



SCHEDE INTEGRATIVE

Questa parte è dedicata ad approfondimenti sui vari capitoli. Si presentano esercizi, problemi integrativi, costruzioni, curiosità, anche con l'ausilio di un computer. I calcoli con i numeri naturali si possono eseguire con ogni tipo di calcolatrice e di software; in particolare, si eseguono con la cosiddetta “precisione infinita”, ossia in forma esatta, senza approssimazioni e con la sola limitazione della memoria dell'hardware usato. Gli ambiti di applicazione del software in aritmetica sono essenzialmente l'esecuzione di calcoli complessi, quali somme di molti addendi, prodotti, potenze, quozienti, naturalmente, ma anche e soprattutto la *divisibilità*, con il cercare se un numero sia *primo* o no, la fattorizzazione, il calcolo di MCD ed mcm. Inoltre, il calcolo combinatorio, che porta spesso a trovare la somma di successioni intere, la risoluzione di *equazioni diofantee*, il calcolo di funzioni aritmetiche. Non ci si limiterà tuttavia all'aritmetica, ma si mostreranno altre costruzioni, con o senza software, nonché risoluzione di equazioni, ecc. In questa sede si farà uso di “Geogebra”, di “Excel”, degli ambienti di calcolo (C.A.S.) e di programmazione delle calcolatrici TI-92 Plus. Quest'ultima è collegabile al computer attraverso il software TI-Connect.

Scheda A1	Le successioni
Scheda A2	Numeri primi
Scheda A3	Monoidi e gruppi
Scheda A4	Calcolo combinatorio e funzioni aritmetiche
Scheda A5	Matrici multiformi
Scheda A6	Algebra delle relazioni
Scheda A7	Insiemi ordinati
Scheda A8	La divisibilità

SCHEDA A1. – Le successioni

Siano X un insieme non vuoto, $n_0 \in \mathbf{N}$ ed $M = \{n \in \mathbf{N} \mid n \geq n_0\}$. Chiamiamo *successione in X* ogni funzione $f: M \rightarrow X$. Se $M = \mathbf{N}$, posto $a_n = f(n)$ per ogni $n \in M$, la successione si denota anche con $(a_n)_{n \in \mathbf{N}}$. In generale sarà $n_0 = 0$ oppure $= 1$.

Una successione $f: M \rightarrow X$ si può definire *esplicitamente*, assegnando ad ogni $n \in M$ il valore $f(n)$, ma anche *ricorsivamente*, assegnando esplicitamente le *condizioni iniziali* $f(0), \dots, f(h-1)$ e facendo dipendere $f(n+h)$ da $f(n), f(n+1), \dots, f(n+h-1)$.

In generale sarà $X = \mathbf{R}$, campo dei numeri reali, di cui $\mathbf{N}$ è pensato come sottoinsieme, ed $n_0 = 0$ oppure $n_0 = 1$. Nel costruire le tabelle si è fatto uso di Excel, il cui riempimento automatico è perfetto per le definizioni ricorsive.

Esempi A1.1. a) Il fattoriale $n!$: $\begin{cases} 0! = 1 \\ (n+1)! = n!(n+1) \end{cases} \Rightarrow \begin{array}{c|cccccc} n & 0 & 1 & 2 & 3 & 4 & 5 & 6 \\ \hline n! & 1 & 1 & 2 & 6 & 24 & 120 & 720 \end{array} \dots$

b) I numeri di Fibonacci. Ecco un esempio di successione che dipende da due valori iniziali:

$$\begin{cases} a_0 = 1 \\ a_1 = 1 \\ a_{n+2} = a_n + a_{n+1} \end{cases} \Rightarrow \begin{array}{c|cccccccc} n & 0 & 1 & 2 & 3 & 4 & 5 & 6 & 7 \\ \hline a_n & 1 & 1 & 2 & 3 & 5 & 8 & 13 & 21 \end{array} \dots$$

c) - Le progressioni geometriche. Fissati $a, q \in \mathbf{R} \setminus \{0\}$, si pone: $\begin{cases} a_0 = a \\ a_{n+1} = a_n \cdot q \end{cases}$. Si ottiene la

successione: $\begin{array}{c|cccccc} n & 0 & 1 & 2 & 3 & \dots & n & \dots \\ \hline a_n & a & a \cdot q & a \cdot q^2 & a \cdot q^3 & \dots & a \cdot q^n & \dots \end{array}$

d) Le serie. Un argomento importante è la somma degli elementi di una successione. Data

$$(a_n)_{n \in \mathbf{N}}, \text{ si pone: } \begin{cases} \sum_{k=0}^0 a_k = a_0 \\ \sum_{k=0}^{n+1} a_k = \sum_{k=0}^n a_k + a_{n+1} \end{cases} \quad . \text{ Questa successione si chiama } \textit{serie} \text{ di } (a_n)_{n \in \mathbf{N}}.$$

Il problema analitico legato alle successioni $(a_n)_{n \in \mathbf{N}}$ è la loro *convergenza*, ossia la determinazione del $\lim_{n \rightarrow \infty} a_n$. Per le serie tale limite si denota anche con

$\sum_{k=0}^{\infty} a_k$ ed è detto somma della serie. Questo problema è risolto in generale nei corsi

di Analisi Matematica. Il problema algebrico è invece il seguente: data la successione $(a_n)_{n \in \mathbf{N}}$ mediante una definizione ricorsiva, trovarne un'espressione esplicita in termini di operazioni aritmetiche, potenze, fattoriali ecc. In particolare,

data un'espressione esplicita di $(a_n)_{n \in \mathbf{N}}$, trovarne una esplicita per $\left(\sum_{k=0}^n a_k \right)_{n \in \mathbf{N}}$.

Esempi A1.2. a) Posto $f(n) = n$, si ha $\sum_{k=0}^n k = \frac{n \cdot (n+1)}{2}$, e la dimostrazione si fa per induzione su n .

b) La serie geometrica. L'espressione analitica per le progressioni geometriche è: $a_n = a \cdot q^n$. Il numero q si dice *ragione* e si suppone $\neq 0$. Se $q = 1$, si ha $a_n = a$ per ogni n e quindi la sua somma è $\left((n+1) \cdot a \right)_{n \in \mathbf{N}}$. Se $q \neq 1$, per le proprietà dei prodotti notevoli si ottiene:

$$\sum_{k=0}^n a \cdot q^k = a \cdot \sum_{k=0}^n q^k = \frac{a \cdot (q^n - 1)}{q - 1}.$$

Pertanto, se $0 < |q| < 1$ si ottiene facilmente la somma della serie geometrica:

$$\sum_{k=0}^{\infty} a \cdot q^k = \lim_{n \rightarrow \infty} \sum_{k=0}^n a \cdot q^k = \lim_{n \rightarrow \infty} \frac{a \cdot (q^n - 1)}{q - 1} = \frac{a}{1 - q}$$

Per esempio, consideriamo il numero decimale periodico

$$7, \overline{245} = 7,245454545... = 7 + \frac{2}{10} + \frac{45}{10^3} + \frac{45}{10^5} + ... = \frac{72}{10} + \frac{45}{10^3} \cdot \left(1 + \frac{1}{10^2} + \frac{1}{10^4} + ... \right).$$

Si tratta della somma del numero $7,2 = \frac{72}{10}$ con la serie geometrica di punto iniziale $\frac{45}{1.000}$ e

ragione $\frac{1}{100}$. La somma è allora:

$$\frac{72}{10} + \frac{45/1000}{1-1/100} = \frac{72}{10} + \frac{45}{990} = \frac{72 \cdot (100-1) + 45}{990} = \frac{7245-72}{990} = \frac{797}{110}$$

(Si è voluta riprodurre nei passaggi finali la nota regola per trovare la frazione generatrice di un numero periodico).

c) I numeri di Fibonacci. La successione di Fibonacci è un caso particolare di questo tipo: si fissino

$$p, q, a_0, a_1 \in \mathbb{R}, q \neq 0; \text{ si ha: } \begin{cases} f(0) = a_0 \\ f(1) = a_1 \\ f(n+2) = p \cdot f(n) - q \cdot f(n+1) \end{cases}.$$

Si consideri l'equazione caratteristica $t^2 - p \cdot t + q = 0$. Dette a e b le radici, eventualmente complesse, esse sono $\neq 0$ perché il loro prodotto è q .

Se sono distinte (cioè se $p^2 - 4q \neq 0$) si ha $a_n = f(n) = A \cdot a^n + B \cdot b^n$, per opportuni A e B da

$$\text{determinare con le condizioni iniziali: } \begin{cases} a_0 = A + B \\ a_1 = A \cdot a + B \cdot b \end{cases}.$$

$$\text{Se invece } b = a, \text{ si ha } a_n = f(n) = (A + n \cdot B) \cdot a^n, \text{ con } A \text{ e } B \text{ tali che } \begin{cases} a_0 = A \\ a_1 = (A + B) \cdot a \end{cases}.$$

Queste formule si dimostrano per induzione rispetto ad n .

Nel caso dei numeri di Fibonacci si ottiene:

$$a = \frac{1+\sqrt{5}}{2} \text{ (il numero aureo)}, b = \frac{1-\sqrt{5}}{2}, f(n) = \frac{a^{n+1} - b^{n+1}}{\sqrt{5}}.$$

Con Excel si verifica facilmente la coincidenza con la serie di Fibonacci, almeno per valori piccoli di n .

Le successioni più semplici sono i polinomi: $f(n) = \sum_{k=0}^n a_k \cdot n^k$, che sono le

restrizioni ad $\mathbf{N}$ dei polinomi a coefficienti reali. Vediamo alcuni esempi particolari.

Esempi A1.3. a). - *Le progressioni aritmetiche.* Siano $a, r \in \mathbb{R} \setminus \{0\}$, si pone: $\begin{cases} a_0 = a \\ a_{n+1} = a_n + r \end{cases}$. Si

$$\text{ottiene la successione: } \begin{array}{c|cccccc} n & 0 & 1 & 2 & 3 & \dots & n & \dots \\ \hline a_n & a & a+r & a+2r & a+3r & \dots & a+n \cdot r & \dots \end{array}$$

La sua espressione esplicita è dunque $a_n = a + n \cdot r$, quindi è un polinomio di primo grado nella variabile naturale n .

b). - I coefficienti binomiali. Si pone, per ogni $n, k \in \mathbb{N}$:

$$C_{n,k} = \frac{n!}{k!(n-k)!} = \frac{1}{k!} \cdot n \cdot (n-1) \cdots (n-k+1).$$

Il significato combinatorio di $C_{n,k}$ è il numero di sottoinsiemi con k elementi di un insieme con n elementi. Da questa formula segue che, fissato k , $C_{n,k}$ è una successione polinomiale di grado k , nulla per ogni $n < k$, ed i cui valori sono tutti numeri naturali. Si usa scrivere $\binom{n}{k}$ in luogo di $C_{n,k}$.

Introduciamo ora il concetto di *derivata (finita)* di una successione: data $(a_n)_{n \in \mathbb{N}}$, per ogni $n \geq 0$ poniamo: $a'_n = a_{n+1} - a_n$. Sia poi $(S_n)_{n \in \mathbb{N}}$ la serie di $(a_n)_{n \in \mathbb{N}}$. Si hanno le seguenti proprietà:

- PROPOSIZIONE A1.4.** a) La derivata finita di una costante è la costante nulla.
 b) La derivata finita di un polinomio di grado $n > 0$ è un polinomio di grado $n-1$.
 c) La derivata finita di $C_{n,k}$ è $C_{n,k-1}$.
 d) $S'_n = a_{n+1}$ per ogni n .

Da quanto precede, segue che eseguendo k volte la derivata finita di un polinomio di grado k si ottiene una successione costante.

Vediamo un esempio con $a_n = n^2$.

n	0	1	2	3	4	5	6	7	8	9	10	11	12	...
a_n	0	1	4	9	16	25	36	49	64	81	100	121	144	...
a'_n		1	3	5	7	9	11	13	15	17	19	21	23	...
a''_n			2	2	2	2	2	2	2	2	2	2	2	...

Denotiamo ora con d_i , $i = 0, 1, \dots, k$ i valori iniziali (per $n = 0$) delle successive derivate finite di un polinomio di grado k . Si osservi che, detti d'_i gli analoghi numeri per la successione (a'_n) , si ha $d'_i = d_{i+1}$ per ogni i . Ciò posto, si ha il seguente risultato:

PROPOSIZIONE A1.5. Ogni successione polinomiale f di grado $k > 0$ si può esprimere univocamente nella forma: $f(n) = \sum_{i=0}^k d_i \binom{n}{i} = \sum_{i=0}^k d_i C_{n,i}$, ossia come combinazione lineare di coefficienti binomiali.

Dimostrazione. Procediamo per induzione rispetto al grado k di f .

Se $k = 1$ allora $f(n) = a + b \cdot n$, ed essendo $d_0 = a$, $d_1 = b$, $C_{n,1} = n$, si ha $f(n) = d_0 + d_1 \cdot C_{n,1}$ per ogni n , come si voleva. Sia ora vero per i polinomi di grado k e dimostriamo che è vero per i polinomi di grado $k+1$. Sia dunque f di grado $k+1$. Procediamo per induzione rispetto ad n : se $n = 0$ si ha:

$f(0) = d_0$, per definizione, ma si ha anche $\sum_{i=0}^{k+1} d_i C_{0,i} = d_0$ perché $C_{0,0} = 1$ e per ogni $i > 0$ si ha

$C_{0,i} = 0$. Sia allora vero che $f(n) = \sum_{i=0}^{k+1} d_i C_{n,i}$ e proviamo che la formula vale anche per $n+1$. Per

definizione di derivata si ha: $f(n+1) = f(n) + f'(n)$, per cui, per l'ipotesi induttiva su n e per quella su k , che è il grado di f' , si ha:

$$\begin{aligned} f(n+1) &= f(n) + f'(n) = \sum_{i=0}^{k+1} d_i C_{n,i} + \sum_{i=0}^k d'_i C_{n,i} = \sum_{i=0}^{k+1} d_i C_{n,i} + \sum_{i=0}^k d_{i+1} C_{n,i} = \\ &= \sum_{i=0}^{k+1} d_i C_{n,i} + \sum_{i=1}^{k+1} d_i C_{n,i-1} = d_0 + \sum_{i=1}^{k+1} d_i (C_{n,i} + C_{n,i-1}) = \\ &= d_0 + C_{n+1,0} + \sum_{i=1}^{k+1} d_i C_{n+1,i} = \sum_{i=0}^{k+1} d_i C_{n+1,i} \end{aligned}$$

come si voleva.

Il vantaggio di questa espressione deriva dalla facilità di ricavare la serie di un coefficiente binomiale. Infatti:

PROPOSIZIONE A1.6. a) $\sum_{i=0}^n \binom{i}{k} = \binom{n+1}{k+1}.$

b) Se f è un polinomio, espresso nella forma: $f(n) = \sum_{i=0}^k d_i \binom{n}{i}$, la sua serie è data

da $\sum_{j=0}^n f(j) = \sum_{i=0}^k d_i \binom{n+1}{i+1}$

Dimostrazione. a) Per $n = 0$, se $k = 0$ si ha $\binom{0}{0} = \binom{1}{1} = 1$; se $k > 0$ si ha $0 = 0$, quindi l'affermazione è vera per $n = 0$ e per ogni k . Sia ora $n \geq 0$ un numero per il quale l'affermazione è vera per ogni k . Ricordiamo che per ogni $m, j > 0$ si ha: $\binom{m-1}{j-1} + \binom{m-1}{j} = \binom{m}{j}$. Allora, per $n+1$ si ha:

$$\sum_{i=0}^{n+1} \binom{i}{k} = \sum_{i=0}^n \binom{i}{k} + \binom{n+1}{k} = \binom{n+1}{k+1} + \binom{n+1}{k} = \binom{n+2}{k+1} = \binom{(n+1)+1}{k+1}$$

b) Segue da a).

Esempi A1.7. a) Sia $f(n) = n^2$ la successione dei quadrati. Si ha: $d_0 = 0$, $d_1 = 1$, $d_2 = 2$, come già calcolato. Pertanto si ha: $n^2 = 1 \cdot \binom{n}{1} + 2 \cdot \binom{n}{2}$, quindi:

$$\sum_{j=0}^n j^2 = 1 \cdot \binom{n+1}{2} + 2 \cdot \binom{n+1}{3} = \frac{(n+1) \cdot n}{2} + 2 \cdot \frac{(n+1) \cdot n \cdot (n-1)}{6} = \frac{2n^3 + 3n^2 + n}{6}$$

b). Sia $a_n = a + n \cdot r$ una progressione aritmetica. Essendo $d_0 = a_0$, $d_1 = r$, $C_{n,1} = n$, si ha $a_n = a \cdot \binom{n}{0} + r \cdot \binom{n}{1}$, quindi la somma S_n dei termini da a_0 ad a_n è la seguente:

$$S_n = a \cdot \binom{n+1}{1} + r \cdot \binom{n+1}{2} = a \cdot (n+1) + r \cdot \frac{(n+1) \cdot n}{2} = \frac{2a+r}{2} \cdot n + \frac{r}{2} \cdot n^2$$

c). Sia $f(n) = 3n^3 - 5n + 1$ e determiniamo S_n . Innanzitutto, calcoliamo le differenze (con Excel):

n	0	1	2	3	4	5	6	7	8	9	10	11	...
f(n)	1	-1	15	67	173	351	619	995	1497	2143	2951	3939	...
f'(n)		-2	16	52	106	178	268	376	502	646	808	988	..
f''(n)			18	36	54	72	90	108	126	144	162	180	...
f'''(n)				18	18	18	18	18	18	18	18	18	..

Pertanto, $f(n) = 1 \cdot \binom{n}{0} - 2 \cdot \binom{n}{1} + 18 \cdot \binom{n}{2} + 18 \cdot \binom{n}{3}$

Ne segue: $S(n) = 1 \cdot \binom{n+1}{1} - 2 \cdot \binom{n+1}{2} + 18 \cdot \binom{n+1}{3} + 18 \cdot \binom{n+1}{4} = \frac{3n^4 + 6n^3 - 7n^2 - 6n + 4}{4}$

d) Si consideri la somma dei primi $n+1$ numeri dispari; a che cosa è uguale?

Costruiamo dapprima una tabella, anche con l'ausilio di Excel, con i primi valori di n :

n	0	1	2	3	4	5	6	7	8	...
$2n+1$	1	3	5	7	9	11	13	15	17	...
$\sum_{k=0}^n 2k+1$	1	4	9	16	25	36	49	64	81	...

Sembra quindi che sia $\sum_{k=0}^n (2k+1) = (n+1)^2$. Per dimostrarlo, procediamo per induzione rispetto

ad n . Per $n = 0$ è vero: entrambi i membri valgono 1. Per ogni $n \geq 0$ per il quale l'uguaglianza è vera, (ipotesi induttiva), dimostriamo che è vera per il successivo $n+1$. Si ha:

$$\begin{aligned} \sum_{k=0}^{n+1} (2k+1) &= \sum_{k=0}^n (2k+1) + (2(n+1)+1) = (n+1)^2 + (2n+3) = n^2 + 2n + 1 + 2n + 3 = \\ &= n^2 + 4n + 4 = (n+2)^2 = ((n+1)+1)^2, \end{aligned}$$

quindi è vero anche per $n+1$. Allora l'uguaglianza è vera per ogni $n \in \mathbf{N}$.

Un altro modo: si ha $a_n = \binom{n}{0} + 2 \cdot \binom{n}{1} \Rightarrow \sum_{j=0}^n a_j = \binom{n+1}{1} + 2 \cdot \binom{n+1}{2} = (n+1)^2$

NOTA. Con il C.A.S. della calcolatrice TI-92 (e a maggior ragione con Derive o Mathematica o simili), si calcola direttamente: $\sum (2k+1, k, 0, n) \rightarrow n^2 + 2n + 1$

SCHEDA A2. – Numeri primi.

Come pro-memoria ricordiamo che un numero naturale $p > 1$ si dice *primo* se ha per divisori solo 1 e se stesso. La proprietà che lo caratterizza è il fatto che ogni volta che divide un prodotto, divide almeno uno dei fattori. Inoltre, ogni numero > 1 e non primo è scomponibile nel prodotto di numeri primi, e a parte l'ordine dei fattori, la scomposizione è unica. Tuttavia, lo scoprire se un numero sia primo o no e trovarne la fattorizzazione è un problema complesso, non appena il numero dato sia grande.

I C.A.S. sono in grado di stabilire se un numero non troppo grande sia primo o no e di trovarne la fattorizzazione attraverso procedimenti assai sofisticati. Tuttavia, i numeri primi sfuggono ai tentativi di esprimerli attraverso successioni esplicite, per la loro almeno apparente estrema casualità. Naturalmente, a parte il 2, tutti i primi sono dispari, quindi spesso i test che ne fanno uso sono estesi ai dispari > 2 .

A2.1 Tra i numeri della forma $2^n - 1$ quali sono primi?

Un metodo per scomporre numeri espressi come valore di espressioni aritmetiche è quello di scomporre la successione con le tecniche dei prodotti notevoli. Se n non è primo si può scrivere nella forma $n = p \cdot q$, con p primo, quindi, ricordando la nota formula algebrica

$$x^q - 1 = (x - 1) \cdot \sum_{k=0}^{q-1} x^k \text{ e posto } x = 2^p, \text{ si ricava la fattorizzazione:}$$

$$2^n - 1 = \left(2^p\right)^q - 1 = \left(2^p - 1\right) \cdot \sum_{k=0}^{q-1} \left(2^p\right)^k$$

Allora l'esponente deve essere un primo p .

Ma $2^p - 1$ sarà sempre primo? Se $p = 2$ sì.

Qui serve un software che sia in grado di verificare se un dato numero naturale dispari $x > 1$ sia primo o no. Nel C.A.S. della TI-92 c'è la funzione "isPrime(x)", che risponde "true" o "false". Cerchiamo allora di costruire una tabella contenente i numeri interi dispari x da 3 in poi e la verifica se x sia primo e se lo sia $2^x - 1$.

Definiamo $y1(x) = \text{isPrime}(x)$ e $y2(x) = \text{isPrime}(2^x - 1)$. Possiamo allora ottenere in modo automatico una tabella delle due funzioni, con $\text{TblStart} = 3$ e $\Delta \text{Tbl} = 2$.

Si vede subito che per $x = 11$, che è primo, si ha:

$$2^{11} - 1 = 23 \cdot 89$$

x	y1	y2
3.	true	true
5.	true	true
7.	true	true
9.	false	false
11.	true	false
13.	true	true
15.	false	false
17.	true	true

x=3.
MAIN RAD AUTO FUNC

NOTA. I numeri della forma $2^p - 1$ che sono primi sono detti *primi di Mersenne*. Non si sa se ce ne siano infiniti o no, ma se ne conoscono con migliaia di cifre.

A2.2. Per ogni numero naturale $n \geq 1$ si calcolino i valori della funzione:

$$A(n) = \text{numero dei primi } p \leq n,$$

e si confrontino i due numeri $A(n)$ e $n/\ln(n)$.

Incominciamo con una tabella calcolata con Excel:

n	1	2	3	4	5	6	7	8	9	10
A(n)	0	1	2	2	3	3	4	4	4	4
ln(n)	0,000	0,693	1,099	1,386	1,609	1,792	1,946	2,079	2,197	2,303
$\frac{A(n) \cdot \ln(n)}{n}$	0,000	0,347	0,732	0,693	0,966	0,896	1,112	1,040	0,977	0,921

Per trovare altri valori della successione $A(n)$ ci serve una lista di primi abbastanza estesa e una paziente tabulazione, oppure la possibilità di scrivere un programma in un qualche linguaggio di programmazione. Per questo possiamo usare una calcolatrice TI-92 Plus, che ci fornisca per ogni $n > 0$ il numero $A(n)$. Il software TI-Connect consentirà poi di trasferire la lista dei risultati sul computer e poi in un foglio Excel.

