

L'infinito sotto il tappeto

Come normalizzare le teorie divergenti

di Massimo Passera



a.
Un gruppo di fisici attorno a Feynman (al centro davanti al tavolino), tra cui Schwinger (tutto a destra nella foto), intenti in una discussione durante una conferenza scientifica nel giugno del 1947.

Immaginate una teoria che, in prima approssimazione, dia risultati corretti nel descrivere la realtà. È un buon inizio, ma appena tentate di raffinare il calcolo, per renderlo più preciso, trovate che la correzione che vi aspettavate piccola, in realtà, è infinita! Cosa fareste? Questa è la situazione in cui si trovarono Freeman Dyson, Richard Feynman, Julian Schwinger e Sin-Itiro Tomonaga con l'elettrodinamica quantistica (Qed), "la strana teoria della luce e degli elettroni", verso la fine degli anni '40.

Il problema degli "infiniti" era già noto da tempo in fisica classica. Basti pensare all'energia del campo elettrico di una sferetta carica, il cui raggio r tenda a zero: l'energia tende all'infinito (ovvero "diverge", come si dice in termini

scientifici) come $1/r$. Per la teoria della relatività speciale di Einstein, parte della massa della sferetta deriva dall'energia (divergente!) contenuta nel campo elettrico che la circonda. Si potrebbe pensare che nessuna carica elettrica è in realtà puntiforme e che il problema è semplicemente dovuto a un'astrazione matematica. Ma come formulare quindi una teoria che descriva oggetti elementari, cioè puntiformi? Il problema perdura nelle teorie di campo quantistiche come la Qed.

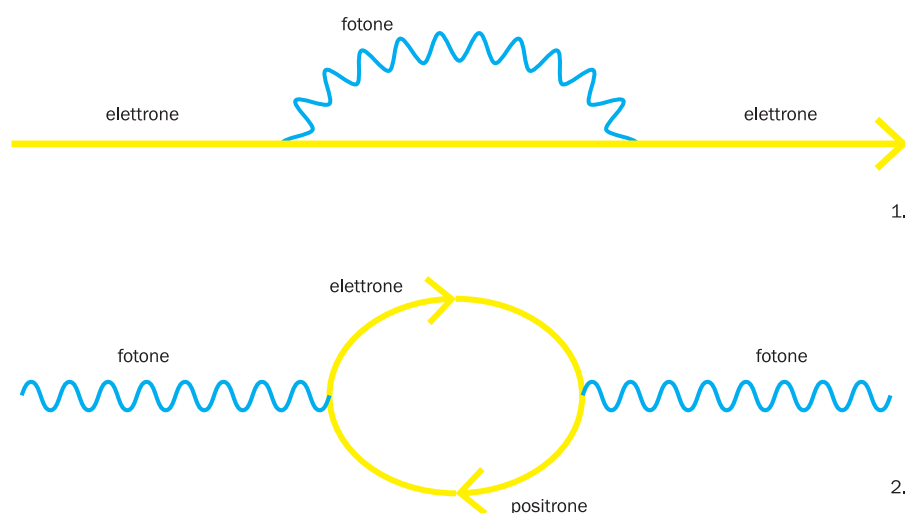
Alle piccolissime scale (vd. in Asimmetrie n. 18 p. 10, ndr) in cui la fisica classica cede il passo a quella quantistica, il principio di indeterminazione di Heisenberg consente, per intervalli di tempo molto piccoli, la presenza di

particelle "virtuali" (vd. anche in Asimmetrie n. 19 p. 19, ndr). Si tratta di particelle con caratteristiche identiche a quelle reali, ma con valori di energia anomali. Un elettrone a riposo, per esempio, ha un'energia pari alla sua massa m moltiplicata per il quadrato della velocità della luce c (è la famosa equazione di Einstein $E = mc^2$). Se l'elettrone è invece in movimento, la sua energia è maggiore di mc^2 , ed è data da una ben precisa funzione di m , c e della quantità di moto p (è l'equazione della relatività speciale $E^2 = p^2c^2 + m^2c^4$, vd. in Asimmetrie n. 19 p. 16, ndr): tanto più veloce è l'elettrone, tanto maggiore è la sua energia. Nessun elettrone può avere quindi energia nulla né, tantomeno, può avere energia negativa! Questo non è più

