

Capitolo B32

algebra lineare sui razionali

Contenuti delle sezioni

- a. spazi vettoriali di sequenze numeriche p. 2
- b. sottospazi p. 9
- c. trasformazioni lineari e loro matrici p. 13
- d. forme lineari e notazioni alla Dirac p. 19
- e. rassegna di trasformazioni lineari e matrici in 2D e 3D p. 22
- f. cambiamenti di base e prodotti di trasformazioni lineari p. 23
- g. determinanti p. 24
- h. invertibilità di una matrice quadrata p. 29

29 pagine

B320.01 In questo capitolo studiamo gli spazi vettoriali costituiti dalle sequenze finite di una determinata lunghezza che chiamiamo “dimensione” le cui componenti sono numeri razionali o analoghi oggetti con valenza numerica.

Queste entità numeriche sono definite come le entità alle quali possono essere applicate due operazioni che godono di molte delle proprietà della somma e del prodotto dei numeri razionali.

Gli insiemi muniti di queste due operazioni, che chiamiamo sempre somma e prodotto, vengono chiamati “campi” e, grazie alle accennate proprietà, vengono considerati insiemi numerici.

Gli spazi vettoriali che affrontiamo sono generalizzazioni della retta razionale $\mathbb{Q}$ e del piano $\mathbb{Q} \times \mathbb{Q}$ visti nei capitoli B30 e B31.

Il campo dei razionali è l’unico campo che abbiamo potuto definire in modo costruttivo; più oltre definiremo altri campi, in particolare il campo dei numeri reali e il campo dei numeri complessi.

Nel seguito di questo capitolo ci serviremo di esempi riguardanti esclusivamente sequenze di numeri razionali.

Tuttavia assumiamo un primo atteggiamento rivolto all’astrazione facendo riferimento a un possibile campo generico che denoteremo con $\mathbb{F}$ e segnalando che le costruzioni e le proprietà che si vanno esponendo si possono applicare a ciascuno dei campi che si possono definire e anticipando che molti risultati di queste pagine sono validi per oggetti e situazioni assai più generali.

B32 a. spazi vettoriali di sequenze numeriche

B32a.01 Introduciamo la scrittura $\mathbb{Q}_{Fld}$ per denotare la struttura della specie campo costituita dall'insieme dei numeri razionali $\mathbb{Q}$ munito delle operazioni di somma, differenza e prodotto tra suoi elementi, dei suoi elementi particolari 0 e 1 e dalla divisione tra un suo elemento qualsiasi e un suo elemento diverso da 0.

Nel seguito con d , e ed f denotiamo numeri interi positivi e rivolgiamo l'attenzione a $\mathbb{Q}^{\times d}$, la potenza cartesiana d -esima di $\mathbb{Q}$.

Prenderemo in considerazione anche una struttura generica, non precisamente individuata, che denotiamo con $\mathbb{F}_{Fld}$ e chiamiamo “campo” che intende generalizzare $\mathbb{Q}_{Fld}$ la quale è costituita da un insieme $\mathbb{F}$ costituente il terreno del campo e da operazioni ed elementi con ruoli formalmente molto simili a quelli che costituiscono il campo dei numeri razionali.

Denotiamo inoltre con $\mathbb{F}^{\times d}$ la potenza cartesiana d -esima della struttura $\mathbb{F}_{Fld}$, ossia l'insieme delle d -uple di elementi di $\mathbb{F}$ (che denotiamo con scritture della forma $\mathbf{a} = \langle a_1, a_2, \dots, a_d \rangle$) munito delle estensioni cartesiane dei componenti del campo $\mathbb{F}_{Fld}$.

Si osserva che se in particolare $\mathbb{F} = \mathbb{Q}$, per $d = 1$ si ha la retta razionale $\mathbb{Q}$, per $d = 2$ il piano $\mathbb{Q} \times \mathbb{Q}$ e per $d = 3$ lo spazio tridimensionale $\mathbb{Q}^{\times 3}$.

Dobbiamo anche avvertire che spesso la distinzione tra un campo e il suo terreno viene trascurata confidando che il contesto annulli i rischi di incomprensioni che possono derivare dalla identificazione di un $\mathbb{F}$ con $\mathbb{F}_{Fld}$.

Aggiungiamo che nel seguito, come spesso accade, chiamiamo **scalari** gli elementi di un campo e diciamo **vettori d -dimensionali** o **punti d -dimensionali** le d -uple di tali scalari.

B32a.02 I vettori si incontrano in una grande varietà di argomentazioni nella matematica e nelle sue applicazioni e vengono denotati in vari modi: utilizzando lettere in grassetto come $\mathbf{x}$ o $\mathbf{b}_j$ e con simboli come $\underline{H}$, $\vec{E}_\perp$ o $|3\rangle$.

Qui presentiamo i vettori con scritture come $\mathbf{v} = \langle v_1, v_2, \dots, v_d \rangle$.

Le componenti di una tale sequenza si dicono anche **coordinate cartesiane del vettore $\mathbf{v}$** .

Consideriamo dunque un campo generico $\mathbb{F}_{Fld} = \langle \mathbb{F}, +, \cdot, -, /, 0, 1 \rangle$ e osserviamo di aver denotato le operazioni di somma, differenza, prodotto e divisione con gli stessi segni usati per i numeri razionali, mentre lo zero e l'unità sono denotati con segni leggermente diversi.

Aggiungiamo tuttavia che in molti contesti certe differenze e certe confusioni di segni non rivestono grande importanza; la precisione delle notazioni è invece importante per la redazione di testi che si vogliono gestire con strumenti automatici.

In ogni campo $\mathbb{F}_{Fld}$ l'elemento zero si distingue dagli altri elementi e conviene servirsi di una notazione particolare per l'insieme degli elementi di $\mathbb{F}$ diversi da zero: scriviamo quindi:

$$(1) \quad \mathbb{F}_{nz} := \mathbb{F} \setminus \{0\}.$$

Tra i vettori di $\mathbb{F}^{\times d}$ gioca invece un ruolo particolare la sequenza di d zeri del campo, ossia $0^{\times d}$ (scrittura abbreviata spesso con 0^d); questa sequenza viene chiamata **vettore zero** o **vettore nullo** di $\mathbb{F}^{\times d}$ e si denota con $\mathbf{0}_d$; come per il piano dei razionali, quando è lecito trascurare di segnalare la dimensione d abbreviamo $0^{\times d}$ con il solo $\mathbf{0}$.

Spesso interessa distinguere il vettore nullo da tutti gli altri vettori e conviene usare una notazione specifica per l'insieme dei vettori di $\mathbb{F}^{\times d}$ diversi da quello nullo: introduciamo quindi la notazione

$$(2) \quad \mathbb{F}^{\times d}_{nz} := \mathbb{F}^{\times d} \setminus \{\mathbf{0}_d\}.$$

B32a.03 A partire dalle operazioni di somma e prodotto del campo $\mathbb{F}$, definiamo l'operazione binaria **somma di due vettori** di $\mathbb{F}^{\times d}$ e la composizione **moltiplicazione di un vettore per uno scalare** di $\mathbb{F}^{\times d}$.

Consideriamo per questo due vettori di $\mathbb{F}^{\times d}$ $\mathbf{v} = \langle v_1, \dots, v_d \rangle$ e $\mathbf{w} = \langle w_1, \dots, w_d \rangle$ e lo scalare $\alpha \in \mathbb{F}$.

Si definisce vettore somma di $\mathbf{v}$ e $\mathbf{w}$ la somma termine a termine delle due sequenze, cioè

$$\mathbf{v} + \mathbf{w} := \langle v_1 + w_1, v_2 + w_2, \dots, v_d + w_d \rangle.$$

Si definisce moltiplicazione per lo scalare α del vettore $\mathbf{v}$ la sequenza

$$\alpha x \mathbf{v} := \mathbf{v} \alpha := \langle \alpha v_1, \alpha v_2, \dots, \alpha v_d \rangle.$$

In genere la precedente espressione si semplifica nella $\alpha \mathbf{v}$.

Si verifica facilmente che valgono le proprietà che seguono.

(1) Commutatività della somma:

$$\forall \mathbf{v}, \mathbf{w} \in \mathbb{F}^{\times d} : \quad \mathbf{v} + \mathbf{w} = \mathbf{w} + \mathbf{v}.$$

(2) Associatività della somma:

$$\forall \mathbf{v}, \mathbf{w}, \mathbf{u} \in \mathbb{F}^{\times d} : \quad \mathbf{v} + (\mathbf{w} + \mathbf{u}) = (\mathbf{v} + \mathbf{w}) + \mathbf{u}.$$

(3) Il vettore nullo $\mathbf{0}_d \in \mathbb{F}^{\times d}$ è elemento neutro per la somma:

$$\forall \mathbf{v} \in \mathbb{F}^{\times d} : \quad \mathbf{v} + \mathbf{0}_d = \mathbf{v} \quad \text{e quindi} \quad \mathbf{0}_d + \mathbf{v} = \mathbf{v}.$$

(4) Ad ogni vettore è associato un **vettore opposto** per la somma:

$$\forall \mathbf{v} \in \mathbb{F}^{\times d} : \quad \exists -\mathbf{v} \in \mathbb{F}^{\times d} \quad \text{tale che} \quad \mathbf{v} + (-\mathbf{v}) = (-\mathbf{v}) + \mathbf{v} = \mathbf{0}_d.$$

(5) Proprietà della moltiplicazione per uno scalare:

$$\begin{aligned} &\forall \mathbf{v}, \mathbf{w} \in \mathbb{F}^{\times d}, \quad \forall \alpha, \beta \in \mathbb{F} : \\ &\alpha(\mathbf{v} + \mathbf{w}) = \alpha \mathbf{v} + \alpha \mathbf{w}, \quad (\alpha + \beta) \mathbf{v} = \alpha \mathbf{v} + \beta \mathbf{v}, \quad (\alpha \beta) \mathbf{v} = \alpha(\beta \mathbf{v}), \quad 1 \mathbf{v} = \mathbf{v}, \quad 0 \mathbf{v} = \mathbf{0}_d. \end{aligned}$$

B32a.04 Dato un qualsiasi campo $\mathbb{F}$ i cui elementi chiamiamo “scalari” si dice **spazio vettoriale** su $\mathbb{F}$ una struttura che consiste in un insieme di entità chiamati “vettori”, nel campo $\mathbb{F}$, in una operazione binaria tra vettori chiamata somma, in una operazione chiamata moltiplicazione dei vettori per uno scalare che a un vettore associa un altro vettore, struttura per le cui componenti valgono le proprietà presentate nel paragrafo precedente.

Un insieme molto ampio di spazi vettoriali si ottiene considerando le funzioni che a ogni elemento di un certo insieme D associano un elemento di un campo $\mathbb{F}$; si definisce come somma di due di tali funzioni f e g la funzione che associa a ogni $x \in D$ la somma dei loro valori $f(x) + g(x)$ e si definisce come moltiplicazione della funzione f per l'elemento del campo $\alpha \in \mathbb{F}$ la funzione che ad x associa la funzione $\alpha f(x)$, proporzionale alla $f(x)$.

Si verifica facilmente che per tale spazio sono valide le proprietà a03(1-5).

Le suddette estensioni delle operazioni somma e moltiplicazione per scalari si dicono **estensioni funzionali delle operazioni** sul campo.

Si osserva che lo spazio vettoriale $\mathbb{F}^{\times d}$ delle sequenze di una data lunghezza d di elementi di $\mathbb{F}$ costituisce un caso particolare di spazio di funzioni con valori su un campo: infatti queste sequenze si possono considerare le funzioni aventi come dominio $\{1, 2, \dots, d\}$, cioè l'insieme delle posizioni delle sue componenti, e come codominio il terreno di $\mathbb{F}$.

B32a.05 Limitiamoci ora ai vettori dello spazio vettoriale d -dimensionale $\mathbb{F}^{\times d}$.

Due vettori nonnulli $\mathbf{v}$ e $\mathbf{w}$ si dicono **vettori proporzionali** o **vettori collineari** sse esiste un $\alpha \in \mathbb{F} \setminus \{0\}$ tale che $\mathbf{v} = \alpha \mathbf{w}$.

La proporzionalità tra vettori è evidentemente un'equivalenza su $\mathbb{F}^{\times d}_{nz}$ e le classi di questa equivalenza sono spesso chiamati **raggi dello spazio vettoriale**.

Per esempio allo stesso raggio di $\mathbb{F}^{\times 3}$ appartengono i vettori $\langle 1, 2, -3 \rangle$, $\langle 2.5, 5, -7.5 \rangle$ e $\langle -0.2, -0.4, 0.6 \rangle$.

Consideriamo l'intero positivo k e l'insieme di k vettori di $\mathbb{F}^{\times d}$ $\{\mathbf{v}_1, \mathbf{v}_2, \dots, \mathbf{v}_k\}$. Si dice **combinazione lineare dei vettori** ogni composizione della forma

$$\alpha_1 \mathbf{v}_1 + \alpha_2 \mathbf{v}_2 + \dots + \alpha_k \mathbf{v}_k =: \sum_{h=1}^k \alpha_h \mathbf{v}_h ,$$

dove per ogni $h = 1, 2, \dots, k$ α_h denota uno scalare.

Gli scalari che intervengono in una combinazione lineare sono detti **coefficienti della combinazione lineare**.

Di due vettori proporzionali si può dire che uno è combinazione lineare dell'altro.

Tra le combinazioni lineari di ogni insieme finito di vettori $\mathbf{w}_1, \dots, \mathbf{w}_k$ si trova il vettore nullo: infatti questo si può ottenere con ogni combinazione lineare della forma $\sum_{h=1}^k 0 \mathbf{w}_h = \mathbf{0}$.