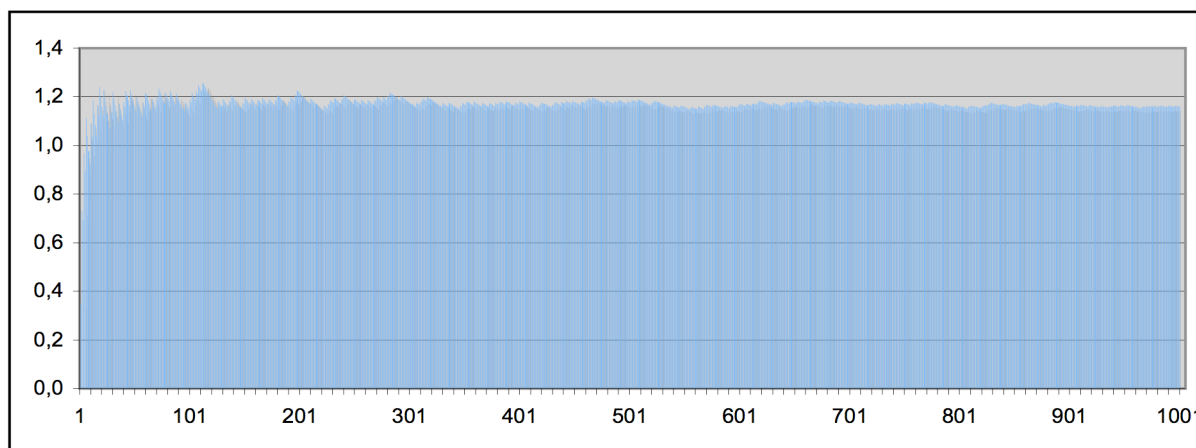
Un semplice algoritmo ricorsivo per il calcolo della successione $A(n)$ è il seguente:

- $A(1) = 0$
- Per ogni $n > 1$, se n è primo allora $A(n) = A(n-1) + 1$; altrimenti, $A(n) = A(n-1)$

Dopo ciò, si divide (in forma approssimata) $A(n)$ per il quoziente $n/\ln(n)$ e si ottiene una successione di numeri decimali, che per un'intuizione di Gauss, dimostrata solo alla fine dell'800, per n tendente all'infinito ha come limite 1.

Calcolati i valori della successione $A(n)$, $1 \leq n \leq 1001$, con un programmino scritto per la calcolatrice TI-92 Plus, poi riportati in una colonna Excel, sono moltiplicati per $\ln(n)$ e divisi per n . Il risultato è mostrato nel grafico ad istogramma riportato nella figura seguente. Si osservi come la convergenza ad 1 sia assai lenta. Per $n = 1001$ si ha infatti $A(1001) = 168$, quindi

$$\frac{A(1001) \cdot \ln(1001)}{1001} \approx 1,1595.$$



A2.3. Sia dato il polinomio $f(n) = n^2 - n + 41$. Se ne calcolino alcuni valori; che cos'hanno di particolare?

Con l'ausilio di Excel, calcoliamo e riportiamo in una tabella i valori di $f(n)$, $1 \leq n \leq 50$ e con l'ausilio della calcolatrice TI-92 Plus e del software TI-Connect riportiamo i valori true/false della funzione $\text{isPrime}(f(n))$. La sorpresa è che per $n \leq 40$ si trovano solo numeri primi. Per $n = 41$ si ha $f(41) = 1.681 = 41^2$.

n	1	2	3	4	5	6	7	8	9	10	11	12
f(n)	41	43	47	53	61	71	83	97	113	131	151	173
isprime(n)	true	true	true	true	true	true	true	true	true	true	true	true
n	13	14	15	16	17	18	19	20	21	22	23	24
f(n)	197	223	251	281	313	347	383	421	461	503	547	593
isprime(n)	true	true	true	true	true	true	true	true	true	true	true	true
n	25	26	27	28	29	30	31	32	33	34	35	36
f(n)	641	691	743	797	853	911	971	1033	1097	1163	1231	1301
isprime(n)	true	true	true	true	true	true	true	true	true	true	true	true
n	37	38	39	40	41	42	43	44	45	46	47	48
f(n)	1373	1447	1523	1601	1681	1763	1847	1933	2021	2111	2203	2297
isprime(n)	true	true	true	true	false	false	true	true	false	true	true	true

A2.4. In quanti modi si può scomporre 100 nella somma di due primi?

Naturalmente si può supporre $100 = p+q$, con $p \leq q$, quindi con $p \leq 47$.

La lista dei primi $p \leq 47$ e la verifica se $100-p$ sia primo o no si può eseguire manualmente.

p	3	5	7	11	13	17	19	23	29	31	37	41	43	47
q	97	95	93	89	87	83	81	77	71	69	63	59	57	53
isPrime(q)	true	false	false	true	false	true	false	false	true	false	false	true	false	true

Dunque, il 100 si scrive in 6 modi diversi come somma di due primi.

NOTA. La congettura di Goldbach, tuttora aperta, afferma che ogni numero pari sia la somma di due primi.

SCHEDA A3. – Monoidi e gruppi.

Ai numeri naturali sono associati vari monoidi, con le operazioni di addizione, moltiplicazione, MCD e mcm. Per costruire altri esempi di monoidi e di gruppi ci serviremo di stuzzicadenti, normali o colorati.

ESEMPIO A3.a) Che cosa si può fare? Fra le tante attività, la più elementare è quella di mettere gli stuzzicadenti in fila uno dietro l'altro e contarli:

$\emptyset$	I	II	III	IIII	IIIII	IIIIII	IIIIIII	IIIIIIII	IIIIIIIIII	...
0	1	2	3	4	5	6	7	8	9	...

Il simbolo del vuoto rappresenta qui una “fila vuota” di stuzzicadenti. Osserviamo che i numeri in basso esprimono la lunghezza della fila di stuzzicadenti sovrastante. Questo primo passo ricorda un poco le tacche che i nostri antenati cavernicoli incidevano su ossa o bastoni per contare. Ed ora, date due file di stuzzicadenti, possiamo ottenerne una nuova collocando la seconda fila dietro la prima. Chiamiamo *concatenazione* questo procedimento. Per rappresentarlo sulla carta ci serve un simbolo, per esempio la “e commerciale” &:

$$\text{IIIII} \& \text{IIIIIII} = \text{IIIIIIIIII}$$

La lunghezza della fila ottenuta è naturalmente la somma delle lunghezze delle due file di partenza. Tutto qui? Sì, perché senza altre regole quasi solo questo si può fare...

Osserviamo solo che se una fila viene concatenata con la fila vuota non cambia.

Nota. L'operazione di concatenazione nell'insieme delle file di stuzzicadenti è automaticamente commutativa, associativa ed ha come elemento neutro la fila vuota. Inoltre, la lunghezza della somma è la somma delle lunghezze. Si tratta infatti di un modo primitivo di rappresentare l'insieme dei numeri naturali e l'operazione di addizione, ossia il monoide $(\mathbb{N}, +, 0)$. Questo monoide è quindi *libero* da regole: la manipolazione degli stuzzicadenti sopra descritta non ha infatti altre regole oltre al modo di costruire la concatenata di due file. Tutto il resto, ossia le proprietà della concatenazione, sono automatiche.

Come trasformare tutto ciò in un gioco un po' più interessante? Proviamo ad inventare regole per complicare il gioco di partenza: al semplice concatenare file di stuzzicadenti, potremmo per esempio aggiungere la regola seguente:

A3.b) Fissato un numero $n > 0$, “mangiare” ogni fila di n stuzzicadenti, ossia sostituirla con la fila vuota. In simboli, $\underbrace{||| \dots |||}_n = \emptyset$

Esempio. Fissiamo $n = 5$. Allora avremo il 4 come lunghezza massima, e le file a disposizione saranno solo cinque: $\emptyset \quad | \quad || \quad ||| \quad ||||$

Che cosa succede quando le concateniamo? Ogni volta che abbiamo 5 stuzzicadenti in una fila, li “mangiamo”, come nella dama, ossia li togliamo dalla fila. In definitiva, una fila di cinque stuzzicadenti equivale alla fila vuota.

Avremo così per esempio:

$$||| \& |||| = \underbrace{||| ||||}_5 = ||.$$

Possiamo anche riassumere in una tabella i 25 risultati possibili. Notiamo che in ogni riga e colonna i risultati sono tutti diversi:

&	$\emptyset$				
$\emptyset$	$\emptyset$				
					$\emptyset$
				$\emptyset$	
			$\emptyset$		
		$\emptyset$			

Nota. Quel che si ottiene, per ogni “tetto” prefissato $n > 0$, è un gruppo con n elementi, detto *gruppo ciclico d'ordine n* . Questo è un altro modo di rappresentare il gruppo delle rotazioni del piano di ampiezza $\frac{2\pi k}{n}$, $0 \leq k \leq n-1$ intorno ad un centro O fissato.

A3.c) Fissati $m, n > 0$, sostituire ad ogni fila di n stuzzicadenti una fila di m stuzzicadenti. In simboli, $\underbrace{||| \dots |||}_n \rightarrow \underbrace{||| \dots |||}_m$ (regola monodirezionale).

Se $m = n$ non succede niente. Se $n < m$, che cosa accade?

Esempio. Prendiamo $n = 4$ ed $m = 5$. La fila vuota e quelle con uno, due o tre stuzzicadenti non cambiano, ma a partire da quella con quattro, si può aumentarne la lunghezza all'infinito, collocando 5 stuzzicadenti al posto di 4 in tutti i modi possibili. Gli informatici scriverebbero “overflow” oppure “out of memory”, ma noi possiamo dominare questa situazione col simbolo ∞ . In definitiva, abbiamo solo cinque file:

$\emptyset$ I II III ∞

Possiamo riassumere in una tabella le 25 possibili concatenazioni. Notiamo che la fila infinita "assorbe" le altre.

&	$\emptyset$	I	II	III	∞
$\emptyset$	$\emptyset$	I	II	III	∞
I	I	II	III	∞	∞
II	II	III	∞	∞	∞
III	III	∞	∞	∞	∞
∞	∞	∞	∞	∞	∞

Nota. Pare che il nostro occhio sappia distinguere, senza contarli, solo fino a 3 oggetti, e dopo ci sia solo una moltitudine indistinguibile. Somiglia un po' a questo nostro gioco. Il monoide ottenuto è commutativo, ha $\emptyset$ come elemento neutro e ∞ come elemento assorbente.

Resta il caso di $n > m$. Qui abbiamo un limite massimo alle lunghezze delle parole, uguale ad $n-1$, perché ogni volta che si arriva ad n , il numero scende ad $m < n$. In questo caso di ottengono monoidi commutativi non regolari, ma privi di elemento assorbente.

Esistono giochi *banali*? Che cosa può significare questa banalità? Potremmo dire che un gioco è banale se fa scomparire tutti gli stuzzicadenti, lasciando solo la fila vuota. In un certo senso, potremmo dire che in tal caso le regole stabilite sono contraddittorie. Per esempio, le regole $III = \emptyset$, $II = \emptyset$ implicano $I = \emptyset$.

Ci procuriamo ora due confezioni di stuzzicadenti e coloriamo in modo diverso il contenuto delle due scatole. Nelle figure sottostanti si useranno i colori nero e rosso, che nella stampa in bianco e nero diventeranno nero e grigio.

A3.d) Il gioco di base, libero da regole, è come quello del caso precedente, ma proprio l'assenza di regole crea subito qualche complicazione, dato che non è lecito per esempio scambiare stuzzicadenti di colori diversi:

$\emptyset$	I	I	II	II	II	II	III	III	III	III	III	III	III	III	...
0	1	2					3								

Le file distinte di data lunghezza $n \geq 0$ sono 2^n . La concatenazione di due file di stuzzicadenti colorati è tale che la fila risultante ha come lunghezza la somma delle

lunghezze delle due file date. La fila vuota non produce variazioni nella sua concatenazione con altre file. Lo scambio dei due termini nella concatenazione in generale produce risultati diversi:

$$||| \& || = |||| \quad || \& ||| = ||||$$

Nota. L'operazione di concatenazione nell'insieme delle file di stuzzicadenti bicolori è automaticamente associativa ed ha come elemento neutro la fila vuota. Inoltre, la lunghezza della somma è la somma delle lunghezze, e vale la legge di cancellazione. Si ottiene il *monoide libero* $(\mathcal{F}_2, \&, \emptyset)$ su due generatori.

Ora introduciamo man mano delle regole.

A3.e). Commutatività: $|| = ||$

Possiamo scambiare di posto in una fila ogni stuzzicadente con ogni altro. In questo caso, per ogni $n \geq 0$ ci sono solo $n+1$ file distinte di lunghezza n . Aggiungiamo la convenzione di collocare sempre prima gli stuzzicadenti neri e poi i rossi. Allora abbiamo un modo “canonico” di mettere in fila gli stuzzicadenti:

$\emptyset$										...
0	1	2	3	4	5	6	7	8	9	

Abbiamo in particolare un modo “canonico” di scrivere il risultato di una concatenazione:

$$||| \& |||| = ||||| = |||||$$

Nota. In questo caso si ottiene un monoide commutativo regolare. Si tratta di un altro modo di rappresentare la *somma diretta* del monoide additivo $(\mathbf{N}, +, 0)$ con se stesso o il monoide moltiplicativo dei monomi monici in due indeterminate.

A3.f) Commutatività e opposti: $|| = ||, || = \emptyset$

Essa implica solo file monocromatiche, una per ciascuna lunghezza possibile. Potremmo disporre le file di stuzzicadenti come segue:

...						$\emptyset$						...
...	-5	-4	-3	-2	-1	0	1	2	3	4	5	...

La concatenazione di due file dello stesso colore ne fa sommare le lunghezze, mentre in caso di colore diverso le lunghezze si sottraggono e resta una fila del colore prevalente. Due file della stessa lunghezza, ma di colore diverso, si distruggono a vicenda lasciando la fila vuota. E' evidente la somiglianza, o meglio l'*isomorfismo*, con il gruppo $(\mathbb{Z}, +)$ dei numeri interi relativi: basta sostituire per esempio ad ogni fila nera il numero positivo che ne esprime la lunghezza, mentre ad ogni fila rossa si sostituisce l'opposto della sua lunghezza.

Imponiamo ora che la massima lunghezza di una fila rossa o di una fila nera siano limitate. Per esempio, stabiliamo di "mangiare" tre stuzzicadenti neri o due rossi consecutivi. Allora abbiamo la seguente situazione:

$\emptyset$	I	I	II	II	II	III	III	III	III	...
0	1	2		3						

Sono possibili quindi file del tipo $IIIIIIIIII$, senza mai due rossi o tre neri consecutivi. Si ottiene un insieme infinito di file che, rispetto alla concatenazione, forma un gruppo non commutativo numerabile. Per esempio,

$$\left(IIIIIIIIIII \right)^{-1} = IIIIIIIIIII.$$

Aggiungiamo regole per potere in qualche modo scambiare gli stuzzicadenti neri con i rossi. Vediamo alcuni esempi interessanti.

A3.g). Commutatività e limitatezza: $II = II$, $II = \emptyset$, $III = \emptyset$.

Non è difficile dedurre che si hanno a disposizione solo sei file di stuzzicadenti:

$\emptyset$ I I II II III

Ed ecco la tavola dei risultati delle 36 concatenazioni: si osservi che ogni fila ha l'*inversa*: $I \& II = III = \emptyset$; $II \& III = IIII = \emptyset$, ecc.

Si ha così un gruppo.

La tavola è stata manipolata con un programma di disegno, per poter mostrare i colori.

&	$\emptyset$	I	II	I	II	III
$\emptyset$	$\emptyset$	I	II	I	II	III
I	I	II	$\emptyset$	II	III	I
II	II	$\emptyset$	I	III	I	II
I	I	II	III	$\emptyset$	I	II
II	II	III	I	I	II	$\emptyset$
III	III	I	II	II	$\emptyset$	I

Nota. L'operazione di concatenazione nell'insieme delle file di stuzzicadenti bicolori è come sempre associativa ed ha come elemento neutro la fila vuota. In questo caso è anche commutativa, ed ogni elemento ha l'inverso, perché la fila vuota compare in ogni riga e colonna. Si ottiene così un gruppo abeliano d'ordine 6, *somma diretta* dei due gruppi ciclici d'ordine 3 (gli stuzzicadenti neri) e 2 (i rossi), quindi *isomorfo* al gruppo ciclico d'ordine 6.

La regola B.3 si generalizza variando le lunghezze massime delle file nere o rosse. Per ogni m, n interi positivi, si ottiene un insieme di $m \cdot n$ file. In definitiva, si ottiene sempre un gruppo abeliano d'ordine $m \cdot n$, *somma diretta* dei due gruppi ciclici d'ordine m ed n .

4.2.h) Inversione e limitatezza: $II = III, II = \emptyset, III = \emptyset$.

In questo caso è un po' più difficile dedurre che si hanno a disposizione solo sei file di stuzzicadenti: $\emptyset \quad I \quad \textcolor{red}{I} \quad II \quad \textcolor{red}{II} \quad III$

Infatti, la regola $II = III$ consente di scambiare rosso-nero con nero-nero-rosso, aumentando la lunghezza. Però, intanto consente di separare sempre gli stuzzicadenti neri dai rossi, quindi le altre due regole provvedono a limitare la lunghezza. Mostriamo un esempio, mettendo un po' di spazio tra gli stuzzicadenti per evidenziare le parti da manipolare:

$$III = II \text{ I} = III \text{ I} = II \text{ II} = II \text{ III} = III \text{ II} = \emptyset \text{ II} = II$$

Ne viene la possibilità di ottenere la tavola di concatenazione: come nel caso precedente, ogni fila ha l'inversa, quindi si ha un gruppo, ma non c'è la proprietà commutativa:

&	$\emptyset$	I	II	$\textcolor{red}{I}$	$\textcolor{red}{II}$	III
$\emptyset$	$\emptyset$	I	II	$\textcolor{red}{I}$	$\textcolor{red}{II}$	III
I	I	II	$\emptyset$	$\textcolor{red}{II}$	III	$\textcolor{red}{I}$
II	II	$\emptyset$	I	III	$\textcolor{red}{I}$	$\textcolor{red}{II}$
$\textcolor{red}{I}$	$\textcolor{red}{I}$	III	$\textcolor{red}{II}$	$\emptyset$	II	I
$\textcolor{red}{II}$	$\textcolor{red}{II}$	$\textcolor{red}{I}$	III	I	$\emptyset$	II
III	III	$\textcolor{red}{II}$	$\textcolor{red}{I}$	II	I	$\emptyset$

La regola $II = III$ è detta *di inversione* perché la fila II è inversa della fila I : lo scambio del rosso e del nero sostituisce al nero la sua fila inversa.

NOTA. In questo caso si ottiene un gruppo non abeliano d'ordine 6, detto *gruppo diedrale* D_3 ed è solo un altro modo di rappresentare il gruppo delle simmetrie di un triangolo equilatero. Se la lunghezza della fila nera che viene mangiata è $m > 2$, e nella regola d'inversione la sequenza rosso-nero diventa la sequenza di $m-1$ neri ed un rosso, si ottengono gruppi non abeliani d'ordine $2m$, detti *gruppi diedrali* D_m . Se $m = 2$ si ottiene invece un gruppo abeliano d'ordine 4.

La regola precedente si generalizza variando le lunghezze massime delle file nere o rosse e la regola d'inversione, ma non in modo arbitrario, per evitare giochi banali, o in cui resta un solo colore. Ma quali di questi insiemi di regole producono giochi banali? Questo è un problema incredibilmente difficile, così come è altrettanto difficile sapere se alla fine avremo un numero finito o infinito di file di stuzzicadenti.

L'uso di lineette verticali per raffigurare gli stuzzicadenti alla fine risulta pesante e rende difficile la lettura di ciò che si sta facendo, oltre alla difficoltà di usare i colori nelle formule e nella stampa. Perché non sostituire al loro posto delle lettere? Ossia, in luogo di  perché non scrivere $bbbaabaaabbb$, sostituendo allo stuzzicadenti nero la lettera a e al posto del rosso la b ? Questo oggetto è una *parola* (o *stringa*) nell'alfabeto $\{a, b\}$. Per migliorarne la leggibilità, si possono anche usare gli esponenti e scrivere $b^3a^2ba^3b^3$, ottenendo una *parola normalizzata*. A causa della somiglianza con i monomi, il solo pericolo è introdurre istintivamente la commutatività delle lettere e trasformare erroneamente $b^3a^2ba^3b^3$ in a^5b^7 . La parola vuota si esprime di solito con il simbolo 1 . Le regole via via introdotte diventano quindi uguaglianze di parole, dette anche *relazioni*.

Il modo di scrivere che si usa in algebra in questi casi è il seguente: entro due specie di parentesi a punta si collocano la lista delle lettere da usare, dette *generatori*, e quella delle regole del gioco, dette *relazioni*. Gli esempi già visti si riscrivono così:

4.2.a $\langle a \mid \emptyset \rangle$	4.2.b $\langle a \mid a^n = a^m \rangle$	4.2.c $\langle a \mid a^n = a^m \rangle$	4.2.d $\langle a \mid a^n = a^m = 1 \rangle$	4.2.e $\langle a, b \mid ab = ba \rangle$
4.2.f $\langle a, b \mid ab = ba = 1 \rangle$	4.2.g $\langle a, b \mid ab = ba, a^3 = 1 = b^2 \rangle$		4.2.h $\langle a, b \mid ba = a^2b, a^3 = 1 = b^2 \rangle$	

Nota. Quel che abbiamo visto è un caso particolare di *presentazione* di un gruppo, ossia del modo usato dagli esperti per descrivere un gruppo astratto come insieme di parole nell'alfabeto costituito dai generatori, manipolate mediante le regole date dalle relazioni.

Per completezza aggiungiamo che gli esponenti nelle parole normalizzate possono essere anche negativi, con la clausola tacita che, denotando con x una lettera qualunque, sia sempre $xx^{-1} = x^{-1}x = 1 = x^0$. Con questa convenzione, per esempio la relazione $ab = ba$ si può esprimere scrivendo $a^{-1}b^{-1}ab = 1$.

Con il pretesto del gioco con gli stuzzicadenti, colorati o no, abbiamo visto esempi di monoidi e gruppi anche assai generali. Ovviamente, aumentando i colori aumenta la complessità delle costruzioni.

NOTA. In questo contesto, il ruolo delle lettere non è quello di variabili su un insieme, ma di *indeterminate*, ossia di oggetti astratti, estranei agli insiemi numerici, manipolati con regole del gioco stabilite di volta in volta. Per evitare questa confusione tra indeterminate e variabili si sono usati inizialmente gli stuzzicadenti colorati al posto delle più comuni lettere dell'alfabeto.

Esempio 4.2.i Sia $G = \langle a, b \mid a^4 = 1, b^2 = a^2, ba = a^3b \rangle$. Questo è un gruppo con 8 elementi, il più sfuggente di tutti. È denotato con Q_8 ed è detto gruppo dei quaternioni. Per costruirlo seguiamo la convenzione che ogni elemento sia scritto in modo che la potenza di a preceda quella di b . Non basta, perché ogni elemento può essere scritto in più modi: $a^3 = a \cdot a^2 = a \cdot b^2$. Inoltre,

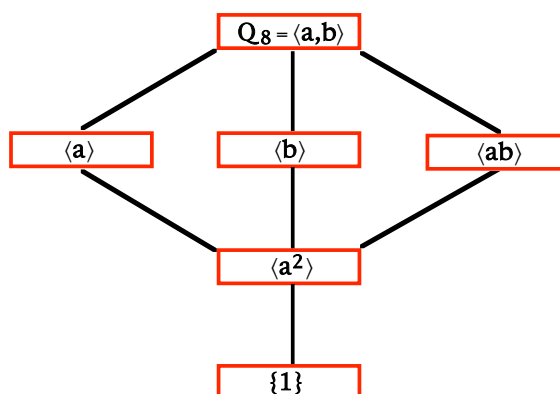
$$b^4 = (b^2)^2 = (a^2)^2 = a^4 = 1, (ab)^2 = abab = a(a^3b)b = a^4b^2 = b^2$$

Ne segue che i sei elementi $a, a^3, b, b^3, ab, a^3b = ab^3 = (ab)^3$ hanno periodo 4.