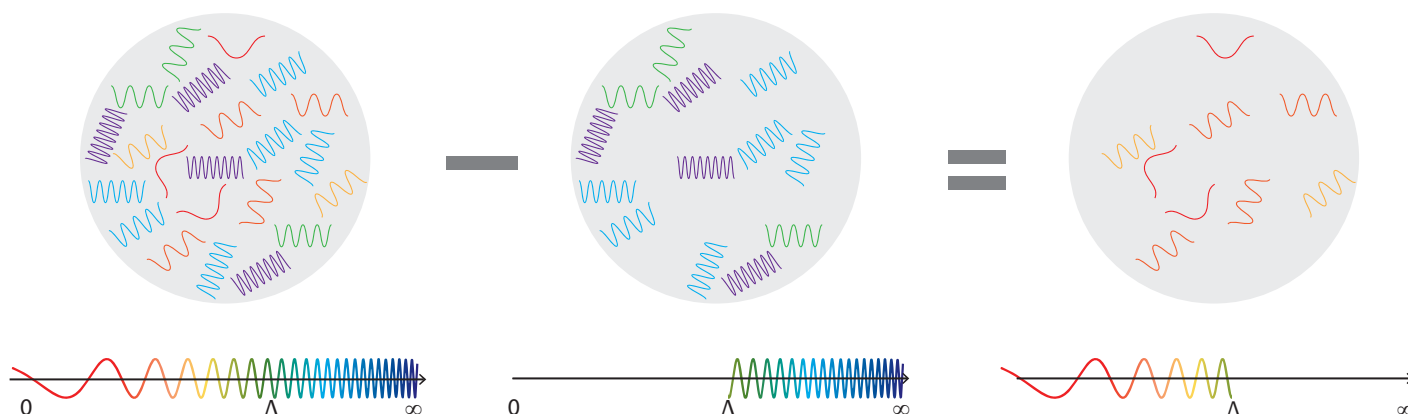
vero per un elettrone virtuale, che può avere invece qualsiasi valore di energia. Ma tanto più anomala è quest'energia (ovvero, tanto più virtuale è l'elettrone), tanto più breve sarà la sua vita. Il concetto di particella virtuale ci permette di comprendere l'origine degli infiniti in teorie di campo quantistiche. Partiamo, come primo esempio, dal fatto che un elettrone libero non può emettere un fotone. Consideriamo infatti un elettrone libero a riposo. Avendo velocità nulla, esso giace nel suo stato di minima energia mc^2 . Se emettesse un fotone, di qualsiasi energia, per la legge di conservazione dell'energia l'elettrone si ritroverebbe ad avere un'energia inferiore al suo valore minimo. Dunque non può emetterlo! Lo stesso vale anche per un elettrone libero in moto, dato che può essere immaginato come un elettrone a riposo visto da un osservatore in moto. Un elettrone libero può tuttavia emettere un fotone virtuale, divenendo esso stesso virtuale, purché lo riassorba rapidamente ritornando così a essere un elettrone reale. Questo processo può essere rappresentato in modo intuitivo tramite i diagrammi di Feynman (vd. fig. b1). Nella fase di transizione, l'elettrone e il fotone virtuali hanno una gamma infinita di energie possibili. Per esempio, il fotone virtuale può essere molto energetico, mentre l'elettrone virtuale può avere energia negativa. L'unico vincolo è che la somma delle loro energie sia pari a quella dell'elettrone iniziale (e finale). Consideriamo un secondo esempio: un fotone, che ha massa a riposo nulla, non può trasformarsi nel vuoto in una coppia costituita da un elettrone e un "positrone" (l'antiparticella dell'elettrone, cioè dotata della stessa massa m dell'elettrone, ma di carica opposta). Questo vale anche se il fotone ha un'energia maggiore di $2mc^2$, maggiore cioè della somma delle

energie dell'elettrone e del positrone liberi a riposo. Infatti, per quanto il fotone (che avendo massa nulla non può mai "stare fermo") sia energetico, se lo osservate in un sistema di riferimento in moto lungo la sua stessa direzione con una velocità sufficiente, esso avrà un'energia inferiore a $2mc^2$, che è il valore minimo necessario a produrre un elettrone e un positrone. Anche in questo secondo esempio, il fotone può però trasformarsi in un elettrone e un positrone virtuali, le cui energie possono variare in una gamma infinita di valori (vd. fig. b2). Dopo un tempo brevissimo, la coppia virtuale svanisce in un lampo di luce, lasciando però una traccia indelebile della sua breve esistenza.

La trattazione dei processi di Qed sopra descritti richiede quindi una specie di somma di tutte le possibili (infinite) configurazioni di energia che possono assumere le particelle virtuali. Il problema è che questa somma talvolta diverge, assume cioè un valore infinito. Ciascuna di queste configurazioni ha un peso, che dipende da quanto l'energia di ciascuna particella virtuale si discosta da quella che avrebbe se fosse reale. Se le configurazioni con particelle virtuali molto energetiche hanno un peso troppo grande, la somma totale è infinita. Queste divergenze, presenti in entrambi gli esempi sopra discussi, sono dette "ultraviolette", come i fotoni di alta energia. Altre divergenze, determinate invece dai contributi virtuali a bassa energia, vengono chiamate "infrarosse". Queste ultime compaiono in teorie con particelle prive di massa (come il fotone) e hanno un'origine fisica; ci indicano che stiamo studiando la quantità sbagliata e scompaiono quando calcoliamo processi fisici che possono essere effettivamente misurati sperimentalmente.



