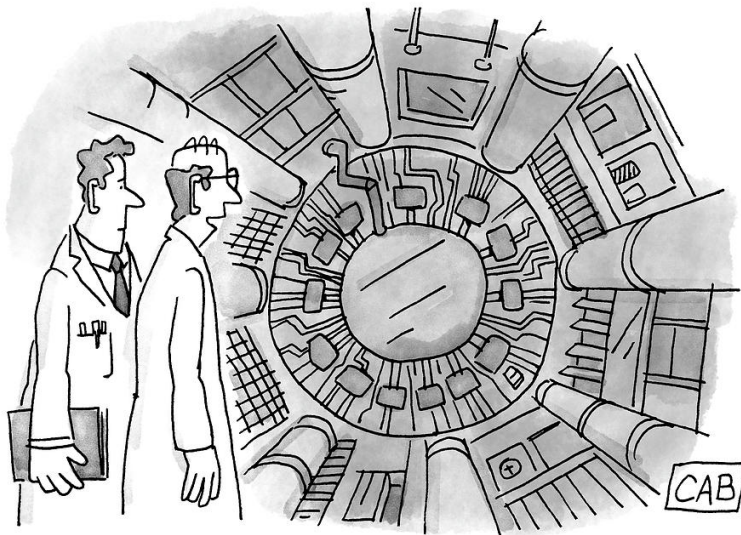


Appunti del corso di:
**Fenomenologia delle Interazioni
Fondamentali**

Gabriele Cembalo

A.A. 2025-2026



"Once you have a collider, every problem starts to look like a particle."

Università degli Studi di Torino
Dipartimento di Fisica
Via Giuria, 1, Torino (TO)

There is a sort of mythology
that grows up about what
happened, which is different
from what really did happen.

Peter Higgs

Prefazione

In questo documento intendo raccogliere i miei appunti basati sul materiale del corso “**Fenomenologia delle Interazioni Fondamentali**”, tenuto dal Prof. P. Torrielli, S. Ucciratie seguito presso *Università degli Studi di Torino* nell’anno accademico 2025-2026. Durante il corso sono stati consigliati diversi libri (elencati in Bibliografia), e cercherò di indicare i riferimenti corrispondenti all’inizio di ogni capitolo.

Questi appunti sono una versione riscritta delle note prese durante le lezioni, quindi la fonte principale è il materiale presentato dal professore. Declino ogni responsabilità e rimando alla mia pagina personale per saperne di più su questi appunti e sul fatto che non hanno natura né accademica né didattica.

Questi appunti vanno chiaramente intesi come note personali, riscritte in maniera ordinata a partire dalle lezioni. Eventuali sviste, errori o imprecisioni sono dovuti ai miei limiti. Inoltre, ho scritto questi appunti principalmente per “spiegare” a me stesso l’argomento, quindi alcune sezioni potrebbero risultare eccessivamente dettagliate o, al contrario, troppo superficiali a seconda del lettore. In ogni caso, spero possano comunque essere utili a qualcuno. Spero anche di essere riuscito a produrre un documento chiaro e ben strutturato.

Una raccolta dei miei appunti è disponibile sulla mia pagina GitHub personale: gCembalo.github.io.

Eventuali errori o refusi possono essere segnalati alla mia email personale: g70362643@gmail.com.

Last update: 29/08/2026

Indice

	Teoria Elettrodebole	1
1	Interazioni fondamentali e particelle elementari	3
1.1	Settore leptónico	5
1.2	Settore adronico	5
1.3	Breve storia delle interazioni deboli	6
2	La teoria di Fermi	9
2.1	Introduzione e intuizioni sulla lagrangiana	9
2.2	Violazione della parità	12
2.2.1	Verifica sperimentale della violazione della parità	16
2.3	Corrente debole leptonica	17
2.3.1	Simmetria $SU(2)$	19
2.4	Corrente debole adronica	22
2.4.1	Corrente con i quark e simmetria $SU(3)$	25
2.4.2	Problema delle FCNC	26
2.4.3	Meccanismo GIM	28
2.5	Larghezze di decadimento e sezioni d'urto	30
2.5.1	Larghezza di decadimento	30
2.5.2	Sezioni d'urto di processi di scattering	31
2.5.3	Decadimento del muone	32
2.6	Unitarietà della matrice S	32
2.6.1	Violazione dell'unitarietà della teoria di Fermi	35
2.7	Rinormalizzazione	41
2.7.1	Parametri della lagrangiana e rinormalizzazione	41
2.7.2	Divergenze ultraviolette	44
2.7.3	Rinormalizzazione e power counting	45
2.7.4	Rinormalizzare la teoria di Fermi	48
3	Teoria di gauge elettrodebole	53
3.1	Piano di lavoro	53
3.1.1	Correnti leptoniche e simmetria $SU(2)$	55
3.2	Interazione elettrodebole per correnti leptoniche	57
3.3	Interazione elettrodebole per correnti adroniche	66

3.4	Invarianza di $\mathcal{L}^{lep} + \mathcal{L}^{had}$ per trasformazioni di $SU(2)_L \times U(1)_Y$	71
3.5	Lagrangiana massless per i campi di gauge	74
3.6	Masse dei bosoni di gauge	76
3.6.1	Massa per il bosone $W^\pm$	77
3.6.2	Massa per il bosone Z^0	79
3.6.3	Come dare massa ai bosoni di gauge	80
3.7	Rottura spontanea di simmetria	84
3.7.1	Meccanismo di Higgs per i bosoni $W^\pm$ e Z	85
3.7.2	Necessità del gauge fixing	89
3.7.3	Propagatori per i bosoni di gauge	93
3.7.4	Parametri liberi della lagrangiana	96
3.7.5	Masse per i fermioni - Quark	97
3.7.6	Masse per i fermioni - Leptoni	104
3.7.7	Ancora sulla matrice CKM	108
3.8	Accoppiamenti del bosone di Higgs	110
3.9	Numero di neutrini leggeri	112
Cromodinamica quantistica		118
4	Introduzione storica	121
4.1	Modello a quark	121
4.2	Dubbi e introduzione del colore	123
4.3	Libertà asintotica e teoria di gauge	125
5	Fondamenti di QCD	129
5.1	Il gruppo $SU(N)$	129
5.2	Lagrangiana della QCD	133
5.2.1	Regole di Feynman	137
5.3	Running coupling e libertà asintotica	140
6	QCD in scattering process	155
6.1	Scattering process $e^+e^- \rightarrow$ Adroni	155
6.1.1	Rappresentazione intuitiva dello scattering	155
6.1.2	$e^+e^- \rightarrow$ adroni al Leading Order (LO)	159
6.1.3	Commenti sui diagrammi del processo (tagli)	163
6.1.4	Nomenclatura: Born e tree-level	167
6.1.5	Calcoli al Next-to-leading-order (NLO) e divergenze IR	169
6.1.6	Correzioni NLO: contributo reale	175
6.1.7	Correzioni NLO: contributo virtuale	186
6.1.8	Correzione del vertice ad un loop	189
6.2	Jets adronici	191
6.2.1	Introduzione	192
6.2.2	Sterman-Weinberg jets	194

6.2.3	Teoria delle Perturbazioni nel regime IR	198
6.3	Distribuzioni cinematiche	202
6.3.1	Sezioni d'urto cumulative	203
6.3.2	Osservabili IR-safe	204
6.3.3	Distribuzione della Thrust in eventi $q\bar{q}g$	208
6.4	Deep Inelastic Scattering	214
6.4.1	Modello a partoni	214
6.4.2	Parton distribution function	217
6.4.3	Deep inelastic neutrino scattering	221
6.4.4	Vincoli sulle PDF	223
6.4.5	Violazione del Bjorken scaling ed equazione DGLAP	228
6.5	Collisioni adroniche	239
6.5.1	Processo di Drell-Yan	240
6.5.2	Produzione di jet (cenni)	245
7	Anomalie	247
7.1	Anomalia assiale	247
7.1.1	Calcolo ad 1-loop	250
7.2	Anomalie nel SM e loro cancellazioni	253
7.2.1	Studio del gruppo $SU(2) \times U(1)$	254
7.2.2	Studio del gruppo $SU(3)$	256
	Appendici	258
A	Relazioni utili per accoppiamenti elettrodeboli	261
B	Scattering $e^+e^- \rightarrow f\bar{f}$ massless	263
C	Fattore di colore e flusso	271
D	Note sui contributi reale e virtuale a $e^+e^- \rightarrow$ adroni	277
D.1	Regolatori dimensionali IR e UV	277
D.2	Contributo reale in gauge di cono di luce	279
E	Algoritmo JADE per Jet da $e^+e^- \rightarrow$ adroni	283
F	Conti Thrust	287
F.1	Conti valore T in $e^+e^- \rightarrow q\bar{q}g$	287
F.2	Conti per l'integrale	288
F.3	Approfondimento Thrust all'endpoint	290
G	Approfondimenti DIS	293
G.1	Momentum sum rule	293
G.2	Prodotto di convoluzione	294

Teoria Elettrodebole

Capitolo 1

Interazioni fondamentali e particelle elementari

Come già sappiamo le interazioni fondamentali, attualmente note, sono:

- **Debole**, mediata dai bosoni vettori, di spin = 1, W^+ , W^- , Z massivi. In particolare si sono misurate le masse:

$$m_{W^\pm} = 80.377 \pm 0.012 \text{ GeV}$$

$$m_Z = 91.1876 \pm 0.0021 \text{ GeV}.$$

- **Elettromagnetica**, mediata dal fotone, di spin = 1, a massa nulla.
- **Forte**, mediata dai gluoni g^a (con $a = 1, \dots, 8$), a massa nulla.
- **Gravitazionale**, mediata dal gravitone (forse), di spin = 2, a massa nulla $m_G = 0$.

Sappiamo che tutte le particelle massive partecipano all'interazione gravitazionale, tuttavia, nei processi considerati in questo corso, l'effetto dell'interazione gravitazionale fra le particelle è del tutto trascurabile rispetto all'intensità delle altre interazioni.

I bosoni mediatori delle interazioni sono tutte particelle elementari, cioè puntiformi e non dotate di struttura interna.

Ovviamente possiamo classificare anche le particelle, e lo facciamo in base alle interazioni a cui partecipano:

- **Leptoni**: non hanno interazioni forti.
- **Adroni**: hanno interazioni forti.

Come sappiamo da corsi precedenti il Modello Standard delle particelle elementari è organizzato come nello schema [1.1](#).

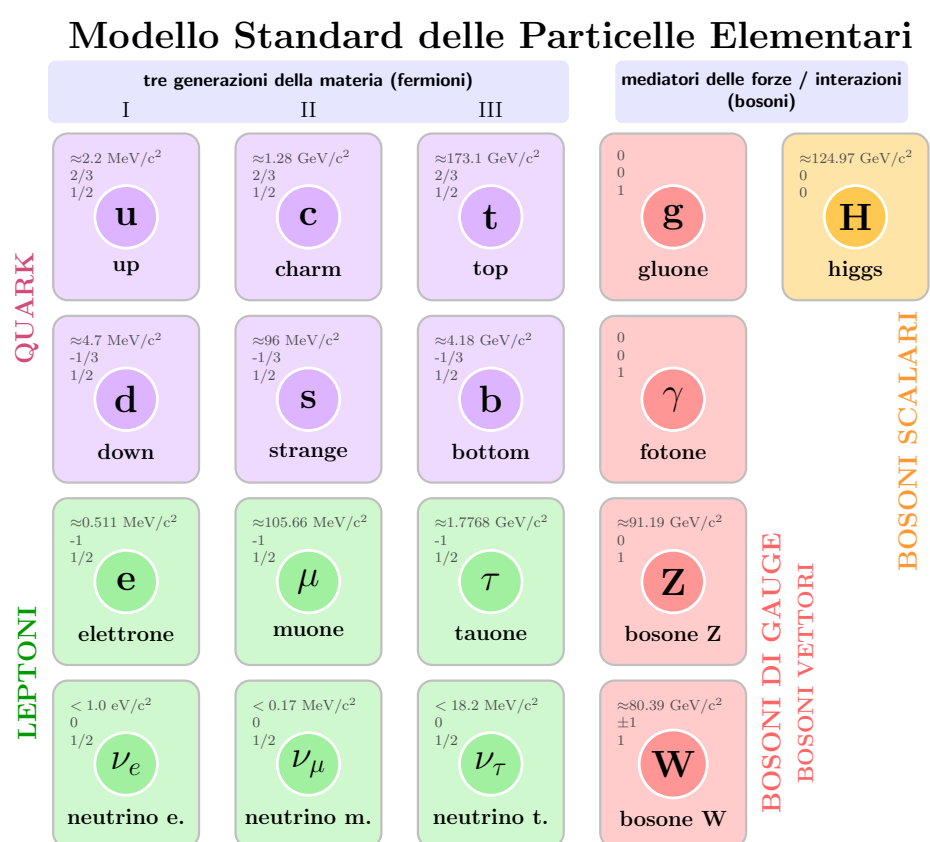


Figura 1.1

1.1 Settore leptonico

I leptoni conosciuti sono:

$$\begin{array}{llll} e^- & , & \nu_e & ; \quad \mu^- & , & \nu_\mu & ; \quad \tau^- & , & \nu_\tau \\ e^+ & , & \bar{\nu}_e & ; \quad \mu^+ & , & \bar{\nu}_\mu & ; \quad \tau^+ & , & \bar{\nu}_\tau. \end{array}$$

I leptoni hanno tutti spin 1/2 e sono considerati puntiformi, senza struttura interna. Sappiamo anche che, per essere consistente, la teoria deve contenere 3 *famiglie* (o generazioni) leptoniche:

$$\begin{array}{ccc} \begin{pmatrix} \nu_e \\ e^- \end{pmatrix} & , & \begin{pmatrix} \nu_\mu \\ \mu^- \end{pmatrix} & , & \begin{pmatrix} \nu_\tau \\ \tau^- \end{pmatrix} & + \text{ antiparticelle.} \\ L_e = +1 & & L_\mu = +1 & & L_\tau = +1 \end{array}$$

Le particelle leptoniche, ossia elettrone, muone e tau, hanno stessi numeri quantici, ma masse diverse:

$$\begin{aligned} m_e &= 511 \text{ KeV} \\ m_\mu &= 106 \text{ MeV} \\ m_\tau &= 1,777 \text{ GeV} \end{aligned}$$

mentre i neutrini, nel modello che sviluppiamo noi, sono massless ed esistono solo nella chiralità left.

1.2 Settore adronico

Nel corso del tempo si sono cominciate ad osservare, nei primi acceleratori di particelle, un numero molto alto di particelle adroniche, ovvero che interagivano anche tramite la forza forte. Con il crescere del numero di rivelazioni si è cominciato a pensare di non riuscire a catalogarle tutte, ma con la piacevole scoperta che in realtà esse non fossero elementari, bensì formate da particelle più piccole, i *quark*, elementari.

Gli adroni li raggruppiamo in:

- **Barioni**, hanno spin semi-intero, sono dunque fermioni, hanno un numero barionico non nullo e sono formati da 3 quark qqq (sono ad esempio p, n, Δ, Λ).
- **Mesoni**, hanno spin intero, sono bosoni, hanno numero barionico nullo, $B = 0$ e sono formati da quark e anti-quark $q\bar{q}$ (sono ad esempio $\pi, \rho, K, D, B, \eta, J/\psi$).

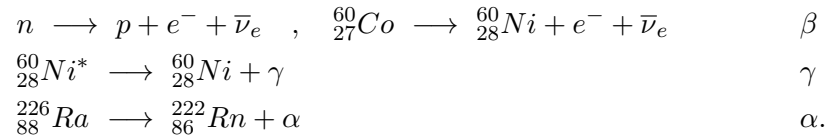
I barioni potremmo considerarli come degli *atomi di quark*, i quali, oltre ad essere a spin = 1/2, possiamo raggrupparli in 3 doppietti (similmente a quelli leptonici):

$$\begin{pmatrix} u \\ d \end{pmatrix} \quad , \quad \begin{pmatrix} c \\ s \end{pmatrix} \quad , \quad \begin{pmatrix} t \\ b \end{pmatrix} \quad \begin{array}{l} \longleftarrow Q = 2/3 \\ \longleftarrow Q = -1/3 \end{array} \quad + \text{ antiquark.}$$

I quark u, d, s sono i tre con massa più piccola, e a causa del fenomeno del confinamento, è molto complicato ricavare le loro masse (in realtà anche degli altri) a partire dalle masse delle particelle non elementari.

1.3 Breve storia delle interazioni deboli

nel 1914 Chadwick dimostrò che lo spettro del decadimento β era continuo, a differenza di quanto pensato e a differenza dello spettro di decadimento α e γ . Alcuni esempi di decadimenti sono:



Inizialmente l'esistenza del neutrino ν non era conosciuta, anche perché per via della debolezza delle sue interazioni era molto difficile rivelarlo, per cui non si capiva molto bene perché lo spettro del decadimento β dovesse essere continuo e non discreto. Nel 1930 W. Pauli, per risolvere il problema dello spettro, e della statistica degli spin del decadimento, ipotizza l'esistenza di una particella neutra, fermionica, che sente l'interazione debole e che partecipa al processo spartendosi l'impulso dello stato finale con l' e^- , rendendo lo spettro continuo. Inizialmente diede il nome *neutrone* a queste particelle, ma nel 1932 Chadwick scoprì il neutrone come lo conosciamo oggi, così fermi rinominò *neutrino* la particella ipotizzata da Pauli. La prima volta che fu nominato il neutrino è stata alla conferenza di Solvay nel 1933, e nello stesso anno Fermi e Perrin, indipendentemente, ipotizzarono che i neutrini fossero a massa nulla. Nel 1934 ci fu la vera milestone della teoria, perché E. Fermi elaborò la teoria dei decadimenti β , la prima teoria delle interazioni deboli e chiamata *teoria di Fermi*, che studieremo tra pochissimo, e la formulò in analogia all'elettrodinamica quantistica (QED).

L'interazione debole è responsabile, in generale, dei processi di decadimento:

- Decadimento del neutrone:

$$n \longrightarrow p + e^- + \bar{\nu}_e. \quad (1.3.1)$$

- Decadimento del muone:

$$\mu^- \longrightarrow e^- + \nu_\mu + \bar{\nu}_e. \quad (1.3.2)$$

- Decadimento dei pioni carichi:

$$\pi^+ \longrightarrow \mu^+ + \nu_\mu. \quad (1.3.3)$$

- Decadimenti β -nucleari, ossia decadimenti β , ma all'interno del nucleo, come ad esempio:

$$N(Z, N) \longrightarrow N(Z + 1, N - 1) + e^- + \bar{\nu}_e \quad (1.3.4)$$

chiamato decadimento β^- e in cui all'interno del nucleo avviene (1.3.1). Il decadimento β^+ è:

$$N(Z, N) \longrightarrow N(Z - 1, N + 1) + e^+ + \nu_e \quad (1.3.5)$$

in cui avviene il processo inverso di (1.3.1), ossia:

$$p \longrightarrow n + e^+ + \nu_e \quad (1.3.6)$$

che non potrebbe avvenire per un protone libero per motivi puramente energetici, essendo p più leggero di n e non potendo decadere in una particella più pesante.

L'interazione debole è anche responsabile di processi di scattering, come:

- Le reazioni indotte dai neutrini:

$$\nu_\mu + e^- \longrightarrow \mu^- + \nu_e. \quad (1.3.7)$$

- Cattura del μ^- da parte di un nucleo:

$$\begin{aligned} \mu^- + N(Z, N) &\longrightarrow N(Z - 1, N + 1) + \nu_\mu \\ (\mu^- + p &\longrightarrow n + \nu_\mu \quad \text{nel nucleo}). \end{aligned} \quad (1.3.8)$$

- Formazione di un deuterone (che si intende una particella formata da $d = n + p$) da due protoni:

$$p + p \longrightarrow d + e^+ + \nu_e \quad (1.3.9)$$

la quale innesca il principale ciclo di fusione nel Sole.

- Il processo:

$$e^+ + e^- \longrightarrow \mu^+ + \mu^- \quad (1.3.10)$$

che può essere mediato sia da un fotone A_μ che da un bosone Z_μ , ed è stato il processo con cui si è misurato lo Z_μ .

Notiamo però che per tutti i processi deboli si osserva sempre la conservazione del numero leptonico (L_e, L_μ, L_τ) per famiglia.

Non abbiamo ancora parlato dell'intensità delle interazioni deboli, motivo per cui viene chiamata debole. I processi di decadimento di tipo debole sono

più lunghi di quelli dell'interazione forte ed elettromagnetica, in particolare si ha:

$$\text{int. forte} \quad \tau \sim 10^{-23} \text{ s} \quad (1.3.11)$$

$$[\text{ex. } \rho \longrightarrow 2\pi]$$

$$\text{int. EM} \quad \tau \sim 10^{-16} \text{ s} \quad (1.3.12)$$

$$[\text{ex. } \pi^0 \longrightarrow 2\gamma]$$

$$\text{int. debole} \quad \tau \sim 10^{-8} - 10^{-6} \text{ s} \quad (1.3.13)$$

$$[\text{ex. } \tau(\mu^- \longrightarrow e^- + \nu_\mu + \bar{\nu}_e = 2.2 \cdot 10^{-6} \text{ s})]$$

$$[\text{ex. } \tau(\pi^+ \longrightarrow \mu^+ + \nu_\mu = 2.6 \cdot 10^{-8} \text{ s})]$$

$$(1.3.14)$$

in cui indichiamo con τ la vita media. Vediamo dunque che le interazioni forti ed elettromagnetiche hanno vita media più piccola, e siccome $\tau = \hbar/\Gamma$, in cui Γ è la lunghezza di decadimento e per cui $\Gamma \propto |A|^2$, in cui A è l'ampiezza di decadimento, se abbiamo una vita media piccola, allora avremo una larghezza di decadimento grande. Dunque, deduciamo che Γ_{weak} è molto più piccola di quella forte ed elettromagnetica. In più se pensiamo al fatto che Γ sia proporzionale ad una qualche potenza della costante di accoppiamento della teoria, allora potremmo dire che l'interazione debole ha una costante piccola; però, vedremo che non è la costante di accoppiamento a rendere debole l'interazione, bensì il fatto che la forza sia mediata da bosoni massivi.

La massa dei bosoni di gauge renderà notevolmente diverse la QCD e l'EW, nonostante siano entrambe teorie di gauge non abeliane.

Detto questo ci aspettiamo che le interazioni deboli diventino rilevanti, dunque riusciamo ad apprezzarle bene, a regimi di bassa energia, in cui le altre interazioni non si manifestano. Siamo nel regime della forza debole quando siamo nell'ordine delle masse di W e Z .

Capitolo 2

La teoria di Fermi

Vediamo in questa sezione tutta la trattazione, pregi e difetti, della teoria di Fermi per le interazioni deboli.

2.1 Introduzione e intuizioni sulla lagrangiana

Come già accennato, Fermi si ispirò alla lagrangiana della QED:

$$\mathcal{L}_{QED} = -\frac{1}{4}F_{\mu\nu}F^{\mu\nu} + \bar{\psi}(i\gamma^\mu\partial_\mu - m)\psi - e\bar{\psi}\gamma^\mu\psi A_\mu \quad (2.1.1)$$

in cui sappiamo che per rendere la simmetria di fase $U(1)$ locale, dobbiamo inserire l'ultimo termine in cui accoppiamo la corrente conservata (globalmente) con un campo di gauge A_μ . Ricordiamo che per una trasformazione di $U(1)$ abbiamo:

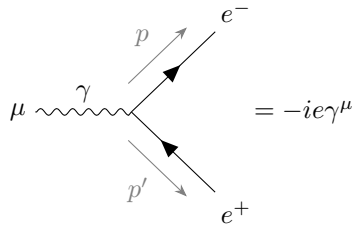
$$\psi(x) \longrightarrow \psi'(x) = e^{ie\alpha(x)}\psi(x) \quad (2.1.2)$$

$$A_\mu(x) \longrightarrow A'_\mu = A_\mu(x) - \partial_\mu\alpha(x) \quad (2.1.3)$$

e notando che per una trasformazione infinitesima si ha $\delta\psi = ie\alpha(x)\psi(x)$, allora la corrente elettromagnetica è:

$$J_{em}^\mu = e\alpha(x)\bar{\psi}\gamma^\mu\psi(x). \quad (2.1.4)$$

Ricordiamoci anche qual è il vertice di interazione della teoria:



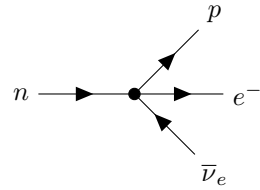
a cui associamo il contributo $-ie\gamma^\mu$ per le regole di Feynman.

Storicamente non c'era ancora l'idea che le interazioni fondamentali fossero determinate dall'invarianza di trasformazioni di gauge locali, però sicuramente scritta una lagrangiana essa dovesse essere invariante per qualche trasformazione. Fermi, quando scrisse la prima lagrangiana della teoria debole nel 1934, basò tutto sulle simmetrie che vedeva nella QED. Infatti notò che J_{em} , contenendo $\bar{\psi}$ e ψ , era invariante di gauge e quindi se si costruiscono termini aggiuntivi, che contengono solo la corrente elettromagnetica, essi non rompono l'invarianza di $\mathcal{L}$.

Inizialmente voleva spiegare il decadimento del neutrone:

$$n \longrightarrow p + e^- + \bar{\nu}_e \quad (2.1.5)$$

e lo fece utilizzando una coppia di (2.1.4), poiché solamente una J_{em}^μ contiene due campi fermionici, ma volendo descrivere un processo con 4 spinori aveva bisogno di 4 campi ψ . Come primo tentativo modellizzò un'interazione di tipo puntiforme (point-like):



$$\mathcal{L}_{int} = \frac{G_F}{\sqrt{2}} \bar{\psi}_p \Gamma'^\mu \psi_n \bar{\psi}_e \Gamma_\mu \psi_{\nu_e} \quad (2.1.6)$$

in cui mise le due correnti:

$$J_L^\mu = \bar{\psi}_{\nu_e} \Gamma^\mu \psi_e \quad (2.1.7)$$

$$J_N^\mu = \bar{\psi}_p \Gamma'^\mu \psi_n \quad (2.1.8)$$

dove le matrici Γ^μ e Γ'^μ , sono in generale diverse, e rappresentano un modo generalizzato di scrivere una corrente¹ di tipo QED, in cui $\Gamma^\mu = \Gamma'^\mu = \gamma^\mu$.

Le matrici Γ^μ sono parte del set completo di forme bilineari per gli spinori di Dirac:

$\bar{\psi}\psi$	(1)	scalare
$\bar{\psi}\gamma^\mu\psi$	(4)	vettore (o vettore polare)
$\bar{\psi}\sigma^{\mu\nu}\psi$	(6)	tensore di rango 2
$\bar{\psi}\gamma^\mu\gamma^5\psi$	(4)	pseudo-vettore (o vettore assiale)
$\bar{\psi}\gamma^5\psi$	(1)	pseudo-scalare.

¹Sono i famosi invarianti bilineari di spinori, che abbiamo visto al corso di *Introduzione alla QFT*.

Sapendo che $[\psi] = M^{3/2}$ concludiamo subito che la dimensione di $\mathcal{L}$ sarebbe M^6 a meno che $[G_F] = M^{-2}$, il quale però sarà la causa della violazione dell'unitarietà e quindi renderà non rinormalizzabile la teoria di Fermi.

Ad ogni modo, il valore di G_F si deduce dai risultati sperimentali, poiché misurando la vita media di n , nel suo decadimento:

$$\tau(n) = 885.7 \pm 0.8 \text{ s} \quad (\sim 15 \text{ minuti}) \quad (2.1.9)$$

si può ricavare:

$$G_F = 1.16639 \cdot 10^{-5} \text{ GeV}^{-2} \simeq 10^{-5} m_p^{-2} \quad (2.1.10)$$

con $m_p = 0.933 \text{ GeV}$. Successivamente con la scoperta del μ , che compie un decadimento β :

$$\mu^- \longrightarrow e^- + \nu_\mu + \bar{\nu}_e \quad (2.1.11)$$

si è riusciti a misurare G_F con un errore più piccolo grazie a:

$$\tau(\mu) = (2.19703 \pm 0.00004) \cdot 10^{-6} \text{ s}. \quad (2.1.12)$$

Dopo molti sforzi teorici, a metà degli anni 50' si arrivò alla scrittura della lagrangiana di Fermi $\mathcal{L}_F$ in una forma generale:

$$\mathcal{L}_F = \frac{G_F}{\sqrt{2}} J^{\mu\dagger}(x) J_\mu(x) \quad (2.1.13)$$

valida per tutti i processi deboli, per i quali si assume stessa costante di accoppiamento. Il fattore $\sqrt{2}$ si mette per ragioni storiche e J^μ è la corrente debole. Abbiamo inserito la *corrente debole* $J^\mu(x)$, in analogia con la corrente EM, che in generale possiamo scrivere come:

$$J^\mu(x) = L^\mu(x) + H^\mu(x) \quad (2.1.14)$$

in cui indichiamo con $L^\mu(x)$ la corrente debole *leptonica* e con $H^\mu(x)$ la corrente debole *adronica*. Questa classificazione per la corrente debole permette di raggruppare i processi deboli in 3 tipologie:

- Processi *puramente leptonici*, controllati da:

$$\mathcal{L}_F^{(L)} = \frac{G_F}{\sqrt{2}} L^{\mu\dagger}(x) L_\mu(x) \quad (2.1.15)$$

come ad esempio il decadimento β del μ :

$$\mu^- \longrightarrow e^- + \bar{\nu}_e + \nu_\mu.$$

- Processi *semi-leptonici*, controllati da:

$$\mathcal{L}_F^{(SL)} = \frac{G_F}{\sqrt{2}} \left(L^{\mu\dagger}(x) H_\mu(x) + H^{\mu\dagger} L_\mu \right) \quad (2.1.16)$$

come ad esempio il decadimento β del neutrone:

$$n \longrightarrow p + e^- + \bar{\nu}_e.$$

- Processi *puramente adronici*, controllati da:

$$\mathcal{L}_F^{(H)} = \frac{G_F}{\sqrt{2}} H^{\mu\dagger}(x) H_\mu(x) \quad (2.1.17)$$

come ad esempio il decadimento della Λ :

$$\Lambda^0 \longrightarrow p + \pi^-.$$

2.2 Violazione della parità

Noi siamo abbastanza soddisfatti di quello trovato, ma non abbiamo idea di come siano fatte le matrici Γ^μ e Γ'^μ . Per capirlo dobbiamo, dopo un rapido ripasso, cosa implichi la violazione della parità nelle interazioni deboli.

Concentriamoci per primo sul settore leptonico² e ricordiamo alcune cose, che conosciamo già sulla chiralità degli spinori di Dirac.

L'operatore di *chiralità* γ_5 è definito da:

$$\gamma_5 = i\gamma^0\gamma^1\gamma^2\gamma^3 = -\frac{i}{4!}\epsilon^{\mu\nu\rho\sigma}\gamma_\mu\gamma_\nu\gamma_\rho\gamma_\sigma \quad (2.2.1)$$

che soddisfa:

$$\gamma_5^\dagger = \gamma_5 \quad ; \quad \gamma_5^2 = \mathbb{1} \quad ; \quad \{\gamma_5, \gamma^\mu\} = 0. \quad (2.2.2)$$

Sappiamo che uno spinore di Dirac si può scomporre in due spinori di chiralità definita, cioè in due autostati di γ_5 con autovalori ± 1 , utilizzando gli operatori di proiezione:

$$P_{L,R} = \frac{1 \mp \gamma_5}{2} \quad (2.2.3)$$

$$P_R^2 = P_R \quad ; \quad P_L^2 = P_L \quad ; \quad P_R P_L = P_L P_R = 0 \quad (2.2.4)$$

che dati gli spinori $\psi_{L,R}$, autostati di γ_5 , si possono scrivere come:

$$\psi_{L,R} = P_{L,R}\psi \quad (2.2.5)$$

²Diamo questa priorità ai processi leptonici visto che sono misurati molto più facilmente e con più precisione; ciò è dato dalla loro maggiore "pulizia", visto che L^μ tratta particelle elementari, a differenza di H^μ che tratta particelle non elementari e sarà un miscuglio di interazioni deboli e forti che sentono i costituenti.

per cui valgono:

$$\gamma_5 \psi_R = \psi_R \quad (2.2.6)$$

$$\gamma_5 \psi_L = -\psi_L \quad (2.2.7)$$

e si ha anche che:

$$P_R \psi_R = \psi_R \quad ; \quad P_L \psi_L = \psi_L \quad (2.2.8)$$

$$P_R \psi_L = 0 \quad ; \quad P_L \psi_R = 0. \quad (2.2.9)$$

Detto tutto questo possiamo vedere che gli spinori di Dirac si possono ridurre a due spinori di Weyl con chiralità definita:

$$\psi = \psi_L + \psi_R. \quad (2.2.10)$$

Oltre l'operatore di chiralità sappiamo esistere anche l'operatore di *elicità*, il quale potrebbe sembrare del tutto scollegato, ma in realtà non è così. L'elicità è definito come la proiezione dello spin lungo la direzione della quantità di moto $\vec{p}$:

$$\hat{h} = \frac{\vec{s} \cdot \vec{p}}{s|\vec{p}|} = \frac{\vec{\Sigma} \cdot \vec{p}}{|\vec{p}|} \quad (2.2.11)$$

dove:

$$\Sigma^k = \gamma^0 \gamma^k \gamma^5 \quad (2.2.12)$$

per il quale vale:

$$[\vec{\Sigma}, \gamma_5] = 0. \quad (2.2.13)$$

Calcoliamo il valor medio dell'elicità (2.2.11) in uno stato di chiralità definita, per esempio lo spinore u_L :

$$\left\langle \frac{\vec{\Sigma} \cdot \vec{p}}{|\vec{p}|} \right\rangle_L = \frac{\sum_r u_L^{\dagger(r)}(p) \frac{\vec{\Sigma} \cdot \vec{p}}{|\vec{p}|} u_L^{(r)}(p)}{\sum_r u_L^{\dagger(r)}(p) u_L^{(r)}(p)} \quad (2.2.14)$$

in cui sommiamo sugli spin r . Per il denominatore abbiamo:

$$\sum_r u_L^{\dagger(r)}(p) u_L^{(r)}(p) = \sum_r u^{\dagger(r)}(p) \left(\frac{1 - \gamma_5}{2} \right) \left(\frac{1 - \gamma_5}{2} \right) u^{(r)}(p) \quad (2.2.15)$$

$$\begin{aligned} & \text{utilizziamo } \left(\frac{1 - \gamma_5}{2} \right) \left(\frac{1 - \gamma_5}{2} \right) = \frac{1 - \gamma_5}{2} \\ & = \sum_r u^{\dagger(r)}(p) \frac{1 - \gamma_5}{2} u^{(r)}(p) \end{aligned} \quad (2.2.16)$$

$$= \sum_r \bar{u}^{(r)}(p) \gamma^0 \frac{1 - \gamma_5}{2} u^{(r)}(p) \quad (2.2.17)$$

$$\text{usiamo l'equazione: } \sum_r u^{(r)}(p) \bar{u}^{(r)}(p) = \not{p} + m$$

$$= \frac{1}{2} \text{Tr}\{(\not{p} + m) \gamma_0 (1 - \gamma_5)\} \quad (2.2.18)$$

$$= \frac{1}{2} \text{Tr}\{(\not{p} \gamma^0) - m(\gamma^0 \gamma_5)\} \quad (2.2.19)$$

a questo punto utilizzando le relazioni:

$$\text{Tr}\{\not{p} \not{q}\} = 4p \cdot q \quad (2.2.20)$$

$$\text{Tr}\{\gamma^\mu \gamma_5\} = -\text{Tr}\{\gamma_5 \gamma^\mu\} = 0 \quad (2.2.21)$$

possiamo concludere che il denominatore vale:

$$\sum_r u_L^{\dagger(r)}(p) u_L^{(r)}(p) = 2E. \quad (2.2.22)$$

Per quanto riguarda il numeratore il conto è simile:

$$\sum_r u_L^{\dagger(r)}(p) \frac{\vec{\Sigma} \cdot \vec{p}}{|\vec{p}|} u_L^{(r)}(p) = \sum_r \bar{u}^{(r)}(p) \gamma^0 \frac{1 - \gamma_5}{2} \frac{\vec{\Sigma} \cdot \vec{p}}{|\vec{p}|} \frac{1 - \gamma_5}{2} u^{(r)}(p) \quad (2.2.23)$$

$$= \sum_r \bar{u}^{(r)}(p) \gamma^0 \frac{\vec{\Sigma} \cdot \vec{p}}{|\vec{p}|} \frac{1 - \gamma_5}{2} u^{(r)}(p) \quad (2.2.24)$$

$$= \frac{1}{2} \text{Tr} \left\{ (\not{p} + m) \gamma^0 \frac{\vec{\Sigma} \cdot \vec{p}}{|\vec{p}|} \frac{1 - \gamma_5}{2} \right\} \quad (2.2.25)$$

$$= \frac{1}{2} \text{Tr} \left\{ (\not{p} + m) \gamma^0 \frac{\gamma^0 \vec{\gamma} \gamma^5 \cdot \vec{p}}{|\vec{p}|} \frac{1 - \gamma_5}{2} \right\} \quad (2.2.26)$$

$$= \frac{\vec{p}}{2|\vec{p}|} \text{Tr} \{ (\not{p} + m) \vec{\gamma} \gamma_5 (1 - \gamma_5) \} \quad (2.2.27)$$

$$= -\frac{\vec{p}}{2|\vec{p}|} \text{Tr} \{ (\not{p} + m) \vec{\gamma} (1 - \gamma_5) \} \quad (2.2.28)$$

$$= -\frac{\vec{p}}{2|\vec{p}|} \text{Tr} \{ \not{p} \vec{\gamma} - m \vec{\gamma} \gamma_5 \} \quad (2.2.29)$$

$$= -\frac{\vec{p}}{2|\vec{p}|} (\text{Tr} \{ \not{p} \vec{\gamma} \} - m \text{Tr} \{ \vec{\gamma} \gamma_5 \}) \quad (2.2.30)$$

$$= -\frac{\vec{p}}{2|\vec{p}|} \cdot 4\vec{p} \quad (2.2.31)$$

$$= -2|\vec{p}| \quad (2.2.32)$$

in cui in (2.2.29) abbiamo utilizzato il fatto che solo un numero pari di matrici γ ha traccia non nulla.

Abbiamo quindi ottenuto che il valor medio dell'elicità per uno stato di chiralità definita sinistrorsa è:

$$\left\langle \frac{\vec{\Sigma} \cdot \vec{p}}{|\vec{p}|} \right\rangle_L = -\frac{|\vec{p}|}{E} = -\beta \quad (2.2.33)$$

e in modo analogo si ottiene che il valor medio dell'elicità per uno stato di chiralità definita destrorsa (cambiarebbe solamente il fatto che in (2.2.28) ci sarebbe un termine $(1 + \gamma_5)$ e dunque non cambierebbe un segno globale.) è:

$$\left\langle \frac{\vec{\Sigma} \cdot \vec{p}}{|\vec{p}|} \right\rangle_R = +\frac{|\vec{p}|}{E} = +\beta. \quad (2.2.34)$$

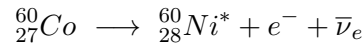
Per particelle a massa nulla vale $\beta = 1$, quindi gli autovalori di chiralità (γ_5) coincidono con quelli di elicità ($\gamma_5 u_L = -1$ e $\gamma_5 u_R = +1$).

La differenza di segno tra (2.2.33) e (2.2.34) pose la domanda se le interazioni deboli conservassero o meno la parità.

2.2.1 Verifica sperimentale della violazione della parità

La proposta che la parità fosse violata nelle interazioni deboli (si era osservato che fosse conservata nelle interazioni forti ed elettromagnetiche) fu proposta nel 1956 da Lee e Young. I due scienziati convinsero Madame Wu ad intraprendere un esperimento criogenico, con un decadimento β^- , per provare la violazione di parità. L'esperimento andò a buon fine ed immediatamente Lee e Young (ma non Madame Wu) vinsero il premio Nobel.

L'esperimento voleva misurare l'asimmetria della distribuzione angolare dei prodotti del decadimento β^- con nuclei polarizzati di ^{60}Co :



in cui $^{60}_{27}\text{Co}$ ha spin 5, mentre $^{60}_{28}\text{Ni}^*$ spin 4. Si ipotizzavano gli elementi del processo fermi, per cui che non ci fosse momento angolare orbitale, e quindi che si dovesse conservare solo lo spin, oltre che la quantità di moto (motivo per cui e^- e $\bar{\nu}_e$ hanno impulsi back-to-back). Si faccia riferimento alla figura 2.1. Il decadimento avviene preferibilmente con la configurazione di sinistra per $\vec{p}_e, \vec{j}_e$ e $\vec{p}_{\nu}, \vec{j}_{\nu}$ (con $\vec{j}$ lo spin). Quindi viene emesso un e^- con valor medio di elicità $-v_e/c$ ed un $\bar{\nu}_e$ con elicità $+1$. La configurazione di sinistra però non è invariante per trasformazioni di parità, cioè per inversioni spaziali. Infatti, se applichiamo $\mathbb{P}$ alla figura di sinistra otteniamo la figura di destra. L'elicità del $\bar{\nu}_e$ diventa -1 e la proiezione di $\vec{j}_e$ su $\vec{p}_e$ (dunque l'elicità) è $+\beta = v_e/c$. Infatti, per inversione spaziale i vettori impulso cambiano segno, mentre i vettori assiali degli spin restano invarianti.

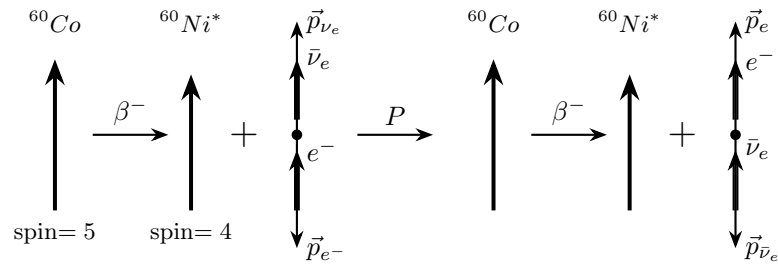


Figura 2.1

Sperimentalmente si osserva la configurazione della figura a sinistra, ma non quella di destra. Dalle osservazioni di numerosi processi di decadimento β si ha che:

- Nei decadimenti β^- vengono sempre emessi un elettrone con valore medio dell'elicità $-v/c$ ed un antineutrino con elicità $+1$.
- Nei decadimenti β^+ vengono sempre emessi un positrone con valore medio dell'elicità $+v/c$ ed un neutrino con elicità -1 .

Queste proprietà implicano che nei decadimenti β la *violazione della parità* è massima³, in quanto l'elettrone e il neutrino partecipano a questi processi con lo spinore ψ_L :

$$\psi_L = \frac{1 - \gamma_5}{2} \psi. \quad (2.2.35)$$

Se la parità fosse conservata, l'elettrone parteciperebbe al decadimento β con lo spinore:

$$\psi = \psi_L + \psi_R \quad (2.2.36)$$

ed il valor medio osservato dell'elicità sarebbe zero, visto che i due estremi sono $\pm\beta$. Invece si osserva asimmetria, dunque *violazione di parità*. Il fatto che si misuri $h = -|\vec{p}|/E = -\beta$ indica che la violazione è massima.

Le interazioni deboli *non* si accoppiano con elettroni destrorsi o antineutrini sinistrorsi. ψ_R non partecipa all'interazione debole leptonica.

NB. Se la parità fosse conservata noi ci aspetteremmo di trovare con il 50% di probabilità entrambe le configurazioni.

2.3 Corrente debole leptonica

Viste le evidenze per le interazioni deboli:

- Violazione massima di parità (solo ψ_L partecipa all'interazione debole).
- Conservazione del numero leptonico (nella singola famiglia).

Possiamo provare a capire chi sono le matrici Γ^μ per la corrente debole leptonica. In generale possiamo scrivere⁴:

$$\Gamma^\mu = \gamma^\mu (a + b\gamma_5) \quad (2.3.1)$$

e se riscriviamo le costanti a e b , che quantificano i contributi delle parti left e right:

$$a = \frac{a_R + a_L}{2}, \quad b = \frac{a_R - a_L}{2} \quad (2.3.2)$$

³Infatti, sperimentalmente se (2.2.34) non si misura si ha massima violazione della parità e si vede che ψ_{l_R} non partecipa alle interazioni deboli. Puoi vedere i capitoli §6.6, 18.2, 18.4, 18.5 dell'Iliopoulos There is a sort of mythology that grows up about what happened, which is different from what really did happen. Iliopoulos per tutti i discorsi riguardanti la scrittura della teoria corrente×corrente di Fermi, la più generica forma della corrente leptonica e per un discorso, sperimentale, riguardante il perché tra tutti gli invarianti bilineari che si possono scegliere, solamente γ^μ (Vettore) e $\gamma^\mu\gamma^5$ (vettore Assiale) sono in accordo con la natura. Per questo la chiamiamo teoria V-A. Eventualmente, c'è l'[articolo](#) in cui ci si chiede, non solo quale siano le forme bilineari da utilizzare in una teoria elettrodebole, ma anche perché tutte le famiglie leptoniche sentono la stessa interazione. Si analizza quella che viene chiamata *universalità*.

⁴Per la QED abbiamo semplicemente $\Gamma^\mu = \gamma^\mu$.

allora:

$$\Gamma^\mu = \gamma^\mu \left(a_L \frac{1 - \gamma_5}{2} + a_R \frac{1 + \gamma_5}{2} \right) \quad (2.3.3)$$

per cui:

$$\bar{\psi} \Gamma^\mu \psi = \bar{\psi} \gamma^\mu \left(a_L \frac{1 - \gamma_5}{2} + a_R \frac{1 + \gamma_5}{2} \right) \psi \quad (2.3.4)$$

$$= \bar{\psi} \gamma^\mu (a_L \psi_L + a_R \psi_R) \quad (2.3.5)$$

ma sapendo che lo spinore ψ_R non partecipa alle interazioni abbiamo che:

$$a_R = 0 \quad (2.3.6)$$

dunque troviamo che:

$$\Gamma^\mu = \frac{a_L}{2} \gamma^\mu (1 - \gamma_5) \quad (2.3.7)$$

in cui possiamo riassorbire il fattore $a_L/2$ nella costante di accoppiamento.

Dell'evidenza sulla conservazione del numero leptonico non ne abbiamo ancora parlato, ma ci sono dei processi, non osservati, che non lo conservano. Infatti, storicamente, oltre la violazione totale della parità, si è osservata la conservazione del numero leptonico (separatamente per singola famiglia). Sperimentalmente quello che si è fatto è stato studiare processi possibili e non possibili (esattamente come per la parità). Ad esempio:

$$(Z, A) \longrightarrow (Z + 2, A) + e^- + e^- \quad (2.3.8)$$

sarebbe cinematicamente favorito a parità di forza di interazione (visto lo spazio delle fasi più grande a disposizione per l' e^- nello stato finale), se avvenissero i processi:

$$\begin{aligned} n &\longrightarrow p + e^- + \bar{\nu}_e \\ \bar{\nu}_e + n &\longrightarrow p + e^- \end{aligned}$$

rispetto al processo osservato:

$$(Z, A) \longrightarrow (Z + 2, A) + e^- + e^- + \bar{\nu}_e + \bar{\nu}_e$$

però per via della conservazione del numero leptonico, che riguarda L_e , L_μ ed L_τ singolarmente⁵, il processo (2.3.8) è proibito.

Altre coppie di processi proibiti/permessi sono:

- Il processo:

$$\begin{aligned} \nu_\mu + (Z, A) &\not\rightarrow (Z - 1, A) + \mu^+ \\ \nu_\mu + (Z, A) &\longrightarrow (Z + 1, A) + \mu^- \end{aligned}$$

⁵In realtà evidenze sperimentali recenti mostrano un piccolo mixing dei numeri leptonici.

- Il processo:

$$\begin{aligned}\nu_e + (Z, A) &\not\rightarrow (Z - 1, A) + e^+ \\ \nu_e + (Z, A) &\longrightarrow (Z + 1, A) + e^-.\end{aligned}$$

- Il processo:

$$\begin{aligned}\mu^+ &\not\rightarrow e^+ + \gamma \\ \text{BR}(\mu^+ \rightarrow e^+ \gamma) &\leq 4 \cdot 10^{-11}.\end{aligned}$$

In base a tutto ciò che abbiamo detto otteniamo per la corrente leptonica la forma:

$$L^\mu = \bar{\psi}_e \gamma^\mu (1 - \gamma_5) \psi_e + \bar{\psi}_\mu \gamma^\mu (1 - \gamma_5) \psi_\mu + \bar{\psi}_\tau \gamma^\mu (1 - \gamma_5) \psi_\tau. \quad (2.3.9)$$

2.3.1 Simmetria $SU(2)$

Consideriamo per semplicità una sola famiglia leptonica:

$$\mathcal{L}_F = \frac{G_F}{\sqrt{2}} L_\mu^\dagger L^\mu \quad (2.3.10)$$

con⁶:

$$L^\mu = \bar{\psi}_\nu \gamma^\mu (1 - \gamma_5) \psi_l \quad (2.3.13)$$

$$L_\mu^\dagger = \bar{\psi}_l \gamma_\mu (1 - \gamma_5) \psi_\nu. \quad (2.3.14)$$

Le correnti scritte sopra sono fatte di leptoni diversi ψ_l , ψ_ν e se organizzati in un doppietto:

$$L = \begin{pmatrix} \psi_\nu \\ \psi_l \end{pmatrix} \quad ; \quad \bar{L} = (\bar{\psi}_\nu \quad \bar{\psi}_l) \quad (2.3.15)$$

ci permettono di riscrivere la corrente leptonica come:

$$L^\mu = \bar{\psi}_\nu \gamma^\mu (1 - \gamma_5) \psi_l \quad (2.3.16)$$

$$= \bar{L} \gamma^\mu (1 - \gamma_5) \begin{pmatrix} 0 & 1 \\ 0 & 0 \end{pmatrix} L \quad (2.3.17)$$

$$L_\mu^\dagger = \bar{\psi}_l \gamma_\mu (1 - \gamma_5) \psi_\nu \quad (2.3.18)$$

$$= \bar{L} \gamma_\mu (1 - \gamma_5) \begin{pmatrix} 0 & 0 \\ 1 & 0 \end{pmatrix} L. \quad (2.3.19)$$

⁶Non è immediato vedere la scrittura di $L_\mu^\dagger$:

$$L_\mu^\dagger = (\bar{\psi}_\nu \gamma^\mu (1 - \gamma_5) \psi_l)^\dagger = \psi_l^\dagger (1 - \gamma^5) \gamma^{\mu\dagger} \gamma^0 \psi_\nu \quad (2.3.11)$$

in cui abbiamo utilizzato $\bar{\psi} = \psi \gamma^0$, $\gamma^{0\dagger} = \gamma^0$ e $\gamma^{k\dagger} = -\gamma^k$. Così:

$$= \psi_l^\dagger (1 - \gamma_5) \gamma^0 \gamma^\mu \psi_\nu = \bar{\psi}_l \gamma^\mu (1 - \gamma_5) \psi_\nu \quad (2.3.12)$$

in cui abbiamo utilizzato $\gamma^{0\dagger} \gamma^0 = \gamma^0 \gamma^0$ e $\gamma^{k\dagger} \gamma^0 = -\gamma^k \gamma^0 = \gamma^0 \gamma^k$, che in generale sono $\gamma^{\mu\dagger} \gamma^0 = \gamma^0 \gamma^\mu$.

Le matrici che abbiamo inserito le possiamo riscrivere tramite i generatori della rappresentazione fondamentale di $SU(2)$. Infatti, se indichiamo:

$$U(\theta) = e^{i\theta_i t_i} \quad , \quad t_i = \frac{\tau_i}{2} \quad , \quad i = 1, 2, 3 \quad (2.3.20)$$

con τ_i le matrici di Pauli:

$$\tau_1 = \begin{pmatrix} 0 & 1 \\ 1 & 0 \end{pmatrix} \quad , \quad \tau_2 = \begin{pmatrix} 0 & -i \\ i & 0 \end{pmatrix} \quad , \quad \tau_3 = \begin{pmatrix} 1 & 0 \\ 0 & -1 \end{pmatrix} \quad (2.3.21)$$

per cui vale l'algebra:

$$[t_i, t_j] = i\epsilon_{ijk} t_k. \quad (2.3.22)$$

Se definiamo i soliti operatori di scala:

$$\tau^\pm = \frac{\tau_1 \pm i\tau_2}{2} = t_1 \pm it_2 \quad (2.3.23)$$

allora abbiamo:

$$\tau^+ = \begin{pmatrix} 0 & 1 \\ 0 & 0 \end{pmatrix} \quad , \quad \tau^- = \begin{pmatrix} 0 & 0 \\ 1 & 0 \end{pmatrix} \quad (2.3.24)$$

perciò:

$$L^\mu = \bar{L} \gamma^\mu (1 - \gamma_5) \tau^+ L \quad (2.3.25)$$

$$= 2\bar{L} \gamma^\mu \tau^+ L_L \quad (2.3.26)$$

$$L_\mu^\dagger = \bar{L} \gamma_\mu (1 - \gamma_5) \tau^- L \quad (2.3.27)$$

$$= 2\bar{L} \gamma_\mu \tau^- L_L \quad (2.3.28)$$

dove abbiamo messo i doppietti left:

$$L_L = P_L L = \frac{1 - \gamma_5}{2} L = \frac{1 - \gamma_5}{2} \begin{pmatrix} \psi_\nu \\ \psi_l \end{pmatrix} = \begin{pmatrix} \psi_\nu \\ \psi_l \end{pmatrix}_L. \quad (2.3.29)$$

Se introducessimo le correnti:

$$J_i^\mu = \bar{L} \gamma^\mu t_i L_L \quad i = 1, 2, 3 \quad (2.3.30)$$

avremmo:

$$L^\mu = 2\bar{L} \gamma^\mu (t_1 + it_2) L_L \quad (2.3.31)$$

$$= 2(J_1^\mu + iJ_2^\mu) \quad (2.3.32)$$

$$L_\mu^\dagger = 2\bar{L} \gamma_\mu (t_1 - it_2) L_L \quad (2.3.33)$$

$$= 2(J_{1\mu} - iJ_{2\mu}) \quad (2.3.34)$$

da cui segue:

$$\mathcal{L} = \frac{G_F}{\sqrt{2}} 4 [J_{1\mu} J_1^\mu + J_{2\mu} J_2^\mu]. \quad (2.3.35)$$

Se aggiungessimo il termine $J_{3\mu} J_3^\mu$ otterremmo:

$$\mathcal{L}'_F = 4 \frac{G_F}{\sqrt{2}} \sum_{i=1}^3 J_{i\mu} J_i^\mu \quad (2.3.36)$$

che sarebbe invariante per $SU(2)$. Infatti, se facciamo una trasformazione infinitesima:

$$L_L \xrightarrow{SU(2)} e^{i\theta_j t_j} L_L \xrightarrow{\text{infinitesima}} (1 - i\theta_j t_j) L_L \quad (2.3.37)$$

e analogamente:

$$\bar{L}_L \longrightarrow (1 + i\theta_j t_j) \bar{L}_L \quad (2.3.38)$$

vediamo che le correnti diventano:

$$J_i^\mu = \bar{L}_L \gamma^\mu t_i L_L \xrightarrow{SU(2)} \bar{L}_L \gamma^\mu (1 - i\theta_j t_j) t_i (1 + i\theta_j t_j) L_L \quad (2.3.39)$$

$$= J_i^\mu + i\theta_j \bar{L}_L \gamma^\mu [t_i, t_j] L_L \quad (2.3.40)$$

$$= J_i^\mu + i\theta_j (i\epsilon_{ijk}) J_k^\mu \quad (2.3.41)$$

$$= [\delta_{ik} + i\theta_k (T_j)_{ik}] J_k^\mu \quad (2.3.42)$$

in cui:

$$(T_a)_{ik} = i\epsilon_{iak} = i\epsilon_{aki} \quad , \quad a = 1, 2, 3 \quad (2.3.43)$$

sono i generatori della rappresentazione aggiunta⁷ di $SU(2)$. Trovato ciò abbiamo per trasformazioni infinitesime (di $SU(2)$ nella rappresentazione aggiunta):

$$\begin{pmatrix} J_1^\mu \\ J_2^\mu \\ J_3^\mu \end{pmatrix} \longrightarrow (\mathbb{1} + i\theta_a T_a) \begin{pmatrix} J_1^\mu \\ J_2^\mu \\ J_3^\mu \end{pmatrix} \quad (2.3.44)$$

mentre per trasformazioni finite:

$$\begin{pmatrix} J_1^\mu \\ J_2^\mu \\ J_3^\mu \end{pmatrix} \xrightarrow{SU(2)} e^{i\theta_a T_a} \begin{pmatrix} J_1^\mu \\ J_2^\mu \\ J_3^\mu \end{pmatrix} = U \begin{pmatrix} J_1^\mu \\ J_2^\mu \\ J_3^\mu \end{pmatrix} \quad (2.3.45)$$

dove abbiamo messo U matrice unitaria 3x3, che ci dà la trasformazione generica di $SU(2)$ nella aggiunta. Conseguentemente abbiamo:

$$(J_1^\mu \quad J_2^\mu \quad J_3^\mu) = \begin{pmatrix} J_1^\mu \\ J_2^\mu \\ J_3^\mu \end{pmatrix}^\dagger \xrightarrow{SU(2)} \left[U \begin{pmatrix} J_1^\mu \\ J_2^\mu \\ J_3^\mu \end{pmatrix} \right]^\dagger = (J_1^\mu \quad J_2^\mu \quad J_3^\mu) U^\dagger. \quad (2.3.46)$$

⁷Ricorda che è la rappresentazione che ha dimensione pari a quella del gruppo, e i cui generatori sono legati alle costanti di struttura.

Perciò:

$$\mathcal{L}'_F = 4 \frac{G_F}{\sqrt{2}} \sum_{i=1}^3 J_{i\mu} J_i^\mu \quad (2.3.47)$$

$$= 4 \frac{G_F}{\sqrt{2}} \begin{pmatrix} J_1^\mu & J_2^\mu & J_3^\mu \end{pmatrix} \begin{pmatrix} J_1^\mu \\ J_2^\mu \\ J_3^\mu \end{pmatrix} \quad (2.3.48)$$

$$\xrightarrow{SU(2)} 4 \frac{G_F}{\sqrt{2}} \begin{pmatrix} J_1^\mu & J_2^\mu & J_3^\mu \end{pmatrix} U^\dagger U \begin{pmatrix} J_1^\mu \\ J_2^\mu \\ J_3^\mu \end{pmatrix} \quad (2.3.49)$$

$$= \mathcal{L}'_F. \quad (2.3.50)$$

Dunque, se aggiungessimo J_3^μ avremmo una lagrangiana invariante per $SU(2)$. Il pezzo che vorremmo aggiungere è:

$$J_3^\mu = \bar{L}_L \gamma^\mu t_3 L_L \quad (2.3.51)$$

$$= \frac{1}{2} \bar{\psi}_{\nu,L} \gamma^\mu \psi_{\nu,L} - \frac{1}{2} \bar{\psi}_{l,L} \gamma^\mu \psi_{l,L} \quad (2.3.52)$$

che è una corrente con un pezzo tra spinori *left* carichi (simile a quella che si ha in QED) e tra spinori *left* neutri, dunque non può essere la corrente elettromagnetica. Infatti non possiamo semplicemente aggiungere J_3^μ alla lagrangiana e dire che essa è la corrente elettromagnetica, sia perché ci sarebbero i neutrini tra i piedi, ma anche perché J_3^μ parla solamente con le componenti di chiralità left. Diagrammaticamente abbiamo:

$$\quad (2.3.53)$$

Però potremmo aggiungere J_3^μ in $\mathcal{L}$ dicendo che i processi che predice non sono osservati sperimentalmente perché trascurabili rispetto gli stessi processi elettromagnetici (vedi ad esempio il secondo diagramma che è analogo a Bhabha scattering).

2.4 Corrente debole adronica

Per la corrente adronica, Fermi scrisse solamente quella con protone e neutroni come campi fondamentali. Non potendo prevedere la violazione di parità, ipotizzò una corrente *solo vettoriale* (simile a quella elettromagnetica):

$$H_\mu^{(\text{Fermi})} = \bar{\psi}_p \gamma_\mu \psi_n. \quad (2.4.1)$$

Un modo più corretto di scrivere la corrente adronica neutrone-protone è:

$$H_\mu = \bar{\psi}_p \gamma_\mu (G_V - G_A \gamma_5) \psi_n \quad (2.4.2)$$

con G_V e G_A determinati sperimentalmente:

$$G_V \simeq 1 \quad , \quad G_A \simeq 1.24. \quad (2.4.3)$$

Però ci sono altri decadimenti deboli, adronici, ma simili ai β , che coinvolgono particelle strane. Alcuni esempi sono:

$$K^+ \longrightarrow \mu^+ + \nu_\mu \quad (2.4.4)$$

$$\Lambda^0 \longrightarrow p + e^- + \bar{\nu}_e \quad (2.4.5)$$

$$\Sigma^- \longrightarrow n + e^- + \bar{\nu}_e \quad (2.4.6)$$

Per cui si ha $\Delta s = 1$ nel vertice di adronico. Sperimentalmente però tali processi sono molto più rari rispetto quelli con $\Delta s = 0$.

Anche per questi decadimenti strani gli elementi di matrice della corrente debole devono essere non nulli. Per esempio:

$$H_\mu^{\Delta s=1} = \bar{\psi}_p \gamma_\mu (G_V^{\Delta s=1} - \gamma_5 G_A^{\Delta s=1}) \psi_{\Lambda^0}. \quad (2.4.7)$$

Empiricamente si è trovato:

$$(G_V)^2 + (G_V^{\Delta s=1})^2 \simeq 1 \quad (2.4.8)$$

e fu' proprio a seguito di questo risultato che Cabibbo nel 1963 suggerì che la corrente debole vettoriale adronica si dovesse scrivere:

$$H_\mu^V = \cos \theta_c H_\mu^{V(\Delta s=0)} + \sin \theta_c H_\mu^{V(\Delta s=1)} \quad (2.4.9)$$

in cui identifichiamo:

$$G_V = \cos \theta_c \quad ; \quad G_V^{\Delta s=1} = \sin \theta_c \quad (2.4.10)$$

così che possa valere (2.4.8). Sperimentalmente valgono:

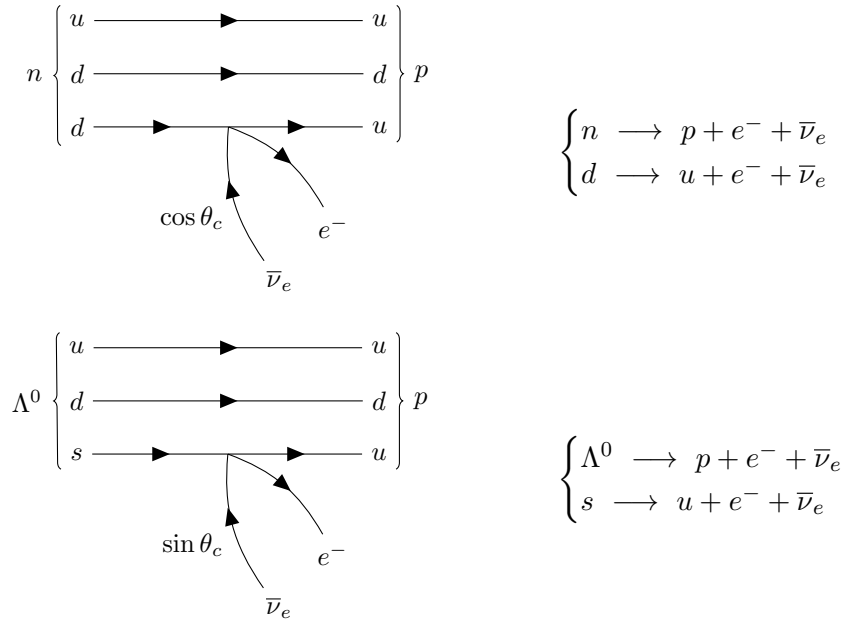
$$\cos \theta_c \simeq 0.9738 \quad ; \quad \sin \theta_c = 0.2273. \quad (2.4.11)$$

Cabibbo in realtà generalizzò ampliando ulteriormente l'algebra delle correnti deboli adroniche.

Il punto di partenza è stato il progressivo rendersi conto che il gran numero di particelle adroniche poteva essere organizzato ipotizzando una struttura composta per gli adroni. Infatti, se non pensassimo gli adroni come elementari, allora dovremmo aggiungere pezzi del tipo (2.4.2) (e poi agire come fece Cabibbo) alla lagrangiana per ogni nuova corrente adronica che descrive il decadimento di ogni particella. Ipotizzando l'esistenza dei

quark (chiaramente poi confermati) si riesce a scrivere una teoria adronica più compatta. Dunque si è spostato il problema dello studio delle interazioni elettrodeboli degli adroni sui quark.

Il principio organizzativo dello spettro adronico fu scoperto da Gell-Mann nei primi anni 60'. Gli adroni (conosciuti allora) sono composti da 3 tipi di particelle più fondamentali: i quark u , d ed s . Successivamente ci si è resi conto che il numero di quark in natura è 6 (organizzati in doppietti). Questi quark sono uniti tra loro dall'interazione nucleare forte. Per esempio si hanno le strutture del protone $p = uud$, del neutrone $n = udd$ e della $\Lambda^0 = uds$. I decadimenti di n e Λ^0 si possono vedere come interazioni tra una corrente di quark ed una leptonica, descritti dai diagrammi:



per cui scriviamo la corrente adronica è:

$$H^\mu = \cos \theta_c \bar{\psi}_u \gamma^\mu (1 - \gamma_5) \psi_d + \sin \theta_c \bar{\psi}_u \gamma^\mu (1 - \gamma_5) \psi_s \quad (2.4.12)$$

la quale descrive tutti i processi che coinvolgono gli adroni che contengono solo u, d, s .

Si ipotizzò un accoppiamento tra correnti left simile anche per i quark.

Si può notare la struttura vettoriale assiale, che non è più $G_V - G_A \gamma_5$, ma è diventata $(1 - \gamma_5)$, che è la stessa valida per i leptoni; questo può funzionare perché, nonostante $G_V \sim 1$ e $G_A \sim 1.24$ ovvero molto diversi, i decadimenti di protone e neutrone devono tenere conto della struttura interna di essi, ossia delle interazioni che i quark hanno tra loro (scambio di gluoni). Ognuna delle interazioni coinvolge una matrice γ^μ e si può vedere che la struttura di Dirac diventa più complessa nel momento in cui accendiamo le interazioni interne, dunque per ricavare l'interazione G_V e G_A per l'adrone non basta

considerare solo l'interazione debole, del decadimento, ma bisogna tenere conto anche dell'interazione forte. Fare tutto ciò non è banale, poiché siamo a livello non perturbativo⁸ della QCD (vedi la sezione di running coupling e libertà asintotica §5.3), la quantità di interazioni è elevatissima, e per studiare la struttura interna degli adroni bisogna fare quella che va sotto il nome di *lattice QCD*. Si riescono a predire in termini teorici (con qualche difficoltà) i valori di G_V e G_A per gli adroni.

Però, essendo solo l'interazione interna a complicare le cose, si può ipotizzare una corrente adronica semplice con una violazione totale di parità, come per i leptoni.

2.4.1 Corrente con i quark e simmetria $SU(3)$

Gell-Mann, nel 1962, ipotizzò che per i quark u, d, s si organizzano per formare adroni secondo le rappresentazioni multidimensionali ($d = 8, 10, \dots$) di $SU(3)$ (di sapore).

Cabibbo, in analogia a quanto succede per i leptoni ed $SU(2)$, e per via dell'organizzazione che si adotta per i quark (ovvero scrivere i quark in un tripletto di $SU(3)$ (2.4.13)), utilizza la simmetria $SU(3)$ per le correnti deboli adroniche. Infatti, utilizzando $SU(3)$ e le matrici di Gell-Mann λ^a ($a = 1, 2, \dots, 8$) ipotizzò un ottetto di correnti conservate. Le matrici di Gell-Mann sono:

$$\begin{aligned} \lambda_1 &= \begin{pmatrix} 0 & 1 & 0 \\ 1 & 0 & 0 \\ 0 & 0 & 0 \end{pmatrix}, & \lambda_2 &= \begin{pmatrix} 0 & -i & 0 \\ i & 0 & 0 \\ 0 & 0 & 0 \end{pmatrix}, & \lambda_3 &= \begin{pmatrix} 1 & 0 & 0 \\ 0 & -1 & 0 \\ 0 & 0 & 0 \end{pmatrix} \\ \lambda_4 &= \begin{pmatrix} 0 & 0 & 1 \\ 0 & 0 & 0 \\ 1 & 0 & 0 \end{pmatrix}, & \lambda_5 &= \begin{pmatrix} 0 & 0 & -i \\ 0 & 0 & 0 \\ i & 0 & 0 \end{pmatrix} \\ \lambda_6 &= \begin{pmatrix} 0 & 0 & 0 \\ 0 & 0 & 1 \\ 0 & 1 & 0 \end{pmatrix}, & \lambda_7 &= \begin{pmatrix} 0 & 0 & 0 \\ 0 & 0 & -i \\ 0 & i & 0 \end{pmatrix}, & \lambda_8 &= \begin{pmatrix} 1 & 0 & 0 \\ 0 & 1 & 0 \\ 0 & 0 & -2 \end{pmatrix}. \end{aligned}$$

In analogia con i doppietti di $SU(2)$ delle correnti leptoniche, Cabibbo raggruppa ψ_u, ψ_d e ψ_s in un tripletto di $SU(3)$:

$$Q = \begin{pmatrix} \psi_u \\ \psi_d \\ \psi_s \end{pmatrix} \quad (2.4.13)$$

⁸In regime perturbativo, cioè ad alte energie, possiamo trascurare l'interazione interna degli adroni e quindi studiare solamente la lagrangiana semplice per l'interazione debole.

da cui segue:

$$\bar{\psi}_u \gamma^\mu (1 - \gamma_5) \psi_d = \bar{Q} \gamma^\mu (1 - \gamma_5) \begin{pmatrix} 0 & 1 & 0 \\ 0 & 0 & 0 \\ 0 & 0 & 0 \end{pmatrix} Q \quad (2.4.14)$$

$$= \bar{Q} \gamma^\mu (1 - \gamma_5) (\lambda_1 + i\lambda_2) Q \quad (2.4.15)$$

$$\bar{\psi}_u \gamma^\mu (1 - \gamma_5) \psi_s = \bar{Q} \gamma^\mu (1 - \gamma_5) \begin{pmatrix} 0 & 0 & 1 \\ 0 & 0 & 0 \\ 0 & 0 & 0 \end{pmatrix} Q \quad (2.4.16)$$

$$= \bar{Q} \gamma^\mu (1 - \gamma_5) (\lambda_4 + i\lambda_5) Q. \quad (2.4.17)$$

Dunque, la corrente adronica è:

$$H^\mu = \cos \theta_c \bar{Q} \gamma^\mu (1 - \gamma_5) (\lambda_1 + i\lambda_2) Q + \sin \theta_c \bar{Q} \gamma^\mu (1 - \gamma_5) (\lambda_4 + i\lambda_5) Q \quad (2.4.18)$$

e anche in questo caso H^μ_μ sarebbe incompleta per generare una lagrangiana invariante per $SU(3)$, poiché mancano $\lambda_3, \lambda_6, \lambda_7$ e λ_8 . In particolare manca una corrente con $\lambda_6 + i\lambda_7$ che è una corrente tra ψ_d e ψ_s , ma che non si osserva sperimentalmente. Questo è il primo esempio di flavour changing.

Oltre a questo si evidenzia una asimmetria tra corrente leptonica e corrente adronica. Prima del '70 si conoscevano solamente due doppietti leptonici, quindi le correnti erano solamente:

$$L^\mu = \bar{\psi}_{\nu_e} \gamma^\mu (1 - \gamma_5) \psi_e + \bar{\psi}_{\nu_\mu} \gamma^\mu (1 - \gamma_5) \psi_\mu \quad (2.4.19)$$

$$H^\mu = \bar{\psi}_u \gamma^\mu (1 - \gamma_5) (\cos \theta_c \psi_d + \sin \theta_c \psi_s). \quad (2.4.20)$$

I leptoni sono raggruppati in doppietti di $SU(2)$, che sono 2 essendoci 2 famiglie, ed essi sono singoletti di $SU(3)$. Invece i quark (3 quark u, d, s) sono organizzati in un tripletto di $SU(3)$, un doppietto di $SU(2)$ (cioè $\begin{pmatrix} u \\ d \end{pmatrix}$) ed un singoletto di $SU(2)$ (s). La combinazione di quark $-\sin \theta_c \psi_d + \cos \theta_c \psi_s$ ortogonale a quella che compare in $H_\mu(x)$, non appare in $\mathcal{L}$ e sembra disaccoppiata dalle interazioni deboli.

2.4.2 Problema delle FCNC

Analizziamo il problema delle *flavour changing neutral currents* (correnti deboli neutre con cambiamento di sapore).

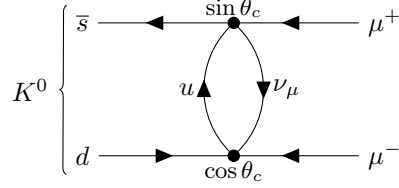
La corrente $H^\mu(x)$ non ammette interazioni al tree-level con la corrente neutra e cambio di sapore:

$$d \leftrightarrow s.$$

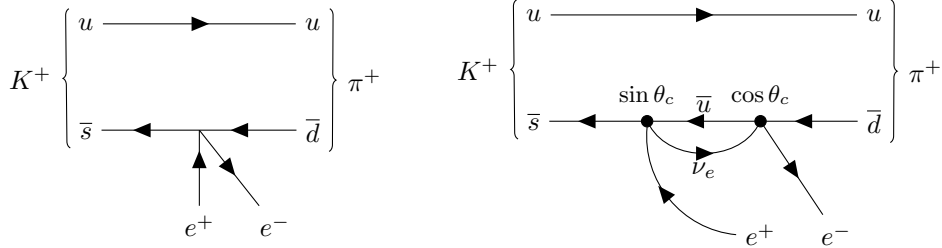
Ad esempio il decadimento:

$$K^0 \longrightarrow \mu^+ + \mu^- \quad (2.4.21)$$

raccontato dal diagramma tree-level:



però, ammette flavour changing ad ordini più elevati (ossia quando compaiono loop), ad esempio i decadimenti:



questi due processi contribuiscono ai propri decadimenti (in particolare il secondo a (2.4.21)) con un coefficiente proporzionale a:

$$\sin \theta_c \cos \theta_c. \quad (2.4.22)$$

Nonostante potremmo dire che i diagrammi visti, contenendo un loop ed essendo proporzionali a $\sin \theta_c \ll 1$, ci aspettiamo essere molto soppressi, rispetto la misura sperimentale il decadimento (2.4.21) è molto sovrastimato:

$$\text{BR}(K^0 \longrightarrow \mu^+ + \mu^-) = (6.84 \pm 0.11) \cdot 10^{-9}$$

il che vuol dire che tra tutti i modi che K^0 ha di decadere, sceglie il modo che stiamo analizzando solamente una volta su 10^9 , il che non è spiegabile solamente dalla presenza di un loop e di $\sin \theta_c$. Dunque, possiamo pensare che esista un ulteriore meccanismo che sopprima le FCNC. Analizziamo quindi il *meccanismo GIM*. Infatti, la corrente del processo che sbriveremmo sarebbe:

$$H^\mu = \bar{\psi}_u(1 - \gamma_5)(\cos \theta_c \psi_d + \sin \theta_c \psi_s) \quad (2.4.23)$$

espressa attraverso un angolo. Quello che ci aspetteremmo è che in questa corrente manchi il contributo ortogonale, cioè:

$$(-\sin \theta_c \psi_d + \cos \theta_c \psi_s) \quad (2.4.24)$$

ciò è quello che intuirono Glasshow, Iliopoulos e Maiani nel 1970.

2.4.3 Meccanismo GIM

Analizziamo meglio il meccanismo di Glassshow-Iliopoulos-Maiani (1970).

La soluzione alla soppressione delle FCNC è stata data da G.I.M., che risolve anche il problema della asimmetria nei numeri quantici che abbiamo visto nella sezione §2.4.1.

L'idea è: introdurre nella corrente la combinazione ortogonale:

$$-\sin \theta_c \psi_d + \cos \theta_c \psi_s \quad (2.4.25)$$

(in cui sta volta l'accoppiamento con ψ_s è dominante) e accoppiarla con un *nuovo sapore di quark*, il quark *charm*, ipotizzato essere più pesante di u . Così la corrente adronica diventa:

$$H^\mu = \bar{\psi}_u \gamma^\mu (1 - \gamma_5) (\cos \theta_c \psi_d + \sin \theta_c \psi_s) + \bar{\psi}_c \gamma^\mu (1 - \gamma_5) (-\sin \theta_c \psi_d + \cos \theta_c \psi_s) \quad (2.4.26)$$

che possiamo scrivere in forma compatta:

$$H^\mu = \sum_{i,j=1}^2 \bar{\psi}_{u_i} \gamma^\mu (1 - \gamma_5) O_{ij} \psi_{d_j} \quad (2.4.27)$$

$$= \sum_{i,j=1}^2 \bar{Q}_i \gamma^\mu (1 - \gamma_5) O_{ij} \begin{pmatrix} 0 & 1 \\ 0 & 0 \end{pmatrix} Q_j \quad (2.4.28)$$

avendo definito i doppietti di quark:

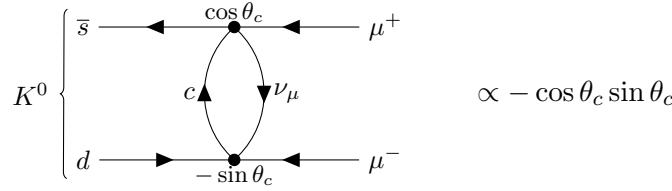
$$Q_1 = \begin{pmatrix} \psi_u \\ \psi_d \end{pmatrix} \quad ; \quad Q_2 = \begin{pmatrix} \psi_c \\ \psi_s \end{pmatrix} \quad (2.4.29)$$

ovvero le componenti $u_1 = u$, $u_2 = c$, $d_1 = d$, $d_2 = s$ e la matrice ortogonale ($OO^T = O^T O = \mathbb{1}$) di Cabibbo:

$$O = \begin{pmatrix} \cos \theta_c & \sin \theta_c \\ -\sin \theta_c & \cos \theta_c \end{pmatrix}. \quad (2.4.30)$$

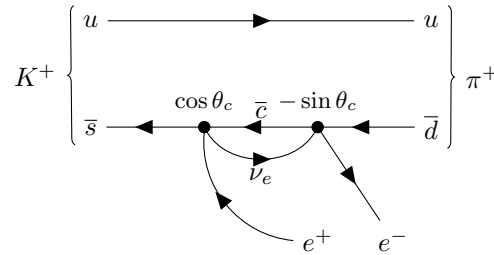
Analizziamo le conseguenze del meccanismo GIM. Per prima cosa esso ci porta ad una soluzione per il problema riguardante la soppressione delle FCNC. Infatti, nell'esempio (2.4.21) si ha:

$$K^0 \left\{ \begin{array}{l} \bar{s} \xleftarrow{\sin \theta_c} \mu^+ \\ u \xrightarrow{\quad} \nu_\mu \\ d \xrightarrow{\cos \theta_c} \mu^- \end{array} \right. \propto \cos \theta_c \sin \theta_c$$



e la cancellazione dei due diagrammi è perfetta se $m_u = m_c$. A noi serve che le due masse siano uguali poiché, comparando tali quark in linee interne, contribuiscono all'ampiezza con dei propagatori che sono uguali per i due quark solo se le masse sono uguali. Nel caso in cui $m_u \neq m_c$ sarebbe proprio tale differenza di massa che ci spiega la forte soppressione (non totale) di questo processo, infatti i contributi dei due diagrammi sarebbero identici ad eccezione della differenza tra le due masse. Di fatto, nonostante la differenza di masse non è così piccola, il grado di cancellazione è molto elevato poiché m_u ed m_c sono paragonabili alla scala di energia $1/\sqrt{G_F}$.

Nota che il discorso si adatta in modo identico anche al processo $K^+ \rightarrow \pi^+ + e^+ + e^-$, che con GIM diventa:



GIM porta anche ad una soluzione dell'asimmetria nella struttura adronica e leptonica della corrente carica debole. Infatti, sia per i leptoni che per gli adroni ora ci sono due (che diventeranno tre) doppietti di $SU(2)$ sinistrorsi:

$$\begin{pmatrix} \psi_{\nu_e} \\ \psi_e \end{pmatrix} \quad \begin{pmatrix} \psi_{\nu_\mu} \\ \psi_\mu \end{pmatrix} \quad ; \quad \begin{pmatrix} \psi_u \\ \psi_d \end{pmatrix} \quad \begin{pmatrix} \psi_c \\ \psi_s \end{pmatrix}. \quad (2.4.31)$$

L'unica differenza è il mixing nel settore dei quark, che è codificato dalla presenza della matrice O di Cabibbo. Vedremo anche il mixing nel settore leptonico.

Anche in questo caso, come nel settore leptonico, manca un pezzo di corrente relativo al terzo generatore di $SU(2)$, che ancora una volta potrebbe esserci, ma senza essere vista per via della soppressione rispetto la corrente elettromagnetica.

Impareremo anche che la struttura in famiglie permette anche la cancellazione delle anomalie (vedi l'ultima parte del corso).

Il quark c è stato scoperto nella sua risonanza $J/\psi = c\bar{c}$ (di massa 3.076 GeV) nel 1974 da due gruppi:

- A SLAC, Richter produce una nuova particella intorno ai 3 GeV (che chiama ψ) nel processo:

$$e^+ + e^- \longrightarrow \psi. \quad (2.4.32)$$

- A Brookhaven, Ting produce una particella con massa simile (che chiama J) nel processo:

$$p + Be \longrightarrow J + X \longrightarrow e^+ + e^- + X \quad (2.4.33)$$

in cui X sono altre particelle non identificate.

Entrambi capiscono di aver scoperto la stessa particella $J/\psi = c\bar{c}$ e ricevono il Nobel nel 1976.

Immaginando una debole energia di legame dei quark si può scrivere:

$$m_c \sim 1.5 \text{ GeV} \quad (2.4.34)$$

che ci spiega la forte soppressione delle FCNC.

2.5 Larghezze di decadimento e sezioni d'urto

Vediamo in questa sezione la larghezza di decadimento e come calcoliamo le sezioni d'urto. Noi chiamiamo $\mathcal{M}_{fi}$ l'ampiezza di probabilità di transizione da uno stato iniziale i ad uno finale f , che si ottiene dai diagrammi di Feynman. Inoltre, n_i conta il numero di particelle dello stato iniziale ed n_f dello stato finale.

2.5.1 Larghezza di decadimento

Nel caso in cui $n_i = 1$ si hanno processi di decadimento del tipo:

$$a \longrightarrow 1 + 2 + \dots + n_f$$

indichiamo la larghezza parziale di decadimento come:

$$\Gamma_{a \rightarrow f} = \frac{1}{2E_a(2\pi)^{3n_f-4}} \int \frac{d^3p_1 \dots d^3p_{n_f}}{2E_1 \dots 2E_{n_f}} \delta^4(p_f - p_a) |\mathcal{M}_{fi}|^2 \quad (2.5.1)$$

e quindi la **larghezza totale di decadimento** come:

$$\Gamma_a = \sum_f \Gamma_{a \rightarrow f}. \quad (2.5.2)$$

Definiamo la **vita media della particella** come:

$$\tau_a = \frac{1}{\Gamma_a} \quad (2.5.3)$$

la quale è sempre misurata nel sistema di riferimento della particella ed è sempre la più piccola possibile.

Ricordiamoci sempre che la larghezza di decadimento non è invariante di Lorentz, ma dipende dall'impulso.

2.5.2 Sezioni d'urto di processi di scattering

Nei casi in cui $n_i = 2$ abbiamo processi di scattering del tipo:

$$a + b \longrightarrow 1 + \cdots + n_f$$

e si può valutare la sezione d'urto:

$$\sigma = \frac{\# \text{ di transizioni per unità di tempo}}{\# \text{ di particelle incidenti per unità di superficie e di tempo}}$$

che dal punto di vista teorico:

$$\sigma_{i \rightarrow f} = \frac{1}{4\phi(2\pi)^{3n_f-4}} \int \prod_{j=1}^{n_f} \frac{d^3 p_j}{2E_j} \delta^4 \left(\sum_{j=1}^{n_f} p_j - p_a - p_b \right) |\mathcal{M}_{fi}|^2 \quad (2.5.4)$$

in cui abbiamo un invariante di Lorentz, che chiamiamo *fattore di flusso di Moller* ϕ :

$$\phi = E_a E_b v_{ab} \quad (2.5.5)$$

$$= \sqrt{(p_a \cdot p_b)^2 - m_a^2 m_b^2} \quad (2.5.6)$$

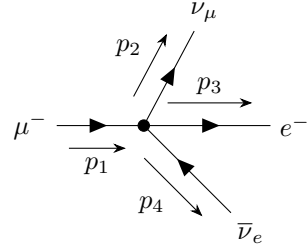
$$= \frac{[s - (m_a + m_b)^2][s - (m_a - m_b)^2]}{2} \quad (2.5.7)$$

$$\xrightarrow{s \gg m_a^2, m_b^2} \frac{s}{2} \quad (2.5.8)$$

da cui $\sigma_{i \rightarrow f}$ è manifestamente Lorentz invariante in quanto dipende solo da invarianti di Lorentz. Infatti, ora (2.5.4) è manifestamente invariante di Lorentz.

2.5.3 Decadimento del muone

Come esempio notevole si è visto il decadimento del muone:



$$\mu^- \longrightarrow \nu_\mu + e^- + \bar{\nu}_e \quad (2.5.9)$$

e ricordando:

$$\mathcal{L}_{int} = \frac{G_F}{\sqrt{2}} \bar{\psi}_{\nu_\mu} \gamma^\mu (1 - \gamma_5) \psi_\mu \bar{\psi}_{\nu_e} \gamma_\mu (1 - \gamma_5) \psi_e \quad (2.5.10)$$

ci siamo calcolati il modulo quadro dell'elemento di matrice:

$$i\mathcal{M} = -i \frac{G_F}{\sqrt{2}} \bar{u}(p_2) \gamma^\mu (1 - \gamma_5) u(p_1) \bar{u}(p_3) \gamma_\mu (1 - \gamma_5) v(p_4) \quad (2.5.11)$$

per un processo non polarizzato. Per poi determinare la larghezza totale di decadimento Γ .

Il risultato è:

$$|\mathcal{M}|^2 = 128 G_F^2 (p_1 p_4) (p_2 p_3) \quad (2.5.12)$$

$$\Gamma = \frac{m_\mu^2 G_F^2}{192 \pi^3}. \quad (2.5.13)$$

I conti, essendo molto simili ad un tipico esercizio di *Fondamenti di Teoria di Campi*, non li riporto in queste note, ma si possono trovare linkate nella mia pagina.

2.6 Unitarietà della matrice S

Come vedremo nelle sezioni che seguiranno riguardo la rinormalizzazione, la teoria di Fermi non è la teoria definitiva per l'interazione elettrodebole e ciò è strettamente legato alla costante di accoppiamento dimensionale. In particolare vedremo che la teoria non è unitaria e non è rinormalizzabile.

Per studiare l'unitarietà dobbiamo analizzare la matrice S . L'unitarietà della teoria è legata alla conservazione delle probabilità. Infatti, la matrice S emerge nel momento in cui studiamo la probabilità di transizione da uno stato iniziale ad uno finale:

$$|i\rangle \longrightarrow |f\rangle \quad (2.6.1)$$

e per farlo dobbiamo calcolare la quantità:

$$|S_{if}|^2 = |\langle f | S | i \rangle|^2 \quad (2.6.2)$$

e sommando su tutti i possibili stati finali⁹ abbiamo:

$$\sum_f |S_{fi}|^2 = 1 \quad (2.6.3)$$

e inserendoci dentro una completezza per gli stati finali:

$$\mathbb{1} = \sum_f |f\rangle \langle f| \quad (2.6.4)$$

$$= \sum_f \prod_{j=1}^{n_f} \int \frac{d^3 p_j^f}{(2\pi)^3 2E_j^f} |f\rangle \langle f| \quad (2.6.5)$$

otteniamo:

$$1 = \sum_f \langle i | S | f \rangle \langle f | S | i \rangle \quad (2.6.6)$$

$$= \langle i | S^\dagger S | i \rangle \quad (2.6.7)$$

da cui vediamo:

$$S^\dagger S = \mathbb{1} \quad (2.6.8)$$

cioè S è unitaria. Questo ha come conseguenza una relazione per le ampiezze di probabilità $\mathcal{M}$ dei processi. Per trovare la relazione tra S ed $\mathcal{M}$ scriviamo:

$$S = \mathbb{1} + iT \quad (2.6.9)$$

in cui $\mathbb{1}$ rappresenta il contributo ad S dato dall'assenza di interazioni, ovvero il contributo $\langle i | i \rangle$, mentre T rappresenta le interazioni tra le particelle, cioè il risultato del calcolo diagrammatico. Da (2.6.9) discende:

$$(\mathbb{1} + iT)^\dagger (\mathbb{1} + iT) = \mathbb{1} \quad (2.6.10)$$

$$i(T - T^\dagger) + T^\dagger T = 0 \quad (2.6.11)$$

$$\implies T^\dagger T = -i(T - T^\dagger) \quad (2.6.12)$$

⁹Sommiamo su tutte le particelle prodotte e tutti i possibili stati di esse, mentre integriamo su tutte le variabili cinematiche.

che è una relazione che connette ordini perturbativi diversi. Sappiamo anche che T e $\mathcal{M}$ sono legate dalla relazione¹⁰:

$$\langle f|T|i\rangle = (2\pi)^4 \delta^4(p_f - p_i) \mathcal{M}_{fi}. \quad (2.6.15)$$

Utilizzando stati $|a\rangle$ e $|b\rangle$ generici, possiamo scrivere:

$$\langle a|T^\dagger T|b\rangle = \sum_f \prod_{j=1}^{n_f} \int \frac{d^3 p_j^f}{(2\pi)^3 2E_j^f} \langle a|T^\dagger|f\rangle \langle f|T|b\rangle \quad (2.6.16)$$

$$= \sum_f \prod_{j=1}^{n_f} \int \frac{d^3 p_j^f}{(2\pi)^3 2E_j^f} (2\pi)^4 \delta^4(p_f - p_a) (2\pi)^4 \delta^4(p_f - p_b) \mathcal{M}_{fb} \mathcal{M}_{fa}^* \quad (2.6.17)$$

in cui abbiamo utilizzato la relazione esplicita di completezza (2.6.5) e notiamo $\mathcal{M}_{fa}^* = \mathcal{M}_{af}^\dagger$. Il membro di sinistra dell'ultima relazione la scriviamo come:

$$-i \langle a|(T - T^\dagger)|b\rangle = -i(M_{ab} - M_{ba}^*) (2\pi)^4 \delta^4(p_a - p_b) \quad (2.6.18)$$

e, dopo aver utilizzato una delle due δ^4 per fissare $p_a = p_f$ e dunque $\delta^4(p_f - p_b) = \delta^4(p_a - p_b)$, possiamo mettere tutto insieme:

$$\sum_f \prod_{j=1}^{n_f} \int \frac{d^3 p_j^f}{(2\pi)^3 2E_j^f} (2\pi)^4 \delta^4(p_a - p_b) \mathcal{M}_{fb} \mathcal{M}_{fa}^* = -i(\mathcal{M}_{ab} - \mathcal{M}_{ba}^*) \quad (2.6.19)$$

che è la **relazione di unitarietà** che cercavamo, valida per qualsiasi stati a, b . La relazione di unitarietà (2.6.19) è scritta in modo più semplice se prendiamo $|a\rangle = |b\rangle$; infatti in questo caso essa diventa:

$$\sum_f \prod_{j=1}^{n_f} \int \frac{d^3 p_j^f}{(2\pi)^3 2E_j^f} (2\pi)^4 \delta^4\left(p_a - \sum_i p_i^f\right) |\mathcal{M}_{fa}|^2 = -i(\mathcal{M}_{aa} - \mathcal{M}_{aa}^*) \quad (2.6.20)$$

$$= 2 \text{Im}\{\mathcal{M}_{aa}\} \quad (2.6.21)$$

che ci dice che la parte immaginaria di un'ampiezza d'urto in avanti è uguale alla somma dei contributi di tutti i possibili stati intermedi di particelle.

¹⁰Puoi vedere il capitolo §4.5 del Peskin There is a sort of mythology that grows up about what happened, which is different from what really did happen. Peskin in cui definisce la matrice S come:

$${}_{\text{out}}\langle p_1, p_2, \dots | k_A, k_B \rangle_{\text{in}} = \langle p_1, p_2, \dots | S | k_A, k_B \rangle \quad (2.6.13)$$

e l'elemento di matrice invariante $\mathcal{M}$:

$$\langle p_1, p_2, \dots | iT | k_A, k_B \rangle = (2\pi)^4 \delta^{(4)}(k_A + k_B - \sum_f p_f) i \mathcal{M}(k_A, k_B \rightarrow p_f). \quad (2.6.14)$$

La relazione (2.6.21) è quello che va sotto il nome di **teorema ottico**, mentre $\mathcal{M}_{aa}$ si chiama **ampiezza d'urto in avanti**.

Ora, il membro a sinistra di (2.6.21) somiglia molto alla sezione d'urto (2.5.4), infatti possiamo riscriverla come:

$$4\phi\sigma_{tot}(a \rightarrow f) = 2 \operatorname{Im}\{\mathcal{M}_{aa}\}. \quad (2.6.22)$$

La relazione (2.6.21) possiamo anche scriverla in termini di diagrammi di Feynman come in figura 2.2.

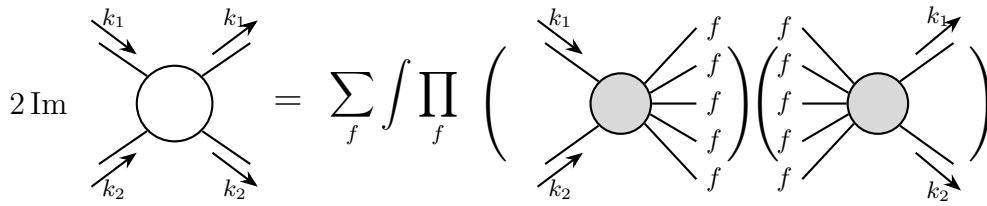


Figura 2.2

Solitamente queste relazioni sono utili poiché per determinare la sezione d'urto totale ci basta calcolare il termine legato al diagramma di sinistra della figura 2.2. Ovviamente questo va fatto per ogni ordine perturbativo, poiché le bolle rappresentano tutti i possibili diagrammi di Feynman che contribuiscono al processo considerato. Tieni a mente però che l'uguaglianza vale solo quando a sinistra e a destra ci sono termini con stessa potenza della costante di accoppiamento.

Di fatto, la relazione (2.6.21) lega relazioni con diverso numero di loop.

2.6.1 Violazione dell'unitarietà della teoria di Fermi

Cerchiamo uno stato a per poter testare il teorema ottico (2.6.21) nella teoria di Fermi per le interazioni deboli.

In particolare, dobbiamo scegliere le particelle dello stato a in modo che possano interagire debolmente in maniera elastica, cioè che esista il processo debole $a \rightarrow a$, così che l'ampiezza $\mathcal{M}_{aa}$ non sia nulla.

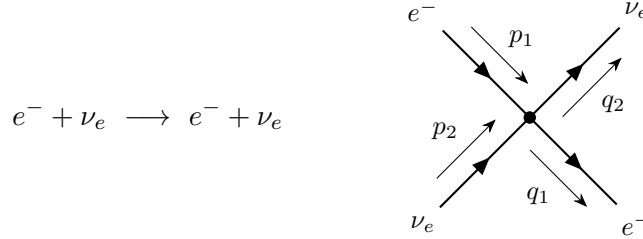
Se scegliamo lo stato:

$$a = e^- + \nu_e \quad (2.6.23)$$

abbiamo il vantaggio che esso ammette un'unico¹¹ processo debole $2 \rightarrow 2$ al

¹¹Gli altri processi con tale stato iniziale hanno più di due particelle nello stato finale, dunque saranno meno probabili e in prima approssimazione possiamo considerare solo il processo $2 \rightarrow 2$. Infatti, nella relazione (2.6.21) l'integrale sullo spazio delle fasi ci uccide i contributi $2 \rightarrow n$, con $n > 2$, ad una prima approssimazione, rispetto il contributo con 2 particelle finali.

tree-level:



Il calcolo dell'ampiezza non polarizzata si svolge in modo identico a quella calcolata per il decadimento del muone, e si ottiene (puoi seguire i conti presenti sulla mia pagina):

$$\sum_{spin} |\mathcal{M}|^2 = 128 G_F^2 (p_1 p_2) (q_1 q_2) \quad (2.6.24)$$

$$= 32 G_F^2 s^2 \quad (2.6.25)$$

ma se trascuriamo le masse, ovvero:

$$s \gg m_e^2, m_\nu^2 \quad (2.6.26)$$

allora abbiamo:

$$s = (p_1 + p_2)^2 \simeq 2 p_1 \cdot p_2 \quad (2.6.27)$$

$$s = (q_1 + q_2)^2 \simeq 2 q_1 \cdot q_2 \quad (2.6.28)$$

che ci portano a dire che nell'espressione trovata:

$$|\mathcal{M}|^2 = \sum_s |\mathcal{M}|^2 = 32 G_F^2 s^2 \quad (2.6.29)$$

abbiamo che solamente contributi left contribuiscono, quindi un solo valore di elicità possibile. Dunque, $\mathcal{M}_{aa}$, in prima approssimazione, è dato da questo processo nel caso di forward scattering (gli stati iniziali e finali sono uguali), per cui:

$$q_1 = p_1 \quad , \quad q_2 = p_2 \quad (2.6.30)$$

ma anche $\mathcal{M}_{fa}$, in prima approssimazione, è dato da tale processo, ma in cui dobbiamo considerare l'integrale su tutto lo spazio delle fasi q_1, q_2 permessi dalla conservazione dell'impulso.

La cinematica del processo, nel sistema del centro di massa, è raffigurata in figura 2.3b. Trascurando le masse avremo:

$$|\vec{p}| = |\vec{q}| = E \quad (2.6.31)$$

e quindi:

$$s = (2E)^2 \implies 2E = \sqrt{s}. \quad (2.6.32)$$

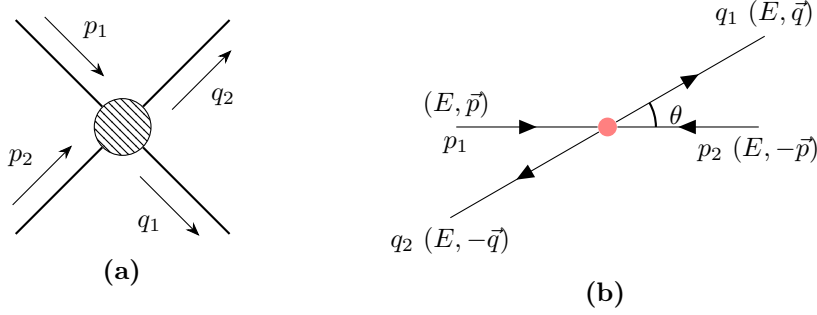


Figura 2.3

Dunque, in generale per un processo $2 \rightarrow 2$ nel sistema del riferimento del centro di massa, avremo che l'ampiezza $\mathcal{M}$ dipende solo da due parametri $\mathcal{M} = \mathcal{M}(s, \cos \theta)$.