Infine, l'elemento $x = a^2 = b^2 = (ab)^2$ commuta con tutti gli altri ed è il quadrato dei sei precedenti. Allora,

$$Q_8 = \{id, a, a^2, a^3, b, b^3, ab, a^3b\}.$$

Qui accanto c'è il diagramma di Hasse dei sei sottogruppi di Q_8 , ordinati per inclusione.



SCHEMA A4. Calcolo combinatorio e funzioni aritmetiche

Numerose parti della Matematica sono riconducibili al calcolo del numero di elementi di un insieme finito (calcolo combinatorio) o allo studio di famiglie di sottoinsiemi di un insieme (geometria combinatoria) o di insiemi di permutazioni (combinatoria algebrica). Collegati a questi problemi ci sono aspetti interessanti di aritmetica, quali le partizioni intere di un numero naturale o le funzioni aritmetiche (teoria dei numeri).

A4.1. Sia X un insieme con n elementi. Calcoliamo il numero $\pi(n)$ di partizioni di X . Il numero cercato è detto *numero di Bell*. Sia $X = \{x_1, \dots, x_n\}$. Se $\pi = \{B_1, \dots, B_k\}$ (dove i B_i sono sottoinsiemi non vuoti e a due a due disgiunti di X) è una partizione con k elementi, la chiamiamo *k-partizione* di X . Sia $v_{n,k}$ il numero di k -partizioni di X

(*numero di Stirling*). Si ha subito $v_{n,1} = v_{n,n} = 1$, inoltre $\pi(n) = \sum_{k=1}^n v_{n,k}$. Calcoliamo

ricorsivamente i numeri di Stirling.

Supposto $n > 1$, poniamo $X^* = X \setminus \{x_n\}$. Sia $\pi = \{B_1, \dots, B_k\}$ una k -partizione di X . Allora $x_n \in B_i$ per un certo $i \in \{1, \dots, k\}$. Distinguiamo due casi:

a) $B_i = \{x_n\}$. Allora $\pi^* = \{B_1, \dots, B_{i-1}, B_{i+1}, \dots, B_k\}$ è una $(k-1)$ -partizione di X^* .

Inversamente da ogni $(k-1)$ -partizione π^* di X^* si ottiene una ed una sola k -partizione π di X aggiungendole $\{x_n\}$. Pertanto X ha $v_{n-1,k-1}$ k -partizioni di questo tipo.

b) B_i contiene altri elementi oltre x_n . Sia $B'_j = B_j$ per ogni $j \neq i$ e sia $B'_i = B_i \setminus \{x_n\}$. L'insieme $\pi' = \{B'_1, \dots, B'_k\}$ è una k -partizione di X^* . Inversamente, ogni k -partizione di X^* produce k diverse k -partizioni di X di questo secondo tipo aggiungendo l'elemento x_n ad uno qualunque dei B'_i , $i = 1, \dots, k$. Pertanto X ha $k \cdot v_{n-1,k}$ k -partizioni di questo tipo.

In definitiva, $v_{n,k} = v_{n-1,k-1} + k \cdot v_{n-1,k}$ per ogni $k = 1, \dots, n$.

Si può allora costruire, anche con l'ausilio di Excel, un triangolo, simile a quello di Tartaglia, per calcolare $v_{n,k}$ per ogni n, k e quindi per calcolare $\pi(n)$: nella prima riga c'è $v_{1,1}$, nella seconda $v_{2,1}$ e $v_{2,2}$ e così via.

$n \backslash k$	1	2	3	4	5	6	7	8	9	10	$\pi(n)$
1	1	0	0	0	0	0	0	0	0	0	1
2	1	1	0	0	0	0	0	0	0	0	2
3	1	3	1	0	0	0	0	0	0	0	5
4	1	7	6	1	0	0	0	0	0	0	15
5	1	15	25	10	1	0	0	0	0	0	52
6	1	31	90	65	15	1	0	0	0	0	203
7	1	63	301	350	140	21	1	0	0	0	877
8	1	127	966	1.701	1.050	266	28	1	0	0	4.140
9	1	255	3.025	7.770	6.951	2.646	462	36	1	0	21.147
10	1	511	9.330	34.105	42.525	22.827	5.880	750	45	1	115.975

A4.2. Collegato a questo problema c'è quello più aritmetico di contare le partizioni di un numero naturale n , ossia di tutte le liste (non crescenti) di addendi interi positivi aventi per somma n . C'è una formula, di Ramanujan, per determinare direttamente tutte le partizioni di n , ma è complicata. Tra i procedimenti ricorsivi il più intuitivo è forse il seguente. Siano note le partizioni di $n-1$. Con i due procedimenti A) e B) costruiamo due liste di partizioni di n .

- A) si aggiunge 1 in coda a ciascuna partizione di $n-1$ e si ha una partizione di n ;
- B) per ogni partizione di $n-1$, si somma 1 a turno ad ogni addendo e si riordina in senso non crescente ogni partizione di n così ottenuta.
- C) Si ordinano in ordine *lessicografico inverso* le partizioni di n così trovate e si eliminano infine i doppi, un lavoraccio ...

Vediamo gli esempi piccoli. Le graffe denotano qui insiemi ordinati, quindi con possibili ripetizioni.

- L'unica partizione di 1 è la banale: $\{1\}$

- Per trovare le partizioni di 2,

- A) aggiungiamo 1 in fondo all'unica lista di 1: $\{1,1\}$;
- B) sommiamo 1 all'unico elemento di $\{1\}$ ottenendo $\{2\}$;
- C) le partizioni di 2 sono: $\{2\}, \{1,1\}$.

- Per trovare le partizioni di 3,

- A) aggiungiamo 1 in fondo alle due liste del 2: $\{2,1\}, \{1,1,1\}$;
- B) sommiamo 1 all'unico elemento di $\{2\}$ ottenendo $\{3\}$, e a ciascuno degli elementi di $\{1,1\}$ ottenendo $\{2,1\}$ e $\{1,2\}$;
- C) riordiniamo ciascuna lista in senso non crescente, poi elenchiamole in ordine lessicografico inverso e sfoltiamo: $\{3\}, \{2,1\}, \{1,1,1\}$.

- Troviamo le partizioni di 4, senza commenti

- A) $\{3,1\}, \{2,1,1\}, \{1,1,1,1\}$
- B) $\{4\}, \{3,1\}, \{2,2\}, \{2,1,1\}, \{1,2,1\}, \{1,1,2\}$

C) in definitiva, le partizioni di 4 sono: $\{4\}$, $\{3,1\}$, $\{2,2\}$, $\{2,1,1\}$, $\{1,1,1,1\}$.

- Concludiamo con le partizioni di 5

A) $\{4,1\}$, $\{3,1,1\}$, $\{2,2,1\}$, $\{2,1,1,1\}$, $\{1,1,1,1,1\}$

B) $\{5\}$, $\{4,1\}$, $\{3,2\}$, $\{3,1,1\}$, $\{2,2,1\}$, $\{2,1,1,1\}$, $\{1,1,1,1,1\}$, ...

C) in definitiva: $\{5\}$, $\{4,1\}$, $\{3,2\}$, $\{3,1,1\}$, $\{2,2,1\}$, $\{2,1,1,1\}$, $\{1,1,1,1,1\}$

La crescita è molto più rapida di quel che sembra. Le partizioni di 6 sono 11:

$\{6\}$, $\{5,1\}$, $\{4,2\}$, $\{4,1,1\}$, $\{3,3\}$, $\{3,2,1\}$, $\{3,1,1,1\}$, $\{2,2,2\}$, $\{2,2,1,1\}$, $\{2,1,1,1,1\}$, $\{1,1,1,1,1,1\}$

Quelle di 7 sono 15 e quelle di 8 sono 22.

Vediamo la seguente tabella, che nella III riga riporta le differenze.

n	1	2	3	4	5	6	7	8	9	10	11	12
p_n	1	2	3	5	7	11	15	22	30	42	56	77

Nota. Ad ogni partizione di un insieme X con n oggetti è associata una partizione di n , formata dalle lunghezze delle classi della partizione. Ma, naturalmente, ogni partizione di n corrisponde a più partizioni di X . A quante? Ecco un altro bel problema combinatorio.

A4.3. Alcune funzioni aritmetiche. Si tratta di particolari funzioni di dominio e codominio $\mathbf{N}^+$ o $\mathbf{Z}$, che esprimono proprietà dei numeri interi. Un esempio è la *funzione di Eulero*, $\varphi: \mathbf{N}^+ \rightarrow \mathbf{N}$, definita da:

$$\varphi(n) = \text{numero di elementi dell'insieme } \{k \in \mathbf{N} \mid 1 \leq k \leq n, \text{MCD}(n,k) = 1\}$$

Questa funzione ha varie applicazioni, sia in Algebra (gruppi ciclici), sia in Teoria dei Numeri e nella Crittografia, e si calcola anche usando la scomposizione in

fattori primi. Posto infatti $n = \prod_{i=1}^r p_i^{a_i}$, dove i p_i sono primi distinti ed $a_i \geq 1$, si ha:

$$\varphi(n) = n \cdot \prod_{i=1}^r \left(1 - \frac{1}{p_i}\right)$$

Analoga è la funzione di Möbius $\mu: \mathbf{N}^+ \rightarrow \mathbf{Z}$, definita da: $\mu(n) = \begin{cases} (-1)^r & \text{se } a_i = 1 \forall i \\ 0 & \text{altrimenti} \end{cases}$

Si noti che $\mu(n) = 0$ se e solo se n è multiplo di un quadrato, perché in tal caso almeno uno degli a_i è maggiore di 1. Vediamo alcuni esempi. Si denoti con p un numero primo.

n	1	p	$p^k, k > 1$	12	30
$\varphi(n)$	1	$p-1$	$p^{k-1} \cdot (p-1)$	4	8
$\mu(n)$	1	-1	0	0	-1

Nella Teoria dei Numeri sono stabilite relazioni fra le varie funzioni aritmetiche e in questo contesto la funzione di Möbius ha un ruolo centrale. Infatti nel calcolo della funzione φ di Eulero si ha:

$$\prod_{i=1}^r \left(1 - \frac{1}{p_i}\right) = 1 - \frac{1}{p_1} - \frac{1}{p_2} - \dots - \frac{1}{p_r} + \frac{1}{p_1 p_2} + \frac{1}{p_1 p_3} + \dots + \frac{1}{p_{r-1} p_r} - \frac{1}{p_1 p_2 p_3} - \frac{1}{p_1 p_2 p_4} - \dots + (-1)^r \frac{1}{p_1 p_2 \dots p_r}$$

Si noti che i segni davanti ad ogni addendo sono precisamente i valori della funzione di Möbius del denominatore, e che i denominatori sono tutti e soli i divisori di n che non sono multipli di un quadrato. Possiamo allora aggiungere anche le frazioni $\mu(d)/d$ corrispondenti agli altri divisori di n , quelli con $\mu(d) = 0$, e potremo scrivere:

$$\prod_{i=1}^r \left(1 - \frac{1}{p_i}\right) = \sum_{d|n} \frac{\mu(d)}{d}$$

Pertanto, per ogni $n \in \mathbf{N}^+$, posto per ogni divisore d di n , $d' = n/d$, si ottiene:

$$\varphi(n) = n \cdot \sum_{d|n} \frac{\mu(d)}{d} = \sum_{d \cdot d' = n} d' \cdot \mu(d)$$

Esempio. Sia $n = 12$. I divisori di 12 sono: 1, 2, 3, 4, 6, 12.

d	1	2	3	4	6	12
d'	12	6	4	3	2	1
$\mu(d)$	1	-1	-1	0	1	0
$d' \cdot \mu(d)$	12	-6	-4	0	2	0

$\Rightarrow \varphi(12) = 12 - 6 - 4 + 0 + 2 + 0 = 4$

Una applicazione importante della funzione di Möbius è la *formula d'inversione delle successioni*: Siano $f, g: \mathbf{N}^+ \rightarrow \mathbf{R}$. Se per ogni $n \in \mathbf{N}^+$ si ha $g(n) = \sum_{d|n} f(d)$, allora $f(n) = \sum_{d \cdot d' = n} g(d') \cdot \mu(d)$.

Vediamone un caso particolare: sia $f(n) = \begin{cases} 1 & \text{se } n = 1 \\ 0 & \text{se } n > 1 \end{cases}$. Allora per ogni $n \in \mathbf{N}^+$ si ha

ovviamente $g(n) = \sum_{d|n} f(d) = 1$. Dalla formula d'inversione segue allora:

$$f(n) = \sum_{d \cdot d' = n} g(d') \cdot \mu(d) = f(n) = \sum_{d \cdot d' = n} \mu(d),$$

ossia, $\sum_{d|n} \mu(d) = \begin{cases} 1 & \text{se } n = 1 \\ 0 & \text{se } n > 1 \end{cases}$.

Da quest'ultima formula segue l'invertibilità della formula d'inversione, ossia se

$$f(n) = \sum_{d \cdot d' = n} g(d') \cdot \mu(d) \text{ allora } g(n) = \sum_{d|n} f(d).$$

In particolare, posto $f = \varphi$ e $g = \text{identità}$, ne segue la nota formula: $n = \sum_{d|n} \varphi(d)$.

Esempio. Sia $n = 30$. I suoi divisori sono: 1, 2, 3, 5, 6, 10, 15, 30. Allora:

$$\varphi(1) + \varphi(2) + \varphi(3) + \varphi(5) + \varphi(6) + \varphi(10) + \varphi(15) + \varphi(30) = 1 + 1 + 2 + 4 + 2 + 4 + 8 + 8 = 30.$$

Il *teorema di Eulero* afferma che dati $a, n \in \mathbf{N}^+$ se $\text{MCD}(a, n) = 1$ allora n divide $a^{\varphi(n)} - 1$. In particolare, Se p è primo ed a è un intero, allora p divide $a^p - a$, e questo è noto come *teorema di Fermat*.

Esempio. Sia $p = 6$: non è primo perché non divide $2^6 - 2 = 62$. Invece, $p = 101$ sarà primo? Con

l'ausilio di Excel si può calcolare $\frac{a^{101} - a}{101}$ per $a = 2, 3, 4, 5$ e scoprire che sono tutti numeri interi,

ma talmente lunghi che non ci stanno in una riga. Però anche solo $a = 2$ fornisce una buona condizione sufficiente a scartare un numero enorme di numeri dispari candidati ad essere primi:

n	3	5	7	9	11
$2^n - 2$	6	30	126	510	2.046
resto	0	0	0	6	0
n	13	15	17	19	21
$2^n - 2$	8.190	32.766	131.070	524.286	2.097.150
resto	0	6	0	0	6
n	23	25	27	29	31
$2^n - 2$	8.388.606	33.554.430	134.217.726	536.870.910	2.147.483.646
resto	0	5	24	0	0

SCHEDA A5. Matrici multiformi.

Nei corsi di Algebra Lineare inseriti nei programmi universitari di Geometria, di istituzioni di Matematiche o di Analisi Matematica si incontrano le matrici associate alle applicazioni lineari tra due spazi vettoriali di dimensioni finite ed ai sistemi lineari. Si apprende così a sommare e moltiplicare matrici, calcolarne il rango, ridurle a scal per righe o, nel caso siano quadrate, trovarne il determinante e gli autovalori, nonché varie scomposizioni canoniche. Tuttavia, le matrici possono essere usate anche per altri scopi, in algebra, geometria combinatoria, grafica, enigmistica, persino in agraria, ed entrano anche in oggetti che fanno parte ormai della nostra vita quotidiana.

A5.1. I quadrati latini. Sono matrici quadrate d'ordine $n \geq 1$, riempite con n oggetti, non necessariamente numeri, senza ripetizioni nelle righe e nelle colonne. Si chiamano così perché gli esempi più antichi erano riempiti da lettere latine maiuscole.

ESEMPIO A5.1.1. Ecco l'unico quadrato latino d'ordine 2 e i due di ordine 3.

$$\begin{bmatrix} A & B \\ B & A \end{bmatrix}, \begin{bmatrix} A & B & C \\ B & C & A \\ C & A & B \end{bmatrix}, \begin{bmatrix} A & B & C \\ C & A & B \\ B & C & A \end{bmatrix}$$

Sono i soli possibili di quegli ordini, con la prima riga in ordine alfabetico.

Detto X l'insieme degli n oggetti distinti che occupano le caselle, possiamo interpretare il quadrato latino come la tavola di una operazione binaria $* : X \times X \rightarrow X$, dotata delle leggi di cancellazione destra e sinistra. Inversamente, ogni operazione in X con le due leggi di cancellazione destra e sinistra fornisce un quadrato latino. La struttura algebrica corrispondente $(X, *)$ si chiama quasi-gruppo; se ha l'elemento neutro, è detta cappio (loop). In un cappio, ogni elemento ha inversi destro e sinistro, dato che in ogni riga e colonna deve trovarsi l'elemento neutro; pertanto, se l'operazione è anche associativa è un gruppo, dato che in tal caso è un monoide ed in un monoide l'inverso destro e l'inverso sinistro di uno stesso elemento coincidono sempre. I primi due quadrati latini dell'esempio A5.1.1.

definiscono gruppi d'ordine 2 e 3 rispettivamente, mentre il terzo non ha l'elemento neutro: infatti, A è elemento neutro solo a destra: $A * A = A$, $A * B = B$, $A * C = C$, ma $B * A = C$.

Mediante l'interpretazione algebrica è possibile definire l'isomorfismo di due quadrati latini considerati come quasi-gruppi $(X, *)$ e $(Y, \bullet)$: sono isomorfi se esiste una funzione $f : X \xrightarrow{1-1} Y$, tale che per ogni $a, b \in X$, $f(a * b) = f(a) \bullet f(b)$. La f si può interpretare come un cambio di nome degli oggetti e dell'operazione, ossia una traduzione da un simbolismo all'altro. Per questo è possibile da un lato codificare ogni quadrato latino d'ordine n mediante l'insieme $\mathbf{N}_n = \{k \in \mathbf{N} \mid 1 \leq k \leq n\}$, ossia impiegando numeri al posto degli oggetti originari; dall'altro, mediante una permutazione di $\mathbf{N}_n$, trasformare il quadrato latino in un altro isomorfo. Di conseguenza, ogni quadrato latino è isomorfo ad un altro la cui prima riga è la successione 1, 2, ..., n .

ESEMPIO A5.1.2. Mediante la sostituzione $f = \begin{cases} C \rightarrow 1 \\ A \rightarrow 2 \\ B \rightarrow 3 \end{cases}$, $\begin{bmatrix} C & A & B \\ A & B & C \\ B & C & A \end{bmatrix}$ diventa $\begin{bmatrix} 1 & 2 & 3 \\ 2 & 3 & 1 \\ 3 & 1 & 2 \end{bmatrix}$

ESEMPIO A5.1.3. Per ogni $n \geq 1$ esistono quadrati latini di ordine n : basta infatti considerare la tavola di moltiplicazione di un gruppo ciclico di ordine n . Per ogni $n = 2m$, $m \geq 2$ ne esiste almeno un altro non isomorfo al precedente: la tavola del *gruppo di Klein* $\langle a, b \mid a^2 = b^2 = 1, ab = ba \rangle$ per $m = 2$ o, per $m > 2$, del *gruppo diedrale* $D_m \langle a, b \mid a^m = b^2 = 1, a^{m-1}b = ba \rangle$. Qui sotto i casi $m = 2, 3$. La seconda tavola codifica quella vista nella sezione degli "stuzzicadenti"

$$\begin{bmatrix} 1 & 2 & 3 & 4 \\ 2 & 1 & 4 & 3 \\ 3 & 4 & 1 & 2 \\ 4 & 3 & 2 & 1 \end{bmatrix}, \begin{bmatrix} 1 & 2 & 3 & 4 & 5 & 6 \\ 2 & 3 & 1 & 5 & 6 & 4 \\ 3 & 1 & 2 & 6 & 4 & 5 \\ 4 & 6 & 5 & 1 & 3 & 2 \\ 5 & 4 & 6 & 2 & 1 & 3 \\ 6 & 5 & 4 & 3 & 2 & 1 \end{bmatrix}$$

I quadrati latini di ordine 9 sono entrati nei nostri passatempi grazie al *sudoku*, che è un gioco divenuto popolare nel 2005, perché contenuto non solo in popolari riviste di enigmistica, ma anche su quotidiani di grande tiratura. Si presenta come una matrice quadrata 9×9 , ripartita in nove sottomatrici 3×3 , dette "riquadri".

Alcune caselle sono occupate da cifre, che costituiscono i dati iniziali, mentre le altre sono vuote. Il giocatore deve compilare tutte le 81 caselle inserendovi le cifre 1, 2, ..., 9, rispettando quelle già occupate, in modo che in ogni riga, colonna e riquadro ci siano una ed una sola volta tutte le nove cifre. La matrice finale è quindi un quadrato latino, sia pure di tipo particolare per via della suddivisione in 9 riquadri ciascuno dei quali contiene tutte le nove cifre.

Come detto, con un'opportuna sostituzione delle nove cifre in tutto lo schema si può sempre far sì che la prima riga sia la *sequenza base* 1, 2, 3, 4, 5, 6, 7, 8, 9. Ne segue che per ogni schema se ne trovano $9! = 362.880$ isomorfi. Per calcolare quanti sono i possibili schemi sudoku, si tratta allora di trovare quanti sono quelli con la sequenza base nella prima riga. Buon divertimento!

NOTA. Un quadrato latino riempito con gli elementi di $\mathbf{N}_n$ ha la proprietà che la somma di ogni riga ed ogni colonna è sempre costante, $\frac{n \cdot (n+1)}{2}$.

16	3	2	13
5	10	11	8
9	6	7	12
4	15	14	1

Ci sono altri tipi di quadrati di numeri con questa proprietà; sono detti *magici* se gli n^2 numeri che li riempiono sono tutti diversi e in successione, e se anche le due diagonali hanno la stessa somma. Ce ne sono di molto antichi. Qui vediamo quello raffigurato da A. Dürer in una sua celebre incisione, *Melancolia*, del 1514; la somma costante è 34.

A5.2. Matrici booleane. Sono matrici riempite con due soli simboli, solitamente 0 ed 1, che possono essere interpretati come Falso/Vero (matrici *booleane*) oppure come gli elementi neutri delle due operazioni di un campo.

Una relazione $(A, B, \mathfrak{R})$ fra due insiemi finiti, con $A = \{a_1, \dots, a_m\}$, $B = \{b_1, \dots, b_n\}$ ed $\mathfrak{R} \subseteq A \times B$, si può rappresentare mediante una matrice $R = [r_{ij}]$, $1 \leq i \leq m$, $1 \leq j \leq n$, detta *matrice d'incidenza* della relazione. Le sue caselle sono definite da:

$$r_{ij} = \begin{cases} 1 & \text{se } a_i \mathfrak{R} b_j \\ 0 & \text{altrimenti} \end{cases}$$

Si tratta quindi di una matrice booleana. (*)

Esempio A5.2.1. Siano $A = B = \{1, 2, 3, 4, 6, 12\}$ e sia $x \mathfrak{R} y$ se x è divisore di y . Allora avremo:

$$\mathfrak{R} = \{(1, 1), (1, 2), (1, 3), (1, 4), (1, 6), (1, 12), (2, 2), (2, 4), (2, 6), (2, 12), (3, 3), (3, 6), (3, 12), (4, 4), (4, 12), (6, 6), (6, 12), (12, 12)\}.$$

Allora la matrice d'incidenza è:

$$\begin{bmatrix} 1 & 1 & 1 & 1 & 1 & 1 \\ 0 & 1 & 0 & 1 & 1 & 1 \\ 0 & 0 & 1 & 0 & 1 & 1 \\ 0 & 0 & 0 & 1 & 0 & 1 \\ 0 & 0 & 0 & 0 & 0 & 1 \\ 0 & 0 & 0 & 0 & 0 & 1 \end{bmatrix}.$$

Data la relazione $\mathfrak{R}$ fra A e B , possiamo definire una relazione $\mathfrak{R}^T$ fra B ed A , che chiameremo *trasposta* di $\mathfrak{R}$, definita come l'insieme delle coppie (b, a) tali che $(a, b) \in \mathfrak{R}$. Chiaramente, se A e B sono finiti, la matrice d'incidenza di $\mathfrak{R}^T$ è la trasposta della matrice d'incidenza di $\mathfrak{R}$.

