

Analisi Matematica II

Appunti delle lezioni tenute dalla Prof.ssa R. Toader

Università di Trieste, CdL Ingegneria Industriale e CdL
Ingegneria Navale

a.a. 2024/2025

1 Lo spazio $\mathbb{R}^N$

Consideriamo l'insieme $\mathbb{R}^N$, costituito dalle N -uple $(x_1, x_2, \dots, x_N)$, dove $x_1, x_2, \dots, x_N$ sono numeri reali. Indicheremo i suoi elementi con i simboli

$$\mathbf{x}, \mathbf{x}', \mathbf{x}'', \dots$$

Cominciamo con l'introdurre un'operazione di addizione in $\mathbb{R}^N$: dati due elementi $\mathbf{x} = (x_1, x_2, \dots, x_N)$ e $\mathbf{x}' = (x'_1, x'_2, \dots, x'_N)$, si definisce $\mathbf{x} + \mathbf{x}'$ in questo modo:

$$\mathbf{x} + \mathbf{x}' = (x_1 + x'_1, x_2 + x'_2, \dots, x_N + x'_N).$$

Valgono le seguenti proprietà:

- a) (associativa) $(\mathbf{x} + \mathbf{x}') + \mathbf{x}'' = \mathbf{x} + (\mathbf{x}' + \mathbf{x}'')$;
- b) esiste un "elemento neutro" $\mathbf{0} = (0, 0, \dots, 0)$: si ha $\mathbf{x} + \mathbf{0} = \mathbf{x} = \mathbf{0} + \mathbf{x}$;
- c) ogni elemento $\mathbf{x} = (x_1, x_2, \dots, x_N)$ ha un "opposto"
 $(-\mathbf{x}) = (-x_1, -x_2, \dots, -x_N)$: si ha $\mathbf{x} + (-\mathbf{x}) = \mathbf{0} = (-\mathbf{x}) + \mathbf{x}$;
- d) (commutativa) $\mathbf{x} + \mathbf{x}' = \mathbf{x}' + \mathbf{x}$.

Pertanto, $(\mathbb{R}^N, +)$ è un "gruppo abeliano". Normalmente, si usa scrivere $\mathbf{x} - \mathbf{x}'$ per indicare $\mathbf{x} + (-\mathbf{x}')$.

Definiamo ora la moltiplicazione di un elemento di $\mathbb{R}^N$ per un numero reale: considerati $\mathbf{x} = (x_1, x_2, \dots, x_N) \in \mathbb{R}^N$ e un numero reale $\alpha \in \mathbb{R}$, si definisce $\alpha\mathbf{x}$ in questo modo:

$$\alpha\mathbf{x} = (\alpha x_1, \alpha x_2, \dots, \alpha x_N).$$

Valgono le seguenti proprietà:

- a) $\alpha(\beta \mathbf{x}) = (\alpha\beta)\mathbf{x}$;
- b) $(\alpha + \beta)\mathbf{x} = (\alpha\mathbf{x}) + (\beta\mathbf{x})$;
- c) $\alpha(\mathbf{x} + \mathbf{x}') = (\alpha\mathbf{x}) + (\alpha\mathbf{x}')$;
- d) $1\mathbf{x} = \mathbf{x}$.

Pertanto, con le operazioni introdotte, $\mathbb{R}^N$ è uno “spazio vettoriale”. Chiameremo i suoi elementi “vettori”; i numeri reali, in questo ambito, verranno chiamati “scalari”.

È utile introdurre il “prodotto scalare” tra due vettori: dati $\mathbf{x} = (x_1, x_2, \dots, x_N)$ e $\mathbf{x}' = (x'_1, x'_2, \dots, x'_N)$, si definisce il numero reale $\mathbf{x} \cdot \mathbf{x}'$ in questo modo:

$$\mathbf{x} \cdot \mathbf{x}' = \sum_{k=1}^N x_k x'_k.$$

Il prodotto scalare è spesso indicato con simboli diversi, quali ad esempio

$$\langle \mathbf{x} | \mathbf{x}' \rangle, \quad \langle \mathbf{x}, \mathbf{x}' \rangle, \quad (\mathbf{x} | \mathbf{x}'), \quad (\mathbf{x}, \mathbf{x}').$$

Valgono le seguenti proprietà:

- a) $\mathbf{x} \cdot \mathbf{x} \geq 0$;
- b) $\mathbf{x} \cdot \mathbf{x} = 0 \Leftrightarrow \mathbf{x} = \mathbf{0}$;
- c) $(\mathbf{x} + \mathbf{x}') \cdot \mathbf{x}'' = (\mathbf{x} \cdot \mathbf{x}'') + (\mathbf{x}' \cdot \mathbf{x}'')$;
- d) $(\alpha\mathbf{x}) \cdot \mathbf{x}' = \alpha(\mathbf{x} \cdot \mathbf{x}')$;
- e) $\mathbf{x} \cdot \mathbf{x}' = \mathbf{x}' \cdot \mathbf{x}$;

Se $\mathbf{x} \cdot \mathbf{x}' = 0$, si dice che i due vettori $\mathbf{x}$ e $\mathbf{x}'$ sono “ortogonali”.

Definiamo infine, solo nello spazio tridimensionale $\mathbb{R}^3$, il “prodotto vettoriale” tra due vettori: dati $\mathbf{x} = (x_1, x_2, x_3)$ e $\mathbf{x}' = (x'_1, x'_2, x'_3)$, si definisce il vettore $\mathbf{x} \times \mathbf{x}'$ in questo modo:

$$\mathbf{x} \times \mathbf{x}' = (x_2 x'_3 - x_3 x'_2, x_3 x'_1 - x_1 x'_3, x_1 x'_2 - x_2 x'_1).$$

Si può verificare che $\mathbf{x} \times \mathbf{x}'$ è ortogonale sia a $\mathbf{x}$ che a $\mathbf{x}'$:

$$(\mathbf{x} \times \mathbf{x}') \cdot \mathbf{x} = 0, \quad (\mathbf{x} \times \mathbf{x}') \cdot \mathbf{x}' = 0.$$

Inoltre, se $\mathbf{x} \times \mathbf{x}' \neq \mathbf{0}$, siccome

$$\det \begin{pmatrix} x_1 & x_2 & x_3 \\ x'_1 & x'_2 & x'_3 \\ x_2 x'_3 - x_3 x'_2 & x_3 x'_1 - x_1 x'_3 & x_1 x'_2 - x_2 x'_1 \end{pmatrix} > 0,$$

la direzione di $\mathbf{x} \times \mathbf{x}'$ è individuata dalla cosiddetta “regola della mano destra”. Ciò significa che la terna $(\mathbf{x}, \mathbf{x}', \mathbf{x} \times \mathbf{x}')$ ha la stessa orientazione di $(\mathbf{e}_1, \mathbf{e}_2, \mathbf{e}_3)$, la base canonica

$$\mathbf{e}_1 = (1, 0, 0), \quad \mathbf{e}_2 = (0, 1, 0), \quad \mathbf{e}_3 = (0, 0, 1).$$

Elenchiamo alcune proprietà del prodotto vettoriale:

- a) $(\mathbf{x} + \mathbf{x}') \times \mathbf{x}'' = (\mathbf{x} \times \mathbf{x}'') + (\mathbf{x}' \times \mathbf{x}'')$;
- b) $(\alpha \mathbf{x}) \times \mathbf{x}' = \alpha(\mathbf{x} \times \mathbf{x}')$;
- c) $\mathbf{x} \times \mathbf{x}' = -\mathbf{x}' \times \mathbf{x}$.

Si può anche dimostrare l'**Identità di Jacobi**

$$\mathbf{x} \times (\mathbf{x}' \times \mathbf{x}'') + \mathbf{x}' \times (\mathbf{x}'' \times \mathbf{x}) + \mathbf{x}'' \times (\mathbf{x} \times \mathbf{x}') = \mathbf{0}.$$

Si noti infine che

$$\mathbf{e}_1 \times \mathbf{e}_2 = \mathbf{e}_3, \quad \mathbf{e}_2 \times \mathbf{e}_3 = \mathbf{e}_1, \quad \mathbf{e}_3 \times \mathbf{e}_1 = \mathbf{e}_2.$$

1.1 Norma e distanza euclidea in $\mathbb{R}^N$

A partire dal prodotto scalare, possiamo definire la “norma” di un vettore $\mathbf{x} = (x_1, x_2, \dots, x_N)$:

$$\|\mathbf{x}\| = \sqrt{\mathbf{x} \cdot \mathbf{x}} = \sqrt{\sum_{k=1}^N x_k^2}.$$

Valgono le seguenti proprietà:

- a) $\|\mathbf{x}\| \geq 0$;
- b) $\|\mathbf{x}\| = 0 \Leftrightarrow \mathbf{x} = \mathbf{0}$;
- c) $\|\alpha \mathbf{x}\| = |\alpha| \|\mathbf{x}\|$;
- d) $\|\mathbf{x} + \mathbf{x}'\| \leq \|\mathbf{x}\| + \|\mathbf{x}'\|$.

