

Marco Parmigiani



L'esperimento della doppia fenditura

Teoria classica e quantistica
a confronto

*L'esperimento più rappresentativo della meccanica quantistica:
dalle frange di interferenza della luce alla funzione d'onda,
un percorso teorico e sperimentale nel mondo dei quanti per
distinguere e capire ciò che osserviamo, calcoliamo e soprattutto
interpretiamo.*

Fisica x Appassionati

L'esperimento della doppia fenditura

Teoria classica e quantistica a confronto

MARCO PARMIGIANI

m.parmigiani@esagramma.it

Dedicato a Marina e ai miei 5 lettori.

Progetto grafico e foto in copertina di Marina Macchi.

Le figure presenti nel testo sono a cura dell'autore e hanno solo funzione esplicativa e rappresentativa. Non riproducono necessariamente apparati sperimentali reali in scala, ma servono solo a chiarire visivamente i concetti trattati.

© 2026 Marco Parmigiani

Tutti i diritti riservati

Prima edizione digitale: agosto 2026

Edizioni Esagramma

Indice

Prefazione	5
-------------------------	---

È vero che nessuno capisce la meccanica quantistica?

Parte Prima - L'esperimento della doppia fenditura secondo la teoria classica

Prima della doppia fenditura	9
---	---

- La luce tra corpuscoli e onde
- Tra classico e quantistico

In laboratorio con Young	11
---------------------------------------	----

- L'esperimento della doppia fenditura
- Osservare prima di interpretare
- La strumentazione essenziale
- La sorgente luminosa
- Come si esegue l'esperimento
- Che cosa si misura

Dal laboratorio alla teoria	14
--	----

- Che cosa significa spiegare un esperimento

L'interpretazione della teoria classica	15
--	----

- Come spiegare formalmente l'esperimento
- La definizione di onda in fisica
- L'interferenza dei campi elettrici
- Come interpretare l'esperimento

I calcoli secondo la teoria classica	18
---	----

- Applicazione alla doppia fenditura
- Il modello ondulatorio classico
- La condizione dei massimi luminosi
- La posizione delle frange sullo schermo

La luce come onda elettromagnetica	20
---	----

- Le equazioni di Maxwell
- Le onde elettromagnetiche

Parte Seconda - L'esperimento della doppia fenditura secondo la teoria quantistica

La doppia fenditura nella teoria quantistica 25

- Dalla luce come onda alla luce come quanto
- Un comportamento inatteso
- Il problema della doppia fenditura
- Dai molti cammini all'ampiezza totale

In laboratorio con Einstein 28

- L'esperimento dell'effetto fotoelettrico
- Osservare prima di interpretare
- La strumentazione essenziale
- La sorgente luminosa
- Come si esegue l'esperimento
- Che cosa si misura

Perché Einstein vinse il Nobel 31

- Quando la teoria ondulatoria fallisce
- Due grandi idee per una rivoluzione concettuale

L'interpretazione della teoria quantistica 33

- Come spiegare formalmente l'esperimento
- La funzione d'onda o ampiezza di probabilità
- L'equazione d'onda di Schrödinger
- L'interpretazione probabilistica di Born
- Feynman e la somma sui cammini
- Come è definita l'azione S

I calcoli secondo la teoria quantistica 36

- Applicazione alla doppia fenditura
- La figura alternata delle frange
- Il significato fisico dell'esperimento

Approfondimento: un unico stato quantistico 38

Il problema quantistico della misura 39

- La scomparsa dell'interferenza
- Il concetto di decoerenza

Comprendere la meccanica quantistica 40

- La funzione d'onda come campo di informazione
- Quindi aveva ragione Feynman?
- La differenza tra formalismo e interpretazione

Parte Terza - Gli esperimenti classici e quantistici messi in pratica

Una simulazione realistica degli esperimenti. 45

- Dalla teoria alla pratica: come si procede in concreto nelle misurazioni
- Gli errori di misura nel caso classico
- Perché ci limitiamo alla misura di ΔY
- La doppia fenditura nel regime dei singoli fotoni
- Gli errori di misura nel caso quantistico
- Una misura complementare: l'effetto fotoelettrico
- Verifica del calcolo non relativistico

Appendici - Sviluppo di alcuni passaggi matematici richiamati nel testo

Appendice A — Interferenza classica e intensità luminosa 53

- La somma algebrica di due onde armoniche
- Dall'ampiezza del campo all'intensità luminosa
- Le equazioni di Maxwell nel vuoto
- Le equazioni d'onda dei campi E e B
- Che cosa si intende con onda piana
- I campi $E \perp B$ e la direzione di propagazione

Appendice B — Funzione d'onda complessa e probabilità 57

- L'operatore hamiltoniano nell'equazione di Schrödinger
- Il modulo quadro della somma di due ampiezze complesse
- La differenza tra fase classica e quantistica

Appendice C — Lagrangiana, azione e integrale sui cammini 61

- La lagrangiana $\mathcal{L}$
- Il caso relativistico e i fotoni
- L'azione S
- Dall'azione alla fase quantistica
- Come si costruisce la somma sui cammini
- L'integrale sui cammini di Feynman
- Applicazione alla doppia fenditura

Appendice D — Fotoni, campo e.m. quantizzato e limite classico 65

- Apparato ottico lineare e limite classico
- I modi del campo elettromagnetico classico
- Il fotone segue gli stessi modi del campo classico

Alcune tappe della meccanica quantistica 68

Bibliografia essenziale 69

Prefazione

È vero che nessuno capisce la meccanica quantistica?

*Credo di poter dire con sicurezza che nessuno
capisce la meccanica quantistica.*

— Richard P. Feynman

La frase del famoso fisico Richard Feynman è tra le più celebri della fisica moderna. Compare nelle sue lezioni raccolte in *The Character of Physical Law* e viene spesso citata per esprimere il carattere *controintuitivo* della teoria quantistica.

Questo libretto nasce da quella provocazione. Non per sostenere che la meccanica quantistica sia incomprensibile, né per presentarla come un insieme di paradossi misteriosi, ma per seguire con attenzione uno degli esperimenti più semplici, profondi e rappresentativi della fisica classica e quantistica: *l'esperimento della doppia fenditura*.

La scelta ovviamente non è casuale. La doppia fenditura permette di osservare, in una forma essenziale ma esauriente, il passaggio da una descrizione classica della luce a una descrizione quantistica della radiazione e, più in generale, della materia.

Nella sua versione classica, l'esperimento mostra la formazione di frange chiare e scure, interpretabili mediante il modello ondulatorio della luce. Nella sua versione quantistica, eseguita con singoli *fotoni* (ma anche con particelle), lo stesso schema sperimentale rivela un aspetto più profondo: ogni rivelazione avviene come un evento localizzato, ma l'insieme di molti eventi ricostruisce progressivamente una figura di interferenza.

L'obiettivo di queste pagine è distinguere con chiarezza tre *livelli* che spesso vengono confusi.

Il primo è il risultato sperimentale: ciò che si osserva e si misura in laboratorio quale base sperimentale per verificare o confutare un'ipotesi scientifica.

Il secondo è il modello teorico: l'insieme di concetti, relazioni matematiche e teorie utilizzate per descrivere e prevedere con un modello il fenomeno fisico.

Il terzo è l'interpretazione fisica: il significato che attribuiamo al modello e ai suoi risultati.

Per questo motivo il percorso coprirà gli aspetti sperimentali, teorici e anche interpretativi.

Da una parte descriveremo il laboratorio: sorgente, fenditure, schermo, misure e frange osservate. Dall'altra la teoria: onde elettromagnetiche, interferenza, funzione d'onda, ampiezze di probabilità e problema della misura. Infine cercheremo di darne una interpretazione fisica.

La doppia fenditura è dunque un esperimento ideale da trattare: è abbastanza semplice da essere descritto in poche linee teoriche e sperimentali, ma abbastanza profondo da costringerci a chiedere che cosa significhi davvero *comprendere* un fenomeno fisico.

È in questo senso che vogliamo interpretare la frase di Feynman: essa non è un invito alla rinuncia, ma un avvertimento metodologico perché comprendere la meccanica quantistica significa anche imparare a distinguere ciò che osserviamo direttamente, ciò che calcoliamo mediante il nostro modello e soprattutto il significato fisico che attribuiamo ai risultati.

Parte Prima

L'esperimento della doppia fenditura secondo la teoria classica

Prima della doppia fenditura

La luce tra corpuscoli e onde

All'inizio dell'Ottocento la natura della luce era ancora oggetto di una discussione profonda. La fisica newtoniana aveva imposto per lungo tempo un'immagine *corpuscolare*: la luce veniva pensata come un flusso di minuscole particelle emesse dai corpi luminosi e propagate nello spazio secondo traiettorie rettilinee. Questa interpretazione era coerente con molti fenomeni noti, come la propagazione rettilinea della luce, la riflessione e, in parte, la rifrazione.

Accanto a questa concezione esisteva però una tradizione diversa. Già nel Seicento, Christiaan Huygens aveva proposto una *teoria ondulatoria* della luce, secondo la quale la propagazione luminosa poteva essere descritta come la trasmissione di un'onda. Questa idea spiegava in modo naturale alcuni fenomeni, come la riflessione e la rifrazione luminosa, difficili da ricondurre a un semplice modello corpuscolare, ma rimase a lungo in posizione minoritaria, anche per il grande peso scientifico dell'autorità di Newton.

La storia della luce, tra XVII e XIX secolo, può quindi essere letta come un confronto tra due modelli: la luce come particella e la luce come onda.

In questo contesto si inserisce il lavoro di Thomas Young. Medico, fisico e studioso di straordinaria ampiezza culturale, Young presentò alla Royal Society, nei primi anni dell'Ottocento, una serie di ricerche dedicate alla teoria della luce e dei colori. Nel 1801 tenne la Bakerian Lecture *On the Theory of Light and Colours*, nella quale introdusse in modo sistematico il principio di interferenza; nel 1803 presentò una nuova Bakerian Lecture, pubblicata nel 1804 con il titolo *Experiments and Calculations Relative to Physical Optics*.

L'idea decisiva era semplice e radicale: se la luce è un'onda, allora due *porzioni* dello stesso fascio luminoso devono potersi sovrapporre come fanno le onde sull'acqua. In alcuni punti le oscillazioni si rinforzano, producendo un massimo di luminosità; in altri si indeboliscono o si annullano, producendo un minimo. Young chiamò questo fenomeno *interferenza*.

La configurazione che oggi chiamiamo *esperimento della doppia fenditura* consiste nel far passare un fascio di luce attraverso due aperture sottili e ravvicinate, per poi osservare la distribuzione luminosa su uno schermo posto a una certa distanza. Se la luce si comportasse soltanto come un insieme di corpuscoli che attraversano una fenditura *oppure* l'altra, ci aspetteremmo di osservare due regioni luminose, corrispondenti alle due aperture.

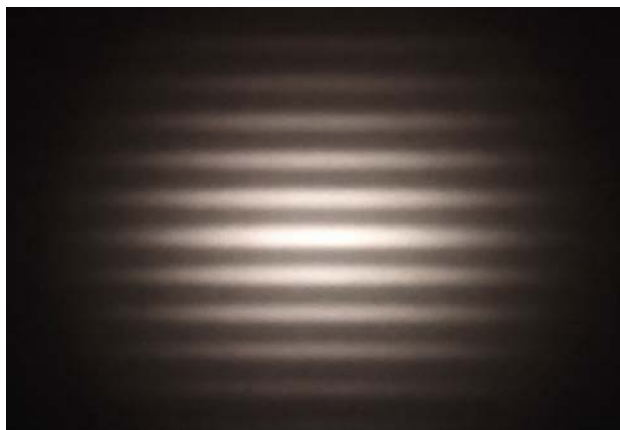
Il risultato osservato è invece molto diverso: sullo schermo compare una successione regolare di frange chiare e scure (come illustrato nelle prossime immagini).

Queste frange costituiscono il dato sperimentale fondamentale. La loro presenza indica che i contributi provenienti dalle due fenditure non si limitano a sommarsi come due immagini indipendenti, ma interferiscono. In termini ondulatori, le frange luminose corrispondono ai punti in cui le onde arrivano in fase, mentre le frange scure corrispondono ai punti in cui giungono in opposizione di fase, come spiegheremo meglio di seguito.

L'esperimento di Young ebbe quindi un'importanza storica enorme: fornì una delle prime dimostrazioni veramente convincenti della natura ondulatoria della luce e contribuì a riaprire il dibattito scientifico sulla teoria corpuscolare allora dominante.

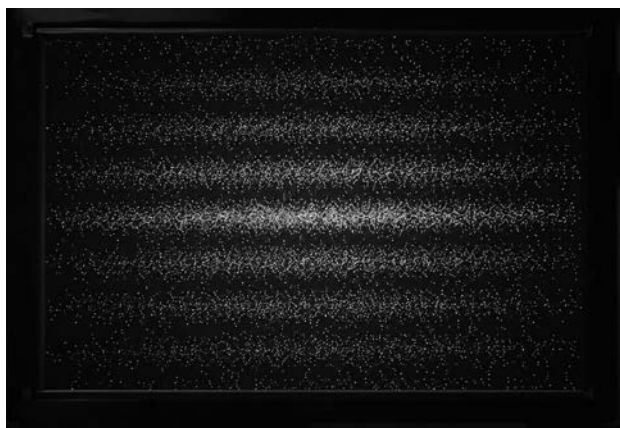
Tra classico e quantistico

L'esperimento della doppia fenditura viene spesso presentato come una prova classica del comportamento ondulatorio della luce, anche se la sua interpretazione moderna è più profonda: nel Novecento, con la meccanica quantistica, la doppia fenditura diventerà il luogo concettuale in cui emergono in modo particolarmente chiaro il *dualismo onda-particella*, il ruolo della misura e il concetto di ampiezza di probabilità (come vedremo più avanti).



Franghe di interferenza sullo schermo di rivelazione: l'alternanza di bande chiare e scure è il risultato osservabile dell'esperimento classico della doppia fenditura.

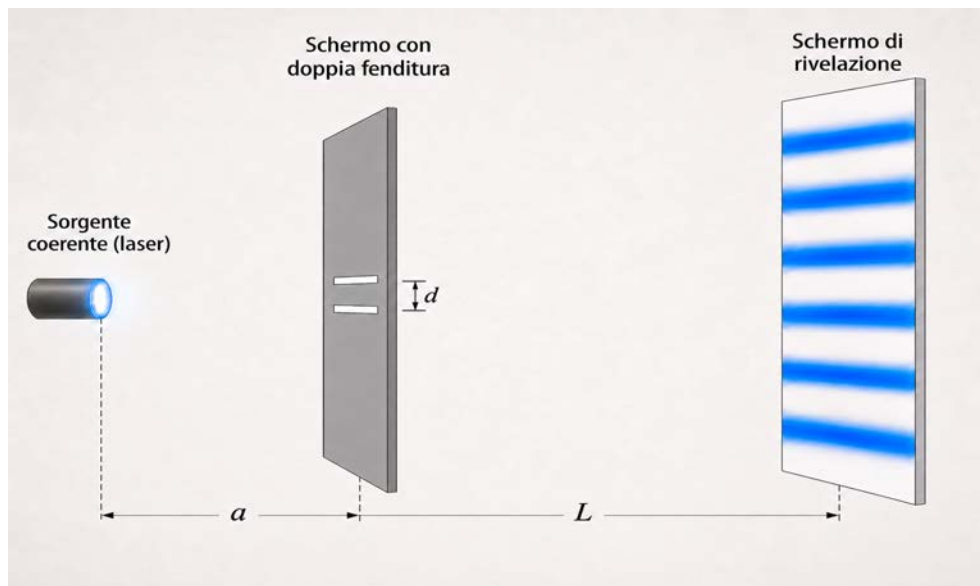
In particolare, come mostrato nella seconda immagine, se si registrano pochi impatti si vedono solo punti isolati; ma man mano che il numero degli impatti aumenta, la loro distribuzione fa emergere progressivamente le frange. Alla fine, con un numero molto grande di eventi, l'insieme dei punti ricostruisce la stessa figura continua della prima immagine.



Distribuzione di eventi quantistici sullo schermo di rivelazione: i singoli impatti luminosi sono eventi puntiformi, tuttavia la figura complessiva, ottenuta accumulando molti eventi, definisce la stessa struttura continua delle frange di interferenza mostrata sopra.

In laboratorio con Young

L'esperimento della doppia fenditura



Schema dell'apparato per l'esperimento della doppia fenditura
(lo schema indicato è fuori scala ed è esclusivamente rappresentativo).

Osservare prima di interpretare

La fisica nasce dagli esperimenti: osserva i fenomeni, li misura, cerca regolarità e costruisce leggi capaci di descrivere ciò che accade. Prima di formulare una teoria, occorre quindi stabilire con precisione che cosa viene osservato, con quali strumenti e in quali condizioni.

L'esperimento della doppia fenditura è, da questo punto di vista, un esempio particolarmente esemplare. L'apparato è semplice nella sua struttura essenziale, ma il risultato che produce è profondo: una sorgente luminosa invia un fascio di luce verso uno schermo opaco nel quale sono praticate due fenditure sottili e ravvicinate. Oltre le fenditure, a una certa distanza, si colloca uno schermo di rivelazione.

La funzione del primo schermo è quella di selezionare due *porzioni* dello stesso fascio luminoso. La funzione del secondo schermo è *registrare* la distribuzione della luce dopo il passaggio attraverso le due fenditure.

È importante ricordare che ciò che si osserva non è il *percorso della luce* nello spazio intermedio, ma la figura luminosa finale prodotta sullo schermo di rivelazione. È proprio questa particolare *distribuzione a frange* l'oggetto sperimentale che il modello teorico deve spiegare.

La strumentazione essenziale

L'apparato sperimentale può essere descritto attraverso quattro elementi principali:

La sorgente luminosa: produce il fascio incidente. Nell'esperimento moderno si usa spesso una sorgente laser, perché fornisce luce quasi monocromatica e coerente. La monocromaticità significa che la luce ha una lunghezza d'onda ben definita; la coerenza significa che le oscillazioni luminose mantengono una relazione di fase stabile, condizione necessaria per osservare frange di interferenza ben contrastate.

La doppia fenditura: è costituita da due aperture strette e molto vicine. La distanza tra i centri delle due fenditure si indica di solito con d . Questa distanza è una grandezza fondamentale, perché determina, insieme alla lunghezza d'onda λ della luce e alla distanza L dallo schermo, la separazione ΔY tra le frange osservate (come mostreremo più avanti).

Lo schermo di rivelazione: è posto a una distanza L dalle fenditure. Su di esso si forma la distribuzione finale dell'intensità luminosa. Nei laboratori didattici può essere un semplice schermo bianco; in configurazioni più sofisticate può essere sostituito da un sensore, da una lastra fotografica o da un rivelatore digitale.

Il sistema di misura: permette di determinare la posizione delle frange, la loro distanza reciproca e l'eventuale variazione della figura osservata al variare dei parametri sperimentali.

La sorgente luminosa

A seconda del *modello teorico* che utilizzeremo per interpretare l'esperimento, cioè un modello classico oppure quantistico, supporremo che la sorgente emetta rispettivamente:

Onde classiche

Nel modello classico la sorgente emette un'onda luminosa, cioè una perturbazione dei campi elettromagnetici che si propaga nello spazio e nel tempo, come spiegheremo meglio in seguito. In questa descrizione, le due fenditure si comportano come *sorgenti secondarie* e le onde che emergono da esse possono sovrapporsi, quindi rinforzarsi o cancellarsi, come previsto dal modello classico per due onde in sovrapposizione (vedi la prossima figura).

Onde quantistiche

Nel modello quantistico, invece, quando si considera l'emissione di singoli *fotoni*, non si deve immaginare necessariamente un'onda materiale che attraversa lo spazio come nel caso classico. Ciò che evolve è lo stato quantistico del fotone, descritto da una *funzione d'onda* che definisce una ampiezza di probabilità. L'interferenza non riguarda quindi due onde fisiche osservabili direttamente, ma le ampiezze associate ai diversi possibili modi in cui il fotone può contribuire al risultato finale (come descriveremo nella *Parte Seconda*).

Nota: nella *Parte Terza* vedremo più concretamente quali sono le grandezze in gioco, come la distanza tra le fenditure, la lunghezza d'onda e la distanza dallo schermo ma anche quali risultati ci si può attendere in un esperimento reale; per ora ci atteniamo alla teoria.

Come si esegue l'esperimento

Per eseguire l'esperimento si allinea anzitutto la sorgente luminosa con la doppia fenditura e con lo schermo di rivelazione. Il fascio deve incidere in modo regolare e simmetrico sulle due aperture, così che entrambe vengano correttamente illuminate. Una volta attraversate le fenditure, la luce raggiunge lo schermo finale.

Se l'allineamento è corretto e la sorgente è sufficientemente coerente (cioè con una fase stabile nel tempo), sullo schermo compare una sequenza di frange. A questo punto si possono misurare la distanza tra gli schermi, la separazione tra le fenditure e la distanza tra le frange. Una volta misurate e conoscendo la lunghezza d'onda λ della sorgente, è possibile verificare quantitativamente la relazione prevista dal modello ondulatorio.

Al fine di *mediare* possibili errori accidentali (come fluttuazioni temporanee, variazioni del cammino ottico, rumore elettronico di fondo) o cercare di *evitare* quelli sistematici (come calibrazioni degli strumenti, disallineamento geometrico, deriva termica del laser) è opportuno ripetere più volte l'esperimento anche smontando e rimontando tutto l'apparato.

