

ALMA MATER STUDIORUM · UNIVERSITÀ DI BOLOGNA

Dipartimento di Fisica e Astronomia, "Augusto Righi"
Corso di Laurea in Fisica

Algebre di Hopf e ZX-Calculus per la semplificazione di circuiti quantistici

Relatore:
Prof. Emanuele Latini

Presentata da:
Sara Maiano

Anno Accademico 2025/2026

Sommario

Il calcolo quantistico rappresenta ad oggi uno dei campi di ricerca più promettenti della fisica contemporanea. I computer quantistici nascono direttamente dalla fisica fondamentale. La loro unità elementare di informazione, il qubit, rappresenta la possibilità di ragionare secondo una logica alternativa alla logica binaria dei bit, sfruttando fenomeni puramente quantistici come la sovrapposizione di stati e l'entanglement. Per questo motivo negli ultimi anni, parallelamente al veloce concretizzarsi della realizzazione di queste nuove tecnologie è emersa inoltre la necessità di costruire un linguaggio teorico che sia maggiormente intuitivo e maneggevole degli operatori lineari propri della meccanica quantistica, allo scopo di rendere meno macchinosi i calcoli degli algoritmi sui qubit che diventano ormai via via sempre più complessi. La soluzione a questa necessità sembra essere ad oggi un nuovo linguaggio diagrammatico sviluppato nell'ambito della *Categorical Quantum Mechanics*, che prende il nome di ZX-Calculus, per il quale è stata recentemente dimostrata la completezza rispetto alle operazioni logiche sui qubit. Obiettivo di questa tesi è fornire al lettore gli strumenti matematici di base per accostarsi, ad un primo livello introduttivo, al linguaggio dello ZX-Calculus e alle strutture algebriche che ne conseguono. Nel primo capitolo presentiamo quindi le definizioni di algebre di Hopf che ritroveremo nelle strutture degli ZX-diagrams. Nel secondo capitolo si forniscono le basi teoriche del Quantum Computing presentando il sistema fisico del qubit a partire dalla descrizione rigorosa dei sistemi quantistici. Nel terzo capitolo infine si affronta la costruzione degli ZX-diagrams e delle regole che costituiscono lo ZX-Calculus.

Indice

1	Algebre di Hopf	7
1.1	Esempio di algebra di Hopf : Funzioni sulle matrici del gruppo $SL(N, \mathbb{C})$	7
1.2	Definizioni di algebra, coalgebra e bialgebra	11
1.3	Algebre di Hopf	23
1.3.1	Coppie duali di algebre di Hopf	27
1.3.2	Esempi di algebre di Hopf	33
1.3.3	Descrizione grafica delle algebre di Hopf	35
2	Quantum Computing	39
2.1	Informatica classica	39
2.1.1	Porte logiche a un bit	40
2.1.2	Porte logiche a due bit	40
2.2	Definizione del qubit nel formalismo della meccanica quantistica	42
2.2.1	Descrizione statistica dei sistemi quantistici	51
2.3	Informatica quantistica	58
2.3.1	Circuiti quantistici	64
2.3.2	I vantaggi dell'informatica quantistica	66
3	Z-X Calculus e algebre di Hopf	72
3.1	Il linguaggio degli ZX-diagrams	72
3.2	Lo ZX-Calculus	75
3.3	Interpretazione dello ZX-Calculus nello spazio di Hilbert	81
3.3.1	Universalità dello ZX-Calculus	87
3.4	Algebre di Hopf nello ZX-Calculus	89
4	Appendice: Algebre di Hopf non normalizzate e F-algebre di Hopf	92

Introduzione

Nel 1981 Richard Feynman ad una celebre conferenza del MIT, propose quella che ancora oggi ricordiamo come la prima idea rivoluzionaria che ci ha condotto a pensare, e realizzare successivamente, quello che oggi chiamiamo un computer quantistico (e che per mia particolare fortuna proprio quest'anno che scrivo questa tesi, è per la prima volta presente nei laboratori del tecnopolo DAMA di Bologna in due versioni differenti NOX e LUX !). Per capire però quale fu l'idea brillante che Feynman propose nel suo lavoro "*Simulating Physics with Computers (1982)*" ripercorriamo velocemente il percorso che la fisica aveva affrontato negli anni precedenti a quella famosa conferenza.

Alla fine degli anni '70 la meccanica quantistica aveva raggiunto risultati straordinari. Si erano ormai sviluppate formalmente, la meccanica matriciale di Heisenberg (1925), la meccanica ondulatoria di Schrodinger (1926), la notazione di Dirac ma anche l'elettrodinamica quantistica (QED), completata negli anni '40 e '50 proprio da Feynman, Schwinger e Tomonaga. Quest'ultima in particolare negli anni '70 era probabilmente la teoria fisica più precisa mai costruita. A quel punto quindi praticamente nessuno dubitava più che la natura fosse fondamentalmente quantistica. Il problema perciò non era capire quali fossero le leggi della natura ma riuscire ad applicarle e qui nacque il vero problema. Le equazioni della meccanica quantistica erano perfettamente note ma la loro risoluzione analitica per i cosiddetti problemi Many-body esplodeva esponenzialmente in difficoltà all'aumentare del numero di particelle coinvolte. Questo rappresentava un problema enorme in ambiti come la chimica quantistica, la fisica della materia condensata, la teoria dei materiali, la superconduttività.

Allo stesso tempo alla fine degli anni '70 anche i computer classici erano ormai diventati una realtà diffusa a cui ci si era abituati e su cui si erano sviluppate teorie formali come quelle di computazione, complessità, macchina di Turing. Per questo era nata una branca della fisica, la fisica computazionale, convinta che saremmo riusciti a simulare qualsiasi sistema con l'aiuto di questa nuova tecnologia. Ma a dare filo da torcere a questo proposito c'erano proprio quei problemi many-body che abbiamo precedentemente citato ed è esattamente qui che nasce l'idea geniale di Richard Feynman, perchè: "*Nature isn't classical, dammit, and if you want to make a simulation of nature, you'd better make it quantum mechanical.*"

Quello che propose Feynman in questa famosa conferenza fu di pensare per la prima volta che potesse esistere un simulatore della natura che ne investigasse il comportamento lavorando proprio nel contesto più profondo dove questa vive, il mondo della meccanica quantistica.

Bisognerà poi aspettare qualche anno per il primo modello matematico rigoroso di computer quantistico universale con il lavoro di David Deutsch nel 1985 "*Quantum Theory, the Church-Turing Principle and the Universal Quantum Computer*", il 1992 per il primo algoritmo con vantaggio quantistico di Deutsch e Jozsa nel lavoro "*Rapid Solution of Problems by Quantum Computation*", David Simon nel 1994 per il primo vantaggio espo-

nenziale fino a Peter Shor col suo algoritmo di fattorizzazione dei numeri interi che mise in discussione la sicurezza dei principali protocolli crittografici a chiave pubblica.

L'intuizione di Feynman insomma diede origine a un nuovo paradigma di calcolo che negli anni successivi conobbe una rapida crescita con l'invenzione dei primi algoritmi e fu Feynman stesso a proporre una prima rappresentazione grafica circuitale come linguaggio standard del calcolo quantistico, circuiti composti da fili, porte logiche e misure. A quasi cinquant'anni dalla prima intuizione di computer quantistico lanciata da Feynman a quella famosa conferenza del MIT, competono oggi tra loro diverse tecnologie sviluppate nell'ambito della ricerca fondamentale, tra cui qubit superconduttivi, ioni intrappolati, atomi neutri e circuiti fotonici, le cui prospettive più promettenti riguardano la simulazione quantistica della materia e delle molecole, lo studio del comportamento della materia a livello atomico, con possibili ricadute nella chimica e nella farmaceutica e per quanto riguarda la fisica, lo studio delle interazioni tra particelle elementari oltre che alla soluzione di problemi di ottimizzazione estremamente complessi. E' ancora presto per capire quale tecnologia si consoliderà e se sarà una soltanto a prevalere. Così come oggi utilizziamo processori diversi per compiti diversi, in futuro potremmo scoprire che alcune architetture quantistiche sono più adatte alla simulazione di sistemi fisici, altre all'ottimizzazione di reti complesse, altre ancora alla progettazione di nuovi materiali o farmaci. Le applicazioni decisive restano in parte ancora da scoprire.

Col crescere però della complessità dei circuiti e degli algoritmi che si riuscivano a teorizzare è nata negli ultimi decenni l'esigenza di sostituire al linguaggio naturale delle matrici unitarie come operazioni logiche sui bit quantistici (i qubit), un linguaggio che fosse maggiormente intuitivo e maneggevole. Ed è proprio su questo proposito che si sviluppa questo lavoro di tesi.

Alla fine degli anni '90 e nei primi anni 2000 il calcolo quantistico era descritto quasi esclusivamente attraverso matrici, circuiti quantistici e notazione di Dirac. Questo formalismo è completo ma presenta un limite: il circuito quantistico rappresenta bene una sequenza di operazioni, ma non sempre rende evidenti le identità algebriche nascoste dietro le trasformazioni quantistiche, questo rende quindi molto difficile verificare quando due circuiti sono equivalenti, molte identità richiedono pagine e pagine di calcoli matriciali. Da qui nasce l'idea: esiste un linguaggio in cui le identità quantistiche siano manipolazioni geometriche invece che manipolazioni matriciali?

Il primo risultato in risposta a questa esigenza arriva nel 2004 con il lavoro di Samson Abramsky e Bob Coecke, "*A Categorical Semantics of Quantum Protocols*" dove Abramsky e Coecke propongono di descrivere protocolli e processi quantistici in termini categoriali, usando strutture composizionali e diagrammatiche. L'idea di fondo è di partire dai processi che coinvolgono i qubit invece di partire dagli stati come vettori e dalle operazioni come matrici. Un processo ha ingressi e uscite, e può essere composto con altri processi in due modi: in sequenza oppure in parallelo. Questa doppia composizione è esattamente ciò che nei diagrammi si rappresenta con fili e nodi. Coecke chiama questo programma *quantum picturalism*: un linguaggio ad alto livello, diagrammatico, che

dovrebbe stare alla meccanica quantistica come i linguaggi di programmazione stanno al codice macchina. Nel suo articolo "*Quantum Picturalism*", Coecke sostiene infatti che il formalismo degli spazi di Hilbert, pur essendo corretto, spesso non aiuta l'intuizione e rende pesanti calcoli che possono diventare molto più trasparenti in forma grafica. Lo ZX-Calculus, oggetto di questa tesi, si inserisce in questo contesto proprio come soluzione alla nostra problematica.

La prima formulazione compare nel lavoro "*Interacting Quantum Observables*" del 2008, poi sviluppato in modo più completo nell'articolo "*Interacting Quantum Observables: Categorical Algebra and Diagrammatics*", pubblicato su New Journal of Physics nel 2011. In quel lavoro gli autori Bob Coecke e Ross Duncan presentano lo ZX-calculus come un calcolo grafico universale per sistemi multi-qubit e come formalizzazione diagrammatica dell'interazione tra osservabili quantistiche complementari, le matrici di Pauli Z ed X (da cui il nome).

Alla base dello ZX-Calculus inoltre emergono naturalmente strutture algebriche come le algebre di Frobenius e le algebre di Hopf e questo costituisce un ponte naturale tra il formalismo dei circuiti quantistici e strutture algebriche più astratte. Questo parallelismo è ciò che si è voluto investigare con il presente lavoro di tesi e nello specifico il testo è organizzato nel modo che segue:

- **Capitolo 1.** In questo capitolo ci proponiamo di introdurre le strutture di algebre di Hopf a partire da quelle di algebra e bialgebra, riportando le dimostrazioni dei risultati più significativi e introducendo due notazioni particolarmente utili alla manipolazione di queste strutture matematiche, la notazione di Sweedler e una notazione grafica che ci avvicinerà all'obiettivo successivo di riconoscere queste strutture all'interno dello ZX-Calculus;
- **Capitolo 2.** in questo capitolo ci proponiamo invece di fornire le basi del formalismo matematico in cui nascono i circuiti quantistici. Definiamo l'unità fondamentale dell'informatica quantistica a partire dai postulati della meccanica quantistica e le operazioni logiche fondamentali sui qubit passando dalla descrizione dei sistemi come vettori su uno spazio di Hilbert alla descrizione come matrici densità. Affrontiamo successivamente un confronto introduttivo tra gli oggetti dell'informatica classica e quelli dell'informatica quantistica evidenziando le differenze ed i vantaggi del passaggio dai bit ai qubit;
- **Capitolo 3.** in questo capitolo ci proponiamo infine di introdurre il formalismo con cui si costruiscono gli ZX-diagrams e successivamente presentare le regole dello ZX-Calculus per poi evidenziare l'emergere delle strutture algebriche di algebre di Frobenius e algebre di Hopf non normalizzate, di cui specifichiamo la forma in appendice.

Capitolo 1

Algebre di Hopf

In questo capitolo vogliamo presentare la nozione di algebre di Hopf; strutture algebriche, dotate di mappe aggiuntive, che ritroveremo in seguito nello sviluppo dello ZX-calculus. Per farlo partiremo da un esempio ben costruito su oggetti matematici che il lettore possa riconoscere come familiari, funzioni su elementi di un gruppo ordinario, e solo successivamente passeremo alla doverosa costruzione assiomatica; per la quale seguiremo il primo capitolo di [5]. Riporteremo le definizioni e proprietà già note di un'algebra, passeremo poi al concetto di coalgebra, e per questa introdurremo una notazione che ci verrà utile nelle dimostrazioni, detta notazione di Sweedler, e infine definiremo una struttura di bialgebra da cui subito arrivare a quella di algebra di Hopf, riscrivendola a sua volta in una seconda notazione grafica particolarmente intuitiva.

1.1 Esempio di algebra di Hopf : Funzioni sulle matrici del gruppo $SL(N, \mathbb{C})$

Partiamo per definire le algebre di Hopf da un esempio costruito sull'algebra delle funzioni definite sul gruppo di matrici $SL(N, \mathbb{C})$.

Il gruppo lineare speciale $SL(N, \mathbb{C})$ è definito come l'insieme delle matrici quadrate di ordine N a coefficienti complessi e determinante unitario, cioè

$$SL(N, \mathbb{C}) = \{A \in M_N(\mathbb{C}) \mid \det(A) = 1\}.$$

L'operazione di gruppo è il prodotto ordinario tra matrici. Poiché il determinante del prodotto è il prodotto dei determinanti,

$$\det(AB) = \det(A) \det(B),$$

il prodotto di due matrici appartenenti a $SL(N, \mathbb{C})$ appartiene ancora a $SL(N, \mathbb{C})$. Inoltre la matrice identità appartiene al gruppo e, per ogni $A \in SL(N, \mathbb{C})$, anche l'inversa A^{-1}

appartiene a $SL(N, \mathbb{C})$, essendo

$$\det(A^{-1}) = \frac{1}{\det(A)} = 1.$$

Ne segue che $SL(N, \mathbb{C})$ costituisce un sottogruppo del gruppo lineare generale $GL(N, \mathbb{C})$.

Indichiamo quindi con G il nostro gruppo e denotiamo con $F(G)$ l'algebra delle funzioni regolari a valori complessi sugli elementi del gruppo stesso.

$F(G)$ è un'algebra con somma e prodotto usuali punto per punto definite come :

$$(f + h)(g) = f(g) + h(g), \quad (f \cdot h)(g) = f(g)h(g), \quad (\lambda f)(g) = \lambda f(g),$$

per $f, h \in F(G)$, $g \in G$, $\lambda \in \mathbb{C}$. L'elemento neutro è I , definito da $I(g) = 1$, $\forall g \in G$.

Usando la struttura di gruppo sugli elementi $g \in G$ possiamo dotare l'algebra $F(G)$ di tre mappe lineari che chiamiamo rispettivamente comoltiplicazione Δ , counità ε e antipodo S :

- $\Delta : \mathcal{F}(G) \rightarrow \mathcal{F}(G \times G)$ definita come $(\Delta(f))(g_1, g_2) := f(g_1 g_2)$,
- $\varepsilon : \mathcal{F}(G) \rightarrow \mathbb{C}$ definita come $\varepsilon(f) := f(e)$,
- $S : \mathcal{F}(G) \rightarrow \mathcal{F}(G)$ definito come $(S(f))(g) := f(g^{-1})$.

Dove e nell'equazione della counità identifica l'unità per le matrici del gruppo $SL(N, \mathbb{C})$. Proviamo a visualizzare ciò che fanno queste tre mappe su un esempio concreto sul gruppo delle matrici di $SL(2, \mathbb{C})$. Queste ultime si scriveranno nella forma generale di :

$$M = \begin{pmatrix} a_{11} & a_{12} \\ a_{21} & a_{22} \end{pmatrix}$$

Possiamo pensare ad $F(G)$ come il funzionale che presa la matrice $\mathbb{U}$, come sopra costruita, ne estrae uno specifico elemento a_{ij} , come per esempio :

$$F_{11}(M) = F_{11} \begin{pmatrix} a_{11} & a_{12} \\ a_{21} & a_{22} \end{pmatrix} = a_{11}$$

che vediamo quindi come il funzionale che prende dalla matrice $\mathbb{U}$ l'ingresso in alto a sinistra. Come agisce la comoltiplicazione in questo nostro esempio ?

$$\begin{aligned} (\Delta(F_{11})) \left(\begin{pmatrix} a_{11} & a_{12} \\ a_{21} & a_{22} \end{pmatrix} \otimes \begin{pmatrix} b_{11} & b_{12} \\ b_{21} & b_{22} \end{pmatrix} \right) &= F_{11} \left(\begin{pmatrix} a_{11} & a_{12} \\ a_{21} & a_{22} \end{pmatrix} \begin{pmatrix} b_{11} & b_{12} \\ b_{21} & b_{22} \end{pmatrix} \right) \\ &= (a_{11} \otimes b_{11} + a_{12} \otimes b_{21}) \end{aligned}$$

La comoltiplicazione in questo caso estrae l'elemento prodotto della prima riga per la prima colonna delle due matrici su cui lavora l'output agendo cioè sul prodotto usuale definito sugli elementi del gruppo.

Come agisce la counità ? La counità $\varepsilon : \mathcal{F}(G) \rightarrow \mathbb{C}$ è definita dall'azione sul funzionale f come la valutazione del funzionale sull'elemento neutro delle matrici del gruppo G , per cui in questo nostro esempio, dove $f = F_{11}$:

$$\varepsilon(F_{11}) = F_{11}(I_{SL(2)}) = F_{11} \begin{pmatrix} 1 & 0 \\ 0 & 1 \end{pmatrix} = 1$$

Come agisce l'antipodo? In un'algebra di Hopf, l'antipodo S svolge un ruolo analogo a quello dell'inverso in un gruppo. Nel nostro caso esso agisce applicando il funzionale all'inverso della matrice data in ingresso. Consideriamo quindi una matrice

$$M = \begin{pmatrix} a & b \\ c & d \end{pmatrix} \in SL(2, \mathbb{C}),$$

poiché il determinante di M è uguale a 1, la sua matrice inversa è

$$U^{-1} = \begin{pmatrix} d & -b \\ -c & a \end{pmatrix}.$$

pertanto

$$(S(F_{11}))(U) = F_{11}(U^{-1}) = d.$$

A questo punto torniamo all'immagine più astratta di G ed $\mathcal{F}(G)$ come gruppo ordinario ed algebra delle funzioni sugli elementi del gruppo. Il solo sguardo alle proprietà che definiscono gli elementi $g \in G$ ci conduce a dimostrare la veridicità delle seguenti affermazioni per le tre mappe che abbiamo appena definito :

1. l'associatività degli elementi del gruppo implica che valga la seguente relazione :

$$(\Delta \otimes \text{id}) \circ \Delta = (\text{id} \otimes \Delta) \circ \Delta$$

dove con id abbiamo indicato la funzione identità sull'algebra delle funzioni $\mathcal{F}(G)$. Infatti ricordando la definizione data per la comoltiplicazione, abbiamo che :

Dimostrazione.

$$(((\Delta \otimes \text{id}) \circ \Delta)f)(g_1, g_2, g_3) = f((g_1 g_2) g_3) \in \mathcal{F}(G \times G \times G)$$

$$(((\text{id} \otimes \Delta) \circ \Delta)f)(g_1, g_2, g_3) = f(g_1 (g_2 g_3)) \in \mathcal{F}(G \times G \times G)$$

e basta quindi osservare l'uguaglianza tra i termini dati in input al funzionale f per dimostrare la validità del risultato :

$$(g_1 g_2) g_3 = g_1 (g_2 g_3) \implies f((g_1 g_2) g_3) = f(g_1 (g_2 g_3))$$

□

2. l'esistenza dell'elemento neutro per gli elementi del gruppo implica che valga la relazione :

$$(\varepsilon \otimes \text{id}) \circ \Delta = (\text{id} \otimes \varepsilon) \circ \Delta = \text{id}$$

infatti, possiamo svolgere il membro di destra e di sinistra dell'equazione, separatamente secondo le definizioni delle mappe di comoltiplicazione e counità, e troviamo:

Dimostrazione.

$$((\varepsilon \otimes \text{id}) \circ \Delta) f(g) = (\varepsilon \otimes \text{id})(f(g_1 g_2)) = f(\varepsilon(g_1) g_2) = f(g_2)$$

$$((\text{id} \otimes \varepsilon) \circ \Delta) f(g) = (\text{id} \otimes \varepsilon)(f(g_1 g_2)) = f(g_1 \varepsilon(g_2)) = f(g_1)$$

dove abbiamo usato la proprietà dell'elemento neutro nell'input al funzionale f

$$f(eg) = f(ge) = f(g)$$

□

per cui quello che troviamo è che questa particolare composizione di mappe risulta agire come la mappa identità, portandoci da un oggetto che vive nello spazio dei funzionali ad un oggetto analogo.

3. l'esistenza dell'elemento inverso per gli elementi del gruppo implica che valga la relazione:

$$m \circ (S \otimes \text{id}) \circ \Delta = \eta \circ \varepsilon = m \circ (\text{id} \otimes S) \circ \Delta$$

dove abbiamo definito altre due mappe addizionali

- $m : \mathcal{F}(G \times G) \rightarrow \mathcal{F}(G)$ tale che $(mh)(g) := h(g, g)$, $h \in \mathcal{F}(G \times G)$
- $\eta : \mathbb{C} \rightarrow \mathcal{F}(G)$ tale che $\eta(1)$ sia l'elemento unità dell'algebra $F(G)$

Infatti anche in questo caso possiamo verificare quanto detto svolgendo separatamente le due uguaglianze di destra e di sinistra della relazione, sostituendo le definizioni date per ciascuna delle mappe:

dimostrazione.

$$\begin{aligned} ((m \circ (S \otimes \text{id}) \circ \Delta)f)(g) &= (m \circ (S \otimes \text{id}))f(g_1, g_2) = (mf)(S(g_1)g_2) = \\ &= (mf)(g_1^{-1}g_2) = f(g^{-1}g) = f(e) = \eta \circ \varepsilon \end{aligned}$$

dove abbiamo utilizzato la relazione nota dell'elemento neutro sugli elementi del gruppo

$$g^{-1}g = gg^{-1} = e$$

□

Quindi gli assiomi che definiscono le proprietà del gruppo G si riflettono sull'algebra $F(G)$ nelle proprietà caratteristiche delle tre mappe aggiuntive di cui l'abbiamo dotata. Possiamo schematizzare questa corrispondenza come segue.

Gruppo G	Algebra $(\mathcal{F}(G), \Delta, \varepsilon, S)$
Proprietà Associativa: $(g_1g_2)g_3 = g_1(g_2g_3)$	$(\Delta \otimes \text{id}) \circ \Delta = (\text{id} \otimes \Delta) \circ \Delta$
Esistenza dell'elemento neutro: $eg = ge = g$	$(\varepsilon \otimes \text{id}) \circ \Delta = (\text{id} \otimes \varepsilon) \circ \Delta = \text{id}$
Esistenza dell'elemento inverso: $gg^{-1} = g^{-1}g = e$	$m \circ (S \otimes \text{id}) \circ \Delta = \eta \circ \varepsilon = m \circ (\text{id} \otimes S) \circ \Delta$

1.2 Definizioni di algebra, coalgebra e bialgebra

Richiamiamo innanzitutto la definizione di un'algebra associativa

Def. 1.2.1.

Preso uno spazio vettoriale A , sul campo $\mathbb{C}$, possiamo costruire un'algebra **associativa** equipaggiando lo spazio vettoriale con una mappa:

$$\begin{aligned} A \times A &\rightarrow A \\ (a, b) &\mapsto ab \end{aligned}$$

tale che per ogni $a, b, c \in A$ e $\lambda \in \mathbb{C}$:

1. $a(bc) = (ab)c$
2. $(a + b)c = ac + bc$
3. $a(b + c) = ab + ac$
4. $\lambda(ab) = (\lambda a)b = a(\lambda b)$

Un elemento 1 di un'algebra è detto *unità di A* se $1a = a1 = a, \forall a \in A$. Se esiste l'elemento neutro, questo è univocamente determinato.

Per algebra da questo momento in poi si intenderà sempre un'algebra associativa munita dell'elemento unità.

Nell'ottica di utilizzare una notazione più simile a quella delle algebre di Hopf possiamo anche riformulare questa definizione scrivendo l'algebra come uno spazio vettoriale A sul campo $\mathbb{C}$ dotato di due mappe lineari :

- la moltiplicazione: $m : A \otimes A \longrightarrow A$, tale che:

$$m \circ (m \otimes \text{id}) = m \circ (\text{id} \otimes m) \quad (1.1)$$

- l'unità: $\eta : \mathbb{C} \longrightarrow A$, tale che:

$$m \circ (\eta \otimes \text{id}) = \text{id} = m \circ (\text{id} \otimes \eta) \quad (1.2)$$

E' facile verificare che le due definizioni sono equivalenti, infatti l'eq (1) definisce la proprietà di associatività degli elementi dell'algebra mentre l'eq (2) definisce $\eta(1)$ come l'elemento unitario di A usando l'identificazione di $\mathbb{C} \otimes A$ e $A \otimes \mathbb{C}$ con A stesso, così come sarà per l'intero capitolo.

Possiamo visualizzare le proprietà che definiscono un'algebra anche attraverso una rappresentazione in diagrammi.

La relazione di associatività diventa per esempio:

$$\begin{array}{ccc}
A \otimes A \otimes A & \xrightarrow{m \otimes \text{id}} & A \otimes A \\
\downarrow \text{id} \otimes m & & \downarrow m \\
A \otimes A & \xrightarrow{m} & A
\end{array}$$

mentre la relazione di unità si traduce nella commutatività di questo secondo diagramma

$$\begin{array}{ccccc}
\mathbb{C} \otimes A & \xrightarrow{\eta \otimes \text{id}} & A \otimes A & \xleftarrow{\text{id} \otimes \eta} & A \otimes \mathbb{C} \\
\downarrow \text{id} & & \downarrow m & & \downarrow \text{id} \\
\mathbb{C} \otimes A & \iff & A & \iff & A \otimes \mathbb{C},
\end{array}$$

dove anche in questo caso abbiamo sottolineato l'identificazione canonica di $\mathbb{C} \otimes A$ e $A \otimes \mathbb{C}$ con A .

Def. 1.2.2.

Siano A e B due algebre. Una mappa lineare $\varphi : A \longrightarrow B$ è detta **omomorfismo di algebre** se:

1. $\varphi(aa') = \varphi(a)\varphi(a') \quad \forall a, a' \in A$
2. $\varphi(1_A) = 1_B$

Tali condizioni possono essere espresse anche come:

$$\varphi \circ m_A = m_B \circ (\varphi \otimes \varphi) \quad \text{e} \quad \varphi \circ \eta_A = \eta_B$$

Questa mappa deve quindi preservare la struttura di algebra (la moltiplicazione e l'unità) passando da A a B .

Def. 1.2.3.

Presi due spazi vettoriali A e B possiamo costruire una **algebra prodotto tensore** $A \otimes B$ il cui spazio vettoriale è il prodotto tensore dei due spazi vettoriali e la cui moltiplicazione è definita come

$$m_{A \otimes B} := (m_A \otimes m_B) \circ (id \otimes \tau \otimes id)$$

dove τ è l'operatore di flip per cui $\tau(v \otimes w) = w \otimes v$

La moltiplicazione dell'algebra prodotto tensore sarà quindi costruita come

$$(a \otimes b)(a' \otimes b') := aa' \otimes bb'$$

dove $a, a' \in A$ e $b, b' \in B$.

Def. 1.2.4.

