

srino le seguenti tautologie:

$$\cdot (\mathcal{D} \Rightarrow \mathcal{D}) \Rightarrow ((\mathcal{D} \Rightarrow \mathcal{D}) \Rightarrow (\mathcal{D} \Rightarrow \mathcal{D})) \quad (\text{proprietà transitiva dell'implicazione}),$$

$$\cdot \mathcal{D} \Rightarrow \mathcal{D},$$

$$\cdot \mathcal{D} \Rightarrow \text{non}(\text{non} \mathcal{D}),$$

$$\cdot (\mathcal{D} \Rightarrow \mathcal{D}) \Rightarrow ((\text{non} \mathcal{D}) \Rightarrow (\text{non} \mathcal{D}));$$

(la scrittura $\mathcal{D} \Rightarrow \mathcal{D}$ deve essere sempre ritenuta equivalente a " $(\text{non} \mathcal{D}) \vee \mathcal{D}$ ").

2. Gli insiemi

In questa esposizione assumiamo come primitiva l'idea di *insieme*, così come ce la presenta la nostra intuizione. Al lettore sono certamente familiari espressioni come "l'insieme degli interi naturali", "l'insieme delle rette del piano" ecc. Consideriamo anche i termini di *aggregato*, *classe*, *famiglia* come sinonimi di *insieme*.

L'espressione

$$x \in T$$

[2.1]

si legge: *x appartiene all'insieme T* o, indifferentemente, *x è un elemento di T*. Ad esempio, se indichiamo con $\mathbb{R}$ l'insieme di tutti i numeri reali, è vero che $\sqrt{2} \in \mathbb{R}$.

Scriveremo $x \notin T$ per negare la [2.1], cioè per affermare che *x non appartiene a T*.

L'intuizione ci suggerisce di considerare uguali due insiemi che abbiano gli stessi elementi; in simboli:

$$(\forall x: x \in A \Leftrightarrow x \in B) \Leftrightarrow A = B.$$

[2.2]

Il significato che abbiamo dato al segno di eguaglianza ci assicura che due insiemi aventi gli stessi elementi devono avere il medesimo ruolo in tutte le enunciazioni che fanno parte della nostra teoria.

Per indicare un insieme basterà allora (quando è possibile) elencarne gli elementi; converremo di elencarli entro parentesi graffe. Ad esempio, $\{2, 5, 6\}$ indica l'insieme i cui elementi sono i numeri 2, 5, 6. La notazione $\{2\}$ indica l'insieme costituito dal solo elemento 2: nella nostra teoria, preso un qualsiasi oggetto a , ci riserviamo il diritto di considerare l'insieme $\{a\}$ che ha come

unico elemento a ; occorre fare bene attenzione a non confondere a con $\{a\}$...

Se ogni elemento di A è un elemento di B , cioè:

$$\forall x: x \in A \Rightarrow x \in B,$$

[2.3]

allora diciamo che *A è contenuto in B* (fig. 2.1), oppure *A è una parte o un sottoinsieme di B*, e scriviamo $A \subset B$. Questa relazione è detta di *inclusione*. Notiamo che la relazione $A \subset B$ è verificata quando è $A = B$. Per esprimere che è $A \subset B$, ma è $A \neq B$, cioè esistono elementi di B che non sono elementi di A , si scrive $A \subsetneq B$ (si dice allora che *A è sottoinsieme proprio di B*, ovvero è *propriamente contenuto in B*).

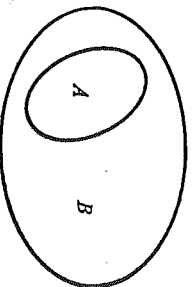


Figura 2.1

Convienne introdurre un particolare insieme privo di elementi, che verrà detto *insieme vuoto*.

Poiché gli insiemi si distinguono solo in base ai loro elementi, e l'insieme vuoto... non ne ha, si deve ritenere che esiste un unico insieme vuoto; esso verrà indicato col simbolo $\emptyset$. L'insieme vuoto è molto utile nello sviluppo formale della teoria degli insiemi, come vedremo tra poco.

Poste queste premesse, il compito che ci spetta in questo paragrafo e nei successivi è quello di esporre alcune importanti costruzioni insiemistiche, cioè procedimenti che, partendo da insiemi assegnati, ci forniscono nuovi insiemi.

***2.1 Osservazione.** Nell'introdurre la nozione di insieme ci siamo basati solo sulla suggestione della nostra intuizione. Accenniamo ora, per sommi capi, al modo con cui si può costruire una teoria assiomatica degli insiemi, cioè una teoria che si basi solo su alcune proposizioni primitive espresse nel linguaggio logico (assiomi).

Oltre ai simboli logici, la nostra teoria deve contenere il simbolo $\in$ che, come risulta dalla [2.1] esprime una relazione, cioè un predicato in due variabili. (Naturalmente, la [2.1] deve essere intesa in modo puramente formale, prescindendo dal significato intuitivo che può esprimere.) Altri predicati si possono derivare da $\in$ mediante un'opportuna definizione (si pensi, ad esempio, alla relazione $A \subset B$, espressa dalla [2.3]). Dunque, la relazione $\in$ può essere posta a fondamento di tutta la teoria degli insiemi; gli assiomi della teoria degli insiemi non fanno che assegnare le proprietà formali della relazione $\in$.

Come abbiamo detto, non esporremo per esteso gli assiomi della teoria degli insiemi, volendo mantenerci a un livello intuitivo: formuleremo solo in modo esplicito, più avanti, un assioma (l'assioma della scelta) che, per il suo carattere speciale, merita di essere segnalato a parte.¹

Notiamo ancora che la relazione [2.1] comporta una gerarchia tra elemento e insieme; ma non si tratta di una gerarchia a ruoli fissi: in altre parole, una stessa lettera può indicare a volte un elemento, a volte un insieme. Ad esempio, una retta di un piano può essere considerata a volte come elemento (dell'insieme di tutte le rette del piano, ad esempio), a volte come insieme (insieme di tutti i suoi punti). Osserviamo poi che nella [2.1] e nella [2.2] si usano le lettere maiuscole per indicare gli insiemi, le lettere minuscole per indicare gli elementi: questo artificio grafico è utile per facilitare la comprensione (e noi lo useremo, in seguito, tutte le volte che potremo) ma è evidente, da quanto detto, che esso non ha alcun fondamento concettuale! Aggiungiamo, infine, che, in una presentazione assiomatica della teoria degli insiemi, è comodo avere a che fare con enti di uno stesso tipo, e cioè solo con insiemi: si fa allora in modo (sia pure un po' artificiosamente) che ogni elemento di un qualsiasi insieme possa, a sua volta, essere considerato come un insieme.

Supponiamo che T sia un insieme e che $\mathscr{P}(x)$ sia una proprietà che ha senso per gli elementi di T ; allora scriviamo

$$\{x: (x \in T) \text{ e } \mathscr{P}(x)\} \quad [2.4]$$

¹ La situazione ha qualche analogia con quella che si trova in geometria elementare per l'assioma delle parallele, che non manca di contenuto intuitivo, ma è, in un certo senso, meno centrale degli altri.

per indicare il sottoinsieme di T formato dagli elementi per cui $\mathscr{P}(x)$ è vera.

Ad esempio, l'espressione

$$\{x: (x \in \mathbb{R}) \text{ e } (\sin x = 0)\}$$

rappresenta l'insieme dei multipli interi di π .

***2.2 Osservazione.** Naturalmente, il fatto che la [2.4] rappresenti un insieme dovrebbe essere postulato, cioè posto come assioma. Lo studioso avrà poi notato che abbiamo parlato di "proprietà che ha senso"; è chiaro che occorre prendere qualche precauzione, dal momento che non avrebbe senso, ad esempio, parlare del "sottoinsieme di $\mathbb{R}$ (retta reale) costituito dai numeri veri". Occorre peraltro rendersi conto che quando si è in una certa teoria, si ha a disposizione solo il linguaggio di quella teoria: qui le proprietà che hanno senso sono quelle espresse attraverso il linguaggio della teoria degli insiemi (e quindi, se si vuole, attraverso i simboli logici e il simbolo $\in$, magari con una lunga serie di manipolazioni, come sarebbe per la definizione della funzione seno).

***2.3 Osservazione.** L'espressione [2.4] permette di definire un insieme, avendo però a disposizione un insieme T già assegnato. Se non prendiamo questa precauzione, cioè se consideriamo semplicemente, in luogo della [2.4] un'espressione del tipo $\{x: \mathscr{P}(x)\}$, non possiamo pretendere che questa rappresenti un insieme, anche se la proprietà $\mathscr{P}(x)$ è espressa nel linguaggio della teoria degli insiemi. Costatiamolo con un esempio: ammettiamo per un momento che l'espressione $\{x: x \notin x\}$ rappresenti un insieme s . (Questo sarebbe dunque "l'insieme di tutti gli insiemi che non contengono sé stessi come elementi".) È chiaro che se è $s \in s$, allora deve essere $s \notin s$ (perché $x \notin x$ è la proprietà caratteristica di s), mentre se è $s \notin s$, allora, per questo stesso fatto, dovrebbe essere $s \in s$. Dunque, l'aver ammesso che $\{x: x \notin x\}$ rappresenti un insieme conduce a un paradosso (è il celebre paradosso di Russell).

Notiamo poi che nelle usuali presentazioni assiomatiche della teoria degli insiemi si esclude che un insieme possa contenere sé stesso come elemento (in conformità di quella gerarchia elemento-insieme a cui abbiamo fatto cenno nell'osservazione 2.1).

Dunque, se la proprietà $x \notin x$ dovesse definire un insieme, questo sarebbe “l'insieme di tutti gli insiemi”: il paradosso di Russell, ci porta allora a riconoscere l'impossibilità di un tale insieme; in altre parole, nella teoria degli insiemi non ci può essere un “universo”.

Infine notiamo che, nei casi concreti, l'indicazione dell'insieme T nella [2.4] può essere sottintesa; in molte questioni di analisi può essere sottinteso, ad esempio, che T coincide con la retta reale, o con il piano complesso ecc.

Esponiamo ora alcune fondamentali operazioni insiemistiche che si possono definire direttamente mediante la relazione di appartenenza.

1) *Unione*. L'unione di due insiemi A, B è l'insieme formato da quegli elementi che appartengono a uno almeno dei due insiemi A, B (fig. 2.2). Esso viene indicato con $A \cup B$. In simboli:

$$A \cup B \stackrel{\text{def}}{=} \{x: (x \in A) \text{ o } (x \in B)\}. \quad [2.5]$$

*2.4 *Osservazione*. Il secondo membro della [2.5] non rientra nella forma dell'espressione [2.4]. Perciò, in una formulazione precisa occorrerebbe anzitutto affermare con un assioma che il secondo membro della [2.5] rappresenta un insieme.