Date le relazioni $(A, B, \mathfrak{R})$ e $(B, C, \mathcal{S})$, si può definire la *relazione composta* $(A, C, \mathcal{T})$, dove $\mathcal{T} = \mathfrak{R} \circ \mathcal{S}$ è definita da: per ogni $a \in A$, $c \in C$,

$$a \mathcal{T} c \Leftrightarrow \exists b \in B \text{ tale che } a \mathfrak{R} b \text{ e } b \mathcal{S} c.$$

Se gli insiemi A , B , C coinvolti sono finiti di ordini rispettivamente m , n , p , possiamo, come detto, rappresentare le relazioni $(A, B, \mathfrak{R})$ e $(B, C, \mathcal{S})$ mediante le matrici R ed S , rispettivamente di dimensioni $m \times n$ e $n \times p$. Utilizziamo per operare sui loro elementi 0 ed 1 le *operazioni booleane* definite come segue:

$$\begin{array}{c|cc} + & 0 & 1 \\ \hline 0 & 0 & 1 \\ 1 & 1 & 1 \end{array} \quad \begin{array}{c|cc} \times & 0 & 1 \\ \hline 0 & 0 & 0 \\ 1 & 0 & 1 \end{array}$$

Abbiamo allora il seguente risultato: la matrice d'incidenza T della relazione composta $\mathcal{T} = \mathfrak{R} \circ \mathcal{S}$ è il prodotto (booleano) righe per colonne delle matrici R ed S .

Tra le relazioni le più importanti sono le funzioni. La matrice d'incidenza R_f di una funzione $f : A \rightarrow B$ è caratterizzata dal fatto che in ogni riga c'è un solo 1. Se poi f è iniettiva, anche in ogni colonna c'è un solo 1, e viceversa. La suriettività

(*) In altri testi la matrice d'incidenza è definita come la trasposta di R .

invece è caratterizzata dal fatto che in ogni colonna ci sia almeno un 1. Infine, f è biiettiva se e solo se la sua matrice è quadrata ed in ogni riga e colonna c'è uno ed un solo 1.

Consideriamo ora il caso $A = B = X$, con n elementi. Tra le relazioni in X c'è l'*identità* id_X , costituita dalle coppie (x, x) , $x \in X$ e contemporaneamente biiezione tra X e se stesso e relazione d'equivalenza in X . La sua matrice sarà denotata al solito con I_n ed avrà la consueta forma: tutti 1 sulla diagonale principale e tutti 0 fuori di essa.

Le biiezioni da X a se stesso sono dette *permutazioni* di X , e le loro matrici d'incidenza *matrici di permutazione*. Ce ne sono $n!$, costituiscono un gruppo isomorfo al gruppo simmetrico su n oggetti ed hanno la proprietà che il prodotto di una di esse per la trasposta dà l'identità I_n . Ossia, l'inversa coincide con la trasposta.

Poiché in ogni riga e colonna c'è un solo 1, nel prodotto di ogni riga per ogni colonna di due matrici di permutazione abbiamo al massimo un prodotto $1 \cdot 1$, mentre gli altri $n-1$ hanno sempre un fattore nullo, quindi c'è una somma di zeri con un unico eventuale addendo = 1. Allora si può pensare che nel calcolo del prodotto di due matrici di permutazione in luogo delle operazioni booleane si

usino quelle di un campo, in cui si avrebbe

$$\begin{array}{c|cc} + & 0 & 1 \\ \hline 0 & 0 & 1 \\ 1 & 1 & 2 \end{array}, \begin{array}{c|cc} \times & 0 & 1 \\ \hline 0 & 0 & 0 \\ 1 & 0 & 1 \end{array}$$

Ne segue che questo gruppo di $n!$ matrici di permutazione si trova come sottogruppo del gruppo generale lineare $\text{GL}_n(K)$ delle matrici invertibili d'ordine n ad elementi in un campo K qualsiasi. Supponiamo ora che la caratteristica di K sia $\neq 2$. Si può calcolare il determinante di ciascuna di queste matrici: si ottiene 1 o -1, e la permutazione si dice in tal caso *pari* o rispettivamente *dispari*. Naturalmente, le matrici corrispondenti alle permutazioni pari formano un sottogruppo isomorfo al gruppo alterno su n oggetti ed incluso nel sottogruppo $\text{SL}_n(K)$ delle matrici a determinante 1.

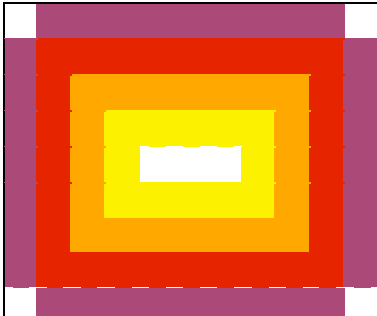
Le matrici booleane dello stesso tipo $m \times n$ si possono sommare con l'addizione booleana, ottenendo un monoide commutativo, in cui la matrice nulla $0_{m,n}$ è l'elemento neutro e quella riempita interamente da 1, che diremo $1_{m,n}$, è l'elemento assorbente. Si possono anche moltiplicare termine a termine: in tal caso

i ruoli di $0_{m,n}$ e $1_{m,n}$ sono scambiati. La somma ed il prodotto di una matrice per se stessa è la matrice stessa (idempotenza); inoltre, le due operazioni sono distributive ed assorbenti l'una rispetto all'altra. Per ogni matrice c'è la *complementare*, con gli zeri e gli 1 scambiati: la somma è la matrice $1_{m,n}$, il prodotto è la matrice $0_{m,n}$. L'insieme di queste 2^{mn} matrici con le due operazioni, gli elementi $0_{m,n}$ ed $1_{m,n}$ e la complementazione è un'algebra di Boole.

Un'applicazione assai diversa delle matrici booleane è la traduzione algebrica di disegni eseguiti su fogli quadrettati, in particolare su schermi di computer, in cui i quadretti sono i pixel, o su insegne luminose, in cui le indicazioni sono ottenute da led o lampade accese o spente. Lo schermo si traduce in una matrice booleana con 1 come ON e 0 come OFF.

1	1	1	1	1	0	1	1	1	1
0	0	0	1	0	0	1	0	0	0
0	0	1	0	0	0	1	1	1	0
0	0	0	1	0	0	0	0	0	1
0	0	0	0	1	0	0	0	0	1
1	0	0	0	1	0	1	0	0	1
0	1	1	1	0	0	0	1	1	0

Nota. Questa idea di matrici per rappresentare disegni si generalizza in vari modi, per introdurre il colore. Il più semplice, ispirato dalla pittura *pointilliste* di Seurat e Signac, consiste nel numerare un insieme di n colori "puri" con i numeri da 1 ad n e sostituire ad ogni puntino (pixel) colorato il suo numero. Allora ogni matrice ad elementi $0, 1, \dots, n$ diventa un disegno colorato, dove 0 è l'assenza di colore. Naturalmente, qui i colori non vengono mischiati.

	<table><tr><td></td><td>0</td></tr><tr><td></td><td>1</td></tr><tr><td></td><td>2</td></tr><tr><td></td><td>3</td></tr><tr><td></td><td>4</td></tr></table>		0		1		2		3		4	<table><tr><td>0</td><td>4</td><td>4</td><td>4</td><td>4</td><td>4</td><td>4</td><td>4</td><td>4</td><td>4</td><td>0</td></tr><tr><td>4</td><td>3</td><td>3</td><td>3</td><td>3</td><td>3</td><td>3</td><td>3</td><td>3</td><td>3</td><td>4</td></tr><tr><td>4</td><td>3</td><td>2</td><td>2</td><td>2</td><td>2</td><td>2</td><td>2</td><td>2</td><td>3</td><td>4</td></tr><tr><td>4</td><td>3</td><td>2</td><td>1</td><td>1</td><td>1</td><td>1</td><td>1</td><td>2</td><td>3</td><td>4</td></tr><tr><td>4</td><td>3</td><td>2</td><td>1</td><td>0</td><td>0</td><td>0</td><td>1</td><td>2</td><td>3</td><td>4</td></tr><tr><td>4</td><td>3</td><td>2</td><td>1</td><td>1</td><td>1</td><td>1</td><td>1</td><td>2</td><td>3</td><td>4</td></tr><tr><td>4</td><td>3</td><td>2</td><td>2</td><td>2</td><td>2</td><td>2</td><td>2</td><td>2</td><td>3</td><td>4</td></tr><tr><td>4</td><td>3</td><td>3</td><td>3</td><td>3</td><td>3</td><td>3</td><td>3</td><td>3</td><td>3</td><td>4</td></tr><tr><td>0</td><td>4</td><td>4</td><td>4</td><td>4</td><td>4</td><td>4</td><td>4</td><td>4</td><td>4</td><td>0</td></tr></table>	0	4	4	4	4	4	4	4	4	4	0	4	3	3	3	3	3	3	3	3	3	4	4	3	2	2	2	2	2	2	2	3	4	4	3	2	1	1	1	1	1	2	3	4	4	3	2	1	0	0	0	1	2	3	4	4	3	2	1	1	1	1	1	2	3	4	4	3	2	2	2	2	2	2	2	3	4	4	3	3	3	3	3	3	3	3	3	4	0	4	4	4	4	4	4	4	4	4	0
	0																																																																																																														
	1																																																																																																														
	2																																																																																																														
	3																																																																																																														
	4																																																																																																														
0	4	4	4	4	4	4	4	4	4	0																																																																																																					
4	3	3	3	3	3	3	3	3	3	4																																																																																																					
4	3	2	2	2	2	2	2	2	3	4																																																																																																					
4	3	2	1	1	1	1	1	2	3	4																																																																																																					
4	3	2	1	0	0	0	1	2	3	4																																																																																																					
4	3	2	1	1	1	1	1	2	3	4																																																																																																					
4	3	2	2	2	2	2	2	2	3	4																																																																																																					
4	3	3	3	3	3	3	3	3	3	4																																																																																																					
0	4	4	4	4	4	4	4	4	4	0																																																																																																					

A5.3. Tavole di moltiplicazione. Si tratta di matrici ad elementi in un insieme finito $X = \{x_1, \dots, x_n\}$. Data un'operazione $*$: $X \times X \rightarrow X$, la tavola di moltiplicazione T di $*$ è definita da: $t_{ij} = x_h \Leftrightarrow x_i * x_j = x_h$. Di solito, se non ci sono ambiguità, si può scrivere: $t_{ij} = h \Leftrightarrow x_i * x_j = x_h$, e così T è una matrice ad elementi interi da 1 ad n . Essa definisce un'operazione $*$ su $\mathbf{N}_n = \{k \in \mathbf{N} \mid 1 \leq k \leq n\}$ che rende la struttura $(\mathbf{N}_n, *)$ isomorfa alla struttura $(X, *)$ di partenza. Quest'ultima codifica è assai utile per rappresentare strutture astratte mediante tavole di moltiplicazione, perché è assai più compatta di quella originale.

NOTA. Alcune delle proprietà di una struttura $(X, *)$ con X finito, si possono leggere direttamente sulla tavola:

- a) La proprietà commutativa si traduce nella simmetria della tavola rispetto alla diagonale che esce dal vertice in alto a sinistra (diagonale principale).
- b) L'elemento neutro dà luogo ad una riga e ad una colonna (nella stessa posizione) uguali rispettivamente alla riga sopra la tavola ed alla colonna a sinistra della tavola.
- c) La legge di cancellazione assicura che in ogni riga e colonna ogni elemento compaia una volta sola. Ossia, la tavola è un quadrato latino.
- d) Ogni elemento ha un simmetrico se e solo se in ogni riga e colonna compare l'elemento neutro e se le posizioni da esso occupate sono simmetriche rispetto alla diagonale principale.
- e) L'idempotenza si traduce nel fatto che la diagonale principale è uguale alla colonna a sinistra della tavola.
- f) L'elemento assorbente ha solo se stesso nella sua riga e nella sua colonna.

Per quanto riguarda la proprietà associativa, che è la più importante, e che non è leggibile sulla tavola, dato che coinvolge tre elementi e non due, se la tavola di moltiplicazione è data come matrice codificata ad elementi interi da 1 ad n , si può scrivere un programmino per verificare se tutte le terne la soddisfino o no. L'algoritmo è il seguente: detti $t[i,j]$ gli elementi della matrice T ,

- se per ogni i, j, k si ha $t[i, t[j, k]] = t[t[i, j], k]$ allora l'operazione è associativa
- se esistono i, j, k tali che $t[i, t[j, k]] \neq t[t[i, j], k]$, l'operazione non è associativa

Questo algoritmo si semplifica se l'operazione ha l'elemento neutro, di solito identificato con 1. Ogni terna che lo coinvolge è associativa, sicché le iterazioni per i, j, k possono partire da 2.

Talora, però, è possibile escludere che un'operazione data mediante la sua tavola sia associativa, mediante alcune proprietà che un gruppo possiede. Supponiamo che T sia un quadrato latino dotato di elemento neutro: la struttura corrispondente è un cappio; se vale l'associatività, si ha un gruppo.

- In un gruppo ogni elemento ha inverso destro ed inverso sinistro uguali, quindi l'unità sarà sempre in posizioni simmetriche rispetto alla diagonale principale. Se non è così, non è un gruppo e quindi l'operazione non è associativa.
- In un gruppo finito G , ogni sottoinsieme H contenente l'unità e chiuso rispetto alla moltiplicazione è un sottogruppo. Allora il suo ordine deve dividere l'ordine di G . Se l'ordine di un tal sottoinsieme non lo fa, G non è un gruppo.
- Il teorema di Cauchy per i gruppi finiti afferma che per ogni divisore primo p dell'ordine di un gruppo finito G esiste almeno un elemento di periodo p . Pertanto, solo nella tavola T di un gruppo d'ordine potenza di 2 l'unità può occupare tutte le caselle della diagonale principale (ossia, tutti gli elementi, eccetto l'unità, hanno periodo 2) In questo caso, inoltre, G è abeliano e T deve essere simmetrica. Altrimenti, non è la tavola di un gruppo.

°	1	2	3	4	5	6
1	1	2	3	4	5	6
2	2	1	5	6	3	4
3	3	4	6	5	2	1
4	4	3	2	1	6	5
5	5	6	4	3	1	2
6	6	5	1	2	4	3

^	1	2	3	4	5	6
1	1	1	1	1	1	1
2	1	2	2	2	2	2
3	1	2	3	3	3	3
4	1	2	3	4	4	4
5	1	2	3	4	5	5
6	1	2	3	4	5	6

*	1	2	3	4	5	6
1	1	2	3	4	5	6
2	2	1	4	5	6	3
3	3	4	1	6	2	5
4	4	5	6	1	3	2
5	5	6	2	3	1	4
6	6	3	5	2	4	1

•	1	2	3	4	5	6
1	1	2	3	4	5	6
2	2	3	1	5	6	4
3	3	1	2	6	4	5
4	4	5	6	3	1	2
5	5	6	4	2	3	1
6	6	4	5	1	2	3

Le quattro tavole qui a lato, complete di intestazioni, forniscono esempi di strutture di ordine 6. Tutte hanno l'unità. Le prime tre sono commutative. La seconda è anche idempotente, ha elemento assorbente, l'1, mentre il 6 è l'elemento neutro; l'operazione è "il minimo tra" ed è associativa, ma non è un gruppo. Le altre sono quadrati latini. La quarta non è associativa, perché l'unità 1 non è sempre in posizione simmetrica rispetto alla diagonale. La terza non è un gruppo, perché tutti gli elementi hanno il quadrato = 1 e quindi non ce ne sono di periodo 3.

Per vedere se la prima delle quattro tavole è associativa, occorrerebbe verificare tutte le $6^3 = 216$ terne, nessuna esclusa. In realtà, dato che ha l'unità 1, ne bastano $(6-1)^3 = 125$, perché le altre lo sono automaticamente. In ogni caso, la verifica con un programma basato su quel semplice algoritmo risponde che è associativa.

Se applichiamo invece il programma alla quarta delle quattro tavole, troveremo:

`non associativa` Infatti, $(2*4)*4 = 5*4 = 2$, mentre $2*(4*5) = 2*3 = 1$.
`{2 4 4}`

Ora sia G un gruppo d'ordine n e sia H un suo sottogruppo d'ordine h . Allora, per il teorema di Lagrange, esiste q tale che $n = h \cdot q$. Il numero q è detto *indice* di H in G ed è il numero di *lateralì destri* $H \cdot g = \{x \cdot g \mid x \in H\}$ di H in G (ma anche il numero di *lateralì sinistri* $g \cdot H$). Se per ogni $g \in G$ si ha $H \cdot g = g \cdot H$, il sottogruppo H si dice *normale* in G . Ciò accade per esempio se G è *abeliano*, ossia se l'operazione è commutativa. Vediamo che cosa accade se ordiniamo gli elementi di G in modo opportuno. Sia G il gruppo ciclico d'ordine 6, che possiamo rappresentare come $\langle a \mid a^6 = 1 \rangle$. Il suo elemento a^3 ha periodo 2, perché al quadrato dà 1. Allora poniamo $H = \langle a^3 \rangle = \{1, a^3\}$, che è un sottogruppo d'ordine 2.

Consideriamo poi i lateralì $H \cdot a = \{1 \cdot a, a^3 \cdot a\} = \{a, a^4\}$ e $H \cdot a^2 = \{1 \cdot a^2, a^3 \cdot a^2\} = \{a^2, a^5\}$. Ordiniamo G nel modo seguente:

$G = \{1, a^3, a, a^4, a^2, a^5\}$ ed eseguiamo la sostituzione $\begin{pmatrix} 1 & a^3 & a & a^4 & a^2 & a^5 \\ 1 & 2 & 3 & 4 & 5 & 6 \end{pmatrix}$. Allora,

la tavola di moltiplicazione ottenuta è scomponibile in blocchi corrispondenti ai tre lateralì di H : numerati da 1 a 3, questi ultimi danno la tavola di un gruppo d'ordine 3, il *gruppo quoziente* G/H , del quale il blocco corrispondente ad H ($= H \cdot 1$) è l'elemento neutro.

$$\begin{bmatrix} 1 & 2 & 3 & 4 & 5 & 6 \\ 2 & 1 & 4 & 3 & 6 & 5 \\ 3 & 4 & 5 & 6 & 2 & 1 \\ 4 & 3 & 6 & 5 & 1 & 2 \\ 5 & 6 & 2 & 1 & 4 & 3 \\ 6 & 5 & 1 & 2 & 3 & 4 \end{bmatrix} \rightarrow \begin{bmatrix} 1 & 2 & 3 & 4 & 5 & 6 \\ 2 & 1 & 4 & 3 & 6 & 5 \\ 3 & 4 & 5 & 6 & 2 & 1 \\ 4 & 3 & 6 & 5 & 1 & 2 \\ 5 & 6 & 2 & 1 & 4 & 3 \\ 6 & 5 & 1 & 2 & 3 & 4 \end{bmatrix} \rightarrow \begin{bmatrix} 1 & 2 & 3 \\ 2 & 3 & 1 \\ 3 & 1 & 2 \end{bmatrix}$$

Se H non è un sottogruppo normale, non funziona. Proviamo con la prima delle quattro tavole d'ordine 6 dell'esempio precedente. L'elemento 2 ha periodo 2, quindi $H = \{1, 2\}$ è un sottogruppo. Si ha poi $H \circ 3 = \{1 \circ 3, 2 \circ 3\} = \{3, 5\}$, ed infine $H \circ 4 = \{1 \circ 4, 2 \circ 4\} = \{4, 6\}$. Allora $G = \{1, 2, 3, 5, 4, 6\}$. Sostituiamo 4 con 5 e viceversa in tutta la tavola e cerchiamo di scomporre in blocchi, ottenendo:

1	2	3	5	4	6
2	1	4	6	3	5
3	5	6	4	2	1
5	3	2	1	6	4
4	6	5	3	1	2
6	4	1	2	5	3

→

1	2	3	5	4	6
2	1	4	6	3	5
3	5	6	4	2	1
5	3	2	1	6	4
4	6	5	3	1	2
6	4	1	2	5	3

Si vede subito che la divisione in blocchi non funziona, perché le sottomatrici non sono tutte costituite dagli elementi di uno stesso laterale.

A5.4. Matrici e campi. Vediamo un procedimento detto *ampliamento quadratico* di un campo, che fa uso di matrici. Sia K un campo e sia $u \in K$ che non sia il quadrato di altri elementi di K . Consideriamo l'insieme

$$F = \left\{ \begin{bmatrix} a & b \cdot u \\ b & a \end{bmatrix} \mid a, b \in K \right\}.$$

Rispetto alle usuali operazioni con le matrici, si vede facilmente che F è un anello con unità e commutativo. Di più, poiché il determinante è $\Delta = a^2 - u \cdot b^2$ ed u non è un quadrato in K , allora solo la matrice

$$\text{nulla è non invertibile. Escluso questo caso, si ha } \begin{bmatrix} a & b \cdot u \\ b & a \end{bmatrix}^{-1} = \frac{1}{\Delta} \begin{bmatrix} a & -b \cdot u \\ -b & a \end{bmatrix} \in F.$$

Pertanto, F è un campo. Si noti che il sottoinsieme delle matrici “scalari”

$$\left\{ \begin{bmatrix} a & 0 \\ 0 & a \end{bmatrix} \mid a \in K \right\}$$

$$\text{forma a sua volta un campo isomorfo a } K \text{ e quindi identificabile con } K. \text{ Posto } i = \begin{bmatrix} 0 & u \\ 1 & 0 \end{bmatrix}, \text{ si ha } i^2 = \begin{bmatrix} u & 0 \\ 0 & u \end{bmatrix} \equiv u \text{ ed } \begin{bmatrix} a & b \cdot u \\ b & a \end{bmatrix} = a + b \cdot i.$$

NOTA. Nel caso di $K = \mathbf{R}$, $u = -1$, si ottiene un campo isomorfo al campo complesso $\mathbf{C}$. Un piccolo vantaggio di questo procedimento si ha perché le operazioni sono quelle naturali delle matrici, derivanti dall'addizione punto per punto e dalla composizione di applicazioni lineari. Un piccolo svantaggio si ha invece dal fatto che le matrici non è ovvio considerarle come numeri, ma del resto anche le coppie (a, b) di numeri reali hanno lo stesso problema.

A5.5. Matrici e applicazioni lineari. Sia K un campo e siano V, W due K -spazi vettoriali. Sia poi $f: V \rightarrow W$ un'applicazione lineare, ossia un K -omomorfismo tra V e W . Esso è un omomorfismo tra i gruppi additivi $(V, +)$ e $(W, +)$ ed inoltre, per ogni $k \in K, v \in V$, si ha $f(kv) = kf(v)$.