Per ogni algebra A si può definire l'**algebra opposta** A^{op} . Un'algebra sullo stesso spazio vettoriale A dotata però di una mappa di moltiplicazione definita come :

$$m_{A^{op}} := m_A \circ \tau$$

Per cui per l'algebra opposta avremo che : $a \cdot_{op} b = b \cdot a$ dove $\cdot_{op}$ e $\cdot$ denotano le mappe di moltiplicazione rispettivamente dell'algebra opposta e dell'algebra di base.

Ora vogliamo costruire una struttura di coalgebra. Per farlo dualizziamo la definizione

di algebra nella sua forma in diagrammi commutativi. Dualizzare l'algebra in questa notazione infatti consiste semplicemente nell'invertire ogni freccia e sostituire ad ogni spazio il suo duale.

$$\begin{array}{ccccc}
 & & A \otimes A \otimes A & \xleftarrow{\Delta \otimes \text{id}} & A \otimes A \\
 & \uparrow \text{id} \otimes \Delta & & & \uparrow \Delta \\
 & A \otimes A & \xleftarrow{\Delta} & & A \\
 \mathbb{C} \otimes A & \xleftarrow{\varepsilon \otimes \text{id}} & A \otimes A & \xrightarrow{\text{id} \otimes \varepsilon} & A \otimes \mathbb{C} \\
 \uparrow \text{id} & & \uparrow \Delta & & \uparrow \text{id} \\
 \mathbb{C} \otimes A & \iff & A & \iff & A \otimes \mathbb{C} .
 \end{array}$$

Possiamo quindi scrivere una definizione assiomatica di algebra duale come:

Def. 1.2.5.

Una **coalgebra** è uno spazio vettoriale A sul campo $\mathbb{C}$, dotato di due mappe lineari:

$$\Delta : A \rightarrow A \otimes A$$

detta **comoltiplicazione** o coprodotto

$$\varepsilon : A \rightarrow \mathbb{C}$$

detta **counità**.

Tali che valgano le due seguenti proprietà:

$$(\Delta \otimes \text{id}) \circ \Delta = (\text{id} \otimes \Delta) \circ \Delta \quad \text{coassociatività} \quad (1.3)$$

$$(\varepsilon \otimes \text{id}) \circ \Delta = \text{id} = (\text{id} \otimes \varepsilon) \circ \Delta \quad \text{counità} \quad (1.4)$$

Le relazioni definite all'interno della coalgebra possono essere pensate come l'analogo dell'associatività e della proprietà di unità dell'algebra A .

La prima uguaglianza nell'eq.1.3 infatti ci sta dicendo che, preso un elemento $a \in A$, applicando la comoltiplicazione lo portiamo nello spazio prodotto tensore di $A \otimes A$, ma dopo questa operazione è equivalente applicare la comoltiplicazione sul primo elemento

del prodotto tensore e l'identità sul secondo ed applicare l'identità sul primo e la comoltiplicazione sul secondo, così come nello spazio dell'algebra era uguale applicare la moltiplicazione ai primi due elementi ed inseguito al terzo oppure agli ultimi due e in seguito al primo. Mentre la seconda uguaglianza nell'eq.1.4 ci sta dicendo che, la mappa della comoltiplicazione seguita dalla counità su uno dei due fattori, restituisce l'elemento originale, così come era con l'applicazione da destra o da sinistra per l'elemento unità nello spazio dell'algebra.

Procedendo alla stessa maniera, molti dei concetti visti per la struttura di algebra possono essere dualizzati per adattarsi alla coalgebra.

Abbiamo infatti :

Def. 1.2.6.

Una mappa lineare su $\mathbb{C}$, $\varphi : A \rightarrow B$ è un **omomorfismo di coalgebre** se

$$\Delta_B \circ \varphi = (\varphi \otimes \varphi) \circ \Delta_A \quad e \quad \varepsilon_A = \varepsilon_B \circ \varphi$$

Def. 1.2.7.

Una **coalgebra prodotto tensore** è una coalgebra costruita sullo spazio vettoriale $A \otimes B$ con mappe di comoltiplicazione e counità definite rispettivamente come :

$$\Delta_{A \otimes B} := (id \otimes \tau \otimes id) \circ (\Delta_A \otimes \Delta_B)$$

$$\varepsilon_{A \otimes B} := \varepsilon_A \otimes \varepsilon_B$$

Dove i due spazi vettoriali A e B sono a loro volta coalgebre.

Def. 1.2.8.

Una **coalgebra opposta** A^{op} è una coalgebra costruita sullo spazio vettoriale A equipaggiata delle mappa di comoltiplicazione e counità definite rispettivamente come :

$$\Delta_{A^{op}} := \tau \circ \Delta_A$$

$$\varepsilon_{A^{op}} := \varepsilon_A$$

Una coalgebra è detta **cocommutativa** se vale che :

$$\tau \circ \Delta = \Delta$$

Def. 1.2.10.

Un sottospazio lineare B di A è una **subcoalgebra** se

$$\Delta(B) \subseteq B \otimes B$$

Prima di passare alla definizione delle bialgebre, ci è utile introdurre una nuova notazione che renda maggiormente maneggevoli le operazioni che coinvolgono la comoltiplicazione. La notazione in questione prende il nome di **notazione di Sweedler**, secondo cui, preso un elemento a appartenente alla coalgebra A , l'applicazione della comoltiplicazione sull'elemento, $\Delta(a) \in A \otimes A$, si scrive come una somma finita :

$$\Delta(a) = \sum_i a_{(1)i} \otimes a_{(2)i}$$

con $a_{1i}, a_{2i} \in A$.

Per semplificare ulteriormente la notazione possiamo anche scrivere :

$$\Delta(a) = \sum a_{(1)} \otimes a_{(2)}$$

sottointendendo la somma sugli indici i .

Da questa notazione è facile definire la mappa induttiva :

$$\Delta^{(n)} : A \rightarrow A^{\otimes(n+1)}$$

come :

$$\Delta^{(n)} := (id \otimes \Delta) \circ \Delta^{(n-1)} \quad \text{per } n > 1 \quad e \quad \Delta^{(1)} = \Delta$$

Dalla coassociatività segue che $\Delta^{(n)}$ è di fatto equivalente a $n - 1$ composizioni di Δ , indipendentemente dal loro ordine, ovvero $\Delta^{(2)} = (id \otimes \Delta) \circ \Delta = (\Delta \otimes id) \circ \Delta$ e così via. Per cui l'elemento $\Delta^{(n)}(a) \in A^{\otimes(n+1)}$ in notazione di Sweedler si scriverà come :

$$\Delta^{(n)}(a) := \sum a_{(1)} \otimes a_{(2)} \otimes \dots \otimes a_{(n+1)}$$

La notazione di Sweedler in alcuni libri di testo potrebbe anche essere riportata in una forma ancora più asciutta, cioè come $\Delta(a) = \sum(a)$, che mette enfasi principalmente sull'elemento in input alla comoltiplicazione. Noi adotteremo maggiormente la versione della notazione che esclude gli indici i ma mantiene i pedici tensoriali. Al fine di prendere confidenza con la notazione potremmo pensare di riscrivere alcune delle proprietà delle coalgebre in notazione di Sweedler.

Coassociatività in notazione di Sweedler:

$$\sum (a_{(1)})_{(1)} \otimes (a_{(1)})_{(2)} \otimes a_{(2)} = \sum a_{(1)} \otimes (a_{(2)})_{(1)} \otimes (a_{(2)})_{(2)} = \Delta^{(2)}(a) = \sum a_{(1)} \otimes a_{(2)} \otimes a_{(3)} \quad (1.5)$$

Dimostrazione. riscriviamo quella che era la proposizione originale per la coassociatività:

$$(\Delta \otimes id) \circ \Delta = (id \otimes \Delta) \circ \Delta$$

Applichiamo la definizione della notazione di Sweedler sul ramo di sinistra:

$$\begin{aligned} ((\Delta \otimes id))(\sum a_{(1)} \otimes a_{(2)}) &= \sum \Delta(a_{(1)}) \otimes id(a_{(2)}) = \sum (\sum (a_{(1)})_1 \otimes (a_{(1)})_2) \otimes a_{(2)} = \\ &= \sum (a_{(1)})_{(1)} \otimes (a_{(1)})_{(2)} \otimes a_{(2)} \end{aligned}$$

ed abbiamo ottenuto il primo ramo di sinistra dell'eq. 1.5.

E' facile poi procedendo allo stesso modo dimostrare anche l'uguaglianza tra la proposizione originale e il secondo membro della proposizione in notazione di sweedler.

Che la proposizione originale sia uguale ad applicare due volte ricorsivamente la comoltiplicazione sull'elemento a è altrettanto facile da verificare ricordando che la notazione di sweedler semplifica gli indici della separazione della comoltiplicazione per privilegiare i pedici delle sole componenti del prodotto tensoriale così come indicato nell'ultimo membro dell'eq.1.5. $\square$

Counità in notazione di Sweedler:

$$\sum \epsilon(a_{(1)})a_{(2)} = a = \sum a_{(1)}\epsilon(a_{(2)})$$

Dimostrazione. riportiamo anche in questo caso la proposizione originale:

$$(\epsilon \otimes id) \circ \Delta = id = (id \otimes \epsilon) \circ \Delta$$

e svolgiamo i passaggi dell'uguaglianza originale utilizzando la notazione di sweedler:

$$((\varepsilon \otimes id)(\Delta(a))) = ((\varepsilon \otimes id))(\sum a_{(1)} \otimes a_{(2)}) = \sum \varepsilon(a_{(1)}) \otimes id(a_{(2)}) = \sum \varepsilon(a_{(1)})a_{(2)}$$

dove il prodotto tensore è stato "assorbito" dall'applicazione della mappa ε che manda l'elemento $a_{(1)}$ del prodotto tensore sul campo $\mathbb{K}$, spostando il prodotto sul campo degli scalari. □

Vogliamo ora estrarre da una coalgebra una struttura di algebra. Il modo più naturale per farlo è attraverso il prodotto di convoluzione.

Proposizione 1.2.1.

Siano A una coalgebra e B un'algebra. Lo spazio vettoriale $\mathcal{L}(A, B)$ di tutte le mappe $\mathbb{C}$ -lineari da A a B , dotato del **prodotto di convoluzione** definito come:

$$(f * g)(a) := (m_B \circ (f \otimes g) \circ \Delta_A)(a) \equiv \sum f(a_{(1)})g(a_{(2)})$$

con $a \in A$ diventa un'algebra con unità $\eta_B \circ \varepsilon_A$.

Dimostrazione. Dalla associatività della moltiplicazione in B e dalla coassociatività della comoltiplicazione in A , otteniamo correttamente a sua volta la proprietà di associatività per il prodotto di convoluzione come :

$$((f * g) * h)(a) = \sum f(a_{(1)*})g(a_{(2)})h(a_{(3)}) = (f * (g * h))(a)$$

dalla definizione di prodotto di convoluzione infatti sappiamo che :

$$(f * g)(a) = \sum f(a_{(1)})g(a_{(2)})$$

allora svolgendo il membro di sinistra della relazione di associatività per il prodotto $*$ troviamo:

$$((f * g) * h)(a) = \sum (f * g)(a_{(1)})h(a_{(2)}) = \sum f(a_{(1)(1)})g(a_{(1)(2)})h(a_{(2)})$$

e allo stesso modo svolgendo il membro di destra :

$$(f * (g * h))(a) = \sum f(a_{(1)})(g * h)(a_{(2)}) = \sum f(a_{(1)})g(a_{(2)(1)})h(a_{(2)(2)})$$

e ricordando la proprietà di coassociatività :

$$\Delta(a_{(1)}) = a_{(1)(1)} \otimes a_{(1)(2)} \quad e \Delta(a_{(2)}) = a_{(2)(1)} \otimes a_{(2)(2)}$$

abbiamo che :

$$a_{(1)(1)} \otimes a_{(1)(2)} \otimes a_{(2)} = a_{(1)} \otimes a_{(2)(1)} \otimes a_{(2)(2)}$$

e con questo abbiamo dimostrato l'associatività del prodotto di convoluzione.

Dalla relazione di counità in notazione di Sweedler è poi immediato dimostrare la relazione

$$((\eta_B \circ \varepsilon_A) * f)(a) = \sum \varepsilon_A(a_{(1)})f(a_{(2)}) = f\left(\sum \varepsilon_A(a_{(1)})a_{(2)}\right) = f(a)$$

che ci dice $\eta_B \circ \varepsilon_A$ è un'elemento unità che agisce da sinistra. Parimenti, è facile verificare che $\eta_B \circ \varepsilon_A$ è un'elemento unità anche agisce da destra, per cui è correttamente definito per essere l'elemento unità dell'algebra indotta col prodotto di convoluzione sulla coalgebra. $\square$

Corollario 1.2.1.

Lo spazio vettoriale duale A' di una coalgebra definita sullo spazio vettoriale A ha struttura di algebra con la mappa di moltiplicazione definita come :

$$(fg)(a) := (f * g)(a) = \sum f(a_{(1)})g(a_{(2)}) \quad a \in A \quad f, g \in A'$$

Lo spazio duale B' dell'algebra B in generale invece non diventa automaticamente una coalgebra dualizzando la moltiplicazione, ma lo diventa nel caso in cui l'algebra B sia di dimensione finita. La ragione è che in generale si verifica sempre che $\mathcal{A}' \otimes \mathcal{A}' \subseteq (\mathcal{A} \otimes \mathcal{A})'$ e quindi $m_{\mathcal{A}'} : \mathcal{A}' \otimes \mathcal{A}' \rightarrow \mathcal{A}'$, mentre $\mathcal{B}' \otimes \mathcal{B}'$ è solo un sottospazio proprio di $(\mathcal{B} \otimes \mathcal{B})'$ se B è di dimensione infinita, cosicché $\Delta_{B'}(f) := f \circ m_B$ potrebbe non finire in $\mathcal{B}' \otimes \mathcal{B}'$. Abbiamo solo due modi per aggirare questo problema: prendere dello spazio duale B' solo un sottospazio oppure allargare il prodotto tensore.

Siamo ora in grado di introdurre la nozione di bialgebra, che introduce sostanzialmente la struttura che ci serve per costruire un'algebra di Hopf, per la quale infatti basterà aggiungere semplicemente una terza mappa, l'antipodo.

Una bialgebra è un'algebra e una coalgebra dove entrambe le strutture sono compatibili nel senso che segue:

Proposizione 1.2.3.

Se A è uno spazio vettoriale che è un'algebra e una coalgebra simultaneamente, allora le seguenti affermazioni sono equivalenti :

1. $\Delta : A \rightarrow A \otimes A$ e $\varepsilon : A \rightarrow \mathbb{C}$ sono omomorfismi di coalgrebre
2. $m : A \otimes A \rightarrow A$ e $\eta : \mathbb{C} \rightarrow A$ sono omomorfismi di algrebre

Dimostrazione. In base alle definizioni precedenti, le equazioni della prima asserzione ci dicono che

$$\Delta \circ m = m_{A \otimes A^\circ} \circ (\Delta \otimes \Delta), \quad \Delta \circ \eta = \eta_{A \otimes A}, \quad \varepsilon \circ m = m_C \circ (\varepsilon \otimes \varepsilon), \quad \varepsilon \circ \eta = \eta_C,$$

mentre le equazioni della seconda possono essere invece espresse come

$$\Delta \circ m = (m \otimes m) \circ \Delta_{A \otimes A}, \quad \varepsilon_{A \otimes A} = \varepsilon \circ m, \quad \Delta \circ \eta = (\eta \otimes \eta) \circ \Delta_C, \quad \varepsilon_C = \varepsilon \circ \eta.$$

I membri di destra delle prime relazioni di (1) e (2) sono entrambi uguali a

$$(m \otimes m) \circ (\text{id} \otimes \tau \otimes \text{id}) \circ (\Delta \otimes \Delta).$$

La seconda uguaglianza di (1) e la terza di (2) sono equivalenti. Le ultime due equazioni significano entrambe che $\varepsilon(1) = 1$. Analogamente, la terza relazione di (1) e la seconda di (2) sono equivalenti. $\square$

Def. 1.2.12.

Una **bialgebra** è uno spazio vettoriale dotato di una struttura algebrica e coalgebrica compatibili secondo la proposizione 1.2.3.

Quindi se A è un'algebra e una coalgebra, sarà una bialgebra se e solo se $\forall a, b \in A$ abbiamo che

$$\Delta(ab) = \Delta(a)\Delta(b) \quad \varepsilon(ab) = \varepsilon(a)\varepsilon(b) \quad \Delta(1) = 1 \otimes 1 \quad \text{e} \quad \varepsilon(1) = 1$$

Prendendo in esame per esempio la prima equazione delle definizioni per una bialgebra, quello che ci dice è che nel caso in cui gli elementi a e b vivano in una bialgebra, allora "spezzare il prodotto" è equivalente a fare il "prodotto dei pezzi spezzati" :

$$(a_{(1)} \otimes a_{(2)})(b_{(1)} \otimes b_{(2)}) \equiv a_{(1)}b_{(1)} \otimes a_{(2)}b_{(2)}.$$

Ora siano A e B due bialgre. Cosa significa essere un omomorfismo di bialgre ? Per omomorfismo di bialgre intendiamo una mappa da A a B che sia simultaneamente un omomorfismo di algre e un omomorfismo di coalgre.

Lo spazio vettoriale $A \otimes B$ dotato dell'algebra prodotto tensore e la struttura di coalgebra diventa una bialgebra.

Ci sono poi tre strutture di bialgre che possiamo estrarre da A prendendo l'opposto della sua struttura di algebra, di coalgebra o di entrambe : A^{op} , A^{cop} , $A^{op,cop}$:

- A^{op} avrà la moltiplicazione opposta $m_{A^{op}}$ e la comoltiplicazione di A

- A^{cop} avrà la moltiplicazione di A e la comoltiplicazione opposta $\Delta_{A^{cop}}$
- $A^{op,cop}$ prenderà sia la moltiplicazione opposta $m_{A^{op}}$ che la comoltiplicazione opposta $\Delta_{A^{cop}}$

Ora introduciamo due elementi distinguibili di uno spazio vettoriale che sia una bialgebra A .

Def. 1.2.13.

Un elemento non nullo $g \in A$ è detto **di tipo gruppo** se $\Delta(g) = g \otimes g$.

Un elemento $x \in A$ è detto **elemento primitivo** se $\Delta(x) = x \otimes 1 + 1 \otimes x$.

E conseguentemente le loro proprietà.

Proposizione 1.2.4.

Sia A una bialgebra. Allora :

1. il prodotto di elementi di tipo gruppo è ancora un elemento di tipo gruppo;
2. se x e y sono elementi primitivi di A allora abbiamo che :

$$\varepsilon(x) = \varepsilon(y) = 0 \quad \text{e} \quad [x, y] := xy - yx \quad \text{è ancora un elemento primitivo}$$

Dimostrazione. La prima asserzione è ovvia. Se g e h sono elementi di tipo gruppo sarà che $\Delta(g) = g \otimes g$ e $\Delta(h) = h \otimes h$ allora $\Delta(gh) = \Delta(g)\Delta(h)$ perchè lavoriamo in una bialgebra, e quindi $\Delta(gh) = gh \otimes gh$, che è ancora un elemento di tipo gruppo.

Per verificare la seconda asserzione invece prendiamo due elementi x e y primitivi.

Allora, siccome $(id \otimes \varepsilon)(\Delta(x)) = x\varepsilon(1) + \varepsilon(x)1$ e $\varepsilon(1) = 1$ per definizione di bialgebra, otteniamo $\varepsilon(x) = 0$.

Sfruttando poi il fatto che Δ è un omomorfismo di algebre otteniamo :

$$\begin{aligned} \Delta(xy) &= (x \otimes 1 + 1 \otimes x)(y \otimes 1 + 1 \otimes y) = xy \otimes 1 + x \otimes y + y \otimes x + 1 \otimes xy \\ \Delta(yx) &= (y \otimes 1 + 1 \otimes y)(x \otimes 1 + 1 \otimes x) = yx \otimes 1 + y \otimes x + x \otimes y + 1 \otimes yx \\ \implies \Delta([x, y]) &= \Delta(xy) - \Delta(yx) = [x, y] \otimes 1 + 1 \otimes [x, y] \end{aligned}$$

□

Dalla Proposizione 1.2.4 quindi capiamo che, il set di elementi di A di tipo gruppo forma un semi gruppo dotato dell'elemento unità e che, gli elementi primitivi di A formano un'algebra di Lie rispetto al commutatore $[\cdot, \cdot]$.

1.3 Algebre di Hopf

Abbiamo ora tutti gli strumenti necessari per costruire la struttura di algebra di Hopf.

Def. 1.3.1.

Una bialgebra A è detta **algebra di Hopf** se esiste una mappa lineare $S : A \rightarrow A$, che prende il nome di **antipodo**, o coinverso, tale che

$$m \circ (S \otimes \text{id}) \circ \Delta = \eta \circ \varepsilon = m \circ (\text{id} \otimes S) \circ \Delta \quad (1.6)$$

L'equazione 1.6 si può tradurre ugualmente nella commutatività del seguente diagramma

$$\begin{array}{ccccc}
 A \otimes A & \xleftarrow{\Delta} & A & \xrightarrow{\Delta} & A \otimes A \\
 \downarrow \text{id} \otimes S & & \downarrow \eta \circ \varepsilon & & \downarrow S \otimes \text{id} \\
 A \otimes A & \xrightarrow{m} & A & \xleftarrow{m} & A \otimes A
 \end{array}$$

mentre in notazione di Sweedler la scriviamo come

$$\sum S(a_1)a_2 = \varepsilon(a)1 = \sum a_1S(a_2)$$

Inoltre l'equazione 1.6 può essere interpretata in modo molto chiaro usando il prodotto di convoluzione dell'algebra $\mathcal{L}(A, A)$, definito come $(f * g)(a) := (m_B \circ (f \otimes g) \circ \Delta_A)(a)$. Confrontando le formule infatti vediamo che l'eq.1.6 altro non ci dice se non che S è l'inverso per convoluzione dell'applicazione identità; cioè

$$S * \text{id} = \text{id} * S = \eta \circ \varepsilon.$$

Dunque, una bialgebra A è un'algebra di Hopf se e solo se la mappa identità di A è invertibile nell'algebra $\mathcal{L}(A, A)$. Questa caratterizzazione implica quindi che l'antipodo di un'algebra di Hopf sia determinato univocamente.

Proposizione 1.3.1.

L'antipodo S di un'algebra di Hopf A è un **anti-omomorfismo di algebre** e un **anti-omomorfismo di coalgebre**.

Questo significa che

$$S : \mathcal{A} \rightarrow \mathcal{A}^{op} \quad \text{è un omomorfismo di algebre}$$

$$\text{e } S : \mathcal{A} \rightarrow \mathcal{A}^{cop} \quad \text{è un omomorfismo di coalgebre.}$$

Per cui abbiamo :

$$S(ab) = S(a)S(b) \quad a, b \in A \quad \text{e} \quad S(1) = 1 \quad (1.7)$$

$$\Delta \circ S = \tau \circ (S \otimes id) \circ \Delta \quad \text{e} \quad \varepsilon \circ S = \varepsilon \quad (1.8)$$

Dimostrazione. Usando gli assiomi di algebra di Hopf e il fatto che Δ e ε siano omomorfismi di coalgebre possiamo dimostrare l'eq.1.7 come segue :

$$\begin{aligned} S(b)S(a) &\stackrel{\text{eq. (10)}}{=} \sum S(b_{(1)}\varepsilon(b_{(2)})) S(a_{(1)}\varepsilon(a_{(2)})) \\ &= \sum S(b_{(1)})S(a_{(1)})\varepsilon(a_{(2)}b_{(2)}) \\ &\stackrel{\text{eq. (13)}}{=} \sum S(b_{(1)})S(a_{(1)}) (a_{(2)}b_{(2)})_{(1)} S((a_{(2)}b_{(2)})_{(2)}) \\ &= \sum S(b_{(1)})(\varepsilon(a_{(1)})1)b_{(2)} S(a_{(2)}b_{(3)}) \\ &= \sum \varepsilon(a_{(1)})\varepsilon(b_{(1)}) S(a_{(2)}b_{(2)}) \\ &= S(ab). \end{aligned}$$

Questo prova la prima uguaglianza dell'eq.1.7. La seconda segue immediatamente da $m \circ (S \otimes id) \circ \Delta(1) = \varepsilon(1)1$ combinata con $\Delta(1) = 1 \otimes 1$ e $\varepsilon(1) = 1$.

Similmente per l'eq.1.8 svolgiamo i calcoli:

$$\begin{aligned}
\sum S(a_{(2)}) \otimes S(a_{(1)}) &= \sum S(a_{(2)}) \varepsilon(a_{(3)}) \otimes S(a_{(1)}) \\
&= \sum (S(a_{(2)}) \otimes S(a_{(1)})) (\varepsilon(a_{(3)}) 1 \otimes 1) \\
&= \sum (S(a_{(2)}) \otimes S(a_{(1)})) (\Delta(a_{(3)}) S(a_{(4)})) \\
&= \sum (S(a_{(2)}) \otimes S(a_{(1)})) (a_{(3)} \otimes a_{(4)}) \Delta(S(a_{(5)})) \\
&= \sum (S(a_{(2)}) a_{(3)} \otimes S(a_{(1)}) a_{(4)}) \Delta(S(a_{(5)})) \\
&= \sum (\varepsilon(a_{(2)}) 1 \otimes S(a_{(1)}) a_{(3)}) \Delta(S(a_{(4)})) \\
&= \sum (1 \otimes S(a_{(1)}) a_{(2)}) \Delta(S(a_{(3)})) \\
&= \sum (1 \otimes \varepsilon(a_{(1)}) 1) \Delta(S(a_{(2)})) = \Delta(S(a)),
\end{aligned}$$

che dimostrano la prima uguaglianza dell'equazione, mentre la seconda si ottiene subito da

$$\varepsilon(S(a)) = \varepsilon\left(S\left(\sum a_{(1)} \varepsilon(a_{(2)})\right)\right) = \varepsilon\left(\sum S(a_{(1)}) a_{(2)}\right) = \varepsilon(\varepsilon(a) 1) = \varepsilon(a).$$

□

In notazione di Sweedler, la prima uguaglianza della 1.8 può essere riscritta come

$$\sum S(a)_{(1)} \otimes S(a)_{(2)} = \sum S(a)_{(2)} \otimes S(a)_{(1)}.$$

Sappiamo come ricavare delle bialgebre con le strutture di A^{op} , A^{cop} e $A^{op,cop}$ prendendo la moltiplicazione opposta, la comoltiplicazione opposta o entrambe. A questo punto ci chiediamo quale sia il completamento esatto perchè queste tre bialgebre possano dirsi essere delle algebre di Hopf. La relazione 1.6 ci dice che l'antipodo mette in relazione le mappe di moltiplicazione e comoltiplicazione, per questo è facile verificare che la bialgebra $A^{op,cop}$ munita sia di moltiplicazione che comoltiplicazione opposta rispetta anch'essa la corretta definizione di antipodo con la stessa mappa S definita per l'algebra di Hopf A . Affinchè le bialgebre A^{op} e A^{cop} siano delle algebre di Hopf invece deve verificarsi che l'antipodo S definito per A sia una mappa invertibile, così che S^{-1} possa completare correttamente le due bialgebre perchè anche loro rispettino la relazione 1.6.

Formalizziamo questo risultato:

Proposizione 1.3.2.

Per ogni algebra di Hopf A , le seguenti condizioni sono equivalenti:

1. L'antipodo S di A è invertibile come applicazione lineare di A .
2. La bialgebra A^{op} è un'algebra di Hopf.
3. La bialgebra A^{cop} è un'algebra di Hopf.

In tal caso, l'inverso S^{-1} di S è l'antipodo di A^{op} e di A^{cop} .

Dimostrazione. (1) $\Rightarrow$ (2): L'inverso S^{-1} di S è un anti-omomorfismo di algebre per l'algebra di Hopf A , come sappiamo dalla Proposizione 1.3.1. Quindi dall'eq.1.7 abbiamo che :

$$\sum S^{-1}(a_{(2)})a_{(1)} = \sum S^{-1}(a_{(2)})(S^{-1} \circ S)(a_{(1)}) = S^{-1}\left(\sum S(a_{(1)})a_{(2)}\right) = S^{-1}(\varepsilon(a)1) = \varepsilon(a)1$$

e similmente si vede che

$$\sum a_{(2)}S^{-1}(a_{(1)}) = \varepsilon(a)1 \quad \text{per ogni } a \in A.$$

Cioè, S^{-1} soddisfa la condizione 1.6 per la bialgebra A^{op} , dunque è un antipodo di A^{op} .

