

ALMA MATER STUDIORUM · UNIVERSITÀ DI BOLOGNA

DIPARTIMENTO DI MATEMATICA

Corso di Laurea in Matematica

L'algebra di Möbius

Tesi di Laurea in Algebra

Relatore:
PAGARIA ROBERTO

Presentata da:
BETTINI MATTEO

Anno Accademico 2024-2025

Introduzione

L'algebra di Möbius è uno strumento sviluppato nell'ambito della matematica discreta, con importanti applicazioni in combinatoria e allo studio degli insiemi parzialmente ordinati. La sua origine risale alla trasformazione di Möbius, introdotta da August Ferdinand Möbius nell'Ottocento, ma la formulazione moderna nasce nel XX secolo, grazie ai lavori di Gian-Carlo Rota sull'algebra di incidenza e la funzione di Möbius su insiemi parzialmente ordinati.

In questa tesi consideriamo particolari poset dotati di join e meet, che sono rispettivamente il minimo dei maggioranti e il massimo dei minoranti. Tali poset sono chiamati reticoli. Andremo a costruire una $\mathbb{K}$ -algebra $A(L)$ generata dagli elementi del reticolo L e con prodotto indotto dal join. Mostriamo identità all'interno dell'algebra $A(L)$ e analizzando il coefficiente di un particolare elemento (ad esempio $\hat{1}$, il massimo) ricaveremo identità notevoli, quali il Teorema Cross-Cut, enunciato da Rota nel 1964 ([Rot64]), e il Teorema di Weisner, dimostrato dal matematico americano Louis Weisner nel 1935.

L'elaborato si articola in quattro capitoli. Il primo capitolo è dedicato alla presentazione dei concetti fondamentali necessari per lo sviluppo della teoria. Si introducono gli insiemi parzialmente ordinati (poset), i relativi diagrammi di Hasse e la definizione di algebra di incidenza. Un ruolo centrale è occupato dalla funzione di Möbius, definita in modo ricorsivo su intervalli di un poset, e dalla relativa formula di inversione di Möbius, risultato con profonde implicazioni in combinatoria.

Il secondo capitolo prosegue con la definizione e l'analisi dell'algebra di Möbius associata a un reticolo finito. Si dimostra come questa possa essere

scomposta in un prodotto di campi, ciascuno associato a un elemento del reticolo. Questa proprietà strutturale ha varie conseguenze notevoli, tra cui i teoremi di Cross-Cut e di Weisner, che consentono di determinare in modo efficace i valori della funzione di Möbius su particolari intervalli del reticolo. Il capitolo si conclude con l'introduzione dei reticoli geometrici e con la definizione del polinomio caratteristico di un poset.

Il terzo capitolo si apre introducendo gli operatori di chiusura, che permettono di associare a un reticolo il sottoreticolo degli elementi chiusi, mantenendo un controllo sulla funzione di Möbius. Nella parte centrale vengono esaminate le condizioni necessarie affinché una mappa tra poset induca un omomorfismo tra le rispettive algebre di Möbius. Infine, viene presentata la classe dei reticoli supersolubili, caratterizzati dalla fattorizzazione completa del polinomio caratteristico, che troveranno applicazione negli esempi del capitolo successivo.

Il quarto capitolo tratta alcuni esempi significativi di poset e il calcolo esplicito della relativa funzione di Möbius. In particolare, analizzeremo:

- il reticolo dei sottospazi di uno spazio vettoriale su un campo finito;
- il reticolo dei sottomoduli di un modulo di torsione finito;
- il reticolo delle partizioni di un insieme.

Questi esempi, scelti per varietà e rilevanza, illustrano le tecniche sviluppate nei capitoli precedenti e ne mostrano applicazioni concrete.

L'algebra di Möbius rappresenta un punto di incontro tra combinatoria e algebra astratta, offrendo un linguaggio potente per descrivere e analizzare strutture discrete complesse. Lungo il percorso seguito in questo elaborato, vedremo come essa permetta non solo di sintetizzare in modo elegante concetti combinatori profondi, ma anche di fornire strumenti concreti per il calcolo e la classificazione.

Indice

Introduzione	i
1 Poset e funzione di Möbius	1
1.1 Poset	1
1.2 Algebra di incidenza e funzione di Möbius	3
2 L'algebra di Möbius	9
2.1 Algebra di Möbius di un reticolo	9
2.2 Idempotenti di $A_V(L, \mathbb{K})$	11
2.3 Conseguenze notevoli	12
2.4 Reticoli geometrici e polinomio caratteristico	14
3 Operatori di chiusura ed estensione di mappe tra poset	21
3.1 Operatori di chiusura	21
3.2 Algebra di Möbius su poset e criteri di estensione	23
3.3 Elementi modulari e reticoli supersolubili	28
4 Funzione di Möbius su poset notevoli	35
4.1 Sottospazi di uno spazio vettoriale finito dimensionale su cam- po finito	35
4.2 Sottomoduli di un modulo di torsione	40
4.3 Reticolo delle partizioni di un insieme	44
Bibliografia	49

Capitolo 1

Poset e funzione di Möbius

1.1 Poset

Definizione 1.1.1. Un *insieme parzialmente ordinato* (abbreviato *poset*) $(P, \leq)$ è un insieme P con una relazione binaria riflessiva, transitiva e anti-simmetrica.

Il termine ‘parzialmente’ indica che non tutti gli elementi sono necessariamente comparabili tra loro.

Definizione 1.1.2. Due poset $(P, \leq_P)$ e $(Q, \leq_Q)$ si dicono *isomorfi* se esiste una biiezione $f : P \rightarrow Q$ tale che per ogni $x, y \in P$, vale $x \leq_P y$ se e solo se $f(x) \leq_Q f(y)$.

Definizione 1.1.3. Dato un poset $(P, \leq)$, chiamiamo *poset opposto a P* , che indichiamo con $(P^{op}, \leq^{op})$, il poset avente gli stessi elementi di P , ma le relazioni d’ordine invertite: dati $x, y \in P$, $x \leq^{op} y$ se e solo se $x \geq y$.

Definizione 1.1.4. Dati i poset $(P, \leq_P)$ e $(Q, \leq_Q)$, diciamo che Q è un *sottoinsieme parzialmente ordinato* di P se $Q \subseteq P$ e per ogni $q_1, q_2 \in Q$, la relazione $q_1 \leq_Q q_2$ vale se e solo se vale $q_1 \leq_P q_2$.

Dato un poset $(A, \leq)$, ogni sottoinsieme $B \subseteq A$ può essere considerato un sottoinsieme parzialmente ordinato di A se dotato della relazione d’ordine parziale indotta da quella su A .

Definizione 1.1.5. Siano $(P, \leq_P)$ e $(Q, \leq_Q)$ due poset. Il poset $Z = (P \times Q, \leq_Z)$ con la relazione d'ordine parziale definita da $(p_1, q_1) \leq_Z (p_2, q_2)$ se e solo se $p_1 \leq_P p_2$ e $q_1 \leq_Q q_2$ è detto *prodotto* dei poset P e Q .

Definizione 1.1.6. Un sottoinsieme C di un poset $(P, \leq)$ è una *catena* se per ogni $x, y \in C$, $x \leq y$ oppure $y \leq x$. Una catena è detta *massimale* se non è contenuta in catene più grandi. Un sottoinsieme B di un poset $(P, \leq)$ è una *anticatena* se per ogni $x, y \in B$ distinti, $x \not\leq y$ e $y \not\leq x$.

Si dice che una catena C ha *lunghezza* n ($\ell(C) = n$) se C ha $n+1$ elementi, e in tal caso la si indica con C_n .

Per semplicità, utilizzeremo la notazione $x < y$ se $x \leq y$ e $x \neq y$.

Definizione 1.1.7. Dato un poset $(P, \leq)$ e $x, y \in P$, diciamo che x *copre* y , e lo indichiamo con $x \succ y$, se $x > y$ e non esiste un altro elemento $z \in P$ tale che $x > z > y$.

Definizione 1.1.8. Dato un poset $(P, \leq)$, si chiama *diagramma di Hasse* associato a P una sua rappresentazione sotto forma di diagramma: ogni elemento di P è rappresentato da un vertice e se y copre x si traccia una linea che collega x , situato più in basso, e y , posto più in alto.

Definizione 1.1.9. Dato un poset $(P, \leq)$, se $x, y \in P$ e $x \leq y$, l'*intervallo* $[x, y]$ è definito come $[x, y] = \{z \in P \mid x \leq z \leq y\}$. Un poset si dice *localmente finito* se ogni suo intervallo è finito.

Esempio 1.1.10. Diamo di seguito alcuni esempi di poset.

- L'insieme dei numeri naturali diversi da zero $\mathbb{N}_+$ con la relazione di divisibilità, cioè $x, y \in \mathbb{N}_+, x \leq y$ se x divide y , indicato $(\mathbb{N}_+, \mid)$;
- L'insieme delle parti di un insieme S di cardinalità finita, con la relazione di inclusione insiemistica, ovvero $A, B \in S, A \leq B$ se $A \subseteq B$, indicato $(P(S), \subseteq)$. Il caso $S = \{a, b, c\}$ è rappresentato in Figura 1.1.

D'ora in avanti, per semplicità di notazione indicheremo con P il poset $(P, \leq)$.

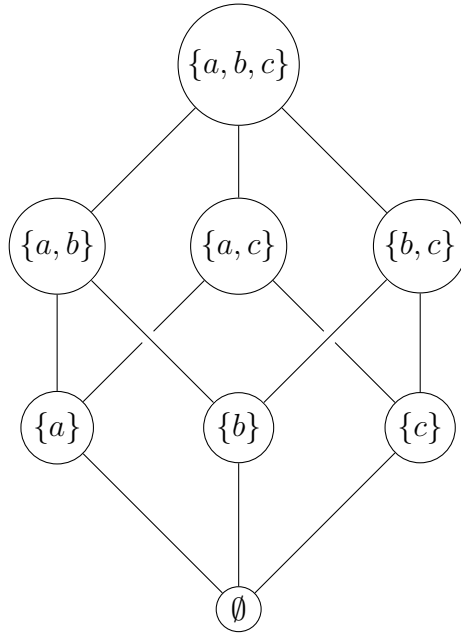


Figura 1.1: Diagramma di Hasse dell'insieme delle parti di $\{a, b, c\}$ con la relazione di inclusione.

1.2 Algebra di incidenza e funzione di Möbius

Dopo aver introdotto le nozioni fondamentali sui poset, possiamo ora definire l'algebra di incidenza e la funzione di Möbius.

Definizione 1.2.1. L'algebra di incidenza $I(P, \mathbb{K})$ di un poset P localmente finito su un campo $\mathbb{K}$ è l'insieme $\{f : P \times P \rightarrow \mathbb{K} \mid f(x, y) = 0 \text{ se } x \not\leq y\}$, con una struttura di $\mathbb{K}$ -algebra associativa con unità data dalle seguenti operazioni:

- $(f + g)(x, y) = f(x, y) + g(x, y)$
- $(f \cdot g)(x, y) = \sum_{x \leq z \leq y} f(x, z)g(z, y)$
- $(r \cdot f)(x, y) = rf(x, y)$

per ogni $f, g \in I(P, \mathbb{K})$, $r \in \mathbb{K}$.

Verifichiamo l'associatività del prodotto: per ogni $f, g, h \in I(P, \mathbb{K})$,

$$\begin{aligned} ((f \cdot g) \cdot h)(x, y) &= \sum_{x \leq z \leq y} (f \cdot g)(x, z) h(z, y) = \sum_{x \leq z \leq y} \sum_{x \leq t \leq z} f(x, t) g(t, z) h(z, y) = \\ &= \sum_{x \leq t \leq y} \sum_{t \leq z \leq y} f(x, t) g(t, z) h(z, y) = \sum_{x \leq t \leq y} f(x, t) (g \cdot h)(t, y) = (f \cdot (g \cdot h))(x, y). \end{aligned}$$

Omettiamo le altre verifiche di $\mathbb{K}$ -algebra perché di routine.

Definiamo ora alcune notevoli funzioni in $I(P, \mathbb{K})$:

- L'elemento neutro moltiplicativo,

$$\delta_P(x, y) = \begin{cases} 1 & \text{se } x = y \\ 0 & \text{altrimenti.} \end{cases}$$

- La *funzione zeta* $\zeta_P(x, y) = 1 \ \forall x \leq y \in P$
- La *funzione di Möbius*, definita dalla seguente ricorrenza:

$$\mu_P(x, x) = 1 \ \forall x \in P \quad (1.1)$$

$$\mu_P(x, y) = - \sum_{x \leq z < y} \mu_P(x, z) \text{ per ogni } x < y \text{ in } P \quad (1.2)$$

Osservazione 1.2.2. La condizione di locale finitezza di P è fondamentale per la definizione di μ_P : garantisce infatti che la somma nell'equazione (1.2) sia finita.

La funzione di Möbius è l'inversa della funzione zeta:

$$\begin{aligned} (\mu_P \cdot \zeta_P)(x, x) &= 1 = \delta_P(x, x) \\ (\mu_P \cdot \zeta_P)(x, y) &= \sum_{x \leq z \leq y} \mu_P(x, z) \zeta_P(z, y) = \sum_{x \leq z \leq y} \mu_P(x, z) = 0 = \delta_P(x, y) \end{aligned}$$

$$\forall x, y \in P, x < y.$$

Un importante risultato che useremo durante la trattazione è la *formula di inversione di Möbius*:

Teorema 1.2.3 (Formula di inversione di Möbius). Sia P poset localmente finito. Allora, date due funzioni $f, g : P \rightarrow \mathbb{K}$, le due seguenti affermazioni sono equivalenti:

1. $f(x) = \sum_{y \geq x} g(y)$;
2. $g(x) = \sum_{y \geq x} \mu_P(x, y) f(y)$.

