

IFT

Instituto de Física Teórica
Universidade Estadual Paulista

DISSERTAÇÃO DE MESTRADO

IFT-D.006/2000

Efeitos Virtuais de Leptoquarks
no Processo de Drell-Yan $p\bar{p} \rightarrow e^+e^-X$

Alexandre Alves



Orientador

Prof. Dr. Oscar José Pinto Éboli

Maio de 2000

Agradecimentos

Em primeiro lugar agradeço ao nosso Senhor Jesus Cristo por nunca ter me abandonado e nunca ter me deixado abater nas horas difíceis. Até aqui o Senhor me trouxe e me deu vitória.

Ao Oscar pela sua imensa dedicação e paciência ao longo desses quase 4 anos de convivência desde os tempos de minha iniciação científica. Todos os teus ensinamentos me fizeram crescer muito, valeu Oscar!

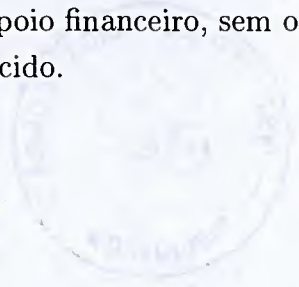
Agradeço muito aos meus pais Márcia e Celso por todo o apoio e carinho ao longo dos anos. Sem vocês jamais teria conseguido coisa alguma, amo-os demais.

Aos meus amigos de batalha Sérgio, Guilherme, Victor e Antônio.

Ao professor Sérgio Novaes pelos valiosos ensinamentos e pelo apoio.

Em especial, quero agradecer à minha bela esposa Elaine por toda a sua dedicação, paciência e amor. Te amo demais!

Por último gostaria de agradecer à FAPESP pelo apoio financeiro, sem o qual o desenvolvimento desse trabalho não poderia ter acontecido.



Resumo

Correções em *Next-to-Leading-Order (NLO)* de *QCD* são calculadas ao processo de Drell-Yan $p\bar{p} \rightarrow e^+e^-X$ com a inclusão de leptquarks escalares virtuais trocados no canal t da reação. As correções de ordem α_s foram calculadas usando um método híbrido que combina resultados analíticos e integração via Monte Carlo para as integrais do espaço de fase de 2 e 3 partículas. As correções de *QCD* são positivas e aumentam as seções de choque totais em cerca de 40% no esquema DIS. A dependência das seções de choque em relação às escalas de renormalização e fatorização não é reduzida em NLO a exemplo de outros processo estudados no Tevatron. Os leptquarks, por sua vez, desacoplam-se rapidamente com o aumento de sua massa, alterando as seções de choque em menos de 1.5% para massas acima de 250 GeV, acoplamentos de Yukawa da ordem da constante de acoplamento eletromagnética e para massas invariantes do par e^+e^- maiores do que 120 GeV.

Palavras Chaves: Correções de *QCD* ; Processo de Drell-Yan ; Leptoquarks Escalares.

Áreas do conhecimento: Física de Partículas Elementares.

Abstract

Next-to-Leading-Order (NLO) QCD corrections are evaluated to the Drell-Yan process $p\bar{p} \rightarrow e^+e^-X$ including the virtual effects of scalar leptoquarks being exchanged in t channel of the reaction. The α_s corrections were obtained utilizing a hybrid method that combines analytical results and Monte Carlo integration of the 2 and 3 particle phase space integrals. The QCD corections are positive and increase the total cross sections about 40% in the DIS scheme. The cross section dependence on the renormalization and fatorization scales do not decrease at NLO just like other processes studied at Tevatron. On the other hand, leptoquarks decouple rapidly with the increasing of their masses , changing less than 1.5% the cross sections for masses above 250 GeV, Yukawa couplings of order of the eletromagnetic one, and e^+e^- invariant masses above 120 GeV.

Índice

1	Introdução	1
2	A Cromodinâmica Quântica	3
2.1	A QCD como um Modelo para as Interações Fortes	3
2.2	O Princípio de Gauge	5
2.3	A Lagrangiana da QCD	8
2.4	Alguns Aspectos da Quantização e Renormalização da QCD	12
2.4.1	Ghosts e Simetria BRST	13
2.4.2	Lagrangiana Renormalizada	18
2.4.3	Liberdade Assintótica	20
3	Correções de QCD ao Processo $e^+e^- \rightarrow$ Hádrons	22
3.1	O Aniquilamento e^+e^- no Modelo Padrão	22
3.2	Cálculos em Ordem Zero	23
3.3	Correções de Ordem α_s : Cálculo Analítico Completo	24
3.3.1	Emissão de Glúon Real	25
3.3.2	Correções Virtuais	28
4	O Formalismo NLO Monte Carlo para Correções de QCD	38
4.1	Descrição do Formalismo	38
4.2	O Aniquilamento $e^+e^- \rightarrow$ Hádrons como Exemplo	41
4.2.1	Processo de Ordem Zero	41
4.2.2	Correções Virtuais	42
4.2.3	Radiação de Glúons Reais	42
4.2.4	Seção de Choque de 2 Corpos	44
4.2.5	Seção de Choque de 3 Corpos	44
5	Correções em <i>NLO</i> de QCD para as Funções de Distribuição de Pártons	46
5.1	O Modelo a Pártons	46

5.2	As Funções de Distribuição de Pártons	48
5.2.1	Universalidade das PDFs	50
5.2.2	Contribuição de Eventos Multipartônicos	52
5.2.3	Correlação entre as funções de onda dos hádrons: soft glúons .	53
5.2.4	Breve Análise dos Higher Twists Através da Expansão em Produtos de Operadores	55
5.3	As Correções da QCD Perturbativa para o DIS	59
5.3.1	Glúons Virtuais	66
5.3.2	Glúon no Estado Inicial	67
5.3.3	Emissão de Glúon Real	68
5.4	PDFs em Ordem α_s	71
6	Correções de Ordem α_s ao Processo de Drell-Yan $p\bar{p} \rightarrow e^+e^-X$	75
6.1	Drell-Yan em Leading-Log Approximation	75
6.2	Contribuições Virtuais ao Nível de 1-Loop	78
6.2.1	Autoenergia dos Quarks	79
6.2.2	Correção do Vértice $(\gamma, Z) q\bar{q}$	81
6.2.3	Contribuição Virtual Total	82
6.3	Radiação de Glúons Reais	83
6.3.1	Emissão de Glúon Soft	85
6.3.2	Emissão de Glúon Colinear	87
6.4	Glúon no Estado Inicial	89
6.5	Seções de Choque Colineares	93
6.6	Fatorização de Massa	95
6.6.1	Termo Hard-Colinear	98
6.6.2	Termo Soft-Colinear	99
6.7	Contribuição de 2 Corpos à Seção de Choque Hadrônica	100
6.8	Seção de Choque de 3 Corpos	101
6.8.1	Glúons Iniciais	101
6.8.2	Emissão de Glúon	101
7	Introdução de Leptoquarks Escalares	103
7.1	Extensões do Modelo Padrão com Acoplamentos Quark-Lépton . . .	103
7.1.1	Motivação de Estudo	103
7.1.2	Lagrangiana de Interação de Leptoquarks	106
7.2	Efeitos Virtuais de Leptoquarks Escalares no Processo Drell-Yan ao Nível de Árvore	109
7.3	Correções Virtuais em Ordem α_s	111

7.3.1	Autoenergia do Leptoquark	112
7.3.2	Correção do Vértice leptoquark–quark–lépton	114
7.3.3	Contribuição Virtual da Caixa	115
7.4	Seção de Choque Virtual Total	121
7.5	Radiação de Glúons Reais	121
7.5.1	Emissão de Glúon Soft	122
7.5.2	Emissão de Glúon Colinear	122
7.6	Estados Finais com Um Quark Não-Massivo Adicional	124
7.7	Termos Hard-Colinear	126
7.8	Termo Soft-Colinear	127
7.9	Seção de Choque Total Hadrônica de 2 Corpos	127
7.10	Seção de Choque Total Hadrônica de 3 Corpos	128
7.10.1	Subtração dos Estados Intermediários On-Shell	129
8	Análises Numéricas	132
8.1	Variação da Seção de Choque Total com os Cut-Offs Teóricos δ_s e δ_c .	132
8.2	Efeitos das Correções de QCD Sobre as Seções de Choque do Processo de Drell-Yan	133
8.3	Estabilidade das Seções de Choque em NLO em Relação à Escolha das Escalas de Renormalização e Fatorização	134
8.4	Distribuições	135
8.5	Região de Exclusão no Espaço de Parâmetros $\lambda \times M_{\ell q}$	137
9	Conclusão	147
A	Regras de Feynman	149
B	Álgebra do Grupo $SU(N)$	153
C	Espaços de Fase em d Dimensões	155
C.1	Definição do Espaço de Fase de N Partículas	155
C.2	Espaço de Fase de 1 Partícula	156
C.3	Espaço de Fase de 2 Partículas Não Massivas	156
C.4	Espaço de Fase de 3 Partículas Não Massivas	157
C.4.1	Espaço de Fase do Decaimento $\gamma^* \rightarrow q\bar{q}g$	159
D	Fórmulas e Integrais Úteis	161
E	Funções Escalares de Passarino-Veltman	163

Referências

168

Capítulo 1

Introdução

O modelo padrão das interações fortes e eletrofracas (SM) tem passado com grande sucesso por todos os testes a que tem sido submetido, explicando qualitativa e quantitativamente todos os dados experimentais obtidos até agora em física de partículas elementares. O grau de concordância teórica e experimental é tão bom que efeitos quânticos de pequena intensidade, devidos a ordens mais altas de teoria de perturbação, podem ser testados com grande precisão em laboratório, exigindo um esforço teórico crescente no cálculo de quantidades observáveis em ordens superiores ao nível de árvore.

Nesse sentido, as correções de ordem superior da Cromodinâmica Quântica (QCD), que corresponde ao setor do modelo padrão das interações fortes entre quarks e glúons, são as mais importantes em processos onde essas partículas estão presentes. Em termos experimentais, a QCD é uma teoria altamente bem sucedida na descrição de fenômenos envolvendo interações fortes no regime de altas energias acessado pelos atuais aceleradores de partículas. Correções de QCD às seções de choque envolvendo hádrons são tipicamente positivas e grandes, chegando, em alguns casos, a ser maiores do que as seções de choque de Born. Além disso, o conhecimento das correções de ordem superior é necessário para reduzir-se incertezas nas previsões teóricas.

Por essas razões, as correções radiativas da QCD são especialmente importantes no cálculo de processos onde investiga-se a existência de nova física além daquela do SM, dado que essas correções podem realçar significativamente as seções de choque devidas a novos fenômenos. Além disso, uma determinação teórica acurada é importante para determinar-se as massas, acoplamentos e outros observáveis físicos de eventuais partículas previstas por extensões do SM.

Entre as diversas extensões estudadas atualmente, existem algumas que descrevem quarks e léptons em um mesmo nível relacionando-os de uma maneira mais fundamental gerando acoplamentos diretos entre essas partículas. Essas interações,

que não ocorrem no modelo padrão, são intermediadas por partículas que carregam números quânticos leptônico e bariônico, e que podem, portanto, acoplar-se a todos os bósons de gauge conhecidos, incluindo o glúon. Tais partículas, se existirem, podem modificar as seções de choque dos processos preditos pelo modelo padrão e também produzir processos novos, ainda não observados experimentalmente.

Nessa dissertação de Mestrado, em especial, analisamos os efeitos virtuais de leptoquarks escalares de primeira geração no processo de Drell-Yan $p\bar{p} \rightarrow e^+e^-X$ com correções de NLO de QCD utilizando um método híbrido que combina resultados analíticos e integração numérica para parte do espaço de fase. Neste tipo de processo, os leptoquarks modificam as seções de choque de forma indireta, ou seja, não são produzidos diretamente, sendo trocados no canal t da reação.

A dissertação está organizada da seguinte maneira.

O capítulo 2 é dedicado a uma revisão sucinta dos aspectos formais da QCD.

No capítulo 3, correções de QCD à produção de hádrons em colisões elétron-pósitron são calculadas analiticamente em NLO visando exemplificar os aspectos básicos da QCD perturbativa.

A descrição precisa do método utilizado para o cálculo das correções de QCD ao processo de Drell-Yan com inclusão de leptoquarks escalares é dada no capítulo 4, onde mostra-se a equivalência dos resultados obtidos com o método e os resultados analíticos obtidos no capítulo 3 para o processo $e^+e^- \rightarrow$ hádrons de uma maneira exata, calculando-se analiticamente todas as integrais do espaço de fase de 2 e 3 partículas.

No capítulo 5, calculamos as correções de QCD às funções de distribuição de quarks no âmbito do espalhamento inelástico profundo, tendo já em mente as correções ao processo de Drell-Yan. Discutimos também alguns aspectos formais a respeito da universalidade das densidades de párons.

Os capítulos 6 e 7 contém o cálculo completo das correções de QCD para o processo de Drell-Yan $p\bar{p} \rightarrow e^+e^-X$ no modelo padrão sem a inclusão de leptoquarks e com a inclusão de leptoquarks respectivamente, utilizando o método descrito no capítulo 4.

Análises numéricas são encontradas no capítulo 8. No capítulo 9 discutimos de forma geral os resultados obtidos em nosso trabalho.

Capítulo 2

A Cromodinâmica Quântica

2.1 A QCD como um Modelo para as Interações Fortes

Há aproximadamente 30 anos atrás a *Cromodinâmica Quântica*, ou *QCD**, como é mais conhecida, foi concebida como um modelo das interações fortes após várias décadas de esforços teóricos e experimentais.

Talvez o primeiro grande avanço na direção da QCD tenha sido o sucesso do *modelo a quarks* [1], proposto por Gell-Mann e Zweig, onde partículas denominadas *quarks* eram postuladas como os constituintes básicos das partículas fortemente interagentes (hádrons). O modelo era capaz de descrever com sucesso o rico e variado padrão da espectroscopia hadrônica, onde o grande número de estados mesônicos e bariônicos claramente revelava a existência de partículas mais fundamentais. Hoje em dia sabemos que os quarks são partículas de spin $1/2$ classificadas em termos de 3 famílias com 2 sabores cada, em um total de 6 tipos diferentes.

Assumindo que mésons são estados formados por pares $q\bar{q}$ enquanto bárions são estados do tipo qqq , podemos verificar que o modelo descreve todo o espectro hadrônico. A confirmação da existência dos sabores mais leves veio na década de 60 com os experimentos do SLAC em colisões inelásticas profundas do tipo elétron-núcleon.

Contudo, o modelo a quarks, como formulado originalmente, não era capaz de descrever a estatística correta de alguns estados hadrônicos. Para obedecer ao *Princípio de Pauli*, os quarks têm de possuir um número quântico extra [2], que foi batizado de *número quântico de cor*. Desse modo, cada sabor de quark pode ter $N = 3$ cores diferentes. As funções de onda dos mésons e bárions são descritos por singletos de cor, levando-se em conta o fato de que quarks nunca foram observados

*Do inglês Quantum Chromodynamics.

isoladamente, e podem ser escritas como

$$|\text{Báron} \rangle = \frac{1}{\sqrt{6}} \varepsilon^{\alpha\beta\gamma} |q_\alpha q_\beta q_\gamma\rangle, \quad |\text{Méson} \rangle = \frac{1}{\sqrt{3}} \delta^{\alpha\beta} |q_\alpha \bar{q}_\beta\rangle, \quad (2.1)$$

onde α, β, γ são índices de cor.

Posteriormente, a medida da razão de decaimento do processo $\pi^0 \rightarrow \gamma\gamma$ e da seção de choque do processo $e^-e^+ \rightarrow$ hádrons, entre outras evidências, confirmaram como sendo 3 o número de cores dos quarks.