2) *Intersezione*. L'intersezione di due insiemi A, B è l'insieme formato dagli elementi che appartengono sia ad A che a B (fig. 2.3). Esso si indica con il simbolo $A \cap B$:

$$A \cap B \stackrel{\text{def}}{=} \{x: (x \in A) \text{ e } (x \in B)\}. \quad [2.6]$$

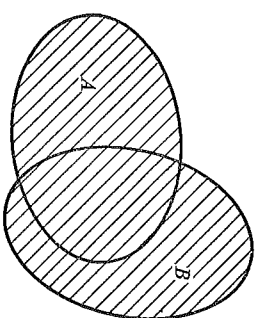


Figura 2.2

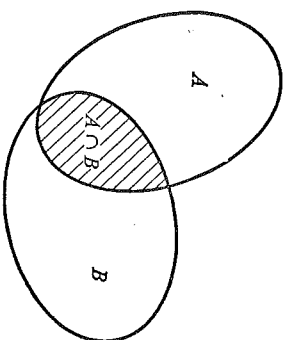


Figura 2.3

Può accadere che i due insiemi A, B siano *disgiunti* (cioè non abbiano alcun elemento in comune). Allora si scrive: $A \cap B = \emptyset$.

3) *Differenza*. La differenza di B da A , che si indica con $A \setminus B$, oppure con $A - B$, è l'insieme degli elementi di A che non appartengono a B (fig. 2.4). In simboli:

$$A \setminus B \stackrel{\text{def}}{=} \{x: (x \in A) \text{ e } (x \notin B)\}. \quad [2.7]$$

4) *Differenza simmetrica*. La differenza simmetrica A, B che si indica col simbolo $A \Delta B$ è l'insieme costituito dagli elementi che appartengono a uno solo degli insiemi A, B (fig. 2.5), cioè

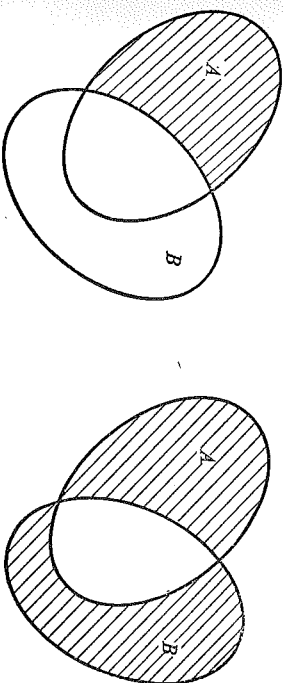


Figura 2.4

Figura 2.5

dagli elementi di A che non appartengono a B e dagli elementi di B che non appartengono ad A . Perciò

$$A \Delta B \stackrel{\text{def}}{=} (A \setminus B) \cup (B \setminus A) = (A \cup B) \setminus (A \cap B). \quad [2.8]$$

Dato un insieme T , ammettiamo di poter considerare un insieme i cui elementi sono tutte le parti (o sottoinsiemi) di T ; questo insieme verrà indicato con $\mathcal{P}(T)$. Ad esempio, se è $T = \{a, b, c\}$, si ha:

$$\mathcal{P}(T) = \{\emptyset, \{a\}, \{b\}, \{c\}, \{a, b\}, \{b, c\}, \{a, c\}, \{a, b, c\}\}.$$

Le operazioni $\cup$ e $\cap$, nell'insieme $\mathcal{P}(T)$, hanno proprietà

formali molto interessanti. Elenchiamone alcune:

$A \cup A = A$	(proprietà di idempotenza)	$A \cap A = A$
$A \cup B = B \cup A$	(proprietà commutativa)	$A \cap B = B \cap A$
$(A \cup B) \cup C =$ $= A \cup (B \cup C)$	(proprietà associativa)	$(A \cap B) \cap C =$ $= A \cap (B \cap C)$
$A \cup (B \cap C) =$ $= (A \cup B) \cap (A \cup C)$	(proprietà distributiva di una operazione rispetto all'altra)	$A \cap (B \cup C) =$ $= (A \cap B) \cup (A \cap C)$
$A \cap (A \cup B) = A$	(proprietà di assorbimento)	$A \cup (A \cap B) = A$

Abbiamo raggruppato queste proprietà su due colonne per mettere in evidenza un aspetto interessante: le proprietà scritte in una colonna si ottengono dalle corrispondenti scritte nell'altra scambiando fra loro i segni $\cup$ e $\cap$. Questa proprietà (che è detta di *dualità*) si spiega facilmente introducendo l'operazione di *complementare*: se $A \in \mathcal{F}(T)$, chiamiamo complementare di A rispetto a T l'insieme $A' = T \setminus A$.

Valgono allora le seguenti proprietà:

$$(A')' = A, \quad [2.10]$$

$$(A \cup B)' = A' \cap B', \quad (A \cap B)' = A' \cup B'. \quad [2.11]$$

Notiamo anzitutto che la [2.10] permette di ricavare una delle [2.11] dall'altra (si sostituisce in luogo di A , A' , in luogo di B , B' , e si prenda il complementare di ambo i membri). Con analogia tecnica si dimostrano le relazioni contenute in una delle colonne della [2.9], partendo dalle corrispondenti dell'altra colonna.

***2.5 Osservazione.** Si dice *algebra di Boole* la struttura che si ottiene assegnando in un insieme (nel nostro caso $\mathcal{F}(T)$) tre operazioni: $\cup$, $\cap$, $'$, aventi le proprietà [2.9], [2.10], [2.11]. È interessante notare che le proprietà [2.9], [2.10], [2.11] valgono anche per la logica delle proposizioni, pur di sostituire il simbolo $\cup$

con il connettivo "o", il simbolo $\cap$ con il connettivo "e", il simbolo $'$ con il connettivo "non", sostituendo poi il simbolo $=$ con il connettivo $\Leftrightarrow$ (di doppia implicazione). Questo, del resto non fa meraviglia se si pensa al modo con cui sono state definite la relazione di eguaglianza fra gli insiemi, e le operazioni insiemistiche (vedi le [2.2], [2.5], [2.6], [2.7]).

Le operazioni di unione e di intersezione si possono estendere al caso di famiglie infinite di insiemi.

Sia $\mathcal{F}$ una famiglia qualunque di insiemi; si indica con $\bigcup_{X \in \mathcal{F}} X$ l'insieme costituito dagli elementi che appartengono a *qualcuno* degli insiemi $X \in \mathcal{F}$. Questo insieme viene detto *insieme-unione* della famiglia $\mathcal{F}$. In simboli:

$$\bigcup_{X \in \mathcal{F}} X = \{x: (\exists X: X \in \mathcal{F} \text{ e } x \in X)\}. \quad [2.12]$$

***2.6 Osservazione.** In realtà, il fatto che la [2.12] definisca un insieme dovrebbe essere un assioma della nostra teoria. Questo assioma contiene come caso particolare quello relativo a due insiemi (a cui abbiamo fatto cenno nell'osservazione 2.3).

Si indica con $\bigcap_{X \in \mathcal{F}} X$ l'insieme costituito dagli elementi che appartengono a *ciascuno* degli insiemi $X \in \mathcal{F}$. Questo insieme viene detto *insieme-intersezione* della famiglia $\mathcal{F}$. In simboli:

$$\bigcap_{X \in \mathcal{F}} X = \{x: (\forall X: X \in \mathcal{F} \Rightarrow x \in X)\}. \quad [2.13]$$

Le considerazioni generali che abbiamo fatto negli insiemi ci porgono l'occasione di una riflessione sui procedimenti che si seguono abitualmente per risolvere le equazioni. Sia T un insieme e sia $\mathcal{P}(x)$ un predicato che abbia senso in T ; possiamo interpretare $\mathcal{P}(x)$ come un'equazione posta in T . L'insieme: $P = \{x: (x \in T) \text{ e } \mathcal{P}(x)\}$ è evidentemente l'insieme delle *soluzioni* di $\mathcal{P}(x)$ in T . Il punto di vista da cui ci poniamo in questa presentazione è molto generale: nel nostro concetto di equazione rientrano, ad esempio, anche quelle che comunemente vengono dette disequazioni, e tante altre cose.

Sia ora $\mathcal{Q}(x)$ un secondo predicato e sia $Q = \{x: (x \in T) \text{ e } \mathcal{Q}(x)\}$; se per ogni $x \in T$ è vera la $\mathcal{Q}(x) \Rightarrow \mathcal{P}(x)$, allora (vedi la [2.3]) si ha $P \subset Q$. In pratica, si cerca $\mathcal{Q}(x)$ tale che l'insieme delle solu-

zioni Q sia facilmente determinabile. Dunque, Q contiene tutte le soluzioni dell'equazione $\mathcal{P}(x)$, ma può contenere altri elementi. Ad esempio supponiamo che T sia la retta reale e che $\mathcal{P}(x)$ significhi $\sqrt{x^2+1}=2x$ e $\mathcal{Q}(x)$ significhi $x^2+1=4x^2$. Allora, per ogni x è vera $\mathcal{P}(x) \Rightarrow \mathcal{Q}(x)$. Si ha $Q = \{1/\sqrt{3}, -1/\sqrt{3}\}$ ma il valore $-1/\sqrt{3}$ non verifica la $\mathcal{P}(x)$.

Il lettore sa già che certe manipolazioni di carattere algebrico possono introdurre "soluzioni estranee", da quanto abbiamo detto sopra risulta chiaramente il carattere generale di questo fenomeno.

Se invece, per ogni $x \in T$ è vera $\mathcal{P}(x) \Leftrightarrow \mathcal{Q}(x)$, allora è $P = Q$; le due equazioni hanno le stesse soluzioni, e si dicono equivalenti.

Concludiamo questo paragrafo con alcune convenzioni a cui ci atterremo in seguito per quello che riguarda l'impiego dei quantificatori nella teoria degli insiemi.