Le notazioni precedenti per le combinazioni lineari sono le più usate, ma qui useremo spesso anche le notazioni equivalenti

$$\mathbf{v}_1 \alpha_1 + \mathbf{v}_2 \alpha_2 + \dots + \mathbf{v}_k \alpha_k = \sum_{h=1}^k \mathbf{v}_h \alpha_h ,$$

in accordo con l'aver considerate equivalenti le scritture $\alpha \mathbf{v}$ e $\mathbf{v} \alpha$.

La notazione che pone i coefficienti a destra dei vettori conduce in maniera più semplice al sistema di espressioni riguardante le rappresentazioni matriciali delle trasformazioni, come vedremo nel paragrafo che segue.

B32a.06 Come vedremo, molto spesso serve esprimere i vettori che si stanno esaminando come combinazioni lineari dei vettori di determinate collezioni finite di $\mathbb{F}^{\times d}$. In effetti risulta conveniente nella pratica di servirsi di determinate sequenze di vettori di riferimento. Vediamo quali di questi insiemi di riferimento sono i più convenienti.

Un insieme di vettori di $\mathbb{F}^{\times d}$ si dice costituire un **insieme di vettori linearmente indipendenti** sse non si può esprimere nessuno di essi come combinazione lineare degli altri.

Viceversa un insieme di vettori tale che almeno uno di essi si può esprimere come combinazione lineare dei rimanenti si dice **insieme di vettori linearmente dipendenti**.

Diamo alcuni esempi di insiemi di vettori linearmente dipendenti.

In $\mathbb{Q} \times \mathbb{Q}$: $\{\langle 0, 1 \rangle, \langle 1, 0 \rangle, \langle 2, -2 \rangle\}$ e $\{\langle 1, 1 \rangle, \langle 1, -1 \rangle, \langle a, b \rangle\}$ per a e b razionali arbitrari.

In $\mathbb{Q}^{\times 3}$: $\{\langle 0, 1, 0 \rangle, \langle 1, 0, 0 \rangle, \langle 2, -2, 0 \rangle, \langle 1, 2, 5 \rangle\}$ e $\{\langle 1, 0, 0 \rangle, \langle 0, 1, 0 \rangle, \langle 1, -1, 0 \rangle, \langle -1, -1, -1 \rangle\}$.

In qualsiasi dimensione $\{\mathbf{u}, \mathbf{u}c\}$, quali che siano $\mathbf{u} \in \mathbb{F}^{\times d}$ e $c \in \mathbb{Q}$, e $\{\mathbf{u}_1, \dots, \mathbf{u}_{k-1}, \sum_{h=1}^{k-1} \mathbf{u}_h d_h\}$, quali che siano $\mathbf{u}_1, \dots, \mathbf{u}_{k-1} \in \mathbb{F}^{\times d}$ e $d_1, \dots, d_{k-1} \in \mathbb{Q}$.

B32a.07 Si osserva che aggiungendo qualche vettore a un sistema di vettori linearmente dipendenti si ottiene ancora un sistema di vettori linearmente dipendenti.

All'opposto si osserva che eliminando qualche elemento da un sistema di vettori linearmente indipendenti si ottiene ancora un sistema di vettori linearmente indipendenti.

(1) Prop.: Un insieme finito di vettori $\{\mathbf{u}_1, \dots, \mathbf{u}_k\} \subset \mathbb{F}^{\times d}$ costituisce un insieme di vettori linearmente indipendenti $\iff$ per ogni k -upla di scalari $\langle c_1, \dots, c_k \rangle$ che soddisfa l'uguaglianza lineare $\sum_{h=1}^k \mathbf{u}_h c_h = \mathbf{0}$ deve essere $c_1 = c_2 = \dots = c_k = 0$.

Dim.: Se esistesse una sequenza- k $\langle c_1, c_2, \dots, c_k \rangle \neq 0^{\times k}$ per la quale $\sum_{h=1}^d \mathbf{u}_h c_h = \mathbf{0}$, posto che sia $c_j \neq 0$ per un opportuno j tale che $1 \leq j \leq k$, sarebbe possibile scrivere $\mathbf{u}_j = -\frac{1}{c_j} \sum_{h=1, \dots, k \wedge h \neq j} \mathbf{u}_h c_h$,
cioè non si avrebbe l'indipendenza lineare di $\{\mathbf{u}_1, \dots, \mathbf{u}_k\}$ ■

In uno spazio vettoriale un insieme di vettori che consente di esprimere ogni vettore come loro combinazione lineare si dice **sistema completo di vettori**.

In $\mathbb{Q}^{\times 3}$ sono sistemi completi di vettori gli insiemi

$$\{\langle 1, 1, 0 \rangle, \langle 1, 0, 1 \rangle, \langle 0, -2, -2 \rangle\} \text{ e } \{\langle 1, 0, -2 \rangle, \langle -1.5, 0, 3.6 \rangle, \langle 2, 0, -2 \rangle\}.$$

All'opposto si constata che sono sistemi incompleti

$$\{\langle 1, 1, 0 \rangle, \langle -0.5, 2.3, 0 \rangle, \langle 2, -2, 0 \rangle\} \text{ e } \{\langle 1, 0, -2 \rangle, \langle -1.5, 0, 3.6 \rangle, \langle 2, 0, -2 \rangle\}.$$

Si trova essere incompleto anche $\{\langle 1, -2, 3 \rangle, \langle 1, 2, -3 \rangle, \langle 0.5, -3, 4.5 \rangle, \langle 1, 6, -9 \rangle\}$.

B32a.08 Dato che togliendo qualche elemento da un sistema di vettori linearmente indipendenti si ha ancora un sistema di vettori linearmente indipendenti, tra questi insiemi di vettori hanno maggiore interesse i più estesi, in quanto tutti gli altri si possono ottenere da questi con semplici eliminazioni.

Viceversa, dato che aggiungendo qualche nuovo vettore a un sistema di vettori linearmente dipendenti si ha ancora un sistema di vettori linearmente dipendenti, tra questi insiemi di vettori hanno maggiore interesse quelli il più possibile ristretti, in quanto molti altri si possono ottenere da questi con ampliamenti effettuabili in completa libertà.

Un sistema di vettori linearmente dipendenti come strumento per esprimere altri vettori mediante loro combinazioni lineari in genere va considerato uno strumento poco economico, in quanto evidentemente ridondante.

Consideriamo per esempio $\left\{ \mathbf{u}_1, \dots, \mathbf{u}_{k-1}, \sum_{h=1}^{k-1} \mathbf{u}_h d_h \right\}$ ed un vettore combinazione lineare dei vettori nel

precedente insieme $\mathbf{v} = \sum_{i=1}^k \mathbf{u}_i a_i$; questo può esprimersi anche come

$$\sum_{h=1}^{k-1} \mathbf{u}_h a_h + a_k \left(\sum_{h=1}^{k-1} \mathbf{u}_h d_h \right) = \sum_{h=1}^{k-1} \mathbf{u}_h (a_h + d_h),$$

cioè come combinazione lineare di soli $k-1$ vettori.

Un sistema di vettori linearmente dipendenti contiene dunque qualche elemento del quale si può fare a meno nelle manovre che si servono di combinazioni lineari.

Quindi, in linea di massima, i sistemi di vettori linearmente indipendenti vanno considerati strumenti più essenziali dei sistemi di vettori linearmente dipendenti.

B32a.09 Si dice **base per uno spazio** $\mathbb{F}^{\times d}$ ogni sistema completo di vettori linearmente indipendenti. Per le considerazioni precedenti le basi di uno spazio vettoriale sono da considerare strumenti convenienti per fornire combinazioni lineari, in quanto economici e capaci di esprimere tutti i vettori.

Consideriamo i d vettori

$$(1) \quad \mathbf{e}_1 := \langle 1, 0, \dots, 0 \rangle, \mathbf{e}_2 := \langle 0, 1, \dots, 0 \rangle, \dots, \mathbf{e}_d := \langle 0, 0, \dots, 1 \rangle.$$

Essi costituiscono un insieme di vettori linearmente indipendenti, in quanto evidentemente

$$\sum_{h=1}^k c_h \mathbf{e}_h = \langle c_1, c_2, \dots, c_k \rangle = \mathbf{0} \text{ implica } c_1 = c_2 = \dots = c_k = 0.$$

Ogni vettore di $\mathbb{F}^{\times d}$ $\mathbf{v} = \langle v_1, v_2, \dots, v_d \rangle$ si può esprimere come combinazione lineare di questi vettori:

$$(2) \quad \langle v_1, v_2, \dots, v_d \rangle = \sum_{i=1}^d \mathbf{e}_i v_i.$$

Quindi l'insieme $\mathfrak{B}_{can} := \{\mathbf{e}_1, \mathbf{e}_2, \dots, \mathbf{e}_d\}$ costituisce una base di $\mathbb{F}^{\times d}$; tale base è detta **base canonica** di $\mathbb{F}^{\times d}$.

Per semplicità espositiva in genere si considerano **basi ordinate**, cioè sequenze di vettori diversi tali che l'insieme dei suoi d componenti costituisca una base per lo spazio.

Quando si stabilisce di fare riferimento ad una determinata base ordinata ogni suo vettore si può individuare con l'intero che fornisce la sua posizione nella sequenza.

Ad esempio se si conviene di fare riferimento a una base ordinata canonica come $\langle \mathbf{e}_1, \mathbf{e}_2, \dots, \mathbf{e}_d \rangle$, può essere comodo denotare i suoi componenti con i simboli $|1\rangle, |2\rangle, \dots, |d\rangle$.

Ciascuno di questi simboli è detto **ket** e fa parte delle cosiddette **notazioni di Dirac** per gli spazi vettoriali, notazioni che riprenderemo in :d .

B32a.10 I vettori di uno spazio $\mathbb{F}^{\times d}$ sono strumenti di lavoro ampiamente usati e conviene prepararsi a richiamarli servendosi di varie notazioni.

In molte formule sui vettori interviene la funzione delta di Kronecker, funzione che può considerarsi del genere $\mathbb{Z} \times \mathbb{Z} \mapsto \{0, 1\}$ o di un genere $\mathbb{Z} \times \mathbb{Z} \mapsto \{0, 1\}$ per un particolare d e le cui componenti sono definite dalla

$$(1) \quad \delta_{i,j} := \begin{cases} 1 & \text{sse } i = j \\ 0 & \text{sse } i \neq j \end{cases}.$$

Per esempio per ogni $i = 1, 2, \dots, d$ il vettore i -esimo della base canonica è dato dall'espressione

$$(2) \quad \mathbf{e}_i = \langle \delta_{i,1}, \delta_{i,2}, \dots, \delta_{i,d} \rangle.$$

I vettori specifici di uno spazio $\mathbb{F}^{\times d}$ in molte circostanze si trattano vantaggiosamente mediante i cosiddetti vettori colonna, matrici di profilo $d \times 1$. Alcuni esempi di vettori colonna sono

$$\begin{bmatrix} a \\ b \end{bmatrix} \quad \begin{bmatrix} x_1 \\ x_2 \\ x_3 \end{bmatrix} \quad \begin{bmatrix} v_1 \\ v_2 \\ \vdots \\ v_d \end{bmatrix}.$$

Un esempio di combinazione lineare viene presentato con la scrittura

$$\forall v_1, v_2, w_1, w_2, \alpha, \beta \in \mathbb{F} \quad : \quad \alpha \begin{bmatrix} v_1 \\ v_2 \end{bmatrix} + \beta \begin{bmatrix} w_1 \\ w_2 \end{bmatrix} = \begin{bmatrix} \alpha v_1 + \beta w_1 \\ \alpha v_2 + \beta w_2 \end{bmatrix}.$$

Molti insiemi di vettori di uno spazio $\mathbb{F}^{\times d}$ si individuano mediante equazioni e disequazioni nelle quali intervengono simboli variabili che esprimono le singole coordinate; tipicamente si denota con x_i la variabile per la componente i -esima, con $i = 1, 2, \dots, d$.

In alternativa se $d = 2, 3, 4$ si trova comodo usare x invece di x_1 , y invece di x_2 , z invece di x_3 e w invece di x_4 .

B32a.11 Si dimostra in generale, ovvero nell'ambito della trattazione assiomatica della teoria degli insiemi servendosi del cosiddetto **lemma di Zorn** (**wi**), che ogni spazio vettoriale non ridotto a $\{0\}$ possiede una base [*Roman: Advanced Linear Algebra* p. 37].

Qui possiamo dimostrare che in uno spazio vettoriale che possiede un insieme spanning finito S , cioè un insieme finito di vettori che consente di esprimere ogni altro vettore come sua combinazione lineare, può condurre a un sistema linearmente indipendente che ha cardinale non superiore ad $|S|$.

(1) Prop.: Sia V uno spazio vettoriale, $\mathcal{I} = \{\mathbf{v}_1, \mathbf{v}_2, \dots, \mathbf{v}_n\}$ un suo insieme di vettori linearmente indipendenti ed $\mathcal{S} = \{\mathbf{w}_1, \mathbf{w}_2, \dots, \mathbf{w}_m\}$ un sottoinsieme spanning (finito) di V . Allora $n \leq m$.

Dim.: Facciamo riferimento a coppie di sequenze di vettori che modifichiamo attraverso successivi stadi che contrassegniamo con $[0]$, $[1]$, ... ed $[n]$.

$$\mathbf{w}_1 \ \mathbf{w}_2 \ \dots \ \mathbf{w}_m \quad \text{---} \quad \mathbf{v}_1 \ \mathbf{v}_2 \ \dots \ \mathbf{v}_n$$

$\mathbf{v}_1$ si può esprimere come combinazione lineare dei $\mathbf{w}_j$ con almeno un coefficiente diverso da zero; da questa combinazione si ricava un'espressione di uno dei $\mathbf{w}_j$ come combinazione dei rimanenti vettori di $\mathcal{S}$ e di $\mathbf{v}_1$.