Possiamo ricordarci le definizioni delle variabili di Mandelstam:

$$s = (p_1 + p_2)^2 = (q_1 + q_2)^2 = (2E)^2 \quad (2.6.33)$$

$$t = (p_1 - q_1)^2 = -\frac{s}{2}(1 - \cos \theta) \quad (2.6.34)$$

per cui potremmo scrivere $\mathcal{M}$ in funzione di s e t , oppure di s e $\cos \theta$. Ad ogni modo, a fissato s , a noi conviene vedere l'ampiezza $\mathcal{M}$ come una funzione di $\cos \theta$, dunque si può sviluppare in termini di polinomi di Legendre (autostati di L^2 in MQ), che costituiscono una base completa ortonormale di funzioni sull'intervallo $-1, +1$:

$$P_l(\cos \theta) = \frac{1}{2^l l!} \frac{d^l}{d \cos \theta^l} (\cos^2 \theta - 1)^l \quad (2.6.35)$$

con la normalizzazione del prodotto scalare (formula di Rodriguez):

$$(P_l, P_n) = \int_{-1}^{+1} dx P_l(x) P_n(x) \quad (2.6.36)$$

$$= \frac{2}{2l+1} \delta_{nl} \quad (2.6.37)$$

che come conseguenza ha:

$$P_l(1) = 1 \quad \forall l. \quad (2.6.38)$$

Possiamo quindi scrivere lo **sviluppo in onde parziali** di $\mathcal{M}$:

$$\mathcal{M}(s, \cos \theta) = \sum_{l=0}^{\infty} (16\pi(2l+1)a_l(s)) P_l(\cos \theta) \quad (2.6.39)$$

con i coefficienti dati da:

$$16\pi(2l+1)a_l(s) = \frac{(P_l, \mathcal{M})}{(P_l, P_l)} \quad (2.6.40)$$

$$= \frac{2l+1}{2} \int_{-1}^1 d\cos\theta P_l(\cos\theta) \mathcal{M} \quad (2.6.41)$$

da cui leggiamo i coefficienti:

$$a_l(s) = \frac{1}{32\pi} \int_{-1}^1 d\cos\theta P_l(\cos\theta) \mathcal{M}(s, \cos\theta). \quad (2.6.42)$$

In questo modo il teorema ottico (in cui mettiamo $\cos\theta = 1$ poiché guardiamo lo scattering in avanti) diventa:

$$2\operatorname{Im}\{\mathcal{M}_{aa}\} = 2\operatorname{Im}\{\mathcal{M}(s, 1)\} \quad (2.6.43)$$

$$= 2\operatorname{Im}\left\{16\pi \sum_{l=0}^{\infty} (2l+1)a_l(s)P_l(1)\right\} \quad (2.6.44)$$

e l'altro membro:

$$\int \frac{d^3q_1}{(2\pi)^3 2E_1} \int \frac{d^3q_2}{(2\pi)^3 2E_2} (2\pi)^4 \delta^4(p_1 + p_2 - q_1 - q_2) |\mathcal{M}(s, \cos\theta)|^2 \quad (2.6.45)$$

$$= \int \frac{dE_1 E_1^2}{E_1 E_2} d\cos\theta d\phi \frac{1}{16\pi^2} \delta(\sqrt{s} - E_1 - E_2) |\mathcal{M}|^2 \quad (2.6.46)$$

$$= \int dE_1 d\cos\theta d\phi \frac{1}{16\pi^2} \delta(\sqrt{s} - 2E_1) |\mathcal{M}|^2 \quad (2.6.47)$$

$$= \frac{1}{16\pi} \int_{-1}^{+1} d\cos\theta |\mathcal{M}(s, \cos\theta)|^2 \quad (2.6.48)$$

in cui per (2.6.46) abbiamo usato la delta spaziale e siamo passati alle coordinate sferiche, mentre per (2.6.47) abbiamo posto $E_1 = E_2$ poiché $|\vec{q}_1| = |\vec{q}_2|$, e per (2.6.48) il fatto che valga:

$$\delta(\sqrt{s} - 2E_1) = \frac{1}{2} \delta\left(E_1 - \frac{\sqrt{s}}{2}\right) \quad (2.6.49)$$

che si dimostra ricordando $\delta(ax) = \frac{1}{|a|} \delta(x)$ insieme alla parità della delta.

Se utilizziamo lo sviluppo in onde parziali di $\mathcal{M}$:

$$= \frac{1}{16\pi} \int_{-1}^{+1} d\cos\theta \left[16\pi \sum_{l=0}^{\infty} (2l+1) a_l(s) P_l(\cos\theta) \right] \times \\ \times \left[16\pi \sum_{n=0}^{\infty} (2n+1) a_n^*(s) P_n(\cos\theta) \right] \quad (2.6.50)$$

$$= \sum_{l,n} 16\pi (2l+1)(2n+1) a_l(s) a_n^*(s) \int_{-1}^{+1} d\cos\theta P_l(\cos\theta) P_n(\cos\theta) \quad (2.6.51)$$

$$= 32\pi \sum_{l=0}^{\infty} (2l+1) |a_l(s)|^2. \quad (2.6.52)$$

Dunque la relazione di unitarietà diventa:

$$2 \operatorname{Im} \left\{ 16\pi \sum_{l=0}^{\infty} (2l+1) a_l(s) \right\} = 32\pi \sum_{l=0}^{\infty} (2l+1) |a_l(s)|^2 \quad (2.6.53)$$

semplificando:

$$\sum_{l=0}^{\infty} (2l+1) \operatorname{Im}\{a_l(s)\} = \sum_{l=0}^{\infty} (2l+1) |a_l(s)|^2. \quad (2.6.54)$$

Abbiamo trovato la relazione (2.6.54), la quale è valida per tutti i processi in cui si trascurano i contributi di processi $2 \rightarrow n$ con $n > 2$.

Nel processo in esame abbiamo l'ampiezza (2.6.29), cioè un oggetto che non dipende da $\cos\theta$. Non dipendendo da $\cos\theta$, tra tutte le onde parziali che abbiamo, l'unica che dà contributo è quella che non dipende da $\cos\theta$, ossia $l = 0$, chiamata anche **onda s**. Abbiamo quindi:

$$a_l = 0 \quad \text{se } l \neq 0. \quad (2.6.55)$$

Dunque, in questo specifico processo la relazione di unitarietà diventa:

$$\operatorname{Im}\{a_0\} = |a_0|^2 \quad (2.6.56)$$

ma se ci scriviamo:

$$a_0(s) = |a_0(s)| e^{i\alpha} = |a_0(s)| (\cos\alpha + i \sin\alpha) \quad (2.6.57)$$

prendendo la parte immaginaria:

$$|a_0(s)| \sin\alpha = |a_0(s)|^2 \implies |a_0(s)| = \sin\alpha \quad (2.6.58)$$

da cui concludiamo:

$$|a_0(s)| \leq 1 \quad (2.6.59)$$

che è la relazione di unitarietà per il particolare processo considerato¹². Inoltre, lo sviluppo in onde parziali di $\mathcal{M}$ si semplifica molto se mettiamo $l = 0$:

$$\mathcal{M}(s, \cos \theta) = 16\pi a_0(s) \implies |\mathcal{M}|^2 = (16\pi)^2 |a_0(s)|^2 \quad (2.6.62)$$

ma conoscendo il risultato (2.6.29):

$$32G_F^2 s^2 = (16\pi)^2 |a_0(s)|^2 \implies |a_0(s)|^2 = \frac{G_F^2 s^2}{8\pi^2} \quad (2.6.63)$$

quindi l'unitarietà è soddisfatta solo se:

$$\frac{G_F s}{2\sqrt{2}\pi} \leq 1. \quad (2.6.64)$$

Dunque, in base a quanto visto, possiamo dire che l'unitarietà è sicuramente violata quando¹³:

$$\frac{G_F s}{2\sqrt{2}\pi} > 1 \implies \sqrt{s} > \sqrt{\frac{2\sqrt{2}\pi}{G_F}} = 875 \text{ GeV}. \quad (2.6.65)$$

Quindi, per il processo che abbiamo studiato la teoria di Fermi viola l'unitarietà per alcuni valori di $\sqrt{s}$ (energia del centro di massa). A bassa energia però siamo sicuri nel dire che non viene violata l'unitarietà. Anche per altri processi si trova che ad alte energie la teoria di Fermi si rompe.

La rottura dell'unitarietà si può anche vedere dal fatto che la sezione d'urto dipende dall'energia, cresce infinitamente con l'energia, e ciò è dovuto alla dimensionalità della costante di accoppiamento:

$$[G_F] = M^{-2}. \quad (2.6.66)$$

¹²In realtà potevamo vedere la relazione (2.6.59) in termini generali, poiché la relazione di unitarietà (2.6.54) deve valere per ogni processo e non può esserci cancellazione tra le onde parziali. L'unitarietà implica:

$$|a_l(s)|^2 = \text{Im}\{a_l(s)\} \quad (2.6.60)$$

ma riscrivendo un numero complesso con la decomposizione polare, e prendendone la parte immaginaria, si vede:

$$|a_l(s)|^2 = |a_l(s)| \sin \alpha \implies |a_l(s)| \leq 1 \quad \forall l. \quad (2.6.61)$$

¹³Limiti simili, ma non identici, si trovano anche per altre ampiezze.

L'ordine di grandezza delle alte/basse energie ce lo dà il valore della costante di Fermi. Infatti, per via del fatto che $[\sigma] = M^{-2} = L^2$, se prendiamo un processo $2 \rightarrow 2$, che contiene solamente l'accoppiamento di Fermi al tree-level, allora per avere le corrette dimensioni per la sezione d'urto:

$$\sigma_{tree} \propto G_F^2 \quad \Rightarrow \quad \sigma_{tree} \propto G_F^2 E^2 \simeq G_F^2 s \quad (2.6.67)$$

in cui l'ultima uguaglianza la scriviamo nel caso di masse trascurabili. Dunque, si ha violazione di unitarietà per:

$$\sqrt{s} \gtrsim \frac{1}{\sqrt{G_F}} \quad (2.6.68)$$

per cui l'ordine di grandezza a cui avviene la violazione di unitarietà è dato proprio dal valore della costante di accoppiamento G_F :

$$\sqrt{s} \gtrsim 10^2 \text{ GeV}. \quad (2.6.69)$$

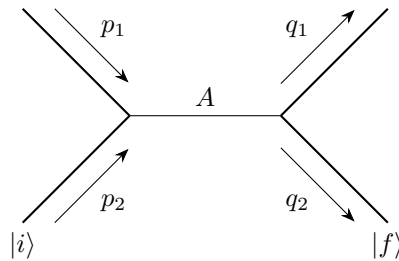
Tuttavia, la teoria di Fermi funziona bene e può essere utilizzata a basse energie, ad esempio predice bene il decadimento del μ che avviene a circa $\sim 100 \text{ MeV}$.

2.7 Rinormalizzazione

Il secondo grande problema della teoria di Fermi è la non rinormalizzabilità, ma prima di addentrarci al caso di Fermi analizziamo cosa ci dicono i parametri della lagrangiana e cosa vuol dire rinormalizzare una teoria.

2.7.1 Parametri della lagrangiana e rinormalizzazione

I parametri della lagrangiana (masse, accoppiamenti) non sono predetti dalla teoria, ma vengono fissati dagli esperimenti. Per esempio, il modo più semplice per fissare la massa di una particella A , è quello di produrre una risonanza, del tipo:



in cui sappiamo che il propagatore di A è proporzionale a:

$$\frac{1}{p^2 - m_A^2} \quad \text{con } p = p_1 + p_2 \quad , \quad p^2 = s. \quad (2.7.1)$$

Studiando il processo:

$$|i\rangle \longrightarrow |f\rangle \quad (2.7.2)$$

si vede che quando $s = p^2$ si avvicina a m_A^2 , allora il contributo di quel diagramma diventa grande, dunque la sezione d'urto del processo tende ad avere un picco in $s = m_A^2$, come nella figura 2.4.

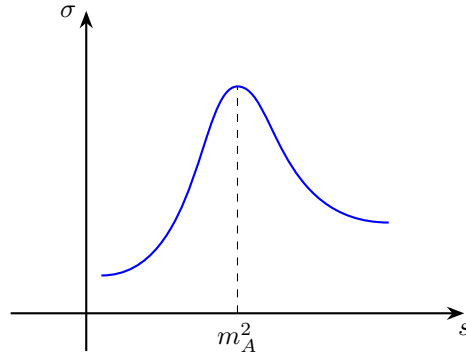
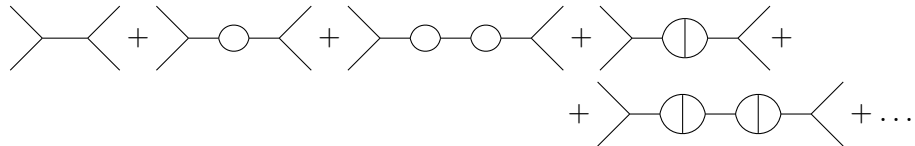


Figura 2.4

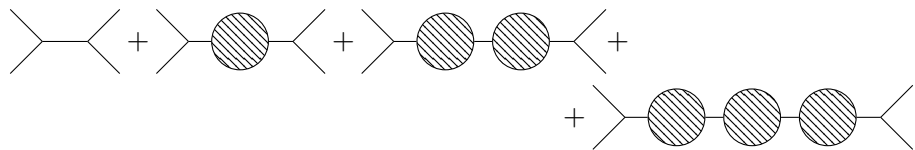
Dunque, nella posizione del picco possiamo determinare il valore di massa. Nota che abbiamo un picco in m_A^2 e non una divergenza, poiché il diagramma che abbiamo disegnato non è l'unico che dà contributo significativo per $s \sim m_A^2$. Infatti, avremo anche diagrammi al tree-level, ad 1-loop, a 2-loop etc. che danno contributo ogni volta che in uno di essi c'è un propagatore nel canale s :



Indicando con $\Sigma(p^2)$, la **self-energy di A**, l'insieme dei diagrammi 1PI, ossia i diagrammi *one-particle-irreducible* (cioè quelli che non possiamo spezzare in due rimuovendo una sola linea), possiamo scrivere:

$$\begin{aligned} \Sigma(p^2) &= \text{diagram with shaded circle} = \text{diagram with empty circle} + \text{diagram with half-filled circle} + \dots \\ &= \Sigma_1 + \Sigma_2 + \dots \end{aligned}$$

Scritta la self-energy di A possiamo riscrivere l'insieme dei diagrammi riso-
nanti come:



che saranno proporzionali a:

$$\frac{1}{p^2 - m_A^2} + \frac{1}{p^2 - m_A^2} \Sigma \frac{1}{p^2 - m_A^2} + \frac{1}{p^2 - m_A^2} \Sigma \frac{1}{p^2 - m_A^2} \Sigma \frac{1}{p^2 - m_A^2} + \dots \quad (2.7.3)$$

$$= \frac{1}{p^2 - m_A^2} \sum_{n=0}^{+\infty} \left(\Sigma \frac{1}{p^2 - m_A^2} \right)^n \quad (2.7.4)$$

$$= \frac{1}{p^2 - m_A^2} \frac{1}{1 - \frac{\Sigma(p^2)}{p^2 - m_A^2}} \quad (2.7.5)$$

$$= \frac{1}{p^2 - m_A^2 - \Sigma(p^2)} \quad (2.7.6)$$

che è parte di quello che si chiama **propagatore risommato alla Dyson**, in cui $\Sigma(p^2)$ è normalmente una quantità complessa. Proprio perché $\Sigma \in \mathbb{C}$ abbiamo che s non potrà mai annullare completamente il denominatore e dare una sezione d'urto divergente. Però vediamo che la risonanza si è spostata in:

$$s = m_A^2 + \Sigma. \quad (2.7.7)$$

Quindi, la posizione della risonanza, che corrisponde alla massa fisica (sperimentale) $m_{A,exp}^2$, non coincide con la massa nella lagrangiana, che è la massa nuda $m_{A,0}^2$, ma massa fisica e massa nuda sono legate da:

$$m_{A,exp}^2 = m_{A,0}^2 + \text{Re}\{\Sigma(m_{A,exp}^2)\}. \quad (2.7.8)$$

Dunque, abbiamo capito che il parametro della lagrangiana $m_{A,0}^2$ non ha diretto significato fisico, ma solo la combinazione (2.7.8) ce l'ha.

Succede una cosa analoga anche per le costanti di accoppiamento, il cui valore viene anche esso fissato sperimentalmente. Infatti, dobbiamo ancora considerare più diagrammi: supponiamo di avere nella lagrangiana un accoppiamento a 3 campi con costante di accoppiamento g_0 . Avremo la costante sperimentale:

$$g_{exp} = g_0 + V_1 + V_2 + \dots \quad (2.7.9)$$

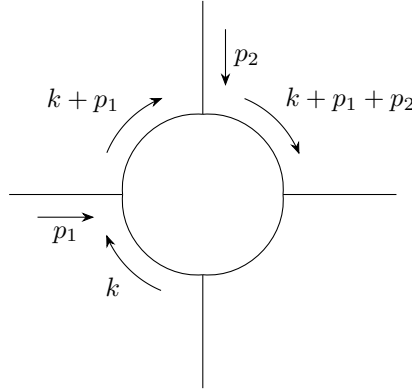
$$= \text{---} \langle \text{---} + \text{---} \bigcirc \text{---} + \text{---} \bigcirc \text{---} + \dots \quad (2.7.10)$$

in cui V_1 contiene tutti i contributi dei vertici a 1-loop con proporzionalità g_0^3 , V_2 tutti quelli a 2-loop, proporzionali a g_0^5 , e così via.

Analogamente a prima leghiamo in qualche modo g_0 a g_{exp} .

2.7.2 Divergenze ultraviolette

Le divergenze ultraviolette nascono in diagrammi contenenti 1 o più loop e sono generate dall'integrale sul quadrimomento del loop. Per esempio, il diagramma:



$$\propto \int d^4k f(k) \sim \int dk |k|^{\alpha+3} \quad (2.7.11)$$

in cui l'ultima uguaglianza è stata scritta nel caso in cui:

$$f(k) \underset{k \rightarrow \infty}{\sim} |k|^\alpha \implies d^4k f(k) \underset{k \rightarrow \infty}{\sim} dk |k|^{\alpha+3}. \quad (2.7.12)$$

Nota che l'integrale scritto in (2.7.11) non è limitato (lo dobbiamo fare) da nessuna delta di conservazione dell'impulso, per cui può raggiungere anche grandi valori di $|k|$, cioè alte energie (infatti alte energie vuol dire alte frequenze, e per questo diciamo divergenze ultraviolette). Da (2.7.11) vediamo che se:

$$\alpha \geq -4 \quad (2.7.13)$$

abbiamo un'integrale divergente. Infatti, noi sappiamo che all'interno di $f(k)$ ci saranno dei propagatori o possibili impulsi derivanti da vertici di interazione, per cui possiamo arrivare a divergenze. Però se in un diagramma ci sono un opportuno numero di propagatori (o altre dipendenze dall'impulso) possiamo rendere il diagramma non divergente.

In pratica gli integrali di loop vanno ad esplorare come si comporta la teoria a tutte le scale di energia, anche le parti alte, in cui l'integrale può divergere.

In realtà, la presenza di queste divergenze nella teoria può essere curata dalla **rinormalizzazione** (che va fatta anche nel caso in cui non si presentino divergenze ultraviolette).

Però, quello che si trova è che proprio i diagrammi che dobbiamo includere all'interno della nostra teoria, cioè le correzioni radiative per legare le masse e accoppiamenti nudi alle quantità fisiche, sono divergenti. Non è così sorprendente visto che nella Σ abbiamo come minimo due propagatori, per cui partiamo già da una potenza $\alpha = -4$ nella nostra funzione integranda $f(k)$ e, se non abbiamo delle potenze di impulsi derivanti da altre parti, quelli

divergono già in partenza. Ciò che succede di solito è che i parametri nudi contengono a loro volta delle divergenze, vanno a compensare quelle UV dei diagrammi a loop, e noi rimaniamo con le quantità fisiche non divergenti. Questi parametri curativi li chiamiamo *contro-termini*.

Il modo di procedere è quello di individuare le divergenze, togliere gli infiniti utilizzando i parametri della lagrangiana e ottenere una teoria senza infiniti e predittiva. Non ci deve disturbare troppo che mettiamo delle divergenze all'interno dei parametri di $\mathcal{L}$, poiché essi non sono le quantità che vediamo negli esperimenti.

La differenza tra teorie rinormalizzabili e non rinormalizzabili è che, nelle prime, una volta terminati i parametri di $\mathcal{L}$ con cui rimuovere gli infiniti non si hanno più diagrammi divergenti, mentre nelle seconde, tra cui la Teoria di Fermi come vedremo, per fare in modo di avere una teoria senza infiniti si devono aggiungere alla lagrangiana un numero infinito di parametri, che rendono la teoria non predittiva.

Vedremo anche che non tutto è perduto nel caso di teorie non rinormalizzabili, poiché con le giuste approssimazioni, saremo in grado di rendere comunque predittiva la teoria.

È più semplice provare la non rinormalizzabilità di una teoria, piuttosto che la rinormalizzabilità, infatti sarà questo che vedremo per la teoria di Fermi.

Esistono diversi *schemes* da poter seguire per rinormalizzare una teoria, che non sono oggetto di questo corso, nonostante per rinormalizzare la QCD si dovrà utilizzare uno di essi, ma puoi vedere le note del corso di *Advanced Quantum Field Theory*.

2.7.3 Rinormalizzazione e power counting

Il power counting è un metodo che fornisce una condizione necessaria, ma non sufficiente, affinché una teoria sia rinormalizzabile. Si basa sull'analisi dimensionale e sostanzialmente contiamo la potenza dei momenti di loop del diagramma.

Considerando un generico diagramma definiamo:

- L : numero di loop.
- V : numero di vertici.
- I_f : numero di linee interne fermioniche.
- I_b : numero di linee interne bosoniche.
- E_f : numero di linee esterne fermioniche.
- E_b : numero di linee esterne bosoniche.

- $n_f^{(i)}$: numero di linee fermioniche che convergono in un vertice i .
- $n_b^{(i)}$: numero di linee bosoniche che convergono in un vertice i .
- $d^{(i)}$: numero delle potenze dell'impulso nel vertice i (per effetto delle derivate dei campi nel corrispondente termine nella lagrangiana).

Notiamo che $d^{(i)}$ è legato alla potenza nulla, o positiva, in un vertice derivante dalla regola di Feynman, e quindi dalla forma di $\mathcal{L}_{int}$.

Nei nostri conti sommeremo sul numero di vertici, per cui vale:

$$\sum_{i=1}^V = \sum_i = V \quad (2.7.14)$$

e valgono le seguenti relazioni:

$$\sum_i n_f^i = 2I_f + E_f \quad (2.7.15)$$

$$\sum_i n_b^i = 2I_b + E_b \quad (2.7.16)$$

che si ottengono guardando quante linee entrano, o escono, in un vertice del diagramma (entrano 2 linee interne ed una esce esterna).

Possiamo vedere come scrivere il numero di loop, cioè il numero di impulsi interni indipendenti, in termini di linee interne e vertici:

$$L = I_f + I_b - (V - 1) \quad (2.7.17)$$

cioè, il numero di momenti indipendenti è il numero di linee interne, meno V condizioni di conservazione (delta), meno un'altra delta di conservazione globale.

Definiamo **grado di divergenza superficiale**:

$$\Delta = dL - I_f - 2I_b + \sum_i d^{(i)} \quad (2.7.18)$$

che rappresenta il numero di potenze dell'impulso a numeratore meno quelle a denominatore. Nell'espressione (2.7.18) indichiamo con d la dimensione dello spazio-tempo (per noi sarà $d = 4$). Nella formula abbiamo messo il termine dL , che è derivante dalla misura di integrazione, il pezzo $I_f + 2I_b$ che tiene conto del fatto che le linee interne contribuiscono con dei propagatori (potenze negative), e poi abbiamo la somma degli impulsi che vengono dai vertici.

Possiamo riscrivere (2.7.18), in $d = 4$, utilizzando le relazioni per i loop (2.7.17) e per i numeri di linee interne (2.7.15), (2.7.16):

$$\Delta = 4L - I_f - 2I_b + \sum_i d^{(i)} \quad (2.7.19)$$

$$= 4(I_f + I_b - (V - 1)) - I_f - 2I_b + \sum_i d^{(i)} \quad (2.7.20)$$

$$= 3I_f + 2I_b + 4 - 4V + \sum_i d^{(i)} \quad (2.7.21)$$

$$= 4 + 3 \frac{\sum_i n_f^i - E_f}{2} + 2 \frac{\sum_i n_b^i - E_b}{2} + \sum_i (d^{(i)} - 4) \quad (2.7.22)$$

$$= 4 - \frac{3}{2}E_f - E_b + \sum_i \left(d^{(i)} - 4 + \frac{3}{2}n_f^i + n_b^i \right). \quad (2.7.23)$$

Per $\Delta \geq 0$ l'ampiezza del diagramma contiene una divergenza ultravioletta, per cui se:

$$d^{(i)} - 4 + \frac{3}{2}n_f^i + n_b^i > 0 \quad (2.7.24)$$

allora Δ cresce al crescere del numero di vertici V . Quindi, anche con un diagramma, che avrà un certo numero di gambe esterne E_f e E_b , con $\Delta < 0$, quando, a fissate gambe esterne, aumentiamo il numero di vertici, il che vuol dire aumentare i loop, allora finiremo sicuramente con un $\Delta \geq 0$ e quindi dobbiamo rinormalizzare l'interazione con quel tale numero di gambe esterne E_f , E_b .

Questo succede per qualunque E_f e E_b . Nella pratica dovremo rinormalizzare (cioé riassorbire le divergenze nei parametri di $\mathcal{L}$) ogni diagramma della teoria, e sicuramente non abbiamo abbastanza parametri nella lagrangiana per farlo.

La condizione di rinormalizzazione diventa dunque:

$$d^i + \frac{3}{2}n_f^i + n_b^i \leq 4 \quad \forall i \quad (2.7.25)$$

in cui il membro di destra rappresenta la dimensionalità dell'operatore che compare nel vertice i . Infatti, sappiamo che i campi fermionici, contati con n_f^i , hanno dimensione di una massa alla $3/2$, mentre i bosonici n_b hanno dimensione 1 e le derivate 1. Però sappiamo che nel vertice c'è la costante di accoppiamento, e infatti tale condizione di rinormalizzazione si può riformulare in termini di costante g_i del vertice. Cominciamo dalla conoscenza

della dimensionalità della lagrangiana $[\mathcal{L}] = M^4$, per cui:

$$\dim(\mathcal{L}) = \dim(g_i) + d^i + \frac{3}{2}n_f^i + n_b^i \quad (2.7.26)$$

$$4 = \dim(g_i) + d^i + \frac{3}{2}n_f^i + n_b^i \quad (2.7.27)$$

$$\dim(g_i) = 4 - d^i - \frac{3}{2}n_f^i - n_b^i \quad (2.7.28)$$

e utilizzando (2.7.25) abbiamo la **condizione di rinormalizzabilità**:

$$\dim(g_i) \geq 0. \quad (2.7.29)$$

Dunque, nel caso in cui $\dim(g_i) < 0$ abbiamo una teoria non rinormalizzabile, e se guardiamo bene noi abbiamo:

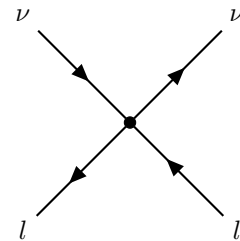
$$\dim(G_F) = -2 \quad (2.7.30)$$

dunque, la teoria di Fermi non è rinormalizzabile e non può essere la teoria definitiva per l'interazione elettrodebole *nel suo complesso*. Ossia, non può essere la teoria dell'interazione debole a tutte le scale di energia. Vedremo però che, nonostante l'essere malata, la teoria di Fermi ci potrà guidare verso una teoria completa, il Modello Standard.

2.7.4 Rinormalizzare la teoria di Fermi

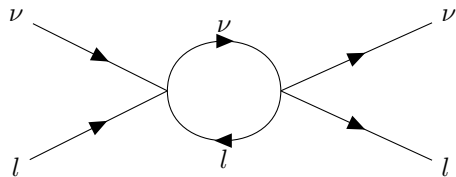
Possiamo vedere nella pratica cosa succede se proviamo, sapendo già di fallire, a rinormalizzare la teoria di Fermi e cerchiamo di capire cosa portarci dietro per la costruzione della teoria elettrodebole completa.

La teoria di fermi sappiamo avere le caratteristiche:

$$\mathcal{L}_F = \frac{G_F}{\sqrt{2}} (\bar{\psi}_\nu \Gamma^\mu \psi_l) (\bar{\psi}_\nu \Gamma^\mu \psi_l)^\dagger \quad \propto G_F$$


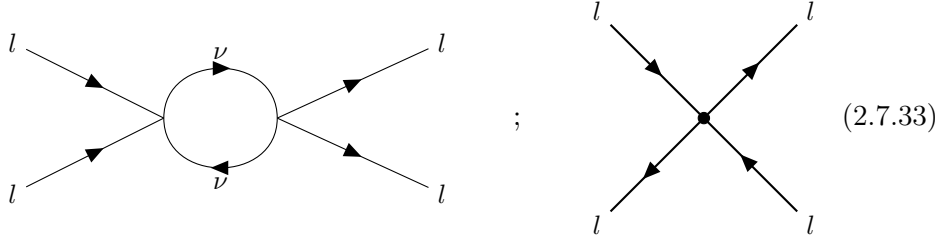
(2.7.31)

con $[G_F] = M^{-2}$. Ad 1-loop abbiamo alcuni diagrammi divergenti; ad esempio:



(2.7.32)

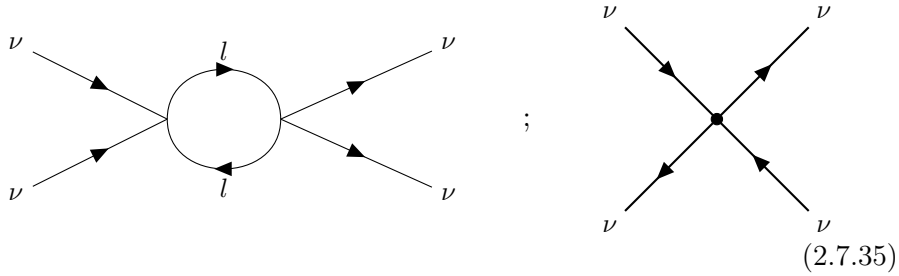
che è per l'appunto divergente e lo riassorbiamo all'interno del vertice (2.7.31). Analogamente il diagramma (di sinistra):


(2.7.33)

essendo divergente può essere riassorbito nel diagramma di destra. Però per riassorbire la divergenza dovremmo rendere possibile nella teoria di Fermi il diagramma di destra, ovvero, dovremmo includere nella lagrangiana un termine di accoppiamento tra correnti neutre:

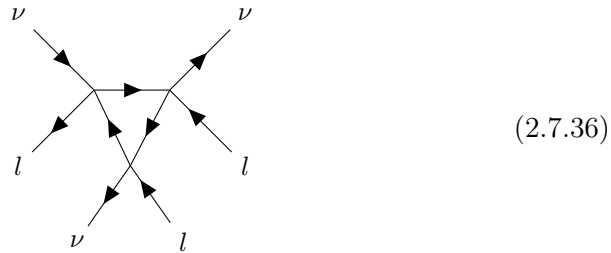
$$\frac{G_N}{\sqrt{2}} (\bar{\psi}_l \Gamma^\mu \psi_l) (\bar{\psi}_\nu \Gamma^\mu \psi_\nu) \quad , \quad [G_N] = M^{-2}. \quad (2.7.34)$$

Stesso ragionamento per il diagramma:

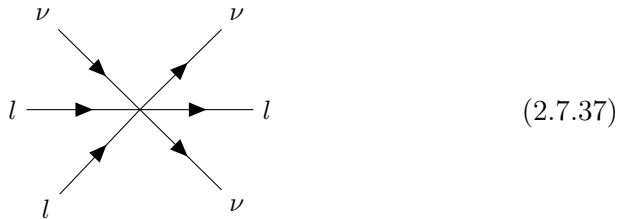

(2.7.35)

Questo è il modo in cui possiamo rinormalizzare i processi $2 \rightarrow 2$.

Però, se prendiamo processi che coinvolgono 6 particelle, ad esempio il diagramma:


(2.7.36)

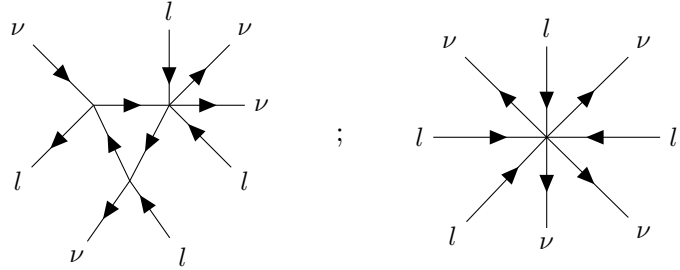
esso risulta ancora divergente e dovremmo riassorbirlo in un vertice di interazione:


(2.7.37)

aggiustando ancora $\mathcal{L}$ con un termine di accoppiamento a tre correnti cariche:

$$\frac{G'}{\sqrt{2}} (\bar{\psi}_\nu \Gamma^\mu \psi_l) (\bar{\psi}_\nu \Gamma^\mu \psi_l) (\bar{\psi}_\nu \Gamma^\mu \psi_l) \quad , \quad [G'] = M^{-5}. \quad (2.7.38)$$

Questo procedimento continua e bisognerebbe via via aggiungere sempre più termini a $\mathcal{L}_F$ e ciò renderebbe questo processo senza fine, poiché i termini che continuiamo ad aggiungere contengono a loro volta divergenze. Ad esempio, avremmo ancora:



$$(2.7.39)$$

che è divergente (quello di sinistra) e lo si può riassorbire nel contro-termine (quello di destra) aggiungendo l'accoppiamento:

$$\frac{G''}{\sqrt{2}} (\bar{\psi}_\nu \Gamma^\mu \psi_l) (\bar{\psi}_\nu \Gamma^\mu \psi_l) (\bar{\psi}_\nu \Gamma^\mu \psi_l) (\bar{\psi}_\nu \Gamma^\mu \psi_l) \quad , \quad [G''] = M^{-8}. \quad (2.7.40)$$

Un modo per limitare il proliferare dei termini da aggiungere alla lagrangiana, affinché essa possa essere utilizzata è il seguente. Scriviamo gli accoppiamenti estraendo la loro parte dimensionale:

$$G_F = \frac{g_F}{M^2} \quad ; \quad G_N = \frac{g_N}{M^2} \quad ; \quad G' = \frac{g'}{M^5} \quad ; \quad G'' = \frac{g''}{M^8} \quad (2.7.41)$$

così che:

$$\dim(g_i) = 0. \quad (2.7.42)$$

In questo modo, qualunque ampiezza ha dimensione zero e se immaginiamo uno scattering $2 \rightarrow 2$, questo sarà del tipo:

$$\begin{aligned} \mathcal{M}(2 \rightarrow 2) &= G_F A_0^F + G_F^2 A_1^F + \cdots + G_N A_0^N + G_N^2 A_1^N + \cdots + G' A_1' + \\ &\quad + \cdots + G'' A'' + \cdots \\ &= \text{diagrammi} + \cdots + \text{diagrammi} + \cdots + \text{diagrammi} + \cdots \\ &= g_F \frac{A_0^F}{M^2} + g_F^2 \frac{A_1^F}{M^4} + \cdots + g_N \frac{A_0^N}{M^2} + g_N^2 \frac{A_1^N}{M^4} + \cdots + g' \frac{A_1'}{M^5} + \\ &\quad + \cdots + g'' \frac{A_1''}{M^8} + \cdots \end{aligned}$$

e sapendo che qualunque ampiezza ha dimensionalità nulla, ovvero $\dim(\mathcal{M}) = 0$, possiamo dire che qualunque rapporto A_j^i/M^k è adimensionale.

Se indichiamo l'energia delle particelle in gioco con:

$$E = \frac{\sqrt{s}}{2} \quad (2.7.43)$$

allora solamente per analisi dimensionale vediamo che:

$$\frac{A_0^F}{M^2}, \frac{A_0^N}{M^2} \propto \frac{E^2}{M^2} \quad ; \quad \frac{A_1^F}{M^4} \propto \frac{E^4}{M^4} \quad (2.7.44)$$

$$\frac{A_1'}{M^5} \propto \frac{E^5}{M^5} \quad ; \quad \frac{A_1''}{M^8} \propto \frac{E^8}{M^8} \quad (2.7.45)$$

il che ci mostra quali processi possiamo descrivere con la nostra teoria. L'ipotesi di base che ci serve fare è che tutte le costanti di accoppiamento G siano dello stesso ordine di grandezza. Infatti, se ci limitiamo¹⁴ a considerare il processo ad energie $E \ll M$, allora i contributi di A_1^F, A', A'' sono trascurabili e quindi possiamo limitare la teoria all'interazione di un solo vertice a due correnti, e avere una teoria predittiva con un errore $\mathcal{O}(E^4/M^4)$.

In generale, se vogliamo una predizione precisa all'ordine $(E/M)^n$, allora possiamo provare ad aggiungere termini di interazione (sempre a patto di poter misurare sperimentalmente i processi necessari) ad $\mathcal{L}$, rinormalizzarli ed escludere solo i termini di ordine superiore.

In questo modo non abbiamo teorie che sono valide in tutto il range energetico, ma valgono (sono rinormalizzabili e predittive) solo fino ad un certo valore. Ciò che otteniamo sono le cosiddette *teorie efficaci*.

In pratica, dunque, anche le teorie non rinormalizzabili (come la teoria di Fermi) possono ancora essere predittive, ma a patto di:

- Limitarsi ad un certo range di energia, per cui $E \ll M$.
- Includere tutti i contributi (rinormalizzare la teoria) fino ad un certo ordine $(E/M)^n$.

Per rendere la teoria di Fermi, almeno a basse energie, abbiamo bisogno di includere nella teoria i vertici con G_F (2.7.31) e G_N (2.7.33) e (2.7.35), ovvero i vertici (2.3.53), che avevamo già trovato parlando della simmetria $SU(2)$ della teoria.

Nota. Abbiamo trovato che la teoria di Fermi non è né unitaria né rinormalizzare in senso assoluto, però può essere unitaria e rinormalizzabile, dunque predittiva, a patto che ci limitiamo ad un certo range energetico, legato alla costante di accoppiamento dimensionale G_F .

¹⁴Nota che il valore di M è in qualche modo fissato dalla conoscenza delle costanti di accoppiamento $G_i = g_i/M^k$.

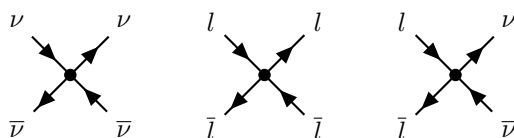
Capitolo 3

Teoria di gauge elettrodebole

Vediamo in questo capitolo come costruiamo una teoria elettrodebole mediata da dei vettori di gauge, e come essa sarà parte del Modello Standard delle particelle elementari.

3.1 Piano di lavoro

Per costruire una teoria elettrodebole non ripartiremo completamente da zero, ma utilizzeremo gli insegnamenti di Fermi. Infatti, abbiamo visto che la teoria di Fermi, pur descrivendo bene i processi deboli alle basse energie, non è una teoria di campo consistente, perché viola l'unitarietà e non è rinormalizzabile. Entrambi i problemi abbiamo visto essere legati alla dimensionalità della costante di accoppiamento G_F . Inoltre, provando a rinormalizzare la teoria e fermandoci nel limite a basse energie, abbiamo notato che i vertici:



cioè un'accoppiamento di correnti neutre, dovessero essere aggiunte alla nostra teoria. I termini di corrente che ci rendono la teoria di Fermi rinormalizzabile a basse energie, sono anche i termini che ci permettevano di completare la simmetria di $SU(2)$ della teoria.

Per costruire la nuova teoria, terremo a mente quali tipi di accoppiamento sicuramente funzionano, ossia quelli di $SU(2)$ della teoria di Fermi, e ispirandoci alla QED proveremo a costruire una teoria di gauge. Dunque, diremo che esiste un campo di gauge W_μ di $SU(2)$, il cui accoppiamento, tramite una costante di accoppiamento adimensionale (così siamo sicuri di rinormalizzare), con i fermioni è dato dalla derivata covariante dell'accoppiamento minimale.

Ricordiamo cosa succedeva in QED. La lagrangiana completa è:

$$\mathcal{L}_{QED} = -\frac{1}{4}F_{\mu\nu}F^{\mu\nu} + \bar{\psi}(i\not{\partial} - m)\psi + eQ\bar{\psi}\gamma^\mu A_\mu\psi \quad (3.1.1)$$

in cui dovremmo inserire la carica elettrica dell'elettrone $Q = -1$. Conosciamo anche le caratteristiche di $\mathcal{L}_{QED}$, ossia:

- La costante di accoppiamento e è adimensionale.
- La lagrangiana è invariante per trasformazioni di gauge $U(1)$ locali del tipo:

$$\psi(x) \longrightarrow e^{-ieQ\alpha(x)}\psi \quad (3.1.2)$$

$$A_\mu(x) \longrightarrow A_\mu(x) - \partial_\mu\alpha(x). \quad (3.1.3)$$

- Il termine di accoppiamento:

$$eQ\bar{\psi}\gamma^\mu A_\mu\psi \quad (3.1.4)$$

si genera a partire dalla lagrangiana di Dirac, promuovendo la derivata parziale ∂_μ ordinaria in derivata covariante:

$$\bar{\psi}(i\not{\partial} - m)\psi \longrightarrow \bar{\psi}(i\not{D} - m)\psi \quad (3.1.5)$$

in cui:

$$\mathcal{D}_\mu = \partial_\mu - ieQA_\mu(x) \quad (3.1.6)$$

che contiene, oltre il campo elettromagnetico, la costante di accoppiamento adimensionale e .

Però potremmo anche essere tentati di riciclare brutalmente la QED e utilizzare la stessa lagrangiana per l'interazione elettrodebole, ovviamente aggiungendo i pezzi di accoppiamento necessari. Però, come si può subito notare dalla QED, per avere una costante adimensionale, non si possono accoppiare tra loro direttamente due correnti fermioniche (ognuna con dimensionalità 3), ma bisogna accoppiare una corrente fermionica ad un campo bosonico (di dimensionalità 1) che è un mediatore dell'interazione:

$$eQ\bar{\psi}\gamma^\mu A_\mu\psi. \quad (3.1.7)$$

Nel caso della QED il gruppo di trasformazioni è $U(1)$, ma noi utilizzeremo il gruppo $SU(2)$, per i motivi già visti.

Inoltre, noi dovremo trovare anche il modo di separare le parti left e right all'interno della nuova lagrangiana, poiché sappiamo che le interazioni deboli parlano solo con gli spinori left. La rottura di parità ci porterà a dover introdurre una derivata covariante per ogni chiralità dei fermioni.

3.1.1 Correnti leptoniche e simmetria $SU(2)$

Possiamo provare a mimare quanto fatto per la QED anche per le interazioni deboli, e possiamo farlo sfruttando la *quasi* simmetria della lagrangiana di Fermi per trasformazioni globali di $SU(2)$. Cominciamo considerando solamente una famiglia leptonica.

Vogliamo includere non solo le correnti cariche, ma anche quella neutra, poiché nella sezione §2.7.4 ci siamo resi conto che J_3^μ fosse necessaria per avere una teoria predittiva.

Ripartiamo dalle correnti deboli puramente left:

$$J_i^\mu = \bar{L}_L \gamma^\mu t_i L_L \quad (3.1.8)$$

in cui abbiamo $t_i = \tau_i/2$, $i = 1, 2, 3$, τ_i le matrici di Pauli e il doppietto left definito come:

$$L_L = \begin{pmatrix} \psi_{\nu_L} \\ \psi_{l_L} \end{pmatrix} = \frac{1 - \gamma_5}{2} \begin{pmatrix} \psi_\nu \\ \psi_l \end{pmatrix}. \quad (3.1.9)$$

Dalla sezione §2.3.1 avevamo visto come trasformassero le correnti, ovvero, con la rappresentazione aggiunta di $SU(2)$:

$$\begin{pmatrix} J_1^\mu \\ J_2^\mu \\ J_3^\mu \end{pmatrix} \longrightarrow U(\theta) \begin{pmatrix} J_1^\mu \\ J_2^\mu \\ J_3^\mu \end{pmatrix} = e^{i\theta_a T_a} \begin{pmatrix} J_1^\mu \\ J_2^\mu \\ J_3^\mu \end{pmatrix} \quad (3.1.10)$$

in cui $(T_a)_{ij} = i\epsilon_{iaj}$ sono i generatori della rappresentazione aggiunta di $SU(2)$ (ovvero le costanti di struttura del gruppo).

Dall'accoppiamento minimale sappiamo che dobbiamo accoppiare (appunto) le correnti, della nostra simmetria globale $SU(2)$, con dei campi di gauge (sono più di uno avendo un gruppo di simmetria di dimensione > 1), associati ciascuno ad un generatore del gruppo. Dunque, introduciamo un tripletto di campi bosonici W_i^μ , con $i = 1, 2, 3$, che si accoppino con J_i^μ . Notiamo però che vogliamo mantenere l'invarianza della lagrangiana, per cui quando accoppiamo bosoni e correnti questo ci fissa il modo di trasformare del tripletto di bosoni di gauge. Si ha che anch'essi trasformino (globalmente) con la rappresentazione aggiunta di $SU(2)$:

$$\begin{pmatrix} W_1^\mu \\ W_2^\mu \\ W_3^\mu \end{pmatrix} \longrightarrow U^\dagger(\theta) \begin{pmatrix} W_1^\mu \\ W_2^\mu \\ W_3^\mu \end{pmatrix} = e^{-i\theta_a T_a} \begin{pmatrix} W_1^\mu \\ W_2^\mu \\ W_3^\mu \end{pmatrix}. \quad (3.1.11)$$

Il termine di interazione debole (minimale), invariante per trasformazioni globali di $SU(2)$, che possiamo scrivere è:

$$\mathcal{L}_{int} = g \sum_{i=1}^3 J_{i,\mu} W_i^\mu \quad (3.1.12)$$

$$= g \sum_{i=1}^3 \bar{L}_L \gamma_\mu t_i L_L W_i^\mu \quad (3.1.13)$$

$$= g \left[\bar{L}_L \gamma_\mu t_1 L_L W_1^\mu + \bar{L}_L \gamma_\mu t_2 L_L W_2^\mu + \bar{L}_L \gamma_\mu t_3 L_L W_3^\mu \right] \quad (3.1.14)$$

in cui notiamo la presenza (che deriva dal completamento della simmetria di $SU(2)$) della corrente neutra:

$$J_3^\mu = \bar{L}_L \gamma_\mu t_3 L_L \quad (3.1.15)$$

$$= \frac{1}{2} (\bar{\psi}_{\nu_L} \quad \bar{\psi}_{l_L}) \gamma^\mu \begin{pmatrix} 1 & 0 \\ 0 & -1 \end{pmatrix} \begin{pmatrix} \psi_{\nu_L} \\ \psi_{l_L} \end{pmatrix} \quad (3.1.16)$$

$$= \frac{1}{2} \bar{\psi}_{\nu_L} \gamma^\mu \psi_{\nu_L} - \frac{1}{2} \bar{\psi}_{l_L} \gamma^\mu \psi_{l_L}. \quad (3.1.17)$$

Abbiamo già ragionato sul perché la corrente J_3^μ non possa essere quella elettromagnetica, anche se ci somiglia molto:

$$J_{QED}^\mu = \bar{\psi}_l \gamma^\mu \psi_l. \quad (3.1.18)$$

Possiamo evidenziare importanti differenze:

- La corrente debole J_3^μ contiene sia particelle cariche (i leptoni l), ma anche particelle neutre (i neutrini ν), che sappiamo bene non accoppiarsi con il fotone.
- La corrente debole J_3^μ contiene solamente la componente left dei campi fermionici, generando quindi un accoppiamento che viola massimamente la parità, ma noi sappiamo che la corrente contiene tutto il campo fermionico e si vede sperimentalmente (vedi il corso di *Particelle 1* per la verifica sperimentale) che conserva la parità.

Dunque, l'accoppiamento $J_3^\mu W_\mu^3$ rappresenta un nuovo accoppiamento rispetto la QED. Questo nuovo accoppiamento che coinvolge correnti neutre non è visibile alle basse energie in cui invece sono evidenti gli accoppiamenti tra correnti cariche (e per questo motivo non è stato inserito nella sua lagrangiana da Fermi). Vedremo in seguito, quando completeremo la lagrangiana del modello standard, che questo fatto è ben spiegabile dalla teoria. La spiegazione si basa proprio sulla somiglianza con l'accoppiamento elettromagnetico, che alle basse energie è predominante rispetto all'accoppiamento debole tra correnti neutre, che risulta quindi schermato a quelle energie (come già in parte avevamo anticipato nella sezione della teoria di Fermi).

3.2 Interazione elettrodebole per correnti leptoniche

Possiamo vedere che l'interazione elettrodebole, cioè il pezzo di interazione (3.1.14) viene fuori anche se seguiamo ciecamente il processo di accoppiamento minimale e introduciamo la derivata covariante agente sugli spinori left:

$$\mathcal{D}_\mu L_L = \partial_\mu L_L - ig \sum_{i=1}^3 W_\mu^i t_i L_L. \quad (3.2.1)$$

Per vederlo esplicitamente dobbiamo riscrivere la lagrangiana per gli spinori, dunque la lagrangiana di Dirac, separando le parti left e right. Consideriamo la lagrangiana di Dirac massless, poiché come ci renderemo presto conto il termine massivo romperebbe tutto ciò che stiamo cercando di costruire e lo potremmo introdurre solo in un secondo momento. Partiamo da $\mathcal{L}_{Dirac}$ e separiamo gli spinori nelle loro due parti:

$$\mathcal{L}_{Dirac} = i\bar{\psi}_\nu \not{\partial} \psi_\nu + i\bar{\psi}_l \not{\partial} \psi_l \quad (3.2.2)$$

$$= i(\bar{\psi}_{\nu L} + \bar{\psi}_{\nu R})\not{\partial}(\psi_{\nu L} + \psi_{\nu R}) + i(\bar{\psi}_{lL} + \bar{\psi}_{lR})\not{\partial}(\psi_{lL} + \psi_{lR}) \quad (3.2.3)$$

possiamo sfruttare l'ortogonalità delle due fasi:

$$\bar{\psi}_L \gamma^\mu \psi_R = \left(\psi \frac{1 - \gamma_5}{2} \right)^\dagger \gamma^0 \gamma^\mu \left(\frac{1 + \gamma_5}{2} \psi \right) \quad (3.2.4)$$

$$= \psi^\dagger \gamma^0 \gamma^\mu \frac{1 - \gamma_5}{2} \frac{1 + \gamma_5}{2} \psi \quad (3.2.5)$$

$$= 0 \quad (3.2.6)$$

di cui ne esiste una analoga della forma $\bar{\psi}_R \gamma^\mu \psi_L$. Se facciamo questo allora abbiamo:

$$\mathcal{L}_{Dirac} = i\bar{\psi}_\nu \not{\partial} \psi_\nu + i\bar{\psi}_l \not{\partial} \psi_l \quad (3.2.7)$$

$$= i\bar{\psi}_{\nu L} \not{\partial} \psi_{\nu L} + i\bar{\psi}_{\nu R} \not{\partial} \psi_{\nu R} + i\bar{\psi}_{lL} \not{\partial} \psi_{lL} + i\bar{\psi}_{lR} \not{\partial} \psi_{lR} \quad (3.2.8)$$

$$= i\bar{L}_L \not{\partial} L_L + i\bar{\psi}_{\nu R} \not{\partial} \psi_{\nu R} + i\bar{\psi}_{lR} \not{\partial} \psi_{lR}. \quad (3.2.9)$$

Infatti, si può vedere che se prendiamo il termine cinetico di $\mathcal{L}_{Dirac}$ per gli spinori left e facciamo accoppiamento minimale:

$$i\bar{L}_L \gamma^\mu \partial_\mu L_L \longrightarrow i\bar{L}_L \gamma^\mu \mathcal{D}_\mu L_L \quad (3.2.10)$$

$$= i\bar{L}_L \gamma^\mu \partial_\mu L_L + g \sum_{i=1}^3 W_\mu^i \bar{L}_L \gamma^\mu t_i L_L \quad (3.2.11)$$

$$= i\bar{L}_L \gamma^\mu \partial_\mu L_L + g \sum_{i=1}^3 W_\mu^i J_i^\mu \quad (3.2.12)$$

otteniamo il termine libero con aggiunto il termine di accoppiamento con i bosoni di gauge, che ci garantisce un'invarianza locale $SU(2)$.

Possiamo dunque sostituire la derivata ordinaria con la derivata covariante per gli spinori left:

$$\mathcal{L} = i\bar{L}_L \not{D} L_L + i\bar{\psi}_{\nu R} \not{D} \psi_{\nu R} + i\bar{\psi}_{lR} \not{D} \psi_{lR}. \quad (3.2.13)$$

e ottenere una teoria con un accoppiamento tra i W_μ e le parti left dei leptoni. Però, notiamo anche che la lagrangiana (3.2.13) contiene degli spinori right che non hanno accoppiamenti deboli, ma che, almeno quelli carichi, dovrebbero accoppiarsi con il campo elettromagnetico.

Dunque, visto che noi vorremmo avere nella stessa teoria sia l'accoppiamento debole che quello elettromagnetico, allora possiamo introdurre un accoppiamento abeliano $U(1)$ (visto che sappiamo che la QED è una teoria di gauge di quel tipo). Possiamo quindi definire le derivate covarianti di tutti gli spinori left e right in modo da generare anche l'accoppiamento (tipo) elettromagnetico. Abbiamo:

$$\mathcal{D}^\mu L_L = \partial^\mu L_L - igW_a^\mu t_a L_L - ig' B^\mu \frac{Y(L_L)}{2} L_L \quad (3.2.14)$$

$$\mathcal{D}^\mu \psi_{\nu R} = \partial^\mu \psi_{\nu R} - ig' B^\mu \frac{Y(\nu_R)}{2} \psi_{\nu R} \quad (3.2.15)$$

$$\mathcal{D}^\mu \psi_{lR} = \partial^\mu \psi_{lR} - ig' B^\mu \frac{Y(l_R)}{2} \psi_{lR} \quad (3.2.16)$$

in cui abbiamo messo gli spinori left accoppiati sia con i bosoni dell'interazione debole, che con il bosone abeliano B^μ , ma gli spinori right solamente con l' $U(1)$, dunque B^μ .

Nota che è noto che i neutrini, essendo considerati massless, e neutri, nel Modello Standard che stiamo costruendo, non avranno parte right accoppiata con nulla, però possiamo continuare a portarci dietro senza fare assunzioni e vedere se si disaccoppiano in modo autonomo.

Come si può notare non abbiamo definito l'accoppiamento $U(1)$ direttamente come l'accoppiamento di QED con il campo A_μ , ma abbiamo introdotto un campo B_μ , per ora non definito. Inoltre, non abbiamo fissato l'accoppiamento $U(1)$ attraverso la carica elettrica Q e la costante d'accoppiamento e , ma di nuovo abbiamo introdotto una nuova costante d'accoppiamento g' e una nuova "carica" Y (normalizzata 1/2), che chiamiamo **ipercarica**. Tutto questo l'abbiamo fatto, perché anche la parte con W_μ^3 di $\mathcal{D}^\mu L_L$ genera un accoppiamento simile alla QED e solo l'unione di questo accoppiamento con quelli $U(1)$, che abbiamo definito, deve riprodurre il corretto accoppiamento elettromagnetico.

Introducendo il nuovo campo B_μ , la costante g' e l'ipercarica (tutti da definire al momento) abbiamo semplicemente introdotto il termine $U(1)$ più

generale possibile, sperando così di poter fissare B_μ , g' e Y in modo da riprodurre l'accoppiamento di QED (vedremo in effetti che la soluzione banale $B_\mu = A_\mu$, $g' = e$ ed $Y = 2Q$ non funziona).

Per quanto riguarda l'accoppiamento con W^μ , come già detto, sappiamo che deve riprodurre quanto insegnato dalla teoria di Fermi.

Introducendo le derivate covarianti appena definite nella lagrangiana di Dirac (3.2.13) (per una sola famiglia leptonica) senza il termine di massa, otteniamo:

$$\mathcal{L} = i\bar{L}_L \not{D} L_L + i\bar{\psi}_{\nu R} \not{D} \psi_{\nu R} + i\bar{\psi}_{lR} \not{D} \psi_{lR} \quad (3.2.17)$$

$$= \mathcal{L}_{Dirac} + \mathcal{L}_C + \mathcal{L}_N \quad (3.2.18)$$

in cui, espandendo i vari accoppiamenti con i campi di gauge, individuiamo il termine di Dirac massless, l'accoppiamento con le correnti cariche e l'accoppiamento con le correnti neutre:

$$\mathcal{L}_{Dirac} = i\bar{L}_L \not{D} L_L + i\bar{\psi}_{\nu R} \not{D} \psi_{\nu R} + i\bar{\psi}_{lR} \not{D} \psi_{lR} \quad (3.2.19)$$

$$= i\bar{\psi}_\nu \not{D} \psi_\nu + i\bar{\psi}_l \not{D} \psi_l \quad (3.2.20)$$

$$\mathcal{L}_C = g \left[\bar{L}_L W_1 t_1 L_L + \bar{L}_L W_2 t_2 L_L \right] \quad (3.2.21)$$

$$\begin{aligned} \mathcal{L}_N = g\bar{L}_L W_3 t_3 L_L + \\ + \frac{g'}{2} \left[Y(L_L) \bar{L}_L \not{B} L_L + Y(\nu_R) \bar{\psi}_{\nu R} \not{B} \psi_{\nu R} + Y(l_R) \bar{\psi}_{lR} \not{B} \psi_{lR} \right]. \end{aligned} \quad (3.2.22)$$

Ricordando l'espressioni delle matrici di scala, costruite a partire dai generatori di $SU(2)$:

$$\tau^+ = t_1 + it_2 = \begin{pmatrix} 0 & 1 \\ 0 & 0 \end{pmatrix} \quad (3.2.23)$$

$$\tau^- = t_1 - it_2 = \begin{pmatrix} 0 & 0 \\ 1 & 0 \end{pmatrix} \quad (3.2.24)$$

introducendo i campi (carichi):

$$(W^\pm)^\mu = \frac{W_1^\mu \mp iW_2^\mu}{\sqrt{2}} \quad (3.2.25)$$

che sono uno l'aggiunto dell'altro ($W^\mp = (W^\pm)^\dagger$), e che se invertiti ci danno:

$$W_1^\mu = \frac{W_\mu^+ + W_\mu^-}{\sqrt{2}}, \quad W_2^\mu = i \frac{W_\mu^+ - W_\mu^-}{\sqrt{2}} \quad (3.2.26)$$

così possiamo riscrivere il termine di corrente carica in termini dei campi $\psi_{\nu L}$ e ψ_{lL} :

$$\mathcal{L}_C = g \left[\bar{L}_L \frac{W^+ + W^-}{\sqrt{2}} t_1 L_L + \bar{L}_L i \frac{W^+ - W^-}{\sqrt{2}} t_2 L_L \right] \quad (3.2.27)$$

$$= \frac{g}{\sqrt{2}} \left[\bar{L}_L W^+ \tau^+ L_L + \bar{L}_L W^- \tau^- L_L \right] \quad (3.2.28)$$

$$= \frac{g}{\sqrt{2}} \left[\bar{\psi}_{\nu L} W^+ \psi_{lL} + \bar{\psi}_{lL} W^- \psi_{\nu L} \right]. \quad (3.2.29)$$

Possiamo anche fare la stessa cosa per il settore neutro:

$$\begin{aligned} \mathcal{L}_N &= g \bar{L}_L W_3 t_3 L_L + \\ &+ \frac{g'}{2} \left[Y(L_L) \bar{L}_L \not{B} L_L + Y(\nu_R) \bar{\psi}_{\nu R} \not{B} \psi_{\nu R} + Y(l_R) \bar{\psi}_{lR} \not{B} \psi_{lR} \right] \end{aligned} \quad (3.2.30)$$

$$\begin{aligned} &= \frac{g}{2} \left[\bar{\psi}_{\nu L} W_3 \psi_{\nu L} - \bar{\psi}_{lL} W_3 \psi_{lL} \right] + \\ &+ \frac{g'}{2} \left[Y(L_L) \bar{\psi}_{\nu L} \not{B} \psi_{\nu L} + Y(L_L) \bar{\psi}_{lL} \not{B} \psi_{lL} + \right. \\ &\left. + Y(\nu_R) \bar{\psi}_{\nu R} \not{B} \psi_{\nu R} + Y(l_R) \bar{\psi}_{lR} \not{B} \psi_{lR} \right]. \end{aligned} \quad (3.2.31)$$

Possiamo introdurre l'isospin I_3 come l'*accoppiamento che hanno i vari campi left e right con W_3* ; guardando bene¹ (3.2.31) vediamo:

$$I_3(\nu_L) = \frac{1}{2} \quad , \quad I_3(l_L) = -\frac{1}{2} \quad , \quad I_3(\nu_R) = I_3(l_R) = 0. \quad (3.2.33)$$

In questo modo possiamo riscrivere $\mathcal{L}_N$ in modo più compatto:

$$\begin{aligned} \mathcal{L}_N &= g \left[I_3(\nu_L) \bar{\psi}_{\nu L} W_3 \psi_{\nu L} + I_3(l_L) \bar{\psi}_{lL} W_3 \psi_{lL} + \right. \\ &+ I_3(\nu_R) \bar{\psi}_{\nu R} W_3 \psi_{\nu R} + I_3(l_R) \bar{\psi}_{lR} W_3 \psi_{lR} \left. \right] + \\ &+ \frac{g'}{2} \left[Y(L_L) \bar{\psi}_{\nu L} \not{B} \psi_{\nu L} + Y(L_L) \bar{\psi}_{lL} \not{B} \psi_{lL} + \right. \\ &\left. + Y(\nu_R) \bar{\psi}_{\nu R} \not{B} \psi_{\nu R} + Y(l_R) \bar{\psi}_{lR} \not{B} \psi_{lR} \right] \end{aligned} \quad (3.2.34)$$

$$= \sum_{\substack{f = \nu_L, l_L \\ \nu_R, l_R}} \left[g I_3(f) \bar{\psi}_f W_3 \psi_f + g' \frac{Y(f)}{2} \bar{\psi}_f \not{B} \psi_f \right] \quad (3.2.35)$$

¹Oppure guardando gli autovalori degli spinori del doppietto L_L rispetto t_3 generatore di $SU(2)$:

$$\frac{1}{2} \begin{pmatrix} 1 & 0 \\ 0 & -1 \end{pmatrix} \begin{pmatrix} \psi_{\nu L} \\ \psi_{lL} \end{pmatrix}. \quad (3.2.32)$$

in cui abbiamo implicitamente definito:

$$Y(\nu_L) = Y(l_L) = Y(L_L). \quad (3.2.36)$$

Notiamo che i fermioni carichi l_L ed l_R si accoppiano sia con W_3^μ sia con il campo B^μ . Ha senso quindi, visto che i fermioni si accoppiano sicuramente con il fotone, pensare che il campo elettromagnetico A^μ sia una combinazione lineare di questi due campi. Introducendo i coefficienti $\cos \theta_W$ e $\sin \theta_W$ (che ci permettono di mantenere la stessa normalizzazione di W_3 e B) scriviamo²:

$$A^\mu = B^\mu \cos \theta_W + W_3^\mu \sin \theta_W \quad (3.2.37)$$

la combinazione lineare ortogonale:

$$Z^\mu = -B^\mu \sin \theta_W + W_3^\mu \cos \theta \quad (3.2.38)$$

e le relative inverse:

$$B^\mu = A^\mu \cos \theta_W - Z^\mu \sin \theta_W \quad (3.2.39)$$

$$W_3^\mu = A^\mu \sin \theta_W + Z^\mu \cos \theta_W. \quad (3.2.40)$$

Sostituendo le relazioni di B^μ e W_3^μ in $\mathcal{L}_N$ otteniamo:

$$\mathcal{L}_N = \mathcal{L}_N^{(A)} + \mathcal{L}_N^{(Z)} \quad (3.2.41)$$

in cui:

$$\mathcal{L}_N^{(A)} = \sum_{\substack{f = \nu_L, l_L \\ \nu_R, l_R}} \left[g \sin \theta_W I_3(f) + g' \cos \theta_W \frac{Y(f)}{2} \right] \bar{\psi}_f A \psi_f \quad (3.2.42)$$

$$\mathcal{L}_N^{(Z)} = \sum_{\substack{f = \nu_L, l_L \\ \nu_R, l_R}} \left[g \cos \theta_W I_3(f) - g' \sin \theta_W \frac{Y(f)}{2} \right] \bar{\psi}_f Z \psi_f. \quad (3.2.43)$$

Tutti questi discorsi sull'accoppiamento elettromagnetico sono molto belli, ma dobbiamo effettivamente richiedere che l'accoppiamento con il fotone sia

²Può essere utile da ricordare che tutto ciò si scrive:

$$\begin{pmatrix} A_\mu \\ Z_\mu \end{pmatrix} = \begin{pmatrix} \cos \theta_w & \sin \theta_w \\ -\sin \theta_w & \cos \theta_w \end{pmatrix} \begin{pmatrix} B_\mu \\ W_\mu^3 \end{pmatrix} \quad ; \quad \begin{pmatrix} B_\mu \\ W_\mu^3 \end{pmatrix} = \begin{pmatrix} \cos \theta_w & -\sin \theta_w \\ \sin \theta_w & \cos \theta_w \end{pmatrix} \begin{pmatrix} A_\mu \\ Z_\mu \end{pmatrix}.$$

quello della QED (ricorda (3.1.18)), il quale lo scriviamo come:

$$\mathcal{L}_N^{(A)} = -e\bar{\psi}_l A\psi_l \quad (3.2.44)$$

$$= -e\left[\bar{\psi}_{lL} A\psi_l + \bar{\psi}_{lR} A\psi_{lR}\right] \quad (3.2.45)$$

$$= e\left[Q(\nu_L)\bar{\psi}_{\nu L} A\psi_{\nu L} + Q(l_L)\bar{\psi}_{lL} A\psi_{lL} + Q(\nu_R)\bar{\psi}_{\nu R} A\psi_{\nu R} + \right. \\ \left. + Q(l_R)\bar{\psi}_{lR} A\psi_{lR}\right] \quad (3.2.46)$$

$$= \sum_{\substack{f = \nu_L, l_L \\ \nu_R, l_R}} eQ(f)\bar{\psi}_f A\psi_f \quad (3.2.47)$$

in cui abbiamo introdotto le cariche elettriche:

$$Q(\nu_L) = Q(\nu_R) = 0 \quad , \quad Q(l_L) = Q(l_R) = -1. \quad (3.2.48)$$

Otteniamo perciò la condizione per cui l'accoppiamento con il fotone riproduce la QED:

$$g \sin \theta_W I_3(f) + g' \cos \theta_W \frac{Y(f)}{2} = eQ(f) \quad (3.2.49)$$

con $f = \nu_L, l_L, \nu_R, l_R$. Se riscriviamo queste relazioni esplicitando i valori di I_3 e Q , abbiamo:

$$\nu_L : \quad \frac{g}{2} \sin \theta_W + \frac{g'}{2} \cos \theta_W Y(L_L) = 0 \quad (3.2.50)$$

$$l_L : \quad -\frac{g}{2} \sin \theta_W + \frac{g'}{2} \cos \theta_W Y(L_L) = -e \quad (3.2.51)$$

$$\nu_R : \quad \frac{g'}{2} \cos \theta_W Y(\nu_R) = 0 \quad (3.2.52)$$

$$l_R : \quad \frac{g'}{2} \cos \theta_W Y(l_R) = -e \quad (3.2.53)$$

da cui otteniamo:

$$(\nu_L) - (l_L) : \quad g \sin \theta_W = e \quad (3.2.54)$$

$$(\nu_L) + (l_L) : \quad g' \cos \theta_W Y(L_L) = -e \quad (3.2.55)$$

$$(\nu_R) : \quad g' \cos \theta_W Y(\nu_R) = 0 \quad (3.2.56)$$

$$(l_R) : \quad g' \cos \theta_W Y(l_R) = -2e. \quad (3.2.57)$$

Dunque si vede bene che le ipercariche del settore right $Y(\nu_R)$ e $Y(l_R)$ sono legate a $Y(L_L)$ dalle relazioni:

$$Y(\nu_R) = 0 \quad , \quad Y(l_R) = 2Y(L_L) \quad (3.2.58)$$

mentre per le costanti di accoppiamento vediamo che:

$$g \sin \theta_W = e \quad , \quad g' \cos \theta_W Y(L_L) = -e \quad (3.2.59)$$

ma visto che $Y(L_L)$ moltiplica sempre g' , allora senza perdita di generalità possiamo definire $Y(L_L) = -1$, a meno di una costante di normalizzazione che viene riassorbita in g' , così da vedere:

$$g \sin \theta_W = g' \cos \theta_W = e \quad (3.2.60)$$

oltre che fissare i valori delle ipercariche:

$$\begin{aligned} Y(L_L) &= Y(\nu_L) = Y(l_L) = -1 \\ Y(\nu_R) &= 0 \\ Y(l_R) &= -2. \end{aligned} \quad (3.2.61)$$

Possiamo fare alcune considerazioni. La relazione $g \sin \theta_W = e$ ci fa subito capire che la soluzione banale $B^\mu = A^\mu$, $g' = e$, che corrisponde a $\sin \theta_W = 0$, non funziona, perché richiederebbe una costante di accoppiamento debole g infinita. Il motivo di questo mixing tra B^μ e W_3^μ per ottenere il campo elettromagnetico A^μ risiede nel differente comportamento dell'interazione debole e di quella elettromagnetica per trasformazioni di parità (la prima la viola massimamente, la seconda la conserva).

La relazione $Y(\nu_R) = 0$, insieme a $I_3(\nu_R) = 0$ e $Q(\nu_R) = 0$, rivela che il neutrino right non si accoppia con niente, né debolmente, né elettromagneticamente. Si dice quindi che ν_R è disaccoppiato, cioè esiste solo come particella libera. Non interagendo con niente, i neutrini right non sono quindi rilevabili e praticamente è come se non esistessero.

Se inseriamo la relazione (3.2.60) all'interno della condizione per riprodurre la QED (3.2.49) otteniamo la *relazione tra isospin, carica e ipercarica*:

$$I_3(f) + \frac{Y(f)}{2} = Q(f) \quad (3.2.62)$$

che è la **relazione di Gell-Mann-Nishijima**.

Utilizzando quanto trovato nella lagrangiana (3.2.43) si ha:

$$\mathcal{L}_N^{(Z)} = \sum_{\substack{f = \nu_L, l_L \\ \nu_R, l_R}} \left[e \frac{\cos \theta_W}{\sin \theta_W} I_3(f) - e \frac{\sin \theta_W}{\cos \theta_W} \frac{Y(f)}{2} \right] \bar{\psi}_f \not{Z} \psi_f \quad (3.2.63)$$

$$= e \frac{\cos \theta_W}{2 \sin \theta_W} \left[\bar{\psi}_{\nu L} \not{Z} \psi_{\nu L} - \bar{\psi}_{l L} \not{Z} \psi_{l L} \right] + \\ + e \frac{\sin \theta_W}{2 \cos \theta_W} \left[\bar{\psi}_{\nu L} \not{Z} \psi_{\nu L} + \bar{\psi}_{l L} \not{Z} \psi_{l L} + 2 \bar{\psi}_{l R} \not{Z} \psi_{l R} \right] \quad (3.2.64)$$

$$= \frac{e}{2 \sin \theta_W \cos \theta_W} \left[(\sin^2 \theta_W + \cos^2 \theta_W) \bar{\psi}_{\nu L} \not{Z} \psi_{\nu L} + \right. \\ \left. + (\sin^2 \theta_W - \cos^2 \theta_W) \bar{\psi}_{l L} \not{Z} \psi_{l L} + 2 \sin^2 \theta_W \bar{\psi}_{l R} \not{Z} \psi_{l R} \right] \quad (3.2.65)$$

$$= \frac{e}{2 \sin \theta_W \cos \theta_W} \left[\bar{\psi}_{\nu L} \not{Z} \psi_{\nu L} + (2 \sin^2 \theta_W - 1) \bar{\psi}_{l L} \not{Z} \psi_{l L} + \right. \\ \left. + 2 \sin^2 \theta_W \bar{\psi}_{l R} \not{Z} \psi_{l R} \right] \quad (3.2.66)$$

possiamo anche riscrivere:

$$\bar{\psi}_{\nu L} \not{Z} \psi_{\nu L} = \frac{1}{2} \bar{\psi}_{\nu} \not{Z} (1 - \gamma_5) \psi_{\nu} \quad (3.2.67)$$

$$\bar{\psi}_{l L} \not{Z} \psi_{l L} = \frac{1}{2} \bar{\psi}_l \not{Z} (1 - \gamma_5) \psi_l \quad (3.2.68)$$

$$\bar{\psi}_{l R} \not{Z} \psi_{l R} = \frac{1}{2} \bar{\psi}_l \not{Z} (1 + \gamma_5) \psi_l \quad (3.2.69)$$

così otteniamo:

$$\mathcal{L}_N^{(Z)} = \frac{e}{4 \sin \theta_W \cos \theta_W} \left[\bar{\psi}_{\nu} \not{Z} (1 - \gamma_5) \psi_{\nu} + \bar{\psi}_l \not{Z} (4 \sin^2 \theta_W - 1 + \gamma_5) \psi_l \right]. \quad (3.2.70)$$

Se ora ci ricordiamo del fatto che i leptoni sono organizzati in tre famiglie, e ipotizziamo assenza di mixing tra famiglie, allora possiamo scrivere la lagrangiana per il settore leptonico come:

$$\mathcal{L}^{lep} = \mathcal{L}_{Dirac}^{lep} + \mathcal{L}_C^{lep} + \mathcal{L}_N^{lep} \quad (3.2.71)$$

in cui abbiamo:

$$\mathcal{L}_{Dirac}^{lep} = \sum_{l=e,\mu,\tau} \left[i\bar{\psi}_{\nu_l} \not{\partial} \psi_{\nu_l} + i\bar{\psi}_l \not{\partial} \psi_l \right] \quad (3.2.72)$$

$$\mathcal{L}_C^{lep} = \frac{g}{\sqrt{2}} \sum_{l=e,\mu,\tau} \left[\bar{\psi}_{\nu_L} W^+ \psi_{l_L} + \bar{\psi}_{l_L} W^- \psi_{\nu_L} \right] \quad (3.2.73)$$

$$= \frac{g}{2\sqrt{2}} \sum_{l=e,\mu,\tau} \left[\bar{\psi}_{\nu} W^+ (1 - \gamma_5) \psi_l + \bar{\psi}_l W^- (1 - \gamma_5) \psi_{\nu} \right] \quad (3.2.74)$$

$$\mathcal{L}_N^{lep} = \sum_{l=e,\mu,\tau} \left\{ -e\bar{\psi}_l A \psi_l + \frac{g}{4\cos\theta_W} \left[\bar{\psi}_{\nu} Z (1 - \gamma_5) \psi_{\nu} - \bar{\psi}_l Z (1 - \gamma_5 - 4\sin^2\theta_W) \psi_l \right] \right\} \quad (3.2.75)$$

con:

$$g = \frac{e}{\sin\theta_W}. \quad (3.2.76)$$

Possiamo notare alcune proprietà di questa lagrangiana:

- Rispetta l'accoppiamento V-A per le correnti cariche.
- Ha solo costanti di accoppiamento adimensionali (ci sono speranze che unitarietà e rinormalizzabilità siano soddisfatte).
- Riproduce l'accoppiamento di QED.
- È basata sulla simmetria $SU(2)_L \times U(1)_Y$ che controlla gli accoppiamenti tra leptoni e bosoni.
- La costante d'accoppiamento debole (3.2.76) è vincolata all'accoppiamento elettromagnetico dall'angolo di Weinberg θ_W , che sperimentalmente vale $\sin\theta_W \sim 0.23$.

Si noti che la costante d'accoppiamento g risulta essere dello stesso ordine di grandezza di e (ad essere precisi, guardando solo (3.2.76), è addirittura $g > e$). Quindi, se dipendesse solo dalla costante d'accoppiamento g , l'interazione debole dovrebbe essere "forte" più o meno come quella elettromagnetica. Alle basse energie risulta invece essere molto più debole. Vedremo più avanti che questo è dovuto ad un altro motivo.

- Sappiamo dagli esperimenti che A^μ conserva la parità e W^μ la viola completamente (infatti nei termini di W^μ abbiamo solo termini con $(1 - \gamma_5)$, cioè parti left); però, dalla lagrangiana leptonica neutra, vediamo che Z^μ è un misto delle due situazioni, poiché contenendo parte dei

neutrini viola la parità, ma con i leptoni, essendo presenti in entrambe le chiralità (cosa che vediamo nei passaggi intermedi per arrivare al risultato finale), non la viola. Dunque, Z^μ viola la parità, ma non completamente.

3.3 Interazione elettrodebole per correnti adroniche

La corrente adronica, dopo aver applicato il meccanismo GIM, abbiamo visto essere:

$$H^\mu = \bar{\psi}_u \gamma^\mu (1 - \gamma_5) \left[\cos \theta_c \psi_d + \sin \theta_c \psi_s \right] + \bar{\psi}_c \gamma^\mu (1 - \gamma_5) \left[-\sin \theta_c \psi_d + \cos \theta_c \psi_s \right] \quad (3.3.1)$$

che con i doppietti di quark (2.4.29) e la matrice di Cabibbo (2.4.30) diventa:

$$H^\mu = \sum_{i=1}^2 O_{ij}^{(c)} \bar{Q}_i \gamma^\mu (1 - \gamma_5) \tau^+ Q_j \quad (3.3.2)$$

la quale è simile alla corrente leptonica per quanto riguarda i doppietti e la presenza della matrice τ^+ di $SU(2)$. Però, a differenza della corrente leptonica, qui abbiamo già una somma sulle famiglie, a colpa del mixing. Infatti, nel caso leptonico non avevamo la matrice $O^{(c)}$, che ci avrebbe fornito il mixing tra le famiglie.

Nota. Per la corrente leptonica non avevamo né la somma sulle famiglie, né il mixing; però, se riuscissimo per la parte adronica ad ottenere un risultato analogo a quello dei leptoni, allora vorrebbe dire che noi potremmo riuscire a scrivere anche un mixing tra i leptoni. Teniamo a mente questa nota e ripensiamoci tra un po' di pagine.

Visto che in (3.3.2) c'è solo il contributo delle correnti cariche, e manca ancora il pezzo per le correnti neutre per completare $SU(2)$, per ripetere quanto fatto per le correnti leptoniche conviene riscrivere i doppietti con il mixing dei quark già all'interno:

$$Q'_1 = \begin{pmatrix} \psi_u \\ \psi_{d'} \end{pmatrix}, \quad Q'_2 = \begin{pmatrix} \psi_c \\ \psi_{s'} \end{pmatrix} \quad (3.3.3)$$

in cui mettiamo:

$$\psi_{d'} = \cos \theta_c \psi_d + \sin \theta_c \psi_s \quad (3.3.4)$$

$$\psi_{s'} = -\sin \theta_c \psi_d + \cos \theta_c \psi_s \quad (3.3.5)$$

cioé facciamo una rotazione di $SU(2)$:

$$\begin{pmatrix} \psi_{d'} \\ \psi_{s'} \end{pmatrix} = V \begin{pmatrix} \psi_d \\ \psi_s \end{pmatrix} \quad , \quad V = O^{(c)} = \begin{pmatrix} \cos \theta_c & \sin \theta_c \\ -\sin \theta_c & \cos \theta_c \end{pmatrix}. \quad (3.3.6)$$

Ridefinendo i doppietti con il mixing la corrente adronica si scrive:

$$H^\mu = \sum_{i=1}^2 \bar{Q}'_i \gamma^\mu (1 - \gamma_5) \tau^+ Q'_i \quad (3.3.7)$$

che è una riscrittura che si può facilmente estendere ad una terza famiglia, come sarà necessario, poiché basta fare la somma da 1 a 3 e ridefinire la matrice V in modo che leghi i quark down delle varie famiglie.

Il complesso coniugato della corrente (3.3.7) risulta:

$$H^\dagger_\mu = \sum_{i=1}^2 \bar{Q}'_i \gamma^\mu (1 - \gamma_5) \tau^- Q'_i. \quad (3.3.8)$$

Vorremmo completare la simmetria $SU(2)$, anche per gli adroni, aggiungendo la corrente debole neutra in analogia al settore leptonic. Per farlo dobbiamo usare gli stessi doppietti usati nelle correnti cariche e definire:

$$H^\mu_3 = \sum_{i=1}^2 \bar{Q}'_i \gamma^\mu (1 - \gamma_5) t_3 Q'_i \quad (3.3.9)$$

in cui, utilizzando i doppietti Q'_i , potremmo aspettarci che H^μ_3 contenga correnti neutre con cambio di flavour, le cosiddette FCNC, e ciò sarebbe in disaccordo con i dati sperimentali. Vediamo anche che, essendoci la matrice t_3 che è diagonale, la corrente neutra lega particelle con la stessa carica. Però, grazie all'unitarietà della matrice V , si può dimostrare che la corrente neutra non genera correnti con cambio di flavour. Per vederlo in maniera compatta conviene introdurre la notazione:

$$u_1, u_2, u_3 = u, c, t \quad (3.3.10)$$

$$d_1, d_2, d_3 = d, s, b \quad (3.3.11)$$

in modo che i doppietti si riscrivano come:

$$Q_i = \begin{pmatrix} \psi_{u_i} \\ \psi_{d_i} \end{pmatrix} \quad , \quad Q'_i = \begin{pmatrix} \psi_{u_i} \\ \psi_{d'_i} \end{pmatrix} \quad (3.3.12)$$

$$\psi_{d'_i} = \sum_j V_{ij} \psi_{d_j} \quad , \quad \bar{\psi}_{d'_i} = \sum_k V_{ik}^* \bar{\psi}_{d_k} = \sum_k V_{ki}^\dagger \bar{\psi}_{d_k}. \quad (3.3.13)$$

Facendo ciò possiamo riscrivere la corrente adronica neutra come:

$$H_3^\mu = \sum_i \bar{Q}'_i \gamma^\mu (1 - \gamma_5) t_3 Q'_i \quad (3.3.14)$$

$$= \sum_i \frac{1}{2} \left[\bar{\psi}_{u_i} \gamma^\mu (1 - \gamma_5) \psi_{u_i} - \bar{\psi}_{d'_i} \gamma^\mu (1 - \gamma_5) \psi_{d'_i} \right] \quad (3.3.15)$$

$$= \sum_i \frac{1}{2} \bar{\psi}_{u_i} \gamma^\mu (1 - \gamma_5) \psi_{u_i} - \sum_{j,k} \frac{1}{2} \underbrace{\sum_i V_{ki}^\dagger V_{ij}}_{=(V^\dagger V)=\delta_{kj}} \bar{\psi}_{d_k} \gamma^\mu (1 - \gamma_5) \psi_{d_j} \quad (3.3.16)$$

$$= \sum_i \frac{1}{2} \bar{\psi}_{u_i} \gamma^\mu (1 - \gamma_5) \psi_{u_i} - \sum_j \frac{1}{2} \bar{\psi}_{d_j} \gamma^\mu (1 - \gamma_5) \psi_{d_j} \quad (3.3.17)$$

$$= \sum_i \bar{Q}_i \gamma^\mu (1 - \gamma_5) t_3 Q_i \quad (3.3.18)$$

in cui vediamo chiaramente che la corrente neutra si può scrivere equivalentemente con i doppietti mescolati Q'_i o con i doppietti Q_i , rivelando che non c'è mixing tra le famiglie nel settore neutro neanche per gli adroni.

Come già sappiamo, il primo passo è riscrivere la lagrangiana massless di Dirac:

$$\mathcal{L}_{Dirac}^{had} = \sum_i \left[i \bar{\psi}_{u_i} \not{\partial} \psi_{u_i} + i \bar{\psi}_{d_i} \not{\partial} \psi_{d_i} \right] \quad (3.3.19)$$

$$= \sum_i \left[i \bar{\psi}_{u_i L} \not{\partial} \psi_{u_i L} + i \bar{\psi}_{d_i L} \not{\partial} \psi_{d_i L} + i \bar{\psi}_{u_i R} \not{\partial} \psi_{u_i R} + i \bar{\psi}_{d_i R} \not{\partial} \psi_{d_i R} \right] \quad (3.3.20)$$

$$= \sum_i \left[i \bar{Q}_{iL} \not{\partial} Q_{iL} + i \bar{\psi}_{u_i R} \not{\partial} \psi_{u_i R} + i \bar{\psi}_{d_i R} \not{\partial} \psi_{d_i R} \right] \quad (3.3.21)$$

$$= \sum_i \left[i \bar{Q}'_{iL} \not{\partial} Q'_{iL} + i \bar{\psi}_{u_i R} \not{\partial} \psi_{u_i R} + i \bar{\psi}_{d'_i R} \not{\partial} \psi_{d'_i R} \right] \quad (3.3.22)$$

in cui nell'ultimo passaggio abbiamo utilizzato quanto visto per la corrente adronica neutra (3.3.18), che può essere scritta indifferentemente con i ψ_{d_i} o $\psi_{d'_i}$.

Ora, possiamo introdurre l'accoppiamento $SU(2)_L \times U(1)_Y$ definendo le derivate covarianti:

$$\mathcal{D}^\mu Q'_{iL} = \partial^\mu Q'_{iL} - ig W_a^\mu t_a Q'_{iL} - ig' B^\mu \frac{Y(Q_L)}{2} Q'_{iL} \quad (3.3.23)$$

$$\mathcal{D}^\mu \psi_{u_i R} = \partial^\mu \psi_{u_i R} - ig' B^\mu \frac{Y(u_R)}{2} \psi_{u_i R} \quad (3.3.24)$$

$$\mathcal{D}^\mu \psi_{d'_i R} = \partial^\mu \psi_{d'_i R} - ig' B^\mu \frac{Y(d_R)}{2} \psi_{d'_i R}. \quad (3.3.25)$$

Notiamo che l'aver introdotto le derivate covarianti per i campi $\psi_{d'_i}$ (invece che per i campi ψ_{d_i}) è essenziale per avere il corretto accoppiamento delle correnti cariche. L'averle considerate anche per le correnti neutre è solo per uniformità, visto che per le correnti neutre è indifferente utilizzare i campi $\psi_{d'_i}$ o ψ_{d_i} . A questo punto possiamo sostituire le derivate ordinarie con le covarianti e individuare i tre termini:

$$\mathcal{L}^{had} = \sum_i \left[i\bar{Q}'_{iL} \not{D} Q'_{iL} + i\bar{\psi}_{u_iR} \not{D} \psi_{u_iR} + i\bar{\psi}_{d'_iR} \not{D} \psi_{d'_iR} \right] \quad (3.3.26)$$

$$= \mathcal{L}_{Dirac}^{had} + \mathcal{L}_C^{had} + \mathcal{L}_N^{had} \quad (3.3.27)$$

dove la lagrangiana cinetica è:

$$\mathcal{L}_{Dirac}^{had} = \sum_i \left[i\bar{\psi}_{u_i} \not{D} \psi_{u_i} + i\bar{\psi}_{d_i} \not{D} \psi_{d_i} \right] \quad (3.3.28)$$

la parte carica risulta:

$$\mathcal{L}_C^{had} = \sum_i g \left[\bar{Q}'_{iL} W_1 t_1 Q'_{iL} + \bar{Q}'_{iL} W_2 t_2 Q'_{iL} \right] \quad (3.3.29)$$

$$= \sum_i \frac{g}{\sqrt{2}} \left[\bar{Q}'_{iL} W^+ \tau^+ Q'_{iL} + \bar{Q}'_{iL} W^- \tau^- Q'_{iL} \right] \quad (3.3.30)$$

$$= \sum_i \frac{g}{\sqrt{2}} \left[\bar{\psi}_{u_iL} W^+ \psi_{d'_iL} + \bar{\psi}_{d'_iL} W^- \psi_{u_iL} \right] \quad (3.3.31)$$

$$= \sum_{i,j} \frac{g}{\sqrt{2}} \left[V_{ij} \bar{\psi}_{u_iL} W^+ \psi_{d_jL} + V_{ij}^* \bar{\psi}_{d_jL} W^- \psi_{u_iL} \right] \quad (3.3.32)$$

$$= \sum_{i,j} \frac{g}{2\sqrt{2}} \left[V_{ij} \bar{\psi}_{u_i} W^+ (1 - \gamma_5) \psi_{d_j} + V_{ij}^* \bar{\psi}_{d_j} W^- (1 - \gamma_5) \psi_{u_i} \right] \quad (3.3.33)$$

mentre per la corrente neutra:

$$\begin{aligned} \mathcal{L}_N^{had} = \sum_i \left\{ g \bar{Q}'_{iL} W_3 t_3 Q'_{iL} + \right. \\ \left. + \frac{g'}{2} \left[Y(Q_L) \bar{Q}'_{iL} \not{B} Q'_{iL} + Y(u_R) \bar{\psi}_{u_i R} \not{B} \psi_{u_i R} + Y(d_R) \bar{\psi}_{d_i R} \not{B} \psi_{d_i R} \right] \right\} \end{aligned} \quad (3.3.34)$$

$$\begin{aligned} = \sum_i \left\{ g \bar{Q}_{iL} W_3 t_3 Q_{iL} + \right. \\ \left. + \frac{g'}{2} \left[Y(Q_L) \bar{Q}_{iL} \not{B} Q_{iL} + Y(u_R) \bar{\psi}_{u_i R} \not{B} \psi_{u_i R} + Y(d_R) \bar{\psi}_{d_i R} \not{B} \psi_{d_i R} \right] \right\} \end{aligned} \quad (3.3.35)$$

$$\begin{aligned} = \sum_i \left\{ \frac{g}{2} \left[\bar{\psi}_{u_i L} W_3 \psi_{u_i L} - \bar{\psi}_{d_i L} W_3 \psi_{d_i L} \right] + \right. \\ \left. + \frac{g'}{2} \left[Y(u_L) \bar{\psi}_{u_i L} \not{B} \psi_{u_i L} + Y(d_L) \bar{\psi}_{d_i L} \not{B} \psi_{d_i L} + \right. \right. \\ \left. \left. + Y(u_R) \bar{\psi}_{u_i R} \not{B} \psi_{u_i R} + Y(d_R) \bar{\psi}_{d_i R} \not{B} \psi_{d_i R} \right] \right\} \end{aligned} \quad (3.3.36)$$

$$= \sum_i \sum_{\substack{f = u_{iL}, d_{iL} \\ u_{iR}, d_{iR}}} \left[g I_3(f) \bar{\psi}_f W_3 \psi_f + g' \frac{Y(f)}{2} \bar{\psi}_f \not{B} \psi_f \right] \quad (3.3.37)$$

in cui possiamo, avendo la stessa struttura della corrispondente leptonica, anche introdurre i campi A^μ e Z^μ , tali da soddisfare (3.2.39) e (3.2.40).

Anche in questo caso possiamo richiedere che l'accoppiamento con A^μ sia quello della QED:

$$\mathcal{L}_{QED}^{had} = \sum_i \left[e Q(u) \bar{\psi}_{u_i} \not{A} \psi_{u_i} + e Q(d) \bar{\psi}_{d_i} \not{A} \psi_{d_i} \right]$$

e seguendo lo stesso procedimento delle correnti leptoniche riotteniamo la relazione tra le costanti di accoppiamento (3.2.60) e la relazione tra isospin, carica e ipercarica (3.2.62) per ogni quark f . Notiamo che in questo caso, avendo:

$$Q(u_L) = Q(u_R) = \frac{2}{3} \quad , \quad Q(d_L) = Q(d_R) = -\frac{1}{3} \quad (3.3.38)$$

$$I_3(u_L) = \frac{1}{2} \quad , \quad I_3(d_L) = -\frac{1}{2} \quad , \quad I_3(u_R) = I_3(d_R) = 0 \quad (3.3.39)$$

otteniamo per l'ipercarica debole:

$$Y(u_L) = Y(d_L) = \frac{1}{3} \quad , \quad Y(u_R) = \frac{4}{3} \quad , \quad Y(d_R) = -\frac{2}{3}. \quad (3.3.40)$$

Vediamo che ora, a differenza del neutrino right ν_R che si disaccoppia per via dei suoi numeri quantici $Y(\nu_R) = I_3(\nu_R) = Q(\nu_R) = 0$, i quark up right non

si disaccoppiano. Infatti, pur non interagendo debolmente, avendo $I_3(u_3) = 0$ e non accoppiandosi con i W^μ , essi interagiscono elettromagneticamente essendo dotati di carica $Q(u_R) = 2/3$.

Dunque, complessivamente otteniamo la corrente neutra:

$$\begin{aligned} \mathcal{L}_N^{had} = & \sum_i \left\{ eQ(u) \bar{\psi}_{u_i} \not{A} \psi_{u_i} + eQ(d) \bar{\psi}_{d_i} \not{A} \psi_{d_i} + \right. \\ & \left. + \sum_{\substack{f = u_{iL}, d_{iL} \\ u_{iR}, d_{iR}}} \left[g \cos \theta_W I_3(f) - g' \sin \theta_W \frac{Y(f)}{2} \right] \bar{\psi}_f \not{Z} \psi_f \right\} \end{aligned} \quad (3.3.41)$$

$$\begin{aligned} = & \sum_i \left\{ \frac{2}{3} e \bar{\psi}_{u_i} \not{A} \psi_{u_i} - \frac{1}{3} e \bar{\psi}_{d_i} \not{A} \psi_{d_i} + \right. \\ & + \frac{g}{4 \cos \theta_W} \left[\bar{\psi}_{u_i} \not{Z} \left(1 - \gamma_5 - \frac{8}{3} \sin^2 \theta_W \right) \psi_{u_i} - \right. \\ & \left. \left. - \bar{\psi}_{d_i} \not{Z} \left(1 - \gamma_5 - \frac{4}{3} \sin^2 \theta_W \right) \psi_{d_i} \right] \right\} \end{aligned} \quad (3.3.42)$$

che, come ci aspettavamo, non contiene mixing tra le famiglie.

3.4 Invarianza di $\mathcal{L}^{lep} + \mathcal{L}^{had}$ per trasformazioni di $SU(2)_L \times U(1)_Y$

In questa sezione vogliamo dimostrare l'invarianza di gauge della lagrangiana:

$$\begin{aligned} \mathcal{L}^{lep} + \mathcal{L}^{had} = & \sum_{i=e,\mu,\tau} \left[i \bar{L}_{iL} \not{D} L_{iL} + i \bar{\psi}_{\nu_{iR}} \not{D} \psi_{\nu_{iR}} + i \bar{\psi}_{l_{iR}} \not{D} \psi_{l_{iR}} \right] + \\ & + \sum_{i=1,2,3} \left[i \bar{Q}'_{iL} \not{D} Q'_{iL} + i \bar{\psi}_{u_{iR}} \not{D} \psi_{u_{iR}} + i \bar{\psi}_{d'_{iR}} \not{D} \psi_{d'_{iR}} \right] \end{aligned} \quad (3.4.1)$$

che contiene l'interazione tra i campi di gauge e i campi fermionici. Ci interessa dimostrare questa invarianza poiché ci servirà in seguito quando parleremo di masse.

Considerando trasformazioni infinitesime³ di $SU(2)_L \times U(1)_Y$ abbiamo:

$$W_i^\mu \longrightarrow W_i^\mu - \partial_\mu \Lambda_i(x) + ig\Lambda_a(x)W_j^\mu(T_a)_{ji} \quad (3.4.3)$$

$$= W_i^\mu - \partial_\mu \Lambda_i(x) - g\epsilon_{jai}W_j^\mu \Lambda_a(x) \quad (3.4.4)$$

$$B^\mu \longrightarrow B^\mu - \partial_\mu \lambda(x) \quad (3.4.5)$$

$$L_{iL} \longrightarrow \left[1 - ig\Lambda_i(x)t_i - ig'\lambda(x)\frac{Y(L_L)}{2} \right] L_{iL} \quad (3.4.6)$$

$$\psi_{\nu_i R} \longrightarrow \left[1 - ig'\lambda(x)\frac{Y(\nu_R)}{2} \right] \psi_{\nu_i R} = \psi_{\nu R} \quad , \quad Y(\nu_R) = 0 \quad (3.4.7)$$

$$\psi_{l_i R} \longrightarrow \left[1 - ig'\lambda(x)\frac{Y(l_R)}{2} \right] \psi_{l_i R} \quad (3.4.8)$$

$$Q'_{iL} \longrightarrow \left[1 - ig\Lambda_k(x)t_k - ig'\lambda(x)\frac{Y(Q_L)}{2} \right] Q'_{iL} \quad (3.4.9)$$

$$\psi_{u_i R} \longrightarrow \left[1 - ig'\lambda(x)\frac{Y(u_R)}{2} \right] \psi_{u_i R} \quad (3.4.10)$$

$$\psi_{d'_i R} \longrightarrow \left[1 - ig'\lambda(x)\frac{Y(d_R)}{2} \right] \psi_{d'_i R}. \quad (3.4.11)$$

Per i vari termini di (3.4.1) si procede in modo analogo, ma noi vediamo il termine $\bar{L}_{iL}\mathcal{D}^\mu L_{iL}$ poiché il più complicato. Partiamo dalla trasformazione della derivata covariante applicata al doppietto:

$$\mathcal{D}^\mu L_{iL} = \partial^\mu L_{iL} - igW_i^\mu t_i L_{iL} - ig'B^\mu \frac{Y(L_L)}{2} L_{iL} \quad (3.4.12)$$

$$\begin{aligned} &\longrightarrow \left[1 - ig\Lambda_i(x)t_i - ig'\lambda(x)\frac{Y(L_L)}{2} \right] \partial^\mu L_{iL} - \\ &\quad - \left[ig\partial^\mu \Lambda_i(x)t_i + ig'\partial^\mu \lambda(x)\frac{Y(L_L)}{2} \right] L_{iL} - \\ &\quad - ig \left[W_i^\mu - \partial_\mu \Lambda_i(x) - g\epsilon_{jai}W_j^\mu \Lambda_a(x) \right] t_i \left[1 - ig\Lambda_j(x)t_j - ig'\lambda(x)\frac{Y(L_L)}{2} \right] L_{iL} - \\ &\quad - ig' \left(B^\mu - \partial^\mu \lambda(x) \right) \frac{Y(L_L)}{2} \left[1 - ig\Lambda_i(x)t_i - ig'\lambda(x)\frac{Y(L_L)}{2} \right] L_{iL} \end{aligned} \quad (3.4.13)$$

limitandoci alle trasformazioni infinitesime e non considerando le potenze

³Per trasformazioni finite avremmo ad esempio:

$$e^{-ig\Lambda_a(x)T_a} W^\mu. \quad (3.4.2)$$

superiori ad uno in Λ_i e λ :

$$\begin{aligned}
 \mathcal{D}^\mu L_{iL} \longrightarrow & \left[1 - ig\Lambda_i(x)t_i - ig'\lambda(x)\frac{Y(L_L)}{2} \right] \partial^\mu L_{iL} - \\
 & - \left[ig\partial^\mu \Lambda_i(x)t_i + ig'\partial^\mu \lambda(x)\frac{Y(L_L)}{2} \right] L_{iL} - \\
 & - igW_i^\mu t_i \left[1 - ig\Lambda_j(x)t_j - ig'\lambda(x)\frac{Y(L_L)}{2} \right] L_{iL} + \\
 & + ig \left[\partial_\mu \Lambda_i(x) + g\epsilon_{jai}W_j^\mu \Lambda_a(x) \right] t_i L_{iL} - \\
 & - igt_i \left[W_i^\mu - igW_i^\mu \Lambda_j(x)t_j - ig'W_i^\mu \lambda(x)\frac{Y(L_L)}{2} - \right. \\
 & \left. - \partial_\mu \Lambda_i(x) - g\epsilon_{jai}W_j^\mu \Lambda_a(x) \right] L_{iL} - \\
 & - ig'B^\mu \frac{Y(L_L)}{2} \left[1 - ig\Lambda_i(x)t_i - ig'\lambda(x)\frac{Y(L_L)}{2} \right] L_{iL} + \\
 & + ig'\partial_\mu \lambda(x)\frac{Y(L_L)}{2} L_{iL}. \tag{3.4.14}
 \end{aligned}$$

in cui le derivate dei campi Λ_i e λ si cancellano e otteniamo:

$$\begin{aligned}
 \mathcal{D}^\mu L_{iL} \longrightarrow & \left[1 - ig\Lambda_i(x)t_i - ig'\lambda(x)\frac{Y(L_L)}{2} \right] \partial^\mu L_{iL} - \\
 & - igW_i^\mu t_i \left[1 - ig\Lambda_j(x)t_j - ig'\lambda(x)\frac{Y(L_L)}{2} \right] L_{iL} + \\
 & + ig^2 W_i^\mu \Lambda_j(x) \epsilon_{ijk} t_k L_{iL} - \\
 & - ig'B^\mu \frac{Y(L_L)}{2} \left[1 - ig\Lambda_i(x)t_i - ig'\lambda(x)\frac{Y(L_L)}{2} \right] L_{iL}. \tag{3.4.15}
 \end{aligned}$$

Vediamo che si può raggruppare davanti a tutto il termine:

$$\left[1 - ig\Lambda_j(x)t_j - ig'\lambda(x)\frac{Y(L_L)}{2} \right]$$

a parte per il termine con ϵ_{ijk} e per l'ordine delle matrici $t_i t_j$ nel secondo termine:

$$\begin{aligned}
 \mathcal{D}^\mu L_{iL} \longrightarrow & \left[1 - ig\Lambda_j(x)t_j - ig'\lambda(x)\frac{Y(L_L)}{2} \right] \times \\
 & \times \left[\partial^\mu L_{iL} - igW_i^\mu t_i L_{iL} - ig'B^\mu \frac{Y(L_L)}{2} L_{iL} \right] - \\
 & - g^2 W_i^\mu \Lambda_j(x) \left[[t_i, t_j] - i\epsilon_{ijk} t_k \right] L_{iL} \tag{3.4.16}
 \end{aligned}$$

in cui: il terzo termine si cancella per via delle regole di commutazione dei generatori del gruppo $SU(2)$, e riusciamo ad individuare la scrittura della derivata covariante del doppietto (3.4.12). Per cui otteniamo:

$$\mathcal{D}^\mu L_{iL} \longrightarrow \left[1 - ig\Lambda_i(x)t_i - ig'\lambda(x)\frac{Y(L_L)}{2} \right] \mathcal{D}^\mu L_{iL} \quad (3.4.17)$$

cioé che $\mathcal{D}^\mu L_{iL}$ trasforma come L_{iL} . Pertanto il termine $\bar{L}_{iL}\mathcal{D}^\mu L_{iL}$ rende $\mathcal{L}^{lep}$ invariante per trasformazioni di gauge di $SU(2)_L \times U(1)_Y$.