Per dire che tutti gli elementi di un insieme T hanno una proprietà $\mathcal{P}(x)$, anziché impiegare la scrittura esatta $\forall x: (x \in T) \Rightarrow \mathcal{P}(x)$, scriveremo più brevemente:

$$\forall x \in T: \mathcal{P}(x). \quad [2.14]$$

Analogamente, per affermare che in T c'è almeno un elemento avente la proprietà $\mathcal{P}(x)$, anziché scrivere $\exists x: (x \in T) \wedge \mathcal{P}(x)$, scriveremo:

$$\exists x \in T: \mathcal{P}(x). \quad [2.15]$$

Esercizi

1. Scrivere, in luogo di ..., il segno $\in$, o il segno $\subset$ (o anche ambedue):
 $a \dots \{a\}, \{a\} \dots \{a, b\}, a \dots \{a, \{a\}, b\}, \emptyset \dots \{a\},$
 $\{a\} \dots \{a, \{a\}, b\}, \{a, b\} \dots \{b, a\}, \emptyset \dots \emptyset.$
2. Di quanti elementi consta l'insieme $\mathcal{P}(\mathcal{P}(\mathcal{P}(\emptyset)))$?
3. Si dimostri che $((A \subseteq B) \wedge (B \subseteq C)) \Rightarrow A \subseteq C$.
4. Si dimostri che $(A \Delta B = \emptyset) \Leftrightarrow (A = B)$.
5. Si rappresentino i seguenti sottoinsiemi di $\mathbb{R}$:
 $\{x: -2 \leq 3x \leq 5\}, \{x: (x-1)^2 \leq (2x-1)^2 - 4\},$
 $\{x: x \geq 8, \sqrt{x-3} + \sqrt{x-8} < 5\}.$

6. È vero che $A \cap B = \emptyset \Rightarrow A \neq B$?

7. Si dimostri che $((A \cap B) \cup C = A \cap (B \cup C)) \Leftrightarrow C \subseteq A$.

8. Si dimostrino le seguenti regole, del tutto evidenti per il buon senso logico, che riguardano la negazione delle [2.14] e [2.15]:

$$\text{non } (\forall x \in T: \mathcal{P}(x)) \Leftrightarrow \exists x \in T: (\text{non } \mathcal{P}(x)),$$

$$\text{non } (\exists x \in T: \mathcal{P}(x)) \Leftrightarrow \forall x \in T: (\text{non } \mathcal{P}(x)),$$

- (Si partirà dalle [1.5] e [1.6], tenendo presente il significato delle [2.14] e [2.15].)

3. L'insieme-prodotto; relazioni; applicazioni

Introduciamo ora un'altra costruzione fondamentale della teoria degli insiemi: quella dell'insieme-prodotto di due insiemi.

Anzitutto, assumiamo come primitiva la nozione di *coppia* (ordinata); per suscitare in modo preciso nella nostra intuizione, pensiamo di avere a disposizione due caselle (la prima e la seconda casella) e di inserire in esse, rispettivamente, due elementi non necessariamente distinti x e y ; per indicare la coppia così ottenuta useremo la notazione (x, y) .

È evidente che bisogna guardarsi dal confondere la coppia (x, y) con l'insieme $\{x, y\}$ costituito dagli elementi x e y . Infatti si ha $\{x, y\} = \{y, x\}$ (per gli insiemi non ha rilevanza l'ordine con cui sono elencati gli elementi), mentre per le coppie si ha:

$$(x, y) = (x', y') \Leftrightarrow x = x' \text{ e } y = y'. \quad [3.1]$$

Siano ora dati due insiemi A e B . Diamo *insieme-prodotto di A per B* , e indicheremo con $A \times B$ l'insieme di tutte le coppie (x, y) , con $x \in A, y \in B$. In simboli:

$$A \times B \stackrel{\text{def}}{=} \{(x, y): x \in A, y \in B\}.$$

Notiamo che non abbiamo richiesto ai due insiemi A e B di essere diversi. Un ben noto esempio di prodotto è $\mathbb{R} \times \mathbb{R}$, dove $\mathbb{R}$ indica la retta reale. La geometria analitica ci insegna a rappresentare biunivocamente il piano mediante l'insieme di tutte le coppie ordinate di numeri reali (fig. 3.1), cioè mediante $\mathbb{R} \times \mathbb{R}$. Perciò l'insieme $\mathbb{R} \times \mathbb{R}$, che comunemente viene scritto $\mathbb{R}^2$, viene detto anche *piano* (numerico).

*3.1 Osservazione. Si può evitare di assumere la nozione di coppia come primitiva, pur di accettare una definizione un po' artificiosa. Si può infatti porre: $(x, y) \stackrel{\text{def}}{=} \{\{x\}, \{x, y\}\}$. Quello che

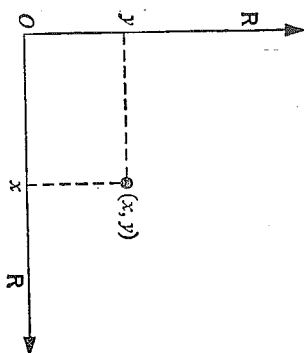


Figura 3.1

importa, è che, con questa definizione, viene soddisfatta la condizione espressa dalla [3.1] (lasciamo al lettore il compito di verificarlo). Una coppia (x, y) , con $x \in A$, $y \in B$ è allora un sottoinsieme di $\mathcal{P}(A \cup B)$ e, perciò, un elemento di $\mathcal{P}(\mathcal{P}(A \cup B))$. Quindi, il prodotto si può definire così:

$$A \times B = \{z: z \in \mathcal{P}(\mathcal{P}(A \cup B)) \text{ e } (\exists x \in A, \exists y \in B: z = \{\{x\}, \{x, y\}\})\}.$$

Si tratta di un'espressione un po' macchinosa, che però ha il vantaggio di presentare $A \times B$ come sottoinsieme di un insieme ben determinato (vedi la [2.4] e l'osservazione 2.3). Dunque, per questa via, non c'è bisogno né d'introdurre un'altra nozione primitiva, né d'introdurre un altro assioma che postuli l'esistenza dell'insieme-prodotto.

Dati, in un certo ordine, tre insiemi A, B, C , si chiama loro prodotto, e si indica con $A \times B \times C$ l'insieme di tutte le terne ordinate (x, y, z) con $x \in A, y \in B, z \in C$. La definizione si estende in modo ovvio al caso di un numero finito qualsiasi di spazi; ad esempio, si indica con $\mathbb{R}^n$ l'insieme di tutte le n -uple ordinate $(x_1, x_2, \dots, x_n)$ di numeri reali. Il prodotto di più insiemi può essere introdotto, peraltro, utilizzando solo il prodotto di due insiemi; è facile vedere infatti che $A \times B \times C$ si può introdurre indifferentemente come $(A \times B) \times C$, oppure come $A \times (B \times C)$.

Abbiamo già usato il termine "relazione" per indicare un predicato in due variabili. Sia ora $\mathcal{R}(x, y)$ una relazione che abbia senso nella teoria degli insiemi. Possiamo allora considerare l'insieme

$$\{(x, y): (x, y) \in A \times B \text{ e } \mathcal{R}(x, y)\} \quad [3.2]$$

cioè il sottoinsieme di $A \times B$ costituito da tutte le coppie per cui la relazione è verificata. Esso viene detto *grafico* della relazione. Dato un sottoinsieme G di $A \times B$, risulta individuata immediatamente una relazione di cui esso è grafico: $(x, y) \in G$.

Nella teoria degli insiemi, spesso il termine "relazione" viene usato nel significato di "grafico", perciò, in futuro, anche noi useremo lo stesso simbolo per indicare una relazione e il suo grafico: scriveremo perciò indifferentemente $\mathcal{R}(x, y)$, oppure $(x, y) \in \mathcal{R}$.

Esempi

a) In $N \times N = N^2$ si consideri la relazione $\mathcal{R}(m, n)$, che significa " m è divisore di n ". L'allievo cerchi di farsi un'idea del grafico di questa relazione con un disegno.

b) In $\mathbb{R}^2$ si consideri la relazione: $x^2 + y^2 \leq 1$. Il suo grafico è rappresentato dal cerchio con centro nell'origine degli assi e raggio 1.

In questo paragrafo e nei due successivi ci proponiamo di studiare particolari tipi di relazioni. Cominceremo con le *applicazioni*.

3.2 DEFINIZIONE Una relazione f definita in $A \times B$ si dice *applicazione* (o *funzione*) di A in B se per ogni $x \in A$ esiste uno e un solo $y \in B$ tale che $(x, y) \in f$.

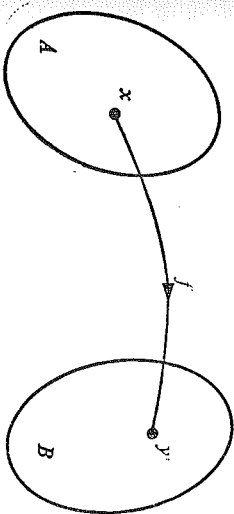


Figura 3.2

Possiamo interpretare intuitivamente la nozione introdotta pensando che f "porti" i punti di A in punti di B ; la coppia (x, y) fa corrispondere al punto $x \in A$ il punto $y \in B$ in cui esso viene portato (fig. 3.2). L'insieme A viene chiamato *dominio* di f , l'insieme B *codominio*. I termini "funzione" e "applicazione" sono equivalenti, tuttavia il termine "funzione" è più tradizionale e lo si impiega di preferenza quando dominio o codominio sono insiemi di numeri reali o complessi. L'espressione:

$$f: A \rightarrow B, \quad \text{scritta anche } A \xrightarrow{f} B,$$

mette in evidenza il dominio e il codominio di f .

Per ogni $x \in A$ l'unico elemento $y \in B$ tale che $(x, y) \in f$ si indica con $f(x)$ e si dice *valore assunto* dalla funzione f in x (fig. 3.3). Partendo dall'espressione di $f(x)$, la notazione completa

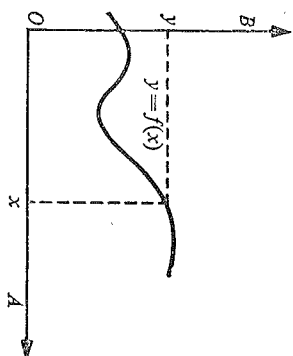


Figura 3.3

che rappresenta l'applicazione f è la seguente:

$$\{x \mapsto f(x): A \rightarrow B\}.$$

In questa espressione, la lettera x è "muta" e può essere sostituita con un'altra lettera qualsiasi. Spesso, però, questa notazione (che è evidentemente ingombrante) può venire più o meno abbreviata, se non vi sono equivoci. Spesso la funzione si indica semplicemente con un'espressione $f(x)$, dove appunto è sottinteso che la lettera x ha il ruolo di "*variable indépendante*". Ad esempio, si potrà parlare della funzione x^2+1 , in luogo di $\{x \mapsto x^2+1: \mathbb{R} \rightarrow \mathbb{R}\}$.