(2) $\Rightarrow$ (1): Sia $\hat{S}$ l'antipodo di A^{op} . Poiché $\hat{S}$ è anch'esso un anti-omomorfismo di algebra di A , otteniamo

$$\begin{aligned} \hat{S}S(a) &= \hat{S}S\left(\sum \varepsilon(a_{(2)})a_{(1)}\right) = \sum \varepsilon(a_{(2)})\hat{S}S(a_{(1)}) \\ &= \sum (a_{(3)}\hat{S}(a_{(2)}))\hat{S}S(a_{(1)}) = \sum a_{(3)}\hat{S}(S(a_{(1)})a_{(2)}) \\ &= \sum a_{(2)}\hat{S}(\varepsilon(a_{(1)})1) = a\hat{S}(1) = a \end{aligned}$$

e analogamente $S\hat{S}(a) = a$ per ogni $a \in A$. Dunque $\hat{S}$ è l'inverso di S .

L'equivalenza tra (1) e (3) si dimostra in modo analogo.

□

Da questo segue:

Corollario 1.3.1.

Se un'algebra di Hopf A è commutativa oppure cocommutativa, allora $S^2 = \text{id}$.

Dimostrazione. Se la nostra algebra di Hopf A è commutativa oppure cocommutativa significa che coincide con A^{op} oppure A^{cop} e questo vuol dire che l'antipodo S coincide con il suo inverso $S^{-1} \implies S^2 = id$

□

Ricordiamo che una coalgebra si dice *cocommutativa* se il coprodotto è invariante rispetto allo scambio dei due fattori tensoriali, cioè se, in notazione di Sweedler, vale la proprietà

$$\sum a_{(1)} \otimes a_{(2)} = \sum a_{(2)} \otimes a_{(1)},$$

per ogni $a \in A$. La cocommutatività rappresenta quindi la proprietà duale della commutatività del prodotto: mentre in un'algebra commutativa l'ordine dei fattori nella moltiplicazione è irrilevante, in una coalgebra cocommutativa è irrilevante l'ordine dei due termini prodotti dal coprodotto.

Inoltre l'antipodo S di un'algebra di Hopf è detto di *ordine finito* se $S^n = id$ per un qualche $n \in \mathbb{N}$. Il più piccolo n è detto *ordine di S* .

Concludiamo questa sottosezione osservando che alcune volte potrebbe risultare più facile verificare le proprietà di un'algebra di Hopf su un set di generatori, per cui potrebbe essere utile il seguente risultato.

Corollario 1.3.2.

Sia A_g un sottoinsieme di un'algebra A che genera A come algebra. Siano

$$\Delta : A \rightarrow A \otimes A \quad \text{e} \quad \varepsilon : A \rightarrow K$$

omomorfismi, e sia

$$S : A \rightarrow A$$

un anti-omomorfismo delle algebre corrispondenti. Se la condizione di coassociatività 1.3 e la condizione di counità 1.4 (e la condizione di antipodo 1.6) sono soddisfatte per gli elementi di A_g , allora esse valgono su tutto A e dunque A è una bialgebra (rispettivamente un'algebra di Hopf).

1.3.1 Coppie duali di algebre di Hopf

Consideriamo ora una algebra di Hopf finito dimensionale A . Il corollario 1.2.1 ci aveva detto che lo spazio vettoriale duale A' sulla struttura di colagebra, diventa un'algebra con la moltiplicazione definita dal prodotto di convoluzione come

$$fg(a) = (f \otimes g)\Delta(a) = \sum f(a_{(1)})g(a_{(2)}).$$

Per $f \in A'$ definiamo un funzionale $\Delta(f) \in (A \otimes A)'$ come

$$\Delta(f)(a \otimes b) := (f \circ m)(a \otimes b) = f(ab) \quad a, b \in A$$

Siccome A è uno spazio vettoriale finito dimensionale

$$(A \otimes A)' \equiv A' \otimes A' \implies \Delta(f) \in A' \otimes A'$$

Non è difficile far vedere che l'algebra A' dotata di questa comoltiplicazione diventa un'algebra di Hopf. L'antipodo, la counità e l'elemento unità saranno definiti come

$$(Sf)(a) = f(s(a)) \quad \varepsilon_{A'}(f) = f(1) \quad \text{e} \quad 1_{A'}(a) = \varepsilon(a).$$

In altre parole, l'algebra di Hopf A' si ottiene dualizzando le mappe della struttura di algebra di Hopf sullo spazio A . Questa situazione si generalizza nel modo seguente.

Def. 1.3.2

Un accoppiamento duale di due bialgebre U e A è una mappa bilineare

$$\langle \cdot, \cdot \rangle : U \times A \rightarrow \mathbb{C}$$

tale che

$$\langle \Delta_U(f), a_1 \otimes a_2 \rangle = \langle f, a_1 a_2 \rangle, \quad \langle f_1 f_2, a \rangle = \langle f_1 \otimes f_2, \Delta_A(a) \rangle, \quad (1.9)$$

$$\langle f, 1_A \rangle = \varepsilon_U(f), \quad \langle 1_U, a \rangle = \varepsilon_A(a), \quad (1.10)$$

per ogni $f, f_1, f_2 \in U$ e $a, a_1, a_2 \in A$. Un accoppiamento duale $\langle \cdot, \cdot \rangle$ tra U e A si dice non degenerare se $\langle f, a \rangle = 0$ per ogni $f \in U$ implica $a = 0$, e se $\langle f, a \rangle = 0$ per ogni $a \in A$ implica $f = 0$.

Proposizione 1.3.3

Se $\langle \cdot, \cdot \rangle$ è un accoppiamento duale tra due algebre di Hopf U e A , allora vale

$$\langle S(f), a \rangle = \langle f, S(a) \rangle, \quad f \in U, \quad a \in A. \quad (1.11)$$

Dimostrazione. Consideriamo le funzioni lineari F , G e H sul prodotto tensoriale dell'algebra di Hopf $U \otimes A$, definite da

$$F(f \otimes a) = \langle f, a \rangle, \quad G(f \otimes a) = \langle S(f), a \rangle, \quad H(f \otimes a) = \langle f, S(a) \rangle,$$

con $f \in U$ e $a \in A$. Mostriamo che $FG = HF = \varepsilon$ nell'algebra $(U \otimes A)'$. Infatti, per 1.9 e 1.10, otteniamo

$$\begin{aligned} FG(f \otimes a) &= \sum \langle f_{(1)}, a_{(1)} \rangle \langle S(f_{(2)}), a_{(2)} \rangle = \sum \langle f_{(1)} S(f_{(2)}), a \rangle \\ &= \langle \varepsilon(f) 1, a \rangle = \varepsilon(f \otimes a), \end{aligned}$$

cioè $FG = \varepsilon$, e analogamente $HF = \varepsilon$. Dunque

$$G = (HF)G = H(FG) = H,$$

il che implica 1.11. □

Dalla discussione precedente alla definizione 1.3.2 abbiamo già mostrato che per ogni algebra di Hopf finito dimensione esiste un accoppiamento duale definito da $\langle f, a \rangle := f(a)$, con $a \in A$, $f \in A'$, che risulta essere un accoppiamento non degenere, quindi che preserva l'intera informazione per entrambe le strutture.

In generale però un accoppiamento di questo tipo non è unico. Per ogni coppia di algebre di Hopf U e A esiste sempre l'accoppiamento banale definito da $\langle f, a \rangle := \varepsilon_U(f) \varepsilon_A(a)$, ed inoltre anche le composizioni di omomorfismi di algebre di Hopf ci restituiscono un ulteriore metodo di costruzione di un corretto accoppiamento.

Dalle formule dell'asserzione 1.9 si vede che ogni accoppiamento duale tra due bialgebre è completamente determinato dai valori della forma bilineare $\langle \cdot, \cdot \rangle$ sui sistemi di generatori delle algebre U e A . Questo ci consente di descrivere in modo sintetico tali accoppiamenti anche mediante matrici.

Un accoppiamento duale non degenere tra due algebre di Hopf U e A racchiude una notevole quantità di informazione sul legame tra le due strutture. Ad esempio, dalla relazione 1.9 segue che U è commutativa se e solo se A è cocommutativa. Nella teoria dei gruppi quantistici si afferma, in generale, che le algebre di Hopf U e A forniscono approcci duali a uno stesso "oggetto quantistico".

Quello che abbiamo costruito quindi con l'accoppiamento duale è una nozione di compatibilità tra due bialgebre che rispetti le mappe di prodotto, coprodotto, unità e counità anche nel caso in cui le due bialgebre siano infinito dimensionali; senza richiedere però strettamente che l'accoppiamento passi dal duale lineare dello spazio vettoriale da cui si parte per derivare la seconda struttura bialgebrica (ed in seguito di Hopf). Il passaggio per lo spazio vettoriale duale dello spazio originale, infatti, non garantisce in generale che da un'algebra infinito dimensionale si possa estrarre una struttura bialgebrica.

Tuttavia, si verifica che A' contiene una massima coalgebra A° e che A° è un'algebra di Hopf quando A lo è.

Supponiamo infatti che A sia un'algebra. Come al solito, consideriamo $A' \otimes A'$ come un sottospazio lineare di $(A \otimes A)'$ identificando $f \otimes g \in A' \otimes A'$ con la funzione lineare su

$A \otimes A$ definita da $(f \otimes g)(a \otimes b) := f(a)g(b)$. Per $f \in A'$ sia $\Delta(f)$ l'elemento di $(A \otimes A)'$ definito da

$$\Delta(f)(a \otimes b) := f(ab), \quad a, b \in A.$$

Sia ora A° l'insieme di tutte le funzionali $f \in A'$ per le quali $\Delta(f)$ appartiene ad $A' \otimes A'$, cioè tali che esistono funzionali $f_1, \dots, f_r, g_1, \dots, g_r \in A'$, con $r \in \mathbb{N}$, tali che

$$f(ab) = \sum f_i(a)g_i(b)$$

per tutti $a, b \in A$. In questo caso, l'elemento $\Delta(f) \in A' \otimes A'$ ha la forma

$$\Delta(f) = \sum f_i \otimes g_i.$$

allora segue che

Proposizione 1.3.4

1. Se A è un'algebra, A° è una coalgebra con comoltiplicazione Δ definita sopra e counità ε data da $\varepsilon(f) = f(1)$.
2. Se A è una bialgebra, allora la coalgebra A° dotata del prodotto per convoluzione (11) è anch'essa una bialgebra. L'unità dell'algebra A° è allora la counità ε della coalgebra A .
3. Se A è un'algebra di Hopf, allora anche A° lo è, con antipodo dato da

$$S(f)(a) = f(S(a)).$$

4. Se A è un'algebra di Hopf $*$, allora l'algebra di Hopf A° dotata dell'involuzione (40) è anch'essa un'algebra di Hopf $*$.

Dimostrazione. 1. Anzitutto dobbiamo assicurarci che $\Delta(A^\circ) \subseteq A^\circ \otimes A^\circ$. Sia $f \in A^\circ$ e

$$\Delta(f) = \sum f_i \otimes g_i,$$

dove le funzionali $\{f_i\}$ sono scelte linearmente indipendenti. Allora possiamo trovare elementi $a_j \in A$ tali che $f_i(a_j) = \delta_{ij}1$. Abbiamo

$$g_j(ab) = \sum_i f_i(a_j)g_i(ab) = f(a_j ab) = \sum_i f_i(a_j a)g_i(b),$$

il che implica che $g_j \in A^\circ$. Analogamente, $f_j \in A^\circ$, e quindi

$$\Delta(f) \in A^\circ \otimes A^\circ.$$

Chiaramente, A° è un sottospazio lineare di A' . La coassociatività di Δ si deduce dall'associatività della moltiplicazione di A :

$$((\Delta \otimes \text{id}) \circ \Delta(f))(a \otimes b \otimes c) = f((ab)c) = f(a(bc)) = ((\text{id} \otimes \Delta) \circ \Delta(f))(a \otimes b \otimes c).$$

L'assioma della counità è ovvio.

2. Vogliamo mostrare che A° è un sott'algebra di A' . Siano $f, g \in A^\circ$ con

$$\Delta(f) = \sum f_i \otimes g_i, \quad \Delta(g) = \sum h_j \otimes k_j.$$

Allora si calcola

$$fg(ab) = \sum_{i,j} f_i h_j(a) g_i k_j(b),$$

per cui $fg \in A^\circ$. La condizione di compatibilità tra comoltiplicazione e moltiplicazione di A° si deduce facilmente da quella di A .

3. Sia

$$\Delta(f) = \sum f_i \otimes g_i$$

per $f \in A^\circ$. Allora

$$S(f)(ab) = \sum_i S(g_i)(a) S(f_i)(b),$$

e quindi $S(f) \in A^\circ$. Verifichiamo una relazione dell'assioma dell'antipodo (l'altra segue in modo analogo):

$$\begin{aligned} (m \circ (S \otimes \text{id}) \circ \Delta(f))(a) &= \sum (S(f_i) g_i)(a) = \sum_{i,(a)} f_i(S(a_{(1)})) g_i(a_{(2)}) \\ &= \sum_i f(S(a_{(1)}) a_{(2)}) = \varepsilon(a) f(1) = 1_{A^\circ} \varepsilon_{A^\circ}(f) = ((\eta \circ \varepsilon)(f))(a). \end{aligned}$$

4. Il punto (4) si verifica facilmente usando il Corollario 1.3.3. □

Dove per *algebra di Hopf ** intendiamo l'algebra di Hopf costruita sulla *-struttura come segue.

Def. 1.3.3

Uno $*$ -spazio vettoriale è uno spazio vettoriale V sul campo $\mathbb{C}$ insieme a una mappa $a \mapsto a^*$, chiamata involuzione e denotata da $*$, tale che

$$(\alpha v + \beta w)^* = \bar{\alpha}v^* + \bar{\beta}w^*, \quad (v^*)^* = v$$

per $v, w \in V$ e $\alpha, \beta \in \mathbb{C}$. Ricordiamo che una $*$ -algebra è un'algebra A con unità insieme a una mappa $a \mapsto a^*$ tale che A diventa uno $*$ -spazio vettoriale e

$$(ab)^* = b^*a^*$$

per tutti $a, b \in A$. Una tale mappa si chiama involuzione algebrica su A . Segue che $1^* = 1$.

Una coalgebra A si dice $*$ -coalgebra se A è dotata di una involuzione $*$ tale che A sia uno $*$ -spazio vettoriale e $\Delta : A \rightarrow A \otimes A$ sia un $*$ -omomorfismo. Ciò significa che

$$\Delta(a^*) = \Delta(a)^*$$

per ogni $a \in A$, dove l'involuzione su $A \otimes A$ è definita da

$$(a \otimes b)^* = a^* \otimes b^*.$$

In una $*$ -coalgebra vale

$$\varepsilon(a^*) = \overline{\varepsilon(a)}, \quad a \in A.$$

Una $*$ -bialgebra è una bialgebra dotata di una involuzione $*$ per la quale sia l'algebra sia la coalgebra siano strutture di $*$ -algebra e $*$ -coalgebra compatibili.

Una $*$ -algebra di Hopf è una bialgebra di Hopf che è anche una $*$ -bialgebra. Sebbene la definizione non imponga alcuna condizione sul comportamento dell'antipodo rispetto all'involuzione, la struttura di algebra di Hopf implica la seguente relazione di compatibilità:

Proposizione 1.3.5

In ogni $*$ -algebra di Hopf A , vale

$$S(S(a^*)^*) = a \quad \text{per ogni } a \in A,$$

cioè

$$S \circ * \circ S \circ * = \text{id}.$$

In particolare, S è invertibile e il suo inverso è dato da

$$S^{-1} = * \circ S \circ *.$$

Corollario 1.3.3

L'algebra duale A' di una $*$ -algebra di Hopf A è una $*$ -algebra con involuzione definita da

$$f^*(a) = \overline{f(S(a)^*)}, \quad f \in A'.$$

1.3.2 Esempi di algebre di Hopf

Abbiamo già visto nel capitolo introduttivo alle algebre di Hopf un esempio costruito estendendo la struttura di gruppo all'insieme dei funzionali sugli elementi del gruppo stesso. In particolare avevamo poi fatto l'esempio di un funzionale che agisce sugli elementi del gruppo di matrici di $SU(2)$, come esportatore del primo elemento in alto a sinistra della matrice data in input. Possiamo formalizzare ora quell'esempio come segue:

1. Algebre di Hopf di coordinate $\mathcal{O}(G)$ di gruppi di Lie matriciali semplici.

Sia G uno dei gruppi di matrici $SL(N, \mathbb{C})$, $SO(N, \mathbb{C})$ oppure $Sp(N, \mathbb{C})$.

Ogni elemento g di G è una matrice complessa $N \times N$, $g = (g_{ij})$. Definiamo le *funzioni di coordinate* u_j^i su G mediante $u_j^i(g) := g_{ij}$, con $g = (g_{ij}) \in G$. Per $g, h \in G$ abbiamo:

$$\Delta(u_j^i)(g, h) = u_j^i(gh) = (gh)_{ij} = \sum_k g_{ik} h_{kj} = \sum_k u_k^i(g) u_j^k(h)$$

e $u_j^i(e) = \delta_{ij}$, da cui

$$\Delta(u_j^i) = \sum_k u_k^i \otimes u_j^k, \quad \varepsilon(u_j^i) = \delta_{ij}. \quad (19)$$

Sia $\mathcal{A} = \mathcal{O}(G)$ la sottoalgebra di $\mathcal{F}(G)$ generata dalle N^2 funzioni u_j^i , per $i, j = 1, 2, \dots, N$. Poiché $\Delta : \mathcal{F}(G) \rightarrow \mathcal{F}(G \times G)$ è un omomorfismo di algebre, abbiamo $\Delta(\mathcal{A}) \subseteq \mathcal{A} \otimes \mathcal{A}$ dalla prima formula di (19). Ogni elemento $g \in G$ ha determinante 1. Ne consegue che la funzione costante uguale a 1 su G appartiene ad $\mathcal{A}$, quindi $\mathcal{A}$ ha un'unità. Inoltre, segue che esistono polinomi p_{ij} in N^2 indeterminate tali che $(g^{-1})_{ij} = p_{ij}(g_{11}, g_{12}, \dots, g_{NN})$, e quindi

$$S(u_j^i)(g) = u_j^i(g^{-1}) = (g^{-1})_{ij} = p_{ij}(u_1^1(g), u_1^2(g), \dots, u_N^N(g)) = p_{ij}(u_1^1, u_1^2, \dots, u_N^N)(g).$$

Pertanto $S(u_j^i) \in \mathcal{A}$ e quindi $S(\mathcal{A}) \subseteq \mathcal{A}$. Di conseguenza, per quanto discusso nell'Esempio 1 precedente, $\mathcal{A} = \mathcal{O}(G)$ è un'algebra di Hopf. Essa è detta *algebra di Hopf di coordinate* di G .

Sia l'algebra astratta delle $F(G)$ che l'algebra appena definita come algebra di coordinate di gruppi di Lie semplici, sono algebre di Hopf dove la moltiplicazione è una mappa commutativa e l'antipodo un omomorfismo di algebre idempotente. Queste condizioni non rimangono in generale vere per un'algebra di Hopf qualsiasi, ma è proprio in questa particolare situazione che nasce l'idea di poter scrivere la natura degli oggetti quantistici su gruppi trascritti in algebre di Hopf *non commutative*. Formalmente questo viene fatto introducendo un parametro di *deformazione* q (indicato anche come e^h dove h assume il significato della costante di Planck) che si caratterizza di un valore specifico per ogni gruppo in particolare. Gli elementi di questa algebra di Hopf *deformata* sono visti allora come funzioni sugli elementi di un *gruppo quantistico*.

Un esempio di algebra di Hopf deformata possiamo costruirlo come :

2. q -deformazioni dei gruppi di matrici $GL_q(n), SL_q(n), \dots$

Ad esempio, $SL_q(2)$ è l'algebra libera generata da $\alpha, \beta, \gamma, \delta$ modulo le relazioni

$$\begin{aligned} \beta\alpha &= q\alpha\beta, & \gamma\alpha &= q\alpha\gamma, & \delta\beta &= q\beta\delta, & \delta\gamma &= q\gamma\delta, \\ \gamma\beta &= \beta\gamma, & \alpha\delta - \delta\alpha &= (q^{-1} - q)\beta\gamma, & \alpha\delta - q^{-1}\beta\gamma &= 1. \end{aligned}$$

Questa è un'algebra di Hopf con

$$\Delta \begin{pmatrix} \alpha & \beta \\ \gamma & \delta \end{pmatrix} = \begin{pmatrix} \alpha & \beta \\ \gamma & \delta \end{pmatrix} \otimes \begin{pmatrix} \alpha & \beta \\ \gamma & \delta \end{pmatrix},$$

counità

$$\varepsilon \begin{pmatrix} \alpha & \beta \\ \gamma & \delta \end{pmatrix} = \begin{pmatrix} 1 & 0 \\ 0 & 1 \end{pmatrix},$$

e antipodo

$$S \begin{pmatrix} \alpha & \beta \\ \gamma & \delta \end{pmatrix} = \begin{pmatrix} \delta & -q\beta \\ -q^{-1}\gamma & \alpha \end{pmatrix}.$$

Per altri esempi di algebre di Hopf si rimanda a [5], di cui riportiamo due esempi così come citati nell'articolo [10].

3. Algebra di gruppo $\mathbb{k}[G]$ di un gruppo G .

$$\Delta(g) = g \otimes g, \quad \varepsilon(g) = 1, \quad S(g) = g^{-1}$$

per $g \in G$, ed estesa linearmente su $\mathbb{k}$.

4. **Algebra di involuppo universale** $U(\mathfrak{g})$ di un'algebra di Lie $\mathfrak{g}$. L'algebra di involuppo universale di un'algebra di Lie $\mathfrak{g}$, indicata con $U(\mathfrak{g})$, è l'algebra associativa generata dagli elementi di $\mathfrak{g}$, nella quale il prodotto è definito in modo che il commutatore riproduca la struttura di Lie, cioè

$$xy - yx = [x, y], \quad x, y \in \mathfrak{g}.$$

Essa contiene $\mathfrak{g}$ come sottospazio vettoriale ed è l'oggetto naturale su cui si definisce una struttura di algebra di Hopf.

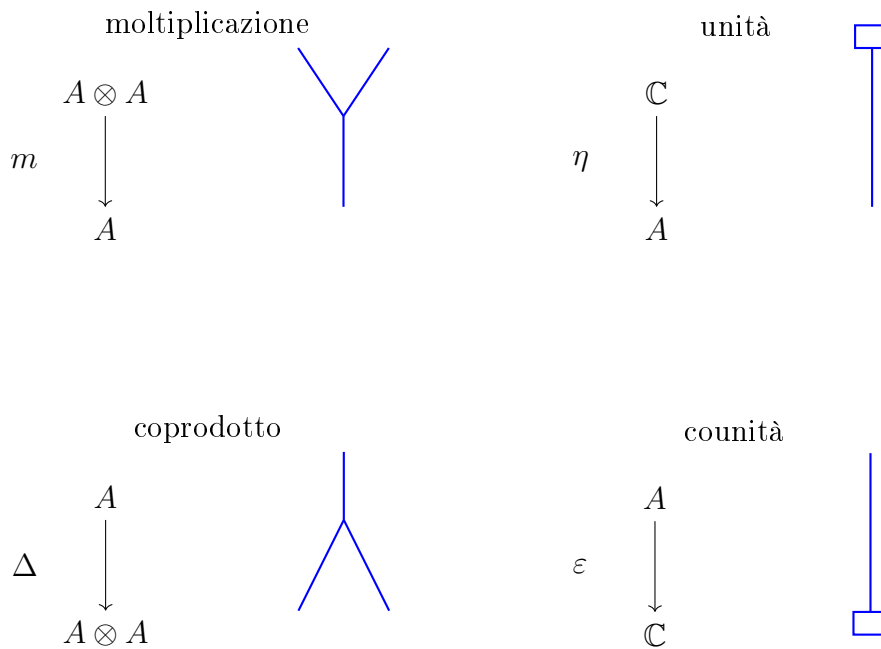
$$\Delta(x) = x \otimes 1 + 1 \otimes x, \quad \varepsilon(x) = 0, \quad S(x) = -x, \quad (x \in \mathfrak{g})$$

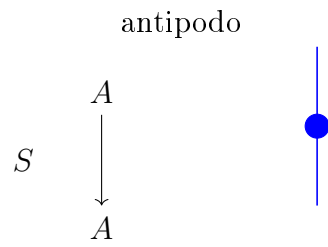
ed estesa come morfismo di (anti-)algebre.

1.3.3 Descrizione grafica delle algebre di Hopf

Possiamo formalmente riformulare la definizione che abbiamo dato di algebra di Hopf in termini di particolari diagrammi che ci aiutino a schematizzare le proprietà che la caratterizzano, passando anche dal ridefinire quelle di algebra, coalgebra e bialgebra.

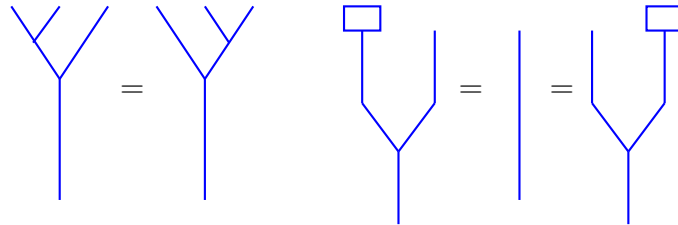
Diciamo quindi che un'algebra di Hopf è uno spazio vettoriale A sull'anello $\mathbb{K}$ (per noi il campo dei numeri complessi $\mathbb{C}$) dotata delle seguenti operazioni:





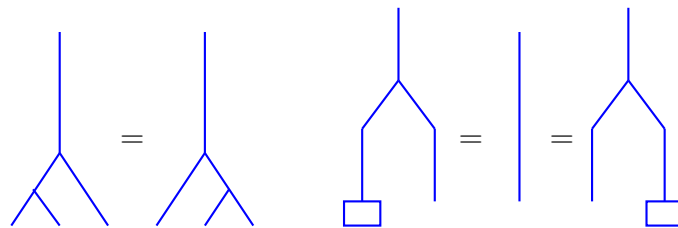
che soddisfano le seguenti relazioni:

(H, m, η) formano un'algebra



$$(xy)z = x(yz) \quad 1x = x = x1$$

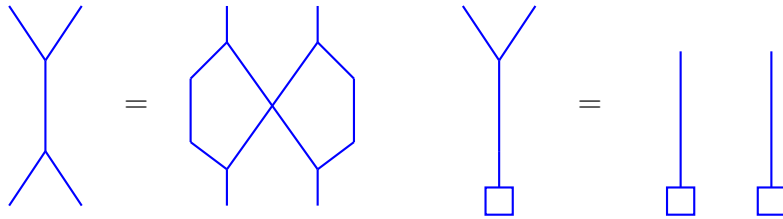
(H, Δ, ε) formano una coalgebra



$$(\Delta \otimes id) \circ \Delta = (id \otimes \Delta) \circ \Delta$$

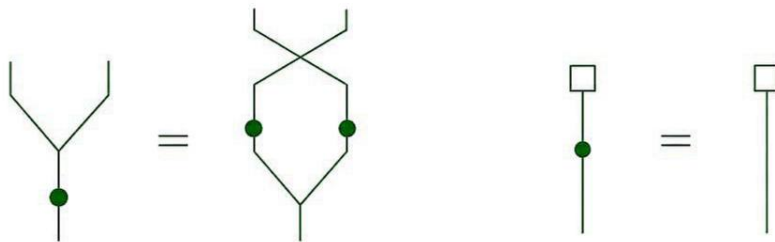
$$(\varepsilon \otimes id) \circ \Delta = id = (id \otimes \varepsilon) \circ \Delta$$

$(H, m, \eta, \Delta, \varepsilon)$ formano una bialgebra



$$\Delta(xy) = \Delta(x)\Delta(y) \quad \varepsilon(xy) = \varepsilon(x)\varepsilon(y)$$

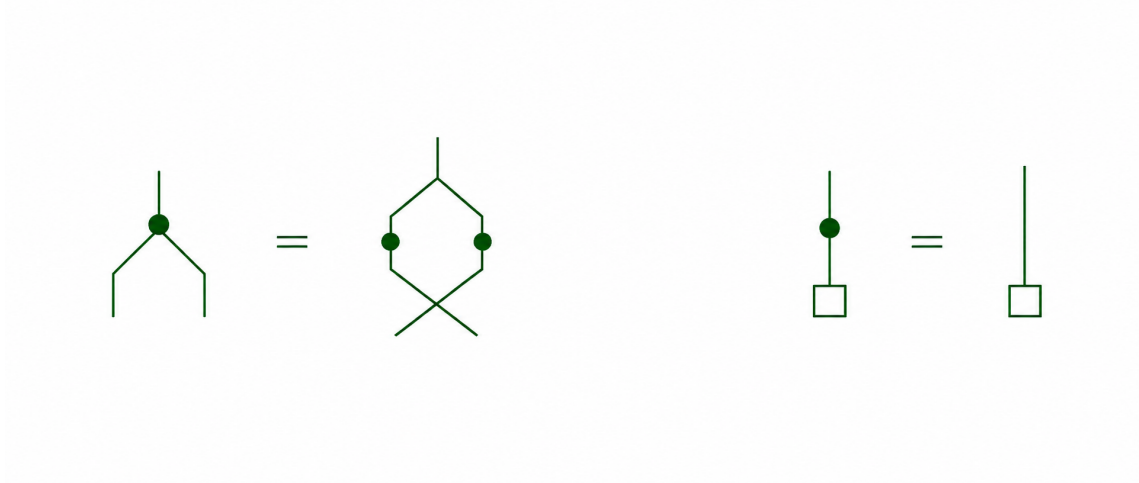
$(H, m, \eta, \Delta, \varepsilon, S)$ forma una Hopf algebra



$$S(xy) = S(y)S(x)$$

$$S(1) = 1$$

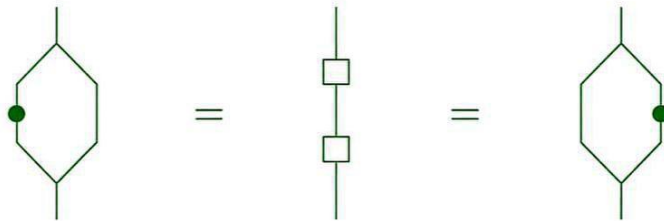
S è un anti-omomorfismo dell'algebra;



$$\Delta \circ S = (S \otimes S) \circ \Delta^{op}$$

$$\varepsilon \circ S = \varepsilon$$

S è un anti-omomorfismo della coalgebra;



$$m \circ (S \otimes id) \circ \Delta = \eta \circ \varepsilon = m \circ (id \otimes S) \circ \Delta$$

Che realizza la relazione fondamentale di un'algebra di Hopf.

