepsilon.$$

TEOREMA 2.4. - **Heine-Cantor** - Sia $f : A \longrightarrow \mathbb{R}^m$, $A \subset \mathbb{R}^n$; se f è una funzione continua su A ed A è un insieme compatto allora f è uniformemente continua su A .

COROLLARIO 2.1. Sia $f : \mathbb{R}^n \rightarrow \mathbb{R}^m$, lineare; allora

- f è uniformemente continua su $\mathbb{R}^n$;
- f trasforma insiemi limitati di $\mathbb{R}^n$ in insiemi limitati di $\mathbb{R}^m$.

3. Differenziabilità e Derivabilità

DEFINIZIONE 2.7. - Sia $f : A \longrightarrow \mathbb{R}^m$, $A \subset \mathbb{R}^n$, A aperto, $x_0 \in A$; diciamo che f è differenziabile in x_0 se esiste una applicazione lineare

$$L : \mathbb{R}^n \longrightarrow \mathbb{R}^m$$

tale che

$$\lim_{h \rightarrow 0} \frac{\|f(x_0 + h) - f(x_0) - L(h)\|}{\|h\|} = 0$$

L'applicazione lineare L si chiama differenziale di f in x_0 e si indica solitamente con $df(x_0)$.

La matrice che la rappresenta si chiama matrice jacobiana di f in x_0 e verrà indicata con $\nabla f(x_0)$. Si ha perciò

$$L(h) = df(x_0)(h) = \nabla f(x_0)h$$

Quando $m = 1$, $\nabla f(x_0)$ si riduce ad un vettore di $\mathbb{R}^n$ e si indica col nome di gradiente di f in x_0 ; faremo uso del nome gradiente anche se $m > 1$.

Osserviamo infine che $\nabla f : A \longrightarrow \mathbb{R}^{m \times n} = \mathcal{M}^{m \times n}$

Sia $f : A \longrightarrow \mathbb{R}^m$, posto

$$\omega(h) = \frac{\|f(x_0 + h) - f(x_0) - df(x_0)(h)\|}{\|h\|}$$

per la definizione di differenziabilità si ha

$$\lim_{h \rightarrow 0} \omega(h) = 0$$

per cui

$$f(x_0 + h) - f(x_0) - df(x_0)(h) = \|h\|\omega(h)$$

Se viceversa vale l'uguaglianza precedente si può verificare che f è differenziabile.

Naturalmente

$$\|f(x_0 + h) - f(x_0)\| \leq [\omega(h) + \|df(x_0)\|] \|h\|$$

e si ha

$$\lim_{h \rightarrow 0} \frac{f(x_0 + h) - f(x_0)}{\|h\|} = df(x_0)$$

per cui

Ogni funzione differenziabile è continua.

DEFINIZIONE 2.8. Sia $f : A \longrightarrow \mathbb{R}$, $A \subset \mathbb{R}^n$, A aperto, $x_0 \in A$; diciamo che f è parzialmente derivabile in x_0 rispetto alla variabile x_i se

$$\lim_{t \rightarrow 0} \frac{f(x_0 + te_i) - f(x_0)}{t}$$

esiste finito.

In tal caso denotiamo il valore di tale limite con il simbolo

$$\frac{\partial f}{\partial x_i}(x_0) \quad \text{oppure con} \quad f_{x_i}(x_0)$$

e lo chiamiamo derivata parziale di f rispetto ad x_i calcolata in x_0 .

(Osserviamo che $te_i = (0, 0, \dots, t, \dots, 0, 0)$).

Se $y \in \mathbb{R}^n$, diciamo che f è derivabile in x_0 rispetto al vettore y se

$$\lim_{t \rightarrow 0^+} \frac{f(x_0 + ty) - f(x_0)}{t}$$

esiste finito.

In tal caso denotiamo il valore di tale limite con $f'(x_0, y)$.

E' facile vedere che f è derivabile rispetto alla i -esima variabile se e solo se $f'(x_0, e_i)$ ed $f'(x_0, -e_i)$ esistono e

$$f'(x_0, e_i) = -f'(x_0, -e_i)$$

In tal caso si ha

$$f'(x_0, e_i) = -f'(x_0, -e_i) = f_{x_i}(x_0).$$

TEOREMA 2.5. Sia $f : A \longrightarrow \mathbb{R}$, $x_0 \in A \subset \mathbb{R}^n$, A aperto e supponiamo che f sia differenziabile in x_0 ; allora f è derivabile in x_0 lungo ogni direzione e si ha

$$f'(x_0, y) = df(x_0)(y) = \langle \nabla f(x_0), y \rangle$$

Se ne deduce in particolare, scegliendo $y = e_i$ ed $y = -e_i$ che f è derivabile in x_0 rispetto ad x_i e che

$$(\nabla f(x_0))_i = f_{x_i}(x_0).$$

DIMOSTRAZIONE. Dal momento che f è differenziabile,

$$\frac{|f(x_0 + h) - f(x_0) - \langle \nabla f(x_0), h \rangle|}{\|h\|} = \omega(h)$$

con $\lim_{h \rightarrow 0} \omega(h) = 0$.

Per $h = ty$ con $t > 0$ si ha

$$|f(x_0 + ty) - f(x_0) - \langle \nabla f(x_0), ty \rangle| = t\|y\|\omega(ty)$$

e

$$\left| \frac{f(x_0 + ty) - f(x_0)}{t} - \langle \nabla f(x_0), y \rangle \right| = \|y\|\omega(ty)$$

Quindi

$$f'(x_0, y) = \langle \nabla f(x_0), y \rangle = df(x_0)(y).$$

Osserviamo che, se f è differenziabile in x_0 , allora si ha

$$f'(x_0, y) = -f'(x_0, -y)$$

e pertanto

$$\lim_{t \rightarrow 0^+} \frac{f(x_0 + ty) - f(x_0)}{t} = \lim_{t \rightarrow 0^-} \frac{f(x_0 + ty) - f(x_0)}{t}$$

□

TEOREMA 2.6. *Sia $f : A \longrightarrow \mathbb{R}^m$, $x_0 \in A \subset \mathbb{R}^n$, A aperto e sia $f = (f_1, f_2, \dots, f_m)$ con $f_j : A \longrightarrow \mathbb{R}$, $j = 1, 2, \dots, m$.*

Allora f è differenziabile in x_0 se e solo se f_j è differenziabile in x_0 per ogni $j = 1, 2, \dots, m$.

Inoltre si ha

$$\nabla f(x_0) = \begin{pmatrix} \nabla f_1(x_0) \\ \nabla f_2(x_0) \\ \vdots \\ \nabla f_m(x_0) \end{pmatrix}$$

TEOREMA 2.7. -Derivazione delle Funzioni Composte - *Sia $f : A \longrightarrow \mathbb{R}^m$, $A \subset \mathbb{R}^n$, e sia $g : B \longrightarrow A$, $B \subset \mathbb{R}^p$.*

Possiamo allora considerare $f(g(\cdot)) : B \rightarrow \mathbb{R}^m$.

Siano $x_0 \in A$, $y_0 \in B$ tali che $g(y_0) = x_0$, f e g siano differenziabili in x_0 ed y_0 , rispettivamente.

Allora $f(g(\cdot))$ è differenziabile in y_0 e si ha

$$\nabla f(g(y_0)) = \nabla f(x_0) \cdot \nabla g(y_0) = \nabla f(g(y_0)) \cdot \nabla g(y_0),$$

essendo il prodotto tra matrici inteso righe per colonne.

DIMOSTRAZIONE. Si ha

$$f(x_0 + h) - f(x_0) - \nabla f(x_0)h = \|h\|\omega_1(h) \text{ con } \lim_{h \rightarrow 0} \omega_1(h) = 0$$

ed anche

$$g(y_0 + k) - g(y_0) - \nabla g(y_0)k = \|k\|\omega_2(k) \text{ con } \lim_{k \rightarrow 0} \omega_2(k) = 0$$

Pertanto posto

$$h(k) = g(y_0 + k) - g(y_0)$$

si ha

$$\lim_{k \rightarrow 0} h(k) = 0$$

e quindi

$$\begin{aligned}
 \frac{f(g(y_0 + k)) - f(g(y_0)) - \nabla f(x_0) \nabla g(y_0) k}{\|k\|} &= \\
 &= \frac{\nabla f(x_0) [g(y_0 + k) - g(y_0) - \nabla g(y_0) k] + \|h(k)\| \omega_1(h(k))}{\|k\|} = \\
 &= \nabla f(x_0) \omega_2(k) + \frac{\|h(k)\| \omega_1(h(k))}{\|k\|} \longrightarrow 0
 \end{aligned}$$

dal momento che

$$\|h(k)\| \leq \|k\|(\omega_2(k) + cost)$$

□

Esplicitiamo in caso semplice il teorema di derivazione delle funzioni composte con lo scopo di illustrarne l'uso.

Sia

$$\begin{aligned}
 f : \mathbb{R}^2 &\rightarrow \mathbb{R} \quad , \quad g : \mathbb{R}^2 \mapsto \mathbb{R}^2 \\
 (x, y) &\mapsto f(x, y) \quad , \quad (t, s) \mapsto (x(t, s), y(t, s))
 \end{aligned}$$

e consideriamo la funzione

$$\phi(t, s) = f(x(t, s), y(t, s)) = f(g(t, s))$$

Utilizzando il teorema possiamo affermare che

$$(\phi_t(t, s), \phi_s(t, s)) = \nabla \phi(t, s) = \nabla f(g(t, s)) \cdot \nabla g(t, s)$$

Ma

$$\nabla f(x, y) = (f_x(x, y), f_y(x, y)), \quad , \quad \nabla g(t, s) = \begin{pmatrix} x_t(t, s) & x_s(t, s) \\ y_t(t, s) & y_s(t, s) \end{pmatrix}$$

per cui

$$\begin{aligned}
 (\phi_t(t, s), \phi_s(t, s)) = \\
 (f_x(x, y), f_y(x, y)) \begin{pmatrix} x_t(t, s) & x_s(t, s) \\ y_t(t, s) & y_s(t, s) \end{pmatrix} = \\
 (f_x(x, y)x_t + f_y(x, y)y_t, f_x(x, y)x_s + f_y(x, y)y_s)
 \end{aligned}$$

Se $f : A \longrightarrow \mathbb{R}^m$, $A \subset \mathbb{R}^n$, A aperto, è differenziabile in A e se definiamo

$$\phi(t) = f(x + th)$$

si ha che ϕ è derivabile per i valori di t tali che $x + th \in A$ (cioè almeno in un intorno $(-\delta, \delta)$ di 0) e si ha

$$\phi'(t) = \nabla f(x + th)h$$

Nel caso in cui f assuma valori reali si ha

$$\phi'(t) = \langle \nabla f(x + th), h \rangle$$

Il teorema di Lagrange applicato alla funzione φ appena introdotta permette di affermare che

Se f assume valori reali ed è differenziabile in A ; allora, allora

$$f(x + h) - f(x) = \langle \nabla f(x + \tau h), h \rangle \quad \tau \in (0, 1)$$

di conseguenza

$$|f(x + h) - f(x)| \leq \|h\| \sup_{t \in (0,1)} \|\nabla f(x + th)\|$$

Tenendo conto che se f assume valori in $\mathbb{R}^m$ allora $f = (f_1, \dots, f_m)$, con f_j a valori reali, si può concludere che

Se $f : A \longrightarrow \mathbb{R}^m$, $A \subset \mathbb{R}^n$, è differenziabile in A ; allora,

$$\|f(x+h) - f(x)\| \leq \|h\| \sup_{t \in (0,1)} \|\nabla f(x+th)\|_0.$$

infatti poichè esiste $p_h \in \mathbb{R}^m$, di norma 1, tale che

$$\begin{aligned}
 (2.1) \quad \|f(x+h) - f(x)\| &= \langle p_h, f(x+h) - f(x) \rangle = \\
 &= \langle p_h, \nabla f(x+\tau h)h \rangle \leq \\
 &\leq \|h\| \sup_{t \in (0,1)} \|\nabla f(x+th)\|_0
 \end{aligned}$$

Quando la funzione f assume valori in $\mathbb{R}^m$, $m > 1$, il precedente risultato può non essere vero. Sia infatti $f : \mathbb{R} \longrightarrow \mathbb{R}^2$ definita da $f(t) = (\cos t, \sin t)$; si ha

$$(0, 0) = f(2\pi) - f(0) \neq 2\pi(-\sin t, \cos t) = 2\pi \nabla f(t) \quad \forall t \in (0, 2\pi)$$

Poichè la definizione non è di immediata verifica, è utile avere condizioni sufficienti che assicurino la differenziabilità di una funzione.

Se $f : A \longrightarrow \mathbb{R}$, $A \subset \mathbb{R}^n$, A aperto, ammette derivate parziali prime continue in A (indicheremo questo dicendo che $f \in C^1$); allora f è differenziabile in A .

Infatti: se ad esempio consideriamo il caso $n = 2$, indichiamo con (x, y) un punto di A e supponiamo che le derivate parziali prime f_x ed f_y siano continue in A ; avremo

$$\begin{aligned}
 & \frac{f(x+h, y+k) - f(x, y) - f_x(x, y)h - f_y(x, y)k}{(h^2 + k^2)} = \\
 &= \frac{f(x+h, y+k) - f(x+h, y) + f(x+h, y) - f(x, y) - f_x(x, y)h - f_y(x, y)k}{\sqrt{h^2 + k^2}} = \\
 & \quad \frac{-f(x, y) - f_x(x, y)h - f_y(x, y)k}{\sqrt{h^2 + k^2}} =
 \end{aligned}$$

per il teorema di Lagrange applicato alle funzioni $f(x+h, \cdot)$ e $f(\cdot, y)$

$$= \frac{(f_x(\xi, y) - f_x(x, y))h + (f_y(x+h, \eta) - f_y(x, y))k}{(h^2 + k^2)}$$

con $|x - \xi| < h$ e $|y - \eta| < k$.

Pertanto, osservando che

$$\left| \frac{h}{h^2 + k^2} \right| \leq 1 \quad \text{e} \quad \left| \frac{k}{h^2 + k^2} \right| \leq 1$$

per la continuità di f_x ed f_y l'ultimo membro tende a 0 quando $(h, k) \rightarrow (0, 0)$.

Possiamo affermare in maniera simile che

Se $f : A \rightarrow \mathbb{R}^m$, $A \subset \mathbb{R}^n$, A aperto, ammette derivate parziali prime continue in A (cioè se ognuna delle m componenti f_j è di classe C^1 : $f_j \in C^1$); allora f è differenziabile in A .

Se $f : A \rightarrow \mathbb{R}$, $A \subset \mathbb{R}^n$, A aperto, e chiamiamo derivata parziale seconda di f rispetto alle variabili x_i ed x_j , calcolata in x , e scriviamo $f_{x_i x_j}(x)$ la derivata rispetto a x_j della Funzione f_{x_i} , calcolata in x .

Nel caso in cui $n = 2$ e le variabili in A si indichino con (x, y) possiamo calcolare 4 derivate parziali seconde:

$$f_{xx}, f_{xy}, f_{yx}, f_{yy}$$

Si può dimostrare che

TEOREMA 2.8. -Schwartz - Sia $f : A \longrightarrow \mathbb{R}$, $A \subset \mathbb{R}^n$, A aperto, e supponiamo che f sia parzialmente derivabile due volte in A e che almeno una tra f_{xy} e f_{yx} sia continua; allora

$$f_{xy}(x, y) = f_{yx}(x, y)$$

Infatti se ad esempio supponiamo che f_{xy} sia continua, posto

$$\omega(h, k) = \frac{f(x + h, y + k) - f(x + h, y) - f(x, y + k) + f(x, y)}{hk}$$

si ha

$$f_{xy}(x) = \lim_{k \rightarrow 0} \lim_{h \rightarrow 0} \omega(h, k)$$

e

$$f_{yx}(x) = \lim_{h \rightarrow 0} \lim_{k \rightarrow 0} \omega(h, k)$$

Pertanto se proviamo, che

$$\lim_{(h,k) \rightarrow (0,0)} \omega(h, k)$$

esiste finito, avremo che

$$\lim_{(h,k) \rightarrow (0,0)} \omega(h,k) = \lim_{k \rightarrow 0} \lim_{h \rightarrow 0} \omega(h,k) = \lim_{h \rightarrow 0} \lim_{k \rightarrow 0} \omega(h,k)$$

e l'uguaglianza delle due derivate seconde

Applicando il teorema di Lagrange alla funzione

$$h \mapsto f(x+h, y+k) - f(x+h, y)$$

si ha

$$\omega(h,k) = \frac{f_x(x+\xi, y+k) - f_x(x+\xi, y)}{k}, \quad -h < \xi < h$$

ed applicando ancora Lagrange alla funzione

$$k \longrightarrow f_x(x+\xi, y+k)$$

si ottiene

$$\omega(h,k) = f_{xy}(x+\xi, y+\eta), \quad -k < \eta < k$$

Ora

$$\lim_{(h,k) \rightarrow (0,0)} (\xi, \eta) = (0, 0)$$

e pertanto, per la continuità di f_{xy}

$$\lim_{(h,k) \rightarrow (0,0)} \omega(h,k) = f_{xy}(x,y).$$

In generale possiamo enunciare il seguente teorema

TEOREMA 2.9. *Sia $f : A \longrightarrow \mathbb{R}^m$, $A \subset \mathbb{R}^n$, A aperto, e supponiamo che f sia parzialmente derivabile due volte in A ; allora*

$$f_{x_i x_j}(x) = f_{x_j x_i}(x)$$

per tutti gli $x \in A$ ove almeno una tra $f_{x_i x_j}$ e $f_{x_j x_i}$ è continua.

Le derivate parziali di ordine superiore si definiscono in maniera del tutto simile.

Le derivate parziali seconde caratterizzano il gradiente della funzione ∇f infatti, se $f : A \longrightarrow \mathbb{R}$, $A \subset \mathbb{R}^n$, A aperto, è differenziabile in A , possiamo considerare la funzione

$$\nabla f : A \longrightarrow \mathbb{R}^n$$

Se ∇f è a sua volta differenziabile (ricordiamo che basta che le derivate parziali prime di ∇f siano continue), possiamo considerare $\nabla(\nabla f)(x)$ e si vede che

$$\nabla(\nabla f)(x) = \begin{pmatrix} \nabla f_{x_1}(x) \\ \cdots \\ \nabla f_{x_n}(x) \end{pmatrix} = \begin{pmatrix} f_{x_1 x_1}(x) & \cdots & f_{x_1 x_n}(x) \\ \cdots & \cdots & \cdots \\ f_{x_n x_1}(x) & \cdots & f_{x_n x_n}(x) \end{pmatrix}$$

La matrice $\nabla(\nabla f)(x)$ si indica solitamente con $Hf(x)$ e si chiama matrice Hessiana di f in x .

La funzione quadratica $g(h) = \langle h, Hf(x)h \rangle$ viene di solito indicata con il nome di forma quadratica hessiana di f in x .

Qualora f ammetta derivate parziali seconde continue in A ($f \in \mathcal{C}^2(A)$), per il teorema di Schwarz, la matrice $Hf(x)$ è simmetrica.

