

Analisi 2

Appunti delle lezioni tenute dal Prof. A. Fonda

Università di Trieste,

CdL Matematica

a.a. 2025/2026

Nota. Questi appunti sono in evoluzione. Si prega di segnalare eventuali errori, per poterli correggere rapidamente.

Premesse: 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) $\mathbf{x} \cdot \mathbf{x}' = \mathbf{x}' \cdot \mathbf{x}$;

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), abbiamo bisogno della seguente **disuguaglianza di Schwarz**.

Teorema. *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. ■

Nota. La disuguaglianza di Schwarz diventa un'uguaglianza se e solo se i due vettori sono linearmente dipendenti. Infatti, vale sicuramente l'uguaglianza se uno dei due vettori è nullo. Se sono non nulli e $\mathbf{x}$ è del tipo $\alpha \mathbf{x}'$, si ha

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

Viceversa, se sono non nulli e vale l'uguaglianza $\mathbf{x} \cdot \mathbf{x}' = \|\mathbf{x}\| \|\mathbf{x}'\|$, allora

$$\left\| \frac{\mathbf{x}}{\|\mathbf{x}\|} - \frac{\mathbf{x}'}{\|\mathbf{x}'\|} \right\|^2 = 1 - 2 \frac{\mathbf{x} \cdot \mathbf{x}'}{\|\mathbf{x}\| \|\mathbf{x}'\|} + 1 = 0,$$

per cui $\frac{\mathbf{x}}{\|\mathbf{x}\|} = \frac{\mathbf{x}'}{\|\mathbf{x}'\|}$. I due vettori hanno la stessa direzione, per cui sono linearmente dipendenti. Analogamente, se $\mathbf{x} \cdot \mathbf{x}' = -\|\mathbf{x}\| \|\mathbf{x}'\|$, allora succede che $\frac{\mathbf{x}}{\|\mathbf{x}\|} = -\frac{\mathbf{x}'}{\|\mathbf{x}'\|}$.

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:

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

Spazi metrici

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 intervallo¹ $]a, b[$ 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

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

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 $\mathring{U}$. Chiaramente, si ha sempre $\mathring{U} \subseteq U$. Si dice che U è un “insieme aperto” se coincide con il suo interno, ossia se $\mathring{U} = U$.

Teorema. *Una palla aperta è un insieme aperto.*

¹Usiamo qui le notazioni $]a, b[= \{x \in \mathbb{R} : a < x < b\}$ e $[a, b] = \{x \in \mathbb{R} : a \leq x \leq b\}$.

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 $\mathring{E} = E$. 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, $\mathring{\emptyset} = \emptyset$, 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. *L'interno di un insieme è un insieme aperto.*

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

Per dimostrare che $B(x_1, \rho) \subseteq \mathring{U}$, 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 $\mathring{U}$. La dimostrazione è completa. ■

Si può dimostrare la seguente implicazione:

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

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

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. 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 $\mathcal{C}U \subseteq \mathcal{C}\overline{U}$.² Prendiamo un $x_1 \in \mathcal{C}U$, 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. 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. ■

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

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

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. *Valgono le seguenti relazioni:*

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

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:

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

Passando ai complementari, si ottiene la seconda uguaglianza. ■

Abbiamo quindi che

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

Come immediato corollario, abbiamo il seguente.

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

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

$$\overline{\mathcal{C}U} = \mathcal{C}\overset{\circ}{U} = \mathcal{C}U,$$

per cui $\mathcal{C}U$ è chiuso.

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

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

per cui $\mathcal{C}U$ è aperto. ■

Si definisce la “frontiera” di un insieme U come differenza tra la sua chiusura e il suo interno:

$$\partial U = \overline{U} \setminus \overset{\circ}{U}.$$

Ad esempio, in $\mathbb{R}$ abbiamo:

$$\overline{\mathbb{Q}} = \mathbb{R}, \quad \overset{\circ}{\mathbb{Q}} = \emptyset, \quad \partial \mathbb{Q} = \mathbb{R}.$$