Fissiamo una K -base $\{e_1, \dots, e_n\}$ di V , ed una K -base $\{\varepsilon_1, \dots, \varepsilon_m\}$ di W . Si tratta di elementi tali che ogni altro $v \in V$ si esprime in uno ed un solo modo nella forma

$$v = \sum_{j=1}^n x_j e_j, \text{ con } x_j \in K \forall j. \text{ Analogamente ogni } w \in W \text{ si esprime in uno ed un solo}$$

modo nella forma $w = \sum_{i=1}^m y_i \varepsilon_i$, con $y_i \in K \forall i$. Pertanto, per ogni $i, 1 \leq i \leq m$, si ha

$$f(e_j) = \sum_{i=1}^m a_{ij} \varepsilon_i, \quad a_{ij} \in K. \text{ Allora, per ogni } v = \sum_{j=1}^n x_j e_j \in V \text{ si ha:}$$

$$f(v) = f\left(\sum_{j=1}^n x_j e_j\right) = \sum_{j=1}^n x_j f(e_j) = \sum_{j=1}^n x_j \sum_{i=1}^m a_{ij} \varepsilon_i = \sum_{i=1}^m \left(\sum_{j=1}^n a_{ij} x_j\right) \varepsilon_i = \sum_{i=1}^m y_i \varepsilon_i$$

Si ha cioè $y_i = \sum_{j=1}^n a_{ij} x_j, 1 \leq i \leq m$. Associamo ora all'applicazione f la *matrice*

$A = [a_{ij}], 1 \leq i \leq m, 1 \leq j \leq n$. Rappresentiamo poi i vettori v e w con le *matrici-*

$$\text{colonna } \begin{bmatrix} x_1 \\ x_2 \\ \dots \\ x_n \end{bmatrix} \text{ e } \begin{bmatrix} y_1 \\ y_2 \\ \dots \\ y_m \end{bmatrix}. \text{ Allora, } \begin{bmatrix} y_1 \\ y_2 \\ \dots \\ y_m \end{bmatrix} = \begin{bmatrix} a_{11} & a_{12} & \dots & a_{1n} \\ a_{21} & a_{22} & \dots & a_{2n} \\ \dots & \dots & \dots & \dots \\ a_{m1} & a_{m2} & \dots & a_{mn} \end{bmatrix} \times \begin{bmatrix} x_1 \\ x_2 \\ \dots \\ x_n \end{bmatrix}, \text{ dove il prodotto è}$$

definito "righe per colonne".

La parte di matematica che studia le matrici sotto l'aspetto di maschere algebriche delle applicazioni lineari è l'Algebra Lineare. Fissate le due basi, le matrici di tipo $m \times n$ ad elementi in K e le applicazioni lineari tra V e W formano, con le operazioni opportune, due K -spazi vettoriali isomorfi di dimensione mn . La composizione di applicazioni lineari è lineare e la sua matrice corrisponde al prodotto righe per colonne delle loro matrici. Pertanto, posto $W = V$ e fissata una sola base, abbiamo così un anello di matrici, isomorfo all'anello degli

endomorfismi. Il gruppo delle unità di questo anello è detto *gruppo generale lineare di grado n sul campo K* e denotato con $GL_n(K)$.

Dal corso di Geometria è noto che $A \in GL_n(K)$ se e solo se è *non-singolare*, ossia se e solo se $\det(A) \neq 0_K$. Equivalentemente, $A \in GL_n(K)$ se e solo se il suo rango è uguale ad n . Più precisamente, se $A \in GL_n(K)$, le sue righe formano una base (ordinata) dello spazio vettoriale K^n ; viceversa, incolonnando i vettori di una base ordinata di K^n si ha una matrice invertibile. Pertanto, c'è una biiezione fra $GL_n(K)$ e l'insieme delle basi ordinate di K^n .

Vediamo il caso di un campo finito: sia $K = GF(q)$ un campo finito d'ordine q ; allora $V = K^n$ ha q^n elementi e non è difficile contare le sue basi ordinate. Se infatti vogliamo costruire una base $\{v_1, \dots, v_n\}$ possiamo scegliere v_1 in $V \setminus \{0\}$, cioè in $q^n - 1$ modi. Scelto v_1 , si deve scegliere v_2 indipendente da v_1 , ossia $v_2 \in V \setminus \text{span}(v_1)$, per cui vi sono per v_2 solo $q^n - q$ scelte. Analogamente, dovrà essere $v_3 \in V \setminus \text{span}(v_1, v_2)$, quindi per lui solo $q^n - q^2$ scelte, e così via. Pertanto, in V ci sono $(q^n - 1)(q^n - q) \dots (q^n - q^{n-1})$ basi ordinate, quindi il gruppo $GL_n(K)$, che si denota in questo

caso con $GL_n(q)$, ha in definitiva $\prod_{i=0}^{n-1} (q^n - q^i) = q^{\binom{n}{2}} \prod_{i=1}^n (q^i - 1)$ elementi.

Sia ora $SL_n(K)$ l'insieme delle matrici quadrate di ordine n con determinante 1_K . Per il teorema di Binét, il determinante è un omomorfismo dal gruppo $GL_n(K)$ al gruppo moltiplicativo K^* del campo K . Il suo nucleo è ovviamente $SL_n(K)$, che risulta così un sottogruppo normale. L'immagine è, come si vede facilmente, tutto K^* , per cui, per il I teorema d'isomorfismo, si ha $GL_n(K)/SL_n(K) \cong K^*$. Se $|K| = q$, allora $|K^*| = q-1$, quindi essendo $|GL_n(q)|/|SL_n(q)| = |K^*|$, ne segue:

$$|SL_n(q)| = \frac{1}{q-1} \prod_{i=0}^{n-1} (q^n - q^i).$$

NOTA. Ci sono quindi per esempio 48 matrici quadrate d'ordine 2 non singolari, ad elementi nel campo $\mathbb{Z}_3$, delle quali 24 hanno determinante 1 $\left(= \begin{bmatrix} 1 \\ 1 \end{bmatrix}_3\right)$ e costituiscono il gruppo $SL_2(3)$. Ci sono altri gruppi d'ordine 24, per esempio il gruppo simmetrico S_4 , il gruppo diedrale D_{12} delle simmetrie di un dodecagono regolare, ecc. Saranno *isomorfi* tra loro?

SCHEMA A6. Strutture relazionali.

NOTA. In algebra per le funzioni $f: X \rightarrow Y$ fra due insiemi X ed Y è frequente l'uso della notazione rovesciata: xf (o x^f) anziché $f(x)$. In particolare, questa notazione è usata per le permutazioni in X . Una conseguenza è che la composizione di due funzioni $f: X \rightarrow Y$ e $g: Y \rightarrow Z$ è denotata con fg in luogo di $g \circ f$. Il vantaggio si ha nel moltiplicare permutazioni e nell'usare la notazione esponenziale: $(x^f)^g = x^{fg}$. Lo svantaggio si ha nel dover usare una notazione meno consueta.

Siano A e B due insiemi e sia $\mathfrak{R} \subseteq A \times B$, cioè sia $\mathfrak{R}$ un insieme di coppie ordinate (a, b) , con $a \in A$ e $b \in B$. Come noto, la terna $(A, B, \mathfrak{R})$ si chiama *relazione* fra A e B . E' comunque consuetudine, qualora non ci siano ambiguità, identificare la relazione con l'insieme $\mathfrak{R}$. Si usa scrivere $a\mathfrak{R}b$ anziché $(a, b) \in \mathfrak{R}$.

Una relazione $\mathfrak{R}$ fra due insiemi finiti $A = \{a_1, \dots, a_m\}$ e $B = \{b_1, \dots, b_n\}$ si può rappresentare mediante una matrice $R = [r_{ij}]$, $1 \leq i \leq m$, $1 \leq j \leq n$, detta *matrice d'incidenza della relazione*. Le sue caselle sono definite da $r_{ij} = \begin{cases} 1 & \text{se } a_i \mathfrak{R} b_j \\ 0 & \text{altrimenti} \end{cases}$. Si tratta quindi di una matrice *booleana*.^(*)

Esempio A6.1. Siano $A = B = \{1, 2, 3, 4, 6, 12\}$ e sia $x\mathfrak{R}y$ se x è divisore di y . Allora avremo:

$\mathfrak{R} = \{(1, 1), (1, 2), (1, 3), (1, 4), (1, 6), (1, 12), (2, 2), (2, 4), (2, 6), (2, 12), (3, 3), (3, 6), (3, 12), (4, 4), (4, 12), (6, 6), (6, 12), (12, 12)\}$.

La sua matrice d'incidenza è quadrata d'ordine 6:

$$R = \begin{bmatrix} 1 & 1 & 1 & 1 & 1 & 1 \\ 0 & 1 & 0 & 1 & 1 & 1 \\ 0 & 0 & 1 & 0 & 1 & 1 \\ 0 & 0 & 0 & 1 & 0 & 1 \\ 0 & 0 & 0 & 0 & 1 & 1 \\ 0 & 0 & 0 & 0 & 0 & 1 \end{bmatrix}$$

Si può anche considerare il *diagramma cartesiano* della relazione: in tal caso conviene scegliere l'asse x verticale ed orientato verso il basso, mentre l'asse y sarà orizzontale ed orientato verso destra. Distribuiti gli elementi di A sull'asse verticale

^(*) In altri testi, in cui per le funzioni si usa la notazione classica $f(x)$, la matrice d'incidenza è definita come la trasposta di R .

e quelli di B nell'altro, le coppie $(a_i, b_j) \in \mathfrak{R}$ daranno luogo a punti corrispondenti anche visivamente agli 1 della matrice.

Data la relazione $\mathfrak{R}$ fra A e B, possiamo definire una relazione $\mathfrak{R}^t$ fra B ed A, che chiameremo *trasposta* di $\mathfrak{R}$, definita come l'insieme delle coppie (b, a) tali che $(a, b) \in \mathfrak{R}$. Chiaramente, se A e B sono finiti, la matrice d'incidenza di $\mathfrak{R}^t$ è la trasposta della matrice d'incidenza di $\mathfrak{R}$. Inoltre, la trasposta della trasposta di $\mathfrak{R}$ è $\mathfrak{R}$ stessa.

Date le relazioni $(A, B, \mathfrak{R})$ e $(B, C, \mathcal{S})$, si può definire la *relazione composta* $(A, C, \mathcal{T})$, dove $\mathcal{T} = \mathcal{S} \circ \mathfrak{R}$ è definita da: per ogni $a \in A, c \in C$,

$$a \mathcal{T} c \in \exists b \in B \text{ tale che } a \mathfrak{R} b \text{ e } b \mathcal{S} c.$$

Esercizio A6.2. Si dimostri che la composizione di relazioni è *associativa*, ossia che date le tre relazioni $(A, B, \mathfrak{R}), (B, C, \mathcal{S}), (C, D, \mathcal{T})$ si ha:

$$\mathcal{T} \circ (\mathcal{S} \circ \mathfrak{R}) = (\mathcal{T} \circ \mathcal{S}) \circ \mathfrak{R}$$

Se gli insiemi A, B, C coinvolti sono finiti di ordini rispettivamente m, n, p, possiamo, come detto, rappresentare le relazioni $(A, B, \mathfrak{R})$ e $(B, C, \mathcal{S})$ mediante le matrici R ed S, rispettivamente di dimensioni $m \times n$ e $n \times p$. Utilizziamo per operare sui loro elementi 0 ed 1 le *operazioni booleane* definite come segue:

$$\begin{array}{c|cc} + & 0 & 1 \\ \hline 0 & 0 & 1 \\ 1 & 1 & 1 \end{array}, \begin{array}{c|cc} \times & 0 & 1 \\ \hline 0 & 0 & 0 \\ 1 & 0 & 1 \end{array}$$

Abbiamo allora il seguente risultato:

PROPOSIZIONE A6.3. La matrice d'incidenza T della relazione composta $\mathcal{T} = \mathcal{S} \circ \mathfrak{R}$ è il prodotto (booleano) righe per colonne delle matrici R ed S.

Dimostrazione. Siano $a_i \in A$ e $c_k \in C$, e distinguiamo due casi:

I. Sia $a_i \mathcal{T} c_k$. Allora $t_{ik} = 1$. Ma inoltre esiste $b_j \in B$ tale che $a_i \mathfrak{R} b_j$ e $b_j \mathcal{S} c_k$, per cui $r_{ij} = 1$ e $s_{jk} = 1$. Ne

$$\text{segue } 1 = r_{ij} \times s_{jk} \Rightarrow \sum_{\lambda=1}^n r_{i\lambda} \times s_{\lambda k} = 1 = t_{ik}.$$

II. Sia $(a_i, c_k) \notin \mathcal{T}$. Allora $t_{ik} = 0$. Ma si ha anche, per ogni $b_l \in B$, $(a_i, b_l) \notin \mathfrak{R}$ oppure $(b_l, c_k) \notin \mathcal{S}$, ossia r_{il}

$$= 0 \text{ oppure } s_{lk} = 0. \text{ Ne viene che tutti i prodotti } r_{il} \times s_{lk} \text{ sono nulli e quindi } \sum_{\lambda=1}^n r_{i\lambda} \times s_{\lambda k} = 0 = t_{ik}.$$

Sia X un insieme e sia $\mathfrak{R} \subseteq X \times X$. Secondo la convenzione sopra detta, identificheremo con $\mathfrak{R}$ la relazione $(X, X, \mathfrak{R})$, detta *relazione in X* . La matrice di una tale relazione, nel caso $|X| = n$, sarà quadrata di ordine n .

Tra le relazioni in X c'è l'*identità* id_X , costituita dalle coppie (x, x) , $x \in X$. Nel caso $|X| = n$ la sua matrice sarà denotata al solito con I_n ed avrà la consueta forma: tutti 1 sulla diagonale principale e tutti 0 fuori di essa. C'è anche lo stesso prodotto cartesiano $X \times X$, la cui matrice ha tutti gli elementi uguali ad 1.

Ricordiamo che un *monoide* è una terna $(M, \cdot, 1_M)$ costituita da un insieme M , una operazione binaria associativa $\cdot$ in M ed un elemento $1_M \in M$, che è il suo elemento neutro. Chiamiamo *elemento assorbente* rispetto all'operazione binaria $\cdot$ un elemento $u \in M$ tale che per ogni $x \in M$ si abbia $u \cdot x = x \cdot u = u$. Chiaramente, se esiste è unico. Sia poi M^* l'insieme dei suoi *elementi invertibili*, ossia di quegli $x \in M$ tali che esiste $x^{-1} \in M$ tale che $x \cdot x^{-1} = x^{-1} \cdot x = 1_M$. Si ha ovviamente $1_M \in M^*$. E' immediato provare che se $x, y \in M^*$ allora $(x \cdot y)^{-1} = y^{-1} \cdot x^{-1}$ ed inoltre $(x^{-1})^{-1} = x$, per cui M^* è un gruppo, detto *gruppo delle unità del monoide*. Per esempio, in $(\mathbf{N}, +, 0)$ si ha $\mathbf{N}^* = \{0\}$, mentre in $(\mathbf{Z}, \cdot, 1)$ si ha $\mathbf{Z}^* = \{1, -1\}$. Ciò posto, denotiamo con $\mathcal{R}(X)$ l'insieme delle relazioni in X .

Esercizio A6.4. Sia X un insieme. Si verifichi che, rispetto alla composizione $\circ$ ed alla identità id_X l'insieme $\mathcal{R}(X)$ delle relazioni in X è un monoide. Tale monoide possiede l'elemento assorbente? Chi sono gli elementi invertibili?

Come in ogni monoide, anche in $(\mathcal{R}(X), \circ, \text{id}_X)$ possiamo definire le *potenze* con esponente naturale, ponendo, per ogni $\mathfrak{R} \in \mathcal{R}(X)$ e per ogni $n \in \mathbf{N}$,

$$\begin{cases} \mathfrak{R}^0 = \text{id}_X \\ \mathfrak{R}^{n+1} = \mathfrak{R}^n \circ \mathfrak{R} \end{cases}.$$

Anche in $(\mathcal{R}(X), \circ, \text{id}_X)$ le potenze possiedono le consuete proprietà: per ogni

$$\mathfrak{R} \in \mathcal{R}(X), m, n \in \mathbf{N} \text{ si ha } \begin{cases} \mathfrak{R}^m \circ \mathfrak{R}^n = \mathfrak{R}^{m+n} \\ (\mathfrak{R}^m)^n = \mathfrak{R}^{mn} \end{cases}.$$

Esercizio A6.5. Sia $\mathfrak{R}$ una relazione nell'insieme X e sia $n \in \mathbf{N}$, $n > 1$. Si provi che:

a) Per ogni $x, y \in X$ si ha: $x\mathfrak{R}^n y \Leftrightarrow \exists x_i \in X, 1 \leq i \leq n-1$, tali che, posto $x_0 = x, x_n = y$, si ha $\forall i, 0 \leq i < n, x_i \mathfrak{R} x_{i+1}$.

b) Sia $\mathcal{S}$ un'altra relazione in X , tale che $\mathfrak{R} \subseteq \mathcal{S}$, allora per ogni $n \in \mathbf{N}$ si ha $\mathfrak{R}^n \subseteq \mathcal{S}^n$.

Siano X un insieme ed $\mathfrak{R}$ una relazione in X . La coppia $(X, \mathfrak{R})$ si chiama *struttura relazionale*. Nel caso finito, cioè per $|X| = n$, prende il nome di *grafo orientato*, ed ha una rappresentazione geometrica come segue: si fissino sul piano n punti distinti, si contrassegni ciascuno col nome di un elemento di X , poi si congiunga il punto x_i col punto x_j con una freccia se si ha $x_i \mathfrak{R} x_j$. Le frecce, non necessariamente rettilinee, prendono il nome di *archi*, mentre i punti sono detti *vertici* o *nod*i. Ciascun arco corrisponde dunque ad una coppia di elementi in relazione.

L'esempio qui a lato traduce la relazione dell'esempio A6.1. I pallini detti *cappi*, intorno ad ogni elemento, senza freccia per praticità, significano che l'elemento è in relazione con se stesso. Si noti che il grafo è *planare* (archi senza incroci).

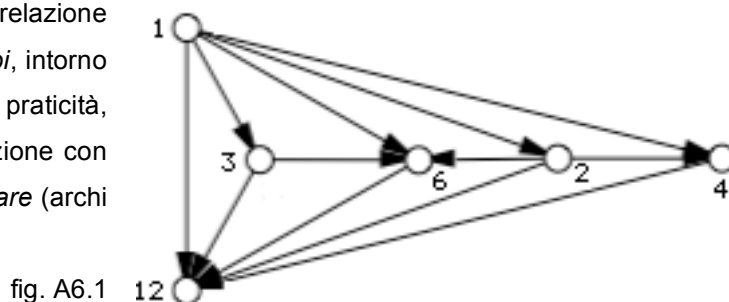


fig. A6.1

Traduciamo ora alcune proprietà delle relazioni in un insieme in termini di operazioni e di inclusione fra le relazioni stesse. Sappiamo che cosa significa per una relazione essere, riflessiva, simmetrica, antisimmetrica o transitiva. Si ha:

PROPOSIZIONE A6.6. Sia data una struttura relazionale $(X, \mathfrak{R})$.

a) $\mathfrak{R}$ è *riflessiva* $\Leftrightarrow \text{id}_X \subseteq \mathfrak{R}$; è *antiriflessiva* $\Leftrightarrow \mathfrak{R} \cap \text{id}_X = \emptyset$

b) $\mathfrak{R}$ è *simmetrica* $\Leftrightarrow \mathfrak{R} = \mathfrak{R}^t$.

c) $\mathfrak{R}$ è *antisimmetrica* $\Leftrightarrow \mathfrak{R} \cap \mathfrak{R}^t \subseteq \text{id}_X$; è *emisimmetrica* $\Leftrightarrow \mathfrak{R} \cap \mathfrak{R}^t = \emptyset$

d) $\mathfrak{R}$ è *transitiva* $\Leftrightarrow \mathfrak{R}^2 \subseteq \mathfrak{R}$.

e) $\mathfrak{R}$ è *totale* $\Leftrightarrow \mathfrak{R} \cup \mathfrak{R}^t = X \times X$. Altrimenti, è parziale.

Una *relazione d'equivalenza* $\mathfrak{R}$ in X è, come noto, una relazione riflessiva, simmetrica e transitiva, ossia tale che $\text{id}_X \subseteq \mathfrak{R}$, $\mathfrak{R} = \mathfrak{R}^t$, $\mathfrak{R}^2 \subseteq \mathfrak{R}$.

Sia data una struttura relazionale $(X, \mathfrak{R})$. La relazione $\overline{\mathfrak{R}} = \bigcup_{n=1}^{\infty} \mathfrak{R}^n$ si chiama

chiusura transitiva della relazione $\mathfrak{R}$. Si ha infatti:

PROPOSIZIONE A6.7. Sia data una struttura relazionale $(X, \mathfrak{R})$ e sia $\overline{\mathfrak{R}}$ la chiusura transitiva di $\mathfrak{R}$.

a) $\overline{\mathfrak{R}}$ è transitiva.

b) $\overline{\mathfrak{R}}$ è inclusa in ogni relazione transitiva contenente $\mathfrak{R}$.

c) Se $\mathfrak{R}$ è simmetrica, la relazione $\sim_{\mathfrak{R}} = \text{id}_X \cup \overline{\mathfrak{R}} = \bigcup_{n=0}^{\infty} \mathfrak{R}^n$ è una relazione

d'equivalenza.

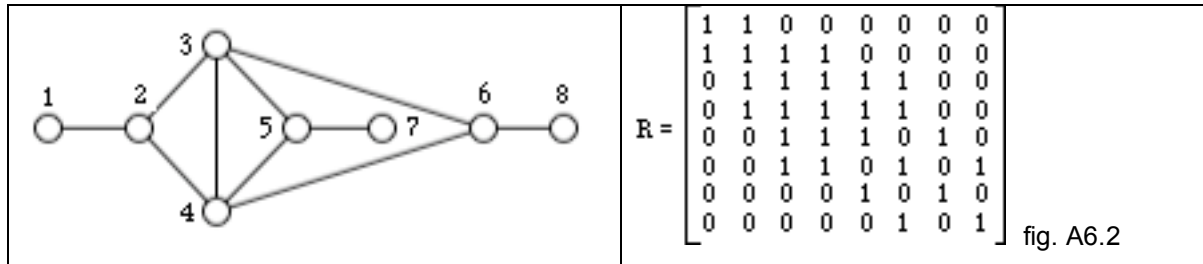
Dimostrazione. a) Siano $a, b, c \in X$, tali che $a \overline{\mathfrak{R}} b$ e $b \overline{\mathfrak{R}} c$. Esistono m, n tali che $a \mathfrak{R}^m b$ e $b \mathfrak{R}^n c$, quindi $a \mathfrak{R}^{m+n} c$. Ne segue $a \overline{\mathfrak{R}} c$, cioè $\overline{\mathfrak{R}}$ è transitiva.

b) Sia $\mathfrak{S}$ una relazione transitiva tale che $\mathfrak{R} \subseteq \mathfrak{S}$. Allora $\mathfrak{R}^2 \subseteq \mathfrak{S}^2 \subseteq \mathfrak{S}$ e, per induzione su n , si ha $\mathfrak{R}^n \subseteq \mathfrak{S}^n \subseteq \mathfrak{S}$ per ogni $n \geq 1$. Pertanto, $\overline{\mathfrak{R}} \subseteq \mathfrak{S}$.

Ad una struttura relazionale simmetrica $(X, \mathfrak{R})$ con $|X| = n$ si può associare un *grafo (non orientato)* costituito da n punti, detti *vertici* o *nodi*, corrispondenti agli elementi di X , e in cui due vertici x_i ed x_j sono congiunti da un segmento (o da un arco di curva non orientato), detto *spigolo*, se e solo se $x_i \mathfrak{R} x_j$. Sostanzialmente, si identificano i due archi orientati opposti che congiungono x_i ed x_j e si toglie l'orientamento.

La matrice R della relazione si chiama in questo caso *matrice d'adiacenza* del grafo, perché il termine *matrice d'incidenza* è usato per indicare la matrice della relazione di appartenenza fra l'insieme dei nodi e quello degli spigoli.

