

Richard P. Feynman

QED

LA STRANA TEORIA DELLA LUCE
E DELLA MATERIA



ADELPHI

Richard P. Feynman

QED

LA STRANA TEORIA DELLA LUCE
E DELLA MATERIA



ADELPHI EDIZIONI

Titolo originale

QED
The Strange Theory
of Light and Matter

Traduzione di Francesco Nicodemi
© 1985 BY RICHARD P. FEYNMAN
© 1989 ADELPHI EDIZIONI S.P.A. MILANO
ISBN 88-459-0719-8

FINITO DI STAMPARE NEL NOVEMBRE 1989
IN AZZATE DAL CONSORZIO ARTIGIANO «L.V.G.»

Printed in Italy

INDICE

Premessa di Leonard Mautner □

Prefazione di Ralph Leighton □

Ringraziamenti □

QED

1. Introduzione □

2. I fotoni: particelle di luce □

3. Gli elettroni e le loro interazioni □

4. Conclusioni indefinite □

Note aggiunte in bozze □

PREMESSA

Le Alix G. Mautner Memorial Lectures sono state concepite per onorare la memoria di mia moglie Alix, morta nel 1982. Oltre ai suoi interessi professionali, che riguardavano la letteratura inglese, Alix coltivò sempre una vivissima curiosità per molti settori della scienza. Sulla spinta del suo esempio, è stato quindi deciso di creare una *fondazione* a suo nome avente come scopo di finanziare un ciclo di conferenze, da tenersi ogni anno, il cui obiettivo sia di comunicare a un pubblico intelligente e avvertito lo spirito e le conquiste della scienza.

Con mio enorme piacere Richard Feynman ha accettato di tenere la prima serie di queste conferenze. L'amicizia tra me e Richard è ormai vecchia di cinquantacinque anni, essendo nata quando eravamo bambini insieme a Far Rockaway, nello Stato di New York. Quella tra lui e Alix cominciò ventidue anni fa. In nome di questa amicizia Alix a lungo cercò di persuaderlo a formulare un'esposizione della fisica delle particelle che fosse comprensibile a lei e a tutti i non fisici.

Un sentito ringraziamento va a tutti coloro che hanno generosamente contribuito alla Fondazione Alix G. Mautner, aiutando così a rendere possibili queste conferenze.

Leonard Mautner
Los Angeles, California
Maggio 1983

PREFAZIONE

Richard Feynman è leggendario nel mondo della fisica per il suo modo di guardare alla realtà: la sua assenza di preconcetti e la sua indipendenza di pensiero gli hanno sovente permesso di raggiungere una comprensione nuova e profonda del comportamento della natura, e di darne descrizioni che sono insieme semplici, eleganti e stimolanti.

Feynman è anche famoso per il suo entusiasmo di maestro. Rifiuta innumerevoli offerte di tenere conferenze per società e organizzazioni prestigiose, ma è sempre disponibile per lo studente che vada a chiedergli di parlare al club di fisica del suo liceo.

Questo libro è un'impresa che non ci risulta sia mai stata tentata prima: esporre in modo piano e senza infingimenti, per un pubblico di non specialisti, un argomento non poco difficile come la teoria dell'elettrodinamica quantistica. Lo scopo è di dare al lettore non specialista un'idea delle forme di pensiero a cui i fisici sono ricorsi per poter spiegare il comportamento della Natura.

Il lettore che stia progettando di studiare fisica, o che lo stia già facendo, lungi dal dover «disimparare» qualcosa leggendo questo libro, vi troverà una descrizione completa e accurata in ogni dettaglio di una struttura teorica alla quale è possibile collegare, senza alcuna modifica, concetti più avanzati. A coloro che hanno già studiato fisica, questo libro rivelerà che cosa in realtà si nascondeva dietro tutti i loro calcoli complicati.

La curiosità per il calcolo infinitesimale venne al giovanissimo Richard Feynman da un libro che si apriva con queste parole: «Ciò che è possibile a uno sciocco è possibile a tutti gli altri». Egli vuole dedicare questo libro ai suoi lettori con parole simili: «Ciò che è comprensibile a uno sciocco, è comprensibile a tutti gli altri».

Ralph Leighton
Pasadena, California
Febbraio 1985

RINGRAZIAMENTI

Questo libro contiene il testo delle lezioni sull'elettrodinamica quantistica da me tenute all'UCLA (University of California, Los Angeles), trascritte e redatte dal mio buon amico Ralph Leighton; alla sua preziosa esperienza di insegnante e di scrittore devo i molti miglioramenti del testo originale che permettono di presentare questa parte fondamentale della fisica a un pubblico più vasto.

Molte «volgarizzazioni» scientifiche raggiungono un'apparente semplicità solo a costo di descrivere qualcosa di diverso da ciò che affermano di descrivere, e anzi di notevolmente distorto. Il rispetto per l'argomento trattato non ci ha permesso di fare altrettanto. Attraverso molte ore di discussione ci siamo sforzati di raggiungere la massima semplicità e trasparenza, rinunciando però a qualsiasi compromesso che portasse a una distorsione della verità.

Richard P. Feynman

QED

LA STRANA TEORIA DELLA LUCE
E DELLA MATERIA

I

INTRODUZIONE

Alix Mautner aveva grande curiosità per la fisica e spesso mi chiedeva di spiegarle questa o quella cosa. All'inizio tutto filava liscio, come nell'ora del giovedì con i miei studenti del Caltech,¹ ma ogni volta finivamo con l'arenarci proprio su quella che secondo me è la parte più interessante: le folli idee della meccanica quantistica. Per spiegarle, le ripetevo, ci voleva assai più tempo che non un'ora o una serata, e così le promisi che un giorno avrei preparato una serie di lezioni su questo argomento.

Quando le lezioni furono pronte, andai a provarle in Nuova Zelanda, la cui lontananza dal resto del mondo mi tranquillizzava: se avessero fatto fiasco, poco male! Invece ebbero successo, sicché mi sento autorizzato a credere che funzionino (se non altro in Nuova Zelanda!). Eccole dunque raccolte in questo libro. Purtroppo Alix, a cui erano destinate, ora non può più ascoltarle.

Contrariamente a quanto si fa di solito in lezioni del genere, la parte della fisica di cui parlerò è già conosciuta. Il pubblico è sempre curioso di sapere cosa c'è di nuovo nell'unificazione di questa teoria con quell'altra teoria. Vuole sempre sapere ciò che non sappiamo, e non ci dà mai la possibilità di parlare delle teorie che conosciamo bene. Così, invece di confondervi le idee con un mucchio di teorie parzialmente elaborate, di teorie ancora 'mezze-crude', ho scelto un argomento che è stato analizzato a fondo e che appartiene a un'area della fisica che io trovo meravigliosa: l'elettrodinamica quantistica o, per brevità, QED.²

In queste lezioni mi propongo innanzitutto di descrivere, con la maggior accuratezza possibile, la strana teoria della luce e della materia o, specificamente, l'interazione tra la luce e gli elettroni. Le cose da spiegare sono molte, ma poiché abbiamo a disposizione quattro lezioni, me la potrò prendere con calma e vedrete che tutto andrà bene.

La fisica, per lunga tradizione, ama sintetizzare molti fenomeni in poche teorie. Per esempio, nei primi tempi della fisica moderna i fenomeni del moto e quelli del calore erano considerati distinti tra loro, così come i fenomeni del suono, della luce e della gravità. Ma dopo che Sir Isaac Newton ebbe spiegato le leggi del moto, si scoprì ben presto che alcuni di questi fenomeni in apparenza diversi erano in realtà aspetti della stessa cosa. Per esempio, il suono poteva essere completamente spiegato come moto degli atomi dell'aria, sicché non venne più considerato come qualcosa di diverso dal moto. Poi si scoprì che anche i fenomeni del calore sono facilmente comprensibili sulla base delle leggi del moto, e fu così che intere sezioni della fisica vennero unificate in una teoria più semplice. La teoria della gravitazione, invece, non era comprensibile a partire dalle leggi del moto e ancor oggi resta separata dalle altre teorie. La gravitazione, almeno sinora, non è spiegabile in termini di altri fenomeni.

Dopo la sintesi del moto, del suono e del calore, vi fu la scoperta di un gran numero di fenomeni chiamati elettrici e magnetici. Nel 1873 James Clerk Maxwell li sintetizzò tutti in un'unica teoria con la luce e con i fenomeni ottici, interpretando la luce come onda elettromagnetica. Si era in tal modo arrivati a tre sistemi distinti di leggi: le leggi del

moto, le leggi dell'elettromagnetismo e le leggi della gravità.

Intorno al 1900, per spiegare la costituzione della materia, venne elaborata una teoria, nota come «teoria dell'elettrone», secondo la quale all'interno degli atomi vi sono delle piccole particelle cariche, appunto gli elettroni. Essa si sviluppò poi gradualmente fino a culminare nell'ipotesi che l'atomo è formato da un nucleo pesante e dagli elettroni che gli girano attorno.

Si trattava, a questo punto, di spiegare il moto degli elettroni attorno al nucleo sulla base delle leggi della meccanica (così come Newton aveva fatto per il moto della Terra attorno al sole). Ma i tentativi in questo senso furono un completo fallimento: tutte le predizioni risultarono sbagliate. (Per inciso, pressappoco nello stesso periodo fu sviluppata la teoria della relatività, che certamente conoscete come una grande rivoluzione della fisica. Se però la si paragona con la scoperta che le leggi del moto di Newton risultano del tutto erronee quando sono applicate agli atomi, la teoria della relatività appare solo una modifica secondaria). Mettere a punto un altro sistema di leggi che sostituisse quelle di Newton richiese parecchio tempo, perché i fenomeni atomici risultavano decisamente strani. Per poter cogliere ciò che accade a livello atomico bisognava rinunciare al comune buon senso. Infine nel 1926 venne sviluppata una teoria «priva di buon senso» che spiegava l'«inusitato comportamento» degli elettroni nella materia. Sembrava una teoria assurda, ma in realtà non lo era: fu chiamata meccanica quantistica. La parola «quanto» si riferisce a questo peculiare aspetto della natura che va contro il buon senso, ed è proprio di questo aspetto che ho intenzione di parlare.

La teoria quantistica ha permesso di spiegare un'infinità di particolari, ad esempio perché un atomo di ossigeno si combina con due di idrogeno per formare l'acqua, e così via. La meccanica quantistica ha dunque fornito la teoria che sottostà alla chimica, sicché si può dire che la chimica teorica fondamentale è in realtà un aspetto della fisica.

La meccanica quantistica ebbe un successo straordinario, perché poteva spiegare tutta la chimica e le diverse proprietà delle sostanze. Rimaneva però ancora il problema dell'interazione della luce con la materia. Occorreva cioè modificare la teoria dell'elettromagnetismo di Maxwell per metterla in accordo con i nuovi principi quantistici. Nel 1929 venne ideata da diversi fisici una nuova teoria che descrive l'interazione quantistica della luce con la materia, e che fu battezzata con l'orribile nome di «elettrodinamica quantistica».

Questa teoria presentava tuttavia dei problemi. Se la si usava per calcolare solo grossolanamente una grandezza fisica, si otteneva una risposta ragionevole. Ma quando si affrontavano calcoli più accurati, si scopriva che correzioni che si presupponevano piccole (termini successivi in una serie, ad esempio) risultavano invece molto grandi, anzi *infinite*! Insomma, al di là di una certa precisione non si riusciva a calcolare *nulla*.

Devo dire che quella che ho appena delineato è una «storia della fisica vista da un fisico», una storia, cioè, che non è mai del tutto veridica, una specie di storia-mito convenzionale che i fisici raccontano ai loro studenti, i quali la raccontano ai loro studenti, e che non è necessariamente collegata all'effettivo sviluppo storico (a me peraltro sconosciuto!).

Ad ogni modo, per continuare con questa 'storia', Paul Dirac, sulla base della teoria

della relatività, formulò una teoria relativistica dell'elettrone nella quale però non si teneva conto degli effetti dell'interazione dell'elettrone con la luce. La teoria di Dirac prevede che l'elettrone si comporti come un piccolo magnete, caratterizzato da un momento magnetico che vale esattamente 1 in certe unità di misura. Intorno al 1948 fu scoperto sperimentalmente che il valore effettivo è invece prossimo a 1,00118 (con un'incertezza di circa 3 sull'ultima cifra). Naturalmente si sapeva che l'elettrone interagisce con la luce, per cui ci si aspettava una piccola correzione. Si riteneva pure che essa fosse spiegabile in base alla nuova teoria dell'elettrodinamica quantistica. Ma quando la si calcolò, invece di 1,00118, si ottenne infinito, il che è sbagliato, sperimentalmente!

Il problema di come calcolare le grandezze fisiche in elettrodinamica quantistica fu infine risolto da Julian Schwinger, da Sin-Itiro Tomonaga e dal sottoscritto intorno al 1948. Schwinger fu il primo a calcolare la correzione al valore del momento magnetico dell'elettrone ricorrendo a un «gioco di bussolotti»; il valore teorico da lui ottenuto fu circa 1,00116, sufficientemente vicino al numero dato dagli esperimenti per dimostrare che eravamo sulla buona pista. Finalmente avevamo una teoria quantistica dell'elettromagnetismo nella quale era possibile portare a termine i conti! Questa è la teoria che vi voglio presentare.

L'elettrodinamica quantistica ha ormai più di cinquant'anni ed è stata verificata con accuratezza sempre maggiore in situazioni sempre più varie. Al momento attuale posso dire con orgoglio che *non vi è discrepanza significativa* tra la teoria e gli esperimenti.

Ecco, ad esempio, alcuni valori recenti che fanno capire quanto la teoria sia stata messa sotto il torchio: gli esperimenti danno per il numero di Dirac il valore di 1,00115965221 (con un'incertezza di circa 4 sull'ultima cifra); la teoria prevede 1,00115965246 (con un'incertezza circa cinque volte maggiore). Per darvi un'idea della precisione di questi numeri, pensate che se si misurasse con la stessa precisione la distanza tra New York e Los Angeles, si avrebbe un risultato approssimato a meno dello spessore in un capello umano. Questo per dirvi con quanta raffinatezza l'elettrodinamica quantistica è stata verificata negli ultimi cinquant'anni, sia sul piano teorico che su quello sperimentale. Ho citato un solo esempio, ma le grandezze misurate con analoga precisione sono molte e tutte risultano in perfetto accordo con le previsioni dell'elettrodinamica quantistica. Sono stati fatti controlli su una scala di distanze che variano da un centinaio di volte le dimensioni della Terra fino a un centesimo di quelle del nucleo atomico. Ho citato questi dati per far colpo e convincervi così che, verosimilmente, la teoria non è del tutto fuori strada. Più avanti vi spiegherò come si fanno questi calcoli.

Vorrei ancora una volta sottolineare l'enorme varietà dei fenomeni descritti dall'elettrodinamica quantistica. Anzi, è più facile dirlo alla rovescia: questa teoria descrive *tutti* i fenomeni del mondo fisico tranne la forza di gravità, cioè la forza che in questo momento vi tiene seduti al vostro posto (a meno che non sia una combinazione di gravità e di cortesia), e i fenomeni radioattivi, che riguardano i nuclei e le transizioni tra i loro livelli energetici. Mettendo da parte la gravità e la radioattività, che cosa rimane? La combustione della benzina nelle automobili, la schiuma e le bolle, la durezza del sale o del

rame, la rigidità dell'acciaio. Inoltre, i biologi stanno cercando di interpretare anche i fenomeni biologici in termini chimici, e, come ho già spiegato, la chimica poggia sull'elettrodinamica quantistica.

A questo punto è necessario chiarire una cosa: ho detto che questa teoria può spiegare tutti i fenomeni del mondo fisico, ma questo non è vero nella pratica.

La maggior parte dei fenomeni a noi noti interessa un numero così grande di elettroni che le nostre povere menti faticano a orientarsi in tale complessità. In simili situazioni possiamo usare la teoria per capire grosso modo ciò che dovrebbe accadere, e in effetti ciò che accade è appunto quanto previsto, grosso modo. Se però si organizza un esperimento di laboratorio che coinvolge solo *pochi* elettroni in circostanze *semplici*, si è in grado di calcolare con grande accuratezza quanto può accadere, e anche di eseguire misure molto precise. Ogni volta che sono stati eseguiti esperimenti di questo genere, l'elettrodinamica quantistica ha funzionato perfettamente.

Noi fisici stiamo sempre attenti a verificare se c'è qualche intoppo in una teoria, e il bello sta proprio qui, perché un eventuale intoppo è una cosa quanto mai interessante. Ma finora la teoria dell'elettrodinamica quantistica non ha rivelato alcuna pecca, ed è quindi, direi, il gioiello della fisica, la creatura di cui andare più fieri.

L'elettrodinamica quantistica è anche il prototipo per le teorie più recenti che cercano di spiegare i fenomeni nucleari, cioè quello che accade all'interno dei nuclei atomici. Se si immagina il mondo fisico come un palcoscenico, gli attori non sono solo gli elettroni, che negli atomi stanno all'esterno del nucleo, ma anche i quark, i gluoni e le dozzine di particelle che si trovano al suo interno. E questi 'attori', sebbene diversissimi l'uno dall'altro, hanno tutti lo stesso stile, uno stile strano e peculiare, lo stile «quantistico». Alla fine di queste lezioni racconterò qualcosa anche sulle particelle nucleari; ma nel frattempo, per rendere le cose più semplici, parlerò solo di fotoni, le particelle di luce, e di elettroni. Perché l'aspetto più importante, e anche il più interessante, è il loro *modo* di comportarsi.

Ho così chiarito l'argomento di queste lezioni. Ma ora sorge un'altra domanda: capirete quello che dirò? Chi va a sentire una conferenza scientifica è già pronto a non capire niente, però, chissà, forse potrà consolarsi ammirando la bella cravatta sgargiante del conferenziere. A voi questo non è possibile. [Feynman è senza cravatta].

Le cose di cui parlerò le insegniamo agli studenti di fisica degli ultimi anni di università: ora, voi pensate che io riuscirò a spiegarle in modo da farvele capire? Ebbene, no, non le capirete. Perché, allora, farvi perdere del tempo? Perché tenervi qui seduti, se non sarete in grado di capire ciò che dirò? Per convincervi a non andar via solo perché questa conferenza vi risulterà incomprensibile, vi dirò che anche i miei studenti di fisica non capiscono queste cose. E non le capiscono perché non le capisco nemmeno io. Il fatto è che non le capisce nessuno.

Permettetemi alcune considerazioni su che cosa significhi «capire». Una conferenza può essere incomprensibile per diverse ragioni. Il conferenziere può essere un cattivo oratore, che non sa esprimersi con efficacia o dice le cose nell'ordine sbagliato, oppure può avere una pronuncia poco chiara. Quest'ultimo è un problema presto risolto: farò di tutto per tenere a freno il mio accento newyorkese!

Altre volte l'oratore, soprattutto se è un fisico, usa parole ordinarie in modo curioso. I fisici usano spesso in senso tecnico parole comuni quali «lavoro», «azione», «energia», e anche «luce», come si vedrà. Ma quando si parla di «lavoro» in fisica non si intende la stessa cosa di quando si parla di «lavoro» nella vita di ogni giorno. Nel corso di queste lezioni mi capiterà di adoperare, senza accorgermene, alcune di queste parole in un senso diverso dal consueto. Cercherò di controllarmi (è parte del mio compito), ma è un errore in cui è molto facile incorrere.

Un altro motivo per cui potreste pensare di non seguire quello che racconterò è che mentre io descriverò *come* funziona la Natura, voi non capirete *perché* la Natura funzioni così. Ma questo non lo capisce nessuno, e quindi io non ve lo so spiegare.

C'è infine un'altra possibilità: che alcune delle cose che vi dirò vi sembrino incredibili, inaccettabili, impossibili da mandar giù. In questi casi è come se calasse un sipario: uno smette di ascoltare. Io descriverò il comportamento della Natura, ma se a voi questo comportamento non piace, il vostro processo di comprensione ne risulterà intralciato. I fisici hanno imparato a convivere con questo problema: hanno cioè capito che il punto essenziale non è se una teoria piaccia o non piaccia, ma se fornisca previsioni in accordo con gli esperimenti. La ricchezza filosofica, la facilità, la ragionevolezza di una teoria sono tutte cose che non interessano. Dal punto di vista del buon senso l'elettrodinamica quantistica descrive una natura assurda. Tuttavia è in perfetto accordo con i dati sperimentali. Mi auguro quindi che riuscirete ad accettare la Natura per quello che è: assurda.

Per me parlare di questa assurdità è un divertimento, perché la trovo incantevole. Vi chiedo come favore di non mettervi a pensare ad altro solo perché non riuscite a credere che la Natura sia così strana. Ascoltatemi fino in fondo, e vedrete che alla fine ne sarete incantati anche voi.

Come farò a spiegarvi cose che non spiego ai miei studenti se non alla fine dei loro studi? Proverò a chiarirlo con un'analogia. I Maya, che nutrivano un grande interesse per Venere come 'stella' del mattino e 'stella' della sera, studiarono a fondo il problema di come prevedere accuratamente il sorgere e il tramontare di questo astro. Dopo alcuni anni di osservazioni, notarono che cinque cicli di Venere erano quasi uguali a otto dei loro «anni nominali» di 365 giorni (essi sapevano che il vero anno delle stagioni è differente e avevano fatto il conto anche rispetto ad esso). Per i loro calcoli, i Maya avevano inventato un sistema fatto di punti e di sbarrette che rappresentavano i numeri (zero compreso) e delle regole di calcolo che permettevano loro di prevedere non solo il sorgere e il tramontare di Venere, ma anche altri fenomeni celesti quali le eclissi lunari.

A quei tempi gli unici in grado di fare calcoli così elaborati erano pochi membri della casta sacerdotale. Immaginiamo ora di chiedere a uno di costoro come si esegue un certo passaggio nel calcolo della predizione del prossimo sorgere di Venere come stella del mattino: la sottrazione di due numeri. E supponiamo di non essere mai andati a scuola e quindi di non conoscere la sottrazione. Come farebbe il sacerdote maya per spiegarci cosa è una sottrazione?

Potrebbe scegliere fra due metodi: insegnarci il suo modo di rappresentare i numeri con punti e sbarrette e le regole per fare una «sottrazione», oppure descrivere ciò che sta

facendo in concreto: «Mettiamo che si debba sottrarre 236 da 584. Prima si contano 584 fagioli e li si mette in una ciotola. Poi se ne tolgono 236 e li si mette da parte. Infine si contano i fagioli rimasti nella ciotola. Questo numero è il risultato della sottrazione di 236 da 584».

Al che voi potreste esclamare: «Per Quetzalcoatl, che barba! Contare fagioli, metterli dentro, tirarli fuori... che lavoraccio!».

Allora il sacerdote maya risponderebbe: «Ecco perché si adoperano le regole con le sbarrette e con i punti. Sono complesse, ma forniscono la risposta in modo molto più efficiente del metodo dei fagioli. Ciò che importa, però, è che non c'è nessuna differenza *nella risposta*: le apparizioni di Venere sono calcolabili sia contando i fagioli (sistema lento, ma facile da capire), sia usando le regole complicate (sistema molto più veloce, ma che richiede anni di studio)».

Capire come funzionano le sottrazioni (sempre che non le si debba realmente eseguire) non è poi tanto difficile. E io farò lo stesso: spiegherò quello che i fisici fanno per prevedere il comportamento della Natura, ma non insegnerò alcun trucco per farlo in modo efficiente. Come vedrete tra poco, per fare qualunque predizione ragionevole nell'ambito dell'elettrodinamica quantistica si dovrebbero tracciare un'infinità di piccole frecce su un foglio di carta. Ci vogliono sette anni di università per imparare a fare questi conti con disinvoltura. Questi sette anni voi li salterete bellamente, perché io vi spiegherò l'elettrodinamica quantistica facendovi vedere che cosa *si fa in concreto*. Spero solo che riuscirete a capirlo meglio di quanto facciano alcuni studenti!

Torniamo all'esempio dei Maya e facciamo un passo avanti: potremmo chiedere al sacerdote perché cinque cicli di Venere sono quasi uguali a 2920 giorni, cioè a otto anni. I perché della sua risposta potrebbero essere molti, ad esempio: «20 è un numero importante nel nostro sistema di numerazione, e dividendo 2920 per 20 si trova 146, che è uno più di un numero che può essere rappresentato in due modi diversi come somma di due quadrati», e così via. Ma una simile teoria non avrebbe in realtà nulla a che fare con Venere. Al giorno d'oggi ci si è resi conto che teorie del genere non servono a nulla. Perciò, ripeto, non ci occuperemo del *perché* del comportamento peculiare della Natura: non esistono buone teorie che diano conto di ciò.

Tutto quello che ho detto fin qui voleva mettervi nelle condizioni di spirito giuste per seguirmi. Altrimenti, tanto varrebbe andarcene tutti a casa. E adesso possiamo incominciare!

Partiamo dalla luce. La prima cosa di cui Newton si accorse, quando si mise a osservare i fenomeni luminosi, fu che la luce bianca è una sovrapposizione di vari colori. Egli, infatti, riuscì a separare la luce bianca in vari colori mediante un prisma; tuttavia, quando fece passare luce di un solo colore, ad esempio il rosso, attraverso un secondo prisma, osservò che essa non veniva ulteriormente separata. Ne concluse che la luce bianca è una sovrapposizione di diversi colori, ciascuno dei quali è puro nel senso che non può essere ulteriormente suddiviso.

(In realtà ogni colore può essere suddiviso ancora una volta, ma in modo differente: secondo la sua cosiddetta «polarizzazione». Questa proprietà della luce non è essenziale per capire le idee di base dell'elettrodinamica quantistica e quindi, per semplificare, non

ne parlerò, anche se questo significherebbe non dare una descrizione assolutamente completa della teoria. Questa piccola semplificazione non toglierà nulla alla comprensione di ciò che dirò, e se la menziono è per una questione di scrupolo).

Durante queste lezioni, quando parlerò di «luce», non mi riferirò unicamente alla luce visibile, dal rosso al violetto. La luce che noi vediamo costituisce solo una parte di una lunga scala, analoga a quella musicale, in cui ci sono note più alte e note più basse di quelle udibili dall'orecchio umano. Il tipo di luce può essere caratterizzato da un numero, detto frequenza; all'aumentare di questo numero si passa dal rosso al blu, al violetto, all'ultravioletto. La luce ultravioletta non può essere vista, ma può impressionare una pellicola fotografica. È sempre luce, solo che il numero che la caratterizza è differente. (Guardiamoci da un'ottica troppo angusta: quello che ci rivelano i nostri strumenti personali, gli occhi, non esaurisce tutto ciò che esiste!). Se si continua ad aumentare la frequenza si trovano i raggi X, i raggi gamma e così via. Variando il numero nella direzione opposta, si va dal blu al rosso, alle onde infrarosse (calorifiche), alle onde della televisione, a quelle radio. Tutte queste cose io le chiamo «luce». In quasi tutti gli esempi che seguiranno, parlerò di luce rossa, ma l'elettrodinamica quantistica descrive i fenomeni relativi a tutto l'intervallo suddetto.

Newton pensava che la luce fosse fatta di particelle, da lui chiamate «corpuscoli», e aveva ragione (anche se a quella conclusione era giunto attraverso un ragionamento errato). Oggi si sa che la luce è costituita da particelle perché se si prende uno strumento molto sensibile che produce un ticchettio quando viene colpito dalla luce, i singoli «clic», quando la luce diventa più fioca, restano della stessa intensità e diventano solo meno frequenti. La luce è dunque un po' come la pioggia: è fatta di tante 'gocce' chiamate fotoni, e quando è di un solo colore, le sue 'gocce' hanno tutte la stessa dimensione.

L'occhio umano è uno strumento molto sensibile: bastano cinque o sei fotoni per attivare una cellula nervosa e inviare un messaggio al cervello. Se l'evoluzione ci avesse portato a una sensibilità visiva dieci volte maggiore dell'attuale, non ci sarebbe stato bisogno di questa spiegazione: percepiremmo la luce molto fioca di un dato colore come una serie di piccoli lampi intermittenti, tutti della stessa intensità.

Per rivelare un singolo fotone, ci si serve di uno strumento chiamato fotomoltiplicatore (fig. 1). Colpendo la placca metallica A, un fotone provoca l'emissione di un elettrone da uno dei suoi atomi.

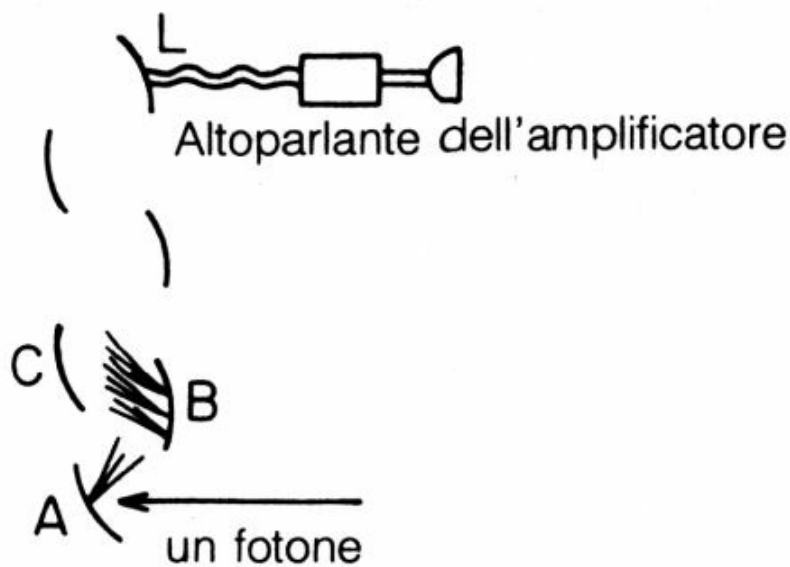


Fig. 1. Il fotomoltiplicatore è in grado di rivelare un singolo fotone. Quando un fotone colpisce la placca A, libera un elettrone che, attratto dalla placca B, carica positivamente, la urta liberando nuovi elettroni. Questo processo si ripete molte volte e miliardi di elettroni arrivano sull'ultima placca, L, dando luogo a una corrente elettrica che viene amplificata da un normale amplificatore. Collegando a quest'ultimo un altoparlante, ogni volta che un fotone di un dato colore colpisce la placca A, si sente un «clic» di intensità uniforme.

L'elettrone liberato è fortemente attratto dalla placca B, carica positivamente, e la urta con forza sufficiente per far emettere tre o quattro nuovi elettroni. Ciascuno di questi è attratto dalla placca C, anch'essa carica positivamente, e il loro impatto su C porta all'emissione di altri elettroni. Questo processo viene ripetuto dieci o dodici volte, così che l'ultima placca, L, viene raggiunta da miliardi di elettroni, sufficienti per costituire una corrente elettrica misurabile. Questa corrente può essere amplificata con un normale amplificatore e inviata a un altoparlante che produce un ticchettio udibile. Ogni volta che un fotone di un dato colore incide sul fotomoltiplicatore si sente un «clic» di intensità uniforme.

Si consideri ora un gran numero di fotomoltiplicatori disposti tutt'intorno a una sorgente luminosa molto fioca che emette luce in tutte le direzioni. Quando arriva su uno qualunque dei rivelatori, la luce produce un «clic» di intensità fissa. È sempre «o tutto o niente»: se un fotomoltiplicatore 'scatta' in un certo istante, non ne scatta nessun altro (tranne nei rari casi in cui due fotoni sono partiti dalla sorgente allo stesso istante). La luce non si suddivide in «mezze particelle» che vanno in direzioni differenti.

Ripeto: la luce si presenta sotto forma di particelle. È molto importante chiarire questo punto, specialmente per chi è andato a scuola, dove probabilmente gli è stato insegnato che la luce si comporta come un'onda. Io vi dico invece che la luce si comporta come particelle.

Si potrebbe obiettare che è il fotomoltiplicatore a rivelare la luce sotto forma di particelle. E invece no: ogni misura fatta con qualsiasi strumento abbastanza sensibile per rivelare luce molto fioca ha sempre portato alla medesima conclusione: la luce è costituita da particelle.

Do per scontato che vi siano note le più comuni proprietà della luce: che un raggio di luce si propaga in linea retta; che si piega quando entra nell'acqua; che quando viene riflesso da una superficie l'angolo di riflessione è uguale a quello di incidenza; che la luce può essere suddivisa nei vari colori; che può essere focalizzata mediante una lente; che in una pozzanghera fangosa cosparsa di petrolio si osservano bellissimi colori, e così via.

Tutti questi fenomeni assai comuni mi serviranno per mostrare il comportamento veramente strano della luce, e li spiegherò in base all'elettrodinamica quantistica. Ho parlato del fotomoltiplicatore per illustrare un aspetto essenziale che forse non vi era altrettanto noto, cioè che la luce è fatta di particelle, ma spero che ora anche su questo non ci siano più dubbi.

Tutti sapete, immagino, che la luce viene parzialmente riflessa da alcune superfici, ad esempio l'acqua. Quanti quadri romantici raffigurano un chiaro di luna che si specchia nelle acque di un lago! (E quante volte ci siamo cacciati nei guai proprio a causa di questo poetico riflesso!). Se si guarda nell'acqua, soprattutto di giorno, si vede quello che c'è sotto la superficie, ma si vede anche quello che la superficie riflette. Altro esempio ben noto è il vetro: se in una stanza c'è una lampada accesa, guardando di giorno attraverso il vetro di una finestra, si vede non solo la scena esterna, ma anche un debole riflesso della lampada. Quindi la luce è parzialmente riflessa dalla superficie del vetro.

Prima di proseguire, voglio avvertirvi di una semplificazione che correggerò più avanti: per il momento farò finta che la riflessione parziale della luce da parte del vetro avvenga solo sulla sua *superficie*. In realtà un pezzo di vetro è un oggetto di complessità mostruosa, in cui un numero enorme di elettroni si agita in modo incredibile. Un fotone che arriva su un pezzo di vetro interagisce con gli elettroni di *tutto* il vetro e non solo sulla superficie. Il fotone e gli elettroni fanno una specie di danza il cui effetto complessivo è lo stesso che si avrebbe se il fotone interagisse unicamente con la superficie. Di qui la semplificazione che ho detto. Quando poi mostrerò quello che accade dentro il vetro, si capirà perché il risultato è lo stesso.

Passo ora a descrivere un esperimento e ad illustrare il suo sorprendente risultato. Inviamo alcuni fotoni di un dato colore, ad esempio rosso, da una sorgente di luce verso un blocco di vetro (fig. 2). Un fotomoltiplicatore posto in A raccoglie i fotoni riflessi dalla superficie frontale del vetro. Un altro fotomoltiplicatore situato in B, all'interno del vetro, misura quanti fotoni attraversano la superficie. Tralasciando l'ovvia difficoltà di inserire un fotomoltiplicatore dentro un blocco di vetro, qual è il risultato di un simile esperimento?

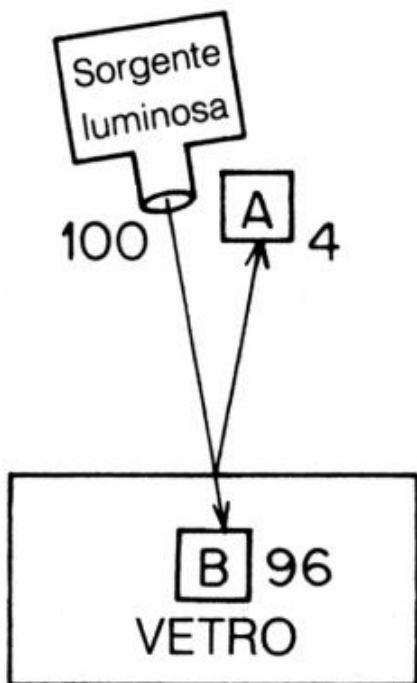


Fig. 2. Esperimento per misurare la riflessione parziale della luce da parte di una singola superficie di vetro. Su 100 fotoni che partono dalla sorgente luminosa, 4 sono riflessi dalla superficie frontale e finiscono nel fotomoltiplicatore in A, mentre gli altri 96 l'attraversano e finiscono nel fotomoltiplicatore in B.

In media, su cento fotoni emessi dalla sorgente che incidono sul vetro in direzione perpendicolare, ne arrivano 4 in A e 96 in B. In questo caso «riflessione parziale» significa dunque che il 4% dei fotoni viene riflesso dalla superficie del vetro, mentre il restante 96% la attraversa. Ed eccoci in grave difficoltà: com'è possibile che la luce venga riflessa *parzialmente*? Ogni fotone finisce o in A o in B, ma come fa a «decidere» se andare in A o in B? [Risate]. Può sembrare una battuta, ma ridere non basta: dobbiamo trovare una teoria che spieghi questo comportamento! Quindi la riflessione parziale è già un profondo mistero, e fu infatti un problema molto difficile per Newton.

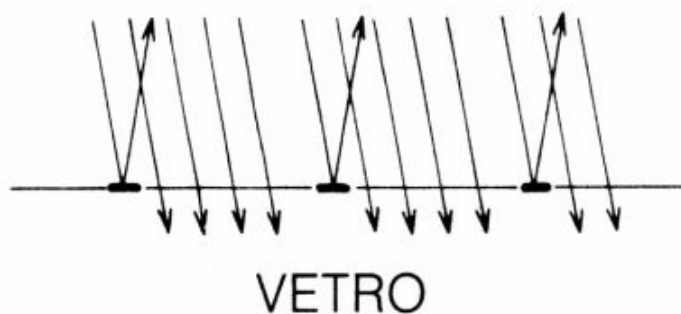


Fig. 3. Per spiegare la riflessione parziale da parte di una singola superficie, questa teoria ipotizza che la superficie sia costituita prevalentemente da «buchi» che lasciano passare la luce e da alcune «chiazze» che invece la riflettono.

Si possono avanzare diverse congetture. Una è che il 96% della superficie del vetro sia fatto di «buchi» che lasciano passare la luce, mentre il rimanente 4% sia fatto di «chiazze» di materiale riflettente (fig. 3). Ma Newton si rese conto che questa non è una spiegazione possibile.³ Discuterò tra breve uno strano aspetto della riflessione parziale che farebbe uscire di senno chiunque insistesse con la teoria dei buchi e delle chiazze (o con qualunque altra ipotesi ragionevole!).

Un'altra teoria potrebbe essere che il fotone ha un meccanismo interno, un qualche congegno, che gli permette di passare attraverso il vetro solo quando è «mirato» a

puntino. Per mettere alla prova questa teoria, possiamo cercare di filtrare i fotoni che non lo sono, interponendo alcuni strati di vetro tra la sorgente e la superficie del nostro blocco. I fotoni che arrivano sul blocco, avendo attraversato i filtri, dovrebbero essere tutti mirati correttamente e quindi non dovrebbero venir riflessi. Purtroppo questa teoria non è in accordo con gli esperimenti: anche dopo aver attraversato diversi filtri il 4% dei fotoni che incidono su una superficie di vetro viene riflesso.

Abbiamo un bel lambiccarci il cervello per inventare una teoria ragionevole che spieghi come fa un fotone a «decidere» se attraversare il vetro o rimbalzare su di esso: è impossibile prevedere che cosa farà il singolo fotone. Alcuni filosofi hanno sostenuto che se le stesse circostanze non producono sempre lo stesso effetto, le previsioni risultano impossibili e la scienza stessa viene a crollare. Qui ci troviamo di fronte proprio a un caso del genere: la situazione è sempre la stessa (fotoni identici che incidono tutti nella stessa direzione sullo stesso pezzo di vetro), ma i risultati sono variabili. Non siamo in grado di prevedere se un dato fotone arriverà in A o in B; possiamo solo prevedere che, in media, su 100 fotoni che incidono sul vetro, 4 saranno riflessi dalla sua superficie frontale. Bisogna concluderne che la fisica, scienza profondamente esatta, è ridotta a calcolare la sola *probabilità* di un evento, invece di prevedere che cosa accade in ciascun caso singolo? Ebbene, sì. È un ripiegamento, ma le cose stanno proprio così: la Natura ci permette di calcolare soltanto delle probabilità. Con tutto ciò la scienza è ancora in piedi. Se la riflessione parziale su una singola superficie è un bel mistero e un problema difficile, la riflessione parziale su due o più superfici è un vero rompicapo. Per spiegare come mai, descriverò un secondo esperimento in cui si misura la riflessione parziale della luce da parte di due superfici. Sostituiamo il blocco di vetro con una lamina di vetro molto sottile le cui due superfici siano esattamente parallele tra di loro, e mettiamo il fotomoltiplicatore sotto la lamina e allineato con la sorgente di luce. Questa volta i fotoni che arrivano in A possono essere stati riflessi da una qualunque delle due superfici; tutti gli altri fotoni finiscono in B (fig. 4). Ci si potrebbe aspettare che la superficie frontale rifletta il 4% della luce e quella posteriore rifletta il 4% del restante 96% per un totale di circa 8%. Si dovrebbe quindi trovare che su 100 fotoni partiti dalla sorgente, circa 8 arrivano in A.

In queste condizioni sperimentali perfettamente controllabili, si osserva che solo in particolari situazioni in A arrivano 8 fotoni su 100. Con alcune lamine di vetro arrivano sempre 15 o 16 fotoni, il doppio del valore atteso, mentre con altre ne arrivano uno o due! Alcune lamine producono una riflessione del 10%, con altre la riflessione parziale scompare del tutto! Come spiegare questi risultati assurdi? Dopo esserci assicurati che le varie lamine sono tutte uniformi e della stessa qualità, scopriamo che esse differiscono solo lievemente per lo spessore.

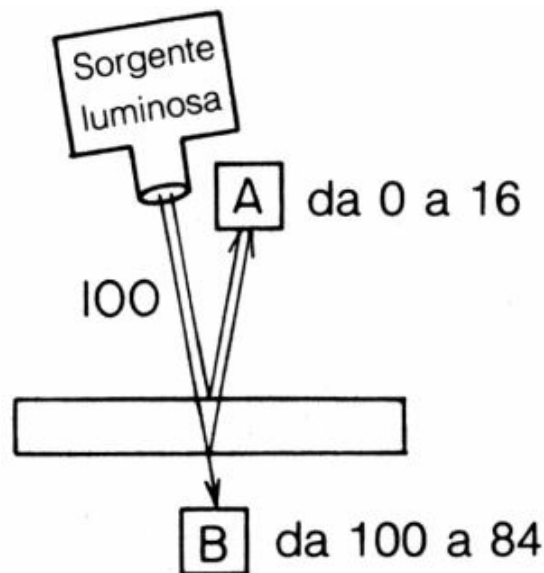


Fig. 4. Esperimento per misurare la riflessione parziale della luce da parte di due superfici di vetro. I fotoni possono arrivare nel fotomoltiplicatore in A riflettendosi sia sulla superficie frontale sia su quella posteriore della lamina; oppure possono attraversare entrambe le superfici unendo nel fotomoltiplicatore in B. Secondo lo spessore della lamina, nel rivelatore in A arrivano da zero a 16 fotoni ogni 100. Questo risultato mette in difficoltà qualunque teoria ragionevole, compresa quella della fig. 3. La riflessione parziale, a quanto sembra, può essere «smorzata» o «amplificata» dalla presenza di altre superfici.

Per mettere alla prova l'ipotesi che la quantità di luce riflessa da due superfici dipende dallo spessore del vetro, si può fare una serie di esperimenti: iniziando con la lamina di vetro più sottile possibile si contano quanti fotoni arrivano nel fotomoltiplicatore A per ogni 100 fotoni emessi dalla sorgente. Poi si sostituisce la lamina con una un po' più spessa e si ripete il conteggio. Che risultati si hanno dopo aver ripetuto una dozzina di volte questo procedimento?

Con la lamina di vetro più sottile possibile si trova che il numero di fotoni che arrivano in A è quasi sempre zero, solo qualche rara volta ne arriva uno. Sostituendo la lamina con una leggermente più spessa, si trova che la quantità di luce riflessa è maggiore, più vicina all'atteso 8%. Dopo alcune altre sostituzioni il numero dei fotoni che arrivano in A supera l'8%. Continuando ad aumentare lo «spessore» - ormai siamo a circa due milionesimi di centimetro, — la quantità di luce riflessa dalle due superfici arriva a un valore massimo del 16%, poi incomincia a diminuire, ripassa per l'8% e scende nuovamente a zero. Se la lamina di vetro ha lo spessore giusto non si ha nessuna riflessione. (E la teoria delle chiazze è servita!).

Aumentando ancora lo spessore, la riflessione parziale torna ad aumentare fino al 16% per poi scendere di nuovo a zero e così di seguito (fig. 5).

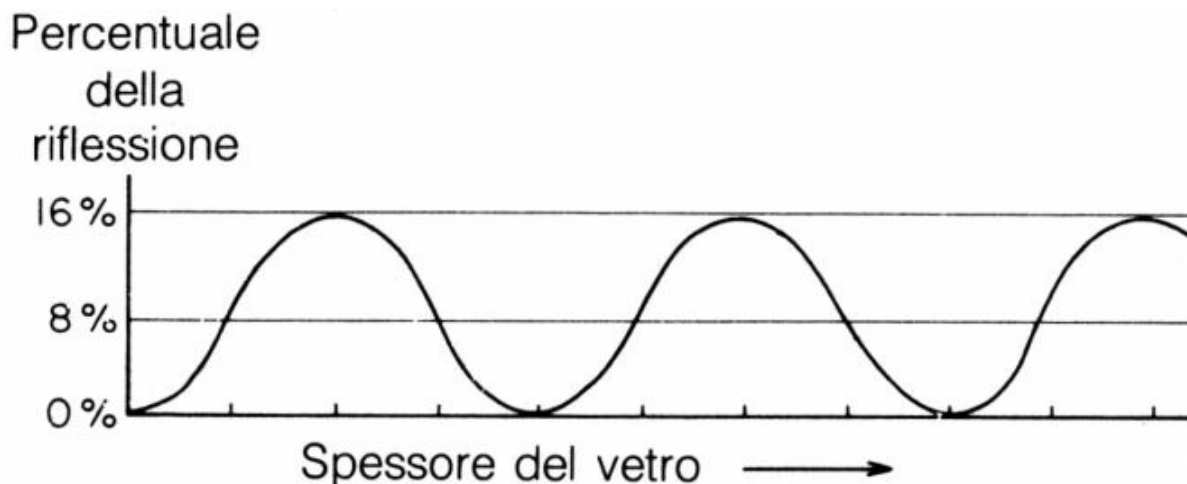


Fig. 5. I risultati di un esperimento in cui si misura attentamente la relazione tra lo spessore della lamina di vetro e la riflessione parziale mostrano il fenomeno noto come «interferenza». Con l'aumentare dello spessore del vetro, la riflessione parziale varia ciclicamente tra zero e il 16%.

Newton scoprì queste oscillazioni e fece un esperimento che per essere correttamente interpretato richiede che il ciclo si ripeta almeno 34.000 volte! Oggi, grazie al laser, che produce una luce molto prossima alla monocromaticità pura, è possibile osservare ancora queste oscillazioni dopo 100.000.000 di cicli, corrispondenti a un vetro spesso oltre 50 metri. (Questo fenomeno non viene osservato in condizioni usuali perché le normali sorgenti di luce non sono monocromatiche).

La previsione dell'8% di luce riflessa risulta quindi giusta come media globale, poiché la quantità riflessa varia con andamento regolare tra zero e il 16%, ma è esatta solo due volte per ogni ciclo, così come un orologio fermo segna l'ora esatta solo due volte al giorno. Come spiegare questa strana dipendenza della riflessione parziale dallo spessore del vetro? Come può il 4% della luce riflettersi sulla superficie frontale (come si è visto nel primo esperimento descritto), se con una seconda superficie posta alla distanza giusta si riesce a eliminare del tutto la riflessione? Non solo: ponendo la seconda superficie a una distanza solo di poco diversa, si può «amplificare» la riflessione fino al 16%! Che la superficie posteriore influisca in qualche modo sulla capacità di quella frontale di riflettere la luce? Che succede se si inserisce una *terza* superficie?

Con una terza superficie, o una quarta, o una quinta, ecc., la quantità di luce riflessa cambia di nuovo, e con questa teoria noi ci troviamo costretti a inseguire sempre nuove superfici, alla ricerca della misteriosa superficie ultima. E questo che deve fare un fotone per «decidere» se riflettersi sulla superficie frontale?

Newton elaborò varie ipotesi ingegnose per spiegare questo problema,⁴ ma alla fine si rese conto di non essere riuscito a sviluppare una teoria soddisfacente. Per un lungo periodo, dopo i vani tentativi di Newton, la riflessione parziale della luce da parte di due superfici venne felicemente spiegata dalla teoria ondulatoria.⁵ Questa tuttavia crollò quando vennero fatti esperimenti in cui luce molto fioca veniva inviata su un fotomoltiplicatore: si osservò infatti che, anche rendendo sempre più fioca la luce, il fotomoltiplicatore continuava a ticchettare con la stessa intensità, solo che i singoli scatti erano meno frequenti. La luce si comportava come un flusso di particelle.

Anche al giorno d'oggi non c'è modello intuitivo che spieghi la riflessione parziale della luce da due superfici; tutto quello che possiamo fare è calcolare la probabilità che un dato fotomoltiplicatore venga colpito da un fotone riflesso da un lamina di vetro. Questo calcolo sarà il mio primo esempio del metodo fornito dalla teoria dell'elettrodinamica quantistica. Vi mostrerò come «si contano i fagioli», cioè come fanno i fisici a ottenere la risposta esatta. Quello che non spiegherò è come fa il fotone a «decidere» se rimbalzare o se passare: non lo spiegherò perché è una domanda a cui non sappiamo dare risposta. (E che probabilmente non ha senso). Mostrerò solamente come calcolare il valore corretto della *probabilità* che la luce sia riflessa da un vetro di dato spessore, perché questa è la sola cosa che i fisici sanno fare! Il procedimento usato per risolvere questo problema specifico è analogo ai procedimenti necessari per risolvere *qualsiasi* altro problema spiegato dall'elettrodinamica quantistica.

E adesso tenetevi saldi, non perché quello che si deve fare sia difficile, ma perché è assolutamente ridicolo: si devono tracciare tante piccole frecce su un foglio di carta, e basta!

Cosa c'entra una freccia con la probabilità che un particolare evento si verifichi? Secondo il metodo del «contare i fagioli», la probabilità di un evento è uguale al quadrato della lunghezza della freccia. Per esempio, nel nostro primo esperimento, quello in cui veniva misurata la riflessione parziale prodotta dalla sola superficie frontale, la probabilità che un fotone arrivasse nel fotomoltiplicatore A era del 4%. Questo corrisponde a una freccia di lunghezza 0,2 perché il quadrato di 0,2 è 0,04 (fig. 6).

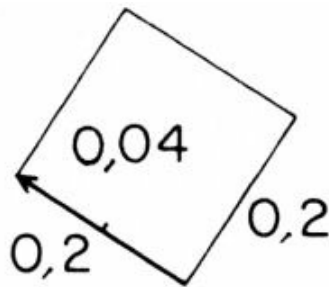


Fig. 6. Le strane proprietà della riflessione parziale da parte di due superfici hanno costretto i fisici ad abbandonare ogni previsione assoluta e a limitarsi a calcolare la probabilità di un evento. Per questi calcoli l'elettrodinamica quantistica fornisce un metodo che consiste nel tracciare piccole frecce su un foglio di carta. La probabilità di un evento è rappresentata dall'area del quadrato costruito su una freccia. Ad esempio, una probabilità dello 0,04 (il 4%) è rappresentata da una freccia lunga 0,2.

Nel secondo esperimento, in cui si usavano lamine di vetro di diversi spessori, i fotoni potevano arrivare in A rimbalzando sia sulla superficie frontale sia su quella posteriore. Come rappresentarlo con una freccia? Per descrivere probabilità che variano tra zero e il 16%, a seconda dello spessore del vetro, la lunghezza della freccia deve variare tra zero e 0,4 (fig. 7).

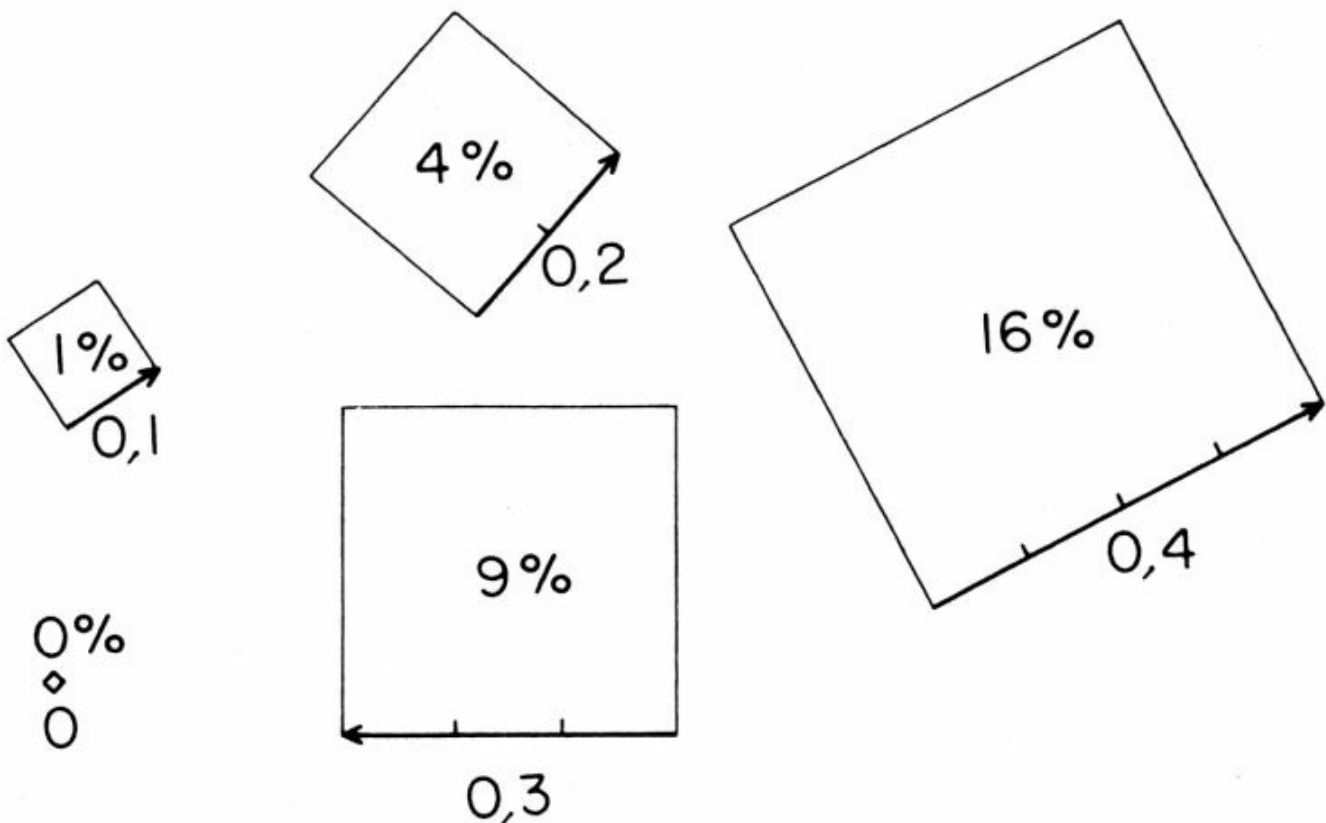


Fig. 7. Frecce che rappresentano probabilità tra zero e il 16% hanno una lunghezza che va da zero a 0,4.