Si giunge quindi alla coppia di sequenze dello stadio $[1]$

$$\mathbf{v}_1 \ \mathbf{w}_{2,1} \ \dots \ \mathbf{w}_{m,1} \quad \text{---} \quad \mathbf{v}_2 \ \dots \ \mathbf{v}_n \quad [1]$$

La prima sequenza riguarda ancora un insieme spanning nel quale il primo vettore dell'insieme $\mathcal{I}$ ha rimpiazzato uno dei vettori di $\mathcal{S}$; la seconda è stata semplicemente decurtata del primo vettore.

Inoltre si è convenuto di usare le notazioni $\mathbf{w}_{j,1}$ per esprimere la prima sequenza decurtata del vettore $\mathbf{w}_j$ rimpiazzato.

Si esegue lo stadio $[2]$ considerando che $\mathbf{v}_2$ si può esprimere come combinazione lineare dei vettori della prima sequenza nella quale almeno uno dei coefficienti dei $\mathbf{w}_{j,1}$ è diverso da 0. Questo rende possibile sostituire uno di questi vettori di $\mathcal{S}$ con $\mathbf{v}_2$ ottenendo la coppia di sequenze

$$\mathbf{v}_1 \ \mathbf{v}_2 \ \mathbf{w}_{3,2} \ \dots \ \mathbf{w}_{m,2} \quad \text{---} \quad \mathbf{v}_3 \ \dots \ \mathbf{v}_n \quad [2]$$

Ancora la prima sequenza costituisce un insieme spanning e risulta possibile esprimere $\mathbf{v}_3$ come combinazione lineare dei primi vettori nella quale almeno uno dei coefficienti dei $\mathbf{w}_{j,2}$ è diverso da 0 (non si può esprimere un vettore $\mathbf{v}_i$ come combinazione lineare dei soli rimanenti vettori di $\mathcal{I}$).

Il procedimento quindi può proseguire fino allo stadio $[n]$ nel quale sono stati eliminati tutti i vettori di $\mathcal{I}$ caratterizzabili con la scrittura:

$$\mathbf{v}_1 \ \dots \ \mathbf{v}_n \ \mathbf{w}_{m-n,n} \ \dots \ \mathbf{w}_{m,n} \quad \text{---} \quad [n]$$

Dunque i vettori dell'insieme spanning $\mathcal{S}$ devono essere in numero non inferiore ad n ■

B32a.12 (1) Coroll.: Ogni spazio vettoriale V che possiede un insieme spanning finito ha tutte le basi con lo stesso cardinale di tale insieme.

Dim.: Siano $\mathfrak{B}_1$ e $\mathfrak{B}_2$ due basi di V . Il procedimento visto in a10 si può effettuare con $\mathcal{B}_1$ nel ruolo di insieme spanning $\mathcal{S}$ e con $\mathcal{B}_2$ nel ruolo di insieme di vettori linearmente indipendenti fino a concludere

che $|\mathcal{B}_1| \geq |\mathcal{B}_2|$. Ma esso si può effettuare anche con $\mathcal{B}_2$ nel ruolo di insieme spanning $\mathcal{S}$ e con $\mathcal{B}_1$ nel ruolo di insieme di vettori linearmente indipendenti fino a stabilire che $|\mathcal{B}_1| \leq |\mathcal{B}_2|$ e conseguentemente concludere che $|\mathcal{B}_1| = |\mathcal{B}_2|$ ■

Ogni spazio vettoriale V che possiede un insieme spanning finito è caratterizzato dall'intero positivo che fornisce il cardinale di tutte le sue basi. Questo si chiama **dimensione dello spazio** V e si denota con $\dim(V)$.

Se $\dim(V) =: n$, si dice che V ha d dimensioni, ovvero che è uno **spazio d -dimensionale**.

Lo spazio $\mathbb{F}^{\times d}$ delle sequenze di lunghezza d ha d dimensioni, in quanto la sua base canonica è costituita dai d vettori $\mathbf{e}_1, \mathbf{e}_2, \dots, \mathbf{e}_d$ e quindi tutte le sue basi hanno cardinale d .

Prescindendo dal particolare valore d , gli spazi vettoriali dotati di un insieme spanning finito si dicono **spazi finitodimensionali** o **spazi di dimensioni finite**.

B32a.13 Consideriamo uno spazio vettoriale d -dimensionale e una sua base ordinata $\langle \mathbf{u}_1, \mathbf{u}_2, \dots, \mathbf{u}_d \rangle$; ciascuno dei suoi vettori si può esprimere in un unico modo come combinazione lineare dei vettori della base.

Quindi ogni spazio finitodimensionale su un campo $\mathbb{F}$ si può trattare come uno spazio vettoriale di sequenze di elementi del campo aventi una unuca determinata lunghezza.

Si usa anche dire che lo spazio $\mathbb{F}^{\times d}$ costituisce un “modello canonico” di tutti gli spazi d -dimensionali sul campo $\mathbb{F}$.

B32 b. sottospazi

B32b.01 Si dice **sottospazio** di $\mathbb{F}^{\times d}$ ogni suo sottoinsieme S tale che ogni combinazione lineare di suoi vettori appartenga allo stesso S . Questa caratteristica di un sottospazio si esprime equivalentemente dicendo che questo insieme di vettori è **insieme di vettori chiuso per combinazione lineare**.

La proprietà di chiusura per combinazione lineare di un sottospazio S significa che effettuando una qualsiasi combinazione lineare di un qualsiasi insieme di vettori di S si può ottenere solo un altro elemento di S .

Due casi molto particolari di sottospazi di $\mathbb{F}^{\times d}$ sono il **sottospazio nullo** $\{0^{\times d}\}$; lo stesso $\mathbb{F}^{\times d}$ viene detto **sottospazio improprio** di se stesso.

Ogni sottospazio di $\mathbb{F}^{\times d}$ sottoinsieme proprio dell'intero spazio si dice **sottospazio proprio**.

Sottospazi semplici da descrivere sono i **sottospazi monodimensionali**, cioè gli insiemi definibili come $\mathbf{v}\mathbb{F} := \{c \in \mathbb{F} : \mathbf{v}c\}$ servendosi di un qualsiasi vettore nonnullo $\mathbf{v}$. Le proprietà dei sottospazi monodimensionali sono in stretta corrispondenza con le proprietà del campo $\mathbb{F}$, fatto che si esprime dicendo che sono **strutture isomorfe** con $\mathbb{F}$.

Da notare che, presi due vettori $\mathbf{v}$ e $\mathbf{w}$ accade che $\mathbf{v}\mathbb{F} = \mathbf{w}\mathbb{F}$ sse i due vettori sono collineari.

B32b.02 Per $d \geq 2$ è utile considerare i sottospazi particolari costituiti da d -uple le quali hanno alcune componenti uguali a 0 e le componenti rimanenti arbitrariamente variabili in $\mathbb{F}$.

Per questi sottospazi usiamo il termine **sottospazi con coordinate azzerate** e la corrispondente abbreviazione **SSCZ**.

Osserviamo esplicitamente che sono sscz di $\mathbb{F}^{\times d}$ questo stesso sottospazio improprio e il suo sottospazio nullo $\{0\}$; gli altri sscz, prevedibilmente, sono chiamati sscz propri.

In $\mathbb{F} \times \mathbb{F}$ sono sscz propri l'asse delle ascisse e l'asse delle ordinate.

In $\mathbb{F}^{\times 3}$ sono sottospazi sscz propri:

$\text{Sscz}_{011} := \{y, z \in \mathbb{Q} : \langle 0, y, z \rangle\}$, insieme individuato dall'equazione $x = 0$, chiamato piano $x = 0$ e denotato con Oyz ;

$\text{Sscz}_{101} := \{x, z \in \mathbb{Q} : \langle x, 0, z \rangle\}$, insieme individuato dall'equazione $y = 0$, chiamato piano $y = 0$ e denotato con Oxz ;

$\text{Sscz}_{110} := \{x, y \in \mathbb{Q} : \langle x, y, 0 \rangle\}$, insieme individuato dall'equazione $z = 0$, chiamato piano $z = 0$ e denotato con Oxy ;

$\text{Sscz}_{100} := \{x \in \mathbb{Q} : \langle x, 0, 0 \rangle\}$, insieme individuato dal sistema di equazioni $y = z = 0$, chiamato **asse della prima coordinata** (x) e denotato con Ox ; $\text{Sscz}_{010} := \{y \in \mathbb{Q} : \langle 0, y, 0 \rangle\}$, insieme individuato dal sistema di equazioni $x = z = 0$, chiamato **asse della seconda coordinata** (y) e denotato con Oy ;

$\text{Sscz}_{001} := \{z \in \mathbb{Q} : \langle 0, 0, z \rangle\}$, insieme individuato dal sistema di equazioni $x = y = 0$, chiamato **asse della terza coordinata** (z) e denotato con Oz .

Le precedenti notazioni si possono generalizzare per gli spazi di dimensioni superiori a 3 con notazioni che presentano d -uple di cifre binarie della forma $\text{Sscz}_{b_1 b_2 \dots b_d}$ con $0 < \sum_{i=1}^d b_i < d$; una tale scrittura individua il sottospazio dei vettori la cui i -esima coordinata vale 0 sse il corrispondente $b_i = 0$ e può assumere qualsiasi valore di $\mathbb{F}$ sse $b_i = 1$.

Per esempio in $\mathbb{F}^{\times 4}$ potrebbero interessare

$$\text{Sscz}_{1101} := \{x, y, w \in \mathbb{F} : \langle x, y, 0, w \rangle\} \text{ e } \text{Sscz}_{1011} := \{x, z, w \in \mathbb{F} : \langle x, 0, z, w \rangle\}.$$

Evidentemente Sscz_{0000} è il sottospazio nullo tetradimensionale $\{\langle 0, 0, 0, 0 \rangle\}$, mentre $\text{Sscz}_{11111} = \mathbb{F}^{\times 5}$.

B32b.03 Un sscz che riguarda l'annullamento di una sola coordinata, ad esempio la i -esima, si può individuare con la semplice equazione $[x_i = 0]$. Un sscz che corrisponde all'annullamento di più coordinate si individua con il sistema delle equazioni di annullamento di tali coordinate: per esempio Sscz_{1001} si può identificare con la scrittura $[x_1 = 0 \wedge x_4 = 0]$ o con l'equivalente sistema di equazioni di annullamento

$$\begin{cases} y = 0 \\ z = 0 \end{cases}.$$

Se si intersecano due sottospazi sscz si ottiene un altro sottospazio dello stesso tipo: per esempio

$$\text{Sscz}_{1101} \cap \text{Sscz}_{1011} = \{x, w \in \mathbb{F} : \langle x, 0, 0, w \rangle\} = \text{Sscz}_{1001}.$$

In generale l'intersezione di due sscz si esprime con il sistema di equazioni di annullamento di coordinate costituito dall'unione delle equazioni di annullamento che individuano i due sottospazi.

Per esempio per l'intersezione dei due sscz

$$\text{Sscz}_{10010} : \begin{cases} x_2 = 0 \\ x_3 = 0 \\ x_5 = 0 \end{cases} \quad \text{e} \quad \text{Sscz}_{01010} : \begin{cases} x_1 = 0 \\ x_3 = 0 \\ x_5 = 0 \end{cases}$$

abbiamo

$$\text{Sscz}_{00010} : \begin{cases} x_1 = 0 \\ x_2 = 0 \\ x_3 = 0 \\ x_5 = 0 \end{cases}$$

Più compattamente

$$[x_2 = 0 \wedge x_3 = 0 \wedge x_5 = 0] \cap [x_1 = 0 \wedge x_3 = 0 \wedge x_5 = 0] = [x_1 = 0 \wedge x_2 = 0 \wedge x_3 = 0 \wedge x_5 = 0].$$

B32b.04 In generale anche l'intersezione di due sottospazi qualsiasi S_1 e S_2 di uno spazio $\mathbb{F}^{\times d}$ è un sottospazio. Infatti scelti come si vuole l'intero positivo k , i vettori $\mathbf{v}_1, \dots, \mathbf{v}_k$ in $S_1 \cap S_2$, e altrettanti scalari $c_1, \dots, c_k$, la combinazione lineare $\sum_{h=1}^k \mathbf{v}_h c_h$ deve appartenere ad entrambi i sottospazi e quindi anche alla loro intersezione.

Se si confrontano due sottospazi S_1 ed S_2 di $\mathbb{F}^{\times d}$ propri e nonnulli possono verificarsi le seguenti situazioni:

- (1) $S_1 = S_2$;
- (2) $S_1 \subset S_2$, relazione equivalente alla $S_1 \cap S_2 = S_1$, a sua volta equivalente alla $S_1 \cup S_2 = S_2$;
- (3) $S_2 \subset S_1$, relazione equivalente alla $S_1 \cap S_2 = S_2$ e alla $S_1 \cup S_2 = S_1$;
- (4) $S_1 \cap S_2$ sottospazio nonnullo;
- (4) $S_1 \cap S_2 = \{\mathbf{0}\}$.

Nei primi tre casi S_1 e S_2 si dicono **sottospazi comparabili**; nei due casi restanti **sottospazi noncomparabili** e in tal caso esistono vettori nonnulli di S_1 che non appartengono ad S_2 e vettori nonnulli di S_2 che non appartengono ad S_1 , ovvero si può affermare $S_1 \setminus S_2 \neq \{\mathbf{0}\}$ e $S_2 \setminus S_1 \neq \{\mathbf{0}\}$.

Semplici esempi delle precedenti situazioni si ottengono con sottospazi sscz:

$$\begin{aligned} \text{Sscz}_{1010} \subset \text{Sscz}_{1011}, \quad \text{Sscz}_{10101} \supset \text{Sscz}_{00100} \\ \{\mathbf{0}^{\times 6}\} \subset \text{Sscz}_{011010} = \text{Sscz}_{011110} \cap \text{Sscz}_{111011} \subset \text{Sscz}_{011110}, \subset \text{Sscz}_{111011}. \end{aligned}$$

B32b.05 Dato un $E \subset \mathbb{F}^{\times d}$ si dice **chiusura lineare** di E l'insieme di tutte le combinazioni lineari di elementi di E ; denotiamo tale insieme con $\text{span}(E)$.