La trascrizione delle algebre di Hopf in diagrammi sarà la notazione che ci permetterà di stabilire l'analogia tra lo Z-X calculus e le algebre di Hopf "*non normalizzate*"; che definiamo in Appendice.

Capitolo 2

Quantum Computing

Vogliamo ora fornire le basi del contesto fisico in cui lo Z-X calculus diventerà di nostro interesse : l'informatica quantistica. Lo Z-X Calculus nasce infatti come metodo grafico di semplificazione dei circuiti logici quantistici, l'equivalente quantistico dei circuiti elettronici classici. Ricorderemo allora innanzitutto la definizione classica di circuito, passando per esempi di porte logiche sull'unità fondamentale dell'informatica classica, il bit, per poi trovare il suo corrispettivo quantistico, il qubit. Del bit quantistico accenneremo le basi formali del contesto matematico in cui vive, passando dalla descrizione dei sistemi quantistici come vettori nello spazio di Hilbert alla descrizione del qubit come vettore della sfera di bloch, per poi sviluppare la descrizione dei sistemi quantistici come matrici densità così da formalizzare l'interazione del bit quantistico con l'ambiente del circuito in cui verrà inserito. Torneremo infine a guardare il qubit come l'analogo informatico del bit classico per analizzare in dettaglio ciò che li rende diversi e stabilire ciò che ci è possibile trasportare dall'informatica classica all'informatica quantistica evidenziandone i vantaggi.

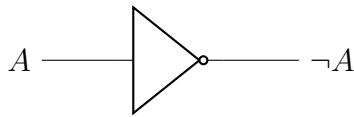
2.1 Informatica classica

Un computer classico è un dispositivo in grado di eseguire programmi, ovvero raccolte di istruzioni scritte ai fini di svolgere dei compiti specifici. Affinché un computer possa eseguire le istruzioni, queste devono essere espresse in un linguaggio che il computer possa comprendere. L'unità fondamentale dell'informazione nei computer classici è il **bit**. Dal punto di vista fisico, il bit è implementato attraverso grandezze elettriche (tipicamente tensione o corrente) opportunamente discretizzate all'interno del circuito entro cui scorrono. I sistemi digitali operano infatti rendendo stabili due soli stati distinguibili, che possono essere interpretati, a seconda del contesto, come "acceso e spento", "vero e falso" oppure "alto e basso". In forma astratta, questi stati vengono rappresentati dai valori 0 e 1, che costituiscono l'alfabeto binario alla base dell'informatica classica e seguono la

matematica dettata dagli assiomi e teoremi dell'algebra di Boole. Un circuito classico è realizzato da un insieme di porte logiche interconnesse progettate per prendere in input segnali binari e restituire il risultato di un'operazione elementare della logica booleana.

2.1.1 Porte logiche a un bit

Stabilito quindi che la variabile fondamentale dell'informatica classica possa assumere solo valori 0 o 1, l'unica operazione logica, secondo l'algebra di Boole, che possiamo eseguire sul singolo bit si realizza nella porta logica NOT :



A	$\neg A$
0	1
1	0

Figura 2.1: Porta logica NOT a un bit

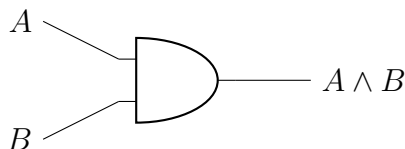
Figura 2.2: Tavola di verità della porta NOT

Quello che abbiamo schematizzato in Figura 2.1 è il primo esempio di porta logica classica. In particolare la porta NOT realizza quello che chiameremo un "bit flip", ovvero un'operazione che manda il bit che si trovi nello stato 0 allo stato 1 e il bit che si trovi nello stato 1 allo stato 0.

2.1.2 Porte logiche a due bit

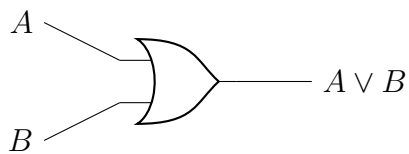
Molta più fantasia si può applicare al caso classico delle porte logiche che prendono in input due bit. Quelle di maggiore rilevanza sono la porta NAND, la porta OR, la porta XOR, e le corrispettive NAND e NOR (porte AND e OR seguite da una porta NOT). Le riportiamo in seguito nella loro forma di simbolo circuitale a sinistra e di tavola di verità sulla destra :

• Porta AND



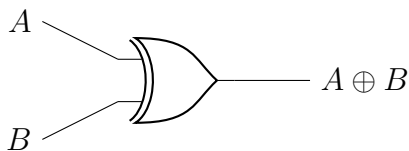
A	B	$A \wedge B$
0	0	0
0	1	0
1	0	0
1	1	1

- Porta OR



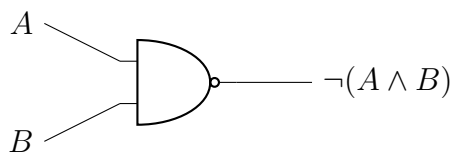
A	B	$A \vee B$
0	0	0
0	1	1
1	0	1
1	1	1

- Porta XOR



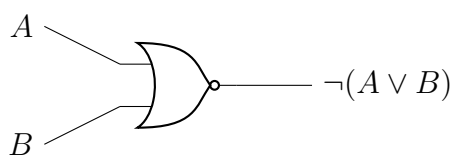
A	B	$A \oplus B$
0	0	0
0	1	1
1	0	1
1	1	0

- Porta NAND



A	B	$\neg(A \wedge B)$
0	0	1
0	1	1
1	0	1
1	1	0

- Porta NOR



A	B	$\neg(A \vee B)$
0	0	1
0	1	0
1	0	0
1	1	0

Ne potremmo elencare tante altre, ma nell'ambito dell'informatica classica esiste un teorema che stabilisce l'"universalità" della porta NAND; un risultato che ci permette di scrivere qualsiasi operazione logica fondamentale in soli termini della porta NAND. Un vantaggio non da poco considerando di dover realizzare queste funzioni logiche in reali componenti circuitali, per cui, in circuiti più robusti, potremo ridurre la nostra attenzione ai soli collegamenti tra componenti se questi sono tutti uguali per costruzione.

A titolo di esempio proviamo a scrivere la funzione NOT in soli termini della porta NAND, utilizzando per le porte logiche definite sopra una notazione più comoda che riportiamo di seguito:

- NOT $x \leftrightarrow 1 - x$;
- x AND $y \leftrightarrow xy$;
- x OR $y \leftrightarrow x + y$;
- x XOR $y \leftrightarrow x \oplus y$;
- x NAND $y = \text{NOT}(x \text{ AND } y) \leftrightarrow 1 - xy$

Allora per scrivere la porta NOT in termini della porta NAND verifichiamo che

$$x \text{ NAND } x = 1 - x^2 = 1 - x = \text{NOT } x$$

□

(l'uguaglianza tra x e x^2 è sempre vera in quanto ci muoviamo nell'ambito di variabili binarie).

Osserviamo che scrivere x^2 significa realizzare una copia della variabile x e nel contesto classico questo ci è sempre permesso. Vedremo invece che in ambito quantistico non ci sarà sempre possibile effettuare la copia di un qubit.

2.2 Definizione del qubit nel formalismo della meccanica quantistica

L'informatica quantistica nasce trasportando ciò che abbiamo visto nella trattazione precedente con il bit, sulla propria unità fondamentale del qubit. Ma cosa è un qubit? Per inserire formalmente il qubit all'interno del framework in cui lavora la meccanica quantistica, abbiamo bisogno di introdurre i postulati fondamentali con cui la teoria definisce come descrivere un qualsiasi sistema quantistico. Abbiamo cioè bisogno di ridefinire in ambito quantistico ciò che nella meccanica classica conosciamo come : *stato*, *osservabile*, *evoluzione* e *misura*, di un sistema. Per farlo seguiremo [9]. Daremo in seguito anche alcuni esempi di come il qubit possa essere realizzato materialmente, ma il vantaggio di introdurre il qubit come un oggetto astratto è di poter sviluppare una

teoria generale del quantum computing e dell'informazione quantistica che non dipenda dalla particolare realizzazione sperimentale del sistema.

Postulato 1. Descrizione degli stati

A ciascun sistema fisico quantistico si fa corrispondere uno spazio di Hilbert $\mathcal{H}$ complesso e separabile. Ogni stato del sistema sarà associato a un raggio del suddetto spazio. Lo stato di un sistema fornisce una descrizione completa del sistema fisico.

Un raggio dello spazio di Hilbert è una classe di equivalenza di vettori che differiscono per un fattore moltiplicativo complesso non nullo. Di conseguenza, gli stati $|\psi\rangle$ ed $e^{i\gamma}|\psi\rangle$, con $|e^{i\gamma}| = 1$, descrivono il medesimo stato fisico, e la loro fase globale non ha alcun significato fisico. È sempre possibile normalizzare un qualsiasi vettore dello spazio $|\psi\rangle$ dividendolo per la sua norma $\|\psi\|$.

Dal primo postulato discende inoltre il *principio di sovrapposizione*: se $|\psi_1\rangle$ e $|\psi_2\rangle$ sono due stati fisici, allora anche ogni loro combinazione lineare $\alpha|\psi_1\rangle + \beta|\psi_2\rangle$, con $\alpha, \beta \in \mathbb{C}$, rappresenta uno stato fisico, purché sia soddisfatta la condizione di normalizzazione

$$|\alpha|^2 + |\beta|^2 = 1 \quad (2.1)$$

A differenza della fase globale, la *fase relativa* tra le componenti della sovrapposizione possiede significato fisico. Infatti, $\alpha|\psi\rangle + \beta|\phi\rangle$ ed $e^{i\gamma}(\alpha|\psi\rangle + \beta|\phi\rangle)$ rappresentano lo stesso stato, mentre $\alpha|\psi\rangle + e^{i\gamma}\beta|\phi\rangle$ descrive, in generale, uno stato fisico differente.

Il più semplice spazio di Hilbert a cui possiamo pensare è proprio quello che definisce il sistema quantistico del nostro qubit: uno spazio di Hilbert bidimensionale generato da due vettori ortonormali che indichiamo come $|0\rangle$ e $|1\rangle$.

Verifichiamo infatti che un qualsiasi sistema quantistico che viva in questo spazio di Hilbert può trovarsi sia nello stato rappresentato dal vettore $|0\rangle$ che nello stato rappresentato dal vettore $|1\rangle$ proprio come il bit classico poteva assumere valori 0 o 1. A differenza del bit classico però il qubit acquista la possibilità di trovarsi in uno stato differente ancora rispetto a quelli che può assumere il bit classico. Il qubit che vive nello spazio di Hilbert oltre che assumere gli stati $|0\rangle$ e $|1\rangle$, può trovarsi in una sovrapposizione lineare di questi ultimi, come :

$$\alpha|0\rangle + \beta|1\rangle$$

con α e β arbitrari, scelti in modo che sia rispettata la condizione di normalizzazione (2.1).

Proprio le condizioni imposte sui coefficienti α e β ci permettono di scrivere per il qubit una rappresentazione particolare, attraverso quella che prende il nome di *sfera di Bloch*.

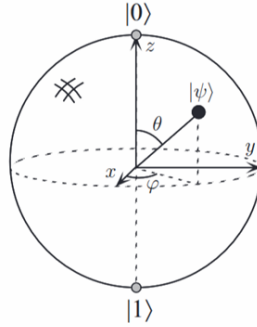


Figura 2.3: Rappresentazione di un qubit mediante la sfera di Bloch.

La forma più generale dello stato del qubit nella rappresentazione della sfera di Bloch sarà quindi ora :

$$|\psi\rangle = \cos\left(\frac{\theta}{2}\right) |0\rangle + e^{i\phi} \sin\left(\frac{\theta}{2}\right) |1\rangle,$$

dove verifichiamo che i coefficienti della sovrapposizione così definiti rispettano a loro volta la condizione di normalizzazione e sono tali per cui : $\theta \in [0, \pi]$ e $\phi \in [0, 2\pi]$.

Specificheremo, dopo aver introdotto la definizione di matrice densità, che attraverso la sfera di Bloch ci sarà possibile individuare se il sistema del qubit in questione si trovi in uno stato puro o in uno stato di miscela (ci basti sapere che per ora noi abbiamo caratterizzato il qubit come un generico stato puro ed è a questi stati che fa riferimento il primo postulato, definiremo in seguito la differenza tra stato puro e stato misto).

Nei nostri circuiti i qubit potrebbero essere realizzati per esempio nello stato di polarizzazione di un fotone oppure nella codifica degli stati $|0\rangle$ e $|1\rangle$ nei livelli iperfini di uno ione.

Postulato 2. Descrizione delle osservabili fisiche

Un'osservabile è una proprietà di un sistema fisico che, in linea di principio, può essere misurata. Nel framework della meccanica quantistica, un'osservabile è rappresentata da un operatore autoaggiunto.

Un operatore è un'applicazione lineare che associa vettori a vettori,

$$A : |\psi\rangle \mapsto A|\psi\rangle, \quad A(a|\psi\rangle + b|\phi\rangle) = aA|\psi\rangle + bA|\phi\rangle.$$

L'aggiunto di un operatore A è definito dalla relazione

$$\langle\varphi|A|\psi\rangle = \langle A^\dagger\varphi|\psi\rangle,$$

valida per tutti i vettori $|\varphi\rangle$ e $|\psi\rangle$ (dove, per semplicità, si è indicato $|A\psi\rangle$ con $A|\psi\rangle$). L'operatore A si dice *autoaggiunto* (o hermitiano) se

$$A^\dagger = A.$$

Un operatore autoaggiunto in uno spazio di Hilbert $\mathcal{H}$ ammette una decomposizione spettrale; in particolare, esiste una base ortonormale completa di $\mathcal{H}$ costituita da autostati dell'operatore. Pertanto un operatore autoaggiunto può essere scritto nella forma

$$A = \sum_n a_n P_n,$$

dove ciascun a_n è un autovalore di A e P_n è il proiettore ortogonale sullo spazio degli autovettori corrispondenti all'autovalore a_n . Nel caso in cui a_n sia non degenere, si ha $P_n = |n\rangle\langle n|$, cioè il proiettore sul corrispondente autovettore. I proiettori soddisfano inoltre le relazioni di ortonormalità e idempotenza

$$P_n P_m = \delta_{nm} P_n, \quad P_n^2 = P_n.$$

Osserviamo quindi che, nella meccanica quantistica, lo stato di un sistema non coincide con le proprietà misurabili ad esso associate, le quali sono rappresentate da oggetti distinti che agiscono sullo spazio di Hilbert del sistema.

Nel caso di un qubit, le osservabili sono rappresentate da operatori hermitiani (autoaggiunti) agenti sullo spazio di Hilbert bidimensionale $\mathbb{C}^2$ e le osservabili fondamentali sono le matrici di Pauli, i cui autostati definiscono le basi di misura utilizzate nel calcolo quantistico:

$$X = \begin{pmatrix} 0 & 1 \\ 1 & 0 \end{pmatrix}, \quad Y = \begin{pmatrix} 0 & -i \\ i & 0 \end{pmatrix}, \quad Z = \begin{pmatrix} 1 & 0 \\ 0 & -1 \end{pmatrix}.$$

Tutte e tre le matrici possiedono gli stessi autovalori $\lambda_1 = +1, \lambda_2 = -1$.

Gli autostati dell'osservabile Z sono i vettori ortonormali con cui abbiamo definito lo spazio bidimensionale del qubit e sono la scelta standard della base computazione nel Quantum Computing

$$|0\rangle = \begin{pmatrix} 1 \\ 0 \end{pmatrix}, \quad |1\rangle = \begin{pmatrix} 0 \\ 1 \end{pmatrix},$$

è sempre possibile però scegliere una diversa base computazionale, utilizzando gli autostati dell'osservabile X

$$|+\rangle = \frac{|0\rangle + |1\rangle}{\sqrt{2}}, \quad |-\rangle = \frac{|0\rangle - |1\rangle}{\sqrt{2}}$$

oppure ancora gli autostati dell'osservabile Y

$$|+i\rangle = \frac{|0\rangle + i|1\rangle}{\sqrt{2}}, \quad |-i\rangle = \frac{|0\rangle - i|1\rangle}{\sqrt{2}}$$

Postulato 3. Misura di osservabili fisiche

Nella meccanica quantistica, il risultato numerico della misura di un'osservabile A su uno stato $|\psi\rangle$ è un autovalore di A . Immediatamente dopo la misura, lo stato quantistico è un autostato di A corrispondente all'autovalore misurato.

Se lo stato del sistema immediatamente prima della misura è $|\psi\rangle$, la probabilità di ottenere come risultato l'autovalore a_n è

$$P(a_n) = \|P_n|\psi\rangle\|^2 = \langle\psi|P_n|\psi\rangle.$$

Se il risultato della misura è a_n , allora il nuovo stato (normalizzato) del sistema dopo la misura diventa

$$\frac{P_n|\psi\rangle}{\sqrt{\langle\psi|P_n|\psi\rangle}}$$

Si osserva quindi che, secondo questa regola, la misura nella meccanica quantistica è un'operazione probabilistica e distruttiva, per cui se si ripete immediatamente dopo nuovamente una misura sullo stato $\frac{P_n|\psi\rangle}{\sqrt{\langle\psi|P_n|\psi\rangle}}$ si otterrà sempre l'autovalore a_n con probabilità unitaria. Si dice che lo stato è *collassato* e non ci è più possibile ricavare le informazioni sullo stato in cui il sistema si trovava prima della misura, se non con la ripetizione della misura stessa su un gran numero di sistemi preparati nello stesso stato iniziale.

Nell'ambito del Quantum Computing ci sarà utile definire diversi tipi di misura oltre quella proiettiva appena presentata, che prende il nome di misura **PVM** (Projector Valued Measurements).

Pensiamo per esempio che la misura nel nostro computer quantistico ci serva a distinguere due stati differenti in cui si possa trovare il sistema. Per esempio potremmo pensare di dover distinguere gli stati $|0\rangle$ e $|1\rangle$ oppure gli stati $|0\rangle$ e $|+\rangle$.

Proviamo a distinguere gli stati di queste due situazioni utilizzando la misura **PVM**: prendiamo la base computazionale della matrice Z per cui la decomposizione spettrale della nostra osservabile sarà la somma dei proiettori

$$P_0 = |0\rangle\langle 0|, \quad P_1 = |1\rangle\langle 1|,$$

tali che

$$\sum_m P_m = I$$

dove con I indichiamo la matrice identità.

Possiamo allora riassumere i possibili risultati della misura come segue :

Stato iniziale	Risultato della misura PVM	Probabilità
$ \psi\rangle = 0\rangle$	$\lambda = +1$	1
	$\lambda = -1$	0
$ \psi\rangle = 1\rangle$	$\lambda = +1$	0
	$\lambda = -1$	1

Tabella 2.1: Risultati della misura PVM dell'osservabile Z sugli stati $|0\rangle$ e $|1\rangle$.

Osserviamo quindi che nel caso di dover distinguere gli stati $|0\rangle$ e $|1\rangle$ questo ci sarebbe possibile con la misura **PVM** in quanto se ottenessimo l'autovalore $+1$ saremmo sicuri che questo venisse dallo stato $|0\rangle$, mentre se ottenessimo l'autovalore -1 saremmo sicuri che venisse dallo stato $|1\rangle$. Lo stesso risultato però non ci è possibile da ottenere nel caso di voler distinguere gli stati $|0\rangle$ e $|+\rangle$.

Stato iniziale	Risultato della misura PVM	Probabilità
$ \psi\rangle = 0\rangle$	$\lambda = +1$	1
	$\lambda = -1$	0
$ \psi\rangle = +\rangle$	$\lambda = +1$	1/2
	$\lambda = -1$	1/2

Tabella 2.2: Risultati della misura PVM dell'osservabile Z sugli stati $|0\rangle$ e $|+\rangle$.

In questo caso infatti notiamo che se ottenessimo l'autovalore $+1$ non saremmo in grado di distinguere se questo sia stato ottenuto dallo stato $|0\rangle$ o dallo $|+\rangle$ e quella misura andrebbe scartata.

La soluzione in questo caso sarebbe utilizzare un diverso tipo di misura, la misura **POVM** (Positive Operator Valued Measurements).

Misura POVM

Le misure **POVM** sono una classe di misure più ampia della meccanica quantistica che contiene come caso particolare anche quello delle misure **PVM**.

Una **POVM** è identificata da una famiglia di operatori $\{M_m\}_m$, che soddisfano unicamente la condizione di completezza

$$\sum_m M_m^\dagger M_m = I.$$

necessaria perchè sia verificata la conservazione delle probabilità. I possibili risultati della misura sono rappresentati da una famiglia di numeri reali $\{\lambda_m\}_m$. Una misura di un sistema nello stato $|\psi\rangle$ presenta in questo caso le seguenti proprietà:

- *Probabilistica*: il risultato della misura può assumere il valore λ_m con probabilità

$$p_m = \langle\psi|M_m^\dagger M_m|\psi\rangle.$$

- *Distruttiva*: dopo una misura con esito λ_m , lo stato del sistema collassa nello stato

$$\frac{M_m|\psi\rangle}{\sqrt{\langle\psi|M_m^\dagger M_m|\psi\rangle}}.$$

così come avevamo visto valere anche per il caso delle misure proiettive.

Nel caso in cui si utilizzi una misura **POVM** non si è più vincolati dalla dimensione dello spazio in cui vive lo stato su cui stiamo effettuando la misura, ci basterà prendere un numero di operatori M_m per cui sia verificata la condizione di completezza. Nel caso del nostro qubit per esempio possiamo provare a risolvere il problema di distinguere i due stati $|0\rangle$ e $|+\rangle$ scegliendo :

$$M_1 = |1\rangle\langle 1| = \begin{pmatrix} 0 & 0 \\ 0 & 1 \end{pmatrix}, \quad M_2 = \left(\frac{|0\rangle - |1\rangle}{\sqrt{2}} \right) \left(\frac{\langle 0| - \langle 1|}{\sqrt{2}} \right) = \frac{1}{2} \begin{pmatrix} 1 & -1 \\ -1 & 1 \end{pmatrix},$$

e infine una matrice M_3 tale per cui sia rispettata la relazione :

$$M_3^\dagger M_3 + M_1^\dagger M_1 + M_2^\dagger M_2 = I.$$

Possiamo riassumere i possibili risultati ancora con una tabella

Stato iniziale	Risultato della misura POVM	Probabilità
$ \psi\rangle = 0\rangle$	λ_1	0
	λ_2	$\neq 0$
	λ_3	$\neq 0$
$ \psi\rangle = +\rangle$	λ_1	1/2
	λ_2	0
	λ_3	$\neq 0$

Tabella 2.3: Risultati della misura POVM applicata agli stati $|0\rangle$ e $|+\rangle$.

Per cui riconosciamo due casi particolari in cui riusciamo ad essere certi di aver misurato sullo stato $|0\rangle$ o sullo stato $|+\rangle$, rispettivamente con λ_2 e con λ_1 come risultati.

Postulato 4. Dinamica ed evoluzione temporale degli stati

L'evoluzione temporale di uno stato quantistico è unitaria; essa è generata da un operatore autoaggiunto, detto Hamiltoniana del sistema.

Nella rappresentazione di Schrödinger della dinamica, il vettore che descrive il sistema evolve nel tempo secondo l'equazione di Schrödinger

$$i\hbar \frac{d}{dt} |\psi(t)\rangle = H |\psi(t)\rangle,$$

dove H è l'operatore Hamiltoniano $\hat{H} = H(\hat{q}, \hat{p})$ ottenuto considerando come operatori le osservabili posizione q e momento coniugato p nella hamiltoniana del sistema. È possibile riscrivere tale equazione, al primo ordine nella quantità infinitesima dt , nella forma

$$|\psi(t + dt)\rangle = (1 - iHdt) |\psi(t)\rangle.$$

L'operatore

$$U(dt) = 1 - iHdt$$

è unitario; infatti, poiché H è autoaggiunto, esso soddisfa

$$U^\dagger U = I$$

al primo ordine in dt .

Poiché il prodotto di operatori unitari è ancora un operatore unitario, l'evoluzione temporale su un intervallo di tempo finito è anch'essa unitaria e può essere scritta nella

forma

$$|\psi(t)\rangle = U(t)|\psi(0)\rangle.$$

Nel caso in cui l'Hamiltoniana sia indipendente dal tempo, l'operatore di evoluzione assume la forma

$$U = e^{-iHt}.$$

Nel contesto del Quantum Computing le matrici unitarie, sono la corretta realizzazione di quelle che defiamo come le porte logiche sui nostri qubit.

Abbiamo già incontrato degli esempi di porte logiche nel commentare il postulato sugli osservabili. Le matrici di Pauli infatti oltre che essere definite correttamente come osservabili risultano soddisfare anche la condizione di unitarietà. Per questo le matrici X, Y e Z così come sopra citate realizzano correttamente la definizione di porta logica quantistica e possono essere riportate anche nella loro forma circuitale :

$$|\psi\rangle \text{ ————— } \boxed{X} \text{ ————— } |\psi'\rangle$$

$$|\psi\rangle \text{ ————— } \boxed{Z} \text{ ————— } |\psi'\rangle$$

$$|\psi\rangle \text{ ————— } \boxed{H} \text{ ————— } |\psi'\rangle$$

Queste ultime, insieme alla porta logica di *Hadamard*,

$$H = \frac{1}{\sqrt{2}} \begin{pmatrix} 1 & 1 \\ 1 & -1 \end{pmatrix} \quad \leftrightarrow \quad |\psi\rangle \text{ ————— } \boxed{H} \text{ ————— } |\psi'\rangle$$

costituiscono il set di porte logiche fondamentali del Quantum Computing.

Postulato 5. Sistemi composti

Lo spazio degli stati di un sistema composto è il prodotto tensoriale degli spazi di Hilbert associati ai singoli sottosistemi. Se n sistemi sono preparati negli stati $|\psi_i\rangle$, lo stato del sistema composto è dato da

$$|\psi\rangle = |\psi_1\rangle \otimes |\psi_2\rangle \otimes \cdots \otimes |\psi_n\rangle.$$

Per semplicità di notazione, il simbolo di prodotto tensoriale viene spesso omesso e si scrive

$$|\psi\rangle = |\psi_1\psi_2\cdots\psi_n\rangle.$$

Questo ci sarà utile nel caso di sistemi a più qubit.

2.2.1 Descrizione statistica dei sistemi quantistici

Avevamo accennato, nel commentare il primo postulato della meccanica quantistica, che gli stati a cui esso fa riferimento sono gli stati definiti come *puri*, ma che esistono anche stati definiti come *misti* (o *di miscela*). Specifichiamo adesso la differenza tra questi due concetti.