Dimostrazione. Consideriamo lo spazio vettoriale $\mathbb{K}^P = \{f : P \rightarrow \mathbb{K}\}$ e l'azione sinistra di $I(P, \mathbb{K})$ su tale spazio data da

$$(\xi f)(x) = \sum_{y \geq x} \xi(x, y) f(y)$$

per ogni $\xi \in I(P, \mathbb{K})$, $f \in \mathbb{K}^P$. Verifichiamo che tale azione è compatibile con la moltiplicazione in $I(P, \mathbb{K})$: dati $\xi, \eta \in I(P, \mathbb{K})$,

$$\begin{aligned} (\eta \xi) f(x) &= \sum_{z \geq x} (\eta \xi)(x, z) f(z) = \sum_{z \geq x} \sum_{x \leq y \leq z} \eta(x, y) \xi(y, z) f(z) = \\ &= \sum_{y \geq x} \eta(x, y) \sum_{z \geq y} \xi(y, z) f(z) = \sum_{y \geq x} \eta(x, y) (\xi f)(y) = \eta(\xi f)(x). \end{aligned}$$

Essendo $\mu_P \zeta_P = \delta_P$, vale $f = \zeta_P g$ se e solo se $\mu_P f = g$, da cui la tesi. $\square$

Talvolta può essere utile una formulazione duale del Teorema 1.2.3:

Teorema 1.2.4 (Formula di inversione di Möbius, forma duale). Sia P poset localmente finito. Allora, date due funzioni $f, g : P \rightarrow \mathbb{K}$, le due seguenti affermazioni sono equivalenti:

1. $f(x) = \sum_{y \leq x} g(y)$;
2. $g(x) = \sum_{y \leq x} f(y) \mu_P(x, y)$.

Dimostrazione. La dimostrazione è analoga a quella del Teorema 1.2.3, ma in questo caso consideriamo l'azione destra di $I(P, \mathbb{K})$ su $\mathbb{K}^P$ data da

$$(f \xi)(x) = \sum_{y \leq x} f(y) \xi(x, y).$$

□

Sfruttando l'algebra di incidenza, si dimostra la seguente relazione tra la funzione di Möbius sul prodotto di due poset P e Q e le funzioni di Möbius su P e su Q :

Proposizione 1.2.5. Siano P e Q due poset localmente finiti e Z loro prodotto, $Z = (P \times Q, \leq_Z)$. Allora per ogni $p_1, p_2 \in P$ e $q_1, q_2 \in Q$ vale

$$\mu_Z((p_1, q_1), (p_2, q_2)) = \mu_P(p_1, p_2)\mu_Q(q_1, q_2).$$

Dimostrazione. Definiamo la funzione $f \in I(Z, \mathbb{K})$ come segue:

$$f((p_1, q_1), (p_2, q_2)) = \mu_P(p_1, p_2)\mu_Q(q_1, q_2).$$

Mostriamo che f è l'inversa di ζ_Z in $I(Z, \mathbb{K})$. Si ha

$$\begin{aligned} (\zeta_Z \cdot f)((p_1, q_1), (p_2, q_2)) &= \sum_{(p_1, q_1) \leq (p, q) \leq (p_2, q_2)} \zeta_Z((p_1, q_1), (p, q))f((p, q), (p_2, q_2)) = \\ &= \sum_{(p_1, q_1) \leq (p, q) \leq (p_2, q_2)} \zeta_P(p_1, p)\zeta_Q(q_1, q)\mu_P(p, p_2)\mu_Q(q, q_2) = \\ &= \sum_{\substack{p_1 \leq p \leq p_2 \\ q_1 \leq q \leq q_2}} \zeta_P(p_1, p)\mu_P(p, p_2)\zeta_Q(q_1, q)\mu_Q(q, q_2) = \\ &= \delta_P(p_1, p_2)\delta_Q(q_1, q_2) = \delta_Z((p_1, q_1), (p_2, q_2)). \end{aligned}$$

Dunque f coincide con l'unica inversa moltiplicativa di ζ_Z , cioè μ_Z . □

Esempio 1.2.6. Calcoliamo il valore della funzione di Möbius associata ai poset presentati nell'Esempio 1.1.10.

- Per $(\mathbb{N}_+, |)$ la formula ricorsiva (1.2) diventa

$$\begin{aligned} \mu_{\mathbb{N}_+}(1, 1) &= 1, \\ \mu_{\mathbb{N}_+}(1, n) &= \begin{cases} (-1)^k & \text{se } n = p_1 p_2 \dots p_k \text{ con i } p_i \text{ primi distinti,} \\ 0 & \text{se } n \text{ non è libero da quadrati.} \end{cases} \end{aligned}$$

- Osserviamo che se $T, U \subseteq S$ e $T \subseteq U$, l'intervallo $[T, U] \subseteq P(S)$ è isomorfo al prodotto $\prod_{i=1}^{|U|-|T|} C_1$, con $C_1 = \{1, 2\}$. Applicando la Proposizione 1.2.5 si ottiene

$$\mu_{P(S)}(T, U) = \prod_{i=1}^{|U|-|T|} \mu_{C_1}(1, 2) = (-1)^{|U|-|T|}.$$

Dunque la formula di inversione di Möbius 1.2.3 può essere riscritta nel modo seguente: siano $f, g : P(S) \rightarrow \mathbb{K}$, allora

$$g(T) = \sum_{U \supseteq T} f(U) \text{ per ogni } U \in P(S)$$

se e solo se

$$f(T) = \sum_{U \supseteq T} (-1)^{|U|-|T|} g(U) \text{ per ogni } U \in P(S).$$

Capitolo 2

L'algebra di Möbius

In questo capitolo ci concentreremo sullo studio dell'algebra di Möbius associata a un reticolo. Dopo averne definito la struttura, analizzeremo le sue proprietà e ne dedurremo importanti conseguenze sulla funzione di Möbius in questo contesto.

2.1 Algebra di Möbius di un reticolo

Introduciamo una particolare tipologia di poset chiamata *reticoli*.

Definizione 2.1.1. Sia P un poset. Un *maggiorante* di $x, y \in P$ è un elemento $z \in P$ tale che $z \geq x, z \geq y$. Il *join* $x \vee y$ di due elementi $x, y \in P$, se esiste, è il minimo dei maggioranti di x e y , cioè vale $x \vee y \leq z'$ per ogni maggiorante z' . Se esiste, il join di due elementi è unico. Analogamente, un *minorante* di $x, y \in P$ è un elemento $w \in P$ tale che $w \leq x, w \leq y$; il *meet* $x \wedge y$ di due elementi $x, y \in P$, se esiste, è il massimo dei minoranti w di x e y , cioè vale $x \wedge y \geq w'$ per ogni minorante w' . Denotiamo $\hat{1}$ l'elemento massimale e $\hat{0}$ quello minimale.

Definizione 2.1.2. Un poset L si dice *reticolo* (in inglese *lattice*) se per ogni coppia di elementi x e y in L esistono un meet e un join.

Esempio 2.1.3. Diamo alcuni esempi di reticoli e non reticoli.

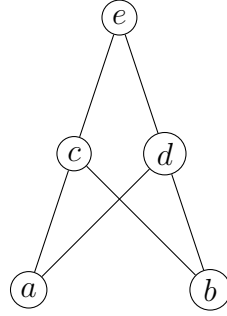


Figura 2.1: Diagramma di Hasse di un poset non reticolo.

- Il poset $P(S)$ dell'Esempio 1.1.10 è un reticolo: infatti, dati $U, V \in P(S)$, si ha $U \vee V = U \cup V$ e $U \wedge V = U \cap V$. In esso sono presenti sia l'elemento minimale $\hat{0} = \emptyset$ sia quello massimale $\hat{1} = S$.
- Il poset $(\mathbb{N}_+, |)$ è un reticolo: se $n, m \in \mathbb{N}_+$, allora $n \vee m = \text{mcm}(n, m)$ e $n \wedge m = \text{MCD}(n, m)$. Questo reticolo ha un elemento minimale $\hat{0} = 1$, ma non ha elemento massimale.
- Il poset rappresentato dal diagramma di Hasse in Figura 2.1 non è un reticolo: infatti c e d sono entrambi maggiori sia di a sia di b , ma a e b non sono confrontabili tra loro; dunque non esiste $c \wedge d$.

Introduciamo adesso una struttura di algebra, detta *algebra di Möbius*, il cui studio porterà a notevoli risultati relativi al calcolo della funzione di Möbius sui reticoli.

Definizione 2.1.4. Dato un reticolo finito L , l'*algebra di Möbius* (relativa al join) di L su $\mathbb{K}$, denotata $A_\vee(L, \mathbb{K})$, è definita come lo spazio vettoriale $\mathbb{K}L$ con base formata dagli elementi di L , munito di un prodotto nel modo seguente: se $x, y \in L$, $x \cdot y = x \vee y$. Tale prodotto è poi esteso per bilinearità. Analogamente si definisce l'algebra di Möbius relativa al meet, $A_\wedge(L, \mathbb{K})$, in cui il prodotto è il meet tra due elementi (nella trattazione tenderemo ad utilizzare, salvo eccezioni, $A_\vee(L, \mathbb{K})$).

2.2 Idempotenti di $A_\vee(L, \mathbb{K})$

Dimostreremo che $A_\vee(L, \mathbb{K})$ è isomorfa a un prodotto di algebre di dimensione 1, ovvero che ammette una base di idempotenti ortogonali. Ciò permette una decomposizione particolarmente utile per il calcolo della funzione di Möbius.

Consideriamo la $\mathbb{K}$ -algebra $A' = \prod_{x \in L} \mathbb{K}\delta'_x$, dove

$$\begin{cases} \delta'_x \delta'_x = \delta'_x \\ \delta'_x \delta'_y = 0, x \neq y \end{cases} \quad (2.1)$$

Essa è dunque generata da idempotenti ortogonali in corrispondenza biunivoca con gli elementi di L . Associamo ora ad ogni $x \in L$ l'elemento $x' \in A'$, $x' = \sum_{y \geq x} \delta'_y$.

L'obiettivo è mostrare il seguente risultato:

Teorema 2.2.1. La mappa $x \mapsto x'$ si estende a un isomorfismo di algebre

$$\phi : A_\vee(L, \mathbb{K}) \rightarrow A'.$$

Dimostrazione. Per dimostrare la suriettività, osserviamo che

$$\delta'_x = \sum_{y \geq x} \mu(x, y) y' = \sum_{y \geq x} \mu(x, y) \phi(y), \quad (2.2)$$

dove la prima uguaglianza segue dal teorema 1.2.3.

Resta da mostrare che ϕ è un morfismo di algebre: dati $x, y \in L$, si ha che

$$\begin{aligned} \phi(x) \cdot \phi(y) &= \left(\sum_{z \geq x} \delta'_z \right) \left(\sum_{w \geq y} \delta'_w \right) = \sum_{\substack{z \geq x \\ w \geq y}} \delta'_z \delta'_w = \sum_{\substack{z \geq x \\ z \geq y}} \delta'_z = \\ &= \sum_{z \geq x \vee y} \delta'_z = \phi(x \vee y) = \phi(x \cdot y). \end{aligned}$$

Poiché $A_\vee(L, \mathbb{K})$ e A' hanno la stessa dimensione come spazi vettoriali, concludiamo che ϕ è un isomorfismo. $\square$

Segue immediatamente dal teorema che $A_\vee(L, \mathbb{K})$ ha una base di idempotenti ortogonali, dati da $\delta_x = \sum_{y \geq x} \mu(x, y)y$.

Osservazione 2.2.2. Vale un risultato analogo per $A_\vee(L, R)$, con R anello commutativo con unità. In questa tesi tratteremo solo il caso di campi.

2.3 Conseguenze notevoli

La decomposizione appena descritta può essere sfruttata per dimostrare alcune importanti proprietà della funzione di Möbius sui reticoli, che verranno discusse in questa sezione.

Definizione 2.3.1. Dati L un reticolo, $z \in L$ e $\alpha \in A_\vee(L, \mathbb{K})$, indichiamo con $[z]\alpha$ il coefficiente di z in α .

Proposizione 2.3.2. Sia L un reticolo finito, e siano $x, y, z \in L$. Allora

$$\sum_{\substack{t \geq y \\ t \vee x = z}} \mu(y, t) = \begin{cases} \mu(y, z) & z \geq y \geq x \\ 0 & \text{altrimenti} \end{cases}$$

Dimostrazione. Consideriamo in $A_\vee(L, \mathbb{K})$ il prodotto $x \cdot \delta_y$ e sviluppiamolo in due modi diversi. Da un lato,

$$x \cdot \delta_y = \left(\sum_{s \geq x} \delta_s \right) \cdot \delta_y = \begin{cases} 0 & \text{se } x \not\leq y \\ \delta_y & \text{se } x \leq y \end{cases}$$

Dall'altro,

$$x \cdot \delta_y = x \cdot \left(\sum_{t \geq y} \mu(y, t)t \right) = \sum_{t \geq y} \mu(y, t)(x \vee t)$$

Se andiamo a considerare il coefficiente di z in $x \cdot \delta_y$, dalla seconda espressione si ha che $[z](x \cdot \delta_y) = \sum_{\substack{t \geq y \\ x \vee t = z}} \mu(y, t)$, mentre dalla prima

$$[z](x \cdot \delta_y) = \begin{cases} \mu(y, z) & \text{se } z \geq y \geq x \\ 0 & \text{altrimenti} \end{cases}$$

Da qui la tesi. $\square$

Denotiamo con $\mu(x)$ il valore $\mu(\hat{0}, x)$.