Per fissare le idee ricordiamo che nel caso di una funzione di due variabili a valori reali

$$Hf(x, y) = \nabla(\nabla f)(x, y) = \begin{pmatrix} \nabla f_x(x, y) \\ \nabla f_y(x, y) \end{pmatrix} = \begin{pmatrix} f_{xx}(x, y) & f_{xy}(x, y) \\ f_{yx}(x, y) & f_{yy}(x, y) \end{pmatrix}$$



4. Formula di Taylor

Consideriamo una funzione $f : A \longrightarrow \mathbb{R}$, $A \subset \mathbb{R}^n$, A aperto, $x_0 \in A$, e sia $h \in S(0, r)$ dove $r > 0$ è scelto in modo che $x_0 + S(0, r) \subset A$.

Definiamo $\phi : (-1, 1) \longrightarrow \mathbb{R}$ mediante la

$$\phi(t) = f(x_0 + th)$$

se supponiamo $f \in \mathcal{C}^k$ (cioè se è derivabile k volte in A), avremo che ϕ è derivabile k volte in $(-1, 1)$ e si ha

$$\phi'(t) = df(x_0 + th)h = \langle h, \nabla f(x_0 + th) \rangle$$

$$\begin{aligned}
 \phi''(t) &= \frac{d}{dt} \sum_{i=1}^n h_i f_{x_i}(x_0 + th) = \\
 &= \sum_{i=1}^n h_i \langle \nabla f_{x_i}(x_0 + th), h \rangle = \langle h, Hf(x_0 + th)h \rangle
 \end{aligned}$$

Possiamo pertanto ottenere una formula di Taylor anche per funzioni di più variabili, sviluppando la funzione φ . Ci limitiamo al secondo ordine in quanto è l'unico di cui abbiamo necessità ed in ogni caso è l'ultimo che possa essere enunciato senza eccessive difficoltà formali.

TEOREMA 2.10. *Se $f : A \longrightarrow \mathbb{R}$, $A \subset \mathbb{R}^n$, A aperto, $x \in A$ e se $f \in C^2(A)$, (e quindi è differenziabile due volte).*

Allora per h abbastanza piccolo si ha

•

$$f(x+h) = f(x) + \langle h, \nabla f(x) \rangle + \langle h, Hf(x+\xi h)h \rangle/2, \quad \xi \in (0,1)$$

(formula di Taylor con il resto di Lagrange)

•

$$f(x+h) = f(x) + \langle h, \nabla f(x) \rangle + \langle h, Hf(x)h \rangle/2 + \|h\|^2 \omega(h)$$

con $\lim_{h \rightarrow 0} \omega(h) = 0$ e $\omega(0) = 0$

(formula di Taylor con resto di Peano).

DIMOSTRAZIONE.

Applicando a φ la formula di McLaurin otteniamo

$$\phi(1) = \phi(0) + \phi'(0) + \phi''(\xi)/2 \quad 0 < \xi < 1$$

da cui tenuto conto che

$$\phi'(t) = \langle h, \nabla f(x + th) \rangle$$

$$\phi''(t) = \langle h, Hf(x + th)h \rangle$$

si ricava la prima affermazione

Inoltre

$$f(x + h) = f(x) + \langle \nabla f(x), h \rangle + \langle h, Hf(x)h \rangle / 2 + \|h\|^2 \omega(h)$$

non appena si definisca

$$\omega(h) = \frac{\langle h, (Hf(x + \xi_h h) - Hf(x))h \rangle}{2\|h\|^2}$$

dove $\lim_{h \rightarrow 0} \omega(h) = 0$ in quanto Hf è continuo e

$$|\omega(h)| \leq \|Hf(x + \xi_h h) - Hf(x)\|/2, \quad \|\xi_h h\| \leq \|h\|.$$

□

5. Massimi e Minimi Relativi

DEFINIZIONE 2.9. Diciamo che x è un punto di minimo (massimo) relativo per f se esiste una sfera $S(x, r)$, $r > 0$, tale che

$$f(y) \geq f(x) \quad (f(y) \leq f(x)) \quad \forall y \in S(x, r) \cap A$$

TEOREMA 2.11. Se x è un punto di minimo (massimo) relativo per f interno al suo dominio, allora

- se f è differenziabile in x si ha $\nabla f(x) = 0$;
- se f ammette derivate seconde continue in x , $Hf(x)$ è semidefinita positiva (negativa).

DIMOSTRAZIONE. Basta osservare che $\varphi(t) = f(x + th)$ ammette un punto di minimo relativo in 0 e che $\forall h \in S(0, r)$

$$0 = \phi'(0) = \langle \nabla f(x), h \rangle$$

ed anche

$$0 \leq \phi''(0) = \langle h, Hf(x)h \rangle$$

La prima condizione assicura che $\nabla f(x) = 0$, mentre la seconda è, per definizione, la semidefinitezza di $Hf(x)$. □

Se $\nabla f(x) = 0$ e $Hf(x)$ è una forma quadratica non definita, allora x non è né punto di massimo relativo, né punto di minimo relativo per f ; un punto siffatto viene solitamente indicato con il nome di 'punto sella'.

TEOREMA 2.12. Se $f \in \mathcal{C}^2(A)$,

- $\nabla f(x) = 0$
- $Hf(x)$ è *definita positiva (negativa)*

allora x è punto di minimo (massimo) relativo per f .

DIMOSTRAZIONE. Si ha

$$f(x+h) - f(x) = \langle h, Hf(x)h \rangle / 2 + \|h\|^2 \omega(h)$$

con $\lim_{h \rightarrow 0} \omega(h) = 0 = \omega(0)$.

Se ne deduce che

$$\begin{aligned} \frac{f(x+h) - f(x)}{\|h\|^2} &= \frac{1}{2} \left\langle \frac{h}{\|h\|}, Hf(x) \frac{h}{\|h\|} \right\rangle + \omega(h) \geq \\ &\geq \min \left\{ \frac{1}{2} \langle u, Hf(x)u \rangle : \|u\| = 1 \right\} + \omega(h) = \frac{M}{2} + \omega(h) \end{aligned}$$

dove

$$M = \min\{\langle u, Hf(x)u \rangle : \|u\| = 1\} = \langle u_0, Hf(x)u_0 \rangle > 0$$

in quanto $\|u_0\| = 1$

(Il minimo esiste per il teorema di Weierstraßed $M > 0$ perché $Hf(x)$ è definita positiva e $u_0 \neq 0$.)

Pertanto, per il teorema della permanenza del segno, si può scegliere $\rho > 0$ in modo che, se $h \in S(0, \rho)$, si abbia

$$\frac{f(x+h) - f(x)}{\|h\|^2} > 0$$

e la tesi. □

6. Convessità

DEFINIZIONE 2.10. Sia $f : A \longrightarrow \mathbb{R}$, $A \subset \mathbb{R}^n$ convesso; diciamo che f è convessa se

$$f(\lambda x + (1 - \lambda)y) \leq \lambda f(x) + (1 - \lambda)f(y) \quad \forall x, y \in A, \forall \lambda \in (0, 1)$$

Inoltre f si dice strettamente convessa se vale la disuguaglianza stretta.

Si prova facilmente che Se f è convessa allora

$$L_{\alpha}^{-} = \{(x \in \mathbb{R}^n : f(x) \leq \alpha\}$$

è un insieme convesso.

Con qualche difficoltà in più si prova che una funzione f convessa su un aperto A è continua.

Inoltre applicando i risultati noti per le funzioni convesse di una variabile alla funzione

$$\varphi(t) = f(x + ty)$$

possiamo dimostrare che

- Se f è convessa su un aperto A , $f'(x, y)$ esiste $\forall x \in A, \forall y \in \mathbb{R}^n$.
- se $f \in \mathcal{C}^2(A)$, allora sono fatti equivalenti:
 - f è convessa

$$f(y) \geq f(x) + \langle \nabla f(x), y - x \rangle \quad \forall x, y \in A$$

- $Hf(x)$ è semidefinita positiva.
- le seguenti condizioni sono ciascuna sufficiente per la successiva:
 - $Hf(x)$ è definita positiva $\forall x \in A$;
 - $f(y) > f(x) + \langle \nabla f(x), y - x \rangle \quad \forall x, y \in A, y \neq x$;
 - f è strettamente convessa.

7. Funzioni Implicite

Se

$$f : A \longrightarrow \mathbb{R} \quad A \subset \mathbb{R}^2$$

è una funzione reale di due variabili reali possiamo considerare l'insieme definito in $\mathbb{R}^2$ da

$$G = \{(x, y) \in A : f(x, y) = 0\}$$

È naturale, per studiare tale insieme, cercare una funzione ϕ il cui grafico coincida localmente con G .

Ciò è equivalente a risolvere rispetto ad y l'equazione $f(x, y) = 0$, ed è il procedimento che si segue quando, per studiare il luogo dei punti del piano in cui

$$x^2 + y^2 = 1$$

si ricava, ad esempio,

$$y = \sqrt{1 - x^2} \quad \text{oppure} \quad y = -\sqrt{1 - x^2}$$

Nel caso in cui non sia facile esplicitare una delle due variabili in funzione della seconda, siamo interessati a sapere se è possibile definire una delle due variabili in funzione dell'altra e a studiare qualche proprietà della funzione che evidentemente non è possibile scrivere esplicitamente in termini di funzioni elementari.

TEOREMA 2.13. - Dini - Sia $A = (x_0 - a, x_0 + a) \times (y_0 - b, y_0 + b)$, $f : A \longrightarrow \mathbb{R}$ e supponiamo che le seguenti condizioni siano verificate:

$$f \in \mathcal{C}^1(A)$$

$$f(x_0, y_0) = 0$$

$$f_y(x_0, y_0) \neq 0$$

Allora esiste $\delta > 0$ ed esiste $\phi : (x_0 - \delta, x_0 + \delta) \longrightarrow (y_0 - b, y_0 + b)$ tale che

$$\phi(x_0) = y_0$$

$$f(x, y) = 0 \Leftrightarrow y = \phi(x), \quad \forall x \in (x_0 - \delta, x_0 + \delta)$$

ϕ è derivabile in $(x_0 - \delta, x_0 + \delta)$ e

$$\phi'(x) = -\frac{f_x(x, \phi(x))}{f_y(x, \phi(x))}.$$

DIMOSTRAZIONE. Sia $f_y(x_0, y_0) > 0$ e siano α, β scelti in modo che $0 < \alpha < a$, $0 < \beta < b$ e $f_y(x, y) > M > 0$ se $|x - x_0| \leq \alpha$ e $|y - y_0| \leq \beta$ (ciò è possibile per la continuità di f_y e per il teorema della permanenza del segno).

Ora, evidentemente, $f(x_0, \cdot)$ è una funzione strettamente crescente in $[y_0 - \beta, y_0 + \beta]$ e pertanto

$$f(x_0, y_0 - \beta) < f(x_0, y_0) = 0 < f(x_0, y_0 + \beta)$$

Ancora per il teorema della permanenza del segno, applicato ad $f(\cdot, y_0 - \beta)$ e ad $f(\cdot, y_0 + \beta)$, si può scegliere $0 < \delta \leq \alpha$, in modo che se

$$|x - x_0| < \delta$$

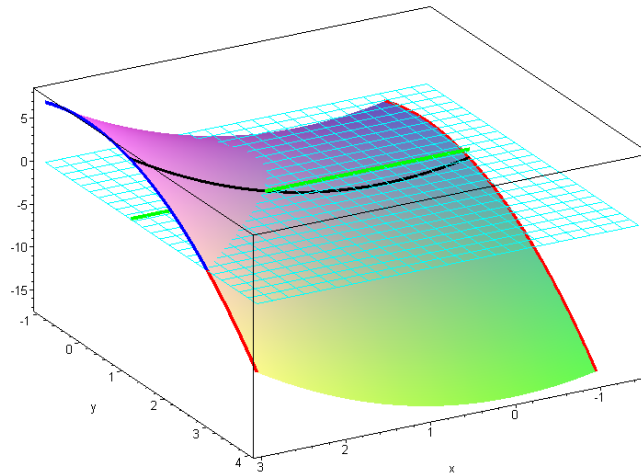


FIGURA 2.2. Il teorema delle funzioni implicite

si abbia

$$f(x, y_0 - \beta) < 0, \quad f(x, y_0 + \beta) > 0$$

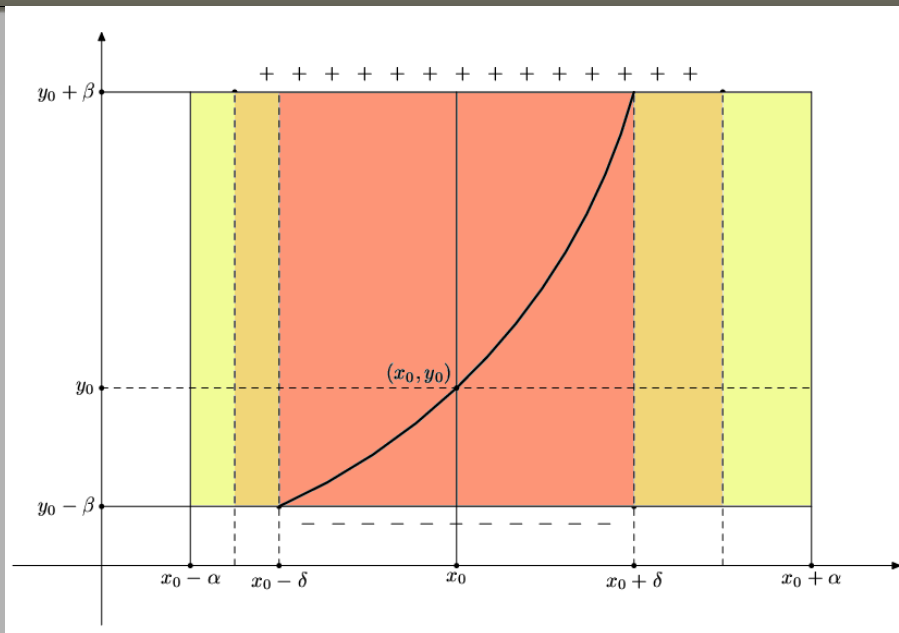


FIGURA 2.3. Il teorema delle funzioni implicite

Pertanto se $|x - x_0| < \delta$, $|y - y_0| < \beta$, si ha

$$f_y(x, y) > 0, \quad f(x, y_0 - \beta) < 0, \quad f(x, y_0 + \beta) > 0$$

e per ogni $x \in (x_0 - \delta, x_0 + \delta)$ si può affermare che esiste uno ed un solo valore $y \in (y_0 - \beta, y_0 + \beta)$ tale che $f(x, y) = 0$ (teorema degli zeri e stretta crescenza di $f(x, \cdot)$).

Possiamo pertanto definire $\phi : (x_0 - \delta, x_0 + \delta) \longrightarrow (y_0 - \beta, y_0 + \beta)$ mediante la $\phi(x) = y$.

Vediamo ora di provare che ϕ è continua e derivabile in $(x_0 - \delta, x_0 + \delta)$.

Siano $x, x + h \in (x_0 - \delta, x_0 + \delta)$, allora

$$f(x + h, \phi(x + h)) - f(x, \phi(x)) = 0$$

e pertanto, se definiamo $k(h) = \phi(x + h) - \phi(x)$, avremo

$$f(x + h, \phi(x) + k(h)) - f(x, \phi(x)) = 0$$

Per il teorema di Lagrange si ha

$$f_x(x + \tau h, \phi(x) + \tau k(h))h + f_y(x + \tau h, \phi(x) + \tau k(h))k(h) = 0$$

con $0 < \tau < 1$, $x + \tau h \in (x_0 - \delta, x_0 + \delta)$ e $\phi(x) + \tau k(h) \in (y_0 - \beta, y_0 + \beta)$, per cui

$$\phi(x + h) - \phi(x) = -h \frac{f_x(x + \tau h, \phi(x) + \tau k(h))}{f_y(x + \tau h, \phi(x) + \tau k(h))}$$

e dal momento che f_x ed f_y sono continue e

$$f_y \geq M > 0 \quad \text{se} \quad (x, y) \in [x_0 - \alpha, x_0 + \alpha] \times [y_0 - \beta, y_0 + \beta]$$

si ha

$$\lim_{h \rightarrow 0} \phi(x + h) - \phi(x) = \lim_{h \rightarrow 0} k(h) = 0$$

Inoltre

$$\frac{\phi(x + h) - \phi(x)}{h} = -\frac{f_x(x + \tau h, \phi(x) + \tau k(h))}{f_y(x + \tau h, \phi(x) + \tau k(h))}$$

e tenuto conto che $(h, k(h)) \rightarrow 0$ per $h \rightarrow 0$ si può concludere che ϕ è derivabile in x e

$$\phi'(x) = -\frac{f_x(x, \phi(x))}{f_y(x, \phi(x))}.$$

□

La dimostrazione fatta è evidentemente valida solo nel caso in cui $A \subset \mathbb{R}^2$ ed f assuma valori reali, ma l'enunciato, con le dovute modifiche, sussiste anche se $A \subset \mathbb{R}^n \times \mathbb{R}^m$ ed f assume valori in $\mathbb{R}^m$.

TEOREMA 2.14. - funzioni implicite - Sia $f : A \times B \longrightarrow \mathbb{R}^m$,

$$A = \{x \in \mathbb{R}^n : \|x - x_0\| < a\} \quad , \quad B = \{y \in \mathbb{R}^m : \|y - y_0\| < b\}$$

e supponiamo che:

- $f \in \mathcal{C}^1(A \times B)$
- $f(x_0, y_0) = 0$
- $\nabla_y f(x_0, y_0)$ sia invertibile.

Allora esistono $\rho, \delta > 0$ ed esiste una funzione

$$\phi : D \longrightarrow E$$

ove

$$D = \{x \in \mathbb{R}^n : \|x - x_0\| < \rho\} \quad \text{ed} \quad E = \{y \in \mathbb{R}^m : \|y - y_0\| < \delta\}$$

tali che

- $\phi(x_0) = y_0$
- $f(x, y) = 0 \Leftrightarrow y = \phi(x), \quad \forall x \in D$
- ϕ è differenziabile in D e si ha

$$\nabla \phi(x) = -[\nabla_y f(x, \phi(x))]^{-1} \nabla_x f(x, \phi(x)) \quad \forall x \in D.$$

8. Massimi e Minimi Vincolati Moltiplicatori di Lagrange

DEFINIZIONE 2.11. Sia $f : A \longrightarrow \mathbb{R}$ e sia $g : A \longrightarrow \mathbb{R}^m$, $A \subset \mathbb{R}^n$; diciamo che $x_0 \in A$ è un punto di massimo (o di minimo) relativo per f vincolato a g se $g(x_0) = 0$ e se esiste $\delta > 0$ tale che

$$f(x) \leq f(x_0) \quad (f(x) \geq f(x_0)) \quad \forall x \in \{x \in A : g(x) = 0\} \cap S(x_0, \delta).$$

A tale proposito possiamo provare il seguente risultato.

TEOREMA 2.15. - **dei moltiplicatori di Lagrange** - Siano $f : A \longrightarrow \mathbb{R}$ e $g : A \longrightarrow \mathbb{R}^m$, $A \subset \mathbb{R}^n$, $m < n$, A aperto, $f, g \in C^1(A)$; supponiamo inoltre che f abbia in $x_0 \in A$ un punto di minimo (o di massimo) relativo vincolato a g .

Allora esistono $\lambda \in \mathbb{R}^m$ e $\mu \in \mathbb{R}$ non contemporaneamente nulli e tali che

$$\mu \nabla f(x_0) + \sum_{i=1}^m \lambda_i \nabla g_i(x_0) = 0.$$

Inoltre, se $\nabla g(x_0)$ ha caratteristica massima ($= m$), allora $\mu \neq 0$ e si può supporre $\mu = 1$.