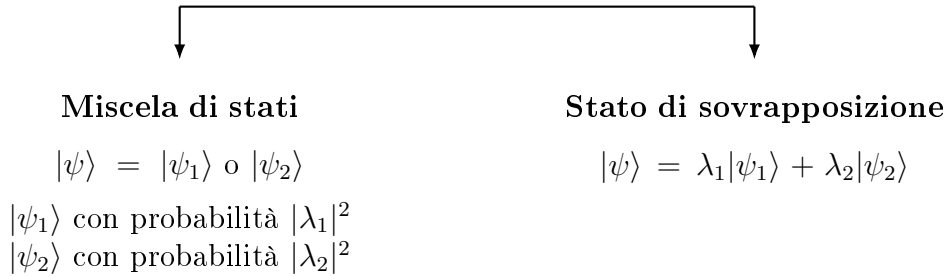
Un sistema in uno stato *puro* è un sistema di cui conosciamo lo stato in cui è stato preparato. Un sistema in uno stato di miscela invece è un sistema la cui conoscenza dello stato in cui è stato preparato contiene a priori una nozione di probabilità.

Spesso nel linguaggio comune confondiamo la nozione di stato di sovrapposizione con quella di probabilità statistica. Se abbiamo uno stato $|\psi\rangle = \alpha|0\rangle + \beta|1\rangle$, dire che il sistema si troverà nello stato $|0\rangle$ con probabilità $|\alpha|^2$ oppure nello stato $|1\rangle$ con probabilità $|\beta|^2$ è parzialmente sbagliato. La formulazione più corretta della nostra frase dovrebbe specificare che la probabilità che il sistema venga trovato nello stato $|0\rangle$ o nello stato $|1\rangle$ è direttamente correlata alla base computazionale su cui proiettiamo il nostro sistema misurandolo. Dovremmo specificare quindi che la nozione di probabilità, nel caso che il nostro sistema si trovi in uno stato di sovrapposizione, entra in gioco solo nell'operazione di misura sul sistema stesso, ma che quest'ultimo senza che noi acquisiamo alcuna informazione su di esso si trova già in una condizione definita. Nel caso invece di un sistema che si trovi in uno stato di miscela, prima ancora di applicare un processo di misura, esiste un livello di ignoranza sullo stato in cui il sistema può trovarsi che assume il significato di probabilità classica così come la conosciamo per la probabilità del lancio di una moneta!

Capire la differenza tra le definizioni di questi due tipi di stati è fondamentale perchè nel lavorare con sistemi quantistici, come il nostro qubit, potremmo trovarci sia nel caso che questo sia stato preparato in uno stato *puro* che nel caso che questo sia stato preparato in uno stato *misto*. Per esempio uno stato di sovrapposizione per un qubit potrebbe essere realizzato come stato di polarizzazione a 45° rispetto agli assi ortogonali, mentre uno stato di miscela si potrebbe ottenere come un fascio di fotoni la cui polarizzazione venga selezionata casualmente come orizzontale o verticale senza che ci venga comunicata la singola selezione.

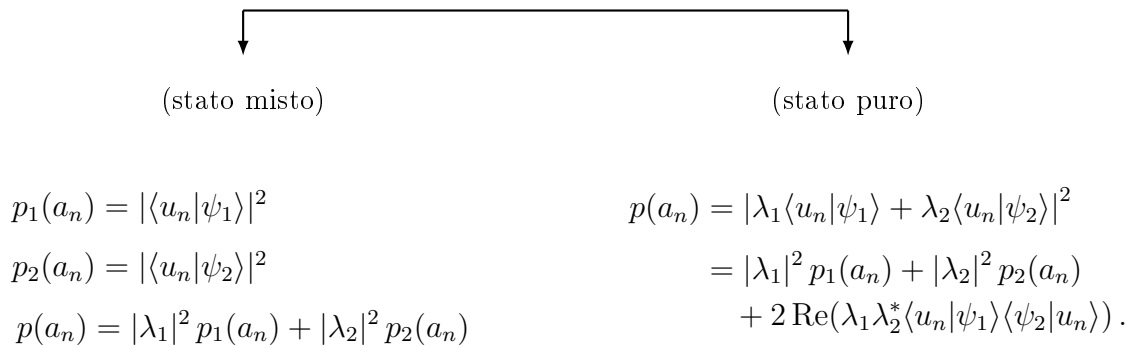
Nei due casi l'applicazione di uno stesso osservabile su l'uno o sull'altro seguirà regole di probabilità molto differenti.

Schematizziamo la situazione come segue :



Scegliamo per entrambi di misurare un'osservabile A utilizzando il suo set di autostati $|u_n\rangle$ tali che : $A|u_n\rangle = a_n|u_n\rangle$.

La probabilità di ottenere come risultato della misura un autovalore a_n di A si calcolerà, nel caso dello stato di miscela come la somma pesata delle singole probabilità proiettive dei due stati in cui è possibile che si trovi il sistema; mentre si calcolerà come unica misura proiettiva dello stato di sovrapposizione nel caso puro:



I due risultati sono significativamente differenti. Notiamo infatti che nel caso dello stato puro la sovrapposizione genera termini di interferenza che nel caso dello stato misto non compaiono. Risulta chiara quindi l'importanza di trattare i due singoli casi coerentemente secondo il proprio formalismo specifico.

Il modo formale di introdurre una trattazione statistica dei sistemi nella meccanica quantistica è quello di definire le *matrici densità*.

Matrici densità

La descrizione dei sistemi quantistici che si trovino in una miscela di stati avviene col passaggio dalla descrizione degli stati come vettori di uno spazio di Hilbert alla descrizione come matrici densità. Introduciamo quindi innanzitutto la definizione di matrice densità nel caso più semplice in cui tutte le probabilità della miscela siano nulle tranne una, che sarà quindi il caso di uno stato puro e trasportiamo tutto ciò che avevamo definito servirci per la descrizione degli stati dai vettori alle matrici. Scopriremo di avere una serie di vantaggi in questo passaggio e che tutte le relazioni definite per il caso puro

si estendono al caso generale.

Caso puro:

Descrizione come vettori		Descrizione come matrici densità
$ \psi(t)\rangle = \sum_n c_n(t) u_n\rangle$	$\longrightarrow$	$\rho(t) = \psi(t)\rangle\langle\psi(t) $
$\sum_n c_n(t) ^2 = 1$	$\longrightarrow$	$\text{Tr } \rho(t) = 1$
$\langle A \rangle(t) = \langle\psi(t) A \psi(t)\rangle$ $= \sum_{n,m} c_n^*(t)c_m(t)A_{nm}$	$\longrightarrow$	$\langle A \rangle(t) = \text{Tr}(\rho(t)A)$
$i\hbar \frac{d}{dt} \psi(t)\rangle = H \psi(t)\rangle$	$\longrightarrow$	$i\hbar \frac{d}{dt}\rho(t) = [H(t), \rho(t)]$

Figura 2.4: Passaggio dalla descrizione tramite vettori di stato alla descrizione tramite matrici densità.

(dove con A_{nm} abbiamo definito gli elementi di matrice dell'osservabile A come $A_{nm} = \langle u_n|A|u_m\rangle$).

Osserviamo quindi che il passaggio dalla descrizione come vettori alla descrizione come matrici densità avviene identificando i coefficienti $c_n^*(t)c_m(t)$ con gli elementi di matrice dell'operatore $\rho(t) = |\psi(t)\rangle\langle\psi(t)|$.

Quello che abbiamo fatto in Fig. 2.4 infatti è stato riscrivere tutte le relazioni esprimibili in termini del vettore di stato $|\psi\rangle$ esclusivamente mediante l'operatore densità $\rho(t)$, mostrando che la conoscenza di $\rho(t)$ è sufficiente a caratterizzare completamente lo stato quantistico del sistema e a ricavare tutte le predizioni fisiche ottenibili dal vettore di stato.

Inoltre, scriveremo la probabilità di ottenere un autovalore a_n misurando l'osservabile A come :

$$P(a_n) = \langle\psi(t)|E_n|\psi(t)\rangle = \text{Tr}[E_n \rho(t)]$$

con

$$E_n = |u_n\rangle\langle u_n|$$

Riassumiamo infine le proprietà che caratterizzano l'operatore densità :

- Gli stati $|\psi(t)\rangle$ ed $e^{i\phi}|\psi(t)\rangle$, che differiscono solo per una fase globale, definiscono lo stesso operatore densità:

$$\rho(t) = |\psi(t)\rangle\langle\psi(t)| = e^{i\phi}|\psi(t)\rangle\langle\psi(t)|e^{-i\phi}.$$

in questo senso l'operatore densità ci evita a priori di dover far attenzione alle fasi globali;

- le formule in $\rho(t)$ sono lineari;
- $\rho^\dagger(t) = \rho(t)$, cioè $\rho(t)$ è hermitiano;
- $\rho^2(t) = \rho(t)$;
- $\text{Tr}[\rho^2(t)] = 1$.

Le ultime due proprietà valgono solo nel caso che si stiano considerando stati puri, in quanto l'operatore densità assume propriamente il significato di proiettore, caratteristica che perderemo nella generalizzazione per gli stati misti, che presentiamo in seguito.

Caso misto:

Consideriamo ora una miscela statistica di stati puri $|\psi_1\rangle, |\psi_2\rangle, \dots, |\psi_k\rangle$, associati a probabilità arbitrarie $p_1, p_2, \dots, p_k$, tali che

$$0 \leq p_k \leq 1, \quad \sum p_k = 1.$$

L'operatore densità della miscela sarà definito come la somma pesata sulle singole probabilità di accadere dei singoli stati puri per i loro rispettivi operatori densità

$$\rho = \sum p_k \rho_k.$$

La probabilità di ottenere l'autovalore a_n misurando l'osservabile A sarà quindi ora

$$P(a_n) = \sum p_k P_k(a_n),$$

dove

$$P_k(a_n) = |\langle u_n | \psi_k \rangle|^2 = \langle \psi_k | E_n | \psi_k \rangle = \text{Tr}(\rho_k E_n).$$

e sostituendo questa espressione nella precedente si ottiene

$$P(a_n) = \sum p_k \text{Tr}(\rho_k E_n) = \sum \text{Tr}(p_k \rho_k E_n) = \text{Tr}(\rho E_n).$$

Come era stato precedentemente accennato, è possibile estendere tutto ciò che si è fatto per il caso puro anche al caso di una miscela statistica.

Riportiamo in seguito i risultati di questo passaggio senza sofferarci sulle dimostrazioni, per cui si rimanda a []

- $\rho^\dagger = \rho$, cioè ρ è hermitiano;
- $\text{Tr}(\rho) = \sum p_k = 1$;
- il valor medio di un osservabile A può essere scritto come

$$\langle A \rangle = \text{Tr}(\rho A) = \sum_n a_n P(a_n) = \text{Tr} \left(\rho \sum_n a_n P_n \right);$$

- l'evoluzione temporale dell'operatore è descritta dall'equazione

$$i\hbar \frac{d\rho(t)}{dt} = [H(t), \rho(t)];$$

- per ogni stato $|u\rangle$ vale la proprietà di positività

$$\langle u | \rho | u \rangle \geq 0.$$

Le proprietà appena discusse caratterizzano completamente l'operatore densità e costituiscono il punto di partenza per la descrizione dei sistemi quantistici aperti. Nei circuiti quantistici, infatti, un qubit reale non può essere considerato un sistema perfettamente isolato: durante la sua evoluzione esso può interagire con l'ambiente circostante, con i dispositivi che implementano le porte logiche o con gli apparati di misura. In tali condizioni lo stato del sistema complessivo è descritto nello spazio di Hilbert del sistema e dell'ambiente e, in generale, non può più essere rappresentato dal solo vettore di stato del qubit.

$$|\psi_{SE}\rangle \in \mathcal{H}_S \otimes \mathcal{H}_E$$

per cui sarà

$$\rho_{SE} = |\psi_{SE}\rangle\langle\psi_{SE}| \quad \text{nel caso puro}$$

$$\rho_{SE} = \sum p_k |\psi_{SE}^k\rangle\langle\psi_{SE}^k| = \sum p_k \rho_{SE}^k \quad \text{nel caso misto}$$

Per distinguere se il sistema totale si trovi in uno stato puro o in uno stato misto riprendiamo la rappresentazione della sfera di Bloch che avevamo introdotto all'inizio del paragrafo.

Una generica matrice densità per un qubit che si trovi in una sovrapposizione di stati potremmo scriverla come:

$$|0\rangle \longrightarrow \rho_0 = |0\rangle\langle 0| = \begin{pmatrix} 1 & 0 \\ 0 & 0 \end{pmatrix} \quad |1\rangle \longrightarrow \rho_1 = |1\rangle\langle 1| = \begin{pmatrix} 0 & 0 \\ 0 & 1 \end{pmatrix}$$

$$|\psi\rangle = \alpha|0\rangle + \beta|1\rangle$$

$$\Rightarrow \rho = (\alpha|0\rangle + \beta|1\rangle)(\alpha^*\langle 0| + \beta^*\langle 1|) = \begin{pmatrix} \alpha \\ \beta \end{pmatrix} \begin{pmatrix} \alpha^* & \beta^* \end{pmatrix} = \begin{pmatrix} |\alpha|^2 & \alpha\beta^* \\ \alpha^*\beta & |\beta|^2 \end{pmatrix}.$$

dove gli elementi fuori dalla diagonale sono detti *termini di coerenza* e sono i responsabili dei fenomeni di interferenza quantistica.

Nel caso invece che il qubit si trovi in uno stato misto scriveremmo :

$$|\psi_1\rangle = |0\rangle, \quad p_1 = |\alpha|^2 \quad \longrightarrow \quad \rho_1 = |0\rangle\langle 0|$$

$$|\psi_2\rangle = |1\rangle, \quad p_2 = |\beta|^2 \quad \longrightarrow \quad \rho_2 = |1\rangle\langle 1|$$

$$\Rightarrow \rho = \sum_k p_k \rho_k = p_1 \rho_1 + p_2 \rho_2 = |\alpha|^2 \rho_1 + |\beta|^2 \rho_2 = \begin{pmatrix} |\alpha|^2 & 0 \\ 0 & |\beta|^2 \end{pmatrix}.$$

Prese due matrici densità potremmo pensare di discriminarne lo stato controllando se sia verificata o meno la condizione di idempotenza, valida solo nel caso di stati puri. Una risposta più veloce però ci arriverebbe dall'osservazione dello stato sulla sfera di Bloch. Scrivendo infatti la più generica matrice densità per un sistema quantistico come

$$\rho = a_0 \mathbb{I} + a_x \sigma_x + a_y \sigma_y + a_z \sigma_z$$

con $a_0, a_x, a_y, a_z \in \mathbb{R}$.

Poiché $\text{Tr}(\rho) = 1$ e $\text{Tr}(\sigma_i) = 0$, segue che $2a_0 = 1$. Pertanto

$$\rho = \frac{\mathbb{I} + \vec{a} \cdot \vec{\sigma}}{2} = \frac{1}{2} \begin{pmatrix} 1 + a_z & a_x - ia_y \\ a_x + ia_y & 1 - a_z \end{pmatrix}.$$

Gli autovalori di ρ sono determinati dall'equazione

$$\det(\rho - \lambda \mathbb{I}) = 0,$$

da cui si ottiene

$$\lambda_{\pm} = \frac{1 \pm |\vec{a}|}{2}.$$

Poiché deve essere $\lambda_{\pm} \geq 0$, si ha immediatamente che $\lambda_+ \geq 0, \forall |\vec{a}|$,

mentre $\lambda_- \geq 0 \iff 1 - |\vec{a}| \geq 0 \iff |\vec{a}| \leq 1$.

Se $|\vec{a}| = 1$, allora $\lambda_- = 0$, e la matrice densità assume la forma

$$\rho = \begin{pmatrix} \lambda_+ & 0 \\ 0 & 0 \end{pmatrix},$$

che corrisponde al caso di uno stato puro. Ci troviamo sulla superficie della sfera di Bloch.

Se invece $|\vec{a}| < 1$, allora $\lambda_+ \neq \lambda_-$, e ci si trova nel caso di uno stato misto. Siamo in questo in uno dei punti interni alla sfera di Bloch.

Sarà importante inoltre nel caso che si abbiano sistema e ambiente in contatto saper distinguere se il sistema totale si trovi in uno stato separabile o in uno stato entangled.

- Stato separabile : $|\psi_{SE}\rangle = |\psi_S\rangle|\psi_E\rangle$

$$\rho_{SE} = \rho_S \otimes \rho_E = (|\psi_S\rangle\langle\psi_S|) \otimes (|\psi_E\rangle\langle\psi_E|)$$

- Stato entangled : $\rho_{SE} \neq \rho_S \otimes \rho_E$
Ad esempio

$$|\psi_{SE}\rangle = \frac{|00\rangle + |11\rangle}{\sqrt{2}}$$

$$\Rightarrow \rho_{SE} = \frac{|00\rangle\langle 00| + |00\rangle\langle 11| + |11\rangle\langle 00| + |11\rangle\langle 11|}{\sqrt{2}} \neq \rho_S \otimes \rho_E$$

Dal punto di vista sperimentale, tuttavia, non si dispone normalmente di informazioni complete sui gradi di libertà dell'ambiente, ma soltanto su quelli del sistema. Per descrivere correttamente il qubit risulta quindi necessario introdurre la *matrice densità ridotta*, ottenuta effettuando la traccia parziale dell'operatore densità del sistema totale rispetto ai gradi di libertà dell'ambiente,

$$\rho_S = \text{Tr}_E(\rho_{SE}).$$

Quest'ultima infatti tiene conto degli effetti dell'interazione tra il sistema e l'ambiente, pur consentendo di descrivere il sistema esclusivamente in termini dei suoi gradi di libertà. Questa forma della matrice densità ridotta risulta particolarmente importante ai fini della nostra trattazione poiché, grazie alla *decomposizione di Schmidt*, si dimostra che se il sistema totale si trova in uno stato entangled allora la matrice densità ridotta ρ_S descrive uno stato misto, mentre se il sistema e l'ambiente si trovano in uno stato separabile la matrice densità ridotta ρ_S rappresenta uno stato puro.

Il formalismo generale con cui verrà descritta l'evoluzione dei qubit nei circuiti quantistici sarà quello delle *quantum operations*, ossia mappe che trasformano matrici densità in matrici densità. Le nozioni introdotte in questo capitolo costituiscono pertanto il fondamento matematico necessario per lo studio dei canali quantistici e della loro rappresentazione.

2.3 Informatica quantistica

Nel paragrafo precedente abbiamo definito le basi matematiche su cui costruire un'informatica quantistica. Abbiamo stabilito per esempio quali caratteristiche debba rispettare una generica porta logica che prenda in input dei qubit, ovvero l'unitarietà.

Se quindi per il nostro qubit le porte logiche sono rappresentate da operatori unitari emerge fin da subito una differenza sostanziale con il bit. Le porte logiche a singolo bit che potevamo pensare seguendo la logica booleana erano di un solo tipo, la porta NOT, mentre le porte logiche associate a un singolo qubit saranno infinite, tante quante sono le matrici 2×2 che possono essere costruite come unitarie.

Riprendiamo allora le matrici che avevamo riconosciuto come porte logiche nel paragrafo precedente per capire quale sia la loro reale azione sul qubit.

La porta logica della matrice di Pauli X realizza il corrispettivo quantistico della porta logica NOT classica, infatti la sua azione sulla base computazionale di

$$|0\rangle = \begin{pmatrix} 1 \\ 0 \end{pmatrix}, \quad |1\rangle = \begin{pmatrix} 0 \\ 1 \end{pmatrix}.$$

porta ai risultati:

$$\begin{aligned} X|0\rangle &= \begin{pmatrix} 0 & 1 \\ 1 & 0 \end{pmatrix} \begin{pmatrix} 1 \\ 0 \end{pmatrix} = \begin{pmatrix} 0 \\ 1 \end{pmatrix} = |1\rangle, \\ X|1\rangle &= \begin{pmatrix} 0 & 1 \\ 1 & 0 \end{pmatrix} \begin{pmatrix} 0 \\ 1 \end{pmatrix} = \begin{pmatrix} 1 \\ 0 \end{pmatrix} = |0\rangle. \end{aligned}$$

La porta X quindi realizza quello che chiamiamo un *Bit Flip* o nel nostro caso un qubit flip:

$$|0\rangle \mapsto |1\rangle, \quad |1\rangle \mapsto |0\rangle.$$

E per linearità, la sua azione su uno stato generico $|\psi\rangle = \alpha|0\rangle + \beta|1\rangle$ sarà

$$X|\psi\rangle = X(\alpha|0\rangle + \beta|1\rangle) = \alpha X|0\rangle + \beta X|1\rangle = \alpha|1\rangle + \beta|0\rangle.$$

L'effetto della porta Z invece sulla base dei vettori $|0\rangle$ e $|1\rangle$ è quello di mandare lo $|0\rangle$ in $|0\rangle$ e $|1\rangle$ in $-|1\rangle$. Mentre prendendo la base computazionale degli autostati della matrice X :

$$|+\rangle = \frac{1}{\sqrt{2}} \begin{pmatrix} 1 \\ 1 \end{pmatrix}, \quad |-\rangle = \frac{1}{\sqrt{2}} \begin{pmatrix} 1 \\ -1 \end{pmatrix}.$$

si ha

$$\begin{aligned} Z|+\rangle &= \begin{pmatrix} 1 & 0 \\ 0 & -1 \end{pmatrix} \left(\frac{1}{\sqrt{2}} \begin{pmatrix} 1 \\ 1 \end{pmatrix} \right) = \frac{1}{\sqrt{2}} \begin{pmatrix} 1 \\ -1 \end{pmatrix} = |-\rangle, \\ Z|-\rangle &= \begin{pmatrix} 1 & 0 \\ 0 & -1 \end{pmatrix} \left(\frac{1}{\sqrt{2}} \begin{pmatrix} 1 \\ -1 \end{pmatrix} \right) = \frac{1}{\sqrt{2}} \begin{pmatrix} 1 \\ 1 \end{pmatrix} = |+\rangle. \end{aligned}$$

Pertanto, nella base costituita dagli stati $|+\rangle$ e $|-\rangle$, anche la porta Z svolge il ruolo di una porta NOT, scambiando tra loro i due stati della base:

$$|+\rangle \mapsto |-\rangle, \quad |-\rangle \mapsto |+\rangle.$$

La porta logica realizzata dalla matrice Y applica contemporaneamente un bit flip e un flip di fase, ma la porta che sarà di nostro maggiore interesse insieme alla porta X e la porta Z è la porta Hadamard. E' facile infatti verificare che la porta Hadamard realizza il passaggio dalla base Z alla base X :

$$H|0\rangle = \frac{1}{\sqrt{2}} \begin{pmatrix} 1 & 1 \\ 1 & -1 \end{pmatrix} \begin{pmatrix} 1 \\ 0 \end{pmatrix} = \frac{1}{\sqrt{2}} \begin{pmatrix} 1 \\ 1 \end{pmatrix} = |+\rangle,$$

$$H|1\rangle = \frac{1}{\sqrt{2}} \begin{pmatrix} 1 & 1 \\ 1 & -1 \end{pmatrix} \begin{pmatrix} 0 \\ 1 \end{pmatrix} = \frac{1}{\sqrt{2}} \begin{pmatrix} 1 \\ -1 \end{pmatrix} = |-\rangle.$$

Pertanto la porta Hadamard trasforma gli stati della base computazionale negli stati della base ruotata:

$$|0\rangle \mapsto |+\rangle, \quad |1\rangle \mapsto |-\rangle.$$

ma anche viceversa e potremmo riguardare all'operazione della porta Hadamard proprio in termini di una sequenza di rotazioni sulla sfera di Bloch:

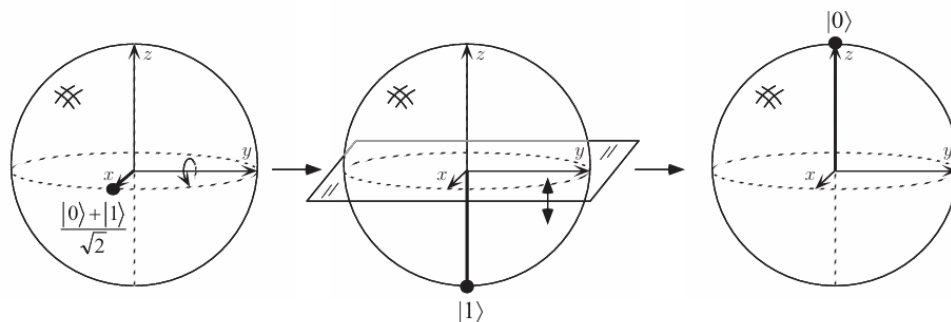
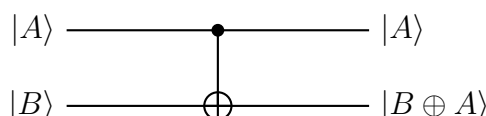


Figura 2.5: Rappresentazione sulla sfera di Bloch dell'azione della porta Hadamard sullo stato $|+\rangle$, [8]

La matrice di Hadamard ci consente di creare sovrapposizioni negli algoritmi.

A questo punto sapendo che le matrici 2×2 unitarie sono potenzialmente infinite potremmo fare ancora tanti esempi, ma nel caso classico avevamo ricordato il teorema di universalità della porta NAND e vogliamo allora costruirne un analogo per il caso quantistico.

Come avviene il passaggio dalle porte logiche a singolo qubit alle porte a più qubit? Nel caso quantistico dobbiamo prestare attenzione che il numero di input sia sempre uguale al numero di output perchè se le nostre porte logiche sono rappresentate da matrici unitarie, deve esistere l'operatore inverso. Il passaggio per costruire una porta NAND quantistica quindi deve tenere conto di questa esigenza. Proviamo con un primo esempio di porta a due qubit, la porta CNOT:



La porta CNOT, agisce applicando al qubit $|B\rangle$ un bit flip a seconda che qubit $|A\rangle$ (che prende il nome di qubit di controllo) sia "acceso o spento":

$$|00\rangle \xrightarrow{CNOT} |00\rangle; \quad |01\rangle \xrightarrow{CNOT} |01\rangle; \quad |10\rangle \xrightarrow{CNOT} |11\rangle; \quad |11\rangle \xrightarrow{CNOT} |10\rangle.$$

Il qubit di controllo agisce come un interruttore. Se esso si trova nello stato $|0\rangle$, la porta non produce alcuna trasformazione sul qubit bersaglio; al contrario, quando il qubit di controllo si trova nello stato $|1\rangle$, sul qubit bersaglio viene applicata la porta X , invertendone lo stato.

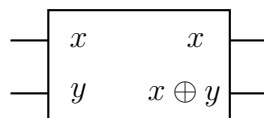
La rappresentazione matriciale della porta CNOT è

$$U_{CN} = \begin{pmatrix} 1 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 \\ 0 & 0 & 0 & 1 \\ 0 & 0 & 1 & 0 \end{pmatrix}.$$

Se l'operazione che viene applicata in modo controllato dalla porta CNOT è la matrice X , possiamo sfruttare la porta Hadamard per costruire una versione della CNOT che utilizzi la porta Z come $X \equiv HZH$.

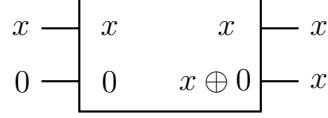
Questo costituisce un vantaggio, perchè in alcuni sistemi fisici è più facile realizzare una porta control- Z .

Tornando al mondo classico, ci accorgiamo che l'analogo circuitale di una porta CNOT sarebbe :



Siamo allora riusciti a trasportare una funzione logica classica nel framework quantistico. Notiamo inoltre che se impostassimo il valore di y a zero, data la corrispondenza tra le

due funzioni logiche, avremmo ottenuto una fotocopiatrice quantistica.



Una delle proprietà fondamentali infatti del bit, discussa nel paragrafo dedicato all'informatica classica, era la possibilità di copiarne liberamente lo stato. Nel caso del qubit, invece, questa proprietà avevamo accennato non essere più garantita. Verifichiamo quindi come si comporta questa presunta "copiatrice quantistica" quando agisce sugli stati della base computazionale $\{|0\rangle, |1\rangle\}$ e, successivamente, su una loro sovrapposizione.