Corollario 2.3.3 (Teorema di Weisner). Sia L un reticolo finito con almeno due elementi e sia $\hat{0} \neq a \in L$. Allora $\sum_{x \vee a = z} \mu(x) = 0$ per ogni $z \in L$.

Dimostrazione. Segue dalla Proposizione 2.3.2, ponendo $x = a$ e $y = \hat{0}$. $\square$

Definizione 2.3.4. Dato un reticolo L , un elemento $a \in L$ si dice *atomo* se $a \succ \hat{0}$ e si dice *coatomo* se $\hat{1} \succ a$.

Teorema 2.3.5 (Teorema Cross-cut). Sia L un reticolo finito e sia X l'insieme dei suoi atomi. Per ogni $k \in \mathbb{N}$, sia N_k il numero di sottoinsiemi di X di cardinalità k il cui join è $\hat{1}$. Allora $\mu(\hat{0}, \hat{1}) = \sum_k (-1)^k N_k$.

Dimostrazione. Per ogni $x \in X$, $\hat{0} - x = \sum_{y \in L} \delta_y - \sum_{y \geq x} \delta_y = \sum_{y \in L, y \not\geq x} \delta_y$. Perciò, per definizione di X , $\prod_{x \in X} (\hat{0} - x) = \delta_{\hat{0}}$, siccome $\delta_{\hat{0}}$ è l'unico idempotente che compare in ogni termine della produttoria. Osserviamo che $[\hat{1}](\delta_{\hat{0}}) = \mu(\hat{0}, \hat{1})$, ma allo stesso tempo $[\hat{1}](\prod_{x \in X} \hat{0} - x) = N_0 - N_1 + N_2 - \dots$, e dunque segue la tesi. $\square$

Osservazione 2.3.6. Il Teorema 2.3.5 ci dà un metodo efficiente per calcolare la funzione di Möbius in $(\hat{0}, \hat{1})$ per alcuni reticoli. Nel reticolo in Figura 2.2 l'insieme degli atomi è $X = \{a, b, c\}$ e $N_0 = 0, N_1 = 0$, mentre $N_2 = 1$ poiché $a \vee c = \hat{1}$, $N_3 = 1$, perciò $\mu(\hat{0}, \hat{1}) = 0$.

Un'altra conseguenza del Teorema 2.3.5 è il seguente risultato:

Corollario 2.3.7. Se L è un reticolo finito in cui $\hat{1}$ non è join di alcun sottoinsieme di atomi, allora $\mu(\hat{0}, \hat{1}) = 0$.

Dimostrazione. Segue immediatamente dal Teorema 2.3.5. $\square$

Tutti i risultati appena mostrati valgono anche nell'algebra duale $A_\wedge(L, \mathbb{K})$, apportando le dovute modifiche, con dimostrazioni analoghe a quelle precedenti, essendo $A_\wedge(L, \mathbb{K}) = A_\vee(L^{op}, \mathbb{K})$.

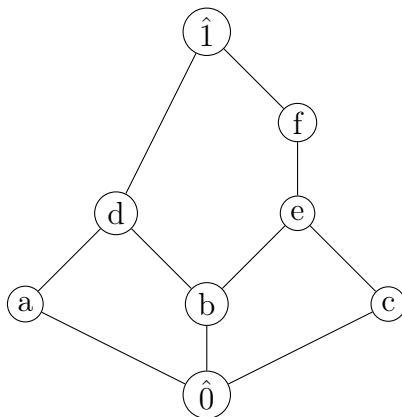


Figura 2.2: Diagramma di Hasse di un reticolo.

2.4 Reticoli geometrici e polinomio caratteristico

Introduciamo ora una classe di reticoli, chiamati reticoli geometrici, la cui algebra di Möbius è di grande interesse. Il loro studio permette infatti di ricavare un vasto numero di identità polinomiali.

Definizione 2.4.1. Un reticolo L si dice *atomico* se ogni suo elemento è join di atomi.

Definizione 2.4.2. Un reticolo L si dice *semimodulare* se vale la seguente proprietà: dati $x, y \in L$ tali che entrambi coprono $x \wedge y$, allora $x \vee y$ copre sia x che y .

Definizione 2.4.3. Un reticolo semimodulare e atomico è detto *geometrico*.

Esempio 2.4.4. Vediamo alcuni esempi:

- $(\mathbb{N}_+, |)$ è semimodulare ma non atomico: gli atomi sono i numeri primi, dunque un numero non libero da quadrati non si scrive come join di atomi;
- il reticolo in Figura 2.3 è atomico ma non semimodulare: $\hat{0} = a \wedge c$ è coperto sia da a sia da c , ma $\hat{1} = a \vee c$ non copre a , dato che $a < d < \hat{1}$;

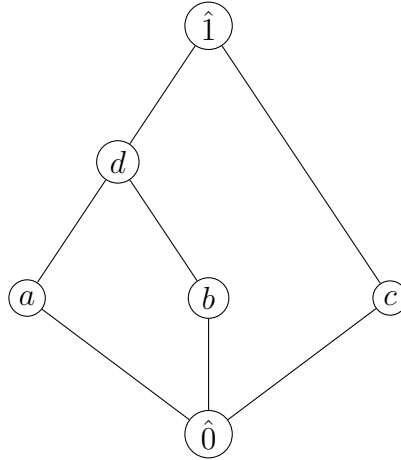


Figura 2.3: Diagramma di Hasse di un reticolo atomico ma non semimodulare.

- i reticoli $V_n(q)$ e $\Pi(S)$ studiati rispettivamente nelle sezioni 4.1 e 4.3 sono entrambi geometrici.

Un'altra interpretazione della semimodularità può essere data introducendo il concetto di grado di un poset.

Definizione 2.4.5. Un poset P si dice *graduato di rango n* se ogni sua catena massimale ha lunghezza n . In tal caso, esiste un'unica *funzione grado* $\rho: P \rightarrow \{0, 1, \dots, n\}$ tale che $\rho(s) = 0$ se $s = \hat{0}$, e $\rho(t) = \rho(s) + 1$ se $t \succ s$. Se $\rho(s) = i$ diciamo che s ha *rango i* .

Proposizione 2.4.6. Sia L un reticolo finito. Le due seguenti affermazioni sono equivalenti:

1. L è graduato, e la funzione grado di L soddisfa

$$\rho(s) + \rho(t) \geq \rho(s \wedge t) + \rho(s \vee t)$$

per ogni $s, t \in L$.

2. se entrambi s e t coprono $s \wedge t$, allora $s \vee t$ copre sia s che t .

Dimostrazione. [Sta12, Proposizione 3.3.2, Capitolo 3].

□

Corollario 2.4.7. Ogni intervallo di un reticolo semimodulare è semimodulare.

Dimostrazione. Poiché ogni intervallo $[s, t] \subseteq L$ è un sottoreticolo di un reticolo semimodulare finito, è anch'esso un reticolo finito graduato. Inoltre vale

$$\rho_{[s,t]}(x) = \rho_L(x) - \rho_L(s).$$

Dunque, se ρ_L soddisfa

$$\rho_L(x) + \rho_L(y) \geq \rho_L(x \wedge y) + \rho_L(x \vee y),$$

allora

$$\begin{aligned} \rho_{[s,t]}(x) + \rho_{[s,t]}(y) &= \rho_L(x) - \rho_L(s) + \rho_L(y) - \rho_L(s) \geq \rho_L(x \wedge y) - \rho_L(s) + \rho_L(x \vee y) - \rho_L(s) = \\ &= \rho_{[s,t]}(x \wedge y) + \rho_{[s,t]}(x \vee y). \end{aligned}$$

Si conclude applicando la Proposizione 2.4.6. $\square$

Dalla Proposizione 2.4.6 si ricava un'applicazione del Corollario 2.3.3 nel caso di reticoli semimodulari finiti.

Proposizione 2.4.8. Sia L un reticolo semimodulare finito di rango n , con funzione grado ρ . Sia $a \in L$ un atomo. Allora

$$\mu(\hat{0}, \hat{1}) = - \sum_{\substack{\text{coatomi } t \\ \text{tali che} \\ t \not\geq a}} \mu(\hat{0}, t).$$

Dimostrazione. Sia $t \in L$ tale che $t \vee a = \hat{1}$. Se $a \leq t$, allora $t = \hat{1}$. Se invece $a \not\leq t$, allora $t \wedge a = \hat{0}$. Per semimodularità, $\rho(t) + \rho(a) \geq \rho(t \wedge a) + \rho(t \vee a)$, quindi se $a \not\leq t$ si ha $\rho(t) + 1 \geq 0 + n$, da cui $\rho(t) = n - 1$ e dunque t è un coatomo. Applicando il Corollario 2.3.3 si ha la tesi. $\square$

Indichiamo con $\ell(s, t) := \ell([s, t])$.

Teorema 2.4.9. Sia L un reticolo semimodulare finito, allora

$$(-1)^{\ell(s,t)} \mu(s,t) \geq 0$$

per ogni $s, t \in L, s \leq t$.

Dimostrazione. Si procede per induzione sul rango di L .

Se L ha rango 1, $L = \{\hat{0}, \hat{1}\}$ e dunque $\mu(\hat{0}, \hat{1}) = -1$. Supponiamo l'ipotesi induttiva sia vera per reticoli di rango strettamente minore di n e consideriamo L reticolo di rango n . Per il Corollario 2.4.7, è sufficiente mostrare che $(-1)^n \mu(\hat{0}, \hat{1}) \geq 0$. Utilizzando l'identità dimostrata nella Proposizione 2.4.8 e moltiplicando da entrambi i lati per $(-1)^n$, si ottiene

$$(-1)^n \mu(\hat{0}, \hat{1}) = (-1)^{n-1} \sum_{\substack{\text{coatomi } t \\ \text{tali che} \\ t \not\leq a}} \mu(\hat{0}, t) = \sum_{\substack{\text{coatomi } t \\ \text{tali che} \\ t \not\leq a}} (-1)^{n-1} \mu(\hat{0}, t)$$

che è somma di quantità non negative e dunque non negativa. $\square$

Introduciamo ora il concetto di polinomio caratteristico di un poset graduato.

Definizione 2.4.10. Sia P un poset con $\hat{0}$, graduato di rango n . Definiamo *polinomio caratteristico* di P il polinomio

$$\chi_P(x) = \sum_{t \in P} \mu(\hat{0}, t) x^{n-\rho(t)}.$$

Definizione 2.4.11. Sia L un reticolo e $A \subseteq L$ un suo sottoinsieme. Indichiamo con $L(A)$ il sottoreticolo di L generato da A , ossia

$$L(A) = \left\{ \bigvee S \mid S \subseteq A \right\}.$$

Teorema 2.4.12. Sia L un reticolo finito e sia X_L l'insieme degli atomi di L . Siano A e B due sottoinsiemi disgiunti di X_L tali che $A \cup B = X_L$, indichiamo con $L(A)$ e $L(B)$ i sottoreticoli di L generati da A e da B e aventi

rispettivamente funzioni di Möbius μ_A e μ_B . Allora

$$\delta_{\hat{0}} = \sum_{x \in L} \mu(\hat{0}, x)x = \left(\sum_{y \in L(A)} \mu_A(\hat{0}, y)y \right) \cdot \left(\sum_{z \in L(B)} \mu_B(\hat{0}, z)z \right).$$

Dimostrazione. Come nella dimostrazione del Teorema 2.3.5,

$$\delta_{\hat{0}} = \prod_{p \in X_L} (\hat{0} - p) = \left(\prod_{p \in A} (\hat{0} - p) \right) \cdot \left(\prod_{q \in B} (\hat{0} - q) \right).$$

Osservando che $\prod_{p \in A} (\hat{0} - p)$ corrisponde all'elemento $\delta_{\hat{0}}$ di $L(A)$ e dunque

$$\prod_{p \in A} (\hat{0} - p) = \sum_{y \in L(A)} \mu_A(\hat{0}, y)y,$$

e che un ragionamento analogo vale per $\prod_{q \in B} (\hat{0} - q)$, si ha la tesi. $\square$

Definizione 2.4.13. Sia L un reticolo e sia X_L l'insieme degli atomi di L . Siano A e B due sottoinsiemi disgiunti di X_L . Diciamo che A e B sono *separatori* di L se $L \simeq L(A) \times L(B)$.

Osservazione 2.4.14. Sia L un reticolo finito graduato e siano A e B separatori di L . Se $x_A \in A$ e $x_B \in B$, allora l'isomorfismo $L \xrightarrow{\sim} L(A) \times L(B)$ manda $x = x_A \vee x_B \in L$ in $(x_A, x_B) \in L(A) \times L(B)$. Segue che

$$\rho_L(x) = \rho_{L(A) \times L(B)}(x_A, x_B) = \rho_L(x_A) + \rho_L(x_B).$$

Inoltre

$$\rho_L(L) = \rho_L(\hat{1}) = \rho_{L(A) \times L(B)}(\hat{1}_A, \hat{1}_B) = \rho_L(\hat{1}_A) + \rho_L(\hat{1}_B) = \rho_L(L(A)) + \rho_L(L(B)).$$

Teorema 2.4.15. Se L è un reticolo geometrico finito e A e B sono separatori di L , allora

$$\chi_L(\lambda) = \sum_{x \in L} \mu(\hat{0}, x) \lambda^{\rho(\hat{1}) - \rho(x)} =$$

$$= \sum_{y \in L(A)} \mu_A(\hat{0}, y) \lambda^{\rho(L(A)) - \rho(y)} \cdot \sum_{z \in L(B)} \mu_B(\hat{0}, z) \lambda^{\rho(L(B)) - \rho(z)} = \chi_{L(A)}(\lambda) \cdot \chi_{L(B)}(\lambda).$$

Dimostrazione. Segue dal Teorema 2.4.12 sostituendo x con $\lambda^{\rho(\hat{1}) - \rho(x)}$, y con $\lambda^{\rho(L(A)) - \rho(y)}$ e z con $\lambda^{\rho(L(B)) - \rho(z)}$; dall'Osservazione 2.4.14 segue infatti

$$\lambda^{\rho(L(A)) - \rho(y)} \cdot \lambda^{\rho(L(B)) - \rho(z)} = \lambda^{\rho(L(A)) + \rho(L(B)) - \rho(y) - \rho(z)} = \lambda^{\rho(\hat{1}) - \rho(x)}.$$

□

Osservazione 2.4.16. L'ipotesi di geometricità di L nel Teorema 2.4.15 assicura che L e i sottoreticoli $L(A)$ e $L(B)$ siano graduati, dunque sono ben definiti i polinomi caratteristici ad essi associati.