TEOREMA 2.16. *Siano $f : A \longrightarrow \mathbb{R}$ e $g : A \longrightarrow \mathbb{R}^m$, $A \subset \mathbb{R}^n$ aperto, $m < n$, $f, g \in \mathcal{C}^1(A)$, $x_0 \in A$; supponiamo che esista $\delta > 0$ tale che*

$$f(x_0) \leq f(x), \forall x \in \{x \in A : g_i(x) \leq 0, i = 1, \dots, m\} \cap S(x_0, \delta)$$

Allora, esistono $\mu \in \mathbb{R}$, $\lambda \in \mathbb{R}^m$ tali che

$$\mu \nabla f(x_0) + \sum_{i=1}^m \lambda_i \nabla g_i(x_0) = 0$$

essendo $\lambda_i = 0$ se $g_i(x_0) < 0$.

Se inoltre $\nabla g(x_0)$ ha caratteristica massima, si può supporre $\mu = 1$ e si ha

$$\lambda_i \geq 0 \quad \text{se } g_i(x_0) = 0.$$

TEOREMA 2.17. - Kuhn-Tucker - *Sia $A \subset \mathbb{R}^n$ aperto, convesso e siano $f, g_i : A \longrightarrow \mathbb{R}$, $i = 1, 2, \dots, m$ funzioni convesse; supponiamo che $f, g_i \in \mathcal{C}^1(A)$ e che $x_0 \in A$ sia scelto in modo che*

$$g_i(x_0) = 0 \quad \text{per } i = 1, 2, \dots, k < m$$

$$g_i(x_0) < 0 \quad \text{per } i = k + 1, \dots, m$$

Supponiamo inoltre che x_0 sia estrema per la funzione

$$F(x) = f(x) + \sum_{i=1}^k \lambda_i g_i(x)$$

essendo $\lambda_i \geq 0$ per $i = 1, 2, \dots, k$; allora

$$f(x_0) \leq f(x) \quad \forall x \in A \text{ tali che } g_i(x) \leq 0.$$

TEOREMA 2.18. Sia $f : A \longrightarrow \mathbb{R}$, $A \subset \mathbb{R}^n$ convesso, chiuso e limitato, f convessa e continua; allora il massimo di f in A è assunto anche in punti che sono sulla frontiera di A .

DIMOSTRAZIONE. Sia

$$f(x) = \max\{f(y) : y \in A\}$$

allora, se x è interno ad A , detti $y, z \in A$ gli estremi del segmento ottenuto intersecando A con una qualunque retta passante per x , si ha

$$x = \lambda y + (1 - \lambda)z$$

e

$$f(x) \leq \lambda f(y) + (1 - \lambda)f(z) \leq \max\{f(y), f(z)\}$$



Osservazione. Nel caso in cui A sia poliedrale, cioè se

$$A = \{x \in \mathbb{R}^n : g_i(x) \leq 0, g_i \text{ lineare}, i = 1, \dots, m\}$$

il massimo si può cercare solo tra i vertici della frontiera.



CAPITOLO 3

INTEGRAZIONE DELLE FUNZIONI DI PIU' VARIABILI.

*** La teoria dell'integrazione per le funzioni reali di più variabili deve tenere conto che si può integrare su sottoinsiemi di dimensione non necessariamente uguale al numero delle variabili.** Ad esempio se f dipende da 3 variabili reali avremo bisogno di definire cosa si intende per integrale di f su un sottoinsieme di $\mathbb{R}^3$, che possiamo intuitivamente definire come un solido (dimensione=3), una superficie (dimensione=2) o una linea (dimensione=1).

Ricordiamo esplicitamente che il concetto di dimensione non è semplice nè univocamente individuato: possiamo parlare di dimensione vettoriale, di dimensione topologica, di dimensione frattale; qui abbiamo fatto semplicemente ricorso ad un concetto intuitivo che si potrebbe precisare, ed in parte si preciserà, parlando di dimensione topologica.

Per semplificare le notazioni e per facilitare la comprensione descriveremo il caso delle funzioni di 3 variabili, essendo facile estendere i concetti al caso delle funzioni con più variabili, a prezzo di una certa complicazione delle notazioni.

1. Integrali Multipli

Cominciamo con il dare la definizione di integrale di una funzione limitata su una classe particolare di sottoinsiemi di $\mathbb{R}^3$ gli intervalli; successivamente estenderemo la definizione ad una più generale classe di insiemi.

1.1. Definizione di Integrale.

DEFINIZIONE 3.1. Siano I_1, I_2, I_3 intervalli chiusi e limitati, $I_i = [a_i, b_i]$, della retta reale. Diciamo che

$$R = I_1 \times I_2 \times I_3$$

è un intervallo chiuso e limitato in $\mathbb{R}^3$.

Nel seguito intenderemo riferirci sempre ad un intervallo chiuso e limitato, anche se queste due proprietà non saranno esplicitamente menzionate.

L'interno di R risulta essere

$$\text{int } R = (a_1, b_1) \times (a_2, b_2) \times (a_3, b_3)$$

DEFINIZIONE 3.2. Sia R un intervallo in $\mathbb{R}^3$; chiamiamo *partizione* di R il prodotto cartesiano $P = P_1 \times P_2 \times P_3$ dove P_i è una partizione dell'intervallo I_i

Denoteremo con $\mathcal{P}(R)$ l'insieme di tutte le partizioni dell'intervallo R .

Se $P \in \mathcal{P}(R)$, i punti di P dividono R in un numero N di intervalli chiusi la cui unione è R . Tali intervalli saranno indicati con

$$\{R_k : k = 1, 2, \dots, N\}$$

DEFINIZIONE 3.3. Sia R un intervallo in $\mathbb{R}^3$ e siano $P, Q \in \mathcal{P}(R)$; diciamo che P è una *partizione più fine* di Q e scriviamo $P < Q$ se $P \supset Q$.

In altre parole P è più fine di Q se e solo se ognuno degli intervalli in cui P suddivide R è contenuto in uno degli intervalli in cui Q suddivide R .

DEFINIZIONE 3.4. Sia R un intervallo in $\mathbb{R}^3$, definiamo *misura* di R il numero

$$\text{mis } R = (b_1 - a_1)(b_2 - a_2)(b_3 - a_3)$$

DEFINIZIONE 3.5. Sia R un intervallo e sia $P \in \mathcal{P}(R)$; siano $R_k, k = 1, 2, \dots, N$, gli intervalli in cui la partizione P suddivide R .

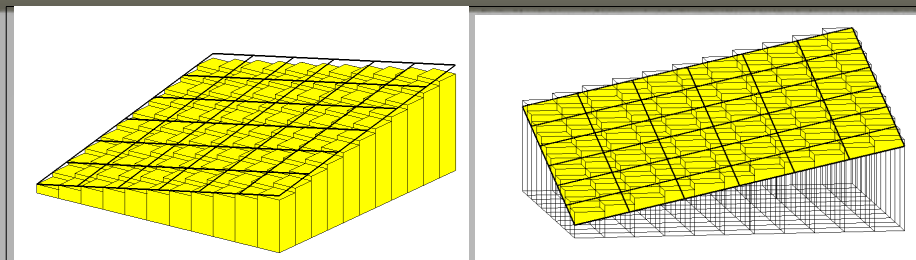


FIGURA 3.1.

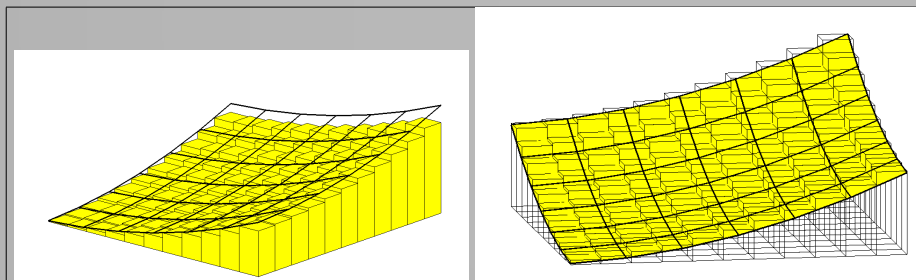


FIGURA 3.2.

Sia $f : R \longrightarrow \mathbb{R}$ una funzione limitata e supponiamo che

$$m \leq f(x) \leq M \quad \forall x \in R$$

Definiamo

$$m_k = \inf\{f(x) : x \in R_k\}$$

$$M_k = \sup\{f(x) : x \in R_k\}$$

definiamo inoltre

$$L(f, P) = \sum_{k=1}^N m_k \text{ mis } R_k$$

$$U(f, P) = \sum_{k=1}^N M_k \text{ mis } R_k$$

$$R(f, P, \Xi) = \sum_{k=1}^N f(\xi_k) \text{ mis } R_k \quad , \quad \xi_k \in R_k$$

essendo Ξ una funzione di scelta che assegna ad ogni intervallo R_k un punto ξ_k .

$L(f, P)$ ed $U(f, P)$ si dicono rispettivamente somme inferiori e somme superiori di f rispetto alla partizione P . $R(f, P, \Xi)$ si dice somma di Riemann di f rispetto alla partizione P e dipende, come è espressamente indicato, anche dalla scelta dei punti ξ_k in R_k .

Esattamente come nel caso di una funzione reale di una variabile reale si può provare che

TEOREMA 3.1. *Siano R un intervallo di $\mathbb{R}^3$, $f : R \longrightarrow \mathbb{R}$ limitata, allora, se $P, Q \in \mathcal{P}(R)$ e se $P < Q$*

$$m \text{ mis } R \leq L(f, Q) \leq L(f, P) \leq R(f, P, \Xi) \leq U(f, P) \leq U(f, Q) \leq M \text{ mis } R$$

e per ogni $P, Q \in \mathcal{P}(R)$

$$m \text{ mis } R \leq L(f, Q) \leq U(f, P) \leq M \text{ mis } R$$

DEFINIZIONE 3.6. *Sia R un intervallo in $\mathbb{R}^3$ e sia $f : R \longrightarrow \mathbb{R}$ una funzione limitata, definiamo*

$$\int_R^{\overline{}} f(x) dx = \inf \{ U(f, P) : P \in \mathcal{P}(R) \}$$

$$\int_R^{\underline{}} f(x) dx = \sup \{ L(f, P) : P \in \mathcal{P}(R) \}$$

essi si dicono rispettivamente, integrale superiore ed integrale inferiore della funzione f sull'intervallo R .

E' immediato provare che

$$m \text{ mis } R \leq \int_R f(x) dx \leq \int_R \bar{f}(x) dx \leq M \text{ mis } R$$

1.2. Condizioni di Integrabilità - Proprietà degli Integrali.

DEFINIZIONE 3.7. Sia R un intervallo in $\mathbb{R}^3$ e sia $f : R \longrightarrow \mathbb{R}$ una funzione limitata; diciamo che:

- f è integrabile se

$$\int_R f(x) dx = \int_R \bar{f}(x) dx$$

ed il valore comune ai due integrali superiore ed inferiore si chiama semplicemente integrale di f su R e si denota

$$\int_R f(x)dx$$

- f soddisfa la condizione di integrabilità se $\forall \varepsilon > 0 \exists P_\varepsilon \in \mathcal{P}(R)$ tale che $0 \leq U(f, P_\varepsilon) - L(f, P_\varepsilon) < \varepsilon$;
- f è integrabile secondo Cauchy-Riemann se $\exists I \in \mathbb{R}$ tale che $\forall \varepsilon > 0 \exists P_\varepsilon \in \mathcal{P}(R)$ tale che se $P \in \mathcal{P}(R)$, $P < P_\varepsilon$ si ha $|R(f, P, \Xi) - I| < \varepsilon \forall \Xi$; il valore I si chiama anche questa volta integrale di f su R .

Osserviamo che se la condizione di integrabilità è soddisfatta se e solo se comunque si scelga $P \in \mathcal{P}(R)$, $P < P_\varepsilon$ si ha

$$0 \leq U(f, P) - L(f, P) \leq U(f, P_\varepsilon) - L(f, P_\varepsilon) < \varepsilon.$$

Come per una sola variabile, si enuncia e si prova che

TEOREMA 3.2. *Sia R un intervallo in $\mathbb{R}^3$ e sia $f : R \longrightarrow \mathbb{R}$ una funzione limitata; sono fatti equivalenti:*

- f è integrabile
- f soddisfa la condizione di integrabilità
- f è integrabile secondo Cauchy-Riemann.

TEOREMA 3.3. *Sia R un intervallo in $\mathbb{R}^3$ e siano $f, g : R \longrightarrow \mathbb{R}$ funzioni limitate ed integrabili su R ; allora*

- $\forall \alpha, \beta > 0, \alpha f + \beta g$ è integrabile su R e

$$\int_R [\alpha f(x) + \beta g(x)] dx = \alpha \int_R f(x) dx + \beta \int_R g(x) dx$$

- fg è integrabile su R ;
- se S e T sono intervalli in $\mathbb{R}^3$ tali che $R = S \cup T$ e $\text{mis}(S \cap T) = 0$,

$$\int_R f(x) dx = \int_S f(x) dx + \int_T f(x) dx$$

- se $f \geq 0$

$$\int_R f(x)dx \geq 0$$

- se $f \geq g$

$$\int_R f(x)dx \geq \int_R g(x)dx$$

- se f è continua, $f \geq 0$,

$$\int_R f(x)dx = 0 \Rightarrow f \equiv 0$$

- $|f|$ è integrabile su R e

$$\left| \int_R f(x)dx \right| \leq \int_R |f(x)|dx$$

se S e T sono intervalli, $S \subset T \subset R$ e se $f \geq 0$,

$$\int_S f(x)dx \leq \int_T f(x)dx$$

TEOREMA 3.4. *Se $f : R \longrightarrow \mathbb{R}$, $R \subset \mathbb{R}^3$ intervallo, è continua, allora f è integrabile.*

1.3. Formule di Riduzione. L'integrale che abbiamo definito non può tuttavia essere calcolato, come per il caso delle funzioni di una variabile reale, facendo uso del concetto di primitiva in $\mathbb{R}^3$; il concetto di primitiva ed il teorema fondamentale del calcolo integrale trovano la loro naturale estensione nell'ambito delle forme differenziali e del teorema di Stokes, di cui parleremo più avanti.

Il calcolo di integrali multipli si può però ricondurre al calcolo di più integrali semplici mediante quelle che si chiamano formule di riduzione.

Se $A \subset \mathbb{R}^3$, la funzione

$$\chi_A : \mathbb{R}^3 \longrightarrow \mathbb{R}$$

definita da

$$\chi_A(x) = \begin{cases} 1 & x \in A \\ 0 & x \notin A \end{cases}$$

si chiama **funzione caratteristica di A** .

TEOREMA 3.5. *Sia R un intervallo in $\mathbb{R}^3$ e sia $f : R \longrightarrow \mathbb{R}$ integrabile.*

Allora si ha

$$\int_R f(x) dx = \int_{a_1}^{b_1} \int_{a_2}^{b_2} \int_{a_3}^{b_3} f(x_1, x_2, x_3) dx_3 dx_2 dx_1$$

ogniquale volta esiste il secondo membro.

DIMOSTRAZIONE. Se $P \in \mathcal{P}(R)$, $R_k = \times [a_{ki}, b_{ki}]$, allora si ha

$$m_k \chi_{R_k}(x) \leq f(x) \leq M_k \chi_{R_k}(x) \quad \forall x \in R_k$$

integrando n volte su $[a_{ki}, b_{ki}]$ e sommando su k si ottiene

$$\sum m_k \text{mis } R_k \leq \int_{a_1}^{b_1} \dots \int_{a_n}^{b_n} f(x_1, \dots, x_n) dx_n \dots dx_1 \leq \sum M_k \text{mis } R_k$$

Poiché f è integrabile, soddisfa il criterio di integrabilità, e si ha la tesi. $\square$

E' necessario estendere la nozione di integrabilità su insiemi che siano più generali di un intervallo in $\mathbb{R}^3$.

A questo scopo occorre precisare la classe dei sottoinsiemi di $\mathbb{R}^3$ sui quali è possibile integrare una funzione.

1.4. Misura di sottoinsiemi di $\mathbb{R}^3$.

DEFINIZIONE 3.8. Sia $A \subset \mathbb{R}^3$ un insieme limitato e sia R un intervallo che contiene A . Definiamo

$$\text{mis}^-(A) = \int_R \chi_A(x) dx \quad , \quad \text{mis}^+(A) = \int_R \bar{\chi}_A(x) dx$$

$\text{mis}^-(A)$ e $\text{mis}^+(A)$ si dicono, rispettivamente *misura interna* e *misura esterna* di A .

Diciamo che A è un sottoinsieme misurabile di $\mathbb{R}^3$ se $\text{mis}^-(A) = \text{mis}^+(A)$; in tal caso definiamo $\text{mis}(A)$, misura di A , il loro comune valore.

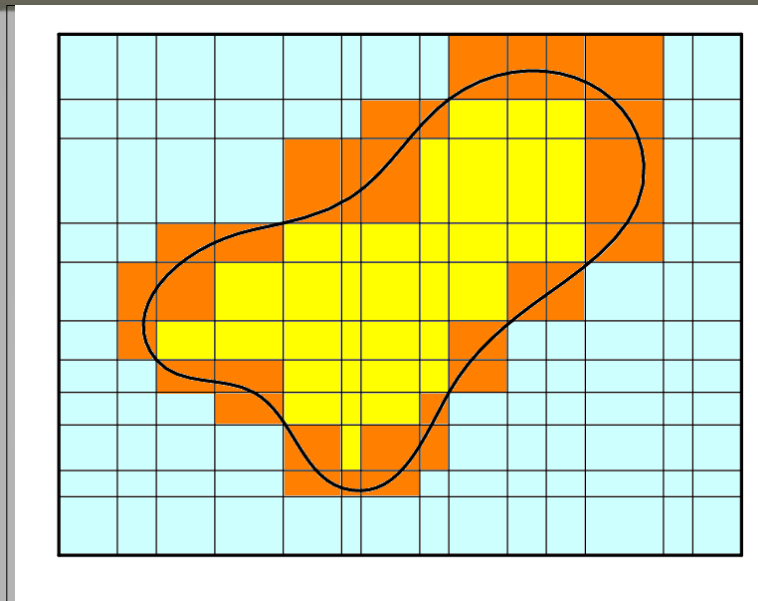


FIGURA 3.3.

E' immediato verificare che la precedente definizione non dipende dalla scelta dell' intervallo R tra tutti quelli che contengono A .

Si può inoltre verificare che

$$\text{mis}^-(A) = \sup_{P \in \mathcal{P}(R)} L(\chi_A, P) \quad \text{mis}^+(A) = \inf_{P \in \mathcal{P}(R)} U(\chi_A, P)$$

In altre parole

- $\text{mis}^-(A)$ è l'estremo superiore delle somme delle misure degli intervalli chiusi che sono contenute in A
- $\text{mis}^+(A)$ è l'estremo inferiore delle somme delle misure degli intervalli chiusi che contengono punti di A

È intuitivamente evidente, si veda la figura 3.3, anche se non immediato da dimostrare che

TEOREMA 3.6. *Sia $A \subset \mathbb{R}^3$ un sottoinsieme limitato, allora*

$$\text{mis}^+(\partial A) = \text{mis}^+(A) - \text{mis}^-(A).$$

Inoltre A è misurabile se e solo se ∂A è misurabile ed ha misura nulla.