Gli stati della base computazionale su un sistema di due qubit sono rappresentati da

$$|00\rangle = \begin{pmatrix} 1 \\ 0 \\ 0 \\ 0 \end{pmatrix}, \quad |01\rangle = \begin{pmatrix} 0 \\ 1 \\ 0 \\ 0 \end{pmatrix}, \quad |10\rangle = \begin{pmatrix} 0 \\ 0 \\ 1 \\ 0 \end{pmatrix}, \quad |11\rangle = \begin{pmatrix} 0 \\ 0 \\ 0 \\ 1 \end{pmatrix}.$$

allora l'azione della porta CNOT, impostata con il qubit bersaglio sullo stato $|0\rangle$, sui vettori della base $\{|0\rangle, |1\rangle\}$ è

$$U_{CN} \begin{pmatrix} 1 \\ 0 \\ 0 \\ 0 \end{pmatrix} = \begin{pmatrix} 1 \\ 0 \\ 0 \\ 0 \end{pmatrix} = |00\rangle$$

$$U_{CN} \begin{pmatrix} 0 \\ 0 \\ 1 \\ 0 \end{pmatrix} = \begin{pmatrix} 0 \\ 0 \\ 0 \\ 1 \end{pmatrix} = |11\rangle$$

Ma nel caso che il qubit di controllo si trovi in uno stato di sovrapposizione $|\psi\rangle = \alpha|0\rangle + \beta|1\rangle$ la porta CNOT agirà in modo lineare producendo uno stato $|\psi'\rangle = \alpha|0\rangle|0\rangle + \beta|1\rangle|1\rangle \neq |\psi\rangle$.

Se avessimo avuto una reale copiatrice quantistica avremmo ottenuto :

$$(\alpha|0\rangle + \beta|1\rangle)_1 |0\rangle_2 \longrightarrow (\alpha|0\rangle + \beta|1\rangle)_1 \otimes (\alpha|0\rangle + \beta|1\rangle)_2 = |\alpha|^2|00\rangle + \alpha\beta|01\rangle + \beta\alpha|10\rangle + |\beta|^2|11\rangle$$

Questa osservazione ci porta ad un importantissimo teorema del Quantum Computing, di cui riportiamo in seguito la dimostrazione.

Teorema di non clonabilità

Non esiste alcuna trasformazione unitaria capace di clonare uno stato quantistico arbitrario. In particolare, non è possibile costruire un operatore unitario U tale che, per ogni stato $|\psi\rangle$,

$$U(|\psi\rangle \otimes |W\rangle) = |\psi\rangle \otimes |\psi\rangle,$$

dove $|W\rangle$ rappresenta uno stato iniziale fissato, detto anche stato “white”.

Dimostrazione. Supponiamo per assurdo che esista un operatore unitario U tale che

$$U(|\psi\rangle \otimes |W\rangle) = |\psi\rangle \otimes |\psi\rangle$$

e

$$U(|\phi\rangle \otimes |W\rangle) = |\phi\rangle \otimes |\phi\rangle.$$

Poiché U è unitario, deve preservare il prodotto scalare. Si ha quindi

$$\langle\phi|\psi\rangle = \langle\phi|\psi\rangle^2.$$

Questa relazione è soddisfatta solo se

$$\langle\phi|\psi\rangle = 0 \quad \text{oppure} \quad \langle\phi|\psi\rangle = 1.$$

Pertanto la clonazione è possibile soltanto per stati ortogonali oppure identici, ma non per stati quantistici arbitrari. Ne segue che non esiste una macchina di clonazione universale per stati quantistici. $\square$

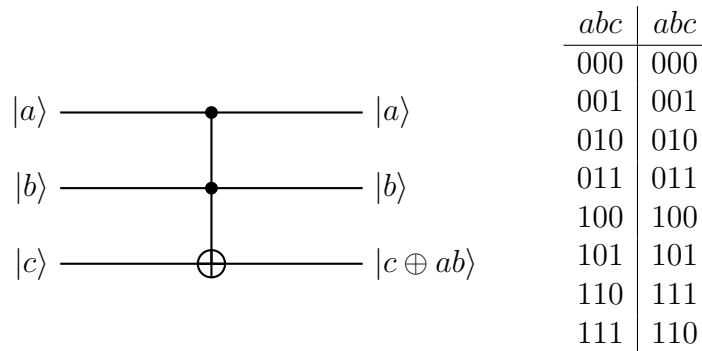
Nonostante la porta CNOT non realizzi una macchina di clonazione, essa introduce una caratteristica esclusivamente quantistica: applicata a uno stato di sovrapposizione, produce infatti uno stato entangled. L’entanglement, piuttosto che la clonazione, costituisce la vera risorsa della computazione quantistica.

L’impossibilità di clonare arbitrariamente un qubit potrebbe far pensare però che non sia possibile costruire un modello di calcolo quantistico universale analogo a quello classico. In un calcolatore classico, infatti, durante l’esecuzione di un algoritmo è spesso necessario duplicare l’informazione e scomporre il problema in una successione di operazioni logiche elementari, riconducibili, grazie al teorema di universalità, alla sola porta NAND.

Nel caso quantistico questo schema deve essere necessariamente modificato. L’assenza di una macchina di clonazione universale impedisce infatti di riprodurre direttamente

i circuiti classici. Ciò nonostante, tale limitazione non preclude la possibilità di costruire un insieme universale di porte quantistiche. Il problema può essere aggirato introducendo un qubit ausiliario (*ancilla*), che permette di realizzare circuiti reversibili capaci di svolgere un ruolo analogo a quello delle porte logiche universali del calcolo classico.

Il modo di realizzare la corrispondenza tra il calcolo classico e quello quantistico quindi è di spostarsi su uno spazio di Hilbert più largo e un esempio importante di questa strategia è quello della porta di Toffoli, o porta CCNOT (*controlled-controlled-NOT*):



Nel passaggio ai qubit notiamo anzitutto che, fissando il terzo qubit nello stato $|0\rangle$, la porta di Toffoli realizza una copia degli stati della base computazionale. Infatti

$$|00\rangle|0\rangle \mapsto |00\rangle|0\rangle,$$

e

$$|11\rangle|0\rangle \mapsto |11\rangle|1\rangle.$$

Mentre, ponendo il terzo ingresso nello stato $c = 1$ si ottiene il corrispettivo di una porta NAND reversibile :

$$c \oplus ab = 1 \oplus ab = \neg(ab),$$

ossia

$$|a, b, 1\rangle \mapsto |a, b, \neg(ab)\rangle.$$

La porta di Toffoli quindi ci permette di trasferire sui qubit i mattoncini fondamentali su cui è costruita l'informatica classica, con la differenza però di avere entanglement lì dove sui bit avevamo una copiatrice. Ci manca un ultimo passaggio per completare la trasposizione dell'informatica classica sui qubit però, dobbiamo avere la possibilità di produrre numeri casuali.

Vogliamo cioè una macchina di Turing non deterministica e questa si realizza sui qubit sfruttando il processo di misura :

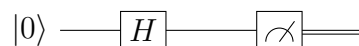


Figura 2.6: Circuito quantistico costituito da una porta Hadamard seguita da una misura.

Applicando la porta Hadamard allo stato iniziale $|0\rangle$, il qubit viene portato nello stato di sovrapposizione

$$|0\rangle \longrightarrow \frac{|0\rangle + |1\rangle}{\sqrt{2}}.$$

La successiva misura proietta il qubit su uno degli stati della base computazionale, restituendo come risultato 0 o 1 con uguale probabilità.

2.3.1 Circuiti quantistici

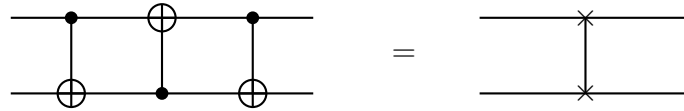
Siamo ora pronti a definire delle regole generali per costruire i circuiti quantistici.

In un circuito quantistico si usano le seguenti convenzioni :

- i circuiti si leggono da sinistra verso destra, i qubit entrano in input a sinistra ed escono in output a destra;
- i fili possono rappresentare il tempo oppure lo spostamento di particelle;
- se non specificato diversamente si assume i qubit siano inizializzati a $|0\rangle$.

Esempio: circuito SWAP

Un primo esempio di circuito quantistico molto semplice è il circuito di SWAP:



L'azione del circuito è

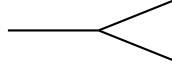
$$|a, b\rangle \xrightarrow{CN_1} |a, a \oplus b\rangle \xrightarrow{CN_2} |a \oplus (a \oplus b), a \oplus b\rangle = |b, a \oplus b\rangle \xrightarrow{CN_3} |b, b \oplus (a \oplus b)\rangle = |b, a\rangle.$$

Pertanto il circuito realizza semplicemente lo scambio dei due qubit:

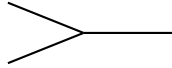
$$|a, b\rangle \longmapsto |b, a\rangle.$$

Dalle proprietà delle trasformazioni unitarie discendono poi alcune operazioni che non possono essere realizzate in un circuito quantistico:

- **Fan-out:** non è possibile sdoppiare arbitrariamente un qubit, poiché ciò richiederebbe la clonazione del suo stato, in contrasto con il teorema di non clonabilità.



- **Fan-in:** non è possibile fondere due qubit in un solo qubit, poiché una tale trasformazione non sarebbe invertibile e violerebbe quindi l'unitarietà dell'evoluzione.



- **Circuiti ciclici:** non è possibile costruire circuiti nei quali l'informazione quantistica ritorni esattamente allo stesso punto del circuito. L'evoluzione di un registro quantistico procede in maniera unidirezionale lungo il circuito.



Un altro esempio di circuito lo era già la "copiatrice quantistica" col CNOT, che è un circuito che genera entanglement. Quest'ultimo inoltre, se costruito con l'aggiunta di una porta Hadamard, produce gli stati che definiscono la *base di Bell*.

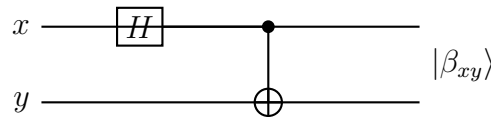


Figura 2.7: Circuito che prepara gli stati di Bell.

L'azione del circuito sugli stati della base computazionale è

$$\begin{aligned}
 |00\rangle &\xrightarrow{H_1} \frac{|0\rangle + |1\rangle}{\sqrt{2}} |0\rangle \xrightarrow{CNOT} \frac{|00\rangle + |11\rangle}{\sqrt{2}} \equiv |\beta_{00}\rangle, \\
 |01\rangle &\xrightarrow{H_1} \frac{|0\rangle + |1\rangle}{\sqrt{2}} |1\rangle \xrightarrow{CNOT} \frac{|01\rangle + |10\rangle}{\sqrt{2}} \equiv |\beta_{01}\rangle, \\
 |10\rangle &\xrightarrow{H_1} \frac{|0\rangle - |1\rangle}{\sqrt{2}} |0\rangle \xrightarrow{CNOT} \frac{|00\rangle - |11\rangle}{\sqrt{2}} \equiv |\beta_{10}\rangle, \\
 |11\rangle &\xrightarrow{H_1} \frac{|0\rangle - |1\rangle}{\sqrt{2}} |1\rangle \xrightarrow{CNOT} \frac{|01\rangle - |10\rangle}{\sqrt{2}} \equiv |\beta_{11}\rangle.
 \end{aligned}$$

Questi quattro stati sono tutti entangled, normalizzati e ortogonali. Chiamiamo stati Φ quelli in cui, in entrambi i rami della sovrapposizione, i qubit si trovano nello stesso stato:

$|\beta_{00}\rangle = |\Phi^+\rangle, |\beta_{10}\rangle = |\Phi^-\rangle$. Chiamiamo invece stati Ψ quelli in cui, in ciascun ramo della funzione d'onda, i due qubit si trovano in stati diversi: $|\beta_{01}\rangle = |\Psi^+\rangle, |\beta_{11}\rangle = |\Psi^-\rangle$. Gli stati

$$\{|\beta_{00}\rangle, |\beta_{01}\rangle, |\beta_{10}\rangle, |\beta_{11}\rangle\} = \{|\Phi^+\rangle, |\Psi^+\rangle, |\Phi^-\rangle, |\Psi^-\rangle\}$$

costituiscono la **base di Bell** e svolgono un ruolo particolarmente importante nel quantum computing perchè sono gli stati che ci permettono di eseguire gli esperimenti di teletrasporto.

2.3.2 I vantaggi dell'informatica quantistica

Per chiudere questa sezione vogliamo introdurre due aspetti dell'informatica quantistica che evidenziano i motivi per cui trasportare la logica dei bit nei qubit ci riserva importanti vantaggi.

- **La complessità quantistica**

Nel calcolo computazionale raggiungere un vantaggio quantistico consiste nel provare ad utilizzare meno risorse di quante ne useremmo con un computer classico. Se per esempio abbiamo un sistema di N qubit lo spazio di Hilbert avrà dimensione 2^N , avremo cioè 2^N vettori della forma $|0100\dots 10\rangle$, con 0 e 1 che spazzano tutte le combinazioni per generare lo spazio di base ed è qui il parallelismo con il calcolo classico. Nel mondo quantistico un qualsiasi stato generico $|\psi\rangle$ si scriverà come combinazione lineare di questi 2^N vettori di base, allora dato che la trasformazione unitaria U che rappresenta la porta logica quantistica rispetta la sovrapposizione, se lo scopo del mio computer dovesse essere confrontare 2^N combinazioni diverse, col computer classico potrei analizzarne una alla volta, col computer quantistico potremmo eseguire in un solo giro il calcolo su ognuna delle combinazioni. Dall'applicazione della trasformazione su tutti gli stati della base in sovrapposizione però potremo estrarre solo un ramo del risultato applicando un processo di misura. Il vero vantaggio quindi consiste nel progettare un algoritmo quantistico che sfrutti interferenze costruttive e distruttive così da amplificare la risposta al nostro problema e sopprimere i risultati che non ci interessano.

Per darci un'idea di quanto concretamente questo vantaggio ci venga in aiuto, diamo in seguito qualche esempio numerico:

se $N = 4$ avremo 16 possibili output

se $N = 100$ avremo $2^{100} \approx 10^{30}$ possibili output

se $N = 300$ avremo 2^{300} possibili output che è un numero maggiore del numero di atomi presenti nell'universo. Se dovessimo allora pensare che eseguire il calcolo sull'ultimo N nel tempo di un solo giro, non ci basterebbero gli atomi dell'universo per implementare abbastanza sistemi da confrontare tra loro.

Un esempio di algoritmo molto semplice che sfrutta la complessità quantistica è l'*Algoritmo di Deutsch-Jozsa*.

- **La comunicazione quantistica**

Un secondo contesto dove il qubit risulta un'alternativa interessante al bit classico è nel caso della sicurezza crittografica. Ad oggi infatti disponiamo di metodi di crittografia classici che garantiscono un livello di sicurezza perfetto, a meno però che sia altrettanto perfetto il meccanismo con cui si scambia la chiave segreta di decodifica. I nostri sistemi classici di trasmissione di una chiave si basano attualmente su algoritmi di fattorizzazione di numeri interi. La moltiplicazione di due numeri interi infatti è un processo banale, mentre l'operazione per tornare indietro ai due fattori originali in generale è un problema esponenzialmente difficile, i migliori algoritmi classici conosciuti richiedono tempo sub-esponenziale per risolverlo. Nel 1994 Peter Shor ha però dimostrato che un computer quantistico può fattorizzare interi in tempo polinomiale accendendo un fanale sulla necessità di rivedere i nostri attuali sistemi di sicurezza. Attualmente, non esistono algoritmi classici noti in grado di fattorizzare numeri in tempi polinomiali. Questo significa che, ad oggi, il processo di fattorizzazione è inefficiente, ma non possiamo prevedere se in futuro la situazione cambierà. Pertanto, affidarsi alla moltiplicazione come meccanismo di sicurezza potrebbe non garantirci protezione per sempre. È proprio qui che interviene la meccanica quantistica. I protocolli di distribuzione delle chiavi che utilizzano i qubit al posto dei bit riescono a identificare con certezza quando la chiave sia stata intercettata e per questo rappresentano l'alternativa efficace per rendere sicure le nostre comunicazioni a lungo raggio. Un esempio di protocollo quantistico che utilizza i qubit per inviare una chiave è il **protocollo BB84**

Il metodo *BB84* consiste nello sfruttare la nozione di probabilità propria del meccanismo di misura sui qubit. Secondo il protocollo Alice e Bob si scambiano inizialmente una sequenza di qubit che Alice ha precedentemente preparato in una certa base (supponiamo Alice abbia potuto scegliere tra le basi X e Z) senza comunicare a Bob quale base Alice abbia scelto per ogni qubit. Bob a quel punto acquisirà una misura sui qubit di Alice utilizzando per ogni qubit una tra le due basi a disposizione. Il protocollo prevede ora che Alice comunichi pubblicamente a Bob le basi in cui aveva preparato i propri qubit (senza ancora comunicare il risultato dei qubit stessi), così che Bob adesso conservi solo le misure dove lui e Alice hanno utilizzato la stessa base. Infine Alice e Bob scelgono un sottoinsieme dei qubit scambiati dei quali si comunicano il risultato. Se non c'è stata intercettazione Alice e Bob troveranno per ogni qubit lo stesso risultato e useranno i qubit del sottoinsieme che non hanno rivelato come chiave segreta. Se invece nella comunicazione c'è stata una intercettazione Alice e Bob troveranno dei qubit discordi e scarteranno questo tentativo.

Una schematizzazione del protocollo *BB84* può essere pensata come segue:

Qubit di A	Valore di A	Base di A	Stato di A	Osservabile di B	Misura di B	Discussione	Autenticazione	Chiave
1	0	σ_X	$ 0\rangle_X$	σ_Z	0	×		
2	1	σ_X	$ 1\rangle_X$	σ_X	1	ok	→	→ 1
3	1	σ_Z	$ 1\rangle_Z$	σ_X	1	×		
4	1	σ_Z	$ 1\rangle_Z$	σ_Z	1	ok	1 ok	
5	0	σ_X	$ 0\rangle_X$	σ_X	0	ok	→	→ 0
6	0	σ_Z	$ 0\rangle_Z$	σ_X	1	×		
7	1	σ_X	$ 1\rangle_X$	σ_X	1	ok	1 ok	
8	0	σ_Z	$ 0\rangle_Z$	σ_Z	0	ok	0 ok	
9	1	σ_Z	$ 1\rangle_Z$	σ_X	1	×		
10	0	σ_Z	$ 0\rangle_Z$	σ_Z	0	ok	→	→ 0

Figura 2.8: Tabella di schematizzazione del protocollo *BB84* di distribuzione di una chiave quantistica

Per ogni singolo qubit esistono diverse possibilità di esito nel caso avvenga un'intercettazione. Se il qubit di Alice venisse intercettato prima di arrivare a Bob, anche l'intercettatore che chiameremo *E* avrebbe da compiere una scelta sulla base da utilizzare per la misura. Nel caso scegliesse la stessa di Alice, A e B non si accorgerebbero del suo intervento. Se invece scegliesse una base diversa da quella di Alice ci sarebbero due possibili risultati e per ognuno Bob misurando sulla stessa base di Alice potrebbe ottenere il risultato corretto oppure sbagliato con uguale probabilità. Quindi ci sarebbero tre casi in cui per singolo qubit A e B non si accorgono di essere stati intercettati, e la probabilità totale che l'intercettazione non venga rivelata su un singolo qubit è

$$p_{\text{non riv.}} = \frac{1}{2} + \frac{1}{8} + \frac{1}{8} = \frac{3}{4}.$$

Se però il processo viene ripetuto per un numero N di qubit molto elevato, questa probabilità tende a zero

$$\left(\frac{3}{4}\right)^N \xrightarrow{N \rightarrow \infty} 0.$$

Un'ulteriore possibilità per trasmettere informazione in modo sicuro è rappresentata dai protocolli di teletrasporto. In seguito ne presentiamo un esempio, seguendo la trattazione di [8].

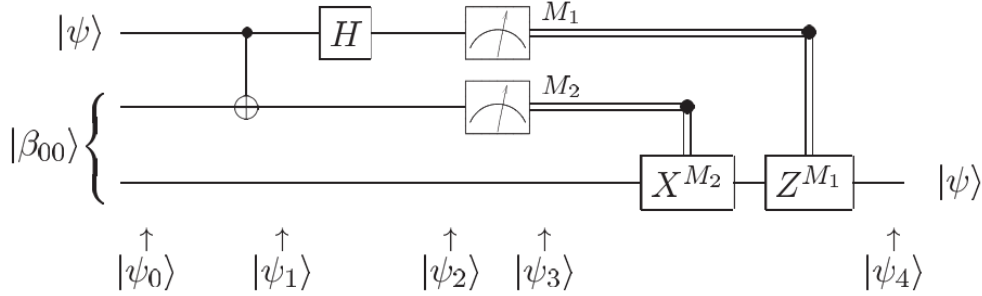


Figura 2.9: Circuito per il teletrasporto dello stato di un qubit. Le due linee superiori rappresentano il sistema di Alice, mentre la linea inferiore rappresenta il sistema di Bob. I misuratori indicano le operazioni di misura e le doppie linee che ne escono trasportano bit classici (le linee singole rappresentano invece i qubit)

Il circuito quantistico mostrato nella Figura 2.9 fornisce una descrizione precisa del protocollo di teletrasporto quantistico. Lo stato da teletrasportare è

$$|\psi\rangle = \alpha|0\rangle + \beta|1\rangle,$$

dove α e β sono ampiezze di probabilità ignote. Lo stato iniziale dell'intero circuito è quindi

$$|\psi_0\rangle = |\psi\rangle|\beta_{00}\rangle.$$

$$|\psi_0\rangle = \frac{1}{\sqrt{2}} [\alpha|0\rangle(|00\rangle + |11\rangle) + \beta|1\rangle(|00\rangle + |11\rangle)].$$

D'ora in avanti adotteremo la convenzione secondo cui i primi due qubit (quelli a sinistra) appartengono ad Alice, mentre il terzo appartiene a Bob. Come illustrato in precedenza, Alice e Bob condividono inizialmente una coppia di Bell. Alice applica quindi una porta CNOT ai propri qubit, ottenendo

$$|\psi_1\rangle = \frac{1}{\sqrt{2}} [\alpha|0\rangle(|00\rangle + |11\rangle) + \beta|1\rangle(|10\rangle + |01\rangle)].$$

Successivamente Alice applica una porta di Hadamard al primo qubit, ottenendo

$$|\psi_2\rangle = \frac{1}{2} [\alpha(|0\rangle + |1\rangle)(|00\rangle + |11\rangle) + \beta(|0\rangle - |1\rangle)(|10\rangle + |01\rangle)].$$

Lo stato così ottenuto può essere riscritto raggruppando opportunamente i termini:

$$|\psi_2\rangle = \frac{1}{2} [|00\rangle(\alpha|0\rangle + \beta|1\rangle) + |01\rangle(\alpha|1\rangle + \beta|0\rangle)]$$

$$+|10\rangle(\alpha|0\rangle - \beta|1\rangle) + |11\rangle(\alpha|1\rangle - \beta|0\rangle)].$$

Questa espressione si scompone naturalmente in quattro termini. Nel primo caso i due qubit di Alice si trovano nello stato $|00\rangle$, mentre il qubit di Bob si trova nello stato

$$\alpha|0\rangle + \beta|1\rangle,$$

che coincide con lo stato iniziale $|\psi\rangle$. Se Alice esegue una misura ottenendo come risultato 00, il qubit di Bob collassa quindi nello stato $|\psi\rangle$.

Analogamente, dalla precedente espressione si ricavano gli stati del qubit di Bob corrispondenti ai possibili esiti della misura di Alice:

$$00 \longrightarrow |\psi_3(00)\rangle = \alpha|0\rangle + \beta|1\rangle,$$

$$01 \longrightarrow |\psi_3(01)\rangle = \alpha|1\rangle + \beta|0\rangle,$$

$$10 \longrightarrow |\psi_3(10)\rangle = \alpha|0\rangle - \beta|1\rangle,$$

$$11 \longrightarrow |\psi_3(11)\rangle = \alpha|1\rangle - \beta|0\rangle.$$

A seconda dell'esito della misura effettuata da Alice, il qubit di Bob si troverà dunque in uno di questi quattro possibili stati. Alice allora conoscendo in quale di questi stati si trovi il qubit di Bob dovrà comunicargli quali operazioni eseguire per ottenere il suo stato iniziale.

Il protocollo di teletrasporto presenta numerose proprietà interessanti, alcune delle quali non approfondiamo e per le quali rimandiamo sempre al [8]. Per il momento è utile soffermarsi su due aspetti fondamentali

Il primo riguarda la possibilità di trasmettere informazione a velocità superiore a quella della luce. A prima vista potrebbe sembrare che il teletrasporto quantistico renda possibile una comunicazione istantanea. Ciò sarebbe tuttavia in contrasto con la teoria della relatività, secondo la quale una trasmissione superluminale consentirebbe anche l'invio di informazione nel passato. In realtà il teletrasporto quantistico non permette alcuna comunicazione più veloce della luce, poiché Alice deve necessariamente comunicare a Bob il risultato della propria misura attraverso un canale classico. Senza questa informazione Bob non è in grado di ricostruire lo stato iniziale e, di conseguenza, il protocollo non trasmette alcuna informazione utile. Poiché la comunicazione classica è limitata dalla velocità della luce, anche il teletrasporto quantistico rispetta tale limite, eliminando ogni apparente paradosso. Un secondo aspetto, altrettanto sorprendente, è che il protocollo sembra realizzare

una copia dello stato quantistico da teletrasportare, in apparente contraddizione con il teorema di non clonazione discusso in precedenza. Anche questa violazione è soltanto apparente: al termine del protocollo lo stato iniziale non esiste più sul qubit di Alice, mentre compare esclusivamente sul qubit di Bob. L'informazione quantistica, dunque, non viene duplicata ma semplicemente trasferita da un sistema fisico a un altro.

Che cosa ci insegna il teletrasporto quantistico? Molto più di quanto possa sembrare a prima vista. Esso non rappresenta soltanto un elegante protocollo per manipolare stati quantistici, ma evidenzia anche come differenti risorse dell'informazione quantistica possano essere convertite l'una nell'altra. In particolare, una coppia entangled condivisa e due bit di comunicazione classica costituiscono una risorsa equivalente, sotto determinati aspetti, alla trasmissione di un qubit. Lo sviluppo dell'informatica quantistica e della teoria dell'informazione quantistica ha mostrato numerosi esempi di questa intercambiabilità tra risorse, molti dei quali si fondano proprio sul protocollo di teletrasporto. Ad esempio, esso può essere impiegato per realizzare porte logiche quantistiche robuste nei confronti del rumore oppure per costruire codici di correzione degli errori. Nonostante tali connessioni con altri ambiti della teoria quantistica, siamo ancora soltanto agli inizi della comprensione delle ragioni profonde per cui il teletrasporto quantistico sia possibile; proprio questo rappresenta uno degli aspetti più affascinanti della disciplina.

Capitolo 3

Z-X Calculus e algebre di Hopf

Introdotti i concetti di porta logica e circuito quantistico vogliamo ora definire la corrispondenza tra i circuiti quantistici e gli ZX -diagrams così che ci sia chiaro come le regole dello ZX -Calculus ci possano essere utili per la costruzione di circuiti equivalenti che risultino semplificati. Nella costruzione del linguaggio degli ZX -diagrams e dello ZX -Calculus seguiremo il secondo capitolo di [1].

3.1 Il linguaggio degli ZX -diagrams

Il linguaggio dello ZX -Calculus consiste in una serie di componenti connesse da fili, simili a circuiti elettronici o diagrammi di flusso di un algoritmo. Il più semplice diagramma non banale che ci possa venire in mente è un filo che scorre dall'alto al basso collegando un input e un output:

$$1_Q = \begin{array}{c} \textit{in} \\ \boxed{} \\ \textit{out} \end{array}$$

Pensiamo ai diagrammi come racchiusi in un'interfaccia fissa che definisce i punti attraversati dai fili. In ogni punto dell'interfaccia deve essere presente esattamente un filo e dobbiamo distinguere quale filo è collegato a quale punto.

Questa è quindi l'unica premessa che rende diversi diagrammi del tipo :

$$1_{\mathcal{Q} \otimes \mathcal{Q}} = \begin{array}{|c|c|} \hline \text{in}_1 & \text{in}_2 \\ \hline \text{out}_1 & \text{out}_2 \\ \hline \end{array} \quad \sigma_{\mathcal{Q}} = \begin{array}{|c|c|} \hline \text{in}_1 & \text{in}_2 \\ \hline \text{out}_1 & \text{out}_2 \\ \hline \end{array}$$

e non è importante quale dei due fili incrociati passi sopra e quale sotto.
I fili possono piegarsi collegando due uscite per formare un *cappuccio* o due ingressi per creare una *tazza*

$$\eta_{\mathcal{Q}} = \begin{array}{|c|c|} \hline \text{out}_1 & \text{out}_2 \\ \hline \end{array} \quad \varepsilon_{\mathcal{Q}} = \begin{array}{|c|c|} \hline \text{in}_1 & \text{in}_2 \\ \hline \end{array}$$

Da ora in avanti gli input e gli output non verranno più etichettati e li distingueremo semplicemente ordinandoli da sinistra verso destra. Vedremo in seguito che l'ordine degli output non ci interesserà nell'applicazione ai circuiti quantistici.