Evidentemente la chiusura lineare di un qualsiasi insieme di vettori di $\mathbb{F}^{\times d}$ è un sottospazio di tale spazio ambiente: infatti la combinazione lineare di due vettori che sono combinazioni lineari di vettori di un dato insieme E si può riscrivere come combinazione lineare di vettori di E .

Questo fatto può essere chiarito dalla seguente formula, nella quale si assume $E = \{\mathbf{v}_1, \mathbf{v}_2, \dots, \mathbf{v}_k\}$

$$(1) \quad \alpha \left(\sum_{h=1}^k c_h \mathbf{v}_h \right) + \beta \left(\sum_{h=1}^k d_h \mathbf{v}_h \right) = \sum_{h=1}^k (\alpha_h c_h + \beta d_h) \mathbf{v}_h .$$

Un insieme di vettori chiusura lineare della forma $\mathbf{span}(E)$ viene chiamato anche **sottospazio sotteso dai vettori** costituenti l'insieme di vettori E .

Il simbolo $\mathbf{span}$ individua una funzione che ha come dominio $\mathfrak{P}(\mathbb{F}^{\times d})$, la collezione di tutti i sottoinsiemi di $\mathbb{F}^{\times d}$, e come codominio l'insieme di tutti i sottospazi di $\mathbb{F}^{\times d}$, insieme che denotiamo con l'espressione $\mathbf{Sbvsp}(\mathbb{F}^{\times d})$.

Questa funzione di insiemi è evidentemente un ampliamento, cioè per ogni $E \subseteq \mathbb{F}^{\times d}$ si ha $\mathbf{span}(E) \supseteq E$. Inoltre essa è una funzione idempotente, ovvero $\mathbf{span}(\mathbf{span}(E)) = \mathbf{span}(E)$. Come vedremo queste due proprietà caratterizzano l'importante insieme delle funzioni chiamate chiusure [B54e].

La chiusura lineare dell'insieme E si può anche individuare come il più ridotto dei sottospazi che contiene tale E ed equivalentemente come intersezione di tutti i sottospazi che contengono E :

$$(2) \quad \mathbf{span}(E) = \bigcap \{S \in \mathbf{Sbvsp}(\mathbb{F}^{\times d}) \mid S \supseteq E\} .$$

B32b.06 Consideriamo due sottospazi S_1 ed S_2 e la loro unione $U := S_1 \cup S_2$. Ovviamente se essi sono confrontabili la loro unione coincide con il più esteso dei sottospazi operandi dell'unione. Se invece essi non sono confrontabili U non è un sottospazio.

Un controesempio evidente si osserva già in due dimensioni: due sottospazi noncomparabili sono due rette diverse passanti per l'origine; queste rette hanno in comune solo l'origine, mentre la somma di due vettori nonnulli, il primo appartenente (solo) alla prima retta, il secondo (solo) alla seconda, non può appartenere a nessuna delle due rette.

Altri semplici controesempi sono forniti da opportune coppie di sscz; consideriamo $\mathbf{Sscz}_{1100}$ e $\mathbf{Sscz}_{0110}$; $\mathbf{e}_1 \in \mathbf{Sscz}_{1100}$, $\mathbf{e}_3 \in \mathbf{Sscz}_{0110}$, mentre $\mathbf{e}_1 + \mathbf{e}_3 \notin \mathbf{Sscz}_{1100} \cup \mathbf{Sscz}_{0110}$.

In generale si possono considerare i due vettori $\mathbf{v}_1 \in S_1 \setminus S_2$ e $\mathbf{v}_2 \in S_2 \setminus S_1$ e si trova che $\mathbf{v}_1 + \mathbf{v}_2$ non può appartenere ad S_1 , perché in tal caso si avrebbe $\mathbf{v}_2 = (\mathbf{v}_1 + \mathbf{v}_2) - \mathbf{v}_1 \in S_1$, contro l'ipotesi; per simmetria tale vettore non può appartenere ad S_2 e dunque non a $S_1 \cup S_2$.

Si deve invece concludere che $\mathbf{v}_1 + \mathbf{v}_2 \notin S_1 \cup S_2$.

Si trova invece che è un sottospazio la chiusura lineare dell'unione di due sottospazi, cioè $\mathbf{span}(S_1 \cup S_2)$. Per esempio accade che

$$\mathbf{e}_1 + \mathbf{e}_3 \in \mathbf{span}(\mathbf{Sscz}_{1100} \cup \mathbf{Sscz}_{0110}) = \mathbf{Sscz}_{1110} .$$

Il sottospazio $\mathbf{span}(S_1 \cup S_2)$ si dice anche **somma dei sottospazi** S_1 ed S_2 in quanto esso si può esprimere come $\{\mathbf{v}_1 \in S_1, \mathbf{v}_2 \in S_2 : \mathbf{v}_1 + \mathbf{v}_2\}$. Esso si denota anche con $S_1 \boxplus S_2$.

Talora è utile considerare la somma di tre e più spazi.

Gli sscz forniscono anche semplici esempi di somme di sottospazi:

$$\mathbf{Sscz}_{0101} \boxplus \mathbf{Sscz}_{0110} = \mathbf{Sscz}_{0111} , \quad \mathbf{Sscz}_{010100} \boxplus \mathbf{Sscz}_{010011} \boxplus \mathbf{Sscz}_{101011} = \mathbb{R}^{\times 6} .$$

Si dimostra facilmente che la somma di sottospazi è un'operazione commutativa, associativa e idempotente.

Inoltre possono servire composizioni di sequenze di sottospazi e per queste proponiamo di servirsi della sommatoria di sottospazi, costruzione espressa da scritte come

$$\boxplus_{i=1}^k V_i .$$

B32 c. trasformazioni lineari e loro matrici

B32c.01 Consideriamo due spazi vettoriali sul campo $\mathbb{F}$ $\mathbf{V}$ ed $\mathbf{X}$; si dice **trasformazione lineare** di $\mathbf{V}$ in $\mathbf{X}$ ogni funzione del genere $T \in [\mathbf{V} \longrightarrow \mathbf{X}]$ tale che,

$$\forall \mathbf{v}, \mathbf{w} \in \mathbf{V}, \alpha, \beta \in \mathbb{F} : T(\mathbf{v}\alpha + \mathbf{w}\beta) = T(\mathbf{v})\alpha + T(\mathbf{w})\beta.$$

La proprietà espressa dalla precedente uguaglianza viene chiamata **rispetto della combinazione lineare**; si dice anche che una trasformazione lineare è una funzione tra due spazi vettoriali che mantiene la combinazione lineare.

Una tale T si chiama anche **omomorfismo di spazi vettoriali** o **omomorfismo lineare** di $\mathbf{V}$ in $\mathbf{X}$.

L'insieme di queste funzioni si denota con $[\mathbf{V} \longrightarrow_{Lin} \mathbf{X}]$.

In talune circostanze può essere opportuno fare riferimento specificamente al campo, parlare di omomorfismo lineare- $\mathbb{F}$ e usare notazioni come $[\mathbf{V} \longrightarrow_{Lin-\mathbb{F}} \mathbf{X}]$.

B32c.02 Le trasformazioni lineari, dette anche **operatori lineari**, si incontrano in moltissimi sviluppi della matematica e delle sue applicazioni e vengono intensamente studiate e utilizzate.

Qui approfondiamo solo lo studio delle trasformazioni lineari tra due spazi di sequenze $\mathbf{V} = \mathbb{F}^{\times d}$ e $\mathbf{X} = \mathbb{F}^{\times e}$, ma è opportuno segnalare che si studiano spazi i cui vettori sono funzioni di variabili reali o complesse con particolari proprietà che risultano particolarmente utili in settori applicativi come meccanica, fisica atomica, chimica strutturistica, biologia molecolare, statistica, economia e ricerca operativa.

Vedremo che negli spazi di funzioni reali derivabili la derivazione è una trasformazione lineare, grazie alla proprietà

$$\frac{d}{dx} [\alpha f(x) + \beta g(x)] = \alpha \frac{d}{dx} f(x) + \beta \frac{d}{dx} g(x).$$

Similmente negli spazi di funzioni reali integrabili in un intervallo $[a, b]$ l'integrazione definita in tale intervallo è una trasformazione lineare che porta allo spazio (monodimensionale) dei numeri reali, grazie alla

$$\int_a^b dx [\alpha f(x) + \beta g(x)] = \alpha \int_a^b dx f(x) + \beta \int_a^b dx g(x).$$

Vi sono inoltre spazi di funzioni per le quali si possono definire delle cosiddette trasformate integrali della forma $\int_a^b dx T(y, x) f(x)$; queste trasformate sono trasformazioni lineari grazie a proprietà come

$$\int_a^b dx T(y, x) [\alpha f(x) + \beta g(x)] = \alpha \int_a^b dx T(y, x) f(x) + \beta \int_a^b dx T(y, x) g(x).$$

B32c.03 Consideriamo ora trasformazioni lineari entro lo spazio $\mathbb{F}^{\times d}$.

Per il loro insieme introduciamo la notazione $\mathbf{Lintr}(\mathbb{F}^{\times d})$. Conveniamo inoltre di scrivere $\mathbb{F}_{nz}$ per l'insieme degli elementi diversi dallo 0 del campo $\mathbb{F}$.

Per ogni $\lambda \in \mathbb{F}_{nz}$ la funzione $[\mathbf{v} \in \mathbf{V} \mapsto \mathbf{v}\lambda]$, cioè la moltiplicazione dei vettori di $\mathbf{V}$ per lo scalare λ , si dice anche **omotetia centrale** di fattore λ .

Chiaramente una tale trasformazione lineare è invertibile e la trasformazione inversa della omotetia centrale di fattore λ è l'omotetia centrale di fattore λ^{-1} .

La moltiplicazione per 1 è l'identità o trasformazione identica di $\mathbf{V}$, cioè $\text{Id}_{\mathbf{V}}$; se si può sottintendere $\mathbf{V}$ l'identità si può denotare con $\hat{1}$ e l'omotetia di fattore λ si denota con $\lambda \cdot \hat{1}$ o con $\hat{\lambda}$.

Se $\lambda > 1$ si parla di **dilatazione** e se $0 < \lambda < 1$ di **contrazione**; per $\lambda = -1$ l'omotetia è la $\lceil \mathbf{v} \mapsto -\mathbf{v} \rceil$, cioè la trasformazione chiamata **simmetria centrale** di centro $\mathbf{0}$ e, con un termine molto usato in fisica, **inversione spaziale**.

Anche la moltiplicazione per il numero 0 può considerarsi una trasformazione lineare, in quanto il suo codominio è costituito dal sottospazio costituito dal solo vettore nullo $0^{\times d}$, entità che si può considerare uno spazio vettoriale a 0 dimensioni.

Essa viene detta **trasformazione nulla** o trasformazione annichilatrice di $\mathbb{F}^{\times d}$; quando si può evitare di segnalare $\mathbf{V}$ essa si denota con $\hat{0}$.

B32c.04 Altre trasformazioni lineari molto semplici si individuano a partire dai vettori della base canonica $\mathbf{e}_1, \mathbf{e}_2, \dots, \mathbf{e}_d$.

Per ogni $i = 1, 2, \dots, d$ si definisce **proiezione sul sottospazio** o **proiettore sul sottospazio**, $\mathbf{e}_i \mathbb{F}$ la trasformazione

$$\mathbf{Prj}_{\mathbf{e}_i} := \left\lceil \sum_{j=1}^d \mathbf{e}_j c_j \mapsto \mathbf{e}_i c_i \right\rceil .$$

Se si presuppone di fare riferimento alla base canonica ordinata, il precedente proiettore si può denotare semplicemente con $\mathbf{Prj}_i$.

Si osserva che $\mathbf{Prj}_{\mathbf{e}_i}(\mathbf{Prj}_{\mathbf{e}_i}(\mathbf{v})) = \mathbf{Prj}_{\mathbf{e}_i}(\mathbf{v})$; questo si esprime dicendo che i proiettori sono operatori lineari idempotenti, cioè operatori che applicati due volte danno lo stesso risultato che si ottiene applicandoli una sola volta.

Inoltre si constata che la composizione di due proiettori relativi a due diversi vettori della base canonica è la trasformazione nulla.

In generale chiamiamo **proiettore su un sottospazio** $\mathbb{F}^{\times d}$ ogni operatore lineare idempotente.

B32c.05 Si osserva che le trasformazioni lineari- $\mathbb{F}$ lineari di un qualsiasi genere $\lceil \mathbf{V} \mapsto \mathbf{X} \rceil$, cioè aventi come dominio e come codominio due spazi vettoriali sullo stesso campo $\mathbb{F}$ possono essere moltiplicate per uno scalare del campo $\mathbb{F}$ e sommate, ovvero possono essere combinate linearmente.

Infatti si può dare la definizione

$$\forall \alpha, \beta \in \mathbb{F}, \forall T, U \in \lceil \mathbf{V} \longrightarrow_{Lin-\mathbb{F}} \mathbf{X} \rceil : (\alpha T + \beta U)(\mathbf{v}) := \alpha \cdot T(\mathbf{v}) + \beta \cdot U(\mathbf{v}) .$$

Si osserva che questa definizione si può applicare a tutte le funzioni T ed U aventi come codominio uno spazio vettoriale lineare- $\mathbb{F}$, quale che sia il loro dominio.

Si osserva inoltre che le trasformazioni lineari, potendo essere combinate linearmente, costituiscono a loro volta uno spazio vettoriale.