È bene essere prudenti su alcune conclusioni che possono esserci suggerite dalla nostra intuizione basata sulla distanza euclidea. Ad esempio, le uguaglianze

$$\overline{B(\mathbf{x}_0, \rho)} = \overline{B}(\mathbf{x}_0, \rho), \quad \partial B(\mathbf{x}_0, \rho) = S(\mathbf{x}_0, \rho).$$

valgono sicuramente in $\mathbb{R}^N$ con la distanza euclidea, ma possono non valere in altri casi. Prendiamo ad esempio la distanza $\hat{d}$ considerata sopra. Allora $B(\mathbf{x}_0, 1) = \{\mathbf{x}_0\}$, che è un insieme chiuso, e $\overline{B}(\mathbf{x}_0, 1) = E$ per cui $\overline{B(\mathbf{x}_0, 1)} \neq \overline{B}(\mathbf{x}_0, 1)$. Inoltre, $\partial B(\mathbf{x}_0, 1) = \emptyset$, mentre $S(\mathbf{x}_0, 1) = E \setminus \{\mathbf{x}_0\}$, per cui $\partial B(\mathbf{x}_0, 1) \neq S(\mathbf{x}_0, 1)$.

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

Funzioni continue

Intuitivamente, una funzione f è “continua” se $f(x)$ varia gradualmente al variare di x nel dominio, cioè quando non si verificano variazioni brusche nei valori della funzione. Per rendere rigorosa questa idea intuitiva, sarà conveniente focalizzare la nostra attenzione fissando un x_0 nel dominio e provando a precisare cosa intendiamo per

$$f \text{ è “continua” in } x_0.$$

Procederemo per gradi.

Primo tentativo. Diremo che f è “continua” in x_0 quando si verifica la cosa seguente:

se x è vicino a x_0 , allora $f(x)$ è vicino a $f(x_0)$.

Osserviamo subito che, sebbene l’idea di continuità vi sia già abbastanza ben formulata, la proposizione precedente non è una definizione accettabile, perché la parola “vicino”, che vi compare due volte, non ha un significato preciso. Innanzitutto, per poter misurare quanto vicino sia x a x_0 e quanto vicino sia $f(x)$ a $f(x_0)$, abbiamo bisogno di introdurre delle distanze. Più precisamente, dovremo supporre che il dominio e il codominio della funzione siano due spazi metrici.

Siano quindi 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. Possiamo riformulare il tentativo di definizione precedente come segue.

Secondo tentativo. Diremo che f è “continua” in x_0 quando si verifica la cosa seguente:

se la distanza $d_E(x, x_0)$ è piccola, allora la distanza $d_F(f(x), f(x_0))$ è piccola.

Ci rendiamo subito conto che il problema riscontrato nel primo tentativo non è stato affatto risolto con questo secondo tentativo, in quanto vi compare ora per due volte la parola “piccola”, che non ha un significato preciso. Ci chiediamo allora: *quanto piccola* vogliamo che sia la distanza $d_F(f(x), f(x_0))$? L’idea che abbiamo in mente è che questa distanza possa essere resa piccola quanto si voglia (purché la distanza $d(x, x_0)$ sia sufficientemente piccola, s’intende). Per poterla misurare, introdurremo quindi un numero reale positivo, che chiameremo ε , e chiederemo che sia $d_F(f(x), f(x_0)) < \varepsilon$, qualora $d(x, x_0)$ sia sufficientemente piccola. L’arbitrarietà di tale ε ci permetterà di prenderlo piccolo quanto si voglia.

Terzo tentativo. Diremo che f è “continua” in x_0 quando si verifica la cosa seguente: preso un qualsiasi numero $\varepsilon > 0$,

se la distanza $d_E(x, x_0)$ è piccola, allora $d_F(f(x), f(x_0)) < \varepsilon$.

Adesso la parola “piccola” compare una sola volta, mentre la distanza $d_F(f(x), f(x_0))$ viene semplicemente controllata dal numero ε . Quindi, almeno la seconda parte della proposizione ha ora un significato ben preciso. Potremmo allora cercare di fare altrettanto con la distanza $d(x, x_0)$, introducendo un nuovo numero reale positivo, che chiameremo δ , che la controlli.

Quarto tentativo (quello buono!). Diremo che f è “continua” in x_0 quando si verifica la cosa seguente: preso un qualsiasi numero $\varepsilon > 0$, è possibile trovare un numero $\delta > 0$ per cui,

se $d_E(x, x_0) < \delta$, allora $d_F(f(x), f(x_0)) < \varepsilon$.

Quest'ultima proposizione, a differenza delle precedenti, non presenta alcun termine impreciso. Le distanze $d_E(x, x_0)$ e $d_F(f(x), f(x_0))$ sono semplicemente controllate da due numeri positivi δ e ε , rispettivamente. Riscriviamola quindi in modo formale.