Scriviamo $\mathcal{D} : m \rightarrow n$ per indicare che il diagramma $\mathcal{D}$ possiede m ingressi e n uscite. Oltre ai fili, lo ZX-calculus contiene quattro tipi di componenti:

- **Vertici Z** (punti verdi ●), etichettati da un angolo $\alpha \in [0, 2\pi)$, detto *fase*, che possono avere un numero arbitrario di ingressi e di uscite, compreso il caso in cui non ne abbiano.
- **Vertici X** (punti rossi ●), anch'essi etichettati da una fase. Possono avere un numero arbitrario di ingressi e di uscite, compreso il caso in cui non ne abbiano.
- **Vertici H** (quadrati gialli ■), devono avere esattamente un ingresso e un'uscita.
- **Vertici $\sqrt{D}$** (rombi ◆), non possono avere né ingressi né uscite.

$$Z_m^n(\alpha) = \begin{array}{|c|} \hline n \\ \hline \text{●} \\ \hline m \\ \hline \end{array} \quad X_m^n(\alpha) = \begin{array}{|c|} \hline n \\ \hline \text{●} \\ \hline m \\ \hline \end{array} \quad H = \begin{array}{|c|} \hline \text{H} \\ \hline \end{array} \quad \sqrt{D} = \blacklozenge$$

Noi ci riferiremo ai vertici X e Z come *spiders* e scegliamo per convenzione di omettere la fase α se questa è nulla. Gli ZX-diagrams possono essere costruiti a partire da questi quattro componenti con fili dritti, incrociati, piegati in due modi diversi :

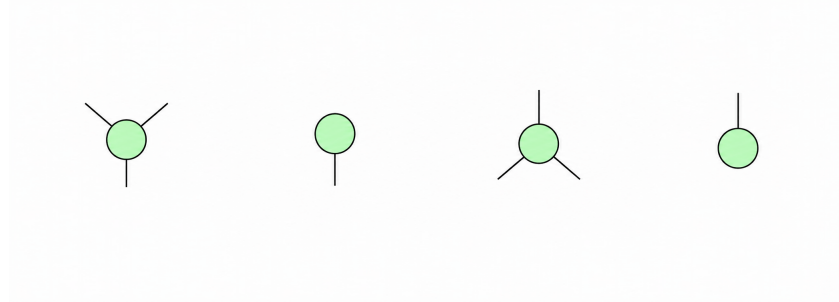
-

-

A collection of various Feynman diagrams, including loops, vertices, and external lines, some labeled with 'H'.

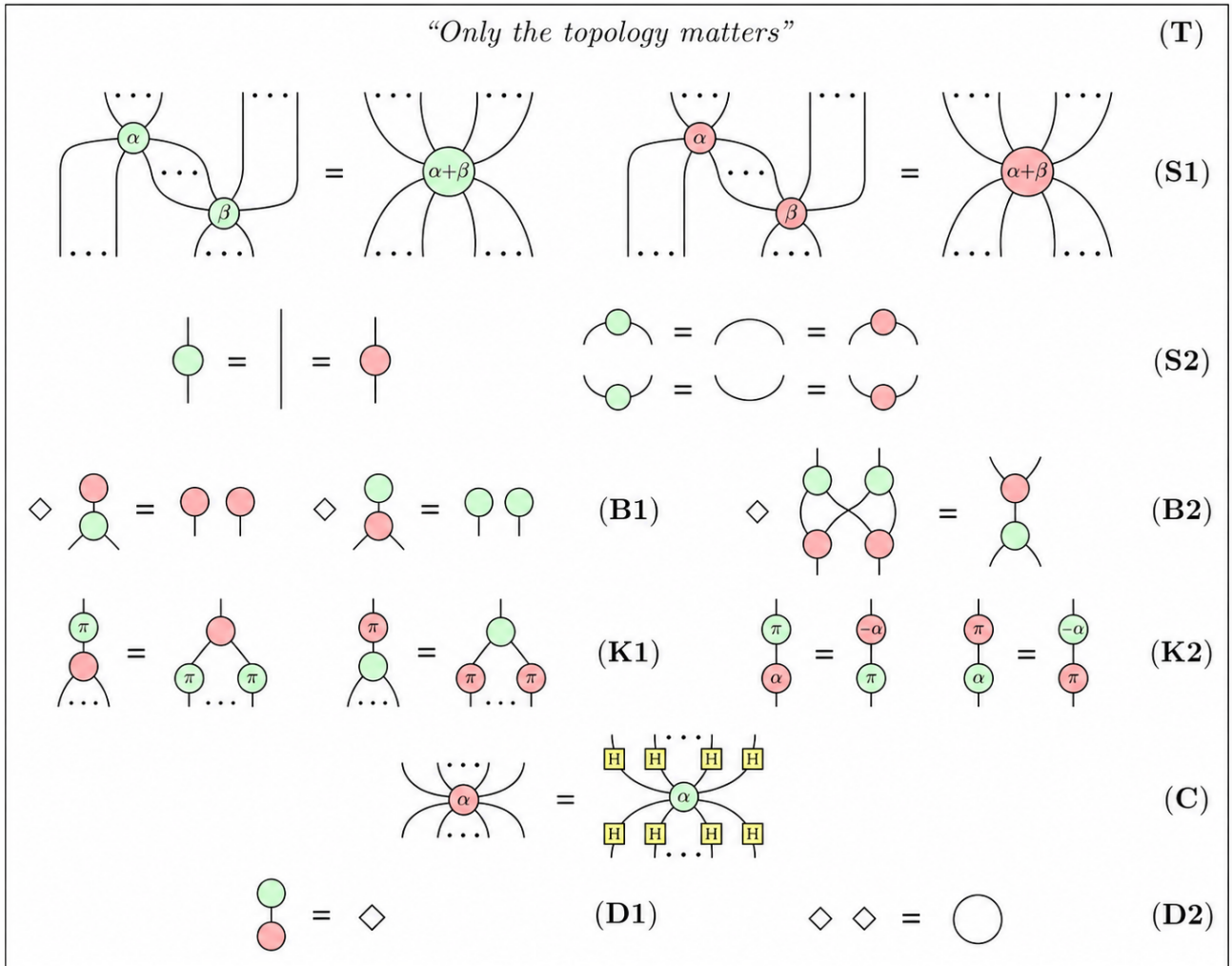
74

ed un'uscita (inizio di un valore), che chiameremo *punti*, oppure ancora un solo ingresso e due uscite (operazione di copia $\mathcal{D} : 1 \longrightarrow 2$) e infine un solo ingresso ma nessuna uscita (operazione di cancellazione):



3.2 Lo ZX -Calculus

In aggiunta alle regole per costruire i diagrammi, lo ZX -Calculus consiste in un set di equazioni che specificano come un diagramma può essere trasformato in un altro. Schematizziamo in seguito il set completo delle regole dello ZX -Calculus e successivamente le commentiamo singolarmente.



T-Rules

In modo sintetico (per una trattazione completa si rimanda al capitolo 4 di [1]) potremmo dire che le T-rules *"only the topology matters"* ci dicono che i fili possono essere stretchati, allungati, girati senza alterare il significato logico del diagramma purchè le connessioni siano mantenute.

Più precisamente, dopo aver identificato (per esempio numerando) gli ingressi e le uscite, qualunque deformazione topologica della struttura interna della rete produce una rete equivalente a quella data.

Due importanti esempi di queste deformazioni *omotopiche* sono:

$$\text{(T}_1\text{)} \quad \text{wavy line} \rightarrow \text{green vertex} \rightarrow \text{four lines} = \text{straight line} \rightarrow \text{green vertex} \rightarrow \text{four lines}$$

$$\text{(T}_2\text{)} \quad \text{wavy line} = \text{straight line}$$

Osservazione 1

Poichè i fili possono essere allungati senza conseguenze logiche sul diagramma, aggiungere un tratto rettilineo di filo all'ingresso o all'uscita di un diagramma non ha alcun effetto. Dunque i fasci di fili rettilinei agiscono come elementi identità nell'algebra dei diagrammi.

Osservazione 2

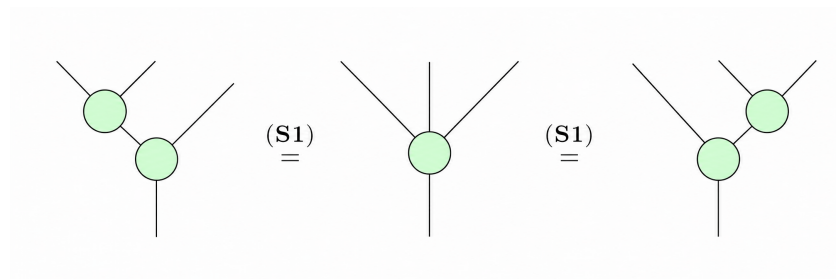
Sebbene le T-rules si riassumano come "*only the topology matters*", questo non significa che la topologia sia sempre preservata. Le altre regole per esempio possono infatti modificare la topologia del diagramma in vari modi, per esempio per rimuovere anelli o per disconnettere vertici precedentemente connessi.

S-Rules

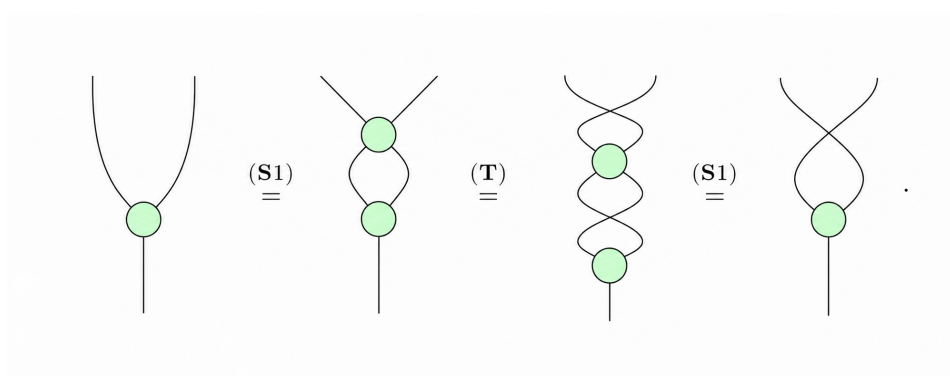
Le S-rules o regole *spiders* riguardano il modo in cui i vertici dello stesso colore interagiscono tra loro. La regola *S1* afferma che vertici connessi dello stesso colore possono essere fusi in un unico vertice la cui fase sia la somma dei due vertici iniziali. Viceversa un vertice può essere decomposto lungo uno o più fili di connessione frammentando la fase iniziale. Si noti che il numero dei fili di connessione è irrilevante.

L'equazione *S2* invece specifica quando gli *spiders* sono banali : vertici di grado 2 con fase $\alpha = 0$ possono essere rimossi o viceversa introdotti in un diagramma senza che la sua logica sia alterata. **Esempio**

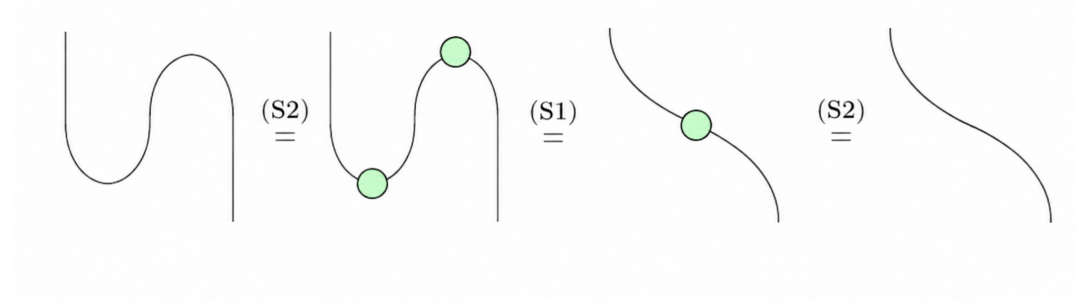
Se consideriamo il vertice $\mathcal{Z} : 2 \longrightarrow 1$ come un'operazione binaria, la regola *S1* ci dice che l'operazione è associativa:



Meno ovvio, la regola $S1$ implica che questa operazione è anche commutativa :



Esempio : Dalle S-Rules possiamo ricavare facilmente la regola $T2$ e similmente anche la regola $T1$

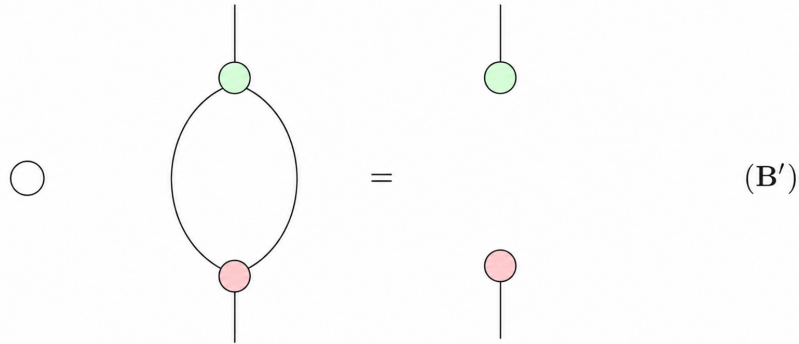


Matematicamente le S-Rules stabiliscono che ogni famiglia di vertici colorati forma una speciale algebra commutativa di Frobenius.

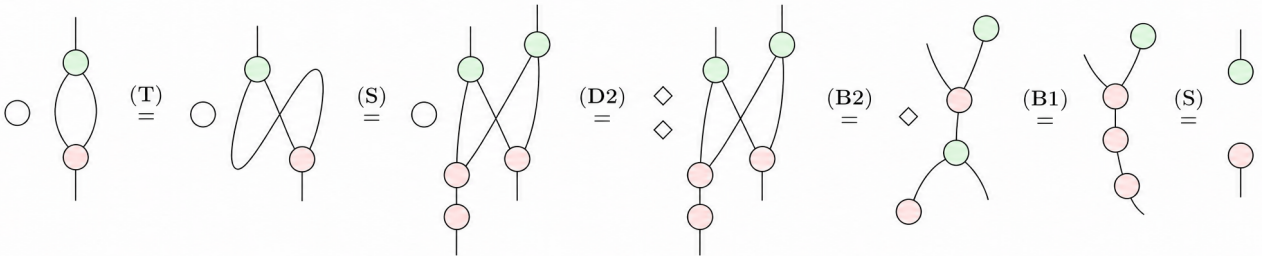
B-Rules

La regola $B1$ può essere riassunta come "il verde copia i punti rossi" e "il rosso copia i punti verdi" in entrambi i casi a meno di un diamante.

La regola $B2$ è un potente principio di commutazione e genera un'intera famiglia di equazioni, permettendo di sostituire cicli alternati di vertici rossi e vertici verdi con diagrammi più semplici. **Esempio:** un'importante equazione derivabile dalle regole B è



che si ottiene come segue:



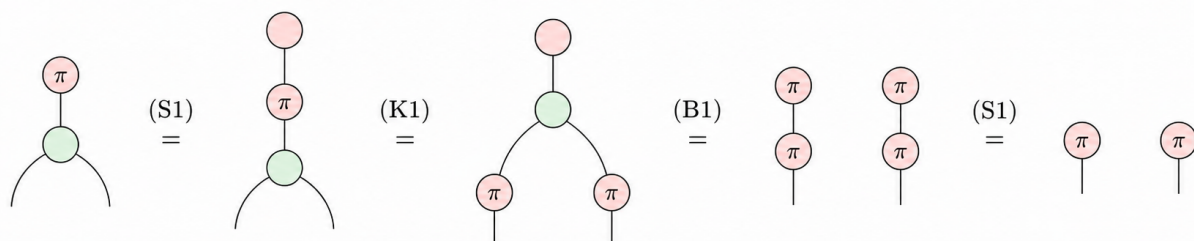
Le regole (B') e (B2) sono conosciute in modo informale come regole di Hopf e regole di bialgebra. Insieme le B-Rules ci dicono che l'interazione di diversi spiders colorati produce una struttura di bialgebra riscalata, che differisce da una bialgebra solo per un fattore di normalizzazione.

Il fatto che queste strutture algebriche emergano naturalmente dalle regole sugli ZX-diagrams è proprio il punto principale del nostro testo e rimarcheremo l'emergere di queste proprietà nel caso specifico del qubit.

K-Rules

Queste regole riguardano un particolare tipo di spiders con fase $\alpha = \pi$. La regola $K1$ dice che i vertici indicizzati α commutano con i vertici di altri colori, che equivale a dire che $X_1^1(\pi)$ è un omomorfismo della comoltiplicazione $Z_1^1(\pi)$ e viceversa.

Esempio: Grazie alla regola $K1$ i vertici con fase π possono essere copiati proprio come



Dato che i vertici con fase π o 0 possono essere copiati, prendono il nome di *punti classici* e quindi di conseguenza $Z_1^1(\pi)$ e $X_1^1(\pi)$ le chiameremo *mappe classiche* (vedremo che nell'interpretazione degli ZX-diagrams come diagrammi dei qubit queste mappe saranno la corretta interpretazione delle porte logiche già introdotte come matrici di Pauli).

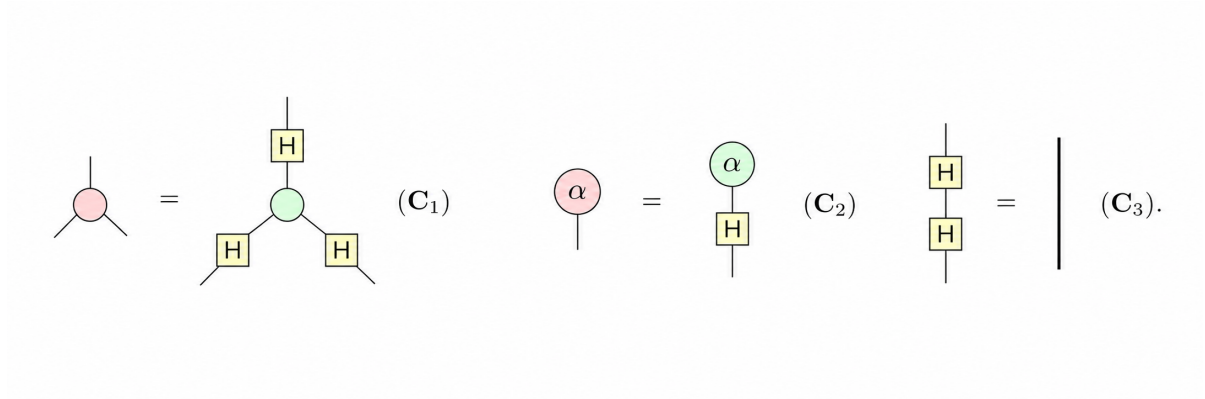
Esempio: Le regole $S1$ e $S2$ ci dicono che lo spider di grado 2 $Z_1^1(\alpha)$ forma un gruppo abeliano e dalla regola $K2$ possiamo dire che il coniugato della mappa $X_1^1(\pi)$ manda ogni elemento nel suo inverso.

C-Rule

Questa regola autorizza il vertice H ad operare come un implicito cambio di colore dei vertici.

(Nella prossima sezione vedremo che la corretta interpretazione di questo vertice sul qubit sarà la porta Hadamard).

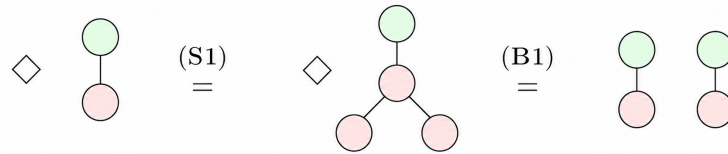
Esempio: Un caso particolare della C-rule è il seguente



D-rules

La regola *D2* ci dice che due diamanti sono equivalenti ad un filo circolare , un loop, ovvero una composizione di un cappello ed una tazza. (Vedremo sempre sui qubit che il doppio diamante rappresenterà la dimensione dello spazio di Hilbert).

La regola *D1* segue invece interamente dalle regole precedenti:



3.3 Interpretazione dello ZX-Calculus nello spazio di Hilbert

Seguendo la notazione di [1], indichiamo con $\mathcal{Q} := \mathbb{C}^2$ lo spazio degli stati di un singolo qubit. La base computazionale, o base *Z*, è data dai vettori

$$|0\rangle, |1\rangle,$$

mentre la base *X* è data da

$$|+\rangle = \frac{1}{\sqrt{2}}(|0\rangle + |1\rangle), \quad |-\rangle = \frac{1}{\sqrt{2}}(|0\rangle - |1\rangle).$$

Di conseguenza, un sistema formato da n qubit è descritto dal prodotto tensoriale di n copie di $\mathcal{Q}$, cioè

$$\mathcal{Q}^{\otimes n} = \underbrace{\mathcal{Q} \otimes \cdots \otimes \mathcal{Q}}_{n \text{ volte}}.$$

Per semplicità di notazione, nello ZX-Calculus si scrive spesso $\mathcal{Q}^n$ al posto di $\mathcal{Q}^{\otimes n}$. Un diagramma $\mathcal{D}$ con n fili in ingresso e m fili in uscita viene quindi interpretato come una mappa lineare tra spazi di Hilbert:

$$\mathcal{D} : \mathcal{Q}^{\otimes n} \longrightarrow \mathcal{Q}^{\otimes m}.$$

In questo senso, gli ingressi del diagramma rappresentano i qubit su cui l'operazione agisce, mentre le uscite rappresentano i qubit prodotti dall'operazione.

L'interpretazione di un diagramma qualsiasi si ottiene assegnando innanzitutto una mappa lineare a ciascun generatore elementare del calcolo. Successivamente, l'interpretazione dei diagrammi composti è determinata in modo naturale dalle due operazioni fondamentali sui diagrammi: la giustapposizione orizzontale corrisponde al prodotto tensoriale di mappe lineari, mentre il collegamento delle uscite di un diagramma con gli ingressi di un altro corrisponde alla composizione di mappe lineari.

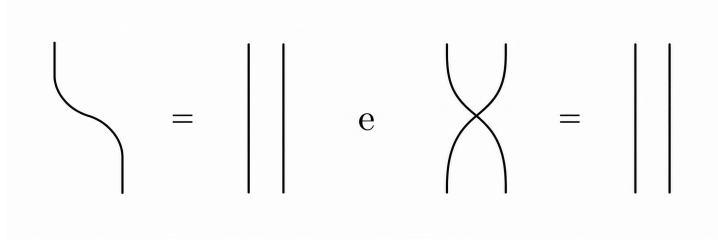
Questa costruzione fornisce dunque una semantica nello spazio di Hilbert per gli ZX-diagrams: ogni diagramma non è soltanto un oggetto grafico, ma rappresenta una precisa applicazione lineare tra spazi di stati quantistici.

Dato quindi un diagramma $\mathcal{D} : n \longrightarrow m$ costruiamo una corrispondente mappa $\mathcal{D} : \mathcal{Q}^n \longrightarrow \mathcal{Q}^m$ come segue: se il diagramma $\mathcal{D}$ è costituito da un solo generatore che è uno tra 1_Q , σ_Q , η_Q , ϵ_Q , $Z_n^m(\alpha)$, $X_n^m(\alpha)$, H , o $\sqrt{D}$, allora le sue corrispondenti mappe lineari sono

$$\begin{array}{c}
\left| \right. = \begin{pmatrix} 1 & 0 \\ 0 & 1 \end{pmatrix} \quad \times = \begin{pmatrix} 1 & 0 & 0 & 0 \\ 0 & 0 & 1 & 0 \\ 0 & 1 & 0 & 0 \\ 0 & 0 & 0 & 1 \end{pmatrix} \\
\cup = |00\rangle + |11\rangle \quad \cap = \langle 00| + \langle 11| \\
\begin{array}{c} \text{diagram with green circle } \alpha \\ \text{diagram with red circle } \alpha \end{array} : \begin{cases} |0 \dots 0\rangle \mapsto |0 \dots 0\rangle \\ |1 \dots 1\rangle \mapsto e^{i\alpha} |1 \dots 1\rangle \\ \text{others} \mapsto 0 \end{cases} \\
\begin{array}{c} \text{diagram with red circle } \alpha \\ \text{diagram with red circle } \alpha \end{array} : \begin{cases} |+\dots+\rangle \mapsto |+\dots+\rangle \\ |-\dots-\rangle \mapsto e^{i\alpha} |-\dots-\rangle \\ \text{others} \mapsto 0 \end{cases} \\
\boxed{H} = \frac{1}{\sqrt{2}} \begin{pmatrix} 1 & 1 \\ 1 & -1 \end{pmatrix} \quad \diamond = \sqrt{2} \\
\begin{array}{c} \text{diagram with red circle } \alpha \end{array} : \mathcal{H}_n \rightarrow \mathcal{H}_m, \quad \begin{cases} |+\dots+\rangle \mapsto |+\dots+\rangle \\ |-\dots-\rangle \mapsto e^{i\alpha} |-\dots-\rangle \\ \text{others} \mapsto 0 \end{cases}
\end{array}$$

Gli ZX-diagrams erano stati definiti come fili che attraversano un'interfaccia fissa sulla quale venivano numerati da sinistra verso destra gli input e gli output. Nel caso dell'interpretazione degli ZX-diagrams nello spazio di Hilbert non importerà in che ordine vengano in output i qubit e per questo definiamo:

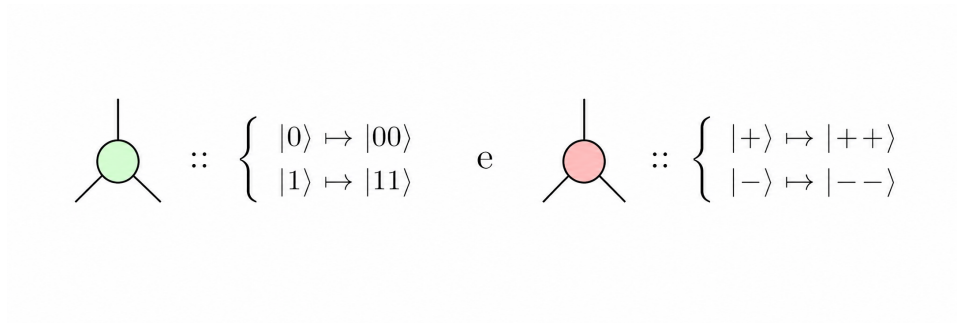
$$\text{id}_{\mathcal{H}} = \begin{pmatrix} 1 & 0 \\ 0 & 1 \end{pmatrix} = \left| \right|, \quad \text{id}_{\mathcal{H}^2} = \begin{pmatrix} 1 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 \\ 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 1 \end{pmatrix} = \left| \right| \left| \right|$$



Esempio: i generatori $Z_1^1(\pi)$, $X_1^1(\pi)$ sono proprio le matrici di Pauli Z e X

$$\text{green circle } \pi = \begin{pmatrix} 1 & 0 \\ 0 & -1 \end{pmatrix} \quad \text{red circle } \pi = \begin{pmatrix} 0 & 1 \\ 1 & 0 \end{pmatrix}$$

Esempio: I generatori Z_2^1 e X_2^1 sono interpretati come segue:



fornendo rispettivamente le mappe che copiano i vettori della base Z e i vettori della base X .

Esempio: considerando la mappa Z_1^0 , notiamo che la sua corrispondente mappa lineare manda $1 \mapsto |0\rangle$ e $1 \mapsto |0\rangle$, quindi per linearità $1 \mapsto |0\rangle + |1\rangle = \sqrt{2}|+\rangle$.
Il set completo delle basi Z e X sarà

$$\begin{array}{c} \text{green circle} \\ \downarrow \end{array} = \sqrt{2} |+\rangle, \quad \begin{array}{c} \text{green circle with } \pi \\ \downarrow \end{array} = \sqrt{2} |-\rangle, \quad \begin{array}{c} \text{red circle} \\ \downarrow \end{array} = \sqrt{2} |0\rangle, \quad \begin{array}{c} \text{red circle with } \pi \\ \downarrow \end{array} = \sqrt{2} |1\rangle.$$

Interpretazione dei diagrammi composti

Se $\mathfrak{D}$ consiste di diversi generatori ci sono due casi:

- se $\mathfrak{D} = \mathfrak{D}_1 \otimes \mathfrak{D}_2$ allora $\mathcal{D} = \mathcal{D}_1 \otimes \mathcal{D}_2$;
- se $\mathfrak{D} = \mathfrak{D}_1 \circ \mathfrak{D}_2$, allora $\mathcal{D} = \mathcal{D}_1 \circ \mathcal{D}_2$.

L'ordine in cui dividiamo il diagramma in pezzi non ha importanza per il risultato finale, a patto che i “tagli” non passino attraverso alcun vertice, né alcun punto in cui i fili si incrociano, né alcun punto di flesso di un filo — più precisamente: solo quei punti di flesso in cui il gradiente del filo cambia segno. (Questi ultimi due possono essere pensati come i “vertici” che definiscono σ_Q , e η_Q e ϵ_Q rispettivamente.)