Il punto essenziale, tuttavia, è concettuale: l'esperimento non mostra direttamente onde che si propagano nello spazio, né traiettorie luminose che partono da una fenditura o dall'altra. Ciò che mostra è esclusivamente una *distribuzione finale di luce* sullo schermo.

È proprio da questa distribuzione e in particolare dall'alternanza regolare di frange chiare e scure, che nasce la necessità di una spiegazione e di un modello teorico.

Che cosa si misura

L'osservabile principale dell'esperimento è la distribuzione dell'intensità luminosa sullo schermo finale. Infatti invece di due sole zone illuminate, corrispondenti alle due fenditure, compare una successione regolare di bande chiare e scure: le *frange di interferenza*.

Le frange luminose come detto sono regioni in cui l'intensità della luce è maggiore; le frange scure sono regioni in cui l'intensità è minore o, idealmente, quasi nulla. La grandezza più importante da misurare è proprio la distanza ΔY tra due frange luminose successive, oppure in modo equivalente, la distanza tra due frange scure successive.

In una configurazione ideale e, come vedremo, per piccoli angoli di osservazione rispetto all'asse centrale, la distanza ΔY tra le frange dipende dalle tre grandezze espresse nella seguente equazione, una relazione che deriveremo in un prossimo paragrafo:

$$\Delta Y \approx \frac{\lambda L}{d}$$

dove λ è la lunghezza d'onda della luce, L è la distanza tra la doppia fenditura e lo schermo di rivelazione e d è la distanza tra le due fenditure, come indicato nella precedente figura.

Questa relazione appartiene al modello teorico. In laboratorio, però, ciò che si registra direttamente è forse più semplice da descrivere: una figura alternata di massimi e minimi di luminosità. La formula serve a definire quantitativamente quella figura e quindi possibilmente a confermare o smentire il nostro modello fisico, con una serie di misure ripetute.

Dal laboratorio alla teoria

Che cosa significa spiegare un esperimento

Un esperimento non *parla* mai da solo. In laboratorio possiamo disporre strumenti, controllare condizioni, misurare grandezze e registrare risultati. Ma per comprendere il significato fisico di ciò che osserviamo è necessario costruire un modello teorico: una rappresentazione del fenomeno, formulata attraverso concetti, relazioni matematiche e ipotesi controllabili.

Nel caso della doppia fenditura, il dato sperimentale è chiaro: quando un fascio di luce attraversa due fenditure sottili e ravvicinate, sullo schermo finale compare una successione regolare di frange luminose e scure. Questo è ciò che si osserva, ma la domanda teorica è: quale modello fisico permette di spiegare questa distribuzione?

Una prima descrizione appartiene alla fisica classica. In questo quadro la luce viene trattata come un'onda elettromagnetica. L'ipotesi è che le due fenditure si comportino come due *sorgenti coerenti*: le onde che da esse si propagano si sovrappongono e producono interferenza. Dove i contributi arrivano in fase si forma una frangia luminosa; dove arrivano in opposizione di fase si forma una frangia scura (vedi la figura seguente). Il modello ondulatorio classico permette quindi di calcolare la posizione delle frange e di collegarla a grandezze misurabili come la lunghezza d'onda, la distanza tra le fenditure e la distanza dallo schermo.

Ma la doppia fenditura ha anche una seconda vita, più profonda, nella fisica del Novecento. Quando l'esperimento viene eseguito con intensità estremamente basse, fino al caso ideale di singoli fotoni, la figura di interferenza riappare progressivamente sullo schermo. Ogni evento di rivelazione è localizzato, come se fosse prodotto da una singola particella; l'insieme di molti eventi, però, ricostruisce una distribuzione classica di interferenza.

È qui che entra in gioco la descrizione quantistica.

In assenza di una misura finalizzata alla rivelazione, il formalismo quantistico non attribuisce al fotone una traiettoria classica ben definita attraverso una delle due fenditure, ma descrive il fenomeno attraverso una funzione d'onda, o meglio attraverso *ampiezze di probabilità* associate ai diversi modi in cui il sistema può evolvere. Come vedremo più avanti le ampiezze si sommano e possono interferire; mentre ciò che si osserva sullo schermo è invece un singolo evento di rivelazione, localizzato in un punto.

È per questo essenziale distinguere fra tre piani:

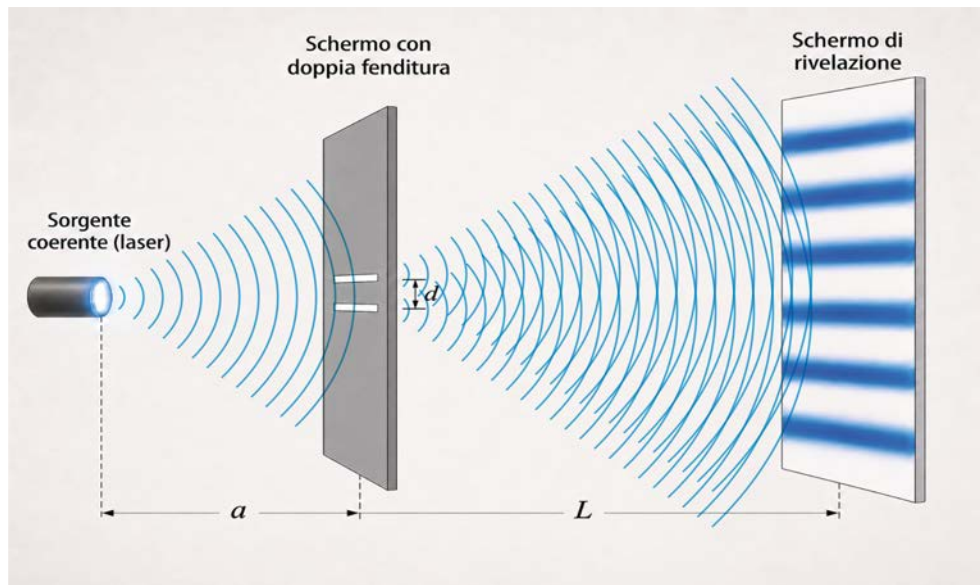
Il risultato sperimentale: la figura osservata sullo schermo, con le sue frange chiare e scure, ma *anche*, nel regime quantistico, i singoli eventi di rivelazione che, accumulandosi, formano una distribuzione di interferenza.

Il modello teorico: la struttura concettuale e matematica usata per descrivere il risultato mediante onde elettromagnetiche classiche *oppure*, nel caso della meccanica quantistica, mediante ampiezze di probabilità.

L'interpretazione fisica: il significato che attribuiamo ai due modelli, cioè che cosa possiamo affermare della luce o della particella prima della misura, che cosa intendiamo per interferenza, quale ruolo svolge il processo di misura e quali limiti presenta il linguaggio classico.

L'interpretazione della teoria classica

Come spiegare formalmente l'esperimento



Interferenza delle onde luminose rappresentate nella descrizione classica
(le due fenditure si comportano come sorgenti d'onda secondarie coerenti).

La definizione di onda in fisica

Prima di affrontare formalmente l'esperimento della doppia fenditura è utile precisare che cosa si intende, in fisica, per *onda*. Un'onda è una *perturbazione* di una grandezza fisica che si propaga nello spazio e nel tempo, trasportando energia senza richiedere necessariamente un trasporto netto di materia. Nel caso della luce, la grandezza che oscilla non è una sostanza materiale, ma il campo elettromagnetico: ciò significa che il *campo elettrico* e il *campo magnetico* variano nel tempo e si propagano nello spazio in modo coordinato (vedi oltre).

In forma semplice, un'onda può essere rappresentata da una funzione come:

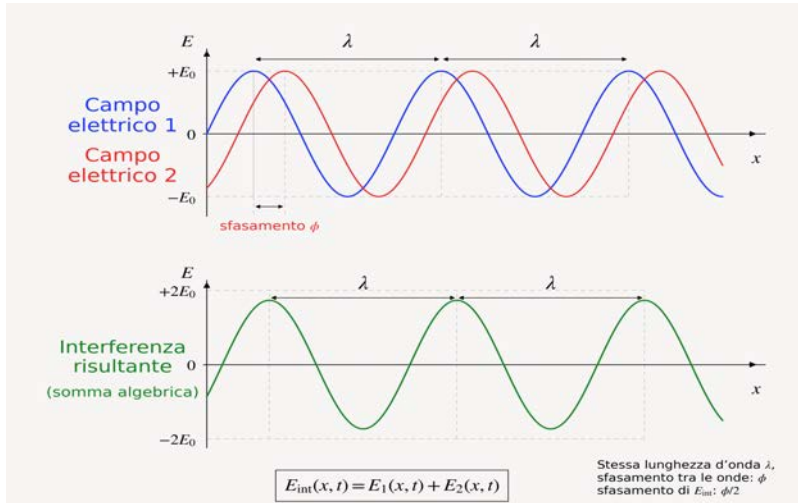
$$E(x, t) = E_0 \cos(kx - \omega t + \phi)$$

dove E_0 è l'ampiezza, $k = \frac{2\pi}{\lambda}$ il numero d'onda, λ la lunghezza d'onda, $\omega = 2\pi/T$ la pulsazione, T il periodo e ϕ la fase iniziale (cioè quando $x = 0$, $t = 0$). La *fase* è particolarmente importante perché, come vedremo, determina il modo in cui due onde si sovrappongono. Inoltre si usa spesso scrivere in forma complessa (essendo $e^{i\varphi} = \cos\varphi + i \sin\varphi$):

$$E(x, t) = \text{Re}[E_0 e^{i(kx - \omega t + \phi)}]$$

dove il campo fisico osservabile è la *parte reale* e dove $(kx - \omega t + \phi)$ è la fase dell'onda.

L'interferenza dei campi elettrici



Schema grafico di due onde del campo elettrico in sovrapposizione
(l'ampiezza del campo risultante è rappresentata con scala diversa).

Per quanto detto prima, nel modello classico i campi elettrici provenienti dalle due fenditure si sommano algebricamente. Se due onde $E_1(x, t)$ ed $E_2(x, t)$ hanno stessa ampiezza E_0 , stessa lunghezza d'onda λ e uno sfasamento relativo ϕ , possiamo scrivere:

$$E_1(x, t) = E_0 \cos(kx - \omega t)$$

$$E_2(x, t) = E_0 \cos(kx - \omega t + \phi).$$

La loro somma algebrica definisce il campo in interazione E_{int} (vedi l'Appendice A):

$$E_{int}(x, t) = E_1(x, t) + E_2(x, t) = 2 E_0 \cos\left(\frac{\phi}{2}\right) \cos\left(kx - \omega t + \frac{\phi}{2}\right).$$

L'ampiezza dell'onda risultante $E_{int}^0 = 2E_0 \cos\left(\frac{\phi}{2}\right)$ dipende quindi dallo sfasamento. In particolare se $\phi = 2n\pi$ (con $n = 0, 1, 2, \dots$) le onde sono in fase e si rinforzano poiché: $E_{int}^0 = 2E_0$. Se $\phi = (2n + 1)\pi$ sono in opposizione di fase e le onde si elidono cioè: $E_{int}^0 = 0$. Pertanto l'ampiezza del campo elettrico risultante varia dal valore nullo al valore massimo $2E_0$.

Ciò si riflette anche sull'*intensità luminosa*, cioè sulla quantità di energia trasportata dall'onda per unità di tempo e per unità di superficie, che risulta proporzionale al valore medio del quadrato del campo elettrico (vedi l'Appendice A):

$$I(\phi) \propto \langle E_{int}^2(x, t) \rangle = 2E_0^2 \cos^2\left(\frac{\phi}{2}\right).$$

Questa dipendenza dell'intensità dallo sfasamento è alla base della formazione delle frange di interferenza, come in effetti si rileva nell'esperimento delle due fenditure (vedi oltre).

Come interpretare l'esperimento

Secondo la teoria classica, come abbiamo visto, la luce che raggiunge la doppia fenditura si comporta come un'onda elettromagnetica. Quando il fronte d'onda incontra lo schermo con le due aperture, ciascuna fenditura diventa, nel modello, una *sorgente secondaria* di onde. Le onde provenienti dalle due fenditure si propagano quindi verso lo schermo finale e si sovrappongono tra loro (vedi la prossima figura).

Questa sovrapposizione può produrre due effetti diversi. Nei punti dello schermo in cui le onde arrivano in fase, le loro ampiezze si sommano e l'intensità luminosa aumenta: si forma quindi una frangia chiara. Nei punti in cui arrivano in opposizione di fase, le ampiezze si compensano parzialmente o totalmente: si forma cioè una frangia scura. Il risultato è una successione regolare di massimi e minimi di luminosità, cioè una *figura di interferenza*.

In questa descrizione, dunque, le frange non sono un effetto misterioso: sono la conseguenza della natura ondulatoria della luce. La teoria classica spiega il risultato osservato assumendo che dalle due fenditure emergano onde coerenti, capaci di interferire tra loro.

Riportiamo di nuovo la formula sintetica (che deriveremo di seguito) che lega la distanza tra due frange successive ΔY ai parametri noti dell'esperimento:

$$\Delta Y \approx \frac{\lambda L}{d}$$

dove λ è la lunghezza d'onda della luce, L è la distanza tra la doppia fenditura e lo schermo di rivelazione e d è la distanza tra le due fenditure.

Invece, lo sfasamento ϕ tra le due onde deriva dalla loro *differenza di cammino* Δl infatti, poiché una lunghezza d'onda completa λ corrisponde a una variazione di fase di 2π , si ha:

$$\phi = \frac{2\pi}{\lambda} \Delta l$$

da cui risulta evidente che per $\Delta l = \lambda$ risulta $\phi = 2\pi$.

Come indicato nella prossima figura, nell'esperimento della doppia fenditura la differenza di cammino tra le onde provenienti dalle due fenditure dipende dalla posizione Y del punto P osservato sullo schermo. Infatti, per piccoli angoli θ e nel caso in cui lo schermo è molto più lontano della distanza tra le due fenditure (cioè per $L \gg d$), si ha come mostreremo:

$$\Delta l \simeq d \sin \theta \simeq \frac{dY}{L}$$

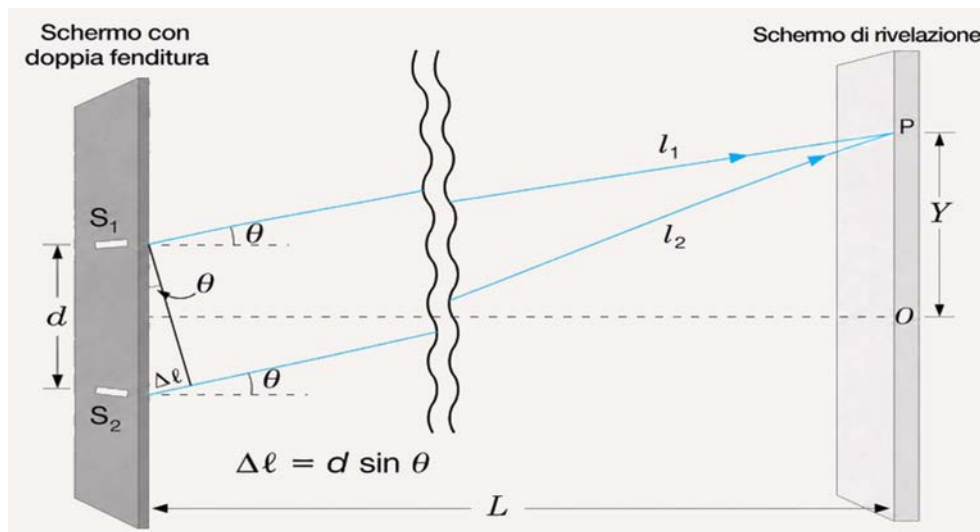
e sostituendo questa espressione nella definizione dello sfasamento di prima si ottiene:

$$\phi(Y) \simeq \frac{2\pi}{\lambda} \frac{dY}{L}.$$

Questa relazione mostra che, a partire dal centro dello schermo, lo sfasamento varia linearmente con la posizione Y del punto P . Poiché l'intensità luminosa dipende, come abbiamo detto, da $\cos^2\left(\frac{\phi}{2}\right)$ è proprio questa variazione regolare dello sfasamento ϕ a produrre l'alternanza periodica di frange luminose e scure.

I calcoli secondo la teoria classica

Applicazione alla doppia fenditura



Differenza di cammino ottico tra due raggi luminosi che convergono nel punto P
(le due linee ondulate indicano il taglio della figura per ragioni di spazio).

Il modello ondulatorio classico

Nel modello ondulatorio classico, la luce è descritta come un'onda elettromagnetica. Tuttavia, quando si vuole calcolare la differenza di cammino ottico tra due contributi ondulatori, si usa spesso una rappresentazione geometrica approssimata mediante raggi luminosi, cioè linee rette che indicano la direzione di propagazione dell'onda.

In questa approssimazione la figura illustra la geometria d'interferenza nella doppia fenditura. I due raggi che emergono dalle due *fenditure/sorgenti* raggiungono lo stesso punto *P* dello schermo di rivelazione, ma i due percorsi non hanno la stessa lunghezza, tra le onde compare una differenza di cammino ottico $\Delta l = |l_1 - l_2|$ responsabile delle frange di interferenza.

Se lo schermo è molto lontano rispetto alla distanza tra le fenditure, cioè se $L \gg d$ i due raggi l_1 e l_2 possono essere considerati quasi paralleli in prossimità dell'uscita dalle fenditure. In questa approssimazione, la differenza di cammino risulta (come si deduce dalla figura):

$$\Delta l \simeq d \sin \theta$$

cioè la differenza di cammino ottico dipende dalla distanza d e dall'angolo θ .

Nota: la figura precedente è puramente rappresentativa e non riproduce l'apparato in scala; distanze e angoli sono stati amplificati per mostrare meglio la geometria del fenomeno.

La condizione dei massimi luminosi

Come si può dedurre dai precedenti paragrafi, si ha interferenza costruttiva quando la differenza di cammino dei raggi luminosi è un multiplo intero della lunghezza d'onda λ :

$$\Delta l \simeq m\lambda$$

dove $m = 0, \pm 1, \pm 2, \pm 3 \dots$ sono gli interi che indicano l'ordine della frangia (sopra e sotto l'asse centrale). Sostituendo $\Delta l = d \sin \theta$ si ottiene subito la *condizione dei massimi luminosi*:

$$d \sin \theta \simeq m\lambda .$$

La posizione delle frange sullo schermo

Per piccoli angoli θ , cioè quando il punto P sullo schermo non è troppo lontano dall'asse centrale, si può usare la seguente approssimazione (essendo $\cos \theta \simeq 1$):

$$\tan \theta = \frac{\sin \theta}{\cos \theta} \simeq \sin \theta .$$

Ora, dalla geometria della figura risulta (per d piccolo rispetto a L , cioè per $d \ll L$):

$$\tan \theta \simeq \frac{Y}{L}$$

da cui segue per l'approssimazione precedente: $\sin \theta \simeq \frac{Y}{L}$ e sostituendo questo risultato nella condizione dei massimi prima ricavata, si ha: $\frac{dY}{L} \simeq m\lambda$ da cui si ottiene infine la distanza ΔY tra le frange (calcolata ad esempio tra $m = 1$ e $m = 0$):

$$\Delta Y \simeq \frac{\lambda L}{d} .$$

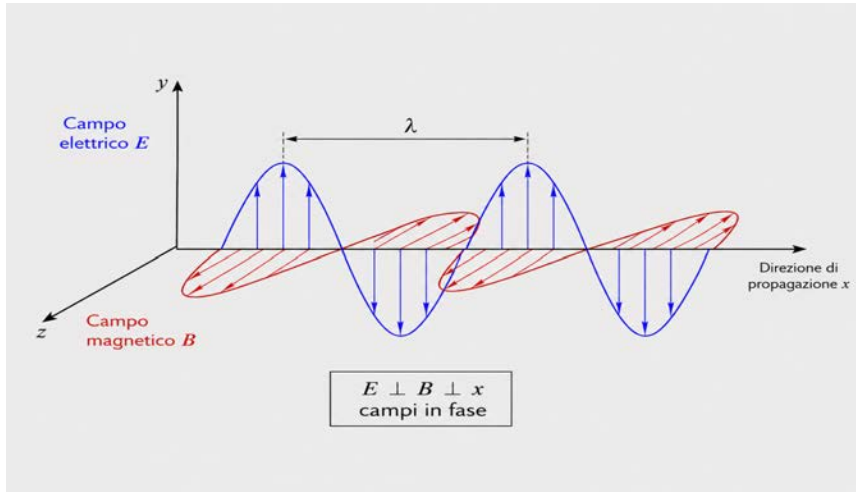
Questa relazione mostra un aspetto fondamentale del modello classico: la figura osservata sullo schermo non dipende solo dalla geometria dell'apparato, ma anche dalla lunghezza d'onda della luce. Se aumenta L , le frange si allargano; se aumenta d , le frange si avvicinano; se cambia λ , cambia anche la spaziatura della figura di interferenza.

Il valore misurato risulta in effetti compatibile, entro le incertezze sperimentali, con quello previsto dalla formula sopra derivata (vedi la *Parte Terza*); ciò indica che le frange osservate, prodotte dall'interferenza tra le onde diffratte dalle due fenditure (che si comportano come sorgenti secondarie), sono ben descritte dal modello ondulatorio classico della luce.