3.3 Osservazione. In certi casi si usa indicare il valore di un'applicazione scrivendo la variabile indipendente a mo' di indice: f_x , anziché $f(x)$. Ad esempio, se il dominio dell'applicazione è l'insieme N degli interi naturali, l'applicazione viene detta *successione*; i valori di una successione a (ma più frequentemente si parla dei *termini* della successione) sono indicati con la notazione: $a_0, a_1, \dots, a_n, \dots$. Le nozioni di funzione di variabile reale e di successione, che nei testi matematici più recenti sono comprese nella nozione generale di applicazione, venivano un tempo considerate come diverse: ciò spiega la diffinitività di notazione, che è rimasta.

Esempi banali ma importanti

- Dato un qualunque insieme A , l'applicazione $\{x \mapsto x: A \rightarrow A\}$ si dice *applicazione identica* di A e si indica con I_A .
- Dato l'insieme A e un suo sottoinsieme B , l'applicazione $\{x \mapsto x: B \rightarrow A\}$ si dice *applicazione di inclusione* di B in A .
- Siano A e B insiemi qualunque; l'applicazione $\{(x, y) \mapsto x: A \times B \rightarrow A\}$ si dice *proiezione canonica* su A , l'applicazione $\{(x, y) \mapsto y: A \times B \rightarrow B\}$ si dice *proiezione canonica* su B .

Data un'applicazione $f: A \rightarrow B$, e dato un sottoinsieme X di A , si dice *immagine* di X il sottoinsieme di B costituito dagli elementi che provengono da qualche elemento di X ; questo sottoinsieme viene indicato con $f(X)$. Pertanto:

$$f(X) \stackrel{\text{def}}{=} \{y: y \in B \text{ e } (\exists x \in X: f(x) = y)\}. \quad [3.3]$$

Si potrà anche scrivere, più brevemente:

$$f(X) = \{f(x): x \in X\}. \quad [3.4]$$

L'insieme $f(A)$ si può denominare senz'altro immagine dell'applicazione f .

Un'applicazione $f: A \rightarrow B$ si dice *costante* se la sua immagine consta di un unico elemento (in altre parole: $\forall x_1 \in A, \forall x_2 \in A: f(x_1) = f(x_2)$).

Diamo ora due importanti definizioni:

3.4 DEFINIZIONE Un'applicazione $f: A \rightarrow B$ si dice surgettiva se $f(A) = B$ (cioè se l'immagine di f coincide con B).

Esempi di applicazioni surgettive sono l'applicazione identica, in un qualunque insieme, le proiezioni canoniche del prodotto $A \times B \rightarrow A$, $A \times B \rightarrow B$ (vedi precedenti esempi a) e c); A e B si suppongono non vuoti). Se $f: A \rightarrow B$ è surgettiva, si usa dire che f è un'applicazione di A su B .

3.5 DEFINIZIONE Un'applicazione $f: A \rightarrow B$ si dice iniettiva se essa porta punti distinti in punti distinti.

In altre parole:

$$f \text{ è iniettiva} \Leftrightarrow (\forall x_1 \in A \quad \forall x_2 \in A: f(x_1) = f(x_2) \Rightarrow x_1 = x_2).$$

Ad esempio, l'applicazione identica in un insieme, l'applicazione di inclusione di un sottoinsieme in un insieme (esempi a) e b)) sono iniettive. La proiezione canonica $A \times B \rightarrow A$ non è iniettiva (a meno che...).

3.6 DEFINIZIONE Un'applicazione che sia iniettiva e surgettiva si dice "biiettiva".

Ad esempio, l'applicazione identica in un insieme qualsiasi è biiettiva. L'applicazione $\{x \mapsto ax + b: \mathbb{R} \rightarrow \mathbb{R}\}$ è biiettiva se $a \neq 0$; se $a = 0$, è un'applicazione costante, e non è né iniettiva né surgettiva.

3.7 Osservazione. In molte questioni non è necessario specificare il dominio e il codominio di un'applicazione, essendo sufficiente conoscere un'espressione che la definisce. Ma in certe questioni la precisazione del dominio e del codominio sono essenziali: come quando ci si chiede se un'applicazione è surgettiva, oppure se è iniettiva (ovviamente!). Ad esempio (indicando con $\mathbb{R}^+$ l'insieme dei numeri reali ≥ 0) l'applicazione: $\{x \mapsto x^2: \mathbb{R} \rightarrow \mathbb{R}\}$ non è né iniettiva né surgettiva; l'applicazione: $\{x \mapsto x^2: \mathbb{R}^+ \rightarrow \mathbb{R}^+\}$ è surgettiva, ma non iniettiva, l'applicazione: $\{x \mapsto x^2: \mathbb{R}^+ \rightarrow \mathbb{R}^+\}$ è iniettiva, ma non surgettiva, e, infine, l'applicazione $\{x \mapsto x^2: \mathbb{R}^+ \rightarrow \mathbb{R}^+\}$ è iniettiva e surgettiva. (Lo studioso compia la verifica aiutandosi anche con i grafici.)

3.8 Osservazione. Spesso è comodo rappresentare gli insiemi come immagini di opportune applicazioni, servendosi di espressioni del tipo [3.3] o [3.4]. Ad esempio, l'insieme degli interi naturali pari si potrà indicare con $\{2n: n \in \mathbb{N}\}$. Sia ora $\mathcal{F}$ una famiglia di insiemi, J un insieme, e supponiamo che esista un'applicazione $\{ \mapsto X_J: J \rightarrow \mathcal{F} \}$ surgettiva. Allora si possono impiegare le seguenti notazioni per indicare l'insieme-unione e l'insieme-intersezione di $\mathcal{F}$:

$$\bigcup_{X \in \mathcal{F}} X \text{ oppure } \bigcup_{J \in J} X_J \text{ oppure } \bigcap_{X \in \mathcal{F}} X \text{ oppure } \bigcap_{J \in J} X_J, \text{ (rispettivamente). [3.5]}$$

3.9 DEFINIZIONE Date due applicazioni $f: A \rightarrow B$ e $g: B \rightarrow C$ si dice applicazione composta di f e g l'applicazione $g \circ f$ così definita:

$$g \circ f \stackrel{\text{def}}{=} \{x \mapsto g(f(x)): A \rightarrow C\}.$$

Si noterà che il codominio di f è stato assunto coincidente con il dominio di g : solo in questo caso si può definire l'applicazione composta (si dice anche: f e g sono componibili). Nella scrittura $g \circ f$ (che rispecchia la scrittura $g(f(x))$) si scrive a destra l'applicazione che viene eseguita per prima (fig. 3.4).

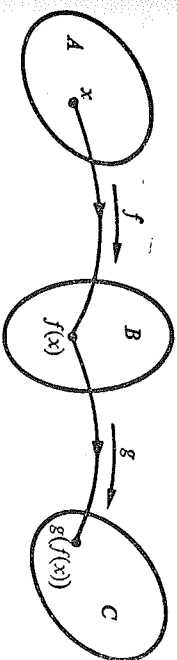


Figura 3.4

Ad esempio: l'applicazione $\{x \mapsto \sqrt{1+x^2}: \mathbb{R} \rightarrow \mathbb{R}\}$ si può considerare composta con l'applicazione $\{x \mapsto 1+x^2: \mathbb{R} \rightarrow \mathbb{R}^+\}$ e l'applicazione $\{y \mapsto \sqrt{y}: \mathbb{R}^+ \rightarrow \mathbb{R}\}$. Se è $f: A \rightarrow B$, si ha

$$f \circ I_A = f, \quad I_B \circ f = f$$

(ricordiamo che I_A e I_B sono le applicazioni identiche di A e di B). Notiamo poi che la composizione delle applicazioni gode della proprietà associativa; cioè: se f e g sono componibili e se g e

h sono componibili, si ha $h \circ (g \circ f) = (h \circ g) \circ f$. Basta infatti notare che si ha $(h \circ (g \circ f))(x) = h(g(f(x))) = ((h \circ g) \circ f)(x)$. Se f è un'applicazione di un insieme A in sé, si possono considerare le iterate di f , cioè le applicazioni $f \circ f, f \circ f \circ f$ ecc.; esse possono venire indicate con i simboli f^2, f^3 ecc. (Però, nel caso di funzioni a valori reali o complessi, o a valori in un gruppo, occorre stare attenti a non confondere queste funzioni con la funzione "quadrato di f ": $x \mapsto f(x)^2$; o "cubo di f " ecc.)

3.10 DEFINIZIONE Data un'applicazione $f: A \rightarrow B$ si chiama inversa della f un'applicazione $g: B \rightarrow A$ tale che:

$$g \circ f = I_A, \quad f \circ g = I_B. \quad [3.6]$$

Il lettore costaterà subito con esempi che non sempre un'applicazione è invertibile (cioè ha un'inversa). Il seguente teorema ci dice chiaramente quando questo può accadere.

3.11 TEOREMA Sia f un'applicazione $A \rightarrow B$. Allora f è invertibile $\Leftrightarrow f$ è biiettiva.

Dimostrazione. Cominciamo con l'implicazione $\Rightarrow$. Supponiamo dunque che f ammetta un'inversa g . Allora:

- a) f è surgettiva; infatti $\forall y \in B$ si ha (per la seconda delle [3.6]) $f(g(y)) = y$; dunque, esiste un elemento di A : $x = g(y)$ tale che $f(x) = y$.
- b) f è iniettiva; infatti sia $f(x_1) = f(x_2) = y$; allora, dalla prima delle [3.6] si ricava $x_1 = g(f(x_1)) = g(y)$, $x_2 = g(f(x_2)) = g(y)$ dunque $x_1 = x_2$.

Dimostriamo ora l'implicazione opposta $\Leftarrow$.

Supponiamo dunque che f sia biiettiva. Essendo f surgettiva, per ogni $y \in B$ esiste qualche x tale che $f(x) = y$; ma, essendo f iniettiva, ne esiste uno solo. Quindi è bene individuata un'applicazione $g: B \rightarrow A$ che fa corrispondere a y l'unico elemento $x \in A$ tale che $f(x) = y$. Dunque, per ogni $y \in B$, è $f(g(y)) = y$, e questa non è che la seconda delle [3.6]. D'altra parte, per ogni $x \in A$, è certamente $g(f(x)) = x$, e questa coincide con la prima delle [3.6]. ■

3.12 TEOREMA L'applicazione inversa di un'applicazione $f: A \rightarrow B$, se esiste è unica.

L'unicità dell'inversa risulta con evidenza dalla seconda parte della dimostrazione del teorema precedente. Riteniamo tuttavia istruttivo darne una dimostrazione autonoma di tipo algebrico. Supponiamo dunque che g e g' siano applicazioni $B \rightarrow A$ soddisfacenti alle relazioni [3.6]. Per la proprietà associativa della composizione, si ha:

$$(g' \circ f) \circ g = g' \circ (f \circ g).$$

D'altra parte, per ipotesi è $g' \circ f = I_A$, $f \circ g = I_B$; si ricava allora subito $g = g'$. ■

L'applicazione inversa di f si indica col simbolo f^{-1} . È evidente che $(f^{-1})^{-1} = f$: dunque, se f è invertibile, è tale anche la sua inversa. Dati due insiemi A, B , se esiste un'applicazione $A \rightarrow B$ invertibile, diciamo che A e B (senza più necessità di considerarli in ordine) si possono mettere in corrispondenza biunivoca. Studieremo più avanti questa importante relazione fra gli insiemi.