Si considerino ora i vari modi in cui un fotone può arrivare dalla sorgente al fotomoltiplicatore A. Dato che, per semplificare, io ho stabilito che la luce rimbalza solo sulla superficie frontale o su quella posteriore, ci sono due possibili percorsi che un fotone può seguire per arrivare in A. In questo caso si devono disegnare *due* frecce, una per ciascun modo in cui l'evento può verificarsi, e poi combinarle in una «freccia finale», il cui quadrato rappresenta la probabilità dell'evento. Se l'evento fosse avvenuto in tre modi differenti, le frecce da combinarsi sarebbero state tre.

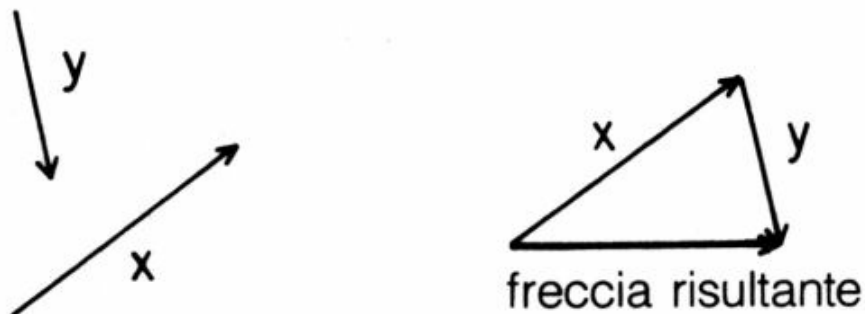


Fig. 8. Le frecce che rappresentano i possibili modi in cui un evento può verificarsi vanno disegnate e poi combinate tra loro («sommate») nel seguente modo: si unisce la coda di una freccia alla punta della precedente (senza cambiare la direzione di nessuna delle due) e infine si traccia la «freccia risultante» dalla coda della prima alla punta dell'ultima.

Vediamo ora come si combinano le frecce. Supponiamo di voler combinare la freccia x con la freccia y (fig. 8): non dobbiamo fare altro che porre la coda di y sulla punta di x (senza cambiare direzione a nessuna delle due) e tracciare la freccia finale dalla coda di x alla punta di y. Tutto qua. Possiamo combinare in questo modo un qualunque numero di frecce (in linguaggio tecnico si parla di «somma delle frecce» e di «risultante»). Ciascuna freccia indica quanto e in che direzione muoversi in questa specie di danza. La freccia risultante descrive lo spostamento *singolo* che si deve fare per andare a finire nello stesso posto (fig. 9).



Fig. 9. Un qualsiasi numero di frecce può esser sommato nel modo descritto alla fig. 8.

Quali regole determinano la lunghezza e la direzione delle singole frecce da combinare per ottenere la freccia risultante? Nel caso specifico delle due superfici di vetro ci sono due frecce da combinare, una che rappresenta la riflessione sulla *superficie frontale* e l'altra che rappresenta quella sulla *superficie posteriore*.

Cominciamo dalla lunghezza. Come si è visto nel primo esperimento, in cui un fotomoltiplicatore era interno al vetro, la superficie frontale riflette circa il 4% dei fotoni

che la colpiscono. Ciò significa che la freccia relativa alla «riflessione frontale» ha lunghezza 0,2. La superficie posteriore del vetro riflette anch'essa il 4%, dunque anche la freccia relativa alla «riflessione posteriore» ha lunghezza 0,2.

Per determinare la direzione di ciascuna freccia, immaginiamo di avere un cronometro con una sola lancetta che gira a grandissima velocità. Quando un fotone lascia la sorgente, facciamo partire il cronometro. Finché il fotone si muove, la lancetta continua a girare, compiendo circa 13.000 rivoluzioni al centimetro nel caso della luce rossa; quando il fotone arriva nel fotomoltiplicatore, fermiamo il cronometro. La direzione in cui si arresta la lancetta è quella lungo la quale va tracciata la freccia.

Per poter ottenere la risposta esatta occorre un'ultima regola: se si considera il percorso di un fotone che rimbalza sulla superficie *frontale* del vetro, si deve invertire il verso della freccia. In altre parole, mentre la freccia relativa alla riflessione *posteriore* va tracciata nello *stesso* verso della lancetta del cronometro, quella relativa alla riflessione *frontale* va tracciata nel verso *opposto*.

Possiamo così disegnare le frecce per il caso della luce che si riflette su una lamina di vetro estremamente sottile. Per ottenere la freccia per la riflessione frontale, immaginiamo un fotone che parte dalla sorgente di luce (la lancetta del cronometro comincia a girare), rimbalza sulla superficie frontale e arriva in A (la lancetta si arresta). A questo punto noi tracciamo una piccola freccia di lunghezza 0,2 nel verso opposto a quello della lancetta (fig. 10).

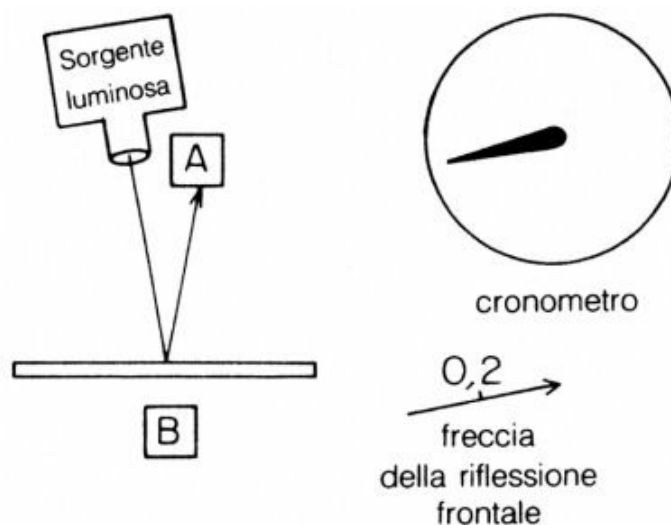


Fig. 10. In un esperimento in cui si misura la riflessione parziale da due superfici si può dire che ogni fotone ha due modi per arrivare in A: riflettendosi o sulla superficie frontale o su quella posteriore. A ciascuno di questi modi corrisponde una freccia di lunghezza 0,2, la cui direzione è determinata dalla lancetta di un «cronometro» che ruota mentre il fotone è in viaggio. La freccia relativa alla riflessione frontale va disegnata nel verso *opposto* a quello in cui si ferma la lancetta del cronometro.

Per ottenere la freccia per la riflessione posteriore, si immagini un fotone che parte dalla sorgente (la lancetta comincia a girare), attraversa la superficie frontale del vetro, rimbalza su quella posteriore e arriva in A (la lancetta si arresta). In questo caso la lancetta sarà orientata quasi nella stessa direzione del caso precedente, perché un fotone che rimbalza sulla superficie posteriore impiega un tempo solo di poco superiore per raggiungere A, dovendo attraversare due volte una lamina di vetro estremamente sottile. Si tracci adesso una piccola freccia di lunghezza 0,2 nello stesso verso della lancetta del cronometro (fig. 11).

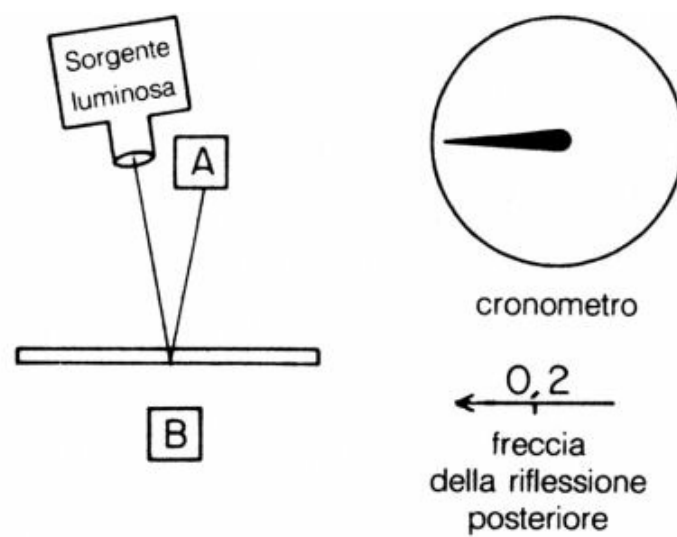


Fig. 11. Un fotone che rimbalza sulla superficie posteriore di una sottile lamina di vetro impiega un tempo leggermente maggiore per arrivare in A. Quindi la lancetta del cronometro che lo segue si trova alla fine in una posizione leggermente diversa da quella in cui si troverebbe seguendo un fotone riflesso dalla superficie frontale. La freccia relativa alla «riflessione» posteriore va tracciata nello *stesso* verso della lancetta del cronometro.

Ora bisogna combinare le due frecce. Poiché hanno entrambe la stessa lunghezza ma direzioni quasi opposte, la freccia risultante ha lunghezza quasi nulla e il suo quadrato sarà ancora più vicino a zero. Dunque la probabilità che la luce si rifletta su una lamina estremamente sottile è praticamente nulla (fig. 12).

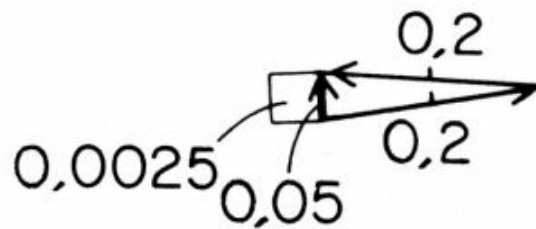


Fig. 12. La freccia risultante, il cui quadrato rappresenta la probabilità di riflessione su una lamina di vetro estremamente sottile, si ottiene sommando la freccia della riflessione frontale a quella della riflessione posteriore. Il risultato è quasi zero.

Se si sostituisce la lamina sottilissima con una leggermente più spessa, il fotone che rimbalza sulla superficie posteriore impiegherà più tempo per arrivare in A; la lancetta del cronometro girerà quindi un po' più a lungo, e le frecce relative alle riflessioni posteriore e frontale formeranno tra loro un angolo un po' più ampio di prima. La freccia finale sarà quindi un po' più lunga, e il suo quadrato proporzionalmente maggiore (fig. 13).

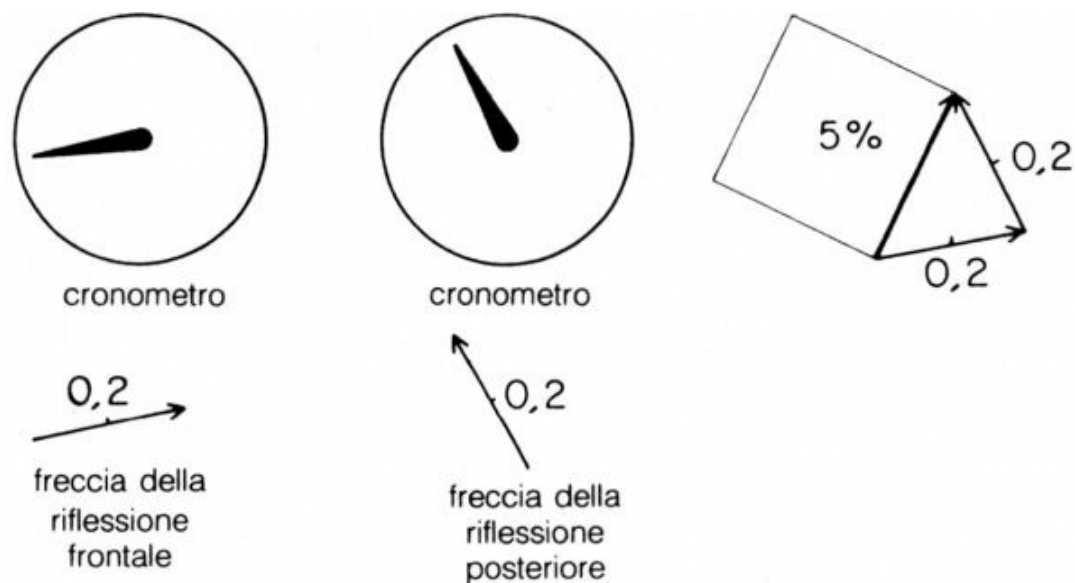


Fig. 13. Se la lamina di vetro è leggermente più spessa, la freccia risultante è un po' più lunga, perché è un po' più ampio l'angolo tra le frecce relative alla riflessione frontale e posteriore. Infatti un fotone che rimbalza sulla superficie posteriore impiega un po' più di tempo per arrivare in A rispetto all'esempio precedente.

Come ulteriore esempio, prendiamo il caso in cui lo spessore del vetro è tale che la lancetta del cronometro percorre esattamente mezzo giro in più quando segue un fotone che rimbalza sulla superficie posteriore. Questa volta la freccia per la riflessione posteriore si troverà diretta esattamente nello stesso verso di quella per la riflessione frontale. Combinandole, si ottiene una freccia finale di lunghezza 0,4, il cui quadrato è 0,16, che rappresenta una probabilità del 16% (fig. 14).

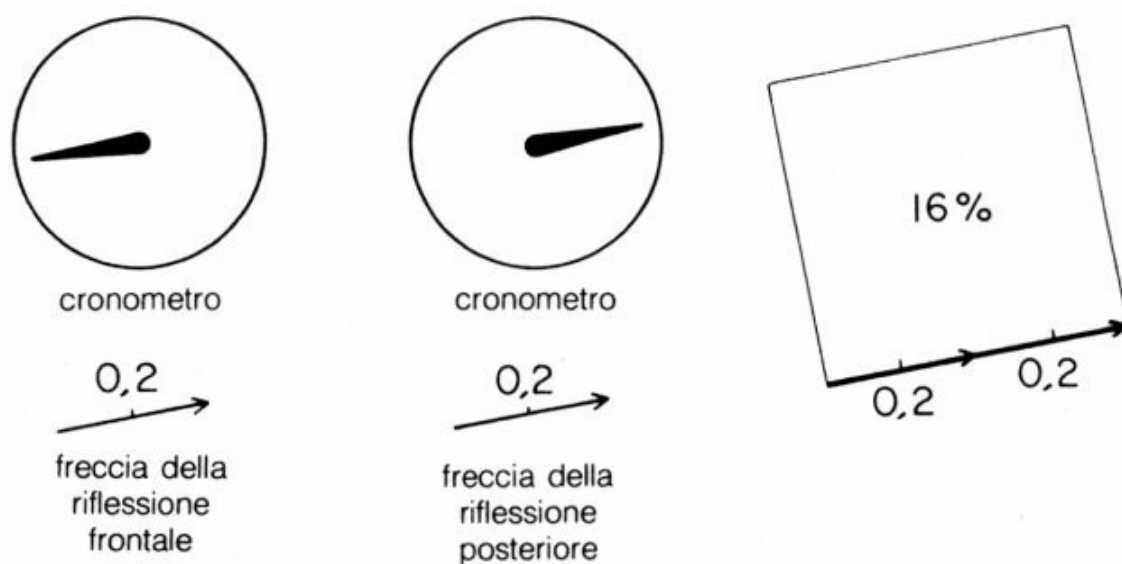


Fig. 14. Quando la lamina di vetro ha uno spessore tale da far compiere alla lancetta del cronometro che segue il fotone riflesso dalla superficie posteriore mezzo giro in più, le frecce relative alla riflessione frontale e posteriore hanno la stessa direzione e lo stesso verso e danno luogo a una freccia risultante di lunghezza 0,4, che rappresenta una probabilità del 16%.

Aumentando ancora lo spessore del vetro finché la lancetta del cronometro che segue il fotone riflesso dalla superficie posteriore compie un *giro completo* in più, le due frecce si troveranno ancora orientate in direzioni esattamente opposte, e la freccia finale sarà zero (fig. 15).

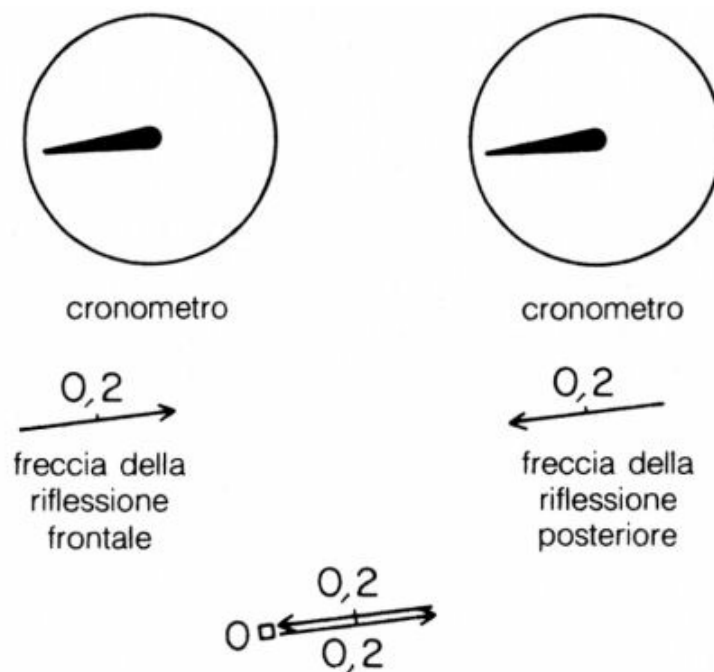


Fig. 15. Quando lo spessore della lamina di vetro è tale da far compiere alla lancetta del cronometro che segue il fotone riflesso dalla superficie posteriore un giro completo in più, la freccia finale è nuovamente zero e non si osserva nessuna riflessione.

Una situazione analoga esisterà ogni volta che lo spessore del vetro sarà tale che la lancetta che segue il fotone riflesso dalla superficie posteriore eseguirà un giro completo in più.

Se lo spessore del vetro è tale che la lancetta che cronometra il tempo della riflessione posteriore esegue esattamente $1/4$ o $3/4$ di giro in più, le due frecce si troveranno a formare un angolo retto tra loro. La freccia risultante sarà dunque l'ipotenusa di un triangolo rettangolo e, in base al teorema di Pitagora, il suo quadrato sarà uguale alla somma dei quadrati dei cateti. Perciò in questo caso si ha il valore esatto «due volte al giorno»: $4\% + 4\% = 8\%$ (fig. 16).

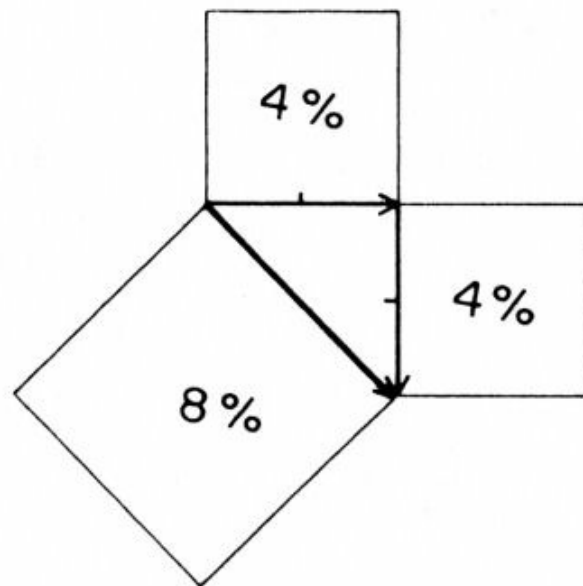


Fig. 16. Quando le frecce relative alla riflessione frontale e posteriore sono ad angolo retto tra loro, la freccia finale è l'ipotenusa di un triangolo rettangolo. Quindi il suo quadrato è la somma dei quadrati delle due frecce, pari all'8%.

Si osservi che via via che si aumenta lo spessore del vetro, la freccia della riflessione frontale conserva sempre la medesima direzione, mentre l'altra cambia gradualmente. Il cambiamento della direzione relativa delle due frecce produce una variazione ciclica della

lunghezza della freccia risultante tra zero e 0,4, cosicché il suo *quadrato* varia ciclicamente tra zero e il 16%, in accordo con le osservazioni sperimentali (fig. 17).

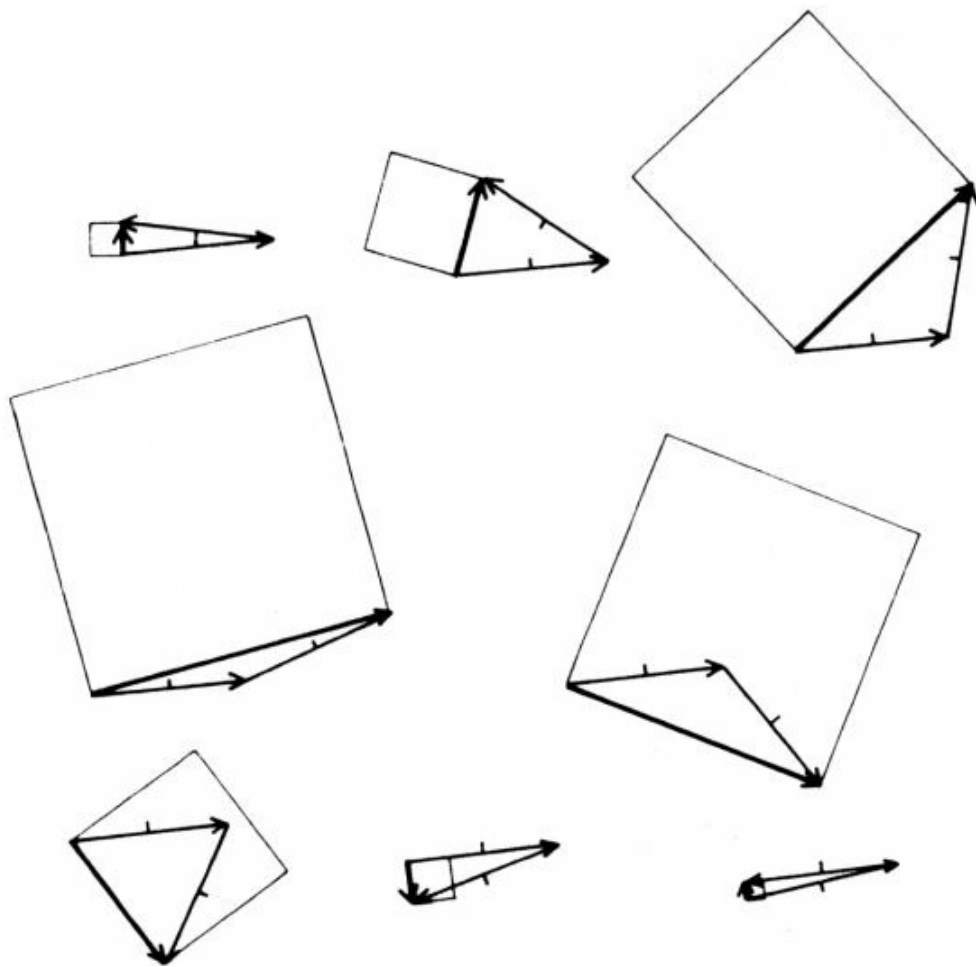


Fig. 17. Via via che le lamine di vetro vengono sostituite con lamine più spesse, la lancetta del cronometro che segue il fotone riflesso dalla superficie posteriore gira di più e l'angolo tra le frecce relative alla riflessione frontale e posteriore cambia. Ciò provoca una variazione ciclica nella lunghezza della freccia finale, il cui quadrato oscilla tra zero e il 16%.

Ho dunque fatto vedere come le strane proprietà della riflessione parziale possono essere esattamente descritte disegnando alcune banalissime frecce su un foglio di carta. In linguaggio tecnico queste frecce vengono chiamate «ampiezze di probabilità», e certo mi fa sentire più importante affermare che sto «calcolando l'ampiezza di probabilità per un evento». Preferisco però essere più onesto e dire che sto considerando le frecce il cui quadrato rappresenta la probabilità che l'evento si verifichi.

Prima di terminare questa lezione voglio parlare dei bellissimi colori che si osservano sulle bolle di sapone oppure sulla superficie scura di una pozzanghera fangosa su cui la vostra automobile ha sgocciolato un po' d'olio. La sottile pellicola oleosa che galleggia sulla pozzanghera agisce come una lamina di vetro molto sottile: riflette la luce di ciascun dato colore da zero fino a un valore massimo, a seconda del suo spessore. Illuminandola con luce rossa monocromatica, poiché lo spessore dello strato d'olio non è uniforme, si osservano chiazze rosse separate da tratti neri dove appunto la riflessione è nulla. Lo stesso accade con la luce blu: chiazze blu separate da tratti neri. Usando luce *sia* rossa *sia* blu, si osservano zone rosse dove lo spessore è tale da riflettere fortemente solo la luce rossa, altre zone che riflettono solo la luce blu, altre ancora che hanno lo spessore adatto per riflettere fortemente sia il rosso sia il blu (mistura percepita come viola), mentre altre zone dove la riflessione si annulla appaiono nere.

Per capire meglio questo punto, bisogna chiarire che nella riflessione parziale da due superfici il ciclo di variazione tra zero e il 16% si ripete più rapidamente per la luce blu che per quella rossa. Quindi per certi spessori sono fortemente riflessi uno solo o entrambi i colori, mentre per altri spessori la riflessione si annulla completamente (fig. 18).

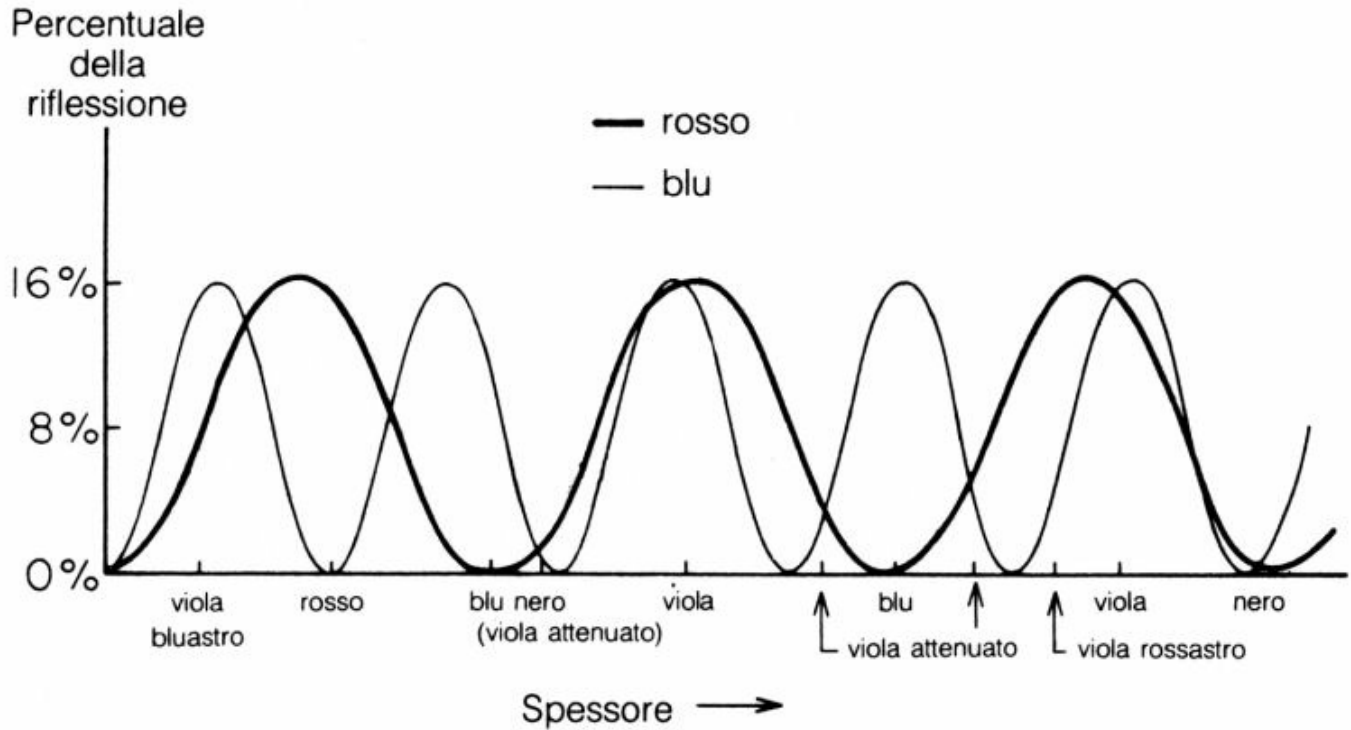


Fig. 18. Le due superfici di una lamina producono una riflessione parziale della luce monocromatica con probabilità che varia ciclicamente tra zero e il 16% all'aumentare dello spessore. Poiché la velocità di rotazione della lancetta del cronometro immaginario varia col colore della luce, il ciclo si ripete a spessori differenti per colori differenti. Quindi, se sulla lamina arriva luce di due colori, ad esempio luce rossa monocromatica e luce blu monocromatica, un dato spessore potrà dare luogo alla riflessione solo del rosso, o solo del blu, o del rosso e del blu in proporzioni differenti, producendo varie sfumature di viola, o potrà anche non riflettere alcun colore (e allora si avrà una zona nera). Se lo spessore della lamina varia da punto a punto, come avviene per una goccia d'olio che si allarga su una pozzanghera, si verificano tutte le situazioni possibili. Con la luce del sole, che contiene tutti i colori, si producono tutte le combinazioni possibili e la gamma dei colori osservabili è ricchissima.

I cicli della riflessione si ripetono con periodicità diversa, perché la lancetta del nostro cronometro immaginario deve girare più rapidamente quando segue un fotone blu che quando ne segue uno rosso. Anzi, la velocità di rotazione della lancetta è l'*unica* differenza tra un fotone rosso e uno blu, o un fotone di qualsiasi altro colore, inclusi i raggi X, le onde radio e così via.

Illuminando con luce rossa e blu una sottile pellicola oleosa, si osservano configurazioni rosse, blu e viola separate da bordi neri. Se la pozzanghera viene illuminata dalla luce del sole, che contiene il rosso, il giallo, il verde e il blu, le zone che riflettono fortemente ciascuno di questi colori si sovrappongono, dando vita alle combinazioni più varie che i nostri occhi percepiscono come colori diversi. Via via che la pellicola d'olio si muove e si spande sulla superficie dell'acqua, mutando spessore nei diversi punti, le configurazioni cromatiche cambiano. (Se però quella stessa pozzanghera voi la guardate di notte, illuminata da un lampione stradale al sodio, vedrete solo strisce giallastre separate da zone scure, perché i lampioni al sodio emettono solo luce gialla).

Questo fenomeno della molteplicità di colori prodotti dalla riflessione parziale della luce bianca su due superfici è chiamato iridescenza, e lo si può osservare in molte situazioni. Ad esempio, vi sarà capitato di domandarvi che cosa produca i variopinti colori

dei colibrì e dei pavoni. Ora lo sapete. Come questi colori sgargianti si siano evoluti è un altro problema interessante. Ogni volta che ammiriamo un pavone, dovremmo rivolgere un pensiero riconoscente a intere generazioni di femmine tutt'altro che fulgide, per essere state così esigenti nella scelta dei loro compagni. (In seguito è intervenuto anche l'uomo, orientando il processo di selezione).

Nella prossima lezione farò vedere come, grazie a questo stravagante procedimento di combinare tante piccole frecce, è possibile ottenere la risposta corretta per altri fenomeni comuni. Spiegherò perché la luce viaggia in linea retta, perché viene riflessa da uno specchio con un angolo uguale a quello con cui vi giunge («l'angolo di riflessione è uguale a quello di incidenza»), perché viene focalizzata da una lente, e così via. Vedrete così che questa nuova impostazione descrive tutte le proprietà della luce a voi note.

II

I FOTONI: PARTICELLE DI LUCE

Questa è la seconda di una serie di conferenze sull'elettrodinamica quantistica, e poiché vedo che il pubblico è completamente diverso (evidentemente perché ho detto che non si sarebbe capito nulla), riassumerò in breve il contenuto della prima lezione.

Abbiamo parlato della luce. La prima caratteristica importante della luce è che si comporta come se fosse formata da particelle: quando una luce monocromatica (cioè di un solo colore) molto debole colpisce un rivelatore, questo emette un ticchettio che diventa meno frequente, ma conserva la stessa intensità, via via che la luce viene resa più fioca.

L'altro aspetto importante del comportamento della luce riguarda la riflessione parziale della luce monocromatica. In media viene riflesso il 4% dei fotoni che colpiscono una *singola* superficie di vetro. Questo è già un mistero, poiché risulta impossibile prevedere quali fotoni rimbalzeranno e quali passeranno. Se poi si inserisce una *seconda* superficie, i risultati sono decisamente strani: invece della riflessione dell'8% che ci si aspetterebbe dalle due superfici, la riflessione parziale può essere amplificata fino al 16% o eliminata del tutto, a seconda dello spessore del vetro.

Questo strano aspetto della riflessione parziale da due superfici può essere spiegato dalla teoria ondulatoria finché si considera luce abbastanza intensa, ma tale teoria non può spiegare come mai il rivelatore continui a ticchettare con la stessa intensità allorché la luce diventa molto fioca. L'elettrodinamica quantistica «risolve» questa dualità ondulatorio-corpuscolare asserendo che la luce è formata da particelle (come sosteneva Newton). Tuttavia questo enorme progresso della scienza ha avuto come prezzo una «ritirata» della fisica, che ora è in grado di calcolare solo la *probabilità* che un dato fotone colpisca il rivelatore, senza offrire un buon modello intuitivo di come ciò avvenga effettivamente.

Nella prima lezione ho descritto come fanno i fisici a calcolare la probabilità che un determinato evento si verifichi. Si disegnano delle frecce su un foglio di carta secondo le seguenti regole.

PRINCIPIO FONDAMENTALE: la probabilità che un evento si verifichi è data dal quadrato della lunghezza di una freccia detta «ampiezza di probabilità». Ad esempio, una freccia lunga 0,4 rappresenta una probabilità di 0,16, cioè del 16%.

REGOLA GENERALE per disegnare la freccia relativa a un evento che può prodursi in più modi alternativi: si tracci una freccia per ciascun modo, e poi si combinino le varie frecce (le si «sommi») unendo la coda di ciascuna alla punta della precedente. Si disegni infine una «freccia risultante», unendo la coda della prima freccia alla punta dell'ultima. Il quadrato di questa freccia finale dà la probabilità complessiva che l'evento considerato si verifichi.

Ho poi illustrato le regole specifiche per disegnare le frecce nel caso della riflessione parziale sul vetro (si veda il capitolo precedente alle pp. 42-44 [intorno alle figure 8-11 in questo eBook]).

Questo per quanto riguarda il contenuto della prima lezione.

Ora vi farò vedere come questo modello del mondo, così radicalmente diverso da ogni descrizione a voi nota (e appunto per questo, magari, preferireste non saperne più niente), è in grado di spiegare tutte le proprietà più conosciute della luce: ad esempio, perché la luce viene riflessa da uno specchio con un angolo di riflessione uguale a quello di incidenza; perché si piega passando dall'aria all'acqua; perché in un mezzo omogeneo si propaga in linea retta; perché viene focalizzata dalle lenti e così via. La teoria descrive correttamente anche proprietà meno conosciute. In effetti, la difficoltà maggiore che ho incontrato nel preparare queste lezioni è stata di resistere alla tentazione di discutere tutte le proprietà della luce che di solito vengono faticosamente studiate a scuola, come il comportamento della luce quando penetra nella zona d'ombra al di là di uno spigolo netto (fenomeno noto come diffrazione). Ma poiché la maggior parte di voi non avrà mai fatto molta attenzione a tali fenomeni, li lascerò perdere. Vi garantisco però (altrimenti gli esempi che sto per fare sarebbero fuorvianti) che *tutti* i fenomeni osservati in cui interviene la luce possono essere spiegati dall'elettrodinamica quantistica; io qui mi limiterò a quelli più semplici e più comuni.

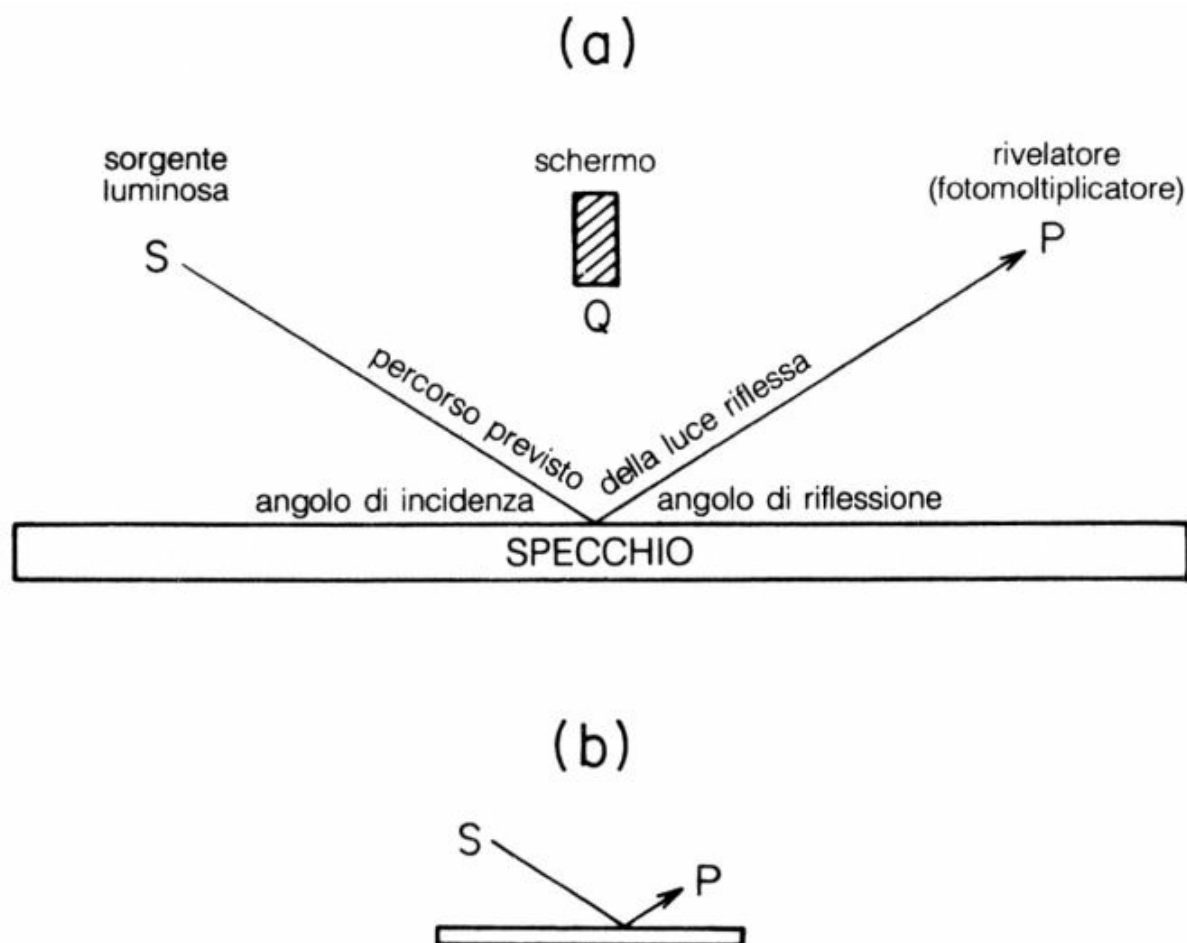


Fig. 19. Secondo la fisica classica, uno specchio riflette la luce solo nel punto in cui l'angolo di riflessione è uguale a quello di incidenza, anche quando sorgente e rivelatore sono a distanze diverse dallo specchio, come nel caso (b).

Comincerò dal problema di determinare come la luce si riflette su uno specchio (fig.

19). In S è situata una sorgente che emette luce monocromatica (mettiamo ancora rossa) di bassissima intensità: un fotone alla volta. In P si trova un fotomoltiplicatore che rivela i fotoni, e poiché è più facile tracciare le frecce se la configurazione è simmetrica, poniamolo, rispetto allo specchio, alla stessa distanza della sorgente. Si vuole calcolare la probabilità che il rivelatore faccia un «clic» in seguito all'emissione di un fotone dalla sorgente. Poiché c'è anche la possibilità che un fotone vada direttamente nel rivelatore, interponiamo uno schermo in Q.

Ci aspetteremmo che tutta la luce che arriva al rivelatore sia stata riflessa dalla parte centrale dello specchio, perché solo in tale zona l'angolo di incidenza è uguale a quello di riflessione, mentre le parti dello specchio vicine alle estremità dovrebbero essere completamente estranee a tutta la faccenda. Giusto?

Così parrebbe, ma vediamo che cosa dice la teoria quantistica. Principio fondamentale: la probabilità che un evento si verifichi è data dal quadrato della freccia risultante ottenuta tracciando una freccia per ciascun modo in cui l'evento può avvenire e poi combinando («sommando») le varie frecce. Nell'esperimento in cui si considerava la riflessione parziale della luce da due superfici, un fotone poteva arrivare dalla sorgente al rivelatore in due modi. Nell'esperimento che stiamo discutendo ora ci sono invece milioni di percorsi che un fotone può seguire: può andare a colpire l'estremità sinistra dello specchio, ad esempio in A o in B, e di lì rimbalzare fino al rivelatore (fig. 20); può rimbalzare nel punto che voi vi aspettereste, cioè al centro, in G; oppure andare fino all'estremità destra dello specchio, in K o in M, e di lì al rivelatore. Forse penserete che sono diventato matto, perché per la maggior parte di questi percorsi gli angoli di incidenza e di riflessione non sono uguali. Nossignori, non sono matto: la luce si propaga proprio così! Ma come può essere?

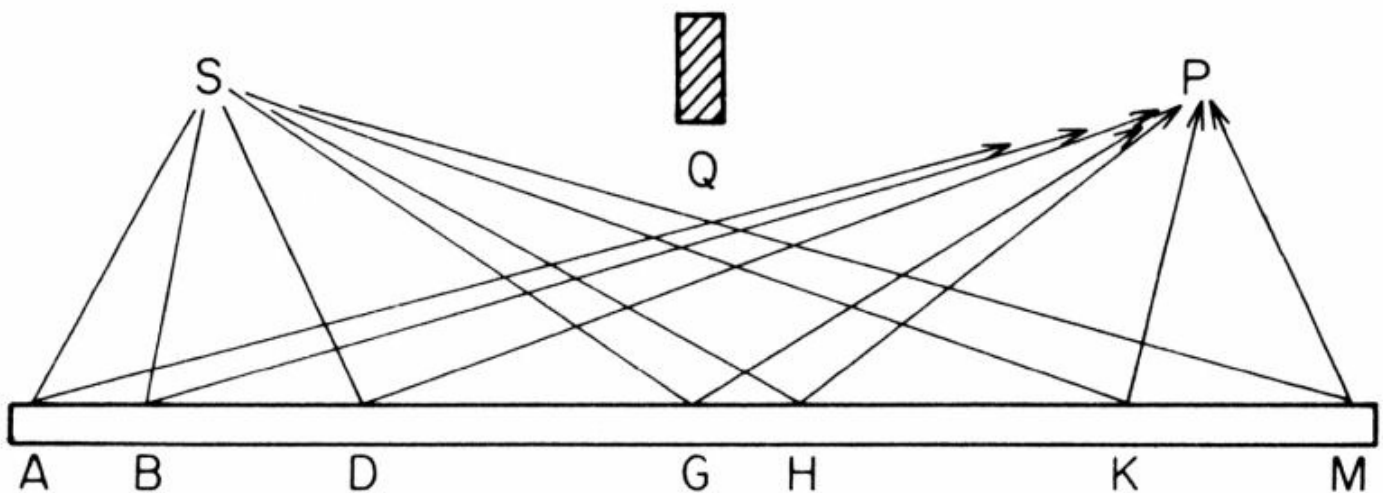


Fig. 20. Secondo l'elettrodinamica quantistica, la luce ha la stessa ampiezza di probabilità di riflettersi in ogni punto dello specchio, da A a M.

Per facilitare le cose, supponiamo che lo specchio sia costituito semplicemente da un lunga striscia che va da sinistra a destra, dimenticando per un momento che esso si estende anche fuori dal piano della carta (fig. 21). Inoltre, mentre la luce può riflettersi in un numero immenso di punti, facciamo per ora l'approssimazione di suddividere lo specchio in un numero finito di piccoli quadrati e di considerare un solo percorso per

ciascun quadrato; più piccoli sono i quadrati, e più numerosi quindi sono i percorsi presi in considerazione, più il calcolo diventa accurato, ma anche complicato.



Fig. 21. Per calcolare più facilmente come si propaga la luce, si consideri momentaneamente solo una striscia di specchio divisa in tanti piccoli quadrati e un solo percorso per ciascun quadrato. Questa semplificazione non pregiudica in alcun modo l'accuratezza dell'analisi.

Per ogni possibile percorso della luce si disegni ora una piccola freccia, che avrà una certa lunghezza e una certa direzione. Cominciamo con la lunghezza. Si potrebbe pensare che la freccia corrispondente al percorso che passa per il punto centrale dello specchio, G, sia decisamente la più lunga, poiché sembra esservi un'elevatissima probabilità che un fotone che arriva nel rivelatore abbia seguito tale cammino, mentre le frecce relative a cammini che passano per le estremità dello specchio dovrebbero essere molto piccole. E invece no: tale regola sarebbe completamente arbitraria. La regola corretta, che descrive ciò che realmente accade, è molto più semplice: la probabilità che ha un fotone di giungere nel rivelatore è quasi la stessa *qualunque* sia il percorso seguito, per cui le nostre frecce hanno tutte quasi la stessa lunghezza. (Vi sono, a dire il vero, alcune piccole differenze dovute alla diversità di angolo e di distanza, ma sono talmente secondarie che le trascurerò). Tutte le frecce avranno dunque la stessa lunghezza arbitraria: scegliamola molto piccola perché le frecce sono molte, corrispondenti a moltissimi percorsi possibili (fig. 22).



Fig. 22. Nei calcoli che verranno discussi, ogni percorso che la luce può seguire sarà rappresentato da una freccia di lunghezza arbitraria fissata, come in figura.

Mentre si può tranquillamente assumere che la lunghezza di tutte le frecce sia più o meno la stessa, le loro direzioni sono chiaramente diverse, essendo diverso il tempo impiegato a seguire i vari percorsi; si ricordi che, come ho spiegato nella prima lezione, la direzione di una data freccia è determinata dalla posizione di arrivo di un'immaginaria lancetta di cronometro, che ruota mentre il fotone segue il suo percorso. Un fotone che vada prima all'estremità sinistra dello specchio, in A, impiega chiaramente più tempo per arrivare nel rivelatore di un altro fotone che si rifletta nella zona centrale dello specchio, in G (fig. 23).

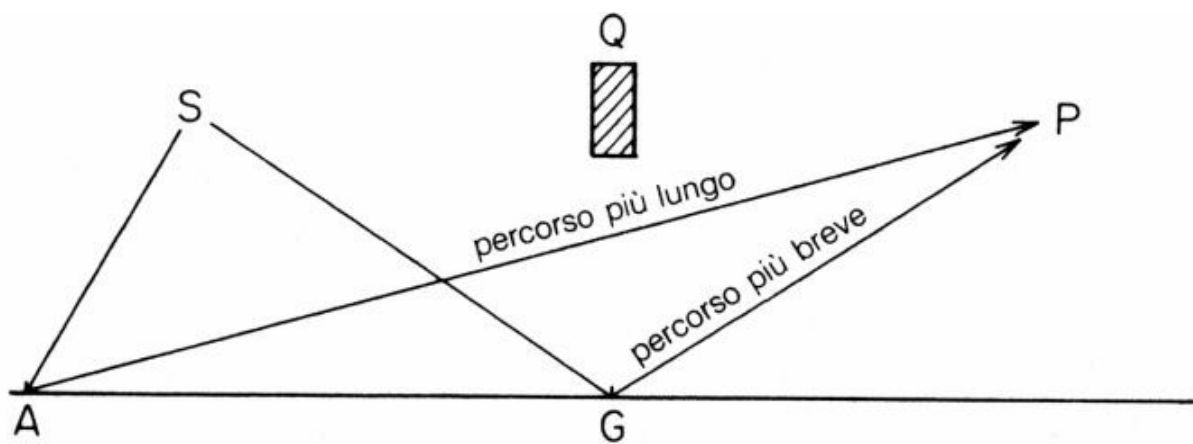


Fig. 23. Mentre la lunghezza delle varie frecce è essenzialmente la stessa, la loro direzione è diversa perché sono diversi i tempi impiegati da un fotone per seguire percorsi diversi. È evidente che occorre più tempo per andare da S in A e poi in P che non da S in G e poi in P.

Immaginate di avere fretta e di dover correre dalla sorgente allo specchio e poi di lì al rivelatore: capireste subito che non è il caso di andare prima fino in A e poi risalire fino al rivelatore; si fa molto prima a toccare lo specchio in un punto della zona centrale.

Per facilitare il calcolo della direzione delle varie frecce, disegnerò un grafico sotto la raffigurazione schematica dello specchio (fig. 24).

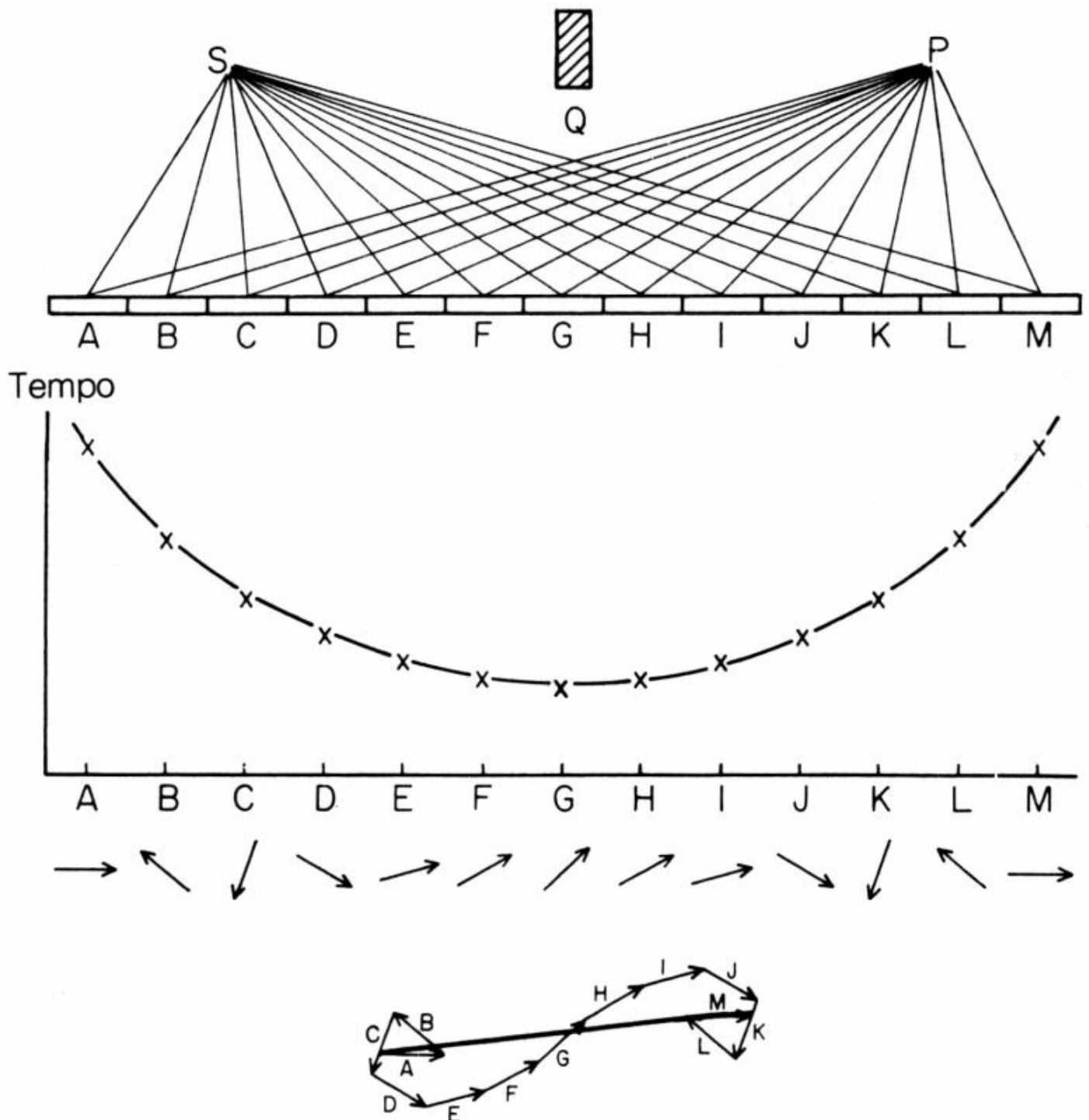


Fig. 24. Nella parte superiore della figura sono mostrati i percorsi che la luce può seguire (nella nostra situazione semplificata), mentre i punti corrispondenti nel grafico sottostante indicano il tempo che un fotone impiega per andare dalla sorgente a quel punto e di lì al fotomoltiplicatore. Subito sotto il grafico sono rappresentate le direzioni delle varie frecce e, sotto ancora, il risultato della loro somma. È evidente che il contributo principale alla lunghezza della freccia finale proviene dalle frecce comprese tra E ed I, le cui direzioni sono quasi le stesse, perché i tempi relativi ai corrispondenti percorsi sono quasi gli stessi. Questa è anche la zona in cui il tempo impiegato è minimo. E perciò approssimativamente corretto dire che la luce segue il percorso che richiede il minimo tempo.

In corrispondenza di ciascun punto dello specchio è indicato, lungo la verticale, il tempo impiegato dalla luce per giungere nel rivelatore seguendo il percorso che passa per quel punto. Quanto maggiore è questo tempo, tanto più in alto nel grafico sarà il punto corrispondente. Incominciamo da sinistra: un fotone che si riflette in A impiegherà un tempo piuttosto lungo, cosicché il punto corrispondente nel grafico sarà piuttosto in alto. Spostandosi verso il centro dello specchio, il tempo impiegato diminuisce, e i punti sul grafico sono via via più in basso. Superato il centro, il tempo necessario a seguire un percorso ricomincia a crescere, e il punto corrispondente nel grafico verrà a trovarsi sempre più in alto. Per aiutare l'occhio, congiungiamo i vari punti: otterremo una curva

simmetrica che inizia dall'alto, scende e poi risale.