In linea di principio, anche la dipendenza dell'intensità I dallo sfasamento ϕ può essere verificata misurando l'intensità luminosa nei diversi punti dello schermo. Tuttavia, qui ci limitiamo alla misura di ΔY , poiché una verifica quantitativa di I richiederebbe un sensore calibrato e sarebbe inoltre influenzata dalla diffrazione dovuta alla larghezza finita delle fenditure.

Nota: ricordiamo che la *diffrazione*, cioè la deviazione e l'allargamento dell'onda, è apprezzabile quando la larghezza della fenditura è dell'ordine della lunghezza d'onda incidente.

La luce come onda elettromagnetica



Schema grafico del campo elettromagnetico di un'onda piana monocromatica
(i campi E , B e la direzione di propagazione x sono tra loro perpendicolari).

Le equazioni di Maxwell

Come abbiamo visto, nella descrizione classica abbiamo rappresentato la luce mediante il solo campo elettrico $E(x, t)$, poiché nel nostro caso è del tutto sufficiente per descrivere la sovrapposizione delle onde e la formazione delle frange di interferenza.

Tuttavia, nella teoria elettromagnetica completa, la luce non è costituita solo da un campo elettrico: è un'onda elettromagnetica, formata da un campo elettrico E e da un campo magnetico B , entrambi variabili nel tempo e nello spazio.

Il punto di partenza, per una descrizione formale, sono le note *equazioni di Maxwell*. Nel vuoto, cioè in assenza di cariche e correnti elettriche o magnetiche, esse possono essere scritte nella seguente forma vettoriale e simmetrica (vedi l'*Appendice A*):

$$\begin{aligned}\nabla \cdot \mathbf{E} &= 0 \quad ; \quad \nabla \times \mathbf{E} = -\frac{\partial \mathbf{B}}{\partial t} \\ \nabla \cdot \mathbf{B} &= 0 \quad ; \quad \nabla \times \mathbf{B} = \mu_0 \epsilon_0 \frac{\partial \mathbf{E}}{\partial t}.\end{aligned}$$

Le due equazioni a sinistra esprimono l'assenza di cariche elettriche e monopoli magnetici, mentre le altre due equazioni a destra affermano rispettivamente che un campo magnetico variabile genera un campo elettrico variabile e che un campo elettrico variabile genera a sua volta un campo magnetico variabile. Questa *reciproca generazione* permette la propagazione di un'onda elettromagnetica anche nel vuoto (dove come vedremo sono definiti ϵ_0 e μ_0).

Le onde elettromagnetiche

Come è noto, a partire dalle equazioni di Maxwell, si ottiene che il campo elettrico $\mathbf{E}$ e il campo magnetico $\mathbf{B}$ soddisfano le seguenti *equazioni d'onda* (vedi l'*Appendice A*):

$$\nabla^2 \mathbf{E} = \mu_0 \varepsilon_0 \frac{\partial^2 \mathbf{E}}{\partial t^2} \quad ; \quad \nabla^2 \mathbf{B} = \mu_0 \varepsilon_0 \frac{\partial^2 \mathbf{B}}{\partial t^2}$$

dove ε_0 è la permittività elettrica e μ_0 è la permeabilità magnetica, entrambe del vuoto.

Confrontandole con l'*equazione canonica* di un'onda $\mathbf{u}$ che si propaga a velocità v :

$$\nabla^2 \mathbf{u} = \frac{1}{v^2} \frac{\partial^2 \mathbf{u}}{\partial t^2}$$

si deduce che la velocità di propagazione dell'onda elettromagnetica nel vuoto è:

$$\frac{1}{v^2} = \mu_0 \varepsilon_0 \quad \text{cioè} \quad v = \frac{1}{\sqrt{\mu_0 \varepsilon_0}} = c.$$

Questo risultato è fondamentale: la velocità prevista dalle equazioni di Maxwell coincide con la velocità della luce c ; da qui nasce l'identificazione della luce come onda elettromagnetica.

Come abbiamo già visto, per un'onda piana il campo elettrico $E(x, t)$ può essere rappresentato in forma semplificata. Consideriamo cioè un'onda monocromatica, sinusoidale e polarizzata lungo una direzione fissata che si propaga lungo l'asse X (a velocità $c = \lambda f$):

$$E(x, t) = E_0 \cos(kx - \omega t + \phi)$$

dove E_0 è l'ampiezza, $k = \frac{2\pi}{\lambda}$ il numero d'onda, λ la lunghezza d'onda, $\omega = 2\pi/T$ la pulsazione, T il periodo e ϕ la fase iniziale. Inoltre il campo magnetico associato $B(x, t)$, ha la stessa struttura ondulatoria, è perpendicolare sia al campo elettrico che alla direzione di propagazione e, nel vuoto, ha ampiezza legata al campo elettrico (vedi l'*Appendice A*):

$$B_0 = \frac{E_0}{c}.$$

In altre parole, nella approssimazione di una semplice onda luminosa monocromatica e piana (condizione che si verifica lontano dalla sorgente) i campi $\mathbf{E}$, $\mathbf{B}$ e la direzione di propagazione $\mathbf{x}$ sono mutuamente perpendicolari, come mostrato nella precedente figura.

Tuttavia, per studiare il fenomeno dell'interferenza si può considerare soltanto il campo elettrico, dato che l'intensità luminosa registrata sullo schermo si dimostra *proporzionale* al valore medio del quadrato del campo elettrico (vedi l'*Appendice A*):

$$I \propto \langle E^2(x, t) \rangle.$$

Il campo magnetico è quindi sempre presente, ma non è necessario rappresentarlo esplicitamente per descrivere la distribuzione di intensità e calcolare la posizione delle frange. La figura di interferenza può essere descritta introducendo solo il campo elettrico, sapendo però che esso fa parte di un'onda elettromagnetica completa.

Parte Seconda

L'esperimento della doppia fenditura secondo la teoria quantistica

La doppia fenditura nella teoria quantistica

Come abbiamo visto nella precedente sezione, la descrizione classica dell'esperimento della doppia fenditura spiega con grande precisione la formazione delle frange quando la luce viene trattata come un'onda elettromagnetica. Tuttavia, questo esperimento assume un significato ancora più profondo quando l'intensità della sorgente viene ridotta fino al regime dei singoli fotoni. In quel caso, la figura di interferenza non scompare ma emerge progressivamente dall'accumulo di molti eventi di rivelazione localizzati.

Questo fatto modifica radicalmente il modo in cui dobbiamo interpretare l'esperimento. Nel modello classico, le frange sono spiegate come il risultato della sovrapposizione di due onde luminose. Nel regime quantistico, invece, ogni fotone viene rivelato sullo schermo come un singolo evento localizzato. Eppure, dopo moltissimi eventi, la distribuzione complessiva ricostruisce la stessa figura di interferenza.

Per comprendere questo passaggio è utile ripercorrere brevemente alcune tappe decisive della nascita della fisica quantistica, per quanto riguarda la natura corpuscolare della luce.

Dalla luce come onda alla luce come quanto

Alla fine dell'Ottocento, la teoria elettromagnetica di Maxwell aveva fornito una descrizione estremamente efficace della luce come onda. La luce era interpretata come una perturbazione dei campi elettrico e magnetico, capace di propagarsi nel vuoto alla velocità:

$$c = \frac{1}{\sqrt{\mu_0 \epsilon_0}}$$

dove ϵ_0 è la permittività elettrica e μ_0 è la permeabilità magnetica, entrambe del vuoto.

Questa descrizione spiegava con grande successo fenomeni ottici come riflessione, rifrazione, diffrazione e interferenza. L'esperimento della doppia fenditura sembrava quindi confermare in modo convincente la natura ondulatoria della luce.

All'inizio del Novecento, però, emersero fenomeni che non potevano essere spiegati pienamente con il solo modello ondulatorio classico. Uno dei più importanti fu l'effetto fotoelettrico: quando una radiazione luminosa colpisce una superficie metallica, può provocare l'emissione di elettroni. Il dato sorprendente era che l'emissione dipendeva dalla frequenza della luce, non semplicemente dalla sua intensità, come previsto dal modello classico.

Nel 1905 Albert Einstein propose una spiegazione radicale. Riprendendo l'ipotesi introdotta da Max Planck nello studio della radiazione di corpo nero, Einstein suggerì che la luce, in alcune interazioni con la materia, si comportasse come se fosse composta da pacchetti discreti di energia. Questi quanti di luce, che più tardi sarebbero stati chiamati *fotoni* (dal chimico Gilbert Lewis nel 1926), hanno per ipotesi energia pari a

$$E = h\nu$$

dove h è la costante di Planck e ν è la frequenza della radiazione.

Un comportamento inatteso

L'effetto fotoelettrico mostrava quindi un aspetto inatteso: la luce, che nell'interferenza si comporta come un'onda, nell'interazione con la materia può manifestarsi come un insieme di eventi discreti. L'energia non viene assorbita in modo continuo, ma in quantità elementari.

Questa duplice natura della luce preparò il terreno alla meccanica quantistica. La domanda che ci si poneva non era più semplicemente: la luce è un'onda oppure una particella?

La domanda diventava più sottile: quale descrizione teorica permette di prevedere correttamente ciò che osserviamo nei diversi esperimenti e in particolare, considerato l'argomento che stiamo trattando, nel classico esperimento della doppia fenditura?

Il problema della doppia fenditura

Nella doppia fenditura, la questione diventa particolarmente evidente. Se la sorgente emette luce con intensità elevata, lo schermo mostra direttamente una figura di interferenza *continua*, come previsto dalla teoria ondulatoria classica. Se però l'intensità viene ridotta fino a emettere un fotone alla volta, ogni fotone rivelato produce un *singolo evento* localizzato.

All'inizio, sullo schermo compaiono punti apparentemente isolati e distribuiti in modo irregolare. Ma aumentando il numero di fotoni rivelati, emerge progressivamente una struttura ordinata: le zone in cui arrivano molti eventi corrispondono alle frange luminose, mentre le zone in cui arrivano pochi eventi corrispondono alle frange scure.

Questo comportamento non può essere descritto dicendo semplicemente che ogni fotone segue una *traiettoria classica* attraverso una fenditura *oppure* l'altra. Se così fosse, ci aspetteremmo una distribuzione simile alla somma di due immagini indipendenti delle fenditure: il risultato osservato è invece una *distribuzione di interferenza*.

La meccanica quantistica introduce allora un nuovo oggetto teorico: la *funzione d'onda* o più in generale l'ampiezza di probabilità. Essa non descrive una traiettoria osservabile del fotone nello spazio, ma permette di *calcolare la probabilità* dei diversi risultati possibili.

Nel caso della doppia fenditura, al fotone sono associate ampiezze di probabilità, relative ai diversi modi in cui esso può contribuire all'evento finale sullo schermo e queste ampiezze si sommano e possono interferire. Ciò che si osserva sperimentalmente, invece, non è l'ampiezza di probabilità in sé, ma un evento localizzato di rivelazione.

In sintesi, per quanto si è detto, la differenza concettuale fondamentale è la seguente:

- ➔ nel modello classico interferiscono campi elettrici;
- ➔ nel modello quantistico interferiscono ampiezze di probabilità.

Ciò significa che nel primo caso l'interferenza riguarda *grandezze fisiche* associate all'onda elettromagnetica; nel secondo invece riguarda *oggetti matematici* che permettono di calcolare la probabilità dei diversi eventi di rivelazione osservabili sullo schermo. E questa, a livello concettuale e quindi anche teorico, è una notevole differenza.

Dai molti cammini all'ampiezza totale

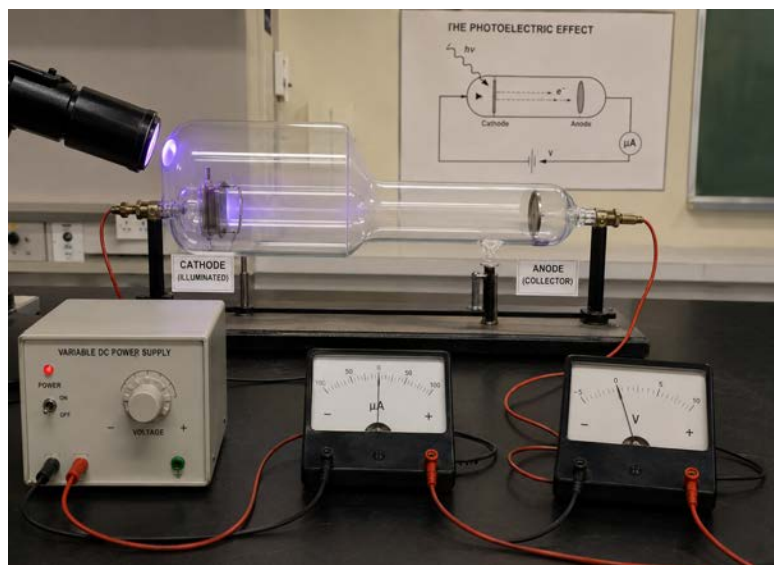
Una formulazione particolarmente elegante di questa idea, come vedremo, fu poi sviluppata da Richard Feynman, attraverso il metodo dell'*integrale sui cammini*. In questa descrizione, per calcolare la probabilità che una particella venga rivelata in un certo punto dello schermo, non si considera un'unica traiettoria classica. Si sommano invece le ampiezze di probabilità associate a *tutti i cammini possibili*, che collegano la sorgente al punto di rivelazione.

Nel limite classico si osserva che i *contributi dei cammini* che risultano molto diversi da quello previsto dalla meccanica classica, tendono a *cancellarsi per interferenza* e quindi resta dominante il cammino classico. Nel regime quantistico, invece, la somma delle diverse ampiezze di probabilità diventa essenziale per descrivere correttamente il risultato dell'esperimento.

Applicata alla doppia fenditura, questa idea mostra perché il fenomeno sia così profondo. Quando non si misura da quale fenditura passa il fotone, le alternative possibili contribuiscono insieme all'ampiezza totale, dato che le ampiezze associate alle due fenditure si sommano e possono interferire. Quando invece si introduce una *misura capace di distinguere* il cammino, la figura di interferenza scompare (come vedremo meglio più avanti).

L'esperimento della doppia fenditura diventa così una finestra privilegiata sulla struttura della teoria quantistica. Non mostra solo che la luce è *a volte onda e a volte particella*. Mostra piuttosto che il linguaggio classico di onde e particelle non basta più: ciò che la teoria calcola sono ampiezze di probabilità, mentre ciò che il laboratorio registra sono eventi osservabili.

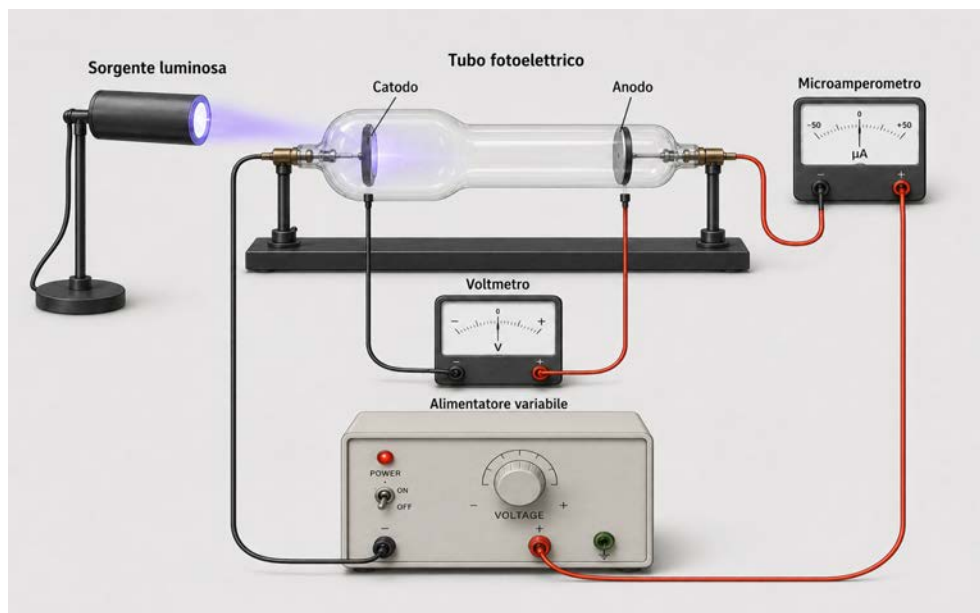
Come già anticipato, nel prossimo paragrafo vedremo come un particolare esperimento sia stato fondamentale per la nascita della realtà dei quanti: l'effetto fotoelettrico.



Strumentazione di laboratorio per l'esperimento dell'effetto fotoelettrico
(lo schema dei collegamenti elettrici è riportato nella pagina seguente).

In laboratorio con Einstein

L'esperimento dell'effetto fotoelettrico



Schema di collegamento dell'apparato per l'esperimento dell'effetto fotoelettrico
(l'alimentatore variabile può applicare una tensione positiva o negativa tra anodo e catodo).

Osservare prima di interpretare

Come abbiamo detto la descrizione ondulatoria classica della luce spiega con grande efficacia fenomeni come interferenza, diffrazione, riflessione e rifrazione. Tuttavia, all'inizio del Novecento emerse un fenomeno sperimentale, abbastanza semplice da eseguire, che mise in difficoltà una descrizione puramente ondulatoria dell'interazione tra luce e materia, che venne poi chiamato: effetto fotoelettrico.

L'esperimento consiste in pratica nell'illuminare una superficie metallica e osservare se dalla superficie vengono emessi elettroni. Questi elettroni, chiamati *fotoelettroni*, possono essere raccolti da un secondo elettrodo, generando una corrente elettrica misurabile.

Il dato sorprendente è che l'emissione non dipende soltanto dall'intensità della luce, cioè dalla quantità complessiva di radiazione incidente, ma soprattutto dalla sua frequenza. Al di sotto di una certa frequenza minima, detta *frequenza di soglia*, non vengono emessi elettroni, anche aumentando l'intensità della radiazione. Al di sopra della soglia, invece, avviene l'emissione e l'energia degli elettroni emessi cresce con la frequenza della luce incidente.

La strumentazione essenziale

L'apparato sperimentale può essere descritto attraverso quattro elementi principali:

La sorgente luminosa: produce la radiazione incidente. È importante poter variare la frequenza della luce e, separatamente, la sua intensità. La frequenza determinerà l'energia di ciascun quanto di luce; l'intensità invece il numero di quanti incidenti per unità di tempo.

Il fotocatodo: è la superficie metallica colpita dalla luce. Se la radiazione possiede frequenza sufficiente, dal metallo vengono estratti elettroni. L'energia minima necessaria per liberare un elettrone dalla superficie prende il nome di *lavoro di estrazione* e lo indicheremo con ϕ .

L'anodo collettore: è un secondo elettrodo posto di fronte al fotocatodo. La sua funzione è raccogliere gli elettroni emessi e inoltre, quando gli elettroni emessi raggiungono l'anodo, nel circuito elettrico preparato esternamente si misura una corrente.

Il circuito di misura: comprende un generatore di tensione variabile, un amperometro sensibile e un voltmetro. Variando la tensione tra catodo e anodo, che sono posti nel tubo fotoelettrico sotto vuoto, si può favorire *oppure* ostacolare il moto degli elettroni emessi.

La sorgente luminosa

La sorgente luminosa ha un ruolo decisivo, perché l'effetto fotoelettrico dipende direttamente dalla frequenza della radiazione incidente. Se la frequenza della luce è inferiore alla frequenza di soglia ν_0 non si osserva emissione di elettroni, qualunque sia l'intensità della radiazione. Ma se la frequenza supera la soglia, l'emissione avviene quasi immediatamente.

L'intensità conserva comunque un ruolo importante: sopra la frequenza di soglia, aumentando l'intensità, aumenta il numero di elettroni emessi per unità di tempo e quindi aumenta la corrente elettrica. Ma l'energia massima dei singoli elettroni emessi dipende dalla frequenza, non dall'intensità, come invece ci aspetteremmo dalla teoria classica.

Come si esegue l'esperimento

Per eseguire l'esperimento si illumina il catodo con una radiazione di frequenza controllata. Gli elettroni eventualmente emessi vengono raccolti dall'anodo e producono una corrente misurabile nel circuito esterno.

Si procede *variando* la frequenza della luce incidente e osservando se compare una corrente elettrica. Quando la frequenza è inferiore alla soglia, non si registra emissione. Appena la frequenza supera la soglia, compare una corrente, dovuta agli elettroni estratti dal metallo.

A frequenza *fissata* sopra la soglia, si può variare l'intensità della luce. In questo caso cambia la corrente elettrica, perché cambia il numero di elettroni emessi, ma non cambia l'energia cinetica massima dei singoli elettroni, che infatti dipende dalla frequenza.

Per misurare questa energia, si applica una *tensione frenante* V_s che cresce fino ad annullare completamente la corrente. Il valore della tensione di arresto consente di ricavare K_{max} cioè l'energia cinetica massima degli elettroni emessi.

In definitiva, esiste una frequenza di soglia sotto la quale non avviene emissione, mentre l'energia massima degli elettroni emessi cresce linearmente con la frequenza della luce.

Che cosa si misura

Le due grandezze principali misurate nell'esperimento dell'effetto fotoelettrico sono quindi la *corrente elettrica* e la *tensione di arresto*.

La corrente elettrica: misura il numero di elettroni che raggiungono l'anodo per unità di tempo. Se la frequenza della luce è superiore alla soglia, aumentando l'intensità della sorgente aumenta il numero di elettroni emessi e quindi aumenta la corrente.