Definizione. 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.

Si può osservare che una o entrambe le disuguaglianze $d_E(x, x_0) < \delta$ e $d_F(f(x), f(x_0)) < \varepsilon$ possono essere sostituite rispettivamente da $d_E(x, x_0) \leq \delta$ e $d_F(f(x), f(x_0)) \leq \varepsilon$, ottenendo definizioni che sono tutte tra loro equivalenti. Questo è dovuto al fatto, da un lato, che ε è un *qualsiasi* 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.

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).

2) Supponiamo che x_0 sia un “punto isolato” di E : esiste cioè un $\rho > 0$ per cui non ci sono punti di E che distino da x_0 per meno di ρ , tranne x_0 stesso. Vediamo che, in questo caso, qualsiasi funzione $f : E \rightarrow F$ risulta continua in x_0 . Infatti, dato $\varepsilon > 0$ qualsiasi, prendendo $\delta = \rho$, avremo che $B(x_0, \delta) = \{x_0\}$, per cui $f(B(x_0, \delta)) = \{f(x_0)\} \subseteq B(f(x_0), \varepsilon)$.

3) 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

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

4) 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

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

5) Siano $E = \mathbb{R}$ e $F = \mathbb{R}$, e consideriamo la “funzione segno” $f : \mathbb{R} \rightarrow \mathbb{R}$ definita da

$$f(x) = \begin{cases} -1 & \text{se } x < 0 \\ 0 & \text{se } x = 0 \\ 1 & \text{se } x > 0. \end{cases}$$

Si può vedere che questa funzione è continua in tutti i punti tranne che in $x_0 = 0$. Infatti, se $x_0 \neq 0$, basterà prendere $\delta < |x_0|$ per avere che f è costante sull'intervallo $]x_0 - \delta, x_0 + \delta[$, quindi continua in x_0 . Per vedere che f non è continua in 0, fissiamo un $\varepsilon \in]0, 1[$; per ogni scelta di $\delta > 0$, è possibile trovare un $x \in]-\delta, \delta[$ tale che $|f(x)| = 1$, per cui $|f(x) - f(0)| > \varepsilon$.

6) La “funzione di Dirichlet” $f : \mathbb{R} \rightarrow \mathbb{R}$ è definita da

$$f(x) = \begin{cases} 1 & \text{se } x \in \mathbb{Q} \\ 0 & \text{se } x \notin \mathbb{Q}. \end{cases}$$

Questa funzione non è continua, in alcun punto x_0 . Infatti, fissato un $\varepsilon \in]0, 1[$, siccome sia $\mathbb{Q}$ che $\mathbb{R} \setminus \mathbb{Q}$ sono densi in $\mathbb{R}$, per ogni x_0 e ogni scelta di $\delta > 0$ ci saranno sicuramente un razionale x' e un irrazionale x'' in $]x_0 - \delta, x_0 + \delta[$; quindi, a seconda che x_0 sia razionale o irrazionale, si avrà che $|f(x'') - f(x_0)| > \varepsilon$ o $|f(x') - f(x_0)| > \varepsilon$.

Enunciamo ora alcune proprietà delle funzioni continue aventi come codominio $F = \mathbb{R}$.

Teorema. Se $f, g : E \rightarrow \mathbb{R}$ sono continue in x_0 , anche $f + g$ lo è.

Dimostrazione. Fissiamo $\varepsilon > 0$. Per la continuità di f e g esistono $\delta_1 > 0$ e $\delta_2 > 0$ tali che

$$\begin{aligned} d(x, x_0) < \delta_1 &\Rightarrow |f(x) - f(x_0)| < \varepsilon, \\ d(x, x_0) < \delta_2 &\Rightarrow |g(x) - g(x_0)| < \varepsilon. \end{aligned}$$

Quindi, se $\delta = \min\{\delta_1, \delta_2\}$, si ha

$$d(x, x_0) < \delta \Rightarrow |(f + g)(x) - (f + g)(x_0)| \leq |f(x) - f(x_0)| + |g(x) - g(x_0)| < 2\varepsilon.$$

Data l'arbitrarietà di ε , ciò dimostra che $f + g$ è continua in x_0 . ■

Teorema. Se $f, g : E \rightarrow \mathbb{R}$ sono continue in x_0 , anche $f \cdot g$ lo è.