Il vettore nullo di questo spazio è dato dalla trasformazione nulla e il vettore opposto della trasformazione $\lceil \mathbf{v} \mapsto T(\mathbf{v}) \rceil$ è la $\lceil \mathbf{v} \mapsto T(-\mathbf{v}) \rceil = \lceil \mathbf{v} \mapsto -T(\mathbf{v}) \rceil$.

Esempi semplici e significativi di somme di trasformazioni lineari di $\mathbf{Lintr}(\mathbb{F}^{\times d})$ sono le somme di proiettori.

Vediamo come si può associare un proiettore ad ogni sscz. Consideriamo la sequenza di d bits $\mathbf{b} = b_1 b_2 \dots b_d$ e il corrispondente sottospazio $\mathbf{Sscz}_{\mathbf{b}}$.

Si dice proiettore corrispondente a tale sscz la trasformazione lineare

$$\left\lceil \sum_{j=1}^d \mathbf{e}_j c_j \in \mathbb{F}^{\times d} \mapsto \sum_{j=1}^d \mathbf{e}_j b_j c_j \right\rceil .$$

Per esempio per $d = 3$ il proiettore relativo alla sscz caratterizzata dall'equazione $\lceil x_2 = 0 \rceil$, cioè il proiettore sul piano $\mathbf{O}x_1x_3$, trasforma il generico vettore $\mathbf{v} = \langle v_1, v_2, v_3 \rangle$ nel vettore $\langle v_1, 0, v_3 \rangle$.

Forse conviene osservare esplicitamente che questo vettore è la somma dei vettori $\mathbf{e}_1 v_1$ ed $\mathbf{e}_3 v_3$ ottenibili applicando a $\mathbf{v}$, risp., il proiettore $\mathbf{Prj}_1$ e il proiettore $\mathbf{Prj}_3$.

Questo induce a considerare il precedente proiettore sul piano $\mathbf{O}x_1x_3$ come la somma dei precedenti proiettori su sottospazi monodimensionali: $\mathbf{Prj}_{101} = \mathbf{Prj}_1 + \mathbf{Prj}_3$.

Evidentemente quindi in generale per il proiettore sopra $\mathbf{Sscz}_{\mathbf{b}}$ si ha

$$\mathbf{Prj}_{b_1b_2\dots b_d} = \sum_{j=1}^d \mathbf{Prj}_j b_j.$$

In effetti vale il seguente enunciato.

(1) Prop.: Se $\mathbf{P}_\mu$ e $\mathbf{P}_\nu$ sono due proiettori e $\mathbf{P}_\mu \circ \mathbf{P}_\nu = \hat{\mathbf{0}}$, allora $\mathbf{P}_\mu + \mathbf{P}_\nu$ è un proiettore.

Due proiettori si dicono **proiettori ortogonali** sse la loro composizione è la trasformazione nulla.

B32c.06 La definizione di trasformazione lineare richiede anche che ogni combinazione lineare di vettori appartenenti a $\text{dom}(T)$ appartiene allo stesso $\text{dom}(T)$, cioè che $\text{span}(\text{dom}(T)) = \text{dom}(T)$ e $\text{dom}(T) \in \mathbf{Sbvsp}(V)$.

Dunque ogni trasformazione lineare ha come dominio un sottospazio.

Inoltre si trova che anche $\text{cod}(T)$ è chiuso rispetto alla combinazione lineare, cioè accade che $\text{cod}(T) \in \mathbf{Sbvsp}(V)$.

Infatti, quali che siano $\mathbf{x}, \mathbf{y} \in \text{cod}(T)$ e $\alpha, \beta \in \mathbb{F}$, se $\mathbf{v}$ e $\mathbf{w}$ sono due vettori di $\text{dom}(T)$ tali che, se $T(\mathbf{v}) = \mathbf{x}$ e $T(\mathbf{w}) = \mathbf{y}$, allora deve essere

$$\mathbf{v}\alpha + \mathbf{w}\beta = T(\mathbf{v}\alpha + \mathbf{w}\beta) \in \text{cod}(T).$$

Dunque le trasformazioni lineari sono funzioni che hanno come dominio e come codominio dei sottospazi.

B32c.07 Introduciamo un sistema di termini che denotano collezioni di trasformazioni lineari tra due spazi V e X , eventualmente coincidenti, definite solo caratterizzando i rispettivi domini e codomini.

Denotiamo con $\lceil V \mapsto_{Lin} X \rceil$ l'insieme degli omomorfismi lineari L dall'intero V in X , cioè degli omomorfismi tali che $\text{dom}(L) = V$; questo insieme di trasformazioni si denota anche con $\mathbf{Lintr}(V, X)$.

Denotiamo con $\lceil V \twoheadrightarrow_{Lin} X \rceil$ l'insieme degli omomorfismi L da V sull'intero X , cioè tali che $\text{cod}(L) = X$; essi sono detti anche **epimorfismi lineari**.

Denotiamo con $\lceil V \twoheadrightarrow_{Lin} X \rceil$ l'insieme degli omomorfismi L dall'intero V sull'intero X , cioè tali che $\text{dom}(L) = V$ e $\text{cod}(L) = X$.

Denotiamo con $\lceil V \hookrightarrow_{Lin} X \rceil$ l'insieme degli omomorfismi L da V in X invertibili, cioè tali che esista $L^{-1} \in \lceil X \hookrightarrow_{Lin} V \rceil$; essi sono detti anche **monomorfismi lineari**.

Denotiamo con $\lceil V \hookrightarrow_{Lin} X \rceil$ l'insieme dei monomorfismi L dall'intero V nell'intero X , cioè tali che $\text{dom}(L) = V$, $\text{cod}(L) = X$ ed esista L^{-1} ; essi sono epimorfismi e sono detti anche **isomorfismi lineari**.

Quando $V=X$ si parla di **endomorfismi lineari** entro V e anche di **operatori lineari** su V .

La loro collezione si denota anche con $\mathbf{Lintr}(V)$.

Gli endomorfismi che sono isomorfismi, e quindi permutazioni del loro dominio, si dicono **automorfismi lineari** di V e anche **operatori lineari invertibili** su V .

B32c.08 Ogni trasformazione lineare T del genere $\lceil \mathbb{F}^{\times d} \mapsto \mathbb{F}^{\times e} \rceil$ risulta determinata dai trasformati dei vettori di una base dello spazio dominio $\mathbb{F}^{\times d}$; in particolare è definita dalle sue azioni sui vettori $\mathbf{e}_i, \mathbf{e}_i, \dots, \mathbf{e}_i$ della base canonica di $\mathbb{F}^{\times d}$.

Infatti se conosciamo i trasformati $T(\mathbf{e}_i)$ per $i = 1, 2, \dots, d$, siamo in grado di determinare il trasformato da T di ogni vettore $\mathbf{v} = \sum_{i=1}^d \mathbf{e}_i v_i$, grazie alla uguaglianza:

$$(1) \quad T(\mathbf{v}) = \left(\sum_{i=1}^d \mathbf{e}_i v_i \right) = \sum_{i=1}^d T(\mathbf{e}_i) v_i .$$

I vettori della base canonica di $\mathbb{F}^{\times e}$ li possiamo denotare con $\mathbf{u}_h$ per $h = 1, 2, \dots, e$ e per ciascuno di essi possiamo scrivere $\mathbf{u}_h := \langle k \in (e) : \delta_{h,k} \rangle$.

Supponiamo di conoscere le espressioni dei vettori $T(\mathbf{e}_i)$ come combinazioni lineari degli $\mathbf{u}_h$ che scriviamo

$$(2) \quad T(\mathbf{e}_i) =: \sum_{h=1}^e \mathbf{u}_h t_{h,i} \quad \text{per } i = 1, 2, \dots, d .$$

Per il trasformato del generico vettore $\mathbf{v}$ abbiamo quindi

$$(3) \quad T(\mathbf{v}) = \sum_{i=1}^d T(\mathbf{e}_i) v_i = \sum_{i=1}^d \sum_{h=1}^e \mathbf{u}_h t_{h,i} v_i = \sum_{h=1}^e \mathbf{u}_h \left(\sum_{i=1}^d t_{h,i} v_i \right) .$$

L'azione della trasformazione T è quindi determinata dalla conoscenza dei $d \cdot e$ scalari $t_{h,i}$ che conviene disporre nelle caselle di una matrice di profilo $e \times d$; in essa l'entrata relativa alla riga h e alla colonna i riguarda il coefficiente del vettore $\mathbf{u}_h$ dello sviluppo nella base canonica dello spazio codominio $\mathbb{F}^{\times e}$ del trasformato del vettore $\mathbf{e}_i$ della base canonica dello spazio dominio $\mathbb{F}^{\times d}$.

La matrice così individuata si dice **matrice rappresentativa della trasformazione nelle basi canoniche** degli spazi codominio e dominio.

Le matrici rappresentative consentono di trattare le trasformazioni lineari con i metodi numerici sviluppati per i calcoli sulle matrici e anche con le procedure che li implementano e con i dispositivi hardware che sono stati sviluppati per poter effettuare con altissima efficienza i calcoli matriciali più impegnativi.

B32c.09 La formula c08(3) mostra che ad ogni trasformazione lineare del genere $\left[\mathbb{F}^{\times d} \mapsto \mathbb{F}^{\times e} \right]$, fissate una base di $\mathbb{F}^{\times d}$ e una base di $\mathbb{F}^{\times e}$, risulta associata una matrice $d \times e$.

Da essa si ricava anche che ogni matrice di $\mathbf{Mat}_{d,e}(\mathbb{F})$, facendo riferimento a una base di $\mathbb{F}^{\times d}$ e una base di $\mathbb{F}^{\times e}$, individua una trasformazione lineare di $\left[\mathbb{F}^{\times d} \mapsto \mathbb{F}^{\times e} \right]$.

Quindi matrici su un campo $\mathbb{F}$ e trasformazioni lineari- $\mathbb{F}$ sono due nozioni equivalenti-eqp.

In termini operativi la stessa c08(3) dice come precisare l'azione di ogni trasformazione lineare mediante calcoli matriciali.

Il vettore da trasformare, $\mathbf{v} \in \mathbb{F}^{\times d}$, si rappresenta con un vettore colonna $d \times 1$ formato dai suoi coefficienti c_i nella sua espressione come combinazione lineare dei vettori della base canonica e la trasformazione si tratta attraverso la sua matrice rappresentativa nelle due basi canoniche di $\mathbb{F}^{\times d}$ ed $\mathbb{F}^{\times e}$. Il vettore trasformato viene allora rappresentato dal vettore colonna $e \times 1$ ottenuto come prodotto righe per colonne della suddetta matrice per il suddetto vettore colonna.

Si osserva anche che le matrici che individuano le trasformazioni di $\left[\mathbb{F}^{\times d} \mapsto \mathbb{F}^{\times e} \right]$ sono in biiezione con le sequenze di $d \cdot e$ elementi del campo $\mathbb{F}$ e quindi consentono di costituire uno spazio vettoriale a $d \cdot e$ -dimensioni.

B32c.10 Prendiamo in esame alcune delle matrici quadrate che rappresentano trasformazioni lineari di $\mathbb{F}^{\times d}$ in se.

Le matrici dell'identità e delle omotetie sono

$$\begin{bmatrix} 1 & 0 & \dots & 0 \\ 0 & 1 & \dots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \dots & 1 \end{bmatrix} \quad \text{e} \quad \begin{bmatrix} \lambda & 0 & \dots & 0 \\ 0 & \lambda & \dots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \dots & \lambda \end{bmatrix}.$$

Le matrici dei proiettori sui vettori della base canonica sono

$$\begin{bmatrix} 1 & 0 & \dots & 0 \\ 0 & 0 & \dots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \dots & 0 \end{bmatrix}, \quad \begin{bmatrix} 0 & 0 & \dots & 0 \\ 0 & 1 & \dots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \dots & 0 \end{bmatrix} \quad \dots \quad \begin{bmatrix} 0 & 0 & \dots & 0 \\ 0 & 0 & \dots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \dots & 1 \end{bmatrix}.$$

Le matrici dei proiettori sopra $SFAC_{101}$, $SFAC_{1011}$, $SFAC_{01110}$ sono, risp.,

$$\begin{bmatrix} 1 & 0 & 0 \\ 0 & 0 & 0 \\ 0 & 0 & 1 \end{bmatrix}, \quad \begin{bmatrix} 1 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 \\ 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 1 \end{bmatrix} \quad \text{e} \quad \begin{bmatrix} 0 & 0 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 & 0 \\ 0 & 0 & 1 & 0 & 0 \\ 0 & 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 0 & 0 \end{bmatrix}.$$

Trasformazioni piuttosto semplici sono anche le combinazioni lineari di proiettori sui sottospazi generati dai vettori della base canonica $D = \sum_{i=1}^d \mathbf{Prj}_i \lambda_i$; l'effetto di tale trasformazione è

$$\sum_{i=1}^d \mathbf{Prj}_i \lambda_i \left(\sum_{j=1}^d \mathbf{e}_j v_j \right) = \sum_{i=1}^d \mathbf{e}_i \lambda_i v_i$$

ed essa è rappresentata nella base canonica dalla matrice diagonale

$$\begin{bmatrix} \lambda_1 & 0 & \dots & 0 \\ 0 & \lambda_2 & \dots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \dots & \lambda_d \end{bmatrix}.$$

B32c.11 Molte trasformazioni lineari del genere $[\mathbb{F}^{\times d} \mapsto \mathbb{F}^{\times e}]$ si trattano convenientemente manipolando le matrici che le rappresentano nelle basi canoniche. Conviene considerare per prime le matrici cosiddette **matrici diadi** che hanno una sola entrata diversa da 0 e uguale a 1.