La tensione di arresto: indicata spesso con V_s , è la tensione frenante necessaria per impedire anche agli elettroni più energetici di raggiungere l'anodo. Essa permette di ricavare l'energia cinetica massima degli elettroni K_{max} (in pratica quelli posti sulla superficie del metallo):

$$K_{max} = eV_s$$

dove con e indichiamo la carica elementare dell'elettrone. Il risultato sperimentale fondamentale è che K_{max} cresce linearmente con la frequenza della luce incidente; per questo motivo Einstein interpretò questo fatto assumendo che la luce venga assorbita in quanti di energia:

$$E = h\nu$$

dove h è la costante di Planck e ν è la frequenza della radiazione.

In particolare, l'energia assorbita serve prima a vincere il lavoro di estrazione Φ , mentre l'energia restante diventa energia cinetica dell'elettrone, cioè descritto in simboli:

$$K_{max} = h\nu - \Phi.$$

Combinando questa ultima relazione con quella precedente $K_{max} = eV_s$ segue subito: $eV_s = h\nu - \Phi$ da cui possiamo ricavare il lavoro di estrazione Φ :

$$\Phi = h\nu - eV_s.$$

Una volta determinato Φ (noti sperimentalmente ν e V_s) la frequenza di soglia ν_0 si ottiene ponendo $K_{max} = 0$ nella relazione precedente $K_{max} = h\nu - \Phi$ da cui segue perciò:

$$h\nu_0 = \Phi \Rightarrow \nu_0 = \frac{\Phi}{h}.$$

L'esperimento dimostra che sotto questa frequenza non si osservano elettroni emessi.

In definitiva, l'effetto fotoelettrico non nega il comportamento ondulatorio della luce osservato nella doppia fenditura. Mostra però che, quando la luce interagisce con la materia, l'energia viene scambiata in quantità discrete.

È questo passaggio sperimentale che rende possibile comprendere la descrizione quantistica della doppia fenditura nel regime dei singoli fotoni: la luce continua a produrre interferenza ma, come vedremo, viene assorbita e rivelata in eventi discreti.

Nota: ricordiamo che in generale l'energia cinetica di ogni singolo elettrone è pari a $K = \frac{1}{2}mv^2$ dove m e v sono rispettivamente la massa e la velocità dell'elettrone; la velocità v varia in funzione della profondità nel metallo e degli urti subiti prima di uscire dalla superficie.

Perché Einstein vinse il Nobel

Quando la teoria ondulatoria fallisce

Nei paragrafi precedenti abbiamo visto come Einstein spiega l'esperimento dell'effetto fotoelettrico. Qui può essere utile riassumere il punto essenziale: dove la teoria ondulatoria classica prova a descrivere il fenomeno, ma fallisce; e dove invece l'ipotesi quantistica riesce a coglierne completamente il significato fisico.

Per cogliere la portata della spiegazione proposta da Einstein, è infatti utile ricordare che cosa ci saremmo dovuti aspettare da una descrizione ondulatoria classica della luce:

Primo: in un'onda elettromagnetica classica l'intensità è proporzionale al valore medio del quadrato del campo elettrico, infatti come abbiamo anticipato risulta (vedi l'Appendice A):

$$I \propto \langle E^2 \rangle.$$

Dato che il campo elettrico esercita per definizione una forza F sulle cariche pari a:

$$F = eE$$

allora, aumentando l'intensità dell'onda, dovrebbe aumentare anche l'energia elementare δW trasferibile agli elettroni (essendo $\delta W = Fdr$). L'esperimento mostra invece che l'energia cinetica degli elettroni emessi dipende dalla frequenza della luce, non dalla sua intensità.

Secondo: seguendo l'immagine classica, una luce sufficientemente intensa dovrebbe poter estrarre elettroni anche a bassa frequenza, se l'energia totale fornita è abbastanza grande. L'esperimento mostra invece l'esistenza di una frequenza di soglia: sotto il valore ν_0 non si ha emissione, qualunque sia l'intensità della radiazione. E infatti nella descrizione di Einstein, la condizione da soddisfare per l'emissione degli elettroni è, come prima descritto:

$$h\nu > \phi$$

dove ϕ è il lavoro di estrazione degli elettroni dal metallo.

Terzo: se l'energia fosse assorbita in modo continuo dall'onda, a bassa intensità ci potremmo aspettare un certo tempo di accumulo prima dell'emissione. L'esperimento non mostra questo ritardo: quando la frequenza è sopra la soglia, l'emissione avviene senza un ritardo apprezzabile. Nell'interpretazione quantistica, l'energia non viene accumulata progressivamente da un'onda continua, ma trasferita in un singolo evento di assorbimento di un *quanto di luce*.

È inoltre opportuno ricordare che l'articolo di Einstein del 1905, intitolato *Über einen die Erzeugung und Verwandlung des Lichtes betreffenden heuristischen Gesichtspunkt* (Un punto di vista euristico sulla emissione e trasformazione della luce), non si occupava solo dell'effetto fotoelettrico, ma anche di emissione e assorbimento della radiazione, di ionizzazione dei gas, di fotoluminescenza e più in generale di scambi energetici tra radiazione e materia.

Il suo obiettivo era proprio quello di mostrare che alcuni fenomeni di emissione e assorbimento della luce, diventano comprensibili solo se si assume che l'energia luminosa non sia distribuita nello spazio in modo continuo, ma possa presentarsi anche in forma discreta (in particolare, quando si considerano i singoli eventi di interazione tra luce e materia).

Due grandi idee per una rivoluzione concettuale

Nell'articolo citato Einstein formula due idee fondamentali, infatti dapprima afferma:

Quando un raggio di luce si espande partendo da un punto, l'energia non si distribuisce su volumi sempre più grandi, bensì rimane costituita da un numero finito di quanti di energia localizzati nello spazio, che si muovono senza suddividersi e che non possono essere assorbiti od emessi parzialmente.

Questa frase introduce il nucleo dell'ipotesi dei *quanti di luce*: l'energia della radiazione, in certi processi, non si trasferisce in modo continuo, ma in unità discrete(!).

La seconda osservazione è altrettanto importante, perché chiarisce il limite del modello della teoria ondulatoria classica:

La teoria ondulatoria della luce, che fa uso di funzioni spaziali continue, si è verificata ottima per quel che riguarda i fenomeni puramente ottici e sembra veramente insostituibile in questo campo. Tuttavia, bisogna tenere presente che le osservazioni ottiche si riferiscono a valori medi nel tempo e non a valori momentanei.

Si noti che Einstein non discredita la teoria ondulatoria. Al contrario, riconosce che essa funziona bene nei fenomeni ottici, come l'interferenza e la diffrazione. Ma osserva correttamente che quei fenomeni riguardano *grandezze medie*, distribuite nello spazio e nel tempo. L'effetto fotoelettrico, invece, mette in evidenza eventi individuali di assorbimento: quando un elettrone riceve un quanto di energia viene emesso, *oppure*, non viene emesso, senza alternative.

In questo senso queste due fondamentali idee, valgono da sole una rivoluzione concettuale. La prima introduce i quanti di energia localizzati; la seconda spiega perché la teoria ondulatoria classica, pur essendo efficace per molti fenomeni ottici, non basta a descrivere l'interazione elementare tra luce e materia.

Ricordiamo infine che Einstein ricevette il Premio Nobel per la Fisica nel 1921 e che gli fu assegnato ufficialmente con la seguente motivazione:

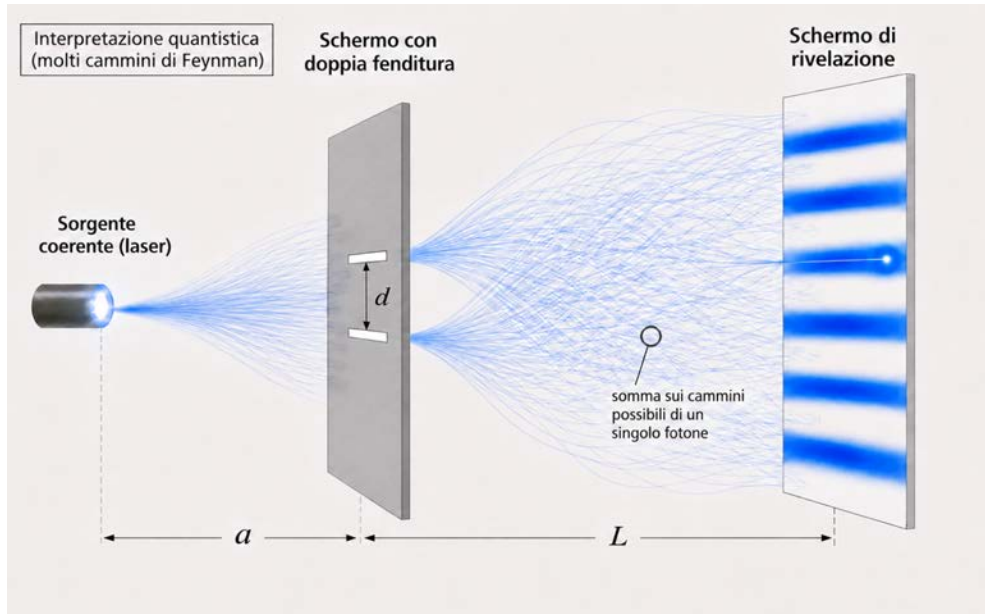
[...] per i suoi contributi alla fisica teorica e, in particolare, per la scoperta della legge dell'effetto fotoelettrico.

Curiosamente non fu premiato per la *Teoria della Relatività*, nonostante oggi essa sia la parte più celebre della sua opera. All'epoca, la relatività era ancora oggetto di discussione e incontrava resistenze teoriche e sperimentali; l'effetto fotoelettrico, invece, offriva una legge verificabile e facilmente misurabile, capace di aprire la strada alla nuova fisica dei quanti.

Per il nostro percorso sulla doppia fenditura, il significato è abbastanza chiaro: la luce si propaga e interferisce come un'onda, ma viene assorbita e rivelata in eventi discreti. È proprio in questa *tensione esplicativa*, tra distribuzione ondulatoria ed evento localizzato, che la meccanica quantistica troverà spazio per costruire una nuova descrizione della realtà fisica.

L'interpretazione della teoria quantistica

Come spiegare formalmente l'esperimento



Interferenza delle ampiezze di probabilità nella descrizione quantistica
(lo schema dell'apparato è fuori scala ed è esclusivamente rappresentativo).

La funzione d'onda o ampiezza di probabilità

Come abbiamo visto, nella descrizione classica dell'esperimento della doppia fenditura abbiamo *interpretato* la figura di interferenza come risultato della sovrapposizione di due onde elettromagnetiche. Nel caso quantistico, invece, il problema cambia natura. Non vogliamo più descrivere soltanto un'onda luminosa continua, ma capire come sia possibile che singoli eventi di rivelazione, accumulandosi nel tempo, producano una figura di interferenza.

Per affrontare questo passaggio, occorre prima introdurre l'oggetto teorico fondamentale della meccanica quantistica: la *funzione d'onda*. Ciò significa che per una particella *non* relativistica, come un elettrone, lo *stato quantistico* può essere descritto da una funzione d'onda:

$$\psi(\mathbf{r}, t)$$

che è soluzione di una particolare equazione d'onda, come descriveremo subito di seguito. Qui $\mathbf{r}$ indica il punto in cui vogliamo valutare l'ampiezza di probabilità al tempo t poiché la funzione d'onda rappresenta appunto una *ampiezza di probabilità* a valori complessi.

L'equazione d'onda di Schrödinger

La funzione d'onda $\psi(\mathbf{r}, t)$ o ampiezza di probabilità non rappresenta una *traiettoria* della particella, ma contiene le *informazioni* necessarie per calcolare i possibili risultati di una misura. La sua evoluzione nel tempo è determinata dalla famosa *equazione di Schrödinger*:

$$i\hbar \frac{\partial \psi(\mathbf{r}, t)}{\partial t} = -\left(\frac{\hbar^2}{2m}\right) \nabla^2 \psi(\mathbf{r}, t) + V(\mathbf{r})\psi(\mathbf{r}, t)$$

dove i è l'unità immaginaria ed $\hbar = \frac{h}{2\pi}$ è la costante di Planck ridotta; questa equazione è valida nel caso di una particella *non* relativistica di massa m soggetta a un potenziale $V(\mathbf{r})$.

Si noti che il primo termine a destra dell'equazione descrive il contributo dell'energia cinetica (vedi l'*Appendice B*); il secondo invece descrive l'interazione con il potenziale: l'equazione stabilisce quindi come evolve nel tempo l'ampiezza quantistica $\psi(\mathbf{r}, t)$ associata al sistema.

Questa formulazione è rigorosamente adatta a particelle massive non relativistiche, come elettroni lenti, neutroni, atomi o molecole. Invece per i fotoni, che sono quanti del campo elettromagnetico e viaggiano nel vuoto alla velocità della luce, la descrizione rigorosa richiede la quantizzazione del campo elettromagnetico (per chiarimenti vedi l'*Appendice D*).

Ad ogni modo, nel seguito useremo il simbolo $\psi(\mathbf{r}, t)$ in senso generale per indicare un'ampiezza di probabilità: nel caso degli elettroni essa è la funzione d'onda di Schrödinger mentre nel caso dei fotoni rappresenta l'ampiezza associata alla rivelazione dei quanti di luce.

Si osservi infine che la soluzione dell'equazione di Schrödinger, come del resto tutte le funzioni complesse, può essere espressa nella seguente forma:

$$\psi(\mathbf{r}, t) = A(\mathbf{r}, t)e^{i\varphi(\mathbf{r}, t)}$$

che ricorda *formalmente* l'onda del campo elettrico classico, introdotto nella precedente sezione, ma in questo caso, come vedremo, il suo *significato fisico* è completamente diverso; come pure quello della *fase* $\varphi(\mathbf{r}, t)$ dell'ampiezza di probabilità è diverso dal caso classico.

L'interpretazione probabilistica di Born

Diciamo subito che la funzione d'onda $\psi(\mathbf{r}, t)$ non è direttamente osservabile. Ciò che si osserva in laboratorio è un evento: per esempio, nel caso della doppia fenditura, un punto luminoso sullo schermo o un segnale in un rivelatore che registra l'evento.

Il collegamento tra la funzione d'onda e il risultato sperimentale fu poi chiarito da Max Born. Secondo la sua interpretazione, la *probabilità* $P(\mathbf{r}, t)$ di rivelare una particella in una certa regione dello spazio, è proporzionale al *modulo quadro* della funzione d'onda:

$$P(\mathbf{r}, t) \propto |\psi(\mathbf{r}, t)|^2.$$

Più precisamente $|\psi(\mathbf{r}, t)|^2$ è una *densità di probabilità* locale: dove questo valore è grande, è più probabile osservare un evento; dove è più piccolo, la probabilità di rivelazione è minore.

Questo è il primo cambiamento concettuale rispetto alla fisica classica. La teoria non assegna alla particella una posizione definita in ogni istante e non descrive una *traiettoria* osservabile *prima* della misura: essa calcola invece le probabilità dei possibili risultati sperimentali.

Feynman e la somma sui cammini

Come già anticipato, Feynman propose una *formulazione intuitiva* della meccanica quantistica, secondo cui la stessa ampiezza di probabilità $\psi(\mathbf{r}, t)$ prima introdotta e associata a un evento, può essere costruita *sommando tutti i contributi* associati ai cammini possibili.

Nel caso non relativistico, questa formulazione è *equivalente* all'equazione di Schrödinger; mentre nella teoria quantistica dei campi il principio può essere esteso anche a campi e fotoni, in modo da poter trattare, ad esempio, il caso della doppia fenditura.

L'idea centrale è che tutte le *alternative indistinguibili* dei possibili cammini, contribuiscono con ampiezze che possono rinforzarsi o cancellarsi; ciò in pratica significa che l'ampiezza totale è la somma dei contributi associati a tutti i cammini possibili (vedi l'*Appendice C*).

Formalmente, se un sistema va da un punto iniziale A a un punto finale B , l'ampiezza complessiva può essere espressa, in forma simbolica, come la seguente somma integrale:

$$\psi_A(B) \propto \int D[\mathbf{r}(t)] e^{\frac{iS[\mathbf{r}(t)]}{\hbar}}$$

dove $S[\mathbf{r}(t)]$ è l'*azione* associata a un determinato cammino, mentre $D[\mathbf{r}(t)]$ indica simbolicamente l'integrazione su tutti i cammini possibili $\mathbf{r}(t)$ che collegano il punto iniziale A al punto finale B . Ogni cammino contribuisce quindi con una fase $\frac{S[\mathbf{r}(t)]}{\hbar}$; i contributi con fasi compatibili si rinforzano; quelli con fasi molto diverse tendono a cancellarsi (vedi l'*Appendice C*).

Questo formalismo, in termini di cammini $\mathbf{r}(t)$, è quello valido per una particella non relativistica. Tuttavia, nei casi relativistici e nella teoria quantistica dei campi l'idea resta la stessa, anche se cambia l'oggetto matematico con cui si sommano i contributi.

Si osservi che questa formulazione *non* significa assolutamente che la particella percorra realmente tutti i cammini, come farebbe un oggetto classico. Significa piuttosto che, per calcolare l'ampiezza del risultato finale, la teoria deve sommare i contributi associati a tutte le alternative possibili, come schematizzato nella figura precedente.

Infine, poiché questo modo di descrivere l'interferenza non vale soltanto per la luce, è implicito che anche particelle materiali, come gli elettroni, possono produrre figure di interferenza quando ad esempio attraversano una doppia fenditura, come dimostrato sperimentalmente.

Come è definita l'azione S

Accenniamo in breve che l'azione S è la grandezza che in meccanica classica caratterizza un cammino dinamico. Si definisce espressamente come (vedi l'*Appendice C*):

$$S = \int \mathcal{L} dt$$

dove $\mathcal{L}$ è la *lagrangiana* che, nel caso non relativistico, è definita semplicemente come: $\mathcal{L} = T - V$ cioè pari all'energia cinetica T meno l'energia potenziale V del sistema considerato. Invece, per le particelle relativistiche e i campi, compreso il campo elettromagnetico, cambia la forma dell'azione S , ma resta comunque il fattore di fase $e^{\frac{iS}{\hbar}}$ che assegna una fase a ciascun contributo associato ai cammini possibili: come al solito è proprio dalla somma di tutti i possibili contributi dei cammini che nasce l'*interferenza quantistica*.

I calcoli secondo la teoria quantistica

Applicazione alla doppia fenditura

Possiamo ora applicare le considerazioni precedenti alla doppia fenditura.

Supponiamo che una particella, o un singolo fotone, venga emesso dalla sorgente e che lo schermo registri un evento nel punto P . Solamente se non viene stabilito, con una misura, da quale fenditura è passato, le due alternative restano *indistinguibili*. In questo caso tuttavia non si sommano le probabilità separate, ma le relative ampiezze di probabilità.

Indichiamo ad esempio con $\psi_1(P)$ l'ampiezza di probabilità associata all'alternativa della prima fenditura (essa rappresenta la somma delle ampiezze di tutti i cammini possibili associati alla prima fenditura) e con $\psi_2(P)$ indichiamo allo stesso modo la somma di tutti i contributi associati alla seconda fenditura, entrambe calcolate nel punto P .

L'ampiezza totale $\psi_A(B)$ nel punto P è, come già anticipato, dato dalla loro somma:

$$\psi_A(B) \equiv \psi(P) = \psi_1(P) + \psi_2(P)$$

(qui A indica la sorgente che omettiamo, mentre $B \equiv P$ è il punto sullo schermo). Poiché la probabilità di rivelazione nel punto P è proporzionale, come sappiamo, al modulo quadro dell'ampiezza totale, cioè $P(P) \propto |\psi(P)|^2$ si ha:

$$P(P) \propto |\psi_1(P) + \psi_2(P)|^2.$$

Sviluppando il modulo quadro dei valori complessi si ottiene infine (vedi l'Appendice B):

$$|\psi_1(P) + \psi_2(P)|^2 = |\psi_1(P)|^2 + |\psi_2(P)|^2 + 2\text{Re}(\psi_1(P)\psi_2^*(P))$$

dove $\psi_2^*(P)$ è il complesso coniugato di $\psi_2(P)$ (vedi la Nota in basso) e dove l'ultimo termine:

$$2\text{Re}(\psi_1(P)\psi_2^*(P)) \propto A\cos(\varphi_1 - \varphi_2)$$

è proprio il *termine di interferenza*, proporzionale ad $A\cos(\varphi_1 - \varphi_2)$ (dove $\varphi_1 - \varphi_2$ è la differenza di fase tra le due funzioni d'onda). Se questo termine è positivo, la probabilità di rivelazione aumenta; se è negativo, la probabilità diminuisce. L'alternanza di massimi e minimi di probabilità produce, evento dopo evento, la figura di interferenza osservata sullo schermo.

Questa è la differenza fondamentale rispetto a una descrizione classica ingenua. Se le due alternative fossero *distinguibili*, cioè se fosse possibile stabilire con una misura da quale fenditura è passato il fotone, ci aspetteremmo una somma di probabilità classica pari a

$$P(P) = P_1(P) + P_2(P).$$

Ma quando la traiettoria non è definita, la meccanica quantistica richiede invece:

$$P(P) \propto |\psi_1(P) + \psi_2(P)|^2$$

e la figura di interferenza nasce proprio da questo modulo quadro della somma di ampiezze.

Nota: un numero complesso può essere espresso come $\psi = a + ib$ (dove $i^2 = -1$) e il suo complesso coniugato è per definizione: $\psi^* = a - ib$ da cui segue $|\psi|^2 = \psi\psi^* = a^2 + b^2$.