Esempio: Il seguente diagramma può essere diviso come segue

The diagram shows a sequence of equalities for a diagram with four vertical lines. The first line has a red circle labeled α , the second a green circle labeled π , the third a red circle, and the fourth a red circle labeled β . The diagram is split by dashed vertical lines. The first equality shows it as a tensor product of three diagrams: a crossing of the first two lines, a crossing of the last two lines, and a vertical line with a red circle. The second equality shows it as a composition of three diagrams: a crossing of the first two lines, a crossing of the last two lines, and a vertical line with a red circle.

Figura 3.1: [1]

da cui si ottiene la seguente applicazione lineare

$$D = \left(\begin{pmatrix} 1 & 0 & 0 & 0 \\ 0 & 0 & 1 & 0 \\ 0 & 1 & 0 & 0 \\ 0 & 0 & 0 & 1 \end{pmatrix} \left(e^{-i\frac{\alpha}{2}} \begin{pmatrix} \cos \frac{\alpha}{2} & i \sin \frac{\alpha}{2} \\ i \sin \frac{\alpha}{2} & \cos \frac{\alpha}{2} \end{pmatrix} \otimes \begin{pmatrix} 1 & 0 \\ 0 & -1 \end{pmatrix} \right) \right)$$

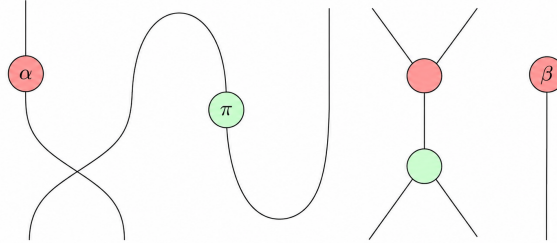
$$\otimes \frac{1}{\sqrt{2}} \left(\begin{pmatrix} 1 & 0 \\ 0 & 0 \\ 0 & 0 \\ 0 & 1 \end{pmatrix} \begin{pmatrix} 1 & 0 & 0 & 1 \\ 0 & 1 & 1 & 0 \end{pmatrix} \right) \otimes e^{-i\frac{\beta}{2}} \begin{pmatrix} i \sin \frac{\beta}{2} \\ \cos \frac{\beta}{2} \end{pmatrix}.$$

A differenza del diagramma, la matrice associata risulta di dimensioni piuttosto elevate (16×32), motivo per cui non viene riportata. Un'altra possibile fattorizzazione del diagramma è, ad esempio,

Figura 3.2: [1]

che riproduce la stessa interpretazione.

In accordo con ciò che ci dice la T-Rule inoltre, il diagramma che abbiamo appena riportato dovrebbe essere equivalente al seguente

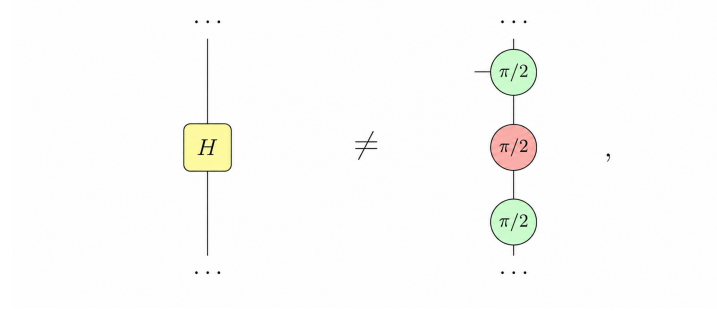


e portarci anch'esso alla stessa interpretazione.

Osservazione Una mappa lineare $f : \mathbb{C} \rightarrow \mathbb{C}$ è completamente determinata dal valore $f(1)$. Per questa ragione, e poiché $\mathcal{Q}^0 = \mathbb{C}$, l'interpretazione nello spazio di Hilbert di un diagramma senza input o output — una mappa da $\mathbb{C}$ in se stesso — è semplicemente un numero complesso.

Proposizione 4.1 (Correttezza). Se i diagrammi $\mathcal{D}_1$ e $\mathcal{D}_2$ sono uguali secondo le regole equazionali dello ZX-Calculus, allora $\mathcal{D}_1 = \mathcal{D}_2$.

Il viceversa della Proposizione è falso: esistono diagrammi che rappresentano la stessa mappa lineare ma che non sono uguali secondo le regole del calcolo ZX. Per esempio, i seguenti diagrammi non sono equivalenti nel calcolo:



La loro interpretazione come applicazioni lineari corrisponde tuttavia alla decomposizione secondo gli angoli di Eulero:

$$H = Z_1^1\left(\frac{\pi}{2}\right) \circ X_1^1\left(\frac{\pi}{2}\right) \circ Z_1^1\left(-\frac{\pi}{2}\right).$$

Vale la pena soffermarsi su questo risultato per due motivi. Innanzitutto, esso evidenzia che non tutte le proprietà valide nello spazio di Hilbert della meccanica quantistica possono essere derivate mediante lo ZX-Calculus, sebbene una grande quantità di equazioni utilizzate nell'informazione quantistica possa essere ottenuta in questo modo.

In secondo luogo, poiché la teoria equazionale dello ZX-Calculus è strettamente più debole di quella degli spazi di Hilbert, essa risulta più generale. Esistono pertanto modelli dello ZX-Calculus che sono distinti dall'usuale interpretazione della meccanica quantistica negli spazi di Hilbert. Tutti questi modelli descrivono un'ampia parte della meccanica quantistica, vista come teoria equazionale, ma risultati come la proporzione precedente non sono necessariamente validi.

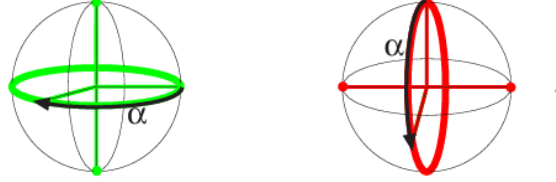
I punti nello ZX-Calculus non sono normalizzati. Questa scelta è motivata da ragioni di semplicità: se si normalizzassero σ_Q e η_Q , allora la regola (T1) richiederebbe ulteriori fattori scalari e, di conseguenza, lo stesso accadrebbe per la regola (S1) e per le altre regole del calcolo.

3.3.1 Universalità dello ZX-Calculus

Affermiamo quindi ora di avere sufficiente potere espressivo per scrivere qualsiasi mappa lineare arbitraria da n qubit a m qubit. Le fasi verde e rossa, rispettivamente

$$\begin{aligned} \text{green circle } \alpha &= Z_1^1(\alpha) = \begin{pmatrix} 1 & 0 \\ 0 & e^{i\alpha} \end{pmatrix} \\ \text{red circle } \alpha &= X_1^1(\alpha) = e^{-i\alpha/2} \begin{pmatrix} \cos \frac{\alpha}{2} & i \sin \frac{\alpha}{2} \\ i \sin \frac{\alpha}{2} & \cos \frac{\alpha}{2} \end{pmatrix} \end{aligned}$$

corrispondono rispettivamente a una rotazione di un angolo α attorno all'asse Z e X della sfera di Bloch



Combinando allora vertici rossi e vertici verdi siamo in grado di scrivere qualsiasi trasformazione unitaria su singolo qubit in termini della sua scomposizione in angoli di Eulero sulla sfera di Bloch

$$\begin{array}{c} \alpha \\ \beta \\ \gamma \end{array} = Z_1^1(\gamma) \circ X_1^1(\beta) \circ Z_1^1(\alpha)$$

La porta **CNOT** è definita come

$$\begin{array}{c} \text{green} \\ \text{red} \end{array} := \begin{array}{c} \text{green} \\ \text{red} \end{array} = \begin{array}{c} \text{green} \\ \text{red} \end{array} = \begin{pmatrix} 1 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 \\ 0 & 0 & 0 & 1 \\ 0 & 0 & 1 & 0 \end{pmatrix}$$

I risultati classici dell'informatica quantistica [8] affermano che la porta ΛX e le trasformazioni unitarie arbitrarie a un qubit sono sufficienti per costruire qualsiasi operatore unitario su n qubit. Lo ZX-Calculus contiene questo insieme universale di porte e, di conseguenza, è in grado di rappresentare qualsiasi operatore unitario su n qubit.

Proposizione 4.2 Sia $A: \mathcal{Q}^n \longrightarrow \mathcal{Q}^m$ un'applicazione lineare. Allora esiste uno ZX-diagram $\mathcal{D}$ la cui interpretazione nello spazio di Hilbert coincide con A .

Osservazione Poiché il viceversa della proposizione non è, in generale, valido, non vi è alcuna ragione per cui lo ZX-diagram $\mathcal{D}$ debba essere unico. Possono infatti esistere più ZX-diagrammi tra loro non equivalenti che rappresentano la stessa applicazione lineare. In ogni caso anche se maggiormente più complesso esiste un metodo formale per passare

da uno ZX-diagram al suo equivalente circuito quantistico, per questo rimandiamo a [2]. Ciò che interessa noi è che se esiste un modo formale di passare dagli ZX-diagrams ai circuiti quantistici e viceversa ed esiste un modo di passare da un diagramma equivalente all'altro allora questo significa che le regole dello ZX-Calculus sono complete e che le possiamo utilizzare per studiare problemi fondamentali come l'ottimizzazione o la correzione degli errori su un circuito quantistico [3].

È importante osservare però che quindi lo ZX-calculus non descrive esclusivamente trasformazioni unitarie. Mentre nel modello circuitale le porte quantistiche che rappresentano l'evoluzione chiusa di un sistema sono operatori unitari $U : H^{\otimes n} \rightarrow H^{\otimes n}$, lo ZX-calculus rappresenta oggetti più generali, processi quantistici lineari espressi dalle mappe $H^{\otimes n} \rightarrow H^{\otimes m}$. Per questo motivo sono ammessi diagrammi con un numero diverso di fili in ingresso e in uscita. Tali diagrammi possono rappresentare, ad esempio, preparazioni di stati, misure, proiezioni o componenti algebriche elementari, come gli *spider*, che non sono necessariamente unitari presi singolarmente.

Le porte quantistiche unitarie costituiscono quindi una sottoclasse dei processi rappresentabili nello ZX-calculus. Quando un diagramma rappresenta un circuito quantistico unitario, esso avrà lo stesso numero di input e output e la mappa lineare associata sarà unitaria, ma la potenza dello ZX-Calculus risiede proprio nella possibilità di scrivere attraverso i suoi diagrammi una classe più ampia di oggetti quantistici così da incorporare nella rappresentazione dei circuiti il comportamento quantistico di questi stessi, che è proprio quello che ci permettere di semplificarli rispettandone la logica.

3.4 Algebre di Hopf nello ZX-Calculus

Scopo ultimo di questo capitolo vuole essere formalizzare la struttura di algebra di Hopf che emerge dallo ZX-Calculus per gli ZX-diagrams.

- i) Una conseguenza delle regole T ed S è che segue naturalmente che lo spazio $\mathcal{H}$ acquista due strutture di algebra e di coalgebra.

Corollario 4.1. Lo spazio a 1 qubit $\mathcal{H}$ è un'algebra, dove

$$m_{\bullet} = \text{diagramma spider rosso} \quad e \quad m_{\bullet} = \text{diagramma spider verde}$$

sono operazioni di moltiplicazione associative con operazioni unità

$$\eta_{\bullet} = \text{diagramma spider rosso} \quad e \quad \eta_{\bullet} = \text{diagramma spider verde},$$

rispettivamente.

Inoltre, $\mathcal{H}$ è una coalgebra, dove

$\Delta_{\bullet} :=$ 
 $e \quad \Delta_{\bullet} :=$ 

sono operazioni di comoltiplicazione coassociative con operazioni counità

$$\varepsilon_{\bullet} := \text{green node with one edge} \quad e \quad \varepsilon_{\bullet} := \text{red node with one edge}$$

rispettivamente.

I corrispettivi assiomi di algebre e coalgebre, che riportiamo in appendice, sono infatti casi speciali delle S-Rules.

- ii) Conseguenza delle B-Rules invece è che $(\mathcal{H}, m_{\bullet}, \eta_{\bullet}, \Delta_{\bullet}, \varepsilon_{\bullet})$ acquista una struttura di algebra di Hopf non normalizzata con antipodo banale.

Dimostrazione. Il primo diagramma precedente afferma che $\Delta_\bullet$ è unitale rispetto a $\eta_\bullet$, a meno della costante di normalizzazione $\diamond$, mentre il secondo diagramma afferma che $m_\bullet$ è counitale rispetto a $\varepsilon_\bullet$, sempre a meno della stessa costante di normalizzazione $\diamond$.

Il terzo diagramma afferma che $\Delta_\bullet$ rispetta il prodotto $m_\bullet$, a meno della costante $\diamond$. In questa situazione, l'assioma dell'antipodo rispetto all'antipodo banale è automaticamente soddisfatto, ottenendo così un'algebra di Hopf non normalizzata. $\square$

Si può mostrare che anche $\Delta_\bullet$ rispetta il prodotto $m_\bullet$, a meno della stessa costante di normalizzazione $\diamond$. Ciò implica che anche

$$(m_{\bullet}, \eta_{\bullet}, \Delta_{\bullet}, \varepsilon_{\bullet})$$

è un'algebra di Hopf non normalizzata.

Una domanda naturale è allora chiedersi se

$$(m_{\bullet}, \eta_{\bullet}, \Delta_{\bullet}, \varepsilon_{\bullet}) \quad \text{oppure} \quad (m_{\bullet}, \eta_{\bullet}, \Delta_{\bullet}, \varepsilon_{\bullet})$$

siano algebre di Hopf non normalizzate. Questo non è vero in generale: infatti accade soltanto in situazioni banali. Invece, l'algebra e la coalgebra verde, oppure

rossa, formano insieme una F -algebra, con la relativa condizione di compatibilità. Insieme alla proposizione precedente, questi dati definiscono una F -algebra di Hopf.

Proposizione La coppia

$$(A, m_{\bullet}, \eta_{\bullet}, \Delta_{\bullet}, \varepsilon_{\bullet}), \quad (A, m_{\bullet}, \eta_{\bullet}, \Delta_{\bullet}, \varepsilon_{\bullet})$$

di F -algebre forma una F -algebra di Hopf.

Capitolo 4

Appendice: Algebre di Hopf non normalizzate e F -algebre di Hopf

In questa appendice richiamiamo alcune nozioni formali riguardanti le algebre di Hopf non normalizzate e le F -algebre nella loro forma grafica più vicina alla trattazione dello Z-X Calculus. Per una trattazione più completa si rimanda ai testi [4], [6] e [7], noi riportiamo le definizioni di [3].

Algebre di Hopf non normalizzate

Richiamiamo la nozione di algebra di Hopf non normalizzata su un campo $\mathbb{C}$. Gli assiomi saranno espressi mediante un linguaggio grafico, utilizzando diagrammi a fili (*string diagrams*) da leggere dal basso verso l'alto.

Un'algebra di Hopf non normalizzata è uno spazio vettoriale A su $\mathbb{C}$, che per noi sarà sempre il campo dei complessi, munito di

i.) una moltiplicazione e unità

$$m_{\bullet} := \begin{array}{c} \diagup \quad \diagdown \\ \bullet \\ | \end{array} : A \otimes A \rightarrow A \qquad \eta_{\bullet} := \begin{array}{c} \bullet \\ | \end{array} : \mathbb{C} \rightarrow A$$

soddisfando le seguenti proprietà:

$$\begin{array}{c} \diagup \quad \diagdown \\ \bullet \\ \diagdown \quad \diagup \\ \bullet \\ | \end{array} = \begin{array}{c} \diagup \quad \diagdown \\ \diagdown \quad \diagup \\ \bullet \\ | \end{array} \qquad \text{Associatività}$$

e l'

$$\text{Diagram} = \text{Line} = \text{Diagram} \quad \text{Unitalità}$$

ii.) una comoltiplicazione e counità

$$\Delta_{\bullet} := \text{Diagram} : A \rightarrow A \otimes A \quad \epsilon_{\bullet} := \text{Diagram} : A \rightarrow \mathbb{C}$$

che soddisfano le seguenti proprietà:

$$\text{Diagram} = \text{Diagram} \quad \text{Coassociatività}$$

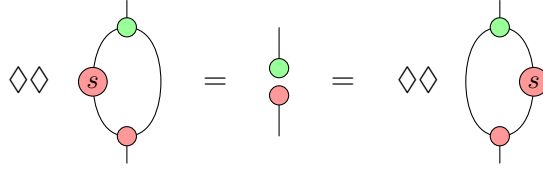
e

$$\text{Diagram} = \text{Line} = \text{Diagram} \quad \text{Counitalità}$$

iii.) una costante di normalizzazione $\diamond \in \mathbb{C} \setminus \{0\}$ tale che Δ e ϵ sono morfismi di algebra non normalizzati

$$\begin{aligned} \diamond \cdot \text{Diagram} &= \text{Diagram} & \diamond \cdot \text{Diagram} &= \text{Diagram} \\ \text{Diagram} &= \diamond & \diamond \cdot \text{Diagram} &= \text{Diagram} \end{aligned}$$

iv.) e un'applicazione lineare $S := s : A \rightarrow A$ che soddisfa gli assiomi dell'antipodo



La nozione usuale di algebra di Hopf normalizzata si ottiene ponendo

$$\diamond = 1.$$

È facile verificare che, data un'algebra di Hopf non normalizzata

$$(A, m_{\bullet}, \eta_{\bullet}, \Delta_{\bullet}, \varepsilon_{\bullet}, S),$$

si ottiene un'algebra di Hopf normalizzata

$$(A, m_{\bullet}, \eta_{\bullet}, \Delta'_{\bullet}, \varepsilon'_{\bullet}, S')$$

ponendo

$$\Delta'_{\bullet} := \diamond^{-1} \Delta_{\bullet}, \quad \varepsilon'_{\bullet} := \diamond^{-1} \varepsilon_{\bullet}, \quad S' := \diamond^{-2} S.$$

Esempio Sia G un gruppo con elemento neutro $e \in G$. Allora l'algebra di gruppo $\mathbb{C}G$, cioè lo spazio vettoriale complesso avente gli elementi di G come base e con prodotto ottenuto estendendo per bilinearità il prodotto del gruppo, diventa un'algebra di Hopf normalizzata con struttura di algebra

$$m_{\bullet}(g \otimes h) := gh, \quad \eta_{\bullet}(1_{\mathbb{C}}) := e,$$

struttura di coalgebra

$$\Delta_{\bullet}(g) := g \otimes g, \quad \varepsilon_{\bullet}(g) := 1,$$

e antipodo

$$S(g) := g^{-1},$$

definiti per ogni $g, h \in G$ ed estesi per linearità a $\mathbb{C}G$.

F-algebre di Hopf

Uno spazio vettoriale A si dice *F-algebra* se è dotato di una struttura di algebra associativa e unitaria

$$(m_{\bullet}, \eta_{\bullet})$$

su A e di una struttura di coalgebra coassociativa e counitaria

$$(\Delta_{\bullet}, \varepsilon_{\bullet})$$

su A , tali che valga la legge di Frobenius

$$(\text{id} \otimes m_{\bullet})(\Delta_{\bullet} \otimes \text{id}) = \Delta_{\bullet} m_{\bullet} = (m_{\bullet} \otimes \text{id})(\text{id} \otimes \Delta_{\bullet}).$$

Una coppia $(A, m_{\bullet}, \eta_{\bullet}, \Delta_{\bullet}, \varepsilon_{\bullet}), (A, m_{\bullet}, \eta_{\bullet}, \Delta_{\bullet}, \varepsilon_{\bullet})$ di F -algebre sullo stesso spazio vettoriale

A si dice F -bialgebra se

$$(A, m_{\bullet}, \eta_{\bullet}, \Delta_{\bullet}, \varepsilon_{\bullet}), \quad (A, m_{\bullet}, \eta_{\bullet}, \Delta_{\bullet}, \varepsilon_{\bullet})$$

sono bialgebre. La coppia si dice F -algebra di Hopf se queste ultime sono algebre di Hopf.

Esempio Sia G un gruppo finito con elemento neutro $e \in G$. Allora $A = \mathbb{C}G$ è un' F -algebra di Hopf con

$$m_{\bullet}(g \otimes h) := gh, \quad \eta_{\bullet}(1_{\mathbb{C}}) := e, \quad \Delta_{\bullet}(g) := \sum_{\substack{h, l \in G \\ hl = g}} h \otimes l, \quad \varepsilon_{\bullet}(g) := \delta_{e, g},$$

e

$$m_{\bullet}(g \otimes h) := \delta_{g, h} g, \quad \eta_{\bullet}(1_{\mathbb{C}}) := \sum_{g \in G} g, \quad \Delta_{\bullet}(g) := g \otimes g, \quad \varepsilon_{\bullet}(g) := 1.$$

Infatti,

$$(A, m_{\bullet}, \eta_{\bullet}, \Delta_{\bullet}, \varepsilon_{\bullet})$$

è l'algebra di Hopf dell'esempio precedente con antipodo

$$S(g) = g^{-1},$$

mentre

$$(A, m_{\bullet}, \eta_{\bullet}, \Delta_{\bullet}, \varepsilon_{\bullet})$$

è una bialgebra. Infatti, utilizzando la notazione

$$m_{\bullet}(g \otimes h) := g \bullet h,$$

si ha

$$\begin{aligned} \Delta_{\bullet}(g) \bullet \Delta_{\bullet}(h) &= \sum_{\substack{k, k', n, n' \in G \\ kk' = g, n'n = h}} k \bullet k' \otimes n \bullet n' \\ &= \sum_{\substack{k, l \in G \\ kl = g}} k \otimes l \\ &= \Delta_{\bullet}(g \bullet h), \end{aligned}$$

e inoltre

$$\varepsilon_{\bullet}(g) \varepsilon_{\bullet}(h) = \delta_{e, g} \delta_{e, h} = \delta_{e, g} \delta_{g, h} = \varepsilon_{\bullet}(g \bullet h).$$

Pertanto anche quest'ultima è dotata dell'antipodo

$$S(g) = g^{-1}.$$

Conclusioni

In questo lavoro di tesi si sono ripercorsi i risultati di costruzione delle algebre di Hopf e dello ZX-Calculus. In particolare nel primo capitolo si è sviluppato interamente il formalismo delle algebre di Hopf a partire dai concetti di struttura algebrica e coalgebrica. Riprendendo l'esempio iniziale del capitolo, costruito sul gruppo di matrici $SL(N, \mathbb{C})$, una algebra di Hopf si può pensare in modo semplice come una struttura che estende in modo naturale il concetto di gruppo all'ambito delle algebre, introducendo operazioni compatibili di natura algebrica e coalgebrica. Usando la struttura di gruppo sugli elementi dell'algebra infatti, che nel nostro caso erano le funzioni a valori complessi sugli elementi del gruppo stesso, si riescono ad estendere su quest'ultima le proprietà di associatività, dell'elemento neutro e dell'elemento inverso, che diventano rispettivamente le proprietà di coassociatività, counità e antipodo, che vivono insieme nella struttura algebrica di Hopf rispettando i criteri di compatibilità stabiliti dalla relazione di struttura bialgebrica che contiene in sé sia una nozione di algebra che una di coalgebra, tali per cui le proprie mappe soddisfino una relazione di omomorfismo tra le stesse.

Nel terzo capitolo abbiamo definito i primi passi per la costruzione del formalismo dello ZX-Calculus a partire dalla definizione degli ZX-diagrams. Il nostro scopo dichiarato era arrivare ad utilizzare lo ZX-Calculus per la semplificazione di circuiti quantistici sostituendo, nella rappresentazione delle operazioni sui qubit, le matrici unitarie con gli ZX-diagrams. Abbiamo quindi ricostruito le principali porte logiche, definite nel secondo capitolo della tesi, in termini degli *spiders* dello ZX-Calculus, ma la potenza di questo calcolo risiede proprio nel fatto di poter trasportare nel suo linguaggio una categoria più ampia di oggetti quantistici oltre che le operazioni logiche sui qubit, come preparazione di stati, misure, componenti algebriche elementari. Questo ci permette di semplificare la composizione dei circuiti quantistici proprio perchè ci consente di lavorare sulla logica degli oggetti e operazioni che li compongono e che rispondono alle regole della meccanica quantistica. Si sono infine evidenziate le strutture algebriche astratte che emergono in modo naturale dalle regole dello ZX-Calculus, confrontando le equazioni degli ZX-diagrams con le rappresentazioni grafiche delle algebre di Hopf non normalizzate e delle algebre di Frobenius. Nel secondo capitolo della tesi si riprendevano le nozioni storiche di qubit definito nel formalismo dei postulati della meccanica quantistica e della rappresentazione degli stati come matrici densità.

Lo ZX-calculus ha ormai raggiunto una notevole maturità teorica, testimoniata dai risultati di completezza ottenuti negli ultimi anni. Quando Coecke e Duncan pubblicarono *Interacting Quantum Observables* (2008, poi esteso nel 2011) infatti, dimostrarono che lo ZX-calculus è universale per il calcolo quantistico a qubit, ovvero che il suo linguaggio grafico è abbastanza espressivo da rappresentare qualsiasi operazione lineare tra sistemi di qubit, almeno a meno di un eventuale fattore scalare globale. Cioè ogni circuito quantistico può essere tradotto in un diagramma ZX. Ma universale non significa completo. La completezza significa una cosa molto più forte: se due circuiti rappresentano la stessa

trasformazione lineare, allora esiste una sequenza finita di regole dello ZX-calculus che permette di trasformare un diagramma nell'altro. All'inizio questo non era noto. Potevano esistere due circuiti matematicamente equivalenti senza che le regole dello ZX riuscissero a dimostrarlo. Il primo risultato di completezza è arrivato con il lavoro "*The ZX-calculus is Complete for Stabilizer Quantum Mechanics*." di Miriam Backens nel 2014 che dimostrava la completezza per la *stabilizer quantum mechanics*, cioè per il frammento della teoria quantistica generato dalle operazioni di Clifford. Questo spostava il significato dello ZX-Calculus da semplice linguaggio grafico a vero e proprio strumento di dimostrazione matematica. Nel 2017 Jeandel, Perdrix e Vilmart, con il lavoro "*A Complete Axiomatisation of the ZX-Calculus for Clifford+T Quantum Mechanics*", hanno successivamente dimostrato che lo ZX-calculus è completo anche per un frammento di circuito che non fosse un circuito Clifford ed è il risultato che si conosce come *Clifford + T* indicando con *T* la particolare porta scelta per la dimostrazione. Questo fu già un risultato enormemente importante perché riguardava praticamente tutti i circuiti realmente utilizzati nell'informatica quantistica. Nel 2020 però Jeandel, Perdrix e Vilmart sono riusciti ad eliminare anche l'ultima limitazione. Con il lavoro "*Complete Axiomatisation of the ZX-Calculus*", hanno dimostrato che lo ZX-calculus è completo per tutta la meccanica quantistica dei qubit. In altre parole questo significa che lo ZX-calculus è un linguaggio matematicamente equivalente al formalismo matriciale e qualsiasi identità tra operatori lineari sui qubit può, almeno in linea di principio, essere dimostrata esclusivamente attraverso le sue regole grafiche.

Bibliografia

- [1] B. Coecke and R. Duncan. Interacting quantum observables. In *Proceedings of the 37th International Colloquium on Automata, Languages and Programming (ICALP)*. Springer, 2008.
- [2] N. de Beaudrap, A. Kissinger, and J. van de Wetering. Circuit extraction for zx-diagrams can be p-hard. In *49th International Colloquium on Automata, Languages, and Programming (ICALP 2022)*, volume 229 of *LIPICs*, pages 119:1–119:19, 2022.
- [3] E. Ercolessi, R. Fioresi, and T. Weber. The geometry of quantum computing. *International Journal of Geometric Methods in Modern Physics*, 2024.
- [4] C. Kassel. *Quantum Groups*, volume 155 of *Graduate Texts in Mathematics*. Springer-Verlag, 1995.
- [5] A. U. Klimyk and K. Schmüdgen. *Quantum Groups and Their Representations*. Springer, Berlin, 1997.
- [6] S. Majid. *Foundations of Quantum Group Theory*. Cambridge University Press, 1995.
- [7] S. Montgomery. *Hopf Algebras and Their Actions on Rings*, volume 82 of *CBMS Regional Conference Series in Mathematics*. American Mathematical Society, Providence, RI, 1993.
- [8] M. A. Nielsen and I. L. Chuang. *Quantum Computation and Quantum Information*. Cambridge University Press, 10th anniversary edition edition, 2010.
- [9] J. Preskill. Lecture notes for physics 229: Quantum information and computation, 1998. California Institute of Technology.
- [10] T. Weber. The zx-calculus and related concepts, 2024. Bologna Quantum Journal Club.