Dimostrazione. Fissiamo $\varepsilon > 0$. Non è restrittivo supporre $\varepsilon \leq 1$, in quanto possiamo sempre porre $\varepsilon' = \min\{\varepsilon, 1\}$ e procedere con ε' al posto di ε . Per la continuità di f e g esistono $\delta_1 > 0$ e $\delta_2 > 0$ tali che

$$\begin{aligned} d(x, x_0) < \delta_1 &\Rightarrow |f(x) - f(x_0)| < \varepsilon, \\ d(x, x_0) < \delta_2 &\Rightarrow |g(x) - g(x_0)| < \varepsilon. \end{aligned}$$

Notiamo che, essendo $\varepsilon \leq 1$, da $|f(x) - f(x_0)| < \varepsilon$ segue che $|f(x)| < |f(x_0)| + 1$. Quindi, se $\delta = \min\{\delta_1, \delta_2\}$, si ha

$$\begin{aligned} d(x, x_0) < \delta &\Rightarrow |(f \cdot g)(x) - (f \cdot g)(x_0)| = \\ &= |f(x)g(x) - f(x)g(x_0) + f(x)g(x_0) - f(x_0)g(x_0)| \\ &\leq |f(x)| \cdot |g(x) - g(x_0)| + |g(x_0)| \cdot |f(x) - f(x_0)| \\ &\leq (|f(x_0)| + 1) \cdot |g(x) - g(x_0)| + |g(x_0)| \cdot |f(x) - f(x_0)| \\ &< (|f(x_0)| + |g(x_0)| + 1)\varepsilon. \end{aligned}$$

Data l'arbitrarietà di ε , ciò dimostra che $f \cdot g$ è continua in x_0 . ■

Teorema. Se $f, g : E \rightarrow \mathbb{R}$ sono continue in x_0 , anche $f - g$ lo è.

Dimostrazione. Segue immediatamente dai due teoremi precedenti e dal fatto che ogni funzione costante è continua, in quanto $f - g = f + (-1) \cdot g$. ■

Teorema (della permanenza del segno). Se $g : E \rightarrow \mathbb{R}$ è continua in x_0 e $g(x_0) > 0$, allora esiste un $\delta > 0$ tale che

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

Dimostrazione. Fissiamo $\varepsilon = g(x_0)$. Per la continuità, esiste un $\delta > 0$ tale che

$$d(x, x_0) < \delta \Rightarrow g(x_0) - \varepsilon < g(x) < g(x_0) + \varepsilon \Rightarrow 0 < g(x) < 2g(x_0).$$

■

Naturalmente, se $g(x_0) < 0$, allora esiste un $\delta > 0$ tale che

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

Teorema. 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 .

Dimostrazione. Si noti che, per la proprietà di permanenza del segno, esiste un $\rho > 0$ tale che il rapporto $\frac{f(x)}{g(x)}$ è definito almeno per tutti gli x di E che distano da x_0 per meno di ρ . Essendo $\frac{f}{g} = f \cdot \frac{1}{g}$, basterà dimostrare che $\frac{1}{g}$ è continua in x_0 . Fissiamo $\varepsilon > 0$; possiamo supporre senza perdita di generalità che $\varepsilon < \frac{|g(x_0)|}{2}$. Per la continuità di g , esiste un $\delta > 0$ tale che

$$d(x, x_0) < \delta \Rightarrow |g(x) - g(x_0)| < \varepsilon.$$

Ma allora, essendo $\varepsilon < \frac{|g(x_0)|}{2}$, anche

$$d(x, x_0) < \delta \Rightarrow |g(x)| > |g(x_0)| - \varepsilon > \frac{|g(x_0)|}{2}.$$

Ne segue che

$$d(x, x_0) < \delta \Rightarrow \left| \frac{1}{g}(x) - \frac{1}{g}(x_0) \right| = \frac{|g(x_0) - g(x)|}{|g(x)g(x_0)|} < \frac{2}{|g(x_0)|^2} \varepsilon.$$

Per l'arbitrarietà di ε , questo dimostra che $\frac{1}{g}$ è continua in x_0 . ■

Sappiamo da quanto sopra che le funzioni costanti sono continue, così come la funzione $f(x) = x$. Usando i teoremi precedenti, abbiamo quindi che tutte le funzioni polinomiali sono continue, così come le funzioni razionali, definite dal rapporto di due polinomi. Più precisamente, esse sono continue sul loro dominio, ossia sull'insieme dei punti in cui il denominatore non si annulla.