Per dimostrare la d), detta “disuguaglianza triangolare” per la norma, abbiamo bisogno della seguente **disuguaglianza di Schwarz**.

Teorema 1 *Presi due vettori $\mathbf{x}, \mathbf{x}'$, si ha*

$$|\mathbf{x} \cdot \mathbf{x}'| \leq \|\mathbf{x}\| \|\mathbf{x}'\|.$$

Dimostrazione. La disuguaglianza è sicuramente verificata se $\mathbf{x}' = \mathbf{0}$, essendo in tal caso $\mathbf{x} \cdot \mathbf{x}' = 0$ e $\|\mathbf{x}'\| = 0$. Supponiamo quindi $\mathbf{x}' \neq \mathbf{0}$. Per ogni $\alpha \in \mathbb{R}$, si ha

$$0 \leq \|\mathbf{x} - \alpha \mathbf{x}'\|^2 = (\mathbf{x} - \alpha \mathbf{x}') \cdot (\mathbf{x} - \alpha \mathbf{x}') = \|\mathbf{x}\|^2 - 2\alpha \mathbf{x} \cdot \mathbf{x}' + \alpha^2 \|\mathbf{x}'\|^2.$$

Prendendo $\alpha = \frac{1}{\|\mathbf{x}'\|^2} \mathbf{x} \cdot \mathbf{x}'$, si ottiene

$$0 \leq \|\mathbf{x}\|^2 - 2 \frac{1}{\|\mathbf{x}'\|^2} (\mathbf{x} \cdot \mathbf{x}')^2 + \frac{1}{\|\mathbf{x}'\|^4} (\mathbf{x} \cdot \mathbf{x}')^2 \|\mathbf{x}'\|^2 = \|\mathbf{x}\|^2 - \frac{1}{\|\mathbf{x}'\|^2} (\mathbf{x} \cdot \mathbf{x}')^2,$$

da cui la tesi.¹ ■

¹Lo stesso risultato si ottiene osservando che, se $a\alpha^2 + b\alpha + c \geq 0$ per ogni $\alpha \in \mathbb{R}$, con $a > 0$, allora $\Delta = b^2 - 4ac \leq 0$.

Dimostriamo ora la proprietà d) della norma, usando la disuguaglianza di Schwarz:

$$\begin{aligned}\|\mathbf{x} + \mathbf{x}'\|^2 &= (\mathbf{x} + \mathbf{x}') \cdot (\mathbf{x} + \mathbf{x}') \\ &= \|\mathbf{x}\|^2 + 2\mathbf{x} \cdot \mathbf{x}' + \|\mathbf{x}'\|^2 \\ &\leq \|\mathbf{x}\|^2 + 2\|\mathbf{x}\| \|\mathbf{x}'\| + \|\mathbf{x}'\|^2 \\ &= (\|\mathbf{x}\| + \|\mathbf{x}'\|)^2,\end{aligned}$$

da cui la disuguaglianza cercata.

Notiamo ancora la seguente **identità del parallelogramma**, di semplice verifica:

$$\|\mathbf{x} + \mathbf{x}'\|^2 + \|\mathbf{x} - \mathbf{x}'\|^2 = 2(\|\mathbf{x}\|^2 + \|\mathbf{x}'\|^2).$$

Definiamo ora, a partire dalla norma, la “distanza euclidea” tra due vettori $\mathbf{x} = (x_1, x_2, \dots, x_N)$ e $\mathbf{x}' = (x'_1, x'_2, \dots, x'_N)$:

$$d(\mathbf{x}, \mathbf{x}') = \|\mathbf{x} - \mathbf{x}'\| = \sqrt{\sum_{k=1}^N (x_k - x'_k)^2}.$$

Valgono le seguenti proprietà:

- a) $d(\mathbf{x}, \mathbf{x}') \geq 0$;
- b) $d(\mathbf{x}, \mathbf{x}') = 0 \Leftrightarrow \mathbf{x} = \mathbf{x}'$;
- c) $d(\mathbf{x}, \mathbf{x}') = d(\mathbf{x}', \mathbf{x})$;
- d) $d(\mathbf{x}, \mathbf{x}'') \leq d(\mathbf{x}, \mathbf{x}') + d(\mathbf{x}', \mathbf{x}'')$.

Quest’ultima viene spesso chiamata “disuguaglianza triangolare”; la dimostriamo:

$$\begin{aligned}d(\mathbf{x}, \mathbf{x}'') &= \|\mathbf{x} - \mathbf{x}''\| \\ &= \|(\mathbf{x} - \mathbf{x}') + (\mathbf{x}' - \mathbf{x}'')\| \\ &\leq \|\mathbf{x} - \mathbf{x}'\| + \|\mathbf{x}' - \mathbf{x}''\| \\ &= d(\mathbf{x}, \mathbf{x}') + d(\mathbf{x}', \mathbf{x}'').\end{aligned}$$

2 Spazi metrici

Definizione 1 Dato un insieme non vuoto E , una funzione $d : E \times E \rightarrow \mathbb{R}$ si chiama “distanza” (su E) se soddisfa alle seguenti proprietà:

- a) $d(x, x') \geq 0$;
- b) $d(x, x') = 0 \Leftrightarrow x = x'$;
- c) $d(x, x') = d(x', x)$;
- d) $d(x, x'') \leq d(x, x') + d(x', x'')$

(la disuguaglianza triangolare). L’insieme E , dotato della distanza d , si dice “spazio metrico”. I suoi elementi verranno spesso chiamati “punti”.

Abbiamo visto che $\mathbb{R}^N$, dotato della distanza euclidea, è uno spazio metrico (nel seguito, parlando dello spazio metrico $\mathbb{R}^N$, se non altrimenti specificato sottintenderemo che la distanza sia sempre quella euclidea). Nel caso $N = 1$, abbiamo la distanza usuale su $\mathbb{R}$: $d(\alpha, \beta) = |\alpha - \beta|$.

È però possibile considerare diverse distanze su uno stesso insieme. Ad esempio, presi due vettori $\mathbf{x} = (x_1, x_2, \dots, x_N)$ e $\mathbf{x}' = (x'_1, x'_2, \dots, x'_N)$, la funzione

$$d_*(\mathbf{x}, \mathbf{x}') = \sum_{k=1}^N |x_k - x'_k|$$

rappresenta anch'essa una distanza in $\mathbb{R}^N$. Lo stesso dicasi per la funzione

$$d_{**}(\mathbf{x}, \mathbf{x}') = \max\{|x_k - x'_k| : k = 1, 2, \dots, N\}.$$

Oppure, si può definire la seguente:

$$\hat{d}(\mathbf{x}, \mathbf{x}') = \begin{cases} 0 & \text{se } \mathbf{x} = \mathbf{x}', \\ 1 & \text{se } \mathbf{x} \neq \mathbf{x}'. \end{cases}$$

Anche questa è una distanza, per quanto strana possa sembrare.

Dati $x_0 \in E$ e un numero $\rho > 0$, definiamo la palla aperta di centro x_0 e raggio ρ :

$$B(x_0, \rho) = \{x \in E : d(x, x_0) < \rho\};$$

analogamente definiamo la palla chiusa

$$\overline{B}(x_0, \rho) = \{x \in E : d(x, x_0) \leq \rho\},$$

e la sfera

$$S(x_0, \rho) = \{x \in E : d(x, x_0) = \rho\}.$$

In $\mathbb{R}$, ogni intervallo $]a, b[$ è una palla aperta e ogni intervallo $[a, b]$ è una palla chiusa: si ha

$$]a, b[= B\left(\frac{a+b}{2}, \frac{b-a}{2}\right), \quad [a, b] = \overline{B}\left(\frac{a+b}{2}, \frac{b-a}{2}\right).$$

Una sfera in $\mathbb{R}$ è quindi costituita da due soli punti.

In $\mathbb{R}^2$, con la distanza euclidea, una palla è un cerchio: la palla aperta non comprende i punti della circonferenza esterna, la palla chiusa sì. Una sfera è semplicemente una circonferenza.

Se in $\mathbb{R}^2$ consideriamo la distanza d_* definita in precedenza, una palla sarà un quadrato, con i lati inclinati di 45 gradi, avente x_0 come punto centrale. Una sfera sarà il perimetro di tale quadrato. Se invece consideriamo la distanza d_{**} , la palla sarà ancora un quadrato, ma con i lati paralleli agli assi cartesiani.

Se invece prendiamo la distanza $\hat{d}$, su un qualsiasi insieme E , allora

$$B(x_0, \rho) = \begin{cases} \{x_0\} & \text{se } \rho \leq 1, \\ E & \text{se } \rho > 1, \end{cases} \quad \overline{B}(x_0, \rho) = \begin{cases} \{x_0\} & \text{se } \rho < 1, \\ E & \text{se } \rho \geq 1, \end{cases}$$

per cui

$$S(x_0, \rho) = \begin{cases} E \setminus \{x_0\} & \text{se } \rho = 1, \\ \emptyset & \text{se } \rho \neq 1. \end{cases}$$

Definizione 2 Un insieme $U \subseteq E$ si dice “intorno” di un punto x_0 se esiste un $\rho > 0$ tale che $B(x_0, \rho) \subseteq U$; in tal caso, il punto x_0 si dice “interno” ad U . L'insieme dei punti interni ad U si chiama “l'interno” di U e si denota con $\overset{\circ}{U}$. Chiaramente, si ha sempre $\overset{\circ}{U} \subseteq U$. Si dice che U è un “insieme aperto” se coincide con il suo interno, ossia se $\overset{\circ}{U} = U$.