Che cosa significa ciò per la direzione delle nostre frecce? La direzione di ciascuna di esse è in corrispondenza col tempo impiegato da un fotone a percorrere un determinato percorso. Disegniamo le varie frecce cominciando da sinistra. Il percorso A richiede il tempo maggiore e la freccia corrispondente sarà orientata in una certa direzione. La freccia per il percorso B sarà orientata lungo una direzione diversa perché il tempo impiegato è diverso. Le frecce relative ai percorsi F, G ed H, che passano per la zona centrale dello specchio, avranno invece direzione pressoché identica perché i tempi corrispondenti sono quasi gli stessi. Superato il centro dello specchio, si vede che ad ogni percorso sul lato destro dello specchio corrisponde un percorso sul lato sinistro che ha lo stesso tempo di percorrenza (questa è una conseguenza della decisione di porre sorgente e rivelatore alla stessa distanza dallo specchio e simmetrici rispetto al suo centro). La freccia per il percorso J, ad esempio, avrà la stessa direzione di quella per il percorso D.

Adesso sommiamo le varie frecce, partendo dalla freccia A e agganciando la coda di ciascuna alla punta della precedente. Se volessimo fare una passeggiata usando le frecce come passi, all'inizio non ci allontaneremmo gran che dal punto di partenza, perché ogni passo porta in una direzione molto diversa dal precedente. Dopo un po', però, le frecce cominciano a rivolgersi tutte pressappoco nella stessa direzione, e allora si progredisce decisamente. Verso la fine della passeggiata ciascun passo è di nuovo in una direzione del tutto diversa dal precedente, per cui non si avanza più.

Non resta che tracciare la freccia risultante. Si congiunge la coda della prima freccia alla punta dell'ultima e si ottiene lo spostamento complessivo della passeggiata. E notate bene che la freccia finale ha lunghezza non trascurabile! L'elettrodinamica quantistica predice che la luce effettivamente si riflette su uno specchio!

Questo risultato richiede alcune considerazioni. Che cosa determina la lunghezza della freccia finale? Vediamo anzitutto che le estremità dello specchio non hanno nessuna importanza: le frecce relative a questi contributi puntano in tutte le direzioni, ma non producono alcuno spostamento complessivo. Se si tagliassero queste estremità, che a voi saranno subito apparse, istintivamente, di nessunissimo interesse, la lunghezza della freccia finale non risulterebbe quasi per niente modificata.

Quale parte dello specchio dà dunque a tale freccia la sua lunghezza apprezzabile? Chiaramente quella le cui frecce hanno tutte quasi la stessa direzione, perché il *tempo* corrispondente è quasi lo *stesso*. Osservando il grafico in cui sono riportati i tempi relativi ai vari percorsi (fig. 24), si vede che le variazioni da un percorso al successivo sono molto ridotte nella zona più bassa della curva, dove il *tempo* impiegato è *minimo*.

Riassumendo, la zona per cui il tempo è minimo è anche quella per cui tempi relativi a percorsi vicini sono quasi gli stessi; le frecce relative a questa zona sono orientate tutte pressappoco nella stessa direzione e si sommano dando un contributo apprezzabile; questa è la zona che determina la probabilità con cui un fotone si riflette su uno specchio. Ed è per questo che, approssimando, possiamo prendere per buona la descrizione del mondo, piuttosto rozza, secondo la quale la luce segue soltanto il percorso che richiede il *minimo tempo*. (È facile dimostrare che per tale percorso l'angolo d'incidenza è uguale all'angolo di riflessione, ma per brevità non lo faremo).

L'elettrodinamica quantistica fornisce dunque la risposta corretta: la zona centrale dello specchio è quella rilevante nel determinare la riflessione; ma questa risposta è stata ottenuta al prezzo di dover credere che la luce si riflette in tutti i punti dello specchio e di dover sommare un gran numero di piccole frecce il cui solo effetto è di cancellarsi tra loro. Tutto ciò può sembrare una perdita di tempo, un giochino stupido buono solo per i matematici. Non sembra «fisica vera» dover introdurre qualcosa solo per poi cancellarlo.

Per verificare se davvero c'è riflessione in tutti i punti dello specchio, facciamo un altro esperimento. Cominciamo con l'eliminare gran parte dello specchio, lasciandone circa un quarto, all'estremità sinistra. Il pezzo rimasto è ancora abbastanza grande, ma è nel posto sbagliato. Si è appena visto (fig. 24) che le frecce relative ai contributi dell'estremità sinistra dello specchio hanno direzioni molto diverse tra loro a causa della notevole differenza di tempo tra percorsi adiacenti. Facciamo adesso un calcolo più accurato, suddividendo la parte rimasta dello specchio in intervalli molto più piccoli, tali che non ci sia molta differenza di tempo tra percorsi vicini (fig. 25).

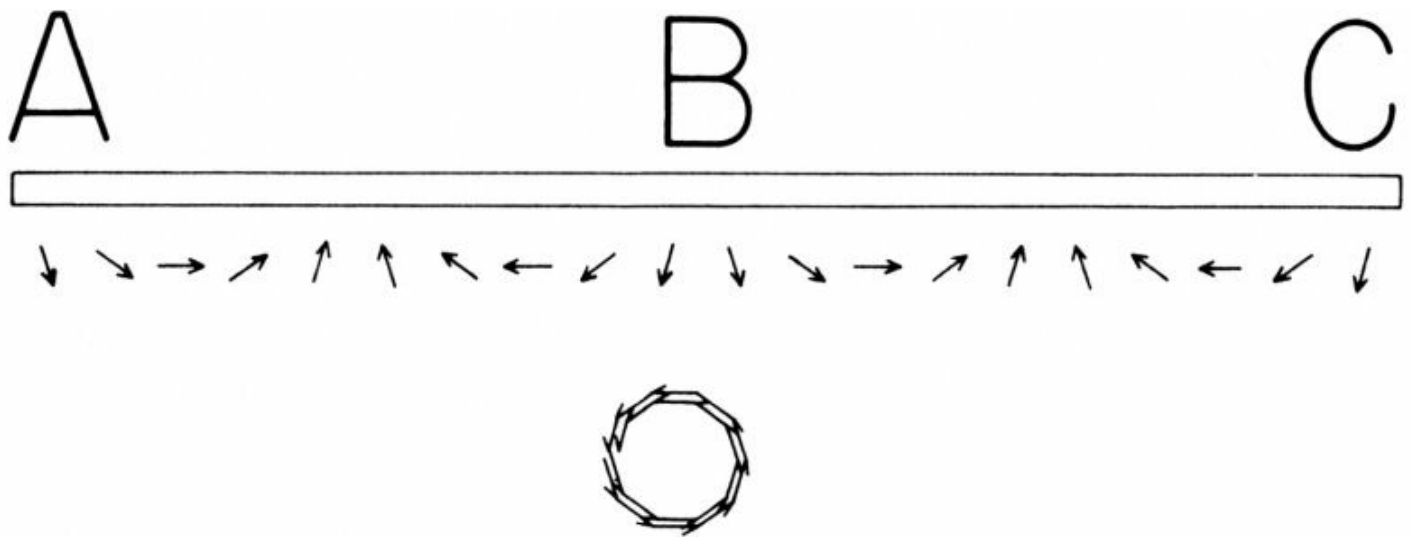


Fig. 25. Per verificare che la riflessione ha luogo effettivamente anche alle estremità dello specchio (anche se poi si elide nella somma finale), si può fare un esperimento con un grande pezzo di specchio disposto in posizione errata per dar luogo alla riflessione da S a P. Si consideri questo pezzo di specchio suddiviso in tratti molto più piccoli, così che il tempo impiegato differisca di poco da un percorso al successivo. Quando tutte le frecce vengono sommate il risultato è nullo: le varie frecce sono disposte in cerchio e la loro somma è praticamente zero.

In questa rappresentazione più dettagliata vediamo che alcune frecce puntano più o meno verso destra mentre le altre puntano più o meno verso sinistra. Quando si sommano *tutte* le frecce si vede che esse sono disposte sostanzialmente lungo una circonferenza, e che quindi non si arriva da nessuna parte. Ma supponiamo di raschiar via per bene dallo specchio le zone a cui corrispondono frecce dirette prevalentemente in una data direzione, ad esempio verso sinistra, in modo che rimangano solo le zone le cui frecce puntano prevalentemente nella direzione opposta (fig. 26).

La somma delle frecce dirette tutte più o meno verso destra darà una serie di semicerchi e produrrà una freccia risultante apprezzabile; in base alla teoria adesso si dovrebbe osservare una riflessione notevole! E infatti così è: la teoria è *corretta*! Uno specchio del genere viene chiamato reticolo di diffrazione e funziona a meraviglia.

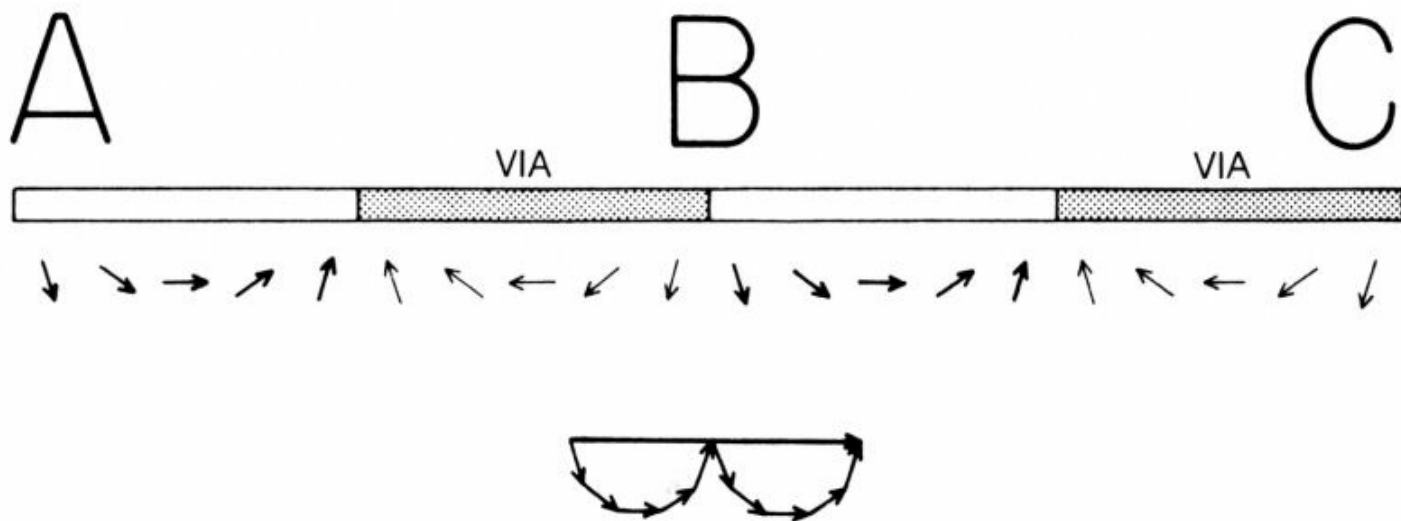


Fig. 26. Se si sommano solo le frecce inclinate verso una particolare direzione, ad esempio verso destra, e si eliminano le altre raschiando il vetro nei punti corrispondenti, si osserva che una notevole quantità di luce viene riflessa dal pezzo di specchio situato nel posto sbagliato. Uno specchio opportunamente raschiato è chiamato reticolo di diffrazione.

Non è fantastico? Si prende un pezzo di specchio su cui non ci si aspetterebbe alcuna riflessione, se ne raschia via una parte ed ecco che lo specchio riflette!⁶

Questo particolare reticolo è fatto su misura per la luce rossa, e non funzionerebbe con luce blu, per la quale occorre invece un reticolo in cui i tratti eliminati siano più vicini tra loro, perché, come ho detto nella prima lezione, la lancetta del cronometro immaginario deve girare più velocemente quando segue un fotone blu che quando ne segue uno rosso. Le strisce opache adatte per la velocità di rotazione del rosso non si trovano nei posti giusti per il blu; le frecce si ingarbugliano e il reticolo non funziona bene. Il caso vuole però che se si sposta il rivelatore in modo che formi un angolo leggermente diverso, il reticolo adatto alla luce rossa diventa adatto a quella blu. Si tratta solamente di un caso fortunato, dovuto alla particolare configurazione geometrica del sistema (fig. 27).

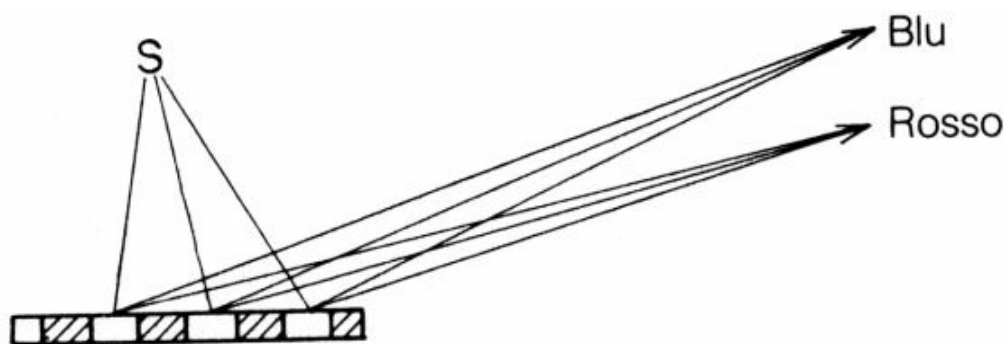


Fig. 27. Un reticolo di diffrazione i cui solchi sono distanziati correttamente per la luce rossa funziona anche per gli altri colori se si pone il rivelatore in posizione differente. Pertanto, a seconda dell'angolo di osservazione, è possibile vedere diversi colori riflessi da una superficie rigata, ad esempio quella di un disco di grammofono.

Illuminando il reticolo con luce bianca, il rosso emergerà ad un certo angolo, l'arancione a un angolo leggermente maggiore, poi verranno il giallo, il verde e il blu, e tutti i colori dell'arcobaleno. Succede comunemente di vedere riflessa luce di vari colori quando si guarda sotto l'angolo giusto un oggetto, illuminato con luce bianca, su cui è incisa una serie di solchi sufficientemente vicini tra loro, ad esempio un disco di grammofono o ancor meglio un videodisco. Avrete certo visto quei meravigliosi adesivi argentati che qui, nell'assolata California, vengono spesso incollati sul retro delle

automobili: quando l'automobile si muove si vedono vivaci colori cangianti dal rosso al blu. Adesso sapete da dove vengono: a causarli è un reticolo di diffrazione, uno specchio raschiato nei punti giusti. Il sole è la sorgente di luce e gli occhi sono i rivelatori. A questo punto potrei spiegare in termini facili come funzionano i laser e gli ologrammi, ma so che non tutti ne hanno conoscenza diretta, e poi ci sono molte altre cose di cui voglio parlare.⁷

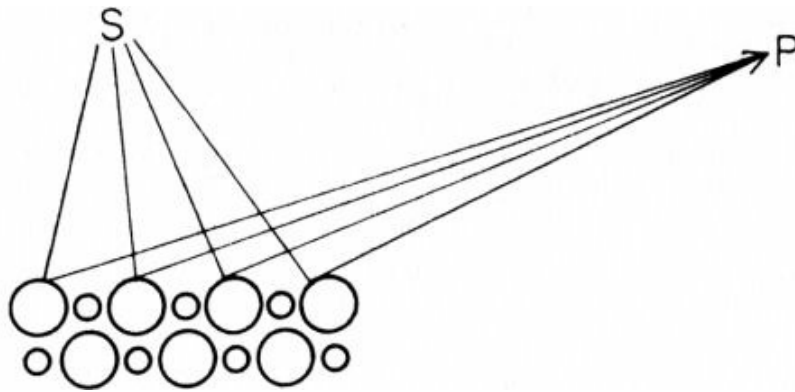


Fig. 28. La Natura ha prodotto molti tipi di reticoli di diffrazione sotto forma di cristalli. Un cristallo di sale riflette solo per alcuni angoli i raggi X, che sono luce per la quale la lancetta del cronometro immaginario ruota a grandissima velocità, anche 10.000 volte maggiore che per la luce visibile. Dagli angoli di riflessione si può risalire all'esatta disposizione spaziale dei singoli atomi.

Un reticolo di diffrazione dimostra quindi che non si possono trascurare le parti di uno specchio che non sembrano contribuire alla riflessione; manipolandolo opportunamente si può mettere in evidenza la realtà della riflessione da tutti i suoi punti e produrre, alcuni fenomeni ottici sorprendenti.

Ma ancor più importante è il fatto che, dimostrando la realtà del contributo di *tutte* le parti dello specchio alla riflessione, si dimostra che c'è un'ampiezza (una freccia) per *ciascun modo* in cui un evento può accadere. E per calcolare correttamente la probabilità di un tale evento nelle diverse circostanze, occorre sommare le frecce relative a *tutti* i modi in cui esso può verificarsi, e non solo a quelli che si ritengono rilevanti!

Parlerò adesso di un fenomeno più familiare dei reticoli di diffrazione, cioè del passaggio della luce dall'aria all'acqua. Immaginiamo che un fotomoltiplicatore possa essere immerso nell'acqua, e collochiamolo nella posizione D, mentre la sorgente si trova nell'aria, in S (fig. 29).

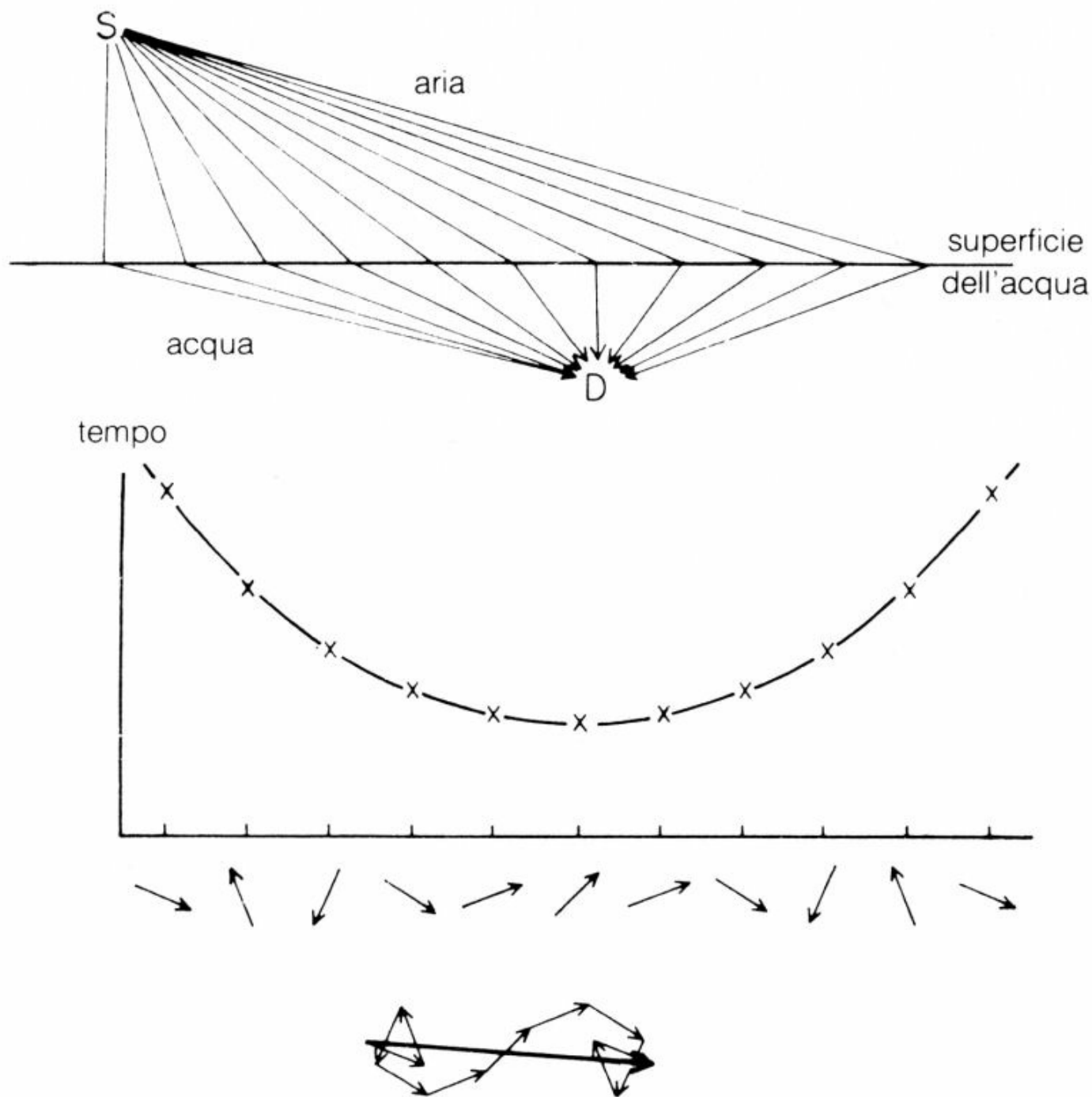


Fig. 29. Secondo la teoria quantistica, la luce può andare da una sorgente nell'aria a un rivelatore nell'acqua seguendo molti percorsi. Se si semplifica il problema, come nel caso dello specchio, si può riportare in un grafico il tempo relativo ai valori percorsi, e sotto di esso tracciare la direzione delle relative frecce. Anche questa volta i contributi maggiori alla lunghezza della freccia risultante vengono da quei percorsi le cui frecce puntano pressappoco nella stessa direzione, perché i tempi corrispondenti sono quasi gli stessi, e anche questa volta ciò avviene quando il tempo impiegato è minimo.

Anche in questo caso vogliamo calcolare la probabilità che un fotone arrivi dalla sorgente al rivelatore. Per farlo dobbiamo prendere in considerazione tutti i possibili percorsi della luce. Ciascuno di questi percorsi aggiunge una piccola freccia e, come nell'esempio precedente, tutte queste frecce hanno quasi la stessa lunghezza. Se facciamo un grafico del tempo necessario a un fotone per compiere ciascuno dei possibili percorsi, la curva che si ottiene è molto simile a quella che si aveva per la luce riflessa da uno specchio: parte dall'alto, scende e poi sale di nuovo; il contributo più importante viene dai luoghi per i quali le frecce puntano pressappoco nella stessa direzione, cioè per i quali il tempo di percorrenza varia poco tra percorsi vicini, ossia dalla parte più bassa della curva.

Questa è anche la zona per cui il tempo è minimo, sicché in realtà basta trovare il percorso di minimo tempo.

Ora, la luce sembra propagarsi più lentamente nell'acqua che nell'aria (il perché lo spiegherò nella prossima lezione), e ciò rende il percorso attraverso l'acqua più 'dispendioso', per così dire, di quello attraverso l'aria. Non è difficile individuare quale percorso richieda il minimo tempo: immaginate di essere un bagnino seduto in S, e che in D ci sia una bella ragazza che sta per affogare (fig. 30). Poiché correre sulla terra è più veloce che nuotare, in che punto entrereste in acqua per raggiungere la bagnante che annega nel più breve tempo possibile?

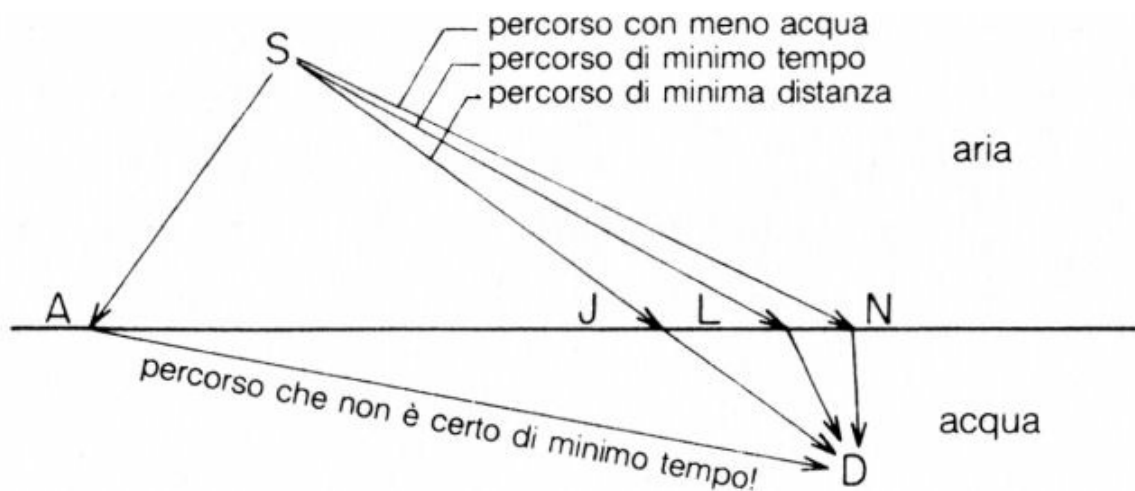


Fig. 30. Trovare il percorso di minimo tempo per la luce è come trovare il percorso che richiede meno tempo a un bagnino che deve attraversare di corsa una spiaggia e poi nuotare per salvare una persona che affoga: lungo il percorso più corto c'è troppa acqua, mentre lungo quello con meno acqua c'è troppa spiaggia; il percorso che richiede il minimo tempo è un compromesso tra i due.

Entrereste in acqua al punto A, per poi mettervi a nuotare come forsennati? Naturalmente no. Ma anche correre direttamente verso la ragazza in pericolo, entrando nell'acqua in J, non è il percorso più rapido. È chiaro che un bagnino non si ferma a fare simili analisi in circostanze d'emergenza, ma è possibile calcolare quale percorso richiede il minimo tempo: si tratta di un compromesso tra il percorso diritto, che passa per J, e quello con meno acqua, che passa per N. La stessa cosa succede per la luce: il percorso che richiede il minimo tempo entra nell'acqua in un punto tra J ed N, ad esempio in L.

Un altro fenomeno cui voglio accennare brevemente è quello del miraggio. Guidando su una strada molto calda, si ha a volte l'impressione di vederla piena di pozzanghere. In realtà quello che si vede è il riflesso del cielo; però quando si vede il cielo sulla strada è perché ci sono delle pozzanghere (riflessione parziale su una singola superficie!). Come mai allora si vede il cielo riflesso sulla strada quando non c'è la minima traccia d'acqua? Bisogna sapere che la luce viaggia più lentamente attraverso l'aria fredda che attraverso quella calda, e per poter vedere il miraggio, l'osservatore deve trovarsi nella zona d'aria più fredda che è sopra lo strato d'aria calda immediatamente vicino alla superficie della strada (fig. 31).



Fig. 31. Trovare il percorso di minimo tempo spiega anche come si produce un miraggio. La luce viaggia più velocemente nell'aria calda che in quella fredda. Quando si vede il riflesso del cielo sulla strada è perché parte della luce del cielo che arriva nell'occhio proviene dalla strada. Poiché l'unica altra situazione in cui si vede il cielo sulla strada si verifica quando c'è acqua che lo riflette, il miraggio viene attribuito alla presenza di acqua.

Come sia possibile vedere il cielo guardando in giù può essere compreso trovando il percorso di minimo tempo. Provate a pensarci da soli: è un problema divertente e abbastanza facile.

Negli esempi finora discussi (riflessione della luce su uno specchio o il suo passaggio dall'aria all'acqua) ho fatto un'approssimazione: per semplicità ho disegnato i percorsi seguiti dalla luce come due tratti rettilinei che formano un angolo tra loro. Ma non c'è bisogno di *assumere* che la luce si propaghi in linea retta quando viaggia attraverso un materiale omogeneo come l'acqua o l'aria; anche questo può essere spiegato in base al principio generale della teoria quantistica: la probabilità di un evento si trova sommando le frecce relative a *tutti* i modi in cui quell'evento può accadere.

Perciò, sommando freccette, farò vedere nel prossimo esempio come si possa avere l'impressione che la luce si propaghi in linea retta. Disponiamo una sorgente e un rivelatore rispettivamente in S e in P (fig. 32) e consideriamo *tutti* i percorsi, comunque contorti, che la luce può seguire per andare dalla sorgente al fotomoltiplicatore. Disegniamo la freccetta relativa a ciascun cammino, e vediamo se si è imparata la lezione!

Per ogni percorso contorto, ad esempio A, c'è un percorso vicino un po' più dritto e decisamente più breve, che cioè prende molto meno tempo. Ma se si considerano percorsi quasi rettilinei, come C, quelli vicini e leggermente più dritti richiedono quasi lo stesso tempo. Per questi ultimi percorsi le frecce si sommano invece di elidersi ed essi sembrano i soli seguiti dalla luce.

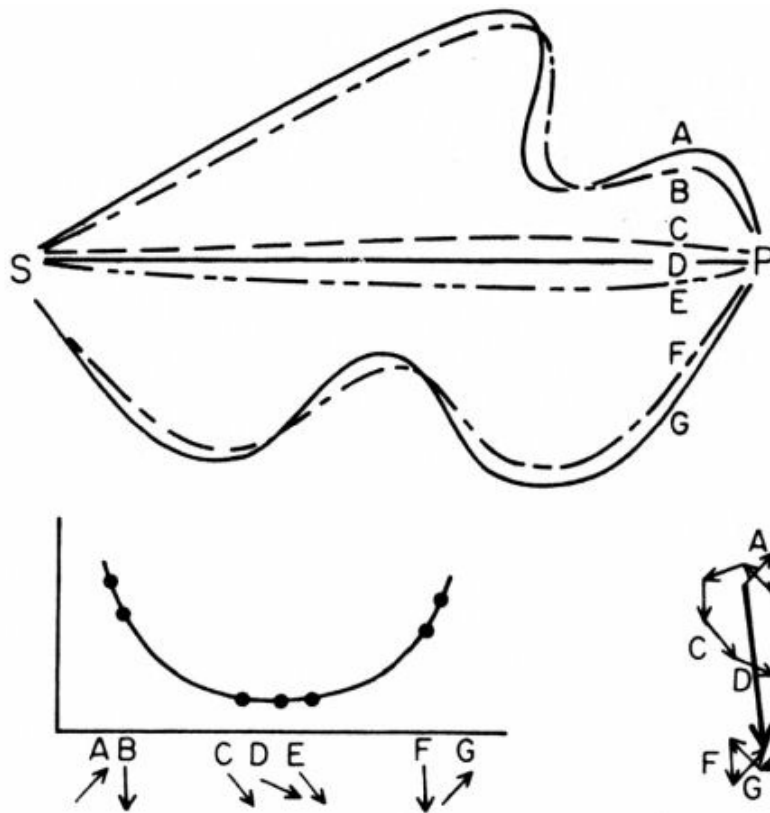


Fig. 32. Usando la teoria quantistica, si può capire perché la luce sembra viaggiare in linea retta. Se si prendono in considerazione tutti i percorsi possibili, si vede che per ogni percorso contorto ce n'è uno vicino di lunghezza notevolmente inferiore, che richiede perciò molto meno tempo e la cui freccia ha una direzione notevolmente diversa. Solo i percorsi vicini a quello rettilineo che passa per D hanno frecce che puntano tutte quasi nella stessa direzione, perché i tempi relativi sono quasi gli stessi. Soltanto tali frecce risultano importanti, perché sono le uniche a dar luogo a una freccia finale apprezzabile.

È importante osservare che la freccia relativa al solo percorso dritto, quello che passa per D, non è sufficiente a dar conto della probabilità con cui la luce arriva dalla sorgente al rivelatore. Anche i percorsi quasi dritti e vicini, attraverso C o E, per esempio, danno contributi importanti. Per cui la luce non viaggia *realmente* lungo una linea retta: essa 'annusa' i percorsi vicini e usa una piccola zona di spazio tutt'attorno. (Analogamente uno specchio deve avere dimensioni abbastanza grandi per dar luogo alla normale riflessione; se è troppo piccolo per poter riflettere tutto il fascio di percorsi importanti, diffonderà la luce in tutte le direzioni dovunque lo si metta).

Per esaminare meglio il comportamento di questo fascio di luce, mettiamo una sorgente in S, un fotomoltiplicatore in P e interponiamo tra loro due blocchi per evitare che i percorsi della luce si allontanino troppo da quello più breve (fig. 33).

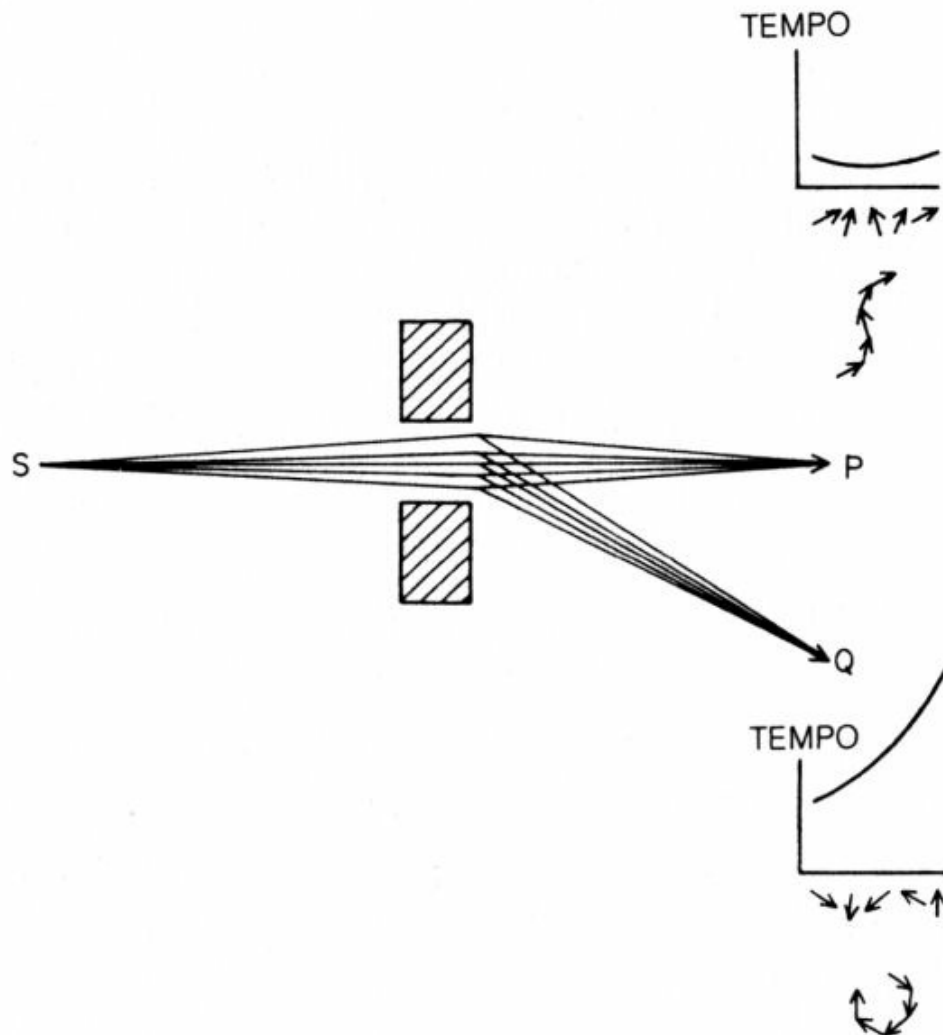


Fig. 33. La luce non segue solo il percorso rettilineo, ma anche tutto il fascio dei percorsi vicini. Se due blocchi sono abbastanza separati da lasciar passare questo fascio, i fotoni vanno normalmente in P e quasi mai in Q.

Collochiamo anche un secondo fotomoltiplicatore in Q, sotto P, e assumiamo di nuovo, per semplicità, che la luce possa andare da S a Q solo seguendo percorsi fatti da due tratti rettilinei. Che cosa osserviamo? Se lo spazio tra i blocchi è sufficiente a permettere molti percorsi vicini, le frecce dei percorsi che arrivano in P si sommano, perché tali percorsi richiedono quasi lo stesso tempo, mentre quelle relative ai percorsi fino a Q si elidono, perché tra di essi vi sono sensibili differenze di tempo. Il fotomoltiplicatore in Q, perciò, non ticchetta.

Ma se avviciniamo i blocchi tra loro, a un certo punto il fotomoltiplicatore in Q inizia a ticchettare! Quando la fenditura tra i due blocchi tende a chiudersi, e vi sono perciò solo pochi percorsi vicini, anche le frecce dei percorsi in direzione di Q cominciano a sommarsi, perché tra loro non vi è quasi più differenza di tempo (fig. 34).

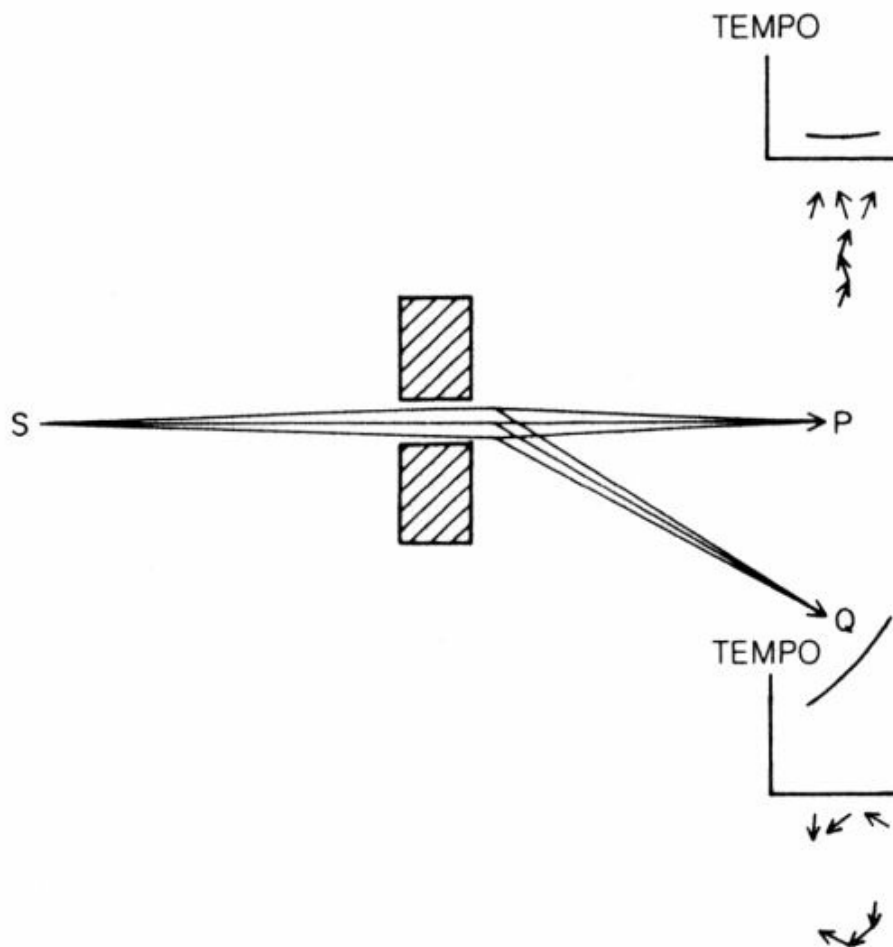


Fig. 34. Se si avvicinano i blocchi in modo che solo pochi percorsi diventino possibili, attraverso la stretta fenditura arriva in Q quasi altrettanta luce che in P, perché non ci sono abbastanza frecce che rappresentano i percorsi fino a Q per elidersi a vicenda.

Naturalmente tutt'e due le frecce risultanti saranno piccole, perché non passa molta luce, essendo la fenditura tanto stretta, ma il rivelatore in Q ticchetta quasi come quello in P! Ecco quindi che se si comprime troppo la luce, per essere sicuri che segua una linea retta, essa rifiuta di collaborare e comincia ad allargarsi.⁸

L'idea che la luce si propaghi in linea retta è dunque solo una comoda approssimazione per descrivere il suo comportamento in situazioni usuali, analoga a quella che si fa dicendo che, quando la luce si riflette su uno specchio, l'angolo di riflessione è uguale a quello di incidenza.

Ho mostrato prima un abile trucco con cui far vedere che la luce si riflette sotto molti angoli da uno specchio. Usando un sistema assai simile, possiamo costringere la luce ad andare da un punto a un altro seguendo parecchi percorsi.

Per semplificare, tratterò prima una linea tratteggiata verticale (fig. 35) tra la sorgente e il rivelatore (linea che non ha alcun significato, è solo artificiale!) e considererò soltanto i percorsi costituiti da due tratti rettilinei. Il tempo relativo a ciascun percorso lo riporto su un grafico simile a quello del caso dello specchio (solo che lo disegno rotato di 90 gradi): la linea dei tempi parte da A, all'estremità superiore, curva verso valori più piccoli, perché i percorsi nella zona centrale sono più corti e richiedono meno tempo, e infine si allontana nuovamente.

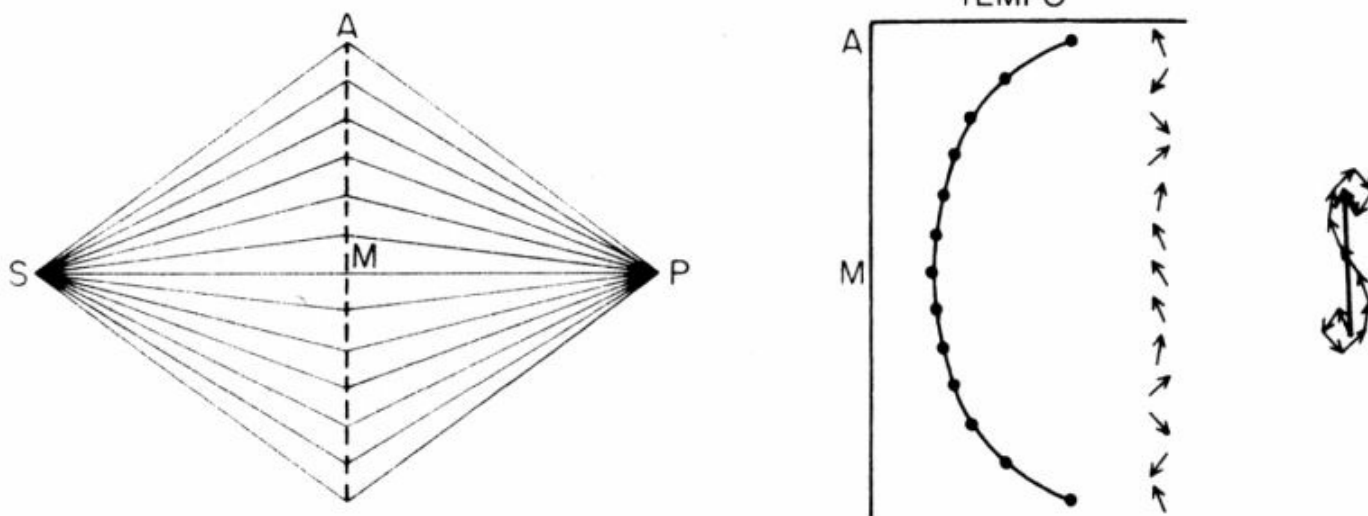


Fig. 35. L'analisi di tutti i possibili percorsi da S a P è più semplice se si considerano solo i percorsi fatti da due tratti rettilinei (in un solo piano). Il risultato è lo stesso che nel caso reale, più complicato: c'è sempre una curva dei tempi con un minimo, da cui proviene il contributo maggiore alla freccia finale.

Adesso divertiamoci un po' e proviamo a 'imbrogliare la luce', facendo in modo che *tutti* i percorsi richiedano lo stesso tempo. Ma come? Come si può ottenere che il percorso più corto, che passa per M, richieda lo stesso tempo del percorso più lungo che passa per A?

Si ricordi che la luce si muove più lentamente nell'acqua che nell'aria; altrettanto succede nel vetro (che però è molto più facile da lavorare!). Quindi, basterà inserire uno spessore adatto di vetro lungo il percorso più breve, che passa per M, e il tempo richiesto per tale percorso sarà esattamente uguale a quello necessario a passare per A. I percorsi vicini a M, che sono un pochino più lunghi, richiederanno un vetro un po' meno spesso (fig. 36).

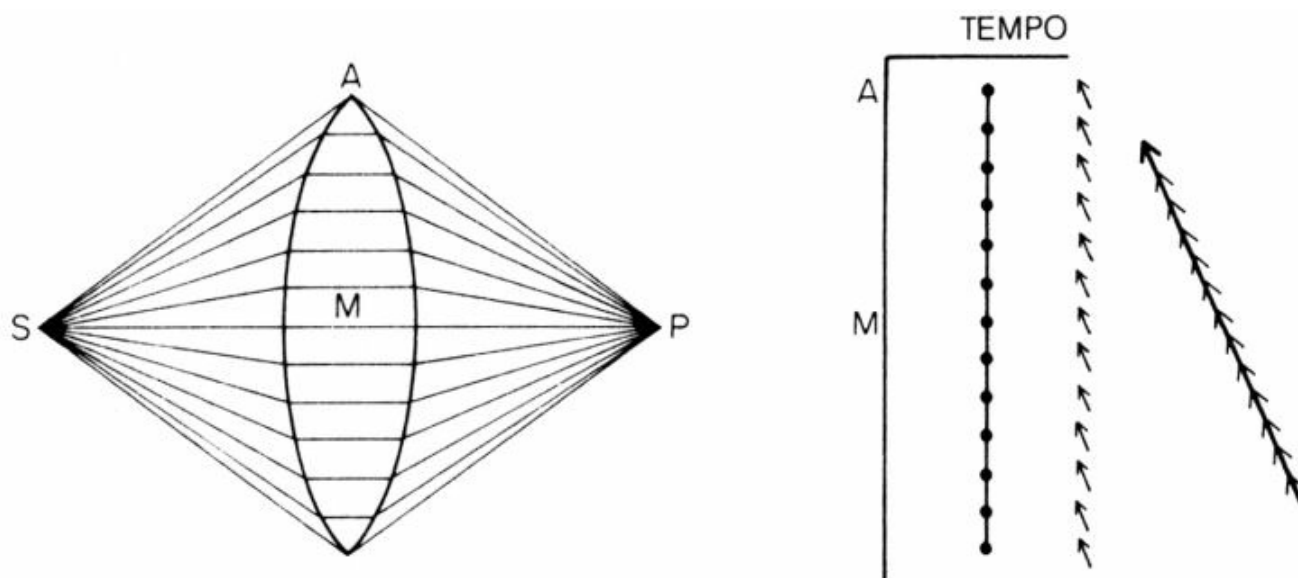


Fig. 36. Si può «imbrogliare» la Natura rallentando la luce che segue i percorsi più brevi: inserendo uno spessore adatto di vetro si fa in modo che tutti i percorsi richiedano lo stesso tempo. Allora tutte le frecce hanno la stessa direzione e si ottiene una freccia risultante favolosa — un'enorme quantità di luce! Un siffatto pezzo di vetro costruito in modo da aumentare la probabilità che la luce arrivi dalla sorgente a un dato punto viene detto lente convergente.

Più ci si avvicina ad A, meno vetro occorre inserire per rallentare la luce. Interponendo lungo ciascun percorso lo spessore di vetro appropriato a compensare il ritardo, si possono rendere uguali tra loro tutti i tempi di percorrenza. Disegnando le frecce per i

vari possibili percorsi della luce, si trova adesso che si è riusciti ad allinearle tutte (e pensare che sono milioni!), per cui il risultato finale è spettacolare: una freccia risultante di lunghezza sorprendente! Naturalmente avrete capito che sto descrivendo una lente convergente. Sistemando le cose in modo che tutti i tempi siano uguali, si può focalizzare la luce, cioè si può rendere molto elevata la probabilità che essa arrivi in un punto determinato e molto bassa la probabilità che arrivi in qualsiasi altro punto.

Questi esempi mostrano come l'elettrodinamica quantistica, che a prima vista sembra una teoria assurda, priva di causalità, senza un funzionamento chiaro, senza alcuna concretezza, riproduce fenomeni ben noti, quali la riflessione su uno specchio, la curvatura dei raggi luminosi nel passaggio dall'aria all'acqua, la focalizzazione delle lenti. Essa prevede anche altri fenomeni forse non altrettanto familiari, come il funzionamento del reticolo di diffrazione, e moltissimi altri. In breve, questa teoria descrive con successo *qualunque* fenomeno in cui interviene la luce.

Gli esempi discussi illustrano come calcolare la probabilità di un evento che può verificarsi in *più modi alternativi* tra loro: si disegna una freccia per ciascun modo e si sommano le frecce. «Sommare le frecce» vuol dire sovrapporre una coda alla punta precedente e tracciare poi la «freccia risultante». Il quadrato di questa freccia finale rappresenta la probabilità dell'evento.

Per farvi sentire meglio il «gusto» quantistico di questa teoria, vi farò vedere ora come i fisici calcolano la probabilità di un evento composto, cioè di un evento che può essere diviso in una serie di passi o che consiste di vari fatti che accadono indipendentemente.

Un esempio di evento composto lo si ottiene modificando il primo esperimento discusso, in cui alcuni fotoni rossi venivano inviati su una superficie di vetro per misurare la riflessione parziale. Al posto del fotomoltiplicatore (fig. 37) mettiamo uno schermo con un foro che lasci passare i fotoni che giungono in A. Mettiamo poi uno strato di vetro in B e piazziamo il fotomoltiplicatore in C. Come si calcola la probabilità che un fotone partito dalla sorgente arrivi in C?

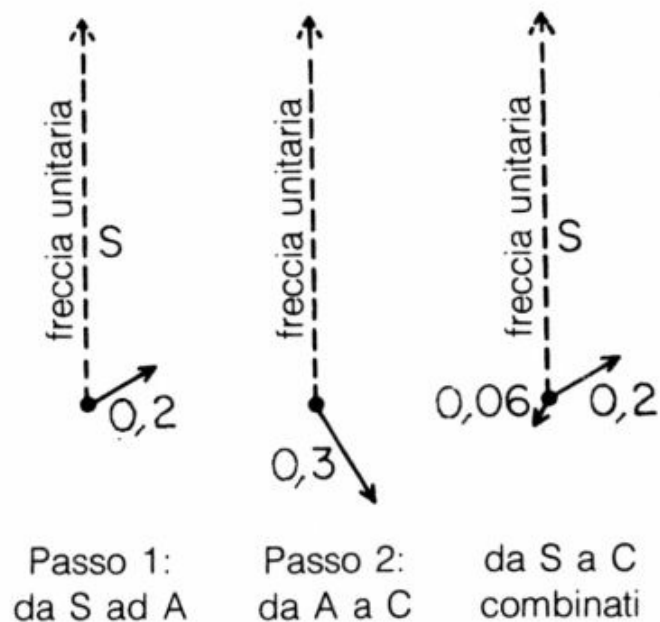
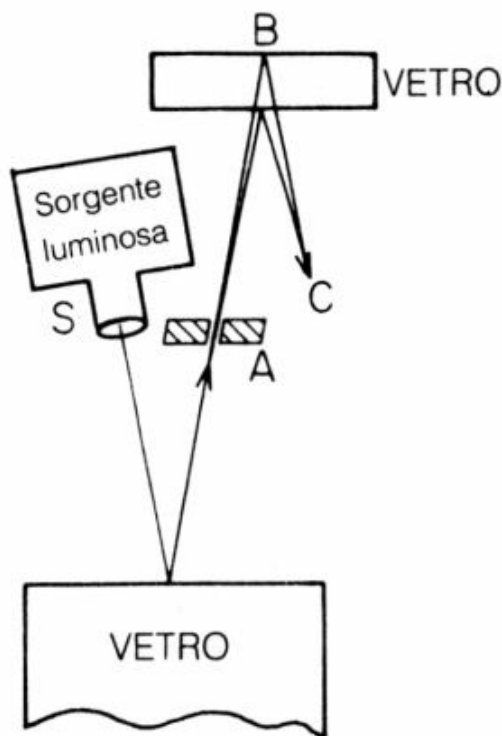


Fig. 37. Un evento composto può essere analizzato come una successione di passi. In questo esempio il percorso seguito da un fotone per andare da S a C può venire diviso in due passi: 1) il fotone va da S ad A; 2) il fotone va da A a C. Ogni passo, analizzato separatamente, dà luogo a una freccia che può essere considerata da un punto di vista nuovo: come una freccia unitaria (cioè di lunghezza 1 e che segna le 12) che ha subito una contrazione e una rotazione. In questo esempio il passo 1 comporta una contrazione a 0,2 e una rotazione fino alle 2; il passo 2 una contrazione a 0,3 e una rotazione fino alle 5. L'ampiezza per due passi consecutivi si ottiene contraendo e ruotando in successione: la freccia unitaria iniziale viene ridotta e ruotata in modo da ottenere una freccia di lunghezza 0,2 che segna le 2, la quale viene a sua volta ridotta a 0,3 e ruotata dalle 12 alle 5, come se fosse la freccia unitaria; si ottiene così una freccia di lunghezza 0,06 diretta verso le 7 sul quadrante dell'orologio. Questo procedimento di contrazioni e rotazioni successive viene chiamato «moltiplicazione» delle frecce.

Questo evento può essere considerato come una successione di due passi. Primo passo: un fotone va dalla sorgente fino al punto A, riflettendosi sull'unica superficie del blocco di vetro. Secondo passo: il fotone va da A fino al fotomoltiplicatore in C, riflettendosi sullo strato di vetro in B. A ciascun passo corrisponde una freccia risultante o, come si dice in linguaggio tecnico, una «ampiezza» (d'ora in avanti userò indifferentemente l'uno o l'altro dei due termini), che può essere calcolata sulla base delle regole già note. L'ampiezza relativa al primo passo ha lunghezza 0,2 (il cui quadrato è 0,04, la probabilità di riflessione su una singola superficie di vetro) ed è orientata secondo una certa direzione, poniamo come una lancetta che segni le 2.

Per calcolare l'ampiezza relativa al secondo passo, poniamo provvisoriamente la sorgente di luce in A in modo che i fotoni vengano diretti verso lo strato di vetro B. Disegniamo le frecce per la riflessione frontale e posteriore e sommiamole, ottenendo (mettiamo) una freccia risultante di lunghezza 0,3 diretta verso le 5 sul quadrante dell'orologio.

Come si devono combinare le due frecce per calcolare l'ampiezza per l'evento complessivo? Consideriamole da un punto di vista nuovo: come istruzioni per una *contrazione* e una *rotazione*.

Nel presente esempio la prima ampiezza ha lunghezza 0,2 ed è diretta verso le 2. Se si inizia con una «freccia unitaria», cioè con una freccia di lunghezza 1 diretta verso l'alto, la si deve *accorciare* da 1 a 0,2 e *ruotare* dalle 12 alle 2. L'ampiezza relativa al secondo passo la si può pensare come una contrazione della freccia unitaria da 1 a 0,3 e una sua

rotazione dalle 12 alle 5.

Per combinare le ampiezze relative a entrambi i passi si deve procedere ad accorciare e ruotare *in successione*. Dapprima si accorcia la freccia unitaria da 1 a 0,2 e la si ruota dalle 12 alle 2, poi si accorcia la freccia così ottenuta da 0,2 a tre decimi di tale valore, e la si ruota dell'angolo compreso tra le 12 e le 5, cioè dalle 2 alle 7. La freccia risultante ha lunghezza 0,06 ed è diretta verso le 7. Essa rappresenta una probabilità data da 0,06 al quadrato, vale a dire 0,0036.