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

Teorema. Siano $f : E \rightarrow F$ 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 $[g \circ f](U) \subseteq W$. ■

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 .

Definizione. Se esiste un $l \in F$ tale che la funzione $\tilde{f} : E \rightarrow F$, definita da

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

risulti continua in x_0 , si dice che l è il “limite di f in x_0 ”, o anche “limite di $f(x)$ per x che tende a x_0 ” e si scrive

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

In altri termini, si ha che l è il limite di f in x_0 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 $x \rightarrow x_0$.

Sappiamo che, se x_0 è un punto isolato, ogni funzione risulterà continua in x_0 . Il problema non presenta pertanto alcun interesse in questo caso. Supporremo quindi che x_0 non sia un punto isolato, ossia che x_0 sia un “punto di accumulazione” di E : ogni intorno di x_0 contiene punti di E distinti da x_0 stesso.³ Nel seguito, supporremo sempre che x_0 sia un punto di accumulazione di E .

³Un semplice ragionamento mostra che, in questo caso, ogni intorno di x_0 contiene infiniti punti di E . Trovatone uno, x_1 , si prende un intorno di x_0 che non lo contenga, nel quale se ne può trovare un secondo, x_2 , e così via...

Per cominciare, verifichiamo l'unicità del limite.

Teorema. *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. ■

Il seguente teorema è una riformulazione del legame stretto che intercorre tra i concetti di limite e di continuità.

Teorema. *Considerata la funzione $f : E \rightarrow F$, si ha che*

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

Dimostrazione. In questo caso, si ha che la funzione $\tilde{f}$ coincide con f . ■

Un'osservazione generale, che potrebbe essere utile in seguito: si ha

$$\lim_{x \rightarrow x_0} f(x) = l \Leftrightarrow \lim_{x \rightarrow x_0} d(f(x), l) = 0.$$

Esempi. 1. Cominciamo con la funzione $f : \mathbb{R} \setminus \{0\} \rightarrow \mathbb{R}$, definita da

$$f(x) = \sin_T\left(\frac{1}{x}\right).$$

Se $x_0 = 0$, il limite di f non esiste, perché in ogni intorno di 0 ci sono valori di x per cui $f(x) = 1$ e valori di x per cui $f(x) = -1$.

2. Dimostriamo invece che

$$\lim_{x \rightarrow 0} x \sin_T\left(\frac{1}{x}\right) = 0.$$

Per far questo, è utile osservare che

$$\left| x \sin_T\left(\frac{1}{x}\right) \right| \leq |x|, \quad \text{per ogni } x \in \mathbb{R} \setminus \{0\}.$$

Ecco allora che, fissato $\varepsilon > 0$, basta prendere $\delta = \varepsilon$ per avere che

$$0 < |x - 0| < \delta \Rightarrow |f(x) - 0| < \varepsilon.$$

Funzioni a valori in $\mathbb{R}$

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

Teorema (della permanenza del segno). *Se*

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

allora esiste un $\delta > 0$ tale che

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

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. *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

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

Naturalmente, si hanno enunciati analoghi qualora g sia di segno opposto.

Teorema. *Se*

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

allora

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

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

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

se $l_2 \neq 0$,

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

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. 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

$$0 < d(x, x_0) < \delta_1 \Rightarrow l - \varepsilon < f_1(x) < l + \varepsilon,$$

$$0 < d(x, x_0) < \delta_2 \Rightarrow l - \varepsilon < f_2(x) < l + \varepsilon.$$

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. ■

Corollario. Se

$$\lim_{x \rightarrow x_0} f(x) = 0,$$

ed esiste un $C > 0$ tale che $|g(x)| \leq C$ per ogni x , allora

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

Dimostrazione. Si ha

$$-C|f(x)| \leq f(x)g(x) \leq C|f(x)|,$$

e il risultato segue dal teorema precedente, tenuto conto che si ha

$$\lim_{x \rightarrow x_0} f(x) = 0 \Leftrightarrow \lim_{x \rightarrow x_0} |f(x)| = 0,$$

prendendo $f_1(x) = -C|f(x)|$ e $f_2(x) = C|f(x)|$. ■

Cambiamento di variabile nel limite

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

Teorema 1. *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. Riguardando la definizione di limite, si ha che $\tilde{f} : E \rightarrow \mathbb{R}$ ivi definita 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 2. *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,

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

■

Alcune considerazioni sull'ultimo teorema dimostrato. Si noti che la sua conclusione 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.