Teorema 2 Una palla aperta è un insieme aperto.

Dimostrazione. Sia $B(x_0, \rho)$ la palla in questione; prendiamo un $x_1 \in B(x_0, \rho)$. Scelto $r > 0$ tale che $r \leq \rho - d(x_0, x_1)$, si ha che $B(x_1, r) \subseteq B(x_0, \rho)$; infatti, se $x \in B(x_1, r)$, allora

$$d(x, x_0) \leq d(x, x_1) + d(x_1, x_0) < r + d(x_1, x_0) \leq \rho,$$

per cui $x \in B(x_0, \rho)$. Abbiamo quindi dimostrato che ogni punto x_1 di $B(x_0, \rho)$ è interno a $B(x_0, \rho)$. ■

Consideriamo ora tre esempi particolari: nel primo, l'insieme U coincide con E ; nel secondo, U è l'insieme vuoto; nel terzo, esso è costituito da un unico punto.

Ogni punto di E è interno all'insieme E stesso, in quanto ogni palla è per definizione contenuta in E . Quindi, l'interno di E coincide con tutto E , ossia $\overset{\circ}{E} = E$. Questo significa che E è un insieme aperto.

L'insieme vuoto non può avere punti interni. Quindi, l'interno di $\emptyset$, non avendo elementi, è vuoto. In altri termini, $\overset{\circ}{\emptyset} = \emptyset$, il che significa che $\emptyset$ è anch'esso un insieme aperto.

L'insieme $U = \{x_0\}$, costituito da un unico punto, in generale non è un insieme aperto (ad esempio in $\mathbb{R}^N$ con la distanza euclidea), ma può esserlo in casi particolari, quando x_0 è un “punto isolato” di E . Ad esempio, se si considera la distanza $\hat{d}$, oppure in $\mathbb{N}$, tutti i punti sono isolati.

Teorema 3 L'interno di un insieme è un insieme aperto.

Dimostrazione. Se $\overset{\circ}{U}$ è vuoto, la tesi è sicuramente vera. Supponiamo allora che $\overset{\circ}{U}$ sia non vuoto. Sia $x_1 \in \overset{\circ}{U}$. Allora esiste un $\rho > 0$ tale che $B(x_1, \rho) \subseteq U$. Se proviamo che $B(x_1, \rho) \subseteq \overset{\circ}{U}$ la dimostrazione sarà completa, in quanto avremo dimostrato che ogni punto x_1 di $\overset{\circ}{U}$ è interno a $\overset{\circ}{U}$.

Per dimostrare che $B(x_1, \rho) \subseteq \overset{\circ}{U}$, sia x un elemento di $B(x_1, \rho)$. Siccome $B(x_1, \rho)$ è un insieme aperto, esiste un $r > 0$ tale che $B(x, r) \subseteq B(x_1, \rho)$. Allora $B(x, r) \subseteq U$, per cui x appartiene a $\overset{\circ}{U}$. La dimostrazione è completa. ■

Si può dimostrare la seguente implicazione:

$$U_1 \subseteq U_2 \quad \Rightarrow \quad \overset{\circ}{U}_1 \subseteq \overset{\circ}{U}_2 .$$

Da essa segue che $\overset{\circ}{U}$ è il più grande insieme aperto contenuto in U : se A è un aperto e $A \subseteq U$, allora $A \subseteq \overset{\circ}{U}$.

Definizione 3 Diremo che il punto x_0 è “aderente” all’insieme U se per ogni $\rho > 0$ si ha che $B(x_0, \rho) \cap U \neq \emptyset$. L’insieme dei punti aderenti ad U si chiama “la chiusura” di U e si denota con $\overline{U}$. Chiaramente, si ha sempre $U \subseteq \overline{U}$. Si dice che U è un “insieme chiuso” se coincide con la sua chiusura, ossia se $U = \overline{U}$.

Teorema 4 Una palla chiusa è un insieme chiuso.

Dimostrazione. Sia $U = \overline{B}(x_0, \rho)$ la palla in questione; voglio dimostrare che $\overline{U} \subseteq U$. A tal fine vedremo che $\mathcal{CU} \subseteq \mathcal{C}\overline{U}$.² Prendiamo un $x_1 \in \mathcal{CU}$, ossia $x_1 \notin \overline{B}(x_0, \rho)$. Scelto $r > 0$ tale che $r \leq d(x_0, x_1) - \rho$, si ha che $B(x_1, r) \cap \overline{B}(x_0, \rho) = \emptyset$; infatti, se per assurdo esistesse un $x \in B(x_1, r) \cap \overline{B}(x_0, \rho)$, allora si avrebbe

$$d(x_0, x_1) \leq d(x_0, x) + d(x, x_1) < \rho + r ,$$

in contrasto con la scelta fatta per r . Quindi, $x_1 \notin \overline{U}$, ossia $x_1 \in \mathcal{C}\overline{U}$. ■

Essendo E il l’insieme universo, ogni punto aderente ad E deve comunque appartenere ad E stesso. Quindi, la chiusura di E coincide con E , ossia $\overline{E} = E$. Questo significa che E è un insieme chiuso.

Notiamo che non esiste alcun punto aderente all’insieme $\emptyset$. Infatti, qualsiasi sia il punto x_0 , per ogni $\rho > 0$ si ha che $B(x_0, \rho) \cap \emptyset = \emptyset$. Quindi, la chiusura di $\emptyset$, non avendo elementi, è vuota. In altri termini, $\overline{\emptyset} = \emptyset$, il che significa che $\emptyset$ è un insieme chiuso.

L’insieme $U = \{x_0\}$, costituito da un unico punto, è sempre un insieme chiuso. Infatti, preso un $x_1 \notin U$, scegliendo $\rho > 0$ tale che $\rho < d(x_0, x_1)$ si ha che $B(x_1, \rho) \cap U = \emptyset$, per cui x_1 non è aderente ad U .

Teorema 5 La chiusura di un insieme è un insieme chiuso.

Dimostrazione. Poniamo $V = \overline{U}$. Se $V = E$, la tesi è verificata. Supponiamo quindi che sia $V \neq E$. Sia $x_1 \notin V$. Allora esiste un $\rho > 0$ tale che $B(x_1, \rho) \cap U = \emptyset$. Vediamo che anche $B(x_1, \rho) \cap V = \emptyset$. Infatti, se per assurdo ci fosse un $x \in B(x_1, \rho) \cap V$, essendo $B(x_1, \rho)$ un insieme aperto, esisterebbe un $r > 0$ tale che $B(x, r) \subseteq B(x_1, \rho)$. Siccome $x \in V = \overline{U}$, dovrebbe essere $B(x, r) \cap U \neq \emptyset$ e quindi anche $B(x_1, \rho) \cap U \neq \emptyset$, in contraddizione con quanto sopra. Quindi, nessun punto x_1 al di fuori di V può essere aderente a V . In altri termini, V contiene tutti i punti ad esso aderenti, pertanto è chiuso. ■

²Denotiamo con $\mathcal{CU}$ il complementare di U in E , ossia l’insieme $E \setminus U$.

Si può dimostrare che

$$U_1 \subseteq U_2 \quad \Rightarrow \quad \overline{U}_1 \subseteq \overline{U}_2.$$

Da questa implicazione segue che $\overline{U}$ è il più piccolo insieme chiuso che contiene U : se C è un chiuso e $C \supseteq U$, allora $C \supseteq \overline{U}$.

Si può dimostrare che l'unione e l'intersezione di due insiemi aperti [chiusi] sono insiemi aperti [chiusi]. Lo stesso vale per un numero finito di insiemi aperti [chiusi]: lo si dimostra per induzione. Se invece si considera un numero infinito di insiemi, la cosa cambia. L'unione di un numero infinito di insiemi aperti è un insieme aperto, l'intersezione in generale non lo è. Ad esempio, in $\mathbb{R}$, prendendo gli aperti

$$A_n = \left] -\frac{1}{n+1}, \frac{1}{n+1} \right[,$$

con $n \in \mathbb{N}$, la loro intersezione è $\{0\}$, che non è un aperto. Analogamente, l'intersezione di un numero infinito di insiemi chiusi è un insieme chiuso, mentre l'unione in generale non lo è. Ad esempio, considerando i chiusi

$$C_n = \left[-1 + \frac{1}{n+1}, 1 - \frac{1}{n+1} \right] ,$$

con $n \in \mathbb{N}$, la loro unione è l'intervallo $] -1, 1[$, che non è un chiuso.

Cercheremo ora di capire le analogie incontrate tra le nozioni di interno e chiusura di un insieme, e quelle di insieme aperto e chiuso.