Osservando attentamente questo procedimento, si vede che il risultato di accorciare e girare in successione due frecce è lo stesso che sommare i rispettivi angoli (quello tra le 12 e le 2 + quello tra le 12 e le 5) e moltiplicare tra loro le lunghezze ($0,2 \times 0,3$). È facile capire perché si sommino gli angoli: essi corrispondono alla rotazione della lancetta di un cronometro immaginario, e la rotazione per i due passi consecutivi è ovviamente la somma dell'angolo di cui la lancetta ha girato nel primo passo e di quello di cui ha girato nel secondo.

Soffermiamoci sul perché questo procedimento venga chiamato «moltiplicazione delle frecce»: è una spiegazione interessante. Proviamo per un momento a considerare la moltiplicazione dal punto di vista dei Greci (questa divagazione non ha nulla a che fare con l'argomento della lezione). Per poter usare numeri che non fossero interi i Greci li rappresentavano con dei tratti di linea. Ogni numero può essere espresso *trasformando* un tratto unitario mediante dilatazioni o contrazioni. Ad esempio, se la linea A rappresenta l'unità (fig. 38), la linea B rappresenta 2 e la linea C rappresenta 3.

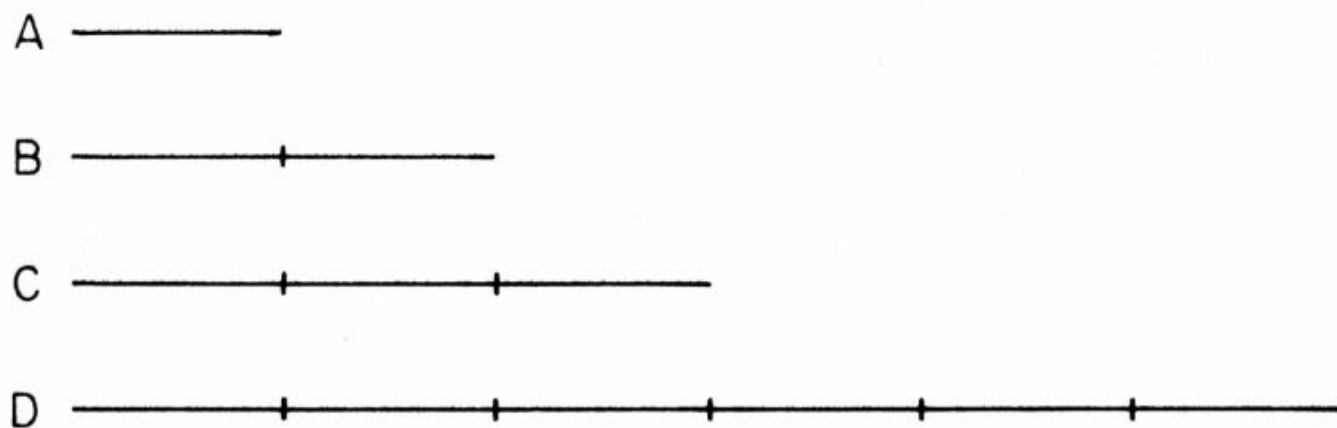


Fig. 38. Qualunque numero può essere ottenuto come trasformazione di un tratto unitario mediante contrazioni o dilatazioni. Se il tratto A rappresenta l'unità, B rappresenta 2 (dilatazione) e C rappresenta 3 (dilatazione). La moltiplicazione si ottiene tramite trasformazioni successive. Ad esempio, moltiplicare 3 per 2 significa dilatare la linea unitaria dapprima 3 volte e poi altre 2, ottenendo come risultato una dilatazione di 6 volte, rappresentata dalla linea D. Se si considera D come tratto unitario, il tratto C rappresenta $1/2$ (contrazione) e il tratto B rappresenta $1/3$ (contrazione). Moltiplicare $1/2$ per $1/3$ significa contrarre il tratto D prima a $1/2$ e poi a $1/3$ di tale risultato, ottenendo come risposta una contrazione a $1/6$, rappresentata dal tratto A.

Come si fa a moltiplicare 3 per 2? Basta applicare le trasformazioni *in successione*: partendo dalla linea unitaria A, la si espande 2 volte e poi 3 volte (o 3 volte e poi 2 volte, l'ordine non fa alcuna differenza). Il risultato è la linea D, la cui lunghezza rappresenta 6. E per moltiplicare $1/3$ per $1/2$? Partendo da D come linea unitaria, la si accorcia prima a $1/2$, ottenendo C, e poi a $1/3$ di questa. Il risultato è la linea A, che rappresenta $1/6$.

La moltiplicazione delle frecce funziona nello stesso modo (fig. 39). Si applicano trasformazioni in successione alla freccia unitaria, solo che le trasformazioni delle frecce

coinvolgono due operazioni: contrazione e rotazione.

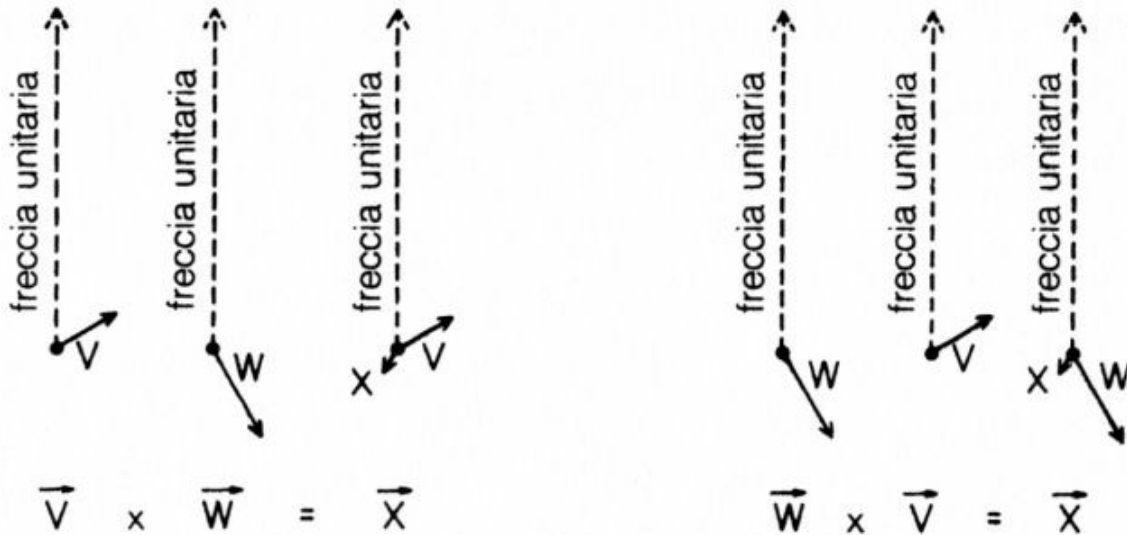


Fig. 39. I matematici hanno osservato che la moltiplicazione tra frecce può essere espressa come successione di trasformazioni di una freccia unitaria, nel nostro caso mediante contrazioni e rotazioni. Come nella comune moltiplicazione, l'ordine non è importante: si ottiene la stessa risposta, freccia X, moltiplicando la freccia V per W o per V.

Per moltiplicare la freccia V per la freccia W, si contrae e si ruota la freccia unitaria dell'ammontare prescritto per V, poi si contrae e si ruota il risultato dell'ammontare prescritto per W; anche in questo caso l'ordine non fa differenza. Dunque la moltiplicazione delle frecce segue la stessa regola di trasformazione sequenziale che si applica ai soliti numeri.⁹

Torniamo al primo esperimento della prima lezione, la riflessione parziale su una singola superficie, avendo in mente questa idea dei passi successivi (fig. 40).

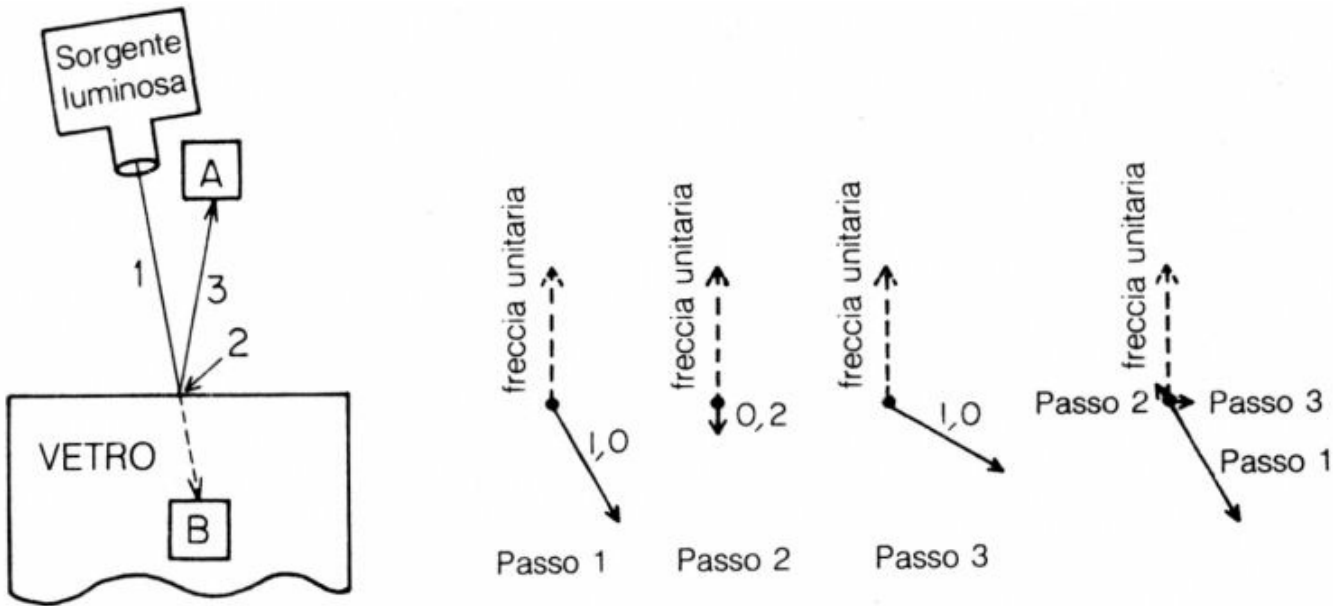


Fig. 40. La riflessione da una superficie singola può essere suddivisa in tre passi, ciascuno dei quali comporta una contrazione e/o una rotazione della freccia unitaria. Il risultato finale è lo stesso di prima: una freccia di lunghezza 0,2 che punta in una certa direzione, ma il nostro metodo di analisi risulta ora più dettagliato.

Possiamo dividere il percorso della luce in tre passi: 1) la luce va dalla sorgente al blocco di vetro, 2) è riflessa dal vetro, 3) va dal vetro al rivelatore. Ogni passo corrisponde a una determinata contrazione e rotazione della freccia unitaria.

Come ricorderete, però, in quella lezione non avevamo preso in considerazione *tutti* i

percorsi che la luce può seguire nel riflettersi sul vetro. Ciò avrebbe richiesto di disegnare e sommare un'enorme quantità di piccole frecce e per evitare troppi dettagli ho dato l'impressione che la luce incida solo in un punto particolare, ossia che non si allarghi. In realtà, nel propagarsi da un punto a un altro, la luce si allarga (a meno che non si usi una lente) e ciò comporta una certa contrazione della freccia unitaria iniziale. Per il momento, tuttavia, manterrò il punto di vista semplificato secondo il quale la luce *non* si allarga, ed è quindi giustificato trascurare tale contrazione. È anche giustificato ammettere che, non essendovi allargamento, ogni fotone che parte dalla sorgente termina in A o in B.

Dunque, nel primo passo non c'è contrazione ma c'è rotazione, corrispondente all'angolo di cui ruota la lancetta del cronometro immaginario mentre il fotone viaggia dalla sorgente alla superficie frontale del vetro. Dopo il primo passo si ha pertanto una freccia di lunghezza 1 diretta, ad esempio, verso le 5.

Il secondo passo consiste nella riflessione del fotone da parte del vetro. In questo caso c'è una notevole contrazione, da 1 a 0,2 e una rotazione di mezzo giro. (Questi numeri sembrano per ora del tutto arbitrari: essi dipendono dal materiale su cui si riflette la luce, e nella terza lezione si vedrà da dove vengono). Il secondo passo è rappresentato da una freccia di lunghezza 0,2 che segna le 6 (mezzo giro).

Nel terzo passo il fotone viaggia dal vetro al rivelatore. Come nel primo passo, non c'è contrazione ma solo rotazione; supponiamo che la distanza percorsa sia leggermente più breve che nel primo passo e che la freccia si porti sulle 4.

Ora «moltiplichiamo» in successione le frecce 1, 2 e 3 (cioè sommiamo gli angoli e moltiplichiamo le lunghezze). Il risultato finale di questi tre passi, ossia 1) rotazione, 2) contrazione e rotazione di mezzo giro, 3) rotazione, è lo stesso di quello trovato nella prima lezione: la rotazione dovuta ai passi 1 e 3 (dalle 12 alle 5 più dalle 12 alle 4) è la stessa di quella trovata lasciando che il cronometro seguisse il fotone per l'intera distanza (dalle 12 alle 9); il mezzo giro in più del passo 2 fa sì che la freccia punti nel verso opposto alla lancetta del cronometro, come accadeva nella prima lezione; infine la contrazione a 0,2 nel secondo passo produce una freccia il cui quadrato rappresenta una probabilità del 4% di riflessione parziale su una superficie singola.

In questo esperimento c'è un problema su cui non ci eravamo soffermati durante la prima lezione: che succede ai fotoni che vanno in B, cioè che attraversano la superficie del vetro? L'ampiezza per questo processo deve avere lunghezza circa 0,98, poiché $0,98 \times 0,98 = 0,9604$, che è ragionevolmente vicino al 96%. Anche questa ampiezza può essere analizzata riducendo il percorso in vari passi (fig. 41).

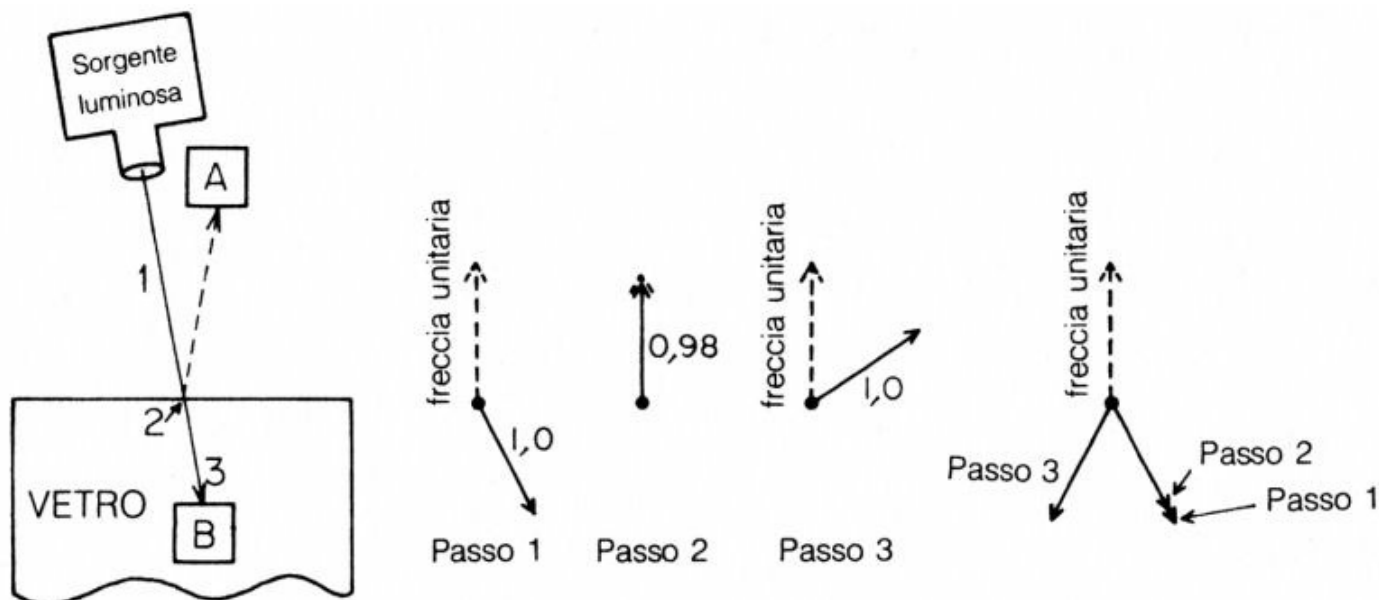


Fig. 41. Anche la trasmissione attraverso una singola superficie può essere suddivisa in tre passi, ciascuno dei quali comporta una contrazione e/o una rotazione. Il quadrato di una freccia di lunghezza 0,98 vale circa 0,96 e rappresenta una probabilità di trasmissione del 96% (che combinata con la probabilità di riflessione del 4% dà conto del 100% della luce).

Il primo di questi è uguale a quello del percorso che finisce in A: il fotone va dalla sorgente al vetro, e la freccia unitaria viene ruotata verso le 5.

Nel secondo passo il fotone attraversa la superficie del vetro: non vi è rotazione associata con tale passo, ma una piccola contrazione a 0,98.

Il terzo passo, la propagazione del fotone all'interno del vetro, comporta una ulteriore rotazione, ma nessuna contrazione.

Il risultato finale consiste in una freccia di lunghezza 0,98, con una direzione imprecisata, il cui quadrato, pari a 0,96, dà la possibilità che il fotone arrivi in B.

Consideriamo adesso nuovamente la riflessione parziale da due superfici. La riflessione sulla superficie frontale è la stessa che per una superficie singola, per cui i tre passi coincidono con quelli appena visti (fig. 40).

La riflessione sulla superficie posteriore può invece essere suddivisa in sette passi (fig. 42).

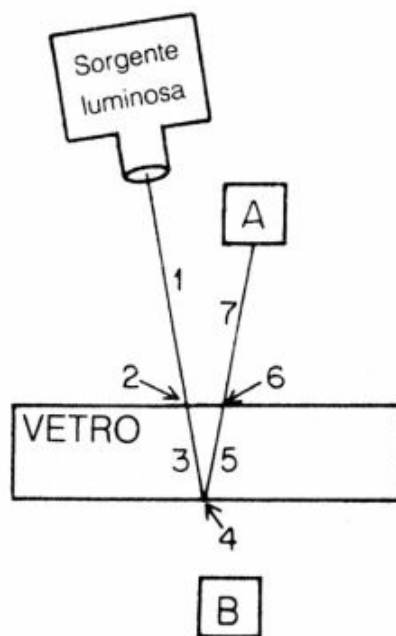


Fig. 42. La riflessione dalla superficie posteriore di una lamina di vetro può essere suddivisa in sette passi. I passi 1, 3, 5 e 7 consistono soltanto in una rotazione; i passi 2 e 6 consistono in una contrazione a 0,98 mentre il 4 consiste in una contrazione a 0,2. Come risultato si ha una freccia di lunghezza 0,192 (approssimata a 0,2 nella prima lezione) ruotata di un angolo corrispondente alla rotazione complessiva della lancetta del cronometro immaginario.

Essi comportano una rotazione pari alla rotazione totale della lancetta del cronometro che segue il fotone lungo l'intero percorso (passi 1, 3, 5 e 7), una contrazione a 0,2 (passo 4) e due riduzioni a 0,98 (passi 2 e 6). La freccia risultante avrà direzione corrispondente alla rotazione totale, come prima, e lunghezza circa 0,192 ($= 0,98 \times 0,2 \times 0,98$), arrotondata a 0,2 nella prima lezione.

Riassumendo, le regole che descrivono la riflessione e la trasmissione della luce da parte del vetro sono: 1) la riflessione dall'aria all'aria (sulla superficie frontale) comporta una contrazione a 0,2 e mezzo giro di rotazione; 2) la riflessione dal vetro nel vetro (sulla superficie posteriore) comporta anch'essa una contrazione a 0,2 ma nessuna rotazione; 3) il passaggio dall'aria al vetro o viceversa comporta una contrazione a 0,98 ma nessuna rotazione.

A rischio di farvi fare indigestione, non so resistere alla tentazione di discutere un altro bellissimo esempio di come si procede all'analisi di un fenomeno mediante tale suddivisione in passi successivi. Spostiamo il rivelatore sotto la lamina di vetro, e consideriamo un fatto non discusso nella prima lezione: la probabilità di *trasmissione* attraverso due superfici di vetro (fig. 43).

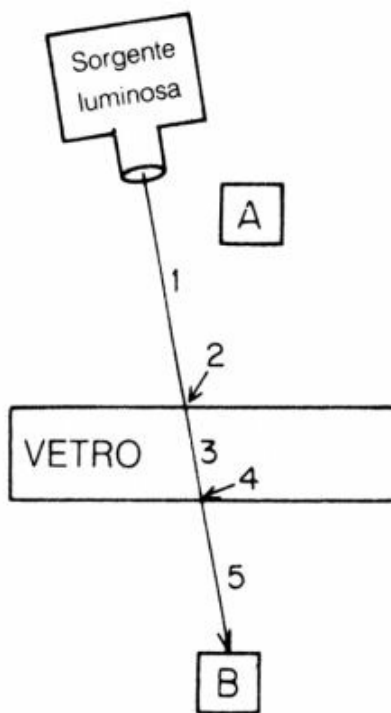


Fig. 43. La trasmissione attraverso due superfici può essere suddivisa in cinque passi. Il passo 2 provoca una contrazione a 0,98 della freccia unitaria, mentre il 4 fa contrarre al 98% la freccia lunga 0,98, ottenendo 0,96; i passi 1, 3 e 5 comportano solo rotazioni. Il quadrato della freccia risultante, lunga 0,96, vale 0,92 e rappresenta una probabilità di trasmissione attraverso due superfici del 92%, corrispondente all'8% di probabilità di riflessione che risulta corretto «due volte al giorno». Allorché lo spessore della lamina è tale da dar luogo a una probabilità di riflessione del 16%, col 92% come probabilità di trasmissione si ottiene il 108% della luce! Deve esserci qualcosa di sbagliato in questa analisi!

Naturalmente la risposta la sapete già: la probabilità che un fotone ha di arrivare in B è semplicemente il 100%, meno quella di arrivare in A, calcolata in precedenza. Così, se si fosse trovata una probabilità del 7% di arrivare in A, quella di arrivare in B sarebbe del 93%. Poiché la probabilità di arrivare in A varia da zero al 16%, a seconda dello spessore

del vetro, quella di arrivare in B deve variare dal 100% all'84%.

Questa, certamente, è la risposta corretta, ma deve essere possibile calcolare *qualunque* probabilità come quadrato di una freccia risultante. Come si calcola la freccia della probabilità per la trasmissione attraverso uno strato di vetro? Come si accorda la sua lunghezza con quella della freccia per la riflessione verso A, così che la probabilità di finire in A più quella di finire in B sommino in ogni caso esattamente al 100%? Discutiamo la cosa un po' più in dettaglio.

Il percorso di un fotone che va dalla sorgente al rivelatore sotto il blocco di vetro, in B, può essere suddiviso in cinque passi. Vediamo come ruota e si contrae la freccia unitaria iniziale mentre il fotone procede.

I primi tre passi sono identici a quelli dell'esempio precedente: il fotone va dalla sorgente al vetro (rotazione ma non contrazione); attraversa la superficie frontale (nessuna rotazione e una contrazione al 98%); si propaga all'interno del vetro (rotazione ma non contrazione).

Il quarto passo, cioè l'attraversamento della superficie posteriore, è analogo al secondo: niente rotazione e ulteriore riduzione al 98%. A questo punto la freccia ha lunghezza 0,96.

Infine il fotone viaggia nuovamente nell'aria fino al rivelatore, e ciò comporta ancora rotazione ma nessuna contrazione. Il risultato finale è una freccia di lunghezza 0,96 diretta nella direzione determinata dalle varie rotazioni della lancetta del cronometro.

Una freccia lunga 0,96 rappresenta una probabilità di circa il 92% (0,96 al quadrato), il che significa una media di 92 fotoni che giungono in B su 100 che partono dalla sorgente. Ciò significa anche che l'8% dei fotoni sono riflessi dalle due superfici e arrivano in A. Ma nella prima lezione si era trovato che la riflessione dell'8% da due superfici si verifica solo in certi casi («due volte al giorno») e che in realtà essa oscilla ciclicamente da zero al 16%, al variare dello spessore della lamina di vetro. Cosa succede quando questo spessore è tale da produrre una riflessione del 16%? Su 100 fotoni emessi dalla sorgente 16 arrivano in A e 92 in B, sicché ci si ritrova col 108% della luce emessa — orrore! È chiaro che c'è qualcosa che non va.

Abbiamo dimenticato di tener conto di *tutti* i percorsi con cui la luce può giungere in B! Ad esempio, può rimbalzare sulla superficie posteriore, ripercorrere il vetro come per andare in A, ma poi riflettersi sulla superficie anteriore e tornare indietro verso B (fig. 44). Questo percorso richiede nove passi. Vediamo cosa accade alla freccia unitaria via via che la luce li compie (niente paura: si tratta solo di contrazioni e rotazioni!).

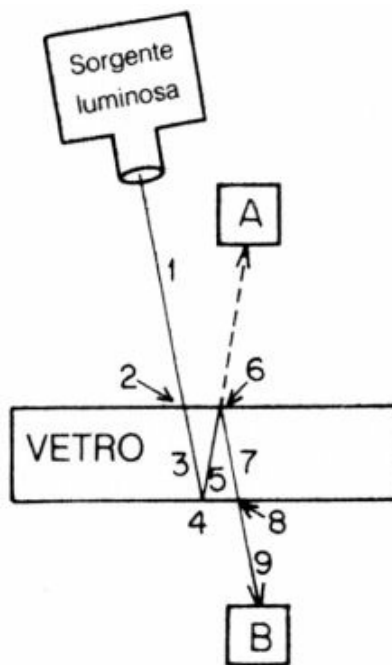


Fig. 44. Per rendere i calcoli più accurati occorre tener conto di un secondo modo in cui la luce può attraversare una lamina con due superfici. Il percorso corrispondente comporta due contrazioni a 0,98 (passi 2 e 8) e due contrazioni a 0,2 (passi 4 e 6) e dà luogo a una freccia di lunghezza 0,0384 (arrotondata a 0,04).

Primo passo: il fotone viaggia nell'aria, si ha rotazione ma non contrazione. Secondo passo: il fotone attraversa la superficie del vetro, niente rotazione ma contrazione al 98%. Terzo: il fotone viaggia nel vetro, rotazione ma non contrazione. Quarto: riflessione sulla superficie posteriore, niente rotazione ma la freccia si contrae allo 0,2 di 0,98, cioè a 0,196. Quinto: il fotone torna indietro nel vetro, solo rotazione. Sesto: il fotone rimbalza sulla superficie frontale, che si comporta come superficie «posteriore», poiché il fotone arriva *dall'interno* del vetro, per cui non si ha rotazione ma una contrazione allo 0,2 di 0,196 cioè a 0,0392. Settimo: il fotone ripercorre il vetro, ancora rotazione. Ottavo: il fotone attraversa la superficie posteriore, niente rotazione ma ulteriore contrazione al 98% di 0,0392 cioè a 0,0384. Infine, nono passo: il fotone si propaga nell'aria e arriva al rivelatore, ancora rotazione ma nessuna contrazione.

Il risultato di tutte queste contrazioni e rotazioni è una freccia di lunghezza 0,0384 (che si può porre uguale a 0,04) orientata nella direzione dell'immaginaria lancetta di cronometro che ha seguito il fotone in questo lungo viaggio. Questa freccia corrisponde a un *secondo* modo in cui la luce può arrivare dalla sorgente al rivelatore. Avendo due percorsi, per trovare la freccia finale si devono ora *sommare* le ampiezze corrispondenti al percorso più diretto, la cui lunghezza è 0,96, e a quello più lungo, che ha lunghezza 0,04.

Le due frecce di solito non hanno la stessa direzione, e il cambiare lo spessore del vetro modifica la loro direzione relativa. Ma si osservi come tutto torna: la rotazione della lancetta che segue il fotone durante i passi 3 e 5 del percorso verso A è esattamente uguale a quella corrispondente ai passi 5 e 7 del percorso verso B. Ciò comporta che quando le due frecce relative alla riflessione si cancellano, dando luogo a una freccia finale che rappresenta riflessione nulla, le frecce relative alla trasmissione si rinforzano dando una freccia finale di lunghezza $0,96 + 0,04 = 1$; quando la probabilità di riflessione è nulla, quella di trasmissione è del 100% (fig. 45).

Viceversa, quando le frecce per la riflessione si rinforzano dando un'ampiezza di 0,4, quelle della trasmissione sono dirette in versi opposti, dando una freccia di lunghezza

$0,96 - 0,04 = 0,92$; calcolando le probabilità (cioè i quadrati), si ottiene 16% per la riflessione e 84% per la trasmissione. E così la Natura, con le sue regole intelligenti, dà sempre ragione del 100% dei fotoni!¹⁰

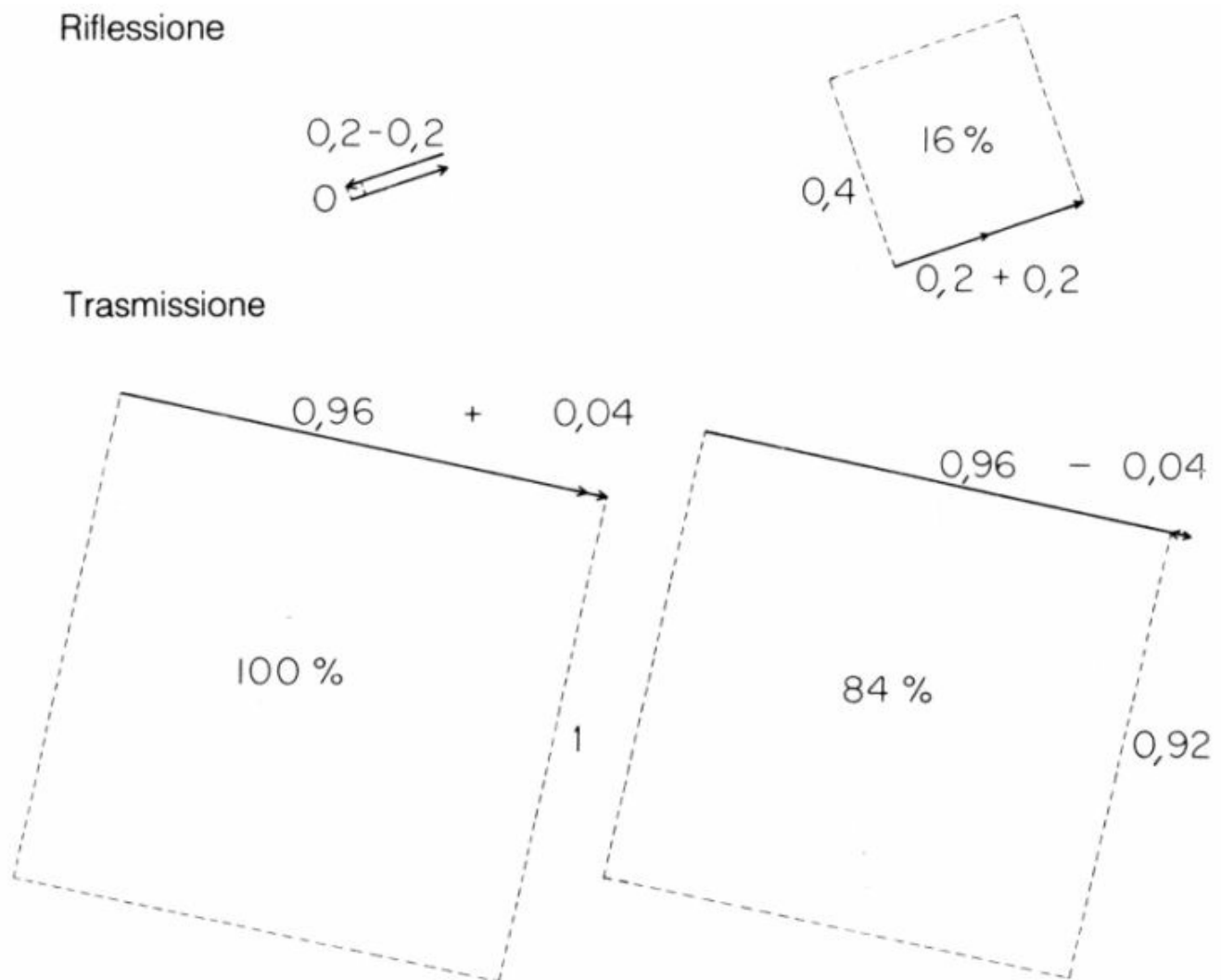


Fig. 45. La Natura fa sempre in modo di dar conto del 100% della luce. Quando lo spessore è tale che le frecce relative alla trasmissione si addizionano, quelle relative alla riflessione si elidono; se le frecce relative alla riflessione si addizionano, quelle della trasmissione sono opposte tra loro.

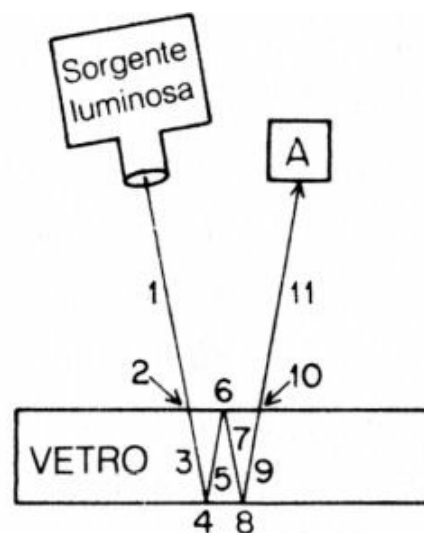


Fig. 46. Per fare un conto ancora più accurato bisognerebbe prendere in considerazione ancora altri percorsi che la luce può seguire per riflettersi. In quello mostrato in figura si ha una contrazione a 0,98 nei passi 2 e 10, e a 0,2 nei passi 4, 6 e 8. Come risultato si ha una freccia lunga circa 0,008 relativa a un possibile percorso che produce riflessione, e che va quindi sommata alle altre frecce relative alla

riflessione (quella «frontale», lunga 0,2, e quella «posteriore», lunga 0,192).

Prima di terminare questa lezione, voglio far osservare che la regola sulla moltiplicazione delle frecce deve essere generalizzata; si devono moltiplicare tra loro le frecce non solo per i fenomeni che consistono di vari passi, ma anche per i fenomeni formati da più eventi che avvengono in modo indipendente e contemporaneo. Supponiamo, ad esempio, di avere in X e in Y due sorgenti e in A e in B due rivelatori (fig. 47), e di voler calcolare la probabilità del seguente fenomeno: X e Y emettono entrambe un fotone e A e B ne acquistano entrambi uno.

I fotoni viaggiano liberamente fino ai rivelatori, senza riflessioni o trasmissioni da un mezzo a un altro; perciò questo è il momento adatto per prendere finalmente in considerazione il fatto che la luce si allarga mentre si propaga. Discuterò pertanto la *regola completa* che descrive la propagazione della luce monocromatica nello spazio vuoto, senza più approssimazioni o semplificazioni. A parte la polarizzazione, questa regola è tutto ciò che occorre sapere sulla luce monocromatica che viaggia nello spazio vuoto. Essa afferma che la *direzione* della freccia è quella dell'immaginaria lancetta di cronometro che ruota un certo numero di volte per centimetro percorso (a seconda del colore del fotone), mentre la sua *lunghezza* è inversamente proporzionale alla distanza percorsa dalla luce, in altre parole la freccia si accorcia man mano che la luce si propaga.¹¹

Supponiamo che la freccia per andare da X ad A abbia lunghezza 0,5 e segni le 5, e che lo stesso valga per la freccia per andare da Y a B (fig. 47). Moltiplicandole tra loro si ottiene una freccia risultante di lunghezza 0,25 orientata verso le 10.

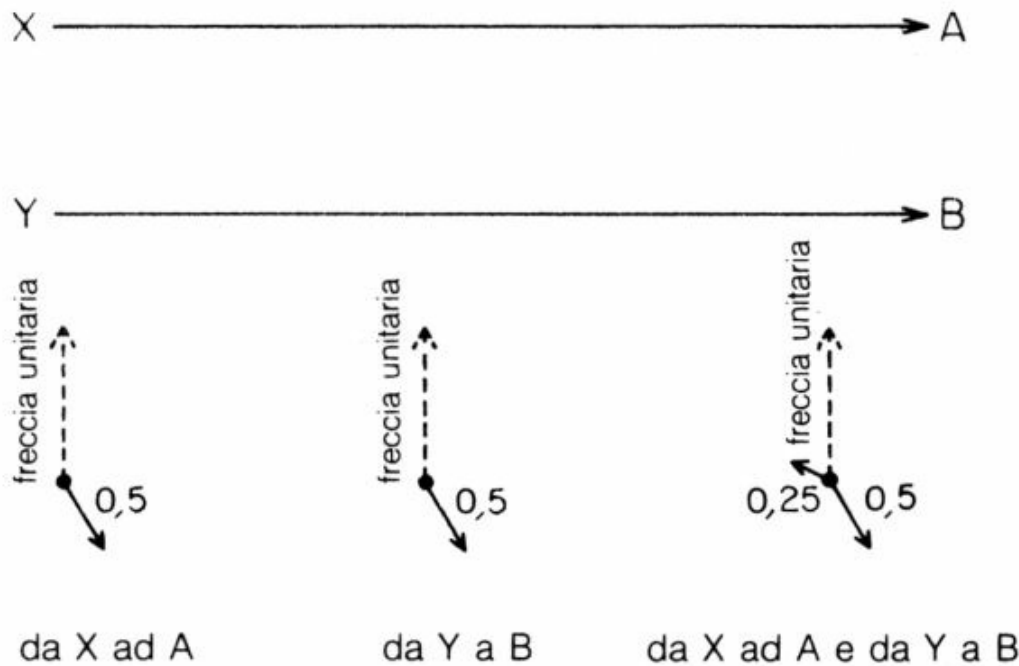


Fig. 47. Se un determinato fenomeno può verificarsi in un modo che dipende da un insieme di eventi indipendenti, l'ampiezza relativa a tale modo va calcolata moltiplicando tra loro le frecce relative ai singoli eventi indipendenti. Nella figura è illustrato il seguente fenomeno: dopo che ciascuna delle due sorgenti X e Y ha emesso un fotone, entrambi i fotomoltiplicatori in A e in B scattano. Tale fenomeno può verificarsi se un fotone va da X ad A mentre un altro va da Y a B (due eventi indipendenti). La probabilità relativa a questo «primo modo» va calcolata moltiplicando tra loro le frecce per ciascun evento indipendente (un fotone da X ad A e uno da Y a B), ottenendo così l'ampiezza relativa a questo modo specifico (l'analisi continua con la fig. 48).

Un momento! Questo fenomeno può avvenire anche in un altro modo: il fotone di X va in B, mentre quello di Y va in A. Ciascuno di questi eventi ha una propria ampiezza, e le relative frecce vanno moltiplicate tra loro per ottenere l'ampiezza corrispondente a questo

altro modo (fig. 48). Le lunghezze delle singole frecce sono molto vicine al caso precedente, essendo le distanze percorse quasi uguali, cioè le frecce per andare da X a B e da Y ad A hanno praticamente lunghezza 0,5; molto diversa da prima sarà invece la loro direzione: la lancetta del cronometro compie 14.000 giri al centimetro per la luce rossa, per cui anche una piccola differenza nel tragitto comporta una notevole differenza nella direzione della freccia.

Le ampiezze relative a questi due modi vanno ora sommate per ottenere la freccia finale. Avendo essenzialmente la stessa lunghezza, le due frecce si cancellano tra loro se hanno direzioni opposte. La loro direzione relativa può essere modificata cambiando la distanza tra le sorgenti o tra i rivelatori: allontanare o avvicinare tra loro un pochino i rivelatori può aumentare la probabilità dell'evento o viceversa renderla nulla, proprio come nel caso della riflessione parziale su due superfici.¹²

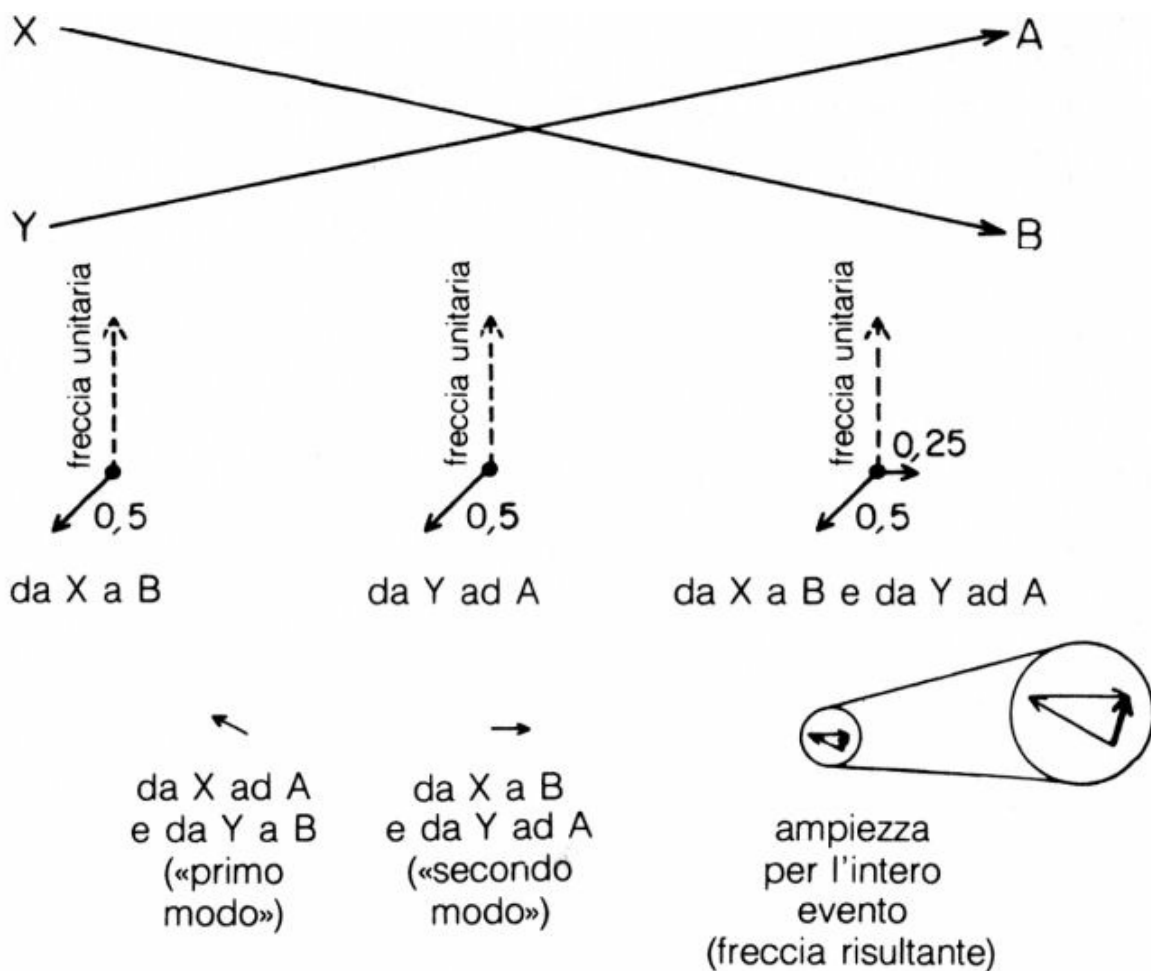


Fig. 48. L'altro modo in cui può verificarsi il fenomeno descritto alla fig. 47: il fotone di X va in B, mentre quello di Y va in A. Anche questo modo dipende dal verificarsi contemporaneo di due eventi indipendenti, per cui anche l'ampiezza relativa a questo «secondo modo» va calcolata moltiplicando tra loro le frecce relative ai due eventi indipendenti. Le frecce per il «primo» e per il «secondo» modo vanno poi sommate tra loro, per avere la freccia risultante per il fenomeno considerato. La probabilità di un evento è sempre rappresentata da una singola *freccia finale*, indipendentemente dal numero delle frecce tracciate, moltiplicate e sommate per ottenerla.

In questo esempio si sono prima moltiplicate delle frecce tra loro, poi si sono sommati i risultati per ottenere la freccia finale (l'ampiezza dell'evento) il cui quadrato dà la probabilità dell'evento.

Occorre sottolineare che qualunque sia il numero di frecce tracciate, moltiplicate o sommate, l'obiettivo finale è sempre quello di calcolare *una sola freccia risultante relativa al fenomeno*. Gli studenti di fisica fanno spesso degli errori all'inizio, perché non

tengono presente questo punto fondamentale. Passano tanto di quel tempo ad analizzare eventi in cui interviene un singolo fotone, che cominciano a pensare che la freccia sia in qualche modo associata al fotone. Ma le frecce sono solo ampiezze di probabilità, e il loro quadrato dà la *probabilità* di un evento completo.¹³

Con la prossima lezione comincerò il processo di semplificazione e spiegazione delle proprietà della materia, in modo da mostrare l'origine dei vari fenomeni discussi: la contrazione a 0,2, il fatto che la luce sembra andare più lenta nel vetro e nell'acqua che nell'aria, eccetera. Perché finora io ho barato: i fotoni in realtà non rimbalzano sulla superficie del vetro, ma interagiscono con gli elettroni *al suo interno*. Farò vedere che i fotoni non fanno altro che viaggiare da elettrone a elettrone e che riflessione e trasmissione altro non sono che il risultato di tante interazioni in cui un elettrone prende un fotone, 'si gratta la testa', per così dire, ed emette un *nuovo* fotone. Questa semplificazione di tutti i fenomeni di cui ho parlato finora è veramente molto bella.

III

GLI ELETTRONI E LE LORO INTERAZIONI

Questa è la terza di quattro lezioni su un argomento abbastanza difficile, la teoria dell'elettrodinamica quantistica. Essendo presente più gente delle volte precedenti è chiaro che alcuni non hanno sentito le altre lezioni, per cui troveranno questa quasi incomprensibile. Anche chi le ha sentite la troverà incomprensibile, ma sa già che la cosa è perfettamente normale; come ho spiegato nella prima lezione il modo di cui disponiamo per descrivere la Natura ci risulta, in generale, incomprensibile.

Lo scopo di queste lezioni è di parlarvi della parte meglio conosciuta della fisica, l'interazione della luce con gli elettroni, interazione che riguarda la maggior parte dei fenomeni usuali, per esempio tutti i fenomeni chimici e tutti quelli biologici. Gli unici a restare fuori da questa teoria sono i fenomeni nucleari e quelli gravitazionali.

Come ho detto nella prima lezione, noi non disponiamo di una spiegazione che descriva in modo soddisfacente anche solo il più semplice dei fenomeni, quale la riflessione parziale della luce sul vetro. Ci è inoltre impossibile prevedere se un dato fotone sarà riflesso o attraverserà il vetro. Sappiamo solo calcolare la *probabilità* che accada un particolare fenomeno, ad esempio che il fotone venga riflesso. (Tale probabilità risulta circa del 4% se la luce incide perpendicolarmente su un blocco di vetro, e aumenta per un'incidenza più obliqua).

In situazioni *abituali* si hanno le seguenti «regole di composizione» per le probabilità: 1) se qualcosa può verificarsi in *modi alternativi*, le varie probabilità vanno *sommate*; 2) se un fenomeno si presenta come *successione di passi*, o dipende dall'avverarsi concomitante di più eventi indipendenti, si devono *moltiplicare* tra loro le probabilità per ciascun passo o per ciascun evento.

Nel mondo misterioso e affascinante della fisica quantistica le probabilità vengono invece calcolate come *quadrati della lunghezza di una freccia*: là dove in circostanze ordinarie ci si aspetterebbe di dover sommare probabilità, ci si trova a dover 'sommare' *freccie*; dove ci si aspetterebbe di moltiplicare probabilità, si devono 'moltiplicare' *freccie*. Le risposte che si ottengono calcolando le probabilità in questo modo, per quanto strane, riproducono esattamente i risultati sperimentali. Personalmente trovo bellissimo dover ricorrere a queste strane regole e a questo strano modo di ragionare per descrivere la Natura, e mi piace parlarne agli altri. È un'analisi che non nasconde nessun congegno segreto: se si vuole capire la Natura, la strada è questa.

Prima di entrare nel vivo della lezione, vi darò un altro esempio del comportamento della luce. Prendiamo una luce monocromatica molto flebile (un fotone alla volta) che viaggia dalla sorgente in S al rivelatore in D (fig. 49).

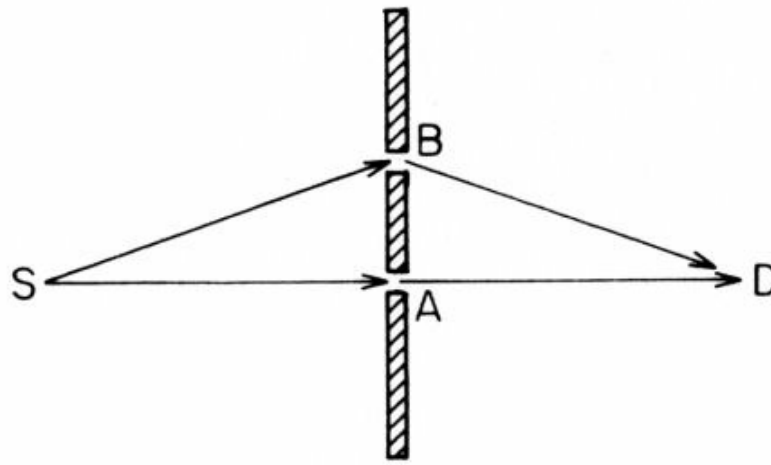


Fig. 49. In uno schermo interposto tra la sorgente S e il rivelatore D vi sono, in A e in B, due fori piccolissimi. Quando sono aperti alternativamente, attraverso ciascuno di essi giunge al rivelatore quasi la stessa quantità di luce (in questo esempio l'1%). Se i fori sono entrambi aperti, c'è «interferenza»: la probabilità che il rivelatore ticchetti varia tra zero e il 4% a seconda della separazione tra A e B (fig. 51 a).

Tra sorgente e rivelatore interponiamo uno schermo in cui vi sono due piccolissimi fori, in A e in B, distanti tra loro alcuni millimetri. (Se sorgente e rivelatore distano 100 centimetri tra loro, il diametro dei fori deve essere inferiore a un decimo di millimetro). Sia A allineato con S e con D, mentre B, trovandosi lateralmente, non è sulla stessa linea.

Quando il foro in B è chiuso, il rivelatore in D ticchetta con una certa frequenza, corrispondente al numero di fotoni che arrivano attraverso A: mettiamo che faccia in media un «clic» ogni 100 fotoni partiti da S. Come abbiamo visto nella seconda lezione, se si chiude il foro in A e si apre quello in B, il contatore ticchetta in media quasi con la stessa frequenza, dato che i fori sono molto piccoli. (Se si 'comprime' troppo la luce non risultano più valide le regole del mondo usuale, ad esempio che la luce viaggia in linea retta). Quando entrambi i fori sono aperti, otteniamo una risposta complicata perché si verifica interferenza: per certe distanze tra i due fori il numero di scatti del rivelatore è maggiore del 2% che ci si aspetta (con un massimo di circa il 4%); per distanze anche pochissimo diverse non si verifica alcuno scatto.

Verrebbe spontaneo pensare che l'apertura di un secondo foro debba *sempre* aumentare la quantità di luce che arriva nel rivelatore, ma nella realtà non accade così. Pertanto è sbagliato dire che la luce viaggia «seguendo questo percorso oppure quest'altro». Io stesso mi sorprendo talvolta a dire: «O andrà di qua o andrà di là», ma nel dire così devo tenere a mente che in realtà sto parlando di ampiezze da sommare: il fotone ha un'ampiezza relativa a un percorso e un'ampiezza relativa all'altro percorso. Se le due ampiezze sono opposte non arriva luce nel punto considerato, anche quando, come nel nostro caso, entrambi i fori sono aperti.

Ed ecco un'ulteriore complicazione nello strano comportamento della Natura. Supponiamo di mettere, in A e in B, opportuni rivelatori che indichino attraverso quale dei due fori passa il fotone quando sono entrambi aperti (è possibile costruire un rivelatore che segnali il passaggio di un fotone senza assorbirlo) (fig. 50). Poiché la probabilità che un fotone vada da S a D dipende dalla distanza tra i fori, il fotone deve per forza dividersi in due furtivamente e poi ricomporsi di nuovo. Giusto? Secondo questa ipotesi, i rivelatori in A e in B dovrebbero sempre scattare insieme (magari con intensità dimezzata), mentre il rivelatore in D dovrebbe scattare con probabilità da zero al 4%, a

seconda della distanza tra i fori A e B.

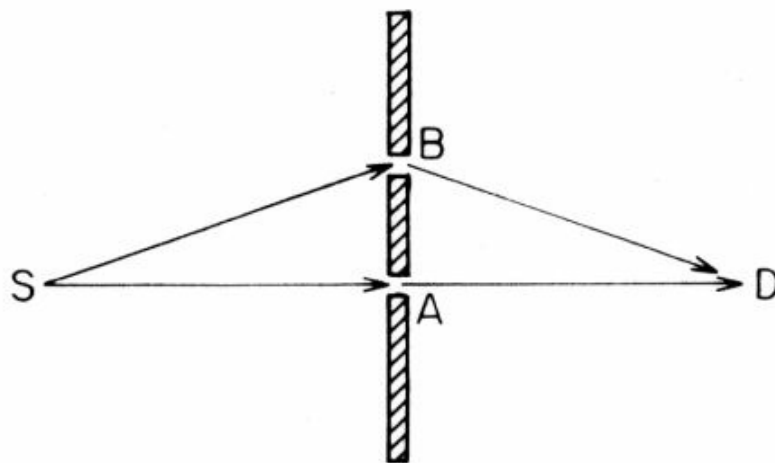


Fig. 50. Se in A e in B si inseriscono speciali rivelatori per determinare attraverso quale dei due fori passa la luce quando sono entrambi aperti, si è cambiato esperimento. Poiché quando i fori sono tenuti sotto controllo il fotone passa sempre o attraverso l'uno o attraverso l'altro, esistono ora due condizioni finali distinguibili: 1) scattano i rivelatori in A e in D, 2) scattano i rivelatori in B e in D. Ciascuno di questi due eventi si verifica con probabilità di circa l'1%. Le probabilità vanno sommate tra loro nel modo normale e si ha una probabilità complessiva del 2% che scatti il rivelatore in D (fig. 51 b).

Invece, nella realtà, i rivelatori in A e in B non scattano *mai* assieme: o scatta A o scatta B. Il fotone non si divide in due: o segue un percorso o segue l'altro.