Osserviamo anche che

$$\text{mis } A = 0 \Leftrightarrow \forall \varepsilon > 0 \exists P_\varepsilon \in \mathcal{P}(R) : U'(\chi_A, P_\varepsilon) < \varepsilon$$

Inoltre, tenuto conto che, se $\text{mis } A = \text{mis } B = 0$ allora $\text{mis } A \cup B = 0$, dal precedente teorema e dal fatto che

$$\partial(A \cup B) \subset \partial A \cup \partial B \quad \partial(A \cap B) \subset \partial A \cup \partial B \quad \partial(A \setminus B) \subset \partial A \cup \partial B$$

si ottiene che, se A e B sono misurabili, allora $A \cup B$, $A \cap B$, $A \setminus B$ sono misurabili.

Infine, tenendo conto che

$$\chi_{A \cup B} = \chi_A + \chi_B - \chi_{A \cap B}$$

si ottiene

$$\text{mis } A \cup B = \text{mis } A + \text{mis } B - \text{mis } A \cap B$$

Abbiamo con ciò che, se $A, B \subset \mathbb{R}^3$ sono misurabili e disgiunti, e se $x \in \mathbb{R}^3$, si ha

- $\text{mis } A \geq 0$
- $\text{mis } A \cup B = \text{mis } A + \text{mis } B$
- $\text{mis}(x + A) = \text{mis } A$
- $\text{mis}(\times_{i=1}^n [0, 1]) = 1$

Si potrebbe anche vedere che tali proprietà sono, da sole, in grado di caratterizzare la misura sui sottoinsiemi di $\mathbb{R}^3$

TEOREMA 3.7. *Sia $R \subset \mathbb{R}^3$ un intervallo e sia $f : R \longrightarrow \mathbb{R}$ limitata; supponiamo inoltre f continua in $R \setminus D$, $\text{mis } D = 0$, allora f è integrabile in R .*

TEOREMA 3.8. *Sia $f : A \longrightarrow \mathbb{R}^m$ continua, $A \subset \mathbb{R}^3$ chiuso e limitato, allora*

$$\text{mis}(\text{gph}(f)) = 0.$$

COROLLARIO 3.1. *Siano $g, f : A \longrightarrow \mathbb{R}$, $A \subset \mathbb{R}^3$ chiuso e limitato; allora, se f e g sono continue*

$$\{(x, y) \in \mathbb{R}^{n+1} : g(x) \leq y \leq f(x)\}$$

è misurabile.

1.5. Integrazione su Domini Normali. Definiamo ora l'integrale di una funzione limitata su un insieme misurabile.

DEFINIZIONE 3.9. Sia $f : R \longrightarrow \mathbb{R}$, $A \subset R \subset \mathbb{R}^3$, A limitato e misurabile, R intervallo in $\mathbb{R}^3$; si definisce

$$\int_A f(x)dx = \int_R \chi_A(x)f(x)dx.$$

E' banale verificare che la definizione non dipende dalla scelta dell'intervallo R che contiene A .

Possiamo dare il seguente criterio di integrabilità.

TEOREMA 3.9. Sia $f : A \longrightarrow \mathbb{R}$, $A \subset \mathbb{R}^3$ chiuso, limitato e misurabile; f continua in $A \setminus D$, $\text{mis } D = 0$. Allora f è integrabile su A .

DEFINIZIONE 3.10. Diciamo che $A \subset \mathbb{R}^{n+1}$ è un dominio normale in $\mathbb{R}^{n+1}$ se esistono un insieme $D \subset \mathbb{R}^n$ chiuso e limitato, e due funzioni continue $g, h : D \longrightarrow \mathbb{R}$ tali che

$$A = \{(x, y) \in \mathbb{R}^n \times \mathbb{R} : x \in D, g(x) \leq y \leq h(x)\}$$

oppure se

$$A = \{(x, y) \in \mathbb{R}^n \times \mathbb{R} : a \leq y \leq b, x \in D_y\}$$

dove D_y è un insieme misurabile in $\mathbb{R}^n$.

e si può verificare che

Ogni dominio normale in $\mathbb{R}^{n+1}$ è un insieme misurabile.

Pertanto è lecito integrare funzioni continue, a meno di insiemi di misura nulla, su domini normali e si ha il seguente

TEOREMA 3.10. *Sia A un dominio normale in $\mathbb{R}^3$ e sia $f : A \longrightarrow \mathbb{R}$ una funzione continua in $A \setminus D$ con $\text{mis } D = 0$.*

Allora f è integrabile su A e si ha

$$\begin{aligned}
 (3.1) \quad \int_A f(x)dx &= \int_R \chi_A(x)f(x)dx = \\
 &= \int_{a_1}^{b_1} \int_{a_2}^{b_2} \int_{a_3}^{b_3} \chi_A(x_1, x_2, x_3, y)f(x_1, x_2, x_3, y)dydx_3dx_2dx_1 = \\
 &= \int_D \int_{g(x_1, x_2, x_3)}^{h(x_1, x_2, x_3)} f(x_1, x_2, \dots, x_n, y)dydx_3dx_2dx_1
 \end{aligned}$$

oppure

$$\begin{aligned}
 (3.2) \quad \int_A f(x)dx &= \int_R \chi_A(x)f(x)dx = \\
 &= \int_{a_1}^{b_1} \int_{a_2}^{b_2} \int_{a_3}^{b_3} \chi_A(x_1, x_2, x_3, y)f(x_1, x_2, x_3, y)dydx_3dx_2dx_1 = \\
 &= \int_a^b \int_{D_y} f(x_1, x_2, \dots, x_n, y)dx_3dx_2dx_1dy
 \end{aligned}$$

1.6. Trasformazione di coordinate in $\mathbb{R}^3$. È spesso utile, per tenere conto delle caratteristiche di un insieme, considerare un cambiamento di variabili in $\mathbb{R}^3$.

Per cambiamento di variabili intendiamo una applicazione

$$V : \mathbb{R}^3 \rightarrow \mathbb{R}^3$$

definita da

$$\mathbb{R}^3 \ni (t, s, r) \mapsto V(t, s, r) = (x(t, s, r), y(t, s, r), z(t, s, r)) \in \mathbb{R}^3$$

che risulti di classe $\mathcal{C}^1$ sia invertibile e sia tale che

$$\frac{\partial(x, y, z)}{\partial(t, s, r)} = \det \begin{pmatrix} x_t & y_t & z_t \\ x_s & y_s & z_s \\ x_r & y_r & z_r \end{pmatrix} \neq 0$$

Sono esempi di trasformazioni di coordinate

- Il cambiamento di variabili lineari

$$\begin{cases} x = a_1u + b_1v + c_1w \\ y = a_2u + b_2v + c_2w \\ z = a_3u + b_3v + c_3w \end{cases}, \quad u, v, w \in \mathbb{R}, \quad z \in \mathbb{R}$$

cioè

$$\begin{pmatrix} x \\ y \\ z \end{pmatrix} = \begin{pmatrix} a_1 & b_1 & c_1 \\ a_2 & b_2 & c_2 \\ a_3 & b_3 & c_3 \end{pmatrix} \begin{pmatrix} u \\ v \\ w \end{pmatrix}$$

- Le **coordinate cilindriche** definite da

$$\begin{cases} x = \rho \cos \theta \\ y = \rho \sin \theta \\ z = z \end{cases}, \quad \rho \in [0, +\infty), \quad \theta \in [0, 2\pi], \quad z \in \mathbb{R}$$

- Le **coordinate sferiche** definite da

$$\begin{cases} x &= \rho \cos \theta \cos \phi \\ y &= \rho \sin \theta \cos \phi \\ z &= \rho \sin \phi \end{cases}, \quad \rho \in [0, +\infty), \theta \in [0, 2\pi], \phi \in [-\pi/2, \pi/2]$$

Si verifica in tali casi che

- Per il cambiamento lineare

$$\frac{\partial(x, y, z)}{\partial(u, v, w)} = \det \begin{pmatrix} a_1 & b_1 & c_1 \\ a_2 & b_2 & c_2 \\ a_3 & b_3 & c_3 \end{pmatrix}$$

- Per le coordinate cilindriche

$$\frac{\partial(x, y, z)}{\partial(\rho, \theta, z)} = \rho$$

- Per le coordinate sferiche

$$\frac{\partial(x, y, z)}{\partial(\rho, \theta, \phi)} = \rho \cos \phi$$

TEOREMA 3.11. - Cambiamento di variabili per integrali multipli - Sia

$$\phi : B \longrightarrow \mathbb{R}^3$$

dove $B \subset \mathbb{R}^3$ aperto e $\phi \in \mathcal{C}^1(B)$.

Supponiamo che A sia un insieme misurabile con $\text{cl } A \subset B$, tale che ϕ è una funzione invertibile e $\nabla \phi$ è una matrice invertibile su $\text{int } A$;

allora se f è limitata su $\phi(A)$ e continua su $\text{int } \phi(A)$, si ha

$$\int_{\phi(A)} f(x) dx = \int_A f(\phi(x)) |\det(\nabla \phi(x))| dx.$$

1.7. Integrali Impropri in $\mathbb{R}^3$. Illustriamo ora per sommi capi il problema di definire l'integrale di una funzione non limitata su un insieme limitato o non limitato.

DEFINIZIONE 3.11. Sia $f : A \longrightarrow \overline{\mathbb{R}}_+$, $A \subset \mathbb{R}^3$, f limitata ed integrabile in ogni compatto misurabile $K \subset A$. Definiamo

$$\int_A f(x)dx = \sup \left\{ \int_K f(x)dx : K \subset A, K \text{ compatto e misurabile} \right\}$$

La definizione si può facilmente estendere a funzioni di segno qualunque, non appena si ricordi che $f = f_+ + f_-$.

Per il calcolo di $\int_A f(x)dx$ è opportuno dare la seguente definizione.

DEFINIZIONE 3.12. Sia $A \subset \mathbb{R}^3$ diciamo che K_i è una successione di domini invadenti A se

- K_i sono insiemi compatti, misurabili, $K_i \subset A$
- $K_{i+1} \supset K_i$
- $\forall K \subset A$, K compatto, misurabile, $\exists i$ tale che $K_i \supset K$.

TEOREMA 3.12. Sia $A \subset \mathbb{R}^3$ misurabile e sia $f : A \longrightarrow \overline{\mathbb{R}}_+$ una funzione integrabile in ogni insieme $K \subset A$, compatto e misurabile.

Allora se K_i è una successione di domini invadenti A , si ha

$$\int_A f(x)dx = \lim_i \int_{K_i} f(x)dx$$

DIMOSTRAZIONE. Si ha

$$\int_{K_i} f(x)dx \leq \int_A f(x)dx$$

e $\int_{K_i} f(x)dx$ è una successione crescente per cui

$$\lim_i \int_{K_i} f(x)dx = \sup \left\{ \int_{K_i} f(x)dx \right\} \leq \int_A f(x)dx$$

D'altra parte, dal momento che, $\forall K \subset A$ esiste $K_i \supset K$, si ha

$$\sup \left\{ \int_{K_i} f(x)dx \right\} \geq \sup \left\{ \int_K f(x)dx : K \subset A \right\} = \int_A f(x)dx$$

□

TEOREMA 3.13. Sia $f : A \longrightarrow \mathbb{R}$, $A \subset \mathbb{R}^2$ misurabile, chiuso e limitato; sia $x_0 \in A$, sia f continua in $A \setminus \{x_0\}$ e

$$\lim_{x \rightarrow x_0} f(x) = +\infty$$

Allora se

$$f(x) \leq \frac{H}{\|x - x_0\|^\alpha}, \quad H \geq 0, \quad \alpha < 2$$

f è integrabile in senso improprio su A .

Se invece

$$f(x) \geq \frac{H}{\|x - x_0\|^\alpha}, \quad H > 0, \quad \alpha \geq 2$$

e se A contiene un cono di vertice x_0 e ampiezza positiva, allora

$$\int_A f(x) dx = +\infty.$$

DIMOSTRAZIONE. Sia $A_k = \text{cl}(A \setminus S(x_0, 1/k))$, A_k è una successione di domini invadenti A ; sia $h \in \mathbb{N}$, si ha, se $k > h$

$$\int_{A_k} f(x)dx = \int_{A_h} f(x)dx - \int_{A_h \setminus A_k} f(x)dx$$

inoltre

$$\int_{A_h \setminus A_k} f(x)dx \leq \int_0^{2\pi} d\theta \int_{1/k}^{1/h} \frac{H}{\rho^\alpha} \rho d\rho$$

non appena si sia convenuto di indicare con ρ e θ le coordinate polari nel piano, centrate in x_0 .

Per quel che riguarda il secondo enunciato, detti θ_0 e θ_1 gli angoli che le semirette delimitanti il settore formano con l'asse x , si ha

$$\int_{A_h \setminus A_k} f(x)dx \geq \int_{\theta_0}^{\theta_1} d\theta \int_{1/k}^{1/h} \frac{H}{\rho^\alpha} \rho d\rho$$

□

In maniera analoga si può provare il seguente

TEOREMA 3.14. Sia $f : A \longrightarrow \mathbb{R}$, $A \subset \mathbb{R}^3$ non limitato; sia f continua.

Se

$$|f(x)| \leq \frac{H}{\|x\|^\alpha} \quad , \quad H \geq 0 \quad , \quad \alpha > 2$$

allora f è integrabile in senso improprio su A .

Se invece

$$f(x) \geq \frac{H}{\|x\|^\alpha} \quad , \quad H > 0 \quad , \quad \alpha \leq 2$$

e se A contiene un cono di ampiezza positiva, allora

$$\int_A f(x) dx = +\infty.$$

2. Integrali dipendenti da un parametro.

Passiamo infine a illustrare brevemente il comportamento di un integrale rispetto a parametri contenuti nella funzione da integrare.

Questo tipo di problematiche si incontra, ad esempio, quando si studiano le trasformazioni integrali (Fourier, Laplace) o nella definizione di funzioni notevoli (come, ad esempio, la funzione Γ).

TEOREMA 3.15. Sia $f : A \times I \longrightarrow \mathbb{R}$, $A \subset \mathbb{R}^3$ chiuso e limitato, $I = [a, b]$. Supponiamo $f \in \mathcal{C}^0(A \times I)$, allora $F : A \times I \times I \longrightarrow \mathbb{R}$ definita da

$$F(x, y, z) = \int_y^z f(x, t) dt$$

è continua in $A \times I \times I$; inoltre F_y ed F_z esistono e sono continue in $A \times I \times I$.

Se $\nabla_x f \in \mathcal{C}^0(A \times I \times I)$, allora F è differenziabile rispetto ad x ,

$$\nabla_x F(x, y, z) = \int_y^z \nabla_x f(x, t) dt$$

e quindi risulta $\nabla_x F$ è continuo in $A \times I \times I$ e $F \in \mathcal{C}^1(A \times I \times I)$.

TEOREMA 3.16. Sia $f : A \times I \longrightarrow \mathbb{R}$, $A \subset \mathbb{R}^3$ chiuso e limitato, $I = [a, +\infty)$, una funzione continua. Consideriamo

$$F(x) = \int_a^{+\infty} f(x, t) dt$$

Se esiste $\phi : I \longrightarrow \mathbb{R}$ tale che

$$|f(x, t)| \leq \phi(t) \quad \forall x \in A ; \quad \int_a^{+\infty} \phi(t) dt < +\infty$$

allora F è definita e continua in A .

Se inoltre $\nabla_x f$ esiste, è continuo in $A \times I$, e se esiste $\psi : I \longrightarrow \mathbb{R}$ tale che

$$\|\nabla_x f(x, t)\| \leq \psi(t) \quad \forall x \in A ; \quad \int_a^{+\infty} \psi(t) dt < +\infty$$

allora $F \in C^1(A)$ e

$$\nabla F(x) = \int_a^{+\infty} \nabla_x f(x, t) dt .$$



CAPITOLO 4

ARCHI E SUPERFICI NELLO SPAZIO.

In questa parte definiamo i concetti di arco, di superficie, di lunghezza di un arco e di area di una superficie nello spazio euclideo a tre dimensioni $\mathbb{R}^3$.

La generalizzazione dei concetti esposti al caso di $\mathbb{R}^n$, $n > 3$, è immediata, come del resto è ovvio che per ottenere una trattazione delle curve in $\mathbb{R}^2$ è sufficiente porre $z = 0$.

1. Linee ed integrali di linea

DEFINIZIONE 4.1. *Chiamiamo curva in $\mathbb{R}^3$ una funzione*

$$\gamma : [a, b] \longrightarrow \mathbb{R}^3, \quad \gamma(t) = (x(t), y(t), z(t)).$$

Chiamiamo traccia di γ , o più raramente supporto di γ , l'insieme

$$\Gamma = R(\gamma) = \{(x, y, z) \in \mathbb{R}^3 : \exists t \in [a, b], (x, y, z) = (x(t), y(t), z(t))\}$$

Indichiamo con $\dot{\gamma} = (\dot{x}(t), \dot{y}(t), \dot{z}(t))$ la derivata di γ .

Una curva γ si dice:

- semplice, se è iniettiva,

- chiusa, se $\gamma(a) = \gamma(b)$,
- regolare, se $\gamma \in C^1([a, b])$ e $\|\dot{\gamma}\| \neq 0$.

Osserviamo che la condizione $\dot{\gamma} \neq 0$ significa che le tre derivate $\dot{x}, \dot{y}, \dot{z}$ non sono mai contemporaneamente nulle ed è spesso espressa nella forma

$$\dot{x}^2 + \dot{y}^2 + \dot{z}^2 > 0$$

DEFINIZIONE 4.2. Sia γ una curva regolare in $\mathbb{R}^3$ definiamo versore tangente alla curva γ nel punto $(x(t), y(t), z(t))$ il versore

$$T_\gamma(t) = \frac{\dot{\gamma}(t)}{\|\dot{\gamma}(t)\|} = \left(\frac{\dot{x}(t)}{\|\dot{\gamma}(t)\|}, \frac{\dot{y}(t)}{\|\dot{\gamma}(t)\|}, \frac{\dot{z}(t)}{\|\dot{\gamma}(t)\|} \right)$$

dove

$$\|\dot{\gamma}(t)\| = \sqrt{\dot{x}^2(t) + \dot{y}^2(t) + \dot{z}^2(t)}.$$

1.1. Lunghezza di una Linea. Passiamo ora a definire la lunghezza di una curva γ .

Sia γ una curva regolare in $\mathbb{R}^3$.

Definiamo poligonale inscritta in Γ , associata alla partizione $P = \{t_0 < t_1 < t_2 < \dots t_n\}$, la spezzata poligonale $\Lambda(\Gamma, P)$ avente per vertici i punti $\gamma(t_i)$.

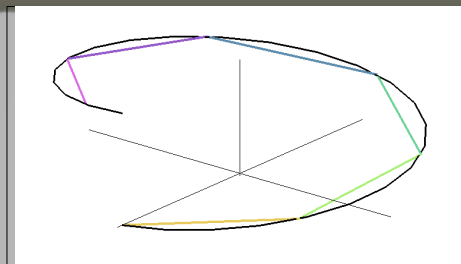


FIGURA 4.1.