Teorema 6 *Valgono le seguenti relazioni:*

$$\overline{\mathcal{C}U} = \mathcal{C}\mathring{U}, \quad (\mathcal{C}U) = \overline{\mathcal{C}U}.$$

Dimostrazione. Vediamo la prima uguaglianza. Se $U = E$, allora $\mathcal{C}U = \emptyset$, per cui $\overline{\mathcal{C}U} = \emptyset$; d'altra parte, $\mathring{U} = E$, per cui $\mathcal{C}\mathring{U} = \emptyset$. L'uguaglianza è così verificata in questo caso. Supponiamo ora che sia $U \neq E$, per cui $\mathcal{C}U \neq \emptyset$. Si ha:

$$\begin{aligned} x \in \overline{\mathcal{C}U} &\Leftrightarrow \forall \rho > 0 \quad B(x, \rho) \cap \mathcal{C}U \neq \emptyset \\ &\Leftrightarrow \forall \rho > 0 \quad B(x, \rho) \not\subseteq U \\ &\Leftrightarrow x \notin \mathring{U} \\ &\Leftrightarrow x \in \mathcal{C}\mathring{U}. \end{aligned}$$

Questo dimostra la prima uguaglianza. Possiamo ora usarla per dedurre la seguente:

$$\mathcal{C}(\mathcal{C}U) = \overline{\mathcal{C}(\mathcal{C}U)} = \overline{U}.$$

Passando ai complementari, si ottiene la seconda uguaglianza. ■

Abbiamo quindi che

$$\overline{U} = \mathcal{C}(\mathring{U}), \quad \mathring{U} = \mathcal{C}(\overline{U}).$$

Come immediato corollario, abbiamo il seguente.

Teorema 7 *Un insieme è aperto [chiuso] se e solo se il suo complementare è chiuso [aperto].*

Dimostrazione. Se U è aperto, $U = \mathring{U}$ e quindi

$$\overline{\mathcal{C}U} = \mathcal{C}\mathring{U} = \mathcal{C}U,$$

per cui $\mathcal{C}U$ è chiuso.

Se U è chiuso, $U = \overline{U}$ e quindi

$$(\mathring{\mathcal{C}U}) = \mathcal{C}\overline{U} = \mathcal{C}U,$$

per cui $\mathcal{C}U$ è aperto. ■

Nota. Nel seguito, qualora non specificato altrimenti, quando parleremo di $\mathbb{R}^N$ come spazio metrico sarà sempre sottinteso che la distanza e la norma su di esso considerate siano quelle euclidee.

3 Funzioni continue

Siano E ed F due spazi metrici, con le loro distanze d_E e d_F , rispettivamente. Sia x_0 un punto di E e $f : E \rightarrow F$ una funzione.

Definizione 4 *Diremo che f è “continua” in x_0 se, comunque preso un numero positivo ε , è possibile trovare un numero positivo δ tale che, se x è un qualsiasi elemento del dominio E che disti da x_0 per meno di δ , allora $f(x)$ dista da $f(x_0)$ per meno di ε . In simboli:*

$$\forall \varepsilon > 0 \quad \exists \delta > 0 : \forall x \in E \quad d_E(x, x_0) < \delta \Rightarrow d_F(f(x), f(x_0)) < \varepsilon.$$

In questa formulazione, spesso la scrittura “ $\forall x \in E$ ” verrà sottintesa.

Osserviamo che una o entrambe le disuguaglianze

$$d_E(x, x_0) < \delta, \quad d_F(f(x), f(x_0)) < \varepsilon$$

possono essere sostituite rispettivamente da

$$d_E(x, x_0) \leq \delta, \quad d_F(f(x), f(x_0)) \leq \varepsilon$$

ottenendo definizioni che sono tutte tra loro equivalenti. Questo è dovuto al fatto, da un lato, che ε è un *qualunque* numero positivo e, dall’altro lato, che se l’implicazione della definizione vale per un certo numero positivo δ , essa vale a maggior ragione prendendo al posto di quel δ un qualsiasi numero positivo più piccolo.

Una rilettura della definizione di continuità ci mostra che f è continua in x_0 se e solo se:

$$\forall \varepsilon > 0 \quad \exists \delta > 0 : f(B(x_0, \delta)) \subseteq B(f(x_0), \varepsilon).$$

Inoltre, è del tutto equivalente considerare una palla chiusa al posto di una palla aperta; risulta inoltre utile la seguente formulazione equivalente, per cui f è continua in x_0 se e solo se:

per ogni intorno V di $f(x_0)$ esiste un intorno U di x_0 tale che $f(U) \subseteq V$.

Nel caso in cui la funzione f sia continua in ogni punto x_0 del dominio E , diremo che “ f è continua su E ”, o semplicemente “ f è continua”.

Vediamo ora alcuni esempi.

Esempio 1. La funzione costante: per un certo $\bar{c} \in F$, si ha che $f(x) = \bar{c}$, per ogni $x \in E$. Essendo $d_F(f(x), f(x_0)) = d_F(\bar{c}, \bar{c}) = 0$ per ogni $x \in E$, tale funzione è chiaramente continua (ogni scelta di $\delta > 0$ va bene).

Esempio 2. Siano $E = \mathbb{R}^N$ e $F = \mathbb{R}^N$. Fissato un numero $\alpha \in \mathbb{R}$, consideriamo la funzione $f : \mathbb{R}^N \rightarrow \mathbb{R}^N$ definita da $f(\mathbf{x}) = \alpha \mathbf{x}$. Vediamo che è continua. Infatti, se $\alpha = 0$, si tratta della funzione costante con valore $\mathbf{0}$, e sappiamo che tale funzione è continua. Sia ora $\alpha \neq 0$. Allora, fissato $\varepsilon > 0$, essendo

$$\|f(\mathbf{x}) - f(\mathbf{x}_0)\| = \|\alpha \mathbf{x} - \alpha \mathbf{x}_0\| = \|\alpha(\mathbf{x} - \mathbf{x}_0)\| = |\alpha| \|\mathbf{x} - \mathbf{x}_0\|,$$

basta prendere $\delta = \frac{\varepsilon}{|\alpha|}$ per avere l’implicazione

$$\|\mathbf{x} - \mathbf{x}_0\| < \delta \Rightarrow \|f(\mathbf{x}) - f(\mathbf{x}_0)\| < \varepsilon.$$

Esempio 3. Siano $E = \mathbb{R}^N$ e $F = \mathbb{R}$. Vediamo che la funzione $f : \mathbb{R}^N \rightarrow \mathbb{R}$ definita da $f(\mathbf{x}) = \|\mathbf{x}\|$ è continua su $\mathbb{R}^N$. Questo seguirà facilmente dalla disuguaglianza

$$\left| \|\mathbf{x}\| - \|\mathbf{x}'\| \right| \leq \|\mathbf{x} - \mathbf{x}'\|,$$

che ora dimostriamo. Si ha:

$$\begin{aligned} \|\mathbf{x}\| &= \|(\mathbf{x} - \mathbf{x}') + \mathbf{x}'\| \leq \|\mathbf{x} - \mathbf{x}'\| + \|\mathbf{x}'\|, \\ \|\mathbf{x}'\| &= \|(\mathbf{x}' - \mathbf{x}) + \mathbf{x}\| \leq \|\mathbf{x}' - \mathbf{x}\| + \|\mathbf{x}\|. \end{aligned}$$

Essendo $\|\mathbf{x} - \mathbf{x}'\| = \|\mathbf{x}' - \mathbf{x}\|$, si ha che

$$\|\mathbf{x}\| - \|\mathbf{x}'\| \leq \|\mathbf{x} - \mathbf{x}'\| \quad \text{e} \quad \|\mathbf{x}'\| - \|\mathbf{x}\| \leq \|\mathbf{x} - \mathbf{x}'\|,$$

da cui la disuguaglianza cercata. A questo punto, considerato un $\mathbf{x}_0 \in \mathbb{R}^N$ e fissato un $\varepsilon > 0$, basta prendere $\delta = \varepsilon$ per avere che

$$\|\mathbf{x} - \mathbf{x}_0\| < \delta \Rightarrow \left| \|\mathbf{x}\| - \|\mathbf{x}_0\| \right| < \varepsilon.$$

I tre teoremi seguenti hanno dimostrazione analoga a quella vista nel corso di Analisi 1, in cui si supposeva che E fosse un sottoinsieme di $\mathbb{R}$.

Teorema 8 Se $f, g : E \rightarrow \mathbb{R}$ sono continue in x_0 e $\alpha \in \mathbb{R}$, anche αf , $f + g$, $f - g$, $f \cdot g$ sono continue in x_0 .

Teorema 9 (della permanenza del segno) Sia $g : E \rightarrow \mathbb{R}$ continua in x_0 . Se $g(x_0) > 0$, allora esiste un intorno U di x_0 per cui $g(x) > 0$ per ogni $x \in U$. Se invece $g(x_0) < 0$, allora esiste un intorno U di x_0 per cui $g(x) < 0$ per ogni $x \in U$.

Teorema 10 Se $f, g : E \rightarrow \mathbb{R}$ sono continue in x_0 e $g(x_0) \neq 0$, anche $\frac{f}{g}$ è continua in x_0 .