Inoltre in presenza dei rivelatori A e B il rivelatore in D scatta esattamente il 2% delle volte, la semplice somma delle probabilità per i percorsi A e B: $1\% + 1\%$. Tale valore non è influenzato dalla distanza tra A e B; l'interferenza *scompare* se si pongono dei rivelatori in A e in B!

La Natura ha congegnato le cose così bene che non riusciremo mai a capire dove sta il trucco: inserendo gli strumenti opportuni possiamo stabilire qual è il percorso seguito dalla luce, ma i bellissimi effetti di interferenza scompaiono. Se non ci sono strumenti rivelatori, gli effetti di interferenza ritornano! Decisamente molto strano!

Per risolvere questo paradosso, vi ricorderò un principio fondamentale: per calcolare correttamente la probabilità di un evento si deve *definire con chiarezza l'evento completo*, in particolare quali siano le condizioni iniziali e finali dell'apparato. Si devono controllare gli strumenti di misura prima e dopo l'esperimento e cercare le eventuali variazioni. Quando prima calcolavamo la probabilità che un fotone vada da S a D senza rivelatori in A e in B, il fenomeno consisteva semplicemente nello scatto del rivelatore in D. Quando tale scatto era l'unico cambiamento nell'apparato di misura, non c'era modo di conoscere il percorso scelto dal fotone e si osservava interferenza.

Mettendo rivelatori in A e in B, abbiamo cambiato il problema. Adesso ci sono *due* eventi completi, cioè due insiemi di condizioni finali, distinguibili: 1) scattano i rivelatori in A e in D, 2) scattano i rivelatori in B e in D. Essendovi diverse condizioni finali possibili, si deve calcolare la probabilità di ciascuna di esse come evento separato e completo.

Per calcolare la probabilità che vi sia uno scatto nei rivelatori in A e in D, si devono moltiplicare le frecce che rappresentano i seguenti passi: un fotone va da S ad A, poi da A a D e il rivelatore in D scattali quadrato della freccia finale dà la probabilità di questo evento. Essa risulta l'1%, cioè la stessa di quando il foro in B è chiuso, perché in ambedue i casi si hanno gli stessi passi. L'altro evento completo è lo scatto dei rivelatori in B e in D.

La relativa probabilità si calcola in modo analogo e anch'essa è circa 14%, la stessa di quando A è chiuso.

Se si vuole sapere con che frequenza ticchetta il contatore in D e non interessa se nel contempo è scattato A o è scattato B, la probabilità è la semplice somma delle probabilità dei due eventi: il 2%. Anche nei casi in cui si trascura qualcosa che *avrebbe potuto* essere osservato per discriminare il percorso seguito dal fotone, siamo di fronte a «stati finali» differenti, cioè a condizioni finali distinguibili, e si devono sommare le *probabilità* relative a ciascuno di essi, non le ampiezze.¹⁴

Ho voluto mettere in evidenza queste cose perché quanto più si vede la stranezza del comportamento della Natura, tanto più si capisce perché sia difficile costruire un modello intuitivo che riesca a descrivere anche i fenomeni più semplici. Perciò la fisica teorica ha rinunciato a farlo.

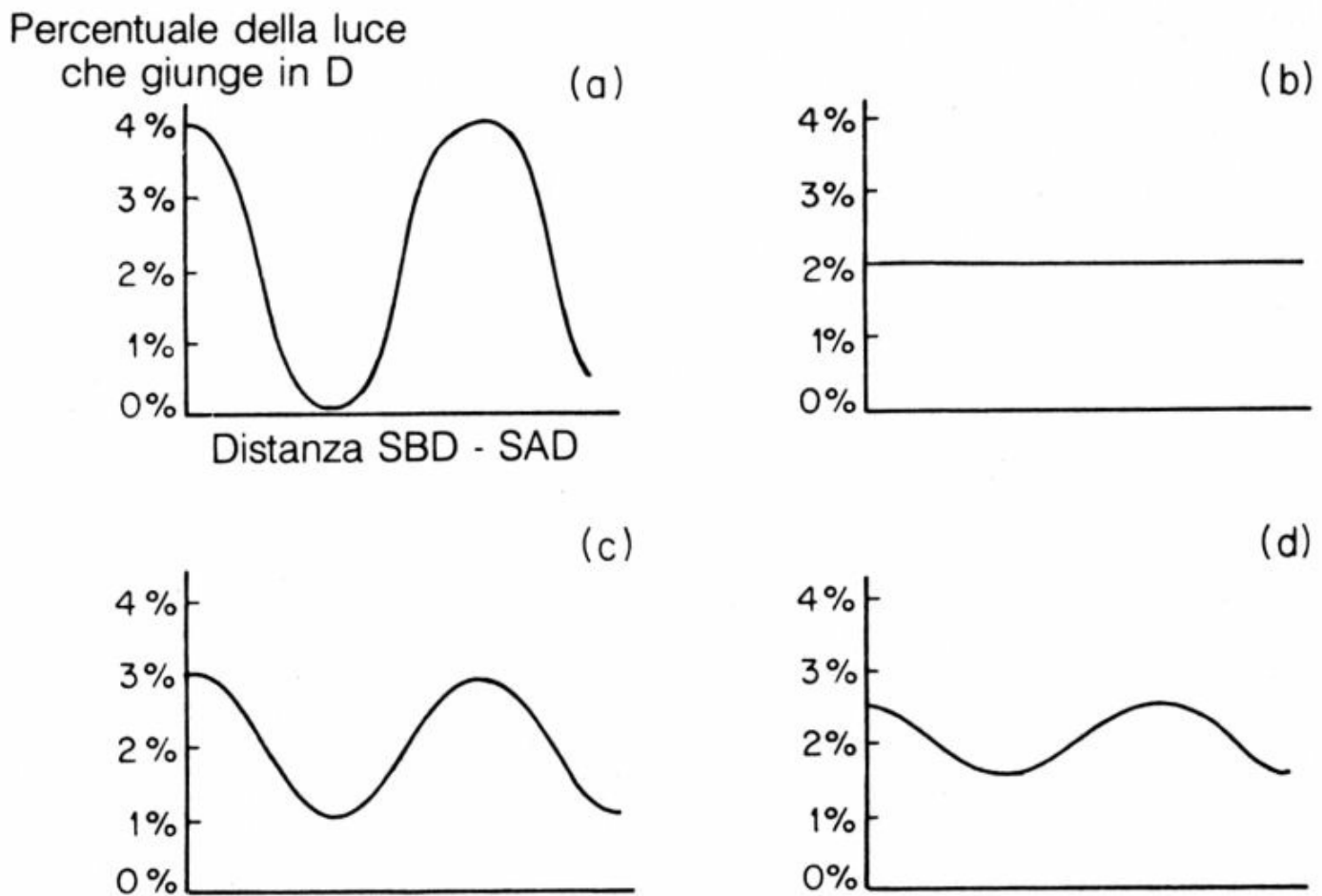


Fig. 51. Se in A e in B non vi sono rivelatori, si produce interferenza: la quantità di luce in D varia tra zero e il 4% (a). Se in A e in B vi sono rivelatori con efficienza del 100%, non si produce interferenza: la quantità di luce che giunge in D è sempre il 2% (b). Se i rivelatori in A e in B non sono completamente affidabili (cioè se qualche volta non viene compiuta nessuna misura in A o in B), vi sono tre condizioni finali distinguibili: scattano A e D, scattano B e D, scatta solamente D. La curva che si ottiene in questo caso è la semplice sovrapposizione dei contributi di ciascuna condizione finale possibile. Meno affidabili sono i rivelatori in A e in B, maggiore è l'interferenza. Ad esempio, i rivelatori nel caso (c) sono meno affidabili che nel caso (d). La quantità di interferenza è determinata dal seguente principio: si deve calcolare indipendentemente la probabilità per ciascuna delle condizioni finali differenti, sommando le relative frecce e prendendo il quadrato della freccia finale; queste probabilità vanno poi sommate tra loro nel modo usuale.

Nella prima lezione ho mostrato come un evento possa essere diviso in più modi alternativi e come vadano 'sommate' le frecce relative a ciascuno di essi. Nella seconda lezione si è poi visto come ogni modo può essere suddiviso in passi successivi, come le frecce per ciascun passo possono considerarsi trasformazioni di una freccia unitaria e come esse vadano 'moltiplicate' tra loro con contrazioni e rotazioni successive.

Conosciamo così tutte le regole necessarie per tracciare e combinare tra loro le frecce (che rappresentano i vari pezzi di un fenomeno) in modo da ottenere la freccia finale il cui quadrato dà la probabilità che esso si verifichi.

È naturale chiedersi fino a che punto si possa spingere tale procedimento di suddivisione degli eventi in sottoeventi via via più semplici. Quali sono le parti più elementari che costituiscono un evento? Esiste un numero limitato di pezzi con cui è possibile ottenere *tutti* i fenomeni che coinvolgono la luce e gli elettroni? C'è un numero limitato di 'lettere' nel linguaggio dell'elettrodinamica quantistica le cui combinazioni formano le 'parole' e le 'frasi' che costituiscono la quasi totalità dei fenomeni naturali?

La risposta è sì, e il numero è tre. Ci sono solo tre eventi elementari di base, necessari per dar luogo a tutti i fenomeni associati con la luce e gli elettroni. Prima di introdurre tali eventi devo presentarvi gli attori, cioè i fotoni e gli elettroni. Dei fotoni, particelle di luce, si è parlato a lungo nelle prime due lezioni. Gli elettroni furono scoperti nel 1895 e classificati come particelle, essendo possibile contarli o aggiungerne uno a una goccia d'olio e misurarne la carica. A poco a poco si comprese che è il movimento di queste particelle a costituire la corrente elettrica che passa nei fili conduttori.

Nei primi tempi successivi alla scoperta degli elettroni si pensò che gli atomi fossero come piccoli sistemi solari, formati da una parte centrale pesante, detta nucleo, e dagli elettroni che le girano attorno seguendo delle 'orbite', come pianeti attorno al sole. Chi pensa che gli atomi siano fatti così è fermo al 1910. Nel 1924 Louis de Broglie suggerì che agli elettroni vada associato un aspetto ondulatorio, e poco tempo dopo C.J. Davisson e L.H. Germer, dei Laboratori Bell, bombardando con elettroni un cristallo di nickel, trovarono che essi rimbalzano con angoli assai curiosi (proprio come fanno i raggi X), e che questi angoli coincidono con quelli previsti dalla formula di de Broglie per la lunghezza d'onda degli elettroni.

I fenomeni osservati quando i fotoni viaggiano su distanze grandi, molto maggiori di quella necessaria per un giro completo della lancetta del cronometro, sono descritti con ottima approssimazione da leggi quali «la luce viaggia in linea retta», perché attorno al percorso di tempo minimo ci sono abbastanza percorsi che si rinforzano tra loro, e insieme c'è un numero sufficiente di altri percorsi che si cancellano tra loro. Ma quando lo spazio a disposizione di un fotone diventa troppo esiguo (come nel caso dei piccoli fori nello schermo), queste leggi non sono più valide: si osserva allora che la luce non viaggia necessariamente in linea retta, che vi è interferenza causata dai due fori e così via. La medesima situazione si presenta con gli elettroni: su grandi distanze si propagano come particelle, lungo percorsi definiti. Ma per piccole distanze, come all'interno di un atomo, lo spazio è così esiguo che non vi è un percorso dominante, non vi è 'orbita'; l'elettrone segue tutti i possibili percorsi, ciascuno con la relativa ampiezza. L'effetto di interferenza diventa molto importante e occorre sommare le frecce per prevedere la probabilità che l'elettrone si trovi in un dato punto.

È interessante osservare che gli elettroni vennero dapprima considerati come particelle, e che il loro aspetto ondulatorio fu scoperto solo in un secondo tempo. La luce invece, a parte il caso di Newton che per errore le attribuiva natura 'corpuscolare', venne ritenuta dapprima un fenomeno ondulatorio e solo più tardi si scoprirono le sue

caratteristiche corpuscolari. In realtà tanto gli elettroni quanto la luce si comportano un po' come onde e un po' come particelle. Piuttosto che introdurre nuovi vocaboli, tipo «partonde», si è scelto di continuare a chiamarle «particelle», tenendo però sempre presente che si comportano secondo le regole di somma e combinazione delle frecce spiegate in queste lezioni. A quanto pare, *tutte* le «particelle» esistenti in Natura — quark, gluoni, neutrini e altre specie di cui parlerò nella prossima lezione — presentano tale comportamento quanto-meccanico.

Ecco dunque i tre eventi elementari di base costituenti tutti i fenomeni in cui intervengono luce ed elettroni:

1. un fotone si propaga da un punto a un altro;
2. un elettrone si propaga da un punto a un altro;
3. un elettrone emette o assorbe un fotone.

A ognuno di questi eventi elementari corrisponde un'ampiezza (una freccia) che può essere determinata in base a precise regole. Discuterò tra un istante queste leggi che permettono di costruire il mondo intero (esclusi, come al solito, i fenomeni nucleari e gravitazionali!).

Debbo premettere che la scena su cui si svolgono i fatti non è solo lo spazio, ma è lo spazio e il tempo. Finora ho trascurato gli aspetti riguardanti il tempo, quali l'istante esatto in cui un fotone lascia la sorgente o in cui arriva al rivelatore. Sebbene lo spazio sia tridimensionale, nel disegnare i grafici¹⁵ considererò solo una dimensione: la collocazione spaziale di un oggetto verrà riportata sull'asse orizzontale e quella temporale sull'asse verticale.

Il primo evento che disegnerò nello spazio e nel tempo, o nello spazio-tempo come potrei inavvertitamente chiamarlo, è una palla ferma (fig. 52). Giovedì mattina, indicato con T_0 , la palla occupa una certa posizione, indicata con X_0 . Alcuni istanti dopo, a T_1 , occuperà ancora la stessa posizione, perché sta ferma. Più tardi ancora, a T_2 , la palla sarà sempre in X_0 . Il grafico relativo alla palla ferma è quindi una striscia che sale verticalmente e che corrisponde alla zona occupata dalla palla.

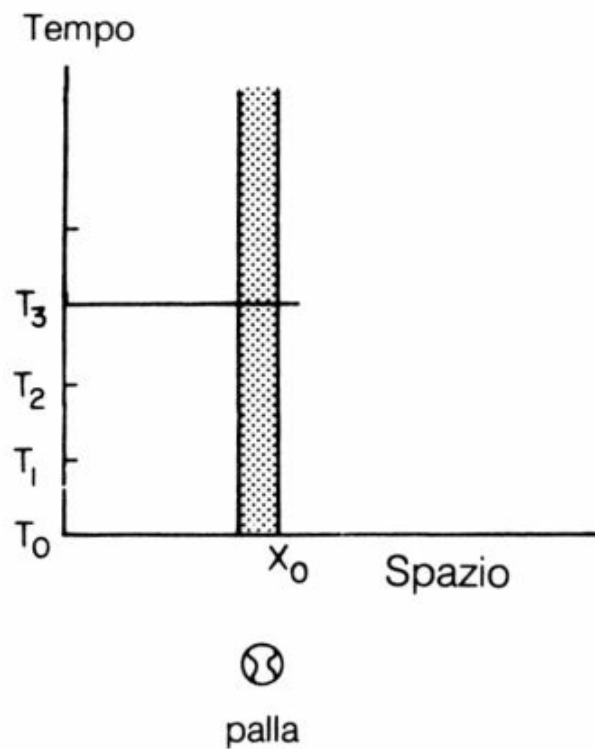


Fig. 52. Tutti gli avvenimenti dell'universo si svolgono nello spazio-tempo. Questo ha quattro dimensioni, tre spaziali e una temporale, ma verrà rappresentato in due sole dimensioni: una spaziale, lungo l'orizzontale, e una temporale, lungo la verticale. In qualunque istante si guardi una palla ferma (ad esempio a T_3), la si trova nello stesso posto. Al passare del tempo si ottiene pertanto una «striscia della palla» diretta verticalmente.

Che cosa succede se la palla si muove dritta verso un muro in assenza di gravità? Giovedì mattina, a T_0 si trova in X_0 (fig. 53), ma dopo un po' non sarà più nello stesso posto: si sarà spostata di un po' e si troverà in X_1 . Via via che la palla si sposta, nel diagramma spaziotemporale si forma una corrispondente «striscia della palla», questa volta inclinata. Urtando il muro, al quale corrisponde una striscia verticale perché è fermo, la palla rimbalza all'indietro e ripassa per X_0 da dove è partita, ma a un istante di tempo differente, T_6 .

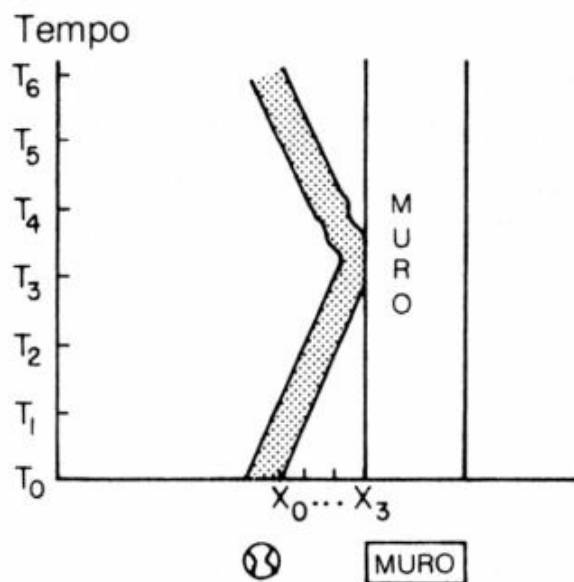


Fig. 53. Una palla in moto perpendicolarmente verso una parete sulla quale rimbalza tornando alla posizione iniziale (indicata sotto il grafico con X_0) si muove in una sola direzione e dà luogo a una «striscia della palla» inclinata. Agli istanti T_1 e T_2 la palla si avvicina alla parete, all'istante T_3 la urta e comincia a tornare indietro.

Per comodità è preferibile scegliere per la scala dei tempi un'unità di misura molto più piccola del secondo. Poiché si avrà a che fare con fotoni ed elettroni, che si muovono molto velocemente, una linea inclinata a 45° rappresenterà qualcosa che si muove alla velocità della luce. Dunque, per una particella che si sposta da X_1T_1 a X_2T_2 con tale velocità, la distanza orizzontale tra X_1 e X_2 è uguale a quella verticale tra T_1 e T_2 (fig. 54).

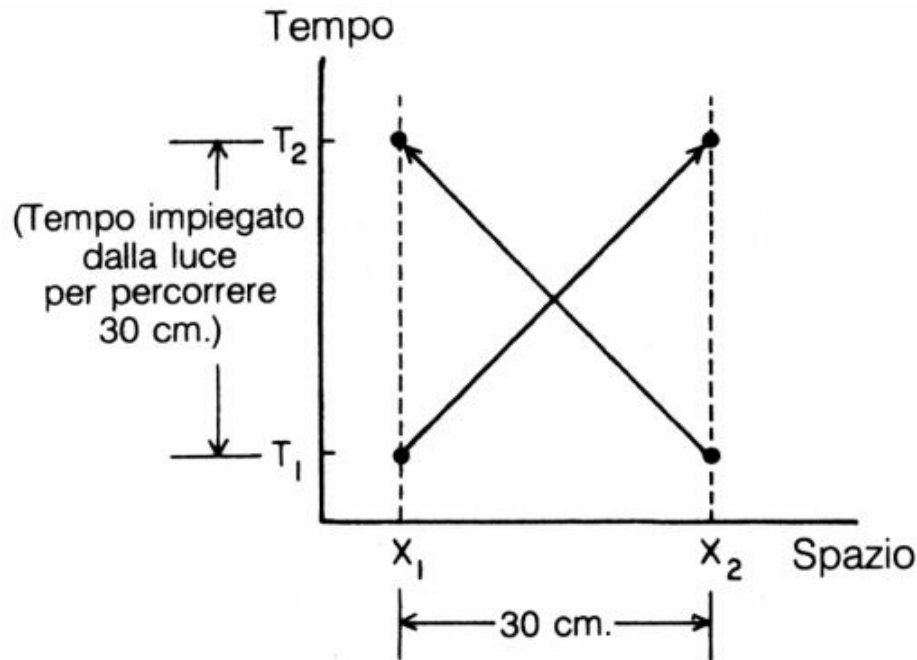


Fig. 54. La scala dei tempi usata nei grafici è tale che particelle che viaggiano alla velocità della luce sono rappresentate da linee inclinate a 45° nello spazio-tempo. Il tempo impiegato dalla luce per percorrere 30 centimetri, ad esempio da X_1 a X_2 o viceversa, è circa un milionesimo di secondo.

Il fattore per cui moltiplicare il tempo (in modo che un angolo di 45° corrisponda alla velocità della luce) viene chiamato c . Nelle formule di Einstein si trovano dappertutto fattori c , dovuti alla infelice scelta, come unità di misura del tempo, del secondo invece che del tempo impiegato dalla luce a percorrere un metro.

Passiamo ora ad analizzare il primo evento elementare: un fotone che viaggia da un punto a un altro. Questo evento verrà rappresentato, senza particolare motivo, da una linea ondulata che va da A a B. Forse, però, dovrei essere più preciso e dire: un fotone che si trova in un dato punto a un dato istante ha una certa ampiezza per arrivare in un altro punto in un altro istante. Nel mio grafico spaziotemporale (fig. 55) un fotone inizialmente nel punto A, cioè in X_1 e T_1 , ha una certa ampiezza di probabilità di comparire nel punto B, in X_2 e T_2 . Indicherò questa ampiezza con F (da A a B).

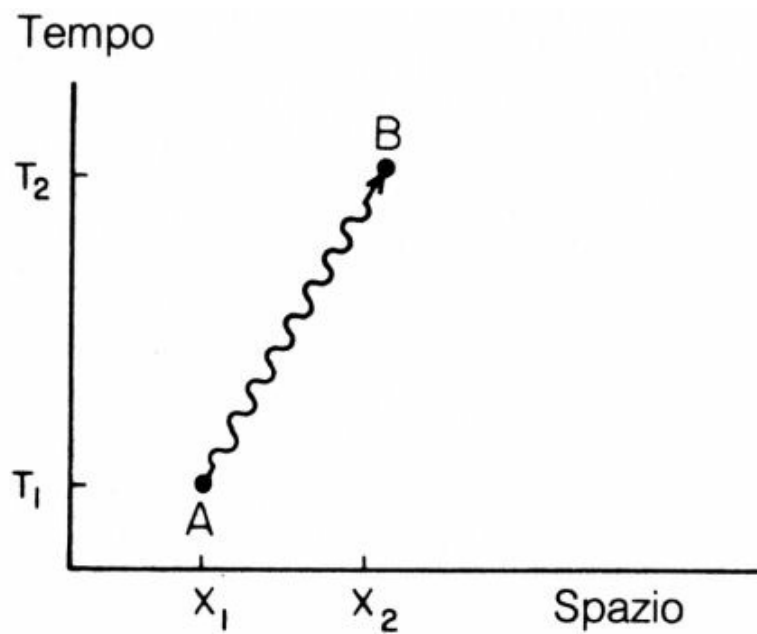


Fig. 55. Un fotone, rappresentato come una linea ondulata, ha un'ampiezza di probabilità per andare da un punto A a un punto B dello spazio-tempo che è indicata con F (da A a B) ed è espressa da una formula che dipende solo dalla differenza di posizione ($X_2 - X_1$) e dalla differenza di tempo ($T_2 - T_1$). In effetti è semplicemente l'inverso della differenza dei loro quadrati, cioè dell'«intervallo» I che può essere scritto $(X_2 - X_1)^2 - (T_2 - T_1)^2$.

C'è una formula precisa per la quantità F (da A a B). È una delle grandi leggi della Natura ed è estremamente semplice: F (da A a B) dipende dalla differenza di *posizione* e dalla differenza di *tempo* tra i due punti. In termini matematici queste due grandezze ¹⁶ possono essere espresse come $(X_2 - X_1)$ e $(T_2 - T_1)$.

Il contributo dominante a F (da A a B) si ha quando la luce si propaga con l'usuale velocità, cioè quando $(X_2 - X_1)$ è uguale a $(T_2 - T_1)$; ci si aspetterebbe che ciò dia l'intera ampiezza, ma c'è un'ampiezza anche per velocità di propagazione maggiori o minori di quella usuale. Nella lezione scorsa abbiamo visto che la luce non si propaga solo in linea retta; adesso troviamo anche che non viaggia alla velocità della luce!

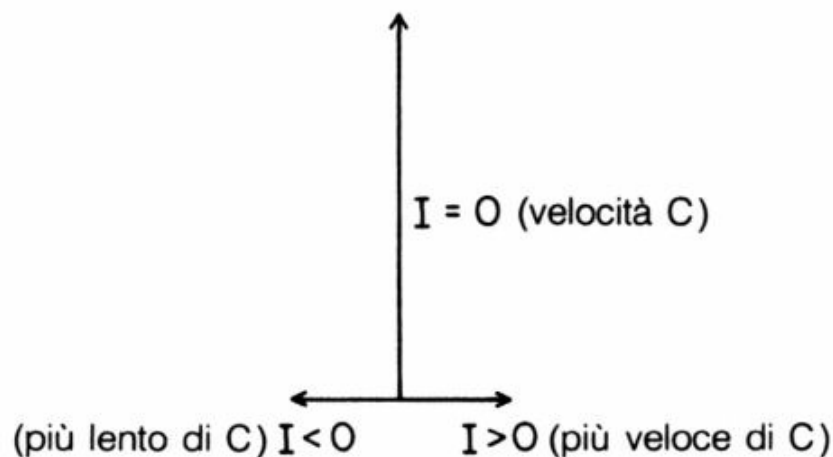


Fig. 56. Se la luce viaggia a velocità c , l'«intervallo» I è nullo e si ha un grosso contributo diretto verso le 12. Se I è maggiore di zero, si ha un piccolo contributo diretto verso le 3 e inversamente proporzionale a I ; se I è minore di zero, si ha un contributo analogo ma diretto verso le 9. Dunque la luce ha ampiezza non nulla per propagarsi con velocità sia maggiore che minore di c , ma tali ampiezze si elidono a vicenda per percorsi lunghi.

Può sorprendere che alla propagazione di un fotone con velocità diverse da quella convenzionale corrispondano ampiezze non nulle. Tali ampiezze sono molto piccole paragonate a quella del contributo con velocità c , anzi si elidono quando la luce viaggia su

lunghe distanze. Ma quando le distanze sono piccole, come in molti casi che vedremo, queste altre possibilità diventano essenziali e bisogna tenerne conto.

Il primo evento elementare, la prima legge fondamentale della fisica, è dunque: un fotone si propaga da un punto a un altro. Tale legge permette di spiegare tutta l'ottica, e costituisce l'intera teoria della luce! Be', non proprio tutta: ho lasciato da parte, come al solito, la polarizzazione, e anche l'interazione della luce con la materia, il che ci porta alla seconda legge.

Il secondo evento elementare fondamentale dell'elettrodinamica quantistica è il seguente: un elettrone si propaga da un punto A a un punto B dello spazio-tempo. (Per il momento lo penseremo come un elettrone semplificato, un elettrone fittizio, senza polarizzazione: ciò che i fisici chiamano particella di «spin zero». L'elettrone reale ha una specie di polarizzazione, che però non aggiunge nulla alle idee principali e rende solo un tantino più complicate le formule).

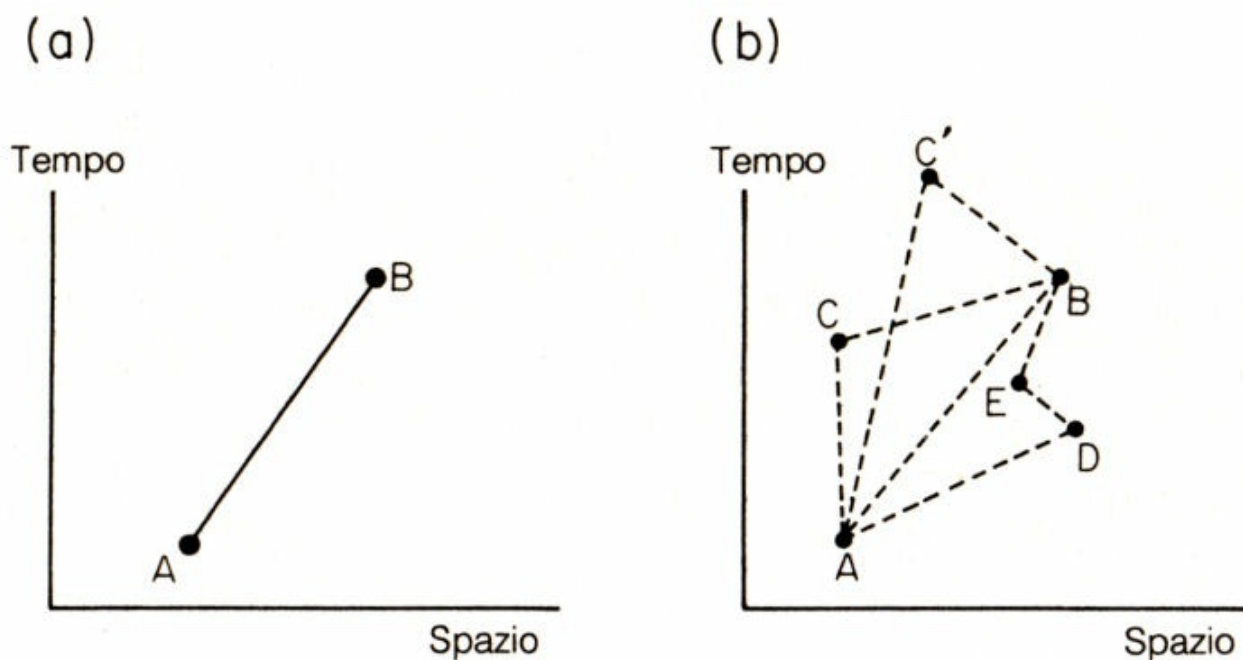


Fig. 57. L'ampiezza di probabilità che un elettrone si propaghi da un punto a un altro dello spazio-tempo sarà indicata con E (da A a B). Rappresento E (da A a B) con una linea retta tra i due punti (a), ma la si può pensare come la risultante di tante ampiezze (b), tra cui quelle per i percorsi con «due salti», in cui l'elettrone cambia direzione nei punti C o C', quelle per i percorsi con «tre salti», in cui cambia direzione in D e in E, oltre all'ampiezza relativa al percorso diretto da A a B. Nel suo viaggio nello spazio-tempo da A a B, l'elettrone può cambiare direzione un qualunque numero di volte (tra zero e infinito), e sono infiniti i punti in cui può farlo. Tutti questi contributi sono inclusi in E (da A a B).

L'ampiezza associata a questo evento base, che indicherò con E (da A a B), dipende anch'essa da $(X_2 - X_1)$ e da $(T_2 - T_1)$, nella stessa combinazione descritta nella nota 16, e inoltre da un numero che chiamerò « n », il quale, una volta determinato, permetterà ai risultati dei nostri calcoli di essere in accordo con i dati sperimentali. (Vedremo più avanti come determinare il valore di n). Purtroppo la formula è alquanto complicata e non saprei come spiegarla in termini semplici. Comunque, è interessante sapere che l'espressione dell'ampiezza per un fotone che va da A a B, F (da A a B), coincide con quella per un elettrone che va da A a B, E (da A a B), se si pone $n = 0$.¹⁷

Infine, il terzo evento elementare è: un elettrone assorbe o emette un fotone (non c'è differenza tra i due casi). Questo evento lo chiamerò «accoppiamento» o «interazione».

Per distinguere nei grafici gli elettroni dai fotoni, rappresenterò un elettrone che

viaggia nello spazio-tempo con un tratto rettilineo. Ogni interazione è quindi descritta come incontro tra due linee rette e una linea ondulata (fig. 58).

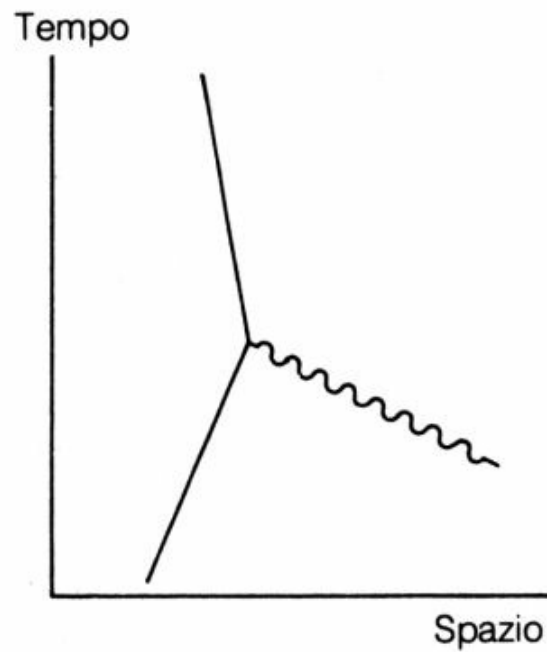


Fig. 58. Un elettrone, rappresentato dal tratto rettilineo, ha una certa ampiezza per assorbire o per emettere un fotone, rappresentato dal tratto ondulato. Poiché l'ampiezza per emissione e quella per assorbimento sono le stesse, entrambe le situazioni verranno chiamate «accoppiamento». L'ampiezza di accoppiamento è un numero, indicato con j , che vale circa $-0,1$ per l'elettrone. Tale numero viene a volte chiamato «carica» della particella.

L'ampiezza corrispondente a questa interazione non è espressa da una formula complicata; anzi, non dipende da nulla, è solo un numero, che io indicherò con j ; esso vale circa $-0,1$, equivalente a una contrazione a circa un decimo e ad una rotazione di mezzo giro.¹⁸

E con questo ho detto tutto quello che c'è da dire sui tre eventi base, a parte le solite piccole complicazioni dovute alla polarizzazione, che abbiamo deciso di trascurare. Adesso dovremo affrontare il problema di combinare insieme questi tre eventi elementari per giungere a rappresentare situazioni un po' più complicate.

Quale primo esempio, calcoleremo la probabilità che due elettroni, inizialmente nei punti 1 e 2 dello spazio-tempo, finiscano nei punti 3 e 4 (fig. 59). Tale evento può verificarsi in più modi. Un primo modo corrisponde al caso in cui l'elettrone inizialmente in 1 va in 3, e a ciò va associata l'ampiezza E (da 1 a 3), ottenuta ponendo 1 e 3 nella formula di E (da A a B), mentre l'elettrone 2 va in 4, con ampiezza E (da 2 a 4). Poiché questi due «sottoeventi» devono essere concomitanti, le due frecce vanno moltiplicate tra loro per ottenere la freccia relativa a questa possibilità. Perciò l'espressione per la freccia del «primo modo» è data da $E(\text{da } 1 \text{ a } 3) \times E(\text{da } 2 \text{ a } 4)$.

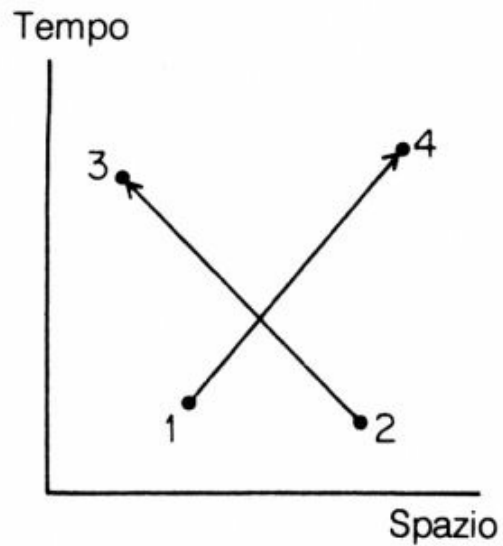
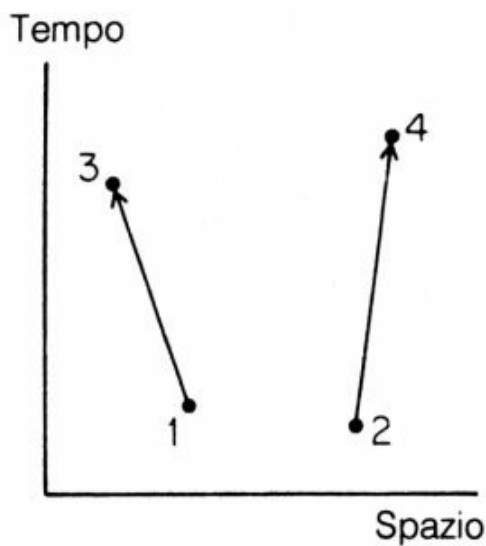


Fig. 59. Per calcolare la probabilità che due elettroni inizialmente nei punti 1 e 2 dello spazio-tempo finiscano nei punti 3 e 4, si deve calcolare, usando l'espressione di $E(\text{da } A \text{ a } B)$, la freccia per il «primo modo», in cui 1 va in 3 e 2 va in 4; poi si calcola la freccia per il «secondo modo», in cui 1 va in 4 mentre 2 va in 3 (una «intersezione»). Sommando queste frecce per il «primo» e per il «secondo modo» si ottiene una buona approssimazione della freccia finale. Ciò è vero per l'elettrone fittizio di «spin zero». Se avessimo incluso la polarizzazione dell'elettrone, avremmo dovuto sottrarre, invece che sommare, le due frecce.

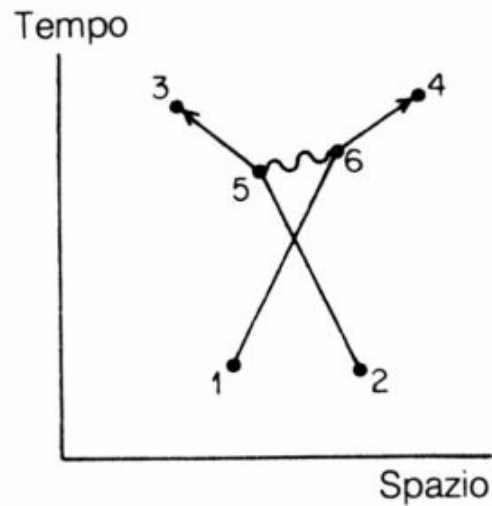
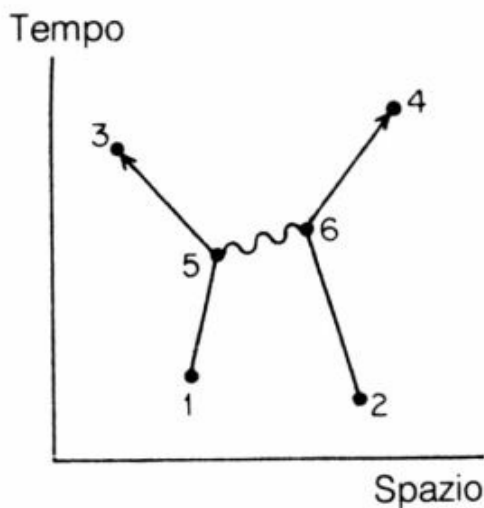


Fig. 60. Due «altri modi» in cui può verificarsi l'evento della fig. 59: in entrambi un fotone viene emesso in 5 e assorbito in 6. Queste alternative hanno condizioni finali identiche ai primi due modi: due elettroni partono e due elettroni arrivano; esse sono perciò indistinguibili dalle precedenti possibilità. Quindi, per ottenere un'approssimazione migliore della freccia risultante che descrive il fenomeno, le frecce relative a questi «altri modi» vanno sommate a quelle dei modi della fig. 59.

Il fenomeno può però avvenire anche in un secondo modo: l'elettrone in 1 va in 4, mentre quello in 2 va in 3. Anche in questo caso si hanno due «sottoeventi» concomitanti. Al «secondo modo» corrisponde perciò un'ampiezza $E(\text{da } 1 \text{ a } 4) \times E(\text{da } 2 \text{ a } 3)$, da sommare alla «prima» freccia.¹⁹

L'ampiezza così ottenuta costituisce una buona approssimazione per la descrizione di questo evento. Per fare un conto più esatto e in migliore accordo con i risultati sperimentali, dobbiamo considerare ancora altri modi in cui l'evento può verificarsi. Ad esempio, in corrispondenza a ciascuno dei due modi principali, un elettrone potrebbe andare prima in un altro punto a suo piacimento e lì emettere un fotone; il secondo elettrone, che nel frattempo se ne è andato in un altro punto, assorbe qui tale fotone (fig. 60). Il calcolo dell'ampiezza per il primo di questi due nuovi modi comporta il prodotto delle ampiezze relative ai seguenti passi: un elettrone va da 1 al nuovo punto 5, dove emette un fotone per poi proseguire fino a 3; l'altro elettrone va da 2 a 6, dove assorbe il

fotone e poi prosegue da 6 a 4. Ricordiamoci inoltre di includere l'ampiezza per il fotone che va da 5 a 6. L'espressione matematica per l'ampiezza relativa a questo modo è dunque data da $E(\text{da } 1 \text{ a } 5) \times j \times E(\text{da } 5 \text{ a } 3) \times E(\text{da } 2 \text{ a } 6) \times j \times E(\text{da } 6 \text{ a } 4) \times F(\text{da } 5 \text{ a } 6)$: un bel po' di contrazioni e rotazioni! (Lascio a voi il piacere di costruire l'espressione per il caso in cui l'elettrone in 1 arriva in 4 mentre quello in 2 arriva in 3).²⁰

Ma un momento: i punti 5 e 6 possono essere dovunque nello spazio e nel tempo — ripeto: dovunque — e per *tutte* queste posizioni si devono calcolare e poi sommare tra loro le frecce relative. Come vedete, sta diventando un lavoro enorme. Non che le regole siano molto difficili; è un po' come giocare a dama: le regole in sé sono semplici, ma si tratta di applicarle ripetutamente. La difficoltà del calcolo proviene dal dover mettere insieme un numero così grande di frecce. Ci vogliono quattro anni di specializzazione per imparare a fare questo conto in modo efficiente, e si tratta di un problema semplice! (Quelli troppo complicati li infiliamo direttamente in un calcolatore!).

Vorrei sottolineare un aspetto connesso con l'emissione e l'assorbimento dei fotoni: se il punto 6 è successivo al 5 si può dire che il fotone è stato emesso nel punto 5 e assorbito in 6 (fig. 61). Viceversa, se 6 è antecedente a 5, potremmo trovare preferibile dire che il fotone è stato emesso in 6 ed assorbito in 5, ma potremmo anche dire che il fotone viaggia all'indietro nel tempo!

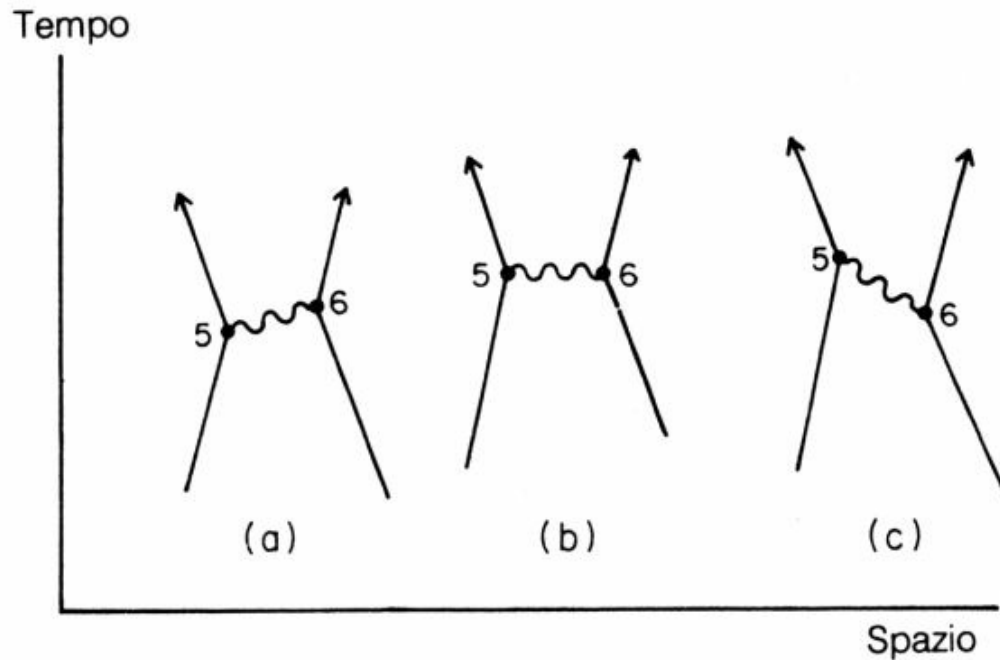


Fig. 61. Poiché la luce ha un'ampiezza non nulla per propagarsi con velocità maggiore o minore di quella convenzionale, nei tre esempi in figura i fotoni possono sempre considerarsi emessi nel punto 5 e assorbiti in 6, anche se il fotone dell'esempio (b) viene emesso allo *stesso istante* in cui viene assorbito, e quello dell'esempio (c) è emesso *dopo* essere stato assorbito — situazione che forse qualcuno preferirebbe descrivere come: fotone emesso in 6 e assorbito in 5; altrimenti il fotone deve viaggiare *all'indietro nel tempo*! Per quanto riguarda i calcoli (e la Natura), le due descrizioni sono assolutamente equivalenti (e possibili), per cui si preferisce semplicemente dire che un fotone viene «scambiato» e inserire le sue posizioni spaziotemporali nell'espressione di $F(\text{da } A \text{ a } B)$.

Comunque, non ci si deve preoccupare della direzione spaziotemporale in cui viaggia il fotone: è già tutto incluso nella formula che esprime $F(\text{da } 5 \text{ a } 6)$ e basta dire che è stato «scambiato» un fotone. Non è meravigliosa la semplicità della Natura?²¹

Oltre al fotone scambiato tra 5 e 6, ne potrebbe essere scambiato un secondo, ad esempio tra i punti 7 e 8 (fig. 62). È troppo faticoso scrivere tutti i passi le cui frecce

vanno moltiplicate tra loro, ma, come avrete notato, a ogni linea retta corrisponde E(da A a B), a ogni linea ondulata F(da A a B), e ad ogni interazione un fattore j . Ci sono sei E(da A a B), due F(da A a B) e quattro j , per *tutti i valori possibili di 5, 6, 7 e 8!* Il che significa dover moltiplicare tra loro e poi sommare miliardi di freccette!

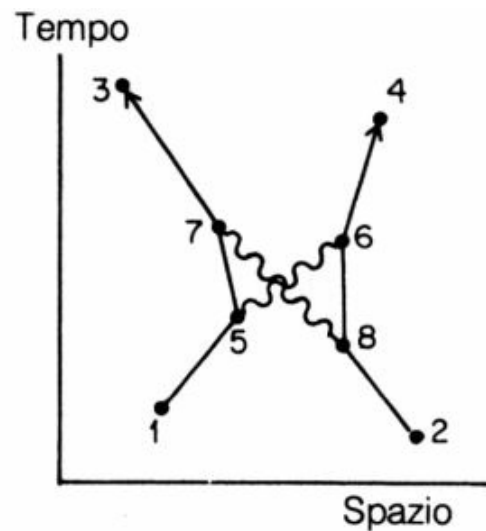


Fig. 62. Il fenomeno della fig. 59 può anche avvenire con lo scambio di *due* fotoni. Sono possibili moltissimi diagrammi di questo tipo (come si vedrà meglio più avanti), e uno di essi è mostrato in figura. La freccia relativa a questo modo va calcolata tenendo conto di tutte le possibili posizioni dei punti 5, 6, 7 e 8, e il conto risulta molto complesso. Poiché j risulta minore di 0,1, la lunghezza di tale freccia, che comprende quattro interazioni, è in generale inferiore a un decimillesimo delle frecce relative al «primo» e al «secondo modo» della fig. 59 che non coinvolgono j .

Insomma, calcolare l'ampiezza per questo semplice evento si direbbe proprio un'impresa disperata, ma quando si è all'università e c'è in ballo la laurea, si tiene duro e si va avanti.

Ma una speranza di successo c'è e si trova nel magico numero j . Abbiamo discusso diversi modi in cui l'evento considerato può verificarsi; ora si noti che i primi due non contengono j nel calcolo dell'ampiezza, i successivi due contengono $j \times j$, e gli ultimi contengono $j \times j \times j \times j$. Poiché $j \times j$ è minore di 0,01, la freccia per i secondi due modi è in linea di massima inferiore all'1% di quella per i primi due, mentre l'ampiezza che contiene $j \times j \times j \times j$ è minore dell'1% dell'1% (cioè di una parte su 10.000) di quella senza j . Se si ha abbastanza tempo a disposizione su un calcolatore, si possono calcolare le ampiezze che coinvolgono sei j (dell'ordine di una parte per milione) e uguagliare la precisione dei dati sperimentali. Ecco come si procede al calcolo di eventi semplici; e questo è tutto!

Si consideri ora un evento diverso: si comincia con un fotone e un elettrone e si finisce con un fotone e un elettrone. Tale evento può avvenire nel modo seguente: il fotone viene assorbito dall'elettrone, il quale continua a viaggiare per un po' e poi emette un nuovo fotone. Questo processo viene chiamato *diffusione della luce*, e nel tracciare i relativi diagrammi si devono includere alcune possibilità bizzarre (fig. 63). Ad esempio, l'elettrone può emettere il nuovo fotone *prima* di assorbire il vecchio (*b*). Oppure, possibilità ancora più strana, l'elettrone emette un fotone, *viaggia all'indietro nel tempo*, assorbe il vecchio fotone e poi si propaga nuovamente in avanti nel tempo (*c*). Il percorso di questi elettroni che viaggiano «all'indietro» può essere tanto lungo da apparire reale negli esperimenti di fisica in laboratorio. Questo comportamento è incluso nei diagrammi considerati e nell'espressione di E(da A a B).

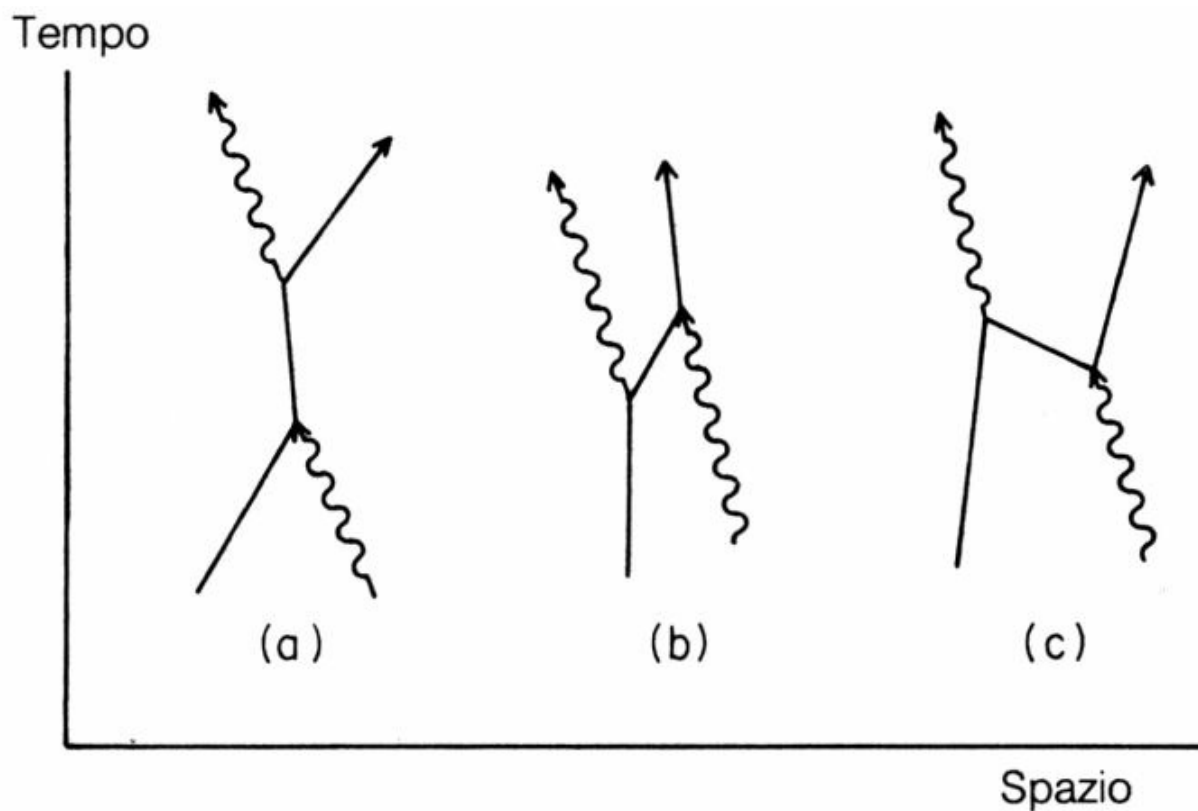


Fig. 63. La diffusione della luce consiste nell'assorbimento e nell'emissione di un fotone da parte di un elettrone. Come evidenziato in (b), questi due eventi non avvengono necessariamente in tale sequenza. L'esempio (c) mostra una possibilità strana ma reale: l'elettrone emette un fotone, poi si precipita all'indietro nel tempo ad assorbire un altro fotone e infine si propaga nuovamente in avanti nel tempo.

Dal punto di vista del tempo che avanza, gli elettroni che viaggiano «all'indietro» appaiono identici a quelli ordinari tranne che vengono attratti da essi, cioè hanno «carica positiva». (Includere gli effetti della polarizzazione renderebbe evidente perché il segno j debba essere rovesciato per gli elettroni che viaggiano all'indietro nel tempo, facendo apparire positiva la loro carica). Per questo motivo sono chiamati «positroni». I positroni sono particelle gemelle degli elettroni e sono un esempio di «antiparticella». ²²

Questa proprietà è generale. Ogni particella esistente in Natura ha un'ampiezza per muoversi all'indietro nel tempo, e perciò una corrispondente antiparticella. Se particella e antiparticella si urtano, possono annichilarsi a vicenda, dando luogo ad altre particelle. Così un elettrone e un positrone che si annichilano producono, di solito, uno o due fotoni. E i fotoni? Come si è visto prima, questi appaiono assolutamente identici quando viaggiano all'indietro nel tempo, per cui coincidono con le proprie antiparticelle. Si noti con quanta astuzia l'eccezione viene inserita nella regola!

Ora cercherò di illustrarvi come appare un elettrone che viaggia all'indietro nel tempo a chi si muove in avanti. A questo scopo dividerò il grafico della fig. 64 in tanti intervalli di tempo, da T_0 a T_{10} , tracciando una serie di linee parallele per aiutare l'occhio. Si inizi a T_0 con un elettrone in moto verso un fotone che viaggia in direzione opposta. All'istante T_3 il fotone si tramuta in due particelle: un elettrone e un positrone. Quest'ultimo non dura a lungo; dopo un po', a T_5 , incontra il primo elettrone e si annichila dando luogo a un nuovo fotone. Nel frattempo, l'elettrone creato dal fotone iniziale continua il suo viaggio nello spazio-tempo.