La figura alternata delle frange

Vediamo quindi come nasce la figura delle frange alternate nel caso quantistico.

Come si può facilmente dimostrare, in un cosiddetto *apparato ottico lineare* come quello della doppia fenditura, lo stato di un fotone si propaga secondo gli stessi *modi spaziali e temporali* del campo elettromagnetico classico (vedi l'*Appendice D*).

In queste condizioni, l'ampiezza di probabilità $\psi(P)$ nel regime di singolo fotone ha la stessa *dipendenza spaziale* del campo elettrico classico $E(P)$, a meno di una costante di normalizzazione. Perciò si può scrivere schematicamente:

$$\psi(P) \propto E(P) .$$

Ne segue che la probabilità $P(P)$ di rivelazione quantistica, proporzionale a

$$P(P) \propto |\psi(P)|^2 \propto |E(P)|^2$$

ha la stessa forma dell'intensità classica, poiché come abbiamo già visto risulta:

$$I(P) \propto |E(P)|^2 .$$

La figura di interferenza è quindi la stessa, anche se cambia completamente il significato fisico: nel caso classico è una distribuzione continua di intensità, nel caso quantistico è una distribuzione di probabilità per eventi localizzati.

Il significato fisico dell'esperimento

Come ormai dovrebbe essere chiaro, nel modello classico l'interferenza riguarda grandezze fisiche ondulatorie, come il campo elettrico. Nel modello quantistico, invece, l'interferenza riguarda ampiezze di probabilità: ciò che appare sullo schermo non è una funzione d'onda, ma una successione di eventi localizzati determinati in modo probabilistico.

La teoria quantistica descrive quindi due aspetti che sembrano difficili da conciliare nel linguaggio classico: da un lato, ogni rivelazione è un evento localizzato; dall'altro, l'insieme di molte rivelazioni segue una distribuzione di interferenza. Questo comportamento *non classico* è ciò che di solito viene indicato come *dualismo onda-particella*.

Le frange osservate sono simili, ma il significato fisico cambia: nella teoria classica derivano dalla sovrapposizione di onde elettromagnetiche; nella teoria quantistica emergono dall'accumulo statistico di eventi individuali, governati dal modulo quadro dell'ampiezza totale.

Esperimenti analoghi sono stati realizzati anche con *particelle dotate di massa*, come elettroni, neutroni, atomi e molecole. Anche in questi casi le frange possono essere descritte dalla funzione d'onda $\psi(\mathbf{r}, t)$ soluzione dell'equazione di Schrödinger, senza dover ricorrere necessariamente alla formulazione equivalente della somma sui cammini di Feynman.

A livello interpretativo, si può forse parlare di *campo di informazione* della funzione d'onda: non è un vero e proprio campo materiale come il campo elettrico, ma più precisamente una struttura matematica che, sia nel caso dei fotoni che nel caso delle particelle massive, *codifica le probabilità* dei risultati osservabili (come chiariremo meglio di seguito).

Approfondimento: un unico stato quantistico

Si osservi che l'ampiezza di probabilità ψ , essendo una funzione complessa, può essere espressa come è noto in forma polare (essendo $e^{i\varphi} = \cos\varphi + i\sin\varphi$):

$$\psi = Ae^{i\varphi}$$

dove $A = |\psi|$ è il modulo dell'ampiezza e dove φ è la *fase quantistica* (vedi l'Appendice B).

Nell'esperimento della doppia fenditura abbiamo introdotto due ampiezze di probabilità, $\psi_1(P)$ e $\psi_2(P)$, associate rispettivamente alla *prima* e alla *seconda* fenditura. Si noti che queste due ampiezze rappresentano esclusivamente dei *contributi* alla medesima ampiezza complessiva $\psi(P)$ e non sono due *stati fisici distinti* posseduti contemporaneamente dalla particella.

D'altra parte, quando entrambe le fenditure sono aperte e non viene effettuata alcuna misura del cammino, le due alternative sono *indistinguibili* e perciò contribuiscono allo stesso evento di rivelazione. In questo caso la probabilità di rivelare la particella, o il singolo fotone, in una piccola regione intorno al punto P è proporzionale a

$$P(P) \propto |\psi_1(P) + \psi_2(P)|^2$$

dove, come dicevamo sopra, possiamo scrivere:

$$\psi_1(P) = A_1 e^{i\varphi_1} \quad \text{e} \quad \psi_2(P) = A_2 e^{i\varphi_2}.$$

Ora, nel caso ideale di fenditure identiche e illuminate simmetricamente, possiamo porre:

$$A_1 = A_2.$$

Ciò significa che, considerate separatamente, le due fenditure forniscono contributi di *uguale modulo* all'ampiezza di probabilità (oltre al termine di interferenza), risultando infatti:

$$|\psi_1(P)| = |\psi_2(P)|$$

dato che i fattori di fase hanno modulo unitario: $|e^{i\varphi_1}| = |e^{i\varphi_2}| = 1$ (vedi la *Nota* in basso).

Tuttavia, la scelta di scomporre l'ampiezza in due soli contributi non è unica. Infatti, potremmo suddividere le aperture, ad esempio, in N piccole zone e associare a ciascuna di esse un contributo $\psi_i(P)$ con $i = 1, 2, 3, \dots, N$ e quindi ottenere una probabilità proporzionale a

$$P(P) \propto |\psi_1(P) + \psi_2(P) + \psi_3(P) + \dots + \psi_N(P)|^2.$$

In definitiva, ciò che non cambia è proprio la funzione d'onda ψ , il cui valore nel punto P è definito come la somma dei diversi contributi. Essa rappresenta un *unico stato quantistico*, anche se può essere espressa come somma di contributi diversi. È esclusivamente da essa che possiamo ottenere, come abbiamo visto, la probabilità di rivelazione nel punto P :

$$P(P) \propto |\psi(P)|^2.$$

Nota: ricordiamo che vale la seguente relazione per il modulo dell'ampiezza di probabilità:

$$|\psi| = |Ae^{i\varphi}| = |A||e^{i\varphi}| = A \text{ poiché } |e^{i\varphi}| = |\cos\varphi + i\sin\varphi| = \sqrt{(\cos^2\varphi + \sin^2\varphi)} = 1.$$

Il problema quantistico della misura

La scomparsa dell'interferenza

Come già osservato, nella doppia fenditura l'interferenza compare quando le alternative associate alle due fenditure restano *indistinguibili*. Se invece un apparato permette di stabilire, con una misura, da quale fenditura è passato il fotone, non si osserva più la figura di interferenza: le frange scompaiono e la distribuzione sullo schermo cambia.

Questo è il punto in cui emerge il cosiddetto *problema della misura*. In meccanica quantistica, prima della rivelazione, il sistema viene descritto mediante ampiezze di probabilità: la teoria non assegna necessariamente una traiettoria definita, ma calcola le probabilità dei possibili risultati. Nel laboratorio, però, ciò che osserviamo è sempre e solo un risultato *determinato*: un punto sullo schermo, un segnale in un rivelatore, un evento localizzato.

La domanda diventa allora: come si passa dalla descrizione quantistica, fatta di possibilità sovrapposte, al singolo risultato concreto registrato dall'apparato?

È importante chiarire che, in fisica, *misura* non significa semplicemente presenza di un osservatore cosciente. Una misura è un'interazione fisica capace di lasciare una traccia, cioè di produrre una informazione accessibile sul sistema. Nel caso della doppia fenditura, se un rivelatore registra quale fenditura è stata attraversata, il sistema quantistico diventa in pratica *correlato* con l'apparato di misura, per questo motivo le due alternative non sono più indistinguibili e la coerenza necessaria all'interferenza viene distrutta.

Tuttavia resta la questione irrisolta di come, in una singola misura, dalla descrizione quantistica fatta di possibilità sovrapposte, emerga un unico risultato misurato dall'apparato.

Il concetto di decoerenza

Questo *processo di correlazione* durante la misura è collegato al concetto di *decoerenza*. Quando un sistema quantistico interagisce con un apparato macroscopico, o più in generale con l'ambiente, le diverse alternative possibili possono diventare di fatto *distinguibili*. La sovrapposizione allora non produce più effetti di interferenza osservabili, perché l'informazione sul cammino è stata, per così dire, *distribuita fisicamente* nell'apparato o nell'ambiente.

Il problema della misura, tuttavia, non è soltanto tecnico. La decoerenza spiega perché l'interferenza diventa praticamente inosservabile nei sistemi *macroscopici*, ma resta aperta la *questione interpretativa* più profonda: perché, in una singola misura, osserviamo un risultato definito e non una sovrapposizione di stati corrispondenti a risultati diversi?

La doppia fenditura rende questo problema particolarmente evidente. Finché non esiste una informazione definita sul cammino del singolo oggetto, il sistema si comporta in modo da produrre interferenza. Quando invece l'informazione sul cammino diventa disponibile, anche solo in linea di principio, l'interferenza scompare: è qui che emerge il problema della misura.

In definitiva, l'esperimento mostra che non possiamo separare completamente il fenomeno osservato dalle condizioni sperimentali. La meccanica quantistica non descrive una traiettoria definita della particella prima della misura, tuttavia fornisce le regole per calcolare le probabilità dei risultati che possono essere osservati, in un determinato apparato sperimentale.

Comprendere la meccanica quantistica

Alla fine di questa breve incursione nel mondo della meccanica quantistica, possiamo essere certi, almeno nel caso dell'esperimento della doppia fenditura, di riuscire finalmente a comprendere questo particolare *fenomeno quantistico*?

La risposta dipende da che cosa intendiamo per *comprendere*. Se comprendere significa ricondurre il fenomeno a immagini familiari — una particella che segue una traiettoria precisa, oppure un'onda materiale che si propaga nello spazio come un'onda sull'acqua — allora la doppia fenditura continua a sfuggire al nostro linguaggio ordinario. Il comportamento quantistico non coincide con nessuna di queste immagini prese separatamente. Ma se comprendere significa *capire e distinguere* con precisione ciò che osserviamo, ciò che calcoliamo e ciò che interpretiamo, allora l'esperimento diventa senz'altro più chiaro.

In definitiva, ciò che *osserviamo* sono eventi localizzati: punti sullo schermo, segnali in un rivelatore, tracce registrate da un apparato. Ciò che *calcoliamo*, invece, non sono traiettorie classiche, ma ampiezze di probabilità. Queste ampiezze si sommano quando le alternative sono indistinguibili: il loro modulo quadro fornisce la probabilità dei risultati osservabili. Ciò che *interpretiamo*, infine, è il significato fisico del formalismo, come chiariremo di seguito.

La funzione d'onda come “campo di informazione”

La funzione d'onda e più in generale l'ampiezza di probabilità, può forse essere interpretata come un *campo di informazione*: non una sostanza materiale distribuita nello spazio, ma una struttura matematica che contiene l'informazione necessaria per prevedere i possibili risultati di una misura. Essa non dice in generale *dove si trova* davvero la particella, come farebbe una mappa classica; dice invece con quale probabilità potremo ottenere un *certo* risultato se predisponiamo un *certo* apparato sperimentale.

Tutto ciò richiede uno sforzo concettuale. Siamo abituati a pensare che la fisica debba descrivere oggetti dotati di proprietà definite in ogni istante: posizione, traiettoria, velocità, cammino seguito. La meccanica quantistica standard ci obbliga invece a cambiare prospettiva. Prima della misura, la teoria non ci autorizza ad attribuire al sistema una storia classica definita. Ci fornisce piuttosto una regola rigorosa per calcolare le probabilità degli eventi osservati.

La doppia fenditura mostra questo cambiamento nel modo più semplice e radicale. Se non misuriamo da quale fenditura passa il sistema, le ampiezze di probabilità associate alle due alternative interferiscono; se invece rendiamo disponibile con una misura l'informazione sul cammino, l'interferenza scompare. Tuttavia, non è una questione di ignoranza soggettiva, come se la particella avesse comunque seguito una traiettoria nascosta che noi non conosciamo. È il formalismo stesso a dirci che alternative indistinguibili devono essere trattate in sovrapposizione, mentre alternative distinguibili producono risultati diversi.

In questo senso, la meccanica quantistica non è incomprensibile: è solo *non classica*. Non ci chiede di rinunciare alla razionalità, ma di rinunciare all'idea che ogni fenomeno debba essere spiegabile con immagini tratte dall'esperienza quotidiana. La sua comprensione passa attraverso un nuovo linguaggio: ampiezze, sovrapposizione, probabilità, misura, informazione.

Quindi aveva ragione Feynman?

La frase di Feynman citata nella *Prefazione*, allora, può essere letta non come una resa, ma come un avvertimento. Nessuno *capisce* la meccanica quantistica se pretende di trasformarla completamente in concetti classici. Tuttavia, possiamo comprenderne molto bene la logica interna, le regole di calcolo, la struttura sperimentale e attribuirle un significato fisico.

L'esperimento della doppia fenditura resta per questo un caso esemplare. Nel modello classico, le frange sono la conseguenza dell'interferenza tra campi elettrici. Nel modello quantistico, le stesse frange emergono dall'interferenza tra ampiezze di probabilità e dall'accumulo statistico di eventi individuali. La figura osservata può essere la stessa: ciò che cambia è il livello di descrizione fisica a cui aspiriamo.

Comprendere la meccanica quantistica significa dunque accettare questa distinzione: il mondo non sempre si lascia descrivere come un insieme di traiettorie definite, ma può essere previsto con straordinaria precisione attraverso una teoria delle ampiezze di probabilità. La doppia fenditura ci insegna proprio questo: il *reale quantistico* non è meno rigoroso del *reale classico*, ma richiede un diverso modo di pensare e di rapportarsi alla realtà fisica.

La differenza tra formalismo e interpretazione

Un punto essenziale è, senz'altro, quello di distinguere tra il formalismo della meccanica quantistica e le sue possibili interpretazioni. Il *formalismo* è l'insieme delle regole matematiche che permettono di calcolare correttamente i risultati degli esperimenti: funzione d'onda, ampiezze di probabilità, operatori, modulo quadro, evoluzione secondo l'equazione di Schrödinger, regole di misura. Questo nucleo formale è straordinariamente preciso e condiviso.

Le *interpretazioni*, invece, cercano di rispondere a una domanda ulteriore: che cosa significa fisicamente quel formalismo? Che cosa rappresenta davvero la funzione d'onda? Che cosa accade durante una misura? Il sistema possiede proprietà definite prima dell'osservazione? La sovrapposizione è una realtà fisica, una descrizione della nostra informazione, oppure qualcosa di ancora diverso?

È per questo che esistono diverse interpretazioni della meccanica quantistica: l'interpretazione di Copenaghen, quella a molti mondi, le teorie a variabili nascoste non locali, le letture relazionali, informative o decoerenti. Esse differiscono nel significato fisico attribuito alla teoria, ma non cambiano le previsioni sperimentali quando usano lo stesso formalismo.

In questo senso possiamo comprendere l'atteggiamento pragmatico spesso riassunto nell'espressione iconica dei fisici: *Zitto e calcola!* La meccanica quantistica funziona, permette di prevedere gli esperimenti con enorme accuratezza e il formalismo può essere usato in modo preciso anche senza risolvere, definitivamente, tutte le questioni interpretative.

Tuttavia, per il nostro desiderio di interpretazione dei fenomeni fisici, non basta soltanto calcolare. Vogliamo anche *capire*, nel senso prima esposto, che cosa stiamo calcolando e perché la doppia fenditura ci costringe a ripensare concetti apparentemente semplici come traiettoria, misura e probabilità, che nella meccanica quantistica assumono un significato diverso da quello della realtà fisica ordinaria, in cui siamo d'altra parte abituati a vivere.

Parte Terza

Gli esperimenti classici e quantistici messi in pratica

Le unità di misura utilizzate

Nei prossimi esperimenti utilizzeremo, per svolgere i relativi calcoli, quattro delle sette unità di base del *Sistema Internazionale*:

m (metro); kg (chilogrammo); s (secondo); A (ampere).

Oltre ad alcune unità di misura da esse derivate:

$Hz = \frac{1}{s}$ (hertz); $N = \frac{kg \cdot m}{s^2}$ (newton); $J = N \cdot m$ (joule); $C = A \cdot s$ (coulomb); $V = \frac{J}{C}$ (volt).

Inoltre, unità di energia molto comune in fisica atomica e quantistica è l'elettronvolt:

$$1 \text{ eV} = 1,60 \cdot 10^{-19} \text{ C} \cdot 1 \frac{J}{C} = 1,60 \cdot 10^{-19} J$$

pari all'energia acquisita da una carica elementare $e = 1,60 \cdot 10^{-19} \text{ C}$ che attraversa una differenza di potenziale di $1V$. Questa unità non fa parte del *Sistema Internazionale* ma viene accettata proprio per esprimere energie molto piccole rispetto al joule.

Una simulazione realistica degli esperimenti

Dalla teoria alla pratica: come si procede in concreto nelle misurazioni

Dopo aver descritto l'esperimento della doppia fenditura dal punto di vista teorico, può essere utile vedere come esso venga effettivamente tradotto in una procedura di laboratorio, anche per capire quali sono le grandezze di misura in gioco.

La fisica, come sappiamo, non consiste soltanto nel formulare equazioni, ma anche e soprattutto nel confrontare quelle equazioni con misure reali al fine di verificare i nostri modelli; misure ottenute con strumenti concreti e inevitabilmente affetti da incertezze.

Consideriamo quindi una configurazione di laboratorio semplice ma realistica.

Utilizziamo allo scopo un laser rosso, con lunghezza d'onda pari a circa

$$\lambda = 6,50 \cdot 10^{-7} m = 650 \text{ nm} \quad (1 \text{ nm} = 10^{-9} m)$$

cioè una radiazione visibile nella regione del rosso ($\text{nm} = \text{nanometro}$).

Il fascio viene inviato su uno schermo con doppia fenditura (comunemente disponibile per esperimenti didattici), ognuna di larghezza dell'ordine di λ per avere l'effetto di diffrazione, dove la distanza d tra i centri delle due fenditure è pari a

$$d = 2,50 \cdot 10^{-4} m = 0,25 \text{ mm}.$$

Mentre lo schermo di rivelazione viene posto a una distanza dalla doppia fenditura pari a

$$L = 2,00 \text{ m}.$$

L'esperimento viene eseguito in un ambiente parzialmente oscurato, in modo da aumentare il contrasto delle frange. Il laser, la doppia fenditura e lo schermo sono montati su supporti regolabili, così da poter allineare con cura l'apparato.

Per misurare la distanza L si può semplicemente usare un metro a nastro, mentre per misurare la distanza tra le frange si può usare un righello millimetrato posto sullo schermo, oppure si può fare una fotografia, nella quale sia presente un riferimento di lunghezza noto.

Abbiamo scelto questi valori poiché da essi nasce un risultato misurabile a occhio nudo o con strumenti semplici dato che la distanza ΔY tra le frange è dell'ordine di qualche millimetro. Infatti, come abbiamo già mostrato nel testo, il modello ondulatorio classico prevede per piccoli angoli di osservazione, la seguente relazione:

$$\Delta Y \simeq \frac{\lambda L}{d}.$$

Sostituendo i valori scelti si ottiene perciò:

$$\Delta Y \simeq \frac{(6,50 \times 10^{-7} m)(2,00 m)}{2,50 \times 10^{-4} m} = 5,20 \times 10^{-3} m = 5,2 \text{ mm}$$

quindi il valore teorico previsto per la distanza tra due frange è pari a circa *5,2 millimetri*.

A questo punto possiamo passare alle misure di laboratorio per verificare il nostro modello.

Gli errori di misura nel caso classico

In laboratorio non conviene misurare direttamente la distanza tra due sole frange consecutive. Una singola lettura può essere influenzata dallo spessore delle frange, dal contrasto non perfetto, dalla risoluzione del righello o da un piccolo errore di allineamento. È preferibile misurare la distanza complessiva tra più frange e poi dividere per il numero di intervalli.

Per esempio, misurando la distanza complessiva corrispondente a dieci intervalli tra frange successive e dividendo poi il valore ottenuto per 10, si ottiene una stima media più affidabile della distanza tra due frange consecutive. Questa procedura riduce l'errore relativo della misura, perché l'incertezza di lettura viene distribuita su più intervalli.

Naturalmente, nessuna misura sperimentale coincide esattamente con il valore teorico. Le principali fonti di incertezza sono diverse, compresi gli errori sistematici:

- ➔ **il laser** può avere una lunghezza d'onda non esattamente pari a 650 nm ;
- ➔ **la doppia fenditura** può avere una tolleranza sulla distanza d ;
- ➔ **la misura della distanza L** può contenere un errore di qualche millimetro;
- ➔ **la posizione delle frange** può non essere determinata con precisione assoluta;
- ➔ **l'allineamento dell'apparato**, la luminosità ambientale e il contrasto delle frange possono influenzare la qualità della lettura.

Il confronto tra teoria ed esperimento consiste quindi nel verificare se il valore misurato risulta compatibile, entro le incertezze, con il valore previsto dalla formula. Se, per esempio, dalla misura su dieci intervalli si ottiene una distanza media tra frange pari a circa 5 mm , il risultato è in buon accordo con la previsione teorica di $5,2\text{ mm}$ (con un errore relativo del 4% circa).

In questo caso il valore misurato è perciò compatibile, entro le incertezze sperimentali, con quello previsto dal modello ondulatorio classico dell'interferenza.