3.13 DEFINIZIONE Data un'applicazione $f: A \rightarrow B$ e dato un sottoinsieme C di A , si dice restrizione di f a C e si indica con $f|_C$, l'applicazione:

$$\{x \mapsto f(x) : C \rightarrow B\}.$$

In altre parole, indicando con j l'applicazione di inclusione:

$$C \rightarrow A, \text{ si ha } f|_C \stackrel{\text{def}}{=} f \circ j.$$

Assegnata un'applicazione $f: A \rightarrow B$, abbiamo già definito l'immagine $f(X)$ di un sottoinsieme X di A . L'applicazione $X \mapsto f(X)$ è un'applicazione $\mathfrak{F}(A) \rightarrow \mathfrak{F}(B)$ che, benché impropriamente, verrà indicata con il medesimo simbolo f . Ovviamente: $f(\emptyset) = \emptyset$.

Analagamente si potrà introdurre un'applicazione $f^{-1}: \mathfrak{F}(B) \rightarrow \mathfrak{F}(A)$ ponendo, per ogni $Y \subset B$:

$$f^{-1}(Y) \stackrel{\text{def}}{=} \{x : x \in A \text{ e } f(x) \in Y\}.$$

L'insieme $f^{-1}(Y)$ si dirà immagine inversa di Y . Osserviamo che questa applicazione, denotata (anch'essa impropriamente) con f^{-1} , esiste anche quando f non è invertibile. È evidente poi che, quando sia $Y \cap f(A) = \emptyset$, risulta $f^{-1}(Y) = \emptyset$.

3.14 TEOREMA Se f è un'applicazione $A \rightarrow B$, e X ed Y sono sottoinsiemi di A , si ha:

$$f(X \cup Y) = f(X) \cup f(Y),$$

$$f(X \cap Y) \subset f(X) \cap f(Y).$$

Lasciamo la semplice dimostrazione allo studioso; notiamo che, nella seconda relazione, il segno $\subset$ non può, in generale, essere sostituito con il segno $=$. Ad esempio, consideriamo la funzione $\{x \mapsto x^2: \mathbb{R} \rightarrow \mathbb{R}\}$, e sia $X = \{x: -1 \leq x \leq 0\}$, $Y = \{x: 0 \leq x \leq 1\}$. Si ha, evidentemente, $X \cap Y = \{0\}$; dunque $f(\{0\}) = \{0\}$, mentre è $f(X) \cap f(Y) = \{x: 0 \leq x \leq 1\}$. Dunque, in questo caso è $f(X \cap Y) \subsetneq f(X) \cap f(Y)$.

L'applicazione $f^{-1}: \mathcal{G}(B) \rightarrow \mathcal{G}(A)$ ha, come ora vediamo, un comportamento migliore di quello della $f: \mathcal{G}(A) \rightarrow \mathcal{G}(B)$, e si mostra molto utile in tante questioni.

3.15 TEOREMA Sia f un'applicazione $A \rightarrow B$, e siano X e Y sottoinsiemi di B . Valgono le seguenti relazioni:

$$f^{-1}(X \cup Y) = f^{-1}(X) \cup f^{-1}(Y),$$

$$f^{-1}(X \cap Y) = f^{-1}(X) \cap f^{-1}(Y),$$

$$f^{-1}(X') = (f^{-1}(X))',$$

dove X' indica il complementare di X in B e, analogamente, $(f^{-1}(X))'$ indica il complementare di $f^{-1}(X)$ in A .

Dimostrazione.

$$f^{-1}(X \cup Y) = \{x: f(x) \in X \cup Y\} = \{x: (f(x) \in X) \vee (f(x) \in Y)\} =$$

$$= \{x: f(x) \in X\} \cup \{x: f(x) \in Y\} = f^{-1}(X) \cup f^{-1}(Y).$$

Analogamente si dimostrano le altre relazioni. ■

È interessante notare che i teoremi 3.14 e 3.15 si estendono al caso di operazioni insiemistiche infinite. Ad esempio: se $\{X_j: j \in J\}$ è una famiglia di parti di B , si ha:

$$f\left(\bigcup_{j \in J} X_j\right) = \bigcup_{j \in J} f(X_j),$$

$$f^{-1}\left(\bigcup_{j \in J} X_j\right) = \bigcup_{j \in J} f^{-1}(X_j).$$

Esercizi

1. Disegnare sommarariamente il grafico delle seguenti relazioni definite in $\mathbb{R} \times \mathbb{R}$ e dire quali di esse rappresentano funzioni $\{x \mapsto y: \mathbb{R} \rightarrow \mathbb{R}\}$:

$$|x| + |y| \leq 1; \quad |x + y| \leq 1; \quad x + y + 2 = 0; \quad y \leq x^2; \quad (x^2 - y)^2 \leq 0.$$

2. Dare un esempio di un'applicazione $\mathbb{R} \rightarrow \mathbb{R}$ che sia: a) iniettiva, ma non surgettiva; b) surgettiva, ma non iniettiva.

3. Dimostrare che un'applicazione $f: A \rightarrow B$ è iniettiva se e soltanto se per ogni $X \subset A$ e per ogni $Y \subset B$ si ha

$$f(X \cap Y) = f(X) \cap f(Y).$$

(Ricordare il teorema 3.14.)

4. Si dimostri che per ogni $f: A \rightarrow B$ e per ogni $X \in \mathcal{G}(A)$ si ha:

$$X \subset f^{-1}(f(X)).$$

Inoltre l'uguaglianza $X = f^{-1}(f(X))$ vale per ogni $X \in \mathcal{G}(A)$ se e soltanto se f è iniettiva.¹

5. Si dimostri che per ogni $f: A \rightarrow B$ e per ogni $Y \in \mathcal{G}(B)$ si ha:

$$f(f^{-1}(Y)) \subset Y.$$

Inoltre l'uguaglianza $f(f^{-1}(Y)) = Y$ per ogni $Y \in \mathcal{G}(B)$ se e solo se f è surgettiva.¹

6. Si calcolino le funzioni composte $f \circ g$ e $g \circ f$ per ciascuna delle seguenti coppie di funzioni reali:

$$f(x) = 1 - 3x, \quad g(x) = x - 2,$$

$$f(x) = x^2 + 1, \quad g(x) = 1/(x^2 + 1),$$

$$f(x) = x + a, \quad g(x) = x - a,$$

$$f(x) = 3x + 2, \quad g(x) = 4x + 3.$$

*7. Si cerchi di caratterizzare le applicazioni f di un insieme A in sé tali che $f^2 = f$. (Si comincerà con l'individuare gli *elementi uniti*, cioè gli elementi $y \in A$ tali che $f(y) = y$.)

*8. Si consideri l'applicazione di $\mathbb{Z}$ (insieme degli interi relativi) in sé così definita: $x \mapsto ax + b$, (con a e b , ovviamente, interi). Si trovi in quali casi essa è biettiva.

¹ Gli esercizi 4 e 5 mettono in guardia dal ritenere che f ed f^{-1} , come applicazioni fra $\mathcal{G}(A)$ e $\mathcal{G}(B)$ siano l'una l'inversa dell'altra (a meno che $f: A \rightarrow B$ non sia biettiva).

*9. Si consideri l'applicazione di $\mathbb{Z} \times \mathbb{Z}$ in sé così definita:

$$(x, y) \mapsto (x', y'),$$

essendo:

$$x' = ax + by,$$

$$y' = cx + dy, \quad \text{con } a, b, c, d \text{ interi.}$$

In quali casi essa è iniettiva? in quali surgettiva?

10. Siano f e g due applicazioni di un insieme E in sé, che commutano (tal, cioè, che: $f \circ g = g \circ f$). Si dimostri che $f \circ g$ è invertibile quando e soltanto quando sia f che g lo sono.

4. Relazioni di equivalenza; l'assioma della scelta

In questo paragrafo e nel successivo studieremo alcune importanti relazioni binarie che si definiscono fra gli elementi di uno stesso insieme.

4.1 DEFINIZIONE Sia dato un insieme E ; diremo che una relazione $\mathcal{R}(x, y)$ definita in $E \times E$ è:

riflessiva se per ogni $x \in E$ è vera $\mathcal{R}(x, x)$,

simmetrica se: $\forall x \in E, \forall y \in E: \mathcal{R}(x, y) \Rightarrow \mathcal{R}(y, x)$,

transitiva se: $\forall x \in E, \forall y \in E, \forall z \in E: \mathcal{R}(x, y) \text{ e } \mathcal{R}(y, z) \Rightarrow \mathcal{R}(x, z)$.

Abbiamo già trovato queste proprietà nella introduzione della relazione di eguaglianza (vedi § 1).

4.2 DEFINIZIONE Una relazione che sia riflessiva, simmetrica, transitiva si dice *relazione di equivalenza*.

Esempi

a) La relazione $x = y$, nell'ambito di un insieme E , è una relazione di equivalenza (vedi § 1).

b) Nell'insieme E possiamo porre equivalenti due elementi qualsiasi; il grafico di questa relazione coincide con tutto l'insieme prodotto $E \times E$.

Le equivalenze introdotte in a) e b) sono, in un certo senso, casi estremi: nella prima ogni elemento risulta equivalente solo a sé stesso, nella seconda risulta equivalente a tutti gli altri.

c) Nella geometria euclidea, il parallelismo tra le rette di un piano è una relazione di equivalenza (a patto di considerare ogni retta parallela a sé stessa; la proprietà transitiva, in questo caso, è un'affermazione equivalente al famoso assioma delle parallele).

Spesso una relazione di equivalenza tra x e y si indica con $x \sim y$; ovviamente, in uno stesso insieme si possono definire più relazioni di equivalenza: allora occorre distinguerle con notazioni opportune.