Il Teorema 2.4.15 mostra che i reticoli che rispettano determinate ipotesi, come i reticoli geometrici, ammettono una fattorizzazione del proprio polinomio caratteristico attraverso separatori.

Capitolo 3

Operatori di chiusura ed estensione di mappe tra poset

In questo capitolo daremo la definizione di operatore di chiusura su un poset, e studieremo la possibilità di estendere una mappa tra poset a una mappa tra le algebre di Möbius ad essi associate. Infine introdurremo il concetto di elemento modulare in un reticolo geometrico ed enunceremo il Teorema di fattorizzazione degli elementi modulari.

3.1 Operatori di chiusura

Definizione 3.1.1. Dato un poset P , si dice *operatore di chiusura* una funzione $f: P \rightarrow P$ tale che:

1. $f(x) \geq x \ \forall x \in P$ (f è crescente);
2. $f(f(x)) = f(x) \ \forall x \in P$ (f è idempotente);
3. $\forall x, y \in P$ tali che $x \leq y$, $f(x) \leq f(y)$ (f preserva l'ordine degli elementi).

Un elemento $x \in P$ si dice *chiuso* se $f(x) = x$.

Il seguente teorema mostra la correlazione tra un reticolo e il sottoreticolo degli elementi chiusi.

Teorema 3.1.2. Sia L un reticolo finito, e sia $x \mapsto \bar{x}$ un operatore di chiusura. Sia $\bar{L}$ il reticolo degli elementi chiusi di L . Allora, per ogni $x, y \in L$, vale

$$\sum_{\substack{z \geq x \\ \bar{y} = \bar{z}}} \mu_L(x, z) = \begin{cases} 0 & \text{se } \bar{x} > x \\ \bar{\mu}(\bar{x}, \bar{y}) & \text{se } \bar{x} = x \end{cases} \quad (3.1)$$

dove $\bar{\mu}$ indica la funzione di Möbius su $\bar{L}$.

Osservazione 3.1.3. Il poset $\bar{L}$ è effettivamente un reticolo. Siano $x, y \in L$, per definizione di join vale $x, y \leq x \vee y$ e applicando la proprietà 3 si ha $\bar{x}, \bar{y} \leq \overline{x \vee y}$; se $z \in \bar{L}$ è tale che $z \geq \bar{x} \geq x$ e $z \geq \bar{y} \geq y$ allora $z \geq x \vee y$; inoltre, essendo z chiuso, sempre per la proprietà 3 vale $z \geq \overline{x \vee y}$. Perciò esiste il join, che denotiamo $\vee_0$, e si ha $\bar{x} \vee_0 \bar{y} = \overline{x \vee y}$. Dall'ultima uguaglianza si nota che la mappa di chiusura preserva il prodotto.

Dimostrazione del Teorema 3.1.2. L'Osservazione 3.1.3 dice che la funzione $x \mapsto \bar{x}$ è un omomorfismo rispetto al join da L a $\bar{L}$, che può essere esteso a un omomorfismo di algebre $A_\vee(L, \mathbb{K}) \rightarrow A_{\vee_0}(\bar{L}, \mathbb{K})$. Consideriamo inoltre l'omomorfismo $\phi : A_\vee(L, \mathbb{K}) \rightarrow A_{\vee_0}(\bar{L}, \mathbb{K})$ dato dalle seguenti relazioni

$$\phi(\delta_x) = \begin{cases} 0 & \text{se } \bar{x} > x \\ \bar{\delta}_{\bar{x}} & \text{se } \bar{x} = x \end{cases} \quad (3.2)$$

per ogni $\delta_x \in A_\vee(L, \mathbb{K})$. Con $\bar{\delta}_{\bar{x}}$ indichiamo l'elemento idempotente in $A_{\vee_0}(\bar{L}, \mathbb{K})$ corrispondente a $\bar{x}$. I due morfismi appena descritti sono in realtà uguali, in quanto

$$\phi(x) = \phi\left(\sum_{y \geq x} \delta_y\right) = \sum_{y \geq x} \phi(\delta_y) = \sum_{\substack{y \geq x \\ y = \bar{y}}} \bar{\delta}_{\bar{y}} = \sum_{\substack{y \geq x \\ y \in \bar{L}}} \bar{\delta}_{\bar{y}} = \bar{x}.$$

Si conclude la dimostrazione osservando che l'identità (3.1) è equivalente alla (3.2): infatti

$$\phi(\delta_x) = \sum_{z \geq x} \mu(x, z) \phi(z) = \sum_{z \geq x} \mu(x, z) \bar{z}.$$

Se $\bar{x} > x$, allora $\sum_{z \geq x} \mu(x, z) \bar{z} = 0$, da cui

$$\sum_{\substack{z \geq x \\ \bar{z} = \bar{y}}} \mu(x, z) = [\bar{y}] \left(\sum_{z \geq x} \mu(x, z) \bar{z} \right) = 0.$$

Se invece $\bar{x} = x$, allora $\sum_{z \geq x} \mu(x, z) \bar{z} = \delta_{\bar{x}} = \sum_{\substack{\bar{z} \in \bar{L} \\ \bar{z} \geq \bar{x}}} \bar{\mu}(\bar{x}, \bar{z}) \bar{z}$, da cui

$$\sum_{\substack{z \geq x \\ \bar{z} = \bar{y}}} \mu(x, z) = [\bar{y}] \left(\sum_{z \geq x} \mu(x, z) \bar{z} \right) = \bar{\mu}(\bar{x}, \bar{y}).$$

□

3.2 Algebra di Möbius su poset e criteri di estensione

Consideriamo adesso un risultato più generale, che dà una caratterizzazione delle mappe (non solo operatori di chiusura) tra poset che possono essere estese a omomorfismi tra le algebre di Möbius associate. Poiché adesso stiamo lavorando con dei poset qualsiasi e non con dei reticoli, è necessario definire cosa sia $A_{\vee}(P, \mathbb{K})$, dato un poset P . La definizione è coerente con quella vista per i reticoli, ma non essendo garantita l'esistenza del join, il prodotto tra due elementi $x, y \in P$ è

$$x \cdot y = \sum_{\substack{t \geq x \\ t \geq y}} \left(\sum_{s \geq t} \mu(t, s) s \right).$$

Le considerazioni su tale algebra sono analoghe a quelle studiate in precedenza. Una trattazione più approfondita è presente su [Gre71].

Diamo alcune definizioni utili per il teorema che seguirà.

Definizione 3.2.1. Una *connessione di Galois* tra due poset P e Q è una coppia di mappe $\{\phi : P \rightarrow Q, \psi : Q \rightarrow P\}$ che invertono l'ordine degli ele-

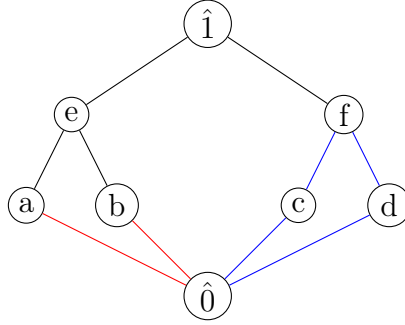


Figura 3.1: Esempi di ideali d'ordine.

menti (se $x, y \in P, x \leq y$, allora $\phi(x) \geq \phi(y)$; lo stesso vale per ψ) e tali che $\forall x \in P, \psi(\phi(x)) \geq x$ e $\forall y \in Q, \phi(\psi(y)) \geq y$.

Osservazione 3.2.2. Data una connessione di Galois $\{\phi, \psi\}$, si ha $\phi\psi\phi = \phi$, infatti per ogni $x \in P$, $\phi(x) \leq \phi(\psi(\phi(x))) \leq \phi(x)$, dove l'ultima disuguaglianza segue da $\psi(\phi(x)) \geq x$ e dal fatto che ϕ scambia l'ordine degli elementi. Analogamente $\psi\phi\psi = \psi$; allora $\psi\phi$ è un operatore di chiusura su P , mentre $\phi\psi$ è un operatore di chiusura su Q . Inoltre ϕ manda chiusi di P in chiusi di Q , infatti se $x \in P$, $\phi(\bar{x}) = \phi\psi\phi(x) = \overline{\phi(x)}$.

Definizione 3.2.3. Un sottoinsieme I di un poset P si dice *ideale d'ordine inferiore* se sono verificate le seguenti condizioni:

1. L'insieme I è non vuoto;
2. per ogni $x \in I$ e $y \in P$, se $y \leq x$ allora $y \in I$.

Definizione 3.2.4. Un ideale d'ordine I di un poset P si dice *principale* se esiste $p \in P$ tale che $I = [\hat{0}, p] = \{x \in P \mid x \leq p\}$.

Esempio 3.2.5. Nel poset in Figura 3.1, sono evidenziati in blu e in rosso due ideali d'ordine; quello in rosso non è principale, mentre quello in blu lo è e corrisponde a $[\hat{0}, f]$.

Enunciamo ora il teorema di estensione:

Teorema 3.2.6. Siano P e Q due poset finiti e sia $\phi : P \rightarrow Q$ una qualsiasi funzione. Allora ϕ estende ad un omomorfismo $A_\vee(P, \mathbb{K}) \rightarrow A_\vee(Q, \mathbb{K})$ se e solo se valgono le seguenti condizioni:

1. se $x, y \in P$ e $x \leq y$, allora $\phi(x) \leq \phi(y)$;
2. per ogni ideale d'ordine principale I , $\phi^{-1}(I)$ è vuoto oppure è un ideale principale (equivalentemente, per ogni $q \in Q$, l'insieme $\{p \in P \mid \phi(p) \leq q\}$ ha un massimo oppure è vuoto).

Dimostrazione. Supponiamo che ϕ estenda ad un omomorfismo. Per ogni $x, y \in P$ tali che $x \leq y$, vale $x \cdot y = x \vee y = y$, ed essendo ϕ omomorfismo $\phi(x) \cdot \phi(y) = \phi(x \cdot y) = \phi(y)$, dunque $\phi(x) \leq \phi(y)$ ed è verificata la proprietà 1.

Mostriamo che vale la proprietà 2. Sia $q \in Q$, allora se $\{p \in P \mid \phi(p) \leq q\}$ è non vuoto, esiste $p \in P$ tale che $\phi(p) \leq q$. Sviluppando $\phi(p)$ in due modi diversi si ha

$$\sum_{x \geq p} \phi(\delta_x) = \phi\left(\sum_{x \geq p} \delta_x\right) = \phi(p) = \sum_{y \geq \phi(p)} \delta_y.$$

Dato che $q \geq \phi(p)$, esiste $x \in P, x \geq p$ tale per cui l'elemento $\phi(\delta_x)$, visto nella sua scrittura in idempotenti di $A_V(Q, \mathbb{K})$, contiene δ_q . Inoltre tale x è unico: se esistesse $z \geq p$ con le stesse caratteristiche si avrebbe $\phi(\delta_x) \cdot \phi(\delta_z) \neq 0$, ma essendo ϕ omomorfismo le immagini di δ_x e δ_z sono ortogonali e si ha una contraddizione. Mostriamo infine che $x = \max\{p \in P \mid \phi(p) \leq q\}$. Per quanto visto in precedenza δ_q compare come addendo in $\phi(\delta_x)$, che a sua volta compare come addendo in $\phi(x) = \sum_{z \geq x} \phi(\delta_z)$, e affinché ciò accada dev'essere $\phi(x) \leq q$. Inoltre, se $\phi(p) \leq q$, allora $x \geq p$.

Vediamo ora il viceversa. Sia $Q_0 = \{q \in Q \mid q \geq \phi(p) \text{ per qualche } p \in P\}$ e definiamo $\psi : Q_0 \rightarrow P$, $\psi(q) = \max\{p \in P \mid \phi(p) \leq q\}$ (ben definita per l'ipotesi 2). Si osserva che ϕ e ψ formano una connessione di Galois tra P e Q_0^{op} : infatti $\phi : P \rightarrow Q_0$ preserva l'ordine degli elementi per la condizione 1, e dunque $\phi : P \rightarrow Q_0^{op}$ lo inverte, mentre dati $q_1, q_2 \in Q_0$ tali che $q_1 \leq^{op} q_2$, vale

$$\psi(q_1) = \max\{p \in P \mid \phi(p) \leq q_1\} \geq \max\{p \in P \mid \phi(p) \leq q_2\} = \psi(q_2);$$

inoltre

$$\forall p' \in P, \quad \psi\phi(p') = \max\{p \in P \mid \phi(p) \leq \phi(p')\} \geq p'$$

e

$$\forall q \in Q_0, \quad \phi\psi(q) = \phi(\max\{p \in P \mid \phi(p) \leq q\}) \geq^{op} q.$$

Indichiamo con $\bar{p} = \psi\phi(p)$ e $\bar{q} = \phi\psi(q)$.