In particolare consideriamo la matrice binaria $MD(j, k)$ avente come unica entrata uguale a 1 nella riga $j \in (d]$ e nella colonna $k \in (e]$ e quindi con le componenti esprimibili come

$$MD_{i,h}(j, k) = \delta_{i,j} \delta_{h,k}, \quad \text{per } i \in (d] \text{ e } h \in (e].$$

Le $d \cdot e$ matrici $MD(j, k)$ costituiscono la base canonica dello spazio vettoriale delle matrici su $\mathbb{F}$ di profilo $d \times e$.

La trasformazione rappresentata dalla matrice diade $MD(j, k)$, associata alla casella $\langle j, k \rangle$, annichila tutti i vettori $\mathbf{e}_i$ della base canonica $\mathfrak{B}_d$ dello spazio dominio $\mathbb{F}^{\times d}$ ad eccezione di $\mathbf{e}_j$ che trasforma nel vettore $\mathbf{u}_k$ della base canonica $\mathfrak{B}_e$ dello spazio codominio $\mathbb{F}^{\times e}$.

Ad esempio, convenendo che per gli indici si abbia $x = 1$, $y = 2$ e $z = 3$, abbiamo

$$MD_{y,z} \left(\sum_{j=1}^3 \mathbf{e}_j v_j \right) = \mathbf{e}_z v_y.$$

B32c.12 Vediamo per esempio come si decompone una matrice 3×2 come combinazione lineare delle 6 matrici diadi 3×2 :

$$\begin{bmatrix} p & q \\ r & s \\ t & u \end{bmatrix} = \begin{bmatrix} 1 & 0 \\ 0 & 0 \\ 0 & 0 \end{bmatrix} p + \begin{bmatrix} 0 & 1 \\ 0 & 0 \\ 0 & 0 \end{bmatrix} q + \begin{bmatrix} 0 & 0 \\ 1 & 0 \\ 0 & 0 \end{bmatrix} r + \begin{bmatrix} 0 & 0 \\ 0 & 1 \\ 0 & 0 \end{bmatrix} s + \begin{bmatrix} 0 & 0 \\ 0 & 0 \\ 1 & 0 \end{bmatrix} t + \begin{bmatrix} 0 & 0 \\ 0 & 0 \\ 0 & 1 \end{bmatrix} u .$$

Nel caso di trasformazione di uno spazio $\mathbb{F}^{\times d}$ in se, la diade corrispondente alla casella $\langle j, k \rangle$ trasforma il versore della base canonica $\mathbf{e}_j$ nel versore $\mathbf{e}_k$ e annichila ogni altro vettore della base; se in particolare $k = j$ coincide con il proiettore su $\mathbf{e}_j$: $MD_{j,j} = \mathbf{Prj}_j$.

B32c.13 Altre interessanti matrici sono le **matrici permutative**, matrici quadrate binarie che corrispondono a permutazioni di insiemi finiti di interi. Definiamo come matrice permutativa associata alla permutazione degli interi $1, 2, \dots, d$ $\mathbf{p} = \langle p_1, p_2, \dots, p_d \rangle \in \mathbf{Perm}_{(d)}$ la matrice binaria di profilo $d \times d$

$$\mathbf{Mprmt}(\mathbf{p}) := [i, j \in (d) : \delta_{i, p_j}] .$$

Per esempio alla permutazione $\langle 2, 3, 4, 1 \rangle$ risulta associata $\mathbf{Mprmt}(\langle 2, 3, 4, 1 \rangle) = \begin{bmatrix} 0 & 1 & 0 & 0 \\ 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 1 \\ 1 & 0 & 0 & 0 \end{bmatrix}$

Si osserva che una matrice permutativa $d \times d$ si esprime naturalmente come somma di d diadi; per la permutazione precedente:

$$\mathbf{Mprmt}(\langle 2, 3, 4, 1 \rangle) = MD_{2,3} + MD_{3,4} + MD_{4,1} + MD_{1,2} .$$

La applicazione della matrice $\mathbf{Mprmt}(\mathbf{p})$ al vettore colonna che rappresenta $\mathbf{e}_k$ fornisce il vettore colonna che costituisce la propria colonna k e questo, per la definizione della matrice, rappresenta il vettore $\mathbf{e}_{p_k}$. Per esempio

$$\begin{bmatrix} 0 & 1 & 0 & 0 \\ 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 1 \\ 1 & 0 & 0 & 0 \end{bmatrix} \begin{bmatrix} 0 \\ 1 \\ 0 \\ 0 \end{bmatrix} = \begin{bmatrix} 1 \\ 0 \\ 0 \\ 0 \end{bmatrix} \quad \text{e} \quad \begin{bmatrix} 0 & 1 & 0 & 0 \\ 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 1 \\ 1 & 0 & 0 & 0 \end{bmatrix} \begin{bmatrix} 0 \\ 0 \\ 0 \\ 1 \end{bmatrix} = \begin{bmatrix} 0 \\ 0 \\ 1 \\ 0 \end{bmatrix} .$$

Dunque la matrice associata a una permutazione di interi effettua la corrispondente permutazione dei vettori della base canonica.

Si vede subito che la permutazione identità è rappresentata dalla matrice $\hat{\mathbf{1}}_d$.

Si trova anche che al prodotto di due permutazioni di interi $\mathbf{p}_1 \cdot \mathbf{p}_2$ corrisponde il prodotto righe per colonne delle matrici permutative rappresentanti le due permutazioni:

$$\mathbf{Mprmt}(\mathbf{p}_1 \cdot \mathbf{p}_2) = \mathbf{Mprmt}(\mathbf{p}_1) \cdot \mathbf{Mprmt}(\mathbf{p}_2) .$$

È inoltre evidente che la permutazione inversa è rappresentata dalla matrice trasposta:

$$\mathbf{Mprmt}(\mathbf{p}^{-1}) = (\mathbf{Mprmt}(\mathbf{p}))^T .$$

B32 d. forme lineari e notazioni alla Dirac

B32d.01 Introduciamo ora le notazioni adottate da Paul Dirac per gli elementi degli spazi vettoriali e le relative trasformazioni lineari, spesso notevolmente vantaggiose.

Nel seguito ci serviremo spesso delle variabili intere positive $i, j \in (\mathbf{d}]$ e della delta di Kronecker, funzione di cui ricordiamo la definizione:

$$\delta_{i,j} := 0 \text{ sse } i \neq j \text{ e } \delta_{i,j} := 1 \text{ sse } i = j .$$

Introduciamo anche la funzione **valore binario di enunciato**, funzione avente come argomento un enunciato $\mathcal{E}$ per il quale si suppone possibile valutare la validità:

$$\text{Bval}(\mathcal{E}) := \begin{cases} 1 & \text{sse } \mathcal{E} \text{ è enunciato valido} \\ 0 & \text{sse } \mathcal{E} \text{ è enunciato nonvalido} \end{cases} .$$

Nel caso di enunciato esprimente la uguaglianza di due parametri interi i e j abbiamo

$$\delta_{i,j} := \text{Bval}(i = j) .$$

Seguendo Dirac identifichiamo ogni vettore $\mathbf{e}_i$ della base canonica con la scrittura equivalente $|i\rangle$ che si dice esprimere un **ket**. Di conseguenza il vettore ottenuto applicandogli la trasformazione lineare T si denota con $T|i\rangle$, semplificazione della scrittura $T(|i\rangle)$.

Si osserva che la componente j -esima del vettore $\mathbf{e}_i$ è $\delta_{i,j}$: quindi si può scrivere

$$\mathbf{e}_i = |i\rangle = \langle j \in (\mathbf{d}] : \delta_{i,j} \rangle =: \delta_{(\mathbf{d}] , i} ,$$

l'ultima scrittura introdotta con il ruolo della abbreviazione.

Si osserva anche che la matrice identità di ordine d , $\mathbf{1}_d$, si può considerare ottenuta dalla sovrapposizione dei d vettori riga rappresentanti i successivi versori della base canonica e, per la dualità della trasposizione, come l'affiancamento dei corrispondenti vettori colonna; quindi si ha la espressione

$$\mathbf{1}_d = [i, j \in (\mathbf{d}] : \delta_{i,j}] .$$

Per lo spazio tridimensionale $\mathbb{F}^{\times 3}$, adottando la solita convenzione per gli indici $x = 1, y = 2$ e $z = 3$, il vettore $\mathbf{v} = \langle v_x, v_y, v_z \rangle = \langle v_1, v_2, v_3 \rangle$ si può scrivere

$$|\mathbf{v}\rangle = \sum_{i=1}^3 v_i |i\rangle = v_x |1\rangle + v_y |2\rangle + v_z |3\rangle .$$

B32d.02 $\langle w|$ Le trasformazioni del genere $[\mathbb{F}^{\times d} \mapsto_{lin} \mathbb{F}]$ sono chiamate tradizionalmente **forme lineari dello spazio** $\mathbb{F}^{\times d}$.

Si osserva che le forme lineari possono essere sottoposte a combinazioni lineari, come accade a tutte le trasformazioni lineari. Quindi le forme lineari, come le trasformazioni lineari di qualsiasi genere $[\mathbb{F}^{\times d} \mapsto_{lin} \mathbb{F}^{\times e}]$, costituiscono uno spazio vettoriale; nel loro caso si ha $\mathbb{F}^{\times e} = \mathbb{F}^{\times d}$.

Le forme lineari su $\mathbb{F}^{\times d}$ si possono quindi rappresentare con matrici $1 \times d$, cioè da d elementi del campo, come i vettori di $\mathbb{F}^{\times d}$.

Questa semplice osservazione induce a stabilire un chiaro collegamento tra forme lineari e vettori che presentano la stessa sequenza di componenti, quale che sia lo spazio $\mathbb{F}^{\times d}$.

In effetti si può associare biunivocamente ad ogni vettore $\mathbf{v}$ dello spazio $\mathbb{F}^{\times d}$ una forma lineare detta **forma duale di un vettore**. Per definire questa associazione facciamo riferimento alla base canonica che ora scriviamo $\mathfrak{B}_d = \{|1\rangle, |2\rangle, \dots, |d\rangle\}$ e seguiamo ancora Dirac assumendo per essa la notazione $\langle i|$, da leggere **bra** di i .

Dunque al generico ket $|i\rangle$ associamo la forma lineare

$$\langle i | := \left[\mathbf{v} = \sum_{j=1}^d |j\rangle v_j \in \mathbb{F}^{\times d} \mapsto v_i \right] .$$

I valori assunti dalla forma $\langle i |$ in corrispondenza dei vari ket sono da scrivere $\langle i | (|h\rangle)$; questa scrittura la semplifichiamo con la più concisa $\langle i | j \rangle$.

Per le definizioni deve essere

$$\forall i, j = 1, 2, \dots, d : \langle i | j \rangle = \delta_{i,j} .$$

La forma lineare duale del generico vettore che riferito alla $\mathfrak{B}_d$ scriviamo $|\mathbf{w}\rangle = \sum_{i=1}^d w_i |i\rangle$, la denotiamo con $\langle \mathbf{w} |$.

Essa si esplicita grazie alla linearità della applicazione dualità:

$$\begin{aligned} \langle \mathbf{w} | &:= \sum_{i=1}^d w_i \langle i | = \left[\sum_{j=1}^d |j\rangle v_j \mapsto \sum_{i=1}^d w_i \langle i | \left(\sum_{j=1}^d |j\rangle v_j \right) \right] \\ &= \left[\sum_{j=1}^d |j\rangle v_j \mapsto \sum_{i,j=1}^d w_i \langle i | j \rangle v_j \right] = \left[\sum_{j=1}^d v_j |j\rangle \mapsto \sum_{i=1}^d w_i v_i \right] . \end{aligned}$$

A questo punto adottiamo la scrittura $\langle \mathbf{w} | \mathbf{v} \rangle$ come abbreviazione di $\langle \mathbf{w} | (|\mathbf{v}\rangle)$, coerentemente a quanto fatto p gli elementi delle basi. Possiamo quindi scrivere in modo compatto:

$$\forall \mathbf{w}, \mathbf{v} \in \mathbb{F}^{\times d} : \langle \mathbf{w} | \mathbf{v} \rangle = \sum_{i=1}^d w_i v_i .$$

B32d.03 La trasformazione lineare che manda i vettori (ket) di $\mathbb{F}^{\times d}$ nei vettori (bra) di $[\mathbb{F}^{\times d} \mapsto \mathbb{F}]$ si chiama **dualità vettoriale** e $[\mathbb{F}^{\times d} \mapsto \mathbb{F}]$ si dice **spazio vettoriale duale** di $\mathbb{F}^{\times d}$.

La dualità vettoriale è evidentemente una trasformazione biunivoca: in effetti essa corrisponde alla trasformazione dei vettori colonna, che rappresentano i ket $|\mathbf{v}\rangle$, nei vettori riga loro trasposti, che rappresentano i bra $\langle \mathbf{v} |$ loro duali.

Come ogni vettore di $\mathbb{F}^{\times d}$ è determinato dalle sue componenti nella base canonica, così ogni forma lineare è determinata dalla sua azione sui vettori della base canonica.

B32d.04 Le trasformazioni della base canonica si possono esprimere con notevole chiarezza mediante i bra e i ket.

Definiamo per ogni $i = 1, 2, \dots, d$ la trasformazione lineare

$$|i\rangle\langle i| := \left[\sum_{h=1}^d |h\rangle v_h \mapsto |i\rangle \sum_{h=1}^d \langle i | h \rangle v_h \right] = \left[\sum_{h=1}^d |h\rangle v_h \mapsto |i\rangle v_i \right] \in [\mathbb{F}^{\times d} \mapsto \mathbb{F} |i\rangle] .$$

Questo operatore quindi è il proiettore relativo ad $\mathbf{e}_i$:

$$|i\rangle\langle i| = \mathbf{Prj}_i .$$

Più in generale, per ogni $i, j = 1, 2, \dots, d$ si definisce l'operatore

$$|i\rangle\langle j| := \left[\sum_{h=1}^d |h\rangle v_h \mapsto |i\rangle \sum_{h=1}^d \langle j | h \rangle v_h \right] = \left[\sum_{h=1}^d |h\rangle v_h \mapsto |j\rangle v_j \right] \in [\mathbb{F}^{\times d} \mapsto \mathbb{F} |j\rangle] ;$$

Per $i \neq j$ questo operatore è quindi lo scambio di vettori $\mathbf{e}_i$ ed $\mathbf{e}_j$ della base canonica, cioè è rappresentato dalla matrice diade relativa alla coppia $\langle i, j \rangle$ che in c11 abbiamo denotata con $MD_{i,j}$.