Data una relazione di equivalenza in un insieme E , e preso un $x \in E$, indichiamo con $[x]$ l'insieme di tutti gli elementi di E equivalenti a x o, come si usa dire, la *classe di equivalenza* di x . In simboli si ha dunque:

$$[x] = \{y: (y \in E) \text{ e } (x \sim y)\}.$$

Vale la seguente importante proposizione:

4.3 TEOREMA Due classi di equivalenza, o sono disgiunte, o coincidono.

Dimostrazione. Siano $[x]$ e $[z]$ due classi di equivalenza, e supponiamo che esse abbiano un elemento w in comune: pertanto è $w \sim x$ e $w \sim z$. Pertanto, è $x \sim z$, per la proprietà transitiva. Sia ora $y \in [x]$, cioè $y \sim x$; per la proprietà transitiva è $y \sim z$, cioè $y \in [z]$; si conclude che è $[x] \subset [z]$. Analogamente si dimostra che è $[z] \subset [x]$. ■

Consideriamo ora l'insieme $\mathcal{F}$ di tutte le classi di equivalenza. $\mathcal{F}$ è un sottoinsieme di $\mathcal{P}(E)$ con le seguenti proprietà:

a) $\forall S \in \mathcal{F}$ è $S \neq \emptyset$ (infatti, la classe di equivalenza $[x]$ contiene almeno x);

b) $\forall S \in \mathcal{F}, \forall T \in \mathcal{F}: S \neq T \Rightarrow S \cap T = \emptyset$ (lo abbiamo visto sopra: due classi di equivalenza diverse sono necessariamente disgiunte);

c) $\bigcup_{S \in \mathcal{F}} S = E$ (infatti un qualunque elemento di E appartiene alla classe di equivalenza che individua).

Un famiglia $\mathcal{F}$ di sottoinsiemi di E aventi le proprietà a), b), c), si dice una *partizione* di E . Dunque, una relazione di equivalenza $\mathcal{R}$ individua una partizione $\mathcal{F}$: questa viene detta *insieme quo-*

ziente di E rispetto alla relazione $\mathcal{R}$ e viene indicata col simbolo $E/\mathcal{R}$.

Reciprocamente, data una partizione $\mathcal{P}$, è subito individuata in modo unico una relazione $\mathcal{R}$ tale che $\mathcal{P} = E/\mathcal{R}$. Questa sarà evidentemente la seguente:

$$x \sim y \Leftrightarrow x \text{ e } y \text{ appartengono a un medesimo elemento di } \mathcal{P}.$$

Data in E una relazione $\mathcal{R}$, risulta individuata l'applicazione $E \rightarrow E/\mathcal{R}$ che porta ogni $x \in E$ nella sua classe di equivalenza $[x]$; si tratta di un'applicazione surgettiva che viene detta *applicazione canonica sul quoziente*.

Prima di lasciare l'argomento, vediamo ancora due esempi di notevole importanza.

d) Consideriamo l'insieme $\mathbb{N}^2$ (costituito dalle coppie di interi naturali). Introduciamo in $\mathbb{N}^2$ la seguente relazione:

$$(m, n) \sim (m', n') \Leftrightarrow m + n' = m' + n. \quad [4.1]$$

Lasciamo al lettore la cura di verificare che si tratta di una relazione di equivalenza; è consigliabile farsi (nel piano, riferito ad assi cartesiani) una rappresentazione grafica di $\mathbb{N}^2$ e delle classi di equivalenza che risultano definite. La relazione di equivalenza [4.1] può essere impiegata per introdurre formalmente l'insieme $\mathbb{Z}$ degli interi relativi, partendo dagli interi naturali: si pone per definizione $\mathbb{Z} = \mathbb{N}^2/\sim$. (Il lettore non si spaventi: il trucco sta nel tenere presente che la coppia (m, n) fa la parte dell'intero relativo $m - n$.) L'addizione si definirà, ovviamente, così:

$$[(m, n)] + [(m_1, n_1)] = [(m + m_1, n + n_1)].$$

Ma, naturalmente, occorre dimostrare che il risultato (che è una classe di equivalenza) dipende solo dalle due classi di equivalenza addende, e non dalle coppie, che le rappresentano. Lasciamo al lettore il compito di dimostrarlo. Il lettore non avrà poi difficoltà (sapendo dove vuole arrivare) a definire la moltiplicazione di due elementi di $\mathbb{Z}$; occorre sempre tener presente che quando si deve operare su classi di equivalenza bisogna dimostrare che il risultato è indipendente dai rappresentanti che le individuano.

e) Sia $\mathbb{Z}$ l'insieme degli interi relativi, e sia $\mathbb{Z}^* = \mathbb{Z} - \{0\}$. Nell'insieme $\mathbb{Z} \times \mathbb{Z}^*$ (che è dunque insieme di tutte le coppie (p, q) di interi relativi, con $q \neq 0$), introduciamo la seguente relazione

$$(p, q) \sim (p', q') \Leftrightarrow pq' = p'q. \quad [4.2]$$

Ad esempio, si ha $(1, 3) \sim (2, 6) \sim (-1, -3) \sim (4, 12) \sim \dots$. Si verifica facilmente che si tratta di una relazione di equivalenza: la proprietà riflessiva e simmetrica sono evidenti; verifichiamo la transitiva. Supponendo $(p, q) \sim (p', q')$ e $(p', q') \sim (p'', q'')$ si ha: $pq' = p'q$ e $p'q'' = p''q'$. Moltiplicando membro a membro la prima per q'' e la seconda per q si ottiene: $pq'q'' = p'q''q$ e $p'q''q = p''q'q$, da cui, essendo $q' \neq 0$, si ricava $pq'' = p''q$, che dice appunto $(p, q) \sim (p'', q'')$.

È istruttivo rappresentare graficamente le classi di equivalenza.

Dopo l'esempio precedente il lettore avrà forse già capito dove si vuole arrivare: l'insieme $\mathbb{Z} \times \mathbb{Z}^*$ può essere considerato come insieme delle frazioni (si interpreti (p, q) come p/q); la relazione di equivalenza introdotta non è altro che l'ordinaria relazione di equivalenza tra frazioni. L'insieme quoziente è dunque l'insieme $\mathbb{Q}$ dei numeri razionali relativi; naturalmente, per completare il discorso, occorrerebbe definire le operazioni e le diseguglianze fra numeri razionali, sempre tenendo presente che, operando su classi di equivalenza, occorre dimostrare l'indipendenza del risultato dai rappresentanti delle classi.

Consideriamo ora un'applicazione f surgettiva di un insieme A in un insieme B (fig. 4.1). Per ogni $y \in B$ l'insieme $f^{-1}(\{y\})$ non

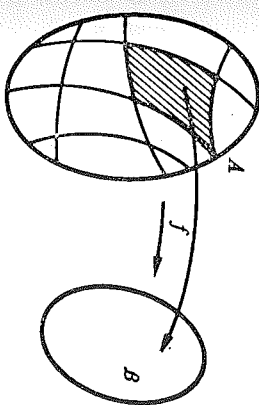


Figura 4.1

è vuoto: la famiglia $\mathcal{P} = \{f^{-1}(\{y\}) : y \in B\}$ è costituita da insiemi disgiunti, e ricopre A ; dunque è una partizione di A . Evidentemente, questa partizione può essere individuata dalla relazione di equivalenza:

$$x_1 \sim x_2 \Leftrightarrow f(x_1) = f(x_2).$$

Definiamo ora un'applicazione $\tilde{f}: \mathcal{P} \rightarrow B$ nel seguente modo: $\tilde{f}(\{x\}) = f(x)$. Si vede subito che questa definizione è legittima perché, se è $[x] = [u]$, allora, essendo $x \sim u$, si ha: $\tilde{f}([u]) = f(u) =$

$=f(x)$. L'applicazione f oltre a essere surgettiva, è iniettiva (infatti, classi diverse vanno a finire in punti diversi, per il modo con cui sono state definite le classi).

Insomma, si è ottenuta un'applicazione iniettiva contraendo in un unico punto tutti i punti aventi uno stesso corrispondente in B ...

Se indichiamo con $\mathcal{F}$ l'applicazione canonica di A sul quoziente $\mathcal{F}$, si ha: $f = \tilde{f} \circ \mathcal{F}$.

Possiamo ora pensare a un altro modo di ricavare dalla f un'applicazione iniettiva, che rimanga ancora surgettiva: quello di sostituire f con una sua restrizione opportuna. Come dovrà essere il dominio della nuova applicazione? Dovrà, evidentemente, contenere uno e un solo punto preso da ciascuno degli insiemi $f^{-1}(\{y\})$, che costituiscono $\mathcal{F}$. Che sia possibile compiere questa scelta è cosa abbastanza accettabile dalla nostra intuizione; tuttavia, nessuna delle costruzioni insiemistiche che abbiamo introdotto ci permette di affermare l'esistenza di un insieme come quello desiderato. Occorre introdurre appositamente un assioma.

4.4 ASSIOMA DELLA SCELTA *Data comunque una partizione $\{X: X \in \mathcal{F}\}$ di un insieme E esiste un sottoinsieme di E che ha uno e un solo elemento in comune con ciascuno degli $X \in \mathcal{F}$.*

4.5 Osservazione. L'assioma della scelta può essere evidentemente anche enunciato in questa forma: data una partizione $\mathcal{F}$ di un insieme E esiste una funzione $\varphi: \mathcal{F} \rightarrow E$ (detta appunto *funzione di scelta*) tale che, per ogni $X \in \mathcal{F}$, è $\varphi(X) \in X$.

4.6 Osservazione. Partendo dall'assioma della scelta non sarebbe difficile dimostrare questa proposizione più generale: data comunque una famiglia $\mathcal{F}$ di sottoinsiemi non vuoti (non necessariamente disgiunti) di un insieme E esiste una funzione $\varphi: \mathcal{F} \rightarrow E$ (funzione di scelta) tale che per ogni $X \in \mathcal{F}$ è $\varphi(X) \in X$.

Esercizi

1. Quante sono le equivalenze possibili in un insieme di 2, o 3, o 4 elementi?

2. Assegnato un intero r , si consideri in $\mathbb{Z}$ questa relazione tra x e y :
esiste un $k \in \mathbb{Z}$ tale che $x - y = kr$.

Si dimostri che questa è una relazione di equivalenza e si determinino le classi di equivalenza.

3. Assegnato un numero reale $T > 0$ si consideri sulla retta reale la seguente relazione tra x e y : esiste un $k \in \mathbb{Z}$ tale che $x - y = kT$.

Si dimostri che è una relazione di equivalenza. Quali sono le funzioni reali che assumono valore costante in ciascuna classe di equivalenza?