Apesar disso, como dissemos, quarks nunca foram observados isoladamente em laboratório, o que levou à hipótese de *confinamento de cor*, onde todos os estados assintóticos hadrônicos são singletos sob rotações no espaço de cores.

É interessante frisar que a hipótese de confinamento implica uma forma totalmente anti-simétrica para as funções dos bárions através das funções de onda de cor. Além disso a hipótese de confinamento só permite a existência de estados ligados de quarks do tipo [3]

$$(3q)^m (q\bar{q})^n, \quad (m, n \geq 0), \quad (2.2)$$

onde $m = 0, n = 1$ correspondem aos mésons, enquanto $m = 1, n = 0$ aos bárions. Outras configurações, apesar de permitidas pela hipótese do confinamento de cor, não são observadas experimentalmente, todos os hádrons observados são previstos pelo modelo a quarks mais simples.

A hipótese de confinamento, contudo, estava firmada em bases fenomenológicas, não havendo na época do surgimento do modelo a quarks nenhuma evidência teórica para sustentar a hipótese. De uma forma geral, a natureza das interações fortes ainda não estava clara.

Motivado pelo sucesso da álgebra de correntes, Gell-Mann sugere que a teoria quântica de campo das interações fortes é uma teoria de gauge abeliana, onde o glúon, o bóson de gauge da interação forte, não possui carga de cor. Nambu, por outro lado, assumiu que os glúons formavam a representação octeta do grupo de cor $SU(3)_C$. Esses conceitos qualitativos, no entanto, não podiam ainda ser colecionados em uma teoria de campo consistente que produzisse resultados quantitativos.

Feynman, então, propõe o seu *modelo a pártons* [4], onde os constituintes dos nucleons são partículas essencialmente pontuais e livres, e que descreve com boa precisão os resultados das experiências de espalhamento inelástico profundo. Os números quânticos dos pártons carregados podiam ser extraídos da experiência, e confirmou-se que eles eram os quarks recém inventados para descrever a espectroscopia hadrônica. Apesar disso, parecia que uma teoria fundamental de interações fortes estava ainda mais longe de ser obtida, pois o modelo a pártons estava baseado na hipótese de estados aproximadamente livres a curtas distâncias dentro de funções

de onda de estados ligados, ou seja, os quarks, apesar de estarem permanentemente ligados dentro dos hádrons, comportavam-se como partículas livres às curtas distâncias em que eram sondados por fótons de grande energia. Como o próprio Feynman ressaltara, aquilo não era consistente com a classe conhecida de teorias de campo renormalizáveis. De fato, naquela época, a única teoria de campo realística renormalizável era a Eletrodinâmica Quântica (QED), a qual sabia-se possuir um comportamento exatamente oposto a grandes energias.

O impasse, no entanto, estava para ser resolvido. Em 1971, 't Hooft prova que as teorias de gauge não-abelianas são renormalizáveis e ajuda os construtores de modelos a colocar juntos os conceitos de quarks, cores e sabores como os ingredientes básicos de uma teoria de gauge para as interações fortes.

Dois anos mais tarde, Gell-Mann *et al.* [5] apontam as vantagens do modelo octeto de cor para os glúons, argumentando que tal modelo tinha uma melhor concordância com as características qualitativas das interações fortes observadas experimentalmente. Alguns dias antes da publicação desse artigo na *Physics Letters*, Gross e Wilczek publicavam um trabalho [6] onde demonstravam que as teorias de Yang e Mills, entre as quais a QCD, possuem uma propriedade notável: a curtas distâncias a interação entre as partículas é cada vez menor e a teoria é bem definida, ao contrário da QED, por exemplo. A QCD, portanto, era uma teoria apta a descrever a liberdade assintótica das interações fortes. Logo em seguida, o *scaling de Bjorken* e o modelo a párons foram compreendidos nesse contexto. Isso tudo permitiu o estudo do regime perturbativo da QCD e suas previsões em processos de espalhamento envolvendo hádrons.

A partir de agora, vamos abordar a QCD de um ponto de vista mais matemático, em busca de uma teoria de campo realista das interações fortes.

2.2 O Princípio de Gauge

O ponto de partida para a construção de uma lagrangiana que descreva as interações fortes atribuindo-lhe as propriedades de liberdade assintótica e confinamento de cor é saber como introduzir interações invariantes por transformações de gauge não-abelianas. Por simplicidade, consideremos o caso da QED, primeiramente.

Considere a lagrangiana que descreve um campo fermiônico de massa m livre de interações

$$\mathcal{L}_0 = \bar{\psi}(x)(i \not{\partial} - m)\psi(x) . \quad (2.3)$$

A lagrangiana $\mathcal{L}_0$ é invariante sob uma transformação de fase global

$$\psi'(x) \rightarrow e^{-iQ\alpha}\psi , \quad (2.4)$$

e que corresponde a uma *transformação de gauge global* do grupo $U(1)$, onde $Q\alpha$ é uma constante real e arbitrária. A fase de $\psi(x)$ é, portanto, uma questão puramente de convenção, e uma vez escolhido o seu valor em um dado ponto do espaço-tempo x , o mesmo valor deve ser atribuído em todos os outros pontos do espaço-tempo, mostrando o caráter global da transformação. Yang e Mills, contudo, já em 1954, escreveram, em um artigo fundamental para o desenvolvimento das teorias de gauge [7],

“It seems that this is not consistent with the localized field concept that underlies the usual physical theories. In the present paper we wish to explore the possibility of requiring all interactions to be invariant under independent rotations of the isotopic isospin at all space-time points.”

A proposta de Yang e Mills consiste na exigência de que a teoria seja invariante sob transformações de fase *locais*, ou seja, de que a convenção da escolha do valor da fase seja feita de maneira independente em cada ponto do espaço-tempo, o que parece ser muito mais natural dada a natureza localizada dos campos. No entanto, permitindo que α varie de ponto a ponto no espaço-tempo, destruimos a invariância de $\mathcal{L}_0$ sob transformações locais do grupo $U(1)$, pois

$$\partial_\mu(e^{-iQ\alpha(x)}\psi(x)) = e^{-iQ\alpha(x)}[\partial_\mu\psi(x) - iQ\partial_\mu\alpha(x)\psi(x)] \quad (2.5)$$

contém, agora, um termo adicional dado pela derivada da fase.

A única maneira de conseguirmos uma invariância sob transformações de fase locais é adicionar um termo a $\mathcal{L}_0$ de modo a cancelar o termo adicional proporcional a $\partial_\mu\alpha(x)$. A modificação necessária está completamente fixada pela lei de transformação de $\psi(x)$: a saída é introduzir um campo vetorial $A_\mu(x)$ (já que $\partial_\mu\alpha(x)$ é um 4-vetor) que se transforme de tal maneira que produza um termo de sinal oposto ao termo adicional de (2.5). A variação de $\mathcal{L}_0$ é dada por

$$\Delta\mathcal{L}_0 = \mathcal{L}'_0 - \mathcal{L}_0 = Q\bar{\psi}(x)\gamma^\mu\psi(x) \cdot \partial_\mu\alpha(x) , \quad (2.6)$$

e isso fixa a forma do termo a ser adicionado à lagrangiana original como sendo

$$-eQ\bar{\psi}(x)\gamma^\mu\psi(x) A_\mu(x) . \quad (2.7)$$

Adicionando esse termo em $\mathcal{L}_0$ temos

$$\mathcal{L} = \mathcal{L}_0 - eQ\bar{\psi}(x)\gamma^\mu\psi(x) A_\mu(x) , \quad (2.8)$$

de modo que a variação de $\mathcal{L}$, dada por

$$\Delta\mathcal{L} = \Delta\mathcal{L}_0 - eQ\bar{\psi}(x)\gamma^\mu\psi(x) \cdot \delta A_\mu(x) , \quad (2.9)$$

onde $\delta A_\mu(x) = A'_\mu(x) - A_\mu(x)$. Para que $\Delta\mathcal{L} = 0$, ou seja, para que a lagrangiana modificada pela introdução do campo vetorial $A_\mu(x)$, que se acopla aos campos fermiônicos de acordo com (2.7), seja invariante sob transformações locais do grupo $U(1)$, $\delta A_\mu(x)$ deve ser da forma

$$\delta A_\mu(x) = \frac{1}{e} \partial_\mu \alpha(x) . \quad (2.10)$$

Impondo somente a invariância sob transformações locais de $U(1)$, a transformação de $A_\mu(x)$ e o seu acoplamento com os campos de matéria estão fixados! Em outras palavras, o *princípio de gauge* determina a dinâmica da teoria forçando o acoplamento dos campos de matéria com campos de gauge, nesse caso um campo vetorial, de forma que a lagrangiana resultante é invariante sob transformações de gauge locais.

A introdução do campo vetorial $A_\mu(x)$ define o que chamamos de *derivada covariante*,

$$D_\mu \equiv \partial_\mu - ieQA_\mu(x) , \quad (2.11)$$

que possui a propriedade fundamental de se transformar exatamente como $\psi(x)$ sob uma *transformação global*

$$(D_\mu \psi(x))' = e^{-iQ\alpha(x)} D_\mu \psi(x) , \quad (2.12)$$

e em termos da qual a lagrangiana $\mathcal{L}$ adquire a forma

$$\mathcal{L} = \bar{\psi}(x)(i \not{D} - m)\psi(x) = \mathcal{L}_0 - eQ\bar{\psi}(x)\gamma^\mu\psi(x) A_\mu(x) . \quad (2.13)$$

A lagrangiana $\mathcal{L}$ descreve justamente a interação férmion-fóton da QED. No entanto, uma teoria completa necessita dos termos que descrevem a propagação do campo A_μ do fóton. Os termos cinéticos são obtidos pela introdução de

$$\mathcal{L}_{cin} = -\frac{1}{4}F_{\mu\nu}F^{\mu\nu} \quad (2.14)$$

onde

$$F_{\mu\nu} = \partial_\mu A_\nu - \partial_\nu A_\mu \quad (2.15)$$

é o tensor de campo eletromagnético usual.

Um fato importante ocorre aqui: um termo de massa para o campo A_μ violaria a invariância de gauge de $\mathcal{L}$. Impondo o *princípio de gauge*, A_μ deve ser um campo vetorial sem massa. Por outro lado, como vimos, um termo de massa para os férmions é permitido.[†]

[†]Isso nem sempre é permitido. No modelo GSW, que descreve o setor eletrofraco do modelo padrão, a lagrangiana não pode conter termos de massa para os férmions antes da quebra espontânea de simetria.

A lagrangiana $\mathcal{L}_{em} = \mathcal{L} + \mathcal{L}_{cin}$ é a lagrangiana completa da QED, a partir da qual qualquer fenômeno eletromagnético envolvendo férmions pode ser descrito. As predições da QED têm sido testadas em um alto grau de acurácia, o que faz dela uma das melhores teorias de campo já criadas e, o que é mais importante neste contexto, é uma teoria cuja dinâmica provém diretamente do *princípio de gauge*.

2.3 A Lagrangiana da QCD

Inspirados no exemplo de QED, vamos estender o *princípio de gauge* às interações fortes em busca de uma lagrangiana invariante de gauge para a QCD.

1. Escolha dos campos de matéria

O primeiro passo na construção de uma teoria de gauge é a escolha dos campos de matéria. No caso da QED, os campos de matéria eram os campos fermiônicos $\psi(x)$, que podem ser atribuídos ao elétron, ou outro férmion carregado. Para a QCD, a escolha desses campos é óbvia: os quarks, que são os únicos tipos de férmions elementares que sofrem interações fortes. Sabemos que cada sabor de quark pode assumir 3 cores diferentes. Vamos, então, denotar um tal campo por $q^\alpha(x)$, onde $q = u, d, s, c, b, t$ indica o seu sabor, $\alpha = 1, 2, 3$ o seu número quântico de cor e m_q sua massa. Para simplificar as equações, vamos adotar uma notação vetorial no espaço das cores: $q^T \equiv (q^1, q^2, q^3)$. A lagrangiana livre é dada, portanto, por

$$\mathcal{L}_0 = \sum_q \bar{q}(x)(i \not{\partial} - m_q)q(x) , \quad (2.16)$$

em outras palavras, os campos de matéria são escolhidos de forma que os quarks pertençam à representação fundamental 3-dimensional do grupo $SU(3)_C$ de cor. Essa lagrangiana é invariante sob a transformação global

$$(q^\alpha)' = U^\alpha_\beta q^\beta , \quad UU^\dagger = U^\dagger U = 1 , \quad \det U = 1 , \quad (2.17)$$

onde as matrizes unitárias U do $SU(3)_C$ podem ser escritas na forma

$$U = \exp\left\{-ig_s \frac{\lambda^a}{2} \theta_a\right\} , \quad (2.18)$$

onde as matrizes λ^a ($a = 1, 2, \dots, 8 = 3^2 - 1$) denotam os geradores da representação da álgebra do $SU(3)_C$, e θ_a são parâmetros arbitrários. Ao contrário da QED, onde o grupo de gauge é o $U(1)_{em}$, que é um grupo abeliano, o grupo de simetria da QCD é um grupo não-abeliano, onde os geradores da álgebra não comutam entre si. A

QCD, assim como a teoria eletrofraca (EW), é descrita por uma *teoria de gauge não-abeliana*. Somente essa classe de teorias renormalizáveis possui o comportamento de liberdade assintótica. As matrizes λ^a , denominadas *matrizes de Gell-Mann* têm traço nulo e satisfazem às seguintes relações de comutação

$$[\lambda^a, \lambda^b] = 2if^{abc}\lambda^c, \quad (2.19)$$

sendo f^{abc} as constantes de estrutura do grupo $SU(3)_C$, que são reais e totalmente anti-simétricas.[†]

2. Aplicação do princípio de gauge

O próximo passo na construção da lagrangiana da QCD é promover a simetria global do grupo $SU(3)_C$ a uma simetria local, onde $\theta_a = \theta_a(x)$. O ponto de partida é a derivada covariante. Como agora temos 8 parâmetros de gauge independentes, 8 bósons de gauge diferentes G_a^μ , os glúons, precisam ser introduzidos. A derivada pode, então, ser escrita como

$$D_\mu = \partial_\mu - ig_s T^a G_a^\mu \equiv \partial_\mu - ig_s G^\mu, \quad (2.20)$$

onde $T^a = \frac{\lambda^a}{2}$ e introduzimos a notação matricial $G^\mu = T^a G_a^\mu$.

A derivada covariante D_μ introduz o acoplamento glúon-quark, de modo que a lagrangiana

$$\mathcal{L} = \sum_q \bar{q}(x)(i \not{D} - m_q)q(x) \quad (2.21)$$

é invariante sob transformações locais do grupo $SU(3)_C$ se $D^\mu q(x)$ varia como o vetor $\partial_\mu q(x)$ sob transformações globais do grupo de gauge. Isso fixa a lei de transformação dos campos de gauge, pois se

$$[\partial_\mu q(x)]' = \partial_\mu q'(x) = U \partial_\mu q(x)$$

sob transformações globais, então, impondo que

$$D'_\mu q'(x) = U(x)[D_\mu q(x)], \quad (2.22)$$

o que é equivalente a

$$D'_\mu = U(x)D_\mu U^\dagger(x), \quad (2.23)$$

implica a lei de transformação

$$G'_\mu = U(x)G_\mu U^\dagger(x) - \frac{i}{g_s}(\partial_\mu U(x))U^\dagger(x). \quad (2.24)$$

[†]Maiores detalhes sobre as propriedades dessas matrizes podem ser encontradas no Apêndice B.