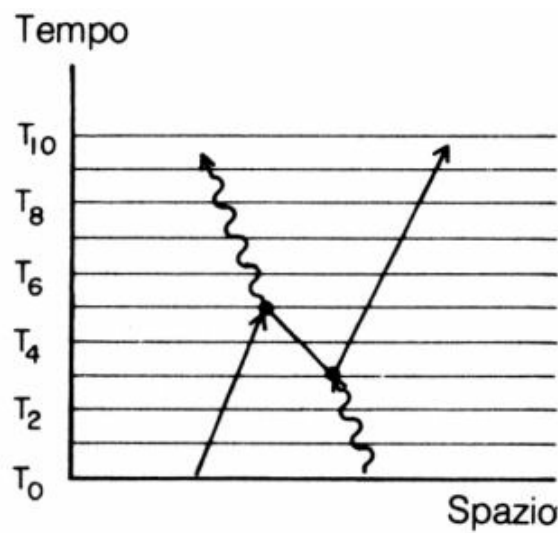


Fig. 64. L'esempio (c) della fig. 63 descritto considerando solo il moto in avanti nel tempo (come si è costretti a fare in laboratorio): da T_0 a T_3 si osservano un elettrone e un fotone che si muovono l'uno verso l'altro. All'improvviso, a T_3 , il fotone si «disintegra» e al suo posto appaiono *due* particelle: un elettrone e un nuovo tipo di particella (chiamata «positrone»). Quest'ultima, che è un elettrone che viaggia all'indietro nel tempo, sembra muoversi verso l'elettrone iniziale (nientemeno!). A T_5 , il positrone incontra questo elettrone e i due si annichilano dando luogo ad un nuovo fotone. Nel frattempo l'elettrone creato dal fotone originale continua il suo moto in avanti nel tempo. Questa successione di avvenimenti è stata osservata in laboratorio ed è automaticamente inclusa nell'espressione di $E(\text{da } A \text{ a } B)$.

Passiamo ora agli elettroni negli atomi. Per poter descrivere il comportamento di questi elettroni occorre introdurre un elemento nuovo: il nucleo, la parte pesante al centro di ciascun atomo, che contiene almeno un protone (quest'ultimo è una sorta di «vaso di Pandora» che apriremo nella prossima lezione).

Non posso spiegare in questa lezione le leggi esatte del comportamento del nucleo, che sono troppo complesse. Ma se lo si considera in quiete, se ne può assimilare il comportamento a quello di una particella la cui ampiezza per andare da un punto a un altro dello spazio-tempo è espressa dalla formula $E(\text{da } A \text{ a } B)$, ma con un valore di n molto più grande; essendo il nucleo tanto più pesante di un elettrone, si può approssimativamente affermare che esso resta di fatto nella stessa posizione spaziale mentre si sposta nel tempo.

L'atomo più semplice, chiamato idrogeno, è formato da un elettrone e da un protone. Tramite lo scambio di fotoni, il protone trattiene vicino a sé l'elettrone che gli danza intorno (fig. 65).²³

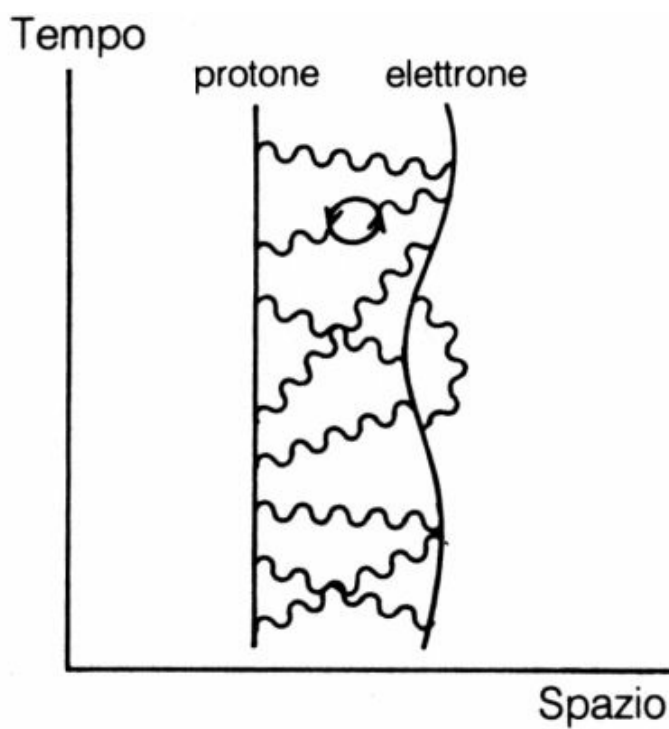


Fig. 65. Un elettrone viene mantenuto entro una certa distanza dal nucleo di un atomo dallo scambio di fotoni col protone (un «vaso di Pandora» in cui guarderemo nell'ultima lezione). Per il momento il protone viene considerato come una particella ferma. In figura è mostrato un atomo di idrogeno, che consiste in un protone e in un elettrone che si scambiano fotoni.

Naturalmente anche gli atomi che contengono più di un protone e un numero uguale di elettroni diffondono la luce (ad esempio, gli atomi dell'aria diffondono la luce solare rendendo azzurro il cielo), ma i loro diagrammi dovrebbero contenere un tale numero di linee diritte e ondulate che non ci si capirebbe più nulla!

Ora vi mostrerò un diagramma relativo alla diffusione della luce da parte dell'elettrone di un atomo di idrogeno (fig. 66). Mentre elettrone e nucleo stanno scambiandosi fotoni, arriva un fotone dall'esterno dell'atomo, urta l'elettrone e ne viene assorbito; dopodiché viene emesso un nuovo fotone. (Ci sono, al solito, anche altre possibilità da considerare, ad esempio che il nuovo fotone venga emesso prima che il vecchio sia stato assorbito). L'ampiezza totale relativa a tutti i modi in cui l'elettrone può diffondere un fotone è rappresentata da una freccia equivalente a una certa contrazione e a una certa rotazione (freccia che in seguito indicherò con «S»). Contrazione e rotazione dipendono dal nucleo e dalla disposizione degli elettroni nell'atomo, e sono differenti per sostanze differenti.

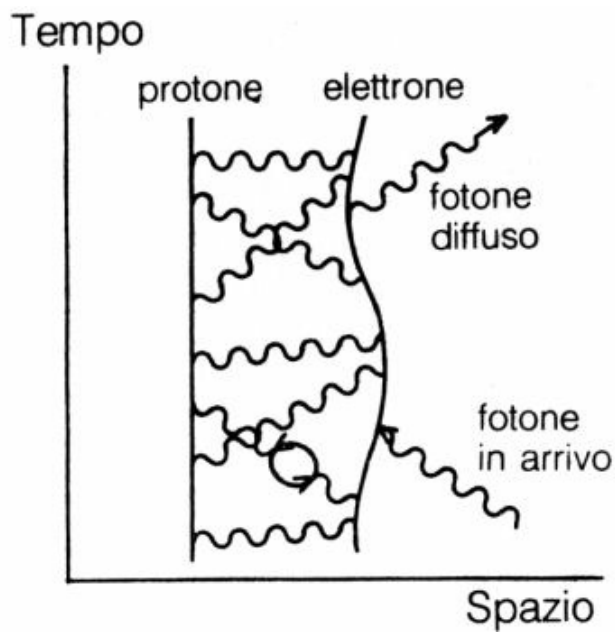


Fig. 66. La diffusione della luce da parte degli elettroni negli atomi è il fenomeno che sta alla base della riflessione parziale in una lamina di vetro. Il diagramma in figura mostra un possibile modo di diffusione in un atomo di idrogeno.

Torniamo alla riflessione parziale della luce su una lamina di vetro. Come funziona? Ho parlato di riflessione sulle superfici frontale e posteriore, ma si è trattato di una semplificazione fatta per facilitare le cose all'inizio. In realtà le superfici non hanno alcun effetto sulla luce. Il fotone incidente viene diffuso dagli elettroni all'interno del vetro, e quello che finisce nel rivelatore è un *nuovo* fotone. La cosa interessante è che, invece di sommare le miriadi di piccole frecce che rappresentano le ampiezze relative alla diffusione del fotone incidente da parte di tutti gli atomi del vetro, si può ottenere lo stesso risultato sommando solo due frecce, corrispondenti alla riflessione sulla «superficie frontale» e sulla «superficie posteriore». Vediamo perché.

Per discutere la riflessione su una lamina dal nuovo punto di vista, occorre prendere in considerazione la dimensione temporale. Finora, parlando della luce proveniente da una sorgente monocromatica, abbiamo usato un cronometro immaginario la cui lancetta gira durante il moto del fotone e, con la sua posizione, determina l'angolo dell'ampiezza relativa a un dato percorso. Nella formula $F(\text{da A a B})$, che esprime l'ampiezza perché il fotone vada da un punto A a un punto B, non si fa menzione di rotazione. Che cosa è successo al cronometro? Che fine ha fatto la rotazione?

Nella prima lezione ho considerato delle «sorgenti monocromatiche», senza ulteriori chiarimenti. Per analizzare correttamente la riflessione parziale su una lamina occorre saperne di più su queste sorgenti monocromatiche. In generale l'ampiezza per l'emissione di un fotone dalla sorgente varia nel *tempo*, e in particolare ciò che varia è il suo angolo. Una sorgente di luce bianca, che è una sovrapposizione di molti colori, emette fotoni in modo caotico; l'angolo dell'ampiezza varia bruscamente e irregolarmente, a sbalzi. Una sorgente *monocromatica* consiste invece in un dispositivo accuratamente costruito in modo che l'ampiezza per l'emissione di un fotone a un certo istante è facilmente calcolata: il suo angolo varia nel tempo con velocità *costante*, come quello di una lancetta di cronometro. (In effetti, questa freccia gira alla stessa velocità della lancetta del cronometro immaginario usato prima, ma in verso opposto: si veda la fig. 67).

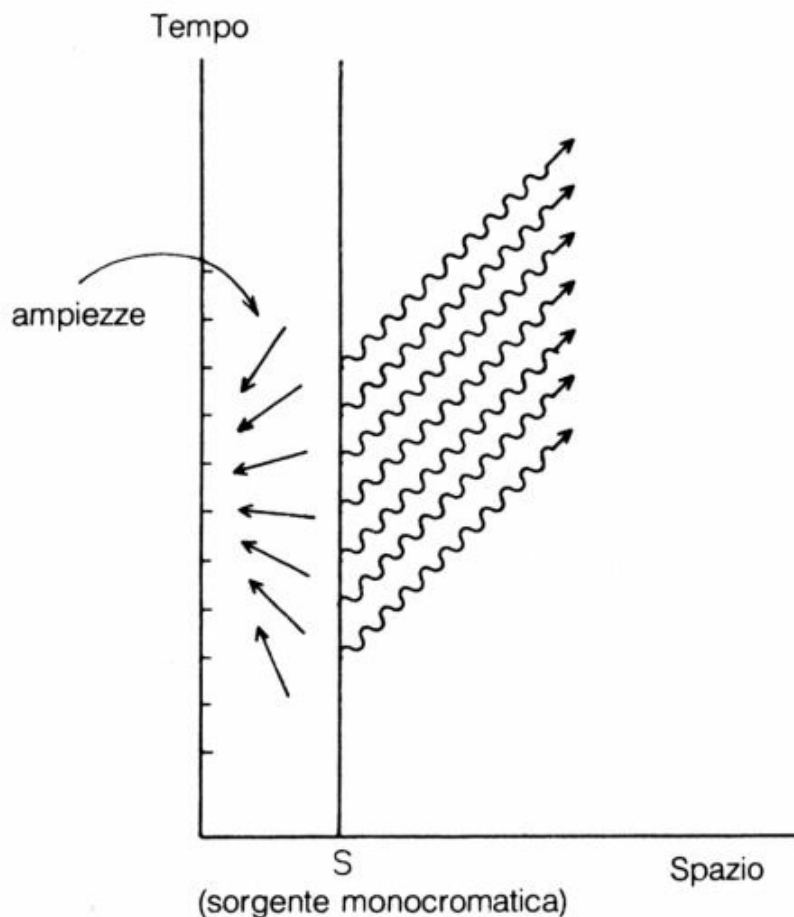


Fig. 67. Una sorgente monocromatica è un bellissimo dispositivo che emette fotoni in modo estremamente prevedibile: l'ampiezza per l'emissione di un fotone a un dato *istante* ruota uniformemente in senso antiorario col passare del tempo. Così l'ampiezza per l'emissione di un fotone a un istante successivo ha un angolo minore. Negli esempi seguenti si assumerà che la luce emessa dalla sorgente viaggi a velocità c , essendo grandi le distanze percorse.

La velocità di rotazione dipende dal colore della luce; la freccia per la luce blu ruota con velocità quasi doppia di quella per la luce rossa. Il «cronometro immaginario» era in realtà la sorgente monocromatica: l'angolo a cui si trova la freccia relativa a un dato percorso dipende dall'*istante* in cui il fotone è stato emesso dalla sorgente.

Una volta emesso il fotone, non vi è ulteriore rotazione della freccia durante il viaggio da un punto a un altro dello spazio-tempo. Sebbene l'espressione F (da A a B) contenga ampiezze per propagazione con velocità diverse da c , nel nostro esperimento la distanza tra sorgente e rivelatore è molto grande (rispetto a quelle atomiche), e quindi gli unici contributi alla lunghezza di F (da A a B) che non si elidono sono quelli che viaggiano con velocità c .

Prima di iniziare a calcolare in questo nuovo modo la riflessione parziale, è opportuno caratterizzare completamente l'evento come segue: il rivelatore in A fa uno scatto all'*istante* T (fig. 68 a). Suddividiamo ora la lamina di vetro in sezioni molto sottili, mettiamo sei. Come abbiamo visto nella seconda lezione, la luce è sostanzialmente riflessa dalla zona centrale di uno specchio, perché, sebbene ogni elettrone diffonda la luce in tutte le direzioni, sommando tra loro tutte le frecce relative a una sezione della lamina l'unica zona i cui contributi *non* si elidono è quella centrale, dove la luce arriva quasi perpendicolarmente dalla sorgente e prosegue lungo una delle due direzioni seguenti: o risale verso il rivelatore oppure va dritta attraverso il vetro. La freccia risultante per l'evento verrà quindi ottenuta sommando le frecce che rappresentano la

riflessione della luce nei sei punti centrali, da X_1 a X_6 , disposti verticalmente all'interno del vetro.

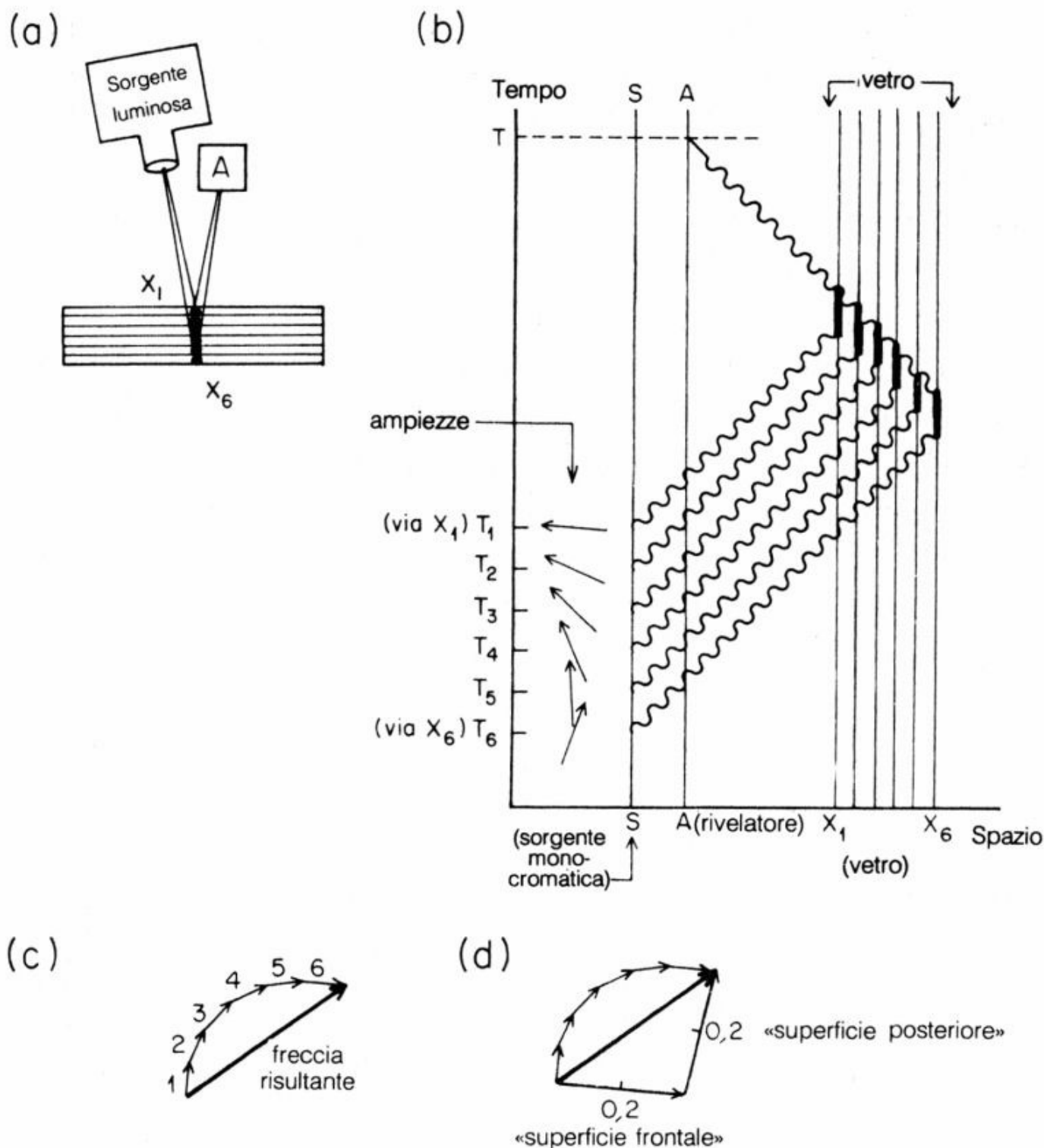


Fig. 68. Per cominciare la nostra nuova analisi della riflessione parziale, dividiamo la lamina di vetro in sezioni (sei, nell'esempio in figura) e consideriamo i possibili modi in cui la luce può andare dalla sorgente al vetro e da lì al rivelatore in A. Gli unici punti importanti della lamina, per i quali le ampiezze di diffusione non si elidono tra loro, sono al centro di ogni sezione; in (a) questi punti, da X_1 a X_6 , sono mostrati nella loro collocazione fisica all'interno del vetro, mentre in (b) sono mostrati come linee verticali in un grafico spaziotemporale. Vogliamo calcolare la probabilità del seguente evento: il rivelatore in A fa uno scatto all'istante T . Nel diagramma spaziotemporale a questo evento corrisponde il punto in cui A e T si intersecano.

Ciascuno dei modi in cui tale evento può verificarsi consiste in quattro passi successivi, per cui occorre moltiplicare quattro frecce. I passi sono illustrati in (b): 1) un fotone parte dalla sorgente a un dato istante (le frecce da T_1 a T_6 rappresentano le ampiezze per l'emissione in tali istanti); 2) il fotone va dalla sorgente a uno dei punti nel vetro (le sei alternative sono disegnate come linee ondulate che salgono verso destra); 3) un elettrone in uno di detti punti diffonde il fotone (mostrato come un breve tratto verticale più spesso); 4) un nuovo fotone (la linea ondolata che sale verso sinistra) si dirige verso il rivelatore dove giunge precisamente all'istante T . Le ampiezze per i passi 2, 3 e 4 sono identiche per tutte le alternative, solo quelle per il passo 1 sono differenti: un fotone diffuso internamente al vetro, ad

esempio in X_2 , deve partire dalla sorgente a T_2 , *prima* di un fotone diffuso da un elettrone alla superficie del vetro, in X_1 .

Dopo aver moltiplicato tra loro le quattro frecce relative a ciascuna alternativa, si ottengono le frecce mostrate in (c), più corte di quelle in (b) e ruotate di 90° in base alle caratteristiche della diffusione della luce da parte degli elettroni del vetro. Sommando nell'ordine queste frecce si ottiene un arco di cerchio, la cui corda è la freccia risultante. La stessa freccia finale è ottenibile disegnando due frecce radiali, mostrate in (d), e «sottraendo» l'una all'altra, cioè girando la freccia «frontale» nella direzione opposta prima di sommarla a quella «posteriore». Questa scorciatoia è stata usata come semplificazione nella prima lezione.

E adesso calcoliamo le frecce relative a questi sei percorsi, ciascuno dei quali comporta quattro passi, ossia quattro frecce da moltiplicare.

1. un fotone viene emesso dalla sorgente a un dato istante.
2. il fotone si propaga dalla sorgente a uno dei punti nel vetro.
3. il fotone è diffuso da un elettrone in quel punto.
4. un nuovo fotone risale fino al rivelatore.

Alle ampiezze relative ai passi 2 e 4 (propagazione di un fotone) non faremo corrispondere alcuna contrazione né alcuna rotazione, perché si può assumere che nel percorso tra la sorgente e il vetro o tra questo e il rivelatore il fascio di luce non si attenui né si allarghi. L'ampiezza per il terzo passo (diffusione di un fotone da parte di un elettrone) vale S , cioè una certa contrazione e una certa rotazione ed è la stessa in ogni punto del vetro. (Come ho detto prima, questo valore di S dipende dal materiale; per il vetro la rotazione risulta di 90°). In conclusione, delle quattro frecce da moltiplicare l'unica a cambiare da un percorso all'altro è quella relativa al primo passo, cioè l'ampiezza per l'emissione di un fotone dalla sorgente a un certo istante.

L'istante in cui un fotone deve essere emesso per arrivare nel rivelatore A al tempo T (fig. 68 b) non è lo stesso per tutti i percorsi. Il fotone che rimbalza in X_2 deve venire emesso leggermente *prima* del fotone che rimbalza in X_1 , perché il suo percorso è più lungo. Dunque la freccia relativa a T_2 sarà ruotata un po' di più di quella a T_1 , perché l'ampiezza per l'emissione di un fotone da parte di una sorgente monocromatica ruota in senso antiorario col passare del tempo. Lo stesso vale per le altre frecce fino a T_6 : hanno tutte lunghezza identica ma sono ruotate di angoli differenti, puntano cioè, in direzioni differenti, in quanto corrispondono a fotoni emessi dalla sorgente in istanti differenti.

Contraendo la freccia relativa a T_1 del fattore prescritto dai passi 2, 3 e 4, e ruotandola dei 90° prescritti dal passo 3, si ottiene la freccia 1 (fig. 68 c). Con lo stesso procedimento si trovano le frecce da 2 a 6. Queste sei frecce hanno tutte la stessa lunghezza (contratta) e formano tra loro esattamente gli stessi angoli delle frecce relative agli istanti di emissione.

Ora dobbiamo sommarle. Congiungendole nell'ordine, da 1 a 6, si ottiene una specie di arco, o parte di un cerchio, la cui corda è appunto la freccia finale. Poiché maggior spessore di vetro significa più sezioni, cioè più frecce e perciò una porzione maggiore di cerchio, la lunghezza della corda aumenta con lo spessore del vetro finché si arriva a un semicerchio e a una freccia finale pari al diametro. Poi la sua lunghezza *diminuisce* con l'aumentare dello spessore, e quando il cerchio è completo si inizia un nuovo ciclo. Il quadrato della lunghezza della freccia finale dà la probabilità dell'evento e varia ciclicamente da zero al 16%.

Un trucco matematico permette di ottenere la stessa risposta in modo diverso (fig. 68 d): se dal centro del «cerchio» tracciamo due frecce, una fino alla coda della freccia 1 e l'altra alla punta della 6, otteniamo due raggi. Se la freccia radiale che va dal centro alla coda della 1 viene ruotata di 180° («cambiata di segno») e sommata all'altra freccia radiale, si ottiene la stessa freccia finale! Questo è quanto avevo fatto nella prima lezione: i due raggi del cerchio sono le frecce che rappresentavano la riflessione sulle superfici «frontale» e «posteriore». Ciascuna di esse ha la famosa lunghezza 0,2. ²⁴

Possiamo quindi ottenere la risposta corretta per la probabilità di riflessione parziale immaginando che tutta la riflessione avvenga solo sulle superfici frontale e posteriore (il che è falso). In questa analisi più intuitiva le frecce relative alla «superficie frontale» e alla «superficie posteriore» sono costruzioni fittizie che permettono di ottenere la risposta esatta, mentre l'analisi discussa adesso, usando un diagramma spaziotemporale dove le varie frecce formano un arco di cerchio, costituisce una rappresentazione più accurata di ciò che succede nella realtà: la riflessione parziale è dovuta alla diffusione della luce da parte degli elettroni *interni* al vetro.

Cosa accade, invece, alla luce che *attraversa* il vetro? Anzitutto c'è da considerare il percorso in cui il fotone attraversa il vetro senza interagire con alcun elettrone, e la relativa ampiezza (fig. 69 a).

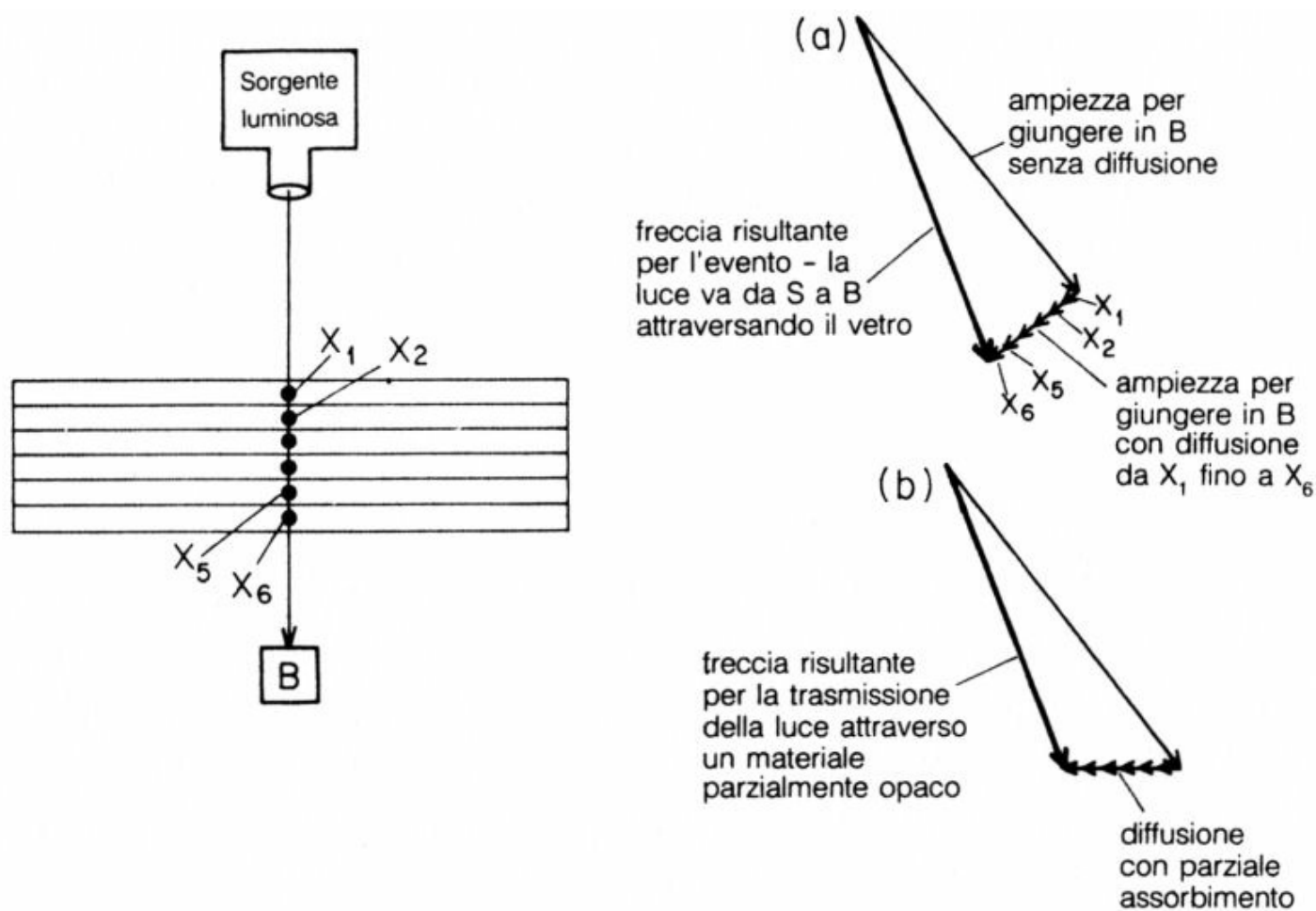


Fig. 69. Nella trasmissione della luce attraverso una lamina di vetro, fino al rivelatore in B, l'ampiezza dominante viene dalla propagazione senza diffusione da parte degli elettroni del vetro, come illustrato in (a). A questa freccia occorre aggiungere quelle più piccole corrispondenti alla diffusione da parte delle singole sezioni, rappresentate dai punti da X_1 a X_6 . Queste sei frecce hanno tutte la stessa lunghezza, perché l'ampiezza di diffusione è la stessa in tutti i punti del vetro, e la stessa direzione, perché la lunghezza dei vari percorsi dalla sorgente al rivelatore in B è la stessa per tutti i punti X. Sommando queste frecce più piccole alla principale, si trova la freccia finale per la trasmissione della luce attraverso il vetro. Essa risulta più ruotata di quanto lo sarebbe se la luce si propagasse senza diffusione. Per questa ragione si ha l'impressione che la luce impieghi più tempo a viaggiare attraverso il vetro che attraverso il vuoto o

l'aria. La rotazione della freccia risultante causata dagli elettroni del materiale è chiamata «indice di rifrazione».

Nei materiali trasparenti le frecce più piccole sono ad angolo retto rispetto alla principale (in realtà, includendo gli effetti della diffusione multipla, si vede che esse formano un trailo di curva così da evitare che la freccia finale risulti più lunga della principale: la Natura ha disposto le cose in modo che non esca mai più luce di quanta ne entri). Nei materiali parzialmente opachi, che assorbono leggermente la luce, le frecce secondarie sono inclinate verso la principale, dando luogo a una freccia risultante sensibilmente più corta, come mostrato in (b). La minore lunghezza di tale freccia rappresenta la ridotta probabilità di trasmissione di un fotone attraverso un materiale parzialmente opaco.

Questa costituisce la freccia più rilevante in termini di lunghezza. Ma ci sono altri sei modi in cui un fotone può giungere nel rivelatore sotto il vetro: il fotone emesso può urtare un elettrone in X_1 e il fotone diffuso finire in B, oppure l'urto può avvenire in X_2 e così via. Queste sei frecce hanno la stessa lunghezza di quelle che formavano l'arco di cerchio dell'esempio precedente, essendo tale lunghezza determinata ancora dall'ampiezza S relativa alla diffusione di un fotone da parte di un elettrone del vetro. Ma questa volta sono rivolte tutte e sei nella stessa direzione, perché i sei percorsi in cui si verifica un unico urto hanno tutti la stessa lunghezza. Per sostanze trasparenti, quali il vetro, la direzione delle frecce secondarie forma un angolo retto con la freccia principale. Sommando tali frecce, si ottiene una freccia risultante che ha la stessa lunghezza di quella principale ma è ruotata in una direzione leggermente diversa. Più spesso è il vetro più frecce secondarie vi sono e maggiore è la rotazione della freccia risultante. Su questo si basa appunto il funzionamento di una lente convergente: inserendo un opportuno spessore di vetro lungo percorsi più corti si ottengono per tutti i percorsi frecce che puntano nella stessa direzione.

Lo stesso effetto si avrebbe se i fotoni andassero più lentamente nel vetro che nell'aria: anche questo darebbe luogo a una maggior rotazione della freccia finale. Ecco perché, come ho detto prima, si ha l'impressione che la luce vada più lenta nel vetro, o nell'acqua, che nell'aria. In realtà, il «rallentamento» della luce non è altro che la maggior rotazione della freccia finale provocata dalla diffusione da parte degli atomi del vetro o dell'acqua. Il grado di rotazione in più della freccia finale, dovuto all'attraversamento di un materiale, viene detto «indice di rifrazione». ²⁵

Per le sostanze che assorbono la luce, le frecce secondarie formano con quella principale un angolo inferiore ai 90° (fig. 69 b). Ne consegue che la freccia risultante è più corta di quella principale, per cui la probabilità che un fotone attraversi un vetro opaco è minore della probabilità che attraversi un vetro trasparente.

In questo modo, a partire da soli tre eventi base, vengono spiegati tutti gli aspetti dei fenomeni dei quali si è parlato nelle prime due lezioni (la riflessione parziale della luce con ampiezza 0,2, il suo «rallentamento» nell'acqua e nel vetro, e così via). Gli stessi eventi elementari permettono altresì di spiegare tutti, o quasi, gli altri fenomeni.

Che l'enorme varietà della Natura sia il risultato di una monotona ripetizione dei tre eventi base in varie combinazioni non è un'idea facile da accettare. Ma è proprio così. Proverò a dare un'idea di come ha luogo tale varietà.

Inizierò dai fotoni (fig. 70).

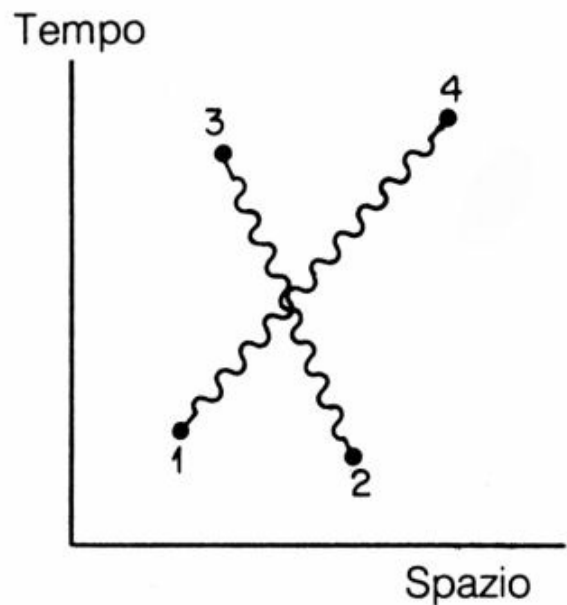
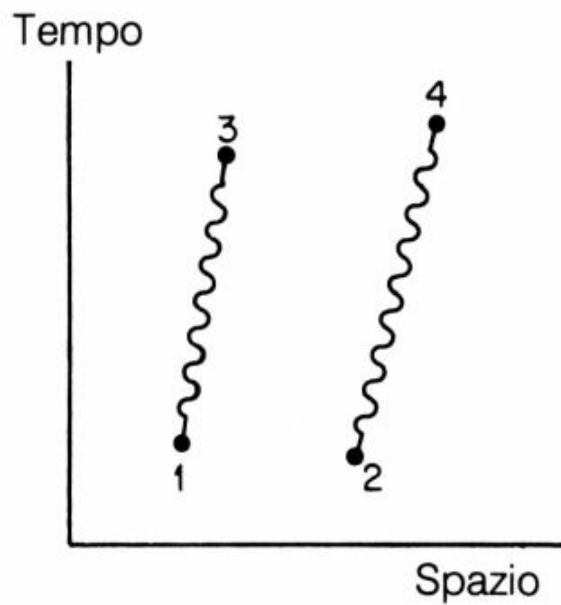


Fig. 70. Due fotoni nei punti 1 e 2 dello spazio-tempo hanno un'ampiezza per giungere nei punti 3 e 4 che può essere approssimata considerando i due modi principali in cui il fenomeno può accadere: $F(\text{da } 1 \text{ a } 3) \times F(\text{da } 2 \text{ a } 4)$ e $F(\text{da } 1 \text{ a } 4) \times F(\text{da } 2 \text{ a } 3)$. A seconda della posizione relativa dei punti 1, 2, 3 e 4, si producono diversi gradi di interferenza.

Qual è la probabilità che due fotoni dai punti 1 e 2 dello spazio-tempo arrivino in due rivelatori nei punti 3 e 4? Per questo fenomeno vi sono due modi principali, ciascuno dei quali dipende dal verificarsi concomitante di due fatti: i due fotoni possono andare direttamente, con ampiezza $F(\text{da } 1 \text{ a } 3) \times F(\text{da } 2 \text{ a } 4)$, oppure possono «incrociarsi», con ampiezza $F(\text{da } 1 \text{ a } 4) \times F(\text{da } 2 \text{ a } 3)$. Le due ampiezze vanno ora sommate e ciò produce interferenza (come si è visto nella seconda lezione), con conseguente variazione della lunghezza della freccia finale, in dipendenza della posizione relativa dei punti nello spazio-tempo.

Cosa succede se si portano a coincidere nello spazio-tempo i punti 3 e 4? (fig. 71). Supponiamo che entrambi i fotoni giungano nel punto 3, e vediamo come ciò influenzi le probabilità dell'evento. In questo caso $F(\text{da } 1 \text{ a } 3) \times F(\text{da } 2 \text{ a } 3)$ e $F(\text{da } 2 \text{ a } 3) \times F(\text{da } 1 \text{ a } 3)$ risultano ampiezze identiche; la loro somma ha quindi lunghezza doppia di quella di ciascuna, e il quadrato della freccia finale è quattro volte quello di ciascuna freccia separatamente. Le due frecce, essendo identiche, sono sempre «allineate». In altre parole, l'interferenza non fluttua al variare della separazione relativa tra i punti 1 e 2: è sempre positiva. Se non si tenesse presente questa interferenza sempre positiva tra i due fotoni, si sarebbe portati a pensare che in media si ha probabilità doppia. Invece la probabilità è sempre quadrupla. Se sono coinvolti parecchi fotoni la probabilità aumenta ulteriormente, rispetto a quanto ci si aspetterebbe.

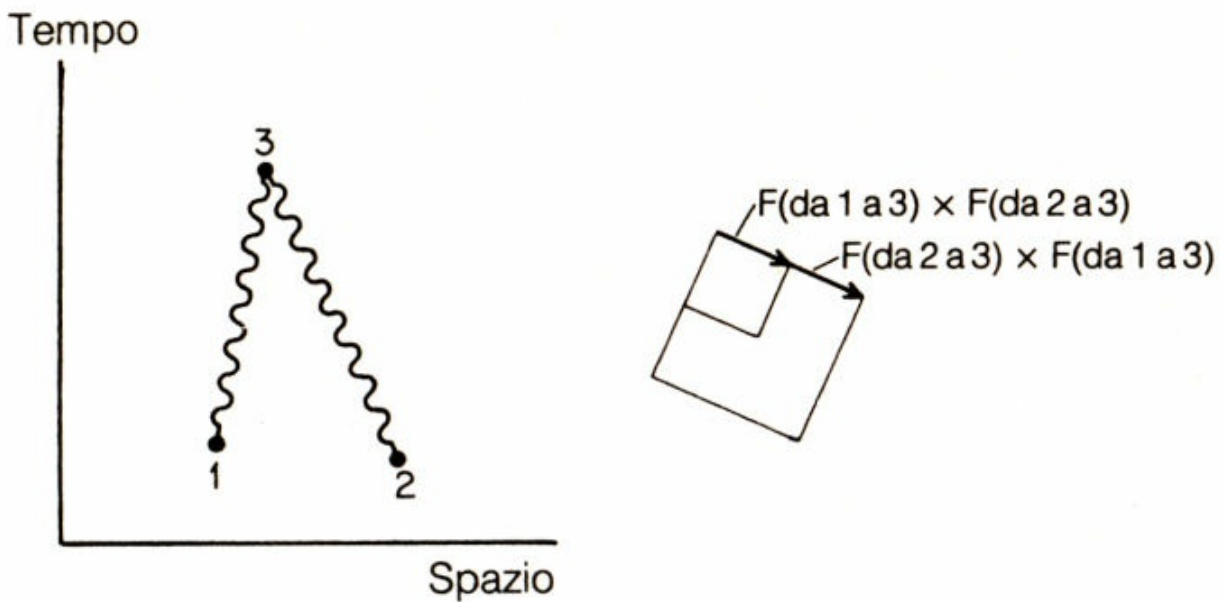


Fig. 71. Portando a coincidere i punti 3 e 4, le due frecce $F(\text{da } 1 \text{ a } 3) \times F(\text{da } 2 \text{ a } 3)$ e $F(\text{da } 2 \text{ a } 3) \times F(\text{da } 1 \text{ a } 3)$ diventano identiche in lunghezza e in direzione. Quando le si somma, essendo sempre «allineate», danno luogo a una freccia lunga il doppio di ciascuna di esse, il cui quadrato è quindi quattro volte maggiore. Pertanto i fotoni tendono a trovarsi nello stesso punto spaziotemporale. L'effetto è ancora più amplificato se sono coinvolti più fotoni. Questa proprietà sta alla base del funzionamento del laser.

Questo fatto ha numerose conseguenze pratiche. In linguaggio tecnico si dice che i fotoni tendono a stare nello stesso «stato», termine che indica in che modo l'ampiezza per trovare un fotone dipenda dalla posizione. La probabilità che un atomo emetta un fotone aumenta se è già presente qualche altro fotone (in uno stato nel quale l'atomo lo può emettere). Einstein scoprì questo fenomeno di «emissione stimolata» formulando il modello quantistico della luce basato sui fotoni. Il laser funziona sulla base di tale principio.

Ripetendo l'analogo conto per il nostro elettrone fittizio di spin zero, si troverebbe lo stesso comportamento. Ma con gli elettroni del mondo reale, che hanno polarizzazione, accade qualcosa di completamente diverso: le due frecce $E(\text{da } 1 \text{ a } 3) \times E(\text{da } 2 \text{ a } 4)$ e $E(\text{da } 1 \text{ a } 4) \times E(\text{da } 2 \text{ a } 3)$ vanno sottratte fra loro, cioè una di esse viene ruotata di 180° prima della somma. Se i punti 3 e 4 coincidono le due frecce hanno stessa lunghezza e direzione, per cui nella sottrazione si elidono (fig. 72). Ciò significa che, a differenza dei fotoni, gli elettroni non amano trovarsi nello stesso posto e si evitano l'un l'altro come la peste: due elettroni con la stessa polarizzazione non potranno mai stare nello stesso punto dello spazio-tempo. Tale proprietà è nota come «principio di esclusione».

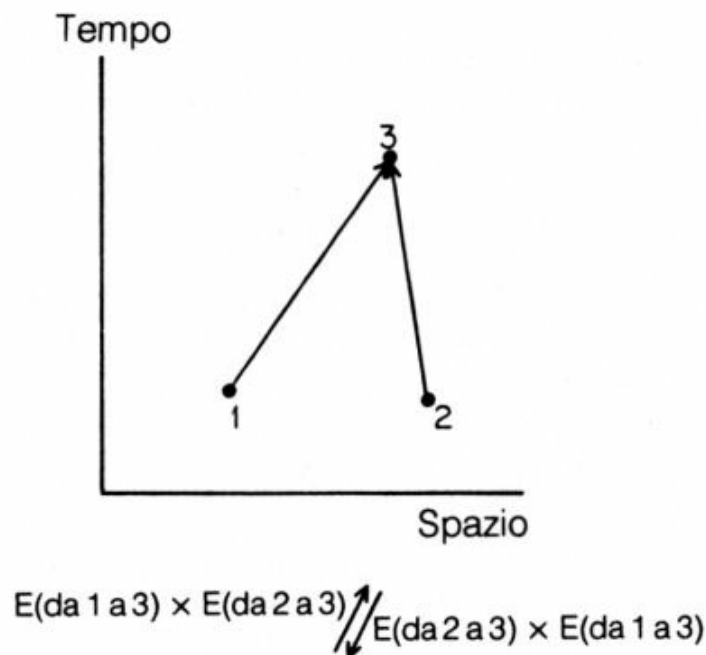


Fig. 72. Se due elettroni (con la stessa polarizzazione) tentano di andare nello stesso punto dello spazio-tempo, si ha sempre interferenza distruttiva per effetto della polarizzazione: le due frecce identiche $E(\text{da } 1 \text{ a } 3) \times E(\text{da } 2 \text{ a } 3)$ e $E(\text{da } 2 \text{ a } 3) \times E(\text{da } 1 \text{ a } 3)$ si sottraggono e si ottiene una freccia finale nulla. L'avversione degli elettroni a occupare la stessa posizione spaziotemporale viene chiamata «principio di esclusione» ed è la ragione della grande varietà degli atomi nell'universo.

Questo principio sta alla base dell'enorme varietà delle proprietà chimiche degli atomi. Un protone che scambia fotoni con un elettrone che gli danza intorno è chiamato atomo di idrogeno. Due protoni in uno stesso nucleo che scambiano fotoni con due elettroni di polarizzazione opposta vengono chiamati atomo di elio. I chimici hanno un loro modo speciale, e complicato, di contare: invece di dire «uno, due, tre, quattro, cinque protoni» dicono: «idrogeno, elio, litio, berillio, boro».

Vi sono solo due stati di polarizzazione possibili per l'elettrone, per cui in un atomo con tre protoni che scambiano fotoni con tre elettroni (situazione chiamata atomo di litio) il terzo elettrone si trova più lontano dal nucleo dei primi due (che si sono presi tutto lo spazio più vicino), e scambia meno fotoni. Perciò, sotto l'effetto di fotoni provenienti da altri atomi, esso si stacca con facilità dal proprio nucleo. Quando gli atomi di litio sono molti e vicini tra loro, questo terzo elettrone viene perduto facilmente e insieme con gli altri elettroni perduti forma un mare di elettroni fluttuante tra un atomo e l'altro e pronto a reagire a qualunque forza elettrica anche piccola (fotoni), producendo una corrente di elettroni: questa è la descrizione del metallo litio, buon conduttore elettrico. Gli atomi di idrogeno e quelli di elio non cedono i loro elettroni ad altri atomi, e sono «isolanti».

Tutti gli atomi, più di un centinaio di tipi diversi, sono fatti da un certo numero di protoni che scambiano fotoni con un egual numero di elettroni. Essi si dispongono secondo configurazioni complesse che determinano un'enorme gamma di proprietà: alcuni materiali sono metalli, altri isolanti; alcuni sono gassosi, altri cristallini; vi sono materiali morbidi e materiali duri, sostanze colorate e sostanze trasparenti; una meravigliosa ed entusiasmante cornucopia di varietà dovuta al principio di esclusione e alla ininterrotta ripetizione di tre semplicissimi eventi: $F(\text{da } A \text{ a } B)$, $E(\text{da } A \text{ a } B)$ e j . (Se gli elettroni reali non avessero polarizzazione, tutti gli atomi avrebbero proprietà molto simili: gli elettroni si addenserebbero insieme, vicini al nucleo del proprio atomo, e non sarebbero attratti facilmente dagli altri atomi per dare luogo a reazioni chimiche).

Vi chiederete come è possibile che tre eventi così semplici portino a un mondo così complesso. Il fatto è che i fenomeni che vediamo intorno a noi sono il risultato del complicatissimo intrecciarsi di un immenso numero di scambi e di interferenze di fotoni. La conoscenza dei tre eventi elementari costituisce solo un primo piccolissimo passo nello studio di qualunque fenomeno *reale*, dove gli scambi di fotoni sono tali e tanti da rendere impossibile ogni tentativo di calcolo; occorre rendersi conto tramite l'esperienza di quali siano le possibilità dominanti. Perciò si introducono grandezze quali «indice di rifrazione», «compressibilità», «valenza», che permettono di fare calcoli approssimati trascurando un'enorme quantità di dettagli. È un po' come conoscere solo le regole del gioco degli scacchi, che sono semplici e fondamentali, rispetto al saper giocare bene, che significa comprendere l'importanza di ogni posizione e la natura delle varie situazioni, ed è molto più difficile e complesso.

I settori della fisica che cercano di scoprire come mai il ferro (che ha 26 protoni) è magnetico mentre il rame (che ne ha 29) non lo è, o perché un gas è trasparente e un altro no, sono chiamati «fisica dello stato solido» o «fisica dello stato liquido», o semplicemente «vera fisica». Il settore che ha trovato quei tre piccoli eventi semplici (la parte più facile) viene chiamato «fisica fondamentale», nome di cui ci siamo impadroniti per mettere a disagio gli altri fisici! I problemi oggi più interessanti, e certamente più rilevanti da un punto di vista pratico, sono ovviamente quelli dello stato solido. Ma è stato detto che non c'è niente di più pratico di una buona teoria, e l'elettrodinamica quantistica è decisamente una buona teoria!

Per finire, torniamo al numero 1,00115965221, così scrupolosamente calcolato e misurato, come vi ho detto nella prima lezione. Questo numero rappresenta la risposta di un elettrone a un campo magnetico esterno, ed è chiamato «momento magnetico». Le leggi per calcolarlo furono discusse per la prima volta da Dirac, il quale, in base all'espressione di E (da A a B), ottenne un risultato molto semplice che, in opportune unità, si può porre uguale a 1. A questo calcolo approssimato del momento magnetico dell'elettrone corrisponde un diagramma molto semplice: nell'andare da un punto a un altro dello spazio-tempo un elettrone interagisce con un fotone dovuto a un magnete (fig. 73).

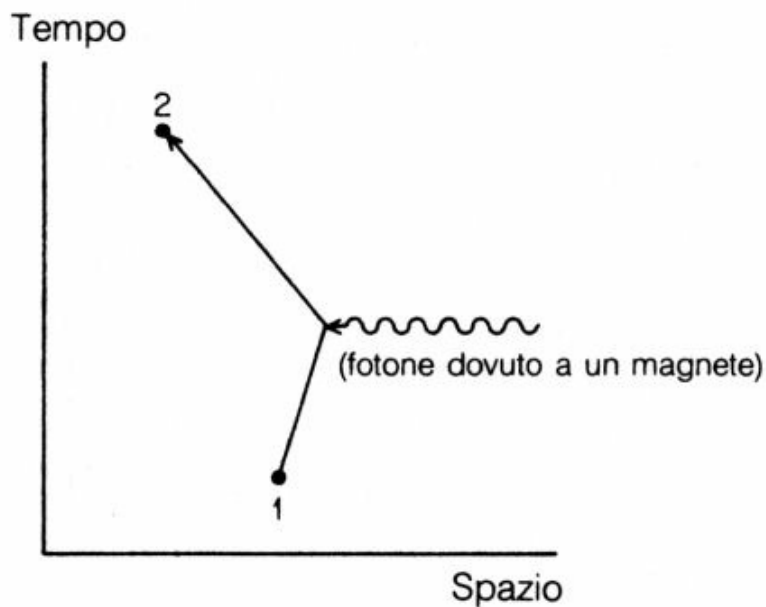


Fig. 73. Il diagramma che rappresenta il calcolo di Dirac del momento magnetico dell'elettrone è molto semplice. Il risultato dei conti rappresentato da questo diagramma verrà posto uguale a 1.

In seguito a esperimenti più precisi si trovò che il valore sperimentale non è esattamente 1, ma leggermente maggiore, vicino a 1,001116. Il valore teorico della correzione venne calcolato per la prima volta nel 1948 da Schwinger, che ottenne $j \times j$ diviso 2π prendendo in considerazione un secondo modo in cui l'elettrone può andare da un punto a un altro: invece di andarvi direttamente, l'elettrone si propaga per un tratto, poi all'improvviso emette un fotone e infine (orrore!) assorbe il fotone da lui stesso emesso (fig. 74): comportamento forse «immorale», ma assolutamente reale! Per calcolare l'ampiezza per questa seconda possibilità, si deve considerare una freccia per ogni punto dello spazio-tempo in cui il fotone può essere emesso e per ogni punto in cui può essere assorbito. Perciò l'espressione conterrà, come fattori moltiplicativi, due termini E (da A a B), un termine F (da A a B) e due termini j in più. Gli studenti di fisica imparano come fare questo semplice conto nel corso introduttivo di elettrodinamica quantistica al secondo anno di specializzazione.

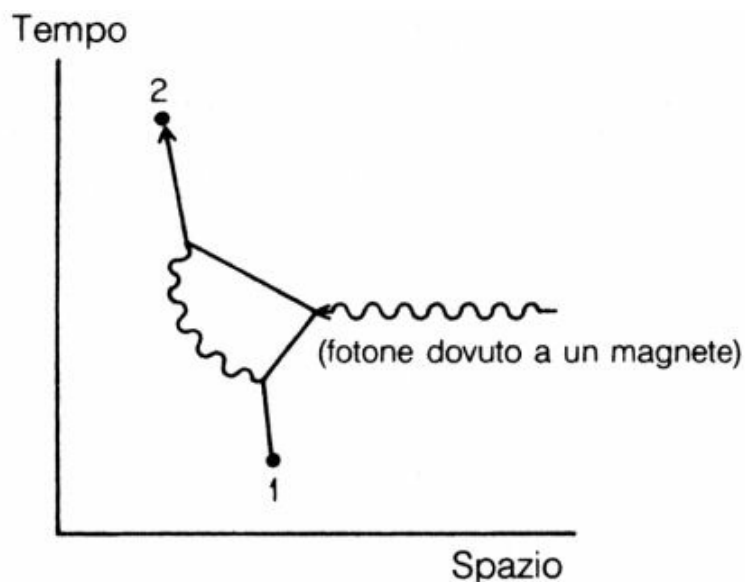


Fig. 74. Gli esperimenti di laboratorio mostrano che il valore reale del momento magnetico dell'elettrone non è 1, ma è leggermente maggiore. Ciò succede perché vi sono altre possibilità: ad esempio, l'elettrone può emettere un fotone e poi riassorbirlo, il che richiede due E (da A a B), una F (da A a B) e due j in più. La correzione che tiene conto di questa possibilità fu calcolata da Schwinger, che trovò $j \times j$ diviso 2π . Questa alternativa è sperimentalmente indistinguibile dal modo discusso prima, in cui l'elettrone va dal punto 1 al punto 2;

pertanto le frecce per le due possibilità vanno sommate e si ha interferenza.

Ma non basta: il comportamento di un elettrone in campo magnetico è stato misurato così accuratamente che nel calcolo teorico si devono prendere in considerazione ancora altre possibilità, ad esempio tutti i modi in cui l'elettrone va da un punto a un altro con *quattro* ulteriori interazioni (fig. 75). Ci sono tre modi in cui l'elettrone può emettere e riassorbire due fotoni. C'è inoltre una nuova possibilità interessante, mostrata sulla destra della fig. 75: l'elettrone emette un fotone che si trasforma in una coppia elettrone-positrone, i quali (con buona pace delle vostre obiezioni «moralì») successivamente si annichilano tra loro, dando luogo a un nuovo fotone che infine viene riassorbito dall'elettrone iniziale. Anche queste sono eventualità di cui tener conto!