Le classi d'equivalenza della relazione $\sim_{\mathfrak{R}}$ sopra definita sono dette *componenti connesse* del grafo. Se c'è n'è una sola, il grafo si dice *connesso*. Il grafo della figura seguente lo è.



Esercizio A6.8. Sia $(X, \mathfrak{R})$ una struttura relazionale simmetrica finita. Si provi che il suo grafo è connesso se e solo se per ogni $x, y \in X$ esiste $n \in \mathbf{N}$ tale che $x \mathfrak{R}^n y$. In tal caso, chi è $\sim_{\mathfrak{R}}$?

Un grafo non orientato connesso si può riguardare come *spazio metrico* nel modo seguente. Siano x, y due vertici distinti: poiché il grafo è connesso, esiste $n \in \mathbf{N}$ tale che $x \mathfrak{R}^n y$. Poniamo $d(x, y) = \min\{n \in \mathbf{N} \mid x \mathfrak{R}^n y\}$. Poniamo poi $d(x, x) = 0$. Allora la funzione $d: X \rightarrow \mathbf{R}$ appena definita è proprio una metrica nel senso usuale. Chiameremo *diametro* del grafo connesso il numero

$$\Delta = \max\{d(x, y) \mid x, y \in X\}.$$

Il grafo della figura A6.2 è connesso ed ha diametro $\Delta = 4$. Si noti che la sua matrice di adiacenza è simmetrica.

PROPOSIZIONE A6.9. Sia $(X, \mathfrak{R})$ una struttura relazionale simmetrica finita con n elementi e sia R la matrice d'incidenza di $\mathfrak{R}$. Detta E la matrice della relazione d'equivalenza $\sim_{\mathfrak{R}}$, esiste $m \in \mathbf{N}$ tale che $(I_n + R)^m = E$. Inoltre, il minimo $m \geq 1$ per cui vale la precedente uguaglianza è il massimo Δ dei diametri delle componenti connesse del grafo.

Dimostrazione. Per induzione su $k \geq 1$, tenuto conto che si opera con le operazioni booleane, si può dimostrare che $(I_n + R)^k = \sum_{i=0}^k R^i$ e che questa è la matrice di $\bigcup_{i=0}^k \mathfrak{R}^i$. Essendo X finito, esiste m tale che $\mathfrak{R}^{m+1} = \mathfrak{R}^m$, quindi $\bigcup_{i=0}^{m+1} \mathfrak{R}^i = \bigcup_{i=0}^m \mathfrak{R}^i$.

Pertanto, $\sim_{\mathfrak{R}} = \bigcup_{i=0}^m \mathfrak{R}^i$ e di conseguenza $E = (I_n + R)^m$. Sia proprio m il minimo intero ≥ 1 che

soddisfa questa uguaglianza. Siano ora x, y distinti e in una stessa componente connessa, per cui esiste $k \geq 1$ tale che $d(x, y) = k$. Ne segue $x \mathfrak{R}^k y$, e k è il minimo che fa questo servizio. Si ha quindi $k \leq m$. Ma allora anche il diametro δ della componente connessa non supera m , e ciò vale per ogni componente connessa, per cui $\Delta \leq m$. Viceversa, siano x, y distinti e tali che $x \sim_{\mathfrak{R}} y$. Allora x ed y sono in una stessa componente connessa e quindi esiste $k \geq 1$ tale che $d(x, y) = k$ e $x \mathfrak{R}^k y$. Ma si ha:

$$x \mathfrak{R}^k y \Rightarrow x \left(\bigcup_{i=0}^k \mathfrak{R}^i \right) y \Rightarrow x \left(\bigcup_{i=0}^{\delta} \mathfrak{R}^i \right) y \Rightarrow x \left(\bigcup_{i=0}^{\Delta} \mathfrak{R}^i \right) y$$

quindi, per ogni $x, y \in X$ si ha $x \left(\bigcup_{i=0}^m \mathfrak{R}^i \right) y \Leftrightarrow x \sim_{\mathfrak{R}} y \Rightarrow x \left(\bigcup_{i=0}^{\Delta} \mathfrak{R}^i \right) y$ e, quindi, per la minimalità di m ,

si ha $m \leq \Delta$. Ciò prova l'asserto.

Due strutture relazionali $(X, \mathfrak{R})$ e $(Y, \mathcal{S})$ si dicono *isomorfe* se esiste una biiezione $f: X \rightarrow Y$ tale che, per ogni $x, x' \in X$ si ha $x \mathfrak{R} x' \Leftrightarrow f(x) \mathcal{S} f(x')$.

In tal caso, la funzione f si chiama *isomorfismo di strutture relazionali*. Chiaramente, due strutture relazionali isomorfe hanno le stesse proprietà. Inoltre, se i due insiemi X e Y sono finiti e se $X = \{x_1, \dots, x_n\}$, ponendo $y_i = f(x_i)$ per $i = 1, \dots, n$ si ha che le matrici d'incidenza di $\mathfrak{R}$ ed $\mathcal{S}$ sono uguali.

Esercizio A6.10. Siano $(X, \mathfrak{R})$, $(Y, \mathcal{S})$ e $(Z, \mathcal{T})$ tre strutture relazionali e siano $f: X \rightarrow Y$ e $g: Y \rightarrow Z$ due isomorfismi di strutture relazionali.

a) Si dimostri che $g \circ f$ è un isomorfismo tra $(X, \mathfrak{R})$ e $(Z, \mathcal{T})$.

b) Si dimostri che f^{-1} è un isomorfismo tra $(Y, \mathcal{S})$ e $(X, \mathfrak{R})$.

Gli isomorfismi tra $(X, \mathfrak{R})$ e se stesso si chiamano *automorfismi* e formano un gruppo rispetto alla composizione. Esso si denota con $\text{Aut}(X, \mathfrak{R})$.

Tra le relazioni in X ci sono anche le funzioni $f: X \rightarrow X$. Si osservi che l'operazione $\circ$ definita tra le relazioni si riduce, nel caso delle funzioni, alla composizione. Poiché la composta di due funzioni è ancora una funzione e l'identità id_X è una funzione, le funzioni formano a loro volta un monoide, il

monoide $(X^X, \circ, \text{id}_X)$ delle funzioni di X . Se f è biiettiva, allora f^t non è altro che l'inversa f^{-1} di f . Se X è finito, come ogni altra relazione in X anche ogni funzione ha una matrice d'incidenza, che avrà in ogni riga uno ed un solo 1. Se f è biiettiva, cioè è una permutazione in X , anche in ogni colonna ci sarà uno ed un solo 1. Una tal matrice è detta *matrice di permutazione*. La sua matrice inversa coinciderà quindi con la sua trasposta. Ovviamente, se $|X| = n$ abbiamo $n!$ matrici di permutazione, che costituiscono un gruppo isomorfo al gruppo simmetrico S_n . Gli elementi di Aut sono particolari permutazioni, cioè $\text{Aut}(X, \mathfrak{R})$ è un sottogruppo di S_n . Si ha:

PROPOSIZIONE A6.11. Siano X un insieme con n elementi, $\mathfrak{R}$ ed $\mathcal{S}$ due relazioni su X , R ed S le rispettive matrici d'incidenza. Sia poi $f: X \rightarrow X$ e sia F la sua matrice.

a) f è un isomorfismo tra $(X, \mathcal{S})$ e $(X, \mathfrak{R})$ se e solo se $S = F^{-1} \times R \times F$.

b) Si ha $f \in \text{Aut}(X, \mathfrak{R})$ se e solo se $R \times F = F \times R$.

Dimostrazione. a) Sia $S = F^{-1} \times R \times F$ e siano y, y' tali che $y \mathfrak{R} y'$. Esistono x, x' tali che $y = f(x), y' = f(x')$. Usando anche per le funzioni la notazione generale per le relazioni, ossia $x f y$ in luogo di $y = f(x)$, segue:

$$x f y, y \mathcal{S} y', y' f^{-1} x' \Rightarrow x (f \circ \mathcal{S} \circ f^{-1}) x'$$

La matrice di $(f \circ \mathcal{S} \circ f^{-1})$ è $F \times S \times F^{-1} = R$, quindi segue $x \mathfrak{R} x'$. Inversamente, poiché i passaggi precedenti sono tutti reversibili, da $x \mathfrak{R} x'$ segue $y \mathcal{S} y'$, per cui f è un isomorfismo. Viceversa, sia f un isomorfismo. Per ogni $x, x' \in X$ tali che $x \mathfrak{R} x'$, posto $y = f(x)$ e $y' = f(x')$ si ha $y \mathcal{S} y'$, ossia

$$x \mathfrak{R} x' \Rightarrow x f y, y \mathcal{S} y', y' f^{-1} x' \Rightarrow x (f \circ \mathcal{S} \circ f^{-1}) x'.$$

Ma vale anche l'implicazione inversa, sicché le due relazioni $\mathfrak{R}$ e $(f \circ \mathcal{S} \circ f^{-1})$ coincidono. Ma allora hanno la stessa matrice, ossia $R = F \times S \times F^{-1}$, da cui l'asserto.

b) Basta porre $\mathcal{S} = \mathfrak{R}$ nella precedente proposizione. Si avrà

$$f \in \text{Aut}(X, \mathfrak{R}) \Leftrightarrow R = F^{-1} \times R \times F \Leftrightarrow F \times R = R \times F.$$

Esercizio A6.12. Date le strutture relazionali $(X, \mathfrak{R}), (Y, \mathcal{S})$, chiamiamo *prodotto diretto* $(X \times Y, \mathcal{T})$ delle strutture date la struttura relazionale definita nel prodotto cartesiano $X \times Y$ ponendo::

$$(x, t) \mathcal{T} (x', y') \Leftrightarrow x \mathfrak{R} x' \text{ e } y \mathcal{S} y'.$$

Si provi che se $\mathfrak{R}$ ed $\mathcal{S}$ sono rispettivamente riflessive, simmetriche, antisimmetriche, transitive anche $\mathcal{S}$ lo è.

APPENDICE. Dati due insiemi disgiunti X ed Y , una relazione $(X, Y, \mathfrak{R})$ non vuota viene anche detta *struttura d'incidenza*. Gli elementi $x \in X$ sono detti *punti*, gli $y \in Y$ sono detti *blocchi* e la relazione $\mathfrak{R}$ viene detta *incidenza* fra punti e blocchi. Possiamo pensare di associare ad ogni $y \in Y$ il sottoinsieme B_y di X formato dagli elementi in relazione con y . Nasce così un'applicazione $b: Y \rightarrow \wp(X)$, $b: y \rightarrow B_y$, tale che:

$$x \mathfrak{R} y \Leftrightarrow x \in b(y)$$

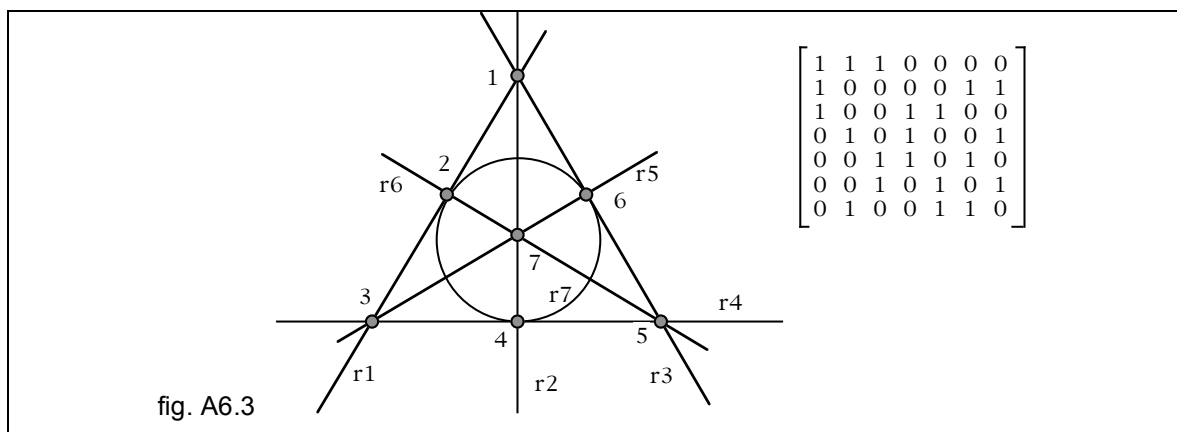
Se tale applicazione è iniettiva, possiamo identificare ogni $y \in Y$ con la sua immagine, cioè considerare Y come sottoinsieme di $\wp(X)$ ed $\mathfrak{R}$ con $\in$.

Questo è il punto di partenza della *Geometria Combinatoria*. Essa prende in considerazione principalmente i casi in cui ci sono delle regolarità, ossia per esempio ogni punto è incidente allo stesso numero di blocchi (*struttura regolare*), oppure ogni blocco è incidente allo stesso numero di punti (*struttura uniforme*), oppure entrambe le condizioni (*configurazione tattica*). Se inoltre ogni stesso numero t di punti è incidente contemporaneamente ad uno stesso numero l di blocchi, abbiamo un *disegno*. Se in particolare ogni coppia di punti distinti ($t = 2$) è incidente ad uno ed un solo blocco ($l = 1$), abbiamo un *piano proiettivo*.

Esempio A6.13. Consideriamo l'insieme X costituito dai vertici, i punti medi dei lati ed il centro di un dato triangolo equilatero. Come Y prendiamo l'insieme costituito dai lati, le mediane e il cerchio inscritto nel triangolo dato. Allora $|X| = |Y| = 7$. Sia poi $\mathfrak{R}$ la usuale appartenenza: otteniamo una struttura d'incidenza in cui:

- due punti distinti sono incidenti ad uno ed un solo blocco;
- due blocchi distinti si incontrano sempre in uno ed un solo punto;
- ogni blocco è incidente a 3 punti;
- ogni punto è incidente a 3 blocchi.

Questa struttura d'incidenza è il più piccolo piano proiettivo, ed è detto *piano di Fano* (v. fig. A6.3.)



SCHEMA A7. Insiemi ordinati.

Siano X un insieme (non vuoto) ed $\mathfrak{R}$ una relazione in X . Diremo che $\mathfrak{R}$ è una relazione d'ordine se è riflessiva, antisimmetrica e transitiva, ossia, usando le nozioni introdotte nella scheda precedente, se

$$\text{I. } \text{id}_X \subseteq \mathfrak{R}$$

$$\text{II. } \mathfrak{R} \cap \mathfrak{R}^t \subseteq \text{id}_X$$

$$\text{III. } \mathfrak{R}^2 \subseteq \mathfrak{R}.$$

Due elementi di $(X, \mathfrak{R})$ si dicono *confrontabili* se si ha $x\mathfrak{R}y$ oppure $y\mathfrak{R}x$. Li diremo *inconfrontabili* in caso contrario. In una relazione d'ordine la proprietà antisimmetrica assicura che se due elementi distinti sono confrontabili, lo sono in un unico modo.

Diremo che $\mathfrak{R}$ è un ordine *totale* se $\mathfrak{R} \cup \mathfrak{R}^t = X \times X$, ossia se tutti gli elementi sono a due a due confrontabili. Altrimenti, è un ordine parziale.

Come noto, gli ordinamenti naturali di $\mathbf{N}$, $\mathbf{Z}$, $\mathbf{Q}$, $\mathbf{R}$ sono totali, e così pure è totale l'ordine alfabetico (o *lessicografico*) fra le parole della lingua italiana. Esempi di ordini non totali, in cui esistono coppie di elementi inconfrontabili, sono l'*inclusione* nell'insieme delle parti di un insieme con almeno due elementi, o la relazione $\mid$ ("è divisore di") in $\mathbf{N}$.

Se $\mathfrak{R}$ è una relazione d'ordine in X , la struttura relazionale $(X, \mathfrak{R})$ verrà detta *insieme (parzialmente) ordinato*, o semplicemente *poset*. Usualmente, per le relazioni d'ordine si usa il simbolo $\leq$. La relazione $<$, detta *ordine stretto*, si definisce ponendo:

$$x < y \text{ se } x \leq y \text{ e } x \neq y.$$

La relazione trasposta di $\leq$ è a sua volta una relazione d'ordine, viene denotata con $\geq$ e viene detta *ordine duale*.

Esercizio A7.1. Sia $(X, \leq)$ un poset e sia $<$ l'ordine stretto associato. Si provi che la relazione $<$ è transitiva e disgiunta dall'identità e dalla sua trasposta.

Inversamente, data una struttura relazionale $(X, \mathfrak{R})$ con $\mathfrak{R}$ transitiva, disgiunta dall'identità e dalla sua trasposta $\mathfrak{R}^t$, posto

$$x \leq y \text{ se } x = y \text{ oppure } x \Re y$$

allora $\leq$ è un ordine in X , di cui $\Re$ è l'ordine stretto associato.

Sia $(X, \leq)$ un poset e siano $x, y \in X$, $x < y$. Una *catena* è un sottoinsieme di X totalmente ordinato. Una catena finita necessariamente ha due *estremi* x ed y :

$$x = x_0 < x_1 < \dots < x_r = y$$

Il numero r si chiama *lunghezza* della catena. Se le catene di X sono tutte finite e di lunghezza non superiore ad un dato numero $m \in \mathbf{N}$, la massima lunghezza delle catene di X è detta *lunghezza* di $(X, \leq)$. Se un tal numero m non esiste, si dice che $(X, \leq)$ ha lunghezza infinita.

Al contrario, una *anticatena* è un sottoinsieme di elementi a due a due inconfrontabili. Se esiste $n \in \mathbf{N}$ tale che ogni anticatena abbia al più n elementi, il numero massimo di elementi delle anticatene è detto *numero di Dilworth* del poset. Chiaramente, un ordine è totale se e solo se tutte le anticatene hanno un solo elemento ciascuna, cioè se il numero di Dilworth del poset è 1.

Gli ordinamenti naturali degli insiemi numerici $\mathbf{Z}$, $\mathbf{Q}$, $\mathbf{R}$ sono ordini totali, ma presentano comunque delle differenze importanti. A titolo di ripasso, ricordiamo alcune di queste. Sia $(X, \leq)$ un insieme ordinato. L'ordine si dice *denso* se per ogni $x, y \in X$, con $x < y$, esiste $z \in X$ tale che $x < z < y$. Come noto dai corsi di Algebra e di Analisi Matematica, gli ordinamenti di $\mathbf{Q}$ e di $\mathbf{R}$ hanno questa proprietà, mentre quello di $\mathbf{Z}$ non ce l'ha. In ogni caso, è immediato verificare che se l'ordine è denso, allora X è infinito.

Sia $(X, \leq)$ un insieme ordinato. Siano poi $x, y \in X$ con $x < y$. L'insieme:

$$[x, y] = \{z \in X \mid x \leq z \leq y\}$$

è detto *intervallo chiuso di estremi x ed y* . L'ordine si dice *localmente finito* se per ogni $x, y \in X$ con $x < y$, l'intervallo $[x, y]$ è finito. Si dice *discreto* se per ogni $x, y \in X$ con $x < y$, l'intervallo $[x, y]$ ha lunghezza finita. Un ordine localmente finito è discreto, ma non viceversa. I poset finiti hanno queste proprietà, ma anche $(\mathbf{Z}, \leq)$ ce le ha, pur essendo infinito, mentre un ordine denso non è localmente finito né discreto.

Sia $(X, \leq)$ un insieme ordinato e siano $x, y \in X$, $x < y$. Si dice che x è *coperto da* y (e si scrive $x \prec y$) se l'intervallo $[x, y]$ contiene solo x ed y , quindi se non esiste alcun altro $z \in X$ tale che $x < z < y$. La relazione $\prec$ così definita viene detta *relazione di copertura* associata all'ordine dato. Chiaramente, se l'ordine è denso, tale relazione è vuota. Si ha:

PROPOSIZIONE A7.2. Sia $(X, \leq)$ un insieme ordinato discreto. Allora la relazione di ordine stretto associata $<$ è la chiusura transitiva $\succsim$ della relazione di copertura $\prec$.

Dimostrazione. Siano $x, y \in X$ tali che $x < y$. Poiché l'ordine è discreto, le catene che li congiungono hanno lunghezza finita non superiore ad un certo k . Esiste pertanto una catena $x_0 = x < x_1 < \dots < x_m = y$ di lunghezza massima $m \leq k$. Chiaramente, per ogni $i = 0, 1, \dots, m-1$, fra x_i ed x_{i+1} non si possono inserire altri elementi di X , perché la catena si allungherebbe. Ne segue che ogni x_i è coperto da x_{i+1} e dunque per ogni i si ha $x \prec^i x_i$. Ma allora si ha $x \prec^m y$, e ciò prova che se $x < y$ allora $x \succsim y$. Essendo però $<$ transitiva, vale anche l'implicazione inversa e quindi l'uguaglianza di $<$ e $\succsim$.

OSSERVAZIONE. Sia $(X, \leq)$ un insieme ordinato finito. L'ordine si rappresenta graficamente mediante il grafo della copertura, detto *diagramma di Hasse*, nel quale per semplicità si tracciano segmenti invece delle frecce, con la convenzione che l'elemento che copre l'altro stia più in alto.

ESERCIZIO A7.3. Sia $(X, \leq)$ un insieme ordinato finito e sia $\prec$ la relazione di copertura. Dette rispettivamente E , I e C le matrici di $\leq$, dell'identità e della copertura, verificare che si ha $E = (I+C)^m$, dove m è la massima lunghezza delle catene ascendenti di X .

Dato un poset $(X, \leq)$ ed un suo sottoinsieme non vuoto A , si chiama *minimo* di A un elemento $a_0 \in A$ tale che per ogni $a \in A$ si abbia $a_0 \leq a$. Chiaramente, se esiste è unico, per la proprietà di antisimmetria della relazione d'ordine, e viene denotato con $\min(A)$. Analogamente, si chiama *massimo* di A un elemento $a_1 \in A$ tale che per ogni $a \in A$ si abbia $a \leq a_1$. Anche il massimo, se esiste, è unico e viene denotato con $\max(A)$.

Se il poset $(X, \leq)$ possiede il minimo, questo viene denotato con 0_X (o con 0, se non ci sono ambiguità). L'eventuale massimo viene denotato con 1_X (o con 1).

Un elemento $b \in X$ si dice *maggiorante* (o *confine superiore*) di A se per ogni $a \in A$ si ha $a \leq b$. Se l'insieme dei maggioranti di A è vuoto, A si dice *superiormente illimitato*; se non è vuoto, A si dice *superiormente limitato*; se non è vuoto ed ha minimo, questo si chiama *estremo superiore* (o *supremo*) di A e viene denotato con $\sup(A)$.

Analogamente, un elemento $c \in X$ si dice *minorante* (o *confine inferiore*) di A se per ogni $a \in A$ si ha $c \leq a$. Se l'insieme dei minoranti non è vuoto ed ha massimo, questo si chiama *estremo inferiore* (o *infimo*) di A e viene denotato con $\inf(A)$.

Se $\inf(A) \in A$ oppure se esiste $\min(A)$ allora si ha $\inf(A) = \min(A)$. Analoga situazione si ha per $\sup(A)$ e $\max(A)$.

L'insieme ordinato $(X, \leq)$ si dice *completo* se ogni sottoinsieme non vuoto che possieda maggioranti ha anche l'estremo superiore.

Esercizio A7.4. Sia $(X, \leq)$ un insieme ordinato. Si provi che è completo se e solo se ogni sottoinsieme non vuoto che possieda minoranti ha anche l'estremo inferiore.

Esercizio A7.5. Sia dato un poset $(X, \leq)$ con il minimo 0_X ed il massimo 1_X .

- a) Si provi che è localmente finito se e solo se è finito.
- b) Si provi che è completo se e solo se ogni sottoinsieme non vuoto ha l'estremo superiore.