Sob uma transformação infinitesimal do $SU(3)_C$, onde

$$U(x) \simeq 1 - ig_s T^a \delta\theta^a(x) , \quad (2.25)$$

as leis de transformação dos campos e da derivada covariante são dadas por

$$\delta q(x) = q'(x) - q(x) = -ig_s T^a \delta\theta^a(x) q(x) , \quad (2.26)$$

$$\delta G_\mu^a = (G_\mu^a)' - G_\mu^a = -\partial_\mu(\delta\theta^a(x)) + g_s f^{abc} \delta\theta^b(x) G_\mu^c , \quad (2.27)$$

$$\delta(D_\mu q) = -ig_s T^a \delta\theta^a(x) \cdot D_\mu q . \quad (2.28)$$

Note que a lei de transformação do campo vetorial do glúon envolve as constantes de estrutura do grupo, o que significa que os campos do glúon pertencem à representação adjunta do grupo $SU(3)_C$ para $\delta\theta^a$ constante. Outro ponto muito importante é o fato de que há uma única constante de acoplamento g_s em todas as leis de transformação. Isso é bem diferente do caso do eletromagnetismo onde cada campo carregado pode se acoplar ao fóton com sua própria carga elétrica ($e, -e, 2e, -3e$, etc). É possível mostrar que, se δG_μ^a é especificado com uma constante g_s , então se δq e a derivada covariante forem definidas com uma outra constante g'_s , só é possível obter $\delta(D_\mu q)$, como dada por (2.28), se $g'_s = g_s$. Todos os campos de gauge se acoplam de uma maneira *universal* aos campos de matéria, não só na QCD, mas em qualquer teoria de gauge não-abeliana.

Resta construir o termo cinético invariante de gauge para os glúons. Para tanto, vamos voltar ao caso do eletromagnetismo. Naquela ocasião, os termos cinéticos em (2.14) do fóton foram construídos com base no tensor de campo eletromagnético $F_{\mu\nu}$ dado em (2.15). A idéia é generalizar aquelas expressões para o caso de glúon. É preciso, então, obter o tensor de campo gluônico e isso é feito se notarmos que $F_{\mu\nu}$ pode ser escrito em termos das derivadas covariantes do eletromagnetismo

$$F_{\mu\nu} = \frac{i}{eQ} [D_\mu, D_\nu] , \quad (2.29)$$

de maneira que o tensor de campo gluônico é escrito como[§]

$$\begin{aligned} G_{\mu\nu} &= \frac{i}{g_s} [D_\mu, D_\nu] = \partial_\mu G_\nu - \partial_\nu G_\mu - ig_s [G_\mu, G_\nu] \equiv T^a G_{\mu\nu}^a , \\ G_{\mu\nu}^a &= \partial_\mu G_\nu^a - \partial_\nu G_\mu^a + g_s f^{abc} G_\mu^b G_\nu^c . \end{aligned} \quad (2.30)$$

Sob uma transformação de gauge infinitesimal,

$$\delta G_{\mu\nu}^a = f^{abc} \delta\theta^b(x) G_{\mu\nu}^c , \quad (2.31)$$

[§]A generalização mais simples possível $G_{\mu\nu}^a = \partial_\mu G_\nu^a - \partial_\nu G_\mu^a$ não é invariante de gauge.

e, consequentemente

$$\begin{aligned}\delta(G_{\mu\nu}^a G^{a\mu\nu}) &= 2G_{\mu\nu}^a \delta G^{a\mu\nu} \\ &= 2f^{abc} \delta\theta^b(x) G_{\mu\nu}^a G^{c\mu\nu} = 0\end{aligned}\quad (2.32)$$

pois envolve a contração de um tensor anti-simétrico com um simétrico nos índices de cor.

Adotando a normalização convencional do termo cinético do glúon, obtemos a lagrangiana clássica da QCD invariante sob transformações do grupo de gauge $SU(3)_C$ de cor

$$\mathcal{L}_{QCD} = -\frac{1}{4} G_{\mu\nu}^a G_a^{\mu\nu} + \sum_q \bar{q}(i \not{D} - m_q)q, \quad (2.33)$$

$$\begin{aligned}\mathcal{L}_{QCD} &= -\frac{1}{4}(\partial_\mu G_\nu^a - \partial_\nu G_\mu^a)(\partial^\mu G_a^\nu - \partial^\nu G_a^\mu) + \sum_q \bar{q}(i \not{\partial} - m_q)q \\ &+ g_s \sum_q \bar{q} \gamma^\mu T^a q G_\mu^a \\ &- \frac{g_s^2}{2} f^{abc}(\partial_\mu G_\nu^a - \partial_\nu G_\mu^a) G_b^\mu G_c^\nu - \frac{g_s^2}{4} f^{abc} f_{ade} G_b^\mu G_c^\nu G_\mu^d G_\nu^e\end{aligned}\quad (2.34)$$

A primeira linha de (2.34) contém os termos cinéticos do glúon e dos quarks de massa m_q , os quais dão origem aos respectivos propagadores. A segunda linha contém a interação quark-glúon, e envolve as matrizes λ^a do $SU(3)_C$. Na última linha temos acoplamentos cúbicos e quárticos entre glúons. Acoplamentos entre bósons de gauge só ocorrem devido à natureza não-abeliana de uma dada teoria e, no caso da QCD, são responsáveis pelos fenômenos de liberdade assintótica e, ao que tudo indica, pelo confinamento de cor.

A constante universal g_s é chamada *constante de acoplamento forte* e tem sido medida a partir de diversos tipos de experiência. A constante de estrutura fina da QCD, definida por

$$\alpha_s = \frac{g_s^2}{4\pi} \quad (2.35)$$

obtida pela média global dos resultados de várias experiências vale $\alpha_s(m_Z) = 0.1173 \pm 0.0020$, onde m_Z é massa do Z [8].[¶]

A intensidade de α_s em relação às constantes de estrutura fina do setor eletrofraco pode chegar a ser duas ordens de magnitude maior na escala de m_Z , o que faz das correções radiativas de ordem superior da QCD as correções quânticas dominantes em qualquer cálculo perturbativo de processos que envolvam quarks, glúons e outras partículas que carreguem cargas de cor, mesmo em extensões do modelo padrão, como é o caso dos leptoquarks e de várias partículas supersimétricas.

[¶]Na próxima seção veremos que α é uma grandeza dependente de escala de energia.

Renormalizabilidade e invariância de gauge permitem, ainda, a adição de um termo em $\mathcal{L}_{QCD}$ do tipo

$$\mathcal{L}_\theta = \frac{\theta\alpha_s}{8\pi} G_{\mu\nu}^a \tilde{G}_a^{\mu\nu}, \quad \tilde{G}_a^{\mu\nu} = \frac{1}{2} \varepsilon^{\mu\nu\lambda\rho} G_{a\lambda\rho}. \quad (2.36)$$

Um termo dessa forma corresponderia a uma interação do tipo $\vec{E} \cdot \vec{B}$ em QED. Tal interação, contudo, violaria as simetrias discretas de paridade e reversão do tempo, em contradição com as propriedades observadas das interações fortes. O termo $\mathcal{L}_\theta$, contudo, pode ser escrito em termos da derivada total de um campo vetorial K_μ que não é invariante de gauge, e portanto, pode ser desprezado no contexto da teoria de perturbação. No entanto, no regime não-perturbativo, $\mathcal{L}_\theta$ pode dar origem a efeitos físicos reais, apesar do fato de ser uma divergência total. O vácuo da QCD tem uma estrutura topológica não-trivial e, neste caso, o termo de superfície não pode ser desprezado, sendo a violação de CP inevitável. Medidas do momento de dipolo elétrico do nêutron, por outro lado, impõem um limite superior para θ dado por $\theta < 10^{-9}$, de forma que, para todos os efeitos podemos desprezar $\mathcal{L}_\theta$. O porquê do parâmetro θ ser tão pequeno ainda não está compreendido totalmente e é denominado na literatura como *the strong CP-problem*.

2.4 Alguns Aspectos da Quantização e Renormalização da QCD

O fato da QCD ser uma teoria quântica de campo invariante sob transformações de gauge do grupo $SU(3)_C$ lhe garante um lugar na seleta classe de teorias renormalizáveis.

Em termos práticos isso quer dizer que todas as divergências ultravioletas (UV) que surgem nos diagramas que contém quarks e glúons em loops, e que aparecem nos cálculos de correções radiativas em teoria de perturbação, podem ser removidas consistentemente a partir de um número finito de contratermos. Em teorias de gauge, as simetrias da teoria impõem fortes restrições sobre a estrutura das divergências, relacionando-as entre si, e reduzindo o número de contratermos que precisam ser calculados para renormalizar a teoria. A renormalização da QCD, ou qualquer outra teoria de gauge realista, não é, no entanto, apenas um artifício matemático para a remoção de infinitos indesejáveis, dado que prediz efeitos físicos mensuráveis como a dependência da constante de acoplamento forte e da massa dos quarks pesados com a escala de energia de uma dada colisão. A razão com que α_s depende dessa escala é determinada pela chamada *função β* e pelas equações do grupo de renormalização. Veremos que a função β da QCD é responsável pelo comportamento de α_s em altas

energias, o que, em última instância, nos leva ao fenômeno da liberdade assintótica.

O cálculo dessa função requer o conhecimento de alguns contratermos necessários à renormalização da teoria e, conseqüentemente, requer o cálculo de amplitudes contendo quarks e glúons em loops. É preciso, então, obter o conjunto das regras de Feynman da lagrangiana quantizada. O procedimento de quantização de teorias de gauge, no entanto, não é nada trivial. Na próxima seção mostramos as principais dificuldades e conseqüências na quantização da QCD utilizando o formalismo canônico, no entanto, sem desenvolvê-lo totalmente.

2.4.1 Ghosts e Simetria BRST

A lagrangiana (2.34) ainda não contém todos os termos necessários à quantização da teoria. Para ilustrar esse fato, consideremos primeiramente o caso do eletromagnetismo. A lagrangiana que descreve a propagação dos fótons é dada por (2.14), e é invariante sob a transformação

$$A'_\mu(x) = A_\mu(x) + \frac{1}{e} \partial_\mu \alpha(x) , \quad (2.37)$$

onde $\alpha(x)$ é uma função arbitrária. A partir daquela lagrangiana obtemos as equações de movimento para A_μ

$$\square A_\mu - \partial^\mu \partial_\nu A^\nu = 0 . \quad (2.38)$$

O primeiro passo na quantização canônica de uma teoria é calcular os momentos canonicamente conjugados às componentes do campo A_μ . Os campos conjugados são dados por

$$\pi^\mu = \frac{\partial \mathcal{L}}{\partial (\partial_0 A_\mu)} = -F^{0\mu} , \quad (2.39)$$

e em seguida postulamos as relações de comutação em tempos iguais, promovendo os campos a operadores, de modo que

$$[A_\mu(x_0, \vec{x}), \pi_\nu(y_0, \vec{y})] = i g_{\mu\nu} \delta^3(\vec{x} - \vec{y}) , \quad (2.40)$$

$$[A_\mu(x_0, \vec{x}), A_\nu(y_0, \vec{y})] = 0 , \quad (2.41)$$

$$[\pi_\mu(x_0, \vec{x}), \pi_\nu(y_0, \vec{y})] = 0 . \quad (2.42)$$

Porém, $F^{00} = 0$, o que implica $\pi^0 = 0$, e dessa forma não é possível postular relações de comutação para a componente temporal de A_μ . Além disso existe um problema ainda mais profundo, que a princípio nos impede de obter as regras de Feynman para a QED. O operador $\square - \partial_\mu \partial_\nu$, no espaço de momentos, é da forma

$$-k^2 g_{\mu\nu} + k_\mu k_\nu , \quad (2.43)$$

cujo operador inverso pode ter a seguinte forma geral $Ag_{\nu\alpha} + Bk_\nu k_\alpha$, que deve satisfazer a seguinte equação

$$(-k^2 g^{\mu\nu} + k^\mu k^\nu)(Ag_{\nu\alpha} + Bk_\nu k_\alpha) = g_\alpha^\mu , \quad (2.44)$$

e que expandida produz a equação

$$-k^2 A \left(g_\alpha^\mu - \frac{1}{k^2} k_\alpha k^\mu \right) = g_\alpha^\mu , \quad (2.45)$$

que não possui solução para nenhum A , ou seja, o operador $\square - \partial_\mu \partial_\nu$ não possui inverso, o que significa dizer que não podemos obter o propagador do fóton.

Esses dois problemas estão relacionados aos graus de liberdade espúrios do fóton. Sendo uma partícula de spin-1 sem massa, o fóton possui apenas dois graus de liberdade de polarização físicos, no entanto, A_μ possui 4 graus de liberdade de Lorentz. A liberdade de escolha dos campos A_μ , a menos de uma transformação de gauge, permite a propagação dos graus de liberdade não-físicos. A maneira de eliminar esses graus de liberdade é a *fixação de um gauge* através de vínculos sobre o campo de gauge, por exemplo, $A_0 = 0$ e $\vec{\nabla} \cdot \vec{A} = 0$. Condições de gauge deste tipo, denominadas *não-covariantes*, eliminam os graus de liberdade espúrios mas complicam algebricamente o formalismo. Ao invés disso, podemos adotar uma *condição covariante*, tal como

$$\partial_\mu A^\mu = 0 , \quad (2.46)$$

denominada *condição de Lorentz*, que simplifica a álgebra do formalismo, e que passaremos a adotar daqui para frente. É possível, dessa forma, adicionar um termo inócuo à lagrangiana original, de maneira que

$$\mathcal{L}' = -\frac{1}{4} F_{\mu\nu} F^{\mu\nu} + \frac{1}{2\eta} (\partial_\mu A^\mu)^2 \quad (2.47)$$

produz para A_μ , a seguinte equação de movimento

$$\left[\square g_{\mu\nu} - \left(1 + \frac{1}{\eta} \right) \partial_\mu \partial_\nu \right] A^\nu = 0 , \quad (2.48)$$

além disso, o campo conjugado a A_0 não é mais identicamente nulo, mas dado por

$$\pi^0 = \frac{1}{\eta} (\partial_\mu A^\mu) , \quad (2.49)$$

onde $\partial_\mu A^\mu$ é um campo livre

$$\frac{1}{\eta} \square (\partial \cdot A) = 0 . \quad (2.50)$$

O propagador do fóton pode, então, ser calculado pelo inverso do operador que age sobre A^ν em (2.48), e é dado por

$$\Pi_{\mu\nu}(k) = \frac{1}{k^2} \left(-g_{\mu\nu} + (1 + \eta) \frac{k_\mu k_\nu}{k^2} \right) . \quad (2.51)$$

Existem ainda algumas sutilezas no procedimento de quantização da QED, dado que a quantização canônica e a condição de Lorentz não são compatíveis [9], contudo a lagrangiana (2.47) é o ponto de partida para uma versão quântica consistente do eletromagnetismo. O que gostaríamos de frisar é que o termo de fixação de gauge, apesar de essencial para a quantização da teoria, quebra a invariância de gauge da lagrangiana original. No entanto, o que é mais interessante, é que a lagrangiana ainda é simétrica o bastante para garantir a validade das identidades de Ward, e dessa forma, garantir a renormalizabilidade da teoria. Isso pode ser constatado reescrevendo (2.47) da seguinte maneira

$$\mathcal{L}' = -\frac{1}{4} F_{\mu\nu} F^{\mu\nu} + B \partial_\mu A^\mu - \frac{\eta}{2} B^2 , \quad (2.52)$$

onde $B(x)$ é um campo escalar. Essa lagrangiana nos permite obter o propagador no gauge de Landau, $\eta = 0$, o que não era possível com (2.47). Podemos mostrar que (2.52) nos dá as seguintes equações de movimento

$$\square A^\mu - \partial^\mu \partial_\nu A^\nu = \partial^\mu B , \quad (2.53)$$

$$\partial_\mu A^\mu - \eta B = 0 . \quad (2.54)$$

Eliminando B de (2.53) usando (2.54), recuperamos (2.48), mostrando que (2.47) e (2.52) produzem as mesmas equações de movimento para $A_\mu(x)$. Além disso, de (2.53) temos que $\square B = 0$, e de (2.54), que $\square(\partial_\mu A^\mu) = 0$, ou seja, B e $\partial_\mu A^\mu$ são campos livres. Por fim, a lagrangiana (2.52) produz $\pi^0 = B$ para o momento conjugado a A^0 , e que é exatamente igual a (2.49).