Possiamo calcolare la lunghezza della poligonale $\Lambda(\Gamma, P)$ mediante la

$$\begin{aligned}
 \ell(\Lambda(\Gamma, P)) &= \sum_{i=1}^n \|\gamma(t_i) - \gamma(t_{i-1})\| = \\
 &= \sum_{i=1}^n \sqrt{(x(t_i) - x(t_{i-1}))^2 + (y(t_i) - y(t_{i-1}))^2 + (z(t_i) - z(t_{i-1}))^2}
 \end{aligned}$$

Usando il teorema di Lagrange, se $P \in \mathcal{P}(a, b)$, si ha

$$\ell(\Lambda(\Gamma, P)) = \sum_{i=1}^n \sqrt{x^2(t_i^1) + y^2(t_i^2) + z^2(t_i^3)}(t_i - t_{i-1})$$

essendo $t_i^1, t_i^2, t_i^3 \in (t_{i-1}, t_i)$, mentre d'altro canto, le somme di Riemann di $\|\dot{\gamma}\|$ sono date da

$$R(\|\dot{\gamma}\|, P, \Xi) = \sum_{i=1}^n \sqrt{\dot{x}^2(\tau_i) + \dot{y}^2(\tau_i) + \dot{z}^2(\tau_i)}(t_i - t_{i-1})$$

$$\begin{aligned} \ell(\Lambda(\Gamma, P)) &= R(\|\dot{\gamma}\|, P, \Xi) - R(\|\dot{\gamma}\|, P, \Xi) + \ell(\Lambda(\Gamma, P)) = \\ &= R(\|\dot{\gamma}\|, P, \Xi) + \sum_{i=1}^n |F(t_i^1, t_i^2, t_i^3) - F(\tau_i, \tau_i, \tau_i)|(t_i - t_{i-1}) \end{aligned}$$

non appena si sia definita $F : [a, b]^3 \longrightarrow \mathbb{R}$ mediante la

$$F(r, s, t) = \sqrt{\dot{x}^2(r) + \dot{y}^2(s) + \dot{z}^2(t)}$$

Dal momento che F è continua sul cubo $[a, b]^3$, si può dimostrare ricorrendo al concetto, non banale, di uniforme continuità che pur di scegliere la partizione P sufficientemente fine si può supporre che

$$|F(t_i^1, t_i^2, t_i^3) - F(\tau_i, \tau_i, \tau_i)| < \frac{\varepsilon}{b-a}$$

poichè al raffinarsi della partizione P si ha

$$R(\|\dot{\gamma}\|, P, \Xi) \rightarrow \int_a^b |\dot{\gamma}(t)| dt$$

e

$$\ell(\Lambda(\Gamma, P)) - R(\|\dot{\gamma}\|, P, \Xi) \rightarrow 0$$

possiamo concludere che se $|\dot{\gamma}|$ è integrabile, allora

$$\ell(\Lambda(\Gamma, P)) \rightarrow \int_a^b |\dot{\gamma}(t)| dt$$

DEFINIZIONE 4.3. Diciamo che due curve γ_1, γ_2 regolari in $\mathbb{R}^3$ sono equivalenti se, essendo

$$\gamma_1 : [a, b] \longrightarrow \mathbb{R}^3, \quad \gamma_1(t) = (x_1(t), y_1(t), z_1(t))$$

$$\gamma_2 : [c, d] \longrightarrow \mathbb{R}^3, \quad \gamma_2(t) = (x_2(t), y_2(t), z_2(t)),$$

esiste una funzione $\phi : [a, b] \longrightarrow [c, d]$ tale che

- $\gamma_1(t) = \gamma_2(\phi(t))$,
- $\phi \in \mathcal{C}^1([a, b])$,
- $\phi(a) = c, \phi(b) = d, \phi(t) > 0 \forall t \in (a, b)$.

Ci riferiremo alla funzione ϕ come ad un cambiamento regolare di parametrizzazione relativo alle curve γ_1, γ_2 .

Ovviamente, dal momento che ϕ è invertibile, anche ϕ^{-1} è un cambiamento regolare di parametrizzazione. Più precisamente mentre ϕ trasforma γ_2 in γ_1 , ϕ^{-1} opera la trasformazione di γ_1 in γ_2 .

1.2. Lunghezza d'Arco.

TEOREMA 4.1. *Sia γ una curva regolare in $\mathbb{R}^3$ e sia γ^* una curva regolare in $\mathbb{R}^3$ ad essa equivalente; si ha*

$$\ell(\gamma) = \ell(\gamma^*)$$

DIMOSTRAZIONE. Sia ϕ un cambiamento regolare di parametrizzazione relativo alle curve γ e γ^* , $\phi : [a, b] \longrightarrow [c, d]$ e sia

$$\gamma(t) = \gamma^*(\phi(t)).$$

Dal momento che $\phi > 0$,

$$\begin{aligned} \ell(\gamma) &= \int_a^b \|\dot{\gamma}(t)\| dt = \int_a^b \|(d/dt)\gamma^*(\phi(t))\| dt = \\ &= \int_a^b \|\dot{\gamma}^*(\phi(t))\dot{\phi}(t)\| dt = \int_a^b \|\dot{\gamma}^*(\phi(t))\|\dot{\phi}(t) dt \end{aligned}$$

e applicando il teorema di integrazione per sostituzione

$$\ell(\gamma) = \int_{\phi^{-1}(c)}^{\phi^{-1}(d)} \|\dot{\gamma}^*(\phi(t))\| \phi(t) dt = \int_c^d \|\dot{\gamma}^*(t)\| dt = \ell(\gamma^*)$$

□

Se γ è una curva regolare in $\mathbb{R}^3$, la funzione

$$s : [a, b] \longrightarrow \mathbb{R}$$

definita da

$$s(t) = \int_a^t \|\gamma(\tau)\| d\tau$$

si chiama lunghezza d'arco della curva γ .

Per le ipotesi fatte, s è una funzione di classe $\mathcal{C}^1([a, b])$ strettamente crescente e $s(t)$ misura la lunghezza del percorso compiuto da un punto che si muova, a partire da $\gamma(a)$, lungo la traccia della curva γ fino al punto $\gamma(t)$.

Inoltre s è una funzione invertibile e detta $t = s^{-1}$ la sua inversa essa pure risulta strettamente crescente e di classe $\mathcal{C}^1([0, \ell(\gamma)])$.

Pertanto è possibile considerare t come un cambiamento regolare di parametrizzazione e, detta γ^* la curva che si ottiene da γ mediante tale cambiamento si ha

$$\gamma^*(s) = \gamma(t(s)).$$

Poichè

$$\frac{dt}{ds}(s_0) = \frac{1}{\frac{ds}{dt}(t(s_0))}$$

si ottiene

$$\dot{\gamma}^*(s) = \frac{\dot{\gamma}(t(s))}{\|\dot{\gamma}(t(s))\|} = T_\gamma(t(s)).$$

DEFINIZIONE 4.4. Sia $f : A \longrightarrow \mathbb{R}$, $A \subset \mathbb{R}^3$ e supponiamo che γ sia una curva regolare in $\mathbb{R}^3$ con traccia contenuta in A (o più brevemente sia γ una curva regolare in A).

Definiamo integrale di linea di f su γ

$$\int_{\gamma} f ds = \int_a^b f(\gamma(t)) ds(t) = \int_a^b f(\gamma(t)) \|\dot{\gamma}(t)\| dt$$

qualora l'ultimo integrale esista.

E' immediato verificare che se f è una funzione continua su A e γ è una curva regolare in A allora f è integrabile su γ .

1.3. Curvatura - Terna Intrinseca. Ricordiamo infine brevemente le definizioni di alcuni utili elementi di una curva.

Abbiamo già definito il vettore tangente ad una curva semplice regolare in $\mathbb{R}^3$ mediante la

$$T_\gamma(t) = \frac{\dot{\gamma}(t)}{\|\dot{\gamma}(t)\|}$$

e abbiamo già provato che

$$T_{\gamma^*}(s) = \left(\frac{d}{ds}x(t(s)), \frac{d}{ds}y(t(s)), \frac{d}{ds}z(t(s)) \right)$$

è un versore.

Definiamo curvatura di γ nel punto $(x(t), y(t), z(t))$ il valore

$$K_\gamma(t) = \frac{\|\dot{\gamma}(t) \wedge \ddot{\gamma}(t)\|}{\|\dot{\gamma}(t)\|^3}$$

(ove con il simbolo $\wedge$ si sia indicato il prodotto vettoriale in $\mathbb{R}^3$).

Definiamo altresì raggio di curvatura di γ nel punto $(x(t), y(t), z(t))$ il valore

$$R_\gamma(t) = \frac{1}{K_\gamma(t)}.$$

La definizione può giustificarsi nella seguente maniera:

Indichiamo con α l'angolo formato dai vettori $\gamma(t)$ e $\gamma(t + \Delta t)$.

E' naturale definire come curvatura media di γ nell'intervallo $[t, t + \Delta t]$ il rapporto

$$\frac{\alpha}{|\Delta s|}$$

essendo $\Delta s = s(t + \Delta t) - s(t)$.

Posto $\Delta\gamma(t) = \gamma(t + \Delta t) - \gamma(t)$, dal momento che si ha

$$\sin \alpha = \frac{\|\dot{\gamma}(t) \wedge \dot{\gamma}(t + \Delta t)\|}{\|\dot{\gamma}(t)\| \|\dot{\gamma}(t + \Delta t)\|}$$

si ottiene

$$\sin \alpha = \frac{\|\dot{\gamma}(t) \wedge \Delta\dot{\gamma}(t)\|}{\|\dot{\gamma}(t)\| \|\dot{\gamma}(t + \Delta t)\|}$$

non appena si tenga conto del fatto che

$$\dot{\gamma}(t + \Delta t) = \dot{\gamma}(t) + \Delta\dot{\gamma}(t)$$

e

$$\dot{\gamma}(t) \wedge \dot{\gamma}(t) = 0$$

Ma allora

$$\frac{\alpha}{|\Delta s|} = \frac{\alpha}{\sin \alpha} \frac{\sin \alpha}{|\Delta s|}$$

e, poiché $\alpha \rightarrow 0$ quando $\Delta t \rightarrow 0$ per quanto abbiamo visto sopra, si ha

$$\lim_{\Delta t \rightarrow 0} \frac{\alpha}{|\Delta s|} = \lim_{\Delta t \rightarrow 0} \frac{\|\dot{\gamma}(t) \wedge \Delta \dot{\gamma}(t)/\Delta t\|}{\|\dot{\gamma}(t)\| \|\dot{\gamma}(t + \Delta t)\| |\Delta s/\Delta t|} = \frac{\|\dot{\gamma}(t) \wedge \ddot{\gamma}(t)\|}{\|\dot{\gamma}(t)\|^3}$$

Chiamiamo piano osculatore a γ nel punto $(x(t), y(t), z(t))$ il piano definito dall'equazione

$$\langle (x, y, z), \dot{\gamma}(t) \wedge \ddot{\gamma}(t) \rangle = 0$$

La definizione di piano osculatore si può interpretare geometricamente come segue.

Consideriamo il piano che è individuato dai vettori $\dot{\gamma}(t)$ e $\dot{\gamma}(t + \Delta t)$; una normale a tale piano sarà data da

$$\dot{\gamma}(t) \wedge \frac{\dot{\gamma}(t + \Delta t) - \dot{\gamma}(t)}{\Delta t} = \dot{\gamma}(t) \wedge \frac{\Delta \dot{\gamma}(t)}{\Delta t}$$

e, se $\Delta t \rightarrow 0$ essa tende a

$$\dot{\gamma}(t) \wedge \ddot{\gamma}(t)$$

Pertanto $\dot{\gamma}(t) \wedge \ddot{\gamma}(t)$ è normale al piano che si ottiene come limite del piano considerato.

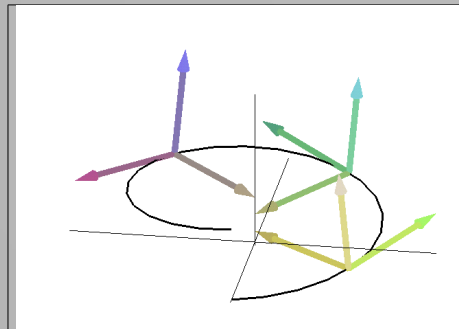


FIGURA 4.2.

In altri termini il piano osculatore alla curva γ nel punto $(x(t), y(t), z(t))$ contiene i vettori $\dot{\gamma}(t)$ e $\ddot{\gamma}(t)$.

Nel caso in cui la curva sia parametrizzata secondo la lunghezza d'arco, dal momento che $\|\dot{\gamma}\| = 1$ si ha

$$\langle \dot{\gamma}, \ddot{\gamma} \rangle = \frac{d}{ds} \|\dot{\gamma}(s)\|^2 / 2 = 0$$

$\dot{\gamma}(s)$ si dice *vettore tangente* a γ in $\gamma(s)$

$\ddot{\gamma}(s)$ si dice *vettore normale principale* a γ in $\gamma(s)$

$\dot{\gamma}(s) \wedge \ddot{\gamma}(s)$ si dice *vettore binormale* a γ in $\gamma(s)$.

I vettori $\dot{\gamma}(s)$, $\ddot{\gamma}(s)$, $\dot{\gamma}(s) \wedge \ddot{\gamma}(s)$ costituiscono il *triedro principale*, o naturale, della curva γ nel punto $(x(s), y(s), z(s))$; il triedro principale è anche noto come *terna intrinseca* della curva.

Ricordiamo infine alcune classiche definizioni senza però entrare nei dettagli.

Se $\gamma : [a, b] \longrightarrow \mathbb{R}^3$ è una curva tale che $z(t) \equiv 0$, diremo che γ è una curva piana.

Il vettore $(-\dot{y}, \dot{x})$ si dice vettore normale alla curva γ .

Sia $R_\gamma(t)$ il raggio di curvatura di γ ; diciamo centro di curvatura di γ il punto $c_\gamma(t)$ centro del cerchio che ha per raggio $R_\gamma(t)$ ed è tangente a γ .

$c_\gamma(t)$ descrive, al variare di t , una curva $\gamma^\#$ che si definisce evoluta della curva γ ; γ , a sua volta, si dice involuta della curva $\gamma^\#$.

Se $\gamma(t) = (x(t), y(t))$, le equazioni della evoluta di γ sono date da

$$x^\#(t) = x(t) - \dot{y}(t) \frac{\dot{x}^2(t) + \dot{y}^2(t)}{\dot{x}(t)\ddot{y}(t) - \dot{y}(t)\ddot{x}(t)}$$

$$y^\#(t) = y(t) + \dot{x}(t) \frac{\dot{x}^2(t) + \dot{y}^2(t)}{\dot{x}(t)\ddot{y}(t) - \dot{y}(t)\ddot{x}(t)}$$

2. Superfici ed Integrali di Superficie

In analogia con quanto fatto per la lunghezza di una linea definiamo

DEFINIZIONE 4.5. Sia $R = [a, b] \times [c, d]$, chiamiamo *superficie parametrica in $\mathbb{R}^3$* una funzione

$$S : R \longrightarrow \mathbb{R}^3, \quad S(u, v) = (x(u, v), y(u, v), z(u, v)).$$

Chiamiamo *traccia o supporto di S* l'insieme

$$\Sigma = \{(x, y, z) \in \mathbb{R}^3 : \exists (u, v) \in R, (x, y, z) = (x(u, v), y(u, v), z(u, v))\}$$

Osserviamo che $\Sigma = R(S)$ è il rango della funzione S . Una *superficie parametrica S in $\mathbb{R}^3$* si dice

- semplice, se è iniettiva,

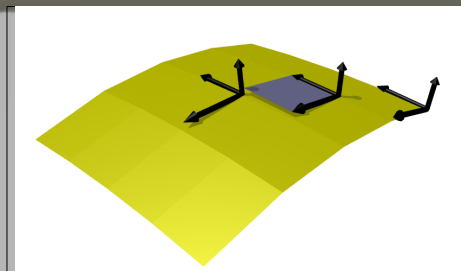


FIGURA 4.3.

- regolare, se $S \in \mathcal{C}^1(R)$ e

$$\nabla S = \begin{pmatrix} S_u \\ S_v \end{pmatrix} = \begin{pmatrix} \nabla_u S \\ \nabla_v S \end{pmatrix}$$

ha caratteristica massima ($= 2$).

Chiamiamo vettore normale ad S in $(x(u, v), y(u, v), z(u, v))$ il vettore

$$\begin{aligned}
 N(u, v) &= S_u(u, v) \wedge S_v(u, v) = \\
 &= \left(\det \begin{pmatrix} y_u & z_u \\ y_v & z_v \end{pmatrix}, -\det \begin{pmatrix} x_u & z_u \\ x_v & z_v \end{pmatrix}, \det \begin{pmatrix} x_u & y_u \\ x_v & y_v \end{pmatrix} \right) = \\
 &= \left(\frac{\partial(y, z)}{\partial(u, v)}, \frac{\partial(z, x)}{\partial(u, v)}, \frac{\partial(x, y)}{\partial(u, v)} \right)
 \end{aligned}$$

E' d'uso indicare con A, B, C le componenti del vettore N . Pertanto

$$N(u, v) = \left(\frac{\partial(y, z)}{\partial(u, v)}, \frac{\partial(z, x)}{\partial(u, v)}, \frac{\partial(x, y)}{\partial(u, v)} \right) = (A, B, C)$$

e

$$\|N\| = \sqrt{A^2 + B^2 + C^2}$$

Per definire l'area di una porzione di superficie possiamo utilizzare le approssimazioni lineari della funzione S che definisce la superficie stessa.

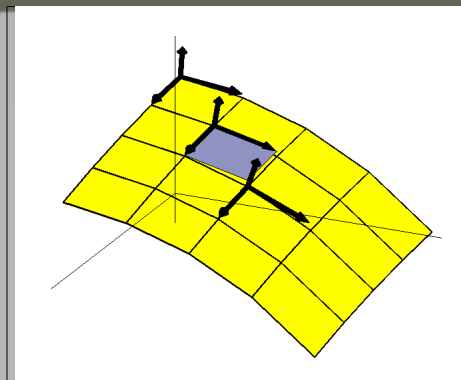


FIGURA 4.4.

Sia $\Omega = P \times Q$ una partizione dell'intervallo R , $P \in \mathcal{P}(a, b)$, $Q \in \mathcal{P}(c, d)$, e siano $\{R_k, k = 0..n\}$ i rettangoli in cui P divide R .