Definiamo ora un omomorfismo $\phi^* : A_V(P, \mathbb{K}) \rightarrow A_V(Q, \mathbb{K})$, nel modo seguente:

$$\phi^*(\delta_p) = \begin{cases} 0 & \text{se } \bar{p} > p \\ \sum_{\substack{q \in Q \\ \bar{q} = \phi(p)}} \delta_q & \text{se } \bar{p} = p. \end{cases}$$

Verifichiamo che ϕ^* è effettivamente un omomorfismo di algebre. Siano $p_1, p_2 \in P$, se $\bar{p}_1 > p_1$ oppure $\bar{p}_2 > p_2$, allora $\phi^*(\delta_{p_1}) \cdot \phi^*(\delta_{p_2}) = 0$. Altrimenti si ha

$$\phi^*(\delta_{p_1}) \cdot \phi^*(\delta_{p_2}) = \left(\sum_{\substack{q \in Q \\ \bar{q} = \phi(p_1)}} \delta_q \right) \left(\sum_{\substack{q \in Q \\ \bar{q} = \phi(p_2)}} \delta_q \right) = \sum_{\substack{q \in Q \\ \bar{q} = \phi(p_1) = \phi(p_2)}} \delta_q. \quad (3.3)$$

Osserviamo che se $\phi(p_1) = \phi(p_2)$ e p_1, p_2 sono chiusi allora

$$p_1 = \bar{p}_1 = \psi\phi(p_1) = \psi\phi(p_2) = \bar{p}_2 = p_2.$$

In tal caso la (3.3) diventa $\phi^*(\delta_{p_1}) \cdot \phi^*(\delta_{p_1}) = \phi^*(\delta_{p_1})$, mentre se ciò non accade $\phi^*(\delta_{p_1}) \cdot \phi^*(\delta_{p_2}) = 0$; dunque ϕ^* è un omomorfismo.

Resta da mostrare che ϕ^* ristretta a P coincide con ϕ . Sia $x \in P$. Allora

$$\phi^*(x) = \sum_{y \geq x} \phi^*(\delta_y) = \sum_{\substack{y \geq x \\ \bar{y} = y}} \sum_{\substack{q \in Q \\ \bar{q} = \phi(y)}} \delta_q = \sum_{\substack{q \in Q \\ \bar{q} \geq \phi(x)}} \delta_q \sum_{\substack{y \geq x \\ \bar{y} = \bar{q} \\ \phi(y) = \bar{q}}} 1.$$

La sommatoria più a destra nell'ultimo termine fa 1, perché $y = \bar{y} = \psi\phi(y) = \psi(\bar{q})$, quindi si somma su un solo termine. Inoltre, essendo $\phi\psi(q) \leq q$, si conclude

$$\phi^*(x) = \sum_{\substack{q \in Q \\ q \geq \phi(x)}} \delta_q = \phi(x).$$

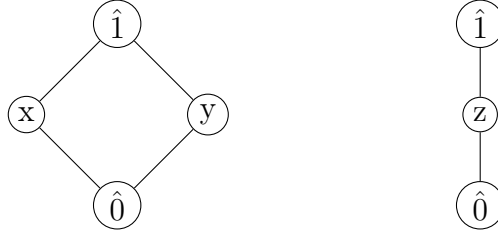


Figura 3.2: Diagramma di Hasse dei poset P (a sinistra) e Q (a destra).

Il teorema è dunque dimostrato. $\square$

Esempio 3.2.7. Consideriamo i poset P e Q in Figura 3.2 e la funzione $\phi : P \rightarrow Q$ tale che $\phi(\hat{0}) = \phi(y) = \hat{0}$ e $\phi(x) = \phi(\hat{1}) = \hat{1}$. Si verifica facilmente che questa rispetta le proprietà 1 e 2 del teorema e dunque si estende ad un omomorfismo $\phi^* : A_\vee(P, \mathbb{K}) \rightarrow A_\vee(Q, \mathbb{K})$. Facciamo il calcolo esplicito: in P , x non è chiuso ($\psi\phi(x) = \psi(\hat{1}) = \hat{1}$), così come $\hat{0}$ ($\psi\phi(\hat{0}) = \psi(\hat{0}) = y$). Allora

$$\phi^*(\delta_{\hat{0}}) = \phi^*(\delta_x) = 0.$$

Sono invece chiusi $\hat{1}$ ($\psi\phi(\hat{1}) = \psi(\hat{1}) = \hat{1}$) e y ($\psi\phi(y) = \psi(\hat{0}) = y$). In Q sono chiusi $\hat{0}$ ($\phi\psi(\hat{0}) = \phi(y) = \hat{0}$) e $\hat{1}$ ($\phi\psi(\hat{1}) = \phi(\hat{1}) = \hat{1}$), mentre z non è chiuso ($\phi\psi(z) = \phi(y) = \hat{0}$). Segue che

$$\begin{aligned} \phi^*(\delta_{\hat{1}}) &= \sum_{\substack{q \in Q \\ \bar{q} = \phi(\hat{1})}} \delta_q = \sum_{\substack{q \in Q \\ \bar{q} = \hat{1}}} \delta_q = \delta_{\hat{1}}; \\ \phi^*(\delta_y) &= \sum_{\substack{q \in Q \\ \bar{q} = \phi(y)}} \delta_q = \sum_{\substack{q \in Q \\ \bar{q} = \hat{0}}} \delta_q = \delta_{\hat{0}} + \delta_z \end{aligned}$$

e dunque $\phi^* = \phi$ su P .

Dal teorema segue un interessante risultato riguardante i sottoinsiemi di un poset e l'algebra di Möbius ad essi associata.

Corollario 3.2.8. Sia P un poset finito e sia $P_0 \subseteq P$. Allora $A_\vee(P_0, \mathbb{K})$ è isomorfa come algebra al sottospazio generato dagli elementi di P_0 in $A_\vee(P, \mathbb{K})$ se e solo se la restrizione di ogni ideale d'ordine principale di P a P_0 è vuota oppure è un ideale d'ordine principale.

Dimostrazione. Consideriamo la funzione $\phi : P_0 \rightarrow P, \phi(p) = p$. Chiaramente ϕ preserva l'ordine degli elementi, inoltre per ipotesi l'insieme

$$\{p \in P_0 \mid \phi(p) \leq q\}$$

ha un massimo oppure è vuoto per ogni $q \in P$. Dunque per il Teorema 3.2.6 ϕ è un omomorfismo iniettivo di algebre e, in particolare, un isomorfismo sull'immagine. $\square$

Osservazione 3.2.9. Nel caso di reticoli, nel Corollario 3.2.8 possiamo sostituire la condizione sugli ideali d'ordine con la richiesta che P_0 sia chiuso rispetto al join, ovvero per ogni $x, y \in P_0$, vale $x \vee y \in P_0$.

3.3 Elementi modulari e reticoli supersolubili

In questa sezione studieremo alcune conseguenze dei teoremi sugli operatori di chiusura e sui criteri di estensione; inoltre introdurremo il concetto di elemento modulare e una categoria di reticoli chiamati reticoli supersolubili.

Cominciamo con l'enunciare un corollario del Teorema 2.4.12.

Teorema 3.3.1. Sia L un reticolo finito e sia $z \in L$. Allora

$$\sum_{x \in L} \mu(\hat{0}, x)x = \left(\sum_{t \leq z} \mu(\hat{0}, t)t \right) \cdot \left(\sum_{y \wedge z = \hat{0}} \mu(\hat{0}, y)y \right).$$

Dimostrazione. Consideriamo la partizione $A \cup B$ di L ottenuta ponendo $A = \{p \in X_L \mid p \leq z\}$ e $B = \{p \in X_L \mid p \not\leq z\}$. Applicando il Teorema 2.4.12 si ha

$$\sum_{x \in L} \mu(\hat{0}, x)x = \left(\sum_{t \leq z} \mu(\hat{0}, t)t \right) \cdot \left(\sum_{y \in L(B)} \mu(\hat{0}, y)y \right).$$

Osserviamo che $L(A) = [\hat{0}, z]$, mentre $L(B) = \{\text{join di atomi} \not\leq z\}$. La dimostrazione consiste nel mostrare che i termini della sommatoria che coinvolgono elementi $y \in L(B)$ sono nulli a meno che $y \wedge z = \hat{0}$. Per ogni

$y \in L$, consideriamo la mappa $t \mapsto t \vee y$ da $[\hat{0}, z]$ a $[y, z \vee y]$. Tale funzione è un omomorfismo rispetto al join; infatti $(t \vee s) \vee y = (t \vee y) \vee (s \vee y)$ per ogni $t, s \in [\hat{0}, z]$. Inoltre ad essa è associato l'operatore di chiusura su $[\hat{0}, z]$ $f(t) = (t \vee y) \wedge z$. Verifichiamo che f è effettivamente un operatore di chiusura:

- f è crescente, infatti $t \leq t \vee y$ e $t \leq z$ perché $t \in [\hat{0}, z]$, dunque $t \leq (t \vee y) \wedge z = f(t)$;
- f preserva l'ordine degli elementi, infatti dati $t, s \in [\hat{0}, z]$, $t \leq s$ vale $t \vee y \leq s \vee y$ e dunque $f(t) = (t \vee y) \wedge z \leq (s \vee y) \wedge z = f(s)$;
- f è idempotente; per ogni $t \in [\hat{0}, z]$, per la crescenza della funzione si ha $f(t) \leq f(f(t))$, ma d'altra parte vale $(t \vee y) \wedge z \leq t \vee y$, da cui $((t \vee y) \wedge z) \vee y \leq (t \vee y) \vee y = t \vee y$ e dunque

$$f(f(t)) = (((t \vee y) \wedge z) \vee y) \wedge z \leq (t \vee y) \wedge z = f(t).$$

Segue che $f(f(t)) = f(t)$.

Sia ora $\delta_{\hat{0}}^{[\hat{0}, z]}$ l'idempotente associato a $\hat{0}$ in $[\hat{0}, z]$, ovvero $\delta_{\hat{0}}^{[\hat{0}, z]} = \sum_{t \leq z} \mu(\hat{0}, t)t$. Se chiamiamo $\bar{f}$ la mappa indotta da f su $A_{\vee}([\hat{0}, z], \mathbb{K})$, applicando il Teorema 3.1.2 si ha che $\bar{f}(\delta_{\hat{0}}^{[\hat{0}, z]})$ è non nullo se e solo se $\hat{0} = \bar{\hat{0}} = (\hat{0} \vee y) \wedge z = y \wedge z$. Se invece $y \wedge z \neq \hat{0}$ allora

$$\bar{f}(\delta_{\hat{0}}^{[\hat{0}, z]}) = \sum_{t \leq z} \mu(\hat{0}, t)((t \vee y) \wedge z) = 0. \quad (3.4)$$

Osserviamo ora che se $t \in [\hat{0}, z]$ allora $t \vee y = ((t \vee y) \wedge z) \vee y$: infatti $t \leq (t \vee y)$ e $t \leq z$, da cui $t \vee y \leq ((t \vee y) \wedge z) \vee y$; d'altra parte, $(t \vee y) \wedge z \leq t \vee y$ e quindi $((t \vee y) \wedge z) \vee y \leq (t \vee y) \vee y = t \vee y$. Sfruttando questa uguaglianza e l'equazione (3.4) si ottiene

$$\left(\sum_{t \leq z} \mu(\hat{0}, t)t \right) \cdot y = \sum_{t \leq z} \mu(\hat{0}, t)(t \vee y) = \sum_{t \leq z} \mu(\hat{0}, t)((t \vee y) \wedge z) \vee y = 0.$$

Questo conclude la dimostrazione. $\square$

Definizione 3.3.2. Sia L un reticolo geometrico. Un elemento $x \in L$ si dice *modulare* se per ogni $y \in L$ vale

$$\rho_L(x) + \rho_L(y) = \rho_L(x \wedge y) + \rho_L(x \vee y). \quad (3.5)$$

Esempio 3.3.3. Sia L un reticolo geometrico finito.

- Gli elementi $\hat{0}$ e $\hat{1}$ sono modulari.
- Ogni atomo a è modulare: se $a \leq y$, allora $a \wedge y = a$ e $a \vee y = y$, quindi l'equazione (3.5) è soddisfatta; se $a \not\leq y$, si ha $\rho(a \wedge y) = \rho(\hat{0}) = 0$, $\rho(a) = 1$ e per semimodularità $\rho(a \vee y) = 1 + \rho(y)$, perciò anche in questo caso vale l'equazione (3.5).
- Se S è un insieme di cardinalità n , allora ogni elemento di $P(S)$ è modulare: infatti se $U \in P(S)$ allora $\rho(U) = |U|$, quindi per ogni $V \in P(S)$ vale $\rho(U) + \rho(V) = |U| + |V| = |U \cup V| + |U \cap V| = \rho(U \cup V) + \rho(U \cap V)$. In questo caso si dice che il reticolo $P(S)$ è modulare.

In un reticolo geometrico, la presenza di un elemento modulare induce una fattorizzazione del polinomio caratteristico, come mostrato dal seguente teorema.