In modo analogo, per tutti gli altri termini, si dimostra che tutta la lagrangiana (3.4.1) è invariante.

3.5 Lagrangiana massless per i campi di gauge

La lagrangiana per i campi di gauge, sia abeliani che non, è una generica lagrangiana di Yang-Mills:

$$\mathcal{L}_{YM} = -\frac{1}{4}B_{\mu\nu}B^{\mu\nu} - \frac{1}{4}W_{\mu\nu}^a W_a^{\mu\nu} \quad (3.5.1)$$

che mantiene l'invarianza di gauge della teoria e ha i tensori:

$$B_{\mu\nu} = \partial_\mu B_\nu - \partial_\nu B_\mu \quad (3.5.2)$$

$$W_{\mu\nu}^a = \partial_\mu W_\nu^a - \partial_\nu W_\mu^a + g\epsilon_{bc}^a W_\mu^b W_\nu^c \quad (3.5.3)$$

con le componenti di $SU(2)$ $a = 1, 2, 3$. Questa lagrangiana genera accoppiamenti a 3 e 4 campi di gauge W_μ^a per $SU(2)$, ma non per il campo B_μ del $U(1)$ abeliano.

Ovviamente a noi interessa scrivere una lagrangiana contenente l'interazione fra i campi *fisici*⁴, ovvero in termini dei tensori:

$$F_{\mu\nu} = \partial_\mu A_\nu - \partial_\nu A_\mu \quad (3.5.4)$$

$$Z_{\mu\nu} = \partial_\mu Z_\nu - \partial_\nu Z_\mu \quad (3.5.5)$$

$$W_{\mu\nu}^\pm = \partial_\mu W_\nu^\pm - \partial_\nu W_\mu^\pm. \quad (3.5.6)$$

La lagrangiana fisica cercata la si può ottenere con le seguenti sostituzioni:

$$W_\mu^1 = \frac{W_\mu^+ + W_\mu^-}{\sqrt{2}} \quad (3.5.7)$$

$$W_\mu^2 = i\frac{W_\mu^+ - W_\mu^-}{\sqrt{2}} \quad (3.5.8)$$

$$W_\mu^a = \sin\theta_W A_\mu + \cos\theta_W Z_\mu \quad (3.5.9)$$

$$B_\mu = \cos\theta_W A_\mu - \sin\theta_W Z_\mu. \quad (3.5.10)$$

⁴I campi fisici sono W_μ^+ , W_μ^- , A_μ , Z_μ .

È un conto lungo e laborioso, non lo abbiamo visto in aula⁵, ma il risultato a cui si arriva è:

$$\mathcal{L}_{YM} = \mathcal{L}_{YM}^{(2)} + \mathcal{L}_{YM}^{(3)} + \mathcal{L}_{YM}^{(4)} \quad (3.5.11)$$

in cui vediamo la *parte libera*:

$$\mathcal{L}_{YM}^{(2)} = -\frac{1}{4}F_{\mu\nu}F^{\mu\nu} - \frac{1}{4}Z_{\mu\nu}Z^{\mu\nu} - \frac{1}{2}W_{\mu\nu}^+W_{\mu\nu}^- \quad (3.5.12)$$

dove individuiamo i contributi di QED, di quella che a volte si può chiamare Z-QED e il pezzo per i $W^\pm$, che ha una diversa normalizzazione poiché sono dei campi complessi. Abbiamo la parte *a tre campi*:

$$\begin{aligned} \mathcal{L}_{YM}^{(3)} = & ig \sin \theta_W \left(W_{\mu\nu}^+ W_{\mu\nu}^- A^\nu - W_{\mu\nu}^- W_{\mu\nu}^+ A^\nu + F_{\mu\nu} W_{\mu\nu}^+ W_{\mu\nu}^- \right) + \\ & + ig \cos \theta_W \left(W_{\mu\nu}^+ W_{\mu\nu}^- Z^\nu - W_{\mu\nu}^- W_{\mu\nu}^+ Z^\nu + Z_{\mu\nu} W_{\mu\nu}^+ W_{\mu\nu}^- \right) \quad (3.5.13) \\ \sim & \sin \theta_w \left(\begin{array}{c} \text{diagramma con } W_+ \text{ e } W_- \text{ che si incontrano in un vertice } \gamma \end{array} \right) + \cos \theta_W \left(\begin{array}{c} \text{diagramma con } W^+ \text{ e } W^- \text{ che si incontrano in un vertice } Z \end{array} \right) \end{aligned}$$

in cui il $\sin \theta_W$ sopprime il contributo dei vertici tipo WWA , mentre $\cos \theta_W$ enfatizza i vertici WWZ . L'ultimo contributo alla lagrangiana è quello *a quattro campi*:

$$\begin{aligned} \mathcal{L}_{YM}^{(4)} = & -\frac{g^2}{2} \left(2g^{\mu\nu}g^{\rho\sigma} - g^{\mu\rho}g^{\nu\sigma} - g^{\mu\sigma}g^{\nu\rho} \right) \times \\ & \times \left[W_\mu^+ W_\nu^- \left(A_\rho A_\sigma \sin^2 \theta_W + Z_\rho Z_\sigma \cos^2 \theta_W + 2A_\rho Z_\sigma \sin \theta_W \cos \theta_W \right) - \right. \\ & \left. - \frac{1}{2} W_\mu^+ W_\nu^+ W_\rho^- W_\sigma^- \right] \quad (3.5.14) \end{aligned}$$

$$\supset \left(\begin{array}{c} \text{diagramma a quattro vertici con } W^+, W^-, \gamma, \gamma \end{array} \right) \quad (3.5.15)$$

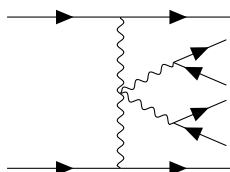
Tutti gli accoppiamenti che abbiamo scritto sono predetti dall'invarianza di gauge, però sperimentalmente i vertici a 3 e 4 bosoni di gauge sono difficili da evidenziare. Infatti, i processi di tale tipo sono eventi rari, cioè hanno piccola sezione d'urto, dunque hanno grande errore statistico. La difficoltà sperimentale ha fatto sì che ci fu' prima la previsione teorica che la verifica negli acceleratori.

⁵Se si vuole svolgere il conto possono essere d'aiuto le relazioni per i vettori di gauge nell'appendice A

Nota che parte della difficoltà sperimentale è data dal fatto che i bosoni W e Z non si misurano direttamente, ma essendo massivi (come vedremo nelle prossime sezioni) decadono e si osservano i prodotti dei decadimenti.

Come ci aspettiamo (essendo Z neutro) negli accoppiamenti a 3 campi non abbiamo un contributo del tipo $Z\gamma\gamma$, però nei vertici a 4 campi compare l'unico contributo in cui Z^μ e A^μ si parlano, cioè $W^+W^-Z\gamma$.

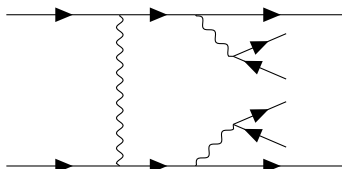
Capiamo meglio il problema sperimentale. Per vedere gli accoppiamenti a 4 campi dovremmo guardare processi con almeno 6 particelle nello stato finale (al tree-level), dunque diagrammi del tipo:



il cui problema è che, avendo 6 particelle finali, lo spazio delle fasi è molto popolato, dunque molto piccolo, e di conseguenza sono eventi molto rari. I processi più frequenti sono processi di scambio.

Dunque, dovremmo produrre moltissime collisioni in modo da avere una statistica sufficiente (un buon errore statistico) per osservare il processo d'interesse. Dobbiamo per questo avere un acceleratore con alta luminosità (molti processi).

Un altro problema è che ci sono contributi al diagramma con 6 particelle, cioè al vertice con 4 bosoni, anche da altri diagrammi senza il vertice che ci interessa; ad esempio un possibile contributo può essere:



I contributi dei diagrammi senza il vertice a 4 bosoni creano quello che chiamiamo *background*. I processi a 3 e 4 bosoni interagenti sono caratterizzati da grande background. Per risolvere il problema, quindi cercare di ridurre il background, ci si mette in situazioni sperimentali (configurazioni di impulso o di spin) in cui si tolgono i diagrammi che non ci interessano (ciò non per forza è facile farlo), che però a sua volta comporta una riduzione ulteriore dello spazio delle fasi.

3.6 Masse dei bosoni di gauge

È immediato rendersi conto che il lavoro non è ancora finito. Infatti, noi abbiamo scritto fin'ora la lagrangiana libera della teoria, ma senza mettere

alcun termine di massa, né per i campi di materia né per i bosoni di gauge (sempre che serva).

A darci delle intuizioni ci sono, come sempre, le evidenze sperimentali. Per quanto riguarda il fotone si vede che:

$$m_\gamma < 2 \cdot 10^{-16} \text{ eV} - 10^{-17} \text{ eV}$$

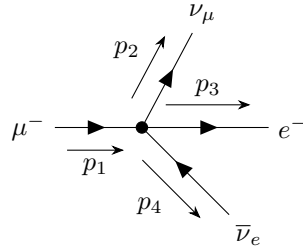
ovvero, gli esperimenti ci dicono che è considerabile massless. Dunque, per quanto riguarda A^μ , la lagrangiana libera che abbiamo scritto è completa e non dobbiamo aggiungerci alcun termine di massa.

Invece, per i $W^\pm$ e Z^0 il discorso è leggermente più sottile. L'idea è che la teoria di gauge massless che abbiamo costruito, a partire dalla Teoria di Fermi, nel caso di limite a bassa energia non riproduce ciò che avevamo trovato in precedenza.

Spiego meglio. Noi conosciamo gli accoppiamenti tra le varie particelle di materia e riusciamo a riprodurre tali accoppiamenti in una teoria di gauge massless; tutto ciò non da problemi. Però, nel momento in cui andiamo a calcolare osservabili fisiche, cioè sezioni d'urto o ampiezze di probabilità, ci aspettiamo di ritrovare gli stessi risultati, nel limite a basse energie, che avevamo nella teoria di Fermi a cui ci siamo ispirati. Nel caso in cui mantenessimo i bosoni $W^\pm$ e Z^0 massless, non riusciremmo a riprodurre la teoria di Fermi. Dunque, abbiamo bisogno di termini di massa per tali campi. Vediamo tutto più esplicitamente nelle sezioni che seguono.

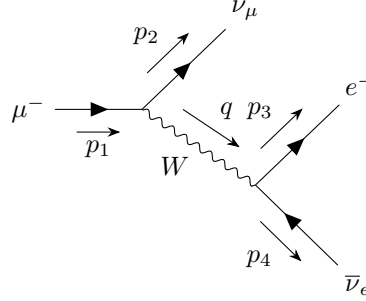
3.6.1 Massa per il bosone $W^\pm$

Prendiamo come esempio il decadimento del muone. Abbiamo già visto che la teoria di Fermi ci dice:



$$\mathcal{M} \propto \frac{G_F}{\sqrt{2}} \bar{u}(p_2) \gamma^\mu (1 - \gamma_5) u(p_1) \bar{u}(p_3) \gamma_\mu (1 - \gamma_5) v(p_4) \quad (3.6.1)$$

mentre la teorie elettrodebole di gauge che abbiamo costruito:



$$\mathcal{M} \propto \left(\frac{g}{2\sqrt{2}} \right)^2 \bar{u}(p_2) \gamma^\mu (1 - \gamma_5) u(p_1) \bar{u}(p_3) \gamma_\mu (1 - \gamma_5) v(p_4) \frac{1}{q^2 - M_W^2} \quad (3.6.2)$$

in cui ipotizziamo la forma del propagatore del bosone W :

$$\sim \frac{g_{\mu\nu}}{q^2 - M_W^2}. \quad (3.6.3)$$

Per riprodurre la teoria di Fermi deve valere:

$$\frac{g^2}{8} \left| \frac{1}{q^2 - M_W^2} \right| \sim \frac{G_F}{\sqrt{2}} \sim 10^{-5} \text{ GeV}^{-2}. \quad (3.6.4)$$

Però abbiamo:

$$q^2 = (p_1 - p_2)^2 \underset{m_\nu \sim 0}{=} m_\mu^2 - 2p_1 p_2 \underset{\text{SR. cm}}{=} m_\mu(m_\mu - 2E_2) \quad (3.6.5)$$

in cui ci siamo messi nel sistema di riferimento in cui il muone che decade è fermo. Abbiamo anche che vale:

$$0 \leq E_2 \leq \frac{m_\mu}{2} \implies 0 \leq q^2 \leq m_\mu^2 \sim (100 \text{ MeV})^2 \quad (3.6.6)$$

che ci dà una condizione per il decadimento.

Vediamo che succede nel caso $M_W = 0$. Dobbiamo analizzare la relazione (3.6.4). Scriviamo:

$$\frac{g^2}{8} \frac{1}{q^2} = \frac{e^2}{8 \sin^2 \theta_W} \frac{1}{q^2} \quad (3.6.7)$$

$$\geq \frac{e^2}{8 m_\mu^2 \sin^2 \theta_W} \quad (3.6.8)$$

$$\geq \frac{e^2}{8 m_\mu^2} \quad (3.6.9)$$

$$= \frac{4\pi\alpha}{8 m_\mu^2} \quad (3.6.10)$$

$$= \frac{\pi}{2} \frac{1}{137} \frac{1}{(0.1 \text{ GeV})^2} \quad (3.6.11)$$

in cui in (3.6.9) abbiamo utilizzato il fatto che $\sin^2 \theta_W \leq 1$. Abbiamo dunque trovato che:

$$\frac{g^2}{8} \frac{1}{g^2} \geq 1 \text{ GeV}^{-2} \gg \frac{G_F}{\sqrt{2}} \sim 10^{-5} \text{ GeV}^{-2}. \quad (3.6.12)$$

Dunque, dal momento che la relazione non è soddisfatta, vediamo che $M_W = 0$ non riproduce la teoria di Fermi a basse energie.

Vediamo ora il caso in cui $M_W \neq 0$. Prendiamo in particolare $M_W^2 \gg m_\mu^2 \geq q^2$, così che:

$$|q^2 - M_W^2| \sim M_W^2 \quad (3.6.13)$$

e dalla condizione (3.6.4) ricaviamo un limite superiore per M_W . Invertendo (3.6.4) possiamo trovare:

$$M_W^2 = \frac{g^2 \sqrt{2}}{8G_F} = \frac{e^2}{\sin^2 \theta_W} \frac{1}{4\sqrt{2}G_F} \quad (3.6.14)$$

in cui se mettiamo che $\sin^2 \theta_W \leq 1$, cioè ignoriamo il fatto di conoscere qualcosa su θ_W , troviamo:

$$M_W^2 \geq \frac{e^2}{4\sqrt{2}G_F} \simeq (37 \text{ GeV})^2 \quad (3.6.15)$$

ma se utilizziamo il valore sperimentale $\sin^2 \theta_W \sim 0.23$, allora troviamo:

$$M_W \sim 80 \text{ GeV}. \quad (3.6.16)$$

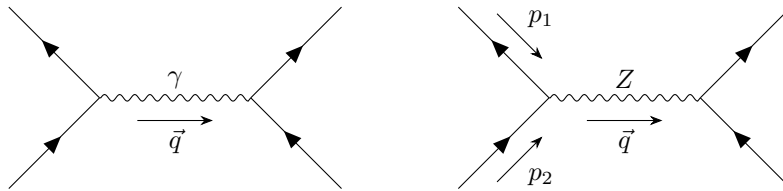
Dunque, i bosoni W devono essere massivi per essere consistenti gauge-Fermi e per le evidenze fenomenologiche.

3.6.2 Massa per il bosone Z^0

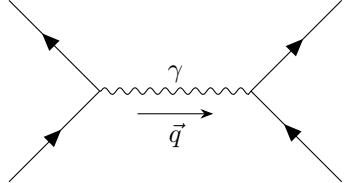
Possiamo fare discorsi analoghi al W anche per il bosone Z . Storicamente, al tempo di Fermi, non c'erano evidenze per Z^0 , cioè non si osservavano correnti neutre alle basse energie. Si riesce a spiegare il contributo soppresso delle correnti neutre tramite la massa del bosone di gauge.

Infatti, prendendo un processo descritto sia dalla QED che dalla teoria elettrodebole, si può mostrare che, nonostante esso prenda contributo da entrambi i diagrammi che come abbiamo visto hanno costanti di accoppiamento $g > e$, il termine debole è soppresso a causa del bosone mediatore massivo.

Consideriamo ad esempio i processi:

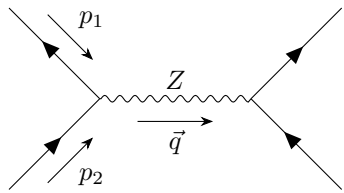


ad ogni valore di $q = p_1 + p_2$ abbiamo il contributo di entrambi i diagrammi. Però, come già anticipato, gli esperimenti ci dicono che a $q^2 = s$ molto piccolo il diagramma con Z è molto soppresso. Vediamo che i due diversi contributi sono:



The diagram shows two incoming fermion lines (arrows pointing right) from the left, meeting at a vertex. A wavy line labeled γ connects this vertex to another vertex on the right, where two outgoing fermion lines (arrows pointing right) emerge. A horizontal arrow labeled $\vec{q}$ points from the left vertex to the right vertex.

$$\propto \frac{e^2}{q^2} \quad (3.6.17)$$



The diagram is similar to the one above, but the wavy line is labeled Z . The incoming fermion lines are labeled p_1 and p_2 with arrows pointing towards the vertex. The horizontal arrow is labeled $\vec{q}$.

$$\propto \frac{e^2}{4 \cos^2 \theta_W \sin^2 \theta_W} \frac{1}{q^2 - M_Z^2} \quad (3.6.18)$$

e se $M_Z = 0$ sono simili a tutte le energie. Per riprodurre il risultato degli esperimenti, per cui a basse energie non si vedono i contributi delle correnti neutre, scriviamo:

$$M_Z^2 \gg q^2 \quad (3.6.19)$$

che implica:

$$\left| \frac{e^2}{4 \cos^2 \theta_W \sin^2 \theta_W} \frac{1}{q^2 - M_Z^2} \right| \sim \frac{e^2}{M_Z^2} \ll \frac{e^2}{q^2}. \quad (3.6.20)$$

Perciò ci aspettiamo che anche il bosone Z abbia una massa, e che questa sia più grande rispetto alle energie in gioco negli esperimenti che non hanno trovato correnti deboli neutre.

Per misurare sperimentalmente (a LEP hanno considerato il processo $e^+e^- \rightarrow \mu^+\mu^-$) la massa dello Z basta aumentare l'energia $q^2 = (p_1 + p_2)^2 = s = 4E_{cm}^2$, così che quando q^2 diventa dell'ordine di M_Z^2 , il fattore $1/(q^2 - M_Z^2)$ diventa molto più grande di $1/q^2$, quindi il contributo dello Z diventa predominante rispetto il fotone e genera un picco nella sezione d'urto; come in figura 3.1.

3.6.3 Come dare massa ai bosoni di gauge

Nelle sezioni precedenti abbiamo capito che i bosoni di gauge dell'interazione debole devono essere massivi, per poter riprodurre i risultati sperimentali rispettati dalla teoria di Fermi. Inoltre, parlando del bosone W abbiamo supposto un propagatore (ispirati da quello del fotone) della forma:

$$\frac{-ig_{\mu\nu}}{p^2 - M_W^2} \quad (3.6.21)$$

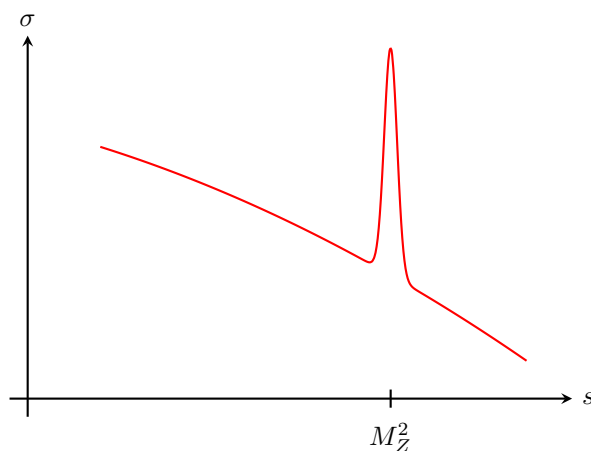


Figura 3.1: Andamento della sezione d'urto per il processo considerato. Individuiamo la prima regione linearmente decrescente poiché in tale regione è dominante il contributo di QED e l'andamento è dato da $1/q^2$. Inoltre siamo tranquilli nel dire che la risonanza della sezione d'urto non porta a divergenze per tutte le considerazioni fatte nella sezione riguardo la rinormalizzazione.

che effettivamente ci riproduce Fermi nel limite $|p^2| \ll M_W^2$. Però vorremmo anche dimostrare che invertendo il termine cinetico della lagrangiana otteniamo effettivamente quella forma.

Storicamente il dover dare massa ai bosoni di gauge fu' un problema. Infatti, se avessimo avuto una teoria di gauge massless avremmo avuto un piano di lavoro ben chiaro, poiché sarebbe bastato mimare quanto fatto per modellizzare la QED e avremmo finito. Invece, il dover mettere un termine di massa nella lagrangiana per i bosoni mediatori rovina tutti i progressi fatti dalla teoria di Fermi. Se con una teoria di gauge massless abbiamo sia l'invarianza per un certo gruppo di simmetria sia la rinormalizzabilità, nel momento in cui aggiungiamo a mano (cioè seguendo il metodo che fin'ora abbiamo sempre utilizzato per dare massa alle particelle) il termine di massa, rompiamo, non solo l'invarianza di gauge (che vedremo nella prossima sezione), ma anche la rinormalizzabilità, rovinando il power counting.

Parliamo un attimo del propagatore e della rinormalizzabilità della teoria.

Nella lagrangiana di Yang-Mills abbiamo il termine cinetico di $W^\pm$:

$$\mathcal{L}_{kin}^W = -\frac{1}{2}W_{\mu\nu}^+W_{-}^{\mu\nu} \quad (3.6.22)$$

$$= -\frac{1}{2}\left(\partial_\mu W_\nu^+ - \partial_\nu W_\mu^+\right)\left(\partial^\mu W_-^\nu - \partial^\nu W_-^\mu\right) \quad (3.6.23)$$

$$= -\partial_\mu W_\nu^+ \partial^\mu W_-^\nu + \partial_\mu W_\nu^+ \partial^\nu W_-^\mu \quad (3.6.24)$$

in cui abbiamo usato:

$$W_{\mu\nu}^+ = \partial_\mu W_\nu^+ - \partial_\nu W_\mu^+.$$

Se aggiungessimo un termine di massa a mano avremmo:

$$\mathcal{L}^W = \mathcal{L}_{kin}^W + M_W^2 W_+^\mu W_\mu^- \quad (3.6.25)$$

$$= W_\alpha^+ X^{\alpha\beta} W_\beta^- \quad (3.6.26)$$

in cui abbiamo scritto l'operatore:

$$X^{\alpha\beta} = -\overleftarrow{\partial}_\mu \overrightarrow{\partial}^\mu g^{\alpha\beta} + \overleftarrow{\partial}^\beta \overrightarrow{\partial}^\alpha + M_W^2 g^{\alpha\beta} \quad (3.6.27)$$

che nello spazio di Fourier⁶ diventa:

$$\overline{X}^{\alpha\beta} = -(-ip_\mu)(ip^\mu)g^{\alpha\beta} + (-ip^\beta)(ip^\alpha) + M_W^2 g^{\alpha\beta} \quad (3.6.30)$$

$$= -(p^2 - M_W^2)g^{\alpha\beta} + p^\alpha p^\beta. \quad (3.6.31)$$

Sappiamo che il propagatore della teoria è l'operatore $i(\overline{X}^{-1})_{\mu\nu}$, tale per cui:

$$(\overline{X})^{\alpha\beta}(\overline{X}^{-1})_{\beta\gamma} = g^\alpha_\gamma. \quad (3.6.32)$$

Facciamo l'ansatz:

$$i(\overline{X}^{-1})_{\beta\gamma} = Ag_{\beta\gamma} + Bp_\beta p_\gamma \quad (3.6.33)$$

ed espandiamo la condizione (3.6.32):

$$A\left[-(p^2 - M_W^2)g^\alpha_\gamma + p^\alpha p_\gamma\right] + B\left[-(p^2 - M_W^2)p^\alpha p_\gamma + p^2 p^\alpha p_\gamma\right] = g^\alpha_\gamma \quad (3.6.34)$$

eguagliando le forme tensoriali si ha:

$$\begin{cases} -A(p^2 - M_W^2)g^\alpha_\gamma = g^\alpha_\gamma \\ [A + M_W^2 B]p^\alpha p_\gamma = 0 \end{cases} \quad (3.6.35)$$

⁶Ricorda che ci possiamo andare operando le sostituzioni:

$$\overleftarrow{\partial}_\mu \longrightarrow -ip_\mu \quad (3.6.28)$$

$$\overrightarrow{\partial}_\mu \longrightarrow ip_\mu. \quad (3.6.29)$$

che portano a:

$$A = -\frac{1}{p^2 - M_W^2} \quad , \quad B = -\frac{A}{M_W^2} = \frac{1}{p^2 - M_W^2} \frac{1}{M_W^2} \quad (3.6.36)$$

e dunque al propagatore:

$$i(\overline{X}^{-1})_{\mu\nu} = \frac{-i}{p^2 - M_W^2} \left(g_{\mu\nu} - \frac{p_\mu p_\nu}{M_W^2} \right). \quad (3.6.37)$$

Nell'espressione trovata, oltre il termine $\frac{-ig_{\mu\nu}}{p^2 - M_W^2}$, c'è anche il contributo di $ip_\mu p_\nu / [M_W^2(p^2 - M_W^2)]$ che, quando $p^\mu \rightarrow \infty$, non va a zero e rovina il power counting per la rinormalizzabilità.

Infatti, il termine $ip_\mu p_\nu / [M_W^2(p^2 - M_W^2)]$ non ci fornisce più una potenza -2 negli impulsi di loop del power counting, bensì una potenza nulla. Un propagatore fatto in questo modo, che discende dal termine massivo introdotto, ci rovina il power counting per la rinormalizzabilità.

Dunque, ci servirà un'ulteriore pezzo nella teoria (il gauge fixing) per avere qualcosa di utilizzabile.

Si può anche vedere che inserendo il termine di massa all'interno della lagrangiana rompiamo l'invarianza di gauge dimostrata nella sezione §3.4.

Infatti, prese le trasformazioni:

$$W_i^\mu \longrightarrow W_i^\mu - \partial_\mu \Lambda_i(x) + ig\Lambda_a(x)W_j^\mu(T_a)_{ji} \quad (3.6.38)$$

per $i = 1, 2$ possiamo vedere (trascurando il termine contenente g nella trasformazione dei W_i^μ , poiché è sufficiente il primo termine a rompere l'invarianza) che:

$$W_\mu^\pm = \frac{1}{\sqrt{2}} (W_\mu^1 \mp iW_\mu^2) \quad (3.6.39)$$

$$\longrightarrow \frac{1}{\sqrt{2}} \left[W_\mu^1 \mp iW_\mu^2 - \partial_\mu (\underbrace{\Lambda_1 \mp i\Lambda_2}_{\Lambda^\pm}) + \dots \right] \quad (3.6.40)$$

$$= W_\mu^\pm - \partial_\mu \Lambda^\pm + \dots \quad (3.6.41)$$

e di conseguenza.

$$M_W^2 W_+^\mu W_\mu^- \longrightarrow M_W^2 W_+^\mu W_\mu^- - M_W^2 \left[W_\mu^+ \partial^\mu \Lambda_- + W_\mu^- \partial^\mu \Lambda_+ \right] + \dots \quad (3.6.42)$$

$$\neq M_W^2 W_+^\mu W_\mu^-. \quad (3.6.43)$$

Dunque, un termine di massa messo a mano nella lagrangiana del nostro modello rompe l'invarianza di gauge (che è, tra l'altro, essenziale per la rinormalizzabilità).

Dobbiamo trovare un altro modo per dare massa, non solo ai bosoni di gauge, ma anche ai campi di materia.

3.7 Rottura spontanea di simmetria

Vediamo ora il meccanismo con cui riusciamo a dare massa alle particelle elementari del Modello Standard. Innanzitutto precisiamo che quando diciamo *rottura* intendiamo dire che si rompe la simmetria del vuoto, ma mai quella di gauge (altrimenti romperemmo la teoria).

Introduciamo un doppietto di $SU(2)$ (visto che già sappiamo che dovremo inserirlo all'interno di una teoria il cui gruppo di simmetria è $SU(2)$):

$$\phi = \begin{pmatrix} \phi_1 \\ \phi_2 \end{pmatrix} \quad (3.7.1)$$

con lagrangiana:

$$\mathcal{L}_\phi = \partial^\mu \phi \partial_\mu \phi - V(\phi) \quad (3.7.2)$$

in cui il potenziale è della classica forma a cappello messicano:

$$V(\phi) = -\mu^2 \phi^\dagger \phi + \lambda (\phi^\dagger \phi)^2 \quad (3.7.3)$$

con il segno "sbagliato" del termine di massa (termine quadratico), dunque $\mu^2 > 0$, che porta $V(\phi)$ ad avere un minimo ϕ_0 , che soddisfa:

$$\left. \frac{\delta V}{\delta \phi} \right|_{\phi=\phi_0} = \left(-\mu^2 + 2\lambda \phi_0^\dagger \phi_0 \right) \phi_0^\dagger = 0 \quad (3.7.4)$$

e dunque tale che:

$$\phi_0^\dagger \phi_0 = \frac{\mu^2}{2\lambda} = \frac{v^2}{2} \quad (3.7.5)$$

in cui abbiamo introdotto:

$$v^2 = \frac{\mu^2}{\lambda}. \quad (3.7.6)$$

Noi sappiamo che i *campi* sono oscillazioni intorno al minimo del potenziale ($\phi_0 \neq 0$):

$$\phi = \phi_0 + \phi'. \quad (3.7.7)$$

Possiamo scrivere la condizione (3.7.5) come:

$$|\phi_1^0|^2 + |\phi_2^0|^2 = \frac{v^2}{2} \quad (3.7.8)$$

che ci dice che il vuoto ha una simmetria $U(2)$, cioè 4 gradi di libertà:

$$(\phi_1^0, \phi_1^{0*}, \phi_2^0, \phi_2^{0*}). \quad (3.7.9)$$

La simmetria $U(2)$ del vuoto viene rotta, in modo **spontaneo**, scegliendo *uno* dei possibili vuoti della teoria, e ci fa passare ad una simmetria $U(1)$,

cioé ad 1 grado di libertà.

Un altro punto molto importante che abbiamo per costruire la teoria è il **teorema di Goldstone**, il quale ci dice che per ogni simmetria del vuoto rotta ($4 - 1 = 3$ simmetrie rotte) esiste uno scalare reale a massa nulla, mentre gli altri bosoni acquistano una massa.

Dunque, nel nostro caso abbiamo 3 bosoni di Goldstone massless ed 1 bosone di Higgs massivo.

Però, a noi questo non serve poiché vorremmo dare massa a 3 bosoni e tenerne uno massless. Possiamo (dobbiamo) utilizzare le interazioni a nostro vantaggio.

Se il campo ϕ è accoppiato a dei bosoni di gauge:

- Tutti i campi di gauge che interagiscono con il campo che ha acquisito massa, anche loro la acquisiscono.
- I bosoni di Goldstone diventano non fisici, cioè esistono dei gauge in cui essi si disaccoppiano (oppure anche se ci si mette in gauge in cui essi interagiscono, essi esistono solamente come particelle virtuali).

Nota. Un bosone massless, ha due gradi di libertà che corrispondono alle polarizzazioni trasverse, ma quando diventa massivo esso guadagna un grado di libertà (polarizzazione longitudinale).

3.7.1 Meccanismo di Higgs per i bosoni $W^\pm$ e Z

Dunque, a noi interessa dare massa ai bosoni $W_\mu^\pm$ e Z_μ del Modello Standard. Abbiamo capito che per poter dare massa ai bosoni di gauge attraverso la rottura spontanea di simmetria, dobbiamo accoppiare il campo ϕ con i vettori di gauge. Lo facciamo come al solito attraverso la derivata covariante $SU(2) \times U(1)$:

$$\partial^\mu \phi \longrightarrow \mathcal{D}^\mu \phi = \partial^\mu \phi - igW_a^\mu t_a \phi - ig' B^\mu \frac{Y(\phi)}{2} \phi \quad (3.7.10)$$

e ricordando le relazioni per le matrici $\tau^\pm$ (3.2.23), (3.2.24), quelle per i $W_\mu^\pm$ (3.2.25) e quelle per B^μ (3.2.39) e W_3^μ (3.2.40) possiamo scrivere:

$$\begin{aligned} \mathcal{D}^\mu \phi &= \partial^\mu \phi - \frac{ig}{\sqrt{2}}(W_+^\mu \tau^+ + W_-^\mu \tau^-)\phi - igW_3^\mu t_3 \phi - ig'B^\mu \frac{Y(\phi)}{2}\phi \quad (3.7.11) \\ &= \begin{pmatrix} \partial^\mu \phi_1 - \frac{ig}{\sqrt{2}}W_+^\mu \phi_2 - ig(A^\mu \sin \theta_W + Z^\mu \cos \theta_W)\frac{1}{2}\phi_1 - \\ \qquad \qquad \qquad - ig'(A^\mu \cos \theta_W - Z^\mu \sin \theta_W)\frac{Y(\phi)}{2}\phi_1 \\ \partial^\mu \phi_2 - \frac{ig}{\sqrt{2}}W_-^\mu \phi_1 + ig(A^\mu \sin \theta_W + Z^\mu \cos \theta_W)\frac{1}{2}\phi_2 - \\ \qquad \qquad \qquad - ig'(A^\mu \cos \theta_W - Z^\mu \sin \theta_W)\frac{Y(\phi)}{2}\phi_2 \end{pmatrix} \quad (3.7.12) \end{aligned}$$

se sostituiamo la relazione per la carica elettrica:

$$g \sin \theta_W = g' \cos \theta_W = e \quad (3.7.13)$$

nell'accoppiamento con A^μ vediamo:

$$\mathcal{D}^\mu \phi = \begin{pmatrix} \partial^\mu \phi_1 - ig \left[\frac{1}{\sqrt{2}}W_+^\mu \phi_2 + \frac{\cos^2 \theta_W - \sin^2 \theta_W Y(\phi)}{2 \cos \theta_W} Z^\mu \phi_1 \right] - ie \frac{1+Y(\phi)}{2} A^\mu \phi_1 \\ \partial^\mu \phi_2 - ig \left[\frac{1}{\sqrt{2}}W_-^\mu \phi_1 - \frac{\cos^2 \theta_W + \sin^2 \theta_W Y(\phi)}{2 \cos \theta_W} Z^\mu \phi_2 \right] + ie \frac{1-Y(\phi)}{2} A^\mu \phi_2 \end{pmatrix}. \quad (3.7.14)$$

A questo punto possiamo capire come rompere la simmetria del vuoto. Infatti, abbiamo quattro campi indipendenti (3.7.9), o equivalentemente $\text{Re } \phi_1, \text{Im } \phi_1, \text{Re } \phi_2, \text{Im } \phi_2$, e con la rottura spontanea di simmetria tre di loro saranno bosoni di Goldstone, rimanendo a massa nulla, mentre il quarto acquisirà un VEV e quindi una massa.

Però sappiamo anche che i campi di gauge che interagiscono con il campo che ha un VEV non nullo, acquisteranno anche loro una massa. Visto che vogliamo che $W_\mu^\pm$ e Z^μ acquistino una massa, mentre A^μ deve rimanere a massa nulla, dobbiamo far sì che in (3.7.14) A^μ non si accoppi con il campo che ha un valore di aspettazione del vuoto diverso da zero. Quindi in pratica abbiamo due possibilità:

- Se ci mettiamo nella situazione in cui $\phi_1^0 = 0$ e $\phi_2^0 \neq 0$, allora A^μ non si deve accoppiare con ϕ_2 , il che vuol dire che dobbiamo imporre $Y(\phi) = 1$.
- Se ci mettiamo nella situazione in cui $\phi_1^0 \neq 0$ e $\phi_2^0 = 0$, allora A^μ non si deve accoppiare con ϕ_1 , il che vuol dire che dobbiamo imporre $Y(\phi) = -1$.

Le due scelte sono del tutto equivalenti, noi scegliamo la prima. Perciò avremo:

$$\mathcal{D}^\mu \phi = \begin{pmatrix} \partial^\mu \phi_1 - ig \left[\frac{1}{\sqrt{2}} W_+^\mu \phi_2 + \frac{\cos^2 \theta_W - \sin^2 \theta_W}{2 \cos \theta_W} Z^\mu \phi_1 \right] - ie A^\mu \phi_1 \\ \partial^\mu \phi_2 - ig \left[\frac{1}{\sqrt{2}} W_-^\mu \phi_1 - \frac{1}{2 \cos \theta_W} Z^\mu \phi_2 \right] \end{pmatrix} \quad (3.7.15)$$

e, avendo scelto $\phi_1^0 = 0$, abbiamo che la condizione (3.7.8) diventa semplicemente:

$$|\phi_2^0| = \frac{v}{\sqrt{2}}. \quad (3.7.16)$$

Ora, per le scelte fatte, abbiamo:

$$\phi_0 = \begin{pmatrix} 0 \\ \phi_2^0 \end{pmatrix} \quad (3.7.17)$$

ma per mettere in evidenza tali scelte, e la conseguenza (3.7.16), conviene riscrivere i campi ϕ_1 e ϕ_2 tramite dei campi complessi $\varphi^\pm$ (uno complesso coniugato dell'altro) e dei campi reali H e φ_0 (che hanno tutto valore di aspettazione del vuoto nullo):

$$\phi_1 = \varphi^+ \quad , \quad \phi_2 = \frac{1}{\sqrt{2}}(v + H + i\varphi_0) \quad (3.7.18)$$

che ovviamente è solo una scelta possibile, per realizzare $\phi_1^0 = 0$ e $|\phi_2^0| = v/\sqrt{2}$, attraverso campi con valore di aspettazione del vuoto nullo, quindi è facile trovare convenzioni diverse nello scrivere il doppietto di Higgs. In questo modo abbiamo:

$$\phi = \begin{pmatrix} \varphi^+ \\ \frac{1}{\sqrt{2}}(v + H + i\varphi_0) \end{pmatrix} \quad (3.7.19)$$

e la derivata covariante:

$$\mathcal{D}^\mu \phi = \begin{pmatrix} \partial^\mu \varphi^+ - \frac{ig}{2} \left[W_+^\mu (v + H + i\varphi_0) + \frac{\cos^2 \theta_W - \sin^2 \theta_W}{\cos \theta_W} Z^\mu \varphi^+ \right] - ie A^\mu \varphi^+ \\ \frac{1}{\sqrt{2}} \partial^\mu (H + i\varphi_0) - \frac{ig}{\sqrt{2}} \left[W_-^\mu \varphi^+ - \frac{1}{2 \cos \theta_W} Z^\mu (v + H + i\varphi_0) \right] \end{pmatrix} \quad (3.7.20)$$

in cui $\varphi^\pm$ e φ_0 sono i tre bosoni di Goldstone senza massa, mentre H è il bosone massivo che prende il nome di **bosone di Higgs** (è la parte reale di ϕ_2 che acquisisce VEV non nullo). Possiamo verificare che i termini massivi sono proprio della forma trovata; infatti, guardando il potenziale $V(\phi)$ esso diventa:

$$V(\phi) = -\mu^2 \phi^\dagger \phi + \lambda (\phi^\dagger \phi)^2 \quad (3.7.21)$$

$$\begin{aligned} &= -\mu^2 \left[\varphi^+ \varphi^- + \frac{v^2}{2} + vH + \frac{1}{2}H^2 + \frac{1}{2}\varphi_0^2 \right] + \\ &\quad + \lambda \left[\varphi^+ \varphi^- + \frac{v^2}{2} + vH + \frac{1}{2}H^2 + \frac{1}{2}\varphi_0^2 \right]^2 \end{aligned} \quad (3.7.22)$$

ma ricordando (3.7.6) ci dice che $v^2 = \mu^2/\lambda$:

$$V(\phi) = \lambda \left[\frac{v^2}{2} + vH + \frac{1}{2}H^2 + \varphi^+\varphi^- + \frac{1}{2}\varphi_0^2 \right] \times \\ \times \left[-\frac{v^2}{2} + vH + \frac{1}{2}H^2 + \varphi^+\varphi^- + \frac{1}{2}\varphi_0^2 \right] \quad (3.7.23)$$

possiamo notare che abbiamo una somma per differenza:

$$V(\phi) = \lambda \left[-\frac{v^4}{4} + H \left(\frac{v^2}{2} - \frac{v^2}{2} \right) + H^2 \left(v^2 + \frac{v^2}{4} - \frac{v^2}{4} \right) + \right. \\ + \varphi^+\varphi^- \left(\frac{v^2}{2} - \frac{v^2}{2} \right) + \varphi_0^2 \left(\frac{v^2}{4} - \frac{v^2}{4} \right) + H^3 \left(\frac{v}{2} + \frac{v}{2} \right) + \\ + H\varphi^+\varphi^-(v+v) + H\varphi_0^2 \left(\frac{v}{2} + \frac{v}{2} \right) + \\ \left. + \frac{H^4}{4} + H^2\varphi^+\varphi^- + \frac{1}{2}H^2\varphi_0^2 + (\varphi^+\varphi^-)^2 + \varphi^+\varphi^-\varphi_0^2 + \frac{\varphi_0^4}{4} \right]. \quad (3.7.24)$$

Come si può notare il termine lineare in H , che corrisponderebbe alla possibilità non fisicamente accettabile di creare dal vuoto il bosone H , si annulla. Questo significa che H è proprio il campo fisico. Similmente si annullano i termini quadratici (cioè di massa) dei campi $\varphi^\pm$ e φ_0 , come atteso. Alla fine otteniamo:

$$V(\phi) = \lambda \left[-\frac{v^4}{4} + v^2H^2 + vH^3 + 2vH\varphi^+\varphi^- + vH\varphi_0^2 + \right. \\ \left. + \frac{H^4}{4} + H^2\varphi^+\varphi^- + \frac{1}{2}H^2\varphi_0^2 + (\varphi^+\varphi^-)^2 + \varphi^+\varphi^-\varphi_0^2 + \frac{\varphi_0^4}{4} \right]. \quad (3.7.25)$$

Si noti la presenza di un valore dell'energia del vuoto non nulla ($-\lambda v^4/4$), che però non ha un particolare significato fisico, essendo l'energia definita a meno di una costante arbitraria, e può essere eliminato.

Ora, considerando tutta la lagrangiana del settore di Higgs $\mathcal{L}^\phi$, in cui abbiamo fatto accoppiamento minimale:

$$\mathcal{L}^{\phi-gauge} = \mathcal{D}_\mu \phi^\dagger \mathcal{D}^\mu \phi - V(\phi) \quad (3.7.26)$$

$$= \mathcal{L}_{(2)}^{\phi-gauge} + \mathcal{L}_{(3)}^{\phi-gauge} + \mathcal{L}_{(4)}^{\phi-gauge} \quad (3.7.27)$$

dove abbiamo indicato a pedice il numero di campi presenti in ogni contributo. Per il termine quadratico, che contiene la parte cinetica di $\varphi^\pm$, φ_0 , H

ed eventuali termini di massa dei bosoni di gauge abbiamo:

$$\begin{aligned}\mathcal{L}_{(2)}^{\phi-gauge} = & \partial_\mu \varphi^+ \partial^\mu \varphi^- + \frac{1}{2} \partial_\mu \varphi_0 \partial^\mu \varphi_0 + \frac{1}{2} \partial_\mu H \partial^\mu H - \\ & - \frac{1}{2} 2\lambda v^2 H^2 + \frac{g^2 v^2}{4} W_+^\mu W_\mu^- + \frac{1}{2} \frac{g^2 v^2}{4 \cos^2 \theta_W} Z^\mu Z_\mu - \\ & - i \frac{gv}{2} W_+^\mu \partial_\mu \varphi^- + i \frac{gv}{2} W_-^\mu \partial_\mu \varphi^+ + \frac{gv}{2 \cos \theta_W} Z^\mu \partial_\mu \varphi_0. \quad (3.7.28)\end{aligned}$$

Nella seconda riga possiamo trovare gli attesi termini di massa per H , $W_\pm^\mu$ e Z^μ che ci permettono di identificare le masse (attenzione alle corrette normalizzazioni):

$$M_H^2 = 2\lambda v^2 = 2\mu^2 \quad (3.7.29)$$

$$M_W^2 = \frac{g^2 v^2}{4} \quad (3.7.30)$$

$$M_Z^2 = \frac{g^2 v^2}{4 \cos^2 \theta_W}. \quad (3.7.31)$$

Notiamo che i termini di massa sono identici ai termini che avevamo aggiunto a mano nella scorsa sezione, però ora non stiamo rompendo la teoria poiché oltre i termini di massa all'interno della lagrangiana sono presenti anche dei termini di interazione che ci sistemano tutto.

Possiamo riscrivere $\mathcal{L}_{(2)}^{\phi-gauge}$ come:

$$\begin{aligned}\mathcal{L}_{(2)}^{\phi-gauge} = & \partial_\mu \varphi^+ \partial^\mu \varphi^- + \frac{1}{2} \partial_\mu \varphi_0 \partial^\mu \varphi_0 + \frac{1}{2} \partial_\mu H \partial^\mu H - \\ & - \frac{1}{2} M_H^2 H^2 + M_W^2 W_+^\mu W_\mu^- + \frac{1}{2} M_Z^2 Z^\mu Z_\mu - \\ & - i M_W W_+^\mu \partial_\mu \varphi^- + i M_W W_-^\mu \partial_\mu \varphi^+ + M_Z Z^\mu \partial_\mu \varphi_0. \quad (3.7.32)\end{aligned}$$

Vedremo nella prossima sottosezione come i termini aggiuntivi, rispetto quelli di massa, ci fanno intuire la strada da seguire per finire di costruire la nostra teoria.

3.7.2 Necessità del gauge fixing

Vediamo che, nella lagrangiana (3.7.32), oltre ai termini di massa per $W_\pm^\mu$ e Z^μ notiamo che compaiono tre accoppiamenti tra i bosoni di gauge e i bosoni di Goldstone (vertici a due campi), che però non contengono nessuna costante di accoppiamento⁷. Il modo più semplice per trattare tali termini

⁷In realtà vediamo che i vertici contengono M_W ed M_Z , che al loro interno hanno la costante di accoppiamento; il problema è che non contengono solamente quella, infatti sappiamo bene che le misure delle masse troviamo dei risultati, che dipendono dalla scala di energia, e che generalmente non sono piccoli (come abbiamo visto dalle scorse sezioni). Per cui, se guardiamo processi ad energie dell'ordine delle masse di W e Z , o più bassi, allora non siamo in regime perturbativo e dobbiamo sommare tutti i contributi.

è di considerarli come dei veri e propri vertici di interazione a due linee e sommare il loro contributo a tutti gli ordini (non dipendendo da una costante di accoppiamento piccola, ogni ordine perturbativo contribuisce con lo stesso peso del precedente). Le regole di Feynman per questi vertici sono:

$$\mu \xrightarrow{p} \text{---} W^+ \text{---} \bullet \text{---} \varphi^- = iM_W p_\mu \quad (3.7.33)$$

$$\mu \xrightarrow{p} \text{---} W^- \text{---} \bullet \text{---} \varphi^+ = -iM_W p_\mu \quad (3.7.34)$$

$$\mu \xrightarrow{p} \text{---} Z \text{---} \bullet \text{---} \varphi^0 = -M_Z p_\mu. \quad (3.7.35)$$

Poichè sono vertici a due linee la loro risommazione entra solo nel propagatore risommato alla Dyson:

$$i \left(\overline{X}_{Dy}^{-1} \right)_{\mu\nu} = \mu \text{---} \text{---} \nu + \mu \text{---} \text{---} \bullet \text{---} \text{---} \nu + \mu \text{---} \text{---} \bullet \text{---} \text{---} \bullet \text{---} \text{---} \nu + \dots \quad (3.7.36)$$

Se cerchiamo ora di calcolare il propagatore risommato per il $W_\pm$ (per lo Z_μ funziona in maniera analoga), ci accorgiamo subito che c'è un problema. Da $\mathcal{L}_{(2)}^{\phi\text{-gauge}}$ si ottengono i seguenti propagatori semplici:

$$\mu \bullet \xrightarrow{p} \text{---} W^\pm \text{---} \bullet \nu = i \left(\overline{X}^{-1} \right)_{\mu\nu} = \frac{-i}{p^2 - M_W^2} \left(g_{\mu\nu} - \frac{p_\mu p_\nu}{M_W^2} \right) \quad (3.7.37)$$

$$\bullet \xrightarrow{p} \text{---} \phi^\pm \text{---} \bullet = i \overline{X}_\varphi^{-1} = \frac{i}{p^2}. \quad (3.7.38)$$

Mettendo tutto nel propagatore risommato (indicando $M = M_W$) otteniamo:

$$\begin{aligned} i \left(\overline{X}_{Dy}^{-1} \right)_{\mu\nu} &= \mu \bullet \xrightarrow{p} \text{---} W^+ \text{---} \bullet \nu + \mu \bullet \xrightarrow{p} \text{---} W^+ \text{---} \bullet \text{---} \text{---} \bullet \xleftarrow{-p} \text{---} W^- \text{---} \bullet \nu + \mu \text{---} \text{---} \bullet \text{---} \text{---} \bullet \text{---} \text{---} \bullet \text{---} \text{---} \nu + \dots \\ &= i \left(\overline{X}^{-1} \right)_{\mu\nu} + i \left(\overline{X}^{-1} \right)_{\mu\alpha} (iM p^\alpha) i \overline{X}_\varphi^{-1} (iM p^\beta) i \left(\overline{X}^{-1} \right)_{\beta\nu} + \\ &\quad + i \left(\overline{X}^{-1} \right)_{\mu\alpha} (iM p^\alpha) i \overline{X}_\varphi^{-1} (iM p^\beta) i \left(\overline{X}^{-1} \right)_{\beta\rho} (iM p^\rho) \times \\ &\quad \times i \overline{X}_\varphi^{-1} (iM p^\sigma) i \left(\overline{X}^{-1} \right)_{\sigma\nu} + \dots \end{aligned} \quad (3.7.39)$$

che se introduciamo:

$$\Pi_\nu^\alpha \equiv (iM p^\alpha) i \overline{X}_\varphi^{-1} (iM p^\beta) i \left(\overline{X}^{-1} \right)_{\beta\nu} \quad (3.7.40)$$

possiamo riscrivere:

$$i \left(\overline{X}_{Dy}^{-1} \right)_{\mu\nu} = i \left(\overline{X}^{-1} \right)_{\mu\nu} + i \left(\overline{X}^{-1} \right)_{\mu\alpha} \Pi_\nu^\alpha + i \left(\overline{X}^{-1} \right)_{\mu\alpha} \Pi_\rho^\alpha \Pi_\nu^\rho + \dots \quad (3.7.41)$$

$$= i \left(\overline{X}^{-1} \right)_{\mu\alpha} \left[\delta_\nu^\alpha + \Pi_\nu^\alpha + \Pi_\rho^\alpha \Pi_\nu^\rho + \dots \right] \quad (3.7.42)$$

$$= i \left(\overline{X}^{-1} \right)_{\mu\alpha} \left[\mathbb{1} + \Pi + \Pi^2 + \Pi^3 + \dots \right]_\nu^\alpha \quad (3.7.43)$$

in cui indichiamo con Π la matrice di elementi Π_β^α . La somma delle matrici $\mathbb{1} + \Pi + \Pi^2 + \Pi^3 + \dots$ è data dalla serie geometrica:

$$\mathbb{1} + \Pi + \Pi^2 + \Pi^3 + \dots = (\mathbb{1} - \Pi)^{-1} \quad (3.7.44)$$

dove abbiamo indicato con $(1 - \Pi)^{-1}$ la matrice inversa di $(1 - \Pi)$, di cui possiamo vedere gli elementi:

$$(1 - \Pi)_\nu^\alpha = \delta_\nu^\alpha - (iMp^\alpha) i \overline{X}_\varphi^{-1} (iMp^\beta) i \left(\overline{X}^{-1} \right)_{\beta\nu} \quad (3.7.45)$$

$$= \delta_\nu^\alpha - (iMp^\alpha) \frac{i}{p^2} (iMp^\beta) \frac{-i}{p^2 - M^2} \left(g_{\beta\nu} - \frac{p_\beta p_\nu}{M^2} \right) \quad (3.7.46)$$

$$= \delta_\nu^\alpha + \frac{M^2}{p^2} \frac{p^\alpha}{p^2 - M^2} \left(p_\nu - \frac{p^2 p_\nu}{M^2} \right) \quad (3.7.47)$$

$$= \delta_\nu^\alpha + \frac{M^2}{p^2} \frac{p^\alpha p_\nu}{p^2 - M^2} \frac{M^2 - p^2}{M^2} \quad (3.7.48)$$

$$= \delta_\nu^\alpha - \frac{p^\alpha p_\nu}{p^2}. \quad (3.7.49)$$

La matrice $(1 - \Pi)$ è singolare, ovvero ha un'autovalore nullo, di conseguenza non è invertibile e non è definibile il propagatore risommato.

Il fatto che non possiamo definire un propagatore per i campi di gauge non ci deve troppo stupire, infatti, il motivo di questo problema è dovuto all'invarianza di gauge. Infatti, in tutte le teorie di gauge non si può definire il propagatore dei campi di gauge prima di fissare il gauge tramite la lagrangiana di gauge-fixing. Lo stesso problema lo si incontra in QED, dove la lagrangiana di gauge $-\frac{1}{4}F_{\mu\nu}F^{\mu\nu}$ risulta singolare e non permette di definire il propagatore del fotone (come abbiamo imparato). Per definire il propagatore del campo A_μ abbiamo visto che si deve prima aggiungere il termine di gauge-fixing (nel gauge di Lorentz è $-\frac{1}{2}(\partial_\mu A^\mu)^2$) e solo dopo possiamo ottenere il propagatore per il fotone ($-ig_{\mu\nu}/p^2$ nel gauge di Lorentz).

Allo stesso modo succede ora per il W_μ e lo Z_μ . Infatti, la rottura spontanea di simmetria ha sì dato una massa ai due bosoni, ma non ha rotto la simmetria di gauge, cosa che possiamo fare solo fissando un gauge. Quindi, anche se il termine di massa nella lagrangiana permette di definire il propagatore di W_μ e Z_μ , la rottura spontanea di simmetria introduce dei termini quadratici aggiuntivi che mantengono la singolarità della lagrangiana di gauge.

Pertanto, anche nel modello $SU(2) \times U(1)$ con rottura spontanea di simmetria bisogna introdurre una lagrangiana di gauge fixing, che permetta di definire i propagatori. La scelta più semplice, e flessibile⁸, è il così detto R_ξ gauge, in cui si introduce la lagrangiana di gauge fixing:

$$\mathcal{L}^{fix} = -\frac{1}{\xi_W} C^+ C^- - \frac{1}{2\xi_Z} C_Z^2 - \frac{1}{2\xi_A} C_A^2 \quad (3.7.50)$$

in cui abbiamo:

$$C^\pm = \partial^\mu W_\mu^\pm \mp i\xi_W M_W \varphi^\pm \quad (3.7.51)$$

$$C_Z = \partial^\mu Z_\mu - \xi_Z M_Z \varphi_0 \quad (3.7.52)$$

$$C_A = \partial^\mu A_\mu. \quad (3.7.53)$$

I parametri di gauge ξ_W, ξ_Z, ξ_A permettono di modificare il gauge, dunque di controllare le quantità che dipendono e quelle che non dipendono dal gauge. Esplicitando la lagrangiana di gauge fixing si ha:

$$\begin{aligned} \mathcal{L}^{fix} = & -\frac{1}{\xi_W} \partial^\mu W_\mu^+ \partial^\nu W_\nu^- - \frac{1}{2\xi_Z} \partial^\mu Z_\mu \partial^\nu Z_\nu - \frac{1}{2\xi_A} \partial^\mu A_\mu \partial^\nu A_\nu - \\ & - \xi_W M_W^2 \varphi^+ \varphi^- - \frac{1}{2} \xi_Z M_Z^2 \varphi_0 \varphi_0 - \\ & - iM_W \partial_\mu W_\mu^+ \varphi^- + iM_W \partial_\mu W_\mu^- \varphi^+ + M_Z \partial_\mu Z^\mu \varphi_0 \end{aligned} \quad (3.7.54)$$

in cui vediamo comparire dei nuovi termini quadratici nei campi $W_\pm^\mu, Z^\mu$ ed A^μ , dei termini di massa per i bosoni di Goldstone (che dipendono dal gauge) e gli stessi vertici a due linee cambiati di segno che ci sono in $\mathcal{L}_{(2)}^{\phi-gauge}$. Dunque questi termini, che dovevamo risommare, si cancellano nella combinazione:

$$\begin{aligned} \mathcal{L}_{(2)}^{\phi-gauge} + \mathcal{L}^{fix} = & \partial_\mu \varphi^+ \partial^\mu \varphi^- - \xi_W M_W^2 \varphi^+ \varphi^- + \frac{1}{2} \partial_\mu \varphi_0 \partial^\mu \varphi_0 - \frac{1}{2} \xi_Z M_Z^2 \varphi_0 \varphi_0 + \\ & + \frac{1}{2} \partial_\mu H \partial^\mu H - \frac{1}{2} M_H^2 H^2 - \frac{1}{2\xi_A} \partial^\mu A_\mu \partial^\nu A_\nu - \\ & - \frac{1}{\xi_W} \partial^\mu W_\mu^+ \partial^\nu W_\nu^- + M_W^2 W_\mu^+ W_\mu^- - \frac{1}{2\xi_Z} \partial^\mu Z_\mu \partial^\nu Z_\nu + \\ & + \frac{1}{2} M_Z^2 Z^\mu Z_\mu \end{aligned} \quad (3.7.55)$$

permettendo di definire propagatori massivi per $W^\pm$ e Z ed un propagatore massless per il fotone. Da questa lagrangiana possiamo vedere che il campo $\varphi^\pm$ ha massa:

$$M_\varphi^2 = \xi_W M_W^2 \quad (3.7.56)$$

⁸Nota che è completamente arbitrario il termine di gauge fixing che si aggiunge, poiché l'unica cosa che conta è che la Fisica sia indipendente da tale scelta. Ovvero, tutte le osservabili fisiche che troviamo, cioè sezioni d'urto, ampiezze o altro, non devono dipendere dai parametri di gauge ξ . Aggiungendo il termine di gauge fixing, il parametro ξ comparirà nei propagatori delle particelle.

che però, dipendendo dal termine di gauge, ci dice che i campi $\varphi^\pm$ non rappresentano particelle fisiche.

3.7.3 Propagatori per i bosoni di gauge

Nella sezione §3.6.3 abbiamo visto come un propagatore per $W^\pm$ non potesse essere definito, con dei termini di massa aggiunti a mano, senza rompere completamente la teoria, mentre nella sezione §3.7.2 abbiamo visto come siano necessari i termini di gauge fixing insieme al meccanismo di Higgs per avere degli operatori cinetici invertibili. Vediamo in questa sezione come trovare delle espressioni per i propagatori dei bosoni di gauge.

Per trovare il propagatore del bosone $W^\pm$, seguendo lo stesso identico procedimento della sezione §3.6.3, dobbiamo vedere i termini quadratici:

$$\mathcal{L}_{(2)}^W = -\partial_\mu W_\nu^+ \partial^\mu W_-^\nu + \partial_\mu W_\nu^+ \partial^\nu W_-^\mu + M_W^2 W_+^\mu W_\mu^- - \frac{1}{\xi_W} \partial_\mu W_+^\mu \partial_\nu W_-^\nu \quad (3.7.57)$$

$$= W_\alpha^+ X^{\alpha\beta} W_\beta^- \quad (3.7.58)$$

in cui abbiamo definito l'operatore:

$$X^{\alpha\beta} = -\overleftarrow{\partial}_\mu \overrightarrow{\partial}^\mu g^{\alpha\beta} + \overleftarrow{\partial}^\beta \overrightarrow{\partial}^\alpha - \frac{1}{\xi_W} \overleftarrow{\partial}^\alpha \overrightarrow{\partial}^\beta + M_W^2 g^{\alpha\beta} \quad (3.7.59)$$

che nello spazio di Fourier⁹ diventa:

$$\overline{X}^{\alpha\beta} = -(-ip_\mu)(ip^\mu)g^{\alpha\beta} + (-ip^\beta)(ip^\alpha) - \frac{1}{\xi_W}(-ip^\alpha)(ip^\beta) + M_W^2 g^{\alpha\beta} \quad (3.7.62)$$

$$= -(p^2 - M_W^2)g^{\alpha\beta} + \left(1 - \frac{1}{\xi_W}\right)p^\alpha p^\beta. \quad (3.7.63)$$

Sappiamo che il propagatore della teoria è l'operatore $i(\overline{X}^{-1})_{\beta\gamma}$, tale per cui:

$$(\overline{X})^{\alpha\beta}(\overline{X}^{-1})_{\beta\gamma} = g_\gamma^\alpha. \quad (3.7.64)$$

Facciamo l'ansatz:

$$i(\overline{X}^{-1})_{\beta\gamma} = Ag_{\beta\gamma} + Bp_\beta p_\gamma \quad (3.7.65)$$

⁹Ricorda che ci possiamo andare operando le sostituzioni:

$$\overleftarrow{\partial}_\mu \longrightarrow -ip_\mu \quad (3.7.60)$$

$$\overrightarrow{\partial}_\mu \longrightarrow ip_\mu. \quad (3.7.61)$$

ed espandiamo la condizione (3.7.64):

$$\begin{aligned} A \left[- (p^2 - M_W^2) g_\gamma^\alpha + \left(1 - \frac{1}{\xi_W} \right) p^\alpha p_\gamma \right] + \\ + B \left[- (p^2 - M_W^2) p^\alpha p_\gamma + \left(1 - \frac{1}{\xi_W} \right) p^2 p^\alpha p_\gamma \right] = g_\gamma^\alpha \end{aligned} \quad (3.7.66)$$

eguagliando le forme tensoriali si ha:

$$\begin{cases} -A(p^2 - M_W^2)g_\gamma^\alpha = g_\gamma^\alpha \\ [A(1 - \frac{1}{\xi_W}) + B(M_W^2 - \frac{p^2}{\xi_W})]p^\alpha p_\gamma = 0 \end{cases} \quad (3.7.67)$$

che portano a:

$$A = -\frac{1}{p^2 - M_W^2} \quad (3.7.68)$$

$$B = \frac{A}{\frac{p^2}{\xi_W} - M_W^2} \left(1 - \frac{1}{\xi_W} \right) \quad (3.7.69)$$

$$= -\frac{1}{p^2 - M_W^2} \frac{\xi_W - 1}{\xi_W} \frac{\xi_W}{p^2 - \xi_W M_W^2} \quad (3.7.70)$$

$$= \frac{1}{p^2 - M_W^2} \frac{1 - \xi_W}{p^2 - \xi_W M_W^2} \quad (3.7.71)$$

e dunque al propagatore:

$$i(\overline{X}^{-1})_{\mu\nu} = \frac{-i}{p^2 - M_W^2} \left(g_{\mu\nu} - (1 - \xi_W) \frac{p_\mu p_\nu}{p^2 - \xi_W M_W^2} \right). \quad (3.7.72)$$

In modo analogo si trova il propagatore per il bosone Z , con il cambio $\xi_W, M_W^2 \rightarrow \xi_Z, M_Z^2$.

Se guardiamo i termini quadratici di $\varphi^\pm$ abbiamo:

$$\mathcal{L}_{(2)}^\varphi = \partial_\mu \varphi^+ \partial^\mu \varphi^- - \xi_W M_W^2 \varphi^+ \varphi^- \quad (3.7.73)$$

$$= \varphi^+ X \varphi^- \quad (3.7.74)$$

dove abbiamo messo:

$$X = \overleftarrow{\partial}_\mu \overrightarrow{\partial}^\mu - \xi_W M_W^2 \quad (3.7.75)$$

che nello spazio degli impulsi è:

$$\overline{X} = p^2 - \xi_W M_W^2 \quad (3.7.76)$$

per cui il propagatore è:

$$i\overline{X}^{-1} = \frac{i}{p^2 - \xi_W M_W^2} \quad (3.7.77)$$

in cui osserviamo che:

$$i\bar{X}^{-1} \xrightarrow{p^2 \rightarrow \infty} \mathcal{O}\left(\frac{1}{p^2}\right) \quad (3.7.78)$$

per ξ finiti, per cui non rovina il power counting. Dunque, scelte di gauge fixing con ξ finito saranno quelle più indicate per verificare la rinormalizzabilità delle teorie.

Vediamo anche che $i\bar{X}^{-1}$ ha un polo non fisico per:

$$p^2 = \xi_W M_W^2 \equiv M_\varphi^2 \quad (3.7.79)$$

che viene cancellato, ordine per ordine, da termini analoghi nel propagatore del W , il quale si può riscrivere, in modo analogo a prima, ma mettendo in evidenza il polo a denominatore, come:

$$i\bar{X}_{\mu\nu}^{-1} = \frac{-i}{p^2 - M_W^2} \left(g_{\mu\nu} - \frac{p_\mu p_\nu}{M_W^2} \right) - \frac{p_\mu p_\nu}{M_W^2} \frac{i}{p^2 - \xi_W M_W^2} \quad (3.7.80)$$

in cui possiamo notare le due componenti del propagatore del W , rispettivamente non longitudinale e longitudinale, i quali corrispondono a due tensori diversi. La combinazione:

$$\frac{-i}{p^2 - M_W^2} \left(g_{\mu\nu} - \frac{p_\mu p_\nu}{M_W^2} \right)$$

fa sì che quando il W è on-shell corrisponde ad un W trasverso¹⁰. Possiamo vedere che il secondo termine, quello longitudinale, è quello che si cancella con il contributo non fisico dato dal propagatore di φ ; dunque, possiamo dire che il φ si comporta come la componente longitudinale del W . Ciò che succede nei diagrammi che hanno un φ dentro è che il loro contributo viene cancellato da un corrispondente diagramma che all'interno ha un W , dunque con parte trasversale e longitudinale.

Notiamo che la cancellazione del polo non fisico è necessario proprio per la non fisicità.

Per fare i conti, ed eventualmente dimostrare alcune cose della teoria, possiamo scegliere ξ a piacere, alcune scelte sono:

- Il *gauge di Feynman* è $\xi = 1$, per cui:

$$i\bar{X}_{\mu\nu}^{-1} = \frac{-ig_{\mu\nu}}{p^2 - M_W^2} \quad ; \quad i\bar{X}^{-1} = \frac{i}{p^2 - M_W^2}. \quad (3.7.82)$$

¹⁰Diciamo che un vettore è *trasverso* se ortogonale al suo quadri-impulso p^μ . Se moltiplichiamo il nostro operatore per p^μ otteniamo:

$$p^\mu \left(g_{\mu\nu} - \frac{p_\mu p_\nu}{M_W^2} \right) = p_\nu - \frac{p^2 p_\nu}{M_W^2} = p_\nu \left(1 - \frac{p^2}{M_W^2} \right) \quad (3.7.81)$$

che si annulla quando il W è on-shell.

- Il *gauge di Landau* è $\xi = 0$, per cui:

$$i\bar{X}_{\mu\nu}^{-1} = \frac{-i}{p^2 - M_W^2} \left(g_{\mu\nu} - \frac{p_\mu p_\nu}{p^2} \right) \quad ; \quad i\bar{X}^{-1} = \frac{i}{p^2}. \quad (3.7.83)$$

- Il *gauge di unitario* è $\xi \rightarrow \infty$, per cui:

$$i\bar{X}_{\mu\nu}^{-1} = \frac{-i}{p^2 - M_W^2} \left(g_{\mu\nu} - \frac{p_\mu p_\nu}{M_W^2} \right) \quad ; \quad i\bar{X}^{-1} = 0 \quad (3.7.84)$$

ovvero, avendo propagatore banale, si ha φ disaccoppiato.

I gauge di Feynman e di Landau sono detti gauge rinormalizzabili, perché è più semplice provare la rinormalizzabilità. Infatti, in tali gauge il propagatore va come $1/p^2$ per $p^2 \rightarrow \infty$.

Nel gauge unitario i bosoni di Goldstone si disaccoppiano e la teoria contiene solo i campi fisici; è quindi più facile provare l'unitarietà della matrice S .

Una prova della rinormalizzabilità e unitarietà del Modello Standard si può trovare al capitolo §2.5.4 del testo "*Gauge Theories of the Strong and Electroweak Interaction*" di M. Bohm, A. Denner, H. Joos.

3.7.4 Parametri liberi della lagrangiana

Dai termini di massa di $W^\pm$ e Z possiamo vedere un legame tra essi e l'angolo di Weinberg:

$$M_W^2 = \frac{g^2 v^2}{4} \quad ; \quad M_Z^2 = \frac{g^2 v^2}{4 \cos^2 \theta_W} \quad (3.7.85)$$

da cui vediamo:

$$\frac{M_W^2}{M_Z^2} = \cos^2 \theta_W. \quad (3.7.86)$$

Dunque, misurate le due masse del W e dello Z , l'angolo di Weinberg è determinato. Per cui dalla relazione che lega g e g' , ovvero (3.2.60), possiamo dire che non dobbiamo determinare due costanti di accoppiamento della teoria, ma in realtà una volta determinato l'angolo θ_W , abbiamo un'unica costante per le interazioni elettro-deboli (che si chiamano così proprio perché identificabili insieme).

La massa del bosone di Higgs è legata a μ^2 e λ :

$$M_H^2 = 2\mu^2 = 2\lambda v^2. \quad (3.7.87)$$

Sperimentalmente si sono misurate nel 2024:

$$\begin{aligned} M_W &= 80.377 \pm 0.012 \text{ GeV} \\ M_Z &= 91.1876 \pm 0.002 \text{ GeV} \quad (LEP) \\ M_H &= 125.25 \pm 0.17 \text{ GeV}. \end{aligned}$$

Pertanto, i parametri liberi della lagrangiana sono:

- Le masse della particelle M_W, M_Z, M_H e in seguito vedremo le masse delle particelle di materia.
- La costante di accoppiamento elettrodebole g , che normalmente viene ricavata da G_F (misurata nel decadimento del muone).

Tutti gli altri parametri, ovvero θ_W , e , μ , λ e v , si ricavano a partire da quelli liberi.

Ad esempio, il valore v di aspettazione del vuoto della componente neutra del doppietto di Higgs si può stimare da:

$$\frac{G_F}{\sqrt{2}} = \left(\frac{g}{2\sqrt{2}} \right)^2 \frac{1}{M_W^2} \quad ; \quad M_W^2 = \frac{g^2 v^2}{4} \quad (3.7.88)$$

per cui si trova:

$$v \simeq \sqrt{\frac{1}{G_F \sqrt{2}}} \simeq 246 \text{ GeV}. \quad (3.7.89)$$

3.7.5 Masse per i fermioni - Quark

Abbiamo risolto quasi tutti i problemi sollevati dalla costruzione del modello di gauge $SU(2)_L \times U(1)_Y$. In particolare siamo riusciti a rendere i bosoni di $SU(2)$ massivi, mentre il fotone a lasciarlo massless. Il problema che ci rimane è che abbiamo ancora i fermioni (campi di materia) *massless*.

Però non possiamo aggiungere a mano un termine convenzionale di massa. Infatti, per motivi diversi rispetto quelli visti per i campi di gauge, un termine di quel tipo violerebbe l'invarianza di gauge poiché accoppia fermioni *left* e *right*, che si trasformano diversamente sotto il gauge di $SU(2)$.

Per la violazione di parità nelle interazioni deboli, il gruppo di invarianza è $SU(2)_L$ e non $SU(2)$.

Infatti, il termine che potremmo provare ad inserire è:

$$\mathcal{L}_{m,Dirac} = -m \bar{\psi} \psi = -m (\bar{\psi}_L + \bar{\psi}_R) (\psi_L + \psi_R) \quad (3.7.90)$$

che, utilizzando le relazioni per gli spinori:

$$\psi_L = \frac{1}{2}(1 - \gamma_5)\psi \quad , \quad \psi_L^\dagger = \psi^\dagger \frac{1}{2}(1 - \gamma_5) \quad , \quad \bar{\psi}_L = \bar{\psi} \frac{1}{2}(1 + \gamma_5)$$

possiamo riscrivere:

$$\mathcal{L}_{m,Dirac} = -m \left(\underbrace{\bar{\psi}_L \psi_L}_{=0} + \bar{\psi}_L \psi_R + \bar{\psi}_R \psi_L + \underbrace{\bar{\psi}_R \psi_R}_{=0} \right) \quad (3.7.91)$$

$$= -m (\bar{\psi}_L \psi_R + \bar{\psi}_R \psi_L) \quad (3.7.92)$$

questo accoppiamento left-right è impedito dall'invarianza di gauge poiché ψ_L e ψ_R si trasformano in modo diverso sotto $SU(2)$, ovvero, rispettivamente come doppietti e singoletti.

La soluzione è, anche per i fermioni, nella rottura spontanea di simmetria con il *meccanismo di Higgs*: i fermioni ricevono massa attraverso il loro accoppiamento con il bosone di Higgs.

La generazione della massa dei quark può avvenire aggiungendo dei termini di interazione fra scalari (Higgs) e fermioni (quark) detti *accoppiamenti di Yukawa*. Questi termini aggiuntivi sono gauge invarianti, Lorentz invarianti e rinormalizzabili (operatori $d = 4$). Questi accoppiamenti sono possibili perché il bosone di Higgs è un doppietto. Per un doppietto left, un singoletto right e il campo di Higgs:

$$\begin{pmatrix} \psi_{1,L} \\ \psi_{2,L} \end{pmatrix}, \quad \psi_{2,R}, \quad \phi = \begin{pmatrix} \phi_1 \\ \phi_2 \end{pmatrix}$$

possiamo scrivere l'accoppiamento:

$$\mathcal{L}_{\text{Yukawa}} = -g_Y (\bar{\psi}_{1,L} \bar{\psi}_{2,L}) \begin{pmatrix} \phi_1 \\ \phi_2 \end{pmatrix} \psi_{2,R} + \text{h.c.} \quad (3.7.93)$$

che è uno scalare di Lorentz con dimensione 4 e conserva: carica, isospin, ipercarica e numero barionico.

Però non abbiamo risolto completamente il problema, poiché se ϕ_2 ha un VEV $v/\sqrt{2} \neq 0$, allora $\mathcal{L}_{\text{Yukawa}}$ contiene un termine di massa solamente per ψ_2 :

$$\mathcal{L}_{\text{Yukawa}} = -g_Y \frac{v}{\sqrt{2}} [\bar{\psi}_{2,L} \psi_{2,R} + \bar{\psi}_{2,R} \psi_{2,L}] + \dots \quad (3.7.94)$$

$$= -g_Y \frac{v}{\sqrt{2}} \bar{\psi}_2 \psi_2. \quad (3.7.95)$$

Dunque, utilizzando il doppietto ϕ si riesce a dare solamente ai quark down e per i leptoni carichi. Se vogliamo dare massa anche ai quark up (e volendo estendere in seguito il ragionamento ai neutrini) abbiamo bisogno di un altro doppietto che trasformi come ϕ per $SU(2)$. Il doppietto che cerchiamo lo chiamiamo ϕ_c , che è definito a partire da (3.7.19) come:

$$\phi_c = i\sigma_2 \phi^* = \begin{pmatrix} 0 & 1 \\ -1 & 0 \end{pmatrix} \begin{pmatrix} \phi_1^* \\ \phi_2^* \end{pmatrix} \quad (3.7.96)$$

$$= \begin{pmatrix} \phi_2^* \\ -\phi_1^* \end{pmatrix} \quad (3.7.97)$$

$$= \begin{pmatrix} \frac{1}{\sqrt{2}}(v + H - i\varphi_0) \\ -\varphi^- \end{pmatrix} \quad (3.7.98)$$

e di conseguenza:

$$\phi_c^\dagger = (\phi_2 \quad -\phi_1) \quad (3.7.99)$$

$$= \left(\frac{1}{\sqrt{2}}(v + H + i\varphi_0) \quad -\varphi^+ \right). \quad (3.7.100)$$

Dobbiamo controllare che ϕ_c trasformi come ϕ per trasformazioni di $SU(2)$. Noi abbiamo già visto che per trasformazioni infinitesime di $SU(2)$:

$$\begin{aligned}\phi &\longrightarrow \left[1 - ig\Lambda_i t_i\right]\phi \\ \phi^* &\longrightarrow \left[1 + ig\Lambda_i t_i^*\right]\phi^*\end{aligned}$$

dove abbiamo i generatori $t_i = \sigma_i/2$. Possiamo vedere:

$$\phi_c \longrightarrow i\sigma_2 \left(1 + ig\Lambda_i t_i^*\right)\phi^* \quad (3.7.101)$$

$$= \phi_c - ig\Lambda_i (-i\sigma_2) \frac{\sigma_i^*}{2} \phi^* \quad (3.7.102)$$

possiamo analizzare solo il secondo termine, ricordando¹¹:

$$\{\sigma_i, \sigma_j\} = 2\delta_{ij} \implies \begin{cases} \sigma_2\sigma_1 = -\sigma_1\sigma_2 \\ \sigma_2\sigma_3 = -\sigma_3\sigma_2 \end{cases} \quad (3.7.103)$$

$$\sigma_1^* = \sigma_1 \quad , \quad \sigma_2^* = -\sigma_2 \quad , \quad \sigma_3^* = \sigma_3 \quad (3.7.104)$$

che ci aiuta a vedere:

$$\Lambda_i (-i\sigma_2) \frac{\sigma_i^*}{2} = -\frac{i}{2} \left[\Lambda_1 \sigma_2 - \Lambda_2 \sigma_2^2 + \Lambda_3 \sigma_2 \sigma_3 \right] \quad (3.7.105)$$

$$= -\frac{i}{2} \left[-\Lambda_1 \sigma_1 \sigma_2 - \Lambda_2 \sigma_2^2 - \Lambda_3 \sigma_3 \sigma_2 \right] \quad (3.7.106)$$

$$= \left(\Lambda_i \frac{\sigma_i}{2} \right) i\sigma_2 \quad (3.7.107)$$

e dunque che complessivamente si ha:

$$\phi_c \longrightarrow \phi_c - ig\Lambda_i t_i \underbrace{i\sigma_2 \phi^*}_{\phi_c} \quad (3.7.108)$$

$$= \left[1 - ig\Lambda_i t_i\right]\phi_c \quad (3.7.109)$$

esattamente come ϕ .

Però il nostro fornire massa ai quark sembra non finire mai. Infatti, sappiamo dal meccanismo GIM (sezione §3.3) e dall'angolo di Cabibbo, che

¹¹Credo sia più semplice utilizzare la proprietà delle matrici di Pauli $\sigma_2 \sigma_i^* \sigma_2 = -\sigma_i$, così che si possa inserire l'identità $\sigma_2 \sigma_2 = 1$ e vedere:

$$\begin{aligned}\phi_c - ig\Lambda_i (-i\sigma_2) \frac{\sigma_i^*}{2} \phi^* &= \phi_c - ig\Lambda_i (-i\sigma_2) \frac{\sigma_i^*}{2} \sigma_2 \sigma_2 \phi^* = \\ &= \phi_c - ig\Lambda_i \left(i \frac{\sigma_i}{2} \right) \sigma_2 \phi^* = \phi_c - ig\Lambda_i \frac{\sigma_i}{2} \phi_c = \phi_c - ig\Lambda_i t_i^i \phi_c.\end{aligned}$$

i doppietti di $SU(2)_L$, dunque gli spinori Q' , e gli autostati di massa, ovvero Q , non coincidono e abbiamo dovuto introdurre una rotazione per mescolare le famiglie. Utilizzando la notazione:

$$\begin{aligned} u_1, u_2, u_3 &= u, c, t \\ d_1, d_2, d_3 &= d, s, b \end{aligned}$$

possiamo generalizzare la relazione tra campi con e senza $'$, introducendo rotazioni indipendenti tra quark up e down, e tra quark left e right:

$$\begin{aligned} \psi_{q'_{i\chi}} &= \sum_j \left(V_{\chi}^{(q)} \right)_{ij} \psi_{q_{j\chi}} \\ \bar{\psi}_{q'_{i\chi}} &= \sum_j \bar{\psi}_{q_{j\chi}} \left(V_{\chi}^{(q)} \right)_{ji}^{\dagger} \end{aligned} \quad , \quad q = u, d \quad \chi = L, R \quad (3.7.110)$$

cioé esplicitamente:

$$\begin{aligned} \psi_{u'_{iL}} &= \sum_j \left(V_L^{(u)} \right)_{ij} \psi_{u_{jL}} \quad , & \bar{\psi}_{u'_{jL}} &= \sum_j \bar{\psi}_{u_{jL}} \left(V_L^{(u)} \right)_{ji}^{\dagger} \\ \psi_{d'_{iL}} &= \sum_j \left(V_L^{(d)} \right)_{ij} \psi_{d_{jL}} \quad , & \bar{\psi}_{d'_{jL}} &= \sum_j \bar{\psi}_{d_{jL}} \left(V_L^{(d)} \right)_{ji}^{\dagger} \\ \psi_{u'_{iR}} &= \sum_j \left(V_R^{(u)} \right)_{ij} \psi_{u_{jR}} \quad , & \bar{\psi}_{u'_{jR}} &= \sum_j \bar{\psi}_{u_{jR}} \left(V_R^{(u)} \right)_{ji}^{\dagger} \\ \psi_{d'_{iR}} &= \sum_j \left(V_R^{(d)} \right)_{ij} \psi_{d_{jR}} \quad , & \bar{\psi}_{d'_{jR}} &= \sum_j \bar{\psi}_{d_{jR}} \left(V_R^{(d)} \right)_{ji}^{\dagger} \end{aligned}$$

in cui le matrici $V_{L,R}^{(u,d)}$ sono per ora generiche matrici complesse. Il più generico termine invariante di $SU(2)$, che accoppia i quark e i doppietti di Higgs, è:

$$\begin{aligned} \mathcal{L}_Y^{had} &= - \sum_{i,j} \left[Y_{ij}^{(d)} \bar{Q}'_{iL} \phi \psi_{d'_{jR}} + Y_{ij}^{(u)} \bar{Q}'_{iL} \phi_c \psi_{u'_{jR}} + \right. \\ &\quad \left. + Y_{ij}^{(d)\dagger} \bar{\psi}_{d'_{iR}} \phi^{\dagger} Q'_{jL} + Y_{ij}^{(u)\dagger} \bar{\psi}_{u'_{iR}} \phi_c^{\dagger} Q'_{jL} \right] \end{aligned} \quad (3.7.111)$$

dove il segno negativo è stato introdotto per dare il corretto segno al termine di massa dei quark che è generato da $\mathcal{L}_Y^{had}$. Inoltre nota che abbiamo introdotto due matrici $Y^{(d,u)}$, i cui elementi sono tutte le costanti di accoppiamento dei vari quark. La seconda riga è l'hermitiano coniugato della prima. Esplicitando i doppietti:

$$Q'_{iL} = \begin{pmatrix} \psi_{u'_{iL}} \\ \psi_{d'_{iL}} \end{pmatrix} \quad , \quad \bar{Q}'_{iL} = \left(\bar{\psi}_{u'_{iL}} \quad \bar{\psi}_{d'_{iL}} \right)$$

abbiamo la lagrangiana (3.7.111):

$$\begin{aligned} \mathcal{L}_Y^{had} = & - \sum_{ij} \left[Y_{ij}^{(d)} \left(\bar{\psi}_{u'_{iL}} \phi_1 + \bar{\psi}_{d'_{iL}} \phi_2 \right) \psi_{d'_{jR}} + Y_{ij}^{(u)} \left(\bar{\psi}_{u'_{iL}} \phi_2^* - \bar{\psi}_{d'_{iL}} \phi_1^* \right) \psi_{u'_{jR}} + \right. \\ & \left. + Y_{ij}^{(d)\dagger} \bar{\psi}_{d'_{iR}} \left(\phi_1^* \psi_{u'_{jL}} + \phi_2^* \psi_{d'_{jL}} \right) + Y_{ij}^{(u)\dagger} \bar{\psi}_{u'_{iR}} \left(\phi_2 \psi_{u'_{jL}} - \phi_1 \psi_{d'_{jL}} \right) \right]. \end{aligned} \quad (3.7.112)$$

Passando dai campi ruotati (quelli con il ') a quelli non ruotati (senza il '), abbiamo:

$$\begin{aligned} \mathcal{L}_Y^{had} = & - \sum_{i,j} \left[\phi_2 \left(V_L^{(d)\dagger} Y^{(d)} V_R^{(d)} \right)_{ij} \bar{\psi}_{d_{iL}} \psi_{d_{jR}} + \phi_2^* \left(V_L^{(d)\dagger} Y^{(d)} V_R^{(d)} \right)_{ij}^\dagger \bar{\psi}_{d_{iR}} \psi_{d_{jL}} + \right. \\ & + \phi_2^* \left(V_L^{(u)\dagger} Y^{(u)} V_R^{(u)} \right)_{ij} \bar{\psi}_{u_{iL}} \psi_{u_{jR}} + \phi_2 \left(V_L^{(u)\dagger} Y^{(u)} V_R^{(u)} \right)_{ij}^\dagger \bar{\psi}_{u_{iR}} \psi_{u_{jL}} + \\ & + \phi_1 \left(V_L^{(u)\dagger} Y^{(d)} V_R^{(d)} \right)_{ij} \bar{\psi}_{u_{iL}} \psi_{d_{jR}} + \phi_1^* \left(V_L^{(u)\dagger} Y^{(d)} V_R^{(d)} \right)_{ij}^\dagger \bar{\psi}_{d_{iR}} \psi_{u_{jL}} - \\ & \left. - \phi_1^* \left(V_L^{(d)\dagger} Y^{(u)} V_R^{(u)} \right)_{ij} \bar{\psi}_{d_{iL}} \psi_{u_{jR}} - \phi_1 \left(V_L^{(d)\dagger} Y^{(u)} V_R^{(u)} \right)_{ij}^\dagger \bar{\psi}_{u_{iR}} \psi_{d_{jL}} \right]. \end{aligned} \quad (3.7.113)$$

Considerando che nelle nostre convenzioni è il campo ϕ_2 ad acquisire un valore di aspettazione del vuoto, tutti i termini di accoppiamento con ϕ_2 e ϕ_2^* sono quelli destinati a generare i termini di massa dei quark¹². Tenendo conto che il termine di VEV di ϕ_2 e ϕ_2^* è uguale a $v/\sqrt{2}$, avremo che $\mathcal{L}_Y^{had}$ genera i corretti termini di massa per i quark up e per i quark down se le matrici:

$$M^{(d)} = V_L^{(d)\dagger} Y^{(d)} V_R^{(d)} \quad , \quad M^{(u)} = V_L^{(u)\dagger} Y^{(u)} V_R^{(u)} \quad (3.7.114)$$

sono diagonali, reali e positive. Visto che una generica matrice complessa può sempre essere diagonalizzata, in modo tale che gli elementi diagonali siano reali e positivi, per mezzo di una trasformazione biunitaria, allora l'unica condizione per avere i corretti termini di massa per tutti i quark è che le matrici $V_{L,R}^{(u,d)}$ siano unitarie.

Avendo a questo punto delle matrici $M^{(d)}$ e $M^{(u)}$ diagonali, reali e positive:

$$M^{(d)} = \begin{pmatrix} M_1^{(d)} & 0 & 0 \\ 0 & M_2^{(d)} & 0 \\ 0 & 0 & M_3^{(d)} \end{pmatrix} \quad , \quad M^{(u)} = \begin{pmatrix} M_1^{(u)} & 0 & 0 \\ 0 & M_2^{(u)} & 0 \\ 0 & 0 & M_3^{(u)} \end{pmatrix}$$

¹²In effetti il termine di accoppiamento scelto nella nostra definizione di $\mathcal{L}^{had}$ tiene già conto del fatto che sia ϕ_2 ad acquisire un valore di aspettazione del vuoto. Se fosse stato ϕ_1 ad avere un VEV, avremmo dovuto definire $\mathcal{L}^{had}$ scambiando ϕ e ϕ_c , altrimenti avremmo generato vertici di Feynman che non conservano la carica elettrica.

i termini di massa generati da $\mathcal{L}_Y^{had}$ sono:

$$\mathcal{L}_{Y,mass}^{had} = -\frac{v}{2} \sum_i \left[M_i^{(d)} \left(\bar{\psi}_{d_{iL}} \psi_{d_{iR}} + \bar{\psi}_{d_{iR}} \psi_{d_{iL}} \right) + M_i^{(u)} \left(\bar{\psi}_{u_{iL}} \psi_{u_{iR}} + \bar{\psi}_{u_{iR}} \psi_{u_{iL}} \right) \right] \quad (3.7.115)$$

$$= -\frac{v}{\sqrt{2}} \sum_i \left[M_i^{(d)} \bar{\psi}_{d_i} \psi_{d_i} + M_i^{(u)} \bar{\psi}_{u_i} \psi_{u_i} \right] \quad (3.7.116)$$

dove si vede chiaramente che le masse dei quark up e down sono date da:

$$m_{d_i} = \frac{v}{\sqrt{2}} M_i^{(d)} \quad , \quad m_{u_i} = \frac{v}{\sqrt{2}} M_i^{(u)}. \quad (3.7.117)$$

In realtà queste relazioni servono per definire le matrici $M^{(d)}$ e $M^{(u)}$ in funzione dei parametri fisici m_{d_i} e m_{u_i} , che si ottengono dagli esperimenti. Ricordando il legame tra v e la massa del bosone W :

$$M_W = \frac{gv}{2} \quad (3.7.118)$$

otteniamo per le matrici $M^{(d)}$ e $M^{(u)}$:

$$M_i^{(d)} = \frac{\sqrt{2}}{v} m_{d_i} = \frac{g}{\sqrt{2}} \frac{m_{d_i}}{M_W} \quad (3.7.119)$$

$$M_i^{(u)} = \frac{\sqrt{2}}{v} m_{u_i} = \frac{g}{\sqrt{2}} \frac{m_{u_i}}{M_W} \quad (3.7.120)$$

e possiamo esplicitare:

$$M^{(d)} = \frac{\sqrt{2}}{v} \begin{pmatrix} m_{d_1} & 0 & 0 \\ 0 & m_{d_2} & 0 \\ 0 & 0 & m_{d_3} \end{pmatrix} \quad , \quad M^{(u)} = \frac{\sqrt{2}}{v} \begin{pmatrix} m_{u_1} & 0 & 0 \\ 0 & m_{u_2} & 0 \\ 0 & 0 & m_{u_3} \end{pmatrix}.$$

A questo punto possiamo usare queste relazioni per riscrivere la lagrangiana di Yukawa per i quark, dunque capire come sono fatti gli accoppiamenti dei quark con i campi scalari. Conviene anche introdurre la matrice unitaria:

$$V = V_L^{(u)\dagger} V_L^{(d)} \quad (3.7.121)$$

che possiamo dire essere unitaria essendole $V_L^{(d)}$ e $V_L^{(u)}$. Usando (3.7.121) possiamo riscrivere:

$$V_L^{(u)\dagger} Y^{(d)} V_R^{(d)} = V M^{(d)} \quad , \quad V_L^{(d)\dagger} Y^{(u)} V_R^{(u)} = V^\dagger M^{(u)} \quad (3.7.122)$$

e ottenere la lagrangiana:

$$\begin{aligned}
\mathcal{L}_Y^{had} = & - \sum_i (m_{d_i} \bar{\psi}_{d_i} \psi_{d_i} + m_{u_i} \bar{\psi}_{u_i} \psi_{u_i}) - \\
& - \frac{g}{2} \sum_i \left[\frac{m_{d_i}}{M_W} \left(H \bar{\psi}_{d_i} \psi_{d_i} + i \phi_0 \bar{\psi}_{d_{iL}} \psi_{d_{iR}} - i \phi_0 \bar{\psi}_{d_{iR}} \psi_{d_{iL}} \right) + \right. \\
& \quad \left. + \frac{m_{u_i}}{M_W} \left(H \bar{\psi}_{u_i} \psi_{u_i} - i \phi_0 \bar{\psi}_{u_{iL}} \psi_{u_{iR}} + i \phi_0 \bar{\psi}_{u_{iR}} \psi_{u_{iL}} \right) \right] - \\
& - \frac{g}{\sqrt{2}} \sum_{i,j} \left[\frac{m_{d_j}}{M_W} \left(V_{ij} \phi^+ \bar{\psi}_{u_{iL}} \psi_{d_{jR}} + V_{ij}^* \phi^- \bar{\psi}_{d_{jR}} \psi_{u_{iL}} \right) - \right. \\
& \quad \left. - \frac{m_{u_j}}{M_W} \left(V_{ij}^* \phi^- \bar{\psi}_{d_{iL}} \psi_{u_{iR}} + V_{ij} \phi^+ \bar{\psi}_{u_{iR}} \psi_{d_{jL}} \right) \right] \quad (3.7.123)
\end{aligned}$$

in cui notiamo che tutti gli accoppiamenti dei quark con gli scalari sono proporzionali alla massa dei fermioni. In particolare vediamo che la parte neutra non mischia i campi e le masse, mentre lo fa per gli accoppiamenti con $\phi^\pm$. Questa è una conseguenza del fatto che si è usato un'accoppiamento con il doppietto di Higgs per dare massa ai quark.

Questa proporzionalità alla massa dei fermioni è ciò che rende l'Higgs molto complicato da misurare. Infatti, negli acceleratori di solito si accelerano i fermioni leggeri, ma ciò vuol dire che da un punto di vista statistico abbiamo un accoppiamento piccolo tra l'Higgs e tali particelle; sono per questo favoriti i processi di formazione dell'Higgs quando si utilizzano fermioni pesanti. Questo è il motivo per cui a LEP dove si usavano e^+e^- non si è riuscito a misurare il bosone di Higgs.

Arrivati a questo punto possiamo riformulare l'espressione della lagrangiana di interazione delle correnti cariche dei quark (non c'è bisogno di considerare le correnti neutre, che sono indifferenti al mixing tra le famiglie) con il bosone W alla luce delle nuove rotazioni introdotte nella lagrangiana di Yukawa. La corrente adronica, che si accoppia con il W^+ , scritta con i nuovi campi ruotati, è:

$$H^\mu = 2 \sum_k \bar{Q}'_{k,L} \gamma^\mu \tau^+ Q'_{k,L} \quad (3.7.124)$$

$$= 2 \sum_k \bar{\psi}_{u'_{k,L}} \gamma^\mu \psi_{d'_{k,L}} \quad (3.7.125)$$

$$= 2 \sum_{i,j,k} \bar{\psi}_{u_{iL}} \left(V_L^{(u)\dagger} \right)_{ik} \gamma^\mu \left(V_L^{(d)} \right)_{kj} \psi_{d_{jL}} \quad (3.7.126)$$

$$= 2 \sum_{i,j} V_{ij} \bar{\psi}_{u_{iL}} \gamma^\mu \psi_{d_{jL}} \quad (3.7.127)$$

in cui compare nuovamente la matrice:

$$V = V_L^{(u)\dagger} V_L^{(d)} \quad (3.7.128)$$

con elementi:

$$V_{ij} = \sum_k \left(V_L^{(u)\dagger} \right)_{ik} \left(V_L^{(d)} \right)_{kj}. \quad (3.7.129)$$

Analogamente la corrente hermitiana coniugata di H^μ sarà:

$$H_\mu^\dagger = 2 \sum_k \bar{Q}'_{k,L} \gamma_\mu \tau^- Q'_{k,L} \quad (3.7.130)$$

$$= 2 \sum_{i,j} V_{ij}^* \bar{\psi}_{d_{jL}} \gamma_\mu \psi_{u_{iL}}. \quad (3.7.131)$$

La lagrangiana per la corrente carica risulta quindi essere come prima:

$$\mathcal{L}_C^{had} = \sum_i \frac{g}{\sqrt{2}} \left[\bar{Q}'_{iL} W^+ \tau^+ Q'_{iL} + \bar{Q}'_{iL} W^- \tau^- Q'_{iL} \right] \quad (3.7.132)$$

$$= \sum_i \frac{g}{\sqrt{2}} \left[\bar{\psi}_{u_{iL}} W^+ \psi_{d'_{iL}} + \bar{\psi}_{d'_{iL}} W^- \psi_{u_{iL}} \right] \quad (3.7.133)$$

$$= \sum_{i,j} \frac{g}{\sqrt{2}} \left[V_{ij} \bar{\psi}_{u_{iL}} W^+ \psi_{d_{jL}} + V_{ij}^* \bar{\psi}_{d_{jL}} W^- \psi_{u_{iL}} \right] \quad (3.7.134)$$

$$= \sum_{i,j} \frac{g}{2\sqrt{2}} \left[V_{ij} \bar{\psi}_{u_i} W^+ (1 - \gamma_5) \psi_{d_j} + V_{ij}^* \bar{\psi}_{d_j} W^- (1 - \gamma_5) \psi_{u_i} \right] \quad (3.7.135)$$

e rispetto a ciò che avevamo scritto l'unica differenza è che, mentre prima la matrice V ruotava i soli quark down, ora è il prodotto delle due rotazioni up e down $V = V_L^{(u)\dagger} V_L^{(d)}$.

Guardando le lagrangiane $\mathcal{L}_Y^{had}$ (3.7.123) e $\mathcal{L}_C^{had}$ (3.7.135) notiamo che, oltre alla costante d'accoppiamento g e alle masse (dei quark e del W), gli unici altri parametri fisici che compaiono sono gli elementi della matrice unitaria V , che viene comunemente chiamata matrice di Cabibbo-Kobayashi-Maskawa (CKM). Le singole rotazioni dei quark up e down $V_{L,R}^{(u,d)}$ non hanno pertanto un diretto significato fisico (non si possono misurare), ma ce l'ha solo la loro combinazione:

$$V_{CKM} = V = V_L^{(u)\dagger} V_L^{(d)}. \quad (3.7.136)$$

Ricordiamo che il fatto che VCKM sia unitaria garantisce la soppressione delle FCNC.

3.7.6 Masse per i fermioni - Leptoni

Per dare massa ai leptoni procediamo come con i quark, introducendo un accoppiamento di Yukawa tra i doppietti scalari di $SU(2)$, ovvero ϕ (3.7.19) e ϕ_c (3.7.98), e i doppietti di leptoni left che trasformano secondo $SU(2)$. Per essere generali, introduciamo la possibilità, anche per i leptoni, che i

doppietti di $SU(2)_L$ (indicati con un ') e gli autostati di massa (indicati senza ') non coincidano, cioè introduciamo anche per i leptoni una rotazione che mescoli le famiglie. Infatti, vista l'evidenza sperimentale che anche i neutrini hanno una massa (seppur piccola rispetto agli altri leptoni), introduciamo un mixing anche nel settore dei leptoni per vedere il legame tra massa dei neutrini e mixing tra le famiglie.