Perché ci limitiamo alla misura di ΔY

Si osservi che, nell'esperimento della doppia fenditura, ci siamo limitati alla misura della distanza ΔY tra massimi consecutivi, senza cioè considerare l'intensità luminosa che varia sullo schermo, poiché si tratta di una grandezza geometrica facilmente determinabile e direttamente confrontabile con la previsione teorica: $\Delta Y \simeq \frac{\lambda L}{d}$.

D'altra parte, una verifica quantitativa dell'intero profilo dell'intensità richiederebbe un sensore con risposta nota e lineare, il controllo del rumore di fondo e dell'efficienza di rivelazione, nonché la considerazione dell'involuppo di diffrazione dovuto alla larghezza delle fenditure.

Inoltre, nel regime a singolo fotone, anche se ogni rivelazione costituisce un evento localizzato, aumentando il numero degli eventi la loro distribuzione tende, entro le fluttuazioni statistiche, allo stesso profilo spaziale normalizzato previsto per l'intensità classica.

Perciò la concordanza del valore misurato di ΔY con quello teorico verifica la periodicità prevista dai due modelli e, insieme alla formazione progressiva delle frange mediante singoli eventi, costituisce una significativa, sebbene parziale, conferma della descrizione quantistica.

La doppia fenditura nel regime dei singoli fotoni

Dopo aver simulato l'esperimento della doppia fenditura con un fascio laser sufficientemente intenso, possiamo chiederci che cosa accade quando la sorgente viene ridotta di intensità fino al regime dei singoli fotoni. La geometria dell'apparato rimane sostanzialmente la stessa, ma cambia il modo in cui il fenomeno viene registrato.

Nel caso classico, lo schermo mostra quasi subito una figura continua di frange chiare e scure. Nel regime quantistico, invece, ogni fotone produce un singolo evento di rivelazione localizzato. Solo accumulando moltissimi eventi individuali emerge progressivamente la stessa distribuzione di interferenza.

Consideriamo quindi la medesima configurazione sperimentale prima utilizzata per il caso classico, con la stessa sorgente luminosa rossa con $\lambda = 650 \text{ nm}$ ma fortemente attenuata.

Il fascio viene inviato sulla stessa doppia fenditura con distanza tra i centri delle due aperture pari a $d = 0,25 \text{ mm}$ e il sistema di rivelazione viene posto alla stessa distanza $L = 2,00 \text{ m}$.

In un esperimento a singolo fotone, però, non basta un normale schermo bianco. Occorre un rivelatore capace di registrare eventi luminosi individuali, per esempio una camera molto sensibile a bassa intensità. L'apparato deve inoltre essere ben schermato dalla luce ambientale, perché anche un debole fondo luminoso può alterare il conteggio degli eventi.

L'attenuazione dell'intensità della sorgente deve essere tale che nell'apparato sia presente, in media, un solo fotone alla volta. Questo *non* significa che possiamo osservare il fotone mentre attraversa l'apparato: ciò che viene registrato è soltanto l'evento finale di rivelazione.

La previsione sulla posizione delle frange resta la stessa del caso classico. Infatti, come abbiamo visto nella sezione dedicata al calcolo quantistico, in un *apparato ottico lineare* l'ampiezza di probabilità del singolo fotone ha la stessa dipendenza spaziale del campo elettrico classico. Cambia però il significato fisico: nel caso classico si misura una distribuzione continua di intensità; nel caso quantistico si misura una distribuzione statistica di eventi localizzati.

La distanza attesa ΔY tra i massimi della distribuzione è quindi ancora pari alla precedente:

$$\Delta Y \simeq \frac{\lambda L}{d}.$$

Perciò, con gli stessi valori sperimentali del caso classico, anche nel regime dei singoli fotoni ci aspettiamo che i massimi della distribuzione siano separati da circa $5,2 \text{ mm}$.

Il dato sperimentale, tuttavia, non appare subito come una frangia continua. All'inizio il sensore registra pochi punti apparentemente irregolari. Aumentando il tempo di acquisizione, cresce il numero di eventi registrati e la distribuzione diventa progressivamente più ordinata.

Semplicemente, le zone in cui i fotoni vengono rivelati più spesso corrispondono ai massimi; le zone in cui vengono rivelati raramente corrispondono ai minimi. Nel limite statistico di un numero molto grande di eventi, la distribuzione può essere considerata praticamente continua, rendendo così possibile individuare le frange e misurare la distanza tra i loro massimi.

Gli errori di misura nel caso quantistico

Per confrontare il risultato con la previsione teorica si misura ancora la distanza tra massimi successivi. Anche in questo caso conviene considerare più intervalli, misurando la distanza complessiva corrispondente a n intervalli tra massimi successivi.

Dividendo la misura complessiva per n si ottiene una stima media della distanza tra due massimi consecutivi. Come nel caso classico, se il valore ottenuto è vicino a $5,2 \text{ mm}$, la distribuzione osservata è compatibile, entro le incertezze sperimentali, con la previsione teorica.

Le fonti di incertezza sono in parte le stesse dell'esperimento classico: lunghezza d'onda λ non perfettamente nota, tolleranza sulla distanza d tra le due fenditure, misura della distanza L tra gli schermi e allineamento dell'apparato.

Nel regime di singolo fotone si aggiungono però altri aspetti da considerare:

- ➔ **il rumore intrinseco** dell'apparato rivelatore;
- ➔ **i conteggi spuri** di fondo del sistema;
- ➔ **l'efficienza** del sensore utilizzato;
- ➔ **il tempo di acquisizione** e quindi il numero totale di eventi registrati.

Se il numero di fotoni rivelati è piccolo, la distribuzione può apparire casuale. Aumentando il numero di eventi, le relative fluttuazioni statistiche si riducono e la figura tende progressivamente alla distribuzione prevista: il singolo evento ovviamente non mostra la frangia; molti eventi, invece, ne ricostruiscono la forma, componendo la figura continua prevista dalla teoria.

Poiché, in generale, la *fluttuazione statistica relativa* dei conteggi diminuisce con l'inverso della radice quadrata del numero N di eventi (cioè va come $\frac{1}{\sqrt{N}}$), maggiore è il numero di dati raccolti, più precisa risulta la stima e quindi più attendibile la verifica del modello teorico.

È inoltre opportuno ripetere più volte l'esperimento, anche smontando e rimontando i componenti o azzerando gli strumenti di misura, sia per *mediare* eventuali errori accidentali ma anche per individuare eventuali errori *sistematici* dell'apparato.

In definitiva, con questo esperimento a bassa intensità, si mette in evidenza:

- ➔ il *dato sperimentale*: cioè la sequenza di eventi localizzati registrati dal sensore e la loro distribuzione spaziale sullo schermo;
- ➔ il *modello teorico*: la distribuzione di probabilità ottenuta dal modulo quadro dell'ampiezza totale, come derivato nel testo;
- ➔ l'*interpretazione fisica*: le frange sullo schermo derivano dall'interferenza delle ampiezze di probabilità, che sono associate alle alternative indistinguibili.

La concordanza tra la distanza misurata e quella prevista tra i massimi mostra che il regime quantistico non elimina la figura di interferenza. La conserva, ma ne cambia il significato fisico: ciò che nel caso classico appare come intensità continua, nel caso quantistico emerge dall'accumulo statistico di molti singoli eventi di rivelazione.

Una misura complementare: l'effetto fotoelettrico

Dopo aver discusso la doppia fenditura, possiamo considerare brevemente anche l'effetto fotoelettrico, non per ripetere la descrizione teorica già svolta, ma per capire come si traduce in pratica la misura dell'energia dei quanti di luce e i valori in gioco.

Come abbiamo già visto l'apparato sperimentale è costituito da una cella fotoelettrica, sotto vuoto, con un fotocatodo metallico e un anodo collettore. Il fotocatodo viene illuminato con una sorgente luminosa di lunghezza d'onda nota; nel circuito esterno si misurano la corrente elettrica prodotta dagli elettroni emessi e la tensione applicata tra catodo e anodo.

Rendendo l'anodo negativo rispetto al fotocatodo, si applica una tensione frenante: aumentando questa tensione, la corrente diminuisce fino ad annullarsi completamente. Il valore per cui la corrente si annulla è la tensione di arresto V_s .

La relazione fondamentale come abbiamo visto è la seguente:

$$h\nu = eV_s + \Phi$$

dove $h\nu$ è l'energia del fotone incidente, Φ è il lavoro di estrazione dal metallo ed eV_s è l'energia cinetica massima degli elettroni emessi ricavata in laboratorio.

Supponiamo quindi di illuminare il fotocatodo con luce violetta di lunghezza d'onda:

$$\lambda = 4,00 \cdot 10^{-7} \text{ m} = 400 \text{ nm} \quad (1 \text{ nm} = 10^{-9} \text{ m}).$$

L'energia del fotone può essere ricavata dalla relazione:

$$E = h\nu \simeq (6,63 \cdot 10^{-34} \text{ J} \cdot \text{s})(7,50 \cdot 10^{14} \text{ Hz}) \simeq 4,97 \cdot 10^{-19} \text{ J}$$

dove $h = 6,63 \cdot 10^{-34} \text{ J} \cdot \text{s}$ è la costante di Planck, mentre la frequenza ν dell'onda è:

$$\nu = \frac{c}{\lambda} \simeq \frac{3,00 \cdot 10^8 \frac{\text{m}}{\text{s}}}{4,00 \cdot 10^{-7} \text{ m}} \simeq 7,50 \cdot 10^{14} \text{ Hz}.$$

A questo punto è utile convertire l'energia da joule in elettronvolt; utilizzando la relazione nota: $1 \text{ eV} = 1,60 \cdot 10^{-19} \text{ J}$ si ottiene:

$$E \simeq \frac{4,97 \cdot 10^{-19}}{1,60 \cdot 10^{-19}} \simeq 3,10 \text{ eV}.$$

Perciò se in laboratorio, realisticamente, si misura una tensione di arresto pari a circa:

$$V_s \simeq 1,10 \text{ V}$$

allora l'energia cinetica massima degli elettroni (cioè quelli posti in superficie) è pari a:

$$K_{max} = eV_s \simeq 1,10 \text{ eV}$$

e il lavoro di estrazione del metallo si ricava dalla relazione $\Phi = E - K_{max}$ perciò:

$$\Phi \simeq 3,10 - 1,10 \simeq 2,00 \text{ eV}.$$

Come si vede sono in gioco valori molto bassi di energia, tipici dei processi atomici.

Una volta noto il lavoro di estrazione Φ , si può ricavare la frequenza di soglia ν_0 , cioè la frequenza minima necessaria per avere emissione di elettroni (cioè per $K_{max} = 0$):

$$\nu_0 = \frac{\Phi}{h}.$$

Inoltre, poiché risulta $\nu_0 = c/\lambda_0$, in termini di lunghezza d'onda λ_0 si ottiene (utilizzando il valore di hc convertito in elettronvolt per metro: $hc = \frac{1,99 \cdot 10^{-25}}{1,60 \cdot 10^{-19}} = 1,24 \cdot 10^{-6} \text{ eV} \cdot \text{m}$):

$$\lambda_0 = \frac{hc}{\Phi} \simeq \frac{1,24 \cdot 10^{-6} \text{ eV} \cdot \text{m}}{2,00 \text{ eV}} \simeq 620 \text{ nm}.$$

Ciò significa che radiazioni con lunghezza d'onda maggiore a circa 620 nm , quindi meno energetiche, non riescono a estrarre elettroni dal metallo (cioè la corrente si annulla).

Per esempio, una luce rossa di 650 nm ha energia:

$$E = \frac{hc}{\lambda} \simeq \frac{1,24 \cdot 10^{-6} \text{ eV} \cdot \text{m}}{650 \cdot 10^{-9} \text{ m}} \simeq 1,91 \text{ eV} < 2 \text{ eV}.$$

Perciò, poiché l'energia è inferiore al lavoro di estrazione prima ricavato, non ci aspettiamo emissione di elettroni, qualunque sia l'intensità della luce (come conferma l'esperimento).

Come è facile prevedere, in questo esperimento le principali incertezze riguardano:

- ➔ **la lunghezza d'onda** della sorgente;
- ➔ **la qualità di vuoto** nella cella fotoelettrica;
- ➔ **la misura della tensione** di arresto;
- ➔ **la sensibilità** dell'amperometro che misura la corrente;
- ➔ **la pulizia della superficie metallica** e l'eventuale presenza di luce ambientale.

Verifica del calcolo non relativistico

Facciamo ora una verifica teorica che ci permetta di capire se il modello di calcolo non relativistico che abbiamo usato è adeguato. A questo scopo ricaviamo la velocità massima degli elettroni emessi, utilizzando la nota relazione: $K_{max} = \frac{1}{2} m_e v_{max}^2$ da cui segue

$$v_{max} = \sqrt{\frac{2K_{max}}{m_e}}.$$

Nel nostro esempio come abbiamo visto $K_{max} = 1,10 \text{ eV} = 1,76 \cdot 10^{-19} \text{ J}$ ed essendo la massa dell'elettrone $m_e = 9,11 \cdot 10^{-31} \text{ kg}$ si ottiene:

$$v_{max} = \sqrt{\left(\frac{2 \cdot 1,76 \cdot 10^{-19}}{9,11 \cdot 10^{-31}}\right)} \simeq 6,2 \cdot 10^5 \frac{\text{m}}{\text{s}}$$

che è circa lo 0,2% della velocità della luce, essendo $c \simeq 3,00 \cdot 10^8 \frac{\text{m}}{\text{s}}$. È sicuramente una velocità molto grande su scala ordinaria, ma ancora molto minore della velocità della luce, perciò il calcolo svolto non relativistico risulta adeguato.

Appendici

Sviluppo di alcuni passaggi matematici richiamati nel testo

Appendice A

Interferenza classica e intensità luminosa

In questa appendice mostriamo come, dalla somma di due onde sfasate, si ottiene il campo elettrico risultante e come da questo si ricava la dipendenza dell'intensità luminosa dalla fase.

Successivamente, richiamiamo il significato delle equazioni di Maxwell nel vuoto e mostriamo in breve come da esse si ottengano le equazioni d'onda del campo elettromagnetico, oltre a definire la direzione dei campi E e B durante la propagazione.

La somma algebrica di due onde armoniche

Nel testo abbiamo descritto formalmente i due campi elettrici provenienti dalle due fenditure (che fungono da sorgenti), nella forma semplice di due *onde piane*:

$$E_1(x, t) = E_0 \cos(kx - \omega t)$$

$$E_2(x, t) = E_0 \cos(kx - \omega t + \phi) .$$

Come risulta evidente le due onde hanno la stessa ampiezza E_0 , la stessa lunghezza d'onda λ e la stessa pulsazione ω , ma sono sfasate di una quantità ϕ . Ora, come già spiegato, il campo elettrico in interazione è la somma algebrica dei due campi:

$$E_{int}(x, t) = E_1(x, t) + E_2(x, t) = E_0 [\cos(kx - \omega t) + \cos(kx - \omega t + \phi)] .$$

Per sviluppare questa somma possiamo usare l'identità trigonometrica:

$$\cos \alpha + \cos \beta = 2 \cos\left(\frac{\alpha + \beta}{2}\right) \cos\left(\frac{\alpha - \beta}{2}\right)$$

dove poniamo per il nostro caso: $\alpha = (kx - \omega t)$ e $\beta = (kx - \omega t + \phi)$.

Quindi sostituendo i valori di α e β si ha: $\frac{\alpha - \beta}{2} = -\frac{\phi}{2}$ e $\frac{\alpha + \beta}{2} = (kx - \omega t + \frac{\phi}{2})$ ed essendo $\cos\left(-\frac{\phi}{2}\right) = \cos\left(\frac{\phi}{2}\right)$ si ottiene infine:

$$E_{int}(x, t) = 2E_0 \cos\left(\frac{\phi}{2}\right) \cos\left(kx - \omega t + \frac{\phi}{2}\right) .$$

Questa è la formula richiamata nel testo; essa mostra che la somma di due onde sinusoidali della stessa pulsazione è ancora un'onda sinusoidale, ma con ampiezza E_{int}^0 dipendente dallo sfasamento ϕ essendo:

$$E_{int}^0 = 2E_0 \cos\left(\frac{\phi}{2}\right)$$

cioè le onde si rafforzano o si elidono in funzione dello sfasamento, come chiarito nel testo.

Nota: si osservi che se il termine $\cos\left(\frac{\phi}{2}\right)$ fosse negativo, significherebbe semplicemente che l'onda risultante è sfasata di π poiché vale la relazione: $-\cos\left(\frac{\phi}{2}\right) = \cos\left(\frac{\phi}{2} + \pi\right)$.

Dall'ampiezza del campo all'intensità luminosa

Nel modello classico, ciò che lo schermo registra non è direttamente il valore istantaneo del campo elettrico, che oscilla molto rapidamente nel tempo, ma l'*intensità luminosa media* associata all'onda. In particolare l'intensità luminosa può essere intesa come la quantità di energia trasportata dall'onda per unità di tempo e per unità di superficie e si misura in $\left[\frac{J}{m^2 \cdot s}\right]$.

Nel caso di un'onda elettromagnetica, questa grandezza è legata al vettore di Poynting (vedi la *Nota* in basso), ma per il nostro scopo, è sufficiente ricordare che l'intensità è proporzionale al valore medio temporale del quadrato del campo elettrico:

$$I \propto \langle E_{int}^2(x, t) \rangle.$$

Nel nostro caso il campo elettrico in interazione, prima ricavato ed elevato al quadrato, è:

$$E_{int}^2(x, t) = 4E_0^2 \cos^2\left(\frac{\phi}{2}\right) \cos^2\left(kx - \omega t + \frac{\phi}{2}\right).$$

Per fare la media temporale, notiamo che il fattore $4E_0^2 \cos^2\left(\frac{\phi}{2}\right)$ non dipende dal tempo, quindi resta fuori dalla media. Rimane da calcolare, su un periodo completo T , l'altro fattore:

$$\langle \cos^2\left(kx - \omega t + \frac{\phi}{2}\right) \rangle = \frac{1}{T} \int_0^T \cos^2\left(kx - \omega t + \frac{\phi}{2}\right) dt.$$

Ora si osservi che, valendo l'identità trigonometrica $\cos^2 u = \frac{1 + \cos 2u}{2}$ si ottiene:

$$\frac{1}{T} \int_0^T \cos^2 u \, dt = \frac{1}{T} \int_0^T \frac{1 + \cos 2u}{2} dt = \frac{1}{2}$$

poiché il primo termine dà $\frac{1}{2}$ mentre il termine oscillante ha *media nulla* su un periodo completo. Quindi posto $u = (kx - \omega t + \frac{\phi}{2})$ segue infine, come riportato nel testo:

$$I(\phi) \propto \langle E_{int}^2(x, t) \rangle = 2E_0^2 \cos^2\left(\frac{\phi}{2}\right).$$

Questa formula mostra che l'intensità osservata sullo schermo dipende dallo sfasamento tra le due onde. Quando $\phi = 0$, le onde sono in fase e l'intensità è massima; quando $\phi = \pi$, le onde sono in opposizione di fase e l'intensità si annulla.

Nota: nel modello elettromagnetico classico, l'intensità luminosa è legata al flusso di energia trasportato dall'onda. Questo flusso è definito dal vettore di Poynting: $\mathbf{S} = \left(\frac{1}{\mu_0}\right) \mathbf{E} \times \mathbf{B} \left[\frac{J}{m^2 \cdot s}\right]$.

Per un'onda piana nel vuoto, i campi $\mathbf{E}$ e $\mathbf{B}$ sono perpendicolari e vale $B = \frac{E}{c}$ quindi il modulo del vettore di Poynting risulta: $S \propto E^2$. Mentre l'intensità sullo schermo è data dal valore *medio temporale* di questo flusso: $I = \langle S \rangle \propto \langle E^2 \rangle$. Perciò, nel nostro caso, dove il campo elettrico da considerare è il campo $E_{int}(x, t)$ si ha come visto sopra: $I \propto \langle E_{int}^2(x, t) \rangle = 2E_0^2 \cos^2\left(\frac{\phi}{2}\right)$.

Le equazioni di Maxwell nel vuoto

È noto che Maxwell riorganizzò in un unico quadro teorico diverse leggi già note dell'elettricità e del magnetismo. Il suo contributo decisivo fu però l'introduzione del termine di *corrente di spostamento* nella legge di Ampère. Grazie a questo termine, le equazioni prevedono che un campo elettrico variabile nel tempo possa generare un campo magnetico, anche nel vuoto: da ciò emerge l'esistenza delle onde elettromagnetiche, identificate da Maxwell con la luce.

Nel testo abbiamo ricordato che la luce, nella teoria classica, è un'onda elettromagnetica, formata da un campo elettrico $\mathbf{E}$ e da un campo magnetico $\mathbf{B}$. Il punto di partenza sono le equazioni di Maxwell nel vuoto (che poi il fisico Oliver Heaviside ridusse alla moderna forma vettoriale e simmetrica); in assenza di cariche e correnti elettriche e magnetiche esse sono:

$$\text{Legge di Gauss per il campo } \mathbf{E}: \quad \nabla \cdot \mathbf{E} = 0$$

$$\text{Legge di Gauss per il campo } \mathbf{B}: \quad \nabla \cdot \mathbf{B} = 0$$

$$\text{Legge di Faraday-Lenz:} \quad \nabla \times \mathbf{E} = -\frac{\partial \mathbf{B}}{\partial t}$$

$$\text{Legge di Ampère-Maxwell:} \quad \nabla \times \mathbf{B} = \mu_0 \varepsilon_0 \frac{\partial \mathbf{E}}{\partial t}.$$

Gli operatori $\nabla \cdot$ e $\nabla \times$, applicati ai campi vettoriali $\mathbf{E}$ o $\mathbf{B}$, definiscono il primo la *divergenza* (che misura la tendenza locale del campo ad *entrare/uscire* da una regione dello spazio) e il secondo il *rotore* (che misura la tendenza locale del campo a *circolare* intorno a un asse).