B32d.05 Le notazioni di Dirac consentono di semplificare varie considerazioni sulle trasformazioni lineari, considerazioni utili in varie circostanze.

Per trattare una generica $T \in [\mathbb{F}^{\times d} \mapsto \mathbb{F}^{\times e}]$ occorre fare riferimento alle basi canoniche dei due spazi.

Mentre per la prima usiamo la solita notazione $\mathfrak{B}_d = \{ \langle 1|, \langle 2|, \dots, \langle d| \}$, per la seconda qui scriviamo $\mathfrak{B}_e = \{ \langle \bar{1}|, \langle \bar{2}|, \dots, \langle \bar{e}| \}$: i kets con l'indice sovrabbarrato sono quindi vettori e -dimensionali.

Sono utili le decomposizioni mediante proiettori delle identità dei due spazi $\mathbb{F}^{\times d}$ e $\mathbb{F}^{\times e}$:

$$(1) \quad \hat{\mathbf{1}}_d = \sum_{i=1}^d |i\rangle\langle i| \quad \text{e} \quad \hat{\mathbf{1}}_e = \sum_{\bar{h}=1}^e |\bar{h}\rangle\langle \bar{h}|.$$

Il formalismo dei bra e dei ket consente utili semplificazioni di scrittura che si avvalgono della eliminazione delle parentesi che delimitano gli argomenti di alcune funzioni.

Dopo aver semplificato la scrittura $T(|\mathbf{v}\rangle)$ nella $T|\mathbf{v}\rangle$, semplifichiamo anche la scrittura $\langle \mathbf{w}|(T(|\mathbf{v}\rangle))$ nella $\langle \mathbf{w}|T|\mathbf{v}\rangle$.

In particolare le entrate della matrice canonica della trasformazione T sono denotate con $\langle i|T|\bar{h}\rangle$.

Con questo formalismo l'azione di una trasformazione T sopra un vettore $\mathbf{v}$ espressa dalla c08(3) si riscrive come

$$T|\mathbf{v}\rangle = \mathbf{1}_e T \mathbf{1}_d |\mathbf{v}\rangle = \sum_{\bar{h}=1}^e |\bar{h}\rangle\langle \bar{h}| T \left(\sum_{i=1}^d |i\rangle\langle i|(|\mathbf{v}\rangle) \right) = \sum_{\bar{h}=1}^e |\bar{h}\rangle \left(\sum_{i=1}^d \langle \bar{h}|T|i\rangle \langle i|\mathbf{v}\rangle \right).$$

Osserviamo che in questo sviluppo compaiono i vettori delle basi canoniche di due spazi e dei loro duali contraddistinti dall'uso di lettere, l'una sovrabbarrata l'altra no, che corrono su insiemi diversi.

B32d.06 Eserc. Si consideri una permutazione dell'insieme $\{1, 2, \dots, d\}$ $\mathbf{p} = \langle p_1, p_2, \dots, p_d \rangle$ e la corrispondente matrice permutativa $\mathbf{Mprmt}(\mathbf{p})$; constatare che

$$\mathbf{Mprmt}\langle p_1, p_2, \dots, p_d \rangle = \sum_{i=1}^d |p_i\rangle\langle i|.$$

B32 e. rassegna di trasformazioni lineari e matrici in 2D e 3D

B32e.01 In questo paragrafo presentiamo una certa gamma di matrici di profilo 2×2 e 3×3 le quali consentono di familiarizzarsi con le trasformazioni lineari che le matrici rappresentano e con i rispettivi significati geometrici e che facilitano la introduzione dei problemi di calcolo basilari dell'algebra lineare. Queste matrici inoltre costituiscono degli strumenti primari per lo studio analitico delle configurazioni del piano e dello spazio tridimensionale.

B32e.02

B32e.03 Determinanti 2×2 e 3×3 . Gruppi S_2 e S_3

B32 f. cambiamenti di base e prodotti di trasformazioni lineari

B32f.01 Consideriamo un automorfismo di uno spazio vettoriale

Esso trasforma una base in una seconda base.

La matrice della trasformazione inversa è l'inversa della matrice della trasformazione diretta.

B32f.02

B32f.03

B32f.04 Fissiamo una base $\mathcal{B} = \langle \mathbf{u}_1, \mathbf{u}_2, \dots, \mathbf{u}_d \rangle$ ed un suo elemento $\mathbf{u}_h$ e consideriamo il vettore $\mathbf{x} = \sum_{j=1}^d x_j \mathbf{u}_j$; definiamo come proiezione di $\mathbf{x}$ sulla direzione $\mathbf{u}_h$ il vettore $x_h \mathbf{u}_h$.

In generale si definisce una **proiezione di un vettore su una direzione**.

Più estesamente si considera un sottoinsieme $\mathcal{C}$ della base, $\mathcal{C} \subseteq \mathcal{B}$, il corrispondente insieme di indici $C \subseteq \{1, 2, \dots, d\}$ e il sottospazio $\text{span}(\mathcal{C})$ sotteso dai vettori di $\mathcal{C}$; si definisce proiezione di $\mathbf{x}$ sul sottospazio $\text{span}(\mathcal{C})$, il vettore $\sum_{h \in C} x_h \mathbf{u}_h$.

In generale si definisce una **proiezione di un vettore su un sottospazio**.

Esempi di endomorfismi sono i proiettori nelle diverse direzioni della base $\mathcal{B}$ e nei diversi sottospazi $\text{span}(\mathcal{C})$.

Vi sono poi le permutazioni delle componenti: anche queste trasformazioni sono degli isomorfismi.

B32f.05 Alla composizione di trasformazioni lineari corrisponde il prodotto di matrici.

B32f.06 Fissate una base per $\mathbb{F}^{\times d}$ e una per $\mathbb{F}^{\times e}$, gli omomorfismi di $\boxed{\mathbb{F}^{\times d} \mapsto \mathbb{F}^{\times e}}$ sono rappresentati fedelmente da matrici di profilo $d \times e$.

B32 g. determinanti

B32g.01 In questo paragrafo introduciamo la nozione generale di determinante di una matrice quadrata di profilo $d \times d$ sopra un campo $\mathbb{F}$, generalizzando quella di determinante di una matrice 2×2 introdotta in B24d06.

Si tratta di una funzione delle entrate delle matrici quadrate che riveste grande importanza per le trasformazioni che le matrici rappresentano, importanza collegata a un suo significato geometrico. Purtroppo la sua definizione generale è piuttosto elaborata, mentre è ben chiara nei casi bidimensionale e tridimensionale, casi nei quali il significato geometrico del determinante risulta intuitivo e chiaramente utilizzabili, in particolare per la comprensione delle soluzioni dei sistemi di equazioni lineari.

B32g.02 Consideriamo la matrice quadrata di profilo $d \times d$

$$A = \begin{bmatrix} a_{1,1} & a_{1,2} & \dots & a_{1,d} \\ a_{2,1} & a_{2,2} & \dots & a_{2,d} \\ \vdots & \vdots & \ddots & \vdots \\ a_{d,1} & a_{d,2} & \dots & a_{d,d} \end{bmatrix}.$$

Consideriamo anche una generica permutazione degli interi $1, 2, \dots, d$ che scriviamo $\mathbf{p} = \langle p_1, p_2, \dots, p_d \rangle$, e la sequenza di n entrate della matrice A associata alla $\mathbf{p}$ della forma

$$\mathbf{a}_{\mathbf{p}} := \langle a_{1,p_1}, a_{2,p_2}, \dots, a_{d,p_d} \rangle.$$

I successivi componenti della $\mathbf{a}_{\mathbf{p}}$ provengono dalle successive righe della matrice e, per il carattere permutativo della $\mathbf{p}$, a ciascuna colonna della matrice appartiene uno e uno solo degli a_{j,p_j} .

Ricordiamo che alla $\mathbf{p}$ è associato il segno $\text{sign}(\mathbf{p})$, numero che vale $+1$ se la permutazione è pari (ovvero se si può ricondurre alla permutazione crescente $\langle 1, 2, \dots, d \rangle$ con un numero pari di scambi) e che vale -1 se essa è dispari (ovvero si riconduce alla permutazione crescente con un numero dispari di scambi).

Consideriamo poi il prodotto associato alla $\mathbf{p}$

$$\pi_{\mathbf{p}} := \text{sign}(\mathbf{p}) \prod_{i=1}^d a_{i,p_i}.$$

Per definire il determinante di A , che denotiamo con $\det(A)$, servono tutti i prodotti $\pi_{\mathbf{p}}$ ottenibili dall'insieme $\mathbf{Perm}_d$ delle $d!$ permutazioni degli interi di $[d]$.

Infatti definiamo **determinante della matrice quadrata A**

$$(1) \quad \det(A) := \sum_{\mathbf{p} \in \mathbf{Perm}_d} \text{sign}(\mathbf{p}) \prod_{j=1}^d a_{j,p_j}.$$

B32g.03 Il secondo membro della g02(1) si dice **sviluppo del determinante** e gli addendi $\pi_{\mathbf{p}}$ di questo sviluppo si chiamano **termini del determinante**.

Le entrate della matrice potrebbero anche essere d^2 variabili nel campo $\mathbb{F}$; in questo caso il determinante si può considerare un'espressione polinomiale in queste variabili.

Per le matrici dei primi ordini le espressioni dei determinanti sono semplici e facili da ricordare:

$$\det[a_{1,1}] = 1 \quad , \quad \det \begin{bmatrix} a_{1,1} & a_{1,2} \\ a_{2,1} & a_{2,2} \end{bmatrix} = a_{1,1}a_{2,2} - a_{1,2}a_{2,1} ,$$

$$\det \begin{bmatrix} a_{1,1} & a_{1,2} & a_{1,3} \\ a_{2,1} & a_{2,2} & a_{2,3} \\ a_{3,1} & a_{3,2} & a_{3,3} \end{bmatrix} = a_{1,1}a_{2,2}a_{3,3} + a_{1,2}a_{2,3}a_{3,1} + a_{1,3}a_{2,1}a_{3,2} - a_{1,3}a_{2,2}a_{3,1} - a_{1,1}a_{2,3}a_{3,2} - a_{1,2}a_{2,1}a_{3,3} .$$

Lo sviluppo del determinante di una matrice di ordine 3 si può ottenere con relativa semplicità con un procedimento chiamato **regola di Sarrus**: si tratta di allargare la matrice affiancandole, nell'ordine, una replica della prima colonna e una replica della seconda: su questo schieramento di elementi del campo si considera la somma dei tre prodotti ottenuti dalla diagonale principale e dalle due linee oblique a essa parallele e a essa si sottraggono i tre prodotti della codiagonale e delle due linee oblique a essa parallele.

Quando cresce l'ordine della matrice il calcolo effettivo del determinante si fa sempre più pesante e pronò agli errori.

In effetti per calcolare determinanti generici in modo efficiente è necessario conoscere meglio varie loro proprietà; in particolare risulta utile riconoscere loro proprietà collegate a simmetrie.

B32g.04 Lo sviluppo del determinante si può esprimere anche come

$$(1) \quad \det(A) := \sum_{\mathbf{p} \in \text{Perm}_d} \text{sign}(\mathbf{p}) \prod_{h=1}^d a_{p_h, h} .$$

I termini di questo sviluppo corrispondono biunivocamente con quelli di g02(1): infatti la commutatività del prodotto in un campo consente di scrivere

$$\text{sign}(\mathbf{p}) \prod_{h=1}^d a_{p_h, h} = \text{sign}(\mathbf{p}^{-1}) \prod_{j=1}^d a_{j, q_j}$$

con $\mathbf{q} = \langle q_1, q_2, \dots, q_d \rangle := \mathbf{p}^{-1}$, avendo anche tenuto conto che il segno di una permutazione coincide con quello della sua inversa.

Il precedente sviluppo si può leggere come sviluppo del determinante della matrice trasposta della A e comporta l'enunciato che segue.

(2) Prop.: Per ogni matrice quadrata A si ha: $\det(A^T) = \det(A)$ ■

In altre parole la trasposizione di una matrice lascia invariato il suo determinante.

Questo enunciato consente di trasformare proprietà del determinante di una matrice A concernenti le righe di questa in proprietà duali concernenti le colonne di A e viceversa.

Coppie di proprietà di questo genere si dicono **proprietà duali per trasposizione di matrici**.

B32g.05 Il determinante di una matrice quadrata A si può utilmente associare alla sequenza dei vettori (ket) costituenti le colonne di A oppure alla sequenza delle forme lineari (bra) costituenti le sue righe.

(1) Prop.: Se si moltiplicano per uno scalare λ tutte le entrate di una linea di una matrice, allora il suo determinante viene moltiplicato per tale λ .

Dim.: Nello sviluppo del determinante della nuova matrice ogni termine viene moltiplicato per λ e lo stesso accade alla loro somma, cioè al determinante ■

(2) Coroll.: Se si moltiplica una matrice $d \times d$ per uno scalare λ , allora il suo determinante viene moltiplicato per λ^d ■

Dim.: Ogni termine dello sviluppo del determinante della nuova matrice viene moltiplicato per λ^d .

(3) Coroll.: Il determinante di una matrice con una linea con tutte le entrate nulle vale 0.