Vediamo ora come si comporta una funzione composta di due funzioni continue.

Teorema 11 Siano E, F, G tre spazi metrici, $f : E \rightarrow F$ continua in x_0 e $g : F \rightarrow G$ continua in $f(x_0)$; allora $g \circ f$ è continua in x_0 .

Dimostrazione. Fissato un intorno W di $[g \circ f](x_0) = g(f(x_0))$, per la continuità di g in $f(x_0)$ esiste un intorno V di $f(x_0)$ tale che $g(V) \subseteq W$. Allora, per la continuità di f in x_0 , esiste un intorno U di x_0 tale che $f(U) \subseteq V$. Ne segue che $[g \circ f](U) \subseteq W$. ■

Consideriamo ora, per ogni $k = 1, 2, \dots, D$, la funzione “ k -esima proiezione” $p_k : \mathbb{R}^D \rightarrow \mathbb{R}$ definita da

$$p_k(x_1, x_2, \dots, x_D) = x_k.$$

Teorema 12 Le funzioni p_k sono continue.

Dimostrazione. Consideriamo un punto $\mathbf{x}_0 = (x_1^0, x_2^0, \dots, x_D^0) \in \mathbb{R}^D$ e fissiamo $\varepsilon > 0$. Notiamo che, per ogni $\mathbf{x} = (x_1, x_2, \dots, x_D) \in \mathbb{R}^D$, si ha

$$|x_k - x_k^0| \leq \sqrt{\sum_{j=1}^D (x_j - x_j^0)^2} = d(\mathbf{x}, \mathbf{x}_0),$$

per cui, prendendo $\delta = \varepsilon$, si ha:

$$d(\mathbf{x}, \mathbf{x}_0) < \delta \Rightarrow |p_k(\mathbf{x}) - p_k(\mathbf{x}_0)| = |x_k - x_k^0| < \varepsilon,$$

il che dimostra che p_k è continua in $\mathbf{x}_0$. ■

Supponiamo ora $f : E \rightarrow \mathbb{R}^M$. Consideriamo le “componenti” della funzione f definite da $f_k = p_k \circ f : E \rightarrow \mathbb{R}$, con $k = 1, 2, \dots, M$, per cui si ha

$$f(x) = (f_1(x), f_2(x), \dots, f_M(x)).$$

Teorema 13 La funzione f è continua in x_0 se e solo se lo sono tutte le sue componenti.

Dimostrazione. Se f è continua in x_0 , lo sono anche le f_k in quanto composte di funzioni continue. Viceversa, supponiamo che le componenti di f siano tutte continue in x_0 . Fissato $\varepsilon > 0$, per ogni $k = 1, 2, \dots, M$ esiste un $\delta_k > 0$ tale che

$$d(x, x_0) < \delta_k \Rightarrow |f_k(x) - f_k(x_0)| < \varepsilon.$$

Posto $\delta = \min\{\delta_1, \delta_2, \dots, \delta_M\}$, si ha

$$d(x, x_0) < \delta \Rightarrow d(f(x), f(x_0)) = \sqrt{\sum_{j=1}^M (f_j(x) - f_j(x_0))^2} < \sqrt{M}\varepsilon,$$

il che, per l'arbitrarietà di ε , completa la dimostrazione. ■

Teorema 14 *Ogni applicazione lineare $\ell : \mathbb{R}^N \rightarrow \mathbb{R}^M$ è continua.*

Dimostrazione. Osserviamo che, essendo le proiezioni p_k lineari, le componenti $\ell_k = p_k \circ \ell$ dell'applicazione lineare ℓ sono anch'esse lineari. Consideriamo la base canonica $[e_1, e_2, \dots, e_N]$ di $\mathbb{R}^N$, con

$$\begin{aligned} e_1 &= (1, 0, 0, \dots, 0), \\ e_2 &= (0, 1, 0, \dots, 0), \\ &\vdots \\ e_N &= (0, 0, 0, \dots, 1). \end{aligned}$$

Ogni vettore $\mathbf{x} = (x_1, x_2, \dots, x_N) \in \mathbb{R}^N$ si può scrivere come

$$\mathbf{x} = x_1 e_1 + x_2 e_2 + \dots + x_N e_N = p_1(\mathbf{x})e_1 + p_2(\mathbf{x})e_2 + \dots + p_N(\mathbf{x})e_N.$$

Quindi, per ogni $k \in \{1, 2, \dots, M\}$,

$$\ell_k(\mathbf{x}) = p_1(\mathbf{x})\ell_k(e_1) + p_2(\mathbf{x})\ell_k(e_2) + \dots + p_N(\mathbf{x})\ell_k(e_N),$$

per cui ℓ_k risulta essere combinazione lineare delle proiezioni $p_1, p_2, \dots, p_N$. Essendo queste ultime continue, anche ℓ_k è continua. Avendo tutte le componenti continue, ℓ è pertanto continua. ■

4 La nozione di limite

Consideriamo due spazi metrici E, F , un punto x_0 di E e una funzione

$$f : E \rightarrow F, \quad \text{oppure} \quad f : E \setminus \{x_0\} \rightarrow F,$$

non necessariamente definita in x_0 . Supporremo inoltre che x_0 sia un “punto di accumulazione” di E : ogni intorno di x_0 contiene infiniti punti di E .

Definizione 5 Si dice che $l \in F$ è il “limite di f in x_0 ”, o anche “limite di $f(x)$ per x che tende a x_0 ” e si scrive

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

se

$$\forall \varepsilon > 0 \quad \exists \delta > 0 : \forall x \in E \quad 0 < d(x, x_0) < \delta \Rightarrow d(f(x), l) < \varepsilon,$$

o equivalentemente,

$$\forall V, \text{ intorno di } l \quad \exists U, \text{ intorno di } x_0 : f(U \setminus \{x_0\}) \subseteq V.$$

Talvolta si scrive anche $f(x) \rightarrow l$ per $x \rightarrow x_0$.

Per cominciare, verifichiamo l'unicità del limite.

Teorema 15 Se esiste, il limite di f in x_0 è unico.

Dimostrazione. Supponiamo per assurdo che ce ne siano due diversi, l e l' . Prendiamo $\varepsilon = \frac{1}{2}d(l, l')$. Allora esiste un $\delta > 0$ tale che

$$0 < d(x, x_0) < \delta \Rightarrow d(f(x), l) < \varepsilon,$$

ed esiste un $\delta' > 0$ tale che

$$0 < d(x, x_0) < \delta' \Rightarrow d(f(x), l') < \varepsilon.$$

Sia $x \neq x_0$ tale che $d(x, x_0) < \delta$ e $d(x, x_0) < \delta'$ (tale x esiste perché x_0 è di accumulazione). Allora

$$d(l', l) \leq d(l, f(x)) + d(f(x), l') < 2\varepsilon = d(l', l),$$

una contraddizione. ■

I seguenti due teoremi evidenziano il legame stretto che intercorre tra i concetti di limite e di continuità.

Teorema 16 Si ha che $\lim_{x \rightarrow x_0} f(x) = l$ se e solo se la funzione $\tilde{f} : E \rightarrow \mathbb{R}$ definita da

$$\tilde{f}(x) = \begin{cases} f(x) & \text{se } x \neq x_0, \\ l & \text{se } x = x_0 \end{cases}$$

è continua in x_0

Teorema 17 Considerata la funzione $f : E \rightarrow F$, definita anche in x_0 , si ha che

$$f \text{ è continua in } x_0 \Leftrightarrow \lim_{x \rightarrow x_0} f(x) = f(x_0).$$

Iniziamo a vedere le proprietà dei limiti che vengono direttamente ereditate dalle funzioni continue. Nei teoremi seguenti, con relativo corollario, le funzioni f e g sono definite su E o su $E \setminus \{x_0\}$, indifferentemente.

Teorema 18 (della permanenza del segno) Sia $g : E \setminus \{x_0\} \rightarrow \mathbb{R}$ tale che

$$\lim_{x \rightarrow x_0} g(x) > 0 ,$$

allora esiste un $\delta > 0$ tale che

$$0 < d(x, x_0) < \delta \Rightarrow g(x) > 0 .$$

Analogamente, se

$$\lim_{x \rightarrow x_0} g(x) < 0 ,$$

allora esiste un $\delta > 0$ tale che

$$0 < d(x, x_0) < \delta \Rightarrow g(x) < 0 .$$

Corollario 1 Se $g(x) \leq 0$ per ogni x in un intorno di x_0 , allora, qualora il limite esista, si ha

$$\lim_{x \rightarrow x_0} g(x) \leq 0 .$$

Analogamente, se $g(x) \geq 0$ per ogni x in un intorno di x_0 , allora, qualora il limite esista, si ha

$$\lim_{x \rightarrow x_0} g(x) \geq 0 .$$

Vediamo ora come si comporta il limite nei confronti delle operazioni in $\mathbb{R}$.