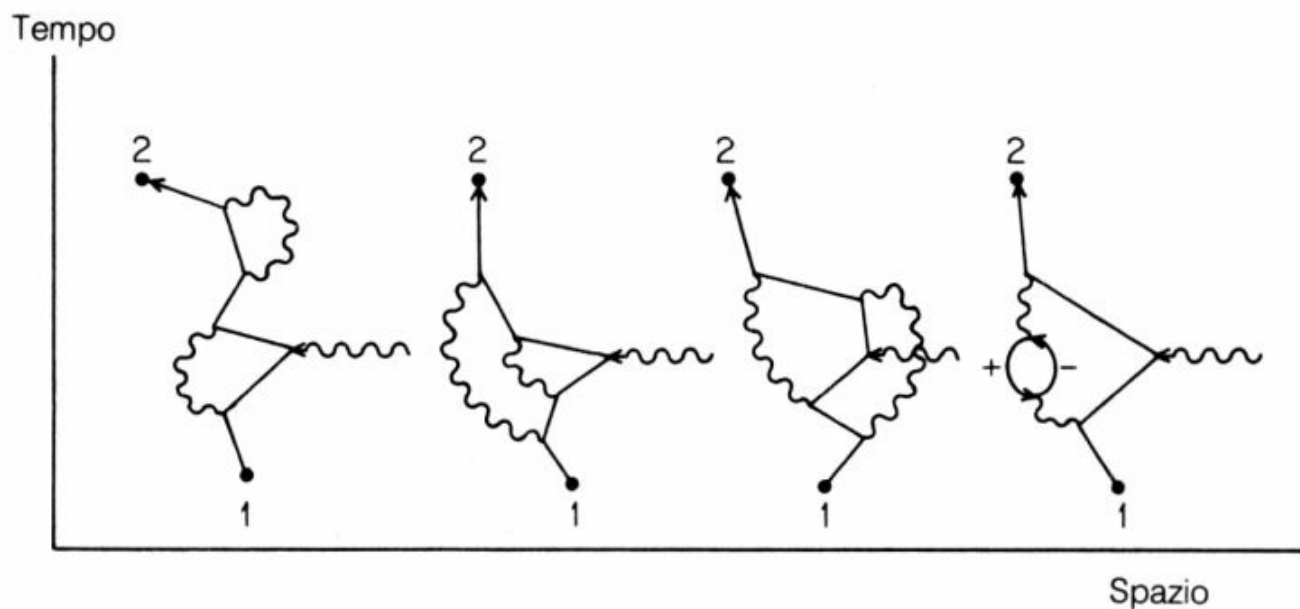


Fig. 75. Gli esperimenti di laboratorio sono diventati così accurati che si sono dovute prendere in considerazione anche alternative che comportano quattro interazioni in più (in tutti i possibili punti nello spazio-tempo). Alcune di queste alternative sono mostrate in figura: quella sulla destra riguarda un fotone che si disintegra in una coppia elettrone-positrone (come descritto alla fig. 64), i quali a loro volta si annichilano producendo un nuovo fotone che infine viene assorbito dall'elettrone.

Ci sono voluti due anni perché due gruppi «indipendenti» di fisici completassero il calcolo di quest'ultimo termine, e poi un altro anno per accorgersi che nel conto c'era un errore. Il valore misurato sperimentalmente risultava leggermente diverso da quello calcolato, e sembrò per un certo periodo che per la prima volta la teoria non fosse in accordo con gli esperimenti; invece si trattava di un semplice errore di aritmetica. Com'è possibile che due gruppi «indipendenti» facciano lo stesso errore? Il fatto è che verso la fine del calcolo i due gruppi confrontarono i propri appunti ed eliminarono le differenze tra i conti. Insomma, erano indipendenti per modo di dire.

L'ampiezza con *sei j* in più comporta il calcolo di un numero di modi ancora maggiore; alcuni sono disegnati alla fig. 76. Ci sono voluti venti anni per calcolare con questo livello di precisione il valore del momento magnetico dell'elettrone. Nel frattempo gli esperimenti continuavano, diventando più accurati e permettendo di aggiungere nuove cifre al valore sperimentale, che continuava ad andare d'accordo con quello teorico.

Per fare questo tipo di calcoli occorre dunque disegnare i relativi diagrammi, scrivere a che cosa equivalgono matematicamente e sommare le corrispondenti ampiezze: un procedimento molto lineare, da «ricetta di cucina».

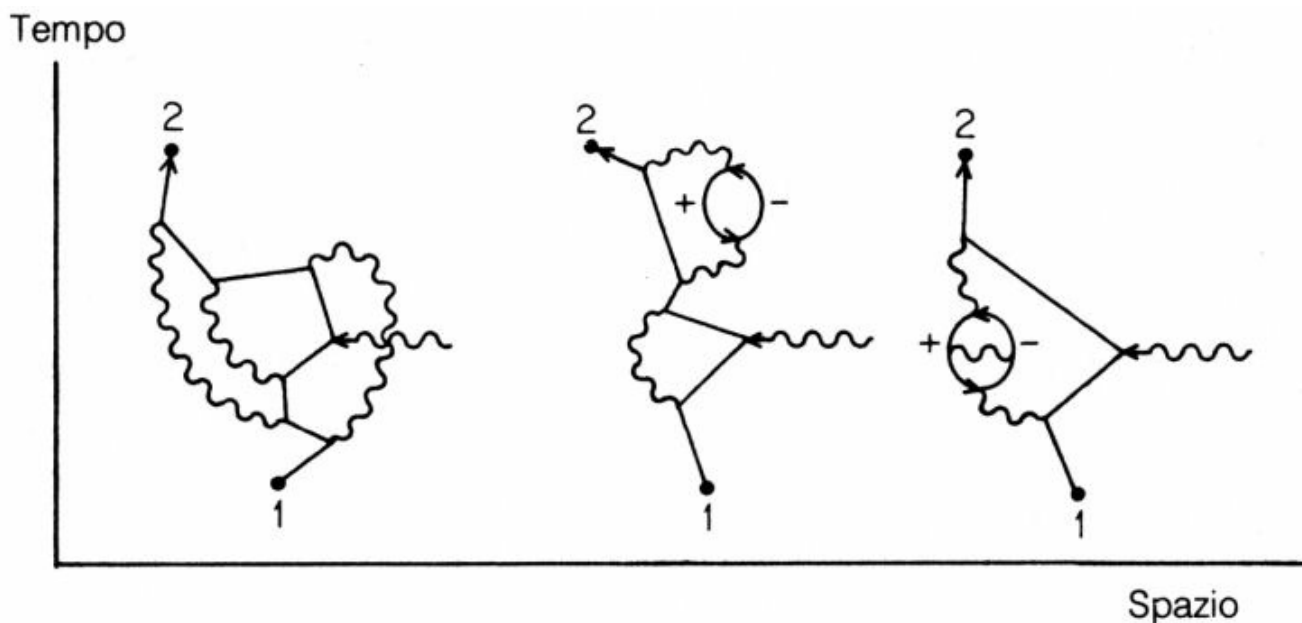


Fig. 76. Si stanno portando a termine i conti per rendere ancora più accurato il valore teorico del momento magnetico dell'elettrone. Il contributo successivo all'ampiezza, dovuto alle alternative con sei interazioni in più, comporta il calcolo di qualcosa come 10.000 diagrammi (ciascuno contenente 500 termini), tre dei quali vengono qui mostrati in figura. Nel 1983 il valore teorico era 1,00115965246, con un'incertezza di circa 20 sulle ultime due cifre, mentre il valore sperimentale era 1,00115965221, con un'incertezza di circa 4 sull'ultima cifra. Per avere un'idea di tale precisione, è come se la distanza tra New York e Los Angeles, di quasi 5.000 chilometri, fosse misurata con l'approssimazione di un capello umano.

Possiamo quindi farlo fare a una macchina. Con i fenomenali calcolatori di cui disponiamo oggi è stato iniziato il conto del termine con otto j in più. Attualmente il valore teorico è 1,00115965246, mentre quello sperimentale è 1,00115965221 con un'incertezza di più o meno quattro sull'ultima cifra decimale. Anche per il valore teorico c'è un'incertezza in parte dovuta all'arrotondamento dei numeri da parte del calcolatore, che produce un'indeterminazione di circa quattro sull'ultima cifra, ma la parte principale (circa 20) è dovuta alla imprecisa conoscenza del valore di j . Questo termine con otto j in più comporta il calcolo di circa diecimila diagrammi, ciascuno comprendente circa cinquecento termini: un conto davvero formidabile, e attualmente in corso.

Sono sicuro che fra alcuni anni il valore teorico e quello sperimentale del momento magnetico dell'elettrone saranno noti con precisione ancora maggiore. Naturalmente non so se i due valori saranno ancora in accordo tra loro. Questo non lo si può mai dire finché non viene fatto il conto e compiuto l'esperimento.

Siamo così tornati al punto da cui eravamo partiti. Avevo parlato di questo numero nella mia prima lezione per «fare colpo»; ora invece spero vi sia chiaro che cosa esso significhi: questo numero rappresenta l'estremo grado di accuratezza con cui abbiamo verificato e continuiamo a verificare l'esattezza della strana teoria dell'elettrodinamica quantistica.

In queste lezioni mi sono spesso divertito a far vedere come l'acquisizione di una teoria così precisa sia stata pagata con una erosione del comune buon senso e con la forzata accettazione di comportamenti molto bizzarri: l'aumento e la soppressione delle probabilità, la riflessione della luce in tutti i punti di uno specchio, la sua propagazione lungo percorsi che non sono linee rette, fotoni che si propagano con velocità maggiore o minore di quella convenzionale per la luce, elettroni che viaggiano all'indietro nel tempo, fotoni che scompaiono trasformandosi in coppie elettrone-positrone, e così via. Tutto ciò è necessario se si vuole capire quale sia il reale comportamento della Natura sotto la

superficie di quasi tutti i fenomeni che vediamo intorno a noi.

Tranne che per il dettaglio tecnico della polarizzazione, ho descritto lo schema nel quale noi inquadrriamo tutti questi fenomeni. Determiniamo le *ampiezze* per ogni modo in cui un evento può verificarsi e poi le sommiamo, laddove in circostanze ordinarie ci si sarebbe aspettati di sommare le probabilità; analogamente moltiplichiamo ampiezze, quando ci si sarebbe aspettati di moltiplicare probabilità. Pensare a qualunque fenomeno in termini di ampiezze può, all'inizio, causare qualche difficoltà dovuta al loro carattere astratto, ma dopo un po' ci si abitua alle stranezze di questo linguaggio. Alla base di un grandissimo numero di fenomeni quotidiani vi sono solo tre eventi fondamentali, uno dei quali è dato da un semplice numero, j , mentre gli altri due sono descritti dalle funzioni $F(\text{da } A \text{ a } B)$ e $E(\text{da } A \text{ a } B)$, molto simili tra loro. Da questi semplicissimi eventi elementari seguono tutte le altre leggi della fisica.

Ma prima di terminare questa lezione voglio fare alcune considerazioni. Per cogliere lo spirito e la fisionomia dell'elettrodinamica quantistica non è necessario includere la polarizzazione. Ma so che voi vi sentirete più tranquilli se vi dirò qualcosa anche su questo argomento. Orbene, i fotoni si presentano in quattro varietà differenti, chiamate polarizzazioni, e connesse da un punto di vista geometrico alle varie direzioni dello spazio-tempo. Ci sono cioè fotoni polarizzati lungo le direzioni X, Y, Z e T. (Avrete forse sentito dire che la luce ha solo due stati di polarizzazione; ad esempio, un fotone che viaggia lungo la direzione Z può essere polarizzato solo lungo le direzioni ortogonali, vale a dire X e Y. In situazioni in cui il fotone percorre lunghe distanze e sembra muoversi con la solita velocità della luce, le ampiezze relative alle polarizzazioni lungo Z e T, come avrete facilmente indovinato, si elidono esattamente tra loro. Ma per i fotoni virtuali, quali quelli scambiati tra un protone e un elettrone in un atomo, è proprio la componente T ad essere la più importante).

Analogamente un elettrone può trovarsi in uno tra quattro possibili stati, pure connessi alla geometria ma in modo alquanto più sottile. Chiamiamoli stati 1, 2, 3 e 4. Il calcolo dell'ampiezza per un elettrone che si propaga da un punto A a un punto B dello spazio-tempo diventa ora un po' più complesso, perché possiamo fare domande come: «Qual è l'ampiezza perché un elettrone partito dal punto A nello stato 2 arrivi nel punto B nello stato 3?». Le sedici possibili combinazioni, dovute ai quattro stati da cui l'elettrone può partire da A e ai quattro in cui può arrivare in B, sono correlate in modo semplice alla già vista formula per $E(\text{da } A \text{ a } B)$.

Per il fotone non è invece necessaria alcuna modifica, perché se parte da A polarizzato nella direzione X, arriverà in B ancora polarizzato lungo la stessa direzione con ampiezza $F(\text{da } A \text{ a } B)$.

La polarizzazione produce un gran numero di accoppiamenti diversi. Ad esempio, c'è un'ampiezza perché un elettrone nello stato 2 assorba un fotone polarizzato lungo X e passi nello stato 3. Non vi è accoppiamento tra tutte le possibili combinazioni della polarizzazione di elettroni e fotoni, ma quelle che si accoppiano lo fanno con ampiezza j , talvolta con un'ulteriore rotazione della freccia di un multiplo di 90° .

I diversi tipi di polarizzazione possibili e i relativi valori dell'accoppiamento possono essere ottenuti in maniera molto bella ed elegante dai principi dell'elettrodinamica

quantistica più le seguenti due ipotesi addizionali: 1) il risultato di un esperimento non cambia se l'apparato viene complessivamente ruotato in una direzione differente; 2) analogamente, non fa alcuna differenza se l'apparato si trova su un'astronave in moto con velocità arbitraria uniforme (questo non è altro che il principio di relatività).

Questa elegante analisi generale mostra che ogni particella deve appartenere a una classe di polarizzazione. Le varie classi si comportano in modo differente e vengono chiamate spin 0, spin $1/2$, spin 1, spin $3/2$, spin 2 e così via. Il comportamento più semplice è quello delle particelle di spin 0 che hanno una sola componente, cioè non hanno alcuna polarizzazione. (Gli elettroni e i fotoni fittizi considerati in queste lezioni sono stati trattati come particelle di spin 0. Tuttavia finora non è stata trovata alcuna particella fondamentale di spin 0). L'elettrone reale costituisce un esempio di particella di spin $1/2$ e il fotone reale un esempio di particella di spin 1. Le particelle di spin $1/2$ e 1 hanno quattro componenti. Gli altri tipi ne avrebbero di più, ad esempio le particelle di spin 2 avrebbero 10 componenti.

La connessione tra relatività e polarizzazione è semplice ed elegante, come ho detto, qualità che certo non riuscirei a raggiungere nella mia spiegazione e che richiederebbe da sola un'altra conferenza. I dettagli legati alla polarizzazione, pur non essendo essenziali per cogliere lo spirito e la fisionomia dell'elettrodinamica quantistica, lo sono per il calcolo corretto di ogni processo reale e spesso hanno effetti rilevanti.

In queste lezioni ho parlato soprattutto di interazioni tra elettroni e fotoni a distanze molto piccole e in condizioni relativamente semplici, in cui intervengono solo poche particelle. Ma vorrei aggiungere qualche considerazione su come appaiono tali interazioni nel mondo macroscopico, dove avviene lo scambio di un numero estremamente elevato di fotoni. In queste situazioni il calcolo delle frecce diventa molto complicato.

Alcune situazioni non sono però particolarmente complesse da analizzare. Ad esempio, in alcuni casi l'ampiezza per l'emissione di un fotone da parte di una sorgente è indipendente dalla precedente emissione di altri fotoni. Ciò si verifica quando la sorgente è molto pesante (un nucleo atomico), o quando vi sono moltissimi elettroni che si muovono tutti nello stesso modo, come avviene nella corrente che va avanti e indietro nell'antenna di una stazione trasmittente o che circola nell'avvolgimento di un elettromagnete. In questi casi viene emesso un gran numero di fotoni, tutti dello stesso tipo. In tali condizioni l'ampiezza per l'assorbimento di un fotone da parte di un elettrone risulta indipendente dal fatto che vi sia stato o no assorbimento di altri fotoni da parte dello stesso o di altri elettroni. Perciò il comportamento dell'elettrone è completamente descritto dall'ampiezza per l'assorbimento di un fotone, che dipende solo dalla posizione spaziotemporale dell'elettrone. (Una situazione di questo genere viene descritta dai fisici usando parole comuni. Essi dicono che l'elettrone si muove in un campo esterno. La parola «campo» viene usata in fisica per descrivere grandezze il cui valore dipende dalla posizione spaziale e temporale. Un buon esempio è dato dalla temperatura dell'aria, che dipende da dove e da quando viene misurata). Se si tiene conto della polarizzazione, il campo può avere più componenti. (Così il campo della luce ha quattro componenti, corrispondenti alle ampiezze per assorbire un fotone con polarizzazione X, Y, Z o T, che in linguaggio tecnico sono chiamate potenziali, vettore e scalare. In fisica classica, a partire

da questi potenziali, si introducono altre grandezze, di uso più comodo, chiamate campi elettrico e magnetico).

In presenza di campi elettrici e magnetici variabili lentamente nel tempo, l'ampiezza relativa alla propagazione di un elettrone tra due punti molto distanti dipende dal percorso seguito. Come si è già visto nel caso della luce, la traiettoria più importante risulta quella per la quale le ampiezze relative ai percorsi vicini hanno tutte quasi la stessa direzione. Di conseguenza la particella non viaggia necessariamente in linea retta.

Questo ci riporta alla fisica classica, secondo la quale l'elettrone si muove attraverso campi di forza seguendo il percorso che rende minima una certa quantità. (I fisici chiamano «azione» tale quantità ed esprimono questa legge col nome di «principio di minima azione»). Questo è un esempio di come le leggi dell'elettrodinamica quantistica descrivono i fenomeni su scala macroscopica. Da qui si potrebbe andare in moltissime direzioni, che purtroppo non avremo tempo di esplorare in questa sede. Ho voluto solo ricordare che i fenomeni osservati nel mondo macroscopico e i bizzarri effetti che avvengono su scala microscopica sono entrambi prodotti dalla interazione di fotoni ed elettroni e, in ultima analisi, sono tutti descritti dalla teoria dell'elettrodinamica quantistica.

IV CONCLUSIONI INDEFINITE

Dividerò questa lezione in due parti. Dapprima parlerò di alcuni problemi collegati alla teoria della elettrodinamica quantistica, supponendo che il mondo sia fatto esclusivamente di fotoni ed elettroni. Poi discuterò le relazioni di questa teoria col resto della fisica.

L'aspetto che fa più impressione nell'elettrodinamica quantistica è la sua assurda descrizione dei fenomeni basata sulle ampiezze, che potrebbe far pensare a qualche difficoltà di fondo. Ma i fisici, che manipolano ampiezze da ormai più di cinquant'anni, si sono completamente abituati al loro uso. Inoltre tutte le nuove particelle e i nuovi fenomeni osservati si inquadrano perfettamente in questo schema concettuale, nel quale la probabilità di un evento è data dal quadrato di una freccia finale la cui lunghezza è determinata dalla combinazione di altre frecce in strani modi (con interferenze e così via). Dunque *non sussiste alcun dubbio di carattere sperimentale* relativamente a questo schema; si possono avere dubbi filosofici a iosa sul significato da dare alle ampiezze (ammesso che ne abbiano uno), ma la fisica è una scienza sperimentale e poiché tale schema da risultati in accordo con gli esperimenti, noi fisici, almeno finora, ne siamo più che soddisfatti.

In elettrodinamica quantistica esiste un insieme di problemi connessi al miglioramento dei metodi di calcolo della somma delle frecce. Sono state ideate varie tecniche da usarsi a seconda delle circostanze e che richiedono agli studenti tre o quattro anni di studio per essere padroneggiate. Trattandosi di dettagli tecnici, non ne discuterò in queste lezioni; vi dico solo che si tratta di migliorare continuamente i metodi per ottenere le previsioni della teoria applicata alle varie situazioni.

Ma c'è un ulteriore problema, caratteristico dell'aspetto quantistico, che ha richiesto venti anni per essere risolto. Esso ha a che fare con gli elettroni e i fotoni ideali, e con i numeri n e j .

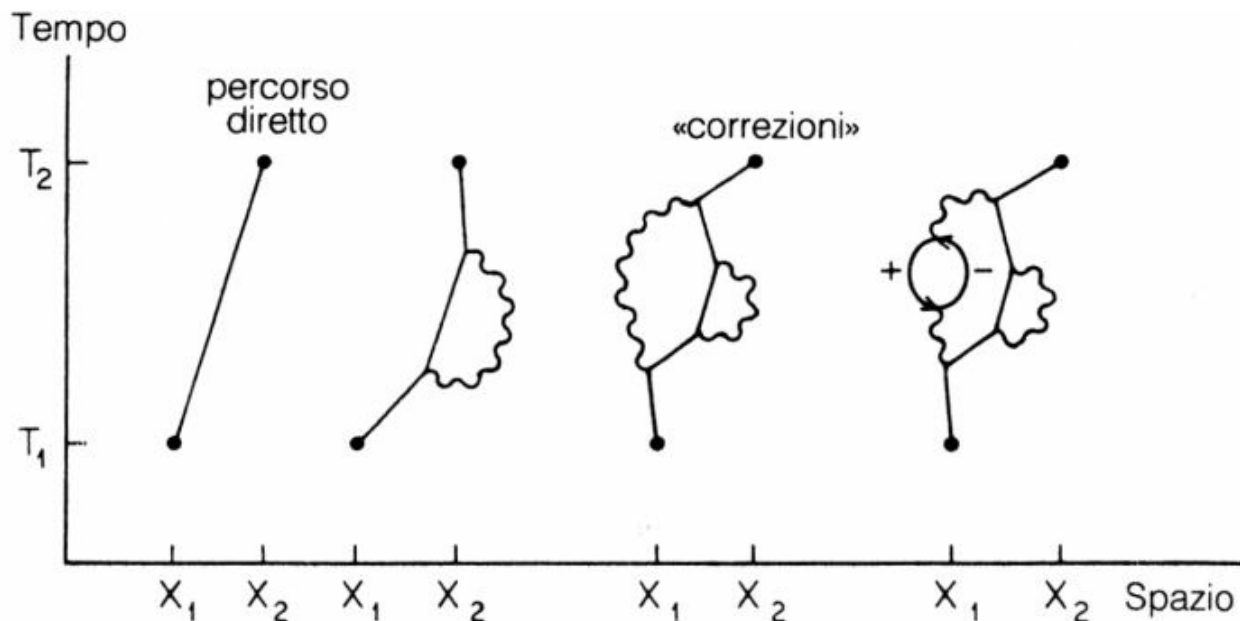


Fig. 77. Per calcolare l'ampiezza che ha un elettrone di propagarsi da un punto a un altro dello spazio-tempo si usa l'espressione di $E(\text{da A a B})$ per il percorso senza interazioni. (Poi si introducono le «correzioni», che consistono nell'emissione e assorbimento di uno o più fotoni). $E(\text{da A a B})$ dipende da $(X_2 - X_1)$, da $(T_2 - T_1)$ e da n , un numero che viene inserito nella formula per ottenere il risultato giusto. Questo numero è chiamato «massa a riposo» dell'elettrone «ideale» e non può essere misurato sperimentalmente perché la massa a riposo di un elettrone reale, m , dipende da tutte le «correzioni». Il calcolo del valore di n da usarsi in $E(\text{da A a B})$ comporta una difficoltà che ha richiesto venti anni per essere superata.

Se gli elettroni fossero ideali e andassero da un punto a un altro dello spazio-tempo seguendo *solo* percorsi diretti (come quello mostrato sulla sinistra della fig. 77), non ci sarebbero problemi: n sarebbe semplicemente la massa dell'elettrone, determinabile mediante l'osservazione, e j sarebbe la sua «carica», cioè l'ampiezza per la sua interazione con un fotone, anch'essa determinabile sperimentalmente.

Ma gli elettroni ideali non esistono e quelli reali di tanto in tanto emettono e assorbono i propri fotoni; perciò, la massa misurata in laboratorio dipende da j , l'ampiezza di interazione con i fotoni. Analogamente la carica osservata riguarda l'interazione tra elettrone reale e fotone reale; quest'ultimo di tanto in tanto si trasforma in una coppia elettrone-positrone, per cui j dipende da $E(\text{da A a B})$ che a sua volta dipende da n (fig. 78). Poiché massa e carica degli elettroni dipendono da tutti i possibili modi di propagazione, i valori sperimentalmente osservati della massa m e della carica e dell'elettrone sono differenti dai numeri n e j introdotti per fare i calcoli.

Se ci fosse una relazione matematica definita tra n e j da una parte, e m ed e dall'altra, non ci sarebbe alcun problema: si dovrebbe semplicemente calcolare da quali valori di n e j partire per ottenere i valori osservati di m e di e . (E se dal calcolo non risultassero i valori corretti di m e di e , si manipolerebbero i valori iniziali di n e j fino a ottenere i valori misurati).

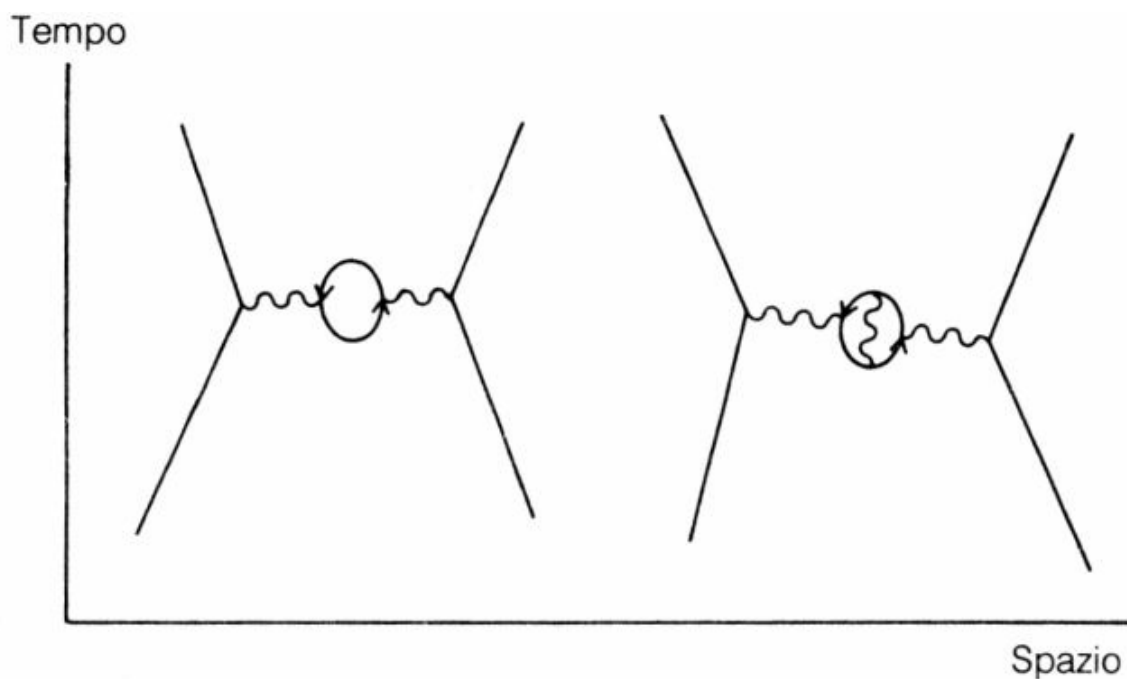


Fig. 78. Il valore sperimentale dell'ampiezza di accoppiamento tra elettrone e fotone, dato dal misterioso numero e , è determinato da esperimenti che includono tutte le possibili «correzioni» alla propagazione di un fotone da un punto a un altro dello spazio-tempo, due delle quali sono mostrate in figura. Per eseguire i calcoli si deve fare uso di un altro numero, j , che non include tali correzioni ma riguarda solo la propagazione diretta di un fotone da un punto a un altro. Per calcolare j si incontra una difficoltà simile a quella presente nel calcolo di n .

Vediamo ora come viene calcolato m in concreto. Si scrive una serie di termini simile a quella discussa per il calcolo del momento magnetico dell'elettrone. Il primo termine, che

non contiene accoppiamenti, è dato da $E(A \rightarrow B)$ e corrisponde a un elettrone ideale che si propaga da un punto a un altro dello spazio-tempo. Il secondo termine contiene due interazioni e rappresenta l'emissione con successivo riassorbimento di un fotone. Poi vi sono termini con quattro, sei, otto interazioni e così via (alcune di queste «correzioni» sono alla fig. 77).

Nel calcolare i termini con interazioni occorre, come sempre, considerare tutti i possibili punti nei quali può avvenire l'interazione, incluso il caso in cui i punti di emissione e di assorbimento si sovrappongono, cioè quando tra loro la distanza è nulla. Ma se si spinge il calcolo fino a distanza nulla, le espressioni ci «esplodono» tra le mani e danno risposte prive di senso, per esempio quantità infinite. Ciò ha causato un sacco di problemi nei primi tempi dell'elettrodinamica quantistica: qualunque grandezza si cercasse di calcolare risultava infinita! (La coerenza del procedimento matematico impone di considerare distanze nulle, ma proprio a questo punto insorge il problema, perché nessun valore di n o di j conduce a risultati sensati).

Ma se invece di includere tutti i possibili punti di interazione fino a distanza zero, si tronca il calcolo quando la distanza tra i punti è molto piccola, ad esempio 10^{-30} cm. (valore che è miliardi e miliardi di volte più piccolo di qualunque distanza osservabile sperimentalmente, il cui minimo è attualmente 10^{-16} cm.), esistono definiti valori di n e di j tali che la massa calcolata coincide col valore di m misurato sperimentalmente, e la carica calcolata coincide col valore sperimentale di e . E qui c'è l'inghippo: se si tronca il calcolo a una distanza minima diversa, ad esempio a 10^{-40} cm., per ottenere gli stessi m ed e si devono usare valori di n e di j differenti!

Dopo vent'anni di congetture, nel 1949 Hans Bethe e Victor Weisskopf si accorsero del seguente fatto: se due ricercatori impegnati a calcolare n e j a partire da valori identici di m ed e , dopo aver troncato i loro calcoli a distanze minime diverse, ottenendo così valori diversi per n e j , usavano poi questi valori per calcolare la risposta a qualche altro problema, ottenevano, una volta prese in considerazione le frecce relative a tutti i possibili termini, un risultato praticamente identico! Anzi, più vicine a zero erano le distanze a cui venivano troncati i calcoli di n e j , più vicine tra loro risultavano le risposte all'altro problema! Schwinger, Tomonaga e io stesso abbiamo inventato indipendentemente metodi di calcolo che confermano questo fatto (vincendo per questo anche dei premi). Finalmente era possibile fare calcoli sensati con la teoria dell'elettrodinamica quantistica!

Sembra dunque che le *uniche* grandezze che dipendono dalla distanza minima tra i punti di interazione sono i valori di n e j , *parametri teorici non direttamente osservabili*, mentre ogni grandezza che può essere misurata ne risulta indipendente.

Questo «gioco di bussolotti» fatto con i valori di n e di j viene chiamato in linguaggio tecnico «rinormalizzazione»: bel nome altisonante per quello che resta pur sempre un procedimento assurdo! L'aver dovuto ricorrere a simili giochi di prestigio ha reso impossibile provare la coerenza interna dell'elettrodinamica quantistica. In effetti è sorprendente che tale coerenza sia tuttora indimostrata e personalmente sospetto che la rinormalizzazione non sia un procedimento matematicamente legittimo. Quello che è certo è che non abbiamo una base di buona matematica per formulare la teoria

dell'elettrodinamica quantistica; non sono buona matematica tutti i giri di parole necessari per descrivere la connessione tra n e j con m ed e . ²⁶

Un interrogativo di grande profondità e bellezza ci viene proposto dal valore osservato della costante di accoppiamento e , cioè dall'ampiezza per emissione o assorbimento di un fotone reale da parte di un elettrone reale. Il valore di tale semplice numero determinato sperimentalmente risulta prossimo a $-0,08542455$. (È un numero che i miei amici fisici non riconosceranno, perché preferiscono ricordarlo come l'inverso del suo quadrato che vale $137,03597$ con un'incertezza di circa 2 sull'ultima cifra decimale. Questo numero costituisce un vero rompicapo fin da quando fu scoperto, oltre cinquant'anni fa, e tutti i migliori fisici teorici lo tengono incorniciato e appeso al muro e ogni giorno ci meditano su).

Vi chiederete subito da dove venga questo valore della costante di accoppiamento: è connesso a π , o magari alla base dei logaritmi naturali? Nessuno lo sa. È uno dei più enigmatici enigmi della fisica, un numero magico che ci viene offerto nel mistero più assoluto. Si potrebbe quasi dire che a scrivere questo numero sia stata la «mano di Dio» e che noi «non sappiamo come Egli abbia mosso la sua matita». Sappiamo perfettamente che cosa fare sperimentalmente per avere una misura accuratissima di questo valore, ma non sappiamo che arzigogolo inventare per farlo venir fuori da un calcolatore, senza avercelo messo dentro di nascosto!

Una buona teoria direbbe qualcosa come: e è uguale a 1 diviso 2π per la radice quadrata di 3. I suggerimenti in questo senso non sono mancati, ma si sono rivelati tutti inutili alla prova dei fatti. Arthur Eddington dimostrò su base puramente logica che il numero cercato dai fisici doveva essere esattamente 136, valore sperimentale di e a quell'epoca. Quando esperimenti più accurati diedero risultati più vicini a 137, Eddington scoprì un piccolo errore nel suo argomento precedente e dimostrò, di nuovo su base puramente logica, che il numero doveva essere l'intero 137! Ogni tanto qualcuno osserva che una certa combinazione di π e di e (la base dei logaritmi naturali), di 2 e di 5 fornisce il misterioso valore della costante di accoppiamento, ma chi fa questo tipo di giochi con l'aritmetica non si rende ben conto di quanti siano i numeri che possono essere riprodotti con π , la base dei logaritmi naturali e così via. La storia della fisica moderna è piena di lavori di persone che hanno ottenuto il valore di e esatto per molte cifre decimali, tranne poi a trovarsi smentiti dalla successiva e più accurata misura.

Anche se oggi si deve ricorrere a un procedimento assurdo per calcolare j , è possibile che un giorno si trovi un'autentica relazione matematica tra j ed e . In tal caso il vero numero misterioso sarebbe j , mentre e sarebbe derivato da esso. Avremmo allora indubbiamente una nuova sfilza di lavori che ci insegnerebbero a costruire j «con le nostre mani», per così dire, mostrandoci come j sia dato da 1 diviso 4π o simili.

E con questo ho svelato tutti i problemi dell'elettrodinamica quantistica.

Nell'organizzare il piano di queste lezioni avevo deciso di limitarmi ad aspetti della fisica ben noti, discuterli a fondo e fermarmi lì. Ma ora che sono arrivato a questo punto, essendo un professore (ossia, per definizione, incapace di smettere di parlare al momento giusto), non posso fare a meno di dirvi qualcosa anche sul resto della fisica.

Debbo precisare immediatamente che le teorie del resto della fisica non fanno

neanche cosa siano i controlli scrupolosi cui è stata sottoposta l'elettrodinamica. Alcune cose che racconterò sono buone ipotesi, altre sono teorie solo parzialmente elaborate, altre ancora sono mere speculazioni. La loro presentazione apparirà perciò alquanto confusa rispetto alle lezioni precedenti: sarà incompleta e carente in molti dettagli. Tuttavia la struttura concettuale della QED serve come base eccellente per la descrizione dei fenomeni nel resto della fisica.

Comincerò con i protoni e i neutroni, che costituiscono i nuclei atomici. Quando vennero scoperti, furono ritenuti particelle elementari, nel senso che l'ampiezza per andare da un punto a un altro potesse essere descritta da E (da A a B) inserendo un opportuno valore di n , ma rapidamente ci si accorse che non era così. Ad esempio, il momento magnetico del protone, calcolato come per l'elettrone, dovrebbe essere vicino a 1, mentre il valore sperimentale risulta tutto diverso: 2,79! Perciò si capì ben presto che nell'interno del protone deve succedere qualcosa che non è spiegato dall'elettrodinamica quantistica. E il neutrone, che, se davvero fosse neutro, non dovrebbe avere interazione magnetica, ha un momento magnetico di circa $-1,93$! Da un pezzo, quindi, si sapeva che dentro il neutrone succedevano cose molto sospette.

C'era inoltre il problema di che cosa tenga insieme protoni e neutroni dentro il nucleo. Non poteva essere lo scambio di fotoni, perché le forze che tengono insieme il nucleo sono molto più intense di quelle elettriche. L'energia necessaria per rompere un nucleo è molto più grande di quella necessaria per strappar via un elettrone da un atomo, nella stessa proporzione in cui la bomba atomica è più distruttiva della dinamite: l'esplosione della dinamite è dovuta a una ridisposizione di elettroni, quella della bomba atomica a una ridisposizione di protoni e neutroni.

Nel tentativo di scoprire che cosa tiene insieme i nuclei, si sono fatti numerosi esperimenti in cui protoni di energia sempre crescente venivano mandati a sbattere contro dei nuclei. Ci si aspettava di osservare solo protoni e neutroni, ma ad energie sufficientemente elevate vennero fuori particelle mai viste prima. I primi furono i pioni, poi le particelle lambda, sigma, rho, e così via fino ad esaurire l'alfabeto. Si cominciò allora a indicare le particelle col valore della massa, come sigma 1190 o sigma 1386, ma era ormai chiaro che il numero delle particelle esistenti al mondo è senza fine, e che la loro varietà dipende solo dall'energia usata per frantumare il nucleo. Attualmente si conoscono più di quattrocento tipi di particelle, ed è inaccettabile che siano tutte elementari: sarebbe troppo complicato! ²⁷

Scienziati dal genio inventivo, come Murray Gell-Mann, uscirono quasi di senno nel tentativo di scoprire le leggi seguite da tutte queste particelle, ma all'inizio degli anni Settanta misero a punto la teoria quantistica delle interazioni forti, detta «cromodinamica quantistica», i cui protagonisti sono i «quark». Le particelle composte di quark si dividono in due classi: alcune, come i protoni e i neutroni, sono fatte da tre quark (e hanno l'orrendo nome di «barioni»); altre, come i pioni, sono costituite da un quark e da un antiquark (e sono chiamate «mesoni»).

particelle di spin 1/2		fotone
		0
nome eletttrone simbolo e 0,511 massa (MeV)	eletttrone e 0,511	-1
	quark d ~10	-1/3
	quark u ~10	+2/3
		accoppiamento

Fig. 79. Il nostro elenco di tutte le particelle esistenti si apre con quelle di «spin 1/2»: l'eletttrone, con massa 0,511 MeV, e i due quark di «sapore» u e d , entrambi con massa di circa 20 MeV. Elettroni e quark hanno «carica», interagiscono cioè coi fotoni con ampiezze date rispettivamente da -1 , $-1/3$ e $+2/3$, in unità della costante di accoppiamento $-j$.

A questo punto conviene avere una tabella delle particelle fondamentali attualmente note (fig. 79). Nel discuterne le proprietà comincerò dalle particelle che si propagano da un punto a un altro con ampiezza E (da A a B), modificata in modo da tener conto della polarizzazione, come nel caso dell'eletttrone; tali particelle sono dette di «spin 1/2». La prima è appunto l'eletttrone, che ha massa 0,511 in unità chiamate MeV, molto usate in fisica. ²⁸

Sotto l'eletttrone lascio uno spazio libero (che verrà riempito in seguito) e ancora sotto elenco due varietà di quark, dette d ed u . La massa di questi quark non è ben conosciuta; una stima ragionevole è di circa 20 MeV per entrambi. In realtà il fatto che il neutrone sia più pesante del protone sembra implicare, come si vedrà tra poco, che il quark d è un po' più pesante del quark u .

A fianco di ciascuna particella è riportata, in unità $-j$, la sua carica, cioè il valore della costante di accoppiamento col fotone cambiato di segno. Con questa scelta la carica dell'eletttrone risulta -1 , coerentemente con la convenzione iniziata da Benjamin Franklin e che non ci siamo mai scrollati di dosso. L'ampiezza dell'accoppiamento a un fotone è $-1/3$ per il quark d e $+2/3$ per il quark u . (Se Franklin avesse conosciuto i quark avrebbe almeno posto la carica di un eletttrone uguale a -3 !).

Ora, la carica di un protone è $+1$ e quella di un neutrone è zero. Giocando un po' con i numeri, si vede che, essendo fatti entrambi da tre quark, il protone deve contenere due u e un d , mentre il neutrone due d e un u (fig. 80).

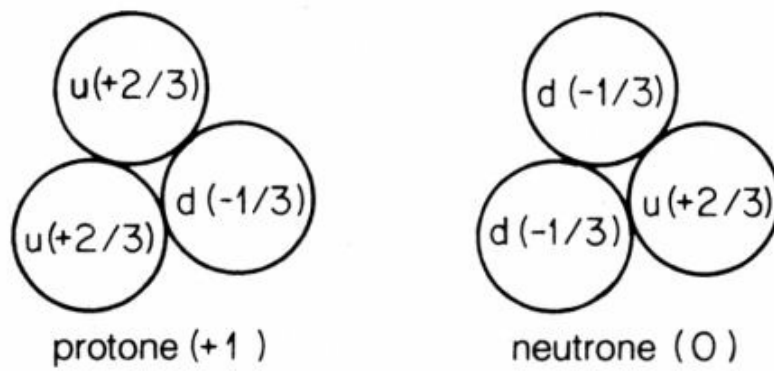


Fig. 80. Tutte le particelle composte da quark appartengono a una di due possibili classi: quelle fatte da un quark e da un antiquark, e quelle fatte da tre quark; protone e neutrone sono gli esempi più comuni di queste ultime. Le cariche dei quark u e d si combinano in modo da dare $+1$ per il protone e zero per il neutrone. Il fatto che protone e neutrone siano fatti di particelle cariche che si muovono al loro interno spiega perché il protone abbia un momento magnetico maggiore di 1 e il neutrone, pur elettricamente neutro, abbia un momento magnetico.

Che cosa mantiene uniti i quark? I fotoni scambiati? (È chiaro che, avendo il quark d carica $-1/3$ e il quark u $+2/3$, essi assorbono ed emettono fotoni come fa l'elettrone). No, le forze elettriche sono troppo deboli. Occorre introdurre qualche altra particella il cui scambio serva a mantenere uniti i quark. Tali particelle sono state chiamate «gluoni» ²⁹ e sono esempi di un altro tipo di particelle, dette di «spin 1», che comprende anche i fotoni; esse viaggiano da un punto a un altro con ampiezza data esattamente dalla stessa espressione che per i fotoni: $F(\text{da } A \text{ a } B)$. L'ampiezza per emissione o assorbimento di un gluone da parte di un quark è data da un altro numero misterioso, g , che risulta molto maggiore di j (fig. 81).

		particelle di spin 1	
		fotone	gluone
	particelle di spin 1/2	0	0
	elettrone e 0,511	-1	0
		0	0
nome	elettrone		
simbolo	e		
massa (MeV)	0,511		
	quark d ~10	-1/3	g
	quark u ~10	+2/3	g
		accoppiamenti	

Fig. 81. I «gluoni» mantengono uniti i quark che costituiscono i protoni e i neutroni, e indirettamente spiegano il fatto che protoni e neutroni restino uniti nei nuclei atomici. I gluoni mantengono uniti i quark con forze molto più intense di quelle elettriche. La costante di accoppiamento per i gluoni, g , è molto maggiore di j . Ciò rende molto più difficile il calcolo di termini che contengono interazioni: la migliore precisione che si riesce ad ottenere per ora è solo del 10%.

I nostri diagrammi che rappresentano lo scambio di gluoni tra quark sono molto simili a quelli che rappresentano lo scambio di fotoni tra elettroni (fig. 82). Così simili che verrebbe da sospettare noi fisici di mancanza di immaginazione: per descrivere le interazioni forti hanno copiato l'elettrodinamica quantistica! E così è, infatti: abbiamo proprio copiato, ma con una piccola variante.

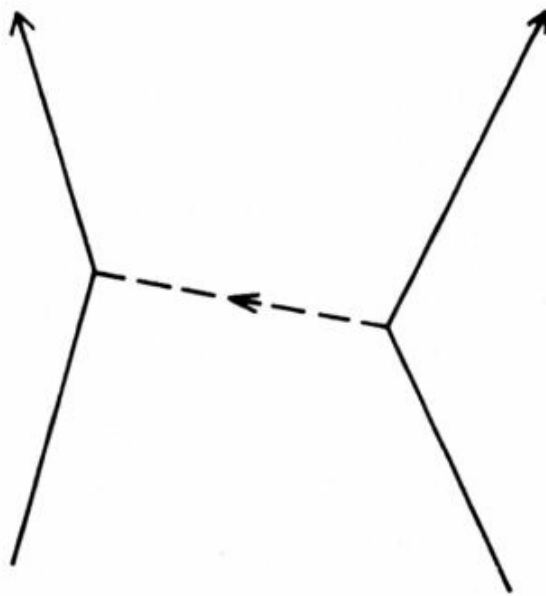


Fig. 82. Il diagramma del modo più semplice in cui due quark possono scambiarsi un gluone è talmente simile a quello che rappresenta lo scambio di un fotone tra due elettroni da far pensare che per descrivere le «interazioni forti», che tengono uniti i quark all'interno dei protoni e neutroni, i fisici abbiano semplicemente copiato la teoria dell'elettrodinamica quantistica. Ebbene, è proprio così - o quasi.

I quark hanno un ulteriore tipo di polarizzazione non legata alla geometria, e i fisici, da quegli sciagurati che sono, invece di darle un bel nome greco, l'hanno battezzata con l'infelice nome di «colore», che naturalmente non ha nulla a che vedere col colore nel senso usuale. Un quark può trovarsi in una fra tre condizioni, o «colori»: R, V o B (iniziali che ci vuol poco a interpretare!). Il «colore» di un quark può cambiare quando esso emette o assorbe un gluone. Questi ultimi esistono in otto varietà differenti, a seconda dei «colori» con cui possono accoppiarsi. Ad esempio, un quark rosso diventa verde emettendo un gluone rosso-antiverde, cioè un gluone che prende il rosso dal quark e gli cede il verde («antiverde» significa che il gluone trasporta il verde in direzione opposta). Un siffatto gluone può essere assorbito da un quark verde, che diventa rosso (fig. 83).

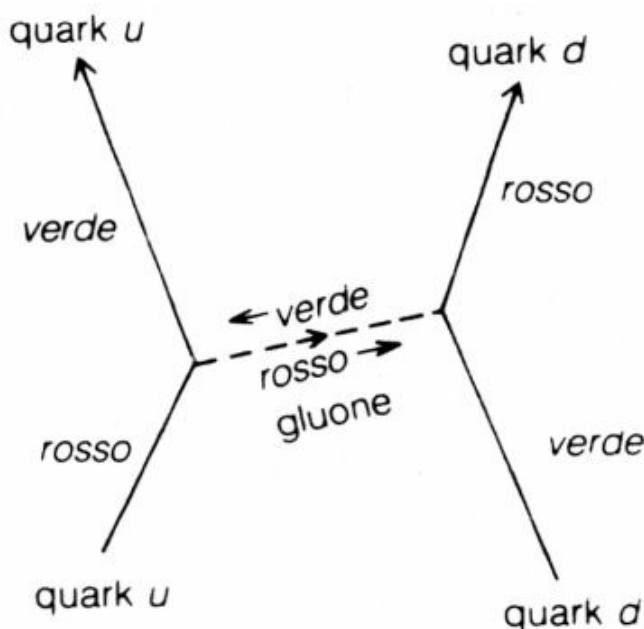


Fig. 83. La teoria che descrive i gluoni differisce dall'elettrodinamica quantistica in quanto i gluoni interagiscono con oggetti che portano un «colore» (i tre stati possibili sono «rosso», «verde» e «blu»). In figura un quark *u* rosso si trasforma in verde emettendo un gluone rosso-antiverde, assorbito poi da un quark *d* verde che diventa rosso. Il prefisso «anti» indica che il «colore» viene trasportato all'indietro nel tempo.

Ci sono otto possibili tipi di gluoni: rosso-antiroso, rosso-antiverde, rosso-antiblu e

così via (sono otto, e non nove come ci si aspetterebbe, perché per motivi tecnici uno è di troppo). La teoria non è eccessivamente complicata. La regola completa per i gluoni è la seguente: i gluoni interagiscono con oggetti «colorati»; si deve solo tenere un po' di contabilità per seguire le tracce del «colore».

Tale regola dà luogo a un'interessante possibilità: i gluoni possono interagire con altri gluoni (fig. 84). Ad esempio, un gluone rosso-antiverde incontrandone uno verde-antiblu si trasforma in gluone rosso-antiblu. La teoria dei gluoni è molto semplice: si devono tracciare i diagrammi possibili avendo cura di seguire il «colore». In qualunque diagramma l'intensità dell'interazione è data da g , costante di accoppiamento per i gluoni.

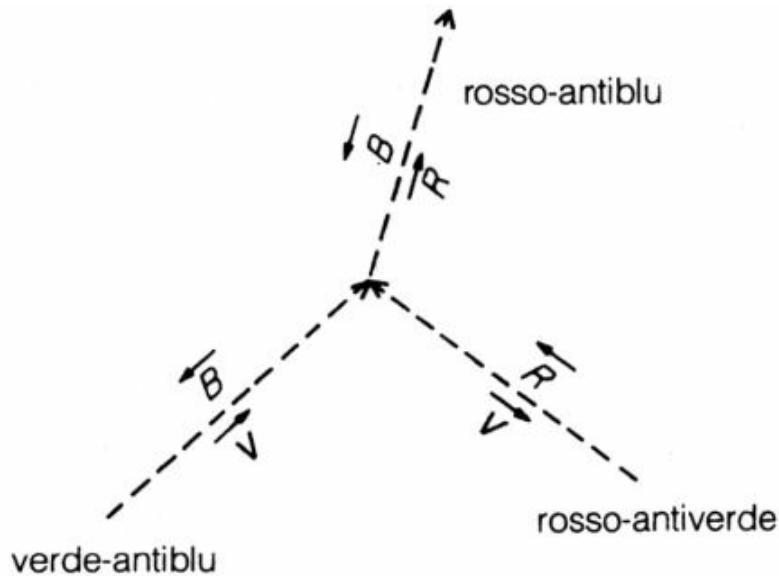


Fig. 84. I gluoni sono anch'essi «colorati», per cui possono interagire tra loro. In figura un gluone verde-antiblu interagisce con uno rosso-antiverde formando un gluone rosso-antiblu. La teoria dei gluoni è facile da capire: basta seguire i «colori».

Questa teoria non è dunque gran che differente dall'elettrodinamica quantistica; come regge il confronto con i dati sperimentali? Ad esempio, come si raffronta il valore osservato del momento magnetico del protone con quello calcolato dalla teoria?

I dati sperimentali, molto accurati, danno un valore di 2,79275. La teoria, invece, riesce a malapena a prevedere 2,7 più o meno 0,3, se si è sufficientemente ottimisti sull'accuratezza dei propri conti: si ha un errore del 10%, cioè una precisione 10.000 volte minore di quella dei dati sperimentali! Abbiamo una teoria semplice e definita, che dovrebbe spiegare tutte le proprietà dei protoni e dei neutroni, eppure non riusciamo a calcolare quasi niente perché la matematica risulta troppo complessa. (Avrete capito a che cosa sto lavorando in questo momento, senza peraltro riuscire a cavarne alcunché). Motivo di questa difficoltà è il valore di g , che risulta molto più grande della costante di accoppiamento tra elettroni e fotoni. I termini con due, quattro o anche sei interazioni non si riducono a piccole correzioni dell'ampiezza principale, ma rappresentano contributi notevoli che non possono essere trascurati. Si hanno pertanto frecce dovute a una miriade di possibilità, e non si è stati finora in grado di sistemarle in modo da trovare, con conti ragionevoli, quale sia la freccia finale.

Si legge nei libri che la scienza è semplice: si inventa una teoria, la si confronta coi dati sperimentali, e se non funziona la si butta e se ne inventa un'altra. Nel nostro caso abbiamo una teoria e abbiamo centinaia di risultati sperimentali, ma non possiamo fare il

raffronto! È una situazione del tutto senza precedenti nella storia della fisica. Siamo temporaneamente intrappolati sotto una valanga di freccette e incapaci di escogitare un metodo adatto per fare i conti.

Nonostante tali difficoltà nell'uso della teoria per fare i calcoli, grazie alla cromodinamica quantistica si ha una comprensione qualitativa di parecchi fenomeni delle interazioni forti tra quark e gluoni. Gli oggetti costituiti da quark che si osservano in natura risultano tutti di «colore» neutro: quelli fatti da tre quark ne contengono uno per ciascun «colore», mentre in quelli fatti da una coppia quark-antiquark vi è la stessa ampiezza perché la coppia sia rossa-antirossa, verde-antiverde o blu-antiblu. Si ha anche una spiegazione del perché non si osservano mai quark come particelle isolate, cioè del fatto che, quale che sia l'energia con cui un protone urta un nucleo, non si osservano mai fuoriuscire quark separati ma sempre fiotti di mesoni (coppie quark-antiquark) e barioni (gruppi di tre quark).

L'elettrodinamica e la cromodinamica quantistiche non esauriscono la descrizione dei fenomeni fisici. In base ad esse un quark non può cambiare «sapore»: un quark u resterà sempre u , e uno d resterà sempre d . Ma la Natura qualche volta si comporta in modo differente. Esiste un processo radioattivo molto lento, chiamato decadimento beta - sono le radiazioni di cui si teme sempre la fuga dai reattori nucleari - nel quale un neutrone si trasforma in protone. Poiché un neutrone è formato da due quark d e da uno u , mentre il protone contiene due u e un d , quel che avviene in realtà è che un quark di tipo d del neutrone si trasforma in quark di tipo u (fig. 85).