A lagrangiana (2.52), com certeza não é mais invariante por

$$A'_\mu(x) = A_\mu(x) + \frac{1}{e} \partial_\mu \alpha(x) ,$$

no entanto, a condição de Lorentz (2.46), embora vincule o campo A_μ , ainda não o fixa de maneira única, pois A_μ pode se transformar por $(1/e)\partial_\mu \omega(x)$, por exemplo, e ainda satisfazer (2.54), desde que $\square \omega(x) = 0$. Essa restrição sobre $\omega(x)$ pode ser introduzida em $\mathcal{L}'$ através de um campo multiplicador $\xi(x)$

$$\mathcal{L}_{em} = -\frac{1}{4} F_{\mu\nu} F^{\mu\nu} + B \partial_\mu A^\mu - \frac{1}{2} \eta B^2 - \xi \square \omega , \quad (2.55)$$

de modo que a equação de Euler-Lagrangre para $\xi(x)$ é justamente $\square \omega = 0$.

Podemos mostrar que a lagrangiana (2.55) é invariante sob as transformações infinitesimais dadas por^{||}

$$A_\mu(x) \rightarrow A_\mu(x) + \epsilon \partial_\mu \omega(x) , \quad (2.56)$$

$$\xi(x) \rightarrow \xi(x) + \epsilon B(x) , \quad (2.57)$$

$$\omega(x) \rightarrow \omega(x) , \quad (2.58)$$

$$B(x) \rightarrow B(x) . \quad (2.59)$$

O papel do termo $\xi \square \omega$ em (2.55) é cancelar a variação do termo de fixação de gauge quando A_μ se transforma de acordo com (2.56). Graças a esse cancelamento, um certo tipo de “invariância de gauge” é restaurada, e é definida pelas transformações (2.56)–(2.59). Note que essas transformações são parametrizadas por uma constante infinitesimal ϵ independente de x , no entanto, ainda há um caráter global devido à presença de $\omega(x)$. As transformações (2.56)–(2.59) representam um tipo de simetria denominada *simetria BRST abeliana*^{**}. Essa simetria remanescente é suficiente para demonstrar as identidades de Ward em QED, logo, (2.55) é uma lagrangiana satisfatória. Na verdade, como os campos ξ e ω não interagem com A_μ , eles podem ser ignorados em cálculos práticos, tornando (2.47) ou (2.52) lagrangianas aceitáveis.

Em teorias de gauge não-abelianas surge ainda um outro tipo de complicação. Vamos tentar generalizar (2.52) para o caso não-abeliano de uma maneira direta e escrever o setor gluônico de $\mathcal{L}_{QCD}$ como

$$\mathcal{L}'_{QCD} = -\frac{1}{4} F_{\mu\nu}^a F_a^{\mu\nu} + B^a (\partial_\mu G_a^\mu) - \frac{1}{2} \eta B_a B^a , \quad (2.60)$$

onde B_a se transforma de acordo com a representação adjunta do $SU(3)_C$, assim como G_a^μ . A equação de movimento para B_a nos revela

$$\partial_\mu G^{a\mu} - \eta B^a = 0 \quad (2.61)$$

como a condição de gauge, e

$$D^{ab\nu} F_{\mu\nu}^b + \partial_\mu B^a = 0 \quad (2.62)$$

como a equação de movimento de G_μ^a , onde

$$D_\mu^{ab} = \delta^{ab} \partial_\mu - g_s f^{abc} A_\mu^c \quad (2.63)$$

é derivada covariante na representação adjunta. Dessa forma, pode-se verificar que

$$D^{ca\mu} D^{ab\nu} F_{\mu\nu}^b = 0 , \quad (2.64)$$

^{||}Quando campos de matéria estiverem envolvidos, transformações apropriadas para eles devem ser definidas.

^{**}Becchi, Rouet, Stora e Tyutin.

o que implica

$$\square B^c = g_s f^{cad} G_\mu^d \partial^\mu B^a , \quad (2.65)$$

mostrando que o campo auxiliar, ao contrário do que ocorria na QED, não é mais um campo livre e, conseqüentemente, $\partial_\mu B_a^\mu$ também não é.

Contudo, a exemplo de $\mathcal{L}_{em}$, podemos introduzir campos auxiliares em (2.60) de forma a compensar a variação do termo de fixação de gauge sob a transformação de G_μ^a dada por

$$\delta G_\mu^a = -D_\mu^{ab} \omega^b . \quad (2.66)$$

O termo $\xi - \omega$ pode ser obtido impondo que (2.61) ainda vale após (2.66), o que implica

$$\partial^\mu [D_\mu^{ab} \omega^b(x)] = 0 . \quad (2.67)$$

Nesse caso, $\omega^a(x)$ não é um campo livre devido ao aspecto não-abeliano do campo do glúon. Note que se $f^{abc} = 0$, como na QED, ou outra teoria de gauge abeliana, $\omega^a(x)$ seria um campo livre. Exatamente como no caso da QED, a condição (2.67) pode ser incorporada à lagrangiana (2.60) por um campo multiplicador $\xi(x)$ de modo que

$$\mathcal{L}_{QCD} = -\frac{1}{4} F_a^{\mu\nu} F_{\mu\nu}^a + B^a (\partial_\mu G_\mu^a) - \frac{1}{2} \eta B_a B^a - \xi^a \partial^\mu D_\mu^{ab} \omega^b . \quad (2.68)$$

A lagrangiana (2.68) possui agora uma invariância residual que é expressa pelas leis de transformação

$$G_\mu^a \rightarrow G_\mu^a + \epsilon D_\mu^{ab} \omega^b \quad (2.69)$$

$$\xi^a \rightarrow \xi^a + \epsilon B^a \quad (2.70)$$

$$B^a \rightarrow B^a \quad (2.71)$$

$$\omega^a \rightarrow \omega^a - \frac{1}{2} \epsilon g_s f^{abc} \omega^b \omega^c , \quad (2.72)$$

A variação do último termo de (2.68), causada por (2.70), cancelaria a variação de $B^a (\partial_\mu G_\mu^a)$ se $\omega^a \rightarrow \omega^a$ simplesmente. Porém o último termo também envolve G_μ^a através da derivada covariante, e sua variação, dada por (2.69) dá origem a um termo adicional que só pode ser cancelado pela transformação não trivial (2.72) de ω . Tal transformação só pode ocorrer se ω anticomutar consigo mesmo, pois o termo $\omega^b \omega^c$ vem contraído por um tensor anti-simétrico nos índices de cor. Os campos $\omega(x)$ e $\xi(x)$ são bastante peculiares, pois são escalares sob transformação de Lorentz obedecendo à estatística de Fermi-Dirac! Na verdade, $\omega(x)$ e $\xi(x)$ (e ϵ também) são denominados matematicamente por *variáveis de Grassmann*. Os campos $\omega(x)$ e $\xi(x)$ são chamados de *ghosts* e têm norma negativa, ao obedecerem à estatística errada, e portanto, produzem probabilidades negativas, o que os torna quantidades

desprovidas de significado físico. Note que os ghosts só se acoplam a glúons, e o seu papel é justamente cancelar a contribuição dos graus de liberdade não físicos do glúon.

O mecanismo que dá origem aos ghosts pode ser entendido de uma forma mais rigorosa e elegante usando o formalismo de integrais de trajetória utilizado por Fadeev e Popov [10]. Cabe enfatizar que esses campos são introduzidos apenas como uma maneira de desenvolver um formalismo covariante simples em comparação com formalismos não-covariantes, tornando as regras de Feynman mais simples e, conseqüentemente, tornando os cálculos mais simples.

A lagrangiana completa da QCD agora é dada por

$$\mathcal{L}_{QCD} = \mathcal{L}_{\text{clássica}} + \mathcal{L}_{\text{fg}} + \mathcal{L}_{\text{ghost}} \quad (2.73)$$

$$\mathcal{L}_{\text{clássica}} = -\frac{1}{4}G_{\mu\nu}^a G_a^{\mu\nu} + \bar{\psi}(i \not{D} - m)\psi \quad (2.74)$$

$$\mathcal{L}_{\text{fg}} = B^a(\partial_\mu G_a^\mu) - \frac{1}{2}\eta B_a B^a \quad (2.75)$$

$$\mathcal{L}_{\text{ghosts}} = -c_1^a \partial^\mu D_\mu^{ab} c_2^b, \quad (2.76)$$

onde $\psi(x)$ representa o campo fermiônico de um dado sabor de quark e ξ e ω foram redefinidos por c_1 e c_2 , respectivamente. O subscrito fg significa *fixação de gauge*. As regras de Feynman obtidas a partir dessa lagrangiana são dadas no Apêndice A.

2.4.2 Lagrangiana Renormalizada

Tendo (2.73) em mãos, vamos agora, redefinir os campos de $\mathcal{L}_{QCD}$, escolhendo $\eta = 0$ para simplificar,

$$G_\mu^a = \sqrt{Z_3} G_{R\mu}^a, \quad (2.77)$$

$$\psi = \sqrt{Z_2} \psi_R, \quad (2.78)$$

$$c_{1,2}^a = \sqrt{Z_2^c} c_{R1,2}^a, \quad (2.79)$$

onde $G_{R\mu}^a$, ψ_R e $c_{R1,2}^a$ correspondem aos campos renormalizados do glúon, do quark e dos ghosts respectivamente. Os contratermos possuem as seguintes definições

$$\delta_2 = Z_2 - 1, \quad (2.80)$$

$$\delta_3 = Z_3 - 1, \quad (2.81)$$

$$\delta_2^c = Z_2^c - 1, \quad (2.82)$$

$$\delta_m = Z_2 m_0 - m, \quad (2.83)$$

$$\delta_1 = \frac{g_{s0}}{g_s} Z_2 \sqrt{Z_3} - 1, \quad (2.84)$$

$$\delta_1^c = \frac{g_{s0}}{g_s} Z_2^c \sqrt{Z_3} - 1 , \quad (2.85)$$

$$\delta_1^{3g} = \frac{g_{s0}}{g_s} Z_3 \sqrt{Z_3} - 1 , \quad (2.86)$$

$$\delta_1^{4g} = \frac{g_{s0}^2}{g_s^2} Z_3^2 - 1 . \quad (2.87)$$

Note que estes oito contratermos dependem de apenas cinco parâmetros básicos: Z_2 , Z_2^c , Z_3 , m_0 e g_{s0} . Existem, portanto, três relações entre eles, ao passo que cinco condições de renormalização apenas são necessárias para remover as divergências UV e determinar todos os oito contratermos. Essas relações entre contratermos são uma consequência direta da simetria de gauge da teoria.

Já ao nível de 1-loop, podemos demonstrar que

$$\delta_1 - \delta_2 = \delta_1^{3g} - \delta_3 = \frac{1}{2}(\delta_1^{4g} - \delta_3) = \delta_1^c - \delta_2^c . \quad (2.88)$$

A lagrangiana da QCD, adquire, portanto, a seguinte forma

$$\mathcal{L}_{QCD} = \mathcal{L}_{ren} + \mathcal{L}_{c.t.} , \quad (2.89)$$

onde $\mathcal{L}_{ren}$ é lagrangiana (2.73) ($\eta = 0$) com os campos renormalizados e $\mathcal{L}_{c.t.}$, a lagrangiana dos contratermos que é dada explicitamente logo abaixo

$$\begin{aligned} \mathcal{L}_{c.t.} = & -\frac{1}{4}\delta_3(\partial_\mu G_\nu^a - \partial_\nu G_\mu^a)^2 + \bar{\psi}(i\delta_2 \not{\partial} - \delta_m)\psi - \delta_2^c \bar{c}^a \partial^2 c^a \\ & + g_s \delta_1 G_\mu^a \bar{\psi} \gamma^\mu \psi - g_s \delta_1^{3g} f^{abc} (\partial_\mu G_\nu^a) G^{b\mu} G^{c\nu} \\ & - g_s^2 \delta_1^{4g} f^{eab} F^{ecd} G_\mu^a G_\nu^b G^{c\mu} G^{d\nu} - g_s \delta_1^c \bar{c}^a f^{abc} \partial^\mu G_\mu^b c^c . \end{aligned} \quad (2.90)$$

A partir desta lagrangiana obtém-se as regras de Feynman dos contratermos necessários à renormalização da QCD, e que podem ser encontradas, por exemplo, em [11, 12].

Os contratermos δ_1 , δ_2 e δ_3 podem ser usados para o cálculo da função β pela fórmula [13]

$$\beta(g_s) = g_s \mu_R \frac{\partial}{\partial \mu_R} (-\delta_1 + \delta_2 + \frac{1}{2}\delta_3) , \quad (2.91)$$

onde μ_R é escala de renormalização. Na QED, $\delta_1 = \delta_2$ pela identidade de Ward-Takahashi, de modo que o *running* de α_{em} é determinado pela renormalização da função de onda do fóton. Na QCD, no entanto, todos os três termos contribuem. O cálculo desses contratermos é bem trabalhoso e pode ser encontrado em diversos livros texto [11, 13, 14], vamos portanto, nos restringir a apresentar os resultados no esquema $\overline{MS}$, onde os termos de subtração contém apenas o pólo

$$\frac{1}{\bar{\epsilon}} \equiv \frac{1}{\epsilon} - \gamma + \ln 4\pi , \quad (2.92)$$

onde $d = 4 + \varepsilon$ é o número de dimensões espaço-temporais no esquema de regularização dimensional [15] e γ é a constante de Euler,

$$\delta_1 = -\frac{g_s^2}{16\pi^2} [C_A + T_R] \frac{2}{\varepsilon} \quad (2.93)$$

$$\delta_2 = -\frac{g_s^2}{16\pi^2} T_R \frac{2}{\varepsilon} \quad (2.94)$$

$$\delta_3 = \frac{g_s^2}{16\pi^2} \left[\frac{2}{3} C_A - \frac{4}{3} n_f T_R \right] \frac{2}{\varepsilon} . \quad (2.95)$$

Os fatores de cor C_A e T_R são dados por 3 e 1/2, respectivamente, para o $SU(3)_C^{\dagger\dagger}$, e n_f é o número de sabores de quarks.

Substituindo (2.93)–(2.95) em (2.91) obtemos a função β da QCD

$$\beta(g_s) = -\frac{g_s^3}{16\pi^2} \left[11 - \frac{2n_f}{3} \right] . \quad (2.96)$$

A solução da equação do grupo de renormalização

$$\frac{d}{d \log(Q/\mu_R)} g_s = \beta(g_s) \quad (2.97)$$

é dada por

$$\alpha_s(Q^2) = \frac{\alpha_s(\mu_R^2)}{1 + \alpha_s(\mu_R^2) b_0 \log(Q^2/\mu_R^2)} , \quad (2.98)$$

onde

$$b_0 \equiv \frac{1}{4\pi} \left[11 - \frac{2n_f}{3} \right] , \quad (2.99)$$

e define o *running* de α_s .

2.4.3 Liberdade Assintótica

Se n_f é menor do que 16 então $b_0 > 0$ e a função β tem sinal negativo. Nesse caso, pelas expressões (2.97)–(2.99) vemos que conforme Q^2 aumenta $\alpha_s(Q^2)$ decresce monotonicamente

$$\lim_{Q^2 \rightarrow \infty} \alpha_s(Q^2) = 0 . \quad (2.100)$$

Na QCD, n_f corresponde ao número de sabores de quarks, ou seja, $n_f \leq 6$. Portanto, a constante de acoplamento da teoria decresce em magnitude conforme aumentamos a escala de energia em uma colisão! Para altíssimas energias, a QCD tende a se comportar como uma teoria de quarks e glúons livres (explicando o sucesso do modelo a párons em descrever colisões de altas energias), esse é o chamado fenômeno de *liberdade assintótica*.