4. In un insieme E siano date due relazioni di equivalenza $\mathcal{R}(x, y)$ ed $\mathcal{R}'(x, y)$; si dimostri che la relazione $(\mathcal{R}(x, y) \text{ e } \mathcal{R}'(x, y))$ è ancora una relazione di equivalenza. Non così, in generale, per la relazione: $(\mathcal{R}(x, y) \text{ o } \mathcal{R}'(x, y))$.

*5. Si dimostri l'affermazione contenuta nell'osservazione 4.6. (Sarà opportuno considerare certi insiemi costituiti da coppie (x, X) , dove è $X \in \mathcal{F}$ ed è $x \in X$...)

5. Relazioni d'ordine; il principio d'induzione

5.1 DEFINIZIONE *Si dice preordinamento una relazione binaria assegnata in un insieme, che goda delle proprietà riflessiva e transitiva.*

Se $\mathcal{R}$ è un preordinamento, occorre tener presente che, non essendo più assicurata la proprietà simmetrica, la scrittura $\mathcal{R}(x, y)$ non è più, in generale, equivalente a $\mathcal{R}(y, x)$.

Esempi

a) Ogni relazione di equivalenza è un preordinamento.

b) La relazione $x \leq y$ nell'ambito degli interi (o razionali, o reali) è un preordinamento.

c) Sia E un insieme qualunque; la relazione $X \subset Y$ è un preordinamento in $\mathfrak{P}(E)$.

d) Nell'insieme dei numeri reali, la relazione $x^2 \leq y^2$ è un preordinamento.

5.2 DEFINIZIONE *Si dice che una relazione binaria $\mathcal{R}$ in un insieme E è antisimmetrica se:*

$$\forall x \in E, \forall y \in E: (\mathcal{R}(x, y) \text{ e } \mathcal{R}(y, x)) \Rightarrow x = y.$$

(In altre parole, il sussistere di entrambe le relazioni $\mathcal{R}(x, y)$ ed $\mathcal{R}(y, x)$ è possibile solo quando x e y coincidono.)

5.3 DEFINIZIONE Si dice ordinamento un preordinamento che goda anche della proprietà antisimmetrica. Pertanto, un ordinamento è una relazione che gode delle proprietà riflessiva, antisimmetrica e transitiva.

Fra gli esempi esposti sopra, si riconosce subito che le relazioni assegnate in b) e c) sono ordinamenti, mentre un'equivalenza (vedi a)) non è un ordinamento, a meno che non si riduca alla relazione di uguaglianza (infatti, se valgono le proprietà simmetrica e antisimmetrica, si deduce subito che: $\mathcal{R}(x, y) \Leftrightarrow x = y$). La relazione dell'esempio d) non è un ordinamento perché se è: $x^2 \leq y^2$ e $y^2 \leq x^2$, si ha: $x^2 = y^2$, e questa vale anche se $x = -y$.

5.4 DEFINIZIONE Dato un preordinamento (rappresentato dal simbolo $\leq$) in un insieme E , si pone:

$$x < y \Leftrightarrow_{a \neq b} (x \leq y) \text{ e } (y \nless x).$$

Potremo chiamare (pre)ordinato un insieme in cui sia stata assegnata una relazione di (pre)ordinamento.

5.5 DEFINIZIONE Un insieme preordinato viene detto insieme diretto (o filtrante) se, dati due suoi elementi qualsiasi x, y esiste almeno un elemento z tale che $x \leq z, y \leq z$.

Si riconosce che negli esempi b), c), d), si definiscono insiemi diretti.

5.6 DEFINIZIONE Un elemento m di un insieme ordinato E si dice massimo se per ogni $x \in E$ si ha $x \leq m$. Si dice massimale se non vi è alcun elemento di E che lo supera, cioè se:

$$\forall y \in E: y \geq m \Rightarrow y = m.$$

Evidentemente, ogni elemento che sia massimo è anche massimale, ma non viceversa. Un insieme ordinato può avere più elementi massimali, mentre non può avere più di un massimo. (Lo si capisce facilmente osservando la figura 5.1; in diagrammi come questo gli elementi sono indicati con punti e si ha $a \leq b$ se a e b sono uniti con un tratto orientato da a verso b ; oltre alle relazioni di ordine che vi sono rappresentate, si devono sottintendere tutte quelle che si ottengono applicando la proprietà transitiva.)

In un insieme ordinato non sempre accade che due elementi x, y siano fra loro confrontabili, cioè sussista una delle due relazioni $x \leq y, y \leq x$. Così, nell'esempio c), se l'insieme E ha più

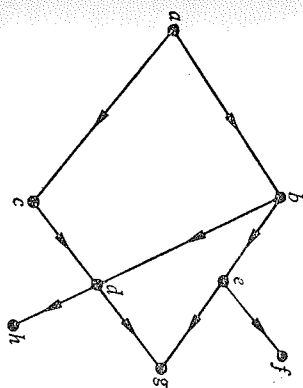


Figura 5.1

di un elemento, si possono trovare coppie X, Y di sottoinsiemi di E tali che $X \not\leq Y$ e $Y \not\leq X$. Nell'esempio b) invece tutti gli elementi sono confrontabili. Risulta allora opportuna la seguente definizione:

5.7 DEFINIZIONE Un ordinamento $\leq$ definito in un insieme E si dice totale (o lineare) se per ogni coppia (x, y) di elementi di E si ha $x \leq y$ oppure $y \leq x$.

Un sottoinsieme T di un insieme ordinato E risulta anch'esso ordinato, con l'ordinamento che gli viene subordinato da E . Se accade che esso risulti totalmente ordinato, si dice che è una catena.

Sia K una catena in E ; supponiamo che esista un elemento $\bar{x} \notin K$ che sia confrontabile con tutti gli elementi di K , allora la catena K può venire ampliata in modo ovvio con l'aggiunta di $\bar{x}$: vogliamo dire con questo che $K \cup \{\bar{x}\}$ è ancora una catena. Vale il seguente teorema:

5.8 TEOREMA In ogni insieme ordinato esistono catene non amplifiabili.

Questa proposizione si dimostra facendo uso dell'assioma (della scelta) 4.4, attraverso la proposizione più generale 4.6, ma la dimostrazione non risulta molto facile. Più facile è dimostrare l'assioma della scelta assumendo come assioma la proposizione 5.8. Può convenire, tutto sommato, prendere come assioma la 5.8,

tanto più che in molte questioni è proprio la 5.8 che viene applicata, direttamente oppure tramite il seguente teorema (detto *lemma di Zorn*).

5.9 TEOREMA Sia E un insieme ordinato non vuoto in cui ogni catena K ammetta una limitazione superiore (cioè esista in E un elemento s tale che $\forall x \in K: x \leq s$). Allora in E esiste almeno un elemento massimale.

Dimostrazione. Il teorema 5.8 ci assicura che in E esiste una catena K non ampliabile; essa è certamente non vuota. Sia s una limitazione superiore per K . Deve essere $s \in K$ perché K non è ampliabile. Dico che s è elemento massimale di E . Infatti sia, per un certo $y \in E$, $y > s$: allora anche y è una limitazione superiore per K , perciò (sempre per il fatto che K non è ampliabile) è $y \in K$, perciò $y \leq s$, da cui finalmente $y = s$. Dunque non c'è alcun y in E tale che $y > s$. Questo ci dice che s è elemento massimale. ■

È importante ora esaminare i rapporti che ci sono tra i preordinamenti e gli ordinamenti; ci proponiamo di costatare che c'è un modo assai naturale per ricavare un ordinamento da un preordinamento.

Sia assegnato in un insieme E un preordinamento; non valendo, in generale, la proprietà antisimmetrica, non si potrà affermare che per ogni coppia (x, y) di elementi di E : $(x \leq y)$ e $(y \leq x) \Rightarrow x = y$. Tuttavia, la relazione $\mathcal{O}(x, y)$: " $(x \leq y)$ e $(y \leq x)$ " è una *relazione di equivalenza* (lasciamo la verifica allo studioso). Allora possiamo considerare l'insieme-quotiente: $E/\mathcal{O}$. La struttura di preordinamento assegnato in E induce una struttura di ordinamento in $E/\mathcal{O}$ ponendo $[x] \leq [y]$, se è $x \leq y$. Naturalmente (come abbiamo già visto in diversi esempi nel paragrafo precedente) occorre verificare che questa relazione che si vuole stabilire per le classi di equivalenza non dipende dal particolare elemento con cui esse sono designate. Sia dunque $x \leq y$ e sia $[x] = [x']$ e $[y] = [y']$, cioè $x \sim x'$, $y \sim y'$. Allora si ha $x' \leq x \leq y \leq y'$, perciò è $x' \leq y'$. Si verifica poi che la relazione introdotta in $E/\mathcal{O}$ è un preordinamento: le proprietà riflessiva e transitiva sono di verifica immediata. Ma vale anche la proprietà antisimmetrica. Sia infatti $[x] \leq [y]$ e $[y] \leq [x]$; ciò significa, per la nostra definizione, che è $x \leq y$ e $y \leq x$; ma allora è $x \sim y$ e perciò $[x] = [y]$. Possiamo concludere con il seguente enunciato.

5.10 TEOREMA Sia assegnato in un insieme E un preordinamento $\leq$, e sia $\mathcal{O}(x, y)$ la relazione: $(x \leq y)$ e $(y \leq x)$. Allora $\mathcal{O}$ è una relazione di equivalenza e il preordinamento di E induce un ordinamento nell'insieme quoziente $E/\mathcal{O}$.

Più avanti, questo teorema ci sarà particolarmente utile.

5.11 Osservazione. Sia E un insieme preordinato (in particolare, ordinato); se indichiamo con $\mathcal{O}(x, y)$ la relazione $x \leq y$ (" y è maggiore o uguale a x "), la relazione $\mathcal{O}(x, y) = \mathcal{O}(y, x)$ (" y è minore o uguale a x ") ha le stesse proprietà formali di $\mathcal{O}$. Si possono allora considerare per $\mathcal{O}$ tutte le nozioni che abbiamo introdotto per $\mathcal{O}$: si può parlare (con senso ovvio) di *minimo*, di *elemento minimale*, di *limitazioni inferiori* ecc.

Svolgiamo infine alcune considerazioni sull'insieme N degli interi naturali, dal punto di vista dell'ordine. L'insieme N ha le seguenti proprietà:

- 1) N è un insieme totalmente ordinato.
- 2) Ogni sottoinsieme non vuoto di N ha un elemento minimo.