Usando una notazione analoga a quella dei quark:

$$\begin{aligned} \nu_1, \nu_2, \nu_3 &= \nu_e, \nu_\mu, \nu_\tau \\ l_1, l_2, l_3 &= e, \mu, \tau \end{aligned}$$

introduciamo le quattro matrici complesse $V_{L,R}^{(\nu,l)}$ per descrivere il legame tra campi con e senza ':

$$\begin{aligned} \psi_{f'_{i\chi}} &= \sum_j \left(V_\chi^{(f)} \right)_{ij} \psi_{f_{j\chi}} \\ \bar{\psi}_{f'_{i\chi}} &= \sum_j \bar{\psi}_{f_{j\chi}} \left(V_\chi^{(f)} \right)_{ji}^\dagger \end{aligned} \quad , \quad f = \nu, l \quad \chi = L, R. \quad (3.7.137)$$

Il più generico termine invariante di $SU(2)$ che accoppia i leptoni e i doppietti di Higgs sarà:

$$\begin{aligned} \mathcal{L}_Y^{lep} &= - \sum_{i,j} \left[Y_{ij}^{(l)} \bar{L}'_{iL} \phi \psi_{l'_{jR}} + Y_{ij}^{(\nu)} \bar{L}'_{iL} \phi_c \psi_{\nu'_{jR}} + \right. \\ &\quad \left. + Y_{ij}^{(l)\dagger} \bar{\psi}_{l'_{iR}} \phi^\dagger L'_{jL} + Y_{ij}^{(\nu)\dagger} \bar{\psi}_{\nu'_{iR}} \phi_c^\dagger L'_{jL} \right] \end{aligned} \quad (3.7.138)$$

che, esplicitando i doppietti e passando dai campi ruotati (con ') a quelli non ruotati (senza '), diventa:

$$\begin{aligned} \mathcal{L}_Y^{lep} &= - \sum_{i,j} \left[\phi_2 \left(V_L^{(l)\dagger} Y^{(l)} V_R^{(l)} \right)_{ij} \bar{\psi}_{l_{iL}} \psi_{l_{jR}} + \phi_2^* \left(V_L^{(l)\dagger} Y^{(l)} V_R^{(l)} \right)_{ij}^\dagger \bar{\psi}_{l_{iR}} \psi_{l_{jL}} + \right. \\ &\quad + \phi_2^* \left(V_L^{(\nu)\dagger} Y^{(\nu)} V_R^{(\nu)} \right)_{ij} \bar{\psi}_{\nu_{iL}} \psi_{\nu_{jR}} + \phi_2 \left(V_L^{(\nu)\dagger} Y^{(\nu)} V_R^{(\nu)} \right)_{ij}^\dagger \bar{\psi}_{\nu_{iR}} \psi_{\nu_{jL}} + \\ &\quad + \phi_1 \left(V_L^{(\nu)\dagger} Y^{(l)} V_R^{(l)} \right)_{ij} \bar{\psi}_{\nu_{iL}} \psi_{l_{jR}} + \phi_1^* \left(V_L^{(\nu)\dagger} Y^{(l)} V_R^{(l)} \right)_{ij}^\dagger \bar{\psi}_{l_{iR}} \psi_{\nu_{jL}} - \\ &\quad \left. - \phi_1^* \left(V_L^{(l)\dagger} Y^{(\nu)} V_R^{(\nu)} \right)_{ij} \bar{\psi}_{l_{iL}} \psi_{\nu_{jR}} - \phi_1 \left(V_L^{(l)\dagger} Y^{(\nu)} V_R^{(\nu)} \right)_{ij}^\dagger \bar{\psi}_{\nu_{iR}} \psi_{l_{jL}} \right]. \end{aligned} \quad (3.7.139)$$

Come visto prima, le masse dei leptoni saranno generate dal valore di aspettazione del vuoto di ϕ_2 e ϕ_2^* che è uguale a $v/\sqrt{2}$. Come per i quark, avremo che $\mathcal{L}_Y^{lep}$ genera i corretti termini di massa per i leptoni carichi e per i neutrini se le matrici:

$$M^{(l)} = V_L^{(l)\dagger} Y^{(l)} V_R^{(l)} \quad , \quad M^{(\nu)} = V_L^{(\nu)\dagger} Y^{(\nu)} V_R^{(\nu)} \quad (3.7.140)$$

sono diagonali, reali e positive. Di nuovo questo è sicuramente possibile se le matrici $V_{L,R}^{(\nu,l)}$ sono *unitarie*. Se pertanto le matrici $M^{(\nu,l)}$ sono diagonali, reali e positive:

$$M^{(l)} = \begin{pmatrix} M_1^{(l)} & 0 & 0 \\ 0 & M_2^{(l)} & 0 \\ 0 & 0 & M_3^{(l)} \end{pmatrix}, \quad M^{(\nu)} = \begin{pmatrix} M_1^{(\nu)} & 0 & 0 \\ 0 & M_2^{(\nu)} & 0 \\ 0 & 0 & M_3^{(\nu)} \end{pmatrix} \quad (3.7.141)$$

i termini di massa generati da $\mathcal{L}_Y^{lep}$ sono:

$$\mathcal{L}_{Y,mass}^{lep} = -\frac{v}{2} \sum_i \left[M_i^{(l)} \left(\bar{\psi}_{l_{iL}} \psi_{l_{iR}} + \bar{\psi}_{l_{iR}} \psi_{l_{iL}} \right) + M_i^{(\nu)} \left(\bar{\psi}_{\nu_{iL}} \psi_{\nu_{iR}} + \bar{\psi}_{\nu_{iR}} \psi_{\nu_{iL}} \right) \right] \quad (3.7.142)$$

$$= -\frac{v}{\sqrt{2}} \sum_i \left[M_i^{(l)} \bar{\psi}_{l_i} \psi_{l_i} + M_i^{(\nu)} \bar{\psi}_{\nu_i} \psi_{\nu_i} \right] \quad (3.7.143)$$

dove si vedono chiaramente le masse dei leptoni carichi e dei neutrini:

$$m_{l_i} = \frac{v}{\sqrt{2}} M_i^{(l)}, \quad m_{\nu_i} = \frac{v}{\sqrt{2}} M_i^{(\nu)}. \quad (3.7.144)$$

Anche per i leptoni queste relazioni servono per definire le matrici $M^{(l)}$ e $M^{(\nu)}$ in funzione dei parametri fisici m_{l_i} e m_{ν_i} :

$$M_i^{(l)} = \frac{\sqrt{2}}{v} m_{l_i} = \frac{g}{\sqrt{2}} \frac{m_{l_i}}{M_W} \quad (3.7.145)$$

$$M_i^{(\nu)} = \frac{\sqrt{2}}{v} m_{\nu_i} = \frac{g}{\sqrt{2}} \frac{m_{\nu_i}}{M_W}. \quad (3.7.146)$$

Usiamo ora queste relazioni per riscrivere la lagrangiana di Yukawa per i quark per capire come sono fatti gli accoppiamenti dei quark con i campi scalari. Introduciamo la matrice unitaria (analogo per i leptoni della matrice V_{CKM}):

$$V^{lep} = V_L^{(\nu)\dagger} V_L^{(l)} \quad (3.7.147)$$

in modo da riscrivere:

$$V_L^{(\nu)\dagger} Y^{(l)} V_R^{(l)} = V M^{(l)}, \quad V_L^{(l)\dagger} Y^{(\nu)} V_R^{(\nu)} = V^\dagger M^{(\nu)} \quad (3.7.148)$$

e ottenere la lagrangiana:

$$\begin{aligned}
\mathcal{L}_Y^{lep} = & - \sum_i (m_{l_i} \bar{\psi}_{l_i} \psi_{l_i} + m_{\nu_i} \bar{\psi}_{\nu_i} \psi_{\nu_i}) - \\
& - \frac{g}{2} \sum_i \left[\frac{m_{l_i}}{M_W} \left(H \bar{\psi}_{l_i} \psi_{l_i} + i \phi_0 \bar{\psi}_{l_{iL}} \psi_{l_{iR}} - i \phi_0 \bar{\psi}_{l_{iR}} \psi_{l_{iL}} \right) + \right. \\
& \quad \left. + \frac{m_{\nu_i}}{M_W} \left(H \bar{\psi}_{\nu_i} \psi_{\nu_i} - i \phi_0 \bar{\psi}_{\nu_{iL}} \psi_{\nu_{iR}} + i \phi_0 \bar{\psi}_{\nu_{iR}} \psi_{\nu_{iL}} \right) \right] - \\
& - \frac{g}{\sqrt{2}} \sum_{i,j} \left[\frac{m_{l_j}}{M_W} \left(V_{ij}^{*lep} \phi^+ \bar{\psi}_{\nu_{iL}} \psi_{l_{jR}} + V_{ij}^{lep} \phi^- \bar{\psi}_{l_{jR}} \psi_{\nu_{iL}} \right) - \right. \\
& \quad \left. - \frac{m_{\nu_j}}{M_W} \left(V_{ij}^{*lep} \phi^- \bar{\psi}_{l_{iL}} \psi_{\nu_{jR}} + V_{ij}^{lep} \phi^+ \bar{\psi}_{\nu_{jR}} \psi_{l_{iL}} \right) \right]. \quad (3.7.149)
\end{aligned}$$

Come conseguenza del mixing, anche la lagrangiana di interazione delle correnti cariche cambia e otteniamo:

$$L^\mu = 2 \sum_k \bar{L}'_{k,L} \gamma^\mu \tau^+ L'_{k,L} \quad (3.7.150)$$

$$= 2 \sum_k \bar{\psi}_{\nu'_{k,L}} \gamma^\mu \psi_{l'_{k,L}} \quad (3.7.151)$$

$$= 2 \sum_{i,j,k} \bar{\psi}_{\nu_{iL}} \left(V_L^{(\nu)\dagger} \right)_{ik} \gamma^\mu \left(V_L^{(l)} \right)_{kj} \psi_{l_{jL}} \quad (3.7.152)$$

$$= 2 \sum_{i,j} V_{ij}^{lep} \bar{\psi}_{\nu_{iL}} \gamma^\mu \psi_{l_{jL}} \quad (3.7.153)$$

$$L_\mu^\dagger = 2 \sum_k \bar{L}'_{k,L} \gamma_\mu \tau^- L'_{k,L} \quad (3.7.154)$$

$$= 2 \sum_{i,j} V_{ij}^{*lep} \bar{\psi}_{l_{jL}} \gamma_\mu \psi_{\nu_{iL}} \quad (3.7.155)$$

e la corrente carica:

$$\mathcal{L}_C^{lep} = \sum_i \frac{g}{\sqrt{2}} \left[\bar{L}'_{iL} W^+ \tau^+ L'_{iL} + \bar{L}'_{iL} W^- \tau^- L'_{iL} \right] \quad (3.7.156)$$

$$= \sum_i \frac{g}{\sqrt{2}} \left[\bar{\psi}_{\nu_{iL}} W^+ \psi_{l'_{iL}} + \bar{\psi}_{l'_{iL}} W^- \psi_{\nu_{iL}} \right] \quad (3.7.157)$$

$$= \sum_{i,j} \frac{g}{\sqrt{2}} \left[V_{ij}^{lep} \bar{\psi}_{\nu_{iL}} W^+ \psi_{l_{jL}} + V_{ij}^{*lep} \bar{\psi}_{l_{jL}} W^- \psi_{\nu_{iL}} \right] \quad (3.7.158)$$

$$= \sum_{i,j} \frac{g}{2\sqrt{2}} \left[V_{ij}^{lep} \bar{\psi}_{\nu_i} W^+ (1 - \gamma_5) \psi_{l_j} + V_{ij}^{*lep} \bar{\psi}_{l_j} W^- (1 - \gamma_5) \psi_{\nu_i} \right]. \quad (3.7.159)$$

Dunque, come per i quark, troviamo un mixing per i leptoni. Inoltre, guardando le lagrangiane L_Y^{lep} (3.7.149) e L_C^{lep} (3.7.159) vediamo che, oltre alla

costante d'accoppiamento g e alle masse, gli unici altri parametri fisici sono gli elementi della matrice unitaria V^{lep} , mentre le singole rotazioni $V_{L,R}^{(\nu,l)}$ non hanno diretto significato fisico.

Notiamo inoltre che se le masse dei neutrini sono nulle ($M^{(\nu)} = Y^{(\nu)} = 0$), allora non si possono distinguere (fisicamente) i neutrini ruotati da $V_{L,R}^{(\nu)}$ da quelli non ruotati, perché nella lagrangiana $\mathcal{L}_Y^{lep} + \mathcal{L}_C^{lep}$ comparirebbero di fatto solo i neutrini ruotati. Infatti, nella lagrangiana $\mathcal{L}_Y^{lep}$ compaiono i campi non ruotati solamente quando proporzionali alle masse dei fermioni, che se messe a zero, eliminano i termini non ruotati. Pertanto, vista l'indistinguibilità di rotazione (o meno) che incontriamo, le matrici di rotazione dei neutrini diventano arbitrarie e possiamo imporre $V_L^{(\nu)} = V_L^{(l)}$ senza cambiare la fisica descritta dalla lagrangiana. Così facendo otteniamo che:

$$\text{se } M^{(\nu)} = Y^{(\nu)} = 0 \quad \Longrightarrow \quad V_L^{(\nu)} = V_L^{(l)} \quad \Longrightarrow \quad V^{lep} = V_L^{(\nu)\dagger} V_L^{(l)} = \mathbb{1}$$

cioè: se le masse dei neutrini sono nulle, allora non c'è mixing nel settore leptonico.

Nel primo Modello Standard formulato da Weinberg-Salam-Glashow, per noi quello analizzato nel corso di *Fondamenti di QFT*, erano stati ipotizzati neutrini a massa nulla.

Avendo introdotto il mixing del settore leptonico siamo stati in grado di vedere il legame tra esso e $M^{(\nu)}$. Sperimentalmente si controlla se i ν sono dotati di massa controllando la presenza o meno del mixing leptonico. La difficoltà successiva è la misura delle masse dei ν .

Nota. Questa osservazione sulle masse dei neutrini e la scomparsa del mixing, in linea di principio, si potrebbe fare anche nel settore adronico. Però, in quel caso non ci è venuto in mente di farlo poiché si osserva già sperimentalmente sia il mixing sia le masse dei quark up e down.

3.7.7 Ancora sulla matrice CKM

Possiamo provare a contare i gradi di libertà contenuti all'interno della matrice V_{CKM} . Sicuramente sappiamo che è una matrice unitaria $n_F \times n_F$, cioè con n_F^2 parametri reali. Degli n_F parametri, alcuni di essi sono delle fasi, che se messe a zero portano V_{CKM} ad essere reale e dunque ortogonale (grazie al fatto che era unitaria). Una matrice $n_F \times n_F$ ortogonale è descritta da $n_F(n_F - 1)/2$ angoli. Perciò il numero di fasi è:

$$n_F^2 - \frac{n_F(n_F - 1)}{2} = \frac{1}{2}n_F(n_F + 1) \quad (3.7.160)$$

e tali fasi parametrizzano la parte immaginaria degli elementi di V . Puoi vedere l'equazione (3.7.166) per un esempio di parametrizzazione.

Però non tutte le fasi sono fisicamente osservabili, poiché alcune possono essere riassorbite nei campi essendo la matrice V_{CKM} sempre pinzata tra

spinori, per questo alcune fasi possono essere eliminate. Infatti, possiamo ridefinire i campi dei quark, di numero:

$$n_F + n_F = 2n_F \quad (3.7.161)$$

cioé n_F up e n_F down, assorbendo una fase nella ψ di ogni quark, che è irrilevante nella lagrangiana.

Però attenzione perché c'è ancora una fase overall delle matrici di rotazione dei quark $V_L^{(u,d)}$ che lascia V_{CKM} invariante. Infatti possiamo vedere che se:

$$V_L^{(u)} \longrightarrow e^{i\alpha} V_L^{(u)} \quad (3.7.162)$$

$$V_L^{(d)} \longrightarrow e^{i\alpha} V_L^{(d)} \quad (3.7.163)$$

allora abbiamo:

$$V_{CKM} = V_L^{(u)\dagger} V_L^{(d)} \longrightarrow V_L^{(u)\dagger} e^{-i\alpha} e^{i\alpha} V_L^{(d)} = V_{CKM}. \quad (3.7.164)$$

Dunque le fasi di V_{CKM} che si possono riassorbire nelle ψ sono $2n_F - 1$. Per questo motivo i parametri fisicamente osservabili sono:

- Angoli: $\frac{1}{2}n_F(n_F - 1)$.
- Fasi: $\frac{1}{2}n_F(n_F + 1) - (2n_F - 1) = \frac{1}{2}(n_F - 1)(n_F - 2)$.

In particolare se $n_F = 2$, allora abbiamo a disposizione un solo angolo, che possiamo prendere essere l'angolo di Cabibbo, e zero fasi; mentre se $n_F = 3$, allora abbiamo bisogno di 3 angoli ed una fase e prendiamo la matrice CKM.

La fase nella matrice CKM è legata alla violazione di CP nel Modello Standard.

La scrittura generica della matrice CKM è:

$$V_{CKM} = V_L^{(u)\dagger} V_L^{(d)} = \begin{pmatrix} V_{ud} & V_{us} & V_{ub} \\ V_{cd} & V_{cs} & V_{cb} \\ V_{td} & V_{ts} & V_{tb} \end{pmatrix} \quad (3.7.165)$$

in cui i vari elementi sono misurati negli esperimenti, poichè gli elementi sulla diagonale misurano gli accoppiamento dei quark senza mixing tra famiglie, mentre quelli off-diagonale misurano il mixing. Una possibile parametrizzazione con 3 angoli ed una fase, indicando $\cos\theta_{ij} = c_{ij}$ e $\sin\theta_{ij} = s_{ij}$, è:

$$V_{CKM} = \begin{pmatrix} c_{12}c_{13} & s_{12}c_{13} & s_{13}e^{-i\delta} \\ -s_{12}c_{23} - c_{12}s_{23}s_{13}e^{i\delta} & c_{12}c_{23} - s_{12}s_{23}s_{13}e^{i\delta} & s_{23}c_{13} \\ s_{12}s_{23} - c_{12}c_{23}s_{13}e^{i\delta} & -c_{12}s_{23} - s_{12}c_{23}s_{13}e^{i\delta} & c_{23}c_{13} \end{pmatrix}. \quad (3.7.166)$$

I valori degli elementi sono misurati sperimentalmente da molti processi. Inoltre, l'unitarietà della matrice CKM, che richiede che la somma dei moduli quadri faccia 1, si può verificare:

$$|V_{ud}|^2 + |V_{us}|^2 + |V_{ub}|^2 = 0.9985 \pm 0.0007 \quad (3.7.167)$$

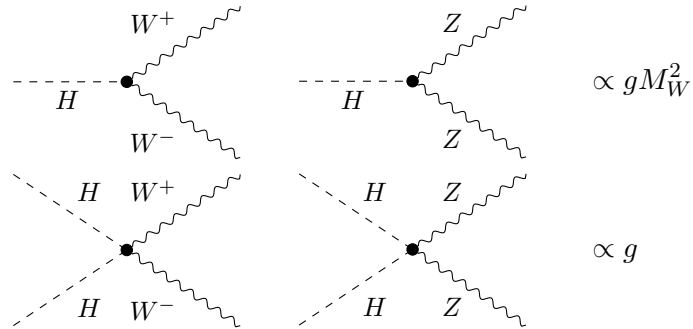
$$|V_{cd}|^2 + |V_{cs}|^2 + |V_{cb}|^2 = 1.001 \pm 0.012 \quad (3.7.168)$$

in cui si ha maggior incertezza sul (3.7.168) a causa della grande incertezza di V_{cs} , poichè difficile da misurare avendo quark che stanno solo nel mare e non dentro protoni o neutroni.

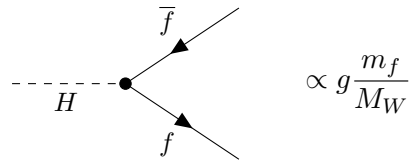
3.8 Accoppiamenti del bosone di Higgs

Per il bosone di Higgs, dalle lagrangiane che abbiamo costruito, possiamo vedere diversi accoppiamenti possibili con le particelle fisiche.

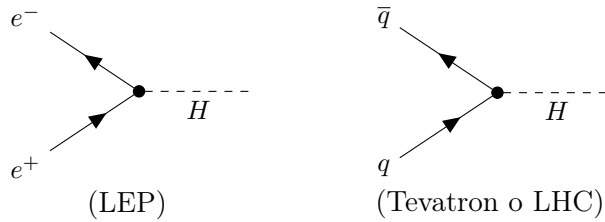
Possiamo avere degli accoppiamenti *Higgs-bosoni* del tipo:



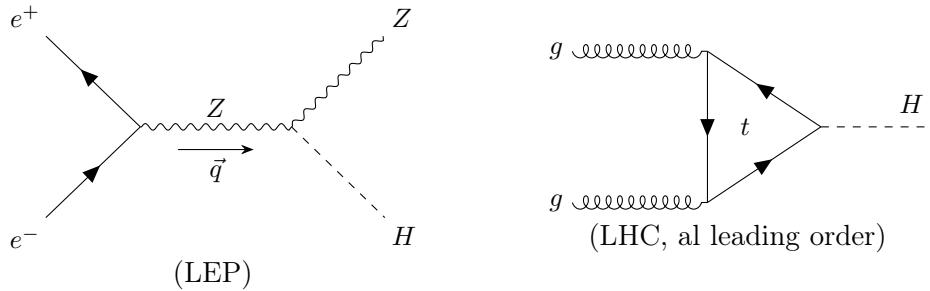
ma potremmo anche avere dei vertici *Higgs-fermioni*:



in cui però l'accoppiamento tra l'Higgs e un fermione di massa $< M_W$ è soppresso. È difficile produrre l'Higgs in uno scattering $f\bar{f}$, del tipo:



proprio poiché soppressi. Nota che LHC è un acceleratore pp , per cui seppur i $\bar{q}$ non sono quark di valenza, ovvero non sono particelle utilizzate direttamente, è possibile in linea teorica avere il vertice disegnato poiché l'antiquark è presente nel mare del protone p . I canali di produzione utilizzati dai vari esperimenti per cercare H sono stati:

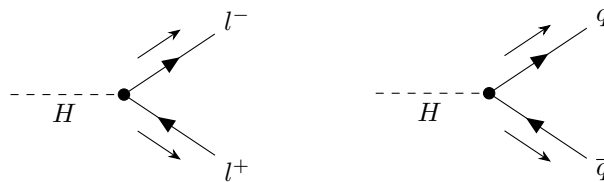


che non sono diagrammi soppressi poiché gli accoppiamenti non sono proporzionali alle masse dei fermioni. Però, non si è osservato l'Higgs a LEP, poiché la massima energia raggiunta (209 GeV) non era sufficiente a produrre H e Z , che sarebbe dovuta essere $s = (M_Z + M_H)^2$. Il diagramma di LHC è sfavorito per colpa del loop, ma è favorito sia dal forte accoppiamento del top t con l'Higgs sia dal fatto che si hanno gluoni, presenti nel protone (mare).

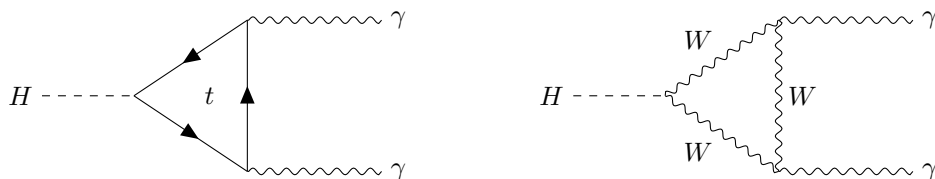
Volendo anche ad LHC si sarebbe potuto cercare lo stesso processo cercato a LEP, ma esso sarebbe sfavorito nello spazio delle fasi essendoci nello stato finale, oltre ad H , anche uno Z .

Da un punto di vista teorico, seppur i diagrammi ad un loop si riescono a calcolare, dunque il diagramma cercato ad LHC, nel momento in cui si vogliono fare conti più precisi, ad ordini superiori, si complicano notevolmente le cose.

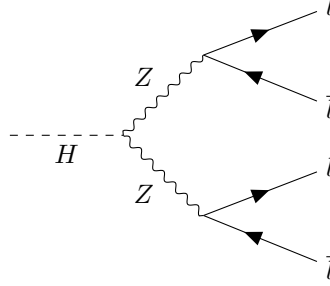
Cosa analoga avviene per i decadimenti:



che sono soppressi, e i cui canali di decadimento negli esperimenti sono:



che però sono soppressi a causa del loop, ma favoriti nello spazio delle fasi e per via della presenza del quark top. Nota, che nel diagramma di destra al posto del loop di $W^\pm$ potrebbe anche esserci il $\varphi^\pm$. Invece, il decadimento:



non ha loop, ma è comunque sfavorito per il suo spazio delle fasi (visto che ci sono molte particelle nello stato finale). Alla fine, per svantaggi di tipo diverso, questi ultimi diagrammi disegnati sono quasi equivalenti.

Ad LHC si misurò la sezione d'urto di:

$$pp \longrightarrow \gamma\gamma$$

e plottandola rispetto $\sqrt{s}$, come in figura 3.2, si è riuscito a trovare il valore risonante, corrispondente alla massa M_H^2 .

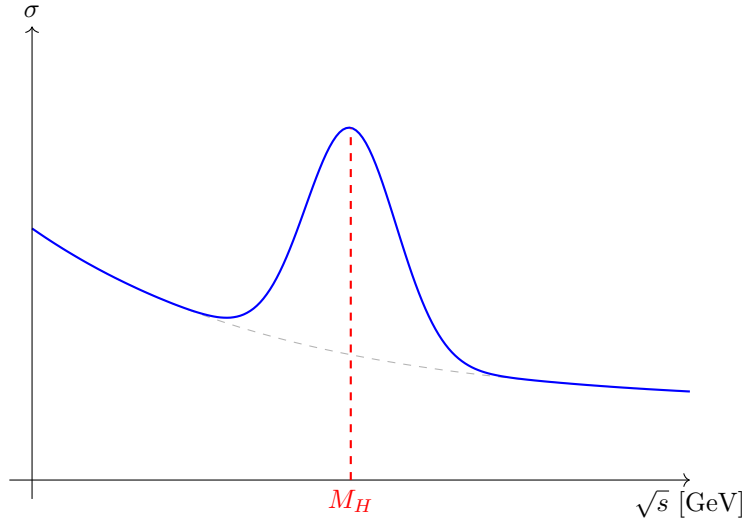


Figura 3.2

3.9 Numero di neutrini leggeri

La misura più accurata del numero di neutrini leggeri ($m_\nu < M_Z/2$, in modo che effettivamente lo Z possa decadere in fermione-antifermione, ovvero due

neutrini) deriva dalla misura della larghezza parziale invisibile Γ_{inv} della produzione di bosoni Z in collisioni e^+e^- (LEP). Il decadimento che si guarda è:

$$e^+ + e^- \longrightarrow Z \longrightarrow f + \bar{f} \quad \begin{array}{c} \mu \xrightarrow{p} \bullet \\ \swarrow \quad \searrow \\ Z \quad Z \\ \searrow \quad \swarrow \\ p_1 \quad p_2 \\ f \quad \bar{f} \end{array} \quad (3.9.1)$$

che in realtà è parte di un processo più grande, ovvero:

$$e^+ + e^- \longrightarrow Z \longrightarrow f + \bar{f}$$

ma nel momento in cui ci si mette con $s \sim M_Z^2$ si può separare la produzione dal decadimento come¹³:

$$\left| \begin{array}{c} e^- \\ \searrow \quad \nearrow \\ \bullet \\ \nearrow \quad \searrow \\ e^+ \end{array} \right|^2 \left(\text{Distribuzione di Breit-Wigner} \right) \left| \begin{array}{c} \searrow \quad \swarrow \\ Z \\ \swarrow \quad \searrow \\ f \quad \bar{f} \end{array} \right|^2.$$

¹³Questo discende dai tagli che si possono fare in un diagramma di Feynman, ed in particolare come si contraggono le gambe esterne dei fermioni. Vedi la sezione §6.1.3. Puoi vedere p.101 del Peskin There is a sort of mythology that grows up about what happened, which is different from what really did happen. Peskin in cui spiega che la distribuzione di Breit-Wigner ci dice che, in Meccanica Quantistica non relativistica, che un atomo instabile si mostra in esperimenti di scattering attraverso delle risonanze. Intorno l'energia di risonanza E_0 l'ampiezza di probabilità è data da:

$$f(E) \propto \frac{1}{E - E_0 + i\Gamma/2} \quad (3.9.2)$$

e la sezione d'urto:

$$\sigma \propto \frac{1}{(E - E_0)^2 + \Gamma^2/4}. \quad (3.9.3)$$

Possiamo calcolare al tree-level (in cui teniamo conto del fattore $1/3$ dato dalle polarizzazioni dello Z) la larghezza del decadimento (3.9.1):

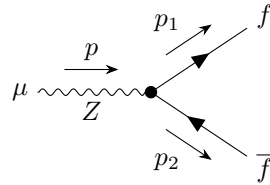
$$\Gamma_{Z \rightarrow f\bar{f}} = \frac{(2\pi)^4}{2M_Z} \int \frac{d^3p_1}{(2\pi)^3 2E_1} \int \frac{d^3p_2}{(2\pi)^3 2E_2} \frac{1}{3} \sum_{spin} |\mathcal{M}|^2 \delta^4(p - p_1 - p_2) \quad (3.9.4)$$

$$= \frac{(2\pi)^{4-6}}{2M_Z} \frac{1}{12} \int \frac{d^3p_1}{E_1 E_2} \sum_{spin} |\mathcal{M}|^2 \delta(M_Z - E_1 - E_2) \quad (3.9.5)$$

$$= \frac{(2\pi)^{-2}}{24M_Z} \int dE_1 E_1^2 \int \frac{d\Omega_1}{4\pi} \frac{1}{E_1^2} \sum_{spin} |\mathcal{M}|^2 \underbrace{\delta(M_Z - 2E_1)}_{\frac{1}{2}\delta\left(E_1 - \frac{M_Z}{2}\right)} \quad (3.9.6)$$

$$= \frac{1}{48\pi} \frac{1}{M_Z} \sum_{spin} |\mathcal{M}|^2. \quad (3.9.7)$$

Possiamo calcolare $\sum_{spin} |\mathcal{M}|^2$ se ricaviamo l'elemento di matrice, dato il contributo del vertice:



$$= i \frac{g}{2 \cos \theta_W} \gamma^\mu (g_V - g_A \gamma_5) \quad (3.9.8)$$

$$\Rightarrow i\mathcal{M} = i \frac{g}{2 \cos \theta_W} \bar{u}(p_1) \gamma^\mu (g_V - g_A \gamma_5) v(p_2) \epsilon_\mu^* \quad (3.9.9)$$

in cui g_V e g_A dipendono dai fermioni f e $\bar{f}$, ed ϵ_μ è il vettore di polarizzazione del bosone Z , per cui vale:

$$\sum_{spin} \epsilon_\mu^* \epsilon_\nu = -g_{\mu\nu} + \frac{p_\mu p_\nu}{M_Z^2}. \quad (3.9.10)$$

In più dobbiamo notare che nel sistema di riferimento di riposo di Z , cioè in cui $\vec{p} = 0$, si ha:

$$\vec{p}_2 = -\vec{p}_1 \quad (3.9.11)$$

e per fermioni mass-less:

$$E_2 = |\vec{p}_2| = |\vec{p}_1| = E_1 \quad (3.9.12)$$

per cui:

$$M_Z = E_1 + E_2 = 2E_1. \quad (3.9.13)$$

Abbiamo:

$$\sum_{spin} |\mathcal{M}|^2 = \frac{g^2}{4 \cos^2 \theta_W} \text{Tr} \left\{ \not{p}_1 \gamma^\mu (g_V - g_A \gamma_5) \not{p}_2 \gamma^\nu (g_V - g_A \gamma_5) \right\} \sum_{spin} \epsilon_\mu^* \epsilon_\nu \quad (3.9.14)$$

$$= \frac{g^2}{4 \cos^2 \theta_W} \text{Tr} \left\{ \not{p}_1 \gamma^\mu \not{p}_2 \gamma^\nu (g_V - g_A \gamma_5)^2 \right\} \left(-g_{\mu\nu} + \frac{p_\mu p_\nu}{M_Z^2} \right) \quad (3.9.15)$$

$$= \frac{g^2}{4 \cos^2 \theta_W} \text{Tr} \left\{ \not{p}_1 \gamma^\mu \not{p}_2 \gamma^\nu (g_V^2 + g_A^2 - 2g_V g_A \gamma_5) \right\} \left(-g_{\mu\nu} + \frac{p_\mu p_\nu}{M_Z^2} \right) \quad (3.9.16)$$

$$= \frac{g^2}{\cos^2 \theta_W} \left[(p_1^\mu p_2^\nu + p_2^\mu p_1^\nu - p_1 p_2 g^{\mu\nu}) (g_V^2 + g_A^2) + \underbrace{2g_V g_A i \epsilon^{\alpha\mu\beta\nu} p_{1\alpha} p_{2\beta}}_{\text{antisimmetrica in } \mu, \nu} \right] \underbrace{\left(-g_{\mu\nu} + \frac{p_\mu p_\nu}{M_Z^2} \right)}_{\text{simmetrica in } \mu, \nu} \quad (3.9.17)$$

$$= \frac{g^2}{\cos^2 \theta_W} \left[-(2-4)p_1 p_2 + \frac{2(p_1 p)(p_2 p) - (p_1 p_2)p^2}{M_Z^2} \right] (g_V^2 + g_A^2) \quad (3.9.18)$$

$$= \frac{g^2}{\cos^2 \theta_W} \left[2p_1 p_2 + 2 \frac{(p_1 p_2)^2}{M_Z^2} - p_1 p_2 \right] \quad (3.9.19)$$

$$= \frac{g^2}{\cos^2 \theta_W} M_Z^2 (g_V^2 + g_A^2) \quad (3.9.20)$$

dove abbiamo utilizzato le relazioni:

$$\begin{aligned} p_1 p &= p_1^2 + p_1 p_2 = p_1 p_2 \\ p_2 p &= p_1 p_2 + p_2^2 = p_1 p_2 \\ p_1 p_2 &= \frac{1}{2}(p_1 + p_2)^2 = \frac{p^2}{2} = \frac{M_Z^2}{2}. \end{aligned}$$

Dunque abbiamo trovato:

$$\Gamma_{Z \rightarrow f \bar{f}} = \frac{1}{4g\pi} \frac{1}{M_Z} \frac{g^2}{\cos^2 \theta_w} M_Z^2 (g_V^2 + g_A^2) \quad (3.9.21)$$

$$= \frac{G_F M_Z^3}{6\sqrt{2}\pi} (g_V^2 + g_A^2) \quad (3.9.22)$$

in cui abbiamo utilizzato:

$$\frac{G_F}{\sqrt{2}} = \frac{g^2}{8M_W^2} \quad , \quad \frac{g^2}{\cos^2 \theta_W} = \frac{g^2}{M_W^2} M_Z^2 = \frac{8}{\sqrt{2}} G_F M_Z^2. \quad (3.9.23)$$

Possiamo indicare le misure sperimentali che si possono fare:

- Γ_{had} la larghezza di decadimento adronico ($Z^0 \rightarrow q\bar{q}$).
- Γ_l la larghezza di decadimento in leptoni carichi (supponendo *lepton universality*, cioè che c'è una larghezza di decadimento comune in leptoni carichi). Ricordiamo che sono presenti 3 famiglie (dunque sarà presente un coefficiente numerico).
- Γ_{inv} la larghezza di decadimento invisibile.

per cui possiamo scrivere:

$$\Gamma_Z = \Gamma_{had} + 3\Gamma_l + \Gamma_{inv} \implies \Gamma_{inv} = \Gamma_Z - \Gamma_{had} - 3\Gamma_l \quad (3.9.24)$$

e assumiamo che Γ_{inv} sia dovuta a tutti i neutrini leggeri del Modello Standard, dal fatto che i neutrini non si rivelano il nome di invisibile, dunque:

$$\Gamma_{inv} = N_\nu \cdot \Gamma_\nu. \quad (3.9.25)$$

Le larghezze di decadimento parziali in fermioni sono date da:

$$\Gamma_f = \Gamma_0 \cos \theta_f (g_V^2 + g_A^2) \quad (3.9.26)$$

$$\text{con } \Gamma_0 = \frac{G_F m_Z^2}{6\sqrt{2}\pi} = 0.332 \text{ GeV} \quad (3.9.27)$$

in cui c_f è il fattore di colore, che vale $c_f = 1$ per i leptoni e $c_f = 3$ per i quark, e in più abbiamo:

$$g_V = I_3 - 2Q \sin^2 \theta_W, \quad g_A = I_3 \quad (3.9.28)$$

con Q carica elettrica e valori:

$$\begin{array}{lll} \nu : & g_V = +\frac{1}{2} & g_A = +\frac{1}{2} \\ l : & g_V = -\frac{1}{2} + 2 \sin^2 \theta_W & g_A = -\frac{1}{2} \\ u : & g_V = +\frac{1}{2} - \frac{4}{3} \sin^2 \theta_W & g_A = +\frac{1}{2} \\ d : & g_V = -\frac{1}{2} + \frac{2}{3} \sin^2 \theta_W & g_A = -\frac{1}{2}. \end{array}$$

Per ciascuno dei canali (adronico, leptonic e invisibile), le predizioni del Modello Standard al tree-level, sono:

$$\Gamma(Z^0 \rightarrow u\bar{u}, c\bar{c}) = \left(\frac{3}{2} - 4 \sin^2 \theta_W + \frac{16}{3} \sin^4 \theta_W \right) \Gamma_0 = 0.286 \text{ GeV}$$

$$\Gamma(Z^0 \rightarrow d\bar{d}, s\bar{s}, b\bar{b}) = \left(\frac{3}{2} - 2 \sin^2 \theta_W + \frac{4}{3} \sin^2 \theta_W \right) \Gamma_0 = 0.369 \text{ GeV}$$

$$\Gamma(Z^0 \rightarrow e^+e^-, \mu^+\mu^-, \tau^+\tau^-) = \left(\frac{1}{2} - 2 \sin^2 \theta_W + 4 \sin^2 \theta_W \right) \Gamma_0 = 0.084 \text{ GeV}$$

$$\Gamma(Z^0 \rightarrow \nu\bar{\nu}) = \frac{1}{2} \Gamma_0 = 0.166 \text{ GeV}.$$

Sommando tutte le ampiezze dei quark si trova $\Gamma_{had} = 1.678$ GeV. Per ridurre la dipendenza del modello teorico conviene considerare rapporti tra le larghezze di decadimento:

$$\Gamma_{inv} = \Gamma_Z - \Gamma_{had} - 3\Gamma_l \quad \Rightarrow \quad \frac{\Gamma_{inv}}{\Gamma_l} = \frac{\Gamma_Z}{\Gamma_l} - \frac{\Gamma_{had}}{\Gamma_l} - 3. \quad (3.9.29)$$

Mettendo un pedice SM alle lunghezze teoriche predette dal Modello Standard, e non mettendo nulla nei valori sperimentali, possiamo vedere che il numero N_ν di neutrini leggeri nel Modello Standard è dato da:

$$N_\nu = \frac{\Gamma_{inv}}{\Gamma_l} \left(\frac{\Gamma_l}{\Gamma_\nu} \right)_{SM} \quad (3.9.30)$$

$$= \left[\frac{\Gamma_Z}{\Gamma_l} - \frac{\Gamma_{had}}{\Gamma_l} - 3 \right] \left(\frac{\Gamma_l}{\Gamma_\nu} \right)_{SM} \quad (3.9.31)$$

che è una forma scelta in modo da diminuire la dipendenza dal modello. Si ottiene:

$$\left(\frac{\Gamma_l}{\Gamma_\nu} \right)_{SM} = 1.99060 \pm 0.00021 \quad (3.9.32)$$

dal calcolo teorico più preciso (Janot-Jadach, 2020, che migliorarono la predizione teorica aumentando il numero di loop), da cui:

$$N_\nu = 2.9963 \pm 0.0074 \quad (3.9.33)$$

in ottimo accordo con l'aspettativa $N_\nu = 3$. Questo non vuol dire che non ci siano altre famiglie leptoniche, ma se ci sono i neutrini associati hanno massa maggiore a $M_Z/2$.

Mettendo la parte hadronica rilevante solamente dal punto di vista sperimentale, poiché nella parte del SM c'è solo la parte leptonica e invisibile, si è ridotta leggermente la dipendenza del conto dal modello (infatti la predizione può essere differente se il modello considerato è differente dal Modello Standard usuale).

Cromodinamica quantistica

Capitolo 4

Introduzione storica

Vedremo in questa parte di corso, come dice il titolo, la trattazione della Cromodinamica Quantistica. In particolare, dopo un'introduzione storica, analizzeremo i suoi fondamenti, la sua struttura matematica, le proprietà caratteristiche delle teorie di gauge non abeliane, quali lo strong coupling e la libertà asintotica. Successivamente vedremo processi di scattering, le divergenze IR associate, i jets adronici e le sezioni d'urto differenziali in teoria delle perturbazioni. Verso la fine analizzeremo il deep inelastic scattering e i di-jet production. Per concludere vedremo la cancellazione delle anomalie nel Modello Standard.

Cominciamo il primo capitolo con una rapida introduzione storica e delle evidenze/necessità che hanno portato alla costruzione della Cromodinamica Quantistica (QCD). Si può seguire il capitolo §17 del Peskin There is a sort of mythology that grows up about what happened, which is different from what really did happen. Peskin.

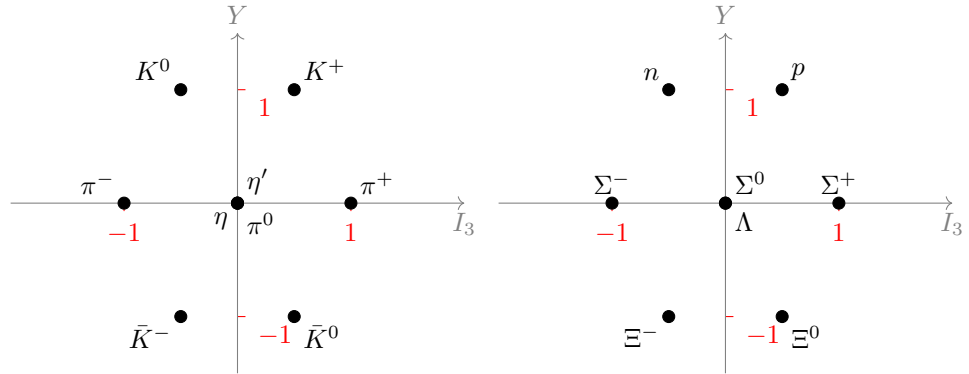
4.1 Modello a quark

Come già accennato nella parte introduttiva §1, abbiamo visto che intorno gli anni '60 cominciarono ad essere rivelate, grazie ai primi collider nel mondo, moltissime particelle adroniche, interagenti tramite la forza forte, che portarono alla preoccupazione di non riuscire a creare un modello di catalogazione di tutte le particelle in natura. Successivamente ci si rese conto che tutte queste nuove particelle in realtà non fossero elementari, bensì che fossero formate da costituenti più piccoli, dunque elementari, e vennero classificati in multipletti. Puoi guardare la figura 4.1. Infatti, la soluzione (in realtà furono due versioni quasi contemporanee) di Gell-Mann e Zweig per giustificare questa struttura ordinata è stato il **modello a quark**, il quale "costruiva" i cosiddetti *mesoni* e *barioni* come degli stati legati di particelle fondamentali chiamate *quark*. Inizialmente i quark fondamentali furono solamente 3 (i più

leggeri):

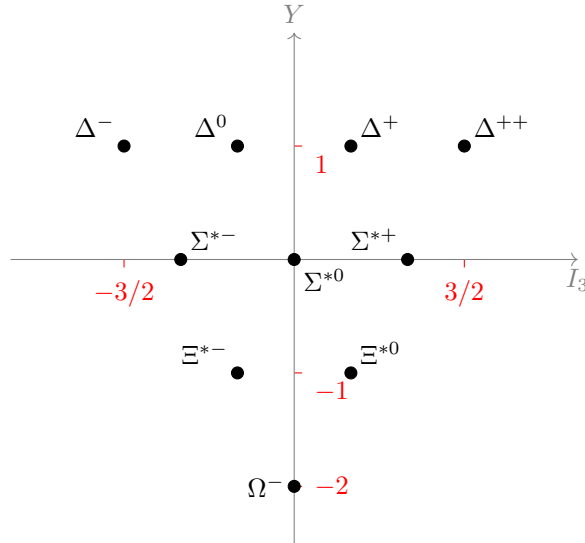
$$u, d, s \quad (4.1.1)$$

e successivamente furono organizzati in 3 famiglie, come accade per i leptoni.



(a) Ottetto di mesoni di spin 0 con masse simili.

(b) Ottetto di barioni di spin 1/2 con masse simili.



(c) Decupletto di barioni di spin 3/2 con masse simili.

Figura 4.1: Multipletti di adroni.

Analogamente al settore leptonic i quark furono postulati con diverso *flavour*, ovvero di diverso tipo (il quark u ha un flavour diverso rispetto il quark d). In particolare, per il modello di Gell-Mann e Zweig si ha:

$$\text{mesoni} = q\bar{q}' \quad (4.1.2)$$

in cui q e q' sono di flavour diverso. Alcuni esempi sono:

$$K^+ = u\bar{s} \quad , \quad K^0 = d\bar{s} \quad , \quad \pi^+ = u\bar{d}.$$

In modo simile si ha:

$$\text{barioni} = qq'q'' \quad (4.1.3)$$

dove ora i tre quark possono anche essere con stesso flavour. Alcuni esempi sono:

$$\begin{aligned} n &= udd \quad , \quad p = uud \\ \Sigma^+ &= uus \quad , \quad \Delta^- = ddd \\ \Delta^{++} &= uuu. \end{aligned}$$

Storicamente si fu' molto entusiasti, dal punto di vista teorico, del modello a quark, però ci furono diversi dubbi riguardo la sua effettiva correttezza. Infatti, passarono diversi anni prima della evidenza sperimentale, e sembrava che il modello fosse più un artificio matematico per caratterizzare le particelle.

Notiamo inoltre, che nonostante i dubbi, l'entusiasmo teorico era dato dal sorgere, dalla costruzione in multipletto, di una simmetria $SU(3)$ globale di flavour. In questo modello si potevano quindi interpretare le tre specie (flavour) di quark come 3 gradi di libertà della teoria, regolati dalle trasformazioni di $SU(3)$ (come spin up e down sono gradi di libertà di $SU(2)$). Fatto questo, i multipletti di Isospin¹ sono delle rappresentazioni del gruppo globale $SU(3)$ di sapore.

Successivamente ci si rese conto che la simmetria $SU(3)$ di flavour era un po' approssimativa, visto che i tre quark u, d, s hanno masse diverse, oltre il fatto che si possano formare stati legati, quali barioni e mesoni, anche utilizzando altre specie. La simmetria $SU(3)_{\text{flavour}}$ è solo un modo per organizzare le particelle adroniche, ma non è in alcun modo dinamica, cioè non genera alcuna forza. Per avere una teoria per la forza forte di tipo dinamico si deve introdurre il gruppo $SU(3)_{\text{colore}}$, la quale è una simmetria esatta, di gauge e di tipo dinamico (contiene 8 gluoni che mediano la forza).

4.2 Dubbi e introduzione del colore

Un forte dubbio del modello era dato dal fatto, che per rispettare le cariche elettriche sperimentali delle particelle (mesoni e barioni), sarebbe stato necessario che i quark avessero avuto carica elettrica frazionaria, in particolare:

$$u = \frac{2}{3} \quad , \quad d = -\frac{1}{3}. \quad (4.2.1)$$

¹I multipletti che abbiamo scritto (come l'ottetto barionico) contengono i multipletti di isospin. Ad esempio, nell'ottetto dei barioni, il protone e il neutrone formano un "sottomultipletto" di isospin all'interno di una struttura più vasta basata sul sapore.

Questa cosa era al quanto strana, poiché si pensava che la carica elettrica elementare, dunque indivisibile, fosse quella dell'elettrone. Questo scetticismo fu' amplificato dai numerosi tentativi sperimentali, fallimentari, di verificare questa carica frazionaria.

Un altro dubbio, un po' più fondamentale in un certo senso, è dato dallo stato legato Δ^{++} . Era considerato leggermente più grave della mancata evidenza sperimentale, poiché quest'ultima poteva essere ritenuta solamente una questione di tempo; però la particella Δ^{++} aveva il problema di rompere il teorema spin-statistica. Infatti, essendo Δ^{++} un fermione, ci si aspettava una funzione d'onda ψ anti-simmetrica; però essendo formato da 3 quark u Δ^{++} aveva la parte $\psi_{flavour}$ simmetrica per scambio, ma non solo, perché essendo gli spin dei tre quark allineati, anche la parte ψ_{spin} è simmetrica per scambio, e non è ancora finita, poiché il momento angolare orbitale è uguale a 0, dunque $\psi_{orbitale}$ è anch'essa simmetrica. Dunque, tutta la funzione d'onda di ψ , con tutti i suoi numeri quantici, crea il problema di rompere il teorema di spin-statistica.

La rottura dell'antisimmetria della funzione d'onda di un fermione venne risolta, in modo un po' grezzo, aggiungendo un grado di libertà, mai stato necessario fino a quel momento. Il ragionamento dietro è stato, circa, che non volendo rinunciare all'antisimmetria dei fermioni, nel momento in cui vediamo una ψ simmetrica, vuol dire che ci stiamo perdendo qualcosa per strada che ci permette di tenere valido lo spin-statistica. Viene per questo introdotto il **colore**. Dunque, questo nuovo numero quantico diventa necessario, con l'introduzione del modello a quark, così da poter antisimmetrizzare la ψ .

In particolare, per poter assegnare questi nuovi gradi di libertà ai quark, si ipotizzò che ci fosse una simmetria (interna ai quark) $SU(3)_c$ di colore (differente da quella di flavour), così che essi potessero essere visti come multipletti di $SU(3)_c$. Si ha che i quark appartengono alla rappresentazione fondamentale (rappresentazione 3) di $SU(3)_c$, ovvero:

$$q^i \quad i = 1, 2, 3 \quad (4.2.2)$$

in cui q ci dice il flavour, mentre i è l'indice di colore. Facendo così, i quark li abbiamo come tripletti di colore, ovvero esistono 3 stati di colore a fissato flavour, e possiamo scrivere:

$$\begin{pmatrix} q^1 \\ q^2 \\ q^3 \end{pmatrix}. \quad (4.2.3)$$

Gli anti-quark appartengono alla rappresentazione antifondamentale (solitamente indicata con un indice basso e nominata $\bar{3}$):

$$\bar{q}_i \quad i = 1, 2, 3. \quad (4.2.4)$$

Il motivo per cui fu' preso proprio $SU(3)$ e non gruppi più grandi² è dettato dal fatto che esso è la scelta più semplice (minimale) per poter vedere gli stati fisici (adroni), come dei singoletti di esso (ossia non colorati), anche se formati da quark *colorati*. Essere un singoletto, ricordiamo, vuol dire che non devono trasformare sotto trasformazioni di $SU(3)_c$, per cui dobbiamo sempre avere gli indici di colore contratti (invarianti). Infatti, il motivo per cui $SU(3)_c$ è stato introdotto è solo per rendere la funzione d'onda antisimmetrica sotto scambio di colore, ma non è stato aggiunto per motivi sperimentali. Vedremo tra pochissimo però come esso sia legato alla libertà asintotica.

Il nostro problema di simmetria sorge nel momento in cui guardiamo i barioni, e per questo vogliamo che scambiando l'indice di colore esso risulti, in qualche modo, antisimmetrico.

I mesoni non ci danno alcun tipo di problema, per cui li vogliamo, non solo singoletti di colore, ma anche simmetrici per scambio. Se prendiamo lo stato:

$$\bar{q}_i q'^i \quad (4.2.5)$$

è un singoletto di colore, come possiamo vedere facendo:

$$SU(3) : \quad q^i \longrightarrow U^i_j q^j \quad (4.2.6)$$

che lascia invariato (4.2.5). Per i barioni, il modo più semplice che abbiamo, per contrarre tutti gli indici, e rendere lo stato antisimmetrico per scambio di colore, è:

$$q^i q^j q^k \epsilon_{ijk} \quad (4.2.7)$$

$$\bar{q}_i \bar{q}_j \bar{q}_k \epsilon^{ijk} \quad (4.2.8)$$

in cui il tensore di Levi-Civita ci permette, non solo di creare singoletti, ma anche oggetti completamente antisimmetrici per scambio.

Sicuramente ci sono gruppi più complicati per costruire singoletti di colore e barioni antisimmetrici, ma la scelta minimale è $SU(3)$.

4.3 Libertà asintotica e teoria di gauge

Il punto però è che la Fisica è una scienza sperimentale, per cui si dovrebbe spiegare il colore anche da un punto di vista dinamico, oltre che quello tecnico (teorico). Occorrerebbe capire cosa sia effettivamente il colore, ma soprattutto capire quale sia il meccanismo dinamico che fa sì che gli adroni siano singoletti di colore, ovvero quale sia il meccanismo per cui gli unici stati asintotici siano singoletti di colore e non triplette (fenomeno del confinamento dei quark).

²Vedremo nel corso del capitolo anche delle evidenze sperimentali, legate al numero di colore, che ci giustificheranno la scelta di $SU(3)_c$ rispetto altri gruppi.

Con i collisori si cominciò a sondare la struttura del protone, tramite *deep inelastic scattering*, per cui si "sparavano" elettroni contro protoni, in modo da provare a vederne la struttura interna, che si sperava essere formata da quark. Ciò che si vide fu' che ad alte energie la struttura di p era composta da moltissime particelle libere, debolmente interagenti, mentre a basse la struttura interna era formata da particelle fortemente legate. L'andamento della costante di accoppiamento, dunque, ha un andamento rappresentato in figura 4.2. Questo andamento viene chiamato **libertà asintotica**. In più, il fatto che ad alte energie si potesse utilizzare la teoria perturbativa, vista la piccolezza dell'accoppiamento, portò alla ricerca per costruire collider sempre più energetici e al voler svolgere calcoli teorici ad ordine perturbativo sempre più alto.

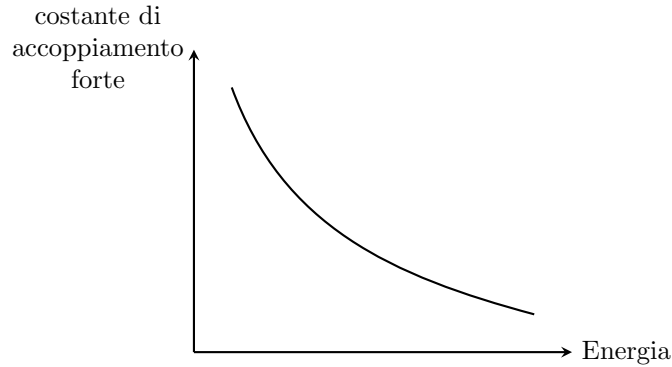


Figura 4.2: Rappresentazione della costante di accoppiamento della forza nucleare forte, che mostra il comportamento di libertà asintotica.

Questa evidenza sperimentale fu' una felice scoperta poiché negli stessi anni si scoprì che le teorie di gauge non abeliane possiedono esattamente questa proprietà di libertà asintotica. Dunque, la simmetria $SU(3)$ globale "fittizia", ipotizzata per sistemare il problema dello spin-statistica, si capì poter essere promossa a *simmetria di gauge per le interazioni forti*:

$$SU(3)_{\text{globale colore}} \longrightarrow SU(3)_{\text{gauge colore}} : QCD.$$

Fatta questa promozione otteniamo che le interazioni forti sono una teoria di gauge non abeliana, che dunque ha la proprietà di libertà asintotica verificata sperimentalmente, di gruppo $SU(3)$. Dunque, si diede un significato dinamico al colore, che non è più un grado di libertà aggiuntivo introdotto ad hoc, ma è proprio un grado di libertà caratteristico delle interazioni forti. Sappiamo anche che, come tutti i gruppi di gauge, ci sono dei bosoni mediatori, in questo caso chiamati **gluoni**.

Inoltre, negli anni successivi Wilson mostrò la *proprietà di confinamento* del colore, ovvero che gli unici stati asintotici (di energia non zero) sono stati

di singoletti di colore. Per questo se prendiamo un mesone $q\bar{q}$ e proviamo a separare i due quark, man mano che li allontaniamo si crea un tubo di flusso di gluoni (che collegano i due quark), il quale, una volta spezzato, non restituisce due quark separati, ma due coppie $q\bar{q}$ (essendo favorito dal punto di vista energetico).

Capitolo 5

Fondamenti di QCD

Vediamo in questo capitolo l'inizio della trattazione della Cromodinamica Quantistica.

5.1 Il gruppo $SU(N)$

Prima di tuffarci nella trattazione della QCD ripassiamo brevemente il gruppo non abeliano $SU(N)$. Si può fare riferimento al capitolo §15.4 del Peskin. There is a sort of mythology that grows up about what happened, which is different from what really did happen. Peskin.

$SU(N)$ è il gruppo di matrici unitarie $N \times N$ con $\det = 1$. Per ogni $U \in SU(N)$, si ha:

$$U_{ij} = e^{i\alpha^a t_{ij}^a}, \quad i, j = 1, \dots, N \quad (5.1.1)$$

dove si intende la somma sugli indici ripetuti (a in questo caso). Le matrici t_{ij}^a sono i generatori di $SU(N)$ nella rappresentazione fondamentale e α^a sono parametri, che caratterizzano la trasformazione.

Dalle caratteristiche delle matrici U discendono delle proprietà dei generatori t^a . Infatti, affinché U sia unitaria, t^a deve essere hermitiana:

$$U^\dagger U = e^{-i\alpha^a (t^a)^\dagger} e^{i\alpha^a t^a} = \mathbb{1} \implies (t^a)^\dagger = t^a. \quad (5.1.2)$$

Affinché U abbia $\det = 1$, i generatori t^a devono essere a traccia nulla:

$$\det(U) = \det(e^{i\alpha^a t^a}) = e^{i\alpha^a \text{tr}(t^a)} = 1 \implies \text{tr}(t^a) = t_{ii}^a = 0. \quad (5.1.3)$$

Però ci si potrebbe chiedere quanti siano i generatori del gruppo. I t^a sono le matrici che generano lo spazio vettoriale delle matrici complesse $N \times N$, il quale ha dimensione $2N^2$. Lo spazio delle matrici hermitiane complesse $N \times N$ ha dimensione N^2 , e lo spazio delle matrici hermitiane a traccia nulla ha dimensione $N^2 - 1$. Ciò significa che ci sono $N^2 - 1$ generatori indipendenti

t^a , ovvero $a = 1, \dots, N^2 - 1$, che è la dimensione della rappresentazione aggiunta del gruppo $SU(N)$.

Notiamo che in (5.1.1) gli indici i, j sono indici della rappresentazione fondamentale ($i, j = 1, 2, 3$), mentre a è un'indice che conta il numero di generatori (nella rappresentazione aggiunta, dunque $a = 1, \dots, N^2 - 1$). Se nel nostro modello di $SU(3)$, prendiamo quark e anti-quark appartenenti rispettivamente alle rappresentazioni fondamentale e anti-fondamentale, entrambe di dimensione N , allora è immediato vedere abbiamo 3 colori diversi di q e $\bar{q}$. Però, se consideriamo i gluoni appartenenti alla rappresentazione aggiunta, di dimensione $N^2 - 1 = 8$, allora ci rendiamo conto che ci sono gluoni a 8 colori differenti.

Sappiamo anche che i generatori soddisfano la relazione di commutazione dell'algebra di Lie:

$$[t^a, t^b] = if^{abc}t^c \quad (5.1.4)$$

dove f^{abc} sono le funzioni di struttura di $SU(N)$. Esse sono uguali a (diverse da) zero per gruppi di Lie abeliani (non-abeliani); per $SU(N)$, la loro presenza (o meno) codifica appieno la natura non-abeliana (abeliano) del gruppo. Senza perdita di generalità, per $SU(N)$ queste possono essere scelte reali e completamente antisimmetriche sotto lo scambio di qualsiasi coppia di indici:

$$f^{abc} = -f^{bac} = -f^{acb} = f^{cab} = \dots \quad (5.1.5)$$

Le costanti di struttura con due o più indici uguali sono quindi nulle:

$$f^{aac} = f^{aba} = f^{abb} = f^{aaa} = 0. \quad (5.1.6)$$

Possiamo prendere le f^{abc} reali e completamente antisimmetriche senza perdere di generalità poiché la fisica si basa sulle sue proprietà e non sulle sue esplicite rappresentazioni.

Le costanti di struttura forniscono i generatori dell'algebra nella rappresentazione aggiunta. Denotiamo con t^a i generatori della rappresentazione fondamentale, di dimensione N , mentre scriviamo t_G^a i generatori della rappresentazione aggiunta, di dimensione $N^2 - 1$. L'etichetta ' G ' nell'espressione di t_G richiama il fatto che la rappresentazione aggiunta è quella in cui "vivono" i gluoni. Infatti, definendo:

$$(t_G^b)_{ac} \equiv if^{abc} \quad (5.1.7)$$

in cui a, c corrono su $1, \dots, N^2 - 1$ cioè sulla dimensione dell'aggiunta, possiamo mostrare che le costanti di struttura soddisfano le relazioni definenti per i generatori (in questo caso della rappresentazione aggiunta, visti gli indici) di $SU(N)$, cioè:

$$(t_G^b)_{ac}^\dagger = (t_G^b)_{ca}^* = -i(f^{cba})^* = -if^{cba} = if^{abc} = (t_G^b)_{ac} \quad (5.1.8)$$

$$\text{Tr}\{t_G^b\} = if^{aba} = 0 \quad (5.1.9)$$

$$[t_G^a, t_G^b]_{cd} = if^{cak}if^{khd} - if^{cbk}if^{kad} = if^{abk}if^{ckd} = if^{abk}(t_G^k)_{cd} \quad (5.1.10)$$

dove per l'ultima equazione abbiamo usato l'identità di Jacobi. Dunque, le costanti di struttura f^{abc} sono una buona rappresentazione dei generatori dell'aggiunta, e possono rappresentare i gluoni.

Possiamo anche calcolare (definire) i cosiddetti **fattori quadratici di Casimir**, definiti dalle tracce dei prodotti di coppie di generatori. Il Casimir quadratico della rappresentazione fondamentale, T_R , è definito come:

$$\text{Tr}\{t^a t^b\} = t_{ij}^a t_{ji}^b = T_R \delta^{ab}, \quad T_R = \frac{1}{2}. \quad (5.1.11)$$

Il valore di $1/2$ è una scelta convenzionale e determina la normalizzazione dei generatori nella rappresentazione fondamentale, come $t^a = T_R \lambda^a$, dove nel caso $SU(3)$ le λ^a sono le otto matrici 3×3 di Gell-Mann. Notiamo che la proprietà (5.1.11) discende dalle proprietà di trasformazione del gruppo, infatti trasformando (nell'aggiunta):

$$t^a \longrightarrow U t^a U^\dagger$$

la traccia è invariante per essa, oltre che per scambio di a e b ; per cui l'unica forma tensoriale che permette di rispettare tale simmetria è una costante moltiplicata per la delta δ^{ab} .

Il Casimir quadratico della rappresentazione aggiunta, indicato con C_A , è definito tramite:

$$\text{Tr}\{t_G^a t_G^b\} = i f^{cak} i f^{kbc} \equiv C_A \delta^{ab}, \quad C_A = N. \quad (5.1.12)$$

Il valore numerico è $C_A = 3$ nel caso della QCD. L'etichetta 'A' richiama ancora i gluoni (tipicamente, i campi quantistici dei gluoni sono indicati con A_μ^a , analogamente ai fotoni). In effetti, C_A è il fattore di Casimir ottenuto da due generatori t_G nella rappresentazione aggiunta di $SU(N)$.

Le relazioni precedenti permettono di riscrivere le funzioni di struttura in termini dei generatori della rappresentazione fondamentale. Infatti possiamo calcolare:

$$\text{Tr}\{[t^a, t^b] t^c\} = \text{Tr}\{i f^{abd} t^d t^c\} \quad (5.1.13)$$

$$= i f^{abd} \text{Tr}\{t^d t^c\} \quad (5.1.14)$$

$$= i f^{abd} T_R \delta^{dc} \quad (5.1.15)$$

$$= i f^{abc} T_R \quad (5.1.16)$$

dunque le f^{abc} non sono completamente indipendenti dai generatori t^a e possiamo scrivere:

$$f^{abc} = -\frac{i}{T_R} \text{Tr}\{[t^a, t^b] t^c\}. \quad (5.1.17)$$

Il terzo fattore di Casimir, C_F , è definito come:

$$(t^a t^a)_{ij} \equiv C_F \delta_{ij} \quad (5.1.18)$$

dove ancora una volta è sottintesa la somma sugli indici ripetuti (in questo caso a) e il pedice F sta a significare "fermioni" poiché C_F riguarda la rappresentazione fondamentale, a cui appartengono i quark, che sono per l'appunto fermioni. Il valore di C_F si deduce direttamente (ricordando che δ^{aa} fornisce la dimensione della dimensionalità dell'aggiunta $= N^2 - 1$, mentre δ_{ii} la dimensionalità della fondamentale $= N$):

$$\text{tr}(t^a t^a) \begin{cases} = T_R \delta^{aa} = T_R(N^2 - 1) \\ = \text{Tr}\{C_F \delta_{ij}\} = C_F \delta_{ii} = C_F N \end{cases} \quad (5.1.19)$$

$$\implies C_F = T_R \frac{N^2 - 1}{N} = \frac{N^2 - 1}{2N}. \quad (5.1.20)$$

Il suo valore numerico è $C_F = 4/3$ nel caso di $SU(3)$, rilevante per le interazioni forti. Si noti inoltre che nel limite $N \rightarrow \infty$ (noto come limite del colore dominante o *leading-colour limit*), si ha $C_F \rightarrow N/2 = C_A/2$. In realtà, per $N = 3$, $C_F = C_A/2$ solo fino a circa il 10%. Delle volte, per semplificare i conti, si considera il limite ad N grande, ma consci del fatto che si sta facendo un'errore del circa 10%.

Un'altra relazione di notevole utilità pratica è l'**identità di Fierz**:

$$t_{ij}^a t_{kl}^a = T_R \left(\delta_{il} \delta_{jk} - \frac{1}{N} \delta_{ij} \delta_{kl} \right) = \frac{1}{2} \left(\delta_{il} \delta_{jk} - \frac{1}{N} \delta_{ij} \delta_{kl} \right). \quad (5.1.21)$$

Possiamo dimostrare questa identità considerando una generica matrice complessa hermitiana $N \times N$ chiamata M :

$$M_{ij} = m_0 \delta_{ij} + m_a t_{ij}^a, \quad i, j = 1, \dots, N \quad (5.1.22)$$

la quale, non essendo trace-less genera lo spazio N^2 dimensionale. I coefficienti di M_{ij} possono essere determinati calcolando le tracce:

$$\text{Tr}\{M\} = m_0 N \implies m_0 = \frac{1}{N} \text{Tr}\{M\}, \quad (5.1.23)$$

$$\text{Tr}\{M t^b\} = m_b T_R \implies m_a = \frac{1}{T_R} \text{Tr}\{M t^a\}. \quad (5.1.24)$$

Pertanto (esplicitando gli indici nel secondo passaggio e poi raccogliendo):

$$M_{ij} = \frac{1}{N} M_{kk} \delta_{ij} + \frac{1}{T_R} M_{lk} t_{kl}^a t_{ij}^a \quad (5.1.25)$$

$$M_{lk} \delta_{il} \delta_{jk} = \frac{1}{N} M_{lk} \delta_{kl} \delta_{ij} + \frac{1}{T_R} M_{lk} t_{kl}^a t_{ij}^a \quad (5.1.26)$$

$$M_{lk} \left[t_{ij}^a t_{kl}^a - T_R \left(\delta_{il} \delta_{jk} - \frac{1}{N} \delta_{ij} \delta_{kl} \right) \right] = 0, \quad \forall M \quad (5.1.27)$$

il che prova l'identità.

5.2 Lagrangiana della QCD

La QCD è una teoria di gauge basata sul gruppo di gauge $SU(3)$ di colore. La sua Lagrangiana classica è:

$$\mathcal{L}_{\text{QCD}} = -\frac{1}{4}F^{\mu\nu,a}F_{\mu\nu}^a + \bar{q}_f^i \left[i\gamma^\mu D_{ij}^\mu - m_f \delta_{ij} \right] \delta_{fg} q_g^j \quad (5.2.1)$$

dove $f, g = u, d, s, \dots$ sono gli indici di sapore (flavour) dei quark, $i, j = 1, 2, 3$ sono gli indici di colore dei quark nella rappresentazione fondamentale di $SU(3)$, $a = 1, \dots, 8$ è l'indice di colore dei gluoni nella rappresentazione aggiunta di $SU(3)$, e μ è un indice di Lorentz. In più, γ_μ sono le matrici di Dirac, mentre m_f e q_f sono rispettivamente la massa e il campo del quark per il sapore f . Il tensore degli sforzi per i campi gluonici A_μ^a è definito come:

$$F_{\mu\nu}^a \equiv \partial_\mu A_\nu^a - \partial_\nu A_\mu^a + g_s f^{abc} A_\mu^b A_\nu^c \quad (5.2.2)$$

$$= \partial_\mu A_\nu^a - \partial_\nu A_\mu^a - ig_s (t_G^b)_{ac} A_\mu^b A_\nu^c \quad (5.2.3)$$

$$= \mathcal{D}_{\mu,ac} A_\nu^c - \partial_\nu A_\mu^a \quad (5.2.4)$$

in cui, oltre a ricordarci la relazione per le costanti di struttura (5.1.7), abbiamo g_s è il parametro della costante di accoppiamento forte, e la derivata covariante nella rappresentazione aggiunta (poiché contiene il generatore dell'aggiunta) è:

$$\mathcal{D}_{\mu,ac} \equiv \partial_\mu \delta_{ac} - ig_s (t_G^b)_{ac} A_\mu^b \quad (5.2.5)$$

$$= \partial_\mu \delta_{ac} + g_s f^{abc} A_\mu^b. \quad (5.2.6)$$

Analogamente, la derivata covariante nella rappresentazione fondamentale (avendo il generatore della rappresentazione fondamentale) è definita come:

$$\mathcal{D}_{\mu,ij} \equiv \partial_\mu \delta_{ij} - ig_s t_{ij}^a A_\mu^a. \quad (5.2.7)$$

Notiamo che in questo caso $F_{\mu\nu}^a$ contiene anche termini proporzionali ad A_μ e A^2 , a differenza del caso abeliano. Possiamo anche notare che il contributo F^2 alla Lagrangiana della QCD contiene il termine cinetico del gluone ($\propto A^2$), mentre i termini di massa per i gluoni sono proibiti come conseguenza dell'invarianza di gauge. In più, poiché il campo di Higgs non è carico sotto il grado di libertà del colore, $SU(3)$ non è rotta spontaneamente e il gluone non acquisisce una massa¹, fornendo una forza forte a lungo raggio.

È importante notare che il termine F^2 nella Lagrangiana presenta anche auto-interazioni gluoniche cubiche (cioè $\propto A^3$) e quartiche (cioè $\propto A^4$). Questi due termini sono la principale novità delle teorie di gauge non-abeliane

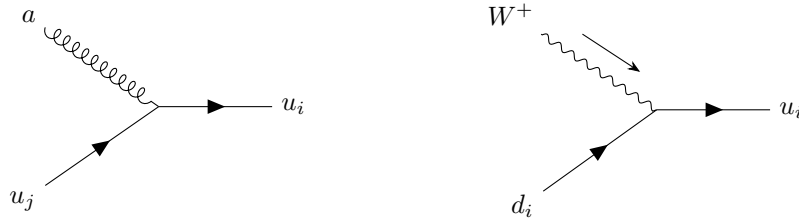
¹Infatti sappiamo che il campo di Higgs fornisce massa solamente ai bosoni di gauge associati alle simmetrie rotte, ma per essere rotta una simmetria deve trasformare in modo non banale il campo H (cioè, in questo caso, H non dev'essere un singoletto di colore).

rispetto alla QED e, come tali, sono governati dalla presenza delle costanti di struttura f^{abc} . Fisicamente, ciò esprime il fatto che i gluoni sono carichi sotto il colore e interagiscono tra loro, a differenza dei fotoni nella QED.

Il termine di massa dei quark ($m_f \delta_{fg} \delta_{ij}$) è diagonale nel sapore ($m_f \delta_{fg}$) e nel colore (δ_{ij}). Vediamo anche che esso dipende dal sapore (m_f), ma non dal colore, ovvero quark di colore diverso ma dello stesso sapore hanno la stessa massa. Il termine cinetico per i quark (che è dentro $\mathcal{D}_\mu$) è indipendente dal sapore e dal colore ($\delta_{fg} \delta_{ij}$).

Le interazioni quark-gluone sono regolate da $g_s t_{ij}^a$, per cui sono diagonali nel sapore, mentre sono non-diagonali nel colore. In particolare, quark con due colori diversi, i e j , possono interagire scambiando un gluone di colore a , essendo t_{ij}^a la loro matrice di mescolamento di colore. In altre parole, un quark di sapore f e colore i che irradia un gluone a mantiene il proprio sapore f , ma cambia il proprio colore in j . La situazione è complementare a quanto accade con le interazioni deboli, dove un quark di sapore f e colore i , irradiando un bosone debole W , transisce in un quark di un sapore diverso g (mixing di famiglia dei quark), ma con lo stesso colore i . Puoi vedere i diagrammi in figura 5.1 per i due esempi.

L'intensità delle interazioni quark-gluone è regolata dal singolo parametro g_s , il che significa che la dinamica delle interazioni forti è insensibile al sapore (*universalità del sapore*). L'unico effetto indiretto del sapore deriva dalle masse, la cui presenza influenza la cinematica.



(a) Osserviamo nelle interazioni forti il color mixing, ma non quello di flavour. (b) Osserviamo nelle interazioni deboli il flavour mixing, ma non quello di colore.

Figura 5.1

Sfruttando l'antisimmetria delle funzioni di struttura, possiamo riscrivere esplicitamente la Lagrangiana della QCD come:

$$\begin{aligned} \mathcal{L}_{\text{QCD}} = & -\frac{1}{4}(\partial_\mu A_\nu^a - \partial_\nu A_\mu^a)^2 - g_s f^{abc}(\partial_\mu A_\nu^a)A_\mu^b A_\nu^c - \\ & -\frac{1}{4}g_s^2 f^{abc}f^{ade}A_\mu^b A_\nu^c A^{d,\mu} A^{e,\nu} \\ & + \bar{q}_f^i(i\gamma^\mu \partial_\mu - m_f)q_f^i + g_s \bar{q}_f^i \gamma^\mu t_{ij}^a q_f^j A_\mu^a. \end{aligned} \quad (5.2.8)$$

Una nota sulle convenzioni. Poiché la Lagrangiana presenta esplicitamente i generatori t^a di $SU(3)$, la cui normalizzazione T_R è convenzionale, ci

si potrebbe chiedere quali siano le conseguenze di scelte diverse per essa. Se si riscalda $t^a \rightarrow \kappa t^a$, il valore di T_R viene riscalato come $T_R \rightarrow \kappa^2 T_R$ (si veda la definizione di T_R (5.1.11)). Per mantenere invariata l'interazione quark-gluone, l'accoppiamento deve essere coerentemente ridefinito come $g_s \rightarrow g_s/\kappa$. A sua volta, per mantenere invariate le auto-interazioni gluoniche, le costanti di struttura devono diventare $f^{abc} \rightarrow \kappa f^{abc}$, il che implica $C_A \rightarrow \kappa^2 C_A$ e $C_F \rightarrow \kappa^2 C_F$. Se facciamo questi riscaldamenti, allora manteniamo tutta la teoria consistente (cioè tutti i termini all'interno di $\mathcal{L}$ non cambiano se riscaldiamo i nostri oggetti) e non modifichiamo gli accoppiamenti anche se cambiamo convenzione. Notiamo anche che possiamo scrivere:

$$C_F = T_R \frac{N^2 - 1}{N} \quad , \quad C_A = 2T_R N. \quad (5.2.9)$$

Scritta la lagrangiana possiamo preoccuparci di ricavare i propagatori delle particelle presenti nella teoria.

Per ottenere il propagatore del gluone, il termine quadratico nella Lagrangiana sopra riportata non è sufficiente. È necessario un ulteriore termine di gauge fixing $\mathcal{L}_{\text{gf}}$. Noi salteremo tutti i conti per arrivare alle conclusioni, ma puoi vedere gli appunti di *Fondamenti di teoria dei campi*. Ciò deriva dal fatto che l'invarianza di gauge implica la presenza di infinite configurazioni di campo gluoniche equivalenti, tutte correlate da trasformazioni di gauge. L'integrale sui cammini della teoria somma su tutte queste, portando a un conteggio eccessivo. A sua volta, questa ridondanza implica l'impossibilità di invertire il termine quadratico della Lagrangiana per ottenere il propagatore del gluone. Il termine di fissaggio del gauge ha l'effetto di selezionare un rappresentante della configurazione di campo per famiglia, rimuovendo la ridondanza. La scelta:

$$\mathcal{L}_{\text{gf}} = -\frac{1}{2\xi} (\partial_\mu A^{a,\mu})^2 \quad (5.2.10)$$

impone effettivamente la condizione $\partial_\mu A^{a,\mu} = 0$, dando origine al seguente propagatore del gluone:

$$\Delta_{\mu\nu}^{ab}(q) = \frac{-i\delta^{ab}}{q^2 + i\epsilon} \left[g_{\mu\nu} - (1 - \xi) \frac{q_\mu q_\nu}{q^2} \right]. \quad (5.2.11)$$

Questo gauge è soprannominato gauge R_ξ , e le scelte comuni per ξ sono $\xi = 1$ (gauge di Feynman) e $\xi = 0$ (gauge di Landau). Il gauge di Feynman è particolarmente semplice, poiché il propagatore è proporzionale a $g_{\mu\nu}$. Nel gauge di Landau il propagatore presenta un termine extra, ma ha la caratteristica interessante di essere trasversale, ovvero $q_\mu \Delta_{\mu\nu}^{ab}(q) = q_\nu \Delta_{\mu\nu}^{ab}(q) = 0$.

Oltre al termine di fissaggio del gauge, che a livello di integrale sui cammini è implementato mediante una funzione delta di Dirac $\delta(G(A))$ (con $G(A) = \partial_\mu A^{a,\mu}$ per il gauge R_ξ), la quantizzazione coerente di Faddeev-Popov delle teorie di gauge non-abeliane richiede l'introduzione di un determinante nell'integrale sui cammini, che tiene conto della variazione del

vincolo di gauge $G(A)$ rispetto ai parametri di una trasformazione di gauge infinitesima. Questo determinante dà origine a un ulteriore termine nella Lagrangiana, che viene scritto in termini di campi ghost c^a :

$$\mathcal{L}_{\text{ghost}} = \bar{c}^a (-\partial^\mu \mathcal{D}_{ac}^\mu) c^c = \bar{c}^a (-\partial^\mu \partial_\mu \delta_{ac} - g_s f^{abc} \partial^\mu A_\mu^b) c^c. \quad (5.2.12)$$

I ghost sono campi scalari anti-commutanti nella rappresentazione aggiunta. Il fatto che siano nella rappresentazione aggiunta ci dice che sono colorati (ossia che sono carichi per $SU(3)_{\text{color}}$). Essi non corrispondono a particelle fisiche propaganti, quanto piuttosto a gradi di libertà "negativi" che cancellano l'effetto delle polarizzazioni longitudinali non fisiche dei bosoni di gauge propaganti (nel gauge R_ξ). Nelle teorie abeliane come la QED, la condizione di gauge $\partial_\mu A^\mu = 0$ è sufficiente a garantire fotoni propaganti trasversali, e i campi ghost non sono necessari. Ciò è coerente con il fatto che nelle teorie abeliane $f^{abc} = 0$, da cui la Lagrangiana dei ghost è un puro termine cinetico, $\mathcal{L}_{\text{ghost}} = -\bar{c} \partial^\mu \partial_\mu c$, ovvero i campi ghost sono completamente disaccoppiati dalla teoria. Nelle teorie di gauge non-abeliane, invece, i campi ghost interagiscono con i bosoni di gauge nel gauge R_ξ .

La presenza dei ghost dipende criticamente dalla scelta del gauge. Infatti, potremmo fare altre scelte per $\mathcal{L}_{gf}$ che potrebbero evitarci la necessità di dover introdurre i campi di ghost nella teoria (che ricordiamo non sono campi fisici, ma solamente oggetti matematici). Per scopi pratici, nella QCD può essere conveniente considerare la seguente Lagrangiana di gauge fixing:

$$\mathcal{L}_{\text{gf}} = -\frac{1}{2\xi} (n^\mu A_\mu^a)^2 \quad (5.2.13)$$

dove n^μ è un quadrivettore *arbitrario*, che definisce la classe dei cosiddetti *gauge assiali*. Il gauge fixing assiale induce il seguente propagatore:

$$\Delta_{\mu\nu}^{ab}(q) = \frac{-i\delta^{ab}}{q^2 + i\epsilon} \left[g_{\mu\nu} - \frac{n_\mu q_\nu + n_\nu q_\mu}{n \cdot q} + \frac{(n^2 + \xi q^2) q_\mu q_\nu}{(n \cdot q)^2} \right]. \quad (5.2.14)$$

Con $n^2 = \xi = 0$ il gauge assiale è chiamato *light-cone gauge*. Nel limite $q^2 \rightarrow 0$, ovvero nel caso del gluone *on-shell*, è facile verificare che il propagatore nel gauge assiale è trasversale, ossia $q^\mu \Delta_{\mu\nu}^{ab}(q) = n^\mu \Delta_{\mu\nu}^{ab}(q) = 0$, il che significa che vengono propagate solo le due polarizzazioni fisiche trasversali del gluone (poiché da 4 gradi di libertà, le condizioni di propagatore trasverso ne tolgono 2). Si può quindi dimostrare che, con tale scelta di gauge, non è necessario introdurre i campi ghost nella Lagrangiana. Infatti, il determinante di Faddeev-Popov, che genera i ghost nell'integrale sui cammini, è proporzionale a $n^\mu A_\mu^a$, che è posto a 0 dal termine di fissaggio del gauge. Questa è una semplificazione che può compensare la maggiore complessità del propagatore rispetto al gauge di Feynman.

5.2.1 Regole di Feynman

Data la lagrangiana completa, dunque (5.2.8) in aggiunta quella di ghost (5.2.12):

$$\begin{aligned}\mathcal{L}_{\text{QCD}} = & -\frac{1}{4}(\partial_\mu A_\nu^a - \partial_\nu A_\mu^a)^2 - g_s f^{abc}(\partial_\mu A_\nu^a)A_\mu^b A_\nu^c - \\ & -\frac{1}{4}g_s^2 f^{abc}f^{ade}A_\mu^b A_\nu^c A^{d,\mu} A^{e,\nu} \\ & + \bar{q}_f^i(i\gamma^\mu \partial_\mu - m_f)q_f^i + g_s \bar{q}_f^i \gamma^\mu t_{ij}^a q_f^j A_\mu^a + \\ & + \bar{c}^a(-\partial^\mu \partial_\mu \delta_{ac} - g_s f^{abc} \partial^\mu A_\mu^b) c^c\end{aligned}\quad (5.2.15)$$

possiamo provare a ricavarne le regole di Feynman nel R_ξ gauge. Notando che per le regole di Feynman per il settore dei quark possono essere dedotte direttamente.

Cominciamo con le parti facili, ovvero i propagatori. Il *propagatore per i quark* si legge:

$$\Delta_{ij}^{fg}(p) = \frac{i(\not{p} + m_f)}{p^2 - m_f^2 + i\epsilon} \delta_{ij} \delta_{fg} \quad (5.2.16)$$

dove i, j sono indici di colore nella rappresentazione fondamentale, mentre f, g sono indici di sapore. Abbiamo già detto che il *propagatore del gluone* nel R_ξ gauge è:

$$\Delta_{\mu\nu}^{ab}(q) = \frac{-i\delta^{ab}}{q^2 + i\epsilon} \left[g_{\mu\nu} - (1 - \xi) \frac{q_\mu q_\nu}{q^2} \right]. \quad (5.2.17)$$

Dalla Lagrangiana dei ghost (se necessaria, a seconda della scelta del gauge), si deduce in modo simile il *propagatore dei ghost*:

$$\Delta_{ab}(p) = \frac{i\delta^{ab}}{p^2 + i\epsilon}. \quad (5.2.18)$$

Possiamo notare che, mentre il propagatore dei gluoni, e dei ghost, sono indipendenti dal segno del momento, si ha che il propagatore dei quark ne dipende.

Ora vediamo i vari termini di interazione. Per dedurre le regole di Feynman dei vertici consideriamo la trasformata di Fourier $\tilde{A}$ del campo gluone:

$$A_\nu^a(x) = \int [dk] e^{ikx} \tilde{A}_\nu^a(k), \quad [dk] \equiv \frac{d^4 k}{(2\pi)^4}. \quad (5.2.19)$$

L'*interazione quark-gluone* dà origine alla regola di Feynman:

$$ig_s \gamma^\mu t_{ij}^a \delta_{fg}. \quad (5.2.20)$$

Si noti che, a causa della presenza del colore, l'interazione quark-fotone è $iQ_f \gamma^\mu \delta_{ij} \delta_{fg}$, con Q_f la carica elettrica del sapore f ($Q_u = 2/3$, $Q_d = -1/3$, e così via). In particolare, il termine δ_{ij} (contrariamente a 1) implementa il fatto che il colore del quark non viene mescolato dalle interazioni fotoni-
che. Lo stesso fattore di colore δ_{ij} modifica tutti i vertici elettrodeboli che coinvolgono i quark nel Modello Standard.

Per quanto riguarda l'accoppiamento a tre gluoni il vertice è (in cui abbiamo tutti i campi gluonici valutati nello stesso punto dello spazio-tempo, x , ma che portano un momento diverso):

$$\begin{aligned} -g_s f^{abc} (\partial_\mu A_\nu^a) A_\mu^b A_\nu^c &= \\ &= \int [dk][dp][dq] e^{i(k+p+q)x} (-g_s) f^{abc} i k_\mu g_{\nu\rho} \tilde{A}_\nu^a(k) \tilde{A}_\mu^b(p) \tilde{A}_\rho^c(q) \end{aligned} \quad (5.2.21)$$

in cui il termine $-i k_\mu g_{\nu\rho} \tilde{A}_\nu^a(k)$ deriva dal passaggio nello spazio di Fourier della derivata $\partial_\mu A_\nu^a$. Questa è solo una delle 3! possibili modalità di scrivere la trasformata di Fourier del vertice, una per ogni permutazione dei campi gluonici. Ad esempio, facendo i cambi nell'ordine $p \leftrightarrow k$, $\nu \leftrightarrow \rho$, $a \leftrightarrow c$ e usando l'antisimmetria delle costanti di struttura, otteniamo tutte espressioni equivalenti:

$$-g_s \int f^{abc} (\partial_\mu A_\nu^a) A_\mu^b A_\nu^c = \quad (5.2.22)$$

$$= \int [dp][dk][dq] e^{i(p+k+q)x} (-g_s) f^{abc} i p_\mu g_{\nu\rho} \tilde{A}_\nu^a(p) \tilde{A}_\mu^b(k) \tilde{A}_\rho^c(q) \quad (5.2.23)$$

$$= \int [dp][dk][dq] e^{i(p+k+q)x} (-g_s) f^{abc} i p_\mu g_{\rho\nu} \tilde{A}_\rho^a(p) \tilde{A}_\mu^b(q) \tilde{A}_\nu^c(k) \quad (5.2.24)$$

$$= \int [dp][dk][dq] e^{i(p+k+q)x} (-g_s) f^{cba} i p_\mu g_{\rho\nu} \tilde{A}_\rho^c(p) \tilde{A}_\mu^b(q) \tilde{A}_\nu^a(k) \quad (5.2.25)$$

$$= - \int [dk][dp][dq] e^{i(k+p+q)x} (-g_s) f^{abc} i p_\mu g_{\nu\rho} \tilde{A}_\nu^a(k) \tilde{A}_\mu^b(q) \tilde{A}_\rho^c(p). \quad (5.2.26)$$

Dunque abbiamo capito che possiamo utilizzare k_μ , oppure $-p_\mu$ in modo equivalente, per cui considerando tutte le possibili permutazioni, dunque scrivendo un'espressione (anti)simmetrica (per via della presenza della costante di struttura) per tutti gli scambi gluonici, otteniamo:

$$\begin{aligned} -g_s f^{abc} (\partial_\mu A_\nu^a) A_\mu^b A_\nu^c &= \int [dk][dp][dq] e^{i(k+p+q)x} \tilde{A}_\nu^a(k) \tilde{A}_\mu^b(q) \tilde{A}_\rho^c(p) \frac{i}{3!} \times \\ &\quad \times g_s f^{abc} [(k-q)_\rho g_{\mu\nu} + (p-k)_\mu g_{\rho\nu} + (q-p)_\nu g_{\rho\mu}] \end{aligned} \quad (5.2.27)$$

dove la seconda riga è la regola di Feynman associata al vertice a tre gluoni. Si noti che l'esponenziale è $e^{i(k+p+q)x}$: integrando in d^4x questo termine

diventa $\delta^{(4)}(k+p+q)$, ovvero i momenti dei tre gluoni sommano a 0. In altre parole, la regola di Feynman sopra riportata è per tre gluoni con momenti k^ν , q^μ e p^ρ tutti entranti nel vertice.

Possiamo dedurre l'auto-interazione quartica dei gluoni in modo simile:

$$-\frac{1}{4}g_s^2 f^{abc} f^{ade} A_\mu^b A_\nu^c A^{d,\mu} A^{e,\nu} = \int [dk][dp][dq][dl] e^{i(k+p+q+l)x} \times \\ \tilde{A}_\mu^b(k) \tilde{A}_\nu^c(p) \tilde{A}_\alpha^d(q) \tilde{A}_\beta^e(l) \left(-\frac{1}{4}\right) g_s^2 f^{abc} f^{ade} g^{\mu\alpha} g^{\nu\beta}. \quad (5.2.28)$$

Permutando i quattro gluoni otteniamo 4! termini equivalenti, dai quali possiamo estrarre il vertice:

$$-\frac{1}{4}g_s^2 f^{abc} f^{ade} A_\mu^b A_\nu^c A^{d,\mu} A^{e,\nu} = \int [dk][dp][dq][dl] e^{i(k+p+q+l)x} \times \\ \times \tilde{A}_\mu^b(k) \tilde{A}_\nu^c(p) \tilde{A}_\alpha^d(q) \tilde{A}_\beta^e(l) \frac{-i}{4!} \times \\ \times (-ig_s^2) \left[f^{abc} f^{ade} (g^{\mu\alpha} g^{\nu\beta} - g^{\nu\alpha} g^{\mu\beta}) + \right. \\ \left. + f^{adc} f^{abe} (g^{\mu\alpha} g^{\nu\beta} - g^{\mu\nu} g^{\alpha\beta}) \right. \\ \left. + f^{abd} f^{ace} (g^{\mu\nu} g^{\alpha\beta} - g^{\nu\alpha} g^{\mu\beta}) \right] \quad (5.2.29)$$

dove le righe dalla terza alla quinta sono la corrispondente regola di Feynman. In cui notiamo ancora che l'esponenziale contiene la somma $(k+p+q+l)x$, cioè integrata ci dà la delta di conservazione nel vertice, e implica tutti i momenti entranti.

Il *vertex ghost-gluone* è relativamente semplice da ricavare (integrando per parti nel primo passaggio):

$$-g_s f^{abc} \bar{c}^a \partial^\mu \left[A_\mu^b c^c \right] = g_s f^{abc} (\partial^\mu \bar{c}^a) A_\mu^b c^c \quad (5.2.30)$$

$$= \int [dk][dp][dq] e^{i(k+q-p)x} i \tilde{c}^a(p) \tilde{A}_\mu^b(k) \tilde{c}^c(q) \left[-g_s f^{abc} p_\mu \right] \quad (5.2.31)$$

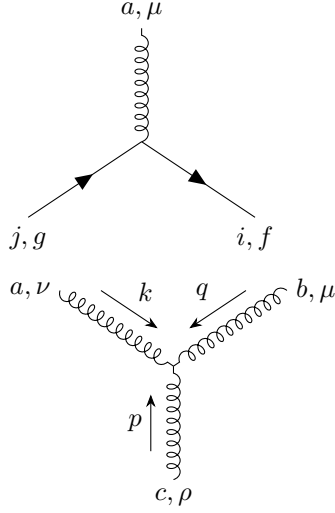
dove la parentesi quadra fornisce la regola di Feynman e l'esponenziale ci dice che ci sono due momenti entranti ed uno (quello associato a $\bar{c}^a$) uscente.

Ricapitolando si hanno le regole:

$$\begin{array}{c} \xrightarrow{p} \\ i, f \quad \quad j, g \end{array} \quad \frac{i(\not{p} + m_f)}{p^2 - m_f^2 + i\epsilon} \delta_{ij} \delta_{fg} \quad (5.2.32)$$

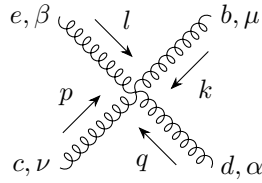
$$\begin{array}{c} \xrightarrow{q} \\ a, \mu \quad \quad b, \nu \end{array} \quad \frac{-i\delta^{ab}}{q^2 + i\epsilon} \left[g_{\mu\nu} - (1 - \xi) \frac{q_\mu q_\nu}{q^2} \right] \quad (5.2.33)$$

$$\begin{array}{c} \xrightarrow{p} \\ a \quad \quad b \end{array} \quad \frac{i\delta^{ab}}{p^2 + i\epsilon} \quad (5.2.34)$$

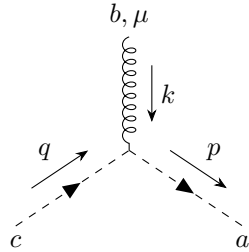


$$ig_s \gamma^\mu t_{ij}^a \delta_{fg} \quad (5.2.35)$$

$$g_s f^{abc} \left[(k - q)_\rho g_{\mu\nu} + (p - k)_\mu g_{\rho\nu} + (q - p)_\nu g_{\rho\mu} \right] \quad (5.2.36)$$



$$-ig_s^2 \left[f^{abc} f^{ade} (g^{\mu\alpha} g^{\nu\beta} - g^{\nu\alpha} g^{\mu\beta}) + f^{adc} f^{abe} (g^{\mu\alpha} g^{\nu\beta} - g^{\mu\nu} g^{\alpha\beta}) + f^{abd} f^{ace} (g^{\mu\nu} g^{\alpha\beta} - g^{\nu\alpha} g^{\mu\beta}) \right] \quad (5.2.37)$$

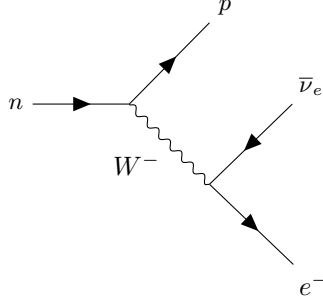


$$-g_s f^{abc} p_\mu. \quad (5.2.38)$$

5.3 Running coupling e libertà asintotica

Al fine di introdurre il concetto di running coupling in modo fenomenologico, seguiamo la discussione di Ellis There is a sort of mythology that grows up about what happened, which is different from what really did happen. Ellis, ma che il capitolo §17 del Peskin There is a sort of mythology that grows up about what happened, which is different from what really did happen. Peskin offre spunti interessanti, e consideriamo un osservabile adimensionale $R(Q^2)$, che assumiamo dipendere da una singola scala di energia fisica Q . R può essere una qualsiasi osservabile, ossia qualcosa di misurabile, come ad esempio una sezione d'urto normalizzata (per esempio una sezione d'urto in una qualche regione dello spazio delle fasi, divisa per la sezione d'urto totale), che assumiamo calcolabile nella teoria perturbativa della QCD. Un esempio

facile che possiamo considerare è il decadimento del neutrone:

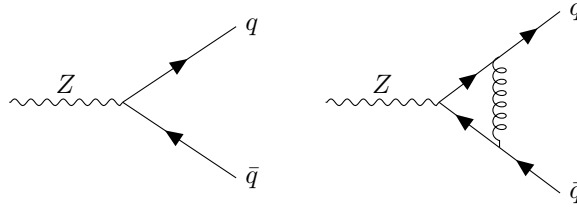


che in teoria elettrodebole non contiene alcuna correzione di QCD, per cui ne è il limite classico. In altri termini, la sezione d'urto del diagramma raffigurato è il primo ordine perturbativo dello stesso processo in QCD.

Un altro esempio è il branching ratio del decadimento dello Z_μ :

$$R = \frac{\Gamma_{Z \rightarrow q\bar{q}}}{\Gamma_Z} \quad (5.3.1)$$

di cui l'ordine 0 e il primo ordine in QCD sono rispettivamente:



in cui il secondo contributo è proporzionale a g_s^2 .

Ad ogni modo, per analisi dimensionale, la predizione classica è $R = \text{costante}$: la dipendenza da Q è infatti irrilevante, poiché non esiste un'altra scala di energia che possa compensarla per dare un risultato adimensionale per R .

A livello quantistico, consideriamo le correzioni QCD alla predizione per R . Tali correzioni coinvolgono, ad esempio, diagrammi a loop che devono essere sottoposti a rinormalizzazione per eliminare le divergenze ultraviolette (UV). Le infinità UV sono contenute nei parametri "nudi" della Lagrangiana, che non sono fisici e devono essere ridefiniti. Si possono sistematicamente suddividere i parametri nudi in una somma di parametri "rinormalizzati" più controtermini UV, dove questi ultimi assorbono tutte le infinità UV, mentre i primi sono quantità fisicamente misurabili. Questo è oggetto del corso di *Advanced QFT*.

Poiché la QCD è una teoria rinormalizzabile, un numero finito di controtermini UV, calcolati ordine per ordine nella teoria perturbativa, è in grado

di assorbire tutte le infinità UV derivanti da diagrammi a più loop. Il prezzo di questa procedura, tuttavia, è che una scala non fisica μ , la *scala di rinormalizzazione*, viene introdotta nel calcolo, il che accade indipendentemente dallo schema di rinormalizzazione adottato². La presenza di μ codifica essenzialmente la quantità di contributi finiti che vengono assorbiti nella definizione dei controtermini UV, insieme alle infinità UV. Il punto fondamentale è che, considerando le correzioni radiative a R in potenze di $\alpha_s \equiv g_s^2/(4\pi)$, la predizione per l'osservabile è $R(Q^2, \mu^2, \alpha_s) \equiv R(Q^2/\mu^2, \alpha_s)$, dove quest'ultima uguaglianza deriva dalla considerazione che Q e μ possono apparire solo come rapporto, al fine di mantenere R adimensionale.

Il considerare α_s al posto di g_s è solo per semplificare la trattazione.

Dunque, abbiamo capito che in una situazione classica, quindi al primo ordine in QCD, abbiamo $R = \text{cost.}$, ma appena consideriamo le correzioni radiative, quindi quantistiche, la nostra osservabile assume la dipendenza $R = R(Q^2/\mu^2, \alpha_s)$.

Ad ogni modo, i risultati fisici non possono dipendere dalla scala non fisica μ , quindi se si fosse in grado di calcolare R a tutti gli ordini nella teoria perturbativa (vedi il proseguio della sezione per capire che succede se ci dovessimo fermare ad un certo ordine), la sua dipendenza da μ svanirebbe esattamente. Esprimiamo formalmente questo fatto richiedendo che:

$$\mu^2 \frac{d}{d\mu^2} R(Q^2/\mu^2, \alpha_s) = \left[\mu^2 \frac{\partial}{\partial \mu^2} + \mu^2 \frac{\partial \alpha_s}{\partial \mu^2} \frac{\partial}{\partial \alpha_s} \right] R(Q^2/\mu^2, \alpha_s) = 0. \quad (5.3.2)$$

Questa equazione è chiamata equazione del gruppo di rinormalizzazione (RGE) o **equazione di Callan-Symanzik**. Poiché R ha acquisito una dipendenza esplicita da μ , così deve fare α_s , per compensare esattamente la prima.³

Al fine di trovare la soluzione dell'equazione precedente, introduciamo le variabili:

$$t = \ln(Q^2/\mu^2) \quad , \quad \beta(\alpha_s) \equiv \mu^2 \frac{\partial \alpha_s}{\partial \mu^2} \quad (5.3.4)$$

in cui $\beta(\alpha_s)$ ci rappresenta quanto velocemente α_s varia con la scala μ^2 . Dopodiché la RGE diventa:

$$\left[-\frac{\partial}{\partial t} + \beta(\alpha_s) \frac{\partial}{\partial \alpha_s} \right] R(e^t, \alpha_s) = 0 \quad (5.3.5)$$

²Si evita lo schema on-shell, utilizzato per la QED, poiché non si vogliono trattare quark on-shell, che ci porterebbero in regime perturbativo.

³Se avessimo $R = (Q^2/\mu^2)$, dunque che non dipendesse da α_s , allora la RGE sarebbe:

$$\mu^2 \frac{\partial}{\partial \mu^2} R(Q^2/\mu^2) = 0 \quad (5.3.3)$$

quindi $R(Q^2/\mu^2)$ sarebbe una costante in una teoria con la scala μ . Il che non avrebbe senso.

dove t e α_s devono essere considerati come due variabili indipendenti.