^{††}Ver Apêndice B.

Sua origem fundamental é a natureza não-abeliana das interações da QCD. Olhando para os contratermos, vemos que, se $C_A = 0$, como seria na QED, por exemplo, a função β seria positiva-definida, esse fator de cor, no entanto surge pela ocasião da inclusão de vértices tríplexes de glúons, o que só ocorre na QCD (ou outras teorias de gauge não-abelianas).

Por outro lado, para Q^2 decrescente α_s cresce em magnitude, o que sugere que quarks e glúons não podem se afastar indefinidamente um dos outros. Isso lembra a propriedade de confinamento de cores, contudo, uma explicação mais rigorosa para o confinamento ainda é necessária [16].

Capítulo 3

Correções de QCD ao Processo $e^+e^- \rightarrow$ Hádrons

3.1 O Aniquilamento e^+e^- no Modelo Padrão

Muitas das idéias e propriedades básicas da QCD perturbativa podem ser ilustradas pelo estudo da produção de hádrons no aniquilamento elétron-pósitron.

Dentre todas as quantidades observáveis da QCD, a seção de choque do processo $e^+e^- \rightarrow$ hádrons talvez seja a mais simples ao cálculo de correções radiativas de ordem superior, prova disso é o fato de que já se conhecem as correções de ordem α_s^3 [17] para $\sigma(e^+e^- \rightarrow q\bar{q})$, quarks leves e α_s^2 [18, 19] para a produção de quarks pesados e event shapes. Além disso, efeitos não-perturbativos têm pouca influência sobre tais correções em relação a processos do tipo hádron-hádron, onde o regime não-perturbativo se faz presente nas funções de distribuição de pártons. Isso tudo faz da medida da seção de choque do aniquilamento e^+e^- um ambiente propício para a determinação precisa da constante de acoplamento forte, além do estudo de várias propriedades da QCD.

A QCD perturbativa também prediz uma rica estrutura de “jatos” para o estado final hadrônico nesse tipo de processo de onde podem ser extraídas informações a respeito da constante de acoplamento α_s , do spin das partículas que formam os jets e da natureza não-abeliana do acoplamento entre glúons, por exemplo.

A seguir apresentamos as correções de QCD em ordem α_s para a produção de 2 jets em colisões elétron-pósitron através de um cálculo analítico completo sem nenhum tipo de aproximação para as regiões singulares do espaço de fase. As divergências ultravioletas que surgem das amplitudes virtuais ao nível de 1-loop e as divergências infravermelhas causadas pela emissão de um glúon soft ou colinear são regularizadas dimensionalmente e se apresentam como pólos do tipo ε^{-i} ($i = 1, 2$).

No capítulo 4 calcularemos novamente a seção de choque desse processo utilizando uma metodologia diferente, baseada em aproximações para as regiões singulares do espaço de fase e que será usada nos cálculos de NLO de QCD para o processo

de Drell-Yan incluindo efeitos virtuais de leptquarks escalares. Mostraremos, nessa ocasião, que a utilização deste método nos dá o mesmo resultado para a seção de choque total inclusiva que passaremos a calcular a seguir.

3.2 Cálculos em Ordem Zero

Para calcular a seção de choque do processo $e^+e^- \rightarrow \gamma, Z \rightarrow q\bar{q}$ corrigida em *next-to-leading order*, vamos utilizar a abordagem normalmente usada no cálculo de correções radiativas de QCD ao decaimento de bósons de gauge eletrofracos. Nessa abordagem a “razão de decaimento” de um fóton virtual em um par $q\bar{q}$ recebe contribuições dos diagramas de Feynman da figura (3.1), que correspondem à amplitude de Born do processo, emissão de um glúon real e contribuições virtuais ao nível de 1-loop. Para obter o processo que realmente estamos interessados em calcular, basta a contração com a corrente leptônica do estado inicial e^+e^- , a correção de NLO de QCD é a mesma do processo $\gamma^* \rightarrow q\bar{q}$ já que o estado inicial é puramente leptônico e não pode ser acoplado a glúons.

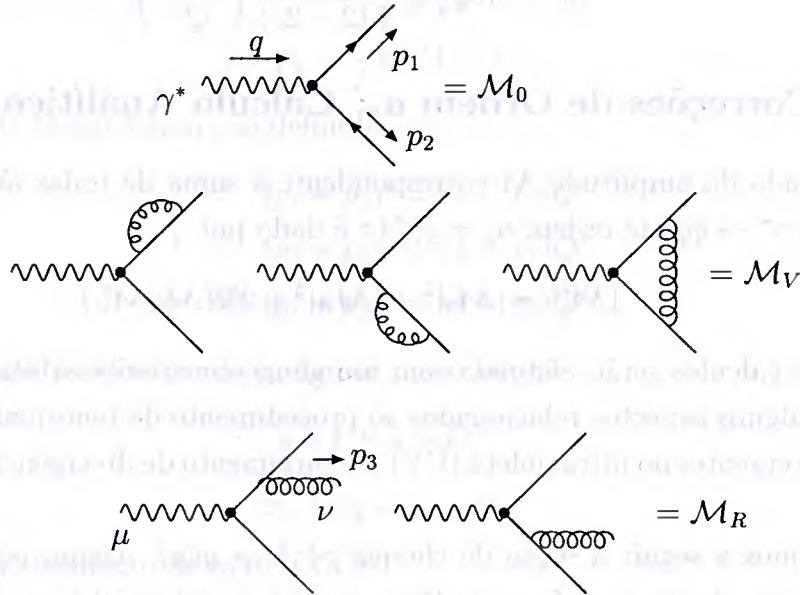


Figura 3.1: Decaimento hadrônico de um fóton virtual em NLO de QCD.

A amplitude de ordem zero pode ser calculada a partir das regras de Feynman do Apêndice A e é dada por

$$\mathcal{M}_0 = -ieQ_q \bar{u}(p_1) \not{\epsilon}(q)v(p_2) , \quad (3.1)$$

onde Q_q é carga elétrica do quark q em unidades da carga do elétron. O quadrado dessa amplitude após soma sobre as polarizações dos quarks e do fóton é dado por

$$|\overline{\mathcal{M}_0}|^2 = -3e^2Q_q^2 \text{Tr}[\not{p}_1\gamma_\mu \not{p}_2\gamma^\mu] , \quad (3.2)$$

onde o fator 3 vem da soma sobre os estados finais de cores dos quarks. Calculando o traço em d dimensões obtemos *

$$|\overline{\mathcal{M}_0}|^2 = (3)2e^2Q_q^2(d-2)Q^2 , \quad (3.3)$$

onde $Q^2 = q^2$ representa a virtualidade do fóton.

Adotando as convenções de [20], a seção de choque do processo $1 \rightarrow N$ é da forma

$$\sigma = \frac{1}{2E_{cm}} |\overline{\mathcal{M}}|^2 \Phi_2^d , \quad (3.4)$$

onde Φ_2^d é dado por (C.14) do Apêndice C. Definindo $d = 4 - 2\varepsilon$, temos

$$\sigma_0 = 3\alpha Q_q^2 Q \frac{\Gamma(2-\varepsilon)}{\Gamma(2-2\varepsilon)} \left(\frac{4\pi\mu^2}{Q^2} \right)^\varepsilon . \quad (3.5)$$

3.3 Correções de Ordem α_s : Cálculo Analítico Completo

O quadrado da amplitude $\mathcal{M}$ correspondente à soma de todas as contribuições ao processo $\gamma^* \rightarrow q\bar{q}$ até ordem $\alpha_s = g_s^2/4\pi$ é dado por

$$|\mathcal{M}|^2 = |\mathcal{M}_0|^2 + |\mathcal{M}_R|^2 + 2\Re(\mathcal{M}_0\mathcal{M}_V^*) . \quad (3.6)$$

Todos os cálculos serão efetuados em um gauge covariante arbitrário η de modo a ilustrar alguns aspectos relacionados ao procedimento de renormalização das amplitudes divergentes no ultravioleta (UV) e o surgimento de divergências infravermelhas (IR).

Obtemos a seguir a seção de choque $\sigma(\gamma^* \rightarrow q\bar{q}g)$. Daqui por diante todas as divergências, tanto as infravermelhas, quanto as ultravioletas serão regularizadas dimensionalmente [20, 21].

*Em regularização dimensional, a constante de acoplamento adquire dimensão de massa

$$\alpha_d = \frac{e^2/4\pi}{(\mu^2)^{\frac{d}{2}-2}} = \frac{\alpha}{(\mu^2)^{\frac{d}{2}-2}} .$$

3.3.1 Emissão de Glúon Real

Utilizando as regras de Feynman do apêndice A obtemos as seguintes amplitudes para a emissão de 1 glúon real

$$\mathcal{A}_R = ig_s T^a e Q_q \bar{u}(p_2) \gamma^\nu \cdot \frac{\not{p}_a}{p_a^2} \cdot \gamma^\mu v(p_1) \epsilon_\mu(q) \epsilon_\nu^*(p_3) , \quad (3.7)$$

$$\mathcal{B}_R = -ig_s T^a e Q_q \bar{u}(p_2) \gamma^\mu \cdot \frac{\not{p}_b}{p_b^2} \cdot \gamma^\nu v(p_1) \epsilon_\mu(q) \epsilon_\nu^*(p_3) , \quad (3.8)$$

de modo que

$$\mathcal{M}_R(\gamma^* \rightarrow q\bar{q}g) = \mathcal{A}_R + \mathcal{B}_R . \quad (3.9)$$

Os 4-momentos p_a e p_b são definidos por

$$p_a = p_2 + p_3 , \quad (3.10)$$

$$p_b = p_1 + p_3 , \quad (3.11)$$

$$q = p_1 + p_2 + p_3 , \quad (3.12)$$

$$E_{cm} = Q = \sqrt{q^2} . \quad (3.13)$$

Introduzindo as frações de energia adimensionais $x_i = \frac{2E_i}{Q}$, $i = 1, 2, 3$, onde E_i corresponde à energia da partícula de 4-momento p_i , e desprezando a massa de todos os pártons (quarks leves e glúons) temos que

$$p_i \cdot p_j = \frac{1}{2} Q^2 (1 - x_k) . \quad (3.14)$$

As variáveis de Mandelstam são definidas por

$$s = (p_1 + p_3)^2 = (1 - x_2) Q^2 , \quad (3.15)$$

$$t = (p_2 + p_3)^2 = (1 - x_1) Q^2 , \quad (3.16)$$

$$u = (p_1 + p_2)^2 = (1 - x_3) Q^2 , \quad (3.17)$$

e a partir dessas relações ainda podemos mostrar que

$$s + t + u = Q^2 , \quad (3.18)$$

$$x_1 + x_2 + x_3 = 2 . \quad (3.19)$$

O quadrado do elemento de matriz (3.9)

$$|\mathcal{M}_R|^2 = |\mathcal{A}_R|^2 + |\mathcal{B}_R|^2 + 2\mathcal{A}_R \mathcal{B}_R^* \quad (3.20)$$

é calculado a partir de (3.7) e (3.8). O resultado após soma sobre as polarizações do fóton, as do glúon e sobre os spins dos quarks, é

$$\overline{|\mathcal{A}_R|^2} = 4 \left(\frac{e Q_q g_s}{p_a^2} \right)^2 \text{Tr}[\not{p}_2 \gamma^\mu \not{p}_a \gamma^\nu \not{p}_1 \gamma_\mu \not{p}_a \gamma_\nu] \quad (3.21)$$

$$|\overline{\mathcal{B}_R}|^2 = 4 \left(\frac{eQ_q g_s}{p_b^2} \right)^2 \text{Tr}[\not{p}_2 \gamma^\mu \not{p}_b \gamma^\nu \not{p}_1 \gamma_\mu \not{p}_b \gamma_\nu] \quad (3.22)$$

$$2\overline{\mathcal{A}_R \mathcal{B}_R^*} = -8 \left(\frac{eQ_q g_s}{p_a \cdot p_b} \right)^2 \text{Tr}[\not{p}_2 \gamma^\mu \not{p}_a \gamma^\nu \not{p}_1 \gamma_\mu \not{p}_b \gamma_\nu] . \quad (3.23)$$

Calculamos os traços com o auxílio do pacote FeynCalc [22] para Mathematica em $d = 4 - 2\varepsilon$. Abaixo está o resultado em termos das variáveis de Mandelstam

$$|\overline{\mathcal{M}_R}|^2 = 8e^2 Q_q^2 g_s^2 \left[(2 - 2\varepsilon)^2 \left(\frac{s}{t} + \frac{t}{s} \right) + (4 - 4\varepsilon) \frac{2uQ^2}{st} - 4\varepsilon(2 - 2\varepsilon) \right] . \quad (3.24)$$

Expressando (3.24) em termos das frações de energia e após um certo exercício algébrico, obtém-se

$$|\overline{\mathcal{M}_R}|^2 = 32e^2 Q_q^2 g_s^2 \mathcal{F}(x_1, x_2) , \quad (3.25)$$

onde

$$\mathcal{F}(x_1, x_2) = (1 - \varepsilon)^2 \frac{x_1^2 + x_2^2}{(1 - x_1)(1 - x_2)} - 2\varepsilon(1 - \varepsilon) \frac{2 - 2x_1 - 2x_2 + x_1 x_2}{(1 - x_1)(1 - x_2)} \quad (3.26)$$

que diverge para $x_1 \rightarrow 1$ e $x_2 \rightarrow 1$!

Quando o glúon é emitido pelo quark do estado final, essa divergência ocorre para $t \rightarrow 0$ ($x_1 \rightarrow 1$), que no limite não-massivo é dada por

$$t = 2p_2 \cdot p_3 = 2E_2 E_3 (1 - \cos \theta_{23}) , \quad (3.27)$$

e que aparece no denominador de $|\mathcal{A}_R|^2$ e do termo de interferência. Dessa forma, quando $E_3 \rightarrow 0$ ou quando $\theta_{23} \rightarrow 0$, então $t \rightarrow 0$ originando um pólo na amplitude.

O caso $E_3 \rightarrow 0$ corresponde à divergência de glúon soft enquanto $\theta_{23} \rightarrow 0$ à situação onde o glúon é emitido na mesma direção de propagação do quark, o que corresponde a uma divergência colinear. As duas são coletivamente denominadas divergências infravermelhas (IR). Se o glúon é emitido pelo antiquark, então $s \rightarrow 0$ ($x_2 \rightarrow 1$) originando também divergências IR para $E_3 \rightarrow 0$ e $\theta_{12} \rightarrow 1$.