Dim.: Ogni termine dello sviluppo del determinante della nuova matrice contiene un fattore 0 ■

(4) Prop.: Il determinante di una matrice avente una linea, cui diamo l'indice $\bar{i}$, esprimibile come somma termine a termine di due d -uple è dato dalla somma dei determinanti delle due matrici che nella linea $\bar{i}$ presentano una delle due d -uple in esame.

Dim.: Nello sviluppo del determinante il termine corrispondente alla permutazione $\mathbf{p} \in \text{Perm}_d$ ha la forma $\text{sign}(\mathbf{p}) (a'_{i,p_i} + a''_{i,p_i}) \prod_{h \neq i} a_{hp_h}$ e si può esprimere come somma della forma

$$\text{sign}(\mathbf{p}) a'_{i,p_i} \prod_{h \neq i} a_{hp_h} + \text{sign}(\mathbf{p}) a''_{i,j} \prod_{h \neq i} a_{hp_h}.$$

Sommando sulle permutazioni $\mathbf{p} \in \text{Perm}_d$ si ottiene la somma dei due determinanti sopra definiti ■

B32g.06 (1) Prop.: Se si scambiano due linee di una matrice quadrata, il suo determinante cambia di segno.

Dim.: Consideriamo il caso della matrice A e della matrice A' ottenuta dalla A scambiando le colonne h e k con $h < k$.

Considerando gli sviluppi g02(1) delle due matrici, al termine $\text{sign}(\mathbf{p}) \cdots a_{h,p_h} \cdots a_{k,p_k} \cdots$ della prima matrice si fa corrispondere biunivocamente nella A' il termine $\text{sign}(\mathbf{p}') \cdots a_{k,p_k} \cdots a_{h,p_h} \cdots$, dove $\mathbf{p}'$ denota una permutazione ottenuta dalla $\mathbf{p}$ con uno scambio e quindi avente $\text{sign}(\mathbf{p}') = -\text{sign}(\mathbf{p})$; questo cambiamento di segno dei termini corrispondenti delle due permutazioni implica $\det(A') = -\det(A)$.

Se si scambiano due righe della A si può condurre un discorso analogo sugli sviluppi della forma g04(1). All'enunciato sullo scambio di due righe si giunge anche considerando che questa trasformazione equivale alla composizione della trasposizione della A , allo scambio di due colonne e a una seconda trasposizione; dato che le trasposizioni non cambiano il determinante, l'effetto dello scambio di due righe coincide con l'effetto allo scambio di due colonne, cioè al cambiamento del segno ■

(2) Prop.: Se una matrice quadrata ha due righe uguali o due colonne uguali, allora il suo determinante è nullo.

Dim.: Il determinante di una tale matrice A deve essere l'opposto del determinante della matrice ottenuta scambiando le due linee uguali, cioè della stessa matrice A : dunque $\det(A) = -\det(A)$ e quindi $\det(A) = 0$ ■

(3) Coroll.: Una matrice quadrata con due linee proporzionali ha il determinante nullo.

Dim.: Moltiplicando per un opportuno scalare λ una delle due linee si ottiene una matrice il cui determinante risulta moltiplicato per λ , ma che presenta due linee uguali e quindi è nullo ■

(4) Prop.: Il determinante di una matrice non viene modificato sommando a una sua linea (riga o colonna) una qualsiasi combinazione lineare di altre linee (risp. righe o colonne).

Dim.: Sommando a una riga i un'altra riga $\bar{i}$ di A , il determinante viene aumentato del determinante della matrice con la riga i sostituita dalla $\bar{i}$, cioè con due righe uguali, cioè di un determinante nullo.

Sommando a una riga i un'altra riga $\bar{i}$ moltiplicata per uno scalare λ , $\det(A)$ viene aumentato del determinante di una matrice con due righe proporzionali, cioè $\det(A)$ non viene modificato.

La modifica generale dell'enunciato si ottiene effettuando più volte le manovre precedenti che non cambiano $\det(A)$ e ripetendo le considerazioni per le colonne, ovvero invocando la invarianza per trasposizione della matrice ■

B32g.07 Il calcolo del determinante di una matrice con il crescere del suo ordine d aumenta rapidamente, in accordo con la rapida crescita del fattoriale $d!$ e quindi del crescere come $\frac{d}{d-1}(d-1)d! = d d!$

del numero delle moltiplicazioni necessarie per effettuare il calcolo seguendo pedissequamente la definizione; si trascura di considerare il costo delle somme, assai meno costose delle moltiplicazioni.

Dato che il calcolo dei determinanti di matrici di ordini elevati (in molte applicazioni si trattano matrici di migliaia e anche milioni di linee) risulta essenziale per risolvere molti problemi applicativi di grande rilievo, risulta importante individuare procedimenti efficienti per queste elaborazioni numeriche.

Evidentemente il determinante di una matrice A diagonale è dato dal semplice prodotto $\prod_{j=1}^d a_{j,j}$ delle entrate sulla diagonale principale che richiede solo $d - 1$ moltiplicazioni.

Il prodotto della forma precedente fornisce anche il determinante di una matrice triangolare inferiore oppure superiore (si osserva in effetti che per trasposizione una matrice triangolare superiore diventa triangolare inferiore e viceversa).

Si è quindi indotti a ricondurre il calcolo del determinante di una matrice A a quello di una matrice A' che abbia poche entrate diverse da 0 al di fuori delle diagonale principale.

Si deve comunque seguire il criterio di ottenere un determinante $\det(A)$ servendosi di determinanti di matrici ottenute dalla A con modifiche che lascino invariato il determinante, oppure con modifiche che si possono controllare facilmente.

B32g.08 Consideriamo due interi $d, e \geq 2$ e una matrice A di profilo $d \times e$.

Ricordiamo che si dice **sottomatrice** della A ogni matrice ottenuta eliminando dalla A un certo numero $r \in [1 : d - 1]$ di righe e un certo numero $s \in [1 : e - 1]$ di colonne.

Consideriamo inoltre $I = \langle i_1, \dots, i_p \rangle$ una sottosequenza propria della $\langle 1, \dots, d \rangle$ e $J = \langle j_1, \dots, j_q \rangle$ una sottosequenza propria della $\langle 1, \dots, e \rangle$.

Con la scrittura $A|_{I \times J}$ denotiamo la sottomatrice della A ottenuta mantenendo di essa solo le righe fornite dalla sottosequenza I e le colonne corrispondenti alla J .

Denotiamo poi con la scrittura $A|_{-I \times J}$ la sottomatrice ottenuta dalla A eliminando le righe fornite dalla I e le colonne relative alla J . Quest'ultima si dice **sottomatrice complementare della sottomatrice** $A|_{I \times J}$.

Delle matrici quadrate $d \times d$ interessano in particolare le sottomatrici quadrate di ordine $d - 1$, cioè le sottomatrici complementari delle sottomatrici costituite da una sola entrata. Per la sottomatrice complementare di quella relativa all'entrata $a_{i,j}$ usiamo la notazione $A|_{-i,j}$, notazione che evidenzia come questa sia associata alla casella $\langle i, j \rangle$.

Chiamiamo poi **minore complementare della matrice** A relativo alla sua casella $\langle i, j \rangle$ il determinante della precedente sottomatrice e scriviamo $\text{Mnrc}_{i,j}(A) := \det(A|_{-i,j})$.

Consideriamo lo sviluppo di $\det(A)$ e isoliamo i termini contenenti $a_{1,1}$; chiaramente la somma di questi termini è data da $a_{1,1} \text{Mnrc}_{1,1}$.

Se per generici $h, k \in (d]$ isoliamo i termini contenenti il fattore $a_{h,k}$, la loro somma è data da $a_{h,k} (-1)^{h+k} \text{Mnrc}_{h,k}(A)$, in quanto $\det(A) = (-1)^{h-1+k-1} \det(A')$ dove A' è la matrice ottenuta dalla A portando la riga h nella prima posizione e la colonna k nella prima posizione e osservando che $(-1)^{h-1+k-1} = (-1)^{h+k}$.

Il prodotto $(-1)^{h+k} \text{Mnrc}_{h,k}(A)$ viene detto **cofattore** o **complemento algebrico** della A relativo all'entrata $a_{h,k}$; qui nel seguito lo denotiamo con $\text{Cftr}_{h,k}(A)$.

Si dice **segno di una casella** $\langle h, k \rangle \in (d] \times (e]$ l'intero $(-1)^{h+k}$.

Si osserva che i due segni delle caselle si distribuiscono similmente alle caselle bianche e nere della classica scacchiera. Per esempio per le matrici 3×4 e 5×5 abbiamo:

$$\begin{bmatrix} + & - & + & - \\ - & + & - & + \\ + & - & + & - \end{bmatrix} \quad \text{e} \quad \begin{bmatrix} + & - & + & - & + \\ - & + & - & + & - \\ + & - & + & - & + \\ - & + & - & + & - \\ + & - & + & - & + \end{bmatrix}$$

B32g.09 Possiamo ora ottenere il cosiddetto sviluppo di un determinante secondo le entrate di una linea.

Consideriamo la riga h della matrice A e separiamo i termini dello sviluppo del suo determinante in d parti, la prima contenente il fattore $a_{h,1}$, la seconda il fattore $a_{h,2}$, ... , la d -esima il fattore $a_{h,d}$.

Come si è visto sommando i termini contenenti $a_{h,k}$ si ottiene $a_{h,k} \text{Cftr}_{h,k}(A)$. Quindi per il determinante

$$(1) \quad \det(A) = \sum_{j=1}^d a_{h,j} \text{Cftr}_{h,j}(A) .$$

Dualmente, scambiando righe e colonne, si ottiene lo sviluppo secondo la colonna k :

$$(2) \quad \det(A) = \sum_{i=1}^d a_{i,k} \text{Cftr}_{i,k}(A) .$$

B32g.10 Consideriamo alcuni esempi.

È utile osservare che le formule precedenti si semplificano quando si incontrano righe o colonne con più entrate nulle, in quanto i relativi addendi si possono eliminare.

B32g.11 Scelti due interi $p, h \in (d]$ con $p \neq h$, se nella (1) per $j = 1, 2, \dots, d$ si sostituiscono, risp., gli $A_{h,j}$ con gli $A_{p,j}$, si ottiene il determinante della matrice ottenuta dalla A sostituendo la riga h con la riga p ; questo determinante di una matrice con due righe uguali è notoriamente nullo.

La conclusione duale-rc si ottiene per lo sviluppo secondo una colonna. Si hanno quindi le due formule

$$(3) \quad \forall h, p \in (d] : \sum_{j=1}^d a_{h,j} A_{p,j} = \delta_{h,p} \det(A) ,$$

$$(4) \quad \forall h, p \in (d] : \sum_{i=1}^d a_{i,k} A_{i,q} = \delta_{k,q} \det(A) .$$

B32 h. invertibilità di una matrice quadrata

B32h.01 La possibilità di stabilire se una matrice quadrata è invertibile e il calcolo effettivo della sua matrice inversa costituiscono operazioni di grande importanza generale e applicativa.

Infatti la conoscenza di una matrice inversa fornisce la possibilità di controllare una trasformazione inversa e molti problemi riguardano la determinazione di punti di uno spazio $\mathbb{F}^{\times d}$ che vengono mandati da una trasformazione lineare in altrettanti punti conosciuti dello stesso spazio o di uno spazio della forma $\mathbb{F}^{\times e}$.

Consideriamo quindi una matrice quadrata A di profilo $d \times d$ e cerchiamo se possiede inversa A^{-1} , cioè la matrice $d \times d$ tale che $A \nabla A^{-1} = A^{-1} \nabla A = \mathbf{1}_d$ e studiamo qualche formula o procedimento che consente di calcolare questa A^{-1} .

B32h.02 Teorema (teorema di Cauchy-Binet)

Se A e B sono due matrici quadrate dello stesso ordine, allora

$$\det(A \cdot B) = \det(A) \cdot \det(B) .$$

Dim.:

B32h.03 Coroll.: Il determinante di una matrice A^{-1} inversa di una data A è il reciproco del suo determinante: $\det(A^{-1}) = \frac{1}{\det(A)}$.

Dim.: Se d è l'ordine delle matrici si ha $A \nabla A^{-1} = \mathbf{1}_d$; dal teorema precedente segue $\det(A) \det(A^{-1}) = \det(\mathbf{1}_d) = 1$ ■

B32h.04 Prop. Una matrice invertibile non è singolare.

Dim.: L'uguaglianza precedente implica che non può essere $\det(A) = 0$ ■

B32h.05 Teorema Una matrice nonsingolare è invertibile ed ha come inversa la matrice

$$A^{-1} = \frac{1}{\det(A)} \begin{bmatrix} A_{1,1} & A_{1,2} & \dots & A_{1,n} \\ A_{2,1} & A_{2,2} & \dots & A_{2,n} \\ \vdots & \vdots & \ddots & \vdots \\ A_{n,1} & A_{n,2} & \dots & A_{n,n} \end{bmatrix}^{\top} ,$$

nella quale interviene la trasposta della matrice le cui entrate $A_{i,j}$ sono i cofattori delle entrate $a_{i,j}$ della matrice data, $A_{i,j} := \text{Cftr}_{i,j}(A)$.

Dim.: Le equazioni g11(3) e g11(4) si possono riscrivere come

$$(1) \quad \forall h, p \in (d] : \sum_{j=1}^d a_{h,j} (A^{-1})_{j,p} = \delta_{h,p} \det(A) ,$$

$$(2) \quad \forall h, p \in (d] : \sum_{i=1}^d (A^{-1})_{q,i} a_{i,k} = \delta_{k,q} \det(A) .$$

Queste costituiscono una formulazione dell'enunciato ■

L'esposizione in <https://www.mi.imati.cnr.it/alberto/> e https://arm.mi.imati.cnr.it/Matexp/matexp_main.php