Un insieme ordinato che ha le proprietà 1) e 2) si dice *bene ordinato*. L'insieme N stesso ha un elemento minimo (lo zero: 0). Dato un $n \in N$, l'insieme $\{x: x \in N, x > n\}$ (cioè l'insieme degli elementi maggiori di n) ha un elemento minimo: questo viene detto *successivo* di n . Lo indicheremo con n^+ .

Allora è chiaro che N ha ancora le seguenti due proprietà:

- 3) Ogni elemento di N ha un successivo.
- 4) (Tenendo conto della 2) basterebbe affermare che, per ogni n l'insieme degli elementi maggiori di n non è vuoto, cioè che N non ha massimo.)

- 4) Ogni elemento di N diverso da 0 è successivo di un elemento di N .

Da queste proprietà si ricava facilmente la seguente importantissima proposizione:

5.12 PRINCIPIO D'INDUZIONE Sia T un sottoinsieme di N avente le seguenti proprietà:

$$\alpha) 0 \in T,$$

$$\beta) \forall n: n \in T \Rightarrow n^+ \in T,$$

$$\text{allora } T = N.$$

Dimostrazione. Indichiamo con T' l'insieme complementare di T ; la tesi equivale, evidentemente, all'affermazione che $T' = \emptyset$. Supponiamo, per assurdo, $T' \neq \emptyset$. Allora (proprietà 2)) T' ha un elemento minimo m , il quale (per la α) è $\neq 0$. Allora, per la 4), esiste n tale che $n^+ = m$. Si ha: $n \in T$, perché m è l'elemento minimo di T' , mentre è: $n^+ \in T'$. Ma questo contrasta con la β . ■

Si deve a G. Peano una presentazione assiomatica degli interi naturali basata sulle nozioni primitive di "numero", "zero" e "successivo"; in essa il principio di induzione figura come assioma.

Il principio d'induzione è uno degli strumenti fondamentali della matematica. Accenniamo a due modi in cui esso viene soprattutto applicato.

a) Definizioni per induzione (o ricorrenza). Supponiamo che siano assegnati un insieme X e un'applicazione $G: N \times X \rightarrow X$; sia poi a un elemento di X . Allora esiste una e una sola applicazione $f: N \rightarrow X$ che soddisfa a queste condizioni

$$f(0) = a,$$

$$f(n^+) = G(n, f(n)).$$

Intuitivamente, l'enunciato è chiaro: la prima relazione assegna il valore di f nello 0, la seconda determina in modo unico il valore di f per l'intero successivo di n , noto che sia il suo valore per n . La dimostrazione di questa proposizione può essere formalmente ottenuta a partire dal principio di induzione; la omettiamo per brevità.

Nella costruzione dell'aritmetica secondo Peano le operazioni sugli interi sono tutte definite per ricorrenza. Ad esempio, la somma $m+n$ di due interi naturali può essere così definita come funzione di n , una volta fissato m :

$$m+0 = m,$$

$$m+n^+ = (m+n)^+.$$

È chiaro che, introdotta l'addizione, si può scrivere $n+1$ in luogo di n^+ (qui 1 è, per definizione, 0^+).

Come altro esempio, consideriamo la funzione $n!$ (fattoriale di n); essa è così definita:

$$0! = 1,$$

$$(n+1)! = (n+1)n!.$$

Con i simboli del nostro enunciato, G è la funzione $(n, x) \rightarrow (n+1)x$. Il lettore, forse, avrebbe compreso più rapidamente se si fosse scritto (per $n > 0$): $n! = 1 \cdot 2 \cdot 3 \cdot \dots \cdot n$. Il fatto è, però, che questa scrittura, con i puntini, ci suggerisce alla buona proprio quel procedimento di ricorrenza che noi abbiamo esposto formalmente. In fondo, anche in discorsi matematici molto semplici viene applicato, più o meno consapevolmente, il principio di induzione. Questo accade, ad esempio, quando parliamo di un'operazione da eseguirsi n "volte": così si scrive: $a^n = a \cdot a \cdot \dots \cdot a$ (n volte). Questa enunciazione sta in luogo della seguente definizione per ricorrenza: $a^0 = 1$, $a^{n^+} = a^n \cdot a$.

b) Dimostrazioni per induzione. Si tratta di dimostrare una proposizione in cui interviene come parametro un intero naturale n . Indichiamo con T l'insieme degli n per cui la proposizione è vera. L'enunciato si suppone vero per $n=0$ e si sa che, se è vero per l'intero n , è vero per $n+1$. Allora l'insieme T gode delle proprietà α) e β) di 5.12: la conclusione è che è $T = N$, cioè l'enunciato è vero per ogni n .

Ad esempio, mostriamo che:

per ogni numero reale $\delta > -1$ e per ogni intero $n \geq 0$ vale la disuguaglianza:

$$(1+\delta)^n \geq 1+n\delta.$$

Possiamo dimostrarla mediante il principio di induzione, in questo modo:

$$(1+\delta)^0 = 1 \geq 1,$$

supponiamo ora che sia $(1+\delta)^n \geq 1+n\delta$; allora è:

$$\begin{aligned} (1+\delta)^{n+1} &= (1+\delta)^n(1+\delta) \geq (1+n\delta)(1+\delta) = \\ &= 1+n\delta+\delta+n\delta^2 \geq 1+(n+1)\delta. \end{aligned}$$

La nostra disuguaglianza vale per $n=0$ e, se vale per un intero n , vale anche per l'intero $n+1$. Dunque risulta dimostrata per ogni $n \in N$.

Esercizi

* 1. In un insieme ordinato, un elemento può essere, al tempo stesso, massimale e minimale? ed essere al tempo stesso massimo e minimo?

* 2. Dato un insieme *finito* E , si descriva una catena massimale di $\mathcal{P}(E)$ (l'ordinamento è sempre quello per inclusione).

* 3. Si dimostri che se in un insieme ordinato finito vi è un unico elemento massimale, esso è massimo.

* 4. Si consideri l'insieme $\mathbb{N}^+$ degli interi positivi, ordinato per divisibilità (cioè $m \leq n$ significa: m è divisore di n). Riconoscere che si tratta di un ordinamento filtrante e che, dati due elementi a e b in $\mathbb{N}^+$ l'insieme dei "seguenti comuni" (cioè l'insieme $\{x: x \geq a, x \geq b\}$) ha un elemento minimo. Come viene chiamato abitualmente questo? Analogamente, si dimostri che l'insieme dei "precedenti" comuni di a e b ha un elemento massimo. Come si possono individuare in $\mathbb{N}^+$, in termini di questo ordinamento, i numeri primi?

5. Si consideri ancora $\mathbb{N}^+$ ordinato per divisibilità. Fra i seguenti sottoinsiemi si ricerchino quelli che sono catene e, fra queste, le catene massimali (cioè non ampliabili).

Per ogni catena non massimale, si cerchi una catena massimale che la contiene:

$$\{2n: n = 1, 2, 3, \dots\}, \quad \{n^2: n = 1, 2, 3, \dots\},$$

$$\{n!: n = 1, 2, 3, \dots\}, \quad \{p^n: n = 1, 2, 3, \dots\} \quad (p \text{ numero primo}),$$

$$\{2^n: n = 1, 2, 3, \dots\}, \quad \{10^n: n = 1, 2, 3, \dots\}.$$

* 6. Il principio d'induzione si applica spesso in questa forma: sia T un sottoinsieme di $\mathbb{N}$ con le seguenti proprietà: $0 \in T$; inoltre: per ogni $n \in \mathbb{N}$, se tutti gli interi $m \leq n$ appartengono a T , anche n^+ appartiene a T . Allora è $T = \mathbb{N}$. Dimostrare questa proposizione, o facendo uso delle proprietà di $\mathbb{N}$, o per mezzo della proposizione 5.12.

* 7. Dimostrare per induzione le seguenti relazioni:

$$1 + 2 + 3 + \dots + n = \frac{n(n+1)}{2},$$

$$1^2 + 2^2 + 3^2 + \dots + n^2 = \frac{n(n+1)(2n+1)}{6},$$

$$1^3 + 2^3 + 3^3 + \dots + n^3 = \left(\frac{n(n+1)}{2}\right)^2,$$

$$\sum_{k=0}^n a^k = \frac{1-a^{n+1}}{1-a} \quad (\text{supponendo } a \neq 1).$$

* 8. Dimostrare la disuguaglianza:

$$(1-a)^n < \frac{1}{1+na} \quad (\text{essendo } 0 < a < 1, n \geq 1).$$

* 9. Sia f la funzione reale così definita $f(x) = ax + b$. Si trovi l'espressione di $f^n = f \circ f \circ \dots \circ f$ (n volte).

* 10. Si consideri la successione $n \mapsto a_n$ (di Fibonacci) così definita: $a_0 = 0, a_1 = 1; a_{n+1} = a_n + a_{n-1}$. Si dimostri che la successione si può così rappresentare: $a_n = Ax^n + By^n$, essendo A, B, x, y numeri reali da determinare.

6. I numeri cardinali; gli insiemi finiti

La nozione di *corrispondenza biunivoca* (introdotta nel § 3) sta alla base della nozione di numero: contare significa stabilire una corrispondenza biunivoca tra un insieme di oggetti e un insieme "campione" (ad esempio: l'insieme delle dita di una mano). Risulta molto naturale la seguente definizione:

6.1 DEFINIZIONE *Dati due insiemi, si dice che essi hanno lo stesso numero cardinale se essi possono essere posti in corrispondenza biunivoca.*

Per indicare che due insiemi X e Y hanno lo stesso numero cardinale, scriveremo $c(X) = c(Y)$. Questa relazione gode delle proprietà riflessiva, simmetrica, transitiva.

* 6.2 Osservazione. Se tutti gli insiemi che si considerano in una certa teoria matematica sono contenuti in uno stesso insieme E (che fa da "universo"), il simbolo $c(X)$ assume il significato di classe di equivalenza in $\mathcal{P}(E)$ (cioè indica la famiglia di tutti i sottoinsiemi di E che si possono porre in corrispondenza biunivoca con X). Ma, se non si fa una tale convenzione, dal momento che, come sappiamo, non esiste "l'insieme di tutti gli insiemi", l'espressione $c(X) = c(Y)$ deve essere intesa, globalmente, come un predicato binario (relazione, ma non in senso insiemistico) tra X e Y .

Cominciamo ad applicare la definizione 6.1 ai casi più semplici. Preso un $n \in \mathbb{N}$, consideriamo l'insieme degli interi che precedono n , cioè l'insieme:

$$I_n = \{0, 1, 2, \dots, n-1\}.$$

Gli insiemi I_n si prestano bene come insiemi-campione.