A seção de choque total do processo é obtida integrando nas variáveis x_1 e x_2 na região física permitida. Os limites de integração decorrem de uma análise do tipo Dalitz plot. Da expressão (3.18) obtemos

$$(1 - x_1)Q^2 + (1 - x_2)Q^2 + 2p_1 \cdot p_2 = Q^2 , \quad (3.28)$$

mas

$$p_1 \cdot p_2 = E_1 E_2 - \vec{p}_1 \cdot \vec{p}_2 = \frac{Q^2}{4} x_1 x_2 (1 - \cos \theta_{12}) , \quad (3.29)$$

logo, inserindo (3.29) em (3.28) e sabendo que $-1 \leq \cos \theta_{12} \leq 1$, então

$$\begin{aligned} 0 &\leq x_1 \leq 1 \\ 1 - x_1 &\leq x_2 \leq 1 . \end{aligned}$$

Dessa forma a seção de choque é dada por[†]

$$\sigma_R = \frac{2\alpha_s}{3\pi} \sigma_0 \left(\frac{4\pi\mu^2}{Q^2} \right)^\varepsilon \frac{1}{\Gamma(2-\varepsilon)} \int_0^1 dx_1 x_1^{-2\varepsilon} \int_{1-x_1}^1 dx_2 x_2^{-2\varepsilon} \left(\frac{1-z^2}{4} \right)^{-\varepsilon} \mathcal{F}(x_1, x_2) , \quad (3.30)$$

onde expressamos o espaço de fase de 3 partículas em d dimensões dada por (C.15) em termos das frações de energia x_1 e x_2 . Essas integrais sobre x_1 e x_2 são bastante complicadas de calcular. Contudo é possível desacoplá-las pela seguinte mudança de variáveis

$$x_2 \equiv 1 - vx_1$$

onde agora, tanto x_1 quanto v variam de 0 até 1. Tal mudança implica

$$\sigma_R = \frac{2\alpha_s}{3\pi} \sigma_0 \left(\frac{4\pi\mu^2}{Q^2} \right)^\varepsilon \frac{1}{\Gamma(2-\varepsilon)} \int_0^1 dx_1 x_1^{-2\varepsilon} (1-x_1)^{-\varepsilon} \int_0^1 dv v^{-\varepsilon} (1-v)^{-\varepsilon} x_1 \mathcal{F}(x_1, v) . \quad (3.31)$$

A função $x_1 \mathcal{F}(v, x_1)$, proveniente do quadrado da amplitude de espalhamento, agora é dada por

$$\begin{aligned} x_1 \mathcal{F}(v, x_1) &= \frac{(v^2 + 1)x_1^2 - 2vx_1 + 1}{v(1-x_1)} - 2 \frac{(v^2 - v + 1)x_1^2 - x_1 + 1}{v(1-x_1)} \varepsilon \\ &\quad + \frac{(v^2 - 2v + 1)x_1^2 + 2(v-1)x_1 + 1}{v(1-x_1)} \varepsilon^2 . \end{aligned} \quad (3.32)$$

É necessário manter termos até ordem ε^2 pois as integrações produzirão termos de ordem até ε^{-2} . Substituindo $x_1 \mathcal{F}(x_1)$ em (3.31) obtemos uma série de produtos de funções B (beta).[‡] Escrevendo (3.31) como

$$\sigma_R = \frac{2\alpha_s}{3\pi} \sigma_0 \left(\frac{4\pi\mu^2}{Q^2} \right)^\varepsilon \frac{1}{\Gamma(2-\varepsilon)} \Delta_R , \quad (3.33)$$

definimos uma função $G(\alpha, \beta)$ dada pelo seguinte produto de funções B (beta)

$$\begin{aligned} G(\alpha, \beta) &\equiv \frac{1}{\Gamma(2-\varepsilon)} B(-2\varepsilon + \alpha, -\varepsilon - 1) B(-\varepsilon + \beta, -\varepsilon) \\ &= \frac{1}{-\varepsilon} \frac{\Gamma(-2\varepsilon + \alpha + 1)}{\Gamma(-2\varepsilon + \beta + 2)} \frac{\Gamma(-\varepsilon + \beta + 1)}{\Gamma(-\varepsilon + 2)} \frac{\Gamma^2(1-\varepsilon)}{\Gamma(-3\varepsilon + \alpha + 1)} , \end{aligned} \quad (3.34)$$

[†] $z = 1 + \frac{2}{x_1 x_2} (1 - x_1 - x_2)$; ver apêndice C.

[‡] A função B (beta) é definida pela sua forma integral

$$B(\alpha, \beta) = \int_0^1 x^\alpha (1-x)^\beta = \frac{\Gamma(1+\alpha)\Gamma(1+\beta)}{\Gamma(2+\alpha+\beta)} .$$

em termos da qual Δ_R tem a expressão

$$\Delta_R = (1 - 2\varepsilon + \varepsilon^2)[G(0, -1) + G(2, 1) + G(2, 1)] + 2(\varepsilon - \varepsilon^2)[G(2, 0) + G(1, -1)] + 2(\varepsilon^2 - 1)G(1, 0) . \quad (3.35)$$

Obtemos, após expandir as funções G usando as propriedades das funções Γ do Apêndice D,

$$\Delta_R = \left(\frac{2}{\varepsilon^2} + \frac{3}{\varepsilon} + \frac{19}{2} + \mathcal{O}(\varepsilon) \right) \frac{\Gamma^2(1 - \varepsilon)}{\Gamma(1 - 3\varepsilon)} . \quad (3.36)$$

Finalmente, usando a expansão da função Γ (D.8) do Apêndice D, temos

$$\frac{\Gamma^2(1 - \varepsilon)}{\Gamma(1 - 3\varepsilon)} = 1 - \gamma\varepsilon + \frac{1}{12}(6\gamma^2 - 7\pi^2)\varepsilon^2 + \mathcal{O}(\varepsilon^3) \quad (3.37)$$

e a expansão

$$\left(\frac{4\pi\mu^2}{Q^2} \right)^\varepsilon = \exp \left[\varepsilon \ln \left(\frac{4\pi\mu^2}{Q^2} \right) \right] = 1 + \ln \left(\frac{4\pi\mu^2}{Q^2} \right) \varepsilon + \frac{1}{2} \ln^2 \left(\frac{4\pi\mu^2}{Q^2} \right) \varepsilon^2 + \mathcal{O}(\varepsilon^3) ,$$

chegamos à expressão da razão de decaimento do processo $\gamma^* \rightarrow q\bar{q}g$ regularizada dimensionalmente

$$\begin{aligned} \sigma_R(\gamma^* \rightarrow q\bar{q}g) &= \frac{2\alpha_s}{3\pi} \sigma_0 \left[\frac{2}{\varepsilon^2} + \frac{1}{\varepsilon} \left[2 \ln \left(\frac{4\pi\mu^2}{Q^2} \right) - 2\gamma + 3 \right] + \ln^2 \left(\frac{4\pi\mu^2}{Q^2} \right) \right. \\ &\quad \left. + (-2\gamma + 3) \ln \left(\frac{4\pi\mu^2}{Q^2} \right) + \gamma^2 - 3\gamma + \frac{7\pi^2}{6} + \frac{19}{2} + \dots \right] . \end{aligned} \quad (3.38)$$

Veja que a seção de choque é divergente no limite $d \rightarrow 4$ ($\varepsilon \rightarrow 0$). No entanto ainda é preciso contabilizar a contribuição dos glúons virtuais, cálculo que é apresentado a seguir.

3.3.2 Correções Virtuais

As correções de glúon virtual ao processo $\gamma^* \rightarrow q\bar{q}$ são calculadas a partir dos diagramas de Feynman da figura (3.1) que correspondem à correção do vértice e autoenergia dos quarks.

As correções virtuais em ordem α_s surgem da interferência de $\mathcal{M}_V$ com a amplitude de Born $\mathcal{M}_0$. A amplitude correspondente à correção do vértice será denotada por $\mathcal{A}_V$ enquanto as autoenergias por $\mathcal{B}_V + \mathcal{C}_V$.

Correções de autoenergia

A autoenergia de um quark sem massa, ao nível de 1-loop, e em um gauge covariante arbitrário η é calculada a partir do diagrama de Feynman da figura (3.2)

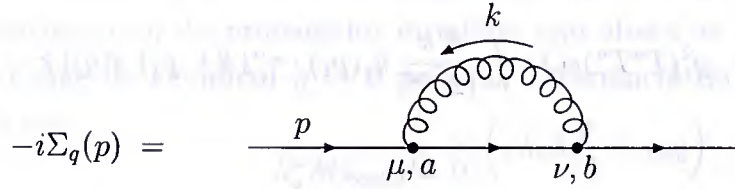


Figura 3.2: Autoenergia do quark.

cuja amplitude é dada por

$$-i\Sigma_q(p) = -g_s^2 C_F \int \frac{d^n k}{(2\pi)^n} \frac{\gamma_\mu (\not{p} - \not{k}) \gamma_\nu}{k^2 (p - k)^2} \times \left(\frac{g^{\mu\nu}}{k^2} + \eta \frac{k^\mu k^\nu}{k^4} \right), \quad (3.39)$$

onde $C_F = 4/3$ é o fator de cor dado no Apêndice B. Usando a álgebra de Dirac em d dimensões, introduzindo a integral paramétrica

$$\frac{k^2}{(p - k)^2} = \int_0^1 dx \frac{1}{[(p - k)^2 x + k^2(1 - x)]^2} \quad (3.40)$$

e definindo

$$l = k - xp, \quad (3.41)$$

$$\Delta = -xp^2(1 - x), \quad (3.42)$$

temos

$$\begin{aligned} -i\Sigma_q(p) = & -g_s^2 C_F \not{p} \left[\int_0^1 dx \int \frac{d^d l}{(2\pi)^d} \frac{(2 - d)(1 - x) - \eta(1 + x)}{[l^2 - \Delta]^2} + \right. \\ & \left. \int_0^1 dx \int \frac{d^d l}{(2\pi)^d} \frac{\eta(1 - x)(4l^2/d + 4x^2 p^2)}{[l^2 - \Delta]^3} \right]. \end{aligned} \quad (3.43)$$

Calculando a integral no momento do loop, obtemos após integração em x

$$-i\Sigma_q(p) = i \not{p} \frac{g_s^2}{16\pi} C_F \left(\frac{-p^2}{4\pi\mu^2} \right)^{-\varepsilon} (1 + \eta) \frac{\Gamma(1 + \varepsilon)\Gamma^2(1 - \varepsilon)}{\Gamma(1 - 2\varepsilon)} \frac{(1 - \varepsilon)}{\varepsilon(1 - 2\varepsilon)}. \quad (3.44)$$

Note que no *gauge de Landau* ($\eta = -1$) a autoenergia dos quarks não contribue à amplitude de glúons virtuais, de modo que somente a correção do vértice precisa ser calculada. Em qualquer outro caso, $\Sigma_q(p)$ precisa ser renormalizada de modo a extrair a divergência UV.

Correção do Vértice

A amplitude virtual $\mathcal{A}_V$ é dada por

$$\begin{aligned} \mathcal{A}_V &= g_s^2 (T^a T^a) e Q_q \int \frac{d^d k}{(2\pi)^d} \bar{u}_q(p_1) \cdot \gamma^\alpha (\not{k} + \not{p}_1) \not{\epsilon}(q) (\not{k} - \not{p}_2) \gamma^\beta \cdot v(p_2) \\ &\times \left(g_{\alpha\beta} + \frac{\eta}{k^2} k_\alpha k_\beta \right) \times \frac{1}{D_3}, \end{aligned} \quad (3.45)$$

onde k é o momento do glúon virtual e

$$D_3 = k^2 \cdot (k + p_1)^2 \cdot (k - p_2)^2.$$

Combinada com a amplitude $\mathcal{M}_0$ e integrada no momento do loop resulta em

$$\begin{aligned} 2\mathcal{M}_0 \mathcal{A}_V^* &= -2ig_s^2 e^2 Q_q^2 C_F \int \frac{d^d k}{(2\pi)^d} \frac{1}{D_3} \left(\text{Tr}[\not{p}_2 \gamma^\mu \not{p}_1 \gamma^\nu (\not{k} + \not{p}_1) \gamma_\mu (\not{k} - \not{p}_2) \gamma_\nu] \right. \\ &\quad \left. + \frac{\eta}{k^2} \text{Tr}[\not{p}_2 \gamma^\mu \not{p}_1 \not{k} (\not{k} + \not{p}_1) \gamma_\mu (\not{k} - \not{p}_2) \not{k}] \right). \end{aligned} \quad (3.46)$$

Antes de prosseguirmos vamos analisar o termo dependente do parâmetro de gauge η , pois estaremos interessados em utilizar o gauge de Landau sendo preciso calcular as contribuições originadas por aquele termo. Para tanto retornemos à expressão da amplitude do vértice dada por (3.45) e isolemos o termo dependente do gauge. Após contração dos índices esse termo torna-se proporcional a

$$\begin{aligned} &\eta \int \frac{d^d k}{(2\pi)^d} \frac{1}{q^2 D_3} \bar{u}_q(p_1) \not{k} (\not{k} + \not{p}_1) \not{\epsilon}(q) (\not{k} - \not{p}_2) \not{k} v_q(p_2) \\ &= \eta \int \frac{d^d k}{(2\pi)^d} \frac{1}{k^2 D_3} \bar{u}_q(p_1) (k^2 + 2k \cdot p_1) \not{\epsilon}(q) (k^2 - 2k \cdot p_2) v_q(p_2), \end{aligned} \quad (3.47)$$

mas, $D_3 = k^2 \cdot (k^2 + 2k \cdot p_1) \cdot (k^2 - 2k \cdot p_2)$, de maneira que a expressão (3.47) acima adquire a forma

$$\eta \int \frac{d^d k}{(2\pi)^d} \frac{1}{k^4} \bar{u}_q(p_1) \not{\epsilon}(q) v_q(p_2). \quad (3.48)$$

Em regularização dimensional, por outro lado, se não é feita distinção entre pólos UV e pólos IR então

$$\int \frac{d^d k}{(2\pi)^d} \frac{1}{k^{2n}} = 0 \quad (3.49)$$

para qualquer $n > 0$ [11], de forma que o termo dependente do gauge não contribui para a expressão do vértice! Em qualquer gauge covariante, os termos finitos da correção de QCD, ao nível de 1-loop, ao vértice eletromagnético $\gamma q\bar{q}$ não-renormalizado são os mesmos.

É, portanto, nesse caso, muito mais vantajoso trabalhar no gauge de Landau, onde os diagramas de autoenergia não contribuem. Essa vantagem é ainda maior quando levamos em conta a simetria de gauge da amplitude. Em primeiro lugar, o termo dependente de η do propagador do glúon não altera as amplitudes reais calculadas no gauge de Feynman $\eta = 0$ pois por invariância do acoplamento de gauge sabemos que

$$k_g^\mu M_{Real,\mu} = 0 \quad , \quad (3.50)$$

onde k_g representa o momento do glúon emitido.

Além disso a função de onda do quark pode ser renormalizada multiplicativamente de modo que o propagador corrigido é dado por

$$S(p) = \frac{iZ_2}{\not{p}} \quad , \quad (3.51)$$

onde a constante de renormalização é obtida através da *condição on shell*

$$Z_2^{-1} = 1 - \frac{d\Sigma_q}{d\not{p}}(p^2 = 0) \quad . \quad (3.52)$$

A identidade de Ward-Takahashi, por outro lado nos assegura que

$$-ik_\mu \Gamma^\mu(p+k, p) = S^{-1}(p+k) - S^{-1}(p) \quad , \quad (3.53)$$

onde k_μ é o momento de um fóton externo que se acopla à corrente eletromagnética do quark e Γ^μ corresponde à amplitude do vértice corrigida, ou seja,

$$\Gamma^\mu(p+k, p) \simeq Z_1^{-1} \gamma^\mu \quad \text{conforme } k \rightarrow 0 \quad .$$

Expandindo a identidade em torno de $k = 0$, temos

$$\begin{aligned} -ik_\mu (Z_1^{-1} \gamma^\mu + \dots) &= S^{-1}(p) + k_\mu \frac{dS^{-1}(p+k)}{dp_\mu} (k=0) + \dots - S^{-1}(p) \\ -iZ_1^{-1} \not{k} &= -iZ_2^{-1} \not{k} \end{aligned}$$

até primeira ordem em k . Dessa forma, mostra-se, em QED, a igualdade

$$Z_1 = Z_2 \quad (3.54)$$

oriunda da identidade de Ward-Takahashi, a qual, na realidade, expressa de forma diagramática a conservação da corrente elétrica, que por sua vez, é uma consequência direta da invariância sob transformações de gauge do grupo $U(1)$ eletromagnético.