Al fine di trovare la soluzione generale della RGE, lavoriamo prima su un esempio più semplice. Per risolvere l'equazione differenziale alle derivate parziali:

$$\left[-\frac{\partial}{\partial x} + b\frac{\partial}{\partial y} \right] F(x, y) = 0 = \vec{v} \cdot \vec{\nabla}_2 F(x, y) \quad (5.3.6)$$

dove $\vec{v} = (-1, b)$, $\vec{\nabla}_2 = (\partial/\partial x, \partial/\partial y)$, si inizia notando che la funzione $Y(x, y) \equiv bx + y$ fornisce una soluzione particolare dell'equazione. Se $F(x, y)$ deve risolvere l'equazione (5.3.6) in generale, le sue dipendenze da x e y non possono essere scorrelate, ma devono piuttosto entrare solo attraverso la combinazione $Y(x, y)$. Chiamando $\gamma(s) = (x(s), y(s))$ la curva caratteristica dell'equazione, definita in modo tale che $d\gamma/ds = (dx/ds, dy/ds) = (-1, b) = \vec{v}$, allora $Y(x, y)$ è invariante lungo γ , ovvero $dY/ds = 0$. $Y(x, y)$ definisce quindi una *famiglia di curve caratteristiche* nel piano (x, y) , le soluzioni di $Y(x, y) = \text{const}$, lungo le quali l'equazione è soddisfatta.

Assumendo per semplicità che F sia regolare all'origine (ma l'argomento può essere generalizzato), abbiamo quindi:

$$F(x, y) = \sum_{k=0}^{\infty} c_k (bx + y)^k \quad (5.3.7)$$

$$= \sum_{k=0}^{\infty} c_k Y(x, y)^k \quad (5.3.8)$$

$$= \sum_{k=0}^{\infty} c_k (b \times 0 + Y(x, y))^k \quad (5.3.9)$$

$$= F(0, Y(x, y)). \quad (5.3.10)$$

La notazione evidenzia che la dipendenza della soluzione generale da uno dei due argomenti è immateriale, e l'intera dipendenza sia da x che da y è codificata attraverso l'invariante $Y(x, y)$.

Notiamo che questo argomento è lo stesso che si segue per risolvere l'equazione d'onda di d'Alembert. Essendo del secondo ordine, se consideriamo il caso $1 + 1$ dimensionale, l'equazione presenta due funzioni invarianti invece di una (avremmo sei funzioni invarianti in $3 + 1$ dimensioni), ovvero $Y_{\pm}(x, t) = x \pm vt$, e la soluzione generale dell'equazione è:

$$F(x, t) = F_+(0, Y_+(x, t)) + F_-(0, Y_-(x, t)).$$

Notiamo inoltre che la funzione invariante può essere trovata imponendo esplicitamente che risolva l'equazione differenziale:

$$-\int_0^x dx' + \int_y^{Y(x,y)} \frac{dy'}{b} = 0. \quad (5.3.11)$$

A questo punto, fatti i commenti sulle equazioni alle derivate parziali tramite un esempio semplificato, possiamo risolvere la RGE fisica (5.3.5) usando la stessa tecnica. Introduciamo una funzione $A_s(t, \alpha_s)$, che viene chiamata **running coupling**, che è implicitamente definita da:

$$-\int_0^t dt' + \int_{\alpha_s}^{A_s(t, \alpha_s)} \frac{da}{\beta(a)} = 0 \quad (5.3.12)$$

$$\implies \ln(Q^2/\mu^2) = t \equiv \int_{\alpha_s}^{A_s(Q^2)} \frac{da}{\beta(a)} \quad (5.3.13)$$

e rappresenta l'invariante di base lungo le curve caratteristiche della RGE. Notiamo che $A_s(t, \alpha_s)$ soddisfa la condizione al contorno $A_s(0, \alpha_s) = \alpha_s$, ovvero la funzione A_s è uguale alla costante di accoppiamento della QCD se $\mu = Q$.

Per costruzione, $A_s(t, \alpha_s)$ soddisfa la RGE. Infatti, se differenziamo l'eq. (5.3.13) rispetto sia a t che a α_s otteniamo:

$$\begin{aligned} \frac{\partial t}{\partial t} = 1 &= \frac{\partial}{\partial t} \int_{\alpha_s}^{A_s(Q^2)} \frac{da}{\beta(a)} = \frac{1}{\beta(A_s(Q^2))} \frac{\partial A_s(Q^2)}{\partial t} \\ \implies \beta(A_s(Q^2)) &= \frac{\partial A_s(Q^2)}{\partial t} \end{aligned} \quad (5.3.14)$$

$$\begin{aligned} \frac{\partial t}{\partial \alpha_s} = 0 &= \frac{\partial}{\partial \alpha_s} \int_{\alpha_s}^{A_s(Q^2)} \frac{da}{\beta(a)} = \frac{1}{\beta(A_s(Q^2))} \frac{\partial A_s(Q^2)}{\partial \alpha_s} - \frac{1}{\beta(\alpha_s)} \\ \implies \frac{\beta(A_s(Q^2))}{\beta(\alpha_s)} &= \frac{\partial A_s(Q^2)}{\partial \alpha_s}. \end{aligned} \quad (5.3.15)$$

In altre parole, la dipendenza di $A_s(Q^2)$ da t è legata alla sua dipendenza da α_s tramite:

$$\frac{\partial A_s(Q^2)}{\partial t} = \beta(\alpha_s) \frac{\partial A_s(Q^2)}{\partial \alpha_s}. \quad (5.3.16)$$

Dunque, la soluzione generale $R(e^t, \alpha_s)$ della RGE è ottenuta mettendo il primo argomento *immateriale*, ovvero $t = 0$, e rimpiazzando il secondo argomento α_s con l'invariante $A_s(t, \alpha_s)$; in questo modo, in termini di $A_s(Q^2)$, è facile mostrare che $R(1, A_s(t, \alpha_s))$ è una soluzione della RGE:

$$\begin{aligned} \left[-\frac{\partial}{\partial t} + \beta(\alpha_s) \frac{\partial}{\partial \alpha_s} \right] R(1, A_s(t, \alpha_s)) &= \\ = \frac{\partial R(1, A_s(t, \alpha_s))}{\partial A_s(t, \alpha_s)} \left[-\frac{\partial A_s(t, \alpha_s)}{\partial t} + \beta(\alpha_s) \frac{\partial A_s(t, \alpha_s)}{\partial \alpha_s} \right] &= 0. \end{aligned} \quad (5.3.17)$$

Notiamo che $R(1, A_s(Q^2))$ non è un'ansatz, bensì abbiamo visto che le variabili t (che contiene al suo interno il rapporto Q^2/μ^2) e α_s sono variabili indipendenti e la dipendenza di $A_s(Q^2)$ da una di esse è strettamente legata

alla dipendenza dall'altra. Dunque, se riesprimiamo $R(e^t, \alpha_s)$ in termini di $A_s(Q^2)$, allora riusciamo ad ottenere un'equazione in cui, la dipendenza di A_s da t e α_s , fa' in modo che R calcolata in questo modo, cioè a diversi valori di t , scorra sulla curva delle soluzioni in modo da mantenere sempre la dipendenza fissata.

Questa è la dichiarazione formale di quanto anticipato sopra: poiché R deve essere strettamente indipendente da μ , dunque l'accoppiamento acquisisce una dipendenza dalla scala per soddisfare la RGE. La funzione $A_s(Q^2)$, che d'ora in poi chiameremo $\alpha_s(Q^2)$ ⁴, ha il significato di un parametro di accoppiamento dipendente dalla scala, ed è, come già detto, chiamato *running coupling*. L'accoppiamento "corre" (cioè varia in funzione della scala di energia) con una velocità logaritmica dettata da $\beta(\alpha_s)$, chiamata *QCD beta function*. Le sue condizioni al contorno sono $A_s(0, \alpha_s) = \alpha_s(\mu^2) = \alpha_s$.

Si può quindi semplicemente calcolare $R(1, \alpha)$ nella teoria perturbativa, cioè con $\mu^2 = Q^2$ e un'accoppiamento generico α , e poi sostituire $\alpha \rightarrow \alpha_s(Q^2)$ per recuperare la piena dipendenza di R dalla scala fisica Q , la quale automaticamente risolve la RGE. Q acquisisce quindi il significato di una scala alla quale l'accoppiamento viene sondato (cioè misurato), ovvero rappresenta la scala tipica del processo di scattering che contribuisce al calcolo di R . Nell'esempio del decadimento dello Z abbiamo $Q^2 = M_Z^2$.

Se si considera una scala di massa esplicita m nel calcolo di R , la corrispondente RGE è più complicata (puoi vedere il Peskin *There is a sort of mythology that grows up about what happened, which is different from what really did happen*. Peskin per dettagli):

$$\left[\mu^2 \frac{\partial}{\partial \mu^2} + \beta(\alpha_s) \frac{\partial}{\partial \alpha_s} - \gamma_m(\alpha_s) m \frac{\partial}{\partial m} \right] R(Q^2/\mu^2, \alpha_s, m/Q) = 0 \quad (5.3.18)$$

la cui soluzione è $R(1, \alpha_s(Q^2), m(Q^2)/Q)$, così come un'equazione alle derivate parziali tridimensionale, richiede l'introduzione di due funzioni (i.e. due invarianti lungo la curva caratteristica dell'equazione), le quali sono il *running coupling* $\alpha_s(Q^2)$, e la *running mass* $m(Q^2)$ (ad esempio, la massa dei quark dipende dalla scala, e dallo schema di rinormalizzazione, scelto), la quale soddisfa l'equazione:

$$Q^2 \frac{\partial m(Q^2)}{\partial Q^2} = -\gamma_m(\alpha_s(Q^2)) m(Q^2) \quad (5.3.19)$$

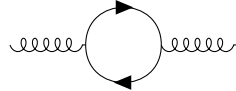
dove γ_m è chiamata *dimensione anomala della massa* e ci dice quanto la massa dipende dalla scala Q^2 , in aggiunta alla sua scala (naturale) E^{-1} .

⁴La notazione temporanea $A_s(Q^2)$ aveva solo lo scopo di rendere chiaro che $A_s(Q^2)$ è una funzione della scala di energia rilevante, mentre α_s è una variabile indipendente rispetto alla quale si sta differenziando.

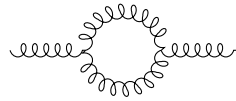
Nella QCD la funzione beta può essere calcolata tramite sviluppo perturbativo:

$$\beta(\alpha_s) = -(\beta_0 \alpha_s^2 + \beta_1 \alpha_s^3 + \dots) \quad (5.3.20)$$

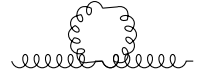
dove il calcolo (vedi il corso di *Advanced QFT*) del coefficiente di ordine più basso β_0 procede attraverso i seguenti diagrammi a 1-loop:



$$(5.3.21)$$



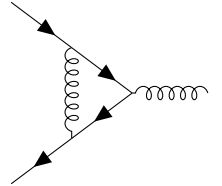
$$(5.3.22)$$



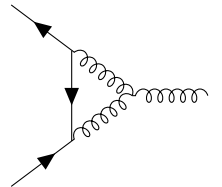
$$(5.3.23)$$



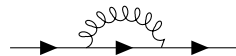
$$(5.3.24)$$



$$(5.3.25)$$



$$(5.3.26)$$



$$(5.3.27)$$

il cui risultato è:

$$\beta_0 = \frac{11C_A - 4T_R N_f}{12\pi}. \quad (5.3.28)$$

La funzione $\beta(\alpha_s)$ ci dice quanto il coupling dipende dalla scala; in particolare ci mostra la velocità con cui α_s cambia rispetto Q^2 . Diciamo che α_s runna sulla scala Q^2 con velocità logaritmica β .

Ricordiamo che $C_A = 3$ è il fattore di Casimir della rappresentazione aggiunta, mentre $T_R = 1/2$ è la normalizzazione dei generatori nella rappresentazione fondamentale di $SU(3)$. N_f rappresenta il numero di sapori di quark "attivi", ovvero quelli che danno un contributo apprezzabile al diagramma (5.3.21). Questi sono i sapori che possono essere considerati privi di massa alla scala Q a cui l'accoppiamento viene sondato, poiché i quark

con una massa $m_q \gg Q$ danno solo contributi trascurabili al diagramma (nel limite $m_q \rightarrow \infty$ essi si disaccoppiano, cioè non contribuiscono alla funzione beta). Risulta quindi $N_f = 4$ per scale $m_c \lesssim Q \lesssim m_b$, $N_f = 5$ se $m_b \lesssim Q \lesssim m_t$, e così via. Il contributo di C_A deriva da una traccia della costante di struttura f^{abc} , ovvero è generato dai diagrammi non-abeliani (5.3.22), (5.3.23), (5.3.24), (5.3.26) sopra riportati.

Per il calcolo del coefficiente β_1 abbiamo bisogno dei diagrammi a due loop.

Se valutiamo β_0 numericamente per $N = 3$ colori, il suo numeratore diventa $33 - 2N_f$, che è positivo per $N_f \leq 16$, come nel caso fisico. La funzione beta della QCD è quindi negativa, il che significa che l'accoppiamento corrente è una funzione decrescente della scala di energia. Questa è la proprietà della *libertà asintotica*.

Possiamo immediatamente notare che per la QED $C_A = 0$ (poiché $f^{abc} = 0$ nel caso abeliano), e la funzione beta è positiva. In altre parole, la QED non è asintoticamente libera. Infatti, poiché è una teoria abeliana, i diagrammi (5.3.22), (5.3.23), (5.3.24) e (5.3.27) non esistono, il che porta a dire $C_A = 0$, che a sua volta comporta $\beta_0^{QED} < 0$ e quindi $\beta^{QED} > 0$. Però, nonostante il coupling della QED runni, esso lo fa' in modo crescente con la scala, ma molto debolmente.

Inoltre, notiamo che β_0 (così come gli altri ordini) è soggetto a convenzione, per via della presenza di C_A e T_R , ma non lo è il suo segno. Infatti, se ricordiamo come cambiano i fattori di Casimir modificando la convenzione, allora è veloce vedere che il segno di β_0 non cambia. Dopotutto è logico che funzioni così poiché è proprio il segno di β ad essere fisicamente rilevante.

Possiamo integrare esplicitamente l'eq. (5.3.4) al più basso ordine (cioè mantenendo solo β_0 nell'eq. (5.3.20)), rinominando rispetto a prima $A_s(Q^2) \equiv \alpha_s(Q_2^2)$ e $\alpha_s = A_s(Q^2) \equiv \alpha_s(Q_1^2)$:

$$\int_{Q_1^2}^{Q_2^2} \frac{d\mu^2}{\mu^2} = \int_{\alpha_s(Q_1^2)}^{\alpha_s(Q_2^2)} \frac{da}{\beta(a)} \quad (5.3.29)$$

$$= \int_{\alpha_s(Q_1^2)}^{\alpha_s(Q_2^2)} \frac{da}{-\beta_0 a^2 + \dots} \quad (5.3.30)$$

$$\beta_0 \ln(Q_2^2/Q_1^2) = \frac{1}{\alpha_s(Q_2^2)} - \frac{1}{\alpha_s(Q_1^2)} \quad (5.3.31)$$

$$\alpha_s(Q_2^2) = \frac{\alpha_s(Q_1^2)}{1 + \alpha_s(Q_1^2)\beta_0 \ln(Q_2^2/Q_1^2)} \quad (5.3.32)$$

dove abbiamo implicitamente assunto che entrambe le scale di energia Q_1 e Q_2 siano nel dominio perturbativo, così che il troncamento della funzione beta al primo ordine sia affidabile. Poiché $\beta_0 > 0$, il denominatore cresce con Q_2 (fissata la scala Q_1^2), da cui $\alpha_s(Q_2^2)$ decresce. Vedi la figura 5.2.

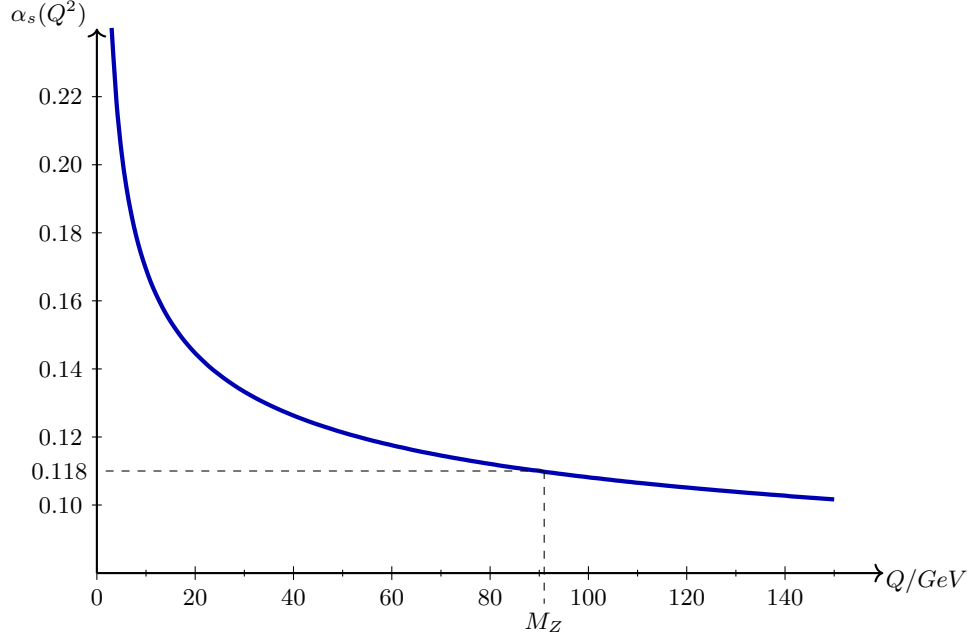


Figura 5.2: Andamento del running coupling rispetto la scala di energia Q . Teniamo a mente che l'andamento $\alpha_s(Q^2)$ è stato trovato dalla teoria, ma i singoli valori sono stati verificati sperimentalmente.

Dall'espressione (5.3.32) sopra riportata, se poniamo $Q_1 = \mu$, $Q_2 = Q$, e $\alpha_s(\mu^2) = \alpha_s$, è immediato verificare che:

$$\frac{\partial \alpha_s(Q^2)}{\partial \alpha_s} = \frac{\beta(\alpha_s(Q^2))}{\beta(\alpha_s)}$$

come prescritto dall'eq. (5.3.14). Inoltre notiamo che l'eq. (5.3.32) può essere riscritta come una serie geometrica:

$$\alpha_s(Q_2^2) = \alpha_s(Q_1^2) \sum_{k=0}^{\infty} \left[-\alpha_s(Q_1^2) \beta_0 \ln(Q_2^2/Q_1^2) \right]^k \quad (5.3.33)$$

il che significa che l'uso dell'accoppiamento $\alpha_s(Q_2)$ in un calcolo induce automaticamente la risommazione di una serie *infinita* di termini logaritmici $\alpha_s^k(Q_1^2) \ln^k(Q_2^2/Q_1^2)$, cioè a tutti gli ordini nell'accoppiamento di riferimento $\alpha_s(Q_1^2)$. Dunque, riscriviamo α_s ad una specifica scala, in teoria delle perturbazioni, in termini di se stessa valutata ad un'altra scala, anch'essa generica, Q_1^2 .

Quando le due scale sono vicine i termini logaritmici sono piccoli e la serie converge perturbativamente, ma se le due scale sono molto lontane, allora i $\ln$ sono grandi, non possiamo fare teoria delle perturbazioni e dobbiamo tenerci tutti gli infiniti termini.

Il valore dell'accoppiamento forte non può essere estratto solo dalla teoria: è necessario un input sperimentale, vale a dire la misurazione dell'accoppiamento a una qualche scala di riferimento Q_1 , a partire dalla quale la relazione nell'equazione (5.3.32) permette di prevedere perturbativamente il valore dell'accoppiamento a qualsiasi altra scala perturbativa Q_2 . Includendo i coefficienti di ordine superiore della funzione beta della QCD, come β_1 e oltre, si ottiene una determinazione più precisa di α_s a qualsiasi scala perturbativa, partendo da un dato valore di input a una scala di riferimento. Scritta come una serie geometrica, l'inclusione di β_1 permette la risommarione dei termini logaritmici sub-leading $\alpha_s^k(Q_1^2) \ln^{k-1}(Q_2^2/Q_1^2)$, e così via per i coefficienti superiori della funzione beta.

L'input sperimentale che viene solitamente scelto è l'accoppiamento valutato alla massa del bosone Z , $M_Z \sim 91.2$ GeV, che sperimentalmente risulta essere $\alpha_s(M_Z^2) \sim 0.118$. L'accoppiamento della QCD al di sotto della scala EW è quindi molto (almeno dieci volte) più grande dell'accoppiamento della QED $\alpha \sim 1/137$. Se l'accoppiamento della QCD viene sondato a scale inferiori, esso cresce significativamente, diventando grande quanto 0.3 a $Q_2 \sim 1$ GeV, dunque facendoci uscire dal regime perturbativo. Viceversa, ad energie più elevate, α_s diventa sempre più piccolo, il che implica che i processi di scattering della QCD che coinvolgono grandi scale di energia sono debolmente accoppiati; ciò ha portato al voler costruire collider sempre più ad alta energia. Le grandi energie di collisione permettono quindi di sondare direttamente la dinamica delle particelle a distanze sempre più piccole, svelando potenzialmente nuova fisica oltre il Modello Standard. D'altra parte, tecnicamente, esse permettono di ottenere predizioni della QCD più affidabili a ordine fissato nella teoria perturbativa.

Siamo ora in grado di commentare la μ -dipendenza dei calcoli teorici troncati a un dato ordine nella teoria perturbativa: poiché l'indipendenza dei risultati fisici da μ è garantita solo tenendo conto di tutti gli ordini perturbativi, e a ogni ordine fissato c'è da aspettarsi una μ -dipendenza.

In effetti, possiamo definire l'osservabile R e la sua indipendenza dalla scala come segue:

$$R \equiv \sum_{k=0}^{\infty} c_k(\mu^2) \alpha_s^k(\mu^2), \quad (5.3.34)$$

$$\mu^2 \frac{dR}{d\mu^2} = 0 = \sum_{k=0}^{\infty} \left(\mu^2 \frac{dc_k(\mu^2)}{d\mu^2} \alpha_s^k(\mu^2) + c_k(\mu^2) k \alpha_s^{k-1}(\mu^2) \beta(\alpha_s(\mu^2)) \right). \quad (5.3.35)$$

Se limitiamo l'analisi al solo scorrimento (*running*) dell'ordine più basso

β_0 , l'equazione precedente risulta soddisfatta se:

$$\frac{dc_0(\mu^2)}{d\mu^2} = 0, \quad \mu^2 \frac{dc_{k+1}(\mu^2)}{d\mu^2} = c_k(\mu) k \beta_0, \quad \forall k \geq 0. \quad (5.3.36)$$

Tuttavia, quando la previsione per l'osservabile R viene troncata a un determinato ordine perturbativo N , ovvero:

$$R_N = \sum_{k=0}^N c_k(\mu^2) \alpha_s^k(\mu^2), \quad (5.3.37)$$

la dipendenza di R_N dalla scala μ è espressa dalla relazione:

$$\begin{aligned} \mu^2 \frac{dR_N}{d\mu^2} &= \sum_{k=0}^N \left(\mu^2 \frac{dc_k(\mu^2)}{d\mu^2} \alpha_s^k(\mu^2) + c_k(\mu^2) k \alpha_s^{k-1}(\mu^2) \beta(\alpha_s(\mu^2)) \right) \\ &= \mu^2 \frac{dc_0(\mu^2)}{d\mu^2} + \sum_{k=0}^{N-1} \mu^2 \frac{dc_{k+1}(\mu^2)}{d\mu^2} \alpha_s^{k+1}(\mu^2) \\ &\quad - \sum_{k=0}^{N-1} c_k(\mu^2) k \alpha_s^{k+1}(\mu^2) \beta_0 - c_N(\mu^2) N \alpha_s^{N+1}(\mu^2) \beta_0 \\ &= -c_N(\mu^2) N \alpha_s^{N+1}(\mu^2) \beta_0, \end{aligned} \quad (5.3.38)$$

dove la somma dei primi tre termini nell'espressione intermedia svanisce in virtù delle condizioni definite precedentemente. L'ultimo termine svanirebbe solo includendo la derivata di c_{N+1} rispetto a μ , ma c_{N+1} rappresenta l'ordine più basso non incluso nel calcolo di R_N .

Il punto cruciale di questa analisi è che la dipendenza dalla scala di un'osservabile R calcolata all'ordine $\mathcal{O}(\alpha_s^N)$ nella teoria perturbativa si manifesta solo all'ordine successivo, $\mathcal{O}(\alpha_s^{N+1})$. Mentre nell'esempio precedente abbiamo considerato solo lo scorrimento all'ordine più basso β_0 , in realtà i coefficienti c_k dipendono da μ anche attraverso β_1, β_2 e così via; in questi casi, la dipendenza dalla scala di R_N governata dai coefficienti β_j entra all'ordine $\mathcal{O}(\alpha_s^{N+j+1})$. Notiamo che se non fosse così non sarebbe possibile ottenere una completa μ -indipendenza includendo tutti gli ordini perturbativi.

Abbiamo appena osservato che, in un calcolo che include correzioni fino all'ordine $\mathcal{O}(\alpha_s^N)$, diverse scelte della scala di rinormalizzazione μ — pur essendo tutte formalmente valide — portano a differenze dell'ordine di $\mathcal{O}(\alpha_s^{N+1})$ nel risultato finale. Tali discrepanze sono dello stesso ordine delle correzioni di ordine superiore mancanti nel calcolo. Di conseguenza, la dispersione dei risultati numerici ottenuti variando μ fornisce una *stima euristica* della potenziale entità dei termini trascurati nella serie perturbativa.

Nella pratica, i calcoli perturbativi vengono solitamente eseguiti impostando $\mu = Q$, dove Q rappresenta una scala di energia caratteristica del

processo o dell'osservabile in esame. Per stimare l'incertezza teorica associata al troncamento dell'espansione, il calcolo viene ripetuto con $\mu = 2Q$ e $\mu = Q/2$. L'involuppo dei tre risultati viene adottato come *fascia di incertezza teorica* associata alla previsione. Vedi la figura 5.3. Ci si aspetta che le correzioni di ordine superiore rientrino nella fascia di incertezza dell'ordine precedente (grafico sulla destra della figura 5.3). Inoltre, poiché l'ampiezza della fascia scala come α_s^{N+1} , essa tende a ridursi man mano che vengono inclusi più ordini, ovvero al crescere di N . Ad ogni modo, per ogni N ci aspettiamo che l'osservabile R_N avrà associato uno spread ai valori calcolati, e tale ampiezza ci darà proprio una stima (teorica) della dipendenza del parametro dalla scelta della scala fisica.

Proprio perché lo spread è dello stesso ordine del primo contributo che stiamo trascurando possiamo dire che una sua stima ci fornisce l'errore teorico della nostra predizione. In questo modo si può intuire se è sensato procedere con i conti, oppure se la stima ottenuta è sufficientemente precisa (anche in base all'esperimento che si sta confrontando).

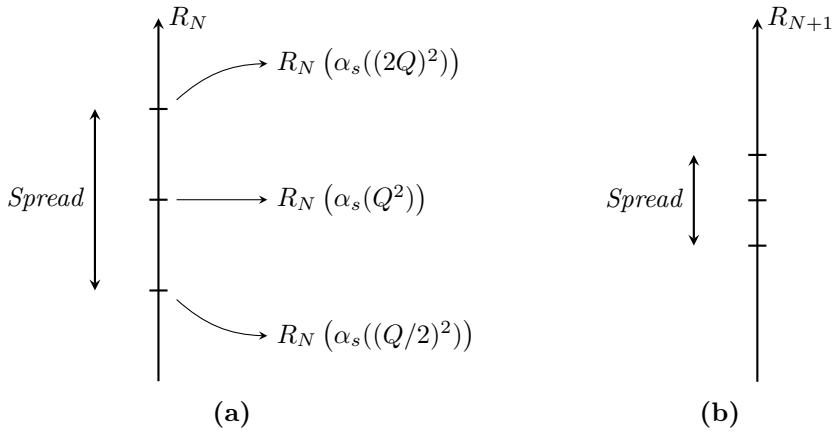


Figura 5.3

Se consideriamo ad esempio la predizione per un osservabile fisico R calcolato al primo ordine nella QCD, dunque $\mathcal{O}(\alpha_s^1)$:

$$R(\alpha_s(\mu^2)) = c_0 + \alpha_s(\mu^2)c_1 + \mathcal{O}(\alpha_s^2) \quad (5.3.39)$$

allora possiamo usare la (5.3.32) per riscrivere $\alpha_s(\mu^2)$, cioè ad una specifica scala, in termini di $\alpha_s(Q^2)$ ad un'altra scala:

$$R(\alpha_s(\mu^2)) = c_0 + \frac{\alpha_s(Q^2)}{1 + \alpha_s(Q^2)\beta_0 \ln(\mu^2/Q^2)} c_1 + \mathcal{O}(\alpha_s^2) \quad (5.3.40)$$

$$= c_0 + \alpha_s(Q^2)c_1 - \alpha_s^2(Q^2)c_1\beta_0 \ln(\mu^2/Q^2) + \mathcal{O}(\alpha_s^2) \quad (5.3.41)$$

$$= c_0 + \alpha_s(Q^2)c_1 + \mathcal{O}(\alpha_s^2) \quad (5.3.42)$$

dove in (5.3.41) abbiamo scritto la frazione come una serie geometrica (troncata all'ordine che ci interessa). All'accuratezza $\mathcal{O}(\alpha_s)$ alla quale stiamo eseguendo il nostro calcolo di R , l'uso di qualsiasi argomento dell'accoppiamento, sia esso μ^2 come in (5.3.43) o Q^2 come nell'ultima riga di (5.3.42), è formalmente permesso: scelte diverse portano solo differenze di $\mathcal{O}(\alpha_s^2)$ nel risultato; dunque non ci preoccupiamo particolarmente di scegliere la scala. Come predetto in precedenza. Infatti, la μ -dipendenza del risultato calcolato a $\mathcal{O}(\alpha_s)$ entra proprio a $\mathcal{O}(\alpha_s^2)$. Questo argomento può essere iterato agli ordini superiori. Per esempio, se consideriamo l'espansione al second'ordine di R :

$$R(\alpha_s(\mu^2)) = c_0 + \alpha_s(\mu^2)c_1 + \alpha_s^2(\mu^2)c_2(\mu) + \mathcal{O}(\alpha_s^3) \quad (5.3.43)$$

il coefficiente c_2 avrà la seguente struttura:

$$c_2(\mu) = \tilde{c}_2 + c_1\beta_0 \ln(\mu^2/Q^2) \quad (5.3.44)$$

con $\tilde{c}_2$ indipendente da μ . Ovvero $c_2(\mu)$ presenterà un'esplicita μ -dipendenza, in modo tale da cancellare esattamente quella indotta a $\mathcal{O}(\alpha_s^2)$ dal running dell'accoppiamento nella seconda riga di (5.3.42). La μ -dipendenza residua nel calcolo a questo punto sarà di $\mathcal{O}(\alpha_s^3)$, e così via per gli ordini superiori. È proprio questa cancellazione di dipendenze tra ordini successivi a rendere possibile la non dipendenza dell'osservabile esatta dalla scala.

Teniamo a mente che l'incertezza viene dalla mancanza di tutti gli ordini perturbativi superiori. Solitamente ci si riferisce a tale incertezza con *Missing Higher Order Uncertainty* (MHOU).

L'accoppiamento della QCD può anche essere convenientemente espresso in termini di una scala di riferimento, Λ , che segna l'energia alla quale l'espressione perturbativa per l'accoppiamento, come l'equazione (5.3.32), formalmente diverge. Nota che per questo motivo Λ varia leggermente a seconda dell'ordine perturbativo che consideriamo. Teniamo a mente che l'accoppiamento fisico non diverge effettivamente, bensì lo fa l'equazione (5.3.32) che è solo un'approssimazione derivata all'interno della teoria perturbativa. Come tale, essa diventa singolare precisamente alla scala in cui la teoria perturbativa stessa viene meno, ovvero a Λ .

Infatti, avevamo già notato, nella figura 5.2, che a grandi energie il coupling è piccolo e possiamo fare teoria delle perturbazioni, ma nel momento in cui scendiamo nella scala, il coupling si ingrandisce e tende a diventare molto grande, facendoci uscire dal regime perturbativo. Per cui, in tale regime, non possiamo più calcolare la funzione β semplicemente fermandoci al primo ordine perturbativo, ma dobbiamo prendere tutti i contributi.

Usando la (5.3.32), il valore di Λ si deduce risolvendo (denominatore uguale a zero):

$$1 + \alpha_s(Q^2)\beta_0 \ln(\Lambda^2/Q^2) = 0 \quad \implies \quad \alpha_s(Q^2) = \frac{1}{\beta_0 \ln(Q^2/\Lambda^2)}. \quad (5.3.45)$$

Λ_{QCD} è la scala a cui α_s diventa grande e Q^2 rimpicciolisce. Infatti, la relazione (5.3.45) ci mostra che $\alpha_s(Q^2)$ diventa piccolo con $\log^{-1}(Q^2/\mu^2)$ per larghi Q^2 .

Con $Q = M_Z$ (che è la tipica scala a cui si estrae il coupling), $\alpha_s(M_Z^2) = 0.118$ e $N_f = 5$ si ottiene $\Lambda \sim 90$ MeV. Se si includono invece gli effetti da β_1 , l'equazione definente per Λ diventa⁵:

$$\frac{1}{\alpha_s(Q^2)} + \beta_1 \ln \left(\frac{\beta_1 \alpha_s(Q^2)}{1 + \beta_1 \alpha_s(Q^2)} \right) = \beta_0 \ln(Q^2/\Lambda^2) \quad (5.3.46)$$

che risulta in $\Lambda \sim 200$ MeV.

Non sorprende che Λ sia dell'ordine della massa degli stati adronici più leggeri (pioni, nucleoni, ...): a un'energia di $\mathcal{O}(\Lambda)$, la QCD diventa una teoria a forte accoppiamento, da cui i gradi di libertà fondamentali di quark e gluoni (che a scale più alte erano gli stati fondamentali) iniziano a essere confinati in stati legati, la cui energia di legame è dell'ordine della loro massa, ovvero $\mathcal{O}(\Lambda)$.

Per le ragioni sopra esposte, il parametro Λ è chiamato *scala non perturbativa* (per sottolineare che a scale di $\mathcal{O}(\Lambda)$ la teoria perturbativa della QCD non è affidabile), o scala di *confinamento* (per evidenziare che è la scala alla quale avviene il confinamento), o *polo di Landau* (per mettere in risalto la singolarità nell'espressione perturbativa dell'accoppiamento corrente).

⁵Nota che l'espressione, come già detto, dipende dall'ordine perturbativo in cui espandiamo la beta function, e in linea di principio, se (fossimo in grado) includessimo tutti i contributi dell'espansione perturbativa, non avremmo una divergenza.

Capitolo 6

QCD in scattering process

Ora che abbiamo visto come storicamente è stata formulata la QCD, e che abbiamo visto i fondamenti di questo pezzo di Modello Standard, possiamo effettivamente buttarci a fare i conti e calcolare sezioni d'urto.

Vedremo in questo capitolo, in particolare, la produzione di adroni dall'annichilazione elettrone-positrone, il quale è uno dei processi di scattering fondamentali che hanno aiutato a stabilire la QCD come teoria delle interazioni forti. In particolare, ha dato fiducia nell'esistenza del grado di libertà del colore, così come della carica frazionaria dei sapori dei quark.

Vedremo prima una visione intuitiva su come analizzare questo processo, successivamente ci immergeremo nei conti di elementi di matrice soffermandoci in particolare sulle divergenze infrarosse che emergono.

Nel corso del capitolo analizzeremo oggetti fondamentali per i processi di scattering di materia adronica, quali il concetto di *jets* e di *distribuzioni cinematiche*.

Nella parte finale del capitolo parleremo del *deep inelastic scattering* e di come questo permetta di sondare, e studiare, la struttura del protone.

6.1 Scattering process $e^+e^- \rightarrow$ Adroni

Vediamo in questa sezione tutti i conti (più o meno) riguardanti il processo:

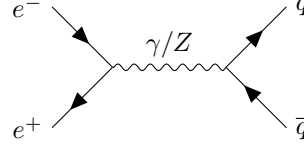
$$e^+ + e^- \longrightarrow \text{adroni} \quad (6.1.1)$$

con tutti i commenti riguardanti eventuali divergenze e rinormalizzazioni.

6.1.1 Rappresentazione intuitiva dello scattering

Abbiamo già detto che il processo avviene tramite l'annichilazione di e^+ ed e^- , che produce un fotone virtuale (o un bosone Z virtuale), il quale

successivamente produce uno stato adronico. Al prim'ordine della teoria perturbativa (LO), il diagramma che rappresenta lo scattering è:



Infatti il modo più semplice (ma non l'unico) di produrre stati adronici, ovviamente, è tramite gli elementi fondamentali della QCD, ovvero quark e gluoni.

Se l'energia nel centro di massa è $\sqrt{s}$, questa è anche l'energia della coppia $q\bar{q}$ prodotta. Possiamo assumere $\sqrt{s} \sim 100$ GeV per fissare la nostra intuizione fisica.

Dunque, se $\sqrt{s} \sim 100$ GeV, allora abbiamo uno spazio fasi grande da riempire e per questo la coppia $q\bar{q}$, prodotta ad alta energia, ha un'alta probabilità di emettere ulteriore radiazione, data principalmente¹ da gluoni (data l'entità della costante di accoppiamento α_s), ognuno dei quali sottrae una certa frazione di energia alla coppia $q\bar{q}$ originale. Diagrammaticamente, queste emissioni successive sono rappresentate schematicamente nella figura 6.1, in cui E_i (tale che $E_{i+1} < E_i$) rappresenta le scale di energia a cui avvengono le emissioni:

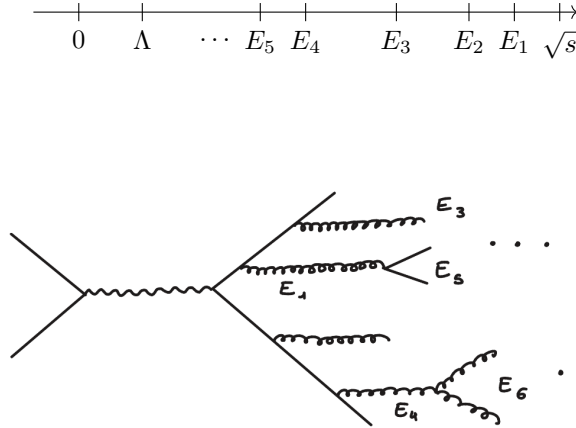


Figura 6.1: Diagramma di Feynman utilizzato per mostrare il parton shower, con etichette energetiche basate sull'evoluzione della scala E_i .

¹Infatti, anche la produzione di fotoni (virtuali o reali) è possibile, ma fortemente soppressa viste le costanti di accoppiamento della QCD e QED. C'è comunque una branca della ricerca che si occupa di fare teoria delle perturbazioni per processi misti, ovvero in cui sono presenti sia termini di QCD, che di QED.

Alcune di queste emissioni possono essere virtuali, il che significa che comportano sia la radiazione che il riassorbimento di partoni (quark e gluoni) all'interno del diagramma.

Lo schema del diagramma di Feynman, disegnato in figura 6.1, descrive la fase perturbativa del processo di scattering, dove $E_i > \Lambda$. Infatti, se $E_{i+1} < E_i$, allora continuiamo a scendere di energia e andiamo sempre di più verso l'energia in cui la QCD diventa forte (forte accoppiamento), dunque dove la teoria non è più perturbativa.

Possiamo intuire l'evoluzione della propagazione del colore durante la produzione di partoni rappresentando graficamente il flusso di colore *leading* associato al diagramma; vedi la figura 6.2a. Da ciò, possiamo vedere come la linea di colore primaria associata alla coppia $q\bar{q}$ (il *dipolo di colore*) ad alta energia $\sqrt{s} \gg \Lambda$ venga frammentata attraverso emissioni successive. Questo processo forma molteplici dipoli di colore a energie più basse, il che è energeticamente favorito.

Raggiunta la scala non perturbativa $\sim \Lambda$, i dipoli di colore si ricombinano (*color reconnection*) per formare singoletti di colore, ovvero adroni, che rappresentano gli stati legati che sopravvivono asintoticamente. Infatti, nel capitolo precedente abbiamo visto come nella regione non perturbativa i costituenti fondamentali della QCD (quark e gluoni) si ricombinano in modo da formare stati legati. Nella precedente immagine del colore, la situazione è rappresentata nella figura 6.2b, in cui i "blobs", ognuno associato a una scala di energia dell'ordine di Λ , rappresentano la formazione di adroni dalle linee di colore lasciate aperte nella fase perturbativa.

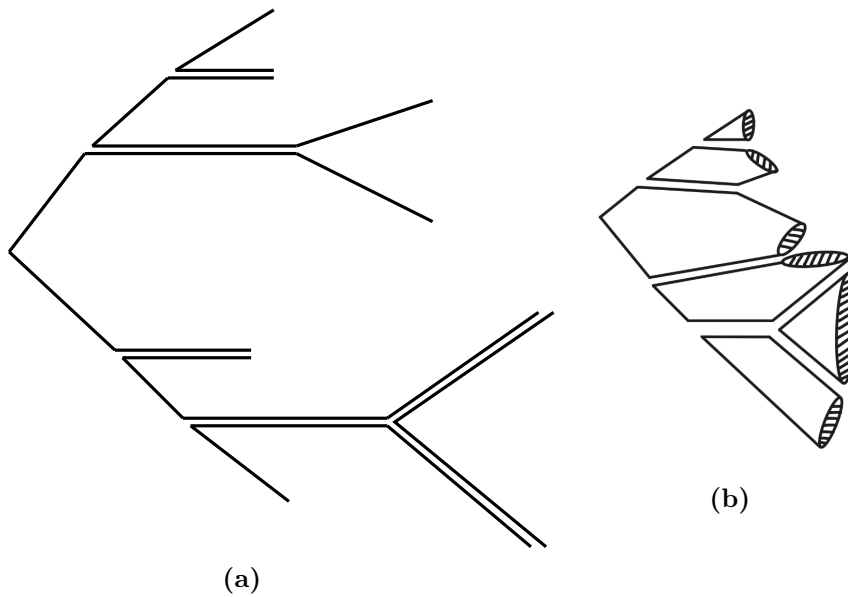
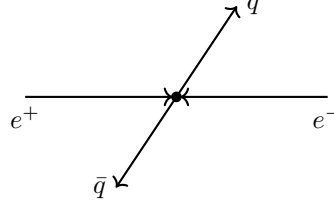


Figura 6.2

Dalla prospettiva dello scattering all'interno di un collisore, possiamo visualizzare geometricamente la produzione primaria della coppia $q\bar{q}$ "hard" (cioé back-to-back):



La radiazione successiva di gluoni a energie $E_i < \sqrt{s}$ modifica la direzione e l'energia di q e $\bar{q}$ (progressivamente meno al diminuire di E_i). Pertanto, un evento di scattering sarà rappresentato più realisticamente dallo schema 6.3a, dove le linee dei gluoni sono più o meno collimate rispetto al quark (o $\bar{q}$) inizialmente prodotto (energetico) a seconda della scala corrispondente E_i .

Dopo la fase di adronizzazione, gli adroni finali vengono rilevati nei calorimetri adronici del rivelatore del collisore sotto forma di depositi di energia, i **jet**, che corrispondono a getti di particelle adroniche piuttosto collimati ed energetici. Vedi la figura 6.3b.

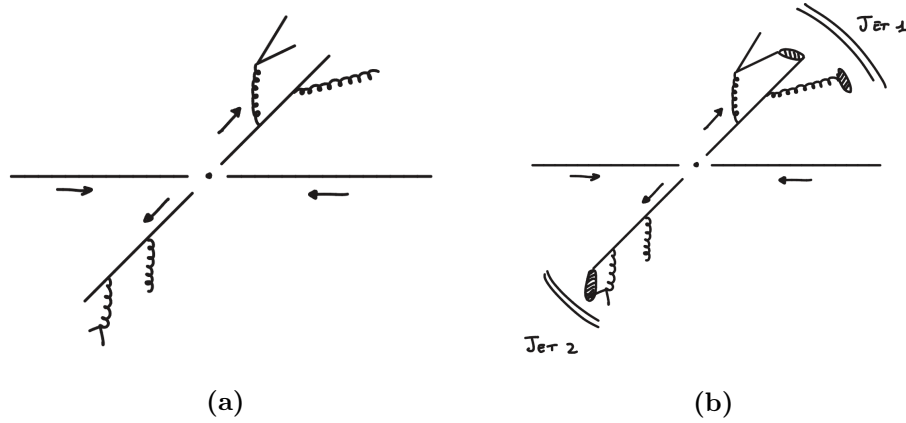


Figura 6.3

Nell'ultima rappresentazione grafica (figura 6.3b), si è assunto per semplicità che le energie E_i fossero basse (o gli angoli di emissione relativi così piccoli) che tutte le emissioni ricadessero in due jet. È possibile che una o più emissioni dalla coppia $q\bar{q}$ siano "hard" e formino altri jet, separati da quelli principali. La probabilità di questo fenomeno può e sarà calcolata in seguito (capitolo §6.2) nella teoria perturbativa della QCD.

Possiamo distinguere le diverse fasi: siamo in regime perturbativo nel momento in cui facciamo scontrare le particelle nello stato iniziale, producono

un bosone virtuale, esso decade in uno stato finale di energia $\sqrt{s}$, il quale può emettere gluoni (o in modo meno probabile altri bosoni) virtuali o reali, che man mano sottraggono energia ai quark "principali". In questa fase abbiamo quelli che si chiamano *partoni* (figura 6.3a). Dopo la formazione di partoni (accumulo di quark, gluoni reali o no), avviene l'*adronizzazione*, cioè passiamo alla fase non perturbativa e i gradi di libertà della QCD devono formare obbligatoriamente degli stati legati (figura 6.3b). Dunque, lo stato finale dell'esperimento è un insieme di adroni.

6.1.2 $e^+e^- \rightarrow \text{adroni}$ al Leading Order (LO)

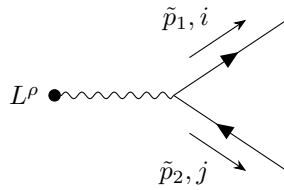
Iniziamo la discussione di questo processo dalla sua predizione all'ordine dominante, ovvero l'approssimazione di Born. Calcoliamo il contributo di Born considerando solo lo scambio di un fotone nel canale s , trascurando cioè lo scambio del bosone Z . Quest'ultimo contributo non può essere trascurato per un'energia nel centro di massa $\sqrt{s} \sim M_Z$, quindi implicitamente stiamo assumendo di eseguire il calcolo ad energie molto più piccole. Infatti, se siamo $\sqrt{s} \ll M_Z$ allora il bosone è molto off-shell e il contributo del suo propagatore è trascurabile. Puoi ricordare le considerazioni fatte nella sezione §3.6.2 in cui abbiamo parlato delle masse dei bosoni di gauge EW, oppure guardare l'Appendice B per il conto completo, con anche il contributo dello Z .

Inoltre, trascuriamo sistematicamente le masse dei quark e dei leptoni, così da poter utilizzare le relazioni ben note:

$$\sum_s u_\alpha \bar{u}_\beta = (\not{p} + m)_{\alpha\beta} \sim \not{p} \quad (6.1.2)$$

$$\sum_s v_\alpha \bar{v}_\beta = (\not{p} - m)_{\alpha\beta} \sim \not{p}. \quad (6.1.3)$$

Siamo dunque nella situazione:



in cui L^ρ rappresenta la corrente leptonica:

$$L^\rho = -\frac{i}{s} \bar{v}(q_2)(-ie\gamma^\rho)u(q_1) \quad (6.1.4)$$

con e è la carica del positrone, q_1 e q_2 sono i momenti dell'elettrone e del positrone, tali che $q_i^2 = 0$, e $s = (q_1 + q_2)^2$ è l'energia nel centro di massa al quadrato. Nel diagramma $\tilde{p}_i$ sono i momenti associati a quark e anti-quark. Nota che in (6.1.4) abbiamo incluso anche il propagatore del fotone. Gli indici i, j sono indici di colore nella rappresentazione fondamentale.

Possiamo fare i conti in modo leggermente diverso rispetto quanto imparato al corso di *Fondamenti di QFT*, ma credo siano più scorrevoli in questo modo. Si può comunque provare a svolgere tutti i conti con tutti gli indici spinoriali espliciti.²

Introduciamo anche il tensore leptónico, in cui mediamo sulle polarizzazioni (essendo i leptoni nello stato iniziale):

$$L^{\rho\sigma} = \frac{1}{4} \sum_{\text{pol. lep.}} L^\rho L^{\sigma*} \quad (6.1.6)$$

$$= \frac{e^2}{4s^2} \text{Tr} \{ \not{q}_1 \gamma^\sigma \not{q}_2 \gamma^\rho \} \quad (6.1.7)$$

$$= \frac{e^2}{s^2} (q_1^\rho q_2^\sigma + q_2^\rho q_1^\sigma - q_1 \cdot q_2 g^{\rho\sigma}) \quad (6.1.8)$$

in cui abbiamo utilizzato la ciclicità della traccia e il risultato, noto, della traccia di 4 matrici γ^μ .

Per la produzione di un quark di sapore f e carica elettrica Q_f (ad es. $Q_u = 2/3, Q_d = -1/3, \dots$), introducendo una δ_{ij} di colore nel vertice, il contributo di Born in $d = 4 - 2\epsilon$ dimensioni (sarà utile aver generalizzato la dimensione per ciò che verrà nel corso del capitolo), sommando sulle polarizzazioni dei quark (essendo lo stato finale), è:

$$|\mathcal{M}_B|^2 = L^{\rho\sigma} \sum_{\text{pol. quark}} \bar{u}(\tilde{p}_1) (ieQ_f \gamma_\rho \delta_{ij}) v(\tilde{p}_2) \bar{v}(\tilde{p}_2) (-ieQ_f \gamma_\sigma \delta_{ji}) u(\tilde{p}_1) \quad (6.1.9)$$

$$= \frac{e^4 Q_f^2}{s^2} N (q_1^\rho q_2^\sigma + q_2^\rho q_1^\sigma - q_1 \cdot q_2 g^{\rho\sigma}) \text{Tr} \{ \not{\tilde{p}}_1 \gamma_\rho \not{\tilde{p}}_2 \gamma_\sigma \} \quad (6.1.10)$$

$$= \frac{4e^4 Q_f^2}{s^2} N (q_1^\rho q_2^\sigma + q_2^\rho q_1^\sigma - q_1 \cdot q_2 g^{\rho\sigma}) \times \\ \times (\tilde{p}_{1\rho} \tilde{p}_{2\sigma} + \tilde{p}_{2\rho} \tilde{p}_{1\sigma} - \tilde{p}_1 \cdot \tilde{p}_2 g_{\rho\sigma}) \quad (6.1.11)$$

$$= \frac{8e^4 Q_f^2}{s^2} N (\tilde{p}_1 \cdot q_1 \tilde{p}_2 \cdot q_2 + \tilde{p}_2 \cdot q_1 \tilde{p}_1 \cdot q_2 - \epsilon q_1 \cdot q_2 \tilde{p}_1 \cdot \tilde{p}_2) \quad (6.1.12)$$

dove abbiamo indicato con $\delta_{ij} \delta_{ji} \equiv N = 3$ il numero di colori dei quark³ e ci siamo ricordati del fatto che valga $g^{\rho\sigma} g_{\rho\sigma} = d$.

²Se si vuole implementare il calcolo si può utilizzare l'equazione di partenza:

$$i\mathcal{M}_B = \bar{v}(q_2) (-ie\gamma^\rho) u(q_1) \frac{i(-ig_{\rho\sigma})}{s} \bar{u}(\tilde{p}_1) (i\gamma^\alpha eQ_f \delta_{ij}) v(\tilde{p}_2). \quad (6.1.5)$$

³Notiamo che il fattore di colore è stato dedotto come segue. Tenendo traccia degli indici di colore per gli spinori dei quark, l'elemento di matrice di Born al quadrato sarebbe:

$$|\mathcal{M}_B|^2 \propto \bar{u}_i(\tilde{p}_1) \delta_{ij} v_j(\tilde{p}_2) \bar{v}_k(\tilde{p}_2) \delta_{kl} u_l(\tilde{p}_1),$$

dove si sottintende la somma sugli indici di colore ripetuti. La contrazione $v_j(\tilde{p}_2) \bar{v}_k(\tilde{p}_2)$ produce una delta di Kronecker di colore δ_{jk} , e analogamente la contrazione $u_l(\tilde{p}_1) \bar{u}_i(\tilde{p}_1)$ produce δ_{li} , quindi il fattore di colore è $\delta_{ij} \delta_{jk} \delta_{kl} \delta_{li} = \delta_{ij} \delta_{ji} = N$.

Parametrizzando:

$$q_1 = \frac{\sqrt{s}}{2}(1, 0, 0, 1) \quad , \quad q_2 = \frac{\sqrt{s}}{2}(1, 0, 0, -1) \quad (6.1.13)$$

$$\tilde{p}_1 = \frac{\sqrt{s}}{2}(1, 0, \sin \tilde{\theta}, \cos \tilde{\theta}) \quad , \quad \tilde{p}_2 = \frac{\sqrt{s}}{2}(1, 0, -\sin \tilde{\theta}, -\cos \tilde{\theta}) \quad (6.1.14)$$

si ottengono le relazioni:

$$q_1 \cdot q_2 = \frac{s}{2} \quad (6.1.15)$$

$$q_i \cdot \tilde{p}_i = s(1 - \cos \tilde{\theta})/4 \quad (6.1.16)$$

$$q_i \cdot \tilde{p}_j = s(1 + \cos \tilde{\theta})/4 \quad (6.1.17)$$

ma soprattutto:

$$|\mathcal{M}_B|^2 = e^4 Q_f^2 N(1 - 2\epsilon + \cos^2 \tilde{\theta}) \quad (6.1.18)$$

Sappiamo che per ottenere la sezione d'urto ci serve integrare nello spazio delle fasi, per cui ci serve la sua misura. Possiamo vedere che lo spazio delle fasi a 2 corpi in $d = 4 - 2\epsilon$ dimensioni, ricordando che il d -momento ha 1 componente energetica e $d - 1$ spaziali, è:

$$d\Phi_2^{(d)}(\tilde{p}_1, \tilde{p}_2; s) = \frac{d^{d-1}\tilde{p}_1}{(2\pi)^{d-1}2\tilde{E}_1} \frac{d^{d-1}\tilde{p}_2}{(2\pi)^{d-1}2\tilde{E}_2} (2\pi)^d \delta^{(d)}(Q - \tilde{p}_1 - \tilde{p}_2) \quad (6.1.19)$$

possiamo utilizzare la decomposizione della delta:

$$\delta^{(d)}(Q - \tilde{p}_1 - \tilde{p}_2) = \delta^{(d-1)}(\vec{q}_1 + \vec{q}_2 - \vec{p}_1 - \vec{p}_2) \delta(E_1 + E_2 - \tilde{E}_1 - \tilde{E}_2) \quad (6.1.20)$$

di cui possiamo utilizzare la prima parte per integrare in $d^{(d-1)}\tilde{p}_2$, ponendoci nel sistema di riferimento del centro di massa ($\vec{q}_1 + \vec{q}_2 = 0$), e dunque settando $\vec{p}_1 = -\vec{p}_2$. In più possiamo anche scrivere $E_1 + E_2 = \sqrt{s}$ e notare che $|\vec{p}_1| = |\vec{p}_2| = \tilde{E}_1 = \tilde{E}_2$ dal momento che stiamo considerando i fermioni massless. Abbiamo quindi:

$$d\Phi_2^{(d)}(\tilde{p}_1, \tilde{p}_2; s) = \frac{\tilde{E}_1^{d-2} d\tilde{E}_1 \sin \tilde{\theta}_1^{d-3} d\cos \tilde{\theta}_1 d\Omega_1^{(d-2)}}{(2\pi)^{d-2} 4\tilde{E}_1^2} \delta(\sqrt{s} - 2\tilde{E}_1) \quad (6.1.21)$$

$$= \frac{(\sqrt{s}/2)^{d-4} \sin \tilde{\theta}_1^{d-4} d\cos \tilde{\theta}_1 (2\pi^{d/2-1})}{(2\pi)^{d-2} 8\Gamma(d/2 - 1)} \quad (6.1.22)$$

$$= \left(\frac{16\pi}{s}\right)^\epsilon \frac{\sin \tilde{\theta}_1^{-2\epsilon} d\cos \tilde{\theta}_1}{\Gamma(1 - \epsilon) 16\pi} \quad (6.1.23)$$

in cui abbiamo integrato in $d\tilde{E}_1$ utilizzando la δ , abbiamo cambiato la dipendenza della dimensione $d = 4 - 2\epsilon$, utilizzato la misura:

$$d^{d-1}p_i = |\vec{p}_i|^{d-2} d|\vec{p}_i| d\Omega_i^{(d-1)}$$

e:

$$\begin{aligned} d\Omega_i^{(d-1)} &= d\Omega_i^{(d-2)} \sin \theta_i^{d-3} d\theta_i \\ &= d\Omega_i^{(d-2)} \sin \theta_i^{d-4} d \cos \theta_i \\ \int d\Omega_i^{(d)} &= \frac{2\pi^{d/2}}{\Gamma(d/2)}. \end{aligned}$$

Quindi per $N = 3$ colori, e sommando sui sapori, la sezione d'urto di Born è:

$$\sigma_B = \frac{N}{2s} \int d\Phi_2^{(d)}(\tilde{p}_1, \tilde{p}_2; s) \sum_f |\mathcal{M}_B|^2 \quad (6.1.24)$$

$$= \frac{4\pi\alpha^2}{3s} \sum_f Q_f^2 \left(\frac{4\pi}{s} \right)^\epsilon \frac{3N(1-\epsilon)\Gamma(2-\epsilon)}{(3-2\epsilon)\Gamma(2-2\epsilon)} \quad (6.1.25)$$

in cui possiamo individuare un termine di flusso ($N/(2s)$) e in cui abbiamo riscritto il risultato in termini della costante di struttura fina $\alpha = e^2/(4\pi)$.

Possiamo anche vedere che (6.1.25) in quattro dimensioni, dunque facendo $\epsilon \rightarrow 0$, si riduce a:

$$\sigma_B(e^+e^- \rightarrow \text{adroni}) = \frac{4\pi\alpha^2}{3s} \sum_f Q_f^2 N \quad (6.1.26)$$

in cui possiamo individuare un fattore cinematico ($4\pi\alpha/(3s)$), la carica di flavour ($\sum_f Q_f^2$) e il numero di colori (N).

Il risultato sopra riportato, che è di pura QED non essendoci nessuna correzione di QCD apparte il fattore N che viene dalle δ_{ij} dei vertici, se verificato sperimentalmente può fornire una conferma del valore di N (o per lo meno la sua esistenza).

Possiamo confrontare (6.1.26) con la sezione d'urto di Born per la produzione di una coppia muone-antimuone (senza massa), che è:

$$\sigma_B(e^+e^- \rightarrow \mu^+\mu^-) = \frac{4\pi\alpha^2}{3s} \quad (6.1.27)$$

in cui non compare nè la carica di flavour, avendo $Q_f^2 = +1$, ma nemmeno nessun fattore di colore ($N = 1$), non essendo presente alcuna delta di colore nel vertice di QED.

Il rapporto R delle due sezioni d'urto, che guarda caso è dipendente solamente da una scala ed è adimensionale (proprio come le caratteristiche analizzate nella sezione §5.3), è:

$$R = \frac{\sigma_B(e^+e^- \rightarrow \text{adroni})}{\sigma_B(e^+e^- \rightarrow \mu^+\mu^-)} = \sum_f Q_f^2 N. \quad (6.1.28)$$

Da un lato, N stabilisce la *normalizzazione globale* del risultato. Dall'altro, la somma su f nell'espressione di R corre sulle specie la cui massa è inferiore a $\sqrt{s}/2$, fornendo la predizione:

$$\sqrt{s} < 2m_c : \quad R = 3 \left[\left(\frac{2}{3} \right)^2 + \left(\frac{1}{3} \right)^2 + \left(\frac{1}{3} \right)^2 \right] = 2, \quad (6.1.29)$$

$$2m_c < \sqrt{s} < 2m_b : \quad R = 2 + 3 \left(\frac{2}{3} \right)^2 = \frac{10}{3}, \quad (6.1.30)$$

$$2m_b < \sqrt{s} : \quad R = \frac{10}{3} + 3 \left(\frac{1}{3} \right)^2 = \frac{11}{3}. \quad (6.1.31)$$

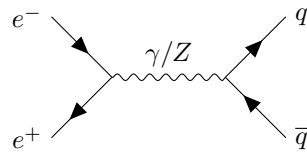
Quest'ultimo valore è comunque valido per $\sqrt{s} \ll M_Z$, per poter trascurare lo scambio del bosone Z . Precisiamo solo che nel caso (6.1.29) i quark considerabili massless sono u, d, s e infatti moltiplichiamo $N = 3$ per la somma delle rispettive cariche al quadrato; in (6.1.30) i quark massless sono u, d, s, c ; nel caso (6.1.31) i fermioni massless sono u, d, s, c, b .

Un confronto accurato di questi risultati con le misurazioni sperimentali (si vedano le immagini tratte da *There is a sort of mythology that grows up about what happened, which is different from what really did happen*. Ellis) ha storicamente dato fiducia nell'esistenza del grado di libertà del colore, oltre a confermare che effettivamente il gruppo di simmetria fosse $SU(3)_c$, dunque con $N = 3$, ma confermò anche la carica frazionaria dei quark, che possiamo trovare dal gap tra i vari "scalini" dell'andamento di R . Vedi la figura 6.4 per i dati sperimentali e la figura 6.5 per l'andamento (abbozzato) di R .

Nota. Noi in tutta questa trattazione abbiamo trascurato il contributo dello scambio del bosone Z , ma nel caso in cui includessimo anche il suo contributo avremmo una modifica al risultato, che cambia leggermente la "shape" del grafico di R .

6.1.3 Commenti sui diagrammi del processo (tagli)

Riassumiamo brevemente come siamo giuti alla situazione presente. All'inizio della sezione §6.1.2 abbiamo disegnato il diagramma del processo:



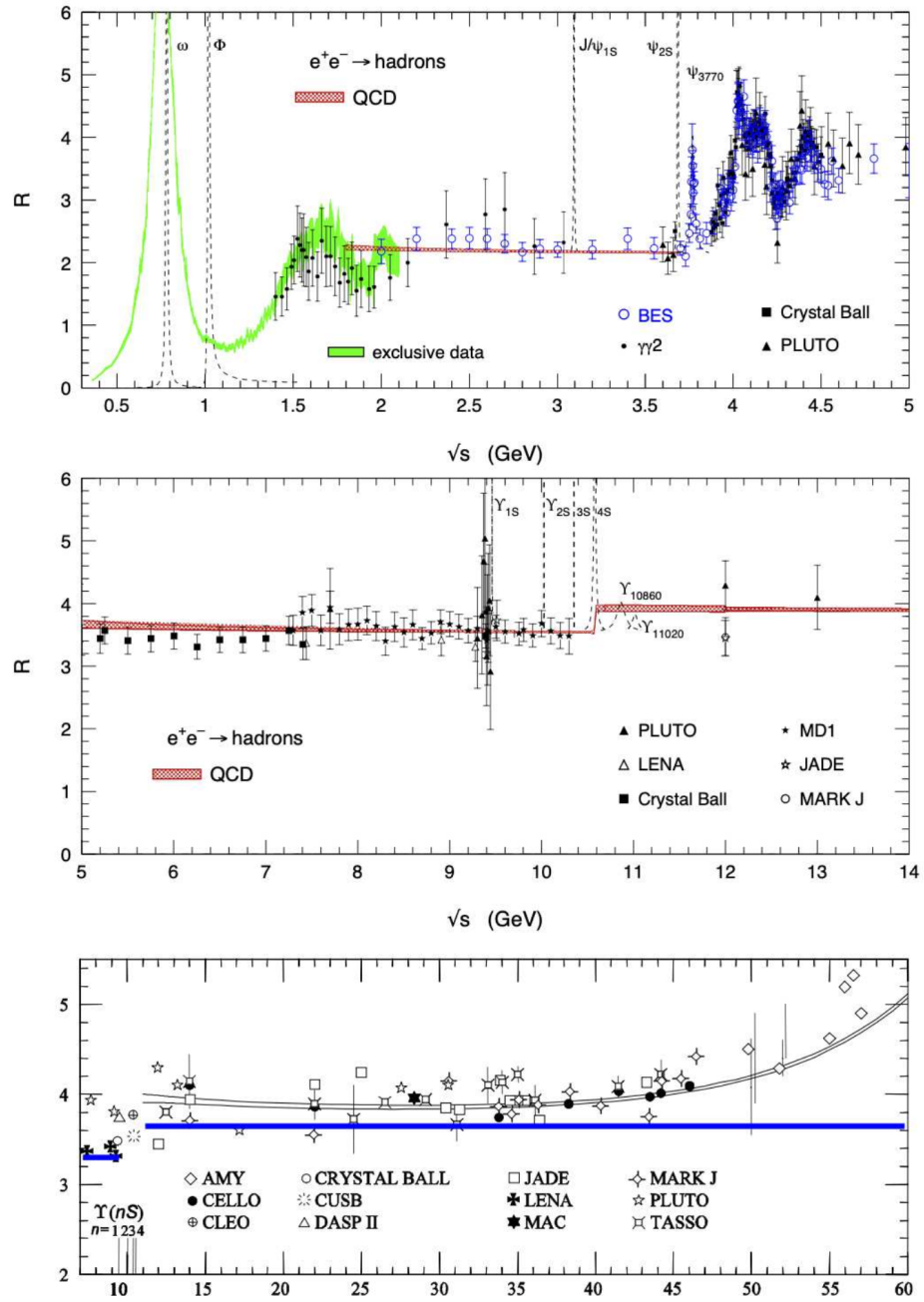
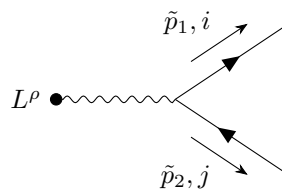
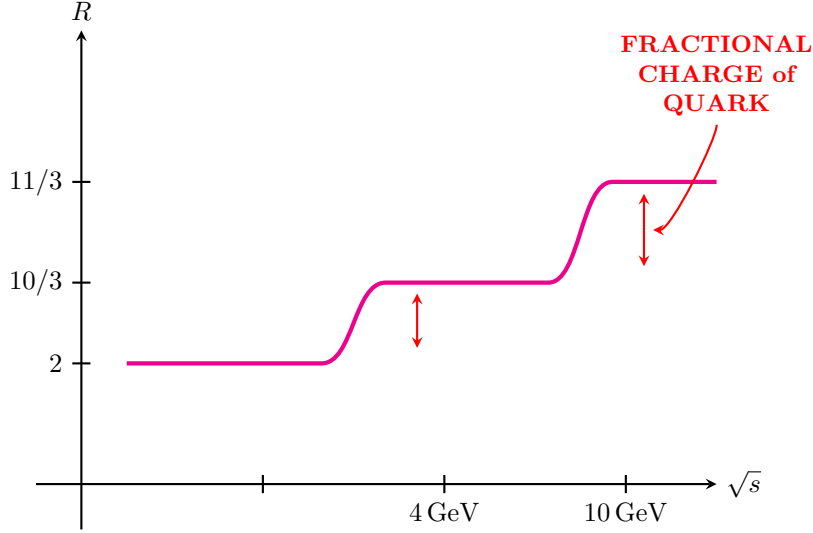


Figura 6.4

che in modo più compatto è:



Figura 6.5: R prediction.

in cui abbiamo considerato la corrente leptonica L^ρ attaccata a sinistra. Indichiamo con un puntino la corrente leptonica (che sarebbero sostanzialmente le gambe esterne che entrano nel vertice) poiché non entrano mai nei conti di QCD, per cui abbreviamo leggermente la notazione. Poi, giustamente, abbiamo contratto $L^{\rho\sigma}$ (cioè il tensore leptonico ottenuto mediando sulle polarizzazioni di due correnti leptoniche) con la parte di corrente di quark.

Tralasciando discorsi sull'unitarietà, che vengono fatti in più dettaglio nel corso di *Advanced QFT*, possiamo fare alcune considerazioni sui diagrammi che stiamo trattando e i risultati a cui portano.

Pensandoci bene, una corrente, leptonica o di quark che sia, è sempre della forma (per la QED che stiamo guardando):

$$L^\rho = -\frac{i}{s} \bar{v}(q_2)(-ie\gamma^\rho)u(q_1) \quad (6.1.32)$$

ossia è un vettore di Lorentz, visto l'indice ρ proveniente dalla matrice di Dirac, e possiede due indici spinoriali, provenienti dagli spinori bidimensionali. Guardando bene la corrente L^ρ è il contributo di due gambe esterne che entrano in un vertice.

Quando costruiamo un tensore leptonico stiamo sostanzialmente guardando i due contributi (uno complesso coniugato) di due vertici. Però, se

guardiamo il risultato che abbiamo ottenuto, ovvero:

$$\begin{aligned}
 L^{\rho\sigma} &= \frac{1}{4} \sum_{\text{pol. lep.}} L^\rho L^{\sigma*} \\
 &= \frac{e^2}{4s^2} \text{Tr} \left\{ \not{q}_1 \gamma^\sigma \not{q}_2 \gamma^\rho \right\} \\
 &= \frac{e^2}{s^2} (q_1^\rho q_2^\sigma + q_2^\rho q_1^\sigma - q_1 \cdot q_2 g^{\rho\sigma})
 \end{aligned}$$

ci rendiamo conto che la costruzione del tensore leptonico (così come sarebbe per i quark) ci fa' contrarre gli indici spinoriali, portandoci alla traccia, che sappiamo svolgere. Notiamo anche che (giustamente essendo un tensore) ci rimangono gli indici di Lorentz.

Possiamo ottenere uno scalare, ossia un'ampiezza di transizione, solo se consideriamo il diagramma nel suo intero, ovvero un tensore leptonico, che tramite un propagatore bosonico, contenente una metrica, è collegato ad un tensore di quark (al cui interno ci sono anche delle delta di colore). In questo modo otteniamo effettivamente uno scalare (tutti gli indici si contraggono) e abbiamo in fin dei conti dei prodotti di tracce da fare. Arriviamo dunque al risultato (6.1.12).

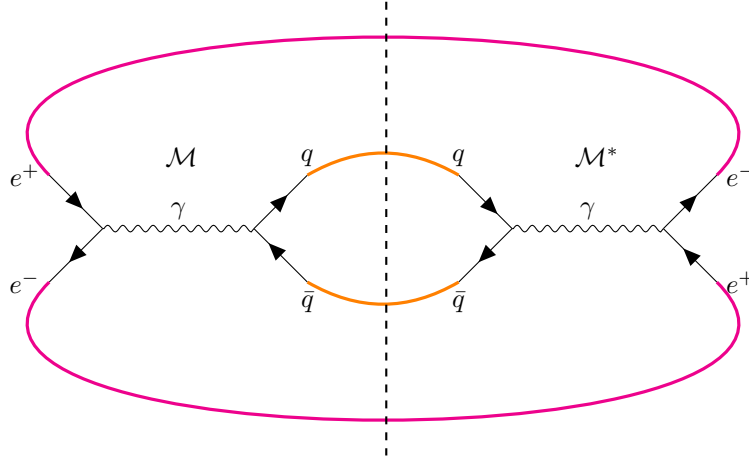
Possiamo notare che in $|\mathcal{M}|^2$ abbiamo tante traccie quante sono le linee fermioniche presenti.

Ciò di cui possiamo renderci conto è che direttamente tramite i diagrammi di Feynman possiamo vedere quali indici contrarre e quindi che tracce scrivere. Noi stiamo calcolando:

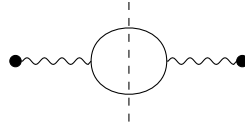
$$\begin{aligned}
 |\mathcal{M}|^2 &= \mathcal{M} \cdot \mathcal{M}^* \\
 &= \left[\begin{array}{c} e^- \\ \nearrow \\ \gamma/Z \\ \nwarrow \\ e^+ \end{array} \right] \left[\begin{array}{c} q \\ \nearrow \\ \gamma/Z \\ \nwarrow \\ \bar{q} \end{array} \right] \left[\begin{array}{c} q \\ \nearrow \\ \gamma/Z \\ \nwarrow \\ \bar{q} \end{array} \right] \left[\begin{array}{c} e^- \\ \nearrow \\ \gamma/Z \\ \nwarrow \\ e^+ \end{array} \right]
 \end{aligned}$$

ma, visto che sappiamo quali indici spinoriali si contraggono (ovvero, dai tensori, leptoni con leptoni e quark con quark) possiamo anche raffigurarlo

come:



e in modo leggermente più compatto:



dove, ancora una volta, mettiamo un "dot" al posto delle correnti leptoniche, e inseriamo un *taglio* a metà del loop di quark. Il *cut* messo a metà del loop ci indica che l'oggetto che stiamo calcolando è un'ampiezza modulo quadro, e che tutte le gambe (particelle) che attraversano il taglio sono contratte tra loro; ovvero, sono delle tracce. Infatti, come abbiamo notato, gli indici dei quark sono contratti tra loro e analogamente gli indici dei leptoni.

Possiamo anche scrivere, con una notazione grafica:

$$\sigma_{LO} = \frac{1}{2s} \int d\Phi_2 \mathcal{M}_B \mathcal{M}_B^* \quad (6.1.33)$$

$$= \frac{1}{2s} \int d\Phi_2 \bullet \text{---} \text{---} \text{---} \text{---} \text{---} \text{---} \text{---} \text{---} \text{---} \bullet \quad (6.1.34)$$

$$= \mathcal{O}(\alpha_s^0). \quad (6.1.35)$$

6.1.4 Nomenclatura: Born e tree-level

Possiamo fare una precisazione riguardo la nomenclatura che utilizziamo e utilizzeremo. Usiamo il termine *livello Born* riferendoci ai diagrammi che contribuiscono al LO di un dato processo, mentre usiamo il termine *livello albero* (tree level) riferendoci a diagrammi senza loop chiusi. In altre parole,

Born è una proprietà del processo e dell'ordine perturbativo al quale è calcolato, mentre *albero* è una proprietà del diagramma. Sebbene spesso (ma non sempre) i diagrammi di livello Born siano diagrammi ad albero, l'essere a livello albero non implica che si stia considerando l'approssimazione di Born, come è evidente ad esempio guardando alle correzioni reali. Lo schema 6.1 dovrebbe aiutare a chiarire la differenza tra i due concetti.

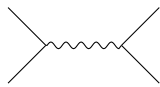
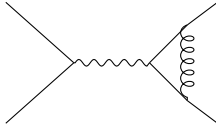
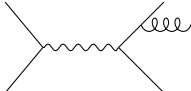
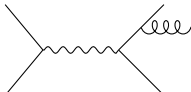
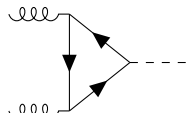
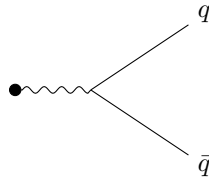
		Tree level	Born	
1.	$e^+e^- \rightarrow 2 \text{ jets LO}$		✓	✓
2.	$e^+e^- \rightarrow 2 \text{ jets NLO}$		x	x
3.	$e^+e^- \rightarrow 2 \text{ jets NLO}$		✓	x
4.	$e^+e^- \rightarrow 3 \text{ jets LO}$		✓	✓
5.	$gg \rightarrow H \text{ LO}$		x	✓

Tabella 6.1

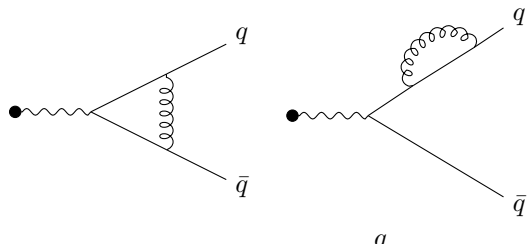
6.1.5 Calcoli al Next-to-leading-order (NLO) e divergenze IR

Si può vedere il capitolo §17 del Peskin There is a sort of mythology that grows up about what happened, which is different from what really did happen. Oltre l'approssimazione di Born, due tipi di correzioni radiative sorgono nel calcolo dei processi di scattering: correzioni virtuali e reali. Le *correzioni virtuali* comportano diagrammi di Feynman con loop chiusi, in cui una particella viene irradiata e riassorbita all'interno del diagramma. Le *correzioni reali* comportano l'irraggiamento di una particella che non viene riassorbita all'interno del diagramma, che quindi compare nello stato finale come una particella extra rispetto alla molteplicità che caratterizza il Born.

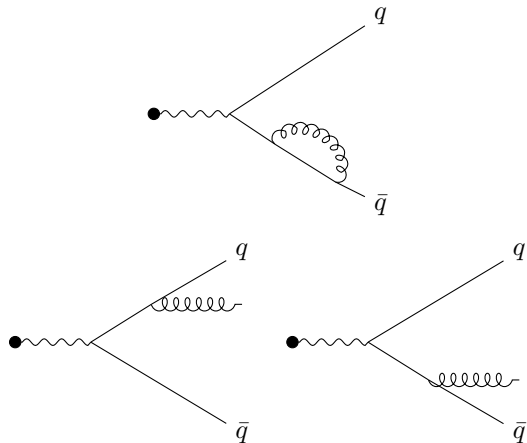
Ad esempio, in $e^+e^- \rightarrow \text{adroni}$, i diagrammi di Born, virtuali e di radiazione reale sono rappresentati di seguito:



Born: $\mathcal{O}(\alpha_s^0)$ (6.1.36)



virtual: $\mathcal{O}(\alpha_s^1)$ (6.1.37)



real: $\mathcal{O}(\alpha_s^{1/2})$ (6.1.38)

in cui ci ricordiamo $\alpha_s = g_s^2/(4\pi)$ e possiamo vedere il contributo di Born dell'ordine α_s^0 , non contenendo alcun vertice di QCD, mentre il contributo virtuale, avendo due vertici di QCD, dà contributo g_s^2 che corrisponde ad una potenza α_s^1 ; allo stesso modo possiamo osservare che i contributi reali, avendo un solo vertice g_s , danno contributo di ordine $\alpha_s^{1/2}$.

Questi diagrammi, se fatti modulo quadro e integrati sullo spazio fasi, ci forniscono i relativi contributi alla sezione d'urto, oltre ai termini di interferenza; però l'ordine diverso dei diagrammi implica che al primo ordine perturbativo noi consideriamo solamente alcune contrazioni. Ripareremo a brevissimo di questa cosa.

Vediamo velocissimamente che in termini dei diagrammi Born, virtuali e reali possiamo scrivere, diagrammaticamente, il processo come:

$$\begin{aligned}
 \text{Shaded Circle} &= \text{Born Diagram} + \left(\text{Virtual Corrections} + \text{Real Corrections} \right) \\
 &+ \left(\text{Virtual Corrections} + \text{Real Corrections} \right). \quad (6.1.39)
 \end{aligned}$$

Ovviamente, avendo più contributi differenti all'ampiezza di probabilità, la sezione d'urto sarà del tipo:

$$\sigma(ee \rightarrow q\bar{q}) = |\mathcal{M}_B + \mathcal{M}_V + \mathcal{M}_R|^2 \quad (6.1.40)$$

$$\begin{aligned}
 &= |\mathcal{M}_B|^2 + |\mathcal{M}_V|^2 + |\mathcal{M}_R|^2 + \\
 &\quad + 2 \text{Re}\{\mathcal{M}_B \mathcal{M}_V^*\} + 2 \text{Re}\{\mathcal{M}_B \mathcal{M}_R^*\} + 2 \text{Re}\{\mathcal{M}_R \mathcal{M}_V^*\} \quad (6.1.41)
 \end{aligned}$$

ma possiamo notare che non tutti i termini danno lo stesso contributo, oltre che non tutti sono possibili. Infatti, la sezione d'urto di Born (LO), ossia il primo termine di (6.1.39), è stata calcolata come:

$$\sigma_{LO} = \frac{1}{2s} \int d\Phi_2 \mathcal{M}_B \mathcal{M}_B^* \quad (6.1.42)$$

$$= \frac{1}{2s} \int d\Phi_2 \text{Born Diagram} \quad (6.1.43)$$

$$= \mathcal{O}(\alpha_s^0). \quad (6.1.44)$$

Analogamente, per le correzioni virtuali e reali, abbiamo:

$$\sigma_V = \frac{1}{2s} \int d\Phi_2 (\mathcal{M}_V \mathcal{M}_B^* + \mathcal{M}_B \mathcal{M}_V^*) \quad (6.1.45)$$

$$= \frac{1}{2s} \int d\Phi_2 2 \text{Re} \text{Virtual Diagram} + \dots \quad (6.1.46)$$

$$= \mathcal{O}(\alpha_s) \quad (6.1.47)$$

reali e virtuali corrispondono a diagrammi molto simili, l'unica differenza è rappresentata dal gluone che attraversa o meno la linea tratteggiata.

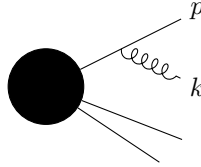
Dunque, fin'ora abbiamo capito che nonostante σ_V e σ_R contengano contributi di diversi numeri di particelle, loro contribuiscono allo stesso ordine della sezione d'urto:

$$\sigma = \sigma_{LO} + \sigma_V + \sigma_R \quad (6.1.52)$$

$$= \mathcal{O}(\alpha_s^0) + \mathcal{O}(\alpha_s^1) + \mathcal{O}(\alpha_s^1). \quad (6.1.53)$$

Però, il calcolo delle correzioni radiative oltre l'approssimazione di Born è delicato, poiché se presi singolarmente, i contributi reali e virtuali sono separatamente divergenti, ovvero $\sigma_R = +\infty$ e $\sigma_V = -\infty$.

Faremo tutti i calcoli per bene nelle prossime sezioni del capitolo, ma per ora al fine di ottenere una comprensione fisica dell'origine di tali divergenze, consideriamo l'emissione di un gluone (potrebbe benissimo essere un fotone) con impulso k da una linea di quark nello stato finale con impulso p , come illustrato di seguito:



Il propagatore del quark che irradia è:

$$\Delta \propto \frac{1}{(p+k)^2 - m_p^2} = \frac{1}{2p \cdot k} = \frac{1}{2E_p E_k (1 - \beta_p \cos \theta_{pk})} \quad (6.1.54)$$

in cui abbiamo parametrizzato:

$$k^\mu = E_k(1, 0, 0, 1) \quad (6.1.55)$$

$$p^\mu = E_p(1, 0, \beta \sin \theta_{pk}, \beta \cos \theta_{pk}) \quad (6.1.56)$$

e dove indichiamo con m_p la massa del quark, E_p, E_k le energie del quark e del gluone in un dato sistema di riferimento, θ_{pk} il loro angolo relativo, e $\beta_p = |\vec{p}|/E_p = \sqrt{E_p^2 - m_p^2}/E_p$ la velocità del quark.

Possiamo vedere che questo propagatore dà dei problemi in alcune situazioni. Infatti, il propagatore diverge ogni volta che l'energia del gluone svanisce, $E_k \rightarrow 0$, il che è soprannominato emissione di un *gluone soft* (sof-fice). Inoltre, se siamo nel caso $\beta_p = 1$, che accade quando $m_p = 0$, ovvero nel limite in cui trascuriamo le masse dei quark rispetto alla scala dura del processo, il propagatore del quark svanisce anche quando $\theta_{pk} \rightarrow 0$, il che è chiamato *emissione collineare*. Potremmo anche dire che il propagatore svanisce quando $E_p \rightarrow 0$, ma in realtà questa singolarità è sistematicamente compensata dagli spinori del quark nel numeratore dell'ampiezza, ed è quindi innocua.

Tali divergenze soffici e collineari non hanno nulla a che fare con quelle di origine UV che sorgono nei calcoli a loop. Esse si verificano all'estremità infrarossa (IR) dello spettro energetico del gluone, e sono quindi collettivamente chiamate *divergenze infrarosse*. Come vedremo in seguito, queste divergenze non sono fisicamente distinguibili.

L'integrazione della radiazione reale sullo spazio delle fasi può essere intrapresa solo previa adeguata regolarizzazione, per esempio mediante la regolarizzazione dimensionale in $d = 4 - 2\epsilon$ dimensioni spazio-temporali. I contributi reali IR-divergenti, una volta integrati, daranno origine a poli nel regolatore ϵ : σ_R divergerà come $1/\epsilon^2$, dove la presenza di un polo doppio deriva dalla possibilità che il gluone irradiato contemporaneamente sia sofficie che collineare rispetto a una delle linee radianti prive di massa. In σ_R sorgeranno poli singoli $1/\epsilon$ quando il gluone è o sofficie o collineare rispetto ad una gamba priva di massa. Analogamente, i contributi virtuali divergeranno come $1/\epsilon^2$ anche dopo la rimozione delle divergenze UV mediante controtermini, poiché il gluone virtuale scambiato può essere sofficie e/o collineare rispetto alle linee esterne prive di massa che compaiono nel diagramma.

Le correzioni reali e virtuali, che sono separatamente IR-divergenti, danno comunque origine a una sezione d'urto fisicamente significativa una volta sommate insieme:

$$\sigma_{NLO} = \sigma_B + \sigma_R + \sigma_V = \text{finita}. \quad (6.1.57)$$

Questo risultato, valido per qualsiasi teoria massless (quindi anche la QED), è una conseguenza di un teorema generale della meccanica quantistica, il *teorema di Kinoshita-Lee-Nauenberg* (KLN), 1962-1964, che è di fondamentale importanza per il calcolo delle correzioni radiative nella fenomenologia moderna. Il teorema afferma in sostanza che gli elementi della matrice di scattering S sono finiti ordine per ordine nella teoria perturbativa una volta che si somma su tutti gli stati degeneri contribuenti.

Seguendo questo teorema possiamo vedere (lo faremo esplicitamente nelle prossime sezioni) che le espressioni per i contributi alla sezione d'urto sono del tipo:

$$\sigma_R \left(\begin{smallmatrix} \text{regolarizzazione} \\ \text{dimensionale} \end{smallmatrix} \right) = \frac{A}{\epsilon^2} + \frac{B}{\epsilon} + C_R \quad (6.1.58)$$

$$\sigma_V \left(\begin{smallmatrix} \text{regolarizzazione} \\ \text{dimensionale} \end{smallmatrix} \right) = -\frac{A}{\epsilon^2} - \frac{B}{\epsilon} + C_V \quad (6.1.59)$$

in cui i coefficienti A e B sono uguali per i processi reali e virtuali grazie al teorema KLN. Possiamo perciò scrivere:

$$\lim_{\epsilon \rightarrow 0} (\sigma_R + \sigma_V) = \text{finito} = C_R + C_V. \quad (6.1.60)$$

Torniamo sul punto dell'indistinguibilità fisica delle divergenze, su cui prima siamo andati veloci.

Dal punto di vista fisico, un gluone reale soffice con $E_k = 0$ non può in alcun modo essere distinto da un gluone virtuale⁶. Analogamente, una coppia collineare $q\bar{q}$ non può essere distinta da un quark che emette un gluone virtuale. Le divergenze IR sono quindi una conseguenza della separazione (che facciamo noi a mano su un foglio di carta) di configurazioni che non possono essere fisicamente distinte, ciascuna delle quali trasporta un peso infinito $1/\epsilon^2$.

Il fatto che le radiazioni soffice e collineari siano associate a contributi divergenti è legato alla tipica scala temporale alla quale esse hanno luogo. Un'emissione soffice con $E_k \rightarrow 0$ avviene a tempi $t \sim 1/E_k \rightarrow \infty$, ovvero molto più tardi rispetto alla scala temporale che caratterizza il processo duro, $t \sim 1/\sqrt{s}$. Pertanto, nella teoria delle perturbazioni dipendente dal tempo, le divergenze IR corrispondono a contributi non soppressi integrati su tempi infiniti, dando origine a risultati infiniti. Le divergenze IR sono quindi legate a fenomeni a lunga distanza, contrariamente alle scale a breve distanza tipiche dello scattering duro sottostante. Una tale chiara separazione delle scale di tempo/distanza tra i contributi duri e IR fornisce anche una motivazione fisica alla base della cancellazione delle divergenze IR (oltre la prova formale del teorema KLN). I contributi radiativi che hanno luogo a tempi lunghi $t \sim 1/E_k \rightarrow \infty$ non possono modificare pesantemente le probabilità di transizione *istantanee* σ_{LO} calcolate a $t \sim 1/\sqrt{s}$, e dunque non possono modificare in modo sostanziale la probabilità che $q\bar{q}$ siano prodotti. Il grosso dell'effetto soffice/collineare dovrebbe cancellarsi, lasciando un residuo finito e tipicamente piccolo (infatti modificano σ_B con un fattore universale moltiplicativo), soppresso da potenze del parametro di accoppiamento (ordine α_s).

Questa immagine, rilevante per le sezioni d'urto inclusive (cioè quando guardiamo al processo totale, ad esempio guardiamo quanti jet vengono prodotti in totale), verrà modificata una volta che avremo a che fare con spettri esclusivi (ossia in esperimenti in cui vogliamo misurare qualcosa di molto specifico, ad esempio ci preoccupiamo di cercare dei gluoni con uno specifico angolo), ovvero con la sezione d'urto per unità di un dato osservabile. Questi saranno ancora finiti per osservabili opportunamente definiti, tuttavia le correzioni radiative indurranno contributi logicamente potenziati dalla radiazione IR.

⁶Nel capitolo §17.2 ci spiega che i due stati sono due configurazioni non distinguibili, cioè degeneri, poiché sperimentalmente abbiamo i rivelatori di dimensione finita, dunque non si è abbastanza precisi per determinare se un quark passa attraverso un calorimetro da solo, o se insieme a lui c'è un quark collineare, e in più i rivelatori hanno una certa energia di soglia, per la quale non si è in grado di rivelare particelle con energia al di sotto, dunque gluoni soffici. In altri termini potremmo dire che guardando lo scattering $e^+e^- \rightarrow$ adroni, in stato finale vediamo degli stati legati, o comunque dei jet, i quali non sono in grado di dirci se essi si sono formati da un gluone virtuale in una gamba fermionica, oppure se sono una combinazione di $q\bar{q}$ collineari. Caso diverso è quando abbiamo $e^+e^- \rightarrow q\bar{q}g$ con momento trasverso del gluone grande, per cui vediamo uno stato finale a 3 jet.

Infatti se il calcolo perturbativo non è limitato alla sezione d'urto totale σ , ma intende sondare osservabili sensibili alla presenza di emissioni soffici e/o collineari, le correzioni alla predizione di Born possono essere grandi (logaritmi grandi). Ciò talvolta richiede il calcolo di ordini superiori⁷ nella teoria perturbativa, o persino la necessità di sommare classi di contributi a tutti gli ordini perturbativi (chiamata *resummation*).

Inoltre, il *teorema KLN* si applica non solo al calcolo della sezione d'urto totale σ , ma di conseguenza, anche alle sezioni d'urto differenziali $d\sigma/dx$, dove x è un'osservabile generica di cui vogliamo calcolare lo spettro nella teoria perturbativa.

Tuttavia, affinché la cancellazione tra contributi reali e virtuali stabilita dal KLN valga per $d\sigma/dx$, x deve essere definita in modo tale da assumere lo stesso valore numerico in presenza di un'emissione reale soffice e/o collineare che in presenza di un'emissione virtuale. In questo modo, per ogni valore di x , $d\sigma/dx$ riceve sia contributi reali che virtuali divergenti, tali che:

$$\frac{d\sigma_R}{dx} + \frac{d\sigma_V}{dx} < \infty \quad (6.1.61)$$

Tale osservabile x è definita come **IR-safe**. Impareremo che non tutte le osservabili si comportano bene, ovvero non sono *safe*.

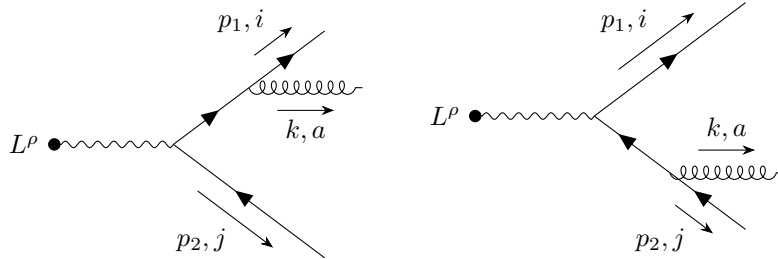
Esempio. Un'osservabile che *non* è IR-safe: il numero di gluoni nello stato finale del processo $e^+e^- \rightarrow 2 \text{ jets}$ al NLO in QCD. Le correzioni virtuali hanno $N_g = 0$, mentre le correzioni reali hanno $N_g = 1$, indipendentemente dalla cinematica. Pertanto:

$$\frac{d\sigma}{dN_g} \propto +\delta(N_g - 1) - \delta(N_g - 0)$$

il che non ha senso come distribuzione fisica. Ripareremo di questo esempio nella sezione §6.3.1.

6.1.6 Correzioni NLO: contributo reale

Prendiamo i diagrammi:



⁷Puoi vedere i contributi di ordine superiore contraendo i vari diagrammi (Born, virtuali e reali) in modo da ottenere ampiezze di ordine $\mathcal{O}(\alpha_s^k)$.

in cui ricordiamo l'espressione della corrente leptonica (6.1.4). Iniziamo calcolando il contributo reale in gauge di Feynman, $\sum_{\mu} \epsilon_{\mu} \epsilon_{\nu}^* = -g_{\mu\nu}$, in $d = 4 - 2\epsilon$, seguendo There is a sort of mythology that grows up about what happened, which is different from what really did happen. Muta, chiamando rispettivamente i diagrammi 1 e 2:

$$\mathcal{M}_{R1} = L^{\rho} \bar{u}(p_1) (i g_s \mu^{\epsilon} \gamma^{\mu} t_{ij}^a) \frac{i(\not{p}_1 + \not{k})}{(p_1 + k)^2} (ie Q_f \gamma_{\rho} \delta_{lj}) v(p_2) \quad (6.1.62)$$

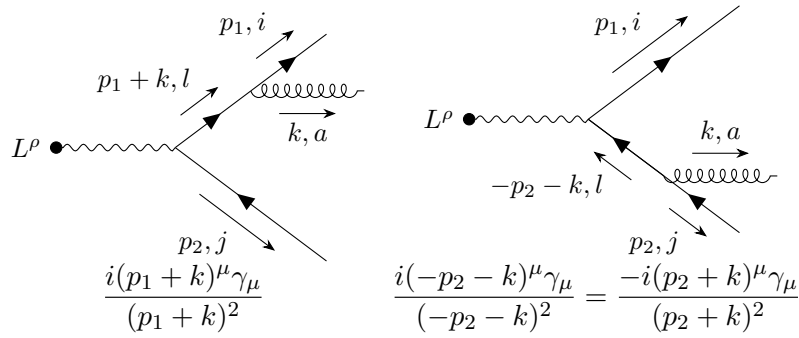
$$\mathcal{M}_{R2} = L^{\rho} \bar{u}(p_1) (ie Q_f \gamma_{\rho} \delta_{il}) \frac{-i(\not{p}_2 + \not{k})}{(p_2 + k)^2} (i g_s \mu^{\epsilon} \gamma^{\mu} t_{lj}^a) v(p_2) \quad (6.1.63)$$

$$\mathcal{M}_R = \mathcal{M}_{R1} + \mathcal{M}_{R2} = -i L^{\rho} g_s \mu^{\epsilon} e Q_f t_{ij}^a \bar{u}(p_1) (\gamma_{\mu} \not{p}_1 \gamma_{\rho} - \gamma_{\rho} \not{p}_2 \gamma_{\mu}) v(p_2) \quad (6.1.64)$$

dove (ricordando che i quark sono massless, per cui $(p_i + k)^2 \sim 2p_i \cdot k$):

$$a_{\mu} = \frac{(p_1 + k)_{\mu}}{2p_1 \cdot k} \quad , \quad b_{\mu} = \frac{(p_2 + k)_{\mu}}{2p_2 \cdot k} \quad (6.1.65)$$

e in cui abbiamo associato un indice ρ al fotone, un indice μ al gluone che vediamo nello stato finale, un indice di colore l al propagatore del quark intermedio tra il vertice $\gamma q \bar{q}$ e il vertice $g q \bar{q}$ e insieme un impulso $p_i + k$ uscente, con rispettivamente $i = 1, 2$. Teniamo solo a mente che, come sappiamo, dobbiamo scrivere il propagatore di un fermione (in questo caso i quark intermedi tra i due vertici) con il segno corretto per il verso dell'impulso nel senso della linea fermionica; spiegato meglio vuol dire che associamo, ai quark intermedi tra i due vertici di QED e QCD, rispettivamente:



in cui possiamo sfruttare il fatto che è equivalente scrivere $p_2 + k$ o $-p_2 - k$ a patto che orientiamo opportunamente le frecce.

Il fattore μ^{ϵ} tiene conto della dimensione della costante di accoppiamento in $d \neq 4$ (ed è il fattore utilizzato nella regolarizzazione dimensionale). Introducendo:

$$S^{\mu\nu} = (\gamma^{\mu} \not{p}_1 \gamma^{\nu} - \gamma^{\nu} \not{p}_2 \gamma^{\mu}) \quad , \quad H_{\rho\sigma} = \text{Tr} \{ \not{p}_1 S^{\mu}_{\rho} \not{p}_2 S_{\sigma\mu} \} \quad (6.1.66)$$

e considerando che gli indici di colore degli spinori ci portano a:

$$\begin{aligned}
 |\mathcal{M}_r|^2 &\propto \bar{u}_i(p_1)t_{ij}^a v_j(p_2)\bar{v}_k(p_2)t_{kl}^a u_l(p_1) \\
 &= \delta_{li}t_{ij}^a \delta_{jk}t_{kl}^{a*} \\
 &= t_{ij}^a t_{ji}^a \\
 &= \delta^{ab} \text{Tr}\{t^a t^b\} \\
 &= \frac{1}{2}\delta^{ab}\delta^{ab} \\
 &= \frac{N_c^2 - 1}{2} \\
 &= C_F N_c
 \end{aligned}$$

e si ottiene:

$$\frac{1}{4} \sum_{\text{pol.}} |\mathcal{M}_R|^2 = -\frac{(g_s \mu^\epsilon e Q_f)^2 C_F N_c}{4s^2} \text{Tr}\{\not{q}_1 \gamma^\sigma \not{q}_2 \gamma^\rho\} \text{Tr}\{\not{p}_1 S^\mu_\rho \not{p}_2 S_{\sigma\mu}\} \quad (6.1.67)$$

$$= -\frac{(g_s \mu^\epsilon e Q_f)^2 C_F N_c}{4s^2} \text{Tr}\{\not{q}_1 \gamma^\sigma \not{q}_2 \gamma^\rho\} H_{\rho\sigma}. \quad (6.1.68)$$

Possiamo notare che la struttura ottenuta è identica al risultato di Born (6.1.10), ma con la sola sostituzione $S^\mu_\rho \leftrightarrow \gamma_\rho$, di cui S^μ_ρ tiene conto del gluone presente nello stato finale. In più, come fatto anche nel calcolo al livello Born, non stiamo considerando delle delta δ_{fg} di flavour, che sarebbero in realtà presenti nei vertici; ciò è lecito farlo poiché, pur considerando δ_{fg} , una volta calcolato il modulo quadro ci uscirebbe una traccia sui flavour, la quale corrisponde alla somma $\sum_{\text{flavour}}$, che noi facciamo esplicitamente dopo. Dunque, non stiamo dimenticando nulla.

Possiamo semplificare il calcolo⁸ considerando che l'espressione per l'elemento di matrice reale al quadrato dovrà essere integrata sullo spazio delle fasi a tre corpi dei momenti p_1, p_2, k . La dipendenza dai momenti p_1, p_2 e k in $|\mathcal{M}_R|^2$ è interamente contenuta nel tensore adronico $H_{\rho\sigma}$, quindi possiamo convenientemente analizzare il tensore:

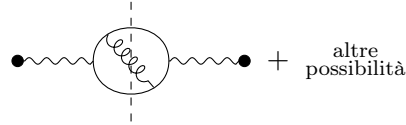
$$I_{\mu\nu} = \int d\Phi_3^{(d)}(p_1, p_2, k; s) H_{\mu\nu} \quad (6.1.69)$$

che dipende solo dal momento $Q_\mu = (q_1 + q_2)_\mu = (p_1 + p_2 + k)_\mu$, dunque dalla somma dei tre momenti e non dai singoli, poiché la dipendenza separata da

⁸Questa semplificazione di conti è fatta seguendo There is a sort of mythology that grows up about what happened, which is different from what really did happen. Muta, ma ad esempio There is a sort of mythology that grows up about what happened, which is different from what really did happen. Ellis svolge tutti i conti, tenendosi le dipendenze di p_1, p_2, k , che sarebbe la cosa corretta per essere inclusivi per tutti i momenti dello stato finale. A noi va bene la semplificazione poiché siamo interessati alla sezione d'urto totale e non a misure inclusive.

p_1, p_2, k è stata integrata. Si noti che $I_{\mu\nu}$ è lo spazio delle fasi integrale del modulo quadro dell'ampiezza di transizione di un fotone nell'insieme $q\bar{q}g$, quindi per il teorema ottico è proporzionale alla parte immaginaria dell'self-energy del fotone. Come tale, per la conservazione della corrente elettromagnetica, esso soddisfa $Q_\mu I^{\mu\nu} = 0$, da cui⁹ $I^{\mu\nu} = \frac{I(s)}{s}(Q^\mu Q^\nu - sg^{\mu\nu})$, con $I(s) = I_\alpha^\alpha/(1-d)$.