Sia $S : R \longrightarrow \mathbb{R}^3$ una superficie semplice, regolare; chiamiamo approssimazione lineare della superficie S relativa alla partizione Ω ed alla scelta Ξ , $\Pi(S, \Omega, \Xi)$ la superficie

$$\Pi(S, \Omega, \Xi)(u, v) = S(\xi_i, \eta_j) + \langle \nabla S(\xi_i, \eta_j), (u - \xi_i, v - \eta_i) \rangle, (u, v) \in R_{ij}$$

Definiamo approssimazione lineare dell'area della superficie S , relativa alla partizione Ω ed alla scelta Ξ ,

$$La(S, \Omega, \Xi) = \sum_{ij} \text{mis}(\Pi(S, \Omega, \Xi)(R_{ij})).$$

Diciamo infine che la superficie S ha area $A(S)$ se $\forall \varepsilon > 0$ esiste $\Omega_\varepsilon \in \mathcal{P}(R)$ tale che $\forall \Omega < \Omega_\varepsilon, \forall \Xi$ si ha

$$|La(S, \Omega, \Xi) - A(S)| < \varepsilon$$

ne segue che se S è una superficie parametrica semplice e regolare in $\mathbb{R}^3$; allora

$$A(S) = \int_R \|n(u, v)\| du dv$$

infatti si ha

$$\text{mis}(\Pi(S, \Omega, \Xi)(R_{ij})) = \|S_u(\xi_i, \eta_j) \wedge S_v(\xi_i, \eta_j)\| (u_i - u_{i-1})(v_j - v_{j-1})$$

da cui

$$\begin{aligned}
 (4.1) \quad \Lambda(S, \Omega, \Xi) &= \sum_{i=1}^n \sum_{j=1}^m \|n(\xi_i, \eta_j)\| (u_i - u_{i-1})(v_j - v_{j-1}) = \\
 &= R(\|n\|, \Omega, \Xi)
 \end{aligned}$$

Ricordando la definizione di integrabilità secondo Cauchy-Riemann si ha che

$$R(\|n\|, \Omega, \Xi) \rightarrow \iint_R \|n\| du dv$$

e quindi

$$A(S) = \int_R \|n(u, v)\| du dv$$

ed indichiamo

$$A(S) = \int_S d\sigma$$

Si può provare, usando il teorema di cambiamento di variabili negli integrali multipli, che se

$$S : R \longrightarrow \mathbb{R}^3 \quad \text{ed} \quad S_1 : R_1 \longrightarrow \mathbb{R}^3$$

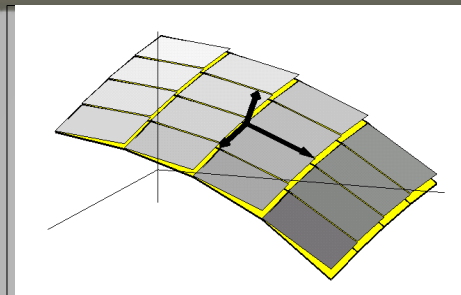


FIGURA 4.5.

sono due rappresentazioni equivalenti della stessa superficie, cioè se

$$S(u, v) = S_1(\phi(u, v))$$

con

$$\phi : R \longrightarrow R_1$$

, invertibile, $\phi \in \mathcal{C}^1(R)$, allora

$$\int_R \|n(u, v)\| du dv = \int_{R_1} \|n_1(u, v)\| du dv.$$

Ciò consente di affermare che la definizione di area di una superficie non dipende dalla parametrizzazione scelta.

DEFINIZIONE 4.6. *Sia S una superficie parametrica semplice e regolare in $\mathbb{R}^3$ e sia $\Sigma \subset A \subset \mathbb{R}^3$; sia inoltre $f : A \rightarrow \mathbb{R}$; definiamo*

$$\int_S f d\sigma = \int_R f(S(u, v)) \|n(u, v)\| du dv.$$

3. Forme Differenziali in $\mathbb{R}^3$

Il linguaggio delle forme differenziali è abbastanza complesso ed astratto nonostante ciò le applicazioni della relativa teoria sono numerose e la teoria stessa consente di formulare l'estensione del teorema fondamentale del calcolo integrale alle funzioni di più variabili.

Ci limitiamo ad illustrare le definizioni ed i risultati fondamentali, senza formalizzare le prime e senza dimostrare i secondi, nel caso di $\mathbb{R}^3$.

In $\mathbb{R}^3$ possiamo considerare

- **forme differenziali di ordine 0 o 0–forme**

$$\omega_0 = f(x, y, z)$$

- **forme differenziali di ordine 1 o 1–forme**

$$\omega_1 = f(x, y, z)dx + g(x, y, z)dy + h(x, y, z)dz$$

- **forme differenziali di ordine 2 o 2–forme**

$$\omega_2 = f(x, y, z)dy \wedge dz + g(x, y, z)dz \wedge dx + h(x, y, z)dx \wedge dy$$

- **forme differenziali di ordine 3 o 3–forme**

$$\omega_3 = f(x, y, z)dx \wedge dy \wedge dz$$

essendo f, g, h funzioni definite in $\mathbb{R}^3$ a valori in $\mathbb{R}$.

Non precisiamo la natura dei simboli dx dy e dz e dell'operazione di prodotto esterno $\wedge$.

Diciamo soltanto che i simboli dx dy e dz ci daranno indicazioni su come trattare ciascuna delle forme mentre il prodotto esterno soddisfa la seguente tabellina

$\wedge$	dx	dy
dx	0	$dx \wedge dy$
dy	$-dx \wedge dy$	0

Simmetricamente

Possiamo considerare in $\mathbb{R}^3$

- **varietà Ω_0 di dimensione 0 o 0-varietà**
cioè **unione finita di punti a ciascuno dei quali è associato un segno (+ o -)**
- **varietà Ω_1 di dimensione 1 o 1-varietà**
cioè **unione finita di linee a ciascuna delle quali è associato un segno (+ o -)**
- **varietà Ω_2 di dimensione 2 o 2-varietà**
cioè **unione finita di superfici a ciascuna delle quali è associato un segno (+ o -)**
- **varietà Ω_3 di dimensione 3 o 3-varietà**
cioè **unione finita di volumi a ciascuno dei quali è associato un segno (+ o -)**

Possiamo poi considerare due operazioni:

- Un'operazione che indichiamo con d che trasforma una k -forma in una $(k + 1)$ -forma
- Un'operazione che indichiamo con ∂ che trasforma una k -varietà in una $(k - 1)$ -varietà

L'operazione d si definisce mediante le seguenti regole che riportiamo omettendo di scrivere esplicitamente la dipendenza da (x, y, z) delle funzioni f, g, h e delle loro derivate.

$$d\omega_0 = f_x(x, y, z)dx + f_y(x, y, z)dy + f_z(x, y, z)dz = f_x dx + f_y dy + f_z dz$$

$$\begin{aligned}
 d\omega_1 &= (f_x dx + f_y dy + f_z dz) \wedge dx + \\
 &\quad + (g_x dx + g_y dy + g_z dz) \wedge dy + \\
 &\quad + (h_x dx + h_y dy + h_z dz) \wedge dz = \\
 &= f_x dx \wedge dx + f_y dy \wedge dx + f_z dz \wedge dx + \\
 &\quad + g_x dx \wedge dy + g_y dy \wedge dy + g_z dz \wedge dy + \\
 &\quad + h_x dx \wedge dz + h_y dy \wedge dz + h_z dz \wedge dz = \\
 &= (f_y - g_x) dy \wedge dx + (f_z - h_x) dz \wedge dx + \\
 &\quad + (h_y - g_z) dy \wedge dz
 \end{aligned}$$

$$\begin{aligned}
 d\omega_2 &= (f_x dx + f_y dy + f_z dz) \wedge dy \wedge dz + \\
 &\quad + (g_x dx + g_y dy + g_z dz) \wedge dz \wedge dx + \\
 &\quad + (h_x dx + h_y dy + h_z dz) \wedge dx \wedge dy = \\
 &= (f_x + g_y + h_z) dx \wedge dy \wedge dz
 \end{aligned}$$

Si ha inoltre che

$$\begin{aligned}
 d\omega_3 &= f_x dx \wedge dx \wedge dy \wedge dz + f_y dy \wedge dx \wedge dy \wedge dz + \\
 &\quad + f_z dz \wedge dx \wedge dy \wedge dz = 0
 \end{aligned}$$

Per quanto riguarda l'operazione ∂ usiamo le seguenti definizioni

Se $V = V(t, s, r)$ è una 3-varietà, cioè se

$$V : [a, b] \times [c, d] \times [\alpha, \beta] \rightarrow \mathbb{R}^3$$

Chiamiamo frontiera di V e scriviamo ∂V la 2-varietà che si ottiene considerando l'unione delle superfici definite dalle seguenti parametrizzazioni a ciascuna delle quali attribuiamo un segno che chiamiamo orientamento della superficie secondo la seguente regola:

Attribuiamo l'indice 0 al primo estremo e l'indice 1 al secondo estremo di ciascuno dei segmenti $[a, b]$, $[c, d]$, $[\alpha, \beta]$ e assegnamo ad ogni parametrizzazione il segno

$$(-1)^{\text{posto della variabile} + \text{indice dell'estremo}}$$

$V(a, s, r)$	con il segno	$-$	$[(-1)^{1+0}]$
$V(b, s, r)$	con il segno	$+$	$[(-1)^{1+1}]$
$V(t, c, r)$	con il segno	$+$	$[(-1)^{2+0}]$
$V(t, d, r)$	con il segno	$-$	$[(-1)^{2+1}]$
$V(t, s, \alpha)$	con il segno	$-$	$[-1]^{3+0}]$
$V(t, s, \beta)$	con il segno	$+$	$[(-1)^{3+1}]$

Se $S = S(u, v)$ è una 2-varietà , cioè se

$$S : [a, b] \times [c, d] \rightarrow \mathbb{R}^3$$

Chiamiamo frontiera di S e scriviamo ∂S la 1–varietà che si ottiene considerando l’unione delle linee definite dalle seguenti trasformazioni a ciascuna delle quali attribuiamo il segno indicato

$S(a, v)$	con il segno–	$[(-1)^{1+0}]$
$S(b, v)$	con il segno+	$[(-1)^{1+1}]$
$S(u, c)$	con il segno–	$[(-1)^{2+0}]$
$S(u, d)$	con il segno+	$[(-1)^{2+1}]$

Se $\gamma = \gamma(t)$ è una 1–varietà , cioè se

$$\gamma : [a, b] \rightarrow \mathbb{R}^3$$

Chiamiamo frontiera di γ e scriviamo $\partial\gamma$ la 0–varietà che si ottiene considerando l’unione dei punti definiti dalle seguenti trasformazioni a ciascuna delle quali attribuiamo il segno indicato

$$\begin{array}{ll}
 \gamma(a) & \text{con il segno} - \quad [(-1)^{1+0}] \\
 \gamma(b) & \text{con il segno} + \quad [(-1)^{1+1}]
 \end{array}$$

Per comprendere il significato dei segni occorre prima chiarire come una k -forma può essere integrata su una k -varietà occorre cioè definire il simbolo

$$\int_{C_k} \omega_k$$

- Se $\omega_0 = f(x, y, z)$ e $C_0 = \pm P$ (cioè è la zero varietà costituita dal punto $P = (x, y, z)$ con il segno $\pm$)

$$\int_{C_0} \omega_0 = \pm f(P)$$

- Se $\omega_1 = f(x, y, z)dx + g(x, y, z)dy + h(x, y, z)dz$ e $C_1 = \pm \gamma$ (cioè è la 1-varietà costituita dalla curva $\gamma = \gamma(t)$ con il segno $\pm$)

$$\int_{C_1} \omega_1 = \pm \int_a^b f(\gamma(t))\dot{x}(t) + g(\gamma(t))\dot{y}(t) + h(\gamma(t))\dot{z}(t)dt$$

- Se $\omega_2 = f(x, y, z)dy \wedge dz + g(x, y, z)dz \wedge dx + h(x, y, z)dx \wedge dy$ e $C_1 = \pm S$ (cioè è la 2-varietà costituita dalla superficie $S = S(u, v)$ con il segno $\pm$

$$\int_{C_2} \omega_2 = \pm \int_a^b \int_c^d f(S(u, v)) \frac{\partial(y, z)}{\partial(u, v)} + g(S(u, v)) \frac{\partial(z, x)}{\partial(u, v)} + h(S(u, v)) \frac{\partial(x, y)}{\partial(u, v)} du dv$$

- Se $\omega_3 = f(x, y, z)dx \wedge dy \wedge dz$ e $C_3 = \pm V = V(t, s, r)$ (cioè è la 3-varietà costituita dal volume V con il segno $\pm$)

$$\int_{C_0} \omega_0 = \pm \int_a^b \int_c^d \int_\alpha^\beta f(V(t, s, r)) \frac{\partial(x, y, z)}{\partial(t, s, r)} dt ds dr$$

A questo punto possiamo anche chiarire il significato del segno che abbiamo attribuito a i vari pezzi di frontiera di una varietà.

3.0.1. Un'idea per rendersi conto dei segni. Consideriamo una 1 - forma

$$\omega_2 = f(x, y, z)dx + g(x, y, z)dy + h(x, y, z)dz$$

ed la 2- varietà S_2 costituita dal quadrato $[0, 1] \times [0, 1]$ nel piano $z = 0$.

Possiamo parametrizzare $S_2 = S_2(u, v)$ nella seguente maniera

$$S_2 \begin{cases} x = u \\ y = v \\ z = 0 \end{cases}, \quad (u, v) \in [0, 1] \times [0, 1]$$

∂S_2 , la frontiera di S_2 , è una curva costituita dalle linee $\gamma_1(u) = S_2(u, 0)$, $\gamma_2(v) = S_2(0, v)$, $\gamma_3(u) = S_2(u, 1)$, $\gamma_4(v) = S_2(1, v)$,

$$\begin{aligned} \gamma_1 \begin{cases} x = u \\ y = 0 \\ z = 0 \end{cases}, \quad u \in [0, 1] & \quad \gamma_2 \begin{cases} x = 1 \\ y = v \\ z = 0 \end{cases}, \quad v \in [0, 1] \\ \gamma_3 \begin{cases} x = u \\ y = 1 \\ z = 0 \end{cases}, \quad u \in [0, 1] & \quad \gamma_4 \begin{cases} x = 0 \\ y = v \\ z = 0 \end{cases}, \quad v \in [0, 1] \end{aligned}$$

Con le regole prima descritte attribuiamo il segno $+$ a γ_1 e γ_2 ed il segno $-$ a γ_3 e γ_4 .
Pertanto

$$\int_{\partial S_2} \omega_2 = + \int_{\gamma_1} \omega_2 + \int_{\gamma_2} \omega_2 - \int_{\gamma_3} \omega_2 - \int_{\gamma_4} \omega_2$$

Se conveniamo che attribuire il segno $+$ ad un lato significa che quel lato è percorso nella direzione positiva dell'asse cui è parallelo mentre il segno $-$ indica che è percorso in direzione opposta, è immediato verificare che con tali scelte di segno si vede che il perimetro del quadrato S_2 , che costituisce la frontiera ∂S_2 , è percorso in senso antiorario, se visto dall'alto (cioè da un punto di vista che giace nel semipiano $z > 0$).

Se non avessimo fatto uso dei segni avremmo percorso i tratti γ_1 e γ_2 in senso antiorario, mentre i tratti γ_3 e γ_4 sarebbero stati percorsi in senso orario.

Ancora se S_1 è la varietà definita da

$$\gamma(t) \begin{cases} x = t \\ y = 0 \\ z = 0 \end{cases} \quad , \quad t \in [0, 1]$$

avremo che la frontiera $\partial\gamma$ è costituita dai due punti

$$P_0 = \gamma(0) \text{ con il segno } -$$

$$P_1 = \gamma(1) \text{ con il segno } +$$

Se $\omega_0 = f(x, y, z)$ è una 0-forma

$$(4.2) \quad \int_{\partial S_1} \omega_0 = - \int_{P_0} \omega_0 + \int_{P_1} \omega_0 = f(P_1) - f(P_0)$$

Si può provare che

TEOREMA 4.2. *Sia ω_n una n -forma in A , allora*

$$d(d\omega_k) = 0.$$

DIMOSTRAZIONE. Dimostriamo il risultato nel caso in cui ω_1 è una 1-forma in $\mathbb{R}^3$

$$\omega_1 = f_1 dx_1 + f_2 dx_2 + f_3 dx_3$$

si ha

$$\begin{aligned}
 d\omega_1 &= \frac{\partial f_1}{\partial x_2} dx_2 \wedge dx_1 + \frac{\partial f_1}{\partial x_3} dx_3 \wedge dx_1 + \\
 &+ \frac{\partial f_2}{\partial x_1} dx_1 \wedge dx_2 + \frac{\partial f_2}{\partial x_3} dx_3 \wedge dx_2 + \\
 &+ \frac{\partial f_3}{\partial x_1} dx_1 \wedge dx_3 + \frac{\partial f_3}{\partial x_2} dx_2 \wedge dx_3
 \end{aligned}$$

e tenendo conto che $dx_i \wedge dx_j = -dx_j \wedge dx_i$,

$$\begin{aligned}
 d(d\omega_1) &= \frac{\partial^2 f_1}{\partial x_3 \partial x_2} dx_3 \wedge dx_2 \wedge dx_1 + \frac{\partial^2 f_1}{\partial x_2 \partial x_3} dx_2 \wedge dx_3 \wedge dx_1 + \\
 &+ \frac{\partial^2 f_2}{\partial x_3 \partial x_1} dx_3 \wedge dx_1 \wedge dx_2 + \frac{\partial^2 f_2}{\partial x_1 \partial x_3} dx_1 \wedge dx_3 \wedge dx_2 + \\
 &+ \frac{\partial^2 f_3}{\partial x_2 \partial x_1} dx_2 \wedge dx_1 \wedge dx_3 + \frac{\partial^2 f_3}{\partial x_1 \partial x_2} dx_1 \wedge dx_2 \wedge dx_3
 \end{aligned}$$

e ogni termine elide il successivo.

□

DEFINIZIONE 4.7. Sia ω_n una $n -$ forma differenziale in A , diciamo che ω_n è **esatta** se esiste una $(n - 1) -$ forma su A , η_{n-1} tale che

$$d\eta_{n-1} = \omega_n$$

η_{n-1} si chiama primitiva di ω_n .

Diciamo che ω_n è **chiusa** se

$$d\omega_n = 0.$$

E' immediata conseguenza del precedente teorema il seguente risultato.

TEOREMA 4.3. Sia ω_n una $n -$ forma su A . Se ω_n è esatta, allora ω_n è chiusa.

DIMOSTRAZIONE. Se ω_n è esatta, $\omega_n = d\eta_{n-1}$ e

$$d\omega_n = d(d\eta_{n-1}) = 0.$$

□

Inoltre

Ogni 3 - forma in $\mathbb{R}^3$ è chiusa

A questo punto è importante riconoscere tra le n – forme su $\mathbb{R}^3$ quelle esatte e determinarne le primitive. Questo problema infatti ha notevoli riscontri sia dal punto di vista matematico che dal punto di vista fisico.

TEOREMA 4.4. *Sia ω_n una n –forma differenziale su un insieme A aperto e stellato; allora se ω_n è chiusa si ha che ω_n è esatta.*

Poichè ci limitiamo a considerare soltanto forme differenziali in $\mathbb{R}^3$, possiamo discutere l’enunciato precedente in ognuno dei tre casi che si presentano e dal momento che i risultati che riguardano le forme differenziali sono strettamente collegati alla teoria dei campi vettoriali, ricordiamo che

è assegnato un campo vettoriale in $\mathbb{R}^3$ se è data una funzione

$$\begin{aligned}
 F : A &\longrightarrow \mathbb{R}^3, \quad A \subset \mathbb{R}^3 \\
 F(x, y, z) &= (f(x, y, z), g(x, y, z), h(x, y, z)) \\
 f, g, h &: A \longrightarrow \mathbb{R}
 \end{aligned}$$

Ad un campo vettoriale F si può associare tanto una 1–forma

$$\omega_1 = f dx + g dy + h dz$$

che una 2–forma

$$\omega_2 = f dy \wedge dz + g dz \wedge dx + h dx \wedge dy$$

e gli integrali di linea di ω_1 e di superficie di ω_2 rappresentano, rispettivamente il lavoro lungo una linea ed il flusso attraverso una superficie del campo vettoriale F .