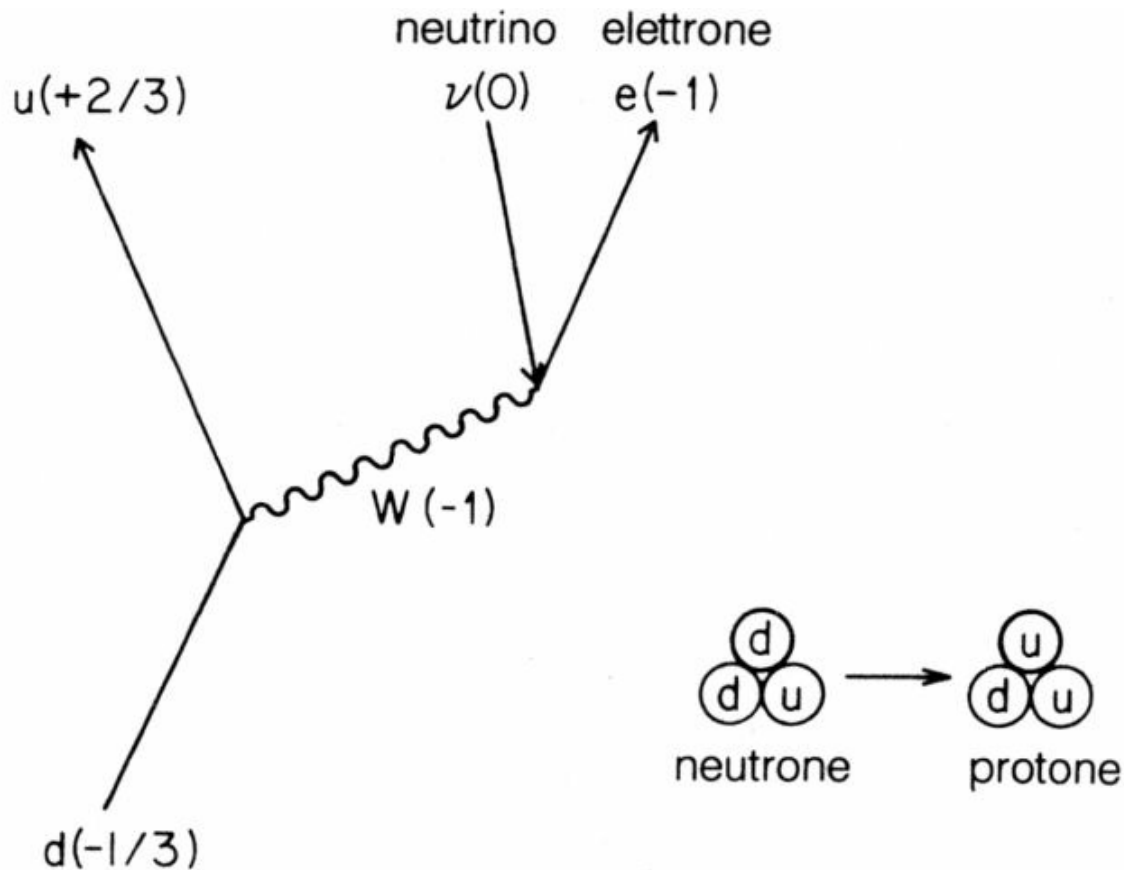


Fig. 85. Quando un neutrone decade in protone, processo chiamato «decadimento beta», l'unica cosa che cambia è il «sapore» di un quark, da d a u , mentre vengono emessi un elettrone e un antineutrino. Per spiegare la relativa lentezza di questo processo è stato proposto che esso avvenga tramite una particella intermedia, chiamata «bosone W intermedio», di massa molto alta (circa 80.000 MeV) e di carica -1.

Ciò si verifica nel modo seguente: il quark d si trasforma in u emettendo una particella tipo fotone, detto W , che si accoppia anche a un elettrone e a una nuova particella detta anti-neutrino, cioè un neutrino che va all'indietro nel tempo. Questo neutrino è una particella di spin $1/2$, come l'elettrone e il quark, ma non ha né massa né carica, per cui non interagisce coi fotoni. Non interagisce neanche coi gluoni, ma solo col W (fig. 86).

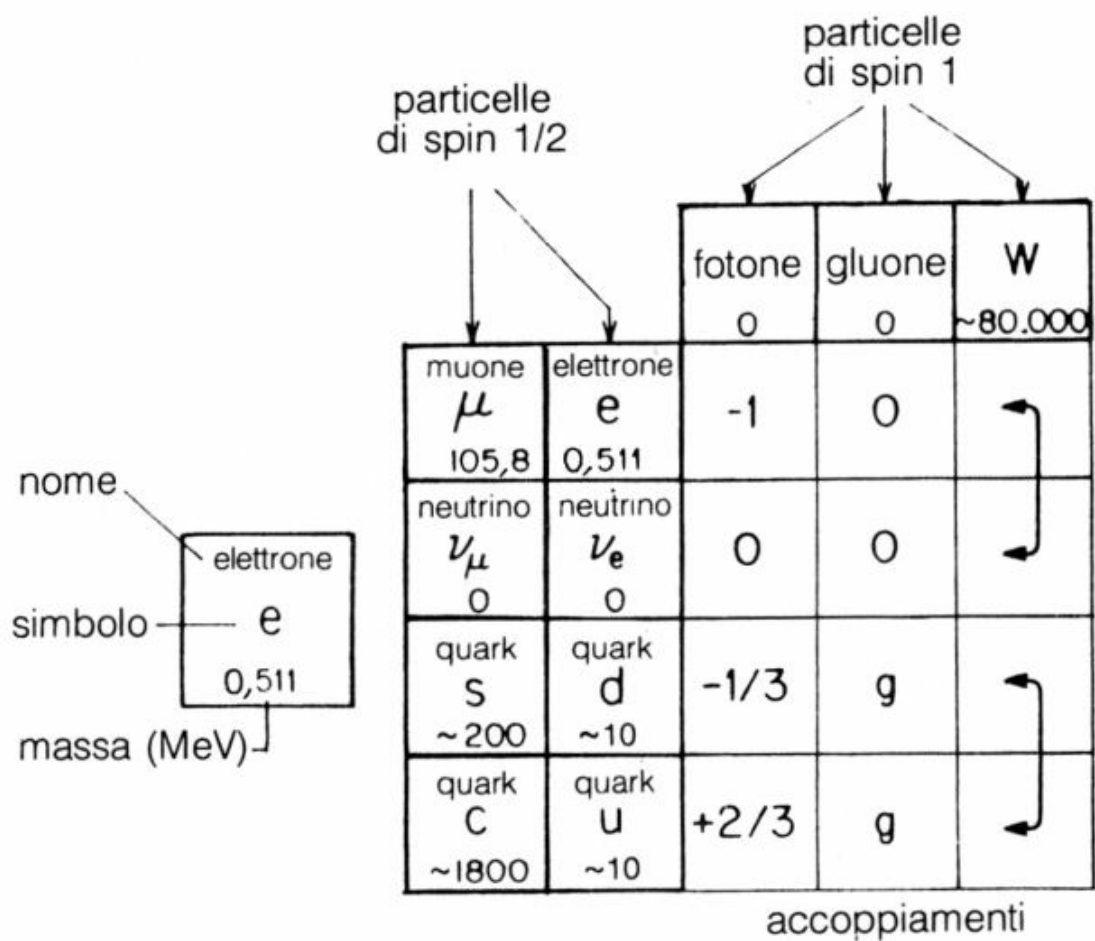


Fig. 86. I W si accoppiano con elettroni e/o neutrini a un estremo e con quark d e/o u all'altro.

I W sono particelle di spin 1, come i fotoni e i gluoni. Essi cambiano il «sapore» e la carica dei quark (il quark d , di carica $-1/3$, viene cambiato in quark u , di carica $+2/3$) ma non il loro «colore». Il W^- trasporta carica -1 e la sua antiparticella, W^+ , porta carica $+1$, per cui essi interagiscono anche coi fotoni. Il decadimento beta impiega un tempo molto maggiore di quello caratteristico per le interazioni tra fotoni ed elettroni, per cui si pensa che, a differenza di fotoni e gluoni, i W devono avere massa molto alta (circa 80.000 MeV). Finora non si è riusciti a vedere un W isolato perché la produzione di particelle di massa così grande richiede energie di collisione estremamente elevate.³⁰

C'è un altro tipo di particelle, chiamate Z_0 , che possono essere considerate come W neutri. Gli Z_0 non modificano la carica e possono interagire con quark u , quark d , elettroni e neutrini (fig. 87).

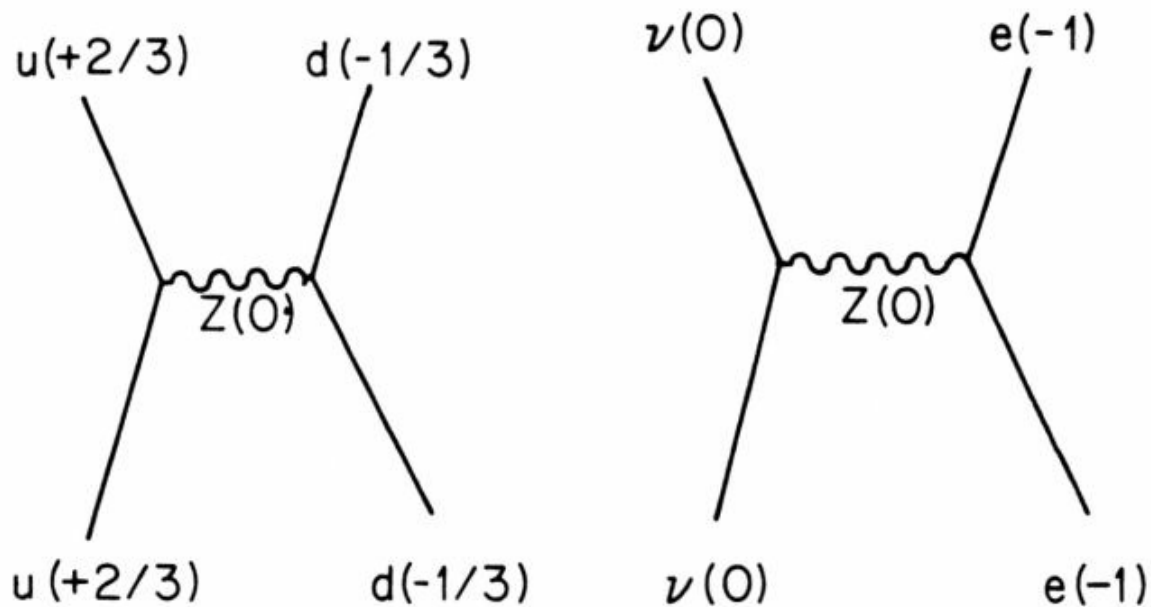


Fig. 87. Quando nessuna delle particelle esterne cambia carica, anche il W deve essere neutro, e in tal caso viene chiamato Z_0 . Questo tipo di interazione è detta «corrente neutra». In figura sono illustrate due possibilità.

Questo tipo di interazione ha avuto il nome improprio di «corrente neutra», e ha suscitato enorme interesse quando è stata osservata alcuni anni or sono.

La teoria dei W risulta molto bella e nitida se si ammette l'esistenza di un'interazione a tre fra i tre tipi di W (fig. 88).

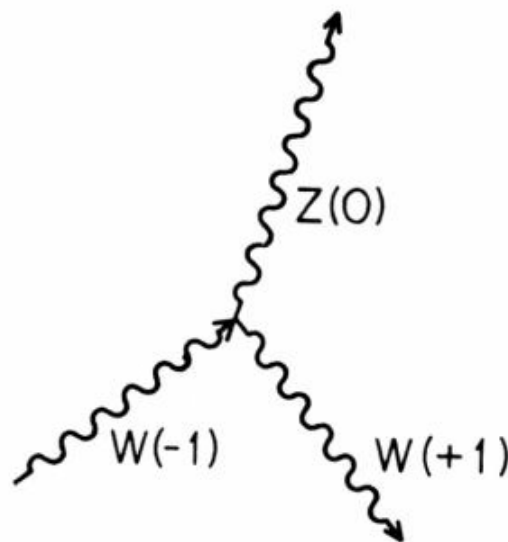


Fig. 88. È anche possibile l'interazione tra un W (-1), la sua antiparticella W (+1) e un W neutro (Z_0). La costante di accoppiamento dei W risulta vicina a j , il che suggerisce l'ipotesi che W e fotoni siano aspetti differenti di una stessa cosa.

La costante di accoppiamento per i W risulta molto vicina a quella dei fotoni, cioè a j . Esiste perciò la possibilità che W e fotoni siano aspetti differenti della stessa cosa. Stephen Weinberg e Abdus Salam sono riusciti a mettere insieme elettrodinamica quantistica e «interazioni deboli» (le interazioni dei W) in un'unica teoria quantistica. Ma nelle espressioni ottenute si vede ancora l'incollatura, per così dire. È evidente che fotoni e W sono in qualche modo connessi tra loro, ma all'attuale livello di comprensione la connessione non emerge con chiarezza; nella teoria si notano ancora delle «giunte», e solo quando saranno state eliminate il collegamento sarà più bello e perciò, probabilmente, più corretto.

Esistono dunque tre tipi di interazioni, tutte descritte da teorie quantistiche: le «interazioni forti» dei quark e gluoni, le «interazioni deboli» dei W e le «interazioni elettriche» dei fotoni. Secondo questo punto di vista le particelle fondamentali dell'universo sono i quark, che hanno «sapore» u o d , ciascuno con tre possibili «colori»; i gluoni, in otto combinazioni di R, V e B; i W, con carica 0 o $+1$; i neutrini, gli elettroni e i fotoni; in totale circa venti particelle di sei varietà differenti, più le loro antiparticelle. Non sarebbe male, ma c'è dell'altro.

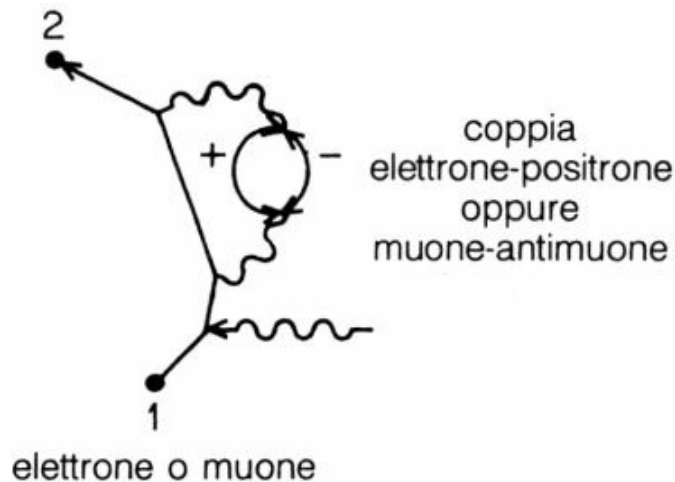


Fig. 89. Bombardando i nuclei con protoni di energia sempre più alta, appaiono via via nuove particelle. Una di queste è il muone, o elettrone pesante. Le interazioni del muone sono descritte dalla stessa teoria che descrive quelle dell'elettrone, tranne che si deve usare un valore di n più grande in E(da A a B). Il momento magnetico del muone dovrebbe risultare leggermente diverso da quello dell'elettrone, grazie alla seguente particolarità: quando un elettrone emette un fotone, questo può disintegrarsi in una coppia elettrone-positrone o muone-antimuone, la cui massa risulta uguale o maggiore di quella dell'elettrone iniziale. Se, invece, è un muone a emettere il fotone che decade in una coppia elettrone-positrone o muone-antimuone, la massa relativa alla coppia è minore o uguale a quella del muone. I dati sperimentali confermano l'esistenza di questa piccola differenza.

Via via che i nuclei venivano bombardati con protoni di energia crescente, venivano fuori nuove particelle. Una delle prime è stato il muone, che è assolutamente identico all'elettrone tranne per la massa molto più alta: 105,8 rispetto a 0,511 MeV. È come se Dio si fosse divertito a provare un diverso valore per la massa! Tutte le proprietà del muone sono correttamente descritte dall'elettrodinamica quantistica, con lo stesso valore di j e la stessa ampiezza E(da A a B), ma con un diverso valore di n .³¹

Avendo il muone massa circa 200 volte maggiore dell'elettrone, la «lancetta del cronometro» associato al muone gira circa 200 volte più rapidamente che per l'elettrone. Ciò ha permesso di verificare che la teoria descrive correttamente i fenomeni elettromagnetici a distanze 200 volte più piccole di prima, distanze peraltro che sono più di otto ordini di grandezza maggiori di quelle a cui potrebbero emergere problemi dovuti agli infiniti che si incontrano nei calcoli.

Il muone può anche intervenire in fenomeni di decadimento beta: il W emesso da un quark d che si trasforma in u può accoppiarsi al muone invece che all'elettrone (fig. 90). In tal caso, un'altra particella, detta neutrino-mu, prende il posto dell'usuale neutrino, che per venire distinto viene chiamato neutrino-e. Così la tavola delle particelle si è arricchita di due nuovi arrivi da porre accanto a elettrone e neutrino: il muone e il neutrino-mu.

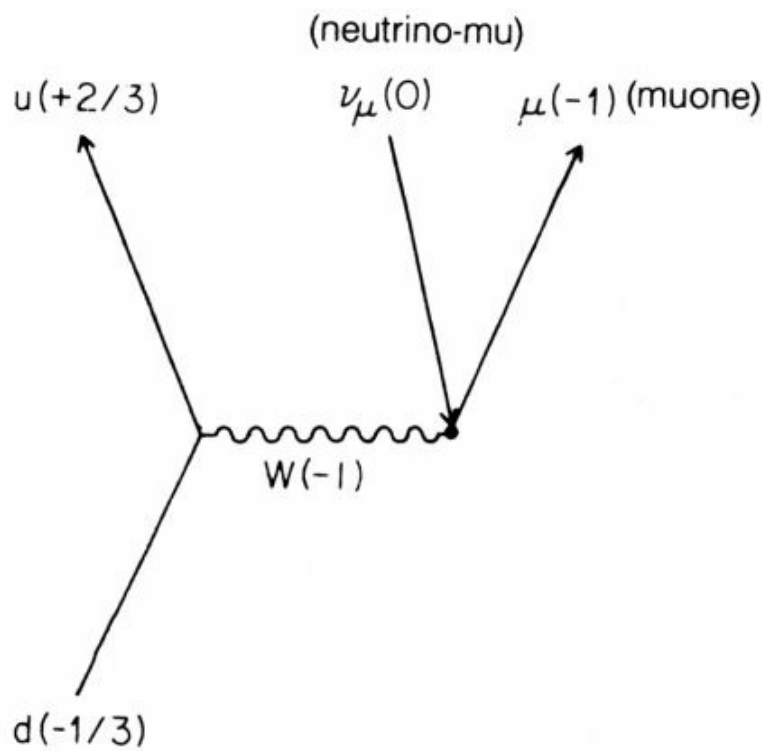


Fig. 90. I W hanno anche un'ampiezza per emettere un muone invece che un elettrone. In questo caso il neutrino-mu prende il posto del neutrino elettronico.

E i quark? È noto da molto tempo che le particelle osservate devono essere costituite anche da un altro tipo di quark, più pesante di u e d . All'elenco delle particelle fondamentali va così aggiunto questo terzo quark, detto s (da «strano»). Il quark s ha una massa di circa 200 MeV, da confrontare con quelle di circa 20 MeV dei quark u e d .

Per molti anni si è pensato che i quark esistessero solo nei tre «sapori» u , d ed s , ma nel 1974 è stata osservata una nuova particella, chiamata mesone ψ , che non può essere costituita da questi quark. C'era anche un ottimo argomento teorico che prevedeva l'esistenza di un quarto tipo di quark, accoppiato al quark s tramite un W , nello stesso modo in cui interagiscono col W i quark u e d (fig. 91). Il «sapore» di questo nuovo quark viene indicato con c . (Non oso dirvi al posto di quale nome stia la c , ma forse l'avrete letto sui giornali. I nomi diventano sempre peggiori!).

Questa ripetizione di particelle con proprietà identiche ma con masse diverse è un completo mistero. A che cosa è dovuta questa strana duplicazione? Come disse il prof. L.I. Rabi quando venne scoperto il muone: «E questo chi l'ha ordinato?».

			particelle di spin 1		
particelle di spin 1/2			fotone	gluone	W
			0	0	~80.000
tau τ ~1860	muone μ 105,8	elettrone e 0,511	-1	0	↕
neutrino ν_τ 0	neutrino ν_μ 0	neutrino ν_e 0	0	0	
quark b ~4800	quark s ~200	quark d ~10	-1/3	g	↕
	quark c ~1800	quark u ~10	+2/3	g	
accoppiamenti					

nome

simbolo

massa (MeV)

elettrone

e

0,511

nome	elettrone
simbolo	e
massa (MeV)	0,511

Fig. 92. Si riparte! A energie ancor più elevate si incontra un'altra ripetizione delle particelle di spin 1/2. La ripetizione sarebbe completa se si trovasse una particella le cui proprietà comportino l'esistenza di un nuovo sapore di quark. Attualmente si stanno preparando esperimenti per cercare le tracce di un'ulteriore ripetizione a energie ancora maggiori. Il motivo di queste ripetizioni è completamente sconosciuto.

Intanto sono stati progettati nuovi esperimenti per vedere se si hanno ulteriori ripetizioni del ciclo. Sono già in costruzione macchine che permetteranno di osservare elettroni ancora più pesanti del tau, fino a una massa di circa 4.000 MeV.

Misteri come questo delle «famiglie» che si ripetono rendono molto interessante il mestiere del fisico teorico: la Natura offre meravigliosi rompicapo! Perché l'elettrone viene ripetuto con massa 206 e 3.640 volte maggiore?

Voglio fare un'ultima precisazione per completare fino in fondo l'esposizione delle proprietà delle particelle. Quando un quark d si accoppia a un W per trasformarsi in u , ha anche una piccola ampiezza di probabilità di trasformarsi in c ; analogamente, c'è una piccola ampiezza che il quark u diventi s anziché d , e un'ampiezza ancora più piccola che si trasformi in b (fig. 93). Dunque i W mescolano le cose tra loro, permettendo ai quark di passare da una colonna all'altra della tabella. Il perché di tali valori delle ampiezze di probabilità per la trasformazione di un quark in un quark di «sapore» diverso è del tutto sconosciuto.

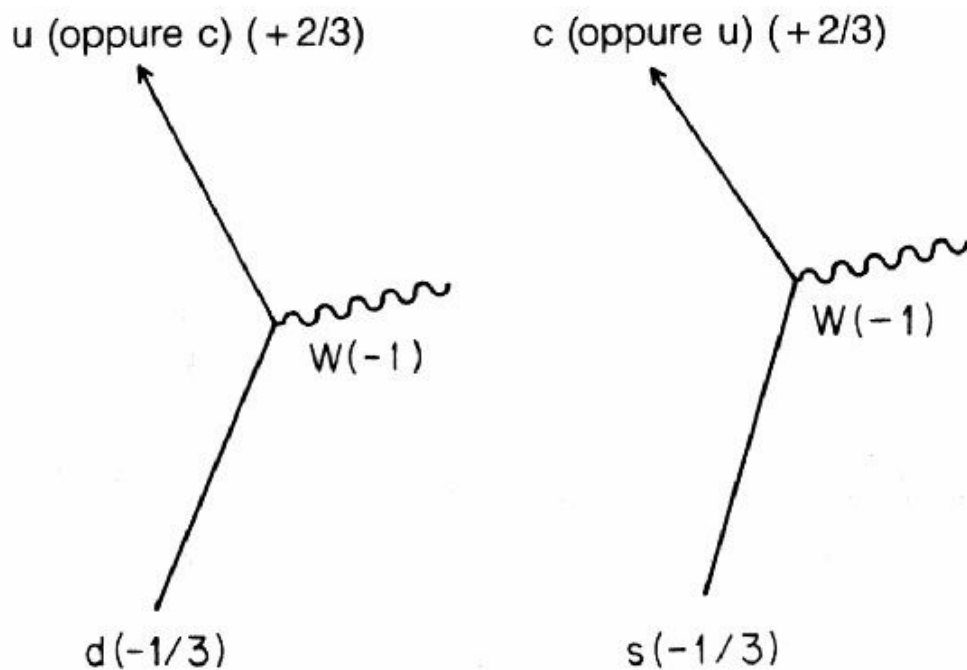


Fig. 93. Il quark d ha una piccola ampiezza per trasformarsi in quark e invece che u , e il quark s ha una piccola ampiezza per trasformarsi in u invece che in e , in entrambi i casi emettendo un W . Pertanto i W sembrano capaci di cambiare il sapore dei quark facendoli passare da una colonna della tabella a un'altra (si veda la fig. 92).

E questo è quanto si sa sulle altre parti della fisica quantistica. Una bella confusione, direte voi, e la fisica si è cacciata in un ginepraio inestricabile. Ma è sempre stato così. La Natura è sempre apparsa intricatissima e caotica finché, andando avanti, cominciamo a distinguere contorni e a formulare teorie; allora emerge un po' di chiarezza e le cose appaiono più semplici. Tutta questa confusione non è niente a confronto di quella che vi avrei mostrato dieci anni fa, quando avrei dovuto parlare di quattrocento tipi di particelle. E pensate alla confusione che c'era agli inizi del secolo, quando si parlava di calore, magnetismo, elettricità, luce, raggi X, raggi ultravioletti, indici di rifrazione, coefficienti di riflessione e di tante altre proprietà dei materiali che oggi vengono descritte mediante l'unica teoria dell'elettrodinamica quantistica.

Una cosa ci tengo a sottolineare: le teorie relative agli altri settori della fisica sono molto simili all'elettrodinamica quantistica, essendo costruite tutte a partire dall'interazione di oggetti di spin $1/2$ (come gli elettroni e i quark) con oggetti di spin 1 (come i fotoni, i gluoni e i W) nell'ambito dello schema basato sulle ampiezze, in cui la probabilità di un evento è data dal quadrato della lunghezza di una freccia. Perché le teorie fisiche hanno tutte struttura così simile?

Ciò può essere dovuto a diversi motivi. Il primo è la scarsa fantasia di noi fisici: quando osserviamo fenomeni nuovi cerchiamo sempre di inquadrarli negli schemi già noti, finché un numero sufficiente di esperimenti non ci costringe a rinunciarvi. Se a un certo punto arriva un fisico che dichiara: «Le cose stanno così e cosà, e notate la meravigliosa somiglianza tra queste teorie», non è perché questa somiglianza esista veramente in Natura, ma perché i fisici non hanno saputo far altro che pensare sempre alla stessa cosa in vesti diverse.

Un'altra possibilità è che sia davvero la stessa cosa, che la Natura conosca un unico modo di operare e ripeta di volta in volta la stessa storia.

Vi è infine la possibilità che le interazioni appaiano simili perché sono aspetti diversi di un'unica interazione; che ci sia una struttura sottostante più ampia che si frammenta

in parti. E queste parti appaiono diverse tra loro come le dita di una stessa mano. Molti fisici stanno dedicando tutti i loro sforzi alla ricerca di una teoria che permetta di unificare tutti i fenomeni noti in un supermodello. È una partita entusiasmante, ma per il momento i giocatori sono in completo disaccordo tra loro su quale sia la struttura unificante. Esagero solo di poco quando dico che queste speculazioni teoriche non sono granché più attendibili della vostra congettura sull'esistenza del quark t , e vi garantisco che i fisici non sanno stimare la massa del t meglio di voi!

Ad esempio, sembra che elettrone, neutrino, quark u e quark d siano da considerarsi un'unica «famiglia»; infatti i primi due sono accoppiati al W come gli ultimi due. Attualmente si pensa che i quark possono solo cambiare «colore» o «sapore». Ma chi ci garantisce che un quark non possa trasformarsi in un neutrino interagendo con particelle per ora sconosciute? Una bella idea: che accadrebbe se fosse vera? Vorrebbe dire che il protone è instabile.

Poniamo che venga proposta una teoria in base alla quale il protone è instabile. Si fanno i conti e si trova che in tal caso non ci sarebbero più protoni nell'universo! Allora si manipolano un po' i numeri, si ipotizza una massa un po' più grande per la nuova particella e dopo molti sforzi si riesce a predire una velocità di decadimento del protone leggermente inferiore al più recente limite sperimentale al di sopra del quale è stato accertato che il protone non decade.

Se nuovi esperimenti abbassano tale valore, le teorie vengono modificate in modo da evitare la difficoltà. Gli esperimenti più recenti hanno dimostrato che la velocità di decadimento del protone è certamente inferiore a un valore cinque volte più piccolo di quello predetto dall'ultima teoria «definitiva». Ebbene, che cosa credete sia successo? La fenice è risorta con un nuovo aggiustamento della teoria, che richiede esperimenti ancora più accurati per essere verificato. Non si sa se il protone decada o meno, e dimostrare che non decade è estremamente difficile.

In queste lezioni non ho mai parlato della gravitazione. Il motivo è che l'attrazione gravitazionale tra due oggetti è estremamente debole: la forza gravitazionale tra due elettroni è inferiore alla forza elettrica tra di loro di un fattore 1 seguito da 40 zeri (o forse 41). Nella materia la forza elettrica serve a tenere gli elettroni vicini al nucleo del loro atomo, dando luogo a un delicato equilibrio di cariche positive e negative che si cancellano tra loro. Invece la forza gravitazionale è puramente attrattiva, e cresce all'aumentare del numero di atomi, finché, quando si considerano corpi con massa sufficientemente ponderosa, come quelli che ci circondano, diventa possibile rivelarne gli effetti: sui pianeti, su noi stessi e così via.

Essendo la forza gravitazionale tanto più debole delle altre interazioni, è impossibile al giorno d'oggi fare esperimenti sufficientemente accurati per misurare effetti la cui spiegazione richieda la precisione di una teoria quantistica della gravitazione.³³ Non vi è dunque modo di sottoporre a verifica sperimentale le teorie quantistiche della gravità, ma questo non impedisce che ne esistano alcune, le quali tutte prevedono l'esistenza dei «gravitoni», particelle appartenenti a una nuova classe di polarizzazione detta «spin 2», e altre particelle fondamentali, alcune con spin $3/2$. La migliore tra queste teorie non è in grado di inquadrare le particelle effettivamente conosciute, mentre prevede l'esistenza di

un gran numero di particelle non osservate. Anche nelle teorie quantistiche della gravitazione compaiono degli infiniti quando si considerano termini con interazioni, ma il procedimento assurdo che permetteva di liberarsi degli infiniti in elettrodinamica quantistica non funziona per la gravità quantistica. Dunque non solo non ci sono dati sperimentali con cui raffrontare una teoria quantistica della gravitazione, ma non ci sono neppure teorie accettabili.

In tutta questa storia vi è un aspetto particolarmente insoddisfacente: i valori osservati delle masse delle particelle. Non vi è nessuna teoria che li spieghi adeguatamente; li si usa di continuo nei conti ma non si ha la minima idea di che cosa siano e da dove vengano. Sono convinto che, a livello fondamentale, l'origine dei valori delle masse costituisce un problema molto serio e interessante.

Spero che tutte queste congetture sull'esistenza di nuove particelle non vi abbiano confuso troppo le idee, ma ho voluto includere questa discussione sulle altre parti della fisica per mostrare come il *carattere* delle leggi fisiche che le descrivono, basato sulle ampiezze e sul calcolo dei diagrammi che rappresentano l'interazione, sia lo stesso che nella teoria dell'elettrodinamica quantistica, che è il nostro miglior modello di una buona teoria fisica.

Note aggiunte in bozze

Nota aggiunta in bozze (novembre 1984)

Successivamente a questa serie di lezioni, strani eventi osservati in alcuni esperimenti fanno sembrare plausibile che si sia alla vigilia della scoperta di qualche altra particella o qualche altro fenomeno, nuovo e inaspettato (e pertanto non discusso in queste lezioni).

Nota aggiunta in bozze (aprile 1985)

Gli «strani eventi» di cui sopra sembrano ormai essere stati un falso allarme. La situazione sarà certamente cambiata di nuovo quando questo libro verrà letto. I tempi dell'editoria sono più lenti di quelli della fisica.

Nota aggiunta dal traduttore (ottobre 1989)

Al momento della stampa di questa edizione gli «strani eventi» di cui alle note precedenti sono definitivamente scomparsi, e la situazione sperimentale non è sostanzialmente diversa da quella descritta da Feynman. Anzi, risultati recentissimi legati alla particella Z_0 sembrano rendere certo che non esistono altre «famiglie» di particelle oltre alle tre di cui si parla in questo libro, se si ammette che tutti i neutrini abbiano masse sufficientemente piccole (< 1 GeV).

In questo libro, con stupefacente chiarezza, un grande fisico ci spiega come tutto ciò che percepiamo dipenda da accadimenti naturali che violano ogni aspettativa del senso comune. La via scelta è la seguente: guidare, come in un vero *tour de force*, ogni testa pensante negli impensabili meandri dell'elettrodinamica quantistica (abbreviata nella sigla QED del titolo). Il punto di partenza è la riflessione della luce. Prendendo le mosse da esperienze elementari, Feynman ci mostra come tale riflessione, lungi dall'essere un semplice mutamento di direzione di un raggio luminoso, sia un accadimento che va contro tutte le concezioni del senso comune. Da ciò una serie inarrestabile di conseguenze, in ogni direzione. E - ciò che più conta per il lettore non specialista - Feynman procede mantenendo sempre la spiegazione in stretto contatto con l'esame di varie esperienze fisiche, così da farci entrare, in certo modo, nella mente dello scienziato che le osserva (e, per certi fenomeni, la prima mente che osservava fu proprio la sua). Mentre al tempo stesso riesce a presentare concetti ben noti in termini sorprendenti per gli stessi fisici. C'è poi una importante novità di atteggiamento, in Feynman. Mentre alcuni fra i suoi illustri predecessori, pur avendo riconosciuto la sconcertante realtà della meccanica quantistica, continuavano a guardare con nostalgia alle sicurezze del senso comune insite nella fisica classica, Feynman è stato forse il primo fisico a vivere senza inibizioni lo shock della quantizzazione. Il suo criterio sembrerebbe il seguente: porsi spudoratamente al di là del buon senso e della intuizione comune, purché i conti tornino e le misure siano esatte.

Intorno a Richard Feynman si è venuta formando una sorta di leggenda. Premio Nobel per la fisica nel 1965 come riconoscimento alle sue ricerche di elettrodinamica quantistica (che sono anche il tema di questo libro), divenne presto una figura molto popolare e molto amata per la sua impressionante versatilità e per lo spirito irriverente, uniti a una rara capacità di rendere accessibili i fatti più complessi della fisica. Questo libro, che apparve nel 1985, ne è la dimostrazione più felice.

1)

Abbreviazione di «California Institute of Technology». [*N.d.T.*] □

2)

Acronimo di «Quantum Electro-Dynamics». [*N.d.T.*] □

3)

Come faceva a saperlo? Newton, che era un genio, annotò: «Perché il vetro può essere lucidato». Chiederete: come diavolo faceva a concludere da ciò che il vetro non è fatto di buchi e di chiazze riflettenti? Ebbene, Newton si faceva da solo le lenti e gli specchi, e sapeva in che cosa consiste questa operazione: si graffia la superficie di un pezzo di vetro con polveri fatte di granuli sempre più fini. Via via che i graffi diventano più sottili, la superficie del vetro cambia aspetto e passa dal grigio spento (dovuto al fatto che i graffi profondi diffondono la luce) a una limpidezza trasparente (perché invece i graffi estremamente sottili lasciano passare la luce). Newton vide quindi che è impossibile ammettere che la luce venga influenzata da irregolarità molto piccole quali graffi, buchi o chiazze; anzi, come lui sapeva, è vero il contrario. Graffi molto sottili, e perciò anche chiazze altrettanto piccole, non disturbano il cammino della luce. La teoria dei buchi e delle chiazze non può dunque funzionare. $\square$

4)

È una vera fortuna per noi che Newton fosse convinto della natura «corpuscolare» della luce, perché questo ci permette di vedere tutte le contorsioni di una mente intelligente e non condizionata da preconconcetti che si sforza di spiegare la riflessione parziale da parte di due o più superfici. (I sostenitori della teoria ondulatoria non hanno mai dovuto fare nessuna fatica per spiegare questo fenomeno). Newton ragionò come segue: nonostante le apparenze, la superficie frontale in realtà non riflette nessuna luce. Infatti, se così fosse, come potrebbe tale luce riflessa essere ricatturata allorché lo spessore è tale da non produrre nessuna riflessione? Quindi la luce deve essere riflessa dalla seconda superficie. Per spiegare il fatto che la quantità di luce riflessa dipende dallo spessore del vetro, Newton propose la seguente idea: la luce che urta la prima superficie produce una specie di onda o di campo che viaggia con essa e ne determina la predisposizione a riflettersi o no sulla seconda superficie. Egli chiamò questo processo «fasi di riflessione facile o di trasmissione facile» e suppose che esso accada ciclicamente, a seconda dello spessore del vetro.

Questa ipotesi incontra due difficoltà; la prima è l'effetto delle superfici addizionali: ogni nuova superficie modifica la riflessione, come ho spiegato nel testo. La seconda è che la luce viene innegabilmente riflessa da un lato, dove non esiste una seconda superficie, quindi la superficie riflettente *deve* essere quella frontale. Nel caso di un'unica superficie, Newton disse che la luce ha una predisposizione a riflettersi. È accettabile una teoria in cui la luce sa che tipo di superficie sta urtando, e sa che si tratta di una superficie unica?

Newton non mise in evidenza queste difficoltà della sua teoria sulle «fasi di trasmissione e di riflessione», pur rendendosi conto, è chiaro, che la teoria non era soddisfacente. Ma ai tempi di Newton si tendeva a sorvolare sui punti carenti di una teoria; oggi invece uno scienziato tende a mettere egli stesso in evidenza i punti in cui la sua teoria non si accorda con i dati dell'osservazione. Dico questo non per criticare Newton, ma per far notare un aspetto positivo del modo di comunicare della scienza d'oggi. $\square$

5)

Questa teoria sfruttava il fatto che le onde si possono sommare o cancellare tra loro, e i calcoli basati su questo modello riproducono i dati delle osservazioni di Newton e di quelle seguite per centinaia d'anni successivi. Poi divenne possibile costruire strumenti abbastanza sensibili da rivelare un singolo fotone: la teoria ondulatoria prevedeva che, rendendo più fioca la luce, il ticchettio del fotomoltiplicatore si sarebbe dovuto indebolire sempre più, mentre in realtà resta della stessa intensità, solo che si dirada. Nessun modello ragionevole era in grado di spiegare questo fatto, e per un certo periodo bisognò far lavorare di più il cervello: per dire se la luce era fatta di onde o di particelle era necessario sapere quale esperimento si stesse analizzando. Questo stato di confusione venne chiamato «dualità ondulatorio-corpuscolare» della luce, e qualche bello spirito dichiarò che la luce era onda di lunedì, mercoledì e venerdì, e particelle di martedì, giovedì e sabato; quanto alla domenica, bisognava pensarci su. Scopo di queste lezioni è di raccontare come è stato «risolto» questo enigma. □

6)

1. Anche le parti dello specchio corrispondenti a frecce dirette prevalentemente verso sinistra danno luogo a una forte riflessione, una volta eliminate le parti le cui frecce puntano nella direzione opposta. Solo quando sono presenti parti a cui corrispondono sia frecce dirette a sinistra sia frecce dirette a destra i contributi si cancellano. La situazione è analoga al caso della riflessione parziale da due superfici: ciascuna riflette la luce, ma la riflessione complessiva si annulla se lo spessore è tale che le corrispondenti frecce hanno direzioni opposte. $\square$

7)

Non resisto alla tentazione di descrivere un reticolo fatto dalla Natura: i cristalli di sale sono costituiti da atomi di sodio e di cloro disposti secondo uno schema regolare. La loro disposizione alternata, come i solchi a cui accennavo prima, agisce come un reticolo se il cristallo è illuminato da luce del colore giusto (raggi X, in questo caso). Individuando con precisione le posizioni in cui al rivelatore arriva la massima quantità della luce che ha subito questo particolare tipo di riflessione (detto diffrazione), è possibile determinare esattamente quanto distano tra loro i solchi, e quindi la distanza tra gli atomi (fig. 28). È un metodo bellissimo, che permette di determinare la struttura di tutti i tipi di cristalli e altresì di confermare che i raggi X hanno la stessa natura della luce. I primi esperimenti di questo tipo furono fatti nel 1914 e furono molto emozionanti: era la prima volta che si vedeva nei particolari la disposizione degli atomi nelle varie sostanze. □

8)

Questo è un esempio del «principio di indeterminazione»: vi è una sorta di «complementarità» tra quello che si sa del percorso della luce tra i due blocchi e quello che si sa del suo percorso successivo: è impossibile avere conoscenza esatta di entrambi. Inquadriamo il principio di indeterminazione nel suo contesto storico: nel periodo in cui si stavano affermando le idee rivoluzionarie della meccanica quantistica, si cercava ancora di interpretarle sulla base di vecchi concetti, quale quello della propagazione rettilinea della luce. Ma, arrivati a un certo punto, i vecchi concetti cessavano di funzionare, per cui fu sviluppato un avvertimento che in sostanza diceva: «I vecchi concetti non servono più quando...». Ma se si buttano via tutti i concetti vecchi e si usano le idee esposte in queste lezioni, sommando le *frecce* per tutti i modi in cui un evento può accadere, non c'è necessità alcuna di un principio di indeterminazione! $\square$

9)

I matematici hanno cercato di individuare tutte le possibili quantità che obbediscono alle regole dell'algebra ($A + B = B + A$, $A \times B = B \times A$, ecc.). In principio queste regole erano state trovate per i numeri interi positivi, usati per contare cose come le persone o le mele. I numeri vennero arricchiti dall'introduzione dello zero, delle frazioni, degli irrazionali (numeri che non si possono esprimere come rapporto tra due interi) e dei numeri negativi, ma sempre continuando a obbedire alle stesse regole dell'algebra. Alcuni dei numeri introdotti dai matematici all'inizio posero delle difficoltà alla gente comune — ad esempio risultava difficile cogliere l'idea di mezza persona — ma oggi non vi sono più problemi e la notizia che in una regione vi è una media di 3,2 persone per km quadrato non desta né indignazione morale né raccapriccio. Nessuno cerca di immaginarsi le 3,2 persone; quel 3,2, come ben si sa, se moltiplicato per 10, dà 32. Pertanto grandezze che verificano le regole dell'algebra possono risultare interessanti anche se non rappresentano sempre situazioni reali. Le frecce su un piano possono essere «sommate» ponendo la coda di una sulla punta di un'altra, e «moltiplicate» con contrazioni e rotazioni successive. Poiché queste frecce verificano le stesse regole dell'algebra dei soliti numeri, i matematici le chiamano numeri. Ma per distinguerle da quelli usuali le chiamano «numeri complessi». Per quelli di voi che hanno studiato abbastanza matematica da sapere che cos'è un numero complesso, avrei potuto dire semplicemente: «La probabilità di un evento è data dal modulo quadrato di un numero complesso. Se un evento può avvenire in più modi alternativi, i rispettivi numeri complessi vanno sommati; se può avvenire come successione di passaggi, i numeri complessi vanno moltiplicati». L'effetto è più solenne, ma non ho detto nulla che non avessi detto anche prima: ho solo usato un linguaggio diverso. $\square$

10)

Avrete notato che per riuscire a dar conto del 100% della luce, 0,0384 è stato cambiato in 0,04 e si è usato 84% come quadrato di 0,92. Ma quando si somma veramente tutto non c'è bisogno di arrotondare 0,0384 e il quadrato di 0,92: tutti i pezzettini di freccia corrispondenti ai vari percorsi possibili per la luce si combinano tra loro dando sempre la risposta esatta. Per coloro che amano questo genere di cose, la fig. 46 illustra un altro possibile percorso della luce dalla sorgente al rivelatore in A: una successione di tre riflessioni e due trasmissioni che danno luogo a una freccia risultante di lunghezza $0,98 \times 0,2 \times 0,2 \times 0,2 \times 0,98$, ossia circa 0,008, una freccia decisamente molto piccola. Perché il calcolo della riflessione parziale su due superfici sia completo, bisogna aggiungere anche questa minuscola freccia, più un'altra ancor più piccola corrispondente a cinque riflessioni, e così via. $\square$

11)

Tale regola concorda con quello che insegnano a scuola, ossia che la quantità di luce che attraversa una superficie è inversamente proporzionale al quadrato della distanza dalla sorgente, perché se una freccia si riduce a metà, il suo quadrato si riduce a un quarto. □

12)

Questo fenomeno, noto come effetto Hanbury-Brown-Twiss, è stato usato per distinguere sorgenti singole da sorgenti doppie di onde radio nello spazio cosmico, anche quando le due sorgenti sono estremamente vicine tra loro. □

13)

Tener chiaro questo principio dovrebbe aiutare gli studenti a non lasciarsi confondere dalla «riduzione del pacchetto d'onde» e da altre magie analoghe. □

14)

L'analisi completa di questa situazione è molto interessante: se i rivelatori in A e in B non sono perfetti e pertanto rivelano i fotoni solo *alcune* volte, vi sono *tre* condizioni finali distinguibili: 1) scattano i rivelatori in A e in D, 2) scattano quelli in B e in D, 3) scatta solo quello in D, mentre A e B restano nello stato iniziale. Per i primi due eventi le probabilità vanno calcolate come illustrato nel testo (a parte un nuovo passo: non essendo i rivelatori perfetti, vi è un'ulteriore contrazione connessa alla probabilità che il rivelatore in A, o in B, non scatti). Quando scatta solo D, non è possibile distinguere i due casi e la Natura gioca con noi facendo intervenire l'interferenza, lo stesso peculiare effetto che si verifica quando non vi sono rivelatori in A e in B (solo che la freccia finale è contratta dell'ampiezza relativa al fatto che i rivelatori *non* scattino). Il risultato finale è una mescolanza, la semplice somma dei tre casi (fig. 51). Al crescere della sensibilità dei rivelatori si produce sempre meno interferenza. $\square$

15)

Questi grafici sono poi diventati uno strumento matematico corrente nella fisica con il nome di «diagrammi di Feynman» [*N.d.T.*]. □

16)

In queste lezioni considero che la posizione di un punto abbia una sola dimensione spaziale. Per individuare un punto nello spazio tridimensionale si deve fare riferimento a una «stanza» e misurare la distanza del punto dal pavimento e da due pareti adiacenti (tutti ad angolo retto tra loro). Si chiamino X_1 , Y_1 e Z_1 questi tre valori. La distanza del nostro punto da un altro con valori X_2 , Y_2 e Z_2 è data dal «teorema di Pitagora tridimensionale» in base al quale il quadrato di tale distanza risulta

$$(X_2 - X_1)^2 + (Y_2 - Y_1)^2 + (Z_2 - Z_1)^2$$

L'eccesso tra tale valore e il quadrato della differenza di tempo, cioè

$$(X_2 - X_1)^2 + (Y_2 - Y_1)^2 + (Z_2 - Z_1)^2 - (T_2 - T_1)^2$$

viene a volte chiamato «intervallo», o I , e la teoria della relatività di Einstein afferma che F (da A a B) dipende solo da I . Il contributo principale a F (da A a B) si ha, come ci si aspetterebbe, quando la distanza spaziale è uguale a quella temporale (cioè quando I è zero). Ma ci sono anche contributi per I non nullo, che sono inversamente proporzionali ad I e puntano nella direzione delle 3 per I positivo (quando la luce va più veloce di c) e verso le 9 per I negativo. Tali contributi si elidono tra loro in molte circostanze (fig. 56). $\square$

17)

La formula che esprime $E(\text{da } A \text{ a } B)$ è complicata, ma c'è un modo interessante per spiegarne la struttura. Si può rappresentare $E(\text{da } A \text{ a } B)$ come somma dell'enorme numero di modi differenti in cui l'elettrone può andare dal punto A al punto B (fig. 57). L'elettrone può fare «un solo salto» andando direttamente da A a B, può fare «due salti» fermandosi prima in C, oppure «tre salti» fermandosi in D e in E e così via. L'ampiezza per un salto diretto da F a G è data da $F(\text{da } F \text{ a } G)$, cioè è la stessa dell'ampiezza per un fotone che vada tra gli stessi due punti. Ogni «fermata» introduce un'ampiezza n^2 , dove n è il numero che abbiamo appena visto e che serve per ottenere il risultato giusto.

L'espressione di $E(\text{da } A \text{ a } B)$ è dunque una somma del tipo: $F(\text{da } A \text{ a } B)$ [cioè «un salto»] + $F(\text{da } A \text{ a } C) \times n^2 \times F(\text{da } C \text{ a } B)$ [«due salti» con fermata in C] + $F(\text{da } A \text{ a } D) \times n^2 \times F(\text{da } D \text{ a } E) \times F(\text{da } E \text{ a } B)$ [«tre salti» con fermata in D e E] + ... per *tutte le possibili posizioni* dei punti C, D, E, ecc.

Si osservi che al crescere di n i percorsi meno diretti danno un contributo maggiore alla freccia finale. Se $n = 0$ (come per il fotone) tutti i termini che contengono n sono nulli e resta solo il primo, cioè $F(\text{da } A \text{ a } B)$. Dunque $E(\text{da } A \text{ a } B)$ e $F(\text{da } A \text{ a } B)$ sono strettamente connesse. $\square$

18)

Questo numero, l'ampiezza per emettere o assorbire un fotone, è talvolta chiamato «carica» della particella. □

19)

Includendo gli effetti della polarizzazione dell'elettrone, la «seconda» freccia andrebbe «sottratta», cioè ruotata di 180° e sommata. (Ma su questo tornerò più avanti). □

20)

Le condizioni finali dell'esperimento in questi due casi più complicati sono le stesse che nei casi più semplici: gli elettroni partono dai punti 1 e 2 ed arrivano in 3 e 4. Non essendo quindi tali alternative distinguibili dalle precedenti, le frecce relative a questi modi vanno sommate a quelle dei primi due. $\square$

21)

Questi fotoni scambiati, che non appaiono mai nelle condizioni iniziali e finali di un esperimento, sono a volte chiamati «fotoni virtuali». $\square$

22)

Dirac propose l'esistenza di «antielettroni» nel 1931; l'anno successivo essi furono rivelati sperimentalmente da Carl Anderson che li battezzò «positroni». Oggi è facile produrre positroni, ad esempio facendo urtare tra loro due fotoni, e conservarli per settimane in un campo magnetico. □

23)

L'ampiezza per lo scambio di un fotone è $(-j) \times F(\text{da } A \text{ a } B) \times j$: due fattori di accoppiamento moltiplicati per l'ampiezza di propagazione del fotone da un punto a un altro. L'ampiezza di accoppiamento di un protone a un fotone è $-j$. $\square$

24)

Il raggio dell'arco dipende evidentemente dalla *lunghezza* della freccia relativa a ciascuna sezione, determinata a sua volta dall'ampiezza di diffusione di un fotone da parte di un elettrone appartenente a un atomo del vetro, che è stata indicata con S. È possibile calcolare tale raggio applicando le formule per i tre eventi base all'enorme numero di scambi di fotoni che si verificano e sommando le ampiezze così ottenute. Si tratta di un problema estremamente complesso, ma nel caso di sostanze molto semplici si è riusciti a calcolare il raggio con notevole successo, e le sue variazioni da sostanza a sostanza sono descritte in maniera ragionevole dall'elettrodinamica quantistica. Tuttavia occorre ammettere che non si è mai tentato questo calcolo a partire dai principi primi per sostanze complesse come il vetro. In tali casi il raggio viene determinato sperimentalmente. Per il vetro gli esperimenti hanno dato il valore di circa 0,2 (per luce che incide perpendicolarmente). □

25)

Le frecce dovute alla riflessione sulle varie sezioni (che formano un «cerchio») hanno la stessa lunghezza di quelle che provocano la rotazione in più della freccia finale relativa alla trasmissione. Vi è dunque una relazione tra coefficiente di riflessione parziale e indice di rifrazione di una sostanza.

Può sembrare che la freccia finale sia più lunga di 1, il che vorrebbe dire che dal vetro esce più luce di quanta ne entri! Tale impressione è però dovuta al non aver considerato l'ampiezza per i percorsi in cui il fotone raggiunge una sezione, un nuovo fotone rimbalza fino a una sezione precedente e infine un terzo fotone attraversa il vetro, o altre possibilità ancora più complicate. Il loro effetto è di ruotare leggermente le frecce secondarie mantenendo la lunghezza di quella finale tra 0,92 e 1, così che la somma delle probabilità di trasmissione e di riflessione è sempre il 100%. $\square$

26)

Tale difficoltà può essere considerata da un altro punto di vista: forse è errata l'idea che due punti possano essere infinitamente vicini, è falsa l'assunzione che la geometria continui a valere invariata fino in fondo. Tuttavia, se si ammette che la minima distanza possibile tra due punti sia 10-100 cm. (la più piccola distanza misurata fino ad oggi in un esperimento è 10-16 cm.), scompaiono gli infiniti ma nascono altre incoerenze: ad esempio la probabilità totale di un evento risulta un po' più grande o un po' più piccola del 100%, oppure compaiono infinitesime quantità di energia negativa. È stato proposto che queste difficoltà siano dovute al non aver preso in considerazione gli effetti della gravità, di solito debolissimi, ma che diventano importanti a distanze così piccole. □

27)

Mentre negli esperimenti di alta energia vengono fuori dai nuclei molte particelle, alle basse energie, cioè in condizioni più normali, i nuclei risultano costituiti solo da protoni e neutroni. □

28)

Un MeV è una quantità molto piccola, adatta quindi a questi tipi di particelle: vale circa $1,78 \times 10^{-27}$ grammi. □

29)

Si osservino i nomi: «fotone» viene dalla parola greca per luce, «elettrone» dalla parola greca per ambra, le cui proprietà elettriche furono le prime ad essere conosciute. Ma con lo sviluppo della fisica moderna l'interesse per il greco antico è andato scemando, e si è arrivati a coniare nomi come «gluoni» (dalla parola inglese glue, «colla»). Anche d e u stanno al posto di nomi, ma non vorrei fare troppa confusione: un quark d non è più *down*, «giù», di quanto un quark u sia *up*, «su». La proprietà di essere d o u di un quark viene chiamata «sapore». $\square$

30)

Dopo la conclusione di questo ciclo di conferenze, sono state raggiunte energie sufficientemente elevate per produrre dei W , e la massa misurata è risultata molto vicina a quella predetta dalla teoria. $\square$

31)

Il momento magnetico del muone è stato misurato con grande accuratezza e risulta 1,001165924 (con un'incertezza di +9 sull'ultima cifra), mentre quello dell'elettrone risulta 1,00115965221 (con un'incertezza di 3 sull'ultima cifra). Ci si chiederà come mai il momento magnetico del muone risulti leggermente maggiore di quello dell'elettrone. In uno dei nostri grafici (fig. 89) un elettrone emette un fotone che si trasforma in una coppia elettrone-positrone. C'è anche la possibilità che il fotone si trasformi in una coppia muone-antimuone, più pesanti dell'elettrone iniziale. Se invece è un muone ad emettere il fotone, questo può sempre trasformarsi in coppia elettrone-positrone, *più leggeri* del muone iniziale, mentre la coppia muone-antimuone avrebbe la stessa massa.

L'elettrodinamica quantistica descrive accuratamente *tutte* le proprietà elettriche del muone, oltre che dell'elettrone. □

32)

Successivamente a queste lezioni è stato osservato qualche indizio dell'esistenza di un quark t con massa molto alta, circa 40.000 MeV. □

33)

Quando Einstein ed altri provarono a unificare la gravitazione con l'elettrodinamica, entrambe le teorie erano ancora approssimazioni classiche. In parole povere erano sbagliate. Nessuna di esse era ancora fondata sullo schema delle ampiezze che è risultato così essenziale. □