Facciamo una breve pausa per vedere meglio questa cosa. $H_{\mu\nu}$ è la parte adronica del modulo quadro $|\mathcal{M}|^2$ e possiamo rappresentarla come:



per cui $I_{\mu\nu}$ non è altro che:

$$I_{\mu\nu}(Q) = \int d\Phi_3^{(d)} \text{ (diagramma) } \propto \text{Im} \left\{ \text{(diagramma)} \right\}$$

in cui abbiamo scritto la proporzionalità per via del teorema ottico (nota che, come richiede il teorema, siamo in una situazione on-shell avendo rappresentato le correnti leptoniche, dunque le particelle esterne). Però $I_{\mu\nu}$, essendo proporzionale alle correzioni a loop (in particolare a due loop) del propagatore del fotone dovrà essere trasverso (come γ), e dunque risolvere $Q_\mu I^{\mu\nu} = 0$.

Capito chi sono, e che forma hanno, $H_{\mu\nu}$ e $I_{\mu\nu}$ possiamo tornare al calcolo

⁹Troviamo che l'unico modo possibile di scrivere un tensore Lorentz covariante con due indici è quello di utilizzare i tensori $g^{\mu\nu}$ e $Q^\mu Q^\nu$, per cui una forma generica che possiamo scrivere è:

$$I^{\mu\nu}(Q) = A(Q^2)g^{\mu\nu} + B(Q^2)Q^\mu Q^\nu \quad (6.1.70)$$

ma imponendo $Q_\mu I^{\mu\nu} = 0$ e contraendo $I^{\mu\nu}g_{\mu\nu} = I^\alpha_\alpha$, possiamo trovare che:

$$A = -BQ^2 = -Bs \quad , \quad B = \frac{I^\alpha_\alpha}{s(1-d)}. \quad (6.1.71)$$

della sezione d'urto, per la quale abbiamo:

$$\sigma_R = \frac{1}{2s} \int d\Phi_3^{(d)}(p_1, p_2, k; s) \frac{1}{4} \sum_{\text{pol.}} |\mathcal{M}_R|^2 \quad (6.1.72)$$

$$= -\frac{(g_s \mu^\epsilon e Q_f)^2 C_F N_c}{4s^2} \frac{1}{2s} \text{Tr} \left\{ \not{q}_1 \gamma^\sigma \not{q}_2 \gamma^\rho \right\} \underbrace{\int d\Phi_3^{(d)}(p_1, p_2, k; s) H_{\rho\sigma}}_{I_{\rho\sigma}} \quad (6.1.73)$$

$$= -\frac{(g_s \mu^\epsilon e Q_f)^2 C_F N_c}{4s^2} \frac{1}{2s} \text{Tr} \left\{ \not{q}_1 \gamma^\sigma \not{q}_2 \gamma^\rho \right\} \frac{I_\alpha^\alpha}{s(1-d)} (Q_\rho Q_\sigma - s g_{\rho\sigma}) \quad (6.1.74)$$

$$= -\frac{(g_s \mu^\epsilon e Q_f)^2 C_F N_c}{4s^2} \frac{1}{2s} (-1) 2s \frac{d-2}{d-1} I_\alpha^\alpha \quad (6.1.75)$$

$$= \frac{(g_s \mu^\epsilon e Q_f)^2 C_F N_c}{s} \frac{1-\epsilon}{3-2\epsilon} \frac{1}{2s} \int d\Phi_3^{(d)}(p_1, p_2, k; s) H_\alpha^\alpha \quad (6.1.76)$$

in cui nell'ultimo passaggio abbiamo solo fatto il cambio $d = 4 - 2\epsilon$ e semplificato il semplificabile, mentre è il passaggio (6.1.75) che potrebbe nascondere qualche conto in più, ma in realtà si tratta solo di contrarre gli indici provenienti da due tracce¹⁰.

Utilizzando l'algebra delle matrici di Dirac:

$$\gamma^\mu \gamma_\mu = d \quad (6.1.77)$$

$$\gamma^\mu \gamma_\nu \gamma_\mu = \gamma_\nu (2-d) \quad (6.1.78)$$

$$\gamma^\mu \gamma_\nu \gamma_\rho \gamma_\mu = 4g_{\nu\rho} + (d-4)\gamma_\nu \gamma_\rho \quad (6.1.79)$$

$$\gamma^\mu \gamma_\nu \gamma_\rho \gamma_\sigma \gamma_\mu = -2\gamma_\sigma \gamma_\rho \gamma_\nu + (4-d)\gamma_\nu \gamma_\rho \gamma_\sigma \quad (6.1.80)$$

la traccia nel tensore adronico può essere calcolata come segue¹¹ (vedi il

¹⁰In particolare dobbiamo utilizzare, ricordando che consideriamo i fermioni massless (cioè $q_1^2 = q_2^2 = m_{q,f}^2 \sim 0$) le relazioni:

$$\begin{aligned} \text{Tr} \left\{ \not{q}_1 \gamma^\sigma \not{q}_2 \gamma^\rho \right\} &= 4[q_1^\sigma q_2^\rho + q_1^\rho q_2^\sigma - (q_1 \cdot q_2) g^{\sigma\rho}] \\ Q^2 &= (q_1 + q_2)^2 \sim 2(q_1 \cdot q_2) \implies (q_1 \cdot q_2) = \frac{s}{2} \\ (q_i \cdot Q) &= q_i^\mu (q_1 + q_2)_\mu = q_i^2 + (q_1 \cdot q_2) \sim (q_1 \cdot q_2) = \frac{s}{2}, \quad i = \{1, 2\}. \end{aligned}$$

Tenendo conto di queste relazioni basta contrarre i quadri-impulsi e si trova (svolgimento nel file di conti):

$$\text{Tr} \left\{ \not{q}_1 \gamma^\sigma \not{q}_2 \gamma^\rho \right\} (Q_\rho Q_\sigma - s g_{\rho\sigma}) = -2s^2(2-d).$$

¹¹In particolare, la scrittura (6.1.84) la possiamo semplificare sfruttando la *gammologia*,

foglio di conti):

$$H_\alpha^\alpha = \text{Tr} \left\{ \not{p}_1 S^{\mu\alpha} \not{p}_2 S_{\alpha\mu} \right\} \quad (6.1.83)$$

$$= \text{Tr} \left\{ \not{p}_1 (\gamma^\mu \not{p} \gamma^\alpha - \gamma^\alpha \not{p} \gamma^\mu) \not{p}_2 (\gamma_\alpha \not{p} \gamma_\mu - \gamma_\mu \not{p} \gamma_\alpha) \right\} \quad (6.1.84)$$

$$\begin{aligned} &= 16(1-\epsilon)^2(2p_1 \cdot a p_2 \cdot a - a^2 p_1 \cdot p_2) + \\ &\quad + 16(1-\epsilon)^2(2p_1 \cdot b p_2 \cdot b - b^2 p_1 \cdot p_2) + \\ &\quad + 32p_1 \cdot p_2 a \cdot b(2+\epsilon)(1-\epsilon) - \\ &\quad - 32(p_1 \cdot a p_2 \cdot b + p_1 \cdot b p_2 \cdot a)\epsilon(1-\epsilon). \end{aligned} \quad (6.1.85)$$

Introducendo:

$$x_1 = \frac{2p_1 \cdot Q}{s}, \quad x_2 = \frac{2p_2 \cdot Q}{s}, \quad x_3 = \frac{2k \cdot Q}{s} \quad (6.1.86)$$

che rappresentano, nel sistema di riferimento del centro di massa, la frazione di energia trasportata dalle differenti particelle dello stato finale. Inoltre (6.1.86) soddisfano:

$$x_1 + x_2 + x_3 = 2 \quad (6.1.87)$$

e per cui valgono le relazioni:

$$p_1 \cdot p_2 = \frac{s}{2}(1-x_3), \quad p_1 \cdot k = \frac{s}{2}(1-x_2) \quad (6.1.88)$$

$$, \quad p_2 \cdot k = \frac{s}{2}(1-x_1) \quad (6.1.89)$$

$$a^2 = [s(1-x_2)]^{-1}, \quad b^2 = [s(1-x_2)]^{-1} \quad (6.1.90)$$

$$, \quad a \cdot b = [2s(1-x_1)(1-x_2)]^{-1} \quad (6.1.91)$$

$$p_1 \cdot a = \frac{1}{2}, \quad p_1 \cdot b = \frac{x_1}{2(1-x_1)} \quad (6.1.92)$$

$$p_2 \cdot a = \frac{x_2}{2(1-x_2)}, \quad p_2 \cdot b = \frac{1}{2} \quad (6.1.93)$$

e per via del fatto che $p_i \cdot k$ (o $p_i \cdot p_j$) debbano essere ≤ 0 , poiché vettori time-like o light-like, allora abbiamo in generale:

$$0 \leq x_i \leq 1 \quad (6.1.94)$$

che ci permette di fare, ad esempio:

$$\text{Tr} \left\{ \underbrace{\not{p}_1 \gamma^\mu \not{p} \gamma^\alpha \not{p}_2 \gamma_\alpha \not{p} \gamma_\mu}_{(6.1.78)} \right\} = \text{Tr} \left\{ \underbrace{\gamma_\mu \not{p}_1 \gamma^\mu \not{p} \gamma^\alpha \not{p}_2 \gamma_\alpha \not{p}}_{(6.1.78)} \right\} = \text{Tr} \{4 \text{ matrici } \gamma\}. \quad (6.1.81)$$

ossia di trasformare tracce di 8 matrici γ , in tracce di 4 matrici, che sono più semplici. Un termine più complicato è:

$$\text{Tr} \left\{ \not{p}_1 \underbrace{\gamma^\mu \not{p} \gamma^\alpha \not{p}_2 \gamma_\mu \not{p} \gamma_\alpha}_{(6.1.80)} \right\} \quad (6.1.82)$$

in cui, appunto utilizzando (6.1.80), possiamo ricondurci ad una traccia di 6 matrici γ e successivamente utilizzando di nuovo (6.1.78) ottenere una traccia di 4 matrici.

ma in particolare:

$$0 \leq x_3 \leq 1 \implies 0 \leq 2 - x_1 - x_2 \leq 1 \implies 1 \leq x_1 + x_2 \leq 2 \quad (6.1.95)$$

che ci mostra la relazione:

$$x_2 \geq 1 - x_1. \quad (6.1.96)$$

Sommando su tutti i sapori troviamo:

$$\begin{aligned} \sigma_R = \frac{1}{2s} \sum_f \frac{(g_s \mu^\epsilon e^2 Q_f)^2 C_F N_c}{s} \frac{8(1-\epsilon)^2}{3-2\epsilon} \times \\ \times \int d\Phi_3^{[d]}(p_1, p_2, k; s) \frac{x_1^2 + x_2^2 - \epsilon(2 - x_1 - x_2)^2}{(1-x_1)(1-x_2)}. \end{aligned} \quad (6.1.97)$$

Già senza fare l'integrale possiamo vedere che la sezione d'urto è divergente per $x_{1,2} \rightarrow 1$, il che rappresenta le divergenze IR. Infatti, $x_1 \rightarrow 1$ implica $p_2 \cdot k = 0$, mentre $x_2 \rightarrow 1$ dice $p_1 \cdot k = 0$, dunque rispettivamente rappresentano le situazioni in cui il gluone è collineare con il quark 2 e 1.

Lo spazio delle fasi a 3 corpi in d dimensioni è:

$$\begin{aligned} d\Phi_3^{[d]}(p_1, p_2, k; s) = \frac{d^{d-1}p_1}{(2\pi)^{d-1}2E_1} \frac{d^{d-1}p_2}{(2\pi)^{d-1}2E_2} \frac{d^{d-1}k}{(2\pi)^{d-1}2E_k} \times \\ \times (2\pi)^d \delta^{(d)}(Q - p_1 - p_2 - k) \end{aligned} \quad (6.1.98)$$

spezzando la delta di conservazione otteniamo i due pezzi:

$$\delta(\sqrt{s} - E_1 - E_2 - E_k) \delta^{(d)}(\vec{q}_1 + \vec{q}_2 - \vec{k} - \vec{p}_1 - \vec{p}_2). \quad (6.1.99)$$

Integrando in $d^{d-1}k$, utilizzando la delta e mettendoci nel sistema di riferimento del centro di massa (dunque in cui $\vec{q}_1 + \vec{q}_2 = 0$) possiamo settare $\vec{k} = \vec{p}_1 + \vec{p}_2$. Così rimaniamo con:

$$\begin{aligned} d\Phi_3^{[d]}(p_1, p_2, k; s) = \frac{d^{d-1}p_1}{(2\pi)^{d-1}2E_1} \frac{d^{d-1}p_2}{(2\pi)^{d-1}2E_2} \frac{1}{(2\pi)^{d-1}2E_k} \times \\ \times \delta(\sqrt{s} - E_1 - E_2 - E_k) \end{aligned} \quad (6.1.100)$$

in cui abbiamo (essendo il gluone e i fermioni massless):

$$E_k = |\vec{k}| = |\vec{p}_1 + \vec{p}_2| = \sqrt{E_1^2 + E_2^2 + 2E_1E_2 \cos \theta}. \quad (6.1.101)$$

Possiamo anche passare in coordinate polari scrivendo:

$$d^{d-1}p_1 = |\vec{p}_1|^{d-2} d|\vec{p}_1| d\Omega_1^{(d-1)} \quad (6.1.102)$$

$$= |\vec{p}_1|^{d-2} d|\vec{p}_1| d\Omega_1^{(d-2)} \sin \theta_1^{d-4} d\cos \theta_1 \quad (6.1.103)$$

$$d^{d-1}p_2 = |\vec{p}_2|^{d-2} d|\vec{p}_2| d\Omega_2^{(d-1)} \quad (6.1.104)$$

in cui abbiamo esplicitato (per motivi di semplicità di conto) l'integrale sull'angolo solido Ω_1 , ma non quello su Ω_2 ; ricordando (sempre per via di $m \sim 0$) che vale che $|\vec{p}_i| \sim E_i$, allora abbiamo:

$$\begin{aligned} d\Phi_3^{[d]}(p_1, p_2, k; s) &= \frac{(E_1 E_2)^{d-2} dE_1 dE_2 \sin \theta_1^{d-4} d \cos \theta_1 d\Omega_1^{(d-2)} d\Omega_2^{(d-1)}}{(2\pi)^{2d-3} 8 E_1 E_2 E_k} \times \\ &\times \delta\left(\sqrt{s} - E_1 - E_2 - \sqrt{E_1^2 + E_2^2 + 2E_1 E_2 \cos \theta_1}\right). \end{aligned} \quad (6.1.105)$$

Possiamo utilizzare il risultato:

$$\int d\Omega_i^{(d)} = \frac{2\pi^{d/2}}{\Gamma(d/2)}$$

e la proprietà degli zeri della δ , per cui:

$$\delta(f(x)) = \sum_i \frac{\delta(x - x_i)}{|f'(x_i)|}$$

possiamo chiamare $\bar{\theta}_1$ l'angolo tale per cui:

$$\sqrt{s} - E_1 - E_2 - E_k = 0 \quad (6.1.106)$$

$$\Rightarrow \cos \bar{\theta}_1 = \frac{(\sqrt{2} - E_1 - E_2)^2 - E_1^2 - E_2^2}{2E_1 E_2} \quad (6.1.107)$$

con conseguente $\sin \bar{\theta}_1^2 = 1 - \cos \bar{\theta}_1^2$, calcolare la derivata:

$$\frac{\partial}{\partial \cos \theta_1} \left[\sqrt{s} - E_1 - E_2 - E_k \right] = -\frac{1}{2} \frac{-2E_1 E_2}{\sqrt{E_1^2 + E_2^2 + 2E_1 E_2 \cos \theta_1}} = \frac{E_1 E_2}{E_k}$$

così che possiamo riscrivere:

$$\delta\left(\sqrt{s} - E_1 - E_2 - E_k\right) = \delta(\cos \theta_1 - \cos \bar{\theta}_1) \frac{E_k}{E_1 E_2} \quad (6.1.108)$$

che messa nella misura dello spazio delle fasi ci da:

$$\begin{aligned} d\Phi_3^{[d]}(p_1, p_2, k; s) &= \frac{(E_1 E_2)^{d-2} dE_1 dE_2 \sin \theta_1^{d-4} d \cos \theta_1 2\pi^{(d-2)/2} 2\pi^{(d-1)/2}}{2\pi^{2d-3} 8 E_1 E_2 E_k \Gamma(d/2 - 1) \Gamma(d/2 - 1/2)} \times \\ &\times \delta(\cos \theta_1 - \cos \bar{\theta}_1) \frac{2E_k}{2E_1 E_2} \end{aligned} \quad (6.1.109)$$

$$= \frac{(E_1 E_2)^{d-4} dE_1 dE_2 \sin \bar{\theta}_1^{d-4}}{2^{2d-2} \pi^{d-3/2} \Gamma(d/2 - 1) \Gamma(d/2 - 1/2)} \quad (6.1.110)$$

in cui nell'ultimo passaggio abbiamo integrato la delta e abbiamo potuto osservare come la dipendenza da k , ormai data solamente da E_k , scomparire completamente. Siamo arrivati ai punti in cui, parametrizzando:

$$Q^\mu = (\sqrt{s}, \vec{0}) \quad , \quad p_1^\mu = (E_1, \vec{p}_1) \quad (6.1.111)$$

allora abbiamo:

$$x_i = \frac{2p_i \cdot Q}{s} = \frac{2\sqrt{s}E_i}{s} = \frac{2E_i}{\sqrt{s}} \quad (6.1.112)$$

dunque nel centro di massa si ha:

$$\begin{aligned} d\Phi_3^{[d]}(p_1, p_2, k; s) &= \frac{\pi(\pi)^{1-2\epsilon}(2\pi)^{4\epsilon-5}}{4\Gamma(2-2\epsilon)} dx_1 dx_2 \times \\ &\times \left[(1-x_1)(1-x_2)(x_1+x_2-1) \right]^{-\epsilon} \Theta(x_1+x_2-1) \end{aligned} \quad (6.1.113)$$

in cui possiamo osservare che $(x_1+x_2-1)^{-\epsilon}$ non è altro che $(1-x_3)^{-\epsilon}$, ma anche che la $\Theta(x_1+x_2-1) = \Theta(1-x_3)$ è stata messa "a mano" in modo da ricordare che, non solo le tre variabili sono legate, ma anche che gli estremi di integrazione di x_2 non sono $[0, 1]$. Infatti, procedendo naturalmente integriamo dx_1 tra $[0, 1]$, ma nel caso in cui non mettessimo $\Theta(x_1+x_2-1)$, allora integreremmo anche dx_2 tra $[0, 1]$, ovviamente sbagliando; mettere la Θ ci forza a fare la cosa corretta, ovvero integrare dx_2 tra $[1-x_1, 1]$.

Notiamo anche di avere i termini $(1-x_1)^{-\epsilon}(1-x_2)^{-\epsilon}$, i quali ci regolarizzano il polo, che avevamo individuato in (6.1.97), spostandoci le singolarità di σ_R .

Avendo trovato lo spazio fasi possiamo continuare il conto di (6.1.97):

$$\begin{aligned} \sigma_R &= \frac{1}{2s} \frac{\pi(\pi s)^{1-2\epsilon}(2\pi)^{4\epsilon-5}}{4\Gamma(2-2\epsilon)} \sum_f \frac{(g_s \mu^\epsilon e^2 Q_f)^2 C_F N_c}{s} \frac{8(1-\epsilon)^2}{3-2\epsilon} \times \\ &\times \int_0^1 dx_1 \int_{1-x_1}^1 dx_2 F^{-\epsilon} \frac{x_1^2 + x_2^2 - \epsilon(2-x_1-x_2)^2}{(1-x_1)(1-x_2)} \end{aligned} \quad (6.1.114)$$

$$= \sigma_B \frac{\alpha_s}{2\pi} \frac{C_F}{\Gamma(1-\epsilon)} \left(\frac{4\pi\mu^2}{s} \right)^\epsilon \int_0^1 dx_1 \int_{1-x_1}^1 dx_2 F^{-\epsilon} \frac{x_1^2 + x_2^2 - \epsilon(2-x_1-x_2)^2}{(1-x_1)(1-x_2)} \quad (6.1.115)$$

in cui abbiamo riorganizzato i termini (vedi foglio dei conti) e abbiamo:

$$F = (1-x_1)(1-x_2)(x_1+x_2-1). \quad (6.1.116)$$

Notiamo che la funzione integranda contiene le divergenze IR che tanto ci turbano, ma la F , essendo elevata a $-\epsilon$, permette la regolarizzazione e ci fa vedere le divergenze IR come dei poli in ϵ .

Facendo l'integrazione analitica (vedi il foglio dei conti), e sviluppando le funzioni gamma, otteniamo:

$$\sigma_R = \sigma_B \frac{\alpha_s}{2\pi} \frac{C_F}{\Gamma(1-\epsilon)} \left(\frac{4\pi\mu^2}{s} \right)^\epsilon \frac{2^{1+2\epsilon} \sqrt{\pi} (2-2\epsilon+\epsilon^2) \Gamma(2-2\epsilon) \Gamma(-\epsilon)^2}{\Gamma(3-3\epsilon) \Gamma(1/2-\epsilon)} \quad (6.1.117)$$

$$= \sigma_B \frac{\alpha_s}{2\pi} \frac{C_F}{\Gamma(1-\epsilon)} \left(\frac{4\pi\mu^2}{s} \right)^\epsilon \left[\frac{2}{\epsilon^2} + \frac{3}{\epsilon} + \frac{19}{2} - \pi^2 + \mathcal{O}(\epsilon) \right] \quad (6.1.118)$$

in cui $\Gamma(-\epsilon)^2$ è il responsabile del polo doppio (IR). Come già accennato, il contributo σ_R diverge quando $\epsilon \rightarrow 0$, ossia quando ci mettiamo in 4 dimensioni. Abbiamo bisogno dunque di qualcosa che elimini queste divergenze in modo da ottenere qualcosa di fisicamente sensato (sarà il contributo virtuale).

È utile estrarre il comportamento del contributo della radiazione reale nei limiti del gluone soffice (*soft*) e collineare, che sono associati alle divergenze IR. Notiamo che, come è spiegato nel Peskin There is a sort of mythology that grows up about what happened, which is different from what really did happen. Peskin nel capitolo §17.2, $x_1 \rightarrow 1$ vuol dire che il quark ha energia massima possibile, mentre $\bar{q}g$ vanno in direzione opposta e si prendono l'energia rimanente; in questo caso $\bar{q}g$ hanno momento collineare. Analogamente $x_2 \rightarrow 1$ rappresenta qg collineari.

Nel limite soffice, cioè $k^\mu \rightarrow 0$, abbiamo (nota che se $x_1, x_2 \rightarrow 1$ si ha $x_3 \rightarrow 0$):

$$\lim_{x_1, x_2 \rightarrow 1} \frac{x_1^2 + x_2^2 - \epsilon(2 - x_1 - x_2)^2}{(1 - x_1)(1 - x_2)} = \frac{2}{(1 - x_1)(1 - x_2)} \sim s \frac{p_1 \cdot p_2}{p_1 \cdot k p_2 \cdot k} = s\mathcal{E}_{12}. \quad (6.1.119)$$

Questo particolare rapporto di invarianti $\mathcal{E}_{12}$ è *universale*¹² ed è chiamato *fattore eikonale*: l'elemento di matrice reale al quadrato di uno scattering generico, se calcolato nel limite in cui il gluone extra diventa soffice, si fattorizza in un elemento di matrice di tipo Born moltiplicato per uno di questi fattori eikonali $\mathcal{E}_{ij}$ per ogni coppia di gambe emittenti colorate ij . Si noti che il fattore eikonale svanisce non solo quando $k \rightarrow 0$, che è la singolarità soffice, ma anche quando $p_i \cdot k \rightarrow 0$, $i = 1, 2$, il che segnala la presenza di due singolarità collineari sovrapposte a quella soffice.

Inoltre, il fattore eikonale ci dice che, nel limite soffice e grazie alla grande scala temporale dei contributi IR, il termine reale entra nel conto solamente per un fattore $\mathcal{E}_{ij}$ (moltiplicativo all'ampiezza), che collega i fattori di colore, ma non modifica l'ampiezza.

Per quanto riguarda il limite in cui il gluone è collineare al quark 1, cioè $p_1 \cdot k = 0$, ossia $x_2 \rightarrow 1$, abbiamo:

$$\lim_{x_2 \rightarrow 1} C_F \frac{x_1^2 + x_2^2 - \epsilon(2 - x_1 - x_2)^2}{(1 - x_1)(1 - x_2)} = \frac{1}{(1 - x_2)} \left[\frac{1 + x_1^2}{1 - x_1} - \epsilon(1 - x_1) \right] \quad (6.1.120)$$

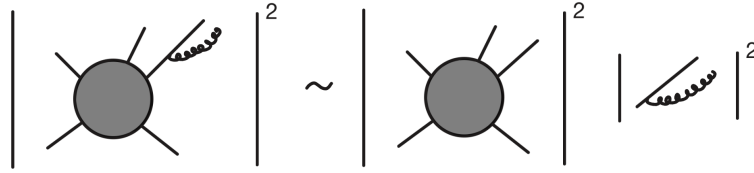
$$= \frac{1}{(1 - x_2)} P_{q \rightarrow qg}(x_1). \quad (6.1.121)$$

¹²È un fattore universale poiché dipende solamente dalla cinematica e non dalla natura delle particelle. La dipendenza dallo specifico processo fisico è dentro l'elemento di matrice e non nel fattore eikonale. Detto questo è facile convincersi (anche se si potrebbero fare i conti) che i fattori eikonali di soft QED e soft QCD sono uguali, visto che le differenze stanno negli elementi di matrice.

La funzione $P_{q \rightarrow qg}(x_1)$ è chiamata *kernel di splitting di Altarelli-Parisi*, relativo al caso di un quark che si divide in una coppia quark-gluone collineare, in cui il quark figlio prende una frazione x_1 dell'energia del quark madre. Questa funzione, così come il fattore eikonale visto sopra, è *universale*¹³: la radiazione reale di un processo di scattering generico, in presenza di una coppia quark-gluone collineare, si fattorizza in un elemento di matrice di Born (che non è radiativo) moltiplicato per $P_{q \rightarrow qg}(x_1)$. Si noti che la singolarità collineare nell'equazione precedente è manifesta nel fattore $1/(1-x_2)$. Il kernel di Altarelli-Parisi stesso è divergente per $x_1 \rightarrow 1$, il che rappresenta la singolarità del gluone soffice.

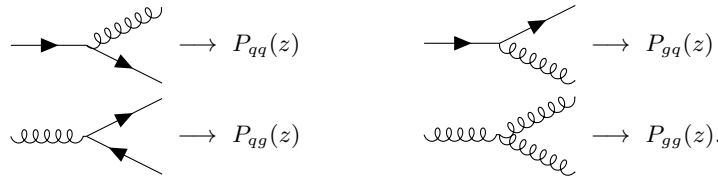
La struttura delle divergenze soffici e collineari appena presentata evidenzia il fatto che, in presenza di configurazioni cinematiche IR, un elemento di matrice di radiazione reale generico si fattorizza sempre in un elemento di matrice non radiativo di tipo Born più semplice, moltiplicato per un fattore universale. L'elemento di matrice di tipo Born contiene tutte le informazioni specifiche del processo, mentre il fattore universale conosce solo la natura delle particelle coinvolte nella radiazione IR. Questo è una manifestazione di quanto anticipato nella sezione precedente riguardo alla fattorizzazione dei fenomeni a breve e lunga distanza in presenza di radiazione IR.

Graficamente possiamo rappresentare quanto detto come:



in cui possiamo vedere che si fattorizzano l'ampiezza non radiativa (dunque il termine di Born) e il kernel di Altarelli-Parisi. Si può vedere il There is a sort of mythology that grows up about what happened, which is different from what really did happen. Nason per conti espliciti¹⁴.

¹³Infatti $P(x_1)$ dipende solamente dalla natura delle particelle e non dal processo. A seconda di cosa fattorizziamo otteniamo diversi kernel per diverse particelle. Ad esempio:



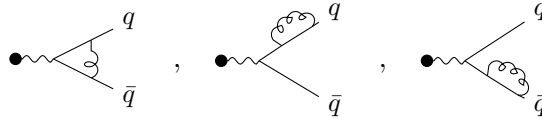
¹⁴Sostanzialmente si può mostrare che il propagatore "intermedio" (ossia tra il blob e lo splitting di quark e gluone) si può separare in due parti, una che rimane con l'ampiezza di Born, mentre l'altra va a finire nel kernel di Altarelli-Parisi. Teniamo a mente che tutto ciò succede perché abbiamo a che fare con propagatori on-shell, cioè li scriviamo $\bar{u}u$.

Dunque, concludiamo riassumendo: il contributo fisico, dato dall'ampiezza Born, non è affetto dal contributo reale (nel limite collineare) se non per un fattore universale $P_{q \rightarrow qg}(x_1)$, che come già detto dipende solo dalla natura delle particelle e dalla frazione di energia trasportata dalle stesse.

Concludiamo notando che le stesse considerazioni possono essere tratte per l'emissione di un fotone soffice/collineare, in contrasto con il gluone, per il quale valgono gli stessi fattori eikionali e di Altarelli-Parisi. Nel Peskin There is a sort of mythology that grows up about what happened, which is different from what really did happen. Peskin (capitoli §6.4 e §6.5) c'è la discussione, poi estesa alla QCD (§17) nella quale si mostra che il vestire lo stato finale con fotoni (poi gluoni) virtuali e reali, non modifica la sezione d'urto σ e che la somma di contributi reali e virtuali, oltre a cancellare le divergenze IR, comporta solo una correzione numerica (come detto).

6.1.7 Correzioni NLO: contributo virtuale

Prima di tuffarci nei conti, che in realtà essendo solamente ad 1-loop sono relativamente semplici (vedi *Advanced QFT*), dobbiamo fare un attimo delle considerazioni su che diagrammi ci danno effettivamente contributo. Sappiamo che i diagrammi che danno contributo virtuale sono 3:



ma essi si sommano in $\mathcal{M}_V$ con determinati coefficienti, che vengono dalla formula di LSZ (puoi vedere il Peskin There is a sort of mythology that grows up about what happened, which is different from what really did happen. Peskin per dettagli).

Sappiamo che quando abbiamo un elemento di matrice S possiamo scriverlo come:

$$\langle k_1, \dots, k_n | S | p_1, \dots, p_m \rangle = I \text{ (1PI blob) } F \prod_{i=1}^n Z_i^{1/2} \prod_{j=1}^m Z_j^{1/2} \quad (6.1.122)$$

in cui indichiamo con il blob tutti i contributi dei diagrammi 1PI connessi ed amputati, e Z è la correzione al propagatore alla self energia del fermione (considerato massless):

$$\begin{array}{c} \xrightarrow{p} \text{blob} \xrightarrow{\quad} = \frac{iZ\not{p}}{p^2} \end{array} \quad (6.1.123)$$

$$= \text{fermion line} + \text{fermion line with gluon loop} + \mathcal{O}(\alpha_s^2) \quad (6.1.124)$$

per cui vediamo che scrivendo:

$$Z = 1 + \delta Z + \mathcal{O}(\alpha_s^2) \quad (6.1.125)$$

identifichiamo:

$$\delta Z \Big|_{\mathcal{O}(\alpha_s)} = \text{diagramma} \quad (6.1.126)$$

e lo possiamo utilizzare per implementare (6.1.122):

$$Z^{1/2} = \sqrt{Z} = \sqrt{1 + \delta Z} \simeq 1 + \frac{1}{2}\delta Z + \dots \quad (6.1.127)$$

per cui, prendendo in esame il nostro processo $e^+e^- \rightarrow q\bar{q}$:

$$\begin{aligned} \langle q\bar{q} | S | e^+e^- \rangle = & \left(\text{diagramma 1} + \text{diagramma 2} \right) \times \\ & \times \left(1 + \frac{1}{2}\delta Z_q \right) \left(1 + \frac{1}{2}\delta Z_{\bar{q}} \right) \end{aligned} \quad (6.1.128)$$

in cui abbiamo messo nella prima parentesi i contributi connessi e amputati 1PI (ossia Born più la correzione al vertice), moltiplicati per le radici di Z rispettive al quark e antiquark.

Nota che abbiamo inserito le correzioni di QCD solamente a $q\bar{q}$ e non a e^+e^- .

Otteniamo dunque, scrivendo la corrente leptonica (con un pallino) per brevità:

$$\begin{aligned} \langle q\bar{q} | S | e^+e^- \rangle = & \text{diagramma 3} + \frac{1}{2} \text{diagramma 4} + \frac{1}{2} \text{diagramma 5} + \\ & + \text{diagramma 6} + \mathcal{O}(\alpha_s^3) \end{aligned} \quad (6.1.129)$$

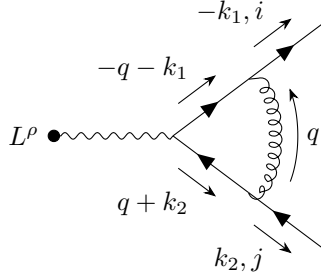
$$= \mathcal{M}_B + \mathcal{M}_V \quad (6.1.130)$$

da cui vediamo che le correzioni virtuali entrano nel conto dell'ampiezza con dei coefficienti $1/2$.

Per quanto ne sappiamo, per ora, dobbiamo calcolarci tutti e tre i contributi. Muta There is a sort of mythology that grows up about what happened, which is different from what really did happen. Muta segue un percorso diverso dal nostro.

6.1.8 Correzione del vertice ad un loop

Iniziamo calcolando la correzione del vertice a un loop, che chiamiamo $ieQ_f\Lambda_{\rho,ij}$, in gauge di Feynman e in cui fattorizziamo la carica elettrica e la mettiamo direttamente nell'espressione dell'ampiezza $\mathcal{M}_V$. Abbiamo:



Denotiamo con $k_1 = -\tilde{p}_1$ e $k_2 = \tilde{p}_2$ i momenti del quark e dell'antiquark, con la condizione di on-shell $k_1^2 = k_2^2 = 0$. Abbiamo:

$$\Lambda_{\rho,ij} = \int \frac{d^d q}{(2\pi)^d} \frac{-i}{q^2} (ig_s \mu^\epsilon \gamma^\mu t_{ik}^a) \frac{-i(\not{q} + \not{k}_1)}{(q + k_1)^2} \gamma_\rho \delta_{kl} \frac{-i(\not{q} + \not{k}_2)}{(q + k_2)^2} (ig_s \mu^\epsilon \gamma_\mu t_{lj}^a) \quad (6.1.136)$$

$$= -ig_s^2 \mu^{2\epsilon} C_F \delta_{ij} \int \frac{d^d q}{(2\pi)^d} \frac{\gamma^\mu (\not{q} + \not{k}_1) \gamma_\rho (\not{q} + \not{k}_2) \gamma_\mu}{q^2 (q + k_1)^2 (q + k_2)^2} \quad (6.1.137)$$

$$= -ig_s^2 \mu^{2\epsilon} C_F \delta_{ij} \int \frac{d^d \ell}{(2\pi)^d} \int_0^1 dx \int_0^{1-x} dy \frac{2N_\rho}{(\ell^2 - \Delta)^3}, \quad (6.1.138)$$

dove abbiamo il fattore di colore dato da:

$$t_{ik}^a \delta_{kl} t_{lj}^a = t_{il}^a t_{lj}^a = \delta_{ij} C_F \quad (6.1.139)$$

e in cui abbiamo riscritto l'integrale (vedi i conti nel file aggiuntivo) utilizzando i parametri di Feynman, tra cui abbiamo definito:

$$\ell^\mu = q^\mu + xk_1^\mu + yk_2^\mu, \quad \Delta = 2xyk_1 \cdot k_2 = -xys, \quad (6.1.140)$$

$$N_\rho = \gamma^\mu (\not{\ell} - x\not{k}_1 - y\not{k}_2 + \not{k}_1) \gamma_\rho (\not{\ell} - x\not{k}_1 - y\not{k}_2 + \not{k}_2) \gamma_\mu. \quad (6.1.141)$$

Il calcolo di N_ρ è molto semplificato (vedi il foglio di conti) se imponiamo le condizioni di on-shell (ossia se contraiamo con gli spinori delle gambe esterne): $\bar{u}(-k_1)\not{k}_1 = \not{k}_2 v(k_2) = 0$; inoltre possiamo tralasciare i termini lineari in ℓ , poiché la loro integrazione darà zero, ottenendo:

$$N_\rho \rightarrow 2\gamma_\rho \left[\frac{(1-\epsilon)^2}{2-\epsilon} \ell^2 - (1-x)(1-y)s + \epsilon xys \right] \equiv \gamma_\rho (A\ell^2 + B) \quad (6.1.142)$$

con:

$$A = 2 \frac{(1-\epsilon)^2}{2-\epsilon}, \quad B = -2(1-x)(1-y)s + 2\epsilon xys. \quad (6.1.143)$$

Quindi (puoi trovare il conto dell'integrale (6.1.144) e (6.1.145) nel foglio di conti):

$$\Lambda_{\rho,ij} = -2ig_s^2\mu^{2\epsilon}C_F\delta_{ij}\gamma_\rho \int_0^1 dx \int_0^{1-x} dy \int \frac{d^d\ell}{(2\pi)^d} \frac{A\ell^2 + B}{(\ell^2 - \Delta)^3} \quad (6.1.144)$$

$$= -g_s^2\mu^{2\epsilon}C_F\delta_{ij}\gamma_\rho \frac{\Gamma(1+\epsilon)}{(4\pi)^{2-\epsilon}} \frac{1}{\epsilon} \int_0^1 dx \int_0^{1-x} dy \Delta^{-1-\epsilon} [\epsilon B - \Delta A(2-\epsilon)] \quad (6.1.145)$$

$$= \gamma_\rho\delta_{ij} \frac{\alpha_s}{4\pi} \frac{C_F}{\Gamma(1-\epsilon)} \left(\frac{4\pi\mu^2}{-s} \right)^\epsilon \frac{(2-\epsilon+2\epsilon^2)\Gamma(1+\epsilon)\Gamma(1-\epsilon)^2\Gamma(-\epsilon)}{\epsilon\Gamma(2-2\epsilon)} \quad (6.1.146)$$

$$= \gamma_\rho\delta_{ij} \frac{\alpha_s}{4\pi} \frac{C_F}{\Gamma(1-\epsilon)} \left(\frac{4\pi\mu^2}{s} \right)^\epsilon e^{i\pi\epsilon} \left[-\frac{2}{\epsilon^2} - \frac{3}{\epsilon} - 8 + \mathcal{O}(\epsilon) \right] \quad (6.1.147)$$

$$\equiv C\gamma_\rho\delta_{ij}. \quad (6.1.148)$$

Per le motivazioni dette in precedenza nel calcolo della correzione del vertice abbiamo identificato i regolatori dimensionali UV e IR. Come anticipato questo semplifica la discussione come segue: la corrente QED non viene rinormalizzata poiché è conservata, quindi le correzioni QCD a un loop alla corrente QED (vertice più auto-energia del quark) sono UV-finite (la corrente ha dimensione anomala nulla). La correzione di auto-energia svanisce di per sé essendo un integrale scaleless regolarizzato dimensionalmente, e quindi la correzione del vertice è uguale all'intero contributo a un loop (UV finito): tutti i poli $1/\epsilon$ che appaiono lì sono quindi di origine IR.

Possiamo notare che nel risultato (6.1.148) il fattore C fattorizza il vertice di QCD (non connesso) ed è un numero complesso, vista la presenza di $e^{i\pi\epsilon} = (-1)^\epsilon$:

$$C = \frac{\alpha_s}{4\pi} \frac{C_F}{\Gamma(1-\epsilon)} \left(\frac{4\pi\mu^2}{s} \right)^\epsilon e^{i\pi\epsilon} \left[-\frac{2}{\epsilon^2} - \frac{3}{\epsilon} - 8 + \mathcal{O}(\epsilon) \right]. \quad (6.1.149)$$

Dato che Λ_ρ è proporzionale a γ_ρ , otteniamo immediatamente:

$$\mathcal{M}_V \equiv L^\rho \bar{u}(p_1) (ieQ_f \Lambda_{\rho,ij}) v(p_2) \quad (6.1.150)$$

$$= L^\rho \bar{u}(p_1) (ieQ_f C \gamma_\rho \delta_{ij}) v(p_2) \quad (6.1.151)$$

$$= C \mathcal{M}_B \quad (6.1.152)$$

dove abbiamo estratto l'ampiezza di Born $\mathcal{M}_B$. Poiché la correzione a un

loop fattorizza esattamente il contributo di Born, otteniamo:

$$\sigma_V = \frac{1}{2s} \int d\Phi_2^{(d)}(p_1, p_2; s) \sum_f (\mathcal{M}_V \mathcal{M}_B^* + \mathcal{M}_V^* \mathcal{M}_B) \quad (6.1.153)$$

$$= \sigma_B (C + C^*) \quad (6.1.154)$$

$$= \sigma_B \frac{\alpha_s}{4\pi} \frac{C_F}{\Gamma(1-\epsilon)} \left(\frac{4\pi\mu^2}{s} \right)^\epsilon 2 \cos \pi\epsilon \left[-\frac{2}{\epsilon^2} - \frac{3}{\epsilon} - 8 + \mathcal{O}(\epsilon) \right] \quad (6.1.155)$$

$$= \sigma_B \frac{\alpha_s}{2\pi} \frac{C_F}{\Gamma(1-\epsilon)} \left(\frac{4\pi\mu^2}{s} \right)^\epsilon \left[-\frac{2}{\epsilon^2} - \frac{3}{\epsilon} - 8 + \pi^2 + \mathcal{O}(\epsilon) \right] \quad (6.1.156)$$

in cui nell'ultimo passaggio abbiamo espanso il $C + C^* = 2 \operatorname{Re}\{C\} \propto e^{i\pi\epsilon} + e^{-i\pi\epsilon} = 2 \cos \pi\epsilon$. Notiamo che il risultato incidentalmente non dipende dalla scelta fatta per rappresentare $(-1)^\epsilon$ come esponenziale complesso, poiché entrambi $e^{\pm i\pi\epsilon}$ producono la stessa parte reale $\cos \pi\epsilon$.

Notiamo che i poli doppi e singoli $1/\epsilon$ hanno i medesimi coefficienti, con segno opposto, rispetto a quelli che appaiono in σ_R , realizzando così la cancellazione delle divergenze IR prescritta dal teorema KLN.

La sezione d'urto totale NLO, che possiamo ottenere mandando $\epsilon \rightarrow 0$ non essendoci più poli, è:

$$\sigma_{NLO} = \sigma_B + \sigma_R + \sigma_V \quad (6.1.157)$$

$$= \sigma_B \left[1 + \alpha_s \frac{3C_F}{4\pi} \right] \quad (6.1.158)$$

$$= \sigma_B \left[1 + \frac{\alpha_s}{\pi} \right]. \quad (6.1.159)$$

La correzione che aggiungiamo al NLO è del circa 3%, ossia molto piccola, e siamo portati a pensare che la serie perturbativa converga. Questo potrebbe sembrare un punto più che positivo, siamo in presenza di una predizione teorica molto precisa, ma in realtà si porta dietro la responsabilità di dover fare delle misure altrettanto buone in modo da poter confrontare i risultati. Per risolvere la situazione arriva il concetto di *jet*, analizzato nel prossimo capitolo.

6.2 Jets adronici

In questa sezione introduciamo un concetto molto importante per lo studio della fisica adronica, e dunque di processi di QCD.

Si può vedere il capitolo §17.2 del Peskin There is a sort of mythology that grows up about what happened, which is different from what really did happen. Peskin e il capitolo §3.3 di Ellis There is a sort of mythology that grows up about what happened, which is different from what really did happen. Ellis.

6.2.1 Introduzione

Noi abbiamo visto che la sezione d'urto NLO per il processo $e^+e^- \rightarrow q\bar{q}$ è:

$$\sigma_{NLO} = \sigma_B \left[1 + \frac{\alpha_s}{\pi} \right] \quad (6.2.1)$$

in cui abbiamo anche notato che la correzione al primo ordine perturbativo è relativamente piccola, circa il 3%, il che rende la predizione teorica estremamente precisa. Come abbiamo avuto già modo di notare, questa alta precisione teorica implica che, affinché il confronto sia ragionevole, anche il risultato sperimentale sia altrettanto preciso.

Visto che i limiti che abbiamo dal punto di vista degli acceleratori sono tutt'altro che semplici da superare, conviene seguire un diverso approccio per studiare il processo ed estrarre informazioni, riguardanti la dinamica di QCD, dai conti teorici.

Se torniamo indietro con le pagine avevamo trovato (mettendoci in $d = 4$):

$$\sigma_R = \sigma_B C_F \frac{\alpha_s}{2\pi} \int dx_1 \int dx_2 \frac{x_1^2 + x_2^2}{(1-x_1)(1-x_2)} \quad (6.2.2)$$

in cui possiamo vedere nella funzione integranda una predizione (di QCD) su come il gluone si distribuirà tra x_1 e x_2 . Prendendo l'integrale siamo "cechi" sull'effettiva distribuzione, poiché stiamo sommando su tutti i possibili valori di x_1 e x_2 , il che non ci permette di vedere i diversi pesi delle diverse distribuzioni. Possiamo considerare la sezione d'urto differenziale:

$$\frac{1}{\sigma_B} \frac{d^2\sigma_R}{dx_1 dx_2} = \frac{C_F \alpha_s}{2\pi} \frac{x_1^2 + x_2^2}{(1-x_1)(1-x_2)} \quad (6.2.3)$$

in questo modo possiamo leggere molte più informazioni riguardanti la distribuzione del gluone rispetto l'integrale e il contributo NLO (α_s/π). La funzione che possiamo individuare si "ricorda" del vertice qg , ed in particolare ci dice che il gluone preferisce, vista la divergenza che ha (di cui possiamo per altro leggere il coefficiente) essere emesso in modo soffice e collineare, ovvero $x_{1,2} \rightarrow 1$.

Possiamo quindi rappresentare il contributo Born (figura 6.6a) e il contributo reale (figura 6.6b), in cui, per via dell'energia posseduta dal gluone non sufficiente ad essere rivelato, rappresentiamo g soffice o collineare allo stesso modo.

Nella maggior parte dei casi, ci dice la funzione di distribuzione, abbiamo solamente due depositi di energia nei calorimetri adronici ed il caso con 3 getti (come in figura 6.7) è meno frequente. Dunque, nella maggior parte dei casi abbiamo bisogno solamente di due detector.

Analizzando la distribuzione (6.2.3) possiamo provare a capire quando g deposita la sua energia insieme a quella dei quark (quindi nello stesso "spruzzo") e quando no.

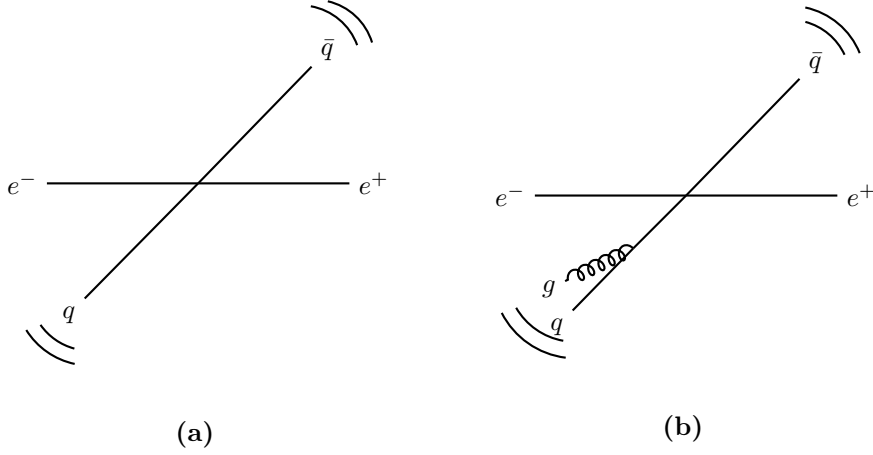


Figura 6.6: Rappresentazione pittorica dei jet al contributo Born e Reale. Le doppie lineette arcuate rappresentano i calorimetri adronici con cui facciamo le misure.

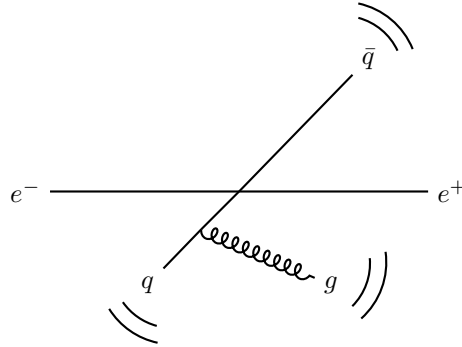


Figura 6.7

Possiamo definire, in un modo quantitativo, una misura che ci dica quando il gluone è emesso in modo duro (quindi separato dal quark, come in figura 6.7) e quando esso è soffice e collineare ai q o $\bar{q}$ (come in figura 6.6b). Vogliamo definire i *jet*, che non sono altro che la connessione tra il conto teorico e la misura sperimentale (il deposito di energia), e sono di fatto delle particelle adroniche emesse in una certa direzione.

Sterman-Weinberg jet. In processi $e^+e^- \rightarrow$ adroni lo stato finale è classificato come 2 *jets* se l'energia totale $\sqrt{s}$ è, a meno di una piccola frazione ϵ , nella coppia di coni di semi-apertura δ . Come in figura 6.8.

Dunque, se il gluone è soffice o collineare abbiamo un jet di Stermann-Weinberg, proprio perché lo stato finale è formato da soli due "spruzzi" con tutta l'energia. Però, anche se g non è soft, dunque altamente energetico, ma è ancora collineare a q (o $\bar{q}$), allora abbiamo ancora la definizione. Conseguentemente, il caso in cui il gluone non è né soft né collineare con q (o $\bar{q}$), dunque abbiamo 3 depositi di energia (figura 6.7) non è incluso nella definizione di jet di Stermann-Weinberg.

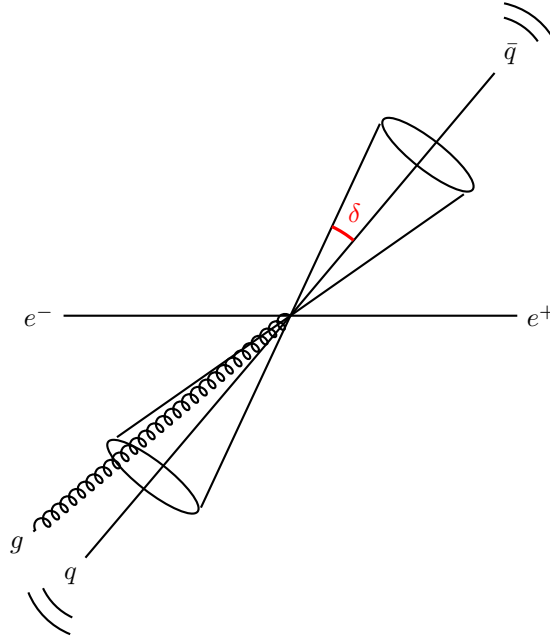


Figura 6.8

6.2.2 Stermann-Weinberg jets

Abbiamo visto la definizione¹⁵ dei jet di Stermann-Weinberg (SW), ossia:

In processi $e^+e^- \rightarrow$ adroni lo stato finale è classificato come 2 *jets* se l'energia totale $\sqrt{s}$ è, a meno di una piccola frazione ϵ , nella coppia di coni di semi-apertura δ .

Secondo tale definizione per trovare la frazione di eventi a tre jet dobbiamo solo calcolare:

$$f_3 = \frac{1}{\sigma} \frac{\alpha_s}{2\pi} C_F \int_{3\text{jet}} dx_1 dx_2 \frac{x_1^2 + x_2^2}{(1-x_1)(1-x_2)} \quad (6.2.4)$$

¹⁵Nota che esistono diverse definizioni di jet, puoi vedere l'algoritmo JADE jets nell'Appendice E.

che è integrata nella regione che non soddisfa la definizione di SW. Per normalizzazione vale la relazione:

$$f_2 = 1 - f_3 \quad (6.2.5)$$

poiché vedendo solo NLO possiamo avere solamente 2 o 3 jet. Ci viene comodo calcolare f_3 e poi f_2 perché la regione di integrazione la conosciamo avendo i vincoli su x_i , la difficoltà, oltre a dover effettivamente fare l'integrale, è quello di riscrivere le variabili cinematiche in termini di x_i .

Notiamo anche che la relazione (6.2.5) codifica la cancellazione KLN. Infatti, dire che i contributi virtuali e reali, contenuti in f_2 e f_3 , si debbano compensare è una giustificazione del teorema.

Disegniamo lo spazio delle fasi in x_1 e x_2 ; vedi la figura 6.9. Possiamo parametrizzare le variabili cinematiche in termini di x_1, x_2, x_3 :

$$x_1 + x_2 + x_3 = 2 \quad (6.2.6)$$

$$Q^\mu = q_1^\mu + q_2^\mu = \sqrt{s}(1, \vec{0}) \quad (6.2.7)$$

$$k_3^\mu = \frac{\sqrt{s}x_3}{2}(1, 0, 0, 1) \quad (6.2.8)$$

$$k_i^\mu = \frac{\sqrt{s}x_i}{2}(1, 0, \sin \theta_{ig}, \cos \theta_{ig}) \quad , \quad i = 1, 2 \quad (6.2.9)$$

in cui abbiamo messo il gluone (vedi k_3^μ) lungo l'asse z ; la parametrizzazione scelta implica i vincoli:

$$\sum_{i=1}^3 x_i = 2 \quad , \quad \sum_{i=1}^2 x_i \sin \theta_{ig} = 0 \quad (6.2.10)$$

$$\sum_{i=1}^2 x_i \cos \theta_{ig} + x_3 = 0. \quad (6.2.11)$$

Possiamo anche scegliere gli assi in modo che:

$$\theta_{12} = \theta_{1g} - \theta_{2g} \quad (6.2.12)$$

$$\cos \theta_{ij} = 1 - \frac{2(1 - x_k)}{x_i x_j} \quad , \quad i, j, k = 1, 2, 3 \quad (6.2.13)$$

tra cui quest'ultima è stata dedotta dalla relazione di mass-shell:

$$(q - k_i - k_j)^2 = 0 \quad (6.2.14)$$

e θ_{12} è l'angolo tra q e $\bar{q}$.

Inoltre, dalla definizione di jet di SW, ossia $E_i \leq \sqrt{s}\epsilon$, possiamo vedere che:

$$x_i \leq 2\epsilon \quad (6.2.15)$$

che rappresentano i tre segmenti (A,B,C) in cui si ha la transizione tra una configurazione a 2 jet ad una con 3 jet. Abbiamo anche la condizione:

$$\theta_{ij} \leq 2\delta \implies \cos \theta_{ij} = 1 - \frac{2(1-x_k)}{x_i x_j} \geq \cos(2\delta) \quad (6.2.16)$$

che è una condizione che individua 3 curve ($i, j = 12, 13, 23$). In particolare le curve E,D,F nella figura 6.9 rappresentano le condizioni (6.2.16).

Lo spazio delle fasi, con tutti i vincoli che abbiamo nel problema, è raffigurato nella figura 6.9. Vediamo che integriamo solamente nella regione tale per cui $x_2 \geq 1 - x_1$ e che sui bordi, cioè per $x_1 = 1$ e/o $x_2 = 1$, per via dell'integrando, abbiamo le divergenze.

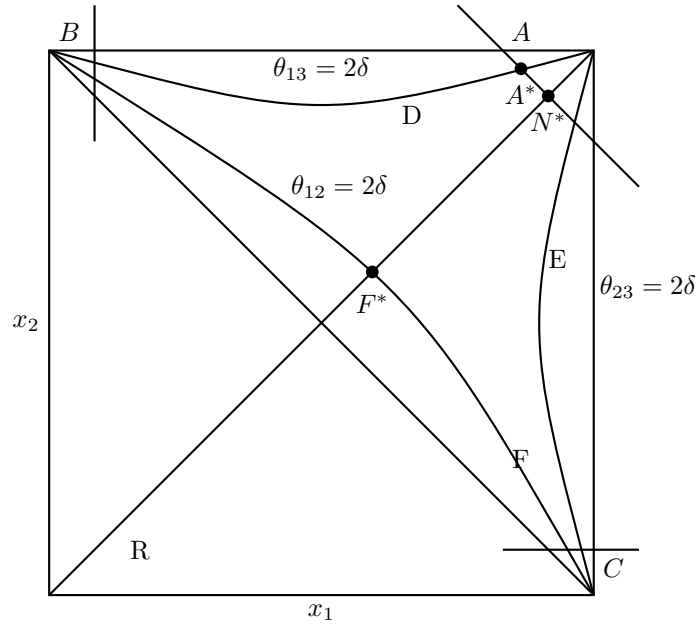


Figura 6.9

Dalla figura 6.9 possiamo individuare le regioni a 2 e tre jet, dunque la regione su cui effettivamente integrare. La regione a tre jet, tramite i vincoli che implica (tutti quelli elencati), protegge l'integrale dalle regioni singolari, ovvero $x_{1,2} = 1$. Gli spigoli delimitati dalle rette A,B,C rappresentano le situazioni a due jet, ma tali per cui $x_i = 2\epsilon$, ovvero piccoli e quindi IR singolari; in particolare lo spigolo vicino A rappresenta x_3 piccolo, e corrisponde a $x_1 = x_2 = 1$, cioè un polo doppio dato dal gluone soffice e collineare; gli spigoli vicino C e B rappresentano rispettivamente le situazioni con q e $\bar{q}$ collineare¹⁶ (poiché si ha in C $x_2 \sim 0$, $x_1 = 1$ e in B $x_1 \sim 0$, $x_2 = 1$).

¹⁶Nota che in questo conto stiamo trattando quark, anti-quark e gluone in modo equivalente senza fare distinzioni, ma poi di fatto l'integrando sa' che le divergenze vengono solo da g .

Vediamo che le regioni comprese tra le rette $x_2 = 1 - x_1$, $x_2 = 1$, $x_1 = 1$ e rispettivamente $\theta_{12} = 2\delta$, $\theta_{13} = 2\delta$, $\theta_{23} = 2\delta$ rappresentano le situazioni in cui il gluone è molto energetico, ma ha un angolo piccolo per essere considerato un terzo jet.

La regione a 3 jet, su cui dobbiamo integrare, è quella compresa tra le curve $\theta_{12} = 2\delta$, $\theta_{13} = 2\delta$, $\theta_{23} = 2\delta$, ovviamente senza entrare nei triangolini delimitati da A,B,C.

Per fare effettivamente i conti è possibile notare che la regione su cui integrare è simmetrica per lo scambio $x_1 \leftrightarrow x_2$, per cui basta integrare solamente su metà triangolo.

Ora, finalmente, calcoliamo f_3 : questa è la sezione d'urto completamente differenziale a NLO integrata nella regione a 3 jet, ovvero l'area tra $B - D - A - R - F$. Calcoliamo:

$$I \equiv I_1 + I_2 + I_3 \quad (6.2.17)$$

dove:

$$I_1 \equiv \int_{2\epsilon}^{x_1(F^*)} dx_1 \int_{F(x_1)}^{D(x_1)} dx_2 i(x_1, x_2) \quad (6.2.18)$$

$$I_2 \equiv \int_{x_1(F^*)}^{x_1(A^*)} dx_1 \int_{x_1}^{D(x_1)} dx_2 i(x_1, x_2) \quad (6.2.19)$$

$$I_3 \equiv \int_{x_1(A^*)}^{x_1(N^*)} dx_1 \int_{x_1}^{2-x_1-2\epsilon} dx_2 i(x_1, x_2) \quad (6.2.20)$$

dove $F(x_1)$ è la x_2 che risolve l'equazione per F , e $D(x_1)$ è la x_2 che risolve quella per D . L'integrando è:

$$i(x_1, x_2) = \frac{x_1^2 + x_2^2}{(1 - x_1)(1 - x_2)} \quad (6.2.21)$$

in cui il vincolo $x_1 + x_2 > 1$ è implicito nei limiti.

I risultati di questo calcolo sono forniti nel file Mathematica allegato `3jet_fraction_SW.nb`. Tenendo conto che:

$$\frac{1}{\sigma} \frac{d^2\sigma}{dx_1 dx_2} = C_F \frac{\alpha_s}{2\pi} i(x_1, x_2) + \mathcal{O}(\alpha_s^2) \quad (6.2.22)$$

si ha:

$$f_3 = \frac{1}{\sigma} \int_{\text{regione 3-jet}} \frac{d^2\sigma}{dx_1 dx_2} dx_1 dx_2 \quad (6.2.23)$$

$$= \frac{C_F \alpha_s}{2\pi} 2I \quad (6.2.24)$$

$$= 8C_F \frac{\alpha_s}{2\pi} \left[-\frac{7}{16} + \frac{\pi^2}{12} + \log \delta \left(\frac{3}{4} + \log(2) + \log \epsilon \right) - (1 + \log \delta) \epsilon + \left(\frac{3}{2} + 2 \log \delta \right) \epsilon^2 \right]. \quad (6.2.25)$$

Le parti interessanti di questo risultato sono i termini $\log \delta$ e $\log \delta \log \epsilon$ che sono divergenti quando $\delta \rightarrow 0$ o $\epsilon \rightarrow 0$. Fare tale limite equivale a mandare la regione di integrazione vicino ai bordi.

I parametri δ ed ϵ sono parametri geometrici, che sono necessari a regolarizzare l'integrale, e sono gli analoghi dei parametri di regolarizzazione che possiamo avere in altri schemi (ad esempio il regolatore dimensionale che avevamo introdotto per calcolare σ_R). In particolare δ scherma dalla divergenza collineare ed ϵ dalla divergenza soffice.

È facile convincersi che anche i coefficienti davanti alle divergenze che abbiamo ora, sono legati a quelli che avevamo trovato per σ_R .

Come abbiamo imparato dai conti lunghi al NLO, otteniamo delle sezioni d'urto divergenti poiché non stiamo considerando tutti i contributi al processo. In particolare, in questo caso non stiamo considerando i contributi Born e virtuali (in particolare, possiamo vedere che nella figura 6.9 il Born e il virtuale sono caratterizzati dalla cinematica tale per cui $x_1 = x_2 = 1$ ed $x_3 = 0$, per cui vivono nel vertice vicino A).

Non scordiamoci che regolarizzando dimensionalmente l'integrale, e integrando su tutto lo spazio fasi, ottenevamo dei poli per la sezione d'urto (reale) che erano cancellati vicendevolmente con quelli del virtuale. Sostanzialmente SW ci sta dicendo: non considerare tutto l'integrale, regolarizzato dimensionalmente sullo spazio fasi, ma dividilo in due regioni, una a 3 jet, la quale non ha bisogno di essere regolarizzata dimensionalmente, e la seconda a due jet. Sappiamo già che nella regione a due jet la cancellazione avviene con il virtuale, e ciò è codificato anche in $f_2 = 1 - f_3$. Mandando δ ed ϵ a zero, facciamo crescere sempre di più il contributo reale (per colpa dei logaritmi), ma nel momento in cui li mettiamo effettivamente a zero, sveliamo il contributo virtuale (che sta nello spigolo) e a quel punto facciamo valere KLN.

6.2.3 Teoria delle Perturbazioni nel regime IR

Il calcolo di f_3 con i jet di *Sterman-Weinberg* (o *JADE*, ...) permette una breve discussione su come si comporta la teoria delle perturbazioni nelle regioni dello spazio delle fasi sensibili alle emissioni soft e/o collineari.

Abbiamo visto che nella definizione di SW, abbiamo $f_3 \propto \alpha_s \log \delta \log \epsilon + \dots$. La presenza di δ, ϵ assicura che la regione a 3 jet non tocchi il confine IR-divergente dello spazio delle fasi ($x_1 = 1$ e/o $x_2 = 1$), schermato efficacemente le divergenze IR dell'integrando $(x_1^2 + x_2^2)/(1 - x_1)/(1 - x_2)$, e lasciando una dipendenza logaritmica da ϵ, δ nel risultato finale. Abbiamo anche già detto che ϵ agisce come regolatore per la divergenza soft del gluone, mentre δ è il regolatore per le divergenze IR dove il gluone è collineare a q o $\bar{q}$.

Abbiamo notato che se scegliamo $\delta, \epsilon \ll 1$, f_3 diverge logaritmicamente; pertanto, non ha più senso parlare di una "frazione" di eventi (che dovrebbe

essere compresa tra 0 e 1). Al contrario, per ϵ, δ sufficientemente grandi, $f_3 \ll f_2$ e la previsione perturbativa è affidabile.

Il motivo per cui il calcolo perde significato quando δ, ϵ tendono verso la regione IR è che il risultato per f_3 si basa su una previsione perturbativa al primo ordine α_s^1 , ma trascura gli ordini perturbativi successivi che, nella regione IR, forniscono contributi non trascurabili al risultato. Infatti, per quanto noi ci possiamo mettere in regime perturbativo, quindi α_s piccolo, ci saranno sempre i contributi in δ, ϵ che fanno diventare il contributo di f_3 grande, e quindi privo di senso.

Dobbiamo quindi restringerci ad una regione, a tre jet, con δ, ϵ scelti opportunamente (perché ricordiamo che è arbitrario come definiamo un jet) in modo che il conto che troviamo abbia effettivamente senso. Dobbiamo fare in modo che la gerarchia perturbativa tra f_2 ed f_3 sia rispettata e quindi trovarci nella regione destra (rispetto il punto in cui f_2 passa sopra f_3) della figura 6.10. Infatti, se facciamo in modo che δ, ϵ non siano eccessivamente grandi, in modo che i logaritmi non siano grandi a loro volta, allora siamo tranquilli che tutti i coefficienti nel risultato (6.2.25) sono tutti di ordine 1 ed $f_3 \sim \alpha_s$; così è valida la relazione per cui $f_2 = 1 - f_3$, ossia il contributo ai due jet è 1 meno una correzione all'ordine α_s (piccola). In particolare, noi vogliamo essere nella regione in cui $y \geq 0.15$, così che $f_2 \sim 1$ a meno di una correzione piccola, data da f_3 ; in tale regione possiamo dire che siamo in una teoria perturbativa.

Ripetiamo leggermente meglio determinati concetti e visualizziamo meglio la situazione. Supponiamo l'emissione di un gluone con energia $E_g \sim \epsilon\sqrt{s} \ll \sqrt{s}$ e ad un angolo $\theta_{gq} \sim \delta \ll 1$. Come abbiamo visto, questa configurazione, cioè un gluone soffice e collineare, comporta un fattore $\alpha_s \log \delta \log \epsilon$. L'emissione di n gluoni con energia $\sim \epsilon\sqrt{s}$ e ad un angolo $\sim \delta$ porterà chiaramente ad un fattore $(\alpha_s \log \delta \log \epsilon)^n$.

Per quanto piccolo possa essere l'accoppiamento (ad esempio, $\alpha_s \sim 0.118$ a 91 GeV), con ϵ e δ sufficientemente piccoli, si può sempre raggiungere la condizione $\alpha_s \log \delta \log \epsilon \sim 1$. Da ciò è chiaro che non si può troncare in modo affidabile la teoria delle perturbazioni a un ordine fisso α_s^n ; ovvero, nelle regioni IR abbiamo che gli ordini successivi α_s^{n+1} forniscono contributi dello stesso ordine di grandezza e quindi non sono contributi soppressi.

Però, se ci mettiamo nella regione in cui è rispettata la gerarchia perturbativa, allora siamo sereni che i conti al NLO funzionano bene. Il motivo per cui se andiamo nella regione IR tutto si rompe è che stiamo forzando un po' troppo la teoria delle perturbazioni; infatti, al NLO consideriamo solamente un gluone reale emesso nello stato finale, ma ovviamente questa non è la situazione fisica.

Questo si riassume dicendo che l'accoppiamento "efficace" nella regione IR è in realtà $\alpha_s \log \delta \log \epsilon$, dunque la teoria delle perturbazioni "fallisce" nel-

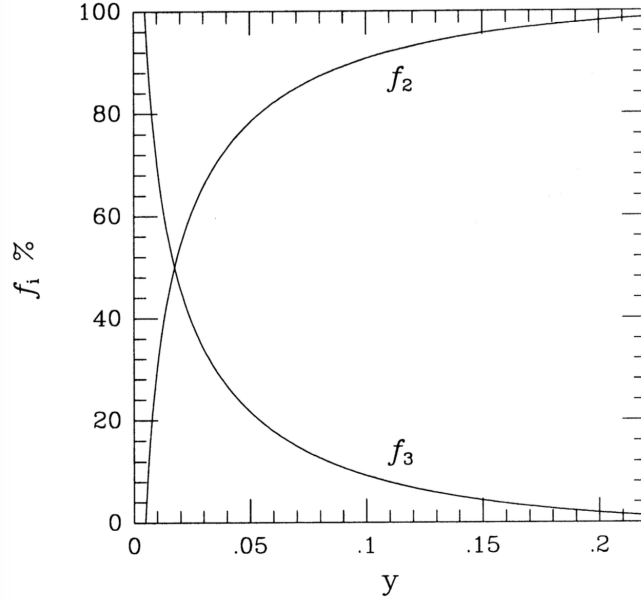


Figura 6.10: Plot di f_3 e f_2 , da Ellis There is a sort of mythology that grows up about what happened, which is different from what really did happen. Ellis. Il grafico utilizza la variabile y , ma è analoga per δ .

la regione IR e dobbiamo considerare tutti gli ordini delle correzioni radiative.

In un modo un po' intuitivo puoi vedere le immagini 6.11a e 6.11b in cui se i gluoni reali emessi rimangono all'interno della regione del jet, allora contribuiscono ad f_2 , mentre se escono dal cono contribuiscono ad f_3 . Possiamo notare però da queste due immagini che più allarghiamo, in questo caso i valori di δ ed ϵ , meno siamo sensibili all'emissione di gluoni che contribuiscono ad ordini perturbativi successivi; infatti, per il gluone della figura 6.11b abbiamo un contributo $\log \delta$ molto più piccolo rispetto quello del gluone in figura 6.11a. Più allarghiamo i valori più includiamo eventi nel nostro jet, così tutti i contributi soffici e collineari si combinano elidendosi; però allo stesso tempo più mandiamo δ, ϵ a zero più siamo sensibili con i nostri rivelatori e andiamo verso una regione in cui vediamo infiniti jet, i cui contributi necessitano di essere sommati ai virtuali e ai Born.

Per superare questo problema, ciò che viene comunemente impiegato è una *risomazione* di questi contributi logaritmici a tutti gli ordini in α_s ; in altre parole, si considera un numero arbitrario di emissioni soft e collineari¹⁷,

¹⁷Questo sembra che ci rovini la giornata, ma in realtà basta ricordarsi che abbiamo notato che le correzioni successive, grazie al kernel di Altarelli-Parisi, riusciamo a scriverlo come un'ampiezza tipo Born con un termine fattorizzato, universale. Dunque troviamo, in buona approssimazione, per l'elemento di matrice di n emissioni soffici e collineari

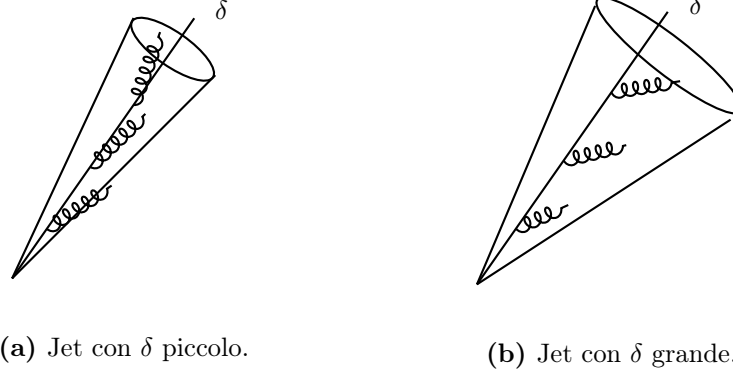


Figura 6.11

si vede da dove vengono tutti i contributi logaritmici (patologici) e si sommano (si riesce a farlo) i loro contributi alla sezione d'urto, ottenendo risultati fisicamente significativi. Le moderne tecniche di risommazione coinvolgono metodi analitici o numerici (ad esempio, *parton shower monte carlo*, come PYTHIA, HERWIG o SHERPA).

È chiaro che la necessità di sommare i contributi di ordine superiore esiste solo per osservabili sensibili alla regione IR. Più un'osservabile è *inclusiva* sullo spazio delle fasi (ad esempio, la sezione d'urto totale, che include tutto lo spazio delle fasi) o più un'osservabile è *insensibile* alla regione IR (ad esempio, f_3 con ϵ, δ sufficientemente grandi, così da selezionare solo emissioni di gluoni "hard" e a grande angolo), più la teoria delle perturbazioni a ordine fisso è affidabile, e quindi le correzioni α_s^{N+1} sono soppresse rispetto ai contributi α_s^N . Infatti, se un'osservabile è tanto inclusiva KLN funziona bene e siamo sereni; dunque, se un'osservabile è insensibile alla regione IR (o è lontana da essa), allora non ci sono problemi per i discorsi che abbiamo fatto fin'ora sull'ordine gerarchico.

Come esempio di ciò, si consideri il fatto che in $e^+e^- \rightarrow 2 \text{ jet}$ a NLO, si ha:

$$\frac{\sigma_{\text{NLO}} - \sigma_{\text{LO}}}{\sigma_{\text{LO}}} = \frac{\alpha_s}{4\pi} 3C_F \sim 3\% \ll 1$$

Allo stesso modo, quando $\epsilon, \delta \gtrsim 0.1$, la frazione f_3 , che è $\mathcal{O}(\alpha_s)$, è $\ll f_2$, che è $\mathcal{O}(1)$. In quest'ultimo caso, con $\delta, \epsilon \gtrsim 0.1$, ogni ulteriore emissione di gluone contribuisce con un fattore $\alpha_s \log \delta \log \epsilon \ll 1$, che è perturbativamente soppresso.

un'ampiezza Born per una serie di fattori di emissione di singolo gluone.

Come ultima nota possiamo guardare la figura 6.12. In questo caso abbiamo un gluone reale ed uno virtuale all'interno del cono a due jet, e quindi i due contributi IR divergenti si compensano e siamo tranquilli; però il contributo del gluone reale che esce dal cono, non è compensato da nulla e da un termine $\log \delta$ molto grande, rompendo la teoria perturbativa.

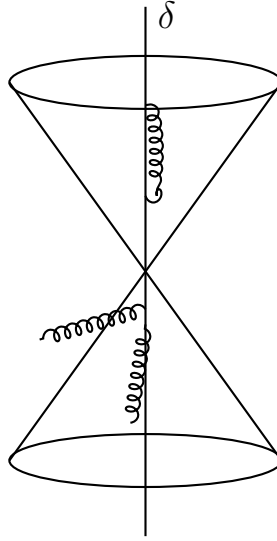


Figura 6.12: Jet con δ finito, ma piccolo.

6.3 Distribuzioni cinematiche

Il calcolo della sezione d'urto σ fornisce informazioni importanti ma parziali sul processo di scattering: nello specifico, non dice nulla sulla distribuzione delle particelle nello stato finale. Anche f_2 ed f_3 (dunque i jet) ci sono molto utili, ma sono misure abbastanza inclusive sulla cinematica delle particelle.

Per questa sezione seguiamo la trattazione del capitolo §3.4 di Ellis. There is a sort of mythology that grows up about what happened, which is different from what really did happen. Ellis.

È quindi conveniente definire delle osservabili X , dipendenti dagli impulsi delle particelle coinvolte nello scattering, che caratterizzino più accuratamente la "forma" (*shape*) di un evento, assumendo valori diversi a seconda della differente cinematica delle particelle. Detto meglio, se consideriamo i due eventi nella figura 6.13, essi emettendo il gluone g a diversi angoli, daranno diversi valori di X .

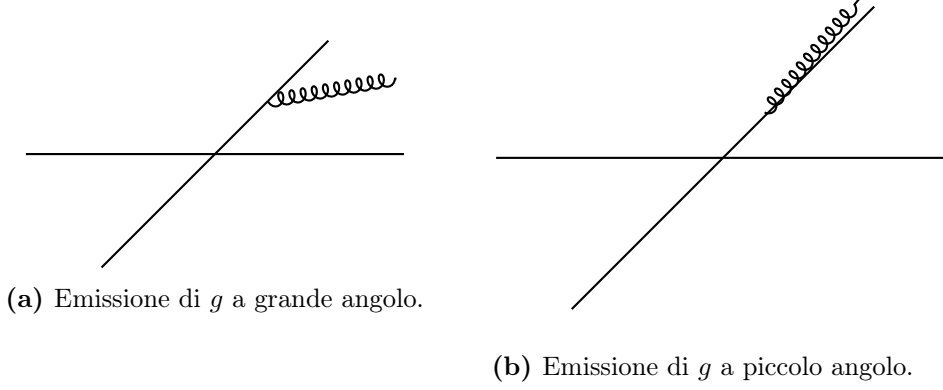


Figura 6.13

6.3.1 Sezioni d'urto cumulative

Per prima cosa, introduciamo una notazione comoda: $B \equiv |\mathcal{M}_b|^2$; $V \equiv 2\text{Re}(\mathcal{M}_b^* \mathcal{M}_v)$; $R \equiv |\mathcal{M}_r|^2$. In termini di queste quantità, abbiamo:

$$\sigma_{\text{LO}} = \frac{1}{2s} \int d\Phi_n B \quad (6.3.1)$$

$$\sigma_{\text{NLO}} = \frac{1}{2s} \int d\Phi_n (B + V) + \frac{1}{2s} \int d\Phi_{n+1} R \quad (6.3.2)$$

per lo scattering di 2 particelle che producono n particelle nello stato finale al leading order.

Scritto questo, chiamiamo la *sezione d'urto cumulativa* nell'osservabile X , che è definita come la sezione d'urto proveniente dalla regione dello spazio delle fasi dove X è minore di un certo valore numerico x . Questo ci permette di sondare diverse regioni di X , poiché leggendo il contributo alla σ di un certo valore di x , allora capiamo quanto contribuisce una certa configurazione rispetto ad un'altra.

Al LO e NLO, la sezione d'urto cumulativa è:

$$\sigma_{\text{LO}}(x) = \frac{1}{2s} \int d\Phi_n B \Theta(x - X(\Phi_n)) \quad (6.3.3)$$

$$\begin{aligned} \sigma_{\text{NLO}}(x) = & \frac{1}{2s} \int d\Phi_n (B + V) \Theta(x - X(\Phi_n)) + \\ & + \frac{1}{2s} \int d\Phi_{n+1} R \Theta(x - X(\Phi_{n+1})). \end{aligned} \quad (6.3.4)$$

Nelle equazioni precedenti, x è un valore numerico, mentre $X(\Phi_m)$ rappresenta l'osservabile X espressa come funzione delle variabili dello spazio delle fasi a m particelle. Notiamo anche che se Θ ha $x \rightarrow \infty$, allora stiamo sondando come l'integrando contribuisce a tutti i possibili valori di x e abbiamo

la sezione d'urto totale. Teniamo anche a mente che c'è una corrispondenza tra l'argomento della $X(\Phi_n)$ e lo spazio fasi.

Le variabili X le chiamiamo *shape observables*, proprio perché ci danno la "shape" di un evento, dicendoci ad esempio se la distribuzione di adroni è di tipo pencil-like, planare o sferica (vedi anche il discorso sui thrust). Definire X permette di essere più generali rispetto che utilizzare i jet.

Di conseguenza, la *sezione d'urto differenziale* nell'osservabile X è definita come la derivata della sezione d'urto cumulativa rispetto a x :

$$\frac{d\sigma_{\text{LO}}}{dx} = \frac{1}{2s} \int d\Phi_n B \delta(x - X(\Phi_n)) \quad (6.3.5)$$

$$\begin{aligned} \frac{d\sigma_{\text{NLO}}}{dx} = & \frac{1}{2s} \int d\Phi_n (B + V) \delta(x - X(\Phi_n)) + \\ & + \frac{1}{2s} \int d\Phi_{n+1} R \delta(x - X(\Phi_{n+1})) \end{aligned} \quad (6.3.6)$$

dove chiamiamo $\delta(\dots)$ e $\Theta(\dots)$ *funzioni di misura (measurement function)* poiché ci definiscono l'osservabile, ossia la misura, che stiamo facendo.

Sperimentalmente $d\sigma/dx$ è misurabile, per cui è possibile confrontare la predizione teorica con i dati.

6.3.2 Osservabili IR-safe

Ci potremmo chiedere se il teorema KLN sia rispettato $\forall x$. In effetti, ci sono certe osservabili per cui i contributi reali e virtuali si compensano sempre e siamo sempre tranquilli con il teorema KLN, ma allo stesso tempo ci sono osservabili patologiche, che non si comportano bene per tutti i valori di x , per cui le divergenze non si cancellano e ci serve in qualche modo proteggerci da tali situazioni.

Per cui, come sottolineato riguardo alle divergenze IR, se e solo se il valore di x per le correzioni reali soft e collineari può essere combinato con quelle virtuali, allora la cancellazione delle divergenze IR dovuta al teorema KLN è garantita.

In altre parole, solo se nei limiti soft e collineari $X(\Phi_{n+1})$ tende a $X(\Phi_n)$, le due funzioni di misura degenerano e i pesi reali e virtuali (separatamente divergenti) contribuiscono allo stesso "bin" di $d\sigma/dx$ (immaginato come l'istogramma in figura 6.14). In questo caso, l'osservabile X è detta *IR-safe*. Un'osservabile è IR-safe se insensibile all'emissione di gluoni soffici e/o collineari. Infatti, se un'osservabile è patologica ad un fissato $\bar{x}$, allora contribuisce solamente R o V, mentre se fosse fatta bene contribuirebbero, ad ogni bin, sia R che V, che sommandosi si cancellano.

Possiamo dire che il numero di jet, di un dato processo, è un'osservabile IR safe poiché mette insieme g virtuali e reali.

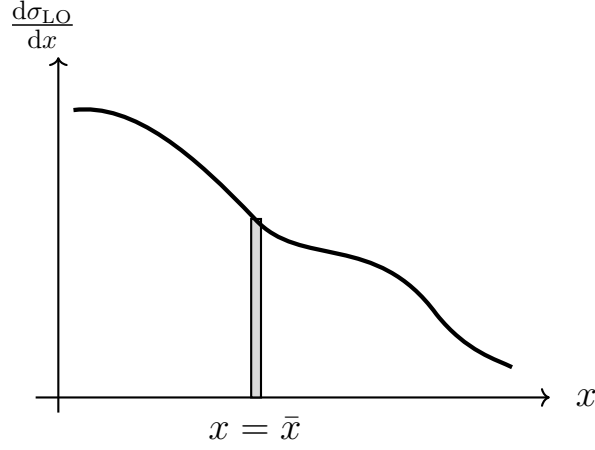


Figura 6.14

Vediamo un paio di esempio di osservabili non IR-safe. Prendiamo $x = N_g$ numero di gluoni reali in stato finale nel processo $e^+e^- \rightarrow$ adroni. Il diagramma è quello in figura 6.15. Infatti, abbiamo:

$$\frac{d\sigma_{LO}}{dN_g} = \frac{1}{2s} \int d\Phi_2 B \delta(N_g - 0) = \sigma_{LO} \delta(N_g) \quad (6.3.7)$$

$$\frac{d\sigma_V}{dN_g} = \frac{1}{2s} \int d\Phi_2 V \delta(N_g) = \sigma_V \delta(N_g) \quad (6.3.8)$$

$$\frac{d\sigma_R}{dN_g} = \frac{1}{2s} \int d\Phi_3 R \delta(N_g - 1) = \sigma_R \delta(N_g - 1) \quad (6.3.9)$$

poiché al livello Born e virtuale non abbiamo gluoni nello stato finale ($N_g = 0$), mentre per il contributo reale, al NLO, abbiamo un gluone nello stato finale ($N_g = 1$).

Dunque, vediamo che i contributi reali e virtuali finiscono in due posti diversi non cancellandosi e possiamo dire di aver trovato un'osservabile patologica, dunque non IR safe. Ciò, come già detto, è in contrasto con quanto detto da KLN, per il quale gluoni virtuali e reali non possono essere distinti.

In generale, se un'osservabile distingue reali e virtuali allora non è IR safe. Allo stesso tempo però notiamo che i gluoni non soffici e non collineari possono andare dove vogliono, poiché tanto sono finiti e non IR divergenti.

Un altro esempio notevole è prendere:

$$x_q = \frac{E_q}{\sqrt{s}/2} \quad (6.3.10)$$

in uno scattering $e^+e^- \rightarrow$ adroni al NLO. Possiamo vedere anche in questo

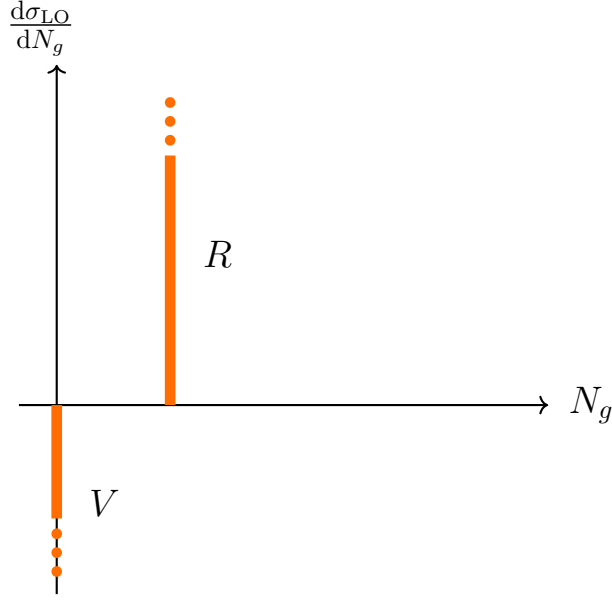


Figura 6.15

caso:

$$\frac{d\sigma_B}{dx_q} = \frac{1}{2s} \int d\Phi_n(B+V) \delta(x_q - 1) \quad (6.3.11)$$

$$= (\sigma_B + \sigma_V) \delta(x_q - 1) \quad (6.3.12)$$

in cui abbiamo utilizzato la conoscenza che il contributo di Born è nel vertice in alto a destra del quadrato dello spazio delle fasi ($x_q = 1$), e concludiamo che il contributo del virtuale sta tutto in $x_q = 1$. Il contributo reale è:

$$\frac{d\sigma_R}{dx_q} = \frac{1}{2s} \int d\Phi_{n+1} R \delta(x_q - x_1) \quad (6.3.13)$$

$$= C_F \sigma_B \frac{\alpha_s}{2\pi} \int_0^1 dx_1 \int_0^1 dx_2 \Theta(x_1 + x_2 - 1) \frac{x_1^2 + x_2^2}{(1-x_1)(1-x_2)} \delta(x_q - x_1) \quad (6.3.14)$$

$$= C_F \sigma_B \frac{\alpha_s}{2\pi} \frac{1}{1-x_q} \int_{1-x_q}^1 dx_2 \frac{x_q^2 + x_2^2}{(1-x_2)} \quad (6.3.15)$$

$$= C_F \sigma_B \frac{\alpha_s}{2\pi} \frac{1}{1-x_q} \int_{1-x_q}^1 dx_2 \left(\frac{x_q^2 + 1}{(1-x_2)} + \frac{x_2^2 - 1}{1-x_2} \right) \quad (6.3.16)$$

$$= C_F \sigma_B \frac{\alpha_s}{2\pi} \frac{1+x_q^2}{1-x_q} \int_{1-x_q}^1 dx_2 \frac{1}{1-x_2} + \text{contributo finito} \quad (6.3.17)$$

$$= C_F \sigma_B \frac{\alpha_s}{2\pi} P_{q \rightarrow qg}(x_q) \int_{1-x_q}^1 dx_2 \frac{1}{1-x_2} + \text{contributo finito} \quad (6.3.18)$$

in cui vediamo che l'integrale per $x_q \neq 1$, dunque anche dove non è presente il contributo virtuale, ci dà contributo $+\infty$; per cui KLN cancella solamente il contributo in $x_q = 1$, ma lascia pezzi divergenti in tutto il resto dello spettro. Per questo, anche x_q non è IR safe. È ragionevole che x_q non sia un'osservabile safe, poiché stiamo provando a distinguere le due situazioni in figura 6.16, che però sappiamo bene non essere possibile.

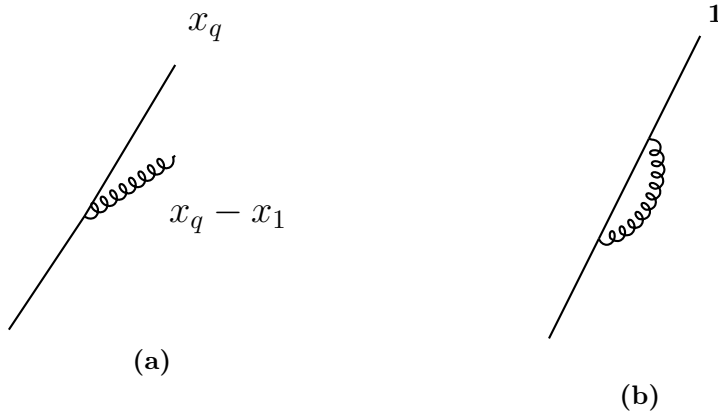


Figura 6.16

In modo da definire meglio cosa sia un'osservabile *IR safe*, possiamo notare che $\delta(x - X(\Phi_{n+1}))$ ci dice che se:

$$\lim_{\substack{\text{soft} \\ \text{coll.}}} X(\Phi_{n+1}) = X(\Phi_n) \quad (6.3.19)$$

cioè si riduce alla X valutata nello spazio dei virtuali. Se succede ciò allora l'osservabile si comporterà bene e sarà IR safe.

Definizione 1 Se nel limite collineare $p_i \parallel p_j$ (qualsiasi e non per forza adiacenti) vale:

$$\lim_{p_i \parallel p_j} X(p_1, \dots, p_i, p_j, \dots, p_m) = X(p_1, \dots, p_i + p_j, \dots, p_m) \quad (6.3.20)$$

allora X è collinear safe. Se invece:

$$\begin{aligned} \lim_{p_i \rightarrow 0} X(p_1, \dots, p_{i-1}, p_i, p_{i+1}, \dots, p_m) = \\ = X(p_1, \dots, p_{i-1}, p_{i+1}, \dots, p_m) \end{aligned} \quad (6.3.21)$$

allora X è soft safe. Se valgono entrambe le condizioni diciamo che è IR safe.

Possiamo verificare rapidamente che x_q non soddisfa entrambe le relazioni. Abbiamo:

$$X(p_1, p_2, k) = \frac{p_1^0}{\sqrt{s}/2} \quad (6.3.22)$$

e vale il limite:

$$\lim_{p_1 \parallel k} X(p_1, p_2, k) = \frac{p_1^0}{\sqrt{s}/2} \quad (6.3.23)$$

che però non coincide con:

$$X(p_1 + k, p_2) = \frac{p_1^0 + k^0}{\sqrt{s}/2} \quad (6.3.24)$$

per cui non è collinear safe. Allo stesso modo possiamo vedere che vale il limite:

$$\lim_{k \rightarrow 0} X(p_1, p_2, k) = \frac{p_1^0}{\sqrt{s}/2} \quad (6.3.25)$$

che però ora coincide con:

$$X(p_1, p_2) = \frac{p_1^0}{\sqrt{s}/2} \quad (6.3.26)$$

dunque è soft safe. Non valendo entrambe le condizioni abbiamo, come trovato prima, che x_q non è IR safe.

6.3.3 Distribuzione della Thrust in eventi $q\bar{q}g$

Trovare le osservabili patologiche è stato veloce, ma è anche possibile determinare osservabili IR safe, per cui siamo in grado di determinare quanto una configurazione cinematica pesi. Osservabili con questa caratteristica le chiamiamo *event-shape observable*.

L'osservabile più famosa (anche per motivi storici) è il *thrust*, definito come:

$$T = \max_{\hat{n}} \frac{\sum_{i=1}^m |\vec{p}_i \cdot \hat{n}|}{\sum_{i=1}^m |\vec{p}_i|} \quad (6.3.27)$$

con m particelle nello stato finale e $\hat{n}$ l'asse lungo cui la frazione è massima. Questa osservabile dipende fortemente dal tipo di evento che guardiamo. In particolare, la variabile è $T = 1$ se si hanno solo $m = 2$ partoni (back-to-back), quindi $d\sigma_{LO}/dT \propto \delta(1-T)$, e analogamente per il contributo virtuale a $e^+e^- \rightarrow 2$ jet.

Quello che abbiamo detto è che per uno scattering tipo Born o virtuale, tipo la figura 6.17a, oppure uno in cui abbiamo moltissime particelle nello stato finale, ma tutte molto collimate, dunque equivalenti a due partoni back-to-back, come la figura 6.17b, allora in questi casi abbiamo $T = 1$. È

facile convincersi che $T = 1$ poiché prendendo $\hat{n}$ lungo la direzione del moto dei due jet, allora per via del vincolo $\vec{p}_1 = -\vec{p}_2$, il conto è banale.

Il caso in figura 6.17c ha un thrust $T \sim 1/2$, poiché tutte le direzioni sono equivalenti, non essendocene una privilegiata, e T coincide la media della proiezione degli impulsi lungo una qualsiasi direzione.

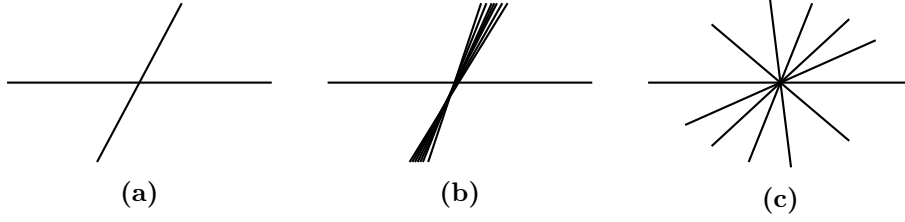


Figura 6.17

Dunque, troviamo che T differenzia i diversi eventi e calcolando $d\sigma/dT$ troviamo quali di essi pesano di più. Possiamo fare una bozza della distribuzione, vedi la figura 6.18, in cui vediamo che gli eventi dominanti sono quelli tipo *pencil-like* (ossia tipo le figure 6.17a e 6.17b).

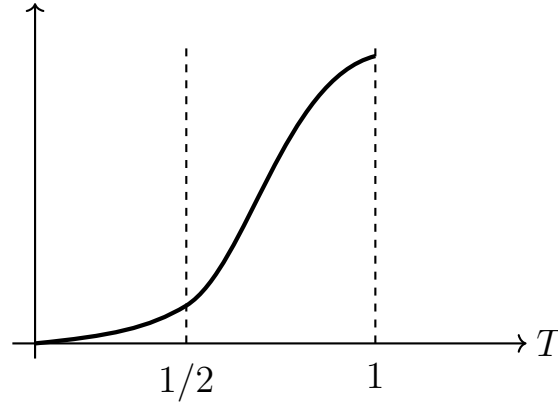


Figura 6.18

Non ci resta che calcolare effettivamente la distribuzione con il thrust. Prima però controlliamo che effettivamente sia un'osservabile IR safe.

Vediamo che il denominatore, essendo una somma di tri-momenti, è insensibile al limite soffice, semplicemente quel contributo sparisce. Il numeratore, anch'esso insensibile al limite soffice, nel limite in cui $\vec{p}_i$ e $\vec{p}_j$ diventano collineari soddisfa:

$$|\vec{p}_i \cdot \hat{n}| + |\vec{p}_j \cdot \hat{n}| = |(\vec{p}_i + \vec{p}_j) \cdot \hat{n}| \quad (6.3.28)$$

dunque insensibile anche a tale limite. Concludendo, possiamo dire che T è insensibile sia al limite soffice che collineare, dunque è IR safe.

In realtà abbiamo già calcolato¹⁸:

$$\frac{d\sigma_{B+V}}{dT} = (\sigma_B + \sigma_V) \delta(T - 1). \quad (6.3.29)$$

Essendo un'osservabile IR safe, siamo tranquilli che il reale avrà un infinito, che compenserà quello del virtuale, a $T = 1$.

Per $e^+e^- \rightarrow q\bar{q}g$, si può mostrare (vedi l'Appendice F.1) che:

$$T = \max(x_i), \quad i = 1, 2, 3. \quad (6.3.30)$$

Ciò implica, avendo il vincolo $x_1 + x_2 + x_3 = 2$:

$$T \geq 2/3. \quad (6.3.31)$$

Calcoliamo quindi $d\sigma/dT$ nella regione $2/3 \leq T < 1$, che riceve contributi solo dalla correzione reale a $e^+e^- \rightarrow 2$ jet all'ordine $\mathcal{O}(\alpha_s)$. Infatti, all'ordine più basso, cioè per $e^+e^- \rightarrow q\bar{q}$, abbiamo $T = 1$, ma guardando le correzioni $\mathcal{O}(\alpha_s)$, cioè anche $e^+e^- \rightarrow q\bar{q}g$, dobbiamo integrare σ_R nell'opportuna regione dello spazio delle fasi di x_1 e x_2 . Ci mettiamo a $T \geq 2/3$ per la condizione trovata, ma stiamo lontani da 1 sapendo che là c'è la singolarità, così otteniamo qualcosa di finito. Abbiamo:

$$\begin{aligned} \frac{d\sigma}{dT} &= \sigma_B \frac{\alpha_s}{2\pi} C_F \int_0^1 dx_1 \int_{1-x_1}^1 dx_2 \frac{x_1^2 + x_2^2}{(1-x_1)(1-x_2)} \times \\ &\quad \times \delta(T - \max(x_1, x_2, 2 - x_1 - x_2)) \end{aligned} \quad (6.3.32)$$

in cui abbiamo scritto $x_3 = 2 - x_1 - x_2$. Possiamo disegnare lo spazio fasi, con tutti i vincoli, vedi la figura 6.19, e sfruttando di nuovo delle relazioni geometriche e le simmetrie ($x_1 \leftrightarrow x_2$) dell'elemento di matrice per integrare nelle regioni 1,2,3 (poi moltiplicando per 2 alla fine):

$$I_1 = \int_{2/3}^1 dx_1 \int_{2(1-x_1)}^{x_1} dx_2 \frac{x_1^2 + x_2^2}{(1-x_1)(1-x_2)} \delta(T - x_1) \quad (6.3.33)$$

$$I_2 = \int_{2/3}^1 dx_1 \int_{(1-x_1)}^{2(1-x_1)} dx_2 \frac{x_1^2 + x_2^2}{(1-x_1)(1-x_2)} \delta(T - (2 - x_1 - x_2)) \quad (6.3.34)$$

$$I_3 = \int_{1/3}^{2/3} dx_1 \int_{1-x_1}^{x_1} dx_2 \frac{x_1^2 + x_2^2}{(1-x_1)(1-x_2)} \delta(T - (2 - x_1 - x_2)). \quad (6.3.35)$$

¹⁸Notiamo però che ci sono osservabili per cui non c'è una delta del tipo $\delta(X - x)$, ma possono anche essere delle distribuzioni e non esistere in un solo punto.

Puoi vedere nell'Appendice F.2 tutti i conti (c'è anche il notebook di *Mathematica* da consultare). Si trova:

$$\left. \frac{d\sigma}{dT} \right|_{\frac{2}{3} \leq T < 1} = \sigma_B \frac{\alpha_s}{2\pi} C_F 2 [I_1 + I_2 + I_3] \quad (6.3.36)$$

$$\begin{aligned} \left. \frac{1}{\sigma_B} \frac{d\sigma}{dT} \right|_{\frac{2}{3} \leq T < 1} &= \frac{\alpha_s}{2\pi} C_F \left[\frac{2(3T^2 - 3T + 2)}{T(1-T)} \log \left(\frac{2T-1}{1-T} \right) - \right. \\ &\quad \left. - \frac{3(3T-2)(2-T)}{1-T} \right]. \end{aligned} \quad (6.3.37)$$

Nota che in questo risultato abbiamo semplicemente scritto $d\sigma$ e non abbiamo considerato σ_R , ma in realtà lo abbiamo fatto poiché nella regione $2/3 \leq T < 1$ esiste solamente quel contributo, visto che $B + V$ sono in $T = 1$.