Esempi. 1. Dimostriamo che

$$\lim_{x \rightarrow 0} \cos\left(x \sin\left(\frac{1}{x}\right)\right) = 1.$$

In effetti, se $f : \mathbb{R} \setminus \{0\} \rightarrow \mathbb{R}$ è definita da $f(x) = x \sin(1/x)$ e $g : \mathbb{R} \rightarrow \mathbb{R}$ è definita da $g(y) = \cos(y)$, sappiamo che $\lim_{x \rightarrow 0} f(x) = 0$, e che g è continua. Per il Teorema 1,

$$\lim_{x \rightarrow 0} g(f(x)) = g\left(\lim_{x \rightarrow 0} f(x)\right) = g(0) = 1.$$

2. Sia ora f come nell'esempio precedente, e sia $g : \mathbb{R} \rightarrow \mathbb{R}$ definita da

$$g(y) = \begin{cases} 1 & \text{se } y \neq 0, \\ 2 & \text{se } y = 0. \end{cases}$$

Si può vedere che, in ogni intorno di $x_0 = 0$, la funzione $g(f(x))$ assume infinite volte il valore 1 e infinite volte il valore 2. Pertanto, in questo caso,

$$\text{il limite } \lim_{x \rightarrow 0} g(f(x)) \text{ non esiste.}$$

Limite delle restrizioni

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 $f : \hat{E} \setminus \{x_0\} \rightarrow F$, con $\hat{E} \subseteq E$, purché x_0 sia di accumulazione per $\hat{E}$: ogni intorno di x_0 deve contenere infiniti punti di $\hat{E}$.

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

Teorema. *Se esiste il limite di f in x_0 e x_0 è di accumulazione anche per $\hat{E}$, 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.

Esempi. 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 $\{(x, y) : x = y\}$.

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 $\frac{1}{2}$.

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

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

Sia ora $E \subseteq \mathbb{R}$. Possiamo considerare le due restrizioni $\hat{f}_1$ e $\hat{f}_2$ agli insiemi $\hat{E}_1 = E \cap]-\infty, x_0]$ e $\hat{E}_2 = E \cap [x_0, +\infty[$. Se x_0 è di accumulazione per $\hat{E}_1$, chiameremo “limite sinistro” di f , quando esiste, il limite di $\hat{f}_1(x)$ per x che tende a x_0 ; lo denoteremo con

$$\lim_{x \rightarrow x_0^-} f(x).$$

Analogamente, se x_0 è di accumulazione per $\hat{E}_2$, chiameremo “limite destro” di f , quando esiste, il limite di $\hat{f}_2(x)$ per x che tende a x_0 ; lo denoteremo con

$$\lim_{x \rightarrow x_0^+} f(x).$$

Si può dimostrare il fatto seguente.

Teorema. *Se x_0 è di accumulazione per $\hat{E}_1$ e per $\hat{E}_2$, il limite di $f(x)$ per x che tende a x_0 esiste se e solo se esistono sia il limite sinistro che il limite destro e hanno lo stesso valore.*

Dimostrazione. Sappiamo già che, se esiste il limite, tutte le restrizioni devono avere lo stesso limite. Viceversa, supponiamo che esistano e coincidano i limiti sinistro e destro, e sia ℓ il loro valore. Fissiamo un $\varepsilon > 0$. Allora esistono $\delta_1 > 0$ e $\delta_2 > 0$ tali che, se $x \in E$,

$$x_0 - \delta_1 < x < x_0 \Rightarrow d(f(x), \ell) < \varepsilon,$$

$$x_0 < x < x_0 + \delta_2 \Rightarrow d(f(x), \ell) < \varepsilon.$$

Preso $\delta = \min\{\delta_1, \delta_2\}$, abbiamo quindi che, se $x \neq x_0$,

$$x_0 - \delta < x < x_0 + \delta \Rightarrow d(f(x), \ell) < \varepsilon,$$

per cui il limite di f in x_0 esiste ed è uguale a ℓ . ■

Esempio. La funzione “segno”, ossia $f : \mathbb{R} \rightarrow \mathbb{R}$ definita da

$$f(x) = \begin{cases} 1 & \text{se } x > 0 \\ 0 & \text{se } x = 0 \\ -1 & \text{se } x < 0 \end{cases}$$

non ha limite in $x_0 = 0$, essendo che $\lim_{x \rightarrow 0^-} f(x) = -1$ e $\lim_{x \rightarrow 0^+} f(x) = 1$.