Teorema 19 Siano $f, g : E \setminus \{x_0\} \rightarrow \mathbb{R}$ tali che

$$l_1 = \lim_{x \rightarrow x_0} f(x), \quad l_2 = \lim_{x \rightarrow x_0} g(x) .$$

Allora

$$\lim_{x \rightarrow x_0} [f(x)g(x)] = l_1 l_2 ;$$

inoltre, se $l_2 \neq 0$,

$$\lim_{x \rightarrow x_0} \frac{f(x)}{g(x)} = \frac{l_1}{l_2} .$$

Consideriamo ora una funzione composta $g \circ f$. Abbiamo due possibili situazioni.

Teorema 20 Sia $f : E \rightarrow F$, oppure $f : E \setminus \{x_0\} \rightarrow F$, tale che

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

Se $g : F \rightarrow G$ è continua in l , allora

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

In altri termini,

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

Dimostrazione. Ricordo che la funzione $\tilde{f} : E \rightarrow \mathbb{R}$ definita da

$$\tilde{f}(x) = \begin{cases} f(x) & \text{se } x \neq x_0, \\ l & \text{se } x = x_0 \end{cases}$$

è continua in x_0 e g è continua in $l = \tilde{f}(x_0)$. Pertanto, $g \circ \tilde{f}$ è continua in x_0 , da cui

$$\lim_{x \rightarrow x_0} g(f(x)) = \lim_{x \rightarrow x_0} g(\tilde{f}(x)) = g(\tilde{f}(x_0)) = g(l).$$

■

Teorema 21 Sia $f : E \rightarrow F$, oppure $f : E \setminus \{x_0\} \rightarrow F$, tale che

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

Supponiamo che l sia un punto di accumulazione di F e che la funzione

$$g : F \rightarrow G, \quad \text{oppure} \quad g : F \setminus \{l\} \rightarrow G,$$

non necessariamente definita in l , sia tale che

$$\lim_{y \rightarrow l} g(y) = L.$$

Se $f(x) \neq l$ per ogni $x \in E \setminus \{x_0\}$, allora

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

Dimostrazione. Consideriamo nuovamente la funzione $\tilde{f} : E \rightarrow F$, continua in x_0 con $\tilde{f}(x_0) = l$. Analogamente, consideriamo la funzione $\tilde{g} : F \rightarrow G$ così definita:

$$\tilde{g}(y) = \begin{cases} g(y) & \text{se } y \neq l, \\ L & \text{se } y = l. \end{cases}$$

Essa è continua in l con $\tilde{g}(l) = L$. Consideriamo la funzione composta $\tilde{g} \circ \tilde{f}$, che per quanto sopra è continua in x_0 con $\tilde{g}(\tilde{f}(x_0)) = \tilde{g}(l) = L$. Essendo $f(x) \neq l$ per ogni x , si ha che, per $x \in E \setminus \{x_0\}$,

$$g(f(x)) = \tilde{g}(f(x)) = \tilde{g}(\tilde{f}(x)),$$

e pertanto,

$$\lim_{x \rightarrow x_0} g(f(x)) = \lim_{x \rightarrow x_0} \tilde{g}(\tilde{f}(x)) = \tilde{g}(\tilde{f}(x_0)) = L.$$

■

Nota. La conclusione del teorema precedente si riassume con la formula

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

Spesso si dice che si è operato il “cambio di variabile $y = f(x)$ ”. Riguardando inoltre le ipotesi dello stesso teorema, si vede subito che è sufficiente richiedere che sia $f(x) \neq l$ per gli x tali che $0 < d(x, x_0) < \delta$. Ciò è dovuto al fatto che la nozione di limite è, in un certo senso, di tipo “locale”. Questa osservazione vale in generale e verrà spesso usata in seguito.

Esempio. Sia ora $f : \mathbb{R}^2 \rightarrow \mathbb{R}$ definita da $f(x, y) = x^2 + y^2$ e $g : \mathbb{R} \setminus \{0\} \rightarrow \mathbb{R}$ definita da $g(z) = z \sin(1/z)$. Il Teorema 21 qui può essere applicato per concludere che

$$\lim_{(x,y) \rightarrow (0,0)} g(f(x, y)) = \lim_{\substack{z \rightarrow \lim_{(x,y) \rightarrow (0,0)} f(x,y)}} g(z) = \lim_{z \rightarrow 0} g(z) = 0.$$

Siamo infine interessati a studiare il limite di una funzione $f : E \setminus \{x_0\} \rightarrow \mathbb{R}^M$, dove x_0 è un punto di accumulazione di uno spazio metrico E . Consideriamo le sue componenti $f_k : E \setminus \{x_0\} \rightarrow \mathbb{R}$ di f , con $k = 1, 2, \dots, M$, per cui si ha:

$$f(x) = (f_1(x), f_2(x), \dots, f_M(x)).$$

Teorema 22 *Il limite $\lim_{x \rightarrow x_0} f(x) = \mathbf{l} \in \mathbb{R}^M$ esiste se e solo se esistono i limiti $\lim_{x \rightarrow x_0} f_k(x) = l_k \in \mathbb{R}$, per ogni $k = 1, 2, \dots, M$. In tal caso, si ha $\mathbf{l} = (l_1, l_2, \dots, l_M)$. Vale quindi la formula*

$$\lim_{x \rightarrow x_0} f(x) = (\lim_{x \rightarrow x_0} f_1(x), \lim_{x \rightarrow x_0} f_2(x), \dots, \lim_{x \rightarrow x_0} f_M(x)).$$

Dimostrazione. Segue direttamente dal teorema sulla continuità delle componenti di una funzione continua. ■

Finora abbiamo considerato due spazi metrici E, F , un punto x_0 di accumulazione per E e una funzione $f : E \rightarrow F$, oppure $f : E \setminus \{x_0\} \rightarrow F$. Siccome l'eventuale valore di f in x_0 è ininfluente ai fini dell'esistenza o meno del limite, nonché del suo effettivo valore, da ora in poi per semplicità considereremo solo il caso $f : E \setminus \{x_0\} \rightarrow F$.

Si può verificare che tutte le considerazioni fatte continuano a valere per una funzione $f : \widehat{E} \setminus \{x_0\} \rightarrow F$, con $\widehat{E} \subseteq E$, purché x_0 sia di accumulazione per $\widehat{E}$: ogni intorno di x_0 deve contenere infiniti punti di $\widehat{E}$.

Sia ora $f : E \setminus \{x_0\} \rightarrow F$, e sia $\widehat{E} \subseteq E$. Possiamo considerare la restrizione di f a $\widehat{E} \setminus \{x_0\}$: è la funzione $\hat{f} : \widehat{E} \setminus \{x_0\} \rightarrow F$ i cui valori coincidono con quelli di f : si ha $\hat{f}(x) = f(x)$ per ogni $x \in \widehat{E} \setminus \{x_0\}$. Talvolta si scrive $\hat{f} = f|_{\widehat{E}}$.

Teorema 23 *Se esiste il limite di f in x_0 e x_0 è di accumulazione anche per $\widehat{E}$, allora esiste anche il limite di $\hat{f}$ in x_0 e ha lo stesso valore:*

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

Dimostrazione. Segue immediatamente dalla definizione di $\hat{f}$. ■

Il teorema precedente viene spesso usato per stabilire la non esistenza del limite per la funzione f : a tal scopo, è sufficiente trovare due diverse restrizioni lungo le quali i valori del limite differiscono.

Esempio 1. La funzione $f : \mathbb{R}^2 \setminus \{(0, 0)\} \rightarrow \mathbb{R}$, definita da

$$f(x, y) = \frac{xy}{x^2 + y^2},$$

non ha limite per $(x, y) \rightarrow (0, 0)$, come si vede considerando le restrizioni alle due rette $\{(x, y) : x = 0\}$ e $\{(x, y) : x = y\}$.

Esempio 2. Più sorprendente è la funzione definita da

$$f(x, y) = \frac{x^2 y}{x^4 + y^2},$$

per la quale le restrizioni a tutte le rette passanti per $(0, 0)$ hanno limite 0, ma la restrizione alla parabola $\{(x, y) : y = x^2\}$ vale costantemente $\frac{1}{2}$.

Esempio 3. Dimostriamo invece che

$$\lim_{(x, y) \rightarrow (0, 0)} \frac{x^2 y^2}{x^2 + y^2} = 0.$$

Fissiamo un $\varepsilon > 0$. Dopo aver verificato che

$$\frac{x^2 y^2}{x^2 + y^2} \leq \frac{1}{2} (x^2 + y^2),$$

risulta naturale prendere $\delta = \sqrt{2\varepsilon}$, per avere che

$$d((x, y), (0, 0)) < \delta \Rightarrow \left| \frac{x^2 y^2}{x^2 + y^2} - 0 \right| < \varepsilon.$$

Consideriamo ora il caso in cui E è uno spazio metrico qualunque ed $F = \mathbb{R}$. Risulterà talvolta utile il seguente “teorema dei due carabinieri”.