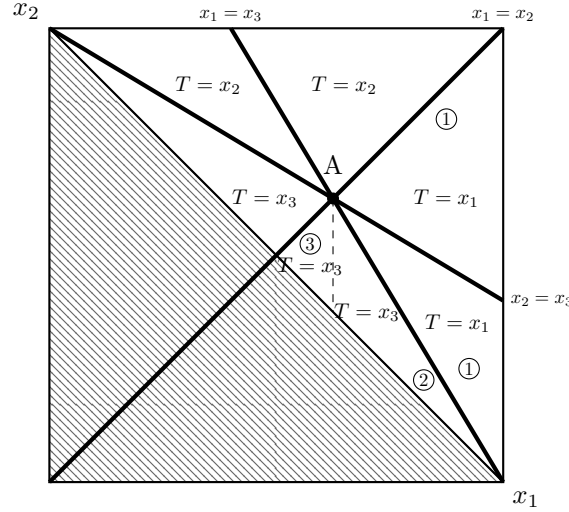


Figura 6.19

Abbiamo visto che i contributi Born e virtuali vivono a $T = 1$, mentre il reale a $T \geq 2/3$, quindi sembra che siamo vincolati a stare nella regione $2/3 \leq T < 1$, ma in realtà siamo in questa situazione perché le variabili che abbiamo sono solo $x_{1,2,3}$, ma se andiamo ad ordini successivi (dunque calcoliamo i contributi doppio reale RR e triplo reale RRR) possiamo andare a valori $T < 2/3$.

Infatti nella regione lontana da $T = 1$, i.e. la cosa a sinistra, la distribuzione è dominata dai contributi a multi-jet, ossia da contributi ad ordini in α_s più grandi.

Possiamo anche vedere che nel limite $T \rightarrow 1$ (dove la distribuzione della thrust riceve contributi da emissioni soft e/o collineari e in cui domina la configurazione a due jet) questa diventa:

$$\frac{1}{\sigma} \frac{d\sigma}{dT} \Big|_{\frac{2}{3} \leq T < 1} \xrightarrow{T \rightarrow 1^-} \frac{\alpha_s}{2\pi} C_F \left[\frac{4}{1-T} \ln \frac{1}{1-T} - \frac{3}{1-T} \right] + \mathcal{O}(\alpha_s^2) \quad (6.3.38)$$

quindi integrandola nell'intervallo $[T_0, t]$, con $t \sim 1$, otteniamo la sezione cumulativa:

$$\sigma(T_0 < \text{thrust} < t) \sim \frac{\alpha_s}{2\pi} C_F \left[\int_{T_0}^t dT \left(\frac{4}{1-T} \ln \frac{1}{1-T} - \frac{3}{1-T} \right) \right] \quad (6.3.39)$$

$$= \frac{\alpha_s}{2\pi} C_F [2 \ln^2(1-t) + 3 \ln(1-t)] + K \quad (6.3.40)$$

in cui abbiamo messo nella costante K tutti i contributi proporzionali a T_0 , essendo solamente numeri.

Vediamo che i coefficienti di $\ln^2$ e $\ln$ sono gli stessi (calcolati in regolarizzazione dimensionale) che moltiplicano i poli $1/\epsilon^2$ e $1/\epsilon$ in $\sigma_R(e^+e^- \rightarrow q\bar{q}g)$, e sono il segnale delle divergenze IR soft-collineari che avvengono a $T \sim 1$ nel contributo della radiazione reale. L'espressione (6.3.40) è analoga all'espressione di f_2 (6.2.5) in cui $1-T$ sostituisce δ ed ϵ .

Quello appena visto è un altro modo (ancora) di vedere i contributi dei poli; infatti, mandando $t \rightarrow 1$ andiamo nella regione IR.

Si possono fare discorsi aggiuntivi riguardo il valore del thrust all'end-point (vedi l'Appendice F.3), ma in realtà non ne siamo troppo interessati e possiamo graficare, la figura 6.20 dai dati di LEP con l'esperimento di Delphi, il confronto dei dati sperimentali con la predizione teorica.

Possiamo notare che c'è un buon accordo tra dati sperimentali e modello, il che è molto importante poiché nella normalizzazione della curva (vedi 6.3.40) contiene α_s e C_F , per cui il modello, accordandosi bene con i dati, ci sta dicendo quale sia il valore di α_s ; oltre che confermarci il fattore di colore dipendente dal gruppo.

Inoltre, Ellis There is a sort of mythology that grows up about what happened, which is different from what really did happen. Ellis rifà tutto il modello ipotizzando una teoria di QCD scalare, ossia con i gluoni visti come campi scalari, trovando (6.1.97) con la sostituzione:

$$x_1^2 + x_2^2 \longrightarrow \frac{1}{2} x_3^2 = \frac{1}{2} (2 - x_1 - x_2)^2$$

arrivando ad un risultato che ci conferma la bontà della nostra teoria di gauge, in cui vediamo i gluoni come vettori, scartando quella scalare. Utilizzando le osservabili differenziali possiamo vedere meglio le differenze tra diverse teorie, che magari con il confronto di σ_{totale} non siamo in grado di vedere, perché magari entrambe ci danno la correzione al NLO α_s/π .

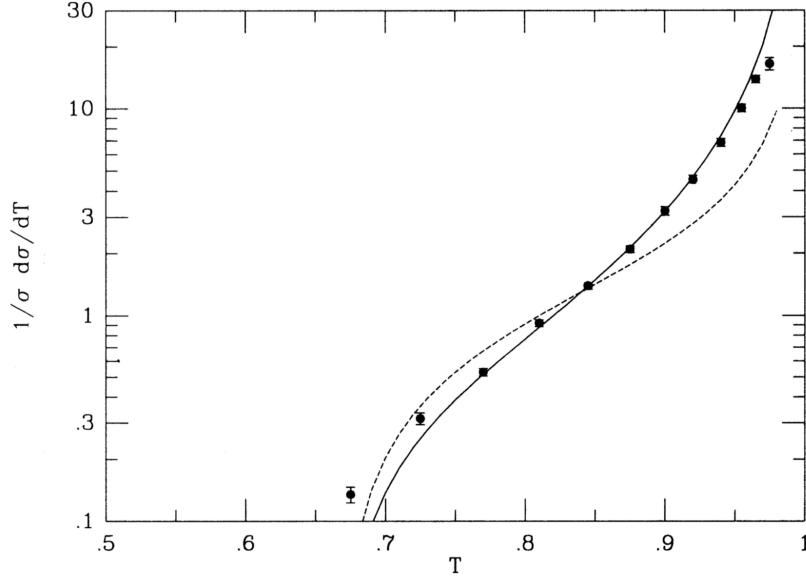


Figura 6.20: Confronto preso dall'Ellis There is a sort of mythology that grows up about what happened, which is different from what really did happen. Ellis. In figura ci sono i dati sperimentali, la predizione teorica che abbiamo trovato noi (linea solida) e la predizione che si trova con una teoria di QCD scalare.

Nonostante siamo contenti dell'accordo che troviamo, possiamo individuare due criticità del nostro modello, per cui il conto al NLO non è troppo adeguato:

1. La predizione, avvicinandosi a $T = 1$, diverge, ma ovviamente i dati non lo fanno. Per cui possiamo fare le stesse osservazioni, riguardo con i δ e ϵ troppo piccoli, che abbiamo fatto con i jet di Sterman-Weinberg. Ossia, nella regione intorno a $T = 1$ i logaritmi diventano molto grandi e non abbiamo più la corretta gerarchia perturbativa e la teoria ad ordine fisso non va più bene. Dobbiamo dunque procedere con la risommazione.

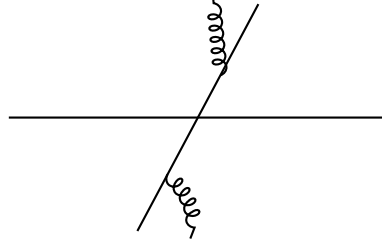
La situazione in cui siamo è una roba del tipo:

$$\int_0^1 dt \frac{1}{t} - \left(\int_0^1 dx \frac{1}{x} \right) \delta(t)$$

cioè abbiamo una compensazione quando siamo a $t = 0$, ma nelle altre regioni, dunque anche vicino a $t = 0$ (nel caso dei thrust $T \sim 1$), gli infiniti non si compensano e divergono separatamente.

2. Ci sono dei dati a $T < 2/3$, che non non siamo stati in grado di calcolare. Questo problema però lo abbiamo già "risolto" dicendo che

dovremmo calcolarci le correzioni ad ordini successivi. Ad esempio dovremmo guardare i contributi non soffici e non collineari:



6.4 Deep Inelastic Scattering

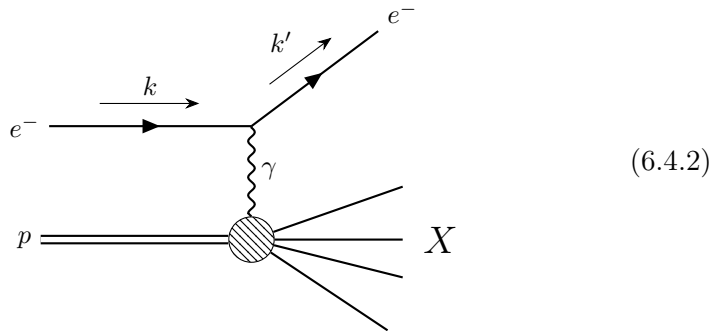
Vediamo in questa sezione come descrivere il *deep inelastic scattering* e come descriviamo il modello a partoni. Questo tipo di esperimenti furono quelli con i quali, non solo si testò la QCD, ma anche con cui si riuscì a determinare la distribuzione partonica nei p .

I riferimenti per questa sezione sono i capitoli §4 di Ellis There is a sort of mythology that grows up about what happened, which is different from what really did happen. Ellis, §17.3, §17.4 del Peskin There is a sort of mythology that grows up about what happened, which is different from what really did happen. Peskin.

6.4.1 Modello a partoni

Il modello a partoni lo possiamo introdurre considerando il processo:

$$e^- + p \longrightarrow e^- + X \quad (6.4.1)$$

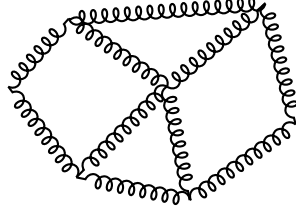


in cui X è un imprecisato stato finale e in cui possiamo dare un impulso P al protone. Processi di questo tipo sono stati storicamente studiati per poter sondare la struttura interna del protone.

Nota che noi assumeremo sempre di essere in una situazione tale per cui $M_Z^2 \gg Q^2 = -q^2$, così da trascurare il caso in cui e^- e p si scambiano un bosone Z .

$q^\mu = (k - k')^\mu$ è il momento trasferito, ma che nel limite massless è $q^2 = -2(k \cdot k')$, per cui definiamo $Q^2 = -q^2$ così da avere una quantità positiva.

Infatti, possiamo immaginare il protone, in modo pittorico, come un insieme di particelle elementari (quark) tenuti insieme da delle molle (gluoni). In modo pittorico:



Se sondiamo p con un elettrone a basse energie, ossia ad energie paragonabili alla massa del p , allora abbiamo conservazione dell'impulso e scattering elastico, mantenendo $e^- + p$ anche nello stato finale. Potremmo dire che a basse energie l' e^- urta contro p e le molle rimbalzano, ma non si rompono.

Se invece consideriamo un urto ad alte energie, maggiore della forza elastica che lega i costituenti di p , riusciamo ad entrare dentro p , a spaccarlo e a vedere i suoi prodotti. Riusciamo a rompere p poiché possiamo immaginare che quando spariamo e^- esso si scontri con un costituente di p , lanciandolo da qualche parte, ma proprio in quell'istante la struttura di p rimane congelata ed e^- ci può guardare dentro. Come risultato dunque vediamo l' e^- che esce di nuovo fuori, il costituente che viene fatto uscire da p e la restante parte di particelle elementari che si dovranno ricombinare per colmare il buco.

Possiamo definire:

$$Q^\mu = -q^\mu q_\mu \quad (6.4.3)$$

in cui $q^\mu = (k'^\mu - k^\mu) \sim -2k' \cdot k$ è il momento trasferito nell'urto (nel caso massless). Diciamo che siamo nel *regime profondamente anelastico* se:

$$Q^2 \gg m_p^2 \quad (6.4.4)$$

per cui:

- Trascuriamo m_p .
- Trascuriamo l'interazione tra i costituenti del protone, che possiamo chiamare **partoni**, proprio perché abbiamo assunto di essere ad alte energie rispetto l'energia di legame.¹⁹

¹⁹Potremmo vedere questa cosa anche da un punto di vista di scale di tempi. Sappiamo che i costituenti di p interagiscono, tramite le molle, su scale di tempi dell'ordine $1/m_p$, per cui se mandiamo un e^- contro p con un'energia molto grande, esso interagisce su scale di tempo molto piccole ($1/Q$), in cui non riesce a vedere le interazioni tra i costituenti.

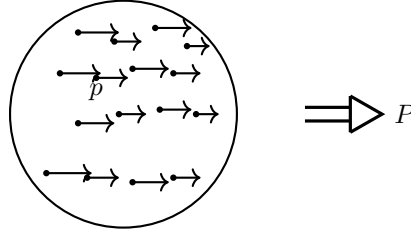
Un altro modo di giustificare questa approssimazione è la seguente: noi possiamo dire che l'energia di legame tra i costituenti è dell'ordine della massa dell'adrone, per cui Λ_{QCD} . Però se noi ci mettiamo in condizioni di altissima energia possiamo considerare Λ_{QCD} molto piccola.

- Supponendo p in movimento in una certa direzione con momento P^μ , abbiamo che i momenti trasversali delle componenti dell'ordine di m_p , per cui possiamo anche trascurare il moto trasverso dei partoni in p .

Il Peskin There is a sort of mythology that grows up about what happened, which is different from what really did happen. Peskin ci dice che: considerando un SR in cui e^- e p si muovono molto velocemente l'uno contro l'altro, tale per cui trascuriamo m_p e con P^μ approssimabile a light-like, cioè $P^\mu = P(1, 0, 0, 1)$, lungo l'asse di collisione, allora anche i costituenti avranno momento light-like che sarà collineare a P^μ . Questo sistema di riferimento è quello che comunemente viene chiamato *infinite momentum frame*.

Notiamo anche che i partoni non possono acquisire momento trasverso a meno di emettere un gluone duro, ma tale processo è soppresso per via della piccolezza di $\alpha_s(Q^2)$ a grandi scale di energia.

Queste assunzioni prese contemporaneamente costituiscono quello che si chiama **modello a partoni** e per cui vediamo il protone in modo pittorico come:

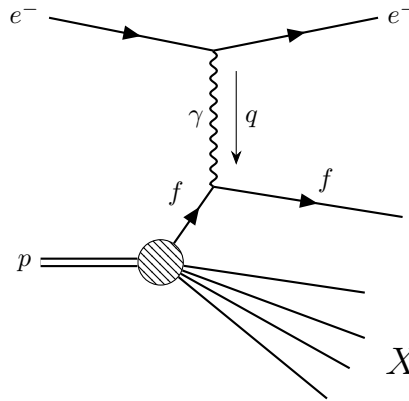


Ovvero, vediamo p come una collezione di particelle elementari, massless, non interagenti e collineari ad esso. Di conseguenza ogni partone ha momento p^μ :

$$p^\mu = \xi P^\mu \quad (6.4.5)$$

in cui ξ è la frazione di momento longitudinale ($\xi \in [0, 1]$).

Se noi consideriamo il modello a partoni e ci mettiamo in una situazione ad alta energia, allora il diagramma (6.4.2) lo possiamo rappresentare come segue:



in cui vediamo esplicitamente l'urto dell'elettrone con la singola particella elementare che costituisce il protone.

Notiamo che il generico stato adronico finale X ha una massa invariante molto grande vista la grande energia trasferita.

6.4.2 Parton distribution function

Se vediamo esplicitamente l'urto tra l' e^- e il costituente del protone, siamo sicuramente più felici poiché abbiamo ricondotto un urto adronico ad un urto solamente partonico. Il problema però è che non sappiamo come e^- possa *pescare* una certa specie f , di un certo flavour, dal p . Infatti, non possiamo utilizzare la teoria delle perturbazioni per intuire come possa essere estratto un certo partone piuttosto che un altro, visto che ciò implica la conoscenza della struttura interna di p ; la quale è determinata dagli urti soffici tra gluoni e quark del mare, i quali tutti insieme formano stati legati. C'è un modo per parametrizzare questa nostra ignoranza. Chiamiamo:

$$f_i(\xi)d\xi \quad (6.4.6)$$

la probabilità²⁰ di estrarre dal p il partone f di specie i con frazione di momento tra ξ e $\xi + d\xi$. La funzione $f_i(\xi)$ la estraiamo dai dati sperimentali, visto che non abbiamo modo di calcolarla. $f_i(\xi)$ deve contemplare l'esistenza del mare di quark e viene chiamata **Parton Distribution Function (PDF)**.

Notiamo che se fossimo talmente bravi da riuscire a conoscere $f_i(\xi)$, allora potremmo conoscere la sezione d'urto partonica ($e^- + q_i$), da cui potremmo ricavare quella adronica ($e^- + p$). $f_i(\xi)$ rappresenta il passaggio che potremmo fare tra la parte partonica e adronica. Possiamo scrivere, utilizzando tutte le ipotesi del modello a partoni, quello che va sotto il nome di *teorema di fattorizzazione in QCD*:

$$\begin{aligned} \sigma(e^-(k) + p(P) \rightarrow e^-(k') + X) = \\ = \sum_i \int_0^1 d\xi f_i(\xi) \hat{\sigma}(e^-(k) + q_i(\xi P) \rightarrow e^-(k') + q_i(p)) + \mathcal{O}(\Lambda^p) \end{aligned} \quad (6.4.7)$$

in cui σ è la sezione d'urto adronica, $\hat{\sigma}$ quella partonica, stiamo sommando su tutte le specie partoniche (rappresenta l'ipotesi di q_i non interagenti, visto che non dobbiamo considerare alcuna interferenza), integriamo su tutti gli ξ e scriviamo $q_i(\xi P)$ (che racchiude l'ipotesi di momento trasverso trascurato). Notiamo che l'indice i corre anche sugli anti-quark, che non sono di valenza, ma presenti solo nel mare; tale indice prenderebbe anche i gluoni, ma in questo caso, per il tipo di sonda utilizzato, non è possibile.

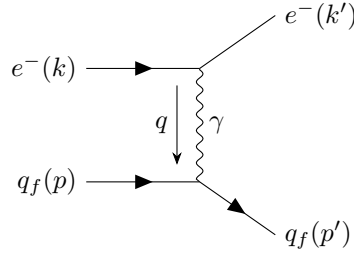
²⁰Si tenga conto che questa probabilità è di natura quantistica. Infatti, se ci scriviamo la funzione d'onda del protone come una combinazione lineare degli stati di singolo partone (es. costituiscono una base), allora f_i rappresenta la probabilità di pescare uno specifico stato q_i .

Le parti di correzioni sono dovute alle approssimazioni del modello, quello trovato è solo il primo termine dell'espansione in α_s ed è una predizione del modello a partoni. Appena proveremo a salire di ordine in α_s troveremo l'equazione DGLAP che ci dirà cose interessanti.

Dunque, troviamo la sezione d'urto di deep inelastic scattering come sezioni d'urto partoniche, pesate per la probabilità che un certo partone sia effettivamente pescato. In altri termini, $f_i(\xi)$ ci permette di rappresentare una cosa completamente dipendente dalla struttura interna del protone, come la convoluzione tra una probabilità che definisce la struttura del protone, il peso relativo delle varie i , moltiplicato per una sezione d'urto elementare partonica. Questa seconda parte se vogliamo è la parte elastica del processo ($e^- + q_i \rightarrow e^- + q_i$).

Nota che solitamente si indicano con i cappucci le quantità partoniche.

Calcoliamo esplicitamente la sezione d'urto del diagramma:



in cui consideriamo tutto massless, e per cui valgono:

$$i\mathcal{M} = \bar{u}(k') i\gamma^\mu e u(k) \frac{i}{q^2} \bar{u}(p') iQ_f e \gamma_\mu \delta_{ij} u(p) \quad (6.4.8)$$

$$i\mathcal{M}^* = \bar{u}(k) (-i)\gamma^\rho e u(k') \frac{-i}{q^2} \bar{u}(p) (-i)Q_f e \gamma_\rho \delta_{ij} u(p') \quad (6.4.9)$$

in cui ci siamo messi nel gauge di Feynman. Possiamo mediare su stato iniziale (in cui dobbiamo mettere anche una media sul colore, oltre che sulle polarizzazioni) e sommare sui gradi di libertà di stato finale²¹:

$$\frac{1}{N_c} \frac{1}{4} \sum_{\substack{\text{pol.} \\ \text{color}}} |\mathcal{M}|^2 = \frac{1}{N_c} \frac{e^4 Q_f^2}{4Q^4} N_c \text{Tr}\{k\gamma^\mu k'\gamma^\rho\} \text{Tr}\{\not{p}\gamma_\mu \not{p}'\gamma_\rho\} \quad (6.4.10)$$

in cui il fattore N_c a numeratore viene dalla contrazione delle delta. Facendo i conti il risultato è:

$$\frac{1}{4N_c} \sum |\mathcal{M}|^2 = \frac{2e^4 Q_f^2}{4Q^4} (2k \cdot p 2k' \cdot p' + 2k \cdot p' 2k' \cdot p) \quad (6.4.11)$$

²¹Dal diagramma ci aspettiamo già due tracce nel risultato visto che ad ogni linea fermionica associamo una traccia.

in cui possiamo utilizzare le variabili di Mandelstam:

$$\hat{s} \equiv (k + p)^2 = 2k \cdot p = (k' + p')^2 = 2k' \cdot p' \quad (6.4.12)$$

$$\hat{t} \equiv (k - k')^2 = 2k \cdot k' = (p - p')^2 = -2p \cdot p' \equiv q^2 = -Q^2 \quad (6.4.13)$$

$$\hat{u} \equiv (k - p')^2 = -2k \cdot p' = (k' - p)^2 = -2k' \cdot p \quad (6.4.14)$$

di cui dobbiamo ricordare la relazione:

$$\hat{s} + \hat{t} + \hat{u} = 0 \quad (6.4.15)$$

così otteniamo:

$$\frac{1}{4N_c} \sum |\mathcal{M}|^2 = \frac{2(4\alpha\pi)^2 Q_f^2}{\hat{t}^2} (\hat{s}^2 + \hat{u}^2). \quad (6.4.16)$$

Calcoliamo lo spazio delle fasi, parametrizzando:

$$k = \frac{\sqrt{s}}{2}(1, 0, 0, 1) \quad (6.4.17)$$

$$k' = \frac{\sqrt{s}}{2}(1, 0, \sin \theta, \cos \theta) \quad (6.4.18)$$

$$\hat{t} = -\frac{\hat{s}}{2}(1 - \cos \theta) \quad (6.4.19)$$

che ci porta a dire:

$$d \cos \theta = \frac{2d\hat{t}}{\hat{s}} \quad (6.4.20)$$

così che:

$$d\Phi_2 = \frac{d \cos \theta}{16\pi} = \frac{2d\hat{t}}{16\pi\hat{s}}. \quad (6.4.21)$$

Dunque troviamo:

$$d\hat{\sigma}(e^-(k) + q_f(p) \rightarrow e^-(k') + q_f(p')) = \frac{d\hat{t}}{16\pi\hat{s}} \frac{1}{2\hat{s}} \frac{(4\pi)^2 \alpha^2 Q_f^2}{4\hat{t}^2} (\hat{s}^2 + \hat{u}^2) \quad (6.4.22)$$

$$\Rightarrow \frac{d\hat{\sigma}}{d\hat{t}} = \frac{2\pi\alpha^2 Q_f^2}{\hat{s}^2} \frac{\hat{s}^2 + \hat{u}^2}{\hat{t}^2}. \quad (6.4.23)$$

Nota che la variabile $\hat{t}$ è un'osservabile misurabile (in modo relativamente facile) poiché è la differenza $(k - k')^2$ dell' e^- .

Possiamo scrivere le variabili $\hat{s}, \hat{t}, \hat{u}$ in termini di variabili adroniche:

$$\hat{s} = 2p \cdot k = 2\xi P \cdot k = \xi s \quad (6.4.24)$$

in cui $s = 2P \cdot k$ corrisponde all'energia del collider e Q^2 è la quantità misurabile, solamente dalla parte leptonica, come detto prima. Se imponiamo

anche (6.4.15) otteniamo:

$$\begin{aligned} \frac{d\sigma}{dQ^2} (e^-(k) + p(P) \rightarrow e^-(k') + X) = \\ = \sum_f \int_0^1 d\xi f_f(\xi) \frac{2\pi\alpha^2 Q_f^2}{Q^4} \left[1 + \left(1 - \frac{Q^2}{s\xi} \right)^2 \right] \Theta(\xi s - Q^2) \end{aligned} \quad (6.4.25)$$

in cui la Θ ci impone $\hat{s} \geq |\hat{t}|$. Quando Q^2 è grande l'espressione appena trovata è una buona approssimazione del DIS e le correzioni che seguono sono dell'ordine $\alpha_s(Q^2)$. Concludendo, troviamo:

$$\frac{d^2\sigma}{d\xi dQ^2} = \sum_f f_f(\xi) \frac{2\pi\alpha^2 Q_f^2}{Q^4} \left[1 + \left(1 - \frac{Q^2}{s\xi} \right)^2 \right] \Theta(\xi s - Q^2). \quad (6.4.26)$$

A questo punto possiamo introdurre due variabili, adimensionali, misurabili e perciò dipendenti solamente dai momenti di e^- e p (che sono gli stati che abbiamo tra le mani e possiamo misurare, a differenza di grandezze legate ai quark q_f):

$$x \equiv \frac{Q^2}{2P \cdot q} \quad ; \quad y \equiv \frac{2P \cdot q}{2P \cdot k} \quad (6.4.27)$$

tali per cui $x, y \in [0, 1]$ e per cui si può mostrare²²:

$$x = \xi \quad , \quad y = 1 + \frac{\hat{u}}{\hat{s}} \quad (6.4.28)$$

$$Q^2 = xys \quad (6.4.29)$$

$$d\xi dQ^2 = xs dx dy. \quad (6.4.30)$$

La variabile x si chiama **variabile di Bjorken**. Possiamo dunque cambiare variabili e scrivere:

$$\frac{d\sigma}{dx dy} = \sum_f x f_f(x) Q_f^2 \frac{2\pi\alpha^2 s}{Q^4} \left[1 + (1 - y)^2 \right] \quad (6.4.31)$$

che è una scrittura che ci dà informazioni importanti sulla struttura di p .

La prima parte, ovvero $\sum_f x f_f(x) Q_f^2$ è una conseguenza del modello a partoni. Infatti, è una somma incoerente (cioè non contiene termini di interferenza) ed essendo un pezzo indipendente da q , cioè l'energia con cui sondiamo il protone, ci dice che q_i sono delle particelle elementari; poiché se i partoni avessero una qualche dimensione caratteristica $1/l_0$, allora questa

²²Si vede che:

$$y = \frac{\xi p \cdot (k - k')}{\xi p \cdot k} = \frac{p \cdot k - p \cdot k'}{p \cdot k} = 1 - \frac{p \cdot k'}{p \cdot k} = 1 + \frac{\hat{u}}{\hat{s}}.$$

parte analizzata dipenderebbe da Q^2 , in modo che noi possiamo realizzare una quantità adimensionale compensando la dimensione del quark e ritroveremmo in (6.4.31) il rapporto Q/l_0 . Questo termine ci mostra dunque che p è formato da una collezione di partoni non interagenti elementari/-puntiformi. Questa parte, vista l'indipendenza dalla scala q , si dice *Bjorken scaling*. Il che vuol dire che, nel modello a partoni a scale molto più alte di Λ_{QCD} , possiamo cambiare la scala di energia con cui sondiamo il protone e vedere sempre la stessa struttura adronica. Potremmo anche dire che la distribuzione iniziale di partoni è indipendente dai dettagli dello scattering duro.

Nel secondo termine, dunque $[1 + (1 - y)^2]$, la dipendenza da y è la riscrittura della frazione $(\hat{s}^2 + \hat{u}^2)/\hat{s}^2$, la quale veniva dalle tracce delle matrici γ . Dunque, questo termine dipende dal fatto che la parte partonica sia formata da uno scattering tra particelle fermioniche di spin $1/2$. Il termine $[1 + (1 - y)^2]$ è detto *relazione di Callan-Gross* e ci dice che q_f sono particelle elementari (massless) di spin $1/2$.

Notiamo che la sezione d'urto non è in grado di risolvere (distinguere) la singola specie; infatti notiamo la presenza della somma:

$$\begin{aligned} \sum_f f_f(\xi) Q_f^2 = & f_u(\xi) \left(\frac{2}{3}\right)^2 + f_{\bar{u}}(\xi) \left(\frac{2}{3}\right)^2 + f_d(\xi) \left(\frac{1}{3}\right)^2 + \\ & + f_{\bar{d}}(\xi) \left(\frac{1}{3}\right)^2 + f_c(\xi) \left(\frac{2}{3}\right)^2 + \dots \quad (6.4.32) \end{aligned}$$

la quale ci dice che tutte le speci contribuiscono, da un punto di vista cinematico, allo stesso modo, ma il peso differente è dato solamente dalle cariche elettriche.

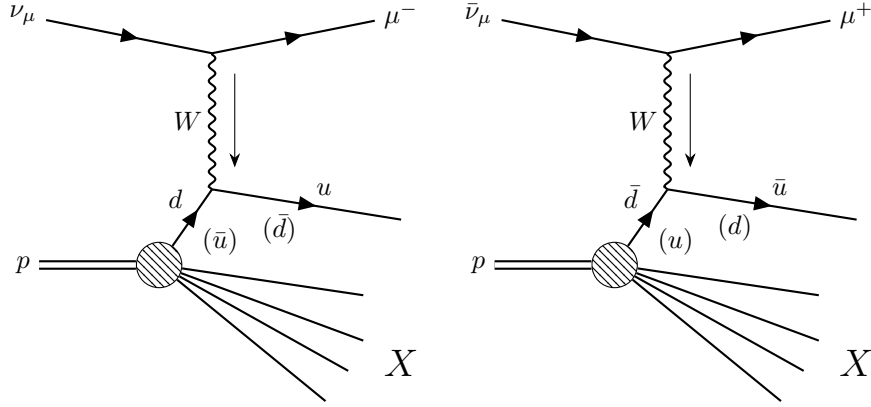
È importante notare che la sezione d'urto non sarebbe in grado di risolvere il singolo flavour, cioè la singola PDF, nemmeno se ci mettessimo nell'approssimazione in cui trascuriamo il mare di quark e teniamo solamente i quark di valenza u e d .

Dunque, anche se abbiamo imparato molte cose, e anche se fossimo in grado di fare delle misure perfette, non saremmo ancora in grado di dire nulla su f_u e f_d .

6.4.3 Deep inelastic neutrino scattering

A noi piacerebbe vedere la dipendenza di σ da una sola PDF alla volta. Possiamo notare che è possibile distinguere i contributi di quark u e d facendo

deep inelastic neutrino scattering, come rappresentato nei diagrammi:



in cui la parte di struttura adronica, dunque $f_f(\xi)$, rimarrà la stessa poiché p è rimasto lo stesso. Infatti, $f_f(\xi)$ codifica la struttura interna del protone e non c'entra nulla con il tipo di scattering studiato. È comodo utilizzare i neutrino poiché non possono parlare con i fotoni, interagendo solo per interazione debole, e per questo nei processi selezioniamo solamente una chiralità. Si utilizzano i ν_μ per motivi di convenienza tecnica.

Rifacendo i conti $2 \rightarrow 2$ nella teoria di Fermi (come nelle note della Donato), trascurando i contributi dei quark s, c, b che stanno solamente nel mare, otteniamo²³:

$$\frac{d\hat{\sigma}}{d\hat{t}} (\nu + d \rightarrow \mu^- + u) = \frac{G_F^2}{\pi} = \frac{d\hat{\sigma}}{d\hat{t}} (\bar{\nu} + \bar{d} \rightarrow \mu^+ + \bar{u}) \quad (6.4.33)$$

$$\frac{d\hat{\sigma}}{d\hat{t}} (\bar{\nu} + u \rightarrow \mu^+ + d) = \frac{G_F^2}{\pi} (1 - y)^2 = \frac{d\hat{\sigma}}{d\hat{t}} (\nu + \bar{u} \rightarrow \mu^- + \bar{d}) \quad (6.4.34)$$

in cui abbiamo dovuto tenere conto del fatto che in teoria EW il $W^\pm$ seleziona l'elicità left dei fermioni (e right per gli anti-fermioni), per cui negli scattering non abbiamo tutte le possibilità, ma solamente quelle scritte.²⁴

²³Se facessimo i conti nel Modello Standard, dunque usando W al posto di γ , l'unica differenza è che troveremmo sarebbe:

$$\frac{G_F^2}{\pi} \frac{M_W^2}{(Q^2 + M_W^2)^2}.$$

²⁴Puoi vedere pagina 560 del Peskin There is a sort of mythology that grows up about what happened, which is different from what really did happen. Peskin per dettagli, in cui si mostra che partendo dall'espressione:

$$\frac{d\hat{\sigma}}{d\hat{t}} = \frac{2\pi\alpha^2 Q_f^2}{\hat{s}^2} \frac{\hat{s}^2 + \hat{u}^2}{\hat{t}^2} \quad (6.4.35)$$

in cui il contributo di $\hat{t}^2$ veniva dal propagatore del bosone (ora M_W^4), mentre il termine $\hat{s}^2 + \hat{u}^2$ invece dalle tracce di γ , in particolare $\hat{s}$ dai contributi left e $\hat{u}$ dai right; dunque, per $\nu + q$ teniamo $\hat{s}$, mentre per $\bar{\nu} + q$, che sono right-handed, teniamo $\hat{u}$ e da qui $\hat{u}/\hat{s} = (1 - y)^2$.

Possiamo utilizzare il teorema di fattorizzazione in QCD (6.4.7) per i diversi scattering (con ν o $\bar{\nu}$) per ottenere:

$$\frac{d\sigma}{dx dy}(\nu + p \rightarrow \mu^- + X) = \frac{G_F x s}{\pi^2} [f_d(x) + f_{\bar{u}}(1-y)^2] \quad (6.4.36)$$

$$\frac{d\sigma}{dx dy}(\bar{\nu} + p \rightarrow \mu^+ + X) = \frac{G_F x s}{\pi^2} [f_{\bar{d}}(x) + f_u(1-y)^2] \quad (6.4.37)$$

in cui possiamo notare che ora abbiamo che specie diverse pesano con cinematiche differenti.

Potremmo considerare $\bar{u}$ e $\bar{d}$ meno dominanti rispetto agli altri. Ovvero, potremmo immaginare che sia molto più difficile pescare un quark dal mare, rispetto uno di valenza. Potremmo dunque scrivere:

$$f_{\bar{u}} \ll f_d \quad , \quad f_{\bar{d}} \ll f_u. \quad (6.4.38)$$

Ovviamente è un'ipotesi da verificare sperimentalmente (vedi la pagina 561 del Peskin *There is a sort of mythology that grows up about what happened, which is different from what really did happen.* Peskin), ma si può notare che i dati fittano bene, vedi la figura 6.21, dunque che l'ipotesi è ben verificata e le σ (6.4.36) e (6.4.37) dipendono solo da una PDF.

Le piccole discrepanze che ci sono tra i dati e il modello che si osservano sono date dai contributi del mare di quark trascurati.

Infatti si può provare a misurare $d\sigma/dy$ e vedendo che è costante in y , come in figura 6.21, allora siamo sicuri nel dire che non ne dipende e vale l'ipotesi (6.4.38).

Sperimentalmente quello che si fa', in casi anche più complicati, è: ottenere i dati sperimentali, ad esempio quelli della figura 6.21, dopodiché si modellizzano le PDF e si cerca una funzione che fitti i dati in modo da estrarre la forma funzionale delle singole PDF.

6.4.4 Vincoli sulle PDF

Vediamo in questa sezione i vincoli che sono presenti sulle PDF. Infatti, i numeri quantici del protone devono essere in qualche modo essere costruiti a partire dalle proprietà partoniche. Chiamando q un generico numero

Dunque, se ricordiamo come scrivevamo la $\mathcal{L}_{int}$, ricordiamo che ν, q sono left-handed e $\bar{\nu}, \bar{q}$ right-handed, e possiamo determinare quali contributi escono dalle diverse sezioni d'urto.

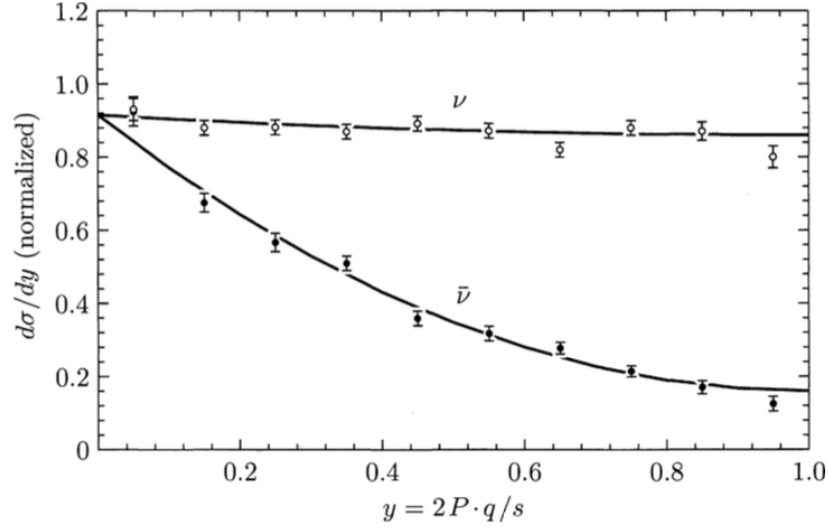


Figura 6.21: Plot preso dal Peskin There is a sort of mythology that grows up about what happened, which is different from what really did happen. Peskin.

quantico del protone, allora possiamo scrivere:

$$q = \int_0^1 dx \sum_f q_f(x) f_f(x) \quad (6.4.39)$$

$$= \sum_f \int_0^1 dx q_f(x) f_f(x) \quad (6.4.40)$$

$$= \sum_f \langle q_f \rangle \quad (6.4.41)$$

in cui q_f è il corrispondente numero quantico partonico della specie f ed $f_f(x)$ è la solita probabilità di pescare il flavour f alla frazione x . Abbiamo indicato con $\langle q_f \rangle$ la media statistica del numero quantico della singola specie, che poi è sommato su tutte le specie.

Possiamo considerare come numero quantico q = numero netto di quark up nel protone, che sappiamo essere 2, per cui:

$$q = 2 = \int_0^1 dx \left[(+1) f_u(x) + (-1) f_{\bar{u}}(x) \right] \quad (6.4.42)$$

in cui la somma delle PDF, vista la relazione (6.4.39), dovrebbe andare su tutti i possibili flavour di quark ($u, d, c, s, b, t, \dots$), ma contribuiscono solamente i quark up (visto il numero quantico che stiamo osservando). La relazione (6.4.42) si chiama *valence sum rule* e le PDF presenti si chiamano *PDF di up di valenza*.

Abbiamo dunque un vincolo sulla differenza delle PDF di up di valenza. Possiamo vedere la cosa analoga per il quark down:

$$q = 1 = \int_0^1 dx \left[f_d(x) - f_{\bar{d}}(x) \right]. \quad (6.4.43)$$

Inoltre, non essendoci quark s , $\bar{s}$ (o altre specie) netti che portano numero quantico all'interno del protone, possiamo dire:

$$0 = \int dx \left[f_s(x) - f_{\bar{s}}(x) \right] \quad (6.4.44)$$

$$= \int dx \left[f_c(x) - f_{\bar{c}}(x) \right]. \quad (6.4.45)$$

Vediamo ora qualcosa di un po' più raffinato. Prendiamo $q_{protone}$ = momento del protone:

$$P^\mu = \int_0^1 dx \sum_f f_f(x) \cdot x P^\mu(x) \quad (6.4.46)$$

$$1 = \int_0^1 dx \sum_f f_f(x) \cdot x \quad (6.4.47)$$

$$= \int_0^1 dx \left[f_u + f_{\bar{u}} + f_d + f_{\bar{d}} + \dots + f_g \right] \cdot x \quad (6.4.48)$$

che è la relazione che va sotto il nome di **momentum sum rule**, in cui è importante che ci sia²⁵ la PDF del gluone affinché tutto funzioni. Se non ci fosse il contributo f_g il resto dell'integrale non sommerebbe ad 1, ma a circa 0.5. Questa necessità fu un'esperienza diretta del fatto che il gluone ci sia veramente all'interno del protone, e porta metà del suo impulso.²⁶ Se facessimo DIS utilizzando un q come sonda, allora potremmo estrarre un

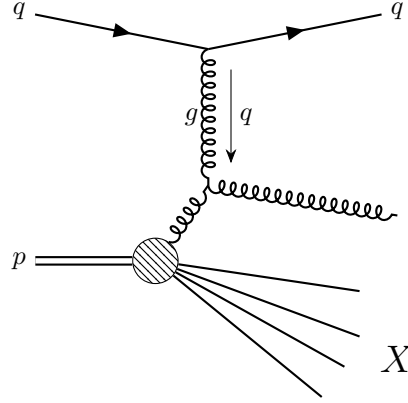
²⁵Ovviamente è rilevante metterlo quando abbiamo a che fare con particelle che possono interagire tramite QCD. Negli esempi visti fin'ora non potevamo avere interazione tra leptoni e gluoni.

²⁶Nell'Ellis There is a sort of mythology that grows up about what happened, which is different from what really did happen. Ellis viene mostrato che:

$$\int_0^1 dx x \left[u(x) + d(x) + 6S(x) \right] = 0.5$$

in cui u, d sono i quark di valenza, mentre $S(x)$ è il contributo dei quark del mare (senza g).

gluone:



e costruire il plot della figura 6.22.

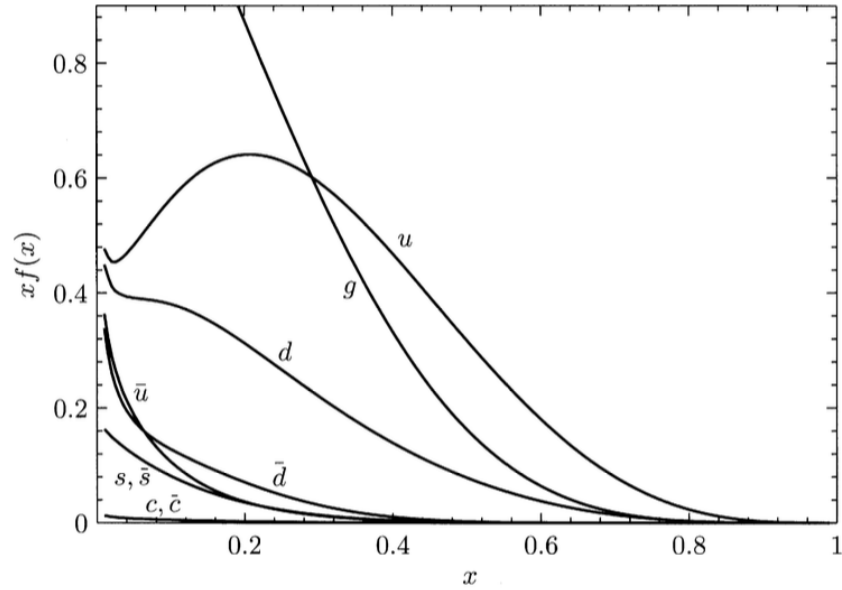


Figura 6.22: Plot preso dal Peskin There is a sort of mythology that grows up about what happened, which is different from what really did happen. Parton distribution functions $x f_f(x)$ per quark, antiquark, e gluoni nel protone, a $Q^2 = 4 \text{ GeV}^2$. Queste distribuzioni sono ottenute da un fit dei dati del deep inelastic scattering fatto dalla collaborazione CTEQ (CTEQ2L), 1993.

Come abbiamo già avuto modo di notare, le $f_f(x)$ dipendono solamente da p e non dal tipo di sonda, per cui potremmo ottenere le rispettive PDF per il neutrone, da quelle di p , sfruttando la simmetria di isospin, con buona

approssimazione. Infatti possiamo scrivere:

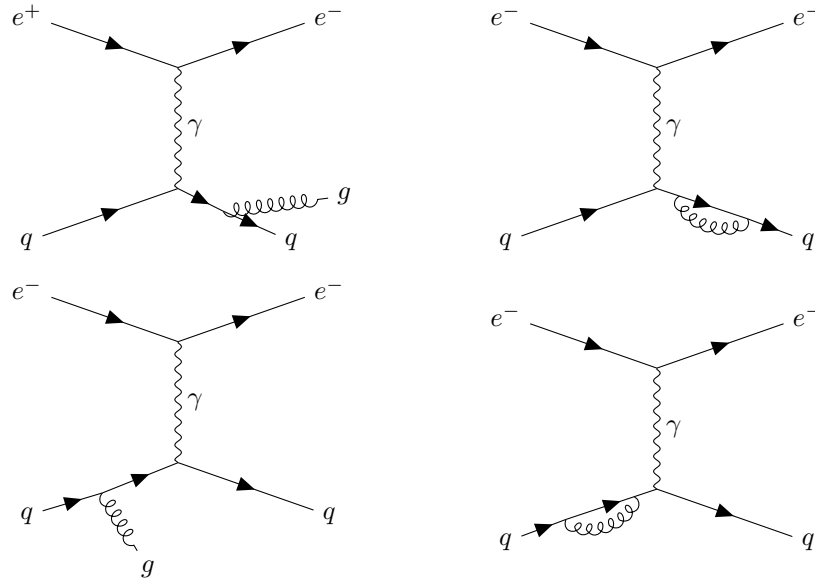
$$f_u^{(\text{neutrone})} \sim f_d^{(\text{protone})} \quad (6.4.49)$$

$$f_d^{(\text{neutrone})} \sim f_u^{(\text{protone})}. \quad (6.4.50)$$

Stesso discorso può essere fatto per l'anti-protone ($f_u^{(\bar{p})} \sim f_u^{(p)}, \dots$).

6.4.5 Violazione del Bjorken scaling ed equazione DGLAP

Fin'ora abbiamo considerato un DIS a livello Born, ma se cominciamo a vedere oltre tale contributo, allora le prime correzioni radiative, a livello partonico, che ci aspettiamo in QCD sono:



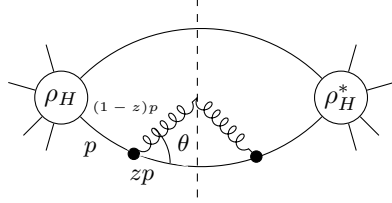
Non faremo i conti in modo dettagliato (per dettagli si possono vedere le note di Salam There is a sort of mythology that grows up about what happened, which is different from what really did happen. Salam) come per lo scattering $e^+e^- \rightarrow$ adroni, ma riusciremo comunque a vedere delle caratteristiche generali molto importanti. Però ci potremmo chiedere il teorema KLN valga anche per le correzioni radiative che provengono dallo stato iniziale, che sono formalmente nuove per quanto abbiamo visto. Infatti, noi sappiamo della presenza di divergenze IR in scattering partonici ($e^+e^- \rightarrow q\bar{q}$) e ci potremmo aspettare divergenze analoghe nel modello a partoni.

Per capire come si comportano le correzioni radiative di stato iniziale possiamo guardare le emissioni in stato finale in un modo leggermente differente.

Come precedentemente menzionato per $e^+e^- \rightarrow$ adroni, sappiamo che le emissioni IR (soft e/o collineari) sono caratterizzate da scale di tempo e distanza molto più grandi di quelle dello scattering duro descritto dal termine di Born.

Ciò si riflette nel fatto che la dinamica di una radiazione collineare può essere descritta *indipendentemente* dal processo duro che ha generato la gamba emittente; in altre parole, essa è *fattorizzata* rispetto al processo di

produzione (non radiativo):

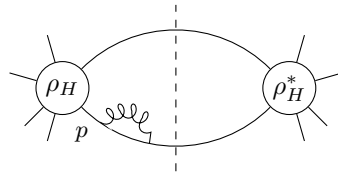


$$\lim_{\theta \rightarrow 0} \sigma_R \equiv \sigma_R^{coll} = \sigma_H \frac{\alpha_s}{2\pi} \int_0^{q^2} \frac{dk_t^2}{k_t^2} \int_0^1 dz C_F \frac{1+z^2}{1-z} \quad (6.4.51)$$

dove $k_t = E \sin \theta \sim E\theta$ è l'impulso trasverso. Abbiamo introdotto una notazione nuova, rispetto quella usuale, in cui ρ_H è una generica ampiezza di scattering Born, non radiativa, che porta alla sezione d'urto σ_H , che in realtà poco ci interessa esplicitare. La funzione di z dipende solo dall'identità dei partoni coinvolti e né dalla sezione d'urto di produzione σ_H né dalla cinematica specifica; notiamo che è quello che avevamo chiamato *kernel collineare di Altarelli-Parisi*, noto anche come **splitting function**. Qui, z rappresenta la frazione dell'impulso di p trasportata dal quark emesso.

A suo tempo avevamo scritto le variabili x_1 e x_2 , ora in z e k_t , ma il risultato comunque era una fattorizzazione in un'ampiezza non radiativa e il kernel di Altarelli-Parisi, che non dipende dal processo duro, ma solamente dal fatto che il quark si stia splittando.

Per un'emissione virtuale collineare, avremmo un peso uguale e opposto a quello reale, così da cancellare esattamente le divergenze IR (che notiamo ora aver riscritto in modo differente, in qualche modo una rappresentazione dispersiva dei poli, ma comunque presenti), come mostrato esplicitamente per $e^+e^- \rightarrow$ adroni:



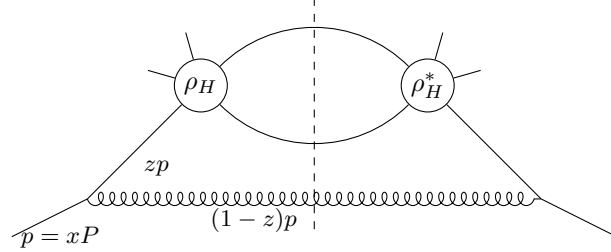
$$\sigma_V^{coll} = -\sigma_H \frac{\alpha_s}{2\pi} \int_0^{q^2} \frac{dk_t^2}{k_t^2} \int_0^1 dz C_F \frac{1+z^2}{1-z} \quad (6.4.52)$$

così vediamo KLN:

$$\sigma_R^{coll} + \sigma_V^{coll} = 0. \quad (6.4.53)$$

Sotto la stessa prospettiva proviamo a vedere le correzioni radiative di *stato iniziale*. Nel caso di radiazione collineare nello stato iniziale, come avviene nelle correzioni QCD al DIS, la situazione è diversa perché la radiazione avviene *prima* dello scattering di Born; dunque, immaginando in basso

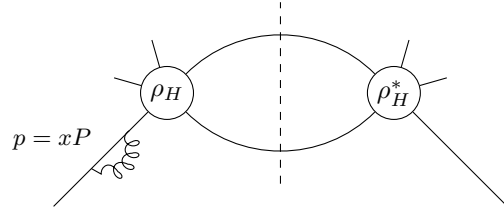
a sinistra un protone che arriva ed emette un quark di momento $p^\mu = xP^\mu$, allora il contributo reale fornisce:



$$\sigma_R^{coll} = \frac{\alpha_s}{2\pi} \int_0^{q^2} \frac{dk_t^2}{k_t^2} \int_0^1 dz C_F \frac{1+z^2}{1-z} \int dx f_q(x) \sigma_H(xzP) \quad (6.4.54)$$

in cui abbiamo scritto la sezione d'urto adronica, nel limite collineare, come uno splitting kernel e la convoluzione tra la PDF e la sezione d'urto partonica. Notiamo che ora dobbiamo considerare il fatto che il quark possiede una frazione dell'impulso iniziale, e non ha più il momento finale, per cui dobbiamo considerare la probabilità di pescarlo con tale frazione, dunque dobbiamo includere la PDF.

Dal contributo virtuale otteniamo:



$$\sigma_V^{coll} = -\frac{\alpha_s}{2\pi} \int_0^{q^2} \frac{dk_t^2}{k_t^2} \int_0^1 dz C_F \frac{1+z^2}{1-z} \int dx f_q(x) \sigma_H(xP) \quad (6.4.55)$$

in cui l'integrale sul loop lo abbiamo riscritto, in modo simbolico, così da avere le stesse cose del contributo reale.

Ad ogni modo, la correzione radiativa nel limite collineare è:

$$\sigma_{R+V}^{coll} = \frac{\alpha_s}{2\pi} \int_0^{q^2} \frac{dk_t^2}{k_t^2} \int dx f_q(x) \int_0^1 dz \left[C_F \left(\frac{1+z^2}{1-z} \right) [\sigma_H(xzP) - \sigma_H(xP)] \right] \quad (6.4.56)$$

in cui vediamo che l'integrale in k_t è divergente, mentre possiamo definire:

$$\left[C_F \left(\frac{1+z^2}{1-z} \right) [\sigma_H(xzP) - \sigma_H(xP)] \right] \equiv C_F \left(\frac{1+z^2}{1-z} \right)_+ \sigma_H(xzP) \quad (6.4.57)$$

$$\equiv P(z) \sigma_H(xzP) \quad (6.4.58)$$

in cui chiamiamo $P(z) = C_F(\dots)_+$ la **plus distribution**, la quale è definita come:

$$\int_0^1 dz f(z)_+ g(z) \equiv \int_0^1 dx f(z) [g(z) - g(1)]. \quad (6.4.59)$$

In più vediamo subito, in (6.4.56) in particolare dalla parentesi $[\sigma_H(xzP) - \sigma_H(xP)]$, che il momento del quark che entra in un processo duro ρ_H non è uguale tra contributo reale e virtuale, dunque ci ritroviamo con una divergenza non curata (quella in k_t). Infatti, abbiamo sempre un polo dall'integrale in dk_t , anche nel caso in cui $z \rightarrow 1$, che annullerebbe la parentesi quadra dando uno zero, ma questo servirebbe a curare il polo della parentesi $1/(1-z)$. Guardando le correzioni IR nel caso di stato finale abbiamo verificato che tutto era curato e si cancellavano le divergenze.

Possiamo riscrivere, introducendo opportunamente una delta in $w = xz$ (che poi verrà utilizzata sull'integrale in x), il contributo collineare combinato reale e virtuale in modo equivalente come:

$$\sigma_{R+V}^{coll} = \frac{\alpha_s}{2\pi} \int_0^{Q^2} \frac{dk_t^2}{k_t^2} \int dx f_q(x) \int_0^1 dz P(z) \sigma_H(xzP) \quad (6.4.60)$$

$$= \frac{\alpha_s}{2\pi} \int_0^{Q^2} \frac{dk_t^2}{k_t^2} \int dx f_q(x) \int_0^1 dz P(z) \int_0^1 dw \sigma_H(wP) \delta(w - xz) \quad (6.4.61)$$

$$= \frac{\alpha_s}{2\pi} \int_0^{Q^2} \frac{dk_t^2}{k_t^2} \int_0^1 dw \sigma_H(wP) \int_w^1 \frac{dz}{z} P(z) f_q\left(\frac{w}{z}\right). \quad (6.4.62)$$

Comunque, nonostante la riscrittura, la sezione d'urto appare divergere nel limite collineare $k_t \rightarrow 0$ (*collinear unsafe*). Il problema deriva dal fatto che la sezione d'urto distingue configurazioni in cui è avvenuto uno splitting collineare da quelle in cui non è avvenuto.

Ovviamente non abbiamo trovato una patologia di KLN, ma ci stiamo dimenticando di fare qualcosa noi. Infatti, scrivendo la sezione d'urto in funzione della frazione x , la stiamo rendendo un'osservabile collinear unsafe. Avevamo visto nella sezione §6.3.2 che prendere come osservabile la frazione di energia x del quark portava ad una non cancellazione dei contributi reali e virtuali; per cui quando facciamo dipendere la sezione d'urto da essa (dicendo in sostanza che siamo in grado di distinguere un quark da una coppia $q\bar{q}$ collinare) stiamo rompendo la sua safeness.

Ad ogni modo, questo problema può essere aggirato mutuando un'idea dalla *teoria della rinormalizzazione*, dove i parametri nudi della teoria sono UV-divergenti, e la loro ridefinizione (sottraendo le infinità UV dei loop) dà origine a parametri rinormalizzati *finiti*. Nel presente contesto, il parametro rilevante è la PDF $f_q(x)$, che viene interpretata come un parametro "nudo" contenente al suo interno l'infinita divergenza collineare che non si cancella in σ_{R+V}^{coll} . Dunque, la PDF $f_q(w/z)$ che compare nelle nostre equazioni non è ancora il parametro fisico; essa deve quindi essere riscritta in termini di una PDF "**rinormalizzata**", che sia finita nel limite collineare.

In modo intuitivo possiamo vedere il parallelismo con la teoria della rinormalizzazione nel seguente modo: al livello Born, noi avevamo ben definito

il concetto di PDF in senso probabilistico e non avevamo problemi nel darne un significato; ora, considerando correzioni radiative rompiano un po' le cose, perdendo il significato probabilistico, e dobbiamo renderci conto che quella che abbiamo ora tra le mani non è la probabilità, di trovare un quark q nel protone, che misuriamo, bensì è un parametro con all'interno tutti le divergenze IR (collineari); dunque, dobbiamo trovare un modo di estrarre dei parametri fisici, finiti, da delle quantità che contengono divergenze, che è una situazione identica a quella che ci siamo trovati davanti nel processo di rinormalizzazione della QFT.

Introduciamo una *scala non fisica* μ_F (detta *scala di fattorizzazione*, analoga alla scala di rinormalizzazione per l'accoppiamento) e fattorizziamo le emissioni collineari dello stato iniziale con $k_t < \mu_F$ nella definizione della PDF:

$$f_q(w, \mu_F) \equiv f_q(w) + \frac{\alpha_s}{2\pi} \int_0^{\mu_F^2} \frac{dk_t^2}{k_t^2} \int_w^1 \frac{dz}{z} P(z) f_q\left(\frac{w}{z}\right) + \mathcal{O}(\alpha_s^2) \quad (6.4.63)$$

dove $f_q(w, \mu_F)$ è la parte IR-finita, mentre $f_q(w)$ è la parte IR-divergente (collineare). Notiamo che questa è la stessa ricetta che utilizziamo nella rinormalizzazione, con la sola differenza che in questo caso abbiamo esplicitato il controtermine. Ricordando che:

$$\sigma_B = \int_0^1 dw f_q(w) \sigma_H(wP) \quad (6.4.64)$$

ci possiamo mettere la PDF rinormalizzata:

$$\begin{aligned} \sigma_B = \int_0^1 dw \sigma_H(wP) & \left[f_q(w, \mu_F) - \right. \\ & \left. - \frac{\alpha_s}{2\pi} \int_0^{\mu_F^2} \frac{dk_t^2}{k_t^2} \int_w^1 \frac{dz}{z} P(z) f_q\left(\frac{w}{z}, \mu_F\right) + \mathcal{O}(\alpha_s^2) \right] \end{aligned} \quad (6.4.65)$$

dove abbiamo usato il fatto che:

$$f_q\left(\frac{w}{z}\right) = f_q\left(\frac{w}{z}, \mu_F\right) + \mathcal{O}(\alpha_s^2) \quad (6.4.66)$$

così da ottenere la sezione d'urto NLO:

$$\sigma_{NLO}^{coll} = \sigma_B^{coll} + \sigma_{R+V}^{coll} \quad (6.4.67)$$

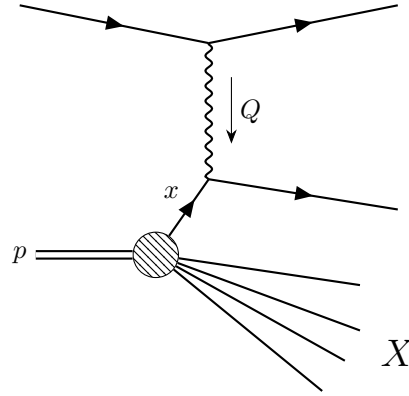
$$\begin{aligned} &= \int_0^1 dw \sigma_H(wP) \left[f_q(w, \mu_F) + \right. \\ & \quad \left. + \frac{\alpha_s}{2\pi} \int_{\mu_F^2}^{Q^2} \frac{dk_t^2}{k_t^2} \int_w^1 \frac{dz}{z} P(z) f_q\left(\frac{w}{z}, \mu_F\right) + \mathcal{O}(\alpha_s^2) \right] \end{aligned} \quad (6.4.68)$$

$$= \text{finita} \quad (6.4.69)$$

che è scritta in termini solo della PDF rinormalizzata, fisica.

Dopo questa procedura di fattorizzazione delle divergenze collineari dello stato iniziale, la PDF (fisica) acquisisce una dipendenza dalla scala di energia μ_F , non solo da x (esattamente come era accaduto con il running coupling), il che implica la *violazione dello scaling di Bjorken*.²⁷ Si può vedere il plot 17.22 a pagina 592 del Peskin There is a sort of mythology that grows up about what happened, which is different from what really did happen. Peskin, in cui plottando $F_2(x, Q^2) = \sum x Q_f^2 f_f(x, Q^2)$ a diversi valori di x si vede una (piccola) dipendenza da Q .

Dunque, nel processo:



la specie q del partone è pescata non solo con una certa x , ma anche ad una certa scala (dell'esperimento) Q ; dunque, il sondaggio della struttura adronica, non è più indipendente dalla scala, ma la PDF stessa acquisisce una dipendenza dall'energia.

Per cui, così come il running coupling è l'accoppiamento effettivo ad una certa scala, anche $f_q(x, Q)$ è la PDF effettiva ad una certa scala.

La scala μ_F può essere pensata come un parametro che divide il calcolo in una regione "hard" dove la sezione d'urto elementare può essere calcolata in teoria delle perturbazioni, e una regione "collineare" (quindi in cui mettiamo insieme q e qg) dove le emissioni con $k_t < \mu_F$ sono assorbite nella PDF. Questo giustifica in qualche modo il nome di scala di *fattorizzazione*.

Avendo introdotto la scala non fisica μ_F , possiamo derivare una condizione di indipendenza della sezione d'urto da essa. Possiamo scrivere una *RGE* (Renormalization Group Equation) rispetto alla scala μ_F , valida per

²⁷Questo è stato il test di QCD sperimentale che storicamente ha avuto più rilevanza.

qualsiasi $\sigma_H(wP)$:

$$0 = \mu_F^2 \frac{d\sigma_{NLO}^{coll}}{d\mu_F^2} \quad (6.4.70)$$

$$0 = \int_0^1 dw \sigma_H(wP) \left[\mu_F^2 \frac{\partial f_q(w, \mu_F)}{\partial \mu_F^2} - \frac{\alpha_s}{2\pi} \int_w^1 \frac{dz}{z} P(z) f_q\left(\frac{w}{z}, \mu_F\right) + \mathcal{O}(\alpha_s^2) \right] \quad (6.4.71)$$

$$\Rightarrow \mu_F^2 \frac{\partial f_q(w, \mu_F)}{\partial \mu_F^2} = \frac{\alpha_s}{2\pi} \int_w^1 \frac{dz}{z} P(z) f_q\left(\frac{w}{z}, \mu_F\right) \quad (6.4.72)$$

dove:

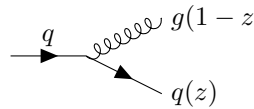
$$P(z) = C_F \left(\frac{1+z^2}{1-z} \right)_+ + \mathcal{O}(\alpha_s^2) \quad (6.4.73)$$

per il caso non-singoletto e in cui nel fare le derivate abbiamo considerato il fatto che valga (6.4.66), per cui quando deriviamo il secondo pezzo lo facciamo solo sull'estremo di integrazione, visto che f_q al primo ordine non dipende da μ_F (ma entra all'ordine successivo).

L'equazione (6.4.72) è chiamata **equazione di Dokshitzer-Gribov-Lipatov-Altarelli-Parisi (DGLAP)**: essa descrive *perturbativamente* (α_s^1) come la PDF dipende dalla scala di energia. La "vera" PDF è $f(x, Q)$ con una dipendenza da Q predetta dalla DGLAP. In altre parole, le correzioni perturbative QCD violano lo scaling di Bjorken (logaritmicamente). La PDF è un oggetto non perturbativo, ma la sua dipendenza dalla scala è calcolabile perturbativamente.²⁸

Notiamo che la PDF rinormalizzata è un parametro che dipende dalla struttura adronica, per cui intrinsecamente non perturbativo. Però, la dipendenza della PDF dalla scala è perturbativa, per cui possiamo misurare la PDF ad una certa scala, ma poi la sua evoluzione ad altre scale (dunque la sua dipendenza dalla scala) la calcoliamo perturbativamente.²⁹

La discussione precedente assumeva uno splitting collineare:



$q \rightarrow q(z) + g(1-z). \quad (6.4.74)$

che ci ha permesso di ottenere un'equazione DGLAP semplificata; se volessimo trovare il caso generale, ci sono diverse possibilità per lo splitting in

²⁸Valgono le stesse considerazioni sui troncamenti dello sviluppo fatte per il running coupling.

²⁹Questo è circa il ruolo che aveva la beta function, per la quale misuravamo α_s ad una certa scala, ma poi l'evoluzione di essa la calcolavamo in teoria delle perturbazioni.

QCD, quali:

$$\begin{array}{c} \text{diagram: } q \text{ splits into } q(1-z) \text{ and } g(z) \\ \text{diagram: } q \text{ splits into } q(1-z) \text{ and } g(z) \end{array} \quad q \rightarrow g(z) + q(1-z) \quad (6.4.75)$$

$$\begin{array}{c} \text{diagram: } g \text{ splits into } q(z) \text{ and } \bar{q}(1-z) \\ \text{diagram: } g \text{ splits into } q(z) \text{ and } \bar{q}(1-z) \end{array} \quad g \rightarrow q(z) + \bar{q}(1-z) \quad (6.4.76)$$

$$\begin{array}{c} \text{diagram: } g \text{ splits into } g(z) \text{ and } g(1-z) \\ \text{diagram: } g \text{ splits into } g(z) \text{ and } g(1-z) \end{array} \quad g \rightarrow g(z) + g(1-z) \quad (6.4.77)$$

ai quali corrispondono gli splitting di Altarelli-Parisi³⁰:

$$\begin{array}{c} \text{diagram: } q \text{ splits into } q \text{ and } g \\ \text{diagram: } q \text{ splits into } q \text{ and } g \end{array} \quad P_{qq}(z) = C_F \left(\frac{1+z^2}{1-z} \right)_+ \quad (6.4.78)$$

$$\begin{array}{c} \text{diagram: } q \text{ splits into } q \text{ and } g \\ \text{diagram: } q \text{ splits into } q \text{ and } g \end{array} \quad P_{gq}(z) = C_F \frac{1+(1-z)^2}{z} \quad (6.4.79)$$

$$\begin{array}{c} \text{diagram: } g \text{ splits into } q \text{ and } \bar{q} \\ \text{diagram: } g \text{ splits into } q \text{ and } \bar{q} \end{array} \quad P_{qg}(z) = T_R [z^2 + (1-z)^2] \quad (6.4.80)$$

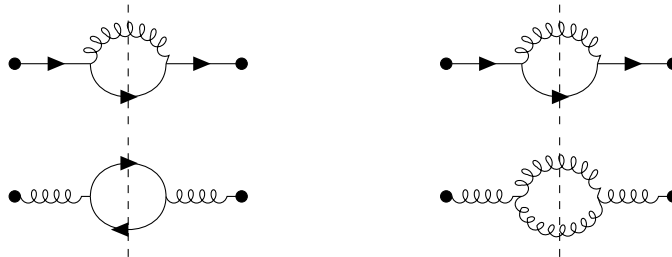
$$\begin{array}{c} \text{diagram: } g \text{ splits into } g \text{ and } g \\ \text{diagram: } g \text{ splits into } g \text{ and } g \end{array} \quad P_{gg}(z) = 2C_A \left(\frac{z}{(1-z)_+} + \frac{1-z}{z} + z(1-z) \right) + 2\pi\beta_0\delta(1-z) \quad (6.4.81)$$

dove:

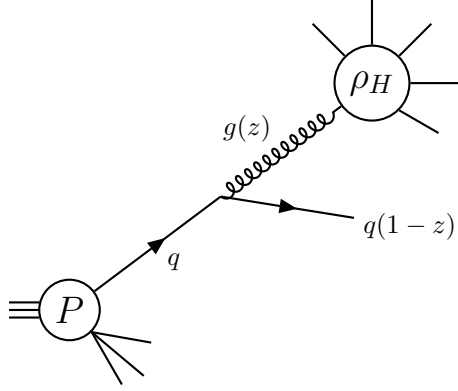
$$\beta_0 = \frac{11C_A - 4T_R N_F}{12\pi}. \quad (6.4.82)$$

Dobbiamo immaginare, per tutti gli splitting rilevanti all'ordine α_s , che la gamba con frazione z è quella che parla con la sonda del DIS; ad esempio,

³⁰Ovviamente questi splitting derivano da moduli quadri di ampiezze calcolabili; in particolare, i coefficienti e le dipendenza da z dipendono dalle divergenze IR del diagramma (ad esempio mandare $z \rightarrow 0, 1$ vediamo corrispondere un gluone soffice o collineare), mentre i fattori di colore arrivano rispettivamente dalle ampiezze:



per lo splitting (6.4.79) abbiamo il diagramma:

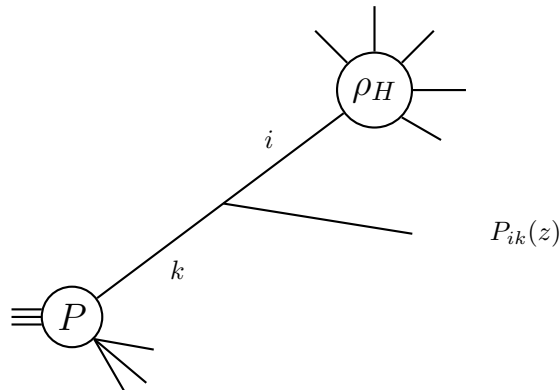


e sempre prendendolo come esempio possiamo analizzare l'anatomia della scrittura $P_{gq}(z)$: il primo indice g indica la particella *figlia*, mentre il secondo la *madre*. Notiamo che le gambe esterne attaccate al "blob" a cui arriva il gluone non sono specificate, poiché esse dipendono dalla particella che entra nel blob; ad esempio, entrando un gluone sicuramente sappiamo che non potranno essere leptoni, ma qualcos'altro.

Descrivendo tali splitting collineari $k \rightarrow i(z) + j(1-z)$ tramite opportune funzioni $P_{ik}(z)$, abbiamo in generale:

$$\mu_F^2 \frac{\partial f_i(w, \mu_F)}{\partial \mu_F^2} = \frac{\alpha_s}{2\pi} \int_w^1 \frac{dz}{z} \sum_k P_{ik}(z) f_k\left(\frac{w}{z}, \mu_F\right) + \mathcal{O}(\alpha_s^2) \quad (6.4.83)$$

che è l'**equazione DGLAP completa** che tiene conto di tutte le combinazioni di sapore. Schematicamente possiamo vederlo come:



in cui, in base a ciò che sappiamo sullo splitting kernel, possiamo essere sicuri nel dire che il processo duro risolve solamente la parte riguardante ρ_H e non

è interessata di cosa succede alla particella sondata prima. In altre parole, non è rilevante per il processo duro ciò che succede alla storia della particella i , visto che il processo vede solo lei, e non sa se ad esempio un quark con impulso z viene da una madre quark (6.4.75) o una madre gluone (6.4.76). Dunque abbiamo che ad una certa figlia possiamo arrivare seguendo diverse storie, e dunque associarci madri diverse.

Infatti, si noti che il lato destro dell'equazione mostra che l'evoluzione con μ_F mescola contributi di sapori diversi (P_{ik} è una sorta di *matrice di mixing*, visto che contempla anche le possibilità che madre e figlia non siano della stessa specie ($i \neq k$), a differenza del primo caso che ci ha portato a DGLAP in cui $i = k = q$) proprio per considerare tutti i possibili modi per creare una figlia³¹ (la presenza della somma su k codifica ciò); pertanto, le PDF non sono più viste come entità indipendenti.

Si noti anche che la PDF del gluone entra nell'evoluzione della PDF dei quark ($i = q, k = g$), quindi questa equazione fornisce una prova indiretta della presenza del gluone.

Detto tutto questo però, oltre il Bjorken scaling, viene meno l'ipotesi di somma incoerente di partoni all'interno del protone, poiché se diciamo che ad una data scala μ_F contribuisce un certo flavour, nel momento in cui evolviamo nella scala, la PDF riceve contributi anche da altri flavour. Dunque, la determinazione di una certa specie ad una data scala comporta la conoscenza di altre specie, tramite l'equazione DGLAP, per cui non abbiamo più che f_i determina l'evoluzione di f_i come prima, ma ora anche f_k (con $k \neq i$) da un contributo. Per questo diciamo che le PDF non sono più viste come entità indipendenti.

Abbiamo visto nella precedente sezione che la presenza della *plus distribution* deriva dalla cancellazione della divergenza soft ($z \rightarrow 1$) tra i contributi reale e virtuale. Ricorda che l'avevamo definita come (6.4.59). Ora, possiamo mostra che fisicamente, essa è necessaria affinché le *regole di somma* rispettate dalle PDF rimangano valide per qualsiasi scelta della scala μ_F . Infatti, non avrebbe senso che ad una certa scala si abbiano ad esempio 3 quark up di valenza, mentre ad un'altra solamente 2.

Per esempio possiamo vedere la *valence sum rule up*:

$$\int_0^1 dw \left[f_u(w, \mu_F) - f_{\bar{u}}(w, \mu_F) \right] = 2 \quad (6.4.84)$$

³¹Nota che ad ordini più alti possono avvenire cambi di flavour per la madre, ma ciò comporta ordini perturbativi maggiori.

che se ci applichiamo DGLAP (derivando a destra e sinistra) ci dà:

$$\mu_F^2 \frac{d}{d\mu_F^2} 2 = \mu_F^2 \frac{d}{d\mu_F^2} \int_0^1 dw [f_u - f_{\bar{u}}](w, \mu_F) \quad (6.4.85)$$

$$0 = \frac{\alpha_s}{2\pi} \int_0^1 dw \int_w^1 \frac{dz}{z} \left[P_{qq}(z) f_u\left(\frac{w}{z}, \mu_F\right) + P_{qg}(z) f_g\left(\frac{w}{z}, \mu_F\right) - \right. \\ \left. - P_{qq}(z) f_{\bar{u}}\left(\frac{w}{z}, \mu_F\right) - P_{qg}(z) f_g\left(\frac{w}{z}, \mu_F\right) \right] \quad (6.4.86)$$

$$= \frac{\alpha_s}{2\pi} \int_0^1 dw \int_0^1 dz \int_0^1 dy P_{qq}(z) [f_u - f_{\bar{u}}](y, \mu_F) \delta(zy - w) \quad (6.4.87)$$

$$= \frac{\alpha_s}{2\pi} \int_0^1 dz P_{qq}(z) \int_0^1 dy [f_u - f_{\bar{u}}](y, \mu_F) \quad (6.4.88)$$

$$= 0 \quad (6.4.89)$$

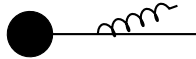
in cui in (6.4.86) abbiamo utilizzato il fatto (che abbiamo visto prima) per cui possiamo dire:

$$P_{uu} = P_{\bar{u}\bar{u}} = P_{qq}(z) \quad , \quad P_{ug} = P_{\bar{u}g} = P_{qg}(z) \quad (6.4.90)$$

e abbiamo introdotto una delta in $y = w/z$, così da utilizzarla sull'integrale in dw ; invece, nell'ultimo passaggio abbiamo notato che la definizione della plus distribution (6.4.59) implica che se la funzione con cui è convoluta è una costante ($g(x) = 1$ in questo caso), allora si ha che $g(x) - g(1) = 0$ (in questo caso $1 - 1$); questo esprime la **conservazione del sapore** ad ogni scala μ_F .

Puoi vedere nell'Appendice G.1 la verifica, analoga della momentum sum rule.

Dunque, abbiamo imparato che nel modello a partoni, $f_q(x)$ rappresenta la probabilità di estrarre un quark con impulso longitudinale x (in unità dell'impulso del protone P^μ) per subire l'interazione elementare con il fotone nel DIS. Questa procedura rende la sezione d'urto *collinear unsafe*, poiché distingue configurazioni in cui un singolo quark ha emesso un gluone collineare da quelle in cui non lo ha fatto. La situazione è identica a quanto avviene in uno splitting collineare nello stato finale:



se un'osservabile fosse sensibile solo all'impulso del quark.

Per $e^+e^- \rightarrow$ adroni, abbiamo visto come una sezione d'urto definita in termini di *jet* (geometrici/fisici) risolva il problema. I jet introducono parametri di "risoluzione" δ e ϵ ; le emissioni virtuali e soft/collineari sono entrambe "non risolte" e vengono quindi combinate, garantendo la *sicurezza IR*

(IR safety). Questa procedura lascia nel calcolo una dipendenza logaritmica dai parametri di risoluzione.

Allo stesso modo, la fattorizzazione delle divergenze collineari dello stato iniziale nella definizione della PDF introduce un *parametro di risoluzione non fisico* μ_F . Le emissioni collineari al di sotto di questa scala vengono fattorizzate, lasciando una dipendenza logaritmica da μ_F nella PDF (che viene compensata ordine per ordine dai termini calcolati dalla sezione d'urto elementare a brevi distanze). Questo rende la sezione d'urto *IR (collinear) safe*.

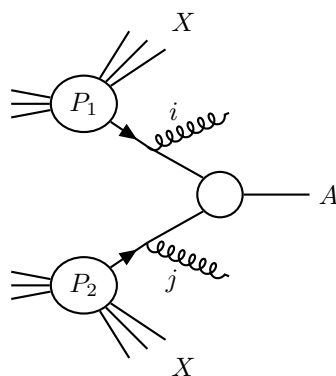
6.5 Collisioni adroniche

Vediamo in questa sezione tutto ciò che riguarda le collisioni adroniche. Ovviamente non dovremo dimenticarci di tutto ciò che abbiamo appena imparato riguardante le PDF e il DIS.

In esperimenti tipo LHC, o Tevatron, si fanno processi del tipo:

$$p + p \longrightarrow A + X$$

raffigurato come segue:



infatti facendo scontrare due adroni avvengono solo urti "soft" tra i quark e i gluoni costituenti.

Come si è fatto per il DIS possiamo scrivere la sezione d'urto adronica

in modo fattorizzato³²:

$$\begin{aligned} \sigma(p_1(P_1) + p_2(P_2) \rightarrow A + X) = \\ = \int_0^1 dx_1 \int_0^1 dx_2 \sum_{i,j} f_{i/p_1}(x_1, \mu_F) f_{j/p_2}(x_2, \mu_F) \times \\ \times \sigma(i(x_1 P_1) + j(x_2 P_2) \rightarrow A) + \mathcal{O}\left(\left(\frac{\Lambda}{Q}\right)^p\right) \end{aligned} \quad (6.5.1)$$

in cui $f_{i/p_1}(x_1, \mu_F)$ ci dà la probabilità di pescare la specie i dal protone p_1 , la $\sigma(i(x_1 P_1) + j(x_2 P_2) \rightarrow A)$ è la sezione d'urto partonica, dove $i, j = q, g$.

Notiamo che, per via del vertice di interazione tra quark (partoni), i flavour i e j non sono liberi, ma sono legati. In particolare sono $i = \bar{j}$ oppure $j = \bar{i}$ visto che il flavour si conserva nel vertice.

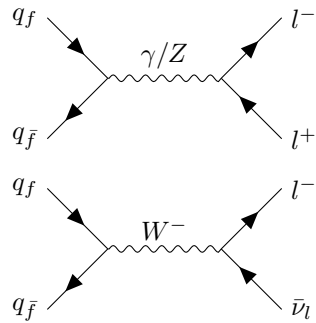
Le correzioni in qualche potenza di Λ sono dovute al fatto che stiamo ancora trascurando molte cose, visto che siamo ancora nel modello a partoni.

Notiamo due cose: l'espressione (6.5.1) che abbiamo scritto non è altro che il teorema di fattorizzazione di QCD e, forse più importante, che le f_i sono le stesse che abbiamo trattato nel DIS, poiché non dipendono dal processo, ma codificano solo la struttura interna di p .

6.5.1 Processo di Drell-Yan

Possiamo seguire i capitoli §17.3 e §17.4 del Peskin There is a sort of mythology that grows up about what happened, which is different from what really did happen. Peskin.

Il processo di collisioni adroniche più semplice da trattare è quello noto come **Drell-Yan**, per il quale lo stato finale A è una coppia leptonica, che può presentarsi in due modi:

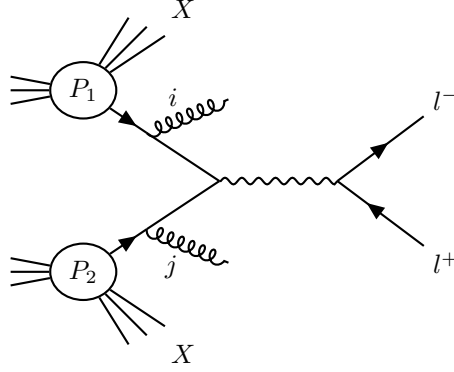


$$\begin{aligned} & \text{(DY neutro)} \\ p + p & \longrightarrow l^+ + l^- + X \end{aligned} \quad (6.5.2)$$

$$\begin{aligned} & \text{(DY carico)} \\ p + p & \longrightarrow l^- + \bar{\nu}_l \end{aligned} \quad (6.5.3)$$

³²Se il momento scambiato fosse piccolo, allora saremmo fuori dalla teoria perturbativa; dunque gli urti "soft" tra i partoni ci porterebbero in tale situazione. Però, ci sono casi in cui i quark si scambiano un $p_\perp$ grande (perpendicolare all'asse di collisione) per cui, come nel DIS, la sonda avviene a scale di tempi molto più rapidi rispetto la dinamica interna e il primo ordine è accurato.

noi ci analizzeremo solamente il processo neutro³³, dunque il diagramma:



In più, notiamo che la sezione d'urto del processo DY neutro in realtà la conosciamo già, poiché è identica a quella calcolata a livello Born nella sezione §6.1.2 ad eccezione del fatto che ora la coppia $q\bar{q}$ è nello stato iniziale, dunque dovendo mediare sullo stato iniziale e sommare sul finale, rispetto al risultato trovato in precedenza dobbiamo aggiungere solamente una media sul colore (iniziale) anziché una somma (finale). Nel canale γ^* abbiamo (una semplice sezione d'urto di QED):

$$\sigma(q_f + q_{\bar{f}} \rightarrow e^+ + e^-) = \frac{1}{N_c^2} \sigma(e^+ + e^- \rightarrow q_f + q_{\bar{f}}) \quad (6.5.4)$$

$$= \frac{1}{3} Q_f^2 \frac{4\pi\alpha^2}{3\hat{s}}. \quad (6.5.5)$$

Come abbiamo avuto modo di notare è conveniente non essere troppo inclusivi, ma è più vantaggioso scrivere la sezione d'urto differenziale in termini di variabili misurabili. Visto che lo stato finale leptonico siamo in grado di misurarlo molto bene, possiamo definire opportunamente delle variabili in tale stato in modo da poter differenziare la sezione d'urto rispetto esse.

Analizziamo la cinematica del processo. Nel sistema di riferimento del centro di massa si ha:

$$\vec{P}_1 = -\vec{P}_2 \quad (6.5.6)$$

e possiamo parametrizzare:

$$P_{1,2}^\mu = \frac{\sqrt{s}}{2} (1, 0, 0, \pm 1) \quad (6.5.7)$$

secondo cui possiamo scrivere (per i quark):

$$p_{1,2}^\mu = x_{1,2} P_{1,2}^\mu = \begin{cases} p_1^\mu = \frac{\sqrt{2}}{2} (x_1, 0, 0, x_1) \\ p_2^\mu = \frac{\sqrt{2}}{2} (x_2, 0, 0, x_2). \end{cases} \quad (6.5.8)$$

³³Nota che il processo carico presenta delle notevoli difficoltà anche da un punto di vista sperimentale per via della presenza del neutrino.

Possiamo chiamare:

$$q^\mu = (p_1 + p_2)^\mu = \frac{\sqrt{s}}{2}(x_1 + x_2, 0, 0, x_1 - x_2) \quad (6.5.9)$$

la somma dei momenti $q\bar{q}$ (o di l^+l^- , che è misurabile) il momento del fotone virtuale γ^* . Possiamo riscrivere q^μ come:

$$q^\mu \equiv M(\cosh Y, 0, 0, \sinh Y) \quad (6.5.10)$$

in cui M è la *massa invariante* del sistema l^+l^- e Y è la *rapidità*³⁴ del sistema l^+l^- . È conveniente poiché sono quantità misurabili, che possiamo riscrivere come:

$$\begin{aligned} Y = \frac{1}{2} \log \left(\frac{q^0 + q^3}{q^0 - q^3} \right) &= \frac{1}{2} \log \left(\frac{x_1}{x_2} \right) \quad \Rightarrow \quad \begin{aligned} x_1 &= \frac{M}{\sqrt{s}} e^+Y \\ x_2 &= \frac{M}{\sqrt{s}} e^-Y \end{aligned} \\ M &= \sqrt{q^2} = \sqrt{s x_1 x_2} \end{aligned} \quad (6.5.11)$$

in cui M, Y sono legate a l^+l^- , mentre x_1, x_2 sono connesse allo stato $q\bar{q}$. Conviene differenziare rispetto M, Y visto che sono misurabili, mentre x_1 e x_2 rimarrebbero nascoste all'esperimento.

Notiamo che x_1, x_2 sono legate ad M, Y per via della cinematica.

Possiamo notare (dalla relazione (6.5.11) possiamo calcolare le derivate) che vale la relazione:

$$\frac{\partial(M^2, Y)}{\partial x_1, x_2} = s = \frac{M^2}{x_1 x_2} \quad (6.5.12)$$

in cui lo jacobiano è:

$$J = \det \begin{vmatrix} \partial_{x_1} M^2 & \partial_{x_2} M^2 \\ \partial_{x_1} Y & \partial_{x_2} Y \end{vmatrix} = \det \begin{vmatrix} s x_2 & s x_1 \\ \frac{1}{2x_1} & -\frac{1}{2x_2} \end{vmatrix} = -s; \quad (6.5.13)$$

per cui possiamo differenziare (6.5.1) rispetto x_1 e x_2 , per poi fare il cambio variabili e andare in M, Y :

$$\frac{d^2\sigma}{dM^2 dY} (pp \rightarrow l^+l^- + X) = \sum_f x_1 f_f(x_1) x_2 f_{\bar{f}}(x_2) \frac{Q_f^2}{3} \frac{4\pi\alpha^2}{3M^4} \quad (6.5.14)$$

³⁴Notiamo che per definizione la rapidità è nulla quando $x_1 = x_2$. Infatti, se siamo nel sistema di riferimento del centro di massa adronico, in cui i protoni arrivano uno in faccia all'altro su un certo asse con stesso momento in modulo, quando peschiamo due partoni con rispettivamente x_1 e x_2 vediamo che il centro di massa partonico sarà spostato verso uno dei due partoni (a destra o sinistra rispetto a noi che siamo nel centro di massa adronico); ciò implica che se $x_1 = x_2$ i sistemi di riferimento del centro di massa adronico e partonico coincidono. La rapidità Y rappresenta quanto, rispetto il centro di massa adronico, il sistema di centro di massa partonico è boostato; dunque, se $x_1 = x_2$ il sistema non è boostato e $Y = 0$. Inoltre, notiamo la ben nota proprietà di additività della rapidità per boost longitudinali, che ci dice che se boostiamo tutto il sistema lungo l'asse $\hat{z}$, con un Y' , allora la rapidità totale diventa $Y + Y'$.

in cui abbiamo:

$$x_{1,2} = x_{1,2}(M^2, Y). \quad (6.5.15)$$

Notiamo che in (6.5.14) abbiamo solamente il flavour f e non due differenti f, f' , poiché in processi di DY, come abbiamo già notato in precedenza, i quark sono legati dal vertice di interazione, per cui il loro sapore non è libero.

Però siamo ancora un po' inclusivi sulla misura che facciamo, poiché potremmo ancora distinguere i momenti dei singoli leptoni anziché guardare il loro totale. Facciamo per questo un po' di cinematica leptonica. Prendiamo i momenti p_1 (quark entrante), p_2 (anti-quark entrante), p_3 (leptone uscente), p_4 (anti-leptone uscente) e scriviamo:

$$p_{1,2}^\mu = \frac{\sqrt{s}}{2}(x_{1,2}, 0, 0, \pm x_{1,2}) \quad (6.5.16)$$

$$p_3^\mu = p_\perp(\cosh y_3, \hat{n}, \sinh y_3) \quad (6.5.17)$$

$$p_4^\mu = p_\perp(\cosh y_4, \hat{n}, \sinh y_4) \quad (6.5.18)$$

in cui abbiamo introdotto le rapidità dei due leptoni y_3 e y_4 , oltre che la variabile di momento trasverso $p_\perp$, il quale è il modulo appunto della parte trasversa, che otteniamo proiettando lungo $\hat{n}$. Teniamo conto che i momenti trasversi (rispetto l'asse dell'urto) sono uguali ed opposti, ma non abbiamo vincoli su quelli longitudinali (facciamo un boost lungo quella direzione). Per come abbiamo parametrizzato abbiamo che i partoni sono nel loro sistema del centro di massa e si muovono lungo l'asse di riferimento $\hat{z}$, quando urtano producono i leptoni back-to-back (sempre nel centro di massa partoni), i quali avranno momento longitudinale lungo $\hat{z}$ e momento trasverso sul piano $\hat{n}$.

Le rapidità y_3 e y_4 rappresentano le componenti longitudinale e temporali dei momenti dei leptoni. y_i rappresenta la rapidità longitudinale data dal boost della particella i dal sistema di riferimento in cui ha momento longitudinale nullo (ricordiamo che $p_\perp$ è invariante per boost longitudinali). Quindi in qualche modo ci dicono quanto i momenti sono proiettati lungo l'asse $\hat{z}$. Ovviamente deve valere delle condizioni su $p_{3,4}^\mu$ per cui si sommino in modo da dare il vettore q^μ definito con (6.5.10).

Possiamo anche considerare l'invariante angolare (che è il momento trasferito al quadrato, ricordando (6.4.13)):

$$\hat{t} = -2p_3 \cdot p_1 = -\sqrt{s}p_\perp x_1 e^{-y_3} \quad (6.5.19)$$

e vedendo la conservazione del momento:

$$(p_1 + p_2)^0 \pm (p_1 + p_2)^3 = (p_3 + p_4)^0 \pm (p_3 + p_4)^3 \quad (6.5.20)$$

possiamo trovare (rispettivamente guardando il segno $+$ e il segno $-$):

$$x_1 = \frac{p_\perp}{\sqrt{s}} (e^{y_3} + e^{y_4}) \quad (6.5.21)$$

$$x_2 = \frac{p_\perp}{\sqrt{s}} (e^{-y_3} + e^{-y_4}). \quad (6.5.22)$$

A questo punto, avendo $p_\perp, y_3, y_4$ misurabili, oltre che con buone proprietà di trasformazioni (di Lorentz) per boost longitudinali, nello stato leptónico, siamo contenti e possiamo scrivere:

$$\frac{d^3\sigma}{dx_1 dx_2 d\hat{t}} = \frac{d^3\sigma}{dp_\perp dy_3 dy_4} \left| \frac{\partial(p_\perp, y_3, y_4)}{\partial(x_1, x_2, \hat{t})} \right| \quad (6.5.23)$$

in cui si può trovare lo jacobiano:

$$\left| \frac{\partial(x_1, x_2, \hat{t})}{\partial(p_\perp, y_3, y_4)} \right| = 2p_\perp x_1 x_2 = \frac{2p_\perp \hat{s}}{s} \quad (6.5.24)$$

che insieme al fatto che possiamo esplicitare:

$$\frac{d^3\sigma}{dx_1 dx_2 d\hat{t}} = f_f(x_1) f_{\bar{f}}(x_2) \frac{d\sigma}{d\hat{t}} \quad (6.5.25)$$

possiamo concludere:

$$\frac{d^3\sigma}{dp_\perp dy_3 dy_4} = 2p_\perp x_1 x_2 f_f(x_1) f_{\bar{f}}(x_2) \frac{d\sigma}{d\hat{t}} \quad (6.5.26)$$

in cui le variabili non sono indipendenti, ma valgono le relazioni trovate in precedenza:

$$x_{1,2} = x_{1,2}(p_\perp, y_3, y_4) \quad (6.5.27)$$

$$\hat{t} = \hat{t}(p_\perp, x_1(p_\perp, y_3, y_4)) \quad (6.5.28)$$

cioé abbiamo tutta l'espressione scritta in termini di $(p_\perp, y_3, y_4)$ che sono le variabili leptóniche. In aggiunta possiamo ricordare che:

$$\frac{d\sigma}{d\hat{t}} (f\bar{f} \rightarrow 3+4) = \frac{Q_f^2}{3} \frac{2\pi\alpha^2}{\hat{s}^2} \left(\frac{\hat{t}^2 + \hat{u}^2}{\hat{s}^2} \right) \quad (6.5.29)$$

sempre con il vincolo sulle variabili di Mandelstam (caso massless):

$$\hat{s} + \hat{t} + \hat{u} = 0$$

così abbiamo una sezione d'urto differenziabile in $\hat{t}$ che è scritta in termini dell'angolo tra il leptone (p_3) e l'asse $\hat{z}$.

Il punto di questa sezione, anche se i conti sono un po' abbozzati, siamo riusciti a vedere che possiamo scrivere la sezione d'urto differenziale in termini di variabili cinematiche dei singoli leptoni in stato finale. Ovviamente, anche se noi abbiamo fatto tutto a livello Born, si possono calcolare correzioni radiative di QCD e rendere più accurato il risultato.

Notiamo per concludere che possiamo ottenere una sezione d'urto differenziale rispetto tante variabili quante lo spazio delle fasi ci permette. Infatti, in questo caso abbiamo: uno spazio fasi a due corpi con $3 \cdot 2$ gradi di libertà, che si riducono a 2 se consideriamo le quattro delta di conservazione (che ci da 4 vincoli), ma che ricrescono a 4 se ci ricordiamo che abbiamo le variabili di Bjorken per le quali ci sono due integrali extra. Dunque, nel caso DY analizzato potremmo avere una sezione d'urto 4 volte differenziale; in effetti, nella pratica quello che si fa' è differenziare in $dp_{\perp}^2$ anzichè solo $dp_{\perp}$. In effetti si può notare che:

$$p_{\perp} dp_{\perp} = \frac{d^2 p_{\perp}}{2\pi}$$

e si ottiene:

$$\frac{d^4}{d^2 p_{\perp} dy_3 dy_4} = x_1 f_1(x_1) x_2 f_2(x_2) \frac{1}{\pi} \frac{d\sigma}{d\hat{t}}. \quad (6.5.30)$$

Ad un certo punto della nostra trattazione abbiamo scritto $d\sigma/(dx_1 dx_2)$ e quello era il massimo che avremmo potuto ottenere poiché avevamo solamente $6 - 4$ gradi di libertà dallo spazio fasi.

6.5.2 Produzione di jet (cenni)

In questa rapida sezione vogliamo solo accennare al fatto che ad LHC si studiano anche processi per cui:

$$\sigma(p + p \rightarrow j + j + X) \quad (6.5.31)$$

la cui grande differenza, rispetto al processo $e^+ e^- \rightarrow \text{adroni}$, è che ora contribuiscono molti più processi, visto che possiamo estrarre molta più roba. I processi partonici che partecipano sono:

$$q\bar{q} \longrightarrow q\bar{q} \quad \begin{array}{c} \begin{array}{c} q \\ \swarrow \\ \text{---} \end{array} \begin{array}{c} \text{---} \\ \searrow \\ q \end{array} \\ \begin{array}{c} \bar{q} \\ \swarrow \\ \text{---} \end{array} \begin{array}{c} \text{---} \\ \searrow \\ \bar{q} \end{array} \end{array} \quad (6.5.32)$$

$$q\bar{q} \longrightarrow q'\bar{q}' \quad \begin{array}{c} \begin{array}{c} q \\ \swarrow \\ \text{---} \end{array} \begin{array}{c} \text{---} \\ \searrow \\ q' \end{array} \\ \begin{array}{c} \bar{q} \\ \swarrow \\ \text{---} \end{array} \begin{array}{c} \text{---} \\ \searrow \\ \bar{q}' \end{array} \end{array} \quad (6.5.33)$$

$$qq \longrightarrow qq \quad \begin{array}{c} q \searrow \quad \nearrow q \\ | \\ q \searrow \quad \nearrow q \\ | \\ q \searrow \quad \nearrow q \end{array} \quad (6.5.34)$$

$$qg \longrightarrow qg \quad \begin{array}{c} q \searrow \quad \nearrow q \\ | \\ g \text{---} \text{---} \text{---} g \end{array} \quad (6.5.35)$$

$$q\bar{q} \longrightarrow gg \quad \begin{array}{c} q \searrow \quad \nearrow g \\ | \\ \bar{q} \searrow \quad \nearrow g \end{array} \quad (6.5.36)$$

$$gg \longrightarrow q\bar{q} \quad \begin{array}{c} g \text{---} \text{---} \text{---} q \\ | \\ g \text{---} \text{---} \text{---} \bar{q} \end{array} \quad (6.5.37)$$

$$(6.5.38)$$

i quali devono essere tutti sommati, dopo essere moltiplicati per le rispettive PDF, poiché non risolvibili sperimentalmente separatamente (ad esempio non distinguiamo un $q\bar{q}$ prodotto da una coppia $q\bar{q}$ o gg). Sostanzialmente i conti in questi processi si complicano poiché non sappiamo cosa entra nel processo, infatti mettiamo le PDF, ma non sappiamo nemmeno nel jet cosa c'è andato.

Nota. Calcolare le ampiezze di processi di questo tipo potrebbe essere un esercizio di esame.

Capitolo 7

Anomalie

Vediamo in questo capitolo conclusivo le anomalie delle teorie di campo. Cominceremo con un discorso generico per poi arrivare a come nel Modello Standard esse si cancellino e come la teoria resti consistente. Si può seguire la trattazione presente al capitolo §12.4.5 sull'Itzykson, Zuber There is a sort of mythology that grows up about what happened, which is different from what really did happen. Zuber.

Per prima cosa dobbiamo definire:

Anomalie := simmetrie classiche (dunque a livello di $\mathcal{L}$) che vengono rotte a livello quantistico (dunque a livello di loop).

Se le simmetrie che si rompono sono accidentali (in qualche modo) o globali, il problema non si pone e la teoria funziona, però ciò diventa problematico se la simmetria che si rompe è di gauge. Infatti, se noi basiamo tutto il nostro modello su tali simmetrie, nel momento in cui esse si rompono perdiamo tutto.

Sotto un'altra veste possiamo dire che a livello di path integral:

$$Z = \int [\mathcal{D}\phi] e^{i\frac{S}{\hbar}} \quad (7.0.1)$$

un'anomalia è tale per cui S (dunque $\mathcal{L}$) è invariante, ma $[\mathcal{D}\phi]$ no!

7.1 Anomalia assiale

Vediamo in questa sezione un esempio di anomalia leggermente semplificato, ma che ci permette di fare un discorso completo.

Prendiamo la teoria definita dalla seguente densità di Lagrangiana:

$$\mathcal{L} = \bar{\psi}(i\not{D} - m)\psi \quad (7.1.1)$$

e definiamo le seguenti correnti:

$$J_\mu^V = \bar{\psi}\gamma_\mu\psi \quad (\text{corrente vettoriale}) \quad (7.1.2)$$

$$J_\mu^A = \bar{\psi}\gamma_\mu\gamma^5\psi \quad (\text{corrente assiale}) \quad (7.1.3)$$

$$J_5 = \bar{\psi}\gamma^5\psi \quad (7.1.4)$$

per le quali si può mostrare (utilizzando le equazioni del moto) che:

$$\partial_\mu J_V^\mu = 0 \quad (7.1.5)$$

$$\partial_\mu J_A^\mu = 2 \text{Im}\{J_5\} \quad (7.1.6)$$

in cui (7.1.5) ce lo aspettiamo poiché facendo $\psi \rightarrow e^{i\alpha}\psi$, vediamo che le componenti left e right trasformano allo stesso modo e dunque $\mathcal{L}$ scritta con $\psi_{L,R}$ è invariante; allo stesso tempo la relazione (7.1.6) la possiamo vedere facendo $\psi \rightarrow e^{i\alpha\gamma^5}\psi$, da cui vediamo che $\psi_{L,R}$ trasformano diversamente e, se c'è il termine di massa, $\mathcal{L}$ non è invariante (poiché $m\bar{\psi}\psi$ mischia i pezzi).

Le conservazioni (o meno) delle due correnti hanno due significati profondamente differenti. La conservazione di J_V^μ è cruciale quando scriviamo la QED con $U(1)$, visto che la trasformazione $\psi \rightarrow e^{i\alpha}\psi$ è scritta a livello locale ed $U(1)$ è la simmetria di gauge del modello.

Allo stesso tempo anche se $\partial_\mu J_A^\mu = 0$ venisse a cadere, pure nel limite massless, non sarebbe un problema, poiché $\psi \rightarrow e^{i\alpha\gamma^5}\psi$ non è una simmetria di gauge.