Le equazioni d'onda dei campi $\mathbf{E}$ e $\mathbf{B}$

Mostriamo ora, in forma sintetica, come dalle equazioni di Maxwell si ottiene l'equazione d'onda per il campo elettrico (qui non mostreremo il caso *analogo* per il campo magnetico).

Partiamo dalla legge di Faraday-Lenz ed applichiamo il rotore $\nabla \times$ a entrambi i membri:

$$\nabla \times (\nabla \times \mathbf{E}) = -\frac{\partial}{\partial t} (\nabla \times \mathbf{B}).$$

Ora, grazie alla legge di Ampère-Maxwell, possiamo sostituire $\nabla \times \mathbf{B}$ con $\mu_0 \varepsilon_0 \frac{\partial \mathbf{E}}{\partial t}$ e quindi:

$$\nabla \times (\nabla \times \mathbf{E}) = -\mu_0 \varepsilon_0 \frac{\partial^2 \mathbf{E}}{\partial t^2}.$$

A questo punto usiamo l'identità vettoriale: $\nabla \times (\nabla \times \mathbf{E}) = \nabla(\nabla \cdot \mathbf{E}) - \nabla^2 \mathbf{E}$ ricordando però che nel vuoto si ha $\nabla \cdot \mathbf{E} = 0$ (per la legge di Gauss) quindi risulta:

$$\nabla \times (\nabla \times \mathbf{E}) = -\nabla^2 \mathbf{E}.$$

Perciò, sostituendo nell'equazione precedente ed eliminando i segni meno, si ottiene:

$$\nabla^2 \mathbf{E} = \mu_0 \varepsilon_0 \frac{\partial^2 \mathbf{E}}{\partial t^2}$$

che è proprio l'*equazione d'onda* per il campo elettrico, come indicato nel testo.

Che cosa si intende con onda piana

Nel testo abbiamo rappresentato il campo elettrico con una forma semplificata:

$$E(x, t) = E_0 \cos(kx - \omega t + \phi)$$

questa è formalmente una *onda piana monocromatica* che si propaga lungo l'asse X .

Si parla di onda piana perché, a ogni istante, i punti che hanno la stessa fase stanno su piani *perpendicolari* alla direzione di propagazione. È una idealizzazione molto utile quando si osserva la luce abbastanza lontano dalla sorgente e i fronti d'onda sono quasi paralleli.

Ricordiamo che il termine $(kx - \omega t + \phi)$ è la fase dell'onda, perciò i piani di fase costante avanzano nello spazio con velocità $v = \frac{\lambda}{T} = \frac{\omega}{k}$ e nel vuoto, per un'onda elettromagnetica:

$$v = \frac{\omega}{k} = c.$$

Quindi l'espressione usata nel testo rappresenta un'onda luminosa idealizzata, piana, sinusoidale, monocromatica e polarizzata (cioè che oscilla lungo una direzione fissata).

I campi $E \perp B$ e la direzione di propagazione

Mostriamo ora perché per un'onda elettromagnetica piana nel vuoto, le equazioni di Maxwell impongono che il campo elettrico E , il campo magnetico B e la direzione di propagazione data dal vettore d'onda k siano tra loro perpendicolari.

Consideriamo un'onda piana del campo elettrico $E(\mathbf{r}, t)$ che si propaga nella direzione del vettore d'onda k al variare del punto $\mathbf{r}$. Possiamo rappresentarla schematicamente come:

$$E(\mathbf{r}, t) = E_0 \cos(\mathbf{k} \cdot \mathbf{r} - \omega t).$$

Poiché nel vuoto $\nabla \cdot E = 0$ si può dimostrare che, applicando l'operatore divergenza, si ha:

$$\nabla \cdot E = -\mathbf{k} \cdot E_0 \sin(\mathbf{k} \cdot \mathbf{r} - \omega t) = 0$$

quindi per il campo elettrico e in modo analogo per il campo magnetico, deve risultare:

$$\mathbf{k} \cdot E_0 = 0 \quad \text{e} \quad \mathbf{k} \cdot B_0 = 0.$$

Queste due relazioni significano che sia E_0 che B_0 devono essere perpendicolari alla direzione di propagazione k quindi i campi oscillano in direzioni trasversali al moto dell'onda.

La perpendicolarità tra E_0 e B_0 deriva invece dalle equazioni di campo espresse con il rotore; cioè si può derivare sia dalla legge di Faraday-Lenz che da quella di Ampère-Maxwell.

Per un'onda piana, esse portano entrambe alla relazione:

$$B_0 = \frac{1}{\omega} \mathbf{k} \times E_0$$

dove k è il vettore d'onda diretto nella direzione di propagazione. Questa relazione mostra che B_0 è perpendicolare sia a k sia a E_0 dato che il prodotto vettoriale $k \times E_0$ è per definizione perpendicolare a entrambi. Inoltre, essendo come abbiamo visto sopra $\frac{\omega}{k} = c$ si ha $B_0 = \frac{E_0}{c}$.

Appendice B

Funzione d'onda complessa e probabilità

In questa appendice richiamiamo alcuni passaggi matematici utilizzati nella sezione quantistica del testo: il significato di operatore nell'equazione di Schrödinger, lo sviluppo del modulo quadro della somma di due ampiezze complesse e la fase nel caso classico e quantistico.

L'operatore hamiltoniano nell'equazione di Schrödinger

Nel testo abbiamo introdotto l'equazione di Schrödinger che, nella forma compatta, si può scrivere in modo equivalente come:

$$i\hbar \frac{\partial \psi(\mathbf{r}, t)}{\partial t} = \hat{H}\psi(\mathbf{r}, t)$$

dove $\psi(\mathbf{r}, t)$ è la funzione d'onda nel punto $\mathbf{r}$ e al tempo t , $\hbar = \frac{h}{2\pi}$ è la costante di Planck ridotta e $\hat{H}$ è l'operatore hamiltoniano.

In matematica, un *operatore* è una regola che agisce su una funzione e produce una nuova funzione. In meccanica quantistica molte grandezze fisiche sono rappresentate da operatori: in particolare l'operatore hamiltoniano $\hat{H}$ è quello associato all'energia totale del sistema.

Ora, nel caso di una particella non relativistica di massa m soggetta a un potenziale $V(\mathbf{r})$, l'operatore hamiltoniano è così definito:

$$\hat{H} = -\left(\frac{\hbar^2}{2m}\right)\nabla^2 + V(\mathbf{r})$$

dove il primo termine rappresenta il contributo dell'energia cinetica (vedi la *Nota* in basso), mentre il secondo termine rappresenta l'energia potenziale. Perciò, quando l'operatore $\hat{H}$ agisce sulla funzione d'onda $\psi(\mathbf{r}, t)$, si ottiene:

$$\hat{H}\psi(\mathbf{r}, t) = -\left(\frac{\hbar^2}{2m}\right)\nabla^2\psi(\mathbf{r}, t) + V(\mathbf{r})\psi(\mathbf{r}, t)$$

e quindi la forma compatta dell'equazione sopra riportata diventa, in forma esplicita:

$$i\hbar \frac{\partial \psi(\mathbf{r}, t)}{\partial t} = -\left(\frac{\hbar^2}{2m}\right)\nabla^2\psi(\mathbf{r}, t) + V(\mathbf{r})\psi(\mathbf{r}, t).$$

La forma compatta è quindi un modo abbreviato per scrivere l'evoluzione temporale della funzione d'onda, raccogliendo nell'operatore $\hat{H}$ i termini che descrivono l'energia totale del sistema, dato dal contributo cinetico sommato a quello potenziale.

Nota: dalla nota relazione di de Broglie per la quantità di moto $\mathbf{p} = \hbar\mathbf{k}$ si ottiene, per un'onda piana $\psi = e^{i\mathbf{k}\cdot\mathbf{r}}$ (che descrive uno stato ideale di particella libera): $-i\hbar\nabla \cdot \psi = \hbar\mathbf{k}\psi = \mathbf{p}\psi$ da cui possiamo definire l'operatore $\hat{\mathbf{p}} = -i\hbar\nabla$ e quindi l'energia cinetica: $\hat{T} = \frac{\hat{\mathbf{p}}^2}{2m} = -\left(\frac{\hbar^2}{2m}\right)\nabla^2$.

Il modulo quadro della somma di due ampiezze complesse

Come spiegato nel testo, nel caso della doppia fenditura, se non si rileva da quale fenditura passa la particella o il singolo fotone, le due alternative restano indistinguibili. In questo caso non si sommano direttamente le probabilità (caso classico), ma le ampiezze di probabilità.

Indichiamo con $\psi_1(P)$ l'ampiezza associata alla prima fenditura e con $\psi_2(P)$ l'ampiezza associata alla seconda fenditura, entrambe valutate nel punto P dello schermo.

Come sappiamo l'ampiezza totale è data dalla somma delle singole ampiezze:

$$\psi(P) = \psi_1(P) + \psi_2(P)$$

mentre la probabilità di rivelazione della particella nel punto P è proporzionale al modulo quadro dell'ampiezza totale, secondo l'interpretazione di Born:

$$P(P) \propto |\psi(P)|^2$$

perciò si ha:

$$P(P) \propto |\psi_1(P) + \psi_2(P)|^2 .$$

Per sviluppare questa espressione, ricordiamo come si calcola il modulo quadro di un numero complesso $\psi = a + ib$ dove i è l'unità immaginaria (con $i^2 = -1$) e dove il suo *coniugato* complesso è così definito: $\psi^* = a - ib$.

Il modulo quadro è proprio espresso dal prodotto di ψ per il suo coniugato ψ^* :

$$|\psi|^2 = \psi\psi^*$$

quindi sostituendo i valori complessi si ha $|\psi|^2 = (a + ib)(a - ib)$ e sviluppando il prodotto:

$$|\psi|^2 = a^2 + iab - iab - i^2b^2 = a^2 + b^2 .$$

Si osservi che il modulo quadro è dunque una quantità reale e non negativa.

Applichiamo ora questa definizione alla somma delle due ampiezze della doppia fenditura e, per semplicità, omettiamo temporaneamente l'indicazione del punto P e scriviamo:

$$|\psi|^2 = |\psi_1 + \psi_2|^2 .$$

Perciò per la definizione di modulo quadro data sopra, cioè $|\psi|^2 = \psi\psi^*$, si ha:

$$|\psi_1 + \psi_2|^2 = (\psi_1 + \psi_2)(\psi_1 + \psi_2)^*$$

e poiché il coniugato della somma è la somma dei coniugati si ottiene:

$$|\psi_1 + \psi_2|^2 = (\psi_1 + \psi_2)(\psi_1^* + \psi_2^*) .$$

Infine sviluppando il prodotto risulta:

$$|\psi_1 + \psi_2|^2 = \psi_1\psi_1^* + \psi_1\psi_2^* + \psi_2\psi_1^* + \psi_2\psi_2^*$$

dove come già detto $\psi_1\psi_1^* = |\psi_1|^2$ e $\psi_2\psi_2^* = |\psi_2|^2$ perciò:

$$|\psi_1 + \psi_2|^2 = |\psi_1|^2 + |\psi_2|^2 + \psi_1\psi_2^* + \psi_2\psi_1^* .$$

Si osservi ora che gli ultimi due termini sono complessi coniugati tra loro cioè

$$(\psi_1 \psi_2^*)^* = \psi_2 \psi_1^*$$

e che in generale, per due numeri complessi coniugati risulta $(a + ib) + (a - ib) = 2a$ (cioè la loro somma è uguale al doppio della parte reale), si ottiene perciò:

$$|\psi_1 + \psi_2|^2 = |\psi_1|^2 + |\psi_2|^2 + 2\text{Re}(\psi_1 \psi_2^*).$$

In definitiva nel caso della doppia fenditura risulta che la probabilità di rivelazione della particella nel punto P , è proporzionale al modulo quadro dell'ampiezza totale:

$$P(P) \propto |\psi_1(P)|^2 + |\psi_2(P)|^2 + 2\text{Re}(\psi_1(P) \psi_2^*(P))$$

dove il termine:

$$2\text{Re}(\psi_1(P) \psi_2^*(P))$$

è proprio il termine di interferenza, come già spiegato nel testo. Ora, se le due alternative fossero distinguibili, si sommerebbero semplicemente le probabilità:

$$P(P) \propto |\psi_1(P)|^2 + |\psi_2(P)|^2$$

quando invece le alternative sono indistinguibili, si sommano prima le ampiezze e solo dopo si calcola il modulo quadro:

$$P(P) \propto |\psi_1(P) + \psi_2(P)|^2.$$

La differenza tra queste due procedure è fondamentale e determina il termine di interferenza.

A questo punto se poniamo: $\psi_1(P) = A_1 e^{i\varphi_1}$ e $\psi_2(P) = A_2 e^{i\varphi_2}$ segue subito:

$$2\text{Re}(\psi_1(P) \psi_2^*(P)) = 2\text{Re}(A_1 e^{i\varphi_1} A_2 e^{-i\varphi_2}) = 2\text{Re}(A_1 A_2 e^{i(\varphi_1 - \varphi_2)})$$

essendo il coniugato $\psi_2^*(P) = A_2 e^{-i\varphi_2}$ e quindi, considerando solo la parte reale, si ha:

$$2\text{Re}(\psi_1(P) \psi_2^*(P)) = 2A_1 A_2 \cos(\varphi_1 - \varphi_2).$$

La differenza tra fase classica e quantistica

Nel testo compaiono due tipi di *fase*, formalmente simili ma fisicamente diversi (vedi anche l'Appendice C). Nel caso classico di un'onda elettromagnetica, la fase compare ad esempio nella forma seguente (dove Re indica la parte reale della forma complessa):

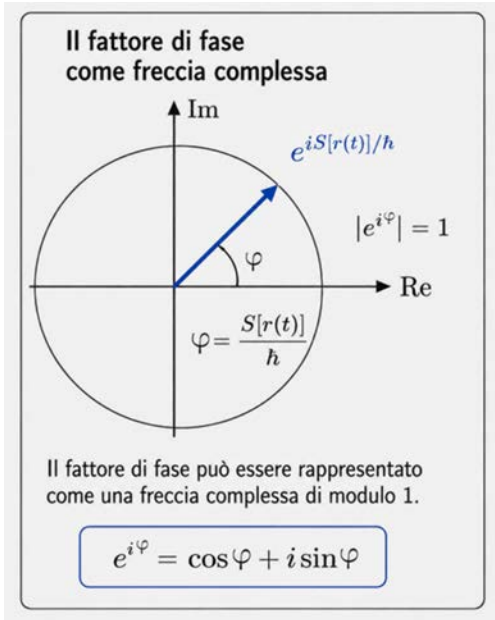
$$E(x, t) = \text{Re}[E_0 e^{i(kx - \omega t + \phi)}].$$

Qui la fase $(kx - \omega t + \phi)$ descrive l'oscillazione del campo elettrico nello spazio e nel tempo. Essa indica in quale punto del ciclo oscillatorio si trova l'onda e determina come due onde classiche si sommano o si cancellano. Nel caso quantistico, invece, nella formulazione di Feynman compare il fattore di fase:

$$e^{\frac{iS[\mathbf{r}(t)]}{\hbar}}.$$

Qui la fase non descrive l'oscillazione di un campo classico osservabile, ma il contributo complesso associato a un possibile cammino. La quantità $S[\mathbf{r}(t)]$ è l'azione del cammino e il rapporto $\frac{S[\mathbf{r}(t)]}{\hbar}$ determina la *fase dell'ampiezza di probabilità*.

La somiglianza matematica è importante: in entrambi i casi le fasi possono produrre interferenza. Ma il significato fisico cambia: nel caso classico interferiscono campi fisici, come il campo elettrico; nel caso quantistico interferiscono ampiezze di probabilità associate a possibilità alternative, come nel caso dei diversi cammini di Feynman, già trattati nel testo.



Il fattore di fase quantistico rappresentato nel piano complesso.

Si noti che il fattore di fase $e^{\frac{iS[r(t)]}{h}}$ può essere rappresentato come una *freccia* nel piano complesso di modulo 1 e direzione determinata dalla fase $\varphi = \frac{S[r(t)]}{h}$. La somma sui cammini consiste allora, schematicamente, nella somma di molte freccette: quelle con fasi simili si rafforzano, mentre quelle con fasi molto diverse tendono a cancellarsi.

A ogni cammino che collega il punto iniziale A al punto finale B corrisponde generalmente un diverso valore dell'azione S e quindi una freccia con fase diversa. La risultante della somma vettoriale di tutte queste frecce rappresenta l'ampiezza complessiva del processo; solo dopo aver effettuato la somma se ne calcola il modulo quadro per ottenere la probabilità.

Nel limite classico, i contributi dei cammini vicini a quello dell'azione stazionaria tendono ad avere fasi simili e quindi a rafforzarsi, mentre quelli più lontani tendono a cancellarsi reciprocamente. Infatti, come spiegato nell'*Appendice C*, la condizione di stazionarietà $\delta S = 0$ individua la traiettoria classica percorsa dalla particella.

Appendice C

Lagrangiana, azione e integrale sui cammini

In questa appendice richiamiamo alcuni strumenti matematici utilizzati nel testo, in particolare per la formulazione di Feynman dell'integrale sui cammini.

Non riprendiamo il significato fisico generale già discusso nel testo, ma ci limitiamo a chiarire come si definiscono la lagrangiana, l'azione e la somma integrale sui cammini.

La lagrangiana $\mathcal{L}$

Nella meccanica classica, lo stato dinamico di una particella può essere descritto attraverso la sua posizione, la sua velocità e le forze a cui è soggetta. Un modo particolarmente generale per formulare le leggi del moto consiste nell'introdurre la *lagrangiana*, indicata con $\mathcal{L}$.

La grandezza $\mathcal{L}$ prende il nome da Joseph-Louis Lagrange, che riformulò la meccanica classica in una forma più generale rispetto alla descrizione newtoniana, nella quale le equazioni del moto vengono ricavate considerando direttamente le forze agenti e la relazione $\mathbf{F} = m\mathbf{a}$. Questa nuova formulazione, esposta nella *Mécanique analytique* del 1788, costituisce uno dei pilastri fondamentali della meccanica analitica.

In particolare, per un sistema classico non relativistico, la lagrangiana è definita semplicemente come la differenza tra l'energia cinetica T e l'energia potenziale V :

$$\mathcal{L} = T - V.$$

È importante osservare che la lagrangiana non coincide con l'energia totale del sistema; ricordiamo infatti che l'energia totale classica, nei casi più semplici, è: $E = T + V$.

Questa differenza non è arbitraria: la lagrangiana è costruita in modo tale che, attraverso il principio di *azione stazionaria* (vedi oltre), permetta di ricavare le equazioni del moto.

In particolare, come è noto, per una particella di massa m , con velocità v , soggetta a un potenziale $V(\mathbf{r})$, l'energia cinetica è:

$$T = \frac{1}{2}mv^2$$

quindi la lagrangiana può essere scritta in funzione di $\mathbf{r}$, v e t :

$$\mathcal{L}(\mathbf{r}, v, t) = \frac{1}{2}mv^2 - V(\mathbf{r}).$$

La formulazione lagrangiana ha il vantaggio di poter essere espressa in coordinate diverse: cartesiane, polari o, più in generale, coordinate generalizzate. La forma esplicita di $\mathcal{L}$ può cambiare, ma le equazioni del moto che si ottengono descrivono lo stesso contenuto fisico.

Questa formulazione mette in evidenza le simmetrie del sistema: da tali simmetrie discendono spesso grandezze conservate, come energia, quantità di moto e momento angolare.

Il caso relativistico e i fotoni

Nel caso relativistico, quando la velocità della particella non è più trascurabile rispetto alla velocità della luce c , la forma della lagrangiana cambia. Per una particella relativistica di massa m sottoposta ad un potenziale scalare $V(\mathbf{r})$ si usa, ad esempio:

$$\mathcal{L} = -mc^2 \sqrt{\left(1 - \frac{v^2}{c^2}\right)} - V(\mathbf{r}) .$$

Questa formula mostra che la lagrangiana dipende dal regime fisico considerato: la forma *non* relativistica vale per velocità piccole rispetto a c , mentre quella relativistica è necessaria quando le velocità diventano confrontabili con quella della luce.

Un fotone tuttavia ha massa nulla e nel vuoto si propaga sempre a velocità c ; non può quindi essere descritto correttamente mediante la lagrangiana relativistica di una particella massiva. La formulazione rigorosa appartiene alla teoria dei campi: la luce è descritta dal campo elettromagnetico e i fotoni sono i *quanti* di tale campo.

Per il nostro scopo, comunque, è sufficiente mantenere l'idea formale generale: una lagrangiana è la grandezza da cui si costruisce l'*azione* e questa è ciò che determina la *fase* dei contributi quantistici nella formulazione di Feynman, come specificheremo di seguito.

L'azione S

Una volta definita la lagrangiana, si introduce l'*azione*, indicata con S : l'azione è definita come l'integrale della lagrangiana lungo un cammino.