Teorema 24 *Supponiamo di avere due funzioni f_1, f_2 per cui*

$$\lim_{x \rightarrow x_0} f_1(x) = \lim_{x \rightarrow x_0} f_2(x) = l.$$

Se $f : E \setminus \{x_0\} \rightarrow \mathbb{R}$ è tale che, per ogni x ,

$$f_1(x) \leq f(x) \leq f_2(x),$$

allora

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

Dimostrazione. Fissato $\varepsilon > 0$, esistono $\delta_1 > 0$ e $\delta_2 > 0$ tali che

$$\begin{aligned} 0 < d(x, x_0) < \delta_1 &\Rightarrow l - \varepsilon < f_1(x) < l + \varepsilon, \\ 0 < d(x, x_0) < \delta_2 &\Rightarrow l - \varepsilon < f_2(x) < l + \varepsilon. \end{aligned}$$

Se $\delta = \min\{\delta_1, \delta_2\}$, allora

$$0 < d(x, x_0) < \delta \Rightarrow l - \varepsilon < f_1(x) \leq f(x) \leq f_2(x) < l + \varepsilon,$$

il che dimostra la tesi. ■

5 La retta ampliata

Consideriamo la funzione $\varphi : \mathbb{R} \rightarrow]-1, 1[$, definita da

$$\varphi(x) = \frac{x}{1 + |x|}.$$

Possiamo definire una nuova distanza su $\mathbb{R}$:

$$\tilde{d}(x, x') = |\varphi(x) - \varphi(x')|.$$

Si può vedere che ogni intorno per la nuova distanza è anche intorno per la vecchia distanza, e viceversa.

Introduciamo ora il nuovo insieme $\tilde{\mathbb{R}}$, definito come unione di $\mathbb{R}$ e di due nuovi elementi, che indicheremo con $-\infty$ e $+\infty$:

$$\tilde{\mathbb{R}} = \mathbb{R} \cup \{-\infty, +\infty\}.$$

L'insieme $\tilde{\mathbb{R}}$ risulta totalmente ordinato se si mantiene l'ordine esistente tra coppie di numeri reali e si pone inoltre, per ogni $x \in \mathbb{R}$,

$$-\infty < x < +\infty.$$

Consideriamo la funzione $\tilde{\varphi} : \tilde{\mathbb{R}} \rightarrow [-1, 1]$, definita da

$$\tilde{\varphi}(x) = \begin{cases} -1 & \text{se } x = -\infty, \\ \varphi(x) & \text{se } x \in \mathbb{R}, \\ 1 & \text{se } x = +\infty. \end{cases}$$

Definiamo, per $x, x' \in \tilde{\mathbb{R}}$,

$$\tilde{d}(x, x') = |\tilde{\varphi}(x) - \tilde{\varphi}(x')|.$$

Si verifica facilmente che $\tilde{d}$ è una distanza su $\tilde{\mathbb{R}}$. In questo modo, $\tilde{\mathbb{R}}$ risulta uno spazio metrico.

Si può verificare che un intorno di $+\infty$ è un insieme che contiene, oltre al punto $+\infty$, un intervallo del tipo $]\alpha, +\infty[$, per un certo $\alpha \in \mathbb{R}$.

Analogamente, un intorno di $-\infty$ è un insieme che contiene, oltre a $-\infty$, un intervallo del tipo $]-\infty, \beta[$, per un certo $\beta \in \mathbb{R}$.

Vediamo ora come si traduce la definizione di limite in alcuni casi in cui compaiono gli elementi $+\infty$ o $-\infty$. Ad esempio, sia $E \subseteq \mathbb{R}$, F uno spazio metrico e $f : E \rightarrow F$ una funzione. Considerando E come sottoinsieme di $\tilde{\mathbb{R}}$, si ha che $+\infty$ è punto di accumulazione per E se e solo se E non è limitato superiormente. In tal caso, si ha:

$$\lim_{x \rightarrow +\infty} f(x) = l \in F \quad \Leftrightarrow \quad \forall V \text{ intorno di } l \quad \exists U \text{ intorno di } +\infty : \\ f(U \cap E) \subseteq V$$

$$\Leftrightarrow \quad \forall \varepsilon > 0 \quad \exists \alpha \in \mathbb{R} : \quad x > \alpha \Rightarrow d(f(x), l) < \varepsilon.$$

Analogamente, se E non è limitato inferiormente, si ha:

$$\lim_{x \rightarrow -\infty} f(x) = l \in F \quad \Leftrightarrow \quad \forall \varepsilon > 0 \quad \exists \beta \in \mathbb{R} : \quad x < \beta \Rightarrow d(f(x), l) < \varepsilon.$$

Si noti che

$$\lim_{x \rightarrow +\infty} f(x) = l \quad \Leftrightarrow \quad \lim_{x \rightarrow -\infty} f(-x) = l.$$

Vediamo ora il caso in cui E sia uno spazio metrico ed $F = \mathbb{R}$, considerato come sottoinsieme di $\widetilde{\mathbb{R}}$. Supponiamo che x_0 sia di accumulazione per E e consideriamo una funzione $f : E \rightarrow \mathbb{R}$, o $f : E \setminus \{x_0\} \rightarrow \mathbb{R}$. Si ha:

$$\begin{aligned} \lim_{x \rightarrow x_0} f(x) = +\infty \quad &\Leftrightarrow \quad \forall V \text{ intorno di } +\infty \quad \exists U \text{ intorno di } x_0 : \\ &f(U \setminus \{x_0\}) \subseteq V \\ &\Leftrightarrow \quad \forall \alpha \in \mathbb{R} \quad \exists \delta > 0 : \quad 0 < d(x, x_0) < \delta \Rightarrow f(x) > \alpha; \end{aligned}$$

analogamente,

$$\lim_{x \rightarrow x_0} f(x) = -\infty \quad \Leftrightarrow \quad \forall \beta \in \mathbb{R} \quad \exists \delta > 0 : \quad 0 < d(x, x_0) < \delta \Rightarrow f(x) < \beta.$$

Si noti che

$$\lim_{x \rightarrow x_0} f(x) = +\infty \quad \Leftrightarrow \quad \lim_{x \rightarrow x_0} (-f(x)) = -\infty.$$

Le situazioni considerate in precedenza possono talvolta presentarsi assieme. Ad esempio, se $E \subseteq \mathbb{R}$ non è limitato superiormente ed $F = \mathbb{R}$, si avrà

$$\begin{aligned} \lim_{x \rightarrow +\infty} f(x) = +\infty \quad &\Leftrightarrow \quad \forall V \text{ intorno di } +\infty \quad \exists U \text{ intorno di } +\infty : \\ &f(U \cap E) \subseteq V \\ &\Leftrightarrow \quad \forall \alpha \in \mathbb{R} \quad \exists \alpha' \in \mathbb{R} : \quad x > \alpha' \Rightarrow f(x) > \alpha; \end{aligned}$$

analogamente,

$$\lim_{x \rightarrow +\infty} f(x) = -\infty \quad \Leftrightarrow \quad \forall \beta \in \mathbb{R} \quad \exists \alpha \in \mathbb{R} : \quad x > \alpha \Rightarrow f(x) < \beta.$$

Se invece $E \subseteq \mathbb{R}$ non è limitato inferiormente ed $F = \mathbb{R}$, si avrà

$$\lim_{x \rightarrow -\infty} f(x) = +\infty \quad \Leftrightarrow \quad \forall \alpha \in \mathbb{R} \quad \exists \beta \in \mathbb{R} : \quad x < \beta \Rightarrow f(x) > \alpha;$$

analogamente,

$$\lim_{x \rightarrow -\infty} f(x) = -\infty \quad \Leftrightarrow \quad \forall \beta \in \mathbb{R} \quad \exists \beta' \in \mathbb{R} : \quad x < \beta' \Rightarrow f(x) < \beta.$$

Vediamo ad esempio il caso di una successione $(a_n)_n$ in uno spazio metrico F . Abbiamo quindi una funzione $f : \mathbb{N} \rightarrow F$ definita da $f(n) = a_n$. Considerando $\mathbb{N}$ come sottoinsieme di $\widehat{\mathbb{R}}$, si vede che l'unico punto di accumulazione è $+\infty$. Adattando la definizione di limite a questo caso, possiamo scrivere:

$$\lim_{n \rightarrow +\infty} a_n = l \in F \quad \Leftrightarrow \quad \forall \varepsilon > 0 \quad \exists \bar{n} \in \mathbb{N} : \quad n \geq \bar{n} \Rightarrow d(a_n, l) < \varepsilon.$$

Pertanto, spesso il limite di una successione si denota semplicemente con $\lim_n a_n$, sottintendendo che $n \rightarrow +\infty$.

Teorema 25 *La funzione f è continua in x_0 se e solo se, presa una successione $(a_n)_n$ in E , si ha*

$$\lim_n a_n = x_0 \quad \Rightarrow \quad \lim_n f(a_n) = f(x_0).$$

Dimostrazione. Supponiamo che f sia continua in x_0 , e sia $(a_n)_n$ una successione in E tale che $\lim_n a_n = x_0$. Per il Teorema 20 sul limite di una funzione composta,

$$\lim_n f(a_n) = f(\lim_n a_n) = f(x_0),$$

cosicchè una delle due implicazioni è dimostrata.