Pois bem, sendo $\Sigma(p) = 0$ no gauge de Landau, segue que $Z_2 = 1$ e, consequentemente, $Z_1 = 1$ até ordem α_s pela identidade de Ward-Takahashi. Trabalhando no

gauge de Landau não precisamos nos preocupar com a renormalização do vértice. Todos os pólos que surgem na expressão de Γ^μ são pólos IR, além disso, temos a garantia de que os termos finitos da expressão do vértice não dependem da escolha do gauge. Em qualquer outro caso, no entanto, ressaltamos que a autoenergia e o vértice precisam ser renormalizados. Nesse caso, $\Sigma_q(p)$ possui uma dependência explícita em η , no entanto, como $Z_1 = Z_2$ no esquema on-shell, essa dependência se cancela quando incluímos Z_1 na expressão do vértice e somamos todas as contribuições virtuais ao nível de 1-loop ao processo $\gamma^* \rightarrow q\bar{q}$, de forma que a seção de choque total inclusiva permanece invariante de gauge, como deve ser.

Talvez devido à tamanha conspiração de fatores que torna o cálculo das correções virtuais de QCD ao nível de 1-loop ao processo $\gamma^* \rightarrow q\bar{q}$, no gauge de Landau, tão simples em relação a outros processos, é que o seu cálculo completo nunca seja apresentado de forma clara, mesmo em livros bastante conhecidos e artigos de revisão de correções de QCD.

No livro “*Applications of Perturbative QCD*” de Richard D. Field [20], no qual grande parte dos cálculos aqui apresentados estão baseados, o autor sugere que, independentemente do gauge utilizado nos cálculos, é possível desprezar as contribuições de autoenergia, desde que trabalhem com os quarks *on shell*, ou seja, desde que $p^2 = 0$ e $\varepsilon > 0$. Note, que isso nos traz de volta ao resultado no gauge de Landau, onde, mesmo para $p^2 \neq 0$, temos $\Sigma(p) = 0$. No entanto, para $\eta \neq -1$, temos que a autoenergia do quark é algo proporcional a

$$\begin{aligned}\Sigma(p) &\sim (-p^2)^\varepsilon \left(\frac{1}{\varepsilon} + \text{finitos} \right) \\ &= \frac{1}{\varepsilon} + \text{finitos} + \ln(-p^2)\end{aligned}$$

após expansão em ε . Aqui existem duas divergências de naturezas distintas: o pólo em ε , que corresponde a uma divergência UV, e $\ln(-p^2)$ no limite $p^2 \rightarrow 0$, que corresponde a uma divergência de massa. Portanto é preciso ter cuidado com essa postura. No capítulo 6, vamos adotar o gauge de Feynman para realizar os cálculos e todo o programa de renormalização será colocado em prática no esquema *on shell*, dado que nesse gauge as autoenergias não podem ser desprezadas.

Esclarecidos tais pontos, retornemos ao cálculo da correção do vértice, agora simplificado pela exclusão dos termos dependentes do parâmetro de gauge.

O numerador da integral, após o cálculo dos traços em d dimensões é da forma

$$\begin{aligned}\mathcal{N} &= \text{Tr}[\not{p}_2 \gamma^\mu \not{p}_1 \gamma^\nu (\not{k} + \not{p}_1) \gamma_\mu (\not{k} - \not{p}_2) \gamma_\nu] \\ &= 4[(-8 + 4d)k \cdot p_1 p_1 \cdot p_2 + (8 - 6d + d^2)k^2 p_1 \cdot p_2 + (8 - 4d)k \cdot p_1 k \cdot p_2 \\ &\quad + (-8 + 4d)k \cdot p_2 p_1 \cdot p_2 + (8 - 4d)(p_1 \cdot p_2)^2]\end{aligned}$$

$$= 4(d-2)[4k \cdot p_1 k \cdot p_2 - q^4 + 2q^2(k \cdot p_2 - k \cdot p_1) + (d-4)q^2 \frac{k^2}{2}] . \quad (3.55)$$

O uso das integrais paramétricas é feito de acordo com o seguinte roteiro:

1. Combinamos os denominadores $(k + p_1)^2$ e $(k - p_2)^2$ usando a fórmula (D.10) e definindo

$$\begin{aligned} p_y &\equiv yp_1 - (1-y)p_2 , \\ D_1 &\equiv k^2 - 2k \cdot p_y . \end{aligned}$$

2. Combinamos o denominador D_1 , gerado pelo item anterior, e o denominador k^2 usando (D.11), identificando

$$\begin{aligned} c &= D_1 \text{ e } d = k^2, \text{ e definindo} \\ l &\equiv k - xp_y \text{ e } \Delta \equiv x^2 p_y^2 , \\ D_2 &\equiv l^2 - \Delta . \end{aligned}$$

3. Finalmente, usamos (D.12) com as mesmas definições do item anterior para calcular os termos que contém o denominador k^4 .

Implementando as mudanças decorrentes das definições acima nos numeradores, obtemos após uma certa álgebra o resultado

$$\Delta_V \equiv \frac{2\mathcal{A}_0\mathcal{A}_V^*}{2E_{cm}} = -\frac{8}{3}ig_s^2e^2Q_q^2Q \int_0^1 dx \int_0^1 dy \frac{d^n l}{(2\pi)^n} \frac{2x\mathcal{N}_1}{[l^2 - \Delta]^3} , \quad (3.56)$$

onde $\mathcal{N}_1$ é dado por

$$\mathcal{N}_1 = -2q^2 + 2xq^2 - 2x^2y(1-y)q^2 + (1-2\varepsilon)\frac{4}{n}l^2 . \quad (3.57)$$

Integrando primeiramente no momento do loop e depois nas variáveis paramétricas x e y , obtemos o seguinte produto de funções B (beta)

$$\begin{aligned} \Delta_V &= \frac{2i}{16\pi^2} \left(\frac{4\pi\mu^2}{-q^2} \right)^\varepsilon \Gamma(1+\varepsilon) [-B(-2\varepsilon-1,0)B(-\varepsilon-1,-\varepsilon-1) \\ &- B(-2\varepsilon+1,0)B(-\varepsilon,-\varepsilon) + B(-2\varepsilon,0)B(-\varepsilon-1,-\varepsilon-1) \\ &+ \frac{1}{\varepsilon}(1-2\varepsilon)B(-2\varepsilon+1,0)B(-\varepsilon,-\varepsilon)] , \end{aligned} \quad (3.58)$$

aonde após expansão em termos de funções Γ e com o auxílio da série

$$\frac{1}{1+x} = 1 - x + x^2 - \dots \quad (3.59)$$

obtemos um resultado promissor

$$\Delta_V = \frac{2\alpha_s}{3\pi} e^2 Q_q^2 \left(\frac{4\pi\mu^2}{-q^2} \right)^\varepsilon \frac{\Gamma(1+\varepsilon)\Gamma^2(1-\varepsilon)}{\Gamma(1-2\varepsilon)} \left(-\frac{2}{\varepsilon^2} - \frac{3}{\varepsilon} - 8 + \mathcal{O}(\varepsilon) \right) . \quad (3.60)$$

Sabendo que

$$\sigma_V = \Delta_V \Phi_2^d , \quad (3.61)$$

onde Φ_2^d é o fator de espaço de fase da equação (C.14), então

$$\sigma_V(\gamma^* \rightarrow q\bar{q}) = \frac{2\alpha_s}{3\pi} \sigma_0 \left(\frac{4\pi\mu^2}{-q^2} \right)^\varepsilon \frac{\Gamma(1+\varepsilon)\Gamma^2(1-\varepsilon)}{\Gamma(1-2\varepsilon)} \left(-\frac{2}{\varepsilon^2} - \frac{3}{\varepsilon} - 8 + \mathcal{O}(\varepsilon) \right) . \quad (3.62)$$

Usando a expansão das funções Γ é fácil mostrar que

$$\frac{\Gamma(1+\varepsilon)\Gamma^2(1-\varepsilon)}{\Gamma(1-2\varepsilon)} = 1 - \gamma\varepsilon + \frac{1}{12}(6\gamma^2 - \pi^2)\varepsilon^2 + \dots$$

Combinando todos os resultados, obtemos

$$\begin{aligned} \sigma_V(\gamma^* \rightarrow q\bar{q}) &= \frac{2\alpha_s}{3\pi} \sigma_0 \left[-\frac{2}{\varepsilon^2} - \frac{1}{\varepsilon} \left[2 \ln \left(\frac{4\pi\mu^2}{-q^2} \right) - 2\gamma + 3 \right] - \ln^2 \left(\frac{4\pi\mu^2}{-q^2} \right) \right. \\ &\quad \left. + (2\gamma - 3) \ln \left(\frac{4\pi\mu^2}{-q^2} \right) - \gamma^2 + 3\gamma + \frac{\pi^2}{6} - 8 + \dots \right] . \end{aligned} \quad (3.63)$$

Para q^2 “timelike”, ou seja, $-q^2 < 0$ é preciso usar a continuação analítica da função *logaritmo*, de modo que possamos combinar os resultados das contribuições reais e virtuais

$$\ln(-q^2) = \ln(q^2) - i\pi ,$$

e dessa forma encontramos o resultado final para a contribuição virtual à seção de choque no gauge de Landau

$$\begin{aligned} \sigma_V(\gamma^* \rightarrow q\bar{q}) &= \frac{2\alpha_s}{3\pi} \sigma_0 \left[-\frac{2}{\varepsilon^2} - \frac{1}{\varepsilon} \left[2 \ln \left(\frac{4\pi\mu^2}{Q^2} \right) - 2\gamma + 3 \right] - \ln^2 \left(\frac{4\pi\mu^2}{Q^2} \right) \right. \\ &\quad \left. + (2\gamma - 3) \ln \left(\frac{4\pi\mu^2}{Q^2} \right) - \gamma^2 + 3\gamma + \frac{\pi^2}{6} - 8 + \pi^2 + \dots \right] , \end{aligned} \quad (3.64)$$

onde mantivemos apenas os termos reais.

Finalmente, somando a contribuição real (3.38) e virtual (3.64), obtemos a seção de choque total do processo em NLO de QCD

$$\begin{aligned} \sigma_{NLO}(\gamma^* \rightarrow q\bar{q}) &= \sigma_0 + \sigma_V + \sigma_R = \frac{2\alpha_s}{3\pi} \sigma_0 \left[-\frac{7\pi^2}{6} + \frac{\pi^2}{6} + \pi^2 + \frac{19}{2} - 8 \right] + \sigma_0 \\ &= \frac{2\alpha_s}{3\pi} \sigma_0 \times \left(\frac{3}{2} \right) + \sigma_0 \\ &= \left(1 + \frac{\alpha_s}{\pi} \right) \times \sigma_0 . \end{aligned} \quad (3.65)$$

Veja que todos os pólos em ε e termos logarítmicos que continham o parâmetro de massa μ cancelaram-se na soma das duas contribuições, de modo que o resultado é finito em 4 dimensões, como esperávamos.

As correções de QCD em NLO ao processo $e^+e^- \rightarrow \gamma, Z \rightarrow q\bar{q}$ são exatamente as mesmas como já argumentamos, logo

$$\sigma_{NLO}(e^+e^- \rightarrow \gamma, Z \rightarrow q\bar{q}) = \left(1 + \frac{\alpha_s(s)}{\pi}\right) \sigma_0(e^+e^- \rightarrow \gamma, Z \rightarrow q\bar{q}) , \quad (3.66)$$

onde já inserimos a constante de acoplamento efetiva da QCD, $\alpha_s(s)$, dada por

$$\alpha_s(s) = \frac{4\pi}{(11 - \frac{2}{3}n_f) \ln(s/\Lambda^2)} ,$$

onde n_f é o número de sabores de quarks e Λ é uma escala de energia que deve ser obtida experimentalmente e que fisicamente corresponde ao ponto onde efeitos não-perturbativos passam a tornar-se importantes.

A seguir apresentamos as distribuições angular e de massa invariante do par $q\bar{q}$, levando em conta as contribuições dos quarks u , d , s e c , em Leading e Next-to-Leading-Order de QCD e do momento transversal do quark em ordem zero, obtidas por integração numérica de Monte Carlo, sem simulação de cortes experimentais. As correções de NLO de QCD, neste caso, são pequenas, dado que

$$\frac{\sigma_{NLO}}{\sigma_{LO}} = 1 + \frac{\alpha_s(s)}{\pi} \approx 1,04 , \quad (3.67)$$

para $s = m_Z^2$.

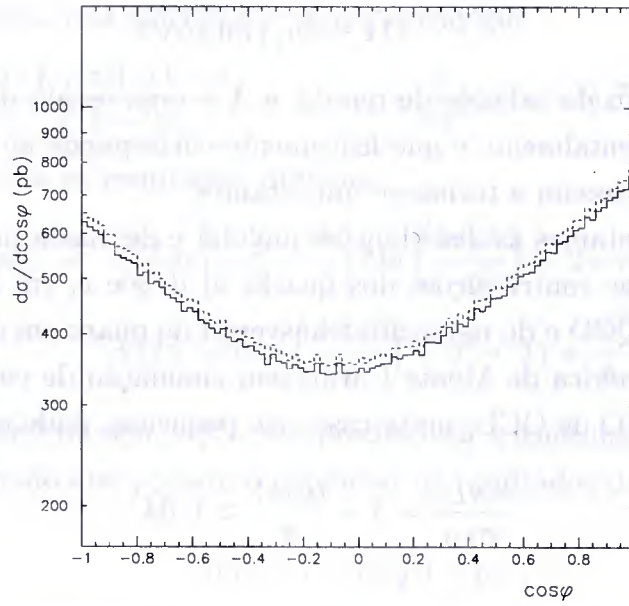


Figura 3.3: Distribuição angular dos quarks do estado final no LEP ($\sqrt{s} = 189$ GeV). O ângulo φ corresponde ao ângulo formado pelo 3-momento do elétron e o 3-momento do quark incidente. A linha sólida corresponde à seção de ordem zero e a linha segmentada corresponde à seção de choque em NLO de QCD.

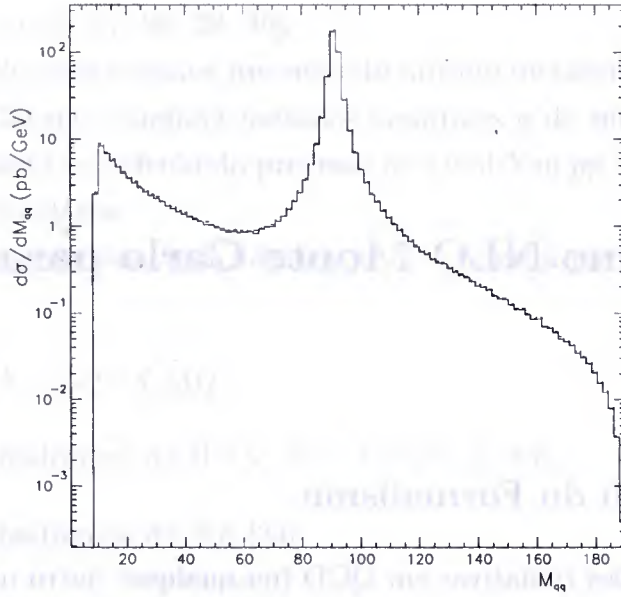


Figura 3.4: Distribuição em termos da massa invariante do par quark-antiquark no LEP. A linha sólida representa a predição de ordem zero e a linha segmentada a seção de choque corrigida em NLO de QCD.