Teorema 3.3.4 (Fattorizzazione dell'elemento modulare). Sia L un reticolo geometrico finito, e sia $z \in L$ un elemento modulare. Allora

$$\begin{aligned} \chi_L(\lambda) &= \sum_{x \in L} \mu(\hat{0}, x) \lambda^{\rho(\hat{1}) - \rho(x)} = \sum_{t \leq z} \mu(\hat{0}, t) \lambda^{\rho(z) - \rho(t)} \cdot \sum_{y \wedge z = \hat{0}} \mu(\hat{0}, y) \lambda^{\rho(\hat{1}) - \rho(z) - \rho(y)} = \\ &= \chi_{[\hat{0}, z](\lambda)} \cdot \sum_{y \wedge z = \hat{0}} \mu(\hat{0}, y) \lambda^{\rho(\hat{1}) - \rho(z) - \rho(y)}. \end{aligned}$$

Dimostrazione. Dividiamo la dimostrazione in più passi.

PASSO 1. Siano $t \leq z$ e y tale che $y \wedge z = \hat{0}$ (e dunque $y \wedge t = \hat{0}$). Mostriamo che $z \wedge (y \vee t) = t$. Chiaramente $z \wedge (y \vee t) \geq t$; d'altra parte, per

la modularità di z si ha

$$\begin{aligned}\rho(z \wedge (y \vee t)) &= \rho(z) + \rho(y \vee t) - \rho(z \vee y) = \\ &= \rho(z) + \rho(y \vee t) - \rho(z) - \rho(y) + \rho(z \wedge y).\end{aligned}$$

Essendo $\rho(z \wedge y) = \rho(\hat{0}) = 0$, utilizzando la semimodularità di L otteniamo

$$\rho(z \wedge (y \vee t)) = \rho(y \vee t) - \rho(y) \leq \rho(y) + \rho(t) - \rho(y \wedge t) - \rho(y) = \rho(t).$$

Dunque $z \wedge (y \vee t) = t$.

PASSO 2. Mostriamo adesso che, se t e y sono come sopra, allora vale $\rho(t \vee y) = \rho(t) + \rho(y)$. Per la modularità di z

$$\rho(z \wedge (y \vee t)) + \rho(z \vee (y \vee t)) = \rho(z) + \rho(y \vee t).$$

Dal Passo 1 sappiamo che $z \wedge (y \vee t) = t$, da cui

$$\rho(z \vee (y \vee t)) = \rho(z) + \rho(y \vee t) - \rho(t).$$

Inoltre, sempre dalla modularità di z segue

$$\rho(z \vee (y \vee t)) = \rho(z \vee y) = \rho(z) + \rho(y) - \rho(z \wedge y) = \rho(z) + \rho(y).$$

Mettendo insieme le due uguaglianze si ottiene $\rho(t \vee y) = \rho(t) + \rho(y)$, come richiesto.

CONCLUSIONE. Adesso possiamo applicare il Teorema 3.3.1 sostituendo x con $\lambda^{\rho(\hat{1})-\rho(x)}$, t con $\lambda^{\rho(z)-\rho(t)}$ e y con $\lambda^{\rho(\hat{1})-\rho(z)-\rho(y)}$. Infatti, dal Passo 2 sappiamo

$$t \cdot y \mapsto \lambda^{\rho(\hat{1})-\rho(t)-\rho(y)} = \lambda^{\rho(\hat{1})-\rho(t \vee y)}$$

e osservando che $t \vee y$ è il prodotto $t \cdot y$ in $A_\vee(L, \mathbb{K})$ la dimostrazione è conclusa. $\square$

Dal Teorema 3.3.4 si evince che se z è un elemento modulare di un reticolo geometrico finito L , il polinomio caratteristico dell'intervallo $[\hat{0}, z]$ divide il

polinomio caratteristico di L . Un corollario di tale teorema è dovuto alla modularità degli atomi.

Corollario 3.3.5. Sia L un reticolo geometrico di rango n e sia a un atomo di L . Allora

$$\chi_L(\lambda) = (\lambda - 1) \sum_{y \wedge a = \hat{0}} \mu(\hat{0}, y) \lambda^{n-1-\rho(y)}.$$

Dimostrazione. Segue applicando il Teorema 3.3.4, essendo a modulare e $\chi_{[\hat{0}, a]}(\lambda) = \lambda - 1$. $\square$

Esiste una classe di reticoli il cui polinomio caratteristico si fattorizza in fattori lineari.

Definizione 3.3.6. Un reticolo geometrico L si dice *supersolubile* se esiste una *catena modulare massimale* ossia una catena massimale

$$\hat{0} = x_0 \leq x_1 \leq \dots \leq x_n = \hat{1}$$

tale che ogni x_i è modulare.

Teorema 3.3.7. Sia L un reticolo geometrico supersolubile di rango n , con catena modulare massimale $\hat{0} = x_0 \leq x_1 \leq \dots \leq x_n = \hat{1}$. Indichiamo con X_L l'insieme degli atomi di L e poniamo

$$e_i = \# \{a \in X_L : a \leq x_i, a \not\leq x_{i-1}\}.$$

Allora $\chi_L(\lambda) = (\lambda - e_1)(\lambda - e_2) \dots (\lambda - e_n)$.

Dimostrazione. Con un ragionamento analogo a quello fatto nella dimostrazione della Proposizione 2.4.8 e sfruttando la modularità di x_{n-1} , abbiamo che $y \wedge x_{n-1} = \hat{0}$ se e solo se $y = \hat{0}$ oppure $y \in X_L$ e $y \not\leq x_{n-1}$. Dunque per il Teorema 3.3.4 si ha

$$\chi_L(\lambda) = \chi_{[\hat{0}, x_{n-1}]}(\lambda) \cdot \left[\sum_{\substack{y \in X_L \\ y \not\leq x_{n-1}}} \mu(\hat{0}, y) \lambda^{n-\rho(x_{n-1})-\rho(y)} + \mu(\hat{0}, \hat{0}) \lambda^{n-\rho(x_{n-1})-\rho(\hat{0})} \right].$$

Siccome $\mu(\hat{0}, y) = -1$, $\mu(\hat{0}, \hat{0}) = 1$, $\rho(x_{n-1}) = n - 1$, $\rho(y) = 1$ e $\rho(\hat{0}) = 0$, l'espressione tra parentesi è uguale a $\lambda - e_n$. Possiamo applicare lo stesso ragionamento al sottoreticolo $[\hat{0}, x_{n-1}]$, dato che x_{n-2} è un elemento modulare di $[\hat{0}, x_{n-1}]$. Iterando questo procedimento si ottiene la tesi. $\square$

Esempio 3.3.8. Vediamo alcuni esempi di reticoli supersolubili.

- Nell'Esempio 3.3.3 abbiamo osservato che se S è un insieme di cardinalità n allora ogni elemento di $P(S)$ è modulare, perciò $P(S)$ è supersolubile. Inoltre, data una catena modulare massimale

$$\hat{0} = S_0 \leq S_1 \leq \dots \leq S_n = \hat{1},$$

è evidente che $e_i = 1$ per ogni i . Si ha quindi $\chi_{P(S)}(\lambda) = (\lambda - 1)^n$.

- Come mostreremo in seguito, i reticoli $V_n(q)$ e $\Pi(S)$ trattati rispettivamente nelle sezioni 4.1 e 4.3 sono entrambi supersolubili.

Capitolo 4

Funzione di Möbius su poset notevoli

In questo capitolo utilizzeremo gli strumenti introdotti precedentemente, in particolare nel Capitolo 2, per calcolare in maniera agevole il valore della funzione di Möbius su alcuni poset di grande interesse. Gli esempi riportati sono tratti da [SO97].

4.1 Sottospazi di uno spazio vettoriale finito dimensionale su campo finito

Definizione 4.1.1. Sia $\mathbb{F}$ un campo finito con q elementi, e sia V un $\mathbb{F}$ -spazio vettoriale di dimensione n . Indichiamo con $V_n(q)$ l'insieme dei sottospazi vettoriali di V , ordinato parzialmente per inclusione: dati U, W sottospazi di V , $U \leq W$ se $U \subseteq W$.

Esempio 4.1.2. In Figura 4.1 è riportato il diagramma di Hasse di $V_2(3)$, il reticolo dei sottospazi vettoriali di $\mathbb{F}_3^2$.

Osservazione 4.1.3. Il join tra due elementi $U, W \in V_n(q)$ corrisponde alla somma dei due spazi vettoriali, $U \vee W = U + W$, mentre il meet alla loro intersezione, $U \wedge W = U \cap W$.

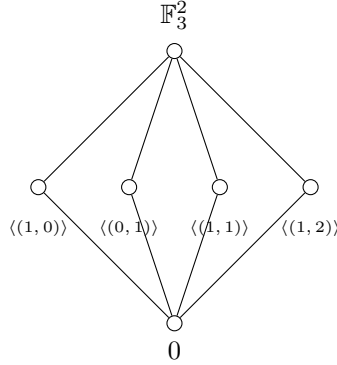


Figura 4.1: Diagramma di Hasse del reticolo $V_2(3)$.

Osservazione 4.1.4. Se $U_1, U_2 \in V_n(q)$, $U_1 \subseteq U_2$, allora vi è una corrispondenza biunivoca tra i sottospazi di U_2 che contengono U_1 e i sottospazi dello spazio quoziente U_2/U_1 . In altre parole, se indichiamo la dimensione di tale spazio con $r = \dim(U_2/U_1)$, l'intervallo $[U_1, U_2] \subseteq V_n(q)$ è isomorfo a $V_r(q)$ e dunque $\mu_{V_n(q)}(U_1, U_2) = \mu_{V_r(q)}(\hat{0}, \hat{1})$.

L'Osservazione 4.1.4 garantisce che è sufficiente calcolare $\mu_{V_r(q)}(\hat{0}, \hat{1})$ per ogni intero $m > 0$ per conoscere il valore della funzione di Möbius in ogni coppia di elementi di $V_n(q)$.

Se $m = 1$, allora $V_1(q)$ è isomorfo alla catena di lunghezza 1, perciò $\mu_{V_1(q)}(\hat{0}, \hat{1}) = -1$. Sia V uno spazio vettoriale su $\mathbb{F}$ di dimensione $m \geq 2$ e sia X_0 un sottospazio 1-dimensionale di V . Per il Teorema 2.3.3, vale

$$\sum_{U \vee X_0 = V} \mu_{V_m(q)}(\hat{0}, U) = 0. \quad (4.1)$$

Un sottospazio U di V soddisfa $X_0 \vee U = V$ se $U = V$ oppure se U ha dimensione $m - 1$ e $X_0 \cap U = \{0\}$, quindi la (4.1) diventa

$$\mu_{V_m(q)}(\hat{0}, \hat{1}) = - \sum_{U \oplus X_0 = V} \mu_{V_m(q)}(\hat{0}, U) = -N(m-1, X_0) \cdot \mu_{V_{m-1}(q)}(\hat{0}, \hat{1}), \quad (4.2)$$

dove $N(m-1, X_0)$ indica il numero di sottospazi $(m-1)$ -dimensionali di V in somma diretta con X_0 . Tale numero può essere calcolato sfruttando il risultato enunciato nella prossima proposizione.

Definizione 4.1.5. Dati due interi non negativi n e k , chiamiamo *coefficiente Gaussiano* o *coefficiente q -binomiale* il numero

$$\begin{bmatrix} n \\ k \end{bmatrix}_q = \frac{(q^n - 1)(q^{n-1} - 1) \dots (q^{n-(k-1)} - 1)}{(q^k - 1)(q^{k-1} - 1) \dots (q - 1)}.$$

Si osserva facilmente che $\begin{bmatrix} n \\ k \end{bmatrix}_q = \begin{bmatrix} n \\ n-k \end{bmatrix}_q$. Una trattazione più approfondita di questo oggetto si trova in [Sta12]

Proposizione 4.1.6. Sia $\mathbb{F}$ il campo con q elementi e sia V un $\mathbb{F}$ -spazio vettoriale di dimensione m . Allora il numero di sottospazi di V di dimensione k è $\begin{bmatrix} m \\ k \end{bmatrix}_q$. Se Y è un sottospazio di V di dimensione r , il numero di sottospazi Z di V di dimensione s tali che $Z \cap Y = \{0\}$ è $q^{rs} \begin{bmatrix} m-r \\ s \end{bmatrix}_q$.

Dimostrazione. Per trovare i sottospazi di V di dimensione k è necessario trovare k vettori linearmente indipendenti. Consideriamo un primo vettore $v_1 \in V$, $v_1 \neq 0$; vi sono $q^n - 1$ possibili scelte per v_1 . Ora cerchiamo un vettore v_2 che sia linearmente indipendente da v_1 , ossia $v_2 \notin \text{Span}\{v_1\}$; vi sono dunque $q^n - q$ possibilità per v_2 . Iterando questo procedimento si ottengono k vettori linearmente indipendenti $v_1, v_2, \dots, v_k$; ciò può essere fatto in

$$(q^n - 1)(q^n - q) \dots (q^n - q^{k-1})$$

modi diversi. Se Z è lo span di $\{v_1, v_2, \dots, v_k\}$, attraverso un calcolo analogo si conclude che esistono

$$(q^k - 1)(q^k - q) \dots (q^k - q^{k-1})$$

basi ordinate di Z . Segue che il numero di sottospazi k -dimensionali di V è proprio $\begin{bmatrix} m \\ k \end{bmatrix}_q$.