Ragioniamo ora per contrapposizione, e supponiamo che f non sia continua in x_0 . Questo significa che esiste un $\varepsilon > 0$ tale che, per ogni $\delta > 0$, esiste almeno un $x \in E$ per cui $d(x, x_0) < \delta$ e $d(f(x), f(x_0)) \geq \varepsilon$. Prendendo $\delta = \frac{1}{n+1}$, per ogni $n \in \mathbb{N}$ esiste pertanto un a_n in E tale che $d(a_n, x_0) < \frac{1}{n+1}$ e $d(f(a_n), f(x_0)) \geq \varepsilon$. Ne segue che $\lim_n a_n = x_0$, ma sicuramente non può essere che $\lim_n f(a_n) = f(x_0)$. ■

Consideriamo ora il caso in cui x_0 sia un punto di accumulazione e $f : E \setminus \{x_0\} \rightarrow F$ una funzione. Come immediata conseguenza del teorema precedente, otteniamo il seguente

Corollario 2 *Avremo che*

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

se e solo se, presa una successione $(a_n)_n$ in $E \setminus \{x_0\}$, si ha

$$\lim_n a_n = x_0 \quad \Rightarrow \quad \lim_n f(a_n) = l.$$

Sia ora U un sottoinsieme dello spazio metrico E . Possiamo caratterizzare la nozione di punto aderente a U facendo uso delle successioni.

Teorema 26 *Un punto $x \in E$ è aderente a U se e solo se esiste una successione $(a_n)_n$ in U tale che $\lim_n a_n = x$.*

Dimostrazione. Se x è aderente a U , allora l'intersezione $B(x, \frac{1}{n+1}) \cap U$ è non vuota, per ogni $n \in \mathbb{N}$, per cui posso sceglierne un elemento, che chiamo a_n . In questo modo, ho costruito una successione $(a_n)_n$ in U , ed è facile vedere che essa ha limite x . Una delle due implicazioni è così dimostrata.

Supponiamo ora che esista una successione $(a_n)_n$ in U tale che $\lim_n a_n = x$. Allora, fissato $\rho > 0$, esiste un $\bar{n} \in \mathbb{N}$ tale che

$$n \geq \bar{n} \Rightarrow d(a_n, x) < \rho,$$

ossia $a_n \in B(x, \rho)$. Quindi, $B(x, \rho) \cap U$ è non vuoto, e questo dimostra che x è aderente a U . ■

6 Insiemi compatti

Data che sia una successione $(a_n)_n$, una sua “sottosuccessione” si ottiene selezionando una successione strettamente crescente di indici $(n_k)_k$ e considerando la funzione composta

$$k \mapsto n_k \mapsto a_{n_k}.$$

Teorema 27 *Se una successione ha limite, allora tutte le sue sottosuccessioni hanno lo stesso limite.*

Dimostrazione. Essendo gli indici n_k in $\mathbb{N}$, dalla $n_{k+1} > n_k$ si deduce che $n_{k+1} \geq n_k + 1$ e, per induzione, che $n_k \geq k$, per ogni k . Ne segue che $\lim_k n_k = +\infty$. Pertanto,

$$\lim_{k \rightarrow +\infty} a_{n_k} = \lim_{n \rightarrow \lim_{k \rightarrow +\infty} n_k} a_n = \lim_{n \rightarrow +\infty} a_n. \quad \blacksquare$$

In uno spazio metrico E , diremo che un sottoinsieme U è “compatto” se ogni successione $(a_n)_n$ in U possiede una sottosuccessione $(a_{n_k})_k$ che ha limite in U . Il Teorema di Bolzano–Weierstrass afferma quindi che, se $E = \mathbb{R}$, gli intervalli del tipo $U = [a, b]$ sono compatti.

Teorema 28 *Ogni insieme compatto di uno spazio metrico E è chiuso e limitato.*

Dimostrazione. Sia U un insieme compatto. Preso un $x \in \bar{U}$, esiste una successione $(a_n)_n$ in U tale che $\lim_n a_n = x$. Essendo U compatto, esiste una sottosuccessione $(a_{n_k})_k$ che ha limite in U . Ma, essendo una sottosuccessione, deve essere $\lim_k a_{n_k} = x$, per cui $x \in U$. Quindi, ogni punto aderente di U appartiene ad U , per cui U è chiuso.

Fissiamo ora un $x_0 \in U$ qualsiasi e dimostriamo che, se $n \in \mathbb{N}$ è sufficientemente grande, allora $U \subseteq B(x_0, n)$. Per assurdo, se così non fosse, potrei trovare una successione $(a_n)_n$ in U tale che $d(a_n, x_0) \geq n$, per ogni $n \in \mathbb{N}$. Ma, essendo U compatto, esiste una sottosuccessione $(a_{n_k})_k$ che ha un certo limite $\bar{x} \in U$. Usando la disuguaglianza triangolare, si ha che

$$|d(a_{n_k}, x_0) - d(\bar{x}, x_0)| \leq d(a_{n_k}, \bar{x}),$$

da cui segue che $\lim_k d(a_{n_k}, x_0) = d(\bar{x}, x_0)$, mentre dovrebbe essere

$$\lim_k d(a_{n_k}, x_0) = +\infty,$$

una contraddizione. Pertanto, U deve essere limitato. ■

Ora vorremmo focalizzare la nostra attenzione sui sottoinsiemi compatti di $\mathbb{R}^M$.

Teorema 29 *Un sottoinsieme $\mathbb{R}^M$ è compatto se e solo se è chiuso e limitato.*

Dimostrazione. Sappiamo già che ogni compatto è chiuso e limitato. Supponiamo ora che U sia un sottoinsieme chiuso e limitato di $\mathbb{R}^M$. Supporremo per semplicità $M = 2$. Allora U è contenuto in un rettangolo

$$I = [a_1, b_1] \times [a_2, b_2].$$

Sia $(\mathbf{a}_n)_n$ una successione in U . Si ha $\mathbf{a}_n = (a_n^1, a_n^2)$, con $a_n^1 \in [a_1, b_1]$ e $a_n^2 \in [a_2, b_2]$. Per la proprietà di Bolzano–Weierstrass, esiste una sottosuccessione $(a_{n_k}^1)_k$ che ha un limite $l_1 \in [a_1, b_1]$. Consideriamo la sottosuccessione $(a_{n_k}^2)_k$, con gli stessi indici di quella appena trovata. Per la proprietà di Bolzano–Weierstrass, esiste una sottosuccessione $(a_{n_{k_j}}^2)_j$ che ha un limite $l_2 \in [a_2, b_2]$. Osservando che

$$d(\mathbf{a}_{n_{k_j}}, (l_1, l_2)) = \sqrt{(a_{n_{k_j}}^1 - l_1)^2 + (a_{n_{k_j}}^2 - l_2)^2},$$

se ne deduce che

$$\lim_j \mathbf{a}_{n_{k_j}} = (l_1, l_2).$$

Il punto $\mathbf{l} = (l_1, l_2)$ è aderente ad U . Essendo U chiuso, $\mathbf{l}$ appartiene ad U . ■

Nel seguito, diremo che una funzione $f : U \rightarrow \mathbb{R}$ è “limitata superiormente” (o “limitata inferiormente”) se lo è la sua immagine $f(U)$. Diremo che f è “limitata” se è sia limitata superiormente che inferiormente. Diremo che “ f ha massimo” (o “ f ha minimo”) se $f(U)$ ce l’ha. Nel caso in cui f abbia massimo, chiameremo “punto di massimo” ogni $\bar{x}$ per cui $f(\bar{x}) = \max f(U)$; analoga definizione per “punto di minimo”.

Teorema 30 (di Weierstrass) *Se U è un insieme compatto e $f : U \rightarrow \mathbb{R}$ è una funzione continua, allora f ha massimo e minimo.*

Dimostrazione. Sia $s = \sup f(U)$. Dimostreremo che esiste un punto di massimo, ossia un $\bar{x} \in U$ tale che $f(\bar{x}) = s$.

Notiamo che è possibile trovare una successione $(y_n)_n$ in $f(U)$ tale che $\lim_n y_n = s$: se $s \in \mathbb{R}$, per ogni $n \geq 1$ possiamo trovare un $y_n \in f(U)$ per cui $s - \frac{1}{n} < y_n \leq s$; se invece $s = +\infty$, per ogni n esiste un $y_n \in f(U)$ tale che $y_n > n$.

In corrispondenza, possiamo trovare una successione $(x_n)_n$ in U tale che $f(x_n) = y_n$. Essendo U compatto, esiste una sottosuccessione $(x_{n_k})_k$ che ha un limite $\bar{x} \in U$. Siccome $\lim_n y_n = s$ e $y_{n_k} = f(x_{n_k})$, la sottosuccessione $(y_{n_k})_k$ ha anch’essa limite s . Allora, per la continuità di f ,

$$f(\bar{x}) = f(\lim_k x_{n_k}) = \lim_k f(x_{n_k}) = \lim_k y_{n_k} = s.$$

Il teorema è così dimostrato, per quanto riguarda l’esistenza del massimo. Per il minimo, si procede in modo analogo (oppure, si considera la funzione continua $g = -f$ e si usa il fatto che g ha massimo). ■