Dai corsi di Algebra e di Analisi è noto che $(\mathbf{R}, \leq)$ e $(\mathbf{Z}, \leq)$ sono completi, mentre $(\mathbf{Q}, \leq)$ non lo è, perché il sottoinsieme $A = \{x \in \mathbf{Q} \mid x \geq 1, x^2 \leq 2\}$ non possiede l'estremo superiore in $(\mathbf{Q}, \leq)$. E' poi facile provare che per ogni insieme X , $(\emptyset(X), \subseteq)$ è completo.

Un poset $(X, \leq)$ si chiama *reticolo* se per ogni coppia $\{x, y\}$ di suoi elementi esistono $\sup\{x, y\}$ ed $\inf\{x, y\}$. Chiaramente, ogni insieme totalmente ordinato ed ogni insieme ordinato completo sono reticoli.

Sia $(X, \leq)$ un insieme ordinato. Un elemento $m \in X$ si dice *massimale* se per ogni $x \in X$, se $m \leq x$ allora $x = m$. Chiaramente, se X ha massimo allora questo è l'unico elemento massimale, ma in generale gli elementi massimali non è detto che esistano né che ce ne sia uno solo. Se un insieme ordinato è finito, ovviamente ha elementi massimali, ma per esempio, $(\mathbb{Z}, \leq)$ non ne ha. L'insieme degli interi compresi fra 1 e 10, ordinato mediante la divisibilità, ha come elementi massimali 6, 7, 8, 9, 10. La seguente proposizione stabilisce una ben nota condizione sufficiente per l'esistenza di elementi massimali:

PROPOSIZIONE A7.6 ("Assioma" o Lemma di Kuratowski-Zorn). Sia $(X, \leq)$ un insieme ordinato. Se ogni catena non vuota di X è limitata superiormente, allora in X esistono elementi massimali.

E' ben noto il problema relativo al Lemma di Zorn: non dipende dagli altri assiomi base della Teoria degli Insiemi, ma è equivalente ad altre proposizioni quali l'*assioma di Zermelo* (o *assioma di scelta*), il quale asserisce la possibilità, data una comunque grande famiglia di insiemi non vuoti $\mathfrak{S}$, di definire una funzione $f : \mathfrak{S} \rightarrow \bigcup_{A \in \mathfrak{S}} A$ tale che per ogni $A \in \mathfrak{S}$ si abbia $f(A) \in A$. In particolare, esso dà la possibilità di scegliere un rappresentante da ogni classe di equivalenza $\sim$ in un insieme X .

Una ulteriore proposizione equivalente al Lemma di Zorn ed all'assioma di scelta è l'assioma del buon ordinamento. Un insieme ordinato $(X, \leq)$ si dice *bene ordinato* se ogni suo sottoinsieme non vuoto possiede il minimo. E' ben noto che $(\mathbb{N}, \leq)$ ha questa proprietà (*principio del minimo*). Si ha:

PROPOSIZIONE A7.7. ("Assioma" del buon ordinamento). In ogni insieme non vuoto X è possibile definire una relazione d'ordine $\leq$ in modo che $(X, \leq)$ sia bene ordinato.

Le catene in $(X, \leq)$ si possono ordinare nel modo seguente: date due catene Σ e Σ' , diremo che Σ' è *più fine* di Σ se ogni termine di Σ lo è anche di Σ' . Sostanzialmente, questo ordinamento coincide con l'inclusione in $\wp(\wp(X))$. Ha

senso quindi parlare di catene massimali, cioè che non possono essere raffinate. Anche la seguente proposizione è equivalente all'assioma di scelta:

PROPOSIZIONE A7.8. ("Assioma" delle catene o di Hausdorff-Birchhoff) Ogni catena di un poset $(X, \leq)$ è inclusa in almeno una catena massimale.

Sia $(X, \leq)$ un poset. Una *catena ascendente* di origine $x_0 \in X$ è una successione:

$$x_0 \leq x_1 \leq x_2 \leq \dots \leq x_n \leq \dots$$

di elementi di X . Diremo che è *finita* se esiste $n \in \mathbf{N}$ tale che per ogni $k > n$ si ha $x_k = x_n$. Similmente si definiscono le catene discendenti e le catene discendenti finite. Diremo che $(X, \leq)$ soddisfa la *condizione sulle catene ascendenti (a.c.c.)* se ogni catena ascendente è finita. Similmente si definisce la condizione sulle catene discendenti (d.c.c.).

Queste condizioni sono importanti in vari settori dell'Algebra. Per esempio, consideriamo un dominio d'integrità A ed ordiniamo i suoi ideali bilateri mediante l'inclusione. Otteniamo un poset $(\mathcal{A}(A), \subseteq)$ con minimo $\{0_A\}$ e massimo A . L'anello A si dice *noetheriano* se il poset $(\mathcal{A}(A), \subseteq)$ soddisfa (a.c.c.), *artiniano* se soddisfa (d.c.c.). La prima situazione è equivalente al fatto che ogni ideale abbia un numero finito di generatori, e si verifica per esempio se A ha gli ideali principali o se è fattoriale. Un celebre teorema di Hilbert afferma che se A è noetheriano anche l'anello dei polinomi $A[x]$ lo è. Queste nozioni sono alla base della Geometria Algebrica.

Esercizio A7.9. Si provi che se un poset $(X, \leq)$ soddisfa la condizione sulle catene ascendenti (a.c.c.), ogni $x \in X$ è contenuto in un elemento massimale.

Esercizio A7.10. Si provi che il poset $(\mathbf{N}, |)$ soddisfa la (d.c.c) ma non la (a.c.c.).

Sia $(X, \leq)$ discreto. Per ogni coppia di elementi $x, y \in X$, con $x < y$, esistono catene massimali di estremi x ed y , ma in generale queste non avranno la stessa lunghezza. Diremo che il poset soddisfa la *condizione di Jordan-Dedekind (J.D.c.)* se per ogni $x, y \in X$, $x < y$, tutte le catene massimali di estremi x ed y hanno la stessa

lunghezza. Un esempio è dato dagli insiemi totalmente ordinati localmente finiti, in cui la catena è unica. Ma anche $(\wp(X), \subseteq)$ (se X è finito) ed $(\mathbf{N} \setminus \{0\}, |)$ soddisfano questa condizione. Esistono però dei reticoli che non la soddisfano, ed un esempio è il cosiddetto reticolo N_5 , visto altrove.

Sia ora $(X, \leq)$ localmente finito, soddisfacente la condizione di Jordan-Dedekind e dotato di minimo. Per ogni $x \in X$, le lunghezze delle catene massimali congiungenti 0_X con x hanno tutte la stessa lunghezza, che chiameremo *altezza* (o *rango* o anche *dimensione*) di x e che denoteremo con $h(x)$. Chiaramente, per ogni $x, y \in X$ si ha:

- i) $h(x) \in \mathbf{N}$
- ii) $h(x) = 0 \Leftrightarrow x = 0_X$
- iii) $x < y \Leftrightarrow x \leq y$ e $h(y) = h(x) + 1$

Se $(X, \leq)$ è dotato anche di massimo, $h(1_X)$ è uguale alla sua lunghezza.

Osservazione. Una funzione $d: X \rightarrow \mathbf{Z}$ che soddisfi la condizione iii) è detta *funzione dimensionale* sul poset. Il concetto di dimensione in uno spazio vettoriale è reinterpretabile da questo punto di vista: il poset è quello dei sottospazi vettoriali, il minimo è il sottospazio nullo 0 ; ed un sottospazio ha dimensione k se e solo se ha altezza k rispetto al minimo: le rette per 0 hanno dimensione 1 perché coprono 0 , i piani per 0 hanno dimensione 2 perché coprono le rette per 0 contenute in essi, ecc.

Dati due insiemi ordinati $(X, \leq)$ ed $(Y, \subseteq)$, un *omomorfismo d'ordine* (o funzione monotona non decrescente) è un'applicazione $f: X \rightarrow Y$ tale che:

$$x \leq x' \Rightarrow f(x) \subseteq f(x')$$

Un *antiomomorfismo d'ordine* (o funzione monotona non crescente) è invece tale che $x \leq x' \Leftrightarrow f(x) \supseteq f(x')$. Si possono anche definire le funzioni strettamente monotone (crescenti o decrescenti), che conservano (o rispettivamente invertono) l'ordine stretto nei due poset. L'interesse per queste funzioni si ha principalmente nell'Analisi Matematica reale, dove le funzioni monotone $f: A(\subseteq \mathbf{R}) \rightarrow \mathbf{R}$ sono oggetto di studio per le loro proprietà quali l'integrabilità sugli intervalli chiusi e limitati. Inoltre, la monotonicità è legata al segno della derivata prima. In campo algebrico sono naturalmente più interessanti

gli isomorfismi, che sono casi particolari degli isomorfismi tra strutture relazionali, esaminati nel capitolo precedente.

Un *isomorfismo d'ordine* tra due insiemi ordinati $(X, \leq)$ ed $(Y, \subseteq)$ è una biiezione $f: X \rightarrow Y$, tale che per ogni $x, x' \in X$ si ha:

$$x \leq x' \Leftrightarrow f(x) \subseteq f(x')$$

Si possono definire anche gli *antiisomorfismi d'ordine*, come gli isomorfismi d'ordine tra $(X, \leq)$ ed $(Y, \subseteq)$. Per esempio, considerando i due insiemi ordinati $(\mathbf{R}, \leq)$ e $(\mathbf{R}^+, \leq)$, per ogni $a \in \mathbf{R}$, $a > 1$ funzione $x \rightarrow a^x$ è un isomorfismo d'ordine, mentre se $0 < a < 1$ la funzione $x \rightarrow a^x$ è un antiisomorfismo d'ordine.

Due insiemi ordinati isomorfi hanno ovviamente le stesse proprietà: se uno è denso, localmente finito, completo, reticolo, ecc. lo è anche l'altro e viceversa. Inoltre, nel caso finito, i due ordini si possono rappresentare con la stessa matrice d'incidenza e con lo stesso diagramma di Hasse, e viceversa.

Per esempio, due insiemi totalmente ordinati finiti con lo stesso numero di elementi sono isomorfi (ed anche antiisomorfi). Non è vero per gli insiemi infiniti: $(\mathbf{Z}, \leq)$ e $(\mathbf{Q}, \leq)$ sono entrambi totalmente ordinati e numerabili, ma il primo è localmente finito ed il secondo è denso.

Un poset $(X, \leq)$ che sia isomorfo al suo duale $(X, \geq)$ è detto *autoduale*. Per esempio, per ogni insieme X , $(\wp(X), \subseteq)$ è autoduale: basta far corrispondere ad ogni $A \in \wp(X)$ il suo complementare $A' = X \setminus A$. Anche ogni insieme totalmente ordinato finito è autoduale. Non è vero nel caso infinito: $(\mathbf{N}, \leq)$ ha minimo, 0, mentre $(\mathbf{N}, \geq)$ non ce l'ha (ossia, non esiste un elemento $u \in \mathbf{N}$ tale che $u \geq n$ per ogni $n \in \mathbf{N}$.)

Gli isomorfismi di un poset $(X, \leq)$ con se stesso formano, come sempre, un gruppo, il gruppo $\text{Aut}(X, \leq)$ detto *l'automorfo di* $(X, \leq)$.

Poiché la trasposta di una relazione d'ordine $\leq$ è ancora una relazione d'ordine, è possibile formulare il cosiddetto *principio di dualità dei poset*. Sia dato un enunciato P sui poset: il suo *duale* P^t è ottenuto scambiando dappertutto il simbolo $\leq$ col simbolo trasposto $\geq$. In tal modo, per esempio, si scambiano fra loro

le nozioni di maggiorante e minorante, massimo e minimo, sup ed inf, catena ascendente e catena discendente. Tale procedimento è detto *dualizzazione* dell'enunciato P.

Si osservi che la nozione duale di intervallo $[x, y]$ non cambia l'insieme:

$$[x, y] = \{z \in X \mid x \leq z \leq y\} = \{z \in X \mid y \geq z \geq x\} = [y, x]^t.$$

E' un esempio di nozione *autoduale*. Si ha:

PROPOSIZIONE A7.11 (Principio di dualità per i poset). Se un enunciato P della teoria dei poset è vero, allora è vero anche il suo duale P^t , e la dimostrazione di P^t si ottiene da quella di P per dualizzazione.

Per finire osserviamo che dati due poset, esistono vari modi per fabbricarne altri. Per esempio, si ha:

Esercizio A7.12. Dati due insiemi ordinati $(X, \leq)$ ed $(Y, \subseteq)$, nell'insieme $X \times Y$ si ponga:

$$(x, y) \Re (x', y') \text{ se } x \leq x' \text{ e } y \subseteq y'.$$

Si provi che $(X \times Y, \Re)$ è un insieme ordinato (*prodotto diretto* dei due poset dati), ma non totalmente.

Esercizio A7.13. Dati due insiemi ordinati $(X, \leq)$ ed $(Y, \subseteq)$, nell'insieme $X \times Y$ si ponga:

$$(x, y) \Re (x', y') \text{ se } x \leq x' \text{ oppure se } x = y \text{ e } y \subseteq y'.$$

Si provi che $(X \times Y, \Re)$ è un insieme totalmente ordinato (ordine *lessicografico*).

SCHEMA A8. – La divisibilità

Dati $a, b \in \mathbf{N}$, si dice che a divide b (o che b è multiplo di a) e esiste $q \in \mathbf{N}$ tale che $b = aq$. Si usa scrivere $a \mid b$. Ne segue che l'unità 1 di $\mathbf{N}$ divide ogni altro numero naturale, mentre lo zero è multiplo di ogni altro numero naturale. Ci proponiamo qui di estendere questa nozione ad altri contesti algebrici.

Per estendere questa definizione ad altri contesti algebrici occorrono un insieme $X \neq \emptyset$ ed una operazione $\cdot$ su X , ossia un *gruppoide* $(X, \cdot)$. La definizione è allora trasferibile così com'è:

per ogni $a, b \in X$, $a \mid b$ se esiste $q \in X$ tale che $b = a \cdot q$

Perché non pretendere invece: $a \mid b$ se esiste $q \in X$ tale che $b = q \cdot a$? Se l'operazione $\cdot$ non è commutativa, ovviamente le cose si complicano. Perciò:

PRIMA RIDUZIONE: $(X, \cdot)$ gruppoide *commutativo*.

Se $(X, \cdot)$ è un gruppo, per ogni $a, b \in X$ esiste sempre $q \in X$ tale che $b = aq$ e si ha quindi $a \mid b$ per ogni $a, b \in X$. Perciò:

SECONDA RIDUZIONE: $(X, \cdot)$ gruppoide *commutativo* ma non un gruppo.

La relazione "divide" ha in $\mathbf{N}$ varie ben note proprietà, sia algebriche sia "relazionali". Vediamone un elenco:

Alcune proprietà "relazionali": siano $a, b, c \in X$:

- i. Riflessiva: $a \mid a$
- ii. Antisimmetrica: se $a \mid b$ e $b \mid a$ allora $a = b$
- iii. Transitiva: se $a \mid b$ e $b \mid c$ allora $a \mid c$

In altre parole, la relazione "divide" è una relazione d'ordine in $\mathbf{N}$, cioè $(\mathbf{N}, \mid)$ è un *insieme ordinato*. Poiché esistono coppie di numeri nessuno dei quali divide l'altro allora tale ordine non è totale. L'unità 1 è il minimo e lo zero è il massimo di $(\mathbf{N}, \mid)$.

Osservazione. L'analogia definizione "additiva": per ogni $a, b \in \mathbf{N}$,

$$a \leq b \text{ se esiste } d \in \mathbf{N} \text{ tale che } b = a + d$$

produce invece un ordine totale in $\mathbf{N}$. In esso, 0 è il minimo e non c'è il massimo.

Vediamo quali proprietà deve avere il gruppoide commutativo $(X, \cdot)$ perché la relazione $\mid$ abbia qui proprietà simili a quelle della relazione "divide" in $\mathbf{N}$.

I) L'*elemento neutro* 1_X assicura la proprietà riflessiva: $a = a \cdot 1_X \Rightarrow a \mid a$

II) La *proprietà associativa* assicura la proprietà transitiva di $\mid$, infatti

$$b = a \cdot q, c = b \cdot p \Rightarrow c = (a \cdot q) \cdot p = a \cdot (q \cdot p)$$

TERZA RIDUZIONE: $(X, \cdot, 1_X)$ monoide commutativo (ma non un gruppo).

Per quanto riguarda la proprietà antisimmetrica, vediamo come esprimerla:

$$a \mid b \text{ e } b \mid a \Leftrightarrow \text{esistono } p, q \in X \text{ tali che } a = b \cdot p \text{ e } b = a \cdot q, \text{ quindi } a = a \cdot (q \cdot p).$$

Tenendo presente che si ha anche $a = a \cdot 1_X$, allora:

III) Se a è *cancellabile*, da $a = b \cdot p$ e $b = a \cdot q$ segue $p \cdot q = 1_X$ e quindi p e q sono invertibili.

QUARTA RIDUZIONE: $(X, \cdot, 1_X)$ monoide commutativo (ma non un gruppo) in cui ogni elemento sia cancellabile, escluso al più l'eventuale elemento assorbente.

Nella maggior parte dei casi, le condizioni poste non bastano ad assicurare la antisimmetria della relazione $\mid$. Per esempio, in $(\mathbf{Z}, \cdot, 1)$ si ha: $-3 \mid 3$ e $3 \mid (-3)$. Si può osservare, a questo punto che:

IV) La proprietà antisimmetrica vale se e solo se *il gruppo delle unità* del monoide contiene solo 1_X .

Quest'ultima condizione, tuttavia, non è verificata negli esempi consueti, come visto per $(\mathbf{Z}, \cdot, 1)$. Lo è invece in $(\mathbf{N}, \cdot, 1)$ (e anche in $(\mathbf{N}, +, 0)$). Dove non è verificata possiamo definire *associati* due elementi a, b se $a \mid b$ e $b \mid a$. La relazione

"essere associati" è di equivalenza nel monoide, come è immediato verificare. Chiaramente, questa relazione è *compatibile* con la relazione $|$, nel senso che se $a | b$ e a', b' sono associati ad a e b rispettivamente, allora $a' | b'$. Ne segue la possibilità di definire la relazione $|$ fra le classi d'equivalenza di elementi associati e allora l'insieme quoziente risulta un insieme ordinato.

Vediamo ora alcune proprietà "algebriche": siano $a, b, c \in X$:

- iv. Unicità del quoziente: se $b = a \cdot c$ e $(a, b) \neq (0, 0)$ allora c è l'unico numero che moltiplicato per a dia b .
- v. Se $a | b$ e $a | c$ allora $a | (b+c)$ ed $a | (bc)$.

La iv) è verificata nel monoide $(X, \cdot, 1_X)$ in cui ora lavoriamo dopo la quarta riduzione. La v) implica la presenza in X di un'altra operazione, $+$, rispetto alla quale l'operazione $\cdot$ sia distributiva. Ma allora:

QUINTA RIDUZIONE: $(X, +, \cdot, 1_X)$ *dominio d'integrità*.

Certamente $(\mathbf{N}, +, \cdot, 1_X)$ non lo è, ma possiamo *simmetrizzarlo* aggiungendo i segni ed ottenere $(\mathbf{Z}, +, \cdot, 1_X)$, che lo è.

Proposizione 8.1. Ogni dominio d'integrità finito $(X, +, \cdot, 1_X)$ è un campo.

Dimostrazione. Sia $a \in X$, $a \neq 0_X$, allora l'applicazione $x \rightarrow a \cdot x$ da X in sé è iniettiva, per cui è anche suriettiva e quindi esiste $x' \in X$ tale che $a \cdot x' = 1_X$. Ne segue che a è invertibile e X è un campo.

In un campo, la moltiplicazione $\cdot$ forma un gruppo sugli elementi non nulli, quindi rende non interessante, come già rilevato, il nostro studio. Pertanto:

SESTA ED ULTIMA RIDUZIONE. $(X, +, \cdot, 1_X)$ *dominio d'integrità infinito e non campo*.

In questo ambiente si sviluppa quindi una teoria della divisibilità, che mantiene le proprietà elementari valide in $\mathbf{N}$, a parte l'antisimmetria della relazione "divide". Per quest'ultima si ha:

Proposizione 8.2. In ogni dominio d'integrità $(X, +, \cdot, 1_X)$, per ogni $a, b \in X$, non entrambi nulli, si ha:

a e b sono associati $\Leftrightarrow$ esiste un elemento invertibile u tale che $b = a \cdot u$.

Dimostrazione. Se $b = a \cdot u$, con u invertibile (rispetto al prodotto), allora $a = b \cdot u^{-1}$ e quindi a e b sono associati. Inversamente, supponiamo che a e b siano associati. Allora, se $a = 0_X$ allora $b = 0_X$ e viceversa. Sia ora $a \neq 0_X$, allora a è cancellabile e quindi, come già osservato, segue $b = a \cdot q$, con q invertibile.

E' interessante quindi sapere chi siano gli elementi invertibili dell'anello. Nel caso di $\mathbf{Z}$ sono solo 1 e -1; nel caso di un anello $K[x]$ di polinomi in una indeterminata a coefficienti in un campo K gli elementi invertibili sono i polinomi di grado zero, identificabili con gli elementi non nulli di K .

In questi due casi è possibile effettuare una scelta "canonica" dei rappresentanti delle classi di elementi associati: nel caso di $\mathbf{Z}$, fra due numeri a e $-a$ si sceglie quello positivo; in $K[x]$ ogni polinomio non nullo è associato ad un polinomio *monico*, cioè avente = 1 il coefficiente del monomio di grado più alto.

Sia dato un dominio d'integrità $(X, +, \cdot, 1_X)$ e siano $a, b \in X$. Un elemento $d \in X$ si dice *massimo comune divisore* di a e b se:

- i. è un divisore di a e di b (*divisore comune*)
- ii. è multiplo di ogni altro divisore comune (*massimo*)

Un elemento $m \in X$ si dice *minimo comune multiplo* di a e b se:

- i. è un multiplo di a e di b (*multiplo comune*)
- ii. è divisore di ogni altro multiplo comune (*minimo*)

In $\mathbf{N}$ per ogni a, b esistono entrambi e sono unici, per cui si scrive $d = \text{MCD}(a, b)$, $m = \text{mcm}(a, b)$ e queste due, MCD e mcm, sono operazioni in $\mathbf{N}$, le quali possiedono varie proprietà.

Proposizione 8.3. Siano $a, b, c \in \mathbf{N}$. Allora valgono le seguenti proprietà

- associativa:

$$\text{MCD}(\text{MCD}(a, b), c) = \text{MCD}(a, \text{MCD}(b, c)),$$

$$\text{mcm}(\text{mcm}(a, b), c) = \text{mcm}(a, \text{mcm}(b, c))$$

- commutativa:

$$\text{MCD}(a, b) = \text{MCD}(b, a),$$

$$\text{mcm}(a, b) = \text{mcm}(b, a)$$

- idempotenza:

$$\text{MCD}(a, a) = a = \text{mcm}(a, a)$$

- assorbimento:

$$\text{MCD}(\text{mcm}(a, b), a) = a = \text{mcm}(\text{MCD}(a, b), a)$$

- distributività reciproca:

$$\text{MCD}(\text{mcm}(a, b), c) = \text{mcm}(\text{MCD}(a, c), \text{MCD}(b, c))$$

$$\text{mcm}(\text{MCD}(a, b), c) = \text{MCD}(\text{mcm}(a, c), \text{mcm}(b, c))$$

- elementi neutri/assorbenti:

$$\text{MCD}(a, 0) = a = \text{mcm}(a, 1)$$

$$\text{MCD}(a, 1) = 1, \text{mcm}(a, 0) = 0$$

La struttura algebrica $(\mathbf{N}, \text{mcm}, \text{MCD})$ così ottenuta è un *reticolo*. Si noti che, rispetto alla relazione d'ordine "divide", si ha:

$$\begin{cases} \text{MCD}(a, b) = \inf\{a, b\} \\ \text{mcm}(a, b) = \sup\{a, b\} \end{cases}$$

Si osservi che 1 e 0 sono, rispettivamente, il minimo ed il massimo di $(\mathbf{N}, |)$. Pertanto, il reticolo così ottenuto su $\mathbf{N}$ è distributivo ed ha minimo e massimo.