Se un sistema va da un punto iniziale A al tempo t_A , a un punto finale B al tempo t_B e segue un determinato cammino $\mathbf{r}(t)$, l'azione associata a quel cammino è:

$$S[\mathbf{r}(t)] = \int_{t_A}^{t_B} \mathcal{L}(\mathbf{r}, v, t) dt .$$

La notazione $S[\mathbf{r}(t)]$ indica che l'azione dipende dall'intero cammino $\mathbf{r}(t)$: cammini *diversi* che collegano gli stessi estremi A e B possono avere valori diversi di posizione, velocità e quindi di lagrangiana a ogni istante; di conseguenza possono avere azioni diverse.

In meccanica classica, il cammino effettivamente seguito dal sistema è quello per cui l'azione è detta *stazionaria*. Questo principio si scrive nella forma variazionale:

$$\delta S = 0 .$$

Si osservi che il simbolo δS indica la variazione dell'azione quando si confronta il cammino considerato con cammini infinitamente vicini. Dire che $\delta S = 0$ significa che, al primo ordine, piccole variazioni del cammino non modificano l'azione S : nel caso quantistico questa è la *traiettoria* più probabile seguita dalla particella (cioè quella dove si ha interferenza costruttiva).

Questo principio viene spesso chiamato *principio di minima azione*, ma il nome è impreciso: l'azione non deve essere necessariamente minima; deve essere stazionaria.

Dall'azione alla fase quantistica

Nella formulazione di Feynman, l'azione non serve soltanto a ricavare il cammino classico, essa entra direttamente nella *fase associata* a ciascun cammino possibile, infatti a ogni cammino $\mathbf{r}(t)$ viene associato un contributo complesso proporzionale a:

$$e^{\frac{iS[\mathbf{r}(t)]}{\hbar}}.$$

In questa espressione $S[\mathbf{r}(t)]$ è l'azione associata al cammino; $\hbar$ è la costante di Planck ridotta (cioè divisa per 2π); i è l'unità immaginaria (con $i^2 = -1$).

Si noti in particolare che il rapporto $\frac{S[\mathbf{r}(t)]}{\hbar}$ è adimensionale, cioè non ha unità di misura e può quindi comparire come argomento dell'esponenziale complesso definendo la *fase quantistica*.

La presenza dell'unità immaginaria i indica che il contributo associato al cammino è un'ampiezza complessa. Cammini con azioni simili hanno fasi simili e possono sommarsi in modo costruttivo; cammini con azioni molto diverse hanno fasi rapidamente variabili e tendono a cancellarsi nella somma complessiva (vedi la precedente figura del fattore di fase).

Come si costruisce la somma sui cammini

Per capire il significato dell'integrale sui cammini, conviene partire da un caso discreto.

Supponiamo che il sistema vada dal punto A al punto B nell'intervallo di tempo compreso tra t_A e t_B . Dividiamo questo intervallo in molti piccoli intervalli temporali:

$$t_A = t_0 < t_1 < t_2 < \dots < t_N = t_B.$$

Un cammino può allora essere approssimato da una successione di punti:

$$A = \mathbf{r}_0, \mathbf{r}_1, \mathbf{r}_2, \dots, \mathbf{r}_N = B$$

dove i punti intermedi $\mathbf{r}_1, \mathbf{r}_2, \dots, \mathbf{r}_{N-1}$ possono assumere molti valori possibili: ogni punto $\mathbf{r}_i = \mathbf{r}(t_i)$ rappresenta una possibile posizione del sistema in corrispondenza dell'istante t_i . Per sommare tutti i cammini, bisogna quindi sommare su tutte le possibili posizioni intermedie.

In modo schematico si può quindi scrivere un particolare integrale multiplo a $(N - 1)$ variabili:

$$\int D[\mathbf{r}(t)] \propto \lim_{N \rightarrow \infty} \int d\mathbf{r}_1 \int d\mathbf{r}_2 \dots \int d\mathbf{r}_{N-1}.$$

Questa espressione significa che l'integrale sui cammini può essere pensato come il *limite* di una somma continua su tutti i possibili punti intermedi attraversati dal sistema.

Nel limite in cui il numero di intervalli tende all'infinito, la successione discreta di punti diventa un cammino continuo $\mathbf{r}(t)$ e la somma su tutte le posizioni intermedie si indica, in simboli:

$$\int D[\mathbf{r}(t)].$$

Il simbolo $D[\mathbf{r}(t)]$ non rappresenta quindi un normale differenziale, ma indica una *somma funzionale*: non si integra su una singola variabile, ma su intere funzioni $\mathbf{r}(t)$ cioè su interi cammini. Infine, come indicato nel testo, a ciascun cammino viene associato un contributo complesso, cioè un fattore di fase dipendente dall'azione del cammino.

L'integrale sui cammini di Feynman

Completiamo ora la notazione introdotta nel precedente paragrafo; in definitiva l'ampiezza di probabilità per andare da A a B può essere così scritta, in forma simbolica:

$$\psi_A(B) \propto \int D[\mathbf{r}(t)] e^{\frac{iS[\mathbf{r}(t)]}{\hbar}}.$$

Questa formula riassume, nel caso di una particella non relativistica, l'idea centrale della formulazione di Feynman: l'ampiezza complessiva si ottiene sommando i contributi complessi associati a tutti i cammini possibili che collegano A a B . Il principio si può estendere anche ai fotoni, ma richiede la formulazione quantistica del campo elettromagnetico.

Riassumendo, i simboli dell'integrale (spesso indicato col termine $K(B, A)$ detto *propagatore*), hanno il seguente significato:

- ➔ $\psi_A(B)$ indica l'ampiezza di probabilità per andare da A a B ;
- ➔ $\int D[\mathbf{r}(t)]$ indica la somma integrale su tutti i cammini possibili $\mathbf{r}(t)$;
- ➔ $e^{\frac{iS[\mathbf{r}(t)]}{\hbar}}$ è il contributo complesso associato a ciascun cammino;
- ➔ $S[\mathbf{r}(t)]$ è l'azione calcolata lungo quel cammino.

Il simbolo $\propto$ indica che la formula è scritta in forma proporzionale: in una trattazione completa compaiono fattori di normalizzazione e dettagli matematici che qui non riteniamo necessari.

Applicazione alla doppia fenditura

Nella doppia fenditura, l'ampiezza di probabilità associata a un punto P dello schermo, può essere vista come somma dei contributi associati ai cammini compatibili con l'apparato.

Quindi, se non si rivela da quale fenditura passa la particella, l'ampiezza $\psi_A(B)$ nel punto P è (qui A indica la sorgente che omettiamo, mentre $B \equiv P$ è il punto sullo schermo):

$$\psi_A(B) \equiv \psi(P) = \psi_1(P) + \psi_2(P)$$

dove $\psi_1(P)$ raccoglie tutti i contributi dei cammini che attraversano la *prima* fenditura (che indichiamo complessivamente con C_1); mentre $\psi_2(P)$ raccoglie tutti i contributi dei cammini che attraversano la *seconda* fenditura (indicati con C_2); formalmente possiamo scrivere:

$$\psi_1(P) \propto \int_{C_1} D[\mathbf{r}(t)] e^{\frac{iS[\mathbf{r}(t)]}{\hbar}} \quad \text{e} \quad \psi_2(P) \propto \int_{C_2} D[\mathbf{r}(t)] e^{\frac{iS[\mathbf{r}(t)]}{\hbar}}.$$

Infine, come già spiegato nell'Appendice B, la probabilità di rivelazione nel punto P si ottiene dal modulo quadro dell'ampiezza totale:

$$P(P) \propto |\psi(P)|^2 = |\psi_1(P) + \psi_2(P)|^2.$$

Questa è la forma matematica che collega la somma sui cammini alla distribuzione di probabilità osservata sullo schermo, come già discusso nel testo.

Nota: in linea di principio, la somma sui cammini potrebbe essere suddivisa anche in classi più fini e non soltanto nei due contributi $\psi_1(P)$ e $\psi_2(P)$, senza modificare l'ampiezza totale.

Appendice D

Fotoni, campo e.m. quantizzato e limite classico

In questa appendice proviamo a chiarire il rapporto tra campo elettromagnetico classico, singolo fotone e modi di propagazione, per arrivare alla figura di interferenza.

Come forse ricorderete nel testo abbiamo affermato:

[...] in un cosiddetto apparato ottico lineare come quello della doppia fenditura, lo stato di un fotone si propaga secondo gli stessi modi del campo elettromagnetico classico.

È proprio ciò che cercheremo di spiegare in questa appendice; ma vediamo innanzitutto cosa si intende con *apparato ottico lineare*.

Apparato ottico lineare e limite classico

Un apparato ottico si dice *lineare* quando rispetta il principio di sovrapposizione: se un primo campo produce un certo effetto e un secondo campo ne produce un altro, allora la somma algebrica dei due campi determina il campo risultante.

Nel regime classico, questo significa che i campi elettrici si sommano linearmente. Per esempio, come già visto nella doppia fenditura, il campo elettrico nel punto P dello schermo può essere scritto come somma dei contributi provenienti dalle due fenditure:

$$E(P) = E_1(P) + E_2(P)$$

mentre l'intensità luminosa osservata è proporzionale al modulo quadro del campo elettrico:

$$I(P) \propto |E(P)|^2 = |E_1(P) + E_2(P)|^2.$$

La linearità riguarda dunque la propagazione e la sovrapposizione dei campi elettrici, mentre la formazione dell'intensità richiede successivamente il modulo al quadrato.

Nel regime quantistico, come abbiamo visto, la stessa idea si traduce nel fatto che le ampiezze di probabilità associate ad alternative indistinguibili si sommano linearmente. Per un singolo fotone nella doppia fenditura, l'ampiezza di rivelazione nel punto P è così definita:

$$\psi(P) = \psi_1(P) + \psi_2(P)$$

mentre la probabilità di rivelazione si ottiene dal modulo quadro:

$$P(P) \propto |\psi(P)|^2 = |\psi_1(P) + \psi_2(P)|^2.$$

Anche qui la linearità riguarda la somma delle ampiezze, mentre la probabilità è data dal calcolo del modulo al quadrato.

Come descritto nel testo, il limite classico si ottiene quando il numero di fotoni è statisticamente molto grande: ogni fotone viene rivelato come evento localizzato, ma l'accumulo di moltissimi eventi segue una distribuzione continua. Questa distribuzione coincide al limite, ed entro le condizioni di validità del modello, con l'intensità prevista dalla teoria classica.

Ma per comprendere perché ciò accade, bisogna chiarire che cosa sono i *modi del campo*.

I modi del campo elettromagnetico classico

In fisica, un *modo del campo elettromagnetico* è una possibile configurazione *spaziale* del campo, che soddisfa le equazioni di Maxwell, insieme alle condizioni imposte dall'apparato sperimentale e, nel nostro caso, è caratterizzata da una ben definita frequenza di oscillazione.

Le condizioni dell'apparato sono importanti: fenditure, specchi, lenti, schermi e bordi determinano quali forme del campo possono propagarsi. Per questo si parla di *modi del campo*: non si tratta di un campo qualsiasi, ma di una forma di propagazione compatibile con l'apparato.

In forma semplificata, il campo elettrico classico sullo schermo può essere descritto come:

$$E(P) = Af(P)$$

dove A è un'ampiezza complessiva, P è un punto dello schermo e $f(P)$ è la *funzione spaziale* del modo considerato, in questo caso quello del campo elettromagnetico.

In pratica, la funzione spaziale $f(P)$ descrive come varia il campo da punto a punto sullo schermo: dove è più intenso, dove è più debole, dove cambia fase. Non indica necessariamente il valore assoluto del campo, ma la sua *forma spaziale*.

Ad esempio, nel caso della doppia fenditura, $f(P)$ contiene la struttura spaziale che produce le frange e il modo finale sullo schermo può essere pensato come la somma di due contributi:

$$f(P) = f_1(P) + f_2(P)$$

quindi il campo classico può essere scritto come:

$$E(P) = Af(P) = A[f_1(P) + f_2(P)] .$$

Il fotone segue gli stessi modi del campo classico

Passiamo ora alla descrizione quantistica. Nella teoria quantistica del campo elettromagnetico, il campo viene quantizzato, ma la quantizzazione avviene con gli stessi *modi normali* che compaiono nella teoria classica.

Questo punto è essenziale: la quantizzazione non cambia la forma spaziale dei modi, cambia invece la modalità in cui tali modi possono essere occupati. In particolare nel campo classico, un modo può avere ampiezza continua, mentre nel campo quantizzato, lo stesso modo può contenere solo un numero discreto di fotoni: $0, 1, 2, 3, \dots n$.

Un singolo fotone è quindi uno *stato quantistico* del campo, in cui un certo modo è occupato da un solo quanto. Possiamo indicare simbolicamente questo stato come:

$$|1_f >$$

che tradotto significa: *uno stato con un fotone nel modo f* .

Se il fotone si trova nel modo f , l'ampiezza di probabilità $\psi(P)$ di rivelarlo nel punto P dello schermo, ha la stessa dipendenza spaziale del modo $f(P)$, possiamo cioè scrivere:

$$\psi(P) = Cf(P)$$

dove C è una costante di normalizzazione.

Questa relazione *non* significa che $\psi(P)$ definisce il campo elettrico classico. Significa invece che $\psi(P)$ ed $E(P)$ hanno la stessa forma spaziale, perché sono entrambi costruiti a partire dallo stesso modo $f(P)$.

In definitiva, per quanto spiegato sopra risulta, nel caso quantistico e classico rispettivamente:

$$\psi(P) = C f(P) \quad \text{e} \quad E(P) = A f(P)$$

e quindi, a meno di una costante di normalizzazione:

$$\psi(P) \propto E(P) .$$

Perciò, tornando al caso della doppia fenditura, possiamo scrivere nel caso quantistico:

$$\psi(P) = C f(P) = C [f_1(P) + f_2(P)]$$

mentre nel caso classico:

$$E(P) = A f(P) = A [f_1(P) + f_2(P)] .$$

Ora, come già spiegato nel testo, la probabilità di rivelazione del singolo fotone è

$$P(P) \propto |\psi(P)|^2$$

quindi:

$$P(P) \propto |f_1(P) + f_2(P)|^2 .$$

Nel caso classico, invece, l'intensità luminosa è proporzionale a

$$I(P) \propto |E(P)|^2$$

perciò:

$$I(P) \propto |f_1(P) + f_2(P)|^2 .$$

Come conseguenza di queste relazioni di proporzionalità si deduce che

$$P(P) \propto I(P) .$$

Questa è la ragione per cui molti fotoni singoli, preparati in successione nello stesso stato, producono sullo schermo la stessa figura spaziale prevista dall'intensità classica.

Ciò ovviamente *non* significa che il fotone sia un'onda classica, significa invece che il *modo spaziale* che governa la probabilità di rivelazione del singolo fotone, è lo stesso modo del campo elettromagnetico classico.

La figura osservata sullo schermo può quindi essere la stessa, anche se cambia il significato fisico: si ha infatti una intensità continua nel modello classico e una distribuzione di probabilità per eventi discreti nel modello quantistico, come descritto nel testo.

// Con questa breve panoramica concludiamo la sezione dedicata alle *Appendici*, auspicando che quanto esposto possa servire da stimolo alla lettura di altri testi di approfondimento, alcuni dei quali sono indicati nella *Bibliografia essenziale* riportata di seguito. //

Alcune tappe della meccanica quantistica

Questa breve e incompleta cronologia mostra come la meccanica quantistica nasca in pratica dal problema dell'energia discreta, passi attraverso la natura corpuscolare della luce e arrivi alla funzione d'onda come strumento probabilistico.

Ciò che si evince è che l'esperimento della doppia fenditura si colloca proprio al centro di questo percorso: esso conserva l'immagine ondulatoria dell'interferenza ma, come abbiamo visto, la riformula in termini di ampiezze di probabilità e singoli eventi di rivelazione.

Studi e scoperte	Personaggi	Descrizione sintetica con equazioni
1900 Quantizzazione dell'energia.	Max Planck	Per spiegare la radiazione di corpo nero, Planck introduce l'idea che l'energia venga scambiata in quantità discrete: $E = h\nu$.
1905 Effetto fotoelettrico e quanti di luce.	Albert Einstein	La luce può essere assorbita in pacchetti discreti di energia, poi chiamati <i>fotoni</i> . L'energia del quanto luminoso γ è: $E_\gamma = h\nu$.
1913 Modello quantizzato dell'atomo.	Niels Bohr	Gli elettroni occupano livelli energetici discreti e l'emissione o assorbimento di luce avviene per salti di energia: $\Delta E = h\nu$.
1924 Onde di materia.	Louis de Broglie	Anche le particelle materiali sono associate a una lunghezza d'onda: $\lambda = \frac{h}{p}$. L'idea ondulatoria viene estesa dalla luce alla materia.
1925–1927 Meccanica delle matrici e principio di indeterminazione.	Werner Heisenberg	La teoria quantistica viene formulata attraverso grandezze osservabili, mentre il principio di indeterminazione stabilisce: $\Delta x \Delta p \geq \frac{h}{2}$.
1926 Meccanica ondulatoria.	Erwin Schrödinger	La dinamica quantistica viene descritta mediante la funzione d'onda ψ governata dall'equazione: $i\hbar \frac{\partial \psi}{\partial t} = \hat{H}\psi$.
1926 Interpretazione probabilistica.	Max Born	La funzione d'onda non è una grandezza osservabile ma il suo modulo quadro fornisce una densità di probabilità: $P \propto \psi ^2$.
1927–1928 Teoria relativistica dell'elettrone.	Paul Dirac	Dirac sviluppa il formalismo della teoria quantistica e formula un'equazione relativistica dell'elettrone, aprendo la strada all'antimateria.
1948 Integrale sui cammini.	Richard Feynman	La funzione d'onda di un evento si calcola sommando le ampiezze associate a tutti i cammini possibili da A a B: $\psi_A(B) \propto \int Dx e^{\frac{is}{\hbar}}$.
1964–2015 Entanglement e test di Bell.	John Bell Alain Aspect John Clauser Anton Zeilinger	Le disuguaglianze di Bell mostrano che la meccanica quantistica non è compatibile con una descrizione classica locale: gli esperimenti confermano le correlazioni quantistiche.

Bibliografia essenziale

I testi indicati di seguito non sono necessari per seguire il percorso del testo, ma possono essere utili per approfondire gli argomenti trattati: onde elettromagnetiche, funzione d'onda, ampiezze di probabilità, integrale sui cammini e campo elettromagnetico quantizzato.

La bibliografia è volutamente snella e privilegia, quando possibile, opere disponibili in italiano o in traduzione italiana, dal livello divulgativo a quello universitario specialistico.

1. Richard P. Feynman

QED. La strana teoria della luce e della materia – Adelphi.

Libro divulgativo-avanzato ma molto pertinente: luce, fotoni, ampiezze e somma dei cammini.

2. Richard P. Feynman, Robert B. Leighton, Matthew Sands

La fisica di Feynman – Zanichelli.

Testo universitario introduttivo, offre il quadro generale: onde, elettromagnetismo, quantistica.

3. David J. Griffiths, Darrell F. Schroeter

Introduzione alla meccanica quantistica – Ambrosiana/Zanichelli.

Testo universitario su funzione d'onda, equazione di Schrödinger, regola di Born e operatori.

4. Paul A. M. Dirac

I principi della meccanica quantistica – Bollati Boringhieri.

Testo avanzato: è un classico teorico, più astratto e concettualmente fondativo.

5. J. J. Sakurai, Jim Napolitano

Meccanica quantistica moderna – Zanichelli.

Testo tecnico sul formalismo moderno, stati, operatori, notazione di Dirac.

6. Francis W. Sears

Ottica - Zanichelli.

Introduttivo-universitario per la parte classica: interferenza, diffrazione, ottica ondulatoria.

7. Richard P. Feynman, Albert R. Hibbs

Quantum Mechanics and Path Integrals – Dover Publications.

Testo di approfondimento avanzato e specifico sull'integrale sui cammini.

8. Rodney Loudon

The Quantum Theory of Light – Oxford University Press.

Testo avanzato specialistico: fotoni, campo e.m. quantizzato, ottica quantistica.

*Credo di poter dire con sicurezza che nessuno
capisce la meccanica quantistica.*
— Richard P. Feynman

La frase del famoso fisico Richard Feynman è tra le più celebri della fisica moderna. Compare nelle sue lezioni raccolte in *The Character of Physical Law* e viene spesso citata per esprimere il carattere *controintuitivo* della teoria quantistica.

Questo libretto nasce da quella provocazione. Non per sostenere che la meccanica quantistica sia incomprensibile, né per presentarla come un insieme di paradossi misteriosi, ma per seguire con attenzione uno degli esperimenti più semplici, profondi e rappresentativi della fisica classica e quantistica: *l'esperimento della doppia fenditura*.

La scelta ovviamente non è casuale. La doppia fenditura permette di osservare, in una forma essenziale ma esauriente, il passaggio da una descrizione classica della luce a una descrizione quantistica della radiazione e, più in generale, della materia.

Avvertenza: anche se il testo contiene varie formule matematiche non preoccupatevi! Queste formule servono soprattutto a far capire che, dietro a spiegazioni di sole parole, ci sono strumenti matematici capaci di dare risposte precise alle nostre domande teoriche e sperimentali: in questo senso vanno interpretate e, se proprio non riuscite a comprenderle, affrontatele come un ostacolo che potete anche saltare... continuando a correre! ;-)

Conoscenze di base: nozioni generali su derivate, integrali e numeri complessi.

Alcuni dei temi trattati li trovate approfonditi nel blog:
www.significatofisico.it