Una 0 – *forma* ed una 3–forma sono semplicemente identificate da una funzione

$$f : A \longrightarrow \mathbb{R}^3, \quad A \subset \mathbb{R}$$

mediante le

$$\omega_0 = f, \quad \omega_3 = f dx \wedge dy \wedge dz$$

Si definisce **divergenza** di un campo vettoriale F la funzione scalare

$$\operatorname{div} F = \frac{\partial f}{\partial x} + \frac{\partial g}{\partial y} + \frac{\partial h}{\partial z}$$

mentre si definisce **rotore** di un campo vettoriale F la funzione vettoriale

$$\text{rot } F = (h_y - g_z, f_z - h_x, g_x - f_y)$$

Se introduciamo il vettore formale

$$D = \left(\frac{\partial}{\partial x}, \frac{\partial}{\partial y}, \frac{\partial}{\partial z} \right)$$

possiamo scrivere che

$$\text{div } F = \langle D, F \rangle$$

$$\text{rot } F = D \wedge F$$

Ricordiamo anche

$$\text{div } \nabla \phi = \frac{\partial^2 \phi}{\partial x^2} + \frac{\partial^2 \phi}{\partial y^2} + \frac{\partial^2 \phi}{\partial z^2} = \Delta \phi$$

$\Delta \phi$ si chiama laplaciano della funzione ϕ ed è di fondamentale importanza in svariati campi della matematica e della fisica.

3.1. Primitive di 1-forme. Consideriamo una 1-forma

$$\omega_1 = f(x, y, z)dx + g(x, y, z)dy + h(x, y, z)dz$$

ed il campo vettoriale

$$F = (f, g, h)$$

ad essa associato. Si ha

$$d\omega_1 = (f_y - g_x)dy \wedge dx + (f_z - h_x)dz \wedge dx + (h_y - g_z)dy \wedge dz$$

e pertanto, affinché ω_1 sia esatta deve essere $d\omega_1 = 0$ e quindi

$$f_y - g_x = 0$$

$$f_z - h_x = 0$$

$$h_y - g_z = 0$$

cioè che

$$\text{rot } F = 0$$

La 0-forma $\omega_0 = \varphi(x, y, z)$ è una primitiva di ω_1 se $d\omega_0 = \omega_1$ cioè se

$$\varphi_x = f$$

$$\varphi_y = g$$

$$\varphi_z = h$$

ed in termini di campo vettoriale se

$$F = \nabla\varphi$$

Possiamo definire

$$A(x, y, z) = \int_{x_0}^x f(t, y, z) dt$$

e scegliere

$$\varphi(x, y, z) = A(x, y, z) + a(y, z)$$

Deve anche essere

$$g(x, y, z) = \varphi_y(t, y, z) = A_y(t, y, z) + a_y(y, z)$$

Poichè a dipende soltanto da (y, z) , affinchè la precedente uguaglianza sia possibile occorre che la funzione

$$g(x, y, z) - A_y(x, y, z)$$

non dipenda da x ; per verificare se questo è vero possiamo derivare rispetto ad x ed otteniamo

$$g_x - A_{yx} = g_x - A_{xy} = g_x - f_y = 0$$

in quanto valgono le condizioni necessarie.

Pertanto l'uguaglianza

$$a_y(y, z) = g(x, y, z) - A_y(x, y, z)$$

è possibile e possiamo concludere che

$$a(y, z) = \int_{y_0}^y \left(g(x, \eta, z) - A_y(x, \eta, z) \right) d\eta + b(z) = B(y, z) + b(z)$$

Ma allora

$$\varphi(x, y, z) = A(x, y, z) + B(y, z) + b(z)$$

e deve infine essere

$$\varphi_z(x, y, z) = A_z(x, y, z) + B_z(y, z) + b_z(z) = h(x, y, z)$$

Affinchè sia possibile trovare b in modo che la precedente uguaglianza sia soddisfatta è necessario che

$$A_z(x, y, z) + B_z(y, z) - h(x, y, z)$$

dipenda solo da z e quindi che

$$\left(h(x, y, z) - A_z(x, y, z) - B_z(y, z) \right)_x = 0$$

$$\left(h(x, y, z) - A_z(x, y, z) - B_z(y, z) \right)_y = 0$$

Si ha

$$(h - A_z - B_z)_x = h_x - A_{zx} - G_{zx} = h_x - A_{xz} = h_x - f_z = 0$$

$$(h - A_z - B_z)_y = h_y - A_{zy} - G_{zy} = h_y - A_{zy} - g_z + F_{yz} = h_y - g_z = 0$$

È allora sufficiente porre

$$b(z) = \int_{z_0}^z (h - A_z - a_z)(x, y, \zeta) d\zeta$$

3.2. Primitive di 2-forme. Consideriamo una 2-forma

$$\omega_2 = f(x, y, z)dy \wedge dz + g(x, y, z)dz \wedge dx + h(x, y, z)dx \wedge dy$$

ed il campo vettoriale

$$F = (f, g, h)$$

ad essa associato. Si ha

$$d\omega_2 = (f_x + g_y + h_z)dx \wedge dy \wedge dz$$

e pertanto, affinché ω_2 sia esatta $d\omega_2 = 0$ e quindi

$$f_x + g_y + h_z = 0$$

cioè che

$$\operatorname{div} F = 0$$

La 1-forma $\omega_1 = \alpha(x, y, z)dx + \beta(x, y, z)dy + \gamma(x, y, z)dz$ cui associamo il campo vettoriale $\Phi = (\alpha, \beta, \gamma)$, è una primitiva di ω_2 se $d\omega_1 = \omega_2$, cioè se

$$\gamma_y - \beta_z = f$$

$$\alpha_z - \gamma_x = g$$

$$\beta_x - \alpha_y = h$$

ed in termini di campo vettoriale se

$$\Phi = \text{rot } F$$

Possiamo scegliere di cercare una primitiva per la quale risulti

$$\gamma(x, y, z) = 0$$

per cui, le prime due condizioni sono soddisfatte se

$$-\beta_z = f$$

$$\alpha_z = g$$

e allo scopo è sufficiente definire

$$\alpha(x, y, z) = \int_{z_0}^z g(x, y, \zeta) d\zeta + a(x, y) = B(x, y, z) + a(x, y)$$

$$\beta(x, y, z) = \int_{z_0}^z -f(x, y, \zeta) d\zeta + b(x, y) = -A(x, y, z) + b(x, y)$$

Per quel che riguarda la terza condizione osserviamo che poichè supponiamo soddisfatta la condizione necessaria per l'esistenza della primitiva di ω_2 , si ha

$$(\beta_x - \alpha_y - h)_z = \beta_{xz} - \alpha_{yz} - h_z =$$

$$= \beta_{zx} - \alpha_{zy} - h_z = -g_y - f_x - h_z = 0$$

ammesso che la condizione necessaria sia soddisfatta, e quindi

$$\beta_x - \alpha_y - h$$

è costante rispetto a z e si ha

$$(\beta_x - \alpha_y - h)(x, y, z) = 0 \quad \Longleftrightarrow \quad (\beta_x - \alpha_y - h)(x, y, z_0) = 0$$

Ma allora, poichè

$$B_x(x, y, z) = \int_{z_0}^z g_x(x, y, \zeta) d\zeta$$

$$A_y(x, y, z) = \int_{z_0}^z f_y(x, y, \zeta) d\zeta$$

avremo

$$\begin{aligned}
 (\beta_x - \alpha_y - h)(x, y, z_0) &= \\
 &= (b_x - a_y - A_y - B_x - h)(x, y, z_0) = (b_x - a_y - h)(x, y, z_0) = 0
 \end{aligned}$$

e l'ultima uguaglianza è verificata se scegliamo, ad esempio,

$$a(x, y) = 0 \quad , \quad b(x, y) = \int_{x_0}^x h(\xi, y, z) d\xi$$

3.3. Primitive di 3-forme. Consideriamo una 3-forma

$$\omega_3 = f(x, y, z) dx \wedge dy \wedge dz$$

si ha sempre

$$d\omega_3 = 0$$

e pertanto ω_3 è sempre esatta

La 2-forma $\omega_2 = \alpha(x, y, z)dx + \beta(x, y, z)dy + \gamma(x, y, z)dz$ cui associamo il campo vettoriale $\Phi = (\alpha, \beta, \gamma)$, è una primitiva di ω_3 se $d\omega_2 = \omega_3$, cioè se

$$\alpha_x + \beta_y + \gamma_z = f$$

ed in termini di campo vettoriale se

$$\operatorname{div} \Phi = f$$

Si verifica subito che possiamo trovare una primitiva definendo

$$\alpha(x, y, z) = \beta(x, y, z) = 0 \qquad \gamma(x, y, z) = \int_{z_0}^z f(x, y, \zeta) d\zeta$$

4. Il Teorema di Stokes

Il teorema di Stokes costituisce la naturale estensione del teorema fondamentale del calcolo integrale, così come il concetto di potenziale di una forma differenziale costituisce la naturale estensione del concetto

di primitiva di una funzione di una variabile reale, ed è di grandissima importanza nella teoria e nelle applicazioni.

Con le definizioni precedenti il teorema di Stokes può essere enunciato come segue

TEOREMA 4.5. - Stokes - Sia ω_{n-1} una $(n-1)$ -forma differenziale su $\mathbb{R}^k$ e sia C una n -varietà in $\mathbb{R}^k$; allora

$$\int_{\partial C} \omega_{n-1} = \int_C d\omega_{n-1}.$$

Il teorema di Stokes è suscettibile di significative conseguenze sia in $\mathbb{R}^2$ che in $\mathbb{R}^3$ soprattutto se si tiene conto del fatto che, ad esempio in $\mathbb{R}^3$, le 1-forme e le 2-forme possono essere identificate con un campo vettoriale. (in $\mathbb{R}^2$ soltanto le 1-forme)

4.1. Il teorema della divergenza in $\mathbb{R}^3$. Consideriamo una 2-forma in $\mathbb{R}^3$

$$\omega_2(x, y, z) = f(x, y, z)dy \wedge dz + g(x, y, z)dz \wedge dx + h(x, y, z)dx \wedge dy$$

ed il campo vettoriale che la identifica

$$F = (f, g, h)$$

Si ha

$$d\omega_2 = (f_x + g_y + h_z)dx \wedge dy \wedge dz$$

Se $V = V(t, s, r)$ è una 3-varietà la cui frontiera è

$$\partial V = +V(1, s, r) \cup -V(0, s, r) \cup -V(t, 1, r) \cup +V(t, 0, r) \cup +V(t, s, 1) \cup -V(t, s, 0)$$

si ha

$$(4.3) \quad \int_{\partial V} f dy \wedge dz + g dz \wedge dx + h dx \wedge dy = \int_V (f_x + g_y + h_z) dx \wedge dy \wedge dz$$

E, se si tiene conto che la normale a ∂V è data da

$$N = \frac{1}{\nu} \left(\frac{\partial(y, z)}{\partial(u, v)}, -\frac{\partial(x, z)}{\partial(u, v)}, \frac{\partial(x, y)}{\partial(u, v)} \right) = \frac{1}{\nu} \left(\frac{\partial(y, z)}{\partial(u, v)}, \frac{\partial(z, x)}{\partial(u, v)}, \frac{\partial(x, y)}{\partial(u, v)} \right)$$

$$\nu = \sqrt{\left(\frac{\partial(y, z)}{\partial(u, v)}\right)^2 + \left(\frac{\partial(x, z)}{\partial(u, v)}\right)^2 + \left(\frac{\partial(x, y)}{\partial(u, v)}\right)^2}$$

si ottiene

$$\begin{aligned} \int_{\partial V} f dy \wedge dz - g dx \wedge dz + h dx \wedge dy &= \\ &= \iint_D \frac{1}{\nu} \left(f \frac{\partial(y, z)}{\partial(u, v)} - g \frac{\partial(x, z)}{\partial(u, v)} + h \frac{\partial(x, y)}{\partial(u, v)} \right) \nu du dv = \\ &= \iint_D \langle F, N \rangle \nu du dv = \iint_{\partial C} \langle F, N \rangle d\sigma \end{aligned}$$

Pertanto si ha

$$(4.4) \quad \iint_{\partial V} \langle F, N_e \rangle d\sigma = \int \iiint_V \operatorname{div} F dx dy dz$$

ove $N_e = N \operatorname{sgn} JV$ è il versore normale a ∂V orientato verso l'esterno di V .

La 4.3 è nota come teorema di Gauss e la 4.4 è la formulazione del teorema della divergenza.

Possiamo dimostrare il teorema della divergenza come segue.

DIMOSTRAZIONE. Proviamo ad esempio che

$$\iiint_V f_x dx \wedge dy \wedge dz = \iint_{\partial V} f dy \wedge dz$$

Si ha

$$\begin{aligned}
 \iint_{\partial V} f dy \wedge dz &= \int_0^1 \int_0^1 f \frac{\partial(y, z)}{\partial(s, r)} \Big|_{t=0}^{t=1} ds dr - \\
 &\quad - \int_0^1 \int_0^1 f \frac{\partial(y, z)}{\partial(t, s)} \Big|_{s=0}^{s=1} dt dr + \int_0^1 \int_0^1 f \frac{\partial(y, z)}{\partial(t, s)} \Big|_{r=0}^{r=1} dt ds = \\
 &= \int_0^1 \int_0^1 \int_0^1 \frac{d}{dt} \left(f \frac{\partial(y, z)}{\partial(s, r)} \right) dt ds dr - \\
 &\quad - \int_0^1 \int_0^1 \int_0^1 \frac{d}{ds} \left(f \frac{\partial(y, z)}{\partial(t, s)} \right) ds dt dr + \\
 &\quad + \int_0^1 \int_0^1 \int_0^1 \frac{d}{dr} \left(f \frac{\partial(y, z)}{\partial(t, s)} \right) dr dt ds
 \end{aligned}$$

e

$$\begin{aligned}
 \iint_{\partial V} f dy \wedge dz &= \\
 &= \int_0^1 \int_0^1 \int_0^1 (f_x x_t + f_y y_t + f_z z_t) \frac{\partial(y, z)}{\partial(s, r)} + f \frac{d}{dt} \frac{\partial(y, z)}{\partial(s, r)} dt ds dr - \\
 &\quad - \int_0^1 \int_0^1 \int_0^1 (f_x x_s + f_y y_s + f_z z_s) \frac{\partial(y, z)}{\partial(t, r)} + f \frac{d}{ds} \frac{\partial(y, z)}{\partial(t, r)} dt ds dr + \\
 &\quad + \int_0^1 \int_0^1 \int_0^1 (f_x x_r + f_y y_r + f_z z_r) \frac{\partial(y, z)}{\partial(t, s)} + f \frac{d}{dr} \frac{\partial(y, z)}{\partial(t, s)} dt ds dr = \\
 &= \int_0^1 \int_0^1 \int_0^1 (f_x x_t + f_y y_t + f_z z_t) \frac{\partial(y, z)}{\partial(s, r)} - \\
 &\quad - (f_x x_s + f_y y_s + f_z z_s) \frac{\partial(y, z)}{\partial(t, r)} + (f_x x_r + f_y y_r + f_z z_r) \frac{\partial(y, z)}{\partial(t, s)} dt ds dr \\
 &= \int_0^1 \int_0^1 \int_0^1 f \left(\frac{d}{dt} \frac{\partial(y, z)}{\partial(s, r)} - \frac{d}{ds} \frac{\partial(y, z)}{\partial(t, r)} + \frac{d}{dr} \frac{\partial(y, z)}{\partial(t, s)} \right) dt ds dr
 \end{aligned}$$

Ma

$$\begin{aligned}
 \left(\frac{d}{dt} \frac{\partial(y, z)}{\partial(s, r)} - \frac{d}{ds} \frac{\partial(y, z)}{\partial(t, r)} + \frac{d}{dr} \frac{\partial(y, z)}{\partial(t, s)} \right) &= \\
 &= \left(\frac{d}{dt} \det \begin{pmatrix} y_s & z_s \\ y_r & z_r \end{pmatrix} - \frac{d}{ds} \det \begin{pmatrix} y_t & z_t \\ y_r & z_r \end{pmatrix} + \frac{d}{dr} \det \begin{pmatrix} y_t & z_t \\ y_s & z_s \end{pmatrix} \right) = \\
 &= \det \begin{pmatrix} y_{st} & z_s \\ y_{rt} & z_r \end{pmatrix} + \det \begin{pmatrix} y_s & z_{st} \\ y_r & z_{rt} \end{pmatrix} - \det \begin{pmatrix} y_{ts} & z_t \\ y_{rs} & z_r \end{pmatrix} - \det \begin{pmatrix} y_t & z_{ts} \\ y_r & z_{rs} \end{pmatrix} + \\
 &\quad + \det \begin{pmatrix} y_{tr} & z_t \\ y_{sr} & z_s \end{pmatrix} + \det \begin{pmatrix} y_t & z_{tr} \\ y_s & z_{sr} \end{pmatrix} = 0
 \end{aligned}$$

per cui

$$\begin{aligned}
 \iint_{\partial V} f dy \wedge dz &= \int_0^1 \int_0^1 \int_0^1 (f_x x_t + f_y y_t + f_z z_t) \frac{\partial(y, z)}{\partial(s, r)} - \\
 &\quad - (f_x x_s + f_y y_s + f_z z_s) \frac{\partial(y, z)}{\partial(t, r)} + (f_x x_r + f_y y_r + f_z z_r) \frac{\partial(y, z)}{\partial(t, s)} dt ds dr \\
 &= \int_0^1 \int_0^1 \int_0^1 f_x \left(x_t \frac{\partial(y, z)}{\partial(s, r)} - x_s \frac{\partial(y, z)}{\partial(t, r)} + x_r \frac{\partial(y, z)}{\partial(t, s)} \right) dt ds dr + \\
 &\quad + \int_0^1 \int_0^1 \int_0^1 f_y \left(y_t \frac{\partial(y, z)}{\partial(s, r)} - y_s \frac{\partial(y, z)}{\partial(t, r)} + y_r \frac{\partial(y, z)}{\partial(t, s)} \right) dt ds dr + \\
 &\quad + \int_0^1 \int_0^1 \int_0^1 f_z \left(z_t \frac{\partial(y, z)}{\partial(s, r)} - z_s \frac{\partial(y, z)}{\partial(t, r)} + z_r \frac{\partial(y, z)}{\partial(t, s)} \right) dt ds dr
 \end{aligned}$$

Poichè

$$\begin{aligned}
 \left(y_t \frac{\partial(y, z)}{\partial(s, r)} - y_s \frac{\partial(y, z)}{\partial(t, r)} + y_r \frac{\partial(y, z)}{\partial t, s} \right) = \\
 \left(y_t \det \begin{pmatrix} y_s & z_s \\ y_r & z_r \end{pmatrix} - y_s \det \begin{pmatrix} y_t & z_t \\ y_r & z_r \end{pmatrix} + y_r \det \begin{pmatrix} y_t & z_t \\ y_s & z_s \end{pmatrix} \right) = 0
 \end{aligned}$$

e

$$\begin{aligned}
 \left(z_t \frac{\partial(y, z)}{\partial(s, r)} - z_s \frac{\partial(y, z)}{\partial(t, r)} + z_r \frac{\partial(y, z)}{\partial t, s} \right) = \\
 \left(z_t \det \begin{pmatrix} y_s & z_s \\ y_r & z_r \end{pmatrix} - z_s \det \begin{pmatrix} y_t & z_t \\ y_r & z_r \end{pmatrix} + z_r \det \begin{pmatrix} y_t & z_t \\ y_s & z_s \end{pmatrix} \right) = 0
 \end{aligned}$$

Possiamo concludere

$$\begin{aligned}
 \iint_{\partial V} f dy \wedge dz &= \\
 &= \int_0^1 \int_0^1 \int_0^1 f_x \left(x_t \frac{\partial(y, z)}{\partial(s, r)} - x_s \frac{\partial(y, z)}{\partial(t, r)} + x_r \frac{\partial(y, z)}{\partial(t, s)} \right) dt ds dr = \\
 &= \int_0^1 \int_0^1 \int_0^1 f_x \det \begin{pmatrix} x_t & y_t & z_t \\ x_s & y_s & z_s \\ x_r & y_r & z_r \end{pmatrix} dt ds dr = \\
 &= \int_0^1 \int_0^1 \int_0^1 f_x \frac{\partial(x, y, z)}{\partial(t, s, r)} = \iiint_V f_x dx \wedge dy \wedge dz dt ds dr
 \end{aligned}$$