Consideriamo adesso un sottospazio Y di V di dimensione r e cerchiamo di contare i sottospazi s -dimensionali Z di V tali che $Z \cap Y = \{0\}$. Cerchiamo dunque s vettori linearmente indipendenti il cui span abbia intersezione banale con Y . Un primo vettore z_1 può essere scelto in $q^m - q^r$ modi. Il secondo vettore z_2 non deve appartenere al sottospazio generato da $Y \cup \{z_1\}$, che ha dimensione $r + 1$; dunque ci sono $q^m - q^{r+1}$ possibili scelte per z_2 . Iterando questo procedimento otteniamo s vettori $z_1, z_2, \dots, z_s$ che soddisfano quanto richiesto, e vi sono

$$(q^m - q^r)(q^m - q^{r+1}) \dots (q^m - q^{r+s-1})$$

modi per sceglierli. Attraverso considerazioni analoghe a quelle del paragrafo precedente si conclude che il numero di sottospazi di V di dimensione s aventi intersezione banale con Y è

$$\frac{(q^m - q^r)(q^m - q^{r+1}) \dots (q^m - q^{r+s-1})}{(q^s - 1)(q^s - q) \dots (q^s - q^{s-1})} = q^{rs} \begin{bmatrix} m - r \\ s \end{bmatrix}_q.$$

□

Dalla Proposizione 4.1.6 segue che

$$N(m - 1, X_0) = q^{m-1} \begin{bmatrix} m - 1 \\ m - 1 \end{bmatrix}_q = q^{m-1},$$

da cui

$$\mu_{V_m(q)}(\hat{0}, \hat{1}) = -q^{m-1} \cdot \mu_{V_{m-1}(q)}(\hat{0}, \hat{1}).$$

Iterando questo procedimento si ottiene il notevole risultato

$$\mu_{V_m(q)}(\hat{0}, \hat{1}) = (-1)^m \cdot q^{1+2+\dots+(m-1)} = (-1)^m \cdot q^{\binom{m}{2}}.$$

In sintesi, presi $U_1, U_2 \in V_n(q)$,

$$\mu_{V_n(q)}(U_1, U_2) = \begin{cases} (-1)^{\dim(U_2/U_1)} \cdot q^{\binom{\dim(U_2/U_1)}{2}} & \text{se } U_1 \subseteq U_2 \\ 0 & \text{altrimenti} \end{cases} \quad (4.3)$$

Mostriamo adesso un esempio di applicazione della formula (4.3).

Esempio 4.1.7. Sia $\mathbb{F}$ il campo con q elementi e sia V uno spazio vettoriale su $\mathbb{F}$. Ci chiediamo quanti sono i sottoinsiemi di V che generano V . Ai fini di questo esempio utilizzeremo la convenzione che l'insieme vuoto $\emptyset$ non genera alcuno spazio vettoriale, mentre l'insieme $\{0\}$ genera lo spazio vettoriale zero dimensionale $\{0\}$.

Se $W \in V_n(q)$, poniamo $f(W)$ uguale al numero di sottoinsiemi di V che generano W , e $g(W)$ il numero dei sottoinsiemi di V il cui span è contenuto in W . Si osserva facilmente che $g(W) = 2^{q^{\dim W}} - 1$, dato che $\emptyset$ non genera nessuno spazio. Inoltre vale

$$g(W) = \sum_{T \leq W} f(T),$$

e per il teorema 1.2.4

$$f(W) = \sum_{T \leq W} g(T) \mu_{V_n(q)}(T, W).$$

Ponendo $W = V$ si ottiene

$$f(V) = \sum_{T \leq V} g(T) \mu_{V_n(q)}(T, V) = \sum_{k=0}^n \binom{n}{k} (-1)^{n-k} q^{\binom{n-k}{2}} (2^{q^k} - 1).$$

Osservazione 4.1.8. Il reticolo $V_n(q)$ è atomico, dato che ogni spazio vettoriale è join dei suoi sottospazi di dimensione 1; inoltre $V_n(q)$ è graduato, con funzione grado data da $\rho_{V_n(q)}(U) = \dim(U)$. Per ogni $U, W \in V_n(q)$, vale

$$\dim(U) + \dim(W) = \dim(U \cup W) + \dim(U \cap W),$$

perciò ogni elemento di $V_n(q)$ è modulare. Dunque $V_n(q)$ è supersolubile e per il Teorema 3.3.7 il suo polinomio caratteristico si scompone in fattori lineari. Consideriamo una catena modulare massimale

$$\hat{0} = V_0 \leq V_1 \leq \dots \leq V_n = \hat{1},$$

indichiamo con $X_{V_n(q)}$ l'insieme degli atomi di $V_n(q)$ e per ogni i poniamo

$$e_i = \# \{a \in X_{V_n(q)} : U \leq V_i, U \not\leq V_{i-1}\}.$$

Per la Proposizione 4.1.6,

$$e_i = \begin{bmatrix} i \\ 1 \end{bmatrix}_q - \begin{bmatrix} i-1 \\ 1 \end{bmatrix}_q = \frac{q^i - 1}{q - 1} - \frac{q^{i-1} - 1}{q - 1} = q^{i-1}$$

e dunque

$$\chi_{V_n(q)}(\lambda) = (\lambda - 1)(\lambda - q)(\lambda - q^2) \dots (\lambda - q^{n-1}).$$

4.2 Sottomoduli di un modulo di torsione

Definizione 4.2.1. Sia R un dominio a ideali principali e sia M un modulo su R . Diciamo che $m \in M$ è un *elemento di torsione* se esiste $r_m \in R$ tale che $r_m \cdot m = 0$. Un modulo su R si dice *di torsione* se ogni suo elemento è di torsione.

Sia R un dominio a ideali principali e sia M un modulo di torsione finito su R . L'insieme

$$\{r \in R : r \cdot m = 0 \quad \forall m \in M\}$$

è un ideale di R , perciò è generato da un certo elemento $r_0 \in R$, unico a meno di moltiplicazione per un invertibile. Chiamiamo r_0 *esponente* di M . Sia $r_0 = p_1^{\alpha_1} p_2^{\alpha_2} \dots p_k^{\alpha_k}$ la fattorizzazione di r_0 in potenze di primi distinti e per $1 \leq i \leq k$ poniamo

$$M_{p_i} = \{m \in M : p_i^{\alpha_i} \cdot m = 0\}.$$

Proposizione 4.2.2. Il modulo M si spezza nella somma diretta

$$M = M_{p_1} \oplus M_{p_2} \oplus \dots \oplus M_{p_k}.$$

Dimostrazione. Vogliamo mostrare che ogni elemento $m \in M$ si scrive in modo unico come $m = m_1 + m_2 + \dots + m_k$ con $m_i \in M_{p_i}$. Per ogni i poniamo

$q_i = \frac{r_0}{p_i^{\alpha_i}}$. Poiché q_i e $p_i^{\alpha_i}$ sono coprimi, esistono $a_i, b_i \in R$ tali che

$$a_i q_i + b_i p_i^{\alpha_i} = 1,$$

dunque, ponendo $e_i = a_i q_i$, si ha

$$e_i \equiv 1 \pmod{p_i^{\alpha_i}}, \quad e_i \equiv 0 \pmod{p_j^{\alpha_j}} \text{ per } j \neq i.$$

Sommando tutti gli e_i si ha

$$\begin{aligned} \sum_i e_i &\equiv 1 \pmod{p_1^{\alpha_1}} \\ \sum_i e_i &\equiv 1 \pmod{p_2^{\alpha_2}} \\ &\dots \\ \sum_i e_i &\equiv 1 \pmod{p_k^{\alpha_k}}, \end{aligned}$$

e per il teorema cinese del resto

$$\sum_i e_i \equiv 1 \pmod{r_0}.$$

Ora, se $m \in M$, allora

$$\left(\sum_i e_i \right) \cdot m = (1 + kr_0) \cdot m = m + kr_0 \cdot m = m.$$

Perciò $m = \sum_i (e_i \cdot m)$ con $e_i \cdot m \in M_{p_i}$, visto che

$$p_i^{\alpha_i} \cdot (e_i \cdot m) = (p_i^{\alpha_i} e_i) \cdot m = (a_i r_0) \cdot m = 0.$$

Per provare l'unicità di tale scrittura mostriamo che $M_{p_j} \cap M_{p_l} = \{0\}$ per ogni $j \neq l$. Sia $m \in M_{p_j} \cap M_{p_l}$, ovvero $p_j^{\alpha_j} \cdot m = 0 = p_l^{\alpha_l} \cdot m$. Ora, poiché $p_j^{\alpha_j}$ e $p_l^{\alpha_l}$ sono coprimi, esistono $a, b \in R$ tali che $ap_j^{\alpha_j} + bp_l^{\alpha_l} = 1$. Moltiplicando

per m si ha

$$m = (ap_j^{\alpha_j} + bp_l^{\alpha_l}) \cdot m = ap_j^{\alpha_j} \cdot m + bp_l^{\alpha_l} \cdot m = 0$$

e ciò conclude la dimostrazione. $\square$

Definizione 4.2.3. Sia R un dominio a ideali principali e sia M un modulo di torsione finito su R . Indichiamo con $L_R(M)$ l'insieme degli R -sottomoduli di M , ordinati parzialmente per inclusione, ovvero $M_1 \leq M_2$ se $M_1 \subseteq M_2$.

Osservazione 4.2.4. Il poset $L_R(M)$ è un reticolo: dati due sottomoduli $M_1, M_2 \in L_R(M)$, si ha $M_1 \wedge M_2 = M_1 \cap M_2$ e $M_1 \vee M_2 = M_1 + M_2$. Inoltre esistono l'elemento minimale $\hat{0} = \{0\}$ e quello massimale $\hat{1} = M$.

Osservazione 4.2.5. La Proposizione 4.2.2 si applica anche ai sottomoduli N di M , ovvero

$$N = N_{p_1} \oplus N_{p_2} \oplus \cdots \oplus N_{p_k},$$

con ogni N_{p_i} sottomodulo di M_{p_i} . Perciò il reticolo $L_R(M)$ è un prodotto di reticoli

$$L_R(M) \simeq L_R(M_{p_1}) \times L_R(M_{p_2}) \times \cdots \times L_R(M_{p_k})$$

e applicando il Teorema 1.2.5 si ha

$$\mu_{L_R(M)}(A, B) = \prod_{i=1}^k \mu_{L_R(M_{p_i})}(A_{p_i}, B_{p_i}).$$

L'Osservazione 4.2.5 ci dice che per calcolare la funzione di Möbius su $L_R(M)$ è sufficiente conoscerne il valore nel caso in cui l'esponente di M è p^n con $p \in R$ primo e n intero positivo. Inoltre, se M è di questo tipo e $A, B \in L_R(M)$ con $A \subseteq B$ allora $[A, B] \simeq L_R(B/A)$, dunque ci basta calcolare $\mu_{L_R(M)}(\hat{0}, \hat{1})$ per ogni modulo finito M con esponente p^n .

Cominciamo considerando il caso $n = 1$. Per ogni $m \in M$ sappiamo che $p \cdot m = 0$, quindi possiamo considerare M come modulo su R/pR . Poiché R è un dominio a ideali principali e p è primo, R/pR è un campo e dunque M è uno spazio vettoriale su tale campo. Essendo M finito, è uno spazio

vettoriale di dimensione finita d su campo finito, poniamo di ordine q . Allora vale l'isomorfismo $L_R(M) \simeq V_d(q)$ e per quanto visto nella sezione 4.1

$$\mu_{L_R(M)}(\hat{0}, \hat{1}) = (-1)^d q^{\binom{q}{2}}.$$

Supponiamo adesso $n \geq 2$. In tal caso, esisterà $m_0 \in M$ tale che $p^2 \cdot m_0 = 0$ ma $p \cdot m_0 \neq 0$. Poniamo $m_1 = p \cdot m_0$ e sia A il sottomodulo di M generato da m_1 . Possiamo supporre che A sia minimale, ovvero non abbia sottomoduli propri. In caso contrario, sia A_1 un sottomodulo minimale non nullo di A . Allora A_1 è generato da un elemento $m_2 = r_2 \cdot m_1$ per un qualche $r_2 \in R$. Sia $m'_0 = r_2 \cdot m_0$. Poiché r_2 e p sono coprimi, si ha $p^2 \cdot m'_0 = 0$ ma $p \cdot m'_0 \neq 0$. Essendo il generatore di A_1 $m_2 = p \cdot m'_0$, possiamo sostituire A con A_1 . Sia dunque A un sottomodulo minimale non nullo di M ; per il Teorema 2.3.3,

$$\sum_{A \vee B = M} \mu_{L_R(M)}(\hat{0}, B) = 0. \quad (4.4)$$

Se B è un sottomodulo proprio di M tale che $A \vee B = A + B = M$, allora per il secondo teorema di omomorfismo di moduli

$$M/B = (A + B)/B \simeq A/(A \cap B) = A,$$

con l'ultima disuguaglianza dovuta alla minimalità di A . In particolare, $p \cdot (M/B) = 0$. D'altra parte, se consideriamo l'omomorfismo di proiezione al quoziente

$$\pi : M \rightarrow M/B$$

abbiamo che $\pi(m_1) = \pi(p \cdot m_0) = p \cdot \pi(m_0) = 0$, da cui $m_1 \in B$. Ma $m_1 \in A$ e ciò contraddice $A \cap B = \{0\}$. Dunque l'equazione (4.4) diventa

$$\mu_{L_R(M)}(\hat{0}, \hat{1}) = 0.$$

In conclusione, abbiamo mostrato che se $A \subseteq B$ sono sottomoduli di M , allora $\mu_{L_R(M)}(A, B) = 0$ a meno che B/A non sia uno spazio vettoriale su R/pR .

4.3 Reticolo delle partizioni di un insieme

Definizione 4.3.1. Sia S un insieme finito. Una *partizione* π di S è una collezione di sottoinsiemi $A_1, A_2, \dots, A_n$ di S non vuoti e a due a due disgiunti, tali che $A_1 \cup A_2 \cup \dots \cup A_n = S$. Tali sottoinsiemi sono chiamati *blocchi* della partizione. Scriviamo $N(\pi) = n$ se π ha n blocchi.