b. Due diagrammi di Feynman (vd in Asimmetrie n. 19 p. 11, ndr), che descrivono un elettrone libero (1) e un fotone libero (2). Un elettrone libero non può emettere un fotone, o meglio lo può fare solo se "riassorbe" il fotone emesso dopo un tempo "brevissimo". In questo tempo brevissimo sia l'elettrone che il fotone sono detti "virtuali" e possono assumere valori qualsiasi di energia, purché la somma delle due energie sia pari a quella dell'elettrone prima e dopo l'emissione/assorbimento del fotone. Allo stesso modo, un fotone libero non può trasformarsi in una coppia elettrone-positrone, se non per un tempo brevissimo. Come nel caso 1) l'elettrone e il positrone sono particelle virtuali con energia qualsiasi e con l'unico vincolo che la somma delle loro energie sia uguale a quella del fotone prima e dopo la produzione/annichilazione della coppia.



Le divergenze ultraviolette della Qed possono essere eliminate mediante la cosiddetta “rinormalizzazione”. È questa la soluzione che trovarono Dyson, Feynman, Schwinger e Tomonaga. Questa procedura consiste nell’assorbire, cioè “spostare”, le divergenze (gli infiniti) contenuti nelle predizioni teoriche, nei parametri liberi della teoria. Questi ultimi diventano essi stessi infiniti e non hanno più alcun significato fisico, ma mantengono una relazione con quantità direttamente misurabili. Quando il risultato del calcolo di un processo fisico viene espresso in termini di quantità misurabili, anziché mediante i parametri liberi del modello teorico (che, come detto sopra, non hanno più alcun significato fisico), le divergenze si cancellano esattamente. Affinché la teoria sia rinormalizzabile, è però necessario che tutti gli infiniti che compaiono nei calcoli possano essere assorbiti in un numero finito di parametri liberi. Questo è proprio ciò che avviene nella Qed, dove è sufficiente esprimere il calcolo di un qualsiasi processo fisico in termini della massa e della carica elettrica dell’elettrone, misurate sperimentalmente, per ottenere un risultato finito. Per quanto bizzarra (o “svitata”, come la definì lo stesso Feynman) possa sembrare, questa procedura, dettata da regole ben precise, funziona perfettamente. Grazie a questa fantastica cancellazione degli infiniti alla fine degli anni ‘40, la Qed acquistò il suo completo potere predittivo e da quasi 70 anni fornisce predizioni in splendido accordo con gli esperimenti. Nel 1971, Gerardus ‘t Hooft e Martinus Veltman riuscirono a dimostrare che

anche le interazioni deboli sono rinormalizzabili, aprendo così la strada del successo alla meravigliosa sintesi del modello standard, l’elegante ed economica teoria delle interazioni elettromagnetiche, forti e deboli. Non tutte le teorie quantistiche dei campi sono però rinormalizzabili. In molti casi, infatti, gli infiniti non possono essere assorbiti in un numero limitato di parametri liberi e la teoria, detta “non rinormalizzabile”, contiene divergenze ultraviolette. Nonostante questo, le teorie non rinormalizzabili hanno un ruolo fondamentale nella nostra descrizione della natura. In particolare, la teoria quantistica della gravità non è rinormalizzabile (vd. p. 13, ndr). Mentre in passato la rinormalizzabilità di una teoria veniva considerata un principio fondamentale, oggi pensiamo che ogni teoria di campo quantistica realistica delle interazioni fondamentali contenga parti rinormalizzabili, come nel modello standard, e altre che non lo sono e che sono divergenti ad alta energia. Per ora, gli infiniti che non riusciamo a spazzare sotto il tappeto parametrizzano la nostra ignoranza della fisica a queste alte energie. Non ci resta che continuare la nostra ricerca verso energie sempre più grandi.

c. Rappresentazione schematica del concetto di rinormalizzazione. La prima sfera rappresenta la somma di tutte le energie che possono assumere le particelle virtuali in un determinato processo. Anche se questa somma è infinita, può essere rinormalizzata sottraendogli la somma di tutte le energie superiori a un certo valore, indicato con λ in figura (seconda sfera). Quello che resta, la terza sfera in figura, è una quantità ben definita. La terza sfera descrive quantità misurabili che accadono nei processi fisici a energie inferiori di λ , mentre la prima e la seconda, che hanno valore infinito, descrivono quantità o parametri non misurabili.

Biografia

Massimo Passera è ricercatore dell’Infn della sezione di Padova. Fisico teorico delle particelle elementari, si occupa principalmente di test di precisione del modello standard.