Proviamo ora a verificare se $\partial_\mu J_V^\mu = \partial_\mu J_A^\mu = 0$ per la teoria di QED (nel caso massless) a livello quantistico (ovvero a loop). Per farlo, ovviamente prendiamo la QED interagente, cioè con il ben noto vertice di interazione, e ci serve introdurre i seguenti tensori:

$$T^{\mu\nu\rho}(k_1, k_2) \equiv i \int d^4x_1 d^4x_2 e^{ik_1 \cdot x_1} e^{ik_2 \cdot x_2} \times \\ \times \langle 0 | T[J_V^\mu(x_1) J_V^\nu(x_2) J_A^\rho(0)] | 0 \rangle \quad (7.1.7)$$

$$T^{\mu\nu}(k_1, k_2) \equiv i \int d^4x_1 d^4x_2 e^{ik_1 \cdot x_1} e^{ik_2 \cdot x_2} \times \\ \times \langle 0 | T[J_V^\mu(x_1) J_V^\nu(x_2) J_5(0)] | 0 \rangle. \quad (7.1.8)$$

Sembra che ci siamo un po' andati ad impelagare in roba complicata, ma in realtà questi tensori ci aiutano nella verifica che vogliamo fare. Infatti, possiamo semplificare la scrittura di $T^{\mu\nu\rho}$ poiché il termine $\langle 0 | J_V^\mu(x_1) J_V^\nu(x_2) | 0 \rangle$ comporta la creazione di un fotone nello stato finale, per cui:

$$T^{\mu\nu\rho}(k_1, k_2) \propto \langle \gamma(k_1, \epsilon_1) \gamma(k_2, \epsilon_2) | J_A^\rho(0) | 0 \rangle \quad (7.1.9)$$

$$= 2ie^2 \epsilon_1^*(k_1)^\mu \epsilon_2^*(k_2)^\nu T_{\mu\nu}^\rho(k_1, k_2). \quad (7.1.10)$$

Dunque abbiamo la struttura ampiezza $T^{\mu\nu\rho}$ espressa graficamente sotto forma di funzioni di correlazione contratte che diagrammaticamente sappiamo calcolare tramite i relativi grafici a loop (il bloob presente nella figura rappresenta, come al solito, tutte le possibili correzioni a loop che possiamo fare). La situazione grafica che abbiamo è:

$$T^{\mu\nu\rho} = J_A^\rho \xrightarrow{q = k_1 + k_2} \text{bloob} \begin{matrix} \nearrow k_1, \mu \\ \searrow k_2, \nu \end{matrix} \quad (7.1.11)$$

che ci rappresenta la scrittura (7.1.10). Effettivamente quello che abbiamo appena concluso può sembrare un po' strano, ma se ci pensiamo un attimo è l'unica cosa che può succedere; infatti, utilizzando LSZ è facile rendersi conto che l'azione delle correnti vettoriali $(\bar{\psi}\gamma^\mu\psi)$ sul vuoto non può fare altro che produrre dei fotoni in stato finale.

Cominciamo a vedere come si traducono le relazioni $\partial_\mu J_V^\mu = \partial_\mu J_A^\mu = 0$, a livello classico, in termini di $T^{\mu\nu\rho}$ così da poter verificare se esse rimangono valide anche a livello quantistico (per il quale è più semplice fare i conti con $T^{\mu\nu\rho}$).

Possiamo vedere che la conservazione della corrente vettoriale $\partial_\mu J_V^\mu = 0$ a livello di $T^{\mu\nu\rho}$ si traduce in:

$$k_{1\mu} T^{\mu\nu\rho}(k_1, k_2) = 0. \quad (7.1.12)$$

Infatti, inserendo l'impulso all'interno dell'integrale nello spazio delle coordinate si ha:

$$k_{1\mu} T^{\mu\nu\rho} = i \int d^4x_1 d^4x_2 k_{1\mu} e^{ik_1 \cdot x_1} (\dots) \quad (7.1.13)$$

in cui la parentesi racchiude tutto ciò che non risiede nell'esponenziale $e^{ik_1 \cdot x_1}$. Riscrivendo il termine $k_{1\mu} e^{ik_1 \cdot x_1}$ come l'azione di una derivata sull'esponenziale (o meglio riscrivendo $k_{1\mu}$ in termini di operatore differenziale):

$$k_{1\mu} T^{\mu\nu\rho} = i \int d^4x_1 d^4x_2 \left(-i \frac{\partial}{\partial x_1^\mu} \right) e^{ik_1 \cdot x_1} (\dots) \quad (7.1.14)$$

e integrando per parti (sottointendendo termini di bordo nullo):

$$k_{1\mu} T^{\mu\nu\rho} = i \int d^4x_1 d^4x_2 e^{ik_1 \cdot x_1} i \frac{\partial}{\partial x_1^\mu} (\dots) \quad (7.1.15)$$

dove l'azione di quest'ultimo operatore differenziale prenderà esplicitamente in considerazione solo la corrente vettoriale $J_V^\mu(x_1)$. Abbiamo dunque (fidandoci del fatto che i termini aggiuntivi derivanti dalle derivate delle Θ

nel prodotto temporalmente ordinato si ricombinino in modo da dare come effetto solo una derivata su $J_V^\mu(x_1)$:

$$k_{1\mu} T^{\mu\nu\rho} = i \int d^4x_1 d^4x_2 e^{ik_1 \cdot x_1} \underbrace{\langle 0 | \frac{\partial}{\partial x_1^\mu} J_V^\mu(x_1) \dots | 0 \rangle}_{=0} \quad (7.1.16)$$

$$\implies k_{1\mu} T^{\mu\nu\rho} = 0. \quad (7.1.17)$$

Analogamente si può vedere che:

$$k_{2\nu} T^{\mu\nu\rho}(k_1, k_2) = 0 \quad (7.1.18)$$

$$(k_1 + k_2)_\rho T^{\mu\nu\rho}(k_1, k_2) = q_\rho T^{\mu\nu\rho}(k_1, k_2) = 2m T^{\mu\nu}. \quad (7.1.19)$$

Si noti che quest'ultimo termine si annulla se consideriamo il caso $m = 0$.

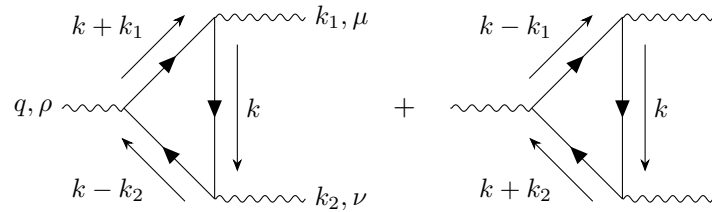
A questo punto possiamo calcolare $T^{\mu\nu\rho}$ ad un loop e vedere se le relazioni (7.1.17), (7.1.18) e (7.1.19) sono mantenute.

Osservazione. Prima di proseguire con i conti ad 1-loop possiamo notare che le relazioni (7.1.17) e (7.1.18) provengono direttamente dalla scrittura (7.1.11), poiché se contraiamo i fotoni, che ricordiamo essere on-shell (visto che abbiamo delle polarizzazioni esplicite), con i loro quadri-impulsi otteniamo effettivamente le identità di Ward (di QED). Dunque, in (7.1.17) e (7.1.18) sono racchiuse le identità di Ward.

7.1.1 Calcolo ad 1-loop

Effettuiamo quindi il calcolo di $k_{1\mu} T^{\mu\nu\rho}$, $k_{2\nu} T^{\mu\nu\rho}$ e $q_\rho T^{\mu\nu\rho}$ a 1-loop per verificare se sono uguali a zero o meno.

Il tensore $T^{\mu\nu\rho}$ a 1-loop è dato dalla somma dei due diagrammi a triangolo (canale diretto e incrociato):



$$+ \quad (7.1.20)$$

in cui facciamo attenzione al fatto che nel vertice di sinistra, in cui arriva la corrente assiale (indice ρ) avremo un contributo di vertice $\gamma^\rho \gamma^5$, mentre negli altri due vertici ci saranno solamente delle γ^μ e γ^ν attaccate alle correnti vettoriali.

Possiamo per semplicità definire le quantità cinematiche associate ai denominatori dei propagatori fermionici nel loop:

$$a \equiv k + k_1 - m \quad , \quad b \equiv k - k_2 - m \quad (7.1.21)$$

$$c \equiv k - m \quad , \quad A \equiv k + k_2 - m \quad (7.1.22)$$

$$B \equiv k - k_1 - m. \quad (7.1.23)$$

e si può verificare facilmente che la traccia associata ai diagrammi risulta:

$$T^{\mu\nu\rho} = - \int [dk] \operatorname{tr} \left\{ \frac{1}{a} \gamma^\rho \gamma^5 \frac{1}{b} \gamma^\nu \frac{1}{c} \gamma^\mu + \frac{1}{A} \gamma^\rho \gamma^5 \frac{1}{B} \gamma^\mu \frac{1}{c} \gamma^\nu \right\}, \quad (7.1.24)$$

dove la misura d'integrazione sul quadrimpulso del loop è definita come $[dk] \equiv d^4k / (2\pi)^4$.

Dato che vogliamo calcolare $k_{1(2)\mu(\nu)} T^{\mu\nu\rho}$, ci rendiamo conto che $T^{\mu\nu\rho}$ è linearmente divergente nell' UV . Dunque, regolarizziamo con il formalismo di *Pauli-Villars*¹:

$$[T^{\mu\nu\rho}]_{\text{REG}} = T^{\mu\nu\rho} - T^{\mu\nu\rho} \Big|_{m \rightarrow M} \quad (7.1.25)$$

e vogliamo verificare se:

$$k_{1\mu} T^{\mu\nu\rho} \Big|_{\text{REG}} = 0 \quad (7.1.26)$$

ad 1-loop. Notiamo che:

$$k_1 = a - c = c - B \quad (7.1.27)$$

e ora, visto che abbiamo regolarizzato l'integrale, possiamo contrarre l'indice μ e cambiare variabile così da ottenere:

$$\begin{aligned} k_{1\mu} T^{\mu\nu\rho} \Big|_{\text{REG}} &= - \int [dk] \operatorname{tr} \left\{ \gamma^\rho \gamma^5 \frac{1}{b} \gamma^\nu \frac{1}{c} - \frac{1}{a} \gamma^\rho \gamma^5 \frac{1}{b} \gamma^\nu + \right. \\ &\quad \left. + \frac{1}{A} \gamma^\rho \gamma^5 \frac{1}{B} \gamma^\nu - \frac{1}{A} \gamma^\rho \gamma^5 \frac{1}{c} \gamma^\nu \right\} - (m \rightarrow M) \end{aligned} \quad (7.1.28)$$

dove di nuovo, ora che l'integrale è regolarizzato, possiamo cambiare variabile all'interno della traccia perturbativa per shiftare gli impulsi del loop (che non

¹Non è oggetto del corso, ma si possono leggere le note del corso di *Advanced QFT* (o forse meglio le dispense del canale telegram del corso di *Complementi di Teoria di Campo*). Il metodo di Pauli-Villars è uno dei molti modi che si hanno di regolarizzare un integrale divergente, il cui modo di procedere è il seguente: prendiamo un integrale divergente e ci sottraiamo la stessa cosa, ma con la massa cut-off M sostituita ad m .

coinvolgono la massa, per cui anche il termine con $m \rightarrow M$ subisce gli stessi cambiamenti):

$$\gamma^\rho \gamma^5 \frac{1}{b} \gamma^\nu \frac{1}{c} \xrightarrow[\substack{(b \rightarrow c) \\ (c \rightarrow A)}]{k^\mu \rightarrow k^\mu + k_2^\mu} \frac{1}{A} \gamma^\rho \gamma^5 \frac{1}{c} \gamma^\nu$$

$$\frac{1}{a} \gamma^\rho \gamma^5 \frac{1}{b} \gamma^\nu \xrightarrow[\substack{(a \rightarrow A) \\ (b \rightarrow B)}]{k^\mu \rightarrow k^\mu - k_1^\mu + k_2^\mu} \frac{1}{A} \gamma^\rho \gamma^5 \frac{1}{B} \gamma^\nu$$

così si cancella tutto (sostituendo i pezzi opportunamente) e otteniamo:

$$k_{1\mu} T^{\mu\nu\rho} \Big|_{\text{REG}} = 0. \quad (7.1.29)$$

Analogamente si mostra che:

$$k_{2\nu} T^{\mu\nu\rho} \Big|_{\text{REG}} = 0. \quad (7.1.30)$$

Possiamo infine vedere la contrazione $q_\rho T^{\mu\nu\rho} \Big|_{\text{REG}}$, in cui scriviamo per convenienza le identità:

$$\begin{aligned} \not{q} \gamma_5 &= 2m \gamma_5 + a \gamma_5 + \gamma_5 b \\ &= 2m \gamma_5 + A \gamma_5 + \gamma_5 B, \end{aligned} \quad (7.1.31)$$

per arrivare alla scomposizione finale:

$$q_\rho T_{\text{REG}}^{\mu\nu\rho} = (2m T^{\mu\nu})_{\text{REG}} + (R^{\mu\nu})_{\text{REG}}. \quad (7.1.32)$$

Dove si può dimostrare:

$$\begin{aligned} (2m T^{\mu\nu})_{\text{REG}} &= -2m \int [dk] \text{Tr} \left\{ \frac{1}{a} \gamma_5 \frac{1}{b} \gamma^\nu \frac{1}{c} \gamma^\mu + \frac{1}{A} \gamma_5 \frac{1}{B} \gamma^\mu \frac{1}{c} \gamma^\nu \right\} - \\ &\quad - (m \rightarrow M) \end{aligned} \quad (7.1.33)$$

$$\begin{aligned} (R^{\mu\nu})_{\text{REG}} &= \int [dk] \text{Tr} \left\{ -\gamma^5 \gamma^\mu \frac{1}{b} \gamma^\nu \frac{1}{c} - \gamma^5 \gamma^\nu \frac{1}{B} \gamma^\mu \frac{1}{c} + \right. \\ &\quad \left. + \gamma^5 \gamma^\nu \frac{1}{c} \gamma^\mu \frac{1}{a} + \gamma^5 \gamma^\mu \frac{1}{c} \gamma^\nu \frac{1}{A} \right\} - (m \rightarrow M). \end{aligned} \quad (7.1.34)$$

Con cambi di variabile simili a quelli utilizzati in precedenza si può mostrare che:

$$(R^{\mu\nu})_{\text{REG}} = 0. \quad (7.1.35)$$

Dunque, dobbiamo solo controllare che $(2mT^{\mu\nu})_{\text{REG}} = 0$ nel caso di $m = 0$ sia verificato. Sfruttando l'integrazione con i parametri di Feynman si trova:

$$[q_\rho T^{\mu\nu\rho}]_{\text{REG}} = \frac{1}{\pi} \epsilon_{\mu\nu\rho\sigma} k_1^\rho k_2^\sigma \int_0^1 dx \int_0^{1-x} dy \left[\frac{m^2}{m^2 - q^2 xy} - \frac{M^2}{M^2 - q^2 xy} \right] \quad (7.1.36)$$

che rimuovendo il regolatore, ovvero mandando $M \rightarrow \infty$, e mettendoci nel caso massless con $m \rightarrow 0$, allora si ottiene:

$$\lim_{\substack{M \rightarrow \infty \\ m \rightarrow 0}} [q_\rho T^{\mu\nu\rho}]_{\text{REG}} = -\frac{1}{\pi} \epsilon_{\mu\nu\rho\sigma} k_1^\rho k_2^\sigma \frac{1}{2} \neq 0. \quad (7.1.37)$$

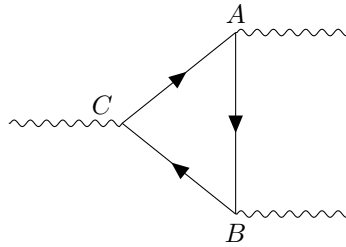
Questo risultato esprime la celebre *Anomalia Assiale*² (o anomalia di Adler-Bell-Jackiw).

Quindi la condizione $\partial_\mu J_A^\mu = 0$, tradotta con $T^{\mu\nu\rho}$, è rispettata classicamente, ma non a loop. Questo si può dimostrare succedere a tutti i loop; cioè non ci sono correzioni a loop successive.

Dunque, abbiamo visto che le correnti assiali vengono rotte a livello quantistico per il modello semplice che abbiamo esaminato. Però, rompere le correnti assiali non è un problema per la QED, visto che la simmetria $U(1)$ non comporta alcuna corrente (di gauge) assiale conservata. Il problema però sorge per il Modello Standard, che è una teoria chirale in cui ci sono correnti assiali di gauge.

7.2 Anomalie nel SM e loro cancellazioni

Nel *Modello Standard* (SM), come già accennato abbiamo dei problemi per l'anomalia assiale. Per far fronte a ciò consideriamo il vertice d'interazione a loop triangolare:



in cui in ogni vertice inseriamo i generatori del gruppo non abeliano, generico, attaccati alle rispettive correnti assiali e vettoriali. Ci poniamo nel caso più generico possibile, ma possiamo vedere che ciò è effettivamente quello che succede con il gruppo $U(1)$, per il quale a seconda della corrente che è

²La chiamiamo *assiale* poiché si rompe la simmetria quando guardiamo la contrazione dei momenti con la corrente assiale J_A^ρ .

attaccata ad un vertice abbiamo un generatore $\mathbb{1}$ (ed un γ^5 per la corrente assiale).

Di conseguenza, nella struttura dell'ampiezza quantistica $T^{\mu\nu\rho}$ compariranno i pezzi gruppali accoppiati tramite la traccia dei generatori (per ora non specificati e che si trovano nella rappresentazione in cui stanno i fermioni):

$$T^{\mu\nu\rho} \propto \text{Tr} \left\{ \left\{ T^a, T^b \right\} T^c \right\}. \quad (7.2.1)$$

Il termine cinetico in $T^{\mu\nu\rho}$ è sempre lo stesso, però a seconda dei fermioni che attacchiamo ai vertici, dunque ai generatori del gruppo che inseriamo, avremo dei contributi alla traccia differenti. Notiamo che l'anticommutatore viene dal fatto che oltre al diagramma disegnato abbiamo anche il diagramma in cui il loop gira al contrario ed A e B sono scambiati.

Si può vedere che, nonostante sia presente l'anomalia assiale, una volta sommato su tutte le possibili particelle che girano nel loop, dunque prendendo tutte le possibili scelte di generatori nei vertici (così da mettere tutti i possibili fermioni attaccati), allora l'anomalia assiale si cura. Dobbiamo quindi verificare che la somma complessiva sulle specie fermioniche sia nulla:

$$\sum_{\substack{\text{specie} \\ \text{fermionica}}} \text{Tr} \left\{ T^a T^b T^c \right\} = 0. \quad (7.2.2)$$

Notiamo prima che il fattore di forma simmetrizzato $\text{Tr} \left\{ \left\{ T^a, T^b \right\} T^c \right\}$ è completamente simmetrico nello scambio di tutti gli indici gruppali; dunque, possiamo considerare solo dei casi prototipici senza vedere tutte le possibili permutazioni simmetriche.

7.2.1 Studio del gruppo $SU(2) \times U(1)$

Consideriamo in prima istanza il sottogruppo di gauge chirale $SU(2) \times U(1)$. Le possibili combinazioni dei generatori danno luogo alle seguenti quattro strutture di traccia da analizzare:

$$\sum_{\text{specie}} \text{Tr} \left\{ \left\{ \tau^a, \tau^b \right\} \tau^c \right\} \quad (7.2.3)$$

$$\sum_{\text{specie}} \text{Tr} \left\{ \left\{ \tau^a, \tau^b \right\} Y \right\} \quad (7.2.4)$$

$$\sum_{\text{specie}} \text{Tr} \left\{ \{Y, Y\} \tau^c \right\} \quad (7.2.5)$$

$$\sum_{\text{specie}} \text{Tr} \{Y^3\} \quad (7.2.6)$$

dove τ^i sono i generatori del gruppo $SU(2)$, mentre Y denota l'operatore associato all'*ipercarica*.

Vediamo esplicitamente il comportamento delle prime due relazioni algebriche.

Per la prima traccia si ottiene immediatamente l'annullamento identico dell'anomalia, valido anche senza sommare sulle diverse specie:

$$(7.2.3) = \sum_{\text{specie}} 2\delta^{ab} \underbrace{\text{Tr}\{\tau^c\}}_{=0} = 0. \quad (7.2.7)$$

Per la seconda traccia, i generatori delle matrici di Pauli $\{\tau^a, \tau^b\}$ selezionano esclusivamente i fermioni *left-handed* (sinistrorsi) poiché i campi *right-handed* sono singoletti sotto $SU(2)$. Sviluppando esplicitamente la somma sulle famiglie di quark (n_q) e di leptoni (n_ℓ), tenendo conto del numero di colori del rispettivo doppietto (rispettivamente 3 ed 1) e dei valori dell'ipercarica Y_L (pari a $\frac{1}{3}$ per i quark e -1 per i leptoni), si ottiene:

$$\begin{aligned} (7.2.4) &= \sum_{\text{specie}} 2\delta^{ab} \text{tr}\{Y_L\} \\ &= \left[n_q \cdot 3 \cdot 2 \cdot \frac{1}{3} + n_\ell \cdot 1 \cdot 2 \cdot (-1) \right] \cdot 2\delta^{ab} \\ &= 2\delta^{ab} [2n_q - 2n_\ell]. \end{aligned} \quad (7.2.8)$$

Dato che nel Modello Standard per ogni generazione il numero di famiglie di quark e di leptoni coincide ($n_q = n_\ell$), questa combinazione si cancella esattamente eliminando il rispettivo contributo anomalo.

Questa relazione potrebbe essere utilizzata anche al contrario, ovvero potremmo dire che affinché si cancelli questo pezzo di anomalia abbiamo bisogno che il numero di famiglie leptoniche sia uguale al numero di famiglie di quark.

Analizziamo ora la terza condizione:

$$(7.2.5) = \sum_{\text{specie}} Y_{\text{specie}}^2 \text{tr}\{\tau^c\} = 0 \quad (7.2.9)$$

che si annulla identicamente in quanto le matrici di Pauli sono a traccia nulla.

Vediamo la quarta relazione, in cui non separiamo le parti left visto che non ci sono i generatori di $SU(2)$, e per cui è necessaria la decomposizione chirale per la corrente assiale:

$$\bar{\psi}\gamma^\mu\gamma^5\psi = \bar{\psi}\gamma^\mu\frac{1}{2}(1+\gamma^5)\psi - \bar{\psi}\gamma^\mu\frac{1}{2}(1-\gamma^5)\psi, \quad (7.2.10)$$

da cui possiamo vedere che le due chiralità contribuiscono con segno opposto all'anomalia. Dunque, in (7.2.6) separiamo i contributi delle componenti

sinistrorse (*left*) e destrorse (*right*):

$$(7.2.6) = \sum_{\text{specie}} \text{tr}(Y^3) = \sum_{\text{specie}} [\text{tr}(Y_L^3) - \text{tr}(Y_R^3)]. \quad (7.2.11)$$

Se sostituiamo esplicitamente i valori delle ipercariche *right-handed* per i singoletti del modello ($Y(u_R) = 4/3$, $Y(d_R) = -2/3$, $Y(e_R) = -2$ e $Y(\nu_R) = 0$), il calcolo esplicito della traccia fornisce (nella prima riga il contributo left e nella seconda il contributo right):

$$(7.2.6) = n_q \cdot 3 \cdot 2 \cdot \left(\frac{1}{3}\right)^3 + n_\ell \cdot 1 \cdot 2 \cdot (-1)^3 - \\ - n_q \cdot 3 \cdot \left[\left(\frac{4}{3}\right)^3 + \left(-\frac{2}{3}\right)^3\right] - n_\ell \cdot 1 \cdot [(-2)^3] \quad (7.2.12)$$

$$= -6(n_q - n_\ell). \quad (7.2.13)$$

Di conseguenza, anche la condizione (7.2.6) si annulla se e solo se $n_q = n_\ell$.

Dunque per il gruppo $SU(2) \times U(1)$ non contiene delle anomalie, poiché esse vengono cancellate se sommiamo su tutti i flavour circolanti in un loop.

7.2.2 Studio del gruppo $SU(3)$

Prendiamo ora in considerazione l'estensione al gruppo di colore $SU(3)$ della QCD. Si presentano i seguenti vari casi per le tracce dei generatori (denotati con t^a): la prima combinazione coinvolge due generatori di $SU(2)$ e uno di $SU(3)$, annullandosi per la traccia di quest'ultimo:

$$1) \quad \sum_{\text{specie}} \text{Tr}\left\{\left\{\tau^a, \tau^b\right\}t^c\right\} = 0. \quad (7.2.14)$$

La seconda combinazione si riduce alla traccia singola del generatore di colore, in quanto l'ipercarica commuta ed è proporzionale all'identità sullo spazio di colore:

$$2) \quad \sum_{\text{specie}} \text{Tr}\left\{\left\{\tau^a, Y\right\}t^c\right\} = \sum_{\text{specie}} \text{Tr}\{Y\tau^a t^c\} = 0. \quad (7.2.15)$$

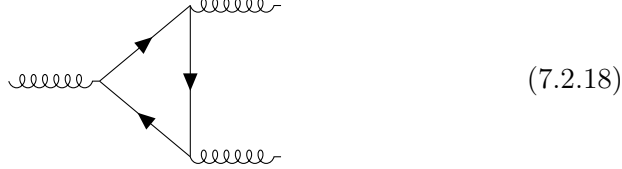
La terza combinazione si annulla poiché la traccia sul doppietto di $SU(2)$ genera una delta di Kronecker moltiplicata per la traccia del singolo generatore di $SU(3)$:

$$3) \quad \sum_{\text{specie}} \text{Tr}\left\{\left\{t^a, t^b\right\}\tau^c\right\} = T_R \delta^{ab} \text{Tr}\{\tau^c\} = 0. \quad (7.2.16)$$

La quarta condizione coinvolge tre generatori di colore puramente vettoriali:

$$4) \quad \sum_{\text{specie}} \text{Tr}\left\{\left\{t^a, t^b\right\}t^c\right\} = 0 \quad (7.2.17)$$

la quale non genera anomalie di gauge. Non c'è anomalia poiché non abbiamo il relativo vertice assiale di gauge accoppiato ai gluoni:



che è il diagramma triangolare misto con gluoni in cui la corrente assiale esterna non è una corrente di gauge (essendo un puro diagramma di QCD).

Infine, l'ultima configurazione con due generatori di $SU(3)$ e l'ipercarica si esprime tramite T_R come:

$$5) \quad \sum_{\text{specie}} \text{Tr} \left\{ \left\{ t^a, t^b \right\} Y \right\} = T_R \delta^{ab} \sum_{\text{specie}} [\text{tr}(Y_L) - \text{tr}(Y_R)]. \quad (7.2.19)$$

Possiamo vedere che contribuiscono solamente i quark, essendo i leptoni dei singoletti di colore, per cui esplicitando le cariche:

$$5) \quad = -T_R \delta^{ab} n_q \left[2 \cdot \left(-\frac{1}{3} \right) + \frac{4}{3} - \frac{2}{3} \right] = 0. \quad (7.2.20)$$

Dunque, avendo esaurito tutte le possibilità e potendo dimostrare che ciò valga a tutti i loop, allora possiamo dire che per tutto il Modello Standard $SU(3) \times SU(2) \times U(1)$ l'anomalia assiale, seppur presente, si riorganizza in modo tale da cancellarsi e da rendere la teoria consistente.

Appendici

Appendice A

Relazioni utili per accoppiamenti elettrodeboli

Riporto in questa Appendice delle relazioni utili per riscrivere la lagrangiana massless per i campi di gauge nella forma:

$$\mathcal{L}_{YM} = \mathcal{L}_{YM}^{(2)} + \mathcal{L}_{YM}^{(3)} + \mathcal{L}_{YM}^{(4)}$$

con il termine cinetico, di accoppiamento a 3 e 4 campi. Per arrivare a tale risultato si dovrebbe partire dalla lagrangiana di Yang-Mills:

$$\mathcal{L}_{YM} = -\frac{1}{4}B_{\mu\nu}B^{\mu\nu} - \frac{1}{4}W_{\mu\nu}^a W_a^{\mu\nu}$$

in cui:

$$\begin{aligned} B_{\mu\nu} &= \partial_\mu B_\nu - \partial_\nu B_\mu \\ W_{\mu\nu}^a &= \partial_\mu W_\nu^a - \partial_\nu W_\mu^a + g\epsilon_{bc}^a W_\mu^b W_\nu^c \end{aligned}$$

e, dopo aver espanso gli indici a, b, c , operare le sostituzioni:

$$\begin{aligned} W_\mu^1 &= \frac{W_\mu^+ + W_\mu^-}{\sqrt{2}} \\ W_\mu^2 &= i \frac{W_\mu^+ - W_\mu^-}{\sqrt{2}} \\ W_\mu^a &= \sin \theta_W A_\mu + \cos \theta_W Z_\mu \\ B_\mu &= \cos \theta_W A_\mu - \sin \theta_W Z_\mu. \end{aligned}$$

Il problema di tutto ciò è che quando si espandono gli indici contratti con il tensore di Levi-Civita ϵ_{abc} i conti possono risultare macchinosi. Per

questo, se qualcuno volesse avventurarsi, possono risultare utili le relazioni:

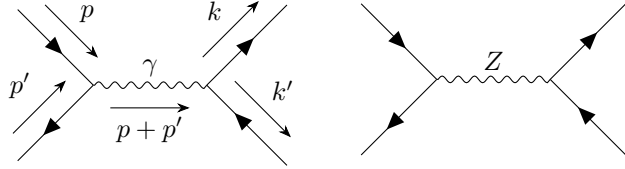
$$\begin{aligned}
\varepsilon_{1bc} W_b^\mu W_c^\nu &= W_2^\mu W_3^\nu - W_3^\mu W_2^\nu \\
&= \frac{i}{\sqrt{2}} \left[\sin \theta_W (W_+^\mu A^\nu - W_-^\mu A^\nu) + \cos \theta_W (W_+^\mu Z^\nu - W_-^\mu Z^\nu) - \right. \\
&\quad \left. - (\mu \leftrightarrow \nu) \right] \\
\varepsilon_{2bc} W_b^\mu W_c^\nu &= W_3^\mu W_1^\nu - W_1^\mu W_3^\nu \\
&= -\frac{1}{\sqrt{2}} \left[\sin \theta_W (W_+^\mu A^\nu + W_-^\mu A^\nu) + \cos \theta_W (W_+^\mu Z^\nu + W_-^\mu Z^\nu) - \right. \\
&\quad \left. - (\mu \leftrightarrow \nu) \right] \\
\varepsilon_{3bc} W_b^\mu W_c^\nu &= W_1^\mu W_2^\nu - W_2^\mu W_1^\nu \\
&= -i [W_+^\mu W_-^\nu - W_-^\mu W_+^\nu] \\
W_1^{\mu\nu} &= \frac{1}{\sqrt{2}} \left\{ W_+^{\mu\nu} + W_-^{\mu\nu} + ig \left[\sin \theta_W (W_+^\mu A^\nu - W_-^\mu A^\nu) + \right. \right. \\
&\quad \left. \left. + \cos \theta_W (W_+^\mu Z^\nu - W_-^\mu Z^\nu) - (\mu \leftrightarrow \nu) \right] \right\} \\
W_2^{\mu\nu} &= \frac{i}{\sqrt{2}} \left\{ W_+^{\mu\nu} - W_-^{\mu\nu} + ig \left[\sin \theta_W (W_+^\mu A^\nu + W_-^\mu A^\nu) + \right. \right. \\
&\quad \left. \left. + \cos \theta_W (W_+^\mu Z^\nu + W_-^\mu Z^\nu) - (\mu \leftrightarrow \nu) \right] \right\} \\
W_3^{\mu\nu} &= \sin \theta_W F^{\mu\nu} + \cos \theta_W Z^{\mu\nu} - ig [W_+^\mu W_-^\nu - W_-^\mu W_+^\nu] \\
B^{\mu\nu} &= \cos \theta_W F^{\mu\nu} - \sin \theta_W Z^{\mu\nu}.
\end{aligned}$$

Appendice B

Scattering $e^+e^- \rightarrow f\bar{f}$ massless

Vediamo in questa Appendice il calcolo al Leading Order dello scattering visto nella sezione §6.1.2, ma a differenza di essa considerando anche il contributo del possibile scambio di un bosone Z nel processo.

Tralasciamo, come al solito, m_e e m_f , utilizziamo il gauge di Feynman e poniamo $m \equiv M_Z$. Possiamo vedere i diagrammi:



Il diagramma con l'Higgs intermedio è proporzionale a m_f , quindi qui è 0. Possiamo enumerare i diagrammi rispettivamente 1 e 2. Abbiamo dunque:

$$i\mathcal{M}_1 = \bar{v}(p')(-ie\gamma_\mu)u(p)\frac{-ig^{\mu\nu}}{q^2}\bar{u}(k)(-ieQ_f\gamma_\nu)v(k')(-1) \quad (\text{B.0.1})$$

dove $Q_u = +2/3$, $Q_d = -1/3$, $q = p + p'$, $q^2 = 2p \cdot p' = s$ e il secondo diagramma:

$$i\mathcal{M}_2 = \bar{v}(p')\left(-i\frac{g_w}{2\cos\theta_w}\gamma_\mu(V_e - A_e\gamma_5)\right)u(p)\frac{-g^{\mu\nu} + \frac{q^\mu q^\nu}{m^2}}{q^2 - m^2}\bar{u}(k) \times \\ \times \left(-i\frac{g_w}{2\cos\theta_w}\gamma_\nu(V_f - A_f\gamma_5)\right)v(k') \quad (\text{B.0.2})$$

dove $q^\mu q^\nu$ si cancellerà tramite l'equazione di Dirac e possiamo ignorare per il momento gli effetti di Breit-Wigner. Saranno introdotti in seguito.

Calcoliamo:

$$\sum_{pol} |\mathcal{M}_1|^2 = \bar{v}(p')\gamma_\mu u(p)\bar{u}(p)\gamma_\rho v(p')\bar{u}(k)\gamma^\mu v(k')\bar{v}(k')\gamma^\rho u(k) \frac{e^4 Q_f^2}{q^4} \quad (\text{B.0.3})$$

$$= \text{Tr}\{\not{p}'\gamma_\mu\not{p}\gamma_\rho\}\text{Tr}\{\not{k}\gamma^\mu\not{k}'\gamma^\rho\} \frac{e^4 Q_f^2}{q^4} \quad (\text{B.0.4})$$

$$= \frac{16e^4 Q_f^2}{q^4} [p'_\mu p_\rho + p'_\rho p_\mu - p \cdot p' g_{\mu\rho}] \times \\ \times [k^\mu k'^\rho + k'^\mu k^\rho - k \cdot k' g^{\mu\rho}] \quad (\text{B.0.5})$$

$$= \frac{16e^4 Q_f^2}{q^4} [2p' \cdot k p \cdot k' + 2p \cdot k p' \cdot k'] \quad (\text{B.0.6})$$

$$= \frac{16e^4 Q_f^2}{2(p \cdot p')^2} [p' \cdot k \cdot p \cdot k' + p \cdot k \cdot p' \cdot k'] \quad (\text{B.0.7})$$

$$= \frac{16e^4 Q_f^2}{2(p \cdot p')^2} [(p \cdot k')^2 + (p \cdot k)^2] \quad (\text{B.0.8})$$

in cui abbiamo $p \cdot k' = p' \cdot k$ e $p \cdot k = p' \cdot k'$. Il secondo termine del modulo quadro è:

$$\sum_{pol} |\mathcal{M}_2|^2 = \bar{v}(p')\gamma_\mu(V_e - A_e\gamma_5)u(p)\bar{u}(p)\gamma_\rho(V_e - A_e\gamma_5)v(p') \times \\ \times \bar{u}(k)\gamma^\mu(V_f - A_f\gamma_5)v(k')\bar{v}(k')\gamma^\rho(V_f - A_f\gamma_5)u(k) \times \\ \times \left(\frac{g_w}{2\cos\theta_w}\right)^4 \frac{(-g^{\mu\nu} + \frac{q^\mu q^\nu}{m^2})(-g^{\rho\sigma} + \frac{q^\rho q^\sigma}{m^2})}{(q^2 - m^2)^2} \quad (\text{B.0.9})$$

$$= \left(\frac{g_w}{2\cos\theta_w}\right)^4 \frac{1}{(q^2 - m^2)^2} \text{Tr}(\not{p}'\gamma_\mu(V_e - A_e\gamma_5)\not{p}\gamma_\rho(V_e - A_e\gamma_5)) \times \quad (\text{B.0.10})$$

$$\times \text{Tr}(\not{k}\gamma^\mu(V_f - A_f\gamma_5)\not{k}'\gamma^\rho(V_f - A_f\gamma_5)) \quad (\text{B.0.11})$$

dove in (B.0.11) ogni termine in $q^\mu q^\nu$ (o $q^\rho q^\sigma$) svanisce poiché: abbiamo $q = p + p'$ e dunque:

$$\begin{aligned} q^\mu q^\nu \bar{v}(p')\gamma_\mu(A + B\gamma_5)u(p) &= q^\nu \bar{v}(p')(\not{p} + \not{p}')(A + B\gamma_5)u(p) \\ &= q^\nu \bar{v}(p')\not{p}(A + B\gamma_5)u(p) \\ &= q^\nu \bar{v}(p')(A - B\gamma_5)\not{p}u(p) \\ &= 0. \end{aligned}$$

quindi il nostro conto (B.0.11) continua come:

$$\sum_{pol} |\mathcal{M}_2|^2 = \underbrace{\left(\frac{g_w}{2 \cos \theta_w} \right)^4 \frac{1}{(q^2 - m^2)^2} \text{Tr} \{ \not{p}' \gamma_\mu \not{p} \gamma_\rho (V_e^2 + A_e^2 - 2V_e A_e \gamma_5) \}}_{\equiv L} \times \text{Tr} \{ \not{k} \gamma^\mu \not{k}' \gamma^\rho (V_f^2 + A_f^2 - 2V_f A_f \gamma_5) \} \quad (\text{B.0.12})$$

$$= L \left[(V_e^2 + A_e^2) 4(p_\mu p_\rho + p_\rho p'_\mu - g_{\mu\rho} p \cdot p') + 2V_e A_e 4i\epsilon_{\mu\rho\sigma\sigma'} \right] \times \\ \times \left[(V_f^2 + A_f^2) 4(k^\mu k'^\rho + k'^\mu k^\rho - k \cdot k' g^{\mu\rho}) + 2V_f A_f 4i\epsilon^{\mu\rho\sigma\sigma'} \right] \quad (\text{B.0.13})$$

$$= L \left[16(V_e^2 + A_e^2)(V_f^2 + A_f^2) 2((p \cdot k)^2 + (p \cdot k')^2) - \right. \\ \left. - 64V_e A_e V_f A_f p^\alpha p'^\beta k^\lambda k'^\delta \epsilon_{\alpha\mu\beta\rho} \epsilon^{\mu\lambda\rho\delta} \right] \quad (\text{B.0.14})$$

$$= L \left[32(V_e^2 + A_e^2)(V_f^2 + A_f^2) ((p \cdot k)^2 + (p \cdot k')^2) + \right. \\ \left. + 128V_e A_e V_f A_f ((p \cdot k)^2 - (p \cdot k')^2) \right]. \quad (\text{B.0.15})$$

in cui in (B.0.13) con gli indici p, p', k, k' del tensore di Levi-Civita intendiamo dire che tali indici sono contratti con i rispettivi dei quadri-impulsi.

Possiamo calcolare il termine di interferenza:

$$\sum_{pol} \mathcal{M}_1^* \mathcal{M}_2 = \sum_{pol} \mathcal{M}_2^* \mathcal{M}_1 \quad (\text{B.0.16})$$

$$= -\bar{u}(p) \gamma_\mu v(p') \frac{Q_f e^2}{q^2} \bar{v}(k') \gamma^\mu u(k) \cdot \bar{v}(p') \frac{g_w^2}{2 \cos^2 \theta_w} \gamma_\rho (V_e - A_e \gamma_5) \times \\ \times u(p) \frac{1}{q^2 - m^2} \bar{u}(k) \gamma^\rho (V_f - A_f \gamma_5) v(k') \quad (\text{B.0.17})$$

$$= \underbrace{\frac{g_w^2 e^2 Q_f}{(2 \cos \theta_w)^2 q^2 (q^2 - m^2)}}_{\equiv N} (-) \text{Tr} \{ \not{p}' \gamma_\mu \not{p} \gamma_\rho (V_e - A_e \gamma_5) \} \times \quad (\text{B.0.18})$$

$$\times \text{Tr} \{ \not{k} \gamma^\mu \not{k}' \gamma^\rho (V_f - A_f \gamma_5) \} \quad (\text{B.0.19})$$

$$= N \cdot 16 \left\{ \left[(p_\mu p_\rho + p_\rho p'_\mu - p \cdot p' g_{\mu\rho}) V_e + A_e p^\alpha p'^\beta i\epsilon_{\alpha\mu\beta\rho} \right] \times \right. \\ \left. \times \left[(k^\mu k'^\rho + k'^\mu k^\rho - k \cdot k' g^{\mu\rho}) V_f + A_f k^\lambda k'^\sigma i\epsilon_{\lambda\mu\sigma\rho} \right] \right\} \quad (\text{B.0.20})$$

$$= 16N \left\{ V_e V_f \left[2((p \cdot k)^2 + (p \cdot k')^2) \right] + \right. \\ \left. + 2A_e A_f \left[(p \cdot k')^2 - (p \cdot k)^2 \right] \right\}. \quad (\text{B.0.21})$$

Quindi complessivamente abbiamo:

$$\begin{aligned}
3 \sum_{pol, colors} |\mathcal{M}|^2 = & \frac{3 \cdot 16e^4 Q_f^2}{2(p \cdot p')^2} [(p \cdot k)^2 + (p \cdot k')^2] + \\
& + \left(\frac{g_w}{\cos \theta_w} \right)^4 \frac{3 \cdot 2}{(q^2 - m^2)^2} \times \\
& \times [(V_e^2 + A_e^2)(V_f^2 + A_f^2) ((p \cdot k)^2 + (p \cdot k')^2) + \\
& + 4V_e V_f A_e A_f ((p \cdot k)^2 - (p \cdot k')^2)] - \\
& - \frac{32g_w^2 e^2 Q_f \cdot 3}{(2 \cos \theta_w)^2 q^2 (q^2 - m^2)} [V_e V_f [(p \cdot k)^2 + (p \cdot k')^2] + \\
& + A_e A_f [(p \cdot k')^2 - (p \cdot k)^2]] . \tag{B.0.22}
\end{aligned}$$

Se ci poniamo nel sistema del Centro di Massa:

$$p^\mu = p(1, 0, 0, 1) \quad ; \quad p'^\mu = p(1, 0, 0, -1) \tag{B.0.23}$$

$$k^\mu = k(1, 0, \sin \theta, \cos \theta) \quad ; \quad k'^\mu = k(1, 0, -\sin \theta, -\cos \theta) \tag{B.0.24}$$

dunque in cui valgono:

$$s = (p + p')^2 = 4p^2 \quad ; \quad s = (k + k')^2 = 4k^2 \implies k = p \tag{B.0.25}$$

$$p \cdot k = p^2(1 - \cos \theta) \quad ; \quad p \cdot k' = p^2(1 + \cos \theta) \quad ; \quad p \cdot p' = 2p^2 \tag{B.0.26}$$

$$\text{con } p^2 = s/4 \tag{B.0.27}$$

allora abbiamo:

$$\begin{aligned}
3 \sum_{pol} |\mathcal{M}|^2 &= \frac{3 \cdot 16e^4 Q_f^2}{s^2/2} (p^4(1 + \cos^2 \theta + 2 \cos \theta) + p^4(1 + \cos^2 \theta - 2 \cos \theta)) + \\
&+ 3 \left(\frac{g_w}{\cos \theta_w} \right)^4 \frac{2}{(s - m^2)^2} [2p^4(V_e^2 + A_e^2)(V_f^2 + A_f^2)(1 + \cos^2 \theta) + \\
&+ 4V_e A_e V_f A_f \cdot 4p^4 \cos \theta] - \\
&- \frac{3 \cdot 16g_w^2 e^2 Q_f}{\cos^2 \theta_w s(s - m^2)} [V_e V_f 2(1 + \cos^2 \theta) + A_e A_f 4 \cos \theta] \cdot p^4
\end{aligned} \tag{B.0.28}$$

$$\begin{aligned}
&= 3(1 + \cos^2 \theta) \left[\frac{64e^4 Q_f^2}{s^2} p^4 + \right. \\
&+ \left(\frac{g_w}{\cos \theta_w} \right)^4 \frac{4}{(s - m^2)^2} p^4 (V_e^2 + A_e^2)(V_f^2 + A_f^2) - \\
&- \frac{32g_w^2 e^2 Q_f p^4}{\cos^2 \theta_w s(s - m^2)} V_e V_f \left. \right] + \\
&+ 3 \cos \theta \left[\left(\frac{g_w}{\cos \theta_w} \right)^4 \frac{32p^4 V_e V_f A_e A_f}{(s - m^2)^2} - \frac{64g_w^2 e^2 Q_f p^4}{\cos^2 \theta_w s(s - m^2)} A_e A_f \right]
\end{aligned} \tag{B.0.29}$$

$$\equiv \left(3 \sum_{pol} |\mathcal{M}|^2 \right)^{TL} . \tag{B.0.30}$$

Lo spazio delle fasi per questo processo è:

$$d\Phi_2 = \frac{d^3 k}{(2\pi)^3 2E} \frac{d^3 k'}{(2\pi)^3 2E'} (2\pi)^4 \delta^{(4)}(p + p' - k - k') \tag{B.0.31}$$

$$= \frac{d^3 k}{(2\pi)^2 4p^2} \delta(2p - 2k) \tag{B.0.32}$$

$$= \frac{k^2 dk d\Omega}{8p^2 (2\pi)^2} \delta(p - k) \tag{B.0.33}$$

$$= \frac{d \cos \theta}{16\pi} \tag{B.0.34}$$

ed essendo il fattore di flusso:

$$\mathcal{F} = 1/2s \tag{B.0.35}$$

si ha dunque:

$$d\sigma/d \cos \theta = \frac{1}{2s} \frac{1}{16\pi} \frac{1}{4} \sum_{pol} |\mathcal{M}|^2 \cdot 3 \tag{B.0.36}$$

dove $1/4$ deriva dalla media sugli spin di e^+e^- .

Se indichiamo:

$$\alpha \equiv \frac{e^2}{4\pi} \quad ; \quad g_w = \frac{e}{\sin \theta_w} \quad (\text{B.0.37})$$

allora abbiamo:

$$\begin{aligned} \frac{d\sigma}{d\cos\theta} = \frac{3}{128\pi s} \frac{s^2}{16} \left\{ \left[\frac{64 \cdot 16\pi^2 \alpha^2 Q_f^2}{s^2} + \right. \right. \\ \left. + \frac{\alpha^2 16\pi^2}{(\sin \theta_w \cos \theta_w)^4} \frac{4}{(s-m^2)^2} (V_e^2 + A_e^2)(V_f^2 + A_f^2) - \right. \\ \left. - \frac{Q_f V_e V_f 32\alpha^2 16\pi^2}{(\sin \theta_w \cos \theta_w)^2 s(s-m^2)} \right] (1 + \cos^2 \theta) + \\ \left. + \left[\frac{\alpha^2 16\pi^2}{(\sin \theta_w \cos \theta_w)^4} \frac{32V_e V_f A_e A_f}{(s-m^2)^2} - \right. \right. \\ \left. - \frac{64Q_f \alpha^2 16\pi^2 A_e A_f}{(\sin \theta_w \cos \theta_w)^2 s(s-m^2)} \right] \cos \theta \left. \right\} \quad (\text{B.0.38}) \end{aligned}$$

$$\begin{aligned} = \frac{3\alpha^2\pi}{2s} \left\{ (1 + \cos^2 \theta) \times \right. \\ \times \left[Q_f^2 + (V_e^2 + A_e^2)(V_f^2 + A_f^2) \left[\frac{s^2}{16} \frac{1}{(s-m^2)^2} \frac{1}{(\sin \theta_w \cos \theta_w)^4} \right] + \right. \\ \left. + 2Q_f V_e V_f \left[\frac{s^2}{4} \frac{1}{s(m^2-s)} \frac{1}{(\sin \theta_w \cos \theta_w)^2} \right] \right] + \\ \left. + (\cos \theta) \left[s^2 \frac{1}{(s-m^2)^2} \frac{1}{(\sin \theta_w \cos \theta_w)^4} Q_f A_e A_f + \right. \right. \\ \left. + 8V_e V_f A_e A_f \frac{s^2}{16} \frac{1}{(s-m^2)^2} \frac{1}{(\sin \theta_w \cos \theta_w)^4} \right] \left. \right\} \quad (\text{B.0.39}) \end{aligned}$$

$$\begin{aligned} = \frac{3\alpha^2\pi}{2s} \left\{ (1 + \cos^2 \theta) [Q_f^2 + (V_e^2 + A_e^2)(V_f^2 + A_f^2)\chi_2(s) - \right. \\ \left. - 2Q_f V_e V_f \chi_1(s)] + \right. \\ \left. + (\cos \theta) [-4Q_f A_e A_f \chi_1(s) + 8V_e V_f A_e A_f \chi_2(s)] \right\}, \quad (\text{B.0.40}) \end{aligned}$$

con le definizioni:

$$\chi_1(s) \equiv \frac{-s}{4(m^2-s)(\sin \theta_w \cos \theta_w)^2} \quad (\text{B.0.41})$$

$$\chi_2(s) \equiv \frac{s^2}{16(m^2-s)^2(\sin \theta_w \cos \theta_w)^4} = \chi_1(s)^2. \quad (\text{B.0.42})$$

Ora, prendendo in considerazione la corretta (si veda ad esempio Peskin pag. 237) **sostituzione di Breit-Wigner**:

$$\frac{1}{s-m^2} \rightarrow \frac{1}{s-m^2+im\Gamma} \quad (\text{B.0.43})$$

nel propagatore della particella bosonica pesante, si ottiene una diversa $d\sigma/d\cos\theta$. (Γ è la larghezza di decadimento totale del bosone Z).

In $|\mathcal{M}_2|^2$ si ha semplicemente $(s - m^2)^2 \rightarrow (s - m^2)^2 + (m\Gamma)^2$, mentre per i termini di interferenza si ha $\mathcal{M}_1\mathcal{M}_2^* + \mathcal{M}_1^*\mathcal{M}_2$ e quindi il contributo di interferenza diventa:

$$\frac{2A}{s - m^2} \rightarrow \frac{2A(s - m^2)}{(s - m^2)^2 + (m\Gamma)^2}. \quad (\text{B.0.44})$$

Quindi:

$$\chi_1(s) \rightarrow \frac{1}{4(\sin\theta_w \cos\theta_w)^2} \underbrace{\frac{s(s - m^2)}{[(s - m^2)^2 + (m\Gamma)^2]}}_{L_1} \quad (\text{B.0.45})$$

e:

$$\chi_2(s) \rightarrow \frac{1}{16(\sin\theta_w \cos\theta_w)^4} \underbrace{\frac{s^2}{[(s - m^2)^2 + (m\Gamma)^2]}}_{L_2}. \quad (\text{B.0.46})$$

Definendo infine:

$$G_F \equiv \frac{g_w^2 \sqrt{2}}{8m^2 \cos^2 \theta_w} = \frac{e^2 \sqrt{2}}{8m^2 \sin^2 \theta_w \cos^2 \theta_w} \quad (\text{B.0.47})$$

abbiamo infine:

$$\chi_1(s) = L_1 \frac{\sqrt{2}G_F m^2}{4\pi\alpha} \quad \text{e} \quad \chi_2(s) = L_2 \left(\frac{\sqrt{2}G_F m^2}{4\pi\alpha} \right)^2. \quad (\text{B.0.48})$$

Con queste notazioni, si ha la sezione d'urto differenziale a tree-level:

$$\begin{aligned} \left(\frac{d\sigma}{d\cos\theta} \right)_{\text{tree level}} &= \frac{3\alpha^2\pi}{2s} [(1 + \cos^2\theta) [Q_f^2 + (V_e^2 + A_e^2)(V_f^2 + A_f^2)\chi_2(s) - \\ &\quad - 2Q_f V_e V_f \chi_1(s)] + \\ &\quad + (\cos\theta) [-4Q_f A_e A_f \chi_1(s) + 8V_e V_f A_e A_f \chi_2(s)]] , \end{aligned} \quad (\text{B.0.49})$$

con le seguenti definizioni per i fattori del propagatore:

$$\chi_1(s) \equiv k \frac{s(s - m^2)}{(s - m^2)^2 + (\Gamma m)^2} \quad ; \quad \chi_2(s) \equiv k^2 \frac{s^2}{(s - m^2)^2 + (\Gamma m)^2} \quad (\text{B.0.50})$$

$$k \equiv \frac{\sqrt{2}G_F m^2}{4\pi\alpha}. \quad (\text{B.0.51})$$

Questo risultato coincide con l'equazione (3.1) di There is a sort of mythology that grows up about what happened, which is different from what really did happen. Ellis, a meno del fattore 3 che tiene conto della *somma sui colori* (numero di colori $N_c = 3$ per i quark).

Appendice C

Fattore di colore e flusso

Possiamo dare una rappresentazione grafica della propagazione del colore. Abbiamo il *propagatore del quark*:

$$j \longrightarrow i \longrightarrow \delta_{ij} \quad (\text{C.0.1})$$

e da esso possiamo vedere:

$$\begin{array}{c} i \\ \boxed{\longrightarrow} \end{array} \begin{array}{c} i \\ \longleftarrow \end{array} = \delta_{ii} = \text{Tr}\{\delta_{ij}\} = N = 3. \quad (\text{C.0.2})$$

Sappiamo di avere il *gluone* nella rappresentazione aggiunta, quindi $3 \otimes \bar{3} = 8 \oplus 1$, dove abbiamo i fattori: 3 per il colore, $\bar{3}$ per l'anti-colore, 8 sono i gluoni e 1 è la traccia. Quindi, il colore si propaga come $3 \otimes \bar{3}$ con traccia nulla:

$$a \text{~~~~~} b \longrightarrow a \begin{array}{c} \longrightarrow \\ \longleftarrow \end{array} b + B \begin{array}{c} \longrightarrow \\ \longleftarrow \end{array} \boxed{\longrightarrow} \begin{array}{c} \longrightarrow \\ \longleftarrow \end{array} b = \delta_{ab}. \quad (\text{C.0.3})$$

Possiamo trovare B calcolando la traccia:

$$\delta_{aa} = N^2 - 1 \quad (\text{C.0.4})$$

$$= \begin{array}{c} \text{---} \\ \text{---} \end{array} + B \begin{array}{c} \text{---} \\ \text{---} \end{array} \quad (\text{C.0.5})$$

$$= N^2 + BN \quad (\text{C.0.6})$$

$$\Rightarrow B = -\frac{1}{N}. \quad (\text{C.0.7})$$

Quindi la (C.0.3) diventa:

$$a \text{~~~~~} b \longrightarrow a \begin{array}{c} \longrightarrow \\ \longleftarrow \end{array} b + \frac{1}{N} \begin{array}{c} \longrightarrow \\ \longleftarrow \end{array} \boxed{\longrightarrow} \begin{array}{c} \longrightarrow \\ \longleftarrow \end{array} b \quad (\text{C.0.8})$$

Di conseguenza possiamo scrivere:

$$t_{ij}^a = A \left[\begin{array}{c} a \\ \uparrow \downarrow \\ j \longrightarrow \quad \longrightarrow i \end{array} - \frac{1}{N} \begin{array}{c} a \\ \uparrow \downarrow \\ j \longrightarrow \quad \longrightarrow i \end{array} \right] \quad (\text{C.0.9})$$

dove il primo termine è il *leading color* (colore dominante) e il secondo è il *subleading color* (colore subdominante). Possiamo anche verificare che $\text{Tr}\{t_{ij}^a\} = 0$:

$$A \left[\begin{array}{c} a \\ \uparrow \downarrow \\ \text{---} \end{array} - \frac{1}{N} \begin{array}{c} a \\ \uparrow \downarrow \\ \text{---} \end{array} \right] = A \begin{array}{c} a \\ \uparrow \downarrow \\ \text{---} \end{array} \left(1 - \frac{1}{N} N \right) = 0. \quad (\text{C.0.10})$$

Possiamo trovare A calcolando:

$$\text{Tr}\{t^a t^b\} = T_R \delta^{ab} \quad (\text{C.0.11})$$

$$= T_R \left[a \text{---} b - \frac{1}{N} a \text{---} b \right] \quad (\text{C.0.12})$$

$$= A^2 \left[\begin{array}{c} a \quad b \\ \uparrow \downarrow \quad \uparrow \downarrow \\ \text{---} \end{array} - \frac{1}{N} \begin{array}{c} a \quad b \\ \uparrow \downarrow \quad \uparrow \downarrow \\ \text{---} \end{array} - \frac{1}{N} \begin{array}{c} a \quad b \\ \uparrow \downarrow \quad \uparrow \downarrow \\ \text{---} \end{array} + \frac{1}{N} \begin{array}{c} a \quad b \\ \uparrow \downarrow \quad \uparrow \downarrow \\ \text{---} \end{array} \right] \quad (\text{C.0.13})$$

$$= A^2 \left[a \text{---} b - \frac{1}{N} a \text{---} b - \frac{1}{N} a \text{---} b + \frac{1}{N^2} N a \text{---} b \right] \quad (\text{C.0.14})$$

$$= A^2 \delta^{ab} \quad (\text{C.0.15})$$

$$\Rightarrow A = \sqrt{T_R}. \quad (\text{C.0.16})$$

Quindi, l'equazione (C.0.9) diventa:

$$t_{ij}^a = \sqrt{T_R} \left[\begin{array}{c} a \\ \uparrow \downarrow \\ j \longrightarrow \quad \longrightarrow i \end{array} - \frac{1}{N} \begin{array}{c} a \\ \uparrow \downarrow \\ j \longrightarrow \quad \longrightarrow i \end{array} \right]. \quad (\text{C.0.17})$$

$$if^{abc} = \frac{1}{T_R} \text{Tr} \left\{ \left[t^a, t^b \right] t^c \right\} \quad (\text{C.0.18})$$

(C.0.18)

(C.0.19)

(C.0.20)

(C.0.21)

(C.0.22)

(C.0.23)

(C.0.24)

(C.0.25)

tra i diagrammi. Alcuni esempi sono mostrati:

Diagrammatic equation (C.0.33): A gluon loop (a circle with two wavy lines labeled γ and two internal gluon lines) is shown to be equal to a sum of diagrams. The first term is $\frac{1}{4}$ times a diagram with four horizontal lines and four curved lines, representing a leading color term. The sum continues with ellipses and a final term $\frac{1}{4}N^3 + \dots$.

$$\text{gluon loop} \rightarrow \frac{1}{4} \text{[diagram]} + \dots = \frac{1}{4}N^3 + \dots \quad (\text{C.0.33})$$

dove i puntini rappresentano termini con potenze inferiori di N . Possiamo anche vedere:

Diagrammatic equation (C.0.34): A gluon loop (a circle with two wavy lines labeled γ and two internal gluon lines) is shown to be equal to a sum of diagrams. The first term is $\frac{1}{4}$ times a diagram with four horizontal lines and four curved lines, representing a leading color term. The sum continues with ellipses and a final term $\frac{1}{4}N + \dots$.

$$\text{gluon loop} \rightarrow \frac{1}{4} \text{[diagram]} + \dots = \frac{1}{4}N + \dots \quad (\text{C.0.34})$$

che è soppresso di un fattore $\frac{1}{N^2} = \frac{1}{9}$ rispetto al precedente.

Quando il calcolo è troppo complicato (più di 2 loop) possiamo calcolarlo al leading color e poi includere i termini $1/N^2$.

Appendice D

Note sui contributi reale e virtuale a $e^+e^- \rightarrow$ adroni

Vediamo in questa sezione prima una discussione, leggermente più quantitativa rispetto quella vista nel capitolo §6, riguardante le divergenze IR e UV per il contributo virtuale; successivamente ho riportato i conti del contributo reale alla sezione d'urto nel light-cone gauge, che pur non costituendo materiale d'esame, può essere considerato interessante.

D.1 Regolatori dimensionali IR e UV

Come abbiamo già detto nel corso delle note, in particolare nel capitolo §6, può essere utile tenere separati i regolatori dimensionali UV e IR. Definendo $\epsilon^{\text{IR}} \equiv \epsilon$ e $\epsilon^{\text{UV}} \equiv \epsilon'$, possiamo prima estrarre la dipendenza da ϵ' dal fattore di forma Λ_ρ : la correzione del vertice è un diagramma con divergenza logaritmica, quindi il suo polo UV può essere calcolato nel limite di invarianti di Mandelstam nulli (qui s), ovvero:

$$\Lambda_\rho|_{\text{UV}} = -2ig_s^2 C_F \gamma_\rho \int_0^1 dx \int_0^{1-x} dy \int \frac{d^d l}{(2\pi)^d} \frac{1}{l^4} \quad (\text{D.1.1})$$

$$= \gamma_\rho \frac{\alpha_s}{4\pi} C_F \frac{1}{\epsilon'} \quad (\text{D.1.2})$$

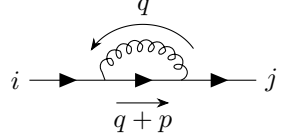
quindi l'equazione (6.1.148) viene sostituita da:

$$\Lambda_\rho = \gamma_\rho \frac{\alpha_s}{4\pi} \frac{C_F}{\Gamma(1-\epsilon)} \left(\frac{4\pi\mu^2}{s} \right)^\epsilon e^{i\pi\epsilon} \left[-\frac{2}{\epsilon^2} - \frac{3}{\epsilon} - 8 + \mathcal{O}(\epsilon) \right] + \gamma_\rho \frac{\alpha_s}{4\pi} C_F \left[\frac{1}{\epsilon'} - \frac{1}{\epsilon} \right] \quad (\text{D.1.3})$$

$$= \gamma_\rho C'. \quad (\text{D.1.4})$$

Lo stesso risultato può essere ottenuto sostituendo $1/\epsilon \rightarrow 1/\epsilon'$ nel coefficiente di A nella seconda riga dell'equazione (6.1.148), poi espandendo in ϵ e considerando $\epsilon/\epsilon' = 1$ nel termine finito.

Distinguendo i due regolatori, la self-energy $i\Sigma(p)$ non svanisce più, quindi la calcoliamo esplicitamente. Abbiamo:


(D.1.5)

per cui:

$$\begin{aligned}
 \Sigma(p) &= -i \int \frac{d^d q}{(2\pi)^d} \frac{-i}{q^2} (ig_s \mu^\epsilon \gamma^\mu t^a) \frac{i(q+p)}{(q+p)^2} (ig_s \mu^\epsilon \gamma_\mu t^a) \\
 &= ig_s^2 \mu^{2\epsilon} C_F \int \frac{d^d q}{(2\pi)^d} \frac{\gamma^\mu (q+p) \gamma_\mu}{q^2 (q+p)^2} \\
 &= -ig_s^2 \mu^{2\epsilon} C_F 2(1-\epsilon) p \int_0^1 dx (1-x) \int \frac{d^d l}{(2\pi)^d} \frac{1}{l^4}
 \end{aligned} \tag{D.1.6}$$

dove $l^\mu = q^\mu + xp^\mu$, e abbiamo ommesso i termini lineari in l per simmetria. Il fattore di colore può essere immediatamente verificato ripristinando gli indici di colore: $\Sigma_{ij}(p) \propto t_{ik}^a t_{kj}^a = C_F \delta_{ij}$, dove come di consueto sottintendiamo la delta di Kronecker di colore.

Il polo UV di $\Sigma(p)$ è:

$$\Sigma(p)|_{\text{UV}} = -2ig_s^2 C_F p \int_0^1 dx (1-x) \int \frac{d^d l}{(2\pi)^d} \frac{1}{l^4} \tag{D.1.7}$$

$$= p \frac{\alpha_s}{4\pi} C_F \frac{1}{\epsilon'} \tag{D.1.8}$$

e poiché con $\epsilon = \epsilon'$ la self-energy svanisce esattamente, si deduce:

$$\Sigma(\not{p}) = \not{p} \frac{\alpha_s}{4\pi} C_F \left[\frac{1}{\epsilon'} - \frac{1}{\epsilon} \right] \tag{D.1.9}$$

$$= \not{p} K. \tag{D.1.10}$$

Lo stesso risultato può essere ottenuto dividendo l'integrale in una regione IR e una UV, ed eseguendo un'integrazione euclidea, si veda ad esempio le pagine 172-173 di Muta There is a sort of mythology that grows up about what happened, which is different from what really did happen. Muta. La correzione del propagatore a un loop è data da:

$$\left(\frac{i\not{p}}{p^2} \right) i\Sigma(\not{p}) \left(\frac{i\not{p}}{p^2} \right) = -iK \frac{\not{p}}{p^2}. \tag{D.1.11}$$

Vediamo che il coefficiente della divergenza UV è lo stesso della correzione del vertice, producendo un contributo di loop UV-finito:

$$\mathcal{M}_v = C' \mathcal{M}_b - \frac{1}{2} K \mathcal{M}_b - \frac{1}{2} K \mathcal{M}_b \tag{D.1.12}$$

$$= C \mathcal{M}_b, \tag{D.1.13}$$

equivalente al risultato precedente.

D.2 Contributo reale in gauge di cono di luce

Supponiamo di trattare una versione semplificata dell'annichilazione e^+e^- , in cui il tensore leptónico è $\hat{L}^{\rho\sigma} \propto g^{\rho\sigma}$. Trascurando le costanti irrilevanti e ponendo $d = 4$, in gauge di Feynman i contributi reali a questo processo sono:

$$\begin{aligned} |m_{r1}|^2 &= \frac{1}{4p_1 \cdot k^2} \bar{u}(p_1) \gamma^\mu (\not{p}_1 + \not{k}) \gamma^\rho v(p_2) \bar{v}(p_2) \gamma_\rho (\not{p}_1 + \not{k}) \gamma^\mu u(p_1) \\ &= 8 \frac{p_2 \cdot k}{p_1 \cdot k}, \end{aligned} \quad (\text{D.2.1})$$

$$|m_{r2}|^2 = 8 \frac{p_1 \cdot k}{p_2 \cdot k}, \quad (\text{D.2.2})$$

$$\begin{aligned} 2m_{r1}m_{r2}^* &= \frac{1}{2p_1 \cdot k p_2 \cdot k} \bar{u}(p_1) \gamma^\mu (\not{p}_1 + \not{k}) \gamma^\rho v(p_2) \bar{v}(p_2) \gamma^\mu (-\not{p}_2 - \not{k}) \gamma_\rho u(p_1) \\ &= \frac{16p_1 \cdot p_2 (p_1 \cdot p_2 + p_1 \cdot k + p_2 \cdot k)}{p_1 \cdot k p_2 \cdot k}. \end{aligned} \quad (\text{D.2.3})$$

Il termine di interferenza non è soppresso rispetto a $|m_{r1}|^2$ nel limite collineare e, inoltre, è l'unico termine con singolarità soft: non è possibile dedurre alcun comportamento IR basandosi sullo studio dei diagrammi. Al contrario, in una gauge fisica esiste una corrispondenza precisa tra i diagrammi e i contributi IR-amplificati. Lo vediamo eseguendo lo stesso calcolo nella gauge del cono di luce (assiale), ovvero quella in cui:

$$\epsilon_\mu \epsilon_\nu^* = -g_{\mu\nu} + \frac{n_\mu k_\nu + n_\nu k_\mu}{n} \cdot k. \quad (\text{D.2.4})$$

Con $d_1 = 2p_1 \cdot k$ e $d_2 = 2p_2 \cdot k$ abbiamo:

$$d_1^2 |m_{r1}|^2 = \bar{u}(p_1) \gamma^\mu (\not{p}_1 + \not{k}) \gamma^\rho v(p_2) \bar{v}(p_2) \gamma_\rho (\not{p}_1 + \not{k}) \gamma^\sigma u(p_1) \left[g_{\mu\sigma} - \frac{n_\mu k_\sigma + n_\sigma k_\mu}{n \cdot k} \right] \quad (\text{D.2.5})$$

$$= \text{Tr} \left\{ \not{p}_1 \gamma^\mu (\not{p}_1 + \not{k}) \gamma^\rho \not{p}_2 \gamma_\rho (\not{p}_1 + \not{k}) \gamma^\sigma \right\} \left[g_{\mu\sigma} - \frac{n_\mu k_\sigma + n_\sigma k_\mu}{n \cdot k} \right] \quad (\text{D.2.6})$$

$$= -2 \text{Tr} \left\{ \not{p}_1 \gamma^\mu (\not{p}_1 + \not{k}) \not{p}_2 (\not{p}_1 + \not{k}) \gamma^\sigma \right\} \left[g_{\mu\sigma} - \frac{n_\mu k_\sigma + n_\sigma k_\mu}{n \cdot k} \right] \quad (\text{D.2.7})$$

$$= 32p_1 \cdot kp_2 \cdot k + \frac{2}{n \cdot k} \text{Tr} \left\{ \not{p}_1 \not{k} \not{p}_1 \not{p}_2 (\not{p}_1 + \not{k}) \not{n} + \not{p}_1 \not{n} (\not{p}_1 + \not{k}) \not{p}_2 \not{p}_1 \not{k} \right\} \quad (\text{D.2.8})$$

$$= 32p_1 \cdot kp_2 \cdot k + \frac{4p_1 \cdot k}{n \cdot k} \text{Tr} \left\{ \not{p}_1 \not{p}_2 (\not{p}_1 + \not{k}) \not{n} + \not{p}_1 \not{n} (\not{p}_1 + \not{k}) \not{p}_2 \right\} \quad (\text{D.2.9})$$

$$= 32p_1 \cdot kp_2 \cdot k + \frac{32p_1 \cdot k}{n \cdot k} (2p_1 \cdot p_2 p_1 \cdot n + p_1 \cdot p_2 n \cdot k + p_1 \cdot np_2 \cdot k - p_1 \cdot kp_2 \cdot n), \quad (\text{D.2.10})$$

quindi:

$$|m_{r1}|^2 = 8 \left[\frac{p_2 \cdot k}{p_1 \cdot k} + \frac{2p_1 \cdot p_2 p_1 \cdot n + p_1 \cdot p_2 n \cdot k + p_1 \cdot np_2 \cdot k - p_1 \cdot kp_2 \cdot n}{p_1 \cdot kn \cdot k} \right], \quad (\text{D.2.11})$$

$$|m_{r2}|^2 = 8 \left[\frac{p_1 \cdot k}{p_2 \cdot k} + \frac{2p_1 \cdot p_2 p_2 \cdot n + p_1 \cdot p_2 n \cdot k + p_2 \cdot np_1 \cdot k - p_2 \cdot kp_1 \cdot n}{p_2 \cdot kn \cdot k} \right]. \quad (\text{D.2.12})$$

L'interferenza risulta:

$$d_1 d_2 m_{r1} m_{r2}^* = \bar{u}(p_1) \gamma^\mu (\not{p}_1 + \not{k}) \gamma^\rho v(p_2) \bar{v}(p_2) \gamma^\sigma (-\not{p}_2 - \not{k}) \gamma_\rho u(p_1) \left[g_{\mu\sigma} - \frac{n_\mu k_\sigma + n_\sigma k_\mu}{n \cdot k} \right] \quad (\text{D.2.13})$$

$$= -\text{Tr} \left\{ \not{p}_1 \gamma^\mu (\not{p}_1 + \not{k}) \gamma^\rho \not{p}_2 \gamma^\sigma (\not{p}_2 + \not{k}) \gamma_\rho \right\} \left[g_{\mu\sigma} - \frac{n_\mu k_\sigma + n_\sigma k_\mu}{n \cdot k} \right] \quad (\text{D.2.14})$$

$$= 2 \text{Tr} \left\{ \not{p}_1 \gamma^\mu (\not{p}_1 + \not{k}) (\not{p}_2 + \not{k}) \gamma^\sigma \not{p}_2 \right\} \left[g_{\mu\sigma} - \frac{n_\mu k_\sigma + n_\sigma k_\mu}{n \cdot k} \right] \quad (\text{D.2.15})$$

$$= 32 p_1 \cdot p_2 (p_1 \cdot p_2 + p_1 \cdot k + p_2 \cdot k) - \frac{2}{n \cdot k} \text{Tr} \left\{ \not{p}_1 \not{k} \not{p}_1 (\not{p}_2 + \not{k}) \not{p}_2 + \not{p}_1 \not{k} (\not{p}_1 + \not{k}) \not{p}_2 \not{k} \not{p}_2 \right\} \quad (\text{D.2.16})$$

$$= 32 p_1 \cdot p_2 (p_1 \cdot p_2 + p_1 \cdot k + p_2 \cdot k) - \frac{4}{n \cdot k} \left\{ p_1 \cdot k \text{Tr} \left\{ \not{p}_1 (\not{p}_2 + \not{k}) \not{p}_2 \right\} + p_2 \cdot k \text{Tr} \left\{ \not{p}_1 \not{k} (\not{p}_1 + \not{k}) \not{p}_2 \right\} \right\} \quad (\text{D.2.17})$$

$$= 32 p_1 \cdot p_2 (p_1 \cdot p_2 + p_1 \cdot k + p_2 \cdot k) - \frac{16}{n \cdot k} [p_1 \cdot n p_2 \cdot k (2 p_1 \cdot p_2 + p_2 \cdot k - p_1 \cdot k) + \quad (\text{D.2.18})$$

$$+ p_2 \cdot n p_1 \cdot k (2 p_1 \cdot p_2 + p_1 \cdot k - p_2 \cdot k) + n \cdot k p_1 \cdot p_2 (p_1 \cdot k + p_2 \cdot k)]. \quad (\text{D.2.19})$$

Prendendo ora il limite collineare $k \rightarrow a p_1 = \frac{1-z}{z} p_1$ si ottiene:

$$|m_{r1}|^2 \rightarrow 8 \frac{p_1 \cdot p_2}{p_1 \cdot k} \frac{a^2 + 2a + 2}{a} = 8 \frac{s}{t} \frac{1 + z^2}{1 - z}, \quad (\text{D.2.20})$$

$$d_1 d_2 m_{r1} m_{r2}^* \rightarrow 0. \quad (\text{D.2.21})$$

Il limite soft $k \rightarrow 0$ è leggermente più complesso e si ha:

$$|m_{r1}|^2 \rightarrow 16 \frac{p_1 \cdot p_2 p_1 \cdot n}{n \cdot k p_1 \cdot k}, \quad (\text{D.2.22})$$

$$|m_{r2}|^2 \rightarrow 16 \frac{p_1 \cdot p_2 p_2 \cdot n}{n \cdot k p_2 \cdot k}, \quad (\text{D.2.23})$$

$$2 m_{r2} m_{r2}^* \rightarrow 16 p_1 \cdot p_2 \frac{p_1 \cdot p_2}{p_1 \cdot k p_2 \cdot k} - 16 \frac{p_1 \cdot p_2 p_1 \cdot n}{p_1 \cdot k n \cdot k} - 16 \frac{p_1 \cdot p_2 p_2 \cdot n}{p_2 \cdot k n \cdot k}, \quad (\text{D.2.24})$$

che fornisce gli stessi risultati di prima.

Nella gauge assiale vediamo che i limiti collineari $\vec{k} \parallel \vec{p}_i$ derivano unicamente da $|m_{ri}|^2$, mentre sia l'interferenza che i diagrammi al quadrato contribuiscono ai limiti soft, coerentemente con il conteggio diagrammatico.

Appendice E

Algoritmo JADE per Jet da $e^+e^- \rightarrow$ adroni

- Parametrizzazione degli impulsi come segue:

$$q = \sqrt{s}(1, \vec{0}) \quad (\text{E.0.1})$$

$$k_3 = \frac{\sqrt{s}}{2} x_3 (1, 0, 0, 1) \quad (\text{E.0.2})$$

$$k_i = \frac{\sqrt{s} x_i}{2} (1, 0, \sin \theta_{ig}, \cos \theta_{ig}) \quad (\text{E.0.3})$$

con i vincoli:

$$\sum_{i=1}^3 x_i = 2 \quad (\text{E.0.4})$$

$$\sum_{i=1}^2 x_i \sin \theta_{ig} = 0 \quad (\text{E.0.5})$$

$$x_3 + \sum_{i=1}^2 x_i \cos \theta_{ig} = 0. \quad (\text{E.0.6})$$

Si ricordi che:

$$\cos \theta_{ij} = 1 - \frac{2(1 - x_k)}{x_i x_j}. \quad (\text{E.0.7})$$

- Si hanno tre jet JADE se $s_{ij} > ys$ per tutti gli $i, j = 1, \dots, 3$, dove $s_{ij} \equiv (k_i + k_j)^2$. Si noti che:

$$s_{ij} = \frac{\sqrt{s}}{2} \frac{\sqrt{s}}{2} 2x_i x_j (1 - \cos \theta_{ij}) = s x_i x_j \frac{1 - \cos \theta_{ij}}{2}. \quad (\text{E.0.8})$$

- Mappatura di queste condizioni sul piano $x_1 - x_2$:

$$x_i x_j (1 - \cos \theta_{ij}) \geq 2y \iff (1 - x_k) \geq y, \quad \forall k = 1, \dots, 3. \quad (\text{E.0.9})$$

Riferendosi al grafico (vedi la figura E.1), le curve A_i rappresentano le equazioni $1 - x_i = y$, e quindi i punti A, B e C hanno coordinate:

$$A = (2y, 1 - y), \quad B = \left(\frac{1+y}{2}, \frac{1+y}{2} \right), \quad (\text{E.0.10})$$

$$C = (1 - y, 1 - y). \quad (\text{E.0.11})$$

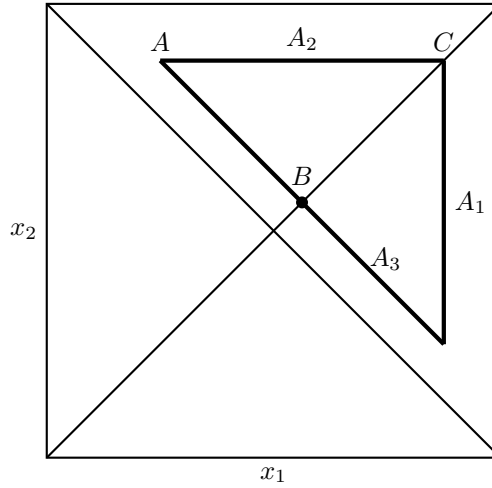


Figura E.1

Ora la procedura è analoga a quanto fatto per i jet SW (Stermann-Weinberg). Si calcoli $2I \equiv I_1 + I_2$, dove:

$$I_1 = \int_{2y}^{\frac{1+y}{2}} dx_1 \int_{1+y-x_1}^{1-y} dx_2 \frac{x_1^2 + x_2^2}{(1-x_1)(1-x_2)} \quad (\text{E.0.12})$$

$$I_2 = \int_{\frac{1+y}{2}}^{1-y} dx_1 \int_{x_1}^{1-y} dx_2 \frac{x_1^2 + x_2^2}{(1-x_1)(1-x_2)} \quad (\text{E.0.13})$$

I risultati per questi integrali sono forniti nel file Mathematica allegato (`3jet_fraction_SW.nb`). Ora, tenendo conto che (nella regione a 3 jet):

$$\frac{1}{\sigma} \frac{d^2\sigma}{dx_1 dx_2} = C_F \frac{\alpha_s}{2\pi} \frac{x_1^2 + x_2^2}{(1-x_1)(1-x_2)} + \mathcal{O}(\alpha_s^2), \quad (\text{E.0.14})$$

si ottiene:

$$f_3 = \frac{1}{\sigma} \int_{3j} \frac{d^2\sigma}{dx_1 dx_2} dx_1 dx_2 \quad (\text{E.0.15})$$

$$= \frac{C_F \alpha_s}{2\pi} 2I \quad (\text{E.0.16})$$

$$= \frac{C_F \alpha_s}{2\pi} \left[\frac{5}{2} - \frac{\pi^2}{3} - 6y - \frac{9}{2}y^2 + 2 \ln \left(\frac{y}{1-y} \right)^2 + \right. \\ \left. + (3 - 6y) \ln \left(\frac{y}{1-2y} \right) + 4\text{Li}_2 \left(\frac{y}{1-y} \right) \right]. \quad (\text{E.0.17})$$

$f_2 = 1 - f_3$. Questo coincide con la (3.35) di There is a sort of mythology that grows up about what happened, which is different from what really did happen. Ellis.

Appendice F

Conti Thrust

Riporto in questa Appendice alcuni conti per la sezione §6.3.3 dei thrust.

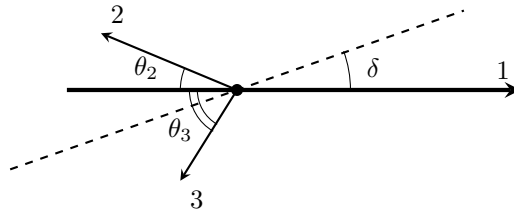
F.1 Conti valore T in $e^+e^- \rightarrow q\bar{q}g$

Vediamo la dimostrazione della relazione:

$$T = \max(x_i), \quad i = 1, 2, 3. \quad (\text{F.1.1})$$

Infine, dobbiamo mostrare che, nel caso di $e^+e^- \rightarrow q\bar{q}g$, la definizione di T coincide con $\max(x_i)$, per $i = 1, 2, 3$. Definiamo y_1, y_2, y_3 come la versione ordinata di x_1, x_2, x_3 , ovvero $y_1 = \max(x_i)$ e $y_3 = \min(x_i)$, cosicché $y_1 > y_2 > y_3$.

Poiché tre particelle con tri-impulso totale $\vec{0}$ giacciono su un piano, possiamo raffigurare un evento generico nel seguente modo:



dove abbiamo tenuto conto graficamente della gerarchia $y_1 > y_2 > y_3$ nella lunghezza delle frecce coinvolte.

Considerando $\hat{\mathbf{n}} = \hat{\mathbf{e}}_1$, per cui si ha:

$$\sum_{i=1}^3 |\mathbf{p}_i \cdot \hat{\mathbf{n}}| = p_1 + p_2 |\cos \theta_2| + p_3 |\cos \theta_3| \quad (\text{F.1.2})$$

dove $p_i \equiv |\mathbf{p}_i|$. La conservazione dell'impulso nella direzione $\hat{\mathbf{e}}_1$ implica che $p_1 = p_2 |\cos \theta_2| + p_3 |\cos \theta_3|$, quindi $\sum_{i=1}^3 |\mathbf{p}_i \cdot \hat{\mathbf{n}}| = 2p_1$ in questo caso.

Di conseguenza, per $\hat{\mathbf{n}} = \hat{\mathbf{e}}_1$:

$$T = \frac{\sum_i |\mathbf{p}_i \cdot \hat{\mathbf{n}}|}{\sum_i |\mathbf{p}_i|} = \frac{2p_1}{\sum_i p_i} = \frac{2y_1}{2} = y_1 = \max(x_i) \quad (\text{F.1.3})$$

dove negli ultimi passaggi abbiamo usato $p_i = \frac{\sqrt{s}}{2} x_i$ e $\sum_i x_i = 2$.

Dobbiamo ancora mostrare che se $\hat{\mathbf{n}} \neq \hat{\mathbf{e}}_1$ il valore della thrust diminuisce, cosicché $\hat{\mathbf{n}} = \hat{\mathbf{e}}_1$ sia la direzione che realizza la condizione di $\max_{\hat{\mathbf{n}}}$ implicita nella definizione di T .

Se consideriamo $\hat{\mathbf{n}}$ che forma un angolo δ rispetto a $\hat{\mathbf{e}}_1$, abbiamo:

$$|\mathbf{p}_1 \cdot \hat{\mathbf{n}}| = p_1 \cos \delta \quad ; \quad |\mathbf{p}_2 \cdot \hat{\mathbf{n}}| = p_2 \cos(\theta_2 - \delta) \quad (\text{F.1.4})$$

$$|\mathbf{p}_3 \cdot \hat{\mathbf{n}}| = p_3 \cos(\theta_3 + \delta). \quad (\text{F.1.5})$$

Sviluppando per piccoli valori di δ e mantenendo i contributi $\mathcal{O}(\delta^2)$, otteniamo:

$$|\mathbf{p}_1 \cdot \hat{\mathbf{n}}| \sim p_1(1 - \delta^2/2) \quad ; \quad |\mathbf{p}_k \cdot \hat{\mathbf{n}}| \sim p_k (\cos \theta_k(1 - \delta^2/2) \pm \sin \theta_k \delta) \quad (\text{F.1.6})$$

per $k = 2, 3$ (con il segno $+$ per $k = 2$ e il segno $-$ per $k = 3$). Quindi:

$$\sum_i |\mathbf{p}_i \cdot \hat{\mathbf{n}}| = 2p_1(1 - \delta^2/2) + \delta(p_2 \sin \theta_2 - p_3 \sin \theta_3) \quad (\text{F.1.7})$$

dove $(p_2 \sin \theta_2 - p_3 \sin \theta_3) = 0$ per la conservazione dell'impulso nella direzione $\perp$ a $\hat{\mathbf{e}}_1$ nel piano di scattering. Quindi:

$$T = y_1(1 - \delta^2/2) < y_1 \quad (\text{F.1.8})$$

cioè, allontanandosi dalla direzione di $\hat{\mathbf{e}}_1$, il valore della thrust diminuisce $\Rightarrow$ il $\max_{\hat{\mathbf{n}}}$ si ha per $\hat{\mathbf{n}} = \hat{\mathbf{e}}_1$ e $T = \max(x_i)$.

F.2 Conti per l'integrale

Possiamo disegnare lo spazio delle fasi [F.1](#), in cui il punto A ha coordinate $(2/3, 2/3)$. Si veda la figura [F.1](#). Si può semplicemente integrare nella regione con $x_1 > x_2$ e moltiplicare per 2 alla fine, a causa della simmetria $x_1 \leftrightarrow x_2$ dell'elemento di matrice.

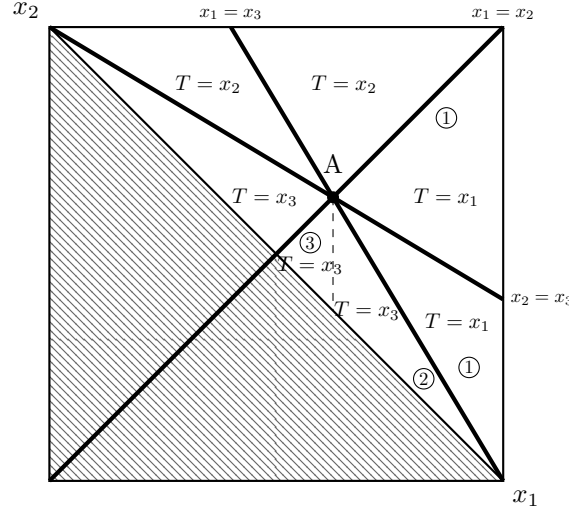


Figura F.1

Valutiamo i contributi rispettivamente dalle regioni 1, 2, 3:

$$I_1 = \int_{2/3}^1 dx_1 \int_{2(1-x_1)}^{x_1} dx_2 \frac{x_1^2 + x_2^2}{(1-x_1)(1-x_2)} \delta(T - x_1) \quad (\text{F.2.1})$$

$$= \int_0^1 dx_1 \int_0^1 dx_2 \frac{x_1^2 + x_2^2}{(1-x_1)(1-x_2)} \delta(T - x_1) \Theta(x_1 > 2/3) \times \\ \times \Theta(x_1 < 1) \Theta(x_1 > x_2) \Theta(x_2 > 2(1-x_1)) \quad (\text{F.2.2})$$

$$= \int_0^1 dx_2 \frac{x_2^2 + T^2}{(1-x_2)(1-T)} \Theta(T > 2/3) \Theta(T < 1) \times \\ \times \Theta(T > x_2) \Theta(x_2 > 2(1-T)) \quad (\text{F.2.3})$$

$$= \int_{2(1-T)}^T dx_2 \frac{x_2^2 + T^2}{(1-x_2)(1-T)} \quad (\text{F.2.4})$$

dove $2/3 < T < 1$. Si ha anche:

$$I_2 = \int_{2/3}^1 dx_1 \int_{1-x_1}^{2(1-x_1)} dx_2 \frac{x_1^2 + x_2^2}{(1-x_1)(1-x_2)} \delta(T - (2 - x_1 - x_2)) \quad (\text{F.2.5})$$

$$= \int_0^1 dx_1 \int_0^1 dx_2 \frac{x_1^2 + x_2^2}{(1-x_1)(1-x_2)} \delta(T - 2 + x_1 + x_2) \Theta(x_1 > 2/3) \times \\ \times \Theta(x_1 < 1) \Theta(x_2 > 1 - x_1) \Theta(x_2 < 2(1 - x_1)) \quad (\text{F.2.6})$$

$$= \int_0^1 dx_1 \frac{x_1^2 + (2 - x_1 - T)^2}{(1-x_1)(-1+x_1+T)} \Theta(x_1 > 2/3) \Theta(x_1 < 1) \times \\ \times \Theta(x_1 < T) \Theta(T < 1) \quad (\text{F.2.7})$$

$$= \int_{2/3}^T dx_1 \frac{x_1^2 + (2 - x_1 - T)^2}{(1-x_1)(-1+x_1+T)}. \quad (\text{F.2.8})$$

Infine:

$$I_3 = \int_{1/2}^{2/3} dx_1 \int_{1-x_1}^{x_1} dx_2 \frac{x_1^2 + x_2^2}{(1-x_1)(1-x_2)} \delta(T - (2 - x_1 - x_2)) \quad (\text{F.2.9})$$

$$= \int_0^1 dx_1 \frac{x_1^2 + (2 - x_1 - T)^2}{(1-x_1)(-1+x_1+T)} \Theta(x_1 < 2/3) \Theta(x_1 > 1/2) \times \\ \times \Theta(x_1 > 2 - x_1 - T) \Theta(2 - T > 1) \quad (\text{F.2.10})$$

$$= \int_{1-T/2}^{2/3} dx_1 \frac{x_1^2 + (2 - x_1 - T)^2}{(1-x_1)(-1+x_1+T)}. \quad (\text{F.2.11})$$

Sommando i tre contributi, otteniamo (vedi notebook `Thrust_qqg.nb` allegato):

$$\frac{1}{\sigma} \frac{d\sigma}{dT} \Big|_{2/3 \leq T < 1} = \frac{\alpha_s}{2\pi} C_F 2(I_1 + I_2 + I_3) + \mathcal{O}(\alpha_s^2) \quad (\text{F.2.12})$$

$$= \frac{\alpha_s}{2\pi} C_F \left[\frac{2(3T^2 - 3T + 2)}{T(1-T)} \ln \left(\frac{2T-1}{1-T} \right) - \frac{3(3T-2)(2-T)}{(1-T)} \right] \quad (\text{F.2.13})$$

in cui in (F.2.12) il fattore 2 viene dalla regione $x_2 > x_1$.

F.3 Approfondimento Thrust all'endpoint

Avendo dedotto l'espressione per $d\sigma_R/dT$, possiamo scrivere direttamente la distribuzione totale NLO, incluso l'endpoint $T = 1$.

A $T = 1$ otteniamo, oltre al reale, un contributo di Born:

$$\frac{d\sigma_B}{dT} = \sigma_B \delta(1 - T) \quad (\text{F.3.1})$$

e un contributo virtuale:

$$\frac{d\sigma_V}{dT} = \sigma_V \delta(1 - T) \quad (\text{F.3.2})$$

tuttavia, poiché $\sigma_V = -\infty$, è necessario dare un senso all'espressione precedente.

Iniziamo riscrivendo il contributo reale come:

$$\frac{d\sigma_R}{dT} = \sigma_B \frac{\alpha_s}{2\pi} C_F \left[4 \frac{\ln(1/(1-T))}{1-T} - 3 \frac{1}{1-T} + g(T) \right] \Theta(T \geq 2/3)$$

dove la funzione $g(T)$ è data da:

$$g(T) = \frac{2}{T(1-T)} \left[(2 - 3T + 3T^2) \ln(2T - 1) - (2 - 5T + 3T^2) \ln(1 - T) \right] \quad (\text{F.3.3})$$

ed è **integrabile** nell'intervallo $T \in [2/3, 1]$.

Per una generica funzione $H(T)$ singolare in $T = 1$, introduciamo la cosiddetta "**prescrizione più**" (o "distribuzione più"), definita come:

$$\int_0^1 dT [H(T)]_+ f(T) \equiv \int_0^1 dT H(T) [f(T) - f(1)] \quad (\text{F.3.4})$$

integrabile in $T = 1$. Essa sostanzialmente prende una distribuzione singolare $H(T)$ e ne sottrae l'integrale (divergente) all'endpoint $T = 1$, come si può vedere riscrivendo quanto sopra nel seguente modo (sebbene matematicamente inadeguato):

$$= \int_0^1 dT \left[H(T) - \delta(1 - T) \left(\int_0^1 dy H(y) \right) \right] f(T). \quad (\text{F.3.5})$$

Vediamo immediatamente che la prescrizione più descrive precisamente la cancellazione fisica delle singolarità IR tra reale e virtuale. Il contributo reale è dato da una distribuzione divergente in $T \sim 1$ (rappresentata da $H(T)$ sopra), ma, a causa della **IR-safety** dell'osservabile T , questa divergenza è esattamente cancellata da un contributo negativo infinito che vive all'endpoint $T = 1$, che è il virtuale. La somma dei due contributi è integrabile in $T \sim 1$.

Queste considerazioni ci permettono finalmente di scrivere la **distribuzione NLO**:

$$\frac{d\sigma_{NLO}}{dT} \equiv \frac{d\sigma_B + d\sigma_R + d\sigma_V}{dT} \quad (\text{F.3.6})$$

$$= \sigma_B \left[\delta(1 - T) + \frac{\alpha_s}{2\pi} C_F \left[4 \left(\frac{\ln(1/(1-T))}{1-T} \right)_+ - 3 \left(\frac{1}{1-T} \right)_+ + g(T) + K \delta(1 - T) \right] \right] \Theta(T \geq 2/3) \quad (\text{F.3.7})$$

dove K è il contributo finito (cioè non IR-divergente) del virtuale, e può essere fissato semplicemente per normalizzazione, ovvero richiedendo che:

$$\int_0^1 dT \frac{d\sigma_{NLO}}{dT} = \sigma_{NLO} = \sigma_B \left[1 + \frac{3}{4} \frac{\alpha_s}{\pi} C_F \right]. \quad (\text{F.3.8})$$

Appendice G

Approfondimenti DIS

Vediamo in questo capitolo un paio di appunti riguardanti il Deep Inelastic Scattering analizzato nel capitolo [6.4](#).

G.1 Momentum sum rule

Vediamo in questa Appendice la verifica del fatto che le plus distribution sono necessarie per rispettare la regola di somma dell'impulso a qualsiasi scala μ_F . Richiediamo:

$$\mu_F^2 \frac{d}{d\mu_F^2} \int_0^1 dw w \left(\sum_{k=-N_F}^{N_F} f_k + f_g \right) (w, \mu_F) = 0 \quad (\text{G.1.1})$$

dove la somma è eseguita su tutti i sapori dei quark (e anti-quark). Sostituendo le equazioni di evoluzione, otteniamo:

$$\begin{aligned} &= \frac{\alpha_s}{2\pi} \int_0^1 dw w \int_w^1 \frac{dz}{z} \left[P_{qq}(z) \sum f_k \left(\frac{w}{z}, \mu_F \right) + 2N_F P_{qg}(z) f_g \left(\frac{w}{z}, \mu_F \right) \right. \\ &\quad \left. + P_{gq}(z) \sum f_k \left(\frac{w}{z}, \mu_F \right) + P_{gg}(z) f_g \left(\frac{w}{z}, \mu_F \right) \right] \quad (\text{G.1.2}) \end{aligned}$$

$$\begin{aligned} &= \frac{\alpha_s}{2\pi} \int_0^1 dw \int_0^1 dz \int_0^1 dy zy (P_{qq} + P_{gq})(z) \sum_{k=-N_F}^{N_F} f_k(y, \mu_F) \delta(zy - w) \\ &\quad + \frac{\alpha_s}{2\pi} \int_0^1 dw \int_0^1 dz \int_0^1 dy zy (2N_F P_{qg} + P_{gg})(z) f_g(y, \mu_F) \delta(zy - w) \quad (\text{G.1.3}) \end{aligned}$$

$$\begin{aligned} &= \frac{\alpha_s}{2\pi} \int_0^1 dz z (P_{qq} + P_{gq})(z) \int_0^1 dy y \sum f_k(y, \mu_F) \\ &\quad + \frac{\alpha_s}{2\pi} \int_0^1 dz z (2N_F P_{qg} + P_{gg})(z) \int_0^1 dy y f_g(y, \mu_F) = 0. \quad (\text{G.1.4}) \end{aligned}$$

Entrambi gli integrali in z sono zero, rispettando la regola di somma dell'impulso a tutte le μ_F ; esplicitamente:

$$\int_0^1 dz z(P_{qq} + P_{gq})(z) = C_F \int_0^1 dz \left[\frac{1+z^2}{1-z}(z-1) + \frac{1+(1-z)^2}{z}z \right] \quad (\text{G.1.5})$$

$$= 0 \quad (\text{G.1.6})$$

$$\int_0^1 dz z(2N_F P_{qg} + P_{gg})(z) = 0 \quad (\text{G.1.7})$$

G.2 Prodotto di convoluzione

In questa sezione vediamo qualcosa che potrebbe benissimo essere saltato, ma allo stesso tempo potrebbe essere utile a qualcuno.

Notazionalmente, si può introdurre un **prodotto di convoluzione** per descrivere le equazioni DGLAP in un modo più compatto e conveniente. Definendo:

$$[A \otimes B](x) \equiv \int_0^1 dz \int_0^1 dy A(z)B(y) \delta(zy - x) \quad (\text{G.2.1})$$

$$= \int_x^1 \frac{dz}{z} A(z) B\left(\frac{x}{z}\right) \quad (\text{G.2.2})$$

$$= \int_x^1 \frac{dt}{t} A\left(\frac{x}{t}\right) B(t) \quad (\text{G.2.3})$$

l'equazione DGLAP può essere riscritta come:

$$\mu_F^2 \frac{\partial f_i}{\partial \mu_F^2}(w, \mu_F) = \frac{\alpha_s}{2\pi} \sum_k [P_{ik} \otimes f_k](w, \mu_F) \quad (\text{G.2.4})$$

Il prodotto di convoluzione così definito è particolarmente utile perché si riduce a un **prodotto ordinario** sotto una trasformata integrale. Introducendo la **trasformata di Mellin**:

$$\hat{f}(N) \equiv \int_0^1 dx x^{N-1} f(x) \quad (\text{G.2.5})$$

(che è essenzialmente una trasformata di Laplace $\mathcal{L}_N[f]$ con $e^{-t} \rightarrow x$), si ottiene:

$$[\widehat{A \otimes B}](N) = \int_0^1 dx x^{N-1} \int_0^1 dz \int_0^1 dy A(z)B(y) \delta(zy - x) \quad (\text{G.2.6})$$

$$= \int_0^1 dz z^{N-1} A(z) \int_0^1 dy y^{N-1} B(y) = \hat{A}(N) \cdot \hat{B}(N). \quad (\text{G.2.7})$$

Questa tecnica costituisce la base per la soluzione analitica delle equazioni DGLAP.

Bibliografia

- [1] R. Barbieri. *Ten Lectures on the ElectroWeak Interactions*. URL: <https://arxiv.org/abs/0706.0684v1>.
- [2] R. Keith Ellis, W. James Stirling e B. R. Webber. *QCD and collider physics*. Cambridge University Press, 2011. ISBN: 978-0-511-82328-2, 978-0-521-54589-1. DOI: [10.1017/CB09780511628788](https://doi.org/10.1017/CB09780511628788). URL: <https://www.cambridge.org/core/books/qcd-and-collider-physics/D0095E6D278BBBC74E9C3636AB4CB80C>.
- [3] John Iliopoulos e Theodore N. Tomaras. *Elementary Particle Physics*. Oxford University Press, 2021.
- [4] Taizo Muta. *Foundations of quantum chromodynamics: an introduction to perturbative methods in gauge theories*. world scientific, 1998.
- [5] P. Nason. *Introduction to perturbative QCD*. URL: https://virgilio.mib.infn.it/~re/extra_material_phenomenology/lecture_notes_QCD_part/lectures_QCD_Nason.pdf.
- [6] Michael E. Peskin e Daniel V. Schroeder. *An Introduction to quantum field theory*. Reading, USA: Addison-Wesley, 1995.
- [7] G. Ridolfi. *An Introduction to the Standard Model of Electroweak interaction*.
- [8] G. P. Salam. *Elements of QCD for hadron colliders*. URL: <https://arxiv.org/abs/1011.5131>.
- [9] M. H. Seymour. *Quantum ChromoDynamics*. URL: <https://arxiv.org/abs/hep-ph/0505192>.
- [10] J.B. Zuber e C. Itzykson. *Quantum Field Theory*. Dover Books on Physics. Dover Publications, 2012. ISBN: 9780486134697.