Se ci poniamo ora in un dominio d'integrità qualsiasi $(X, +, \cdot, 1_X)$, non è detto che ogni coppia di elementi $a, b \in X$ possieda un massimo comune divisore o un minimo comune multiplo.

Tuttavia, se d' è associato a d e d è un massimo comune divisore di a e b , anche d' lo è. Inoltre, d è massimo comune divisore anche di a' , b' , se a' e b' sono associati rispettivamente ad a e b . Lo stesso dicasi per il minimo comune multiplo. Pertanto, l'unicità non è assicurata, ma lo è a meno di elementi associati. Scriveremo perciò $d = \text{MCD}(a, b)$ (o anche $d = (a, b)$) e $m = \text{mcm}(a, b)$ anche se ciò non sarebbe del tutto corretto.

Si pone quindi il problema di trovare condizioni che assicurino l'esistenza di MCD ed mcm. Vediamo in $\mathbf{N}$ come si fa.

I. - In $\mathbf{N}$ sovente si definisce $\text{MCD}(a, b)$ come "il più grande" dei divisori comuni di a e b , dove l'espressione "il più grande" non si riferisce all'ordine parziale $|$ ma a quello totale $\leq$. Quest'ultimo è legato all'altro dal fatto che, per ogni $a, b \in \mathbf{N} \setminus \{0\}$,

$$a | b \text{ implica } a \leq b.$$

Pertanto, ogni divisore comune di a e b è $\leq d$. Ne segue che basta determinare l'insieme $D(a)$ dei divisori di a e quello, $D(b)$, dei divisori di b , entrambi finiti, farne l'intersezione $D(a) \cap D(b)$ e prenderne il massimo:

$$D(12) = \{1, 2, 3, 4, 6, 12\}, D(18) = \{1, 2, 3, 6, 9, 18\},$$

$$D(a) \cap D(b) = \{1, 2, 3, 6\} \Rightarrow \text{MCD}(12, 18) = 6.$$

Si ha inoltre: $\text{mcm}(12, 18) = 12 \cdot 18 / 6 = 36$, come si può verificare.

Ma tutto ciò è esportabile al massimo in $\mathbf{Z}$, in cui c'è l'ordine totale come in $\mathbf{N}$, ma non per esempio fra i polinomi.

II. - Il *procedimento euclideo delle divisioni successive*: in $\mathbf{N}$ si può eseguire la *divisione col resto*: per ogni $a, b \in \mathbf{N}$, $b \neq 0$, esistono e sono unici $q \in \mathbf{N}$ (quoziente) ed $r \in \mathbf{N}$ (resto) tali che:

$$a = b \cdot q + r, 0 \leq r < b.$$

Infatti, $\{d \in \mathbf{N} \mid \exists q \in \mathbf{N}, d = a - b \cdot q\}$ non è vuoto, dato che contiene a , ed ha perciò il minimo, che chiameremo r . Allora $a = b \cdot q + r$, ed inoltre necessariamente $0 \leq r < b$, altrimenti r non sarebbe il minimo.

Si noti che $b | a$ se e solo se $r = 0$. Si può dimostrare facilmente il seguente lemma:

Lemma 8.3. Per ogni $a, b \in \mathbf{N}$, $b \neq 0$, posto $a = b \cdot q + r$, $0 \leq r < b$, si ha:

- a) Ogni divisore comune di a e b è anche divisore comune di b ed r , e viceversa.
 b) $\text{MCD}(a, b) = \text{MCD}(b, r)$

Ne segue il seguente algoritmo: supposto $a \geq b$ e posto $d = \text{MCD}(a, b)$,

$$a = b \cdot q_1 + r_1, \quad 0 \leq r_1 < b: \quad \text{se } r_1 = 0 \text{ allora } d = b$$

$$b = r_1 \cdot q_2 + r_2, \quad 0 \leq r_2 < r_1: \quad \text{se } r_2 = 0 \text{ allora } d = r_1$$

$$r_1 = r_2 \cdot q_3 + r_3, \quad 0 \leq r_3 < r_2: \quad \text{se } r_3 = 0 \text{ allora } d = r_2$$

.....

Il procedimento termina dopo un numero finito di passi, poiché i resti decrescono ad ogni passo.

$$18 = 12 \cdot 1 + 6$$

$$12 = 6 \cdot 2 + 0 \Rightarrow \text{MCD}(18, 12) = 6$$

Anche in alcuni domini d'integrità si può eseguire la divisione col resto, per esempio in $\mathbf{Z}$, in $K[x]$ ed in $\mathbf{Z}[i] = \{a+ib, a, b \in \mathbf{Z}\}$ (*anello degli interi di Gauss*).

- In $\mathbf{Z}$, posto $|x|$ = valore assoluto di x , per ogni a, b , con $b \neq 0$, esistono e sono unici $q, r \in \mathbf{Z}$ tali che:

$$a = b \cdot q + r, \quad 0 \leq r < |b|.$$

- In $K[x]$, K campo, posto $\text{gr}(p)$ = grado del polinomio non nullo p , per ogni coppia a, b di polinomi, b diverso dal polinomio nullo, esistono e sono unici due polinomi q, r tali che:

$$a = b \cdot q + r, \text{ con } r \text{ polinomio nullo oppure } \text{gr}(r) < \text{gr}(b).$$

- In $\mathbf{Z}[i]$, posto $\|a+ib\| = a^2+b^2$, che è un intero non negativo, per ogni $\alpha, \beta \in \mathbf{Z}[i]$, con $\beta \neq 0+i0$, esistono $\theta, \rho \in \mathbf{Z}$ tali che:

$$\alpha = \beta \cdot \theta + \rho, \quad 0 \leq \|\rho\| < \|\beta\|.$$

Con l'analogo algoritmo si trovano quindi MCD ed mcm anche in questi anelli, detti *anelli euclidei*. Molti altri domini d'integrità non hanno però questa possibilità, per esempio, $\mathbf{Z}[x]$, $K[x_1, \dots, x_n]$ con $n > 1$, $\mathbf{Z}[\sqrt{5}]$.

III. - La scomposizione in fattori irriducibili. In $\mathbf{N}$ si definisce *primo* o *irriducibile* un numero p maggiore di 1 e divisibile solo per 1 e per se stesso. Una sua proprietà rilevante, e che lo caratterizza, è:

$$\text{per ogni } a, b \in \mathbf{N}, p \mid a \cdot b \Rightarrow p \mid a \text{ oppure } p \mid b.$$

Esistono infiniti numeri primi, come ha dimostrato Euclide. Inoltre:

Teorema fondamentale dell'aritmetica: ogni numero naturale > 1 si può scrivere in uno ed un solo modo come prodotto di fattori primi, a parte l'ordine dei fattori stessi.

Sia Π l'insieme dei primi. Per ogni $a \in \mathbf{N}$, $a \neq 0$, accorpare con l'uso delle potenze i fattori primi uguali e ponendo $= 0$ gli esponenti dei primi mancanti nella sua scomposizione, possiamo rappresentare a nella forma:

$$a = \prod_{p \in \Pi} p^{\alpha_p}$$

dove gli esponenti sono nulli per tutti i primi p tranne un numero finito.

Ciò posto, se $b = \prod_{p \in \Pi} p^{\beta_p}$, si ha:

$$b \mid a \text{ se e solo se per ogni } p \in \Pi, \beta_p \leq \alpha_p.$$

Questa proprietà è nota come "criterio generale di divisibilità", e dipende dalle proprietà delle potenze: posto $\theta_p = \alpha_p - \beta_p$ e $q = \prod_{p \in \Pi} p^{\theta_p}$, segue $a = b \cdot q$. Il viceversa è parimenti immediato.

Ne segue la regola che usualmente si insegna nella scuola secondaria: per ogni a, b non nulli, scritti come sopra, si ha:

$$\text{MCD}(a, b) = \prod_{p \in \Pi} p^{\min\{\alpha_p, \beta_p\}}, \quad \text{mcm}(a, b) = \prod_{p \in \Pi} p^{\max\{\alpha_p, \beta_p\}}.$$

Esempio: $12 = 2^2 \cdot 3 = 2^2 \cdot 3^1 \cdot 5^0 \cdot 7^0 \dots$, $18 = 3^2 \cdot 2 = 2^1 \cdot 3^2 \cdot 5^0 \cdot 7^0 \dots$

$$\text{MCD}(12, 18) = 2^1 \cdot 3^1 \cdot 5^0 \cdot 7^0 \dots = 6, \quad \text{mcm}(12, 18) = 2^2 \cdot 3^2 \cdot 5^0 \cdot 7^0 \dots = 36$$

Questo procedimento, con qualche variante, funziona anche in altri anelli. Innanzitutto, in un anello ogni elemento è multiplo di tutti gli elementi invertibili e di tutti i suoi associati; lo diremo *irriducibile* se questi sono i soli suoi divisori. Per esempio, -3 in $\mathbf{Z}$ ha come divisori $1, -1, 3, -3$, quindi è irriducibile.

Non avremo quindi l'unicità della fattorizzazione, ma solo una *unicità essenziale*, cioè a meno di elementi associati, oltre che dell'ordine dei fattori:

$$12 = 2 \cdot 2 \cdot 3 = (-2) \cdot 2 \cdot (-3) = (-2) \cdot (-2) \cdot 3$$

Inoltre, la proprietà dei numeri primi di dividere necessariamente un fattore tutte le volte che dividono un prodotto non è sempre vera. In $\mathbf{Z}[i\sqrt{5}]$, per esempio, si dimostra che 3 è irriducibile ma non ha questa proprietà. Un elemento p di un dato dominio $(X, +, \cdot, 1_X)$, non nullo e non invertibile, si dice *primo* se

$$\text{per ogni } a, b \in X, p \mid a \cdot b \Rightarrow p \mid a \text{ oppure } p \mid b.$$

Si vede facilmente che se p è primo in X è anche irriducibile, ma non vale sempre il viceversa, come detto.

Un dominio d'integrità $(X, +, \cdot, 1_X)$ si dice *fattoriale* se ogni $a \in X$, non nullo e non invertibile, è esprimibile come prodotto di fattori irriducibili ed in modo essenzialmente unico, cioè se vale in esso il teorema fondamentale dell'aritmetica. Chiaramente, in un tal dominio ogni coppia di elementi possiede MCD ed mcm. Inoltre, ogni irriducibile è primo.

Tutti gli anelli euclidei sono fattoriali, ma non tutti i domini d'integrità lo sono, ed un esempio è $\mathbf{Z}[i\sqrt{5}]$, in cui esistono elementi esprimibili in modi essenzialmente diversi come prodotto di irriducibili. Si ha invece un teorema importante di Gauss: se A è un dominio fattoriale, anche $A[x]$ lo è. Ne segue che,

essendo $K[x_1, \dots, x_n] = (K[x_1, \dots, x_{n-1}])[x_n]$, e $K[x_1]$ euclideo, quindi fattoriale, allora $K[x_1, x_2]$, $K[x_1, x_2, x_3]$, ... sono fattoriali. Analogamente, è fattoriale $\mathbf{Z}[x]$.

Si può dimostrare che un dominio d'integrità $(X, +, \cdot, 1_X)$ è fattoriale se e solo se ogni $a \in X$, non nullo e non invertibile, è esprimibile come prodotto di fattori irriducibili ed inoltre vale una delle seguenti proprietà, tutte equivalenti fra loro:

- ogni irriducibile è primo
- ogni coppia di elementi possiede MCD
- ogni coppia di elementi possiede mcm.

A questo punto, in un dominio fattoriale si pongono due problemi:

- a) Classificare gli elementi irriducibili (= primi)
- b) Scomporre un dato elemento in un prodotto di irriducibili.

In $\mathbf{N}$ entrambi i problemi sono aperti e sono tuttora oggetto di ricerche, a causa delle applicazioni alla crittografia e quindi alla tutela della segretezza nella trasmissione di informazioni bancarie, militari, politiche. Si hanno numerosi criteri di primalità e di divisibilità, alcuni dei quali sono ben noti ed elementari, ma ce ne sono di ben più sofisticati, alcuni dei quali sono implementati sui programmi matematici più comuni (Mathematica, Maple, Derive, ...)

- (Eratostene?) Se un numero dispari non è divisibile per i numeri dispari minori o uguali alla sua radice quadrata è primo.
- (Fermat-Eulero) Un numero p è primo se e solo se per ogni $a \in \mathbf{N}$, $0 < a < p \Rightarrow a^{p-1} \equiv 1 \pmod{p}$.
- Un numero è divisibile per 2, 3, 5, 11, ... se e solo se ... (criteri di divisibilità)

Invece, in alcuni casi, il primo problema è completamente risolto, mentre il secondo no: per esempio, in $\mathbf{C}[x]$ i polinomi irriducibili sono solo quelli di primo grado; in $\mathbf{R}[x]$ sono solo quelli di primo grado e quelli di secondo grado col discriminante negativo.

Per $\mathbf{Q}[x]$ il problema è riconducibile a $\mathbf{Z}[x]$. Infatti, dato un polinomio $f(x)$ a coefficienti razionali, si possono ridurre questi ultimi allo stesso denominatore D , poi raccogliere a fattor comune il denominatore stesso e il MCD d dei numeratori,

ottenendo: $f(x) = \frac{d}{D} f_1(x)$, dove la frazione è poi ridotta ai minimi termini e il polinomio $f_1(x) = a_n x^n + \dots + a_0$ ha i coefficienti a_i interi e *coprime*, ossia $\text{MCD}(a_n, \dots, a_0) = 1$. Un polinomio come f_1 è detto *primitivo*, e si può dimostrare che il prodotto di polinomi primitivi è primitivo.

Si può anche dimostrare che un polinomio primitivo è irriducibile in $\mathbf{Q}[x]$ se e solo se lo è in $\mathbf{Z}[x]$. Una ovvia condizione necessaria per l'irriducibilità è l'assenza di radici. Tale condizione è anche sufficiente se il grado è al più 3. Una condizione sufficiente di irriducibilità è la seguente:

Proposizione (Eisenstein). Un polinomio primitivo $f_1(x) = a_n x^n + \dots + a_0$ è irriducibile se esiste un primo p che divide tutti i coefficienti tranne a_n , mentre p^2 non divide a_0 .

Per esempio, per il criterio di Eisenstein, per ogni $n \geq 2$ i polinomi $x^n - 3$ sono tutti irriducibili in $\mathbf{Z}[x]$ ed in $\mathbf{Q}[x]$, in quanto $p = 3$ divide tutti i coefficienti tranne quello di x^n , ma 9 non divide il termine noto -3. La condizione è però solo sufficiente: $x^2 + 2x + 4$ è irriducibile, ma non soddisfa il criterio di Eisenstein.

Se consideriamo l'insieme $D(n)$ dei divisori di un dato $n \in \mathbf{N}$ abbiamo dei problemi di tipo numerico o configurazionale.

- a) Quanti elementi ha $D(n)$?
- b) Come trovare questi elementi?
- c) Che cosa ha di particolare un numero n con un numero dispari di divisori?
- d) Per quali n l'insieme ordinato $(D(n), |)$ è totalmente ordinato?
- e) Per quali n il reticolo $D(n)$ è un'algebra di Boole?

Escludiamo subito il caso di $n = 0$, perché $D(0) = \mathbf{N}$. Negli altri casi, $D(n)$ è finito e le domande hanno una facile risposta; innanzitutto, scomponiamo n in fattori primi, $n = p_1^{n_1} \dots p_r^{n_r}$. Applicando il criterio generale di divisibilità si ottiene che ogni divisore k di n ha la forma $k = p_1^{k_1} \dots p_r^{k_r}$, con $k_i \leq n_i$. Pertanto, n ha esattamente $m = (n_1+1)(n_2+1) \dots (n_r+1)$ divisori. Ciò risponde ai quesiti a) e b).

c) Elenchiamo gli elementi di $D(n)$ in senso crescente: possiamo osservare che se $k \in D(n)$, anche $n/k \in D(n)$, per cui, di norma, i divisori vanno a coppie. Pertanto, se n è dispari, c'è un divisore k che coincide con n/k , dunque $n = k^2$ è un quadrato.

d) Se n fosse multiplo di due primi distinti p e q allora $D(n)$ non sarebbe totalmente ordinato, per cui esiste $k \in \mathbb{N}$ tale che $n = p^k$. Ovviamente, viceversa, per un tal n i divisori sono sempre confrontabili, perché $p^h \mid p^k$ se e solo se $h \leq k$.

e) Un'algebra di Boole è un reticolo distributivo con minimo e massimo e in cui ogni elemento a ha un *complemento* a' , tale che

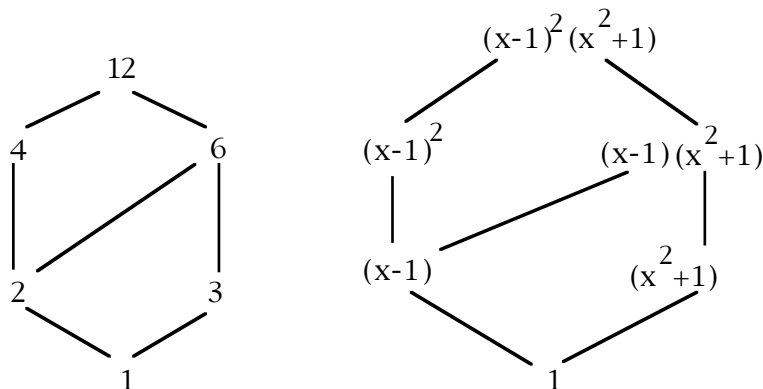
$$\sup\{a, a'\} = \text{massimo}$$

$$\inf\{a, a'\} = \text{minimo}$$

L'esempio più elementare è costituito dall'insieme dei sottoinsiemi di un insieme X , ordinato mediante l'inclusione; in esso $\inf$ e $\sup$ sono l'intersezione e l'unione d'insiemi, il minimo è l'insieme vuoto $\emptyset$ ed il massimo è X . Il complementare di un sottoinsieme Y di X è l'usuale complementare $X \setminus Y$.

In generale, $D(n)$ è un reticolo distributivo con minimo 1 e massimo n , ma non è un'algebra di Boole; lo è se e solo se $n = p_1 p_2 \dots p_r$, con p_i primi distinti. Infatti, in questo caso, per ogni divisore k di n poniamo $k' = n/k$ ed allora $\text{MCD}(k, k') = 1$ e $\text{mcm}(k, k') = k \cdot k' = n$. Se invece $n = p_1^{n_1} \dots p_r^{n_r}$ e, per esempio, $n_1 > 1$, posto $k = p$ allora un eventuale complemento k' , per essere primo con k , dovrebbe essere un divisore di $p_2^{n_2} \dots p_r^{n_r}$, per cui $\text{mcm}(k, k') \mid p p_2^{n_2} \dots p_r^{n_r} \neq n$.

Problemi simili si possono porre anche nei domini fattoriali, naturalmente modulo la relazione "elementi associati". Si scopre così che in ogni dominio fattoriale, per ogni x non nullo la struttura del reticolo dei divisori di x (modulo la detta relazione) dipende unicamente dal numero dei fattori irriducibili distinti che dividono x e dai loro esponenti, indipendentemente non solo dai fattori stessi, ma anche dal dato dominio fattoriale. Se confrontiamo il reticolo dei divisori di 12 o di 18 o dei divisori di $(x-1)^2(x^2+1)$ in $\mathbb{R}[x]$ (modulo il prodotto per numeri reali non nulli) avremo lo stesso numero, 6, di divisori e lo stesso *diagramma di Hasse* (v. fig. seguente). Analogamente, i divisori di 30 o quelli di $(x-1)(x-2)(x+5)$ formeranno algebre di Boole con 2^3 elementi, isomorfe fra loro e con l'algebra dei sottoinsiemi di $\{1, 2, 3\}$.



C'è infine un ulteriore aspetto algebrico da considerare, che lega in un certo senso addizione e MCD:

Proposizione 8.5. (Identità di Bézout). Se $d = \text{MCD}(a, b)$ in $\mathbf{Z}$, esistono $u, v \in \mathbf{Z}$ tali che $au + bv = d$.

Dimostrazione. Questi due coefficienti u, v si ottengono dal procedimento euclideo delle divisioni successive; innanzitutto,

$$a = b \cdot q_1 + r_1 \Rightarrow r_1 = u_1 a + v_1 b, \text{ con } u_1 = 1 \text{ e } v_1 = -q_1;$$

allora, per induzione su i , supposto $r_j = u_j a + v_j b$ per ogni $j \leq i$, e $r_{i-1} = r_i q_{i+1} + r_{i+1}$ allora

$$r_{i+1} = (u_{i-1} - q_{i+1} u_i) a + (v_{i-1} - q_{i+1} v_i) b$$

Pertanto, poiché d è uno di questi resti, vale anche per lui la stessa espressione come combinazione lineare di a e b .

Esempio: $6 = \text{MCD}(12, 18) = -1 \cdot 12 + 1 \cdot 18$.

Osservazione. Dati $a, b, c \in \mathbf{Z}$, l'equazione lineare $ax + by = c$, a due incognite intere, è detta *diophantea*. Le soluzioni esistono se e solo se c è un multiplo di $d = \text{MCD}(a, b)$. Infatti, chiaramente ogni divisore di a e b deve dividere c , e ciò accade anche per d . Viceversa, sia $c = c'd$; per l'identità di Bézout esistono u, v tali che $au + bv = d$, ma allora $a(uc') + b(vc') = c$. In particolare, se $d = 1$ allora l'equazione ha sempre (infinite) soluzioni.

L'analoga equazione in $\mathbf{N}$, ossia $ax + by = c$, con a, b, c interi non negativi dati ed x, y incognite intere non negative, ha altre limitazioni oltre a quella di $d \mid c$. Dividiamo tutto per d , ossia consideriamo il caso di a, b coprimi. Si può dimostrare che esistono soluzioni

per ogni $c > ab - a - b$, mentre non ne esistono per $c = ab - a - b$. Per valori di c minori di $ab - a - b$ occorre provare. In ogni caso, le soluzioni sono in numero finito, perché certamente x ed y non possono superare rispettivamente la parte intera di c/a e c/b

Come detto, l'espressione di $\text{MCD}(a, b)$ a partire da a e b è detta *identità di Bézout*. Non in tutti i domini fattoriali vale questa identità, ed un esempio è $\mathbf{Z}[x]$. Essa è legata alla nozione di *ideale* di un anello: un ideale I di un anello $(A, +, \cdot, 1_A)$ è un sottogruppo del gruppo additivo ed è anche un ideale del monoide moltiplicativo, ossia è un sottoinsieme non vuoto di A tale che, per ogni $x, y \in I$ e per ogni $a \in A$ si ha:

- i) $x - y \in I$
- ii) $x \cdot i$ ed $i \cdot x$ appartengono ad I .

In un anello commutativo un ideale I si dice *principale* se esiste $m \in I$ tale che $I = \{m \cdot a \mid a \in A\}$, cioè è costituito dai multipli di m . In tal caso I si dice *generato* da m . Si usa scrivere $I = \langle m \rangle$. La nozione di divisore si traduce in termini di ideali osservando che, per ogni $m, n \in A$ si ha:

$$m \mid n \text{ se e solo se } \langle n \rangle \subseteq \langle m \rangle.$$

Un dominio d'integrità i cui ideali siano tutti principali è detto *dominio ad ideali principali*. Un esempio è costituito dagli anelli euclidei, come $\mathbf{Z}$ o $K[x]$, ma non tutti i domini ad ideali principali sono euclidei. Si può dimostrare inoltre che un dominio ad ideali principali è sempre fattoriale ed in esso vale l'identità di Bézout, ma se non è euclideo non c'è un algoritmo semplice per trovare i coefficienti u, v .