La retta ampliata

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

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

Si tratta di una funzione invertibile, con inversa $\varphi^{-1} :]-1, 1[\rightarrow \mathbb{R}$, definita da

$$\varphi^{-1}(y) = \frac{y}{1 - |y|}.$$

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

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

Per la nuova distanza, la palla aperta di centro $x_0 \in \mathbb{R}$ e raggio ρ è data da

$$\tilde{B}(x_0, \rho) = \{x : |\varphi(x) - \varphi(x_0)| < \rho\}.$$

È importante notare che gli intorni di un punto $x_0 \in \mathbb{R}$ rimangono gli stessi di quelli definiti dalla distanza usuale in $\mathbb{R}$. Infatti, essendo φ continua in x_0 , per ogni $\rho_1 > 0$ esiste un $\rho_2 > 0$ tale che

$$|x - x_0| < \rho_2 \quad \Rightarrow \quad |\varphi(x) - \varphi(x_0)| < \rho_1,$$

ossia

$$]x_0 - \rho_2, x_0 + \rho_2[\subseteq \tilde{B}(x_0, \rho_1).$$

Viceversa, essendo φ^{-1} continua in $y_0 = \varphi(x_0) \in]-1, 1[$, per ogni $\rho_1 > 0$ esiste un $\rho_2 > 0$ tale che

$$|y - y_0| < \rho_2 \quad \Rightarrow \quad y \in]-1, 1[\quad \text{e} \quad |\varphi^{-1}(y) - \varphi^{-1}(y_0)| < \rho_1;$$

In particolare, prendendo $y = \varphi(x)$,

$$|\varphi(x) - \varphi(x_0)| < \rho_2 \quad \Rightarrow \quad \varphi(x) \in]-1, 1[\quad \text{e} \quad |x - x_0| < \rho_1,$$

ossia

$$\tilde{B}(x_0, \rho_2) \subseteq]x_0 - \rho_1, x_0 + \rho_1[.$$

Da quanto visto, si deduce 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

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

Essa è invertibile, con inversa $\tilde{\varphi}^{-1} : [-1, 1] \rightarrow \tilde{\mathbb{R}}$ definita da

$$\tilde{\varphi}^{-1}(y) = \begin{cases} -\infty & \text{se } y = -1, \\ \varphi^{-1}(y) & \text{se } y \in]-1, 1[, \\ +\infty & \text{se } y = 1. \end{cases}$$

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. Vediamo ad esempio cos'è una palla aperta centrata in $+\infty$:

$$B(+\infty, \rho) = \{x \in \tilde{\mathbb{R}} : |\tilde{\varphi}(x) - 1| < \rho\} = \{x \in \tilde{\mathbb{R}} : \tilde{\varphi}(x) > 1 - \rho\},$$

e quindi

$$B(+\infty, \rho) = \begin{cases} \tilde{\mathbb{R}} & \text{se } \rho > 2, \\]-\infty, +\infty] & \text{se } \rho = 2, \\]\varphi^{-1}(1 - \rho), +\infty] & \text{se } \rho < 2, \end{cases}$$

dove abbiamo usato le notazioni

$$]a, +\infty] = \{x \in \tilde{\mathbb{R}} : x > a\} =]a, +\infty[\cup \{+\infty\}.$$

Possiamo quindi affermare 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 $\widetilde{\mathbb{R}}$, si ha che $+\infty$ è punto di accumulazione per E se e solo se E non è limitato superiormente. In tal caso, si ha:

$$\begin{aligned} \lim_{x \rightarrow +\infty} f(x) = l \in F &\Leftrightarrow \forall V \text{ intorno di } l \ \exists U \text{ intorno di } +\infty : \\ &f(U \cap E) \subseteq V \\ &\Leftrightarrow \forall \varepsilon > 0 \ \exists \alpha \in \mathbb{R} : \ x > \alpha \Rightarrow d(f(x), l) < \varepsilon . \end{aligned}$$

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

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

Si noti che

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

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:

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

analogamente,

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

Si noti che

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

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

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

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,

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