Definizione 4.3.2. Date due partizioni π_1 e π_2 di S , scriviamo $\pi_1 \leq \pi_2$ se ogni blocco di π_1 è contenuto in un blocco di π_2 . L'insieme delle partizioni di S , con questa relazione d'ordine, è un reticolo, che viene indicato con $\Pi(S)$.

Esempio 4.3.3. In Figura 4.2 è riportato il diagramma di Hasse di $\Pi(\{1, 2, 3\})$.

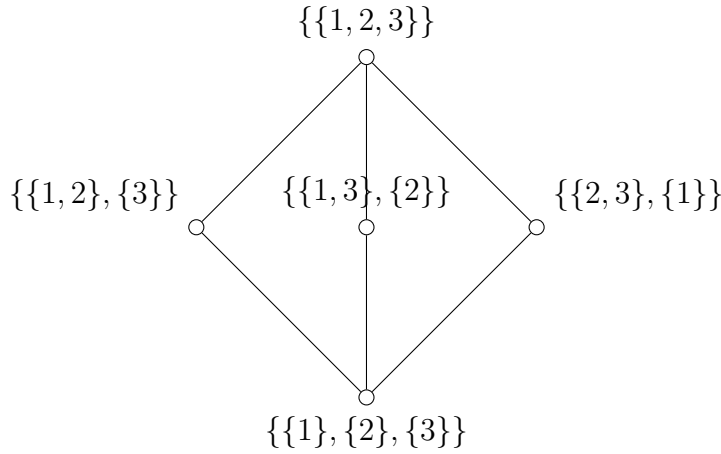


Figura 4.2: Diagramma di Hasse del reticolo delle partizioni di $\{1, 2, 3\}$.

Osservazione 4.3.4. Siano $\pi_1, \pi_2 \in \Pi(S)$, aventi rispettivamente blocchi $A_1, A_2, \dots, A_n$ e $B_1, B_2, \dots, B_m$. Allora $\pi_1 \wedge \pi_2$ è la partizione con blocchi $\{A_i \cap B_j\}_{i=1,2,\dots,n, j=1,2,\dots,m}$. La descrizione del join richiede maggiore attenzione ed è data dalla seguente proposizione.

Proposizione 4.3.5. Siano $\pi_1, \pi_2 \in \Pi(S)$, aventi rispettivamente blocchi $A_1, A_2, \dots, A_n$ e $B_1, B_2, \dots, B_m$. Allora esiste $\pi_1 \vee \pi_2$; inoltre due elementi $x, y \in S$ appartengono allo stesso blocco di $\pi_1 \vee \pi_2$ se e solo se esiste una catena di elementi $x = s_1, s_2, \dots, s_k = y$ tale che per ogni intero $1 \leq i \leq k-1$, s_i e s_{i+1} stanno nello stesso blocco di π_1 o nello stesso blocco di π_2 .

Dimostrazione. Sia λ un maggiorante di π_1 e π_2 e siano $x, y \in S$ tali per cui esiste una catena del tipo descritto sopra. Allora per ogni $1 \leq i \leq k-1$, esiste un blocco W_i di π_1 o di π_2 che contiene s_i e s_{i+1} . Siccome $\pi_1, \pi_2 \leq \lambda$, ciascun W_i è contenuto in un blocco di λ . Ma per ogni i , $\emptyset \neq \{s_{i+1}\} \subseteq W_i \cap W_{i+1}$, ed essendo i blocchi di una partizione a due a due disgiunti, l'unica possibilità è che tutti i W_i siano contenuti nello stesso blocco di λ . Dunque x, y appartengono entrambi a tale blocco.

Consideriamo ora la partizione η definita dalla seguente condizione: due elementi $x, y \in S$ appartengono allo stesso blocco di η se e solo se esiste una catena di elementi $x = s_1, s_2, \dots, s_k = y$ tale che per ogni intero $1 \leq i \leq k-1$, s_i e s_{i+1} stanno nello stesso blocco di π_1 o nello stesso blocco di π_2 . Mostriamo che η è effettivamente una partizione. Siano C, D due blocchi di η con intersezione non vuota, cioè esiste $y \in C \cap D$. Siano $x \in C, z \in D$. Poiché x e y stanno nello stesso blocco di η , esiste una catena di elementi $x = s_1, s_2, \dots, s_k = y$ con le proprietà enunciate sopra; analogamente esiste una catena $y = s_k, s_{k+1}, \dots, s_{k+r} = z$ con le stesse caratteristiche. Otteniamo dunque una catena $x = s_1, s_2, \dots, s_k = y, s_{k+1}, \dots, s_{k+r} = z$ che collega x e z , perciò questi due elementi stanno nello stesso blocco di η . Per l'arbitrarietà di x e z concludiamo che $C = D$.

Si osserva facilmente che $\pi_1, \pi_2 \leq \eta$, ossia che η è un maggiorante di π_1 e π_2 . Da quanto osservato nel primo paragrafo della dimostrazione concludiamo $\eta = \pi_1 \vee \pi_2$. $\square$

L'elemento massimale è la partizione composta da un solo blocco, l'elemento minimale è la partizione composta da $|S|$ singoletti.

Osservazione 4.3.6. Se due insiemi hanno stessa cardinalità, i reticoli delle loro partizioni sono isomorfi. A meno di isomorfismo denotiamo $\Pi(n)$ il reticolo delle partizioni di un insieme con n elementi.

Il nostro obiettivo è quello di determinare la funzione di Möbius per $\Pi(n)$. Cominciamo col calcolare $\mu_{\Pi(n)}(\hat{0}, \hat{1})$. Sia $S = \{s_1, s_2, \dots, s_n\}$ e sia $\pi_0 \in \Pi(S)$ data da $\{\{s_1, s_2, \dots, s_{n-1}\}, \{s_n\}\}$. Allora $\pi \in \Pi(S)$ è tale che $\pi \wedge \pi_0 = \hat{0}$ se e solo se $\pi = \hat{0}$ oppure π ha tutti i blocchi di cardinalità 1 eccetto quello contenente s_n , che ha due elementi. Esistono esattamente $n-1$ partizioni

fatte in questo modo.

Se $\pi_i = \{\{s_1\}, \{s_2\}, \dots, \{s_{i-1}\}, \{s_i, s_n\}, \{s_{i+1}\}, \dots, \{s_{n-1}\}\}$, allora

$$[\pi_i, \hat{1}] \simeq \Pi(n-1).$$

Applicando la versione duale del Teorema 2.3.3, si ottiene

$$\mu_{\Pi(n)}(\hat{0}, \hat{1}) = -(n-1) \cdot \mu_{\Pi(n-1)}(\hat{0}, \hat{1}).$$

Iterando e tenendo a mente che $\mu_{\Pi(2)}(\hat{0}, \hat{1}) = -1$ concludiamo che

$$\mu_{\Pi(n)}(\hat{0}, \hat{1}) = (-1)^{n-1} \cdot (n-1)!.$$

Consideriamo adesso il caso generale. Siano $\pi_1, \pi_2 \in \Pi(S)$ con $\pi_1 \leq \pi_2$. Sia $\pi_2 = \{B_1, B_2, \dots, B_r\}$, allora ogni blocco B_i è unione disgiunta di blocchi di π_1 . Supponiamo che per $i = 1, 2, \dots, r$, B_i sia unione di n_i blocchi di π_1 , e quindi $N(\pi_2) = \sum_{i=1}^r n_i$. Allora

$$[\pi_1, \pi_2] \simeq \Pi(n_1) \times \dots \times \Pi(n_r).$$

Applicando la Proposizione 1.2.5, otteniamo

$$\mu_{\Pi(n)}(\pi_1, \pi_2) = \prod_{i=1}^r \mu_{\Pi(n_i)}(\hat{0}, \hat{1}) = \prod_{i=1}^r (-1)^{n_i-1} (n_i-1)! = (-1)^{N(\pi_1)-N(\pi_2)} \prod_{i=1}^r (n_i-1)!.$$

Concludiamo studiando gli elementi modulari di $\Pi(n)$ e dimostrando che tale reticolo è supersolubile.

Osservazione 4.3.7. Gli atomi del reticolo $\Pi(n)$ sono le partizioni costituite da tutti singoletti, eccetto un unico blocco di cardinalità 2. Ce ne sono esattamente $\binom{n}{2}$. Osserviamo ora che $\Pi(n)$ è un reticolo atomico. Infatti, sia $\pi \in \Pi(n)$ una partizione qualsiasi. Per ogni blocco $B \in \pi$ con $|B| \geq 2$, consideriamo tutti gli atomi che uniscono due elementi di B (e lasciano gli altri come singoletti). Il join dell'insieme di questi atomi, presi su tutti i blocchi di π , è esattamente π . Inoltre $\Pi(n)$ è graduato di rango $n-1$ con

funzione grado $\rho_{\Pi(n)}(\pi) = n - |\pi|$, dove $|\pi|$ indica il numero di blocchi di π .

Proposizione 4.3.8. Una partizione $\pi \in \Pi(n)$ è un elemento modulare di $\Pi(n)$ se e solo se π ha al più un blocco di cardinalità maggiore o uguale a 2. Perciò il numero di elementi modulari di $\Pi(n)$ è $2^n - n$.

Dimostrazione. Se tutti i blocchi di π sono singoletti, allora $\pi = \hat{0}$ che è modulare. Supponiamo che π abbia un blocco A con $r > 1$ elementi e tutti gli altri blocchi siano singoletti. Allora $|\pi| = n - r + 1$. Per ogni $\sigma \in \Pi(n)$, si ha $\rho(\sigma) = n - |\sigma|$. Poniamo $k = |\sigma|$ e

$$j = \# \{B \in \sigma : A \cap B \neq \emptyset\}.$$

Si ha che $|\pi \wedge \sigma| = j + (n - r)$ e $|\pi \vee \sigma| = k - j + 1$, da cui $\rho(\pi) = r - 1$, $\rho(\sigma) = n - k$, $\rho(\pi \wedge \sigma) = r - j$ e $\rho(\pi \vee \sigma) = n - k + j - 1$. Perciò π è modulare.

Viceversa, supponiamo che $\pi = \{B_1, B_2, \dots, B_k\}$ con $|B_1|, |B_2| > 1$. Siano $a \in B_1$, $b \in B_2$ e poniamo

$$\sigma = \{(B_1 \cup \{b\}) \setminus \{a\}, (B_2 \cup \{a\}) \setminus \{b\}, B_3, \dots, B_k\}.$$

Osserviamo che $|\sigma| = |\pi| = k$; inoltre

$$\pi \wedge \sigma = \{\{a\}, \{b\}, B_1 \setminus \{a\}, B_2 \setminus \{b\}, B_3, \dots, B_k\},$$

$$\pi \vee \sigma = \{B_1 \cup B_2, B_3, \dots, B_k\}$$

e dunque $|\pi \wedge \sigma| = k + 2$, $|\pi \vee \sigma| = k - 1$. Ma

$$\rho(\pi) + \rho(\sigma) = 2n - 2k \neq 2n - 2k - 1 = \rho(\pi \wedge \sigma) + \rho(\pi \vee \sigma)$$

e quindi π non è modulare. □

Per la Proposizione 4.3.8, una catena massimale di $\Pi(n)$ è modulare se e solo se è della forma $\hat{0} = \pi_0 \leq \pi_1 \leq \dots \leq \pi_{n-1} = \hat{1}$, dove per ogni $i > 0$ π_i ha esattamente un blocco non singoletto B_i , di cardinalità $i + 1$, con $B_i \subseteq B_{i+1}$. In particolare $\Pi(n)$ è supersolubile, vediamo come si fattorizza il suo

polinomio caratteristico. Data una catena modulare massimale indicata come sopra, gli atomi di $\Pi(n)$ minori o uguali di π_i sono tutte e sole le partizioni con un unico blocco non singoletto $\{a, b\} \subseteq B_i$. Questi sono esattamente $\binom{i+1}{2}$, perciò utilizzando le notazioni del Teorema 3.3.7 si ha

$$e_i = \binom{i+1}{2} - \binom{i}{2} = i.$$

Applicando il Teorema 3.3.7 si ottiene

$$\chi_{\Pi(n)}(\lambda) = (\lambda - 1)(\lambda - 2) \dots (\lambda - n + 1).$$

Bibliografia

- [Gre71] Curtis Greene. “On the Möbius algebra of a partially ordered set”. In: *Möbius algebras (Proc. Conf., Univ. Waterloo, Waterloo, Ont., 1971)*. With two appendices by Henry Crapo and a reprint of an article by Louis Solomon. University of Waterloo, Waterloo, ON, 1971, pp. 3–38.
- [Rot64] Gian-Carlo Rota. “On the foundations of combinatorial theory. I. Theory of Möbius functions”. In: *Z. Wahrscheinlichkeitstheorie und Verw. Gebiete* 2 (1964).
- [SO97] Eugene Spiegel e Christopher J. O’Donnell. *Incidence algebras*. Vol. 206. Monographs and Textbooks in Pure and Applied Mathematics. Marcel Dekker, Inc., New York, 1997.
- [Sta07] Richard P. Stanley. “An introduction to hyperplane arrangements”. In: *Geometric combinatorics*. Vol. 13. IAS/Park City Math. Ser. Amer. Math. Soc., Providence, RI, 2007, pp. 389–496.
- [Sta12] Richard P. Stanley. *Enumerative combinatorics. Volume 1*. Second. Vol. 49. Cambridge Studies in Advanced Mathematics. Cambridge University Press, Cambridge, 2012.