□

4.2. Il Teorema del Rotore. Se consideriamo una 1-forma in $\mathbb{R}^3$

$$\omega_1 = f dx + g dy + h dz$$

e se

$$F = (f, g, h)$$

è il campo vettoriale ad essa associato, si ha

$$d\omega_1 = (g_x - f_y)dx \wedge dy + (h_y - g_z)dy \wedge dz + (h_x - f_z)dx \wedge dz$$

Sia $S = S(u, v)$ una 2-varietà, cioè sia

$$S : [0, 1] \times [0, 1] \rightarrow \mathbb{R}^3$$

e sia ∂S la sua frontiera

$$\partial S = -S(1, v) \cup +S(0, v) \cup +S(u, 1) \cup -S(u, 0)$$

si ha

$$\begin{aligned}
 (4.5) \quad \int_{\partial S} f dx + g dy + h dz &= \\
 &= \int_S (g_x - f_y) dx \wedge dy + (h_y - g_z) dy \wedge dz + (h_x - f_z) dx \wedge dz
 \end{aligned}$$

Ricordando che il versore tangente a ∂S è dato da

$$T = \left(\frac{\dot{x}}{\sqrt{\dot{x}^2 + \dot{y}^2 + \dot{z}^2}}, \frac{\dot{y}}{\sqrt{\dot{x}^2 + \dot{y}^2 + \dot{z}^2}}, \frac{\dot{z}}{\sqrt{\dot{x}^2 + \dot{y}^2 + \dot{z}^2}} \right)$$

si ha

$$\begin{aligned}
 \int_{\partial S} f dx + g dy + h dz &= \int_a^b (f \dot{x} + g \dot{y} + h \dot{z}) dt = \\
 &= \int_a^b \langle F, T \rangle \sqrt{\dot{x}^2 + \dot{y}^2 + \dot{z}^2} dt = \int_{\partial S} \langle F, T \rangle ds
 \end{aligned}$$

mentre

$$\begin{aligned}
 \int_S (g_x - f_y) dx \wedge dy + (h_y - g_z) dy \wedge dz + (h_x - f_z) dx \wedge dz = \\
 = \iint_D \frac{1}{\nu} \left((g_x - f_y) \frac{\partial(x, y)}{\partial(u, v)} + (h_y - g_z) \frac{\partial(y, z)}{\partial(u, v)} + (h_x - f_z) \frac{\partial(x, z)}{\partial(u, v)} \right) \nu dudv = \\
 = \iint_D \langle \text{rot } F, N \rangle \nu dudv = \iint_S \langle \text{rot } F, N \rangle d\sigma
 \end{aligned}$$

Si ottiene perciò che

$$(4.6) \quad \int_{\partial C} \langle F, T \rangle ds = \iint_C \langle \text{rot } F, N \rangle d\sigma$$

La 4.6 è nota come teorema del rotore.

Osserviamo che nella formula compare N e non N_e .

Le formule 4.6 e 4.4 assumono un'interessante aspetto nel caso in cui $F = \nabla\phi$.

Si ha infatti in tal caso che, se C_2 è una 2-varietà e C_3 è una 3-varietà

$$(4.7) \quad \int_{\partial C_2} \langle \nabla\phi, T \rangle ds = 0$$

$$\iint_{\partial C_3} \langle \nabla\phi, N_e \rangle ds = \iiint_{C_3} \Delta\phi dx dy dz$$

Possiamo dimostrare il teorema del rotore come segue.

DIMOSTRAZIONE. Proviamo il teorema nel caso in cui

$$F = (0, 0, h(x, y, z))$$

per cui

$$\text{rot } F = (h_y(x, y, z), -h_x(x, y, z), 0)$$

In tal caso si ha

$$\begin{aligned}
 \int_{\partial S} \langle F, T \rangle ds &= \int_{\partial S} h dz = \int_0^1 h(x(u, 0), y(u, 0), z(u, 0)) z_u(u, 0) du - \\
 &\quad - h(x(u, 1), y(u, 1), z(u, 1)) z_u(u, 1) du + \\
 &+ \int_0^1 h(x(1, v), y(1, v), z(1, v)) z_v(1, v) dv - \\
 &\quad - h(x(0, v), y(0, v), z(0, v)) z_v(0, v) dv = \\
 &= - \int_0^1 \int_0^1 \frac{d}{dv} h(x(u, v), y(u, v), z(u, v)) z_u(u, v) dv du + \\
 &\quad + \int_0^1 \int_0^1 \frac{d}{du} h(x(u, v), y(u, v), z(u, v)) z_v(u, v) du dv = \\
 &= \int_0^1 \int_0^1 -(h_x x_v + h_y y_v + h_z z_v) z_u - h z_{uv} + (h_x x_u + h_y y_u + h_z z_u) z_v + h z_{vu} du dv = \\
 &= \int_0^1 \int_0^1 h_x (x_u z_v - x_v z_u) + h_y (y_u z_v - y_v z_u) du dv = \\
 &= \int_0^1 \int_0^1 h_x \frac{\partial y, z}{\partial(u, v)} + h_y \frac{\partial(x, z)}{\partial(u, v)} du dv = \iint_{\partial S} \langle \text{rot } F, N \rangle d\sigma
 \end{aligned}$$



4.3. La formula di Green nel piano. Consideriamo la 1-forma in $\mathbb{R}^2$

$$\omega_1(x, y) = f(x, y)dx + g(x, y)dy$$

ed il campo vettoriale su $\mathbb{R}^2$

$$F = (f, g)$$

Avremo

$$d\omega_1 = (-f_y + g_x)dx \wedge dy$$

Sia $\gamma = \gamma(t)$ una 1-varietà in $\mathbb{R}^2$, cioè sia

$$\gamma : [a, b] \rightarrow \mathbb{R}^2$$

avremo

$$\partial\gamma = +\gamma(b) \cup -\gamma(a)$$

si ha

$$(4.8) \quad \int_{\partial\gamma} f dx + g dy = \int_{\gamma} (-f_y + g_x) dx \wedge dy.$$

Ora, se N è il versore normale a $\partial\gamma$,

$$N = \left(\frac{\dot{y}}{\sqrt{\dot{x}^2 + \dot{y}^2}}, -\frac{\dot{x}}{\sqrt{\dot{x}^2 + \dot{y}^2}} \right)$$

e pertanto

$$\begin{aligned} \int_{\partial\gamma} f dx - g dy &= \int_a^b (f\dot{x} + g\dot{y}) dt = \\ &= \int_a^b \langle F, N \rangle (\dot{x}^2 + \dot{y}^2) dt = \int_{\partial\gamma} \langle F, N \rangle ds \end{aligned}$$

La 4.8 è nota come teorema di Green e si può dimostrare come segue.

DIMOSTRAZIONE.

$$\begin{aligned}
 \int_{\Omega} g_x(x, y) dx dy &= \int_a^b \int_c^d g_x(x(t, s), y(t, s)) \left| \frac{\partial(x, y)}{\partial(t, s)} \right| dt ds = \\
 &= \int_a^b \int_c^d g_x(x, y) [x_t y_s - x_s y_t] dt ds = \\
 &= \int_a^b \int_c^d \left[g_x(x, y) x_t y_s - g_x(x, y) x_s y_t + \right. \\
 &\quad \left. + g_y(x, y) y_t y_s - g_y(x, y) y_s y_t \right] dt ds = \\
 &= \int_a^b \int_c^d \left[g_x(x, y) x_t + g_y(x, y) y_t \right] y_s - \\
 &\quad - \int_a^b \int_c^d \left[g_x(x, y) x_s + g_y(x, y) y_s \right] y_t dt ds = \\
 &= \int_c^d \left(\int_a^b \left(\frac{d}{dt} g(\phi(t, s)) \right) y_s(t, s) dt \right) ds - \\
 &\quad - \int_a^b \left(\int_c^d \left(\frac{d}{ds} g(\phi(t, s)) \right) y_s(t, s) ds \right) dt = \\
 &= \int_c^d \left([g(\phi(t, s)) y_s(t, s)]_{t=a}^{t=b} - \int_a^b g(\phi(t, s)) y_{st}(t, s) dt \right) ds - \\
 &\quad - \int_a^b \left([g(\phi(t, s)) y_t(t, s)]_{s=c}^{s=d} - \int_c^d g(\phi(t, s)) y_{ts}(t, s) ds \right) dt =
 \end{aligned}$$

□

4.4. Campi conservativi. Concludiamo questo paragrafo precisando i risultati ottenuti per le 1-forme in $\mathbb{R}^3$ nell'ambito dello studio dei campi vettoriali.

DEFINIZIONE 4.8. Sia $A \subset \mathbb{R}^3$ aperto e siano

$$f, g, h : A \longrightarrow \mathbb{R} \quad , \quad f, g, h \in C^2(A)$$

Sia

$$\omega = fdx + gdy + hdz$$

e

$$F = (f, g, h)$$

ω è la 1-forma differenziale corrispondente al campo vettoriale F e reciprocamente F è il campo vettoriale che corrisponde alla 1-forma ω .

Diamo ora alcune definizioni che sono la controparte relativa al campo F delle definizioni date in precedenza, per le 1-forme ω .

DEFINIZIONE 4.9. Sia $F : A \longrightarrow \mathbb{R}^3$, $A \subset \mathbb{R}^3$ aperto, un campo vettoriale $F = (f, g, h)$. Diciamo che F è un campo chiuso se

$$\operatorname{rot} F = (h_y - g_z, f_z - h_x, g_x - f_y) = (0, 0, 0)$$

Diciamo che F è conservativo, oppure che F ammette potenziale, se esiste

$$\phi : A \longrightarrow \mathbb{R}$$

tale che

$$\nabla \phi = F$$

In tal caso ϕ si chiama potenziale di F .

Osserviamo che, evidentemente, un campo vettoriale è chiuso o conservativo se e solo se la corrispondente 1-forma ω è chiusa o esatta, rispettivamente.

E' pertanto conseguenza dei precedenti risultati che

TEOREMA 4.6. Sia $F : A \longrightarrow \mathbb{R}^k$, $A \subset \mathbb{R}^k$ aperto; se F è conservativo allora F è chiuso.

Possiamo osservare che il teorema è immediata conseguenza del teorema di Schwarz.

Si può anche dimostrare che

TEOREMA 4.7. *Sia $F : A \longrightarrow \mathbb{R}^3$, $A \subset \mathbb{R}^3$ aperto, stellato. Se F è chiuso, allora F è conservativo.*

Il campo vettoriale $F : \mathbb{R}^2 \setminus \{(0, 0)\} \longrightarrow \mathbb{R}^2$ definito da

$$F(x, y) = \left(\frac{-y}{x^2 + y^2}, \frac{x}{x^2 + y^2} \right)$$

mostra che la condizione A stellato non è inessenziale. Essa può tuttavia essere un po' attenuata.

A questo scopo definiamo

DEFINIZIONE 4.10. *Sia $F : A \longrightarrow \mathbb{R}^3$, $A \subset \mathbb{R}^3$ aperto, e sia γ una 1-varietà in $\mathbb{R}^k$, $\gamma(t) = (x(t), y(t), z(t))$, $t \in [a, b]$; definiamo*

$$\begin{aligned} \int_{\gamma} F &= \int_{\gamma} \omega = \int_a^b f(\gamma(t))\dot{x}(t) + g(\gamma(t))\dot{y}(t) + h(\gamma(t))\dot{z}(t)dt = \\ &= \int_a^b \left\langle F(\gamma(t)), \frac{\dot{\gamma}(t)}{\|\dot{\gamma}(t)\|} \right\rangle \sqrt{\dot{x}^2(t) + \dot{y}^2(t) + \dot{z}^2(t)} dt = \int_{\gamma} \langle F, T_{\gamma} \rangle ds \end{aligned}$$

dove $T_{\gamma}(t)$ indica il vettore tangente unitario alla curva γ .

Possiamo dimostrare che

TEOREMA 4.8. *Sia $F : A \longrightarrow \mathbb{R}^k$, $A \subset \mathbb{R}^k$ aperto e connesso, un campo vettoriale; sono condizioni equivalenti*

- (1) F è conservativo;
- (2) $\int_{\gamma} F = 0$ su ogni curva chiusa γ contenuta in A ;
- (3) $\int_{\gamma} F$ dipende solo dagli estremi di γ .

(Per semplicità la funzione F e la funzione γ sono supposte di classe $\mathcal{C}^2$, ma è sufficiente $F \in \mathcal{C}^0$ e γ regolare a tratti.)

DIMOSTRAZIONE.

1) $\Rightarrow$ 2)

Se $F = \nabla\phi$ con $\phi : A \longrightarrow \mathbb{R}$, si ha

$$\begin{aligned}
 (4.9) \quad \int_{\gamma} F &= \int_a^b \phi_x(\gamma(t))\dot{x}(t) + \phi_y(\gamma(t))\dot{y}(t) + \phi_z(\gamma(t))\dot{z}(t)dt = \\
 &= \int_a^b \frac{d}{dt}\phi(\gamma(t))dt = \phi(\gamma(b)) - \phi(\gamma(a)) = 0
 \end{aligned}$$

2) $\Rightarrow$ 3). Siano γ_1 e γ_2 due curve con gli stessi estremi; allora $\gamma = \gamma_1 - \gamma_2$ è chiusa e

$$\int_{\gamma} F = \int_{\gamma_1} F - \int_{\gamma_2} F = 0$$

3) $\Rightarrow$ 1). Poiché vale 3), se γ è una qualunque curva avente per estremi $P = (x, y, z) \in A$ e $P_0 = (x_0, y_0, z_0) \in A$ fissato, possiamo definire

$$\phi(x) = \int_{\gamma} F$$

Si ha allora, ad esempio

$$\frac{\phi(x+t, y, z) - \phi(x, y, z)}{t} = \frac{1}{t} \int_{\gamma^*} F$$

essendo $\gamma^*(s) = (x+s, y, z)$, $0 \leq s \leq t$.

Pertanto

$$\frac{\phi(x+t, y, z) - \phi(x, y, z)}{t} = \frac{1}{t} \int_0^t f(x+s, y, z) ds = f(x+\sigma_t, y, z)$$

$0 \leq \sigma_t \leq t$ e

$$\lim_{t \rightarrow 0} \frac{\phi(x+t, y, z) - \phi(x, y, z)}{t} = \lim_{t \rightarrow 0} f(x+\sigma_t) = f(x)$$

□

DEFINIZIONE 4.11. Sia $A \subset \mathbb{R}^3$, diciamo che A è semplicemente connesso se ogni curva chiusa contenuta in A può essere deformata con continuità sino a ridursi ad un punto senza uscire da A .

Per la precisione, se $\gamma_0, \gamma_1 : [a, b] \longrightarrow A$ sono due curve chiuse, diciamo che γ_0 e γ_1 sono omotope se esiste $\psi \in \mathcal{C}^2$,

$$\psi : [0, 1] \times [a, b] \longrightarrow A$$

tale che, se $s \in [0, 1]$ e $t \in [a, b]$

$$\psi(0, t) = \gamma_0(t) , \quad \psi(1, t) = \gamma_1(t), \quad \psi(s, 0) = \psi(s, 1)$$

Diciamo che A è semplicemente connesso se ogni 1-varietà chiusa a valori in A è omotopa ad un punto di A .

TEOREMA 4.9. *Sia $F : A \longrightarrow \mathbb{R}^3$ un campo vettoriale chiuso, $A \subset \mathbb{R}^3$, e siano γ_0 e γ_1 due curve chiuse a valori in A , omotope, allora*

$$\int_{\gamma_0} F = \int_{\gamma_1} F.$$

DIMOSTRAZIONE. Poichè γ_0 e γ_1 sono omotope, esiste S tale che

$$S : [0, 1] \times [a, b] \longrightarrow A$$

tale che, se $u \in [0, 1]$ e $v \in [a, b]$

$$S(0, v) = \gamma_0(v) , \quad S(1, v) = \gamma_1(v), \quad S(u, 0) = S(u, 1)$$

Definiamo

$$\gamma_2(u) = S(u, 0) = S(u, 1) , \quad u \in [0, 1]$$

Per il teorema di Stokes avremo

$$\partial S = \gamma_2 + \gamma_1 - \gamma_2 - \gamma_0$$

per cui

$$\int_{\gamma_2} F + \int_{\gamma_1} F - \int_{\gamma_2} F - \int_{\gamma_0} F = \int_{\partial S} \omega = \int_S d\omega = 0$$

e si può concludere che

$$\int_{\gamma_1} F - \int_{\gamma_2} F = 0$$

□

COROLLARIO 4.1. *Sia $F : A \longrightarrow \mathbb{R}^k$ un campo vettoriale, $A \subset \mathbb{R}^k$ semplicemente connesso, allora F è conservativo se e solo se F è chiuso.*

Elenco delle figure

1.1		22
2.1	Il teorema degli zeri	31
2.2	Il teorema delle funzioni implicite	63
2.3	Il teorema delle funzioni implicite	64
3.1		76
3.2		76
3.3		86
4.1		107
4.2		118
4.3		121
4.4		123

4.5

126

Indice

Capitolo 1. SPAZI EUCLIDEI n -DIMENSIONALI.	3
1. Norma e Prodotto scalare	3
2. Applicazioni Lineari	12
3. Forme Bilineari e Quadratiche	15
4. Proprietà Topologiche	18
Capitolo 2. LE FUNZIONI DI PIÙ VARIABILI	23
1. Limiti	24
2. Continuità	28
3. Differenziabilità e Derivabilità	33
4. Formula di Taylor	53
5. Massimi e Minimi Relativi	56
6. Convessità	59
7. Funzioni Implicite	60

8. Massimi e Minimi Vincolati Moltiplicatori di Lagrange	68
Capitolo 3. INTEGRAZIONE DELLE FUNZIONI DI PIU' VARIABILI.	73
1. Integrali Multipli	74
1.1. Definizione di Integrale	74
1.2. Condizioni di Integrabilità - Proprietà degli Integrali	79
1.3. Formule di Riduzione	83
1.4. Misura di sottoinsiemi di $\mathbb{R}^3$	85
1.5. Integrazione su Domini Normali	90
1.6. Trasformazione di coordinate in $\mathbb{R}^3$	93
1.7. Integrali Impropri in $\mathbb{R}^3$	96
2. Integrali dipendenti da un parametro.	101
Capitolo 4. ARCHI E SUPERFICI NELLO SPAZIO.	105
1. Linee ed integrali di linea	105
1.1. Lunghezza di una Linea	106
1.2. Lunghezza d'Arco	110

1.3. Curvatura - Terna Intrinseca	114
2. Superfici ed Integrali di Superficie	120
3. Forme Differenziali in $\mathbb{R}^3$	127
3.0.1. Un'idea per rendersi conto dei segni	135
3.1. Primitive di 1-forme	144
3.2. Primitive di 2-forme	148
3.3. Primitive di 3-forme	151
4. Il Teorema di Stokes	152
4.1. Il teorema della divergenza in $\mathbb{R}^3$	153
4.2. Il Teorema del Rotore	162
4.3. La formula di Green nel piano	168
4.4. Campi conservativi	171
Elenco delle figure	179