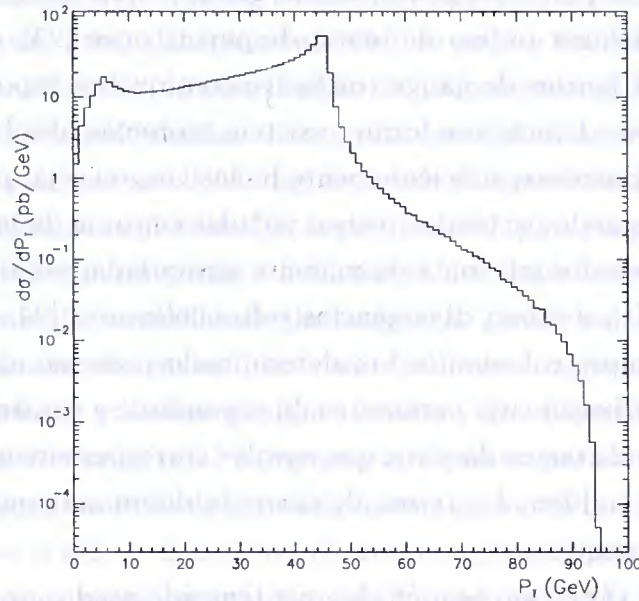


Figura 3.5: Momento transversal P_T do quark em LO.

Capítulo 4

O Formalismo NLO Monte Carlo para Correções de QCD

4.1 Descrição do Formalismo

Cálculos de correções radiativas em QCD (ou qualquer outro modelo realístico das interações fundamentais) requerem uma maneira apropriada de se tratar os diversos tipos de infinitos que, invariavelmente, ocorrem quando lidamos com teoria de perturbação.

Grandes esforços foram feitos nos últimos anos no sentido de justificar a teoria de perturbação como uma poderosa e insubstituível ferramenta de cálculo e investigação dos efeitos quânticos de ordem superior em fenômenos de escala subatômica. Veltman e 't Hooft mostraram, em um trabalho fundamental, que qualquer teoria que possui simetria por transformações de gauge é livre de divergências ultravioletas (UV) em qualquer ordem de teoria de perturbação [23], em outras palavras, mostraram que as teorias de gauge (inclusive com quebra espontânea de simetria) são renormalizáveis. Da mesma forma, existem teoremas absolutamente gerais que mostram que em processos suficientemente inclusivos, ou seja, processos onde todos os efeitos de emissão de partículas reais e virtuais em uma dada ordem de teoria de perturbação são levados em conta de maneira apropriada, são livres de divergências infravermelhas (IR), a saber, divergências soft e colineares [24, 25].

Na prática sempre calculamos um determinado processo até uma ordem finita de teoria de perturbação cujo parâmetro de expansão é a constante de acoplamento forte, onde o segundo termo da série, que envolve correções virtuais ao nível de 1-loop e emissão real de 1 glúon, é o termo de correção dominante quando essa constante é suficientemente pequena.

Em ordem α_s , vários são os métodos que têm sido usados no cálculo de correções radiativas. Todos têm em comum a motivação de, ao mesmo tempo tornar os cálculos analíticos mais simples usando aproximações de forma sistemática para as regiões

singulares do espaço de fase, e utilizar métodos de integração numérica para as partes finitas. Entre eles destacamos o *formalismo de subtração de dipolos*[26] e o formalismo de *cut-off*[27, 28, 29, 30].

Nesse capítulo, descrevemos um método híbrido de cálculo de seções de choque em NLO de QCD que combina métodos analíticos e de integração numérica, via Monte Carlo, usado no cálculo do processo de Drell-Yan $p\bar{p} \rightarrow e^+e^-X$ com inclusão de leptoquarks escalares.

Entre outros processos, essa técnica já foi utilizada para o cálculo de diversas reações como

- $p\bar{p} \rightarrow W^+X \rightarrow e^+\nu X$ [31]
- produção hadrônica de $W^\pm\gamma$, $Z\gamma$ e $W^\pm Z$ [32, 33]
- produção hadrônica de ZZ [34]
- produção direta de fótons [35]
- produção simétrica de dois hádrons [36]

e corresponde a uma extensão do *phase space slicing method* [37, 38] para o cálculo de processos que envolvem hádrons no estado inicial.

A principal motivação para o desenvolvimento e aplicação da técnica foi a necessidade de ampliar a descrição quantitativa de um dado processo, passando a um nível de aproximação superior ao *leading-log approximation (LLA)*, onde a seção de choque de Born, com o acoplamento efetivo para a energia total do processo, é convoluída com as funções de distribuição de pártons corrigidas em NLO de QCD.

Além disso, apesar de muitos desses processos já estarem calculados analiticamente, a informação que se obtém das seções de choque totais nunca é a abordagem ideal para comparação com os dados experimentais, pois na prática, é preciso levar em conta os cortes experimentais e a eficiência dos detectores. Nesse ponto é que o formalismo se vale da versatilidade da integração numérica via Monte Carlo, onde cortes experimentais, eficiências, etc. são facilmente implementados no programa além de nos permitir calcular qualquer número de observáveis simultaneamente, bastando apenas construir um histograma apropriado para a quantidade desejada.

Para se alcançar um nível de acurácia teórica maior do que o *LLA* é preciso tratar os processos $2 \rightarrow 3$ e $2 \rightarrow 2$ ao nível de 1-loop apropriadamente, e nesse sentido é que incorporamos os aspectos analíticos do formalismo *NLO Monte Carlo*.

Ao calcularmos os subprocessos de ordem α_s , divergências soft, colineares e ultravioletas surgem e são tratadas usando regularização dimensional. Adotando um

esquema de renormalização conveniente, subtraímos as divergências ultravioletas que surgem dos diagramas que contém um loop gluônico. Ainda assim restam divergências soft e colineares em todos os subprocessos, $2 \rightarrow 2$ ou $2 \rightarrow 3$.

A idéia básica é dividir o espaço de fase em três regiões distintas: *região soft* (S), *região colinear* (C) e *região hard* (H). Na prática essa divisão é feita introduzindo-se os cut-offs teóricos δ_s , para delimitar a região soft e δ_c para delimitar a região colinear. Se δ_s é escolhido suficientemente pequeno, então o subprocesso $2 \rightarrow 3$ relevante, no caso, o processo de emissão de glúon $q\bar{q} \rightarrow ge^+e^-$ é calculado usando a *aproximação eikonal*, onde se coloca o momento do glúon igual a zero no numerador da amplitude. A expressão resultante é então integrada analiticamente na região soft do espaço de fase e o resultado somado à contribuição dos processos virtuais - as divergências soft se cancelam nesse estágio restando as colineares. As partes finitas, que aliás, já dependem explicitamente de δ_s , que surgem após o cancelamento das divergências soft, são incluídas na contribuição de 2 corpos à seção de choque total. Tais partes finitas serão posteriormente incluídas em um programa de simulação numérica.

Em termos práticos, a divisão do espaço de fase é realizada da seguinte forma

$$\sigma_{total} = \sigma_2 + \sigma_3 \quad (4.1)$$

$$\sigma_3 = \frac{1}{2s} \int_S |\mathcal{M}|^2 d\Phi_3 + \frac{1}{2s} \int_H |\mathcal{M}|^2 d\Phi_3 \quad (4.2)$$

e, posteriormente, a região H é dividida em duas partes

$$\sigma_3 = \frac{1}{2s} \int_S |\mathcal{M}|^2 d\Phi_3 + \frac{1}{2s} \int_{HC} |\mathcal{M}|^2 d\Phi_3 + \frac{1}{2s} \int_{\overline{HC}} |\mathcal{M}|^2 d\Phi_3, \quad (4.3)$$

onde, HC corresponde à região onde existem divergências colineares, e $\overline{HC}$ à região livre de quaisquer divergências. As seções de choque elementares calculadas em S e HC são somadas em σ_2 , enquanto aquelas calculadas em 4 dimensões em $\overline{HC}$ são absorvidas em σ_3 .

Se δ_c é suficientemente pequeno, então, em HC é válido usar a *Leading Pole Approximation* (LPA), onde na amplitude de um determinado processo, mantemos apenas aqueles termos que são divergentes quando algum invariante de Mandelstam vai a zero no regime colinear. Após integração no espaço de fase, as divergências são fatorizadas e incluídas nas funções de distribuição de pártons apropriadas ou se cancelam com termos semelhantes que sobreviveram dos diagramas virtuais. Novamente, os termos finitos que restam após o cancelamento ou fatorização das divergências colineares são incluídos em σ_2 , enquanto que, na região $\overline{HC}$ os elementos de matriz relevantes são calculados exatamente em 4 dimensões.

O cálculo agora consiste em duas partes: um conjunto de contribuições de 3 corpos - σ_3 - e um conjunto de contribuições de 2 corpos - σ_2 , todos finitos, mas

com dependências explícitas nos parâmetros de corte δ_s e δ_c . Contudo, qualquer observável inclusivo recebe contribuição de σ_2 e σ_3 , de modo que, na soma, a dependência em δ_s e δ_c deve ser cancelada. Fisicamente, o resultado não pode depender de nossa escolha precisa da região limite onde realizamos o cálculo analítico e o cálculo numérico. É condição necessária, portanto, que as seções de choque inclusivas numéricas sejam independentes da escolha desses cut-offs, evidentemente respeitando o fato de que devem ser pequenas o suficiente para que as aproximações soft e colinear sejam válidas.

Em todos os processos em que a técnica foi utilizada, checkou-se que os resultados são estáveis em relação a variações de δ_s e δ_c . A checagem numérica de insensibilidade a essas variações em relação ao nosso processo pode ser encontrada no capítulo 8.

4.2 O Aniquilamento $e^+e^- \rightarrow$ Hádrons como Exemplo

Cada passo descrito acima pode ser exemplificado através de sua aplicação no cálculo das correções de ordem α_s ao processo $e^+e^- \rightarrow$ hádrons que discutimos no capítulo anterior, onde apresentamos um cálculo analítico completo.

Nosso interesse maior é mostrar que, pela utilização do formalismo podemos obter o mesmo resultado para a seção de choque em *next-to-leading-order* dada por (3.66) mostrando de maneira explícita o cancelamento entre os termos que contém dependência nos cut-offs δ_s e δ_c . Isso só é possível, a princípio, em reações cujos estados iniciais não contém hádrons, pois não há funções de distribuição de pártons a serem integradas, o que possibilita o cálculo analítico completo.

Não mostraremos todos os detalhes de cada passo dos cálculos, nos restringindo a apresentar os resultados e apenas comentando a obtenção das expressões, isso será feito nos próximos capítulos onde calculamos as correções de QCD para o processo de Drell-Yan com a inclusão de leptquarks escalares. No entanto, em ambos os casos, as técnicas utilizadas são as mesmas, apesar do processo de Drell-Yan ser mais complicado devido à presença de funções de distribuição de pártons.

4.2.1 Processo de Ordem Zero

O primeiro passo na aplicação do formalismo é o cálculo do processo em *leading order* (LO) em $d = 4 - 2\epsilon$ dimensões espaço-temporais. É essencial obter a amplitude e o espaço de fase de 2 partículas em d dimensões, pois suas expressões serão utilizadas em todos os passos do cálculo da seção de choque de 2 corpos σ_2 .

A exemplo do que fizemos no capítulo anterior, vamos calcular, de fato, o decai-

mento de um fóton virtual γ^* em um par $q\bar{q}$ e, depois, no final do dia, simplesmente trocar $\sigma_0(\gamma^* \rightarrow q\bar{q})$ por $\sigma_0(e^+e^- \rightarrow \gamma, Z \rightarrow q\bar{q})$, dado que as correções de QCD são as mesmas.

A amplitude do processo $\gamma^* \rightarrow q\bar{q}$ é dada por (3.1)

$$\mathcal{M}_0 = -ieQ_q \bar{u}(p_1) \not{\epsilon}(q)v(p_2) ,$$

e o seu módulo quadrado, após soma sobre números quânticos do estado final e média dos números quânticos do estado inicial é dado por (3.3) com $d = 4 - 2\varepsilon$

$$|\overline{\mathcal{M}_0}|^2 = 12 \cdot 4\pi\alpha Q_q^2(1 - \varepsilon)Q^2 . \quad (4.4)$$

O espaço de fase do processo $1 \rightarrow 2$ em d dimensões é o da expressão (C.14)

$$\Phi_2^d = \frac{(4\pi)^\varepsilon}{8\pi} \frac{\Gamma(1 - \varepsilon)}{\Gamma(2 - 2\varepsilon)} Q^{-2\varepsilon} ,$$

de modo que

$$\sigma_0(\gamma^* \rightarrow q\bar{q}) = 3\alpha Q_q^2(1 - \varepsilon)Q \frac{\Gamma(1 - \varepsilon)}{\Gamma(2 - 2\varepsilon)} \left(\frac{4\pi\mu^2}{Q^2} \right)^\varepsilon . \quad (4.5)$$

A seguir descrevemos o cálculo das correções em NLO.

4.2.2 Correções Virtuais

As correções virtuais são calculadas no esquema de regularização dimensional, como já dissemos, e no presente caso são dadas por

$$\sigma_V(\gamma^* \rightarrow q\bar{q}) = \frac{2\alpha_s}{3\pi} \sigma_0 \left(\frac{4\pi\mu^2}{Q^2} \right)^\varepsilon \frac{\Gamma(1 - \varepsilon)}{\Gamma(1 - 2\varepsilon)} \left(-\frac{2}{\varepsilon^2} - \frac{3}{\varepsilon} - 8 + \frac{2\pi^2}{3} + \mathcal{O}(\varepsilon) \right) , \quad (4.6)$$

obtidas de (3.64).

4.2.3 Radiação de Glúons Reais

Até aqui nada de novo, correções virtuais envolvem cálculos de integrais de loops através das técnicas padrão e podem, inclusive, ser encontradas na literatura.

A novidade está no cálculo das amplitudes de emissão de glúon real* nas regiões soft e colinear do espaço de fase. Ao contrário do que fizemos no capítulo 3, onde o elemento de matriz em d dimensões do processo $\gamma^* \rightarrow q\bar{q}g$ era integrado em todo o espaço de fase de 3 partículas em d dimensões, estaremos interessados aqui apenas

*No processo de Drell-Yan, além da emissão de glúons reais, existem processos envolvendo glúons no estado inicial, o que não ocorre aqui.

naquelas regiões onde um glúon adicional no estado final contribui para a seção de choque de 2 jatos. Isso ocorre em duas situações diferentes: primeiro, quando o glúon não tem energia suficiente para dar origem a mais um jato identificável, além daqueles provenientes da hadronização do par $q\bar{q}$, ou seja, quando o glúon é suficientemente soft; segundo, quando um hard glúon é emitido na mesma direção definida pelo momento de um dos quarks, ou seja, o glúon pode formar um jato, mas os seus estados hadrônicos estão degenerados com os originados do quark (ou antiquark), são, portanto, glúons colineares. Além disso existe uma região do espaço de fase em que ambas as situações podem ocorrer simultaneamente, ou seja, um glúon soft pode ser emitido colinearmente à direção dos quarks. A todos os casos estão associadas divergências expressas em termos de pólos do tipo ε^{-i} ($i = 1, 2$), para o limite $\varepsilon \rightarrow 0$.

A idéia é, portanto, calcular o elemento de matriz do processo $\gamma^* \rightarrow q\bar{q}g$ e o espaço de fase de 3 partículas de forma aproximada nos regimes soft e colinear do glúon. O objetivo é demonstrar os teoremas de fatorização de divergências soft e colineares em NLO de QCD [26, 38] para o processo em questão. Para divergências soft, demonstra-se que

$$\mathcal{M}_{R,soft} \simeq \mathcal{M}_0 \times 2g_s T^a \left[\frac{p_2^\mu}{t} - \frac{p_1^\mu}{s} \right] \epsilon_\mu^*(p_3) , \quad (4.7)$$

onde os invariantes e momentos estão definidos no capítulo anterior, ao passo que, para as colineares, temos o seguinte

$$|\overline{\mathcal{M}}_R|_{col}^2 \simeq 4\pi\alpha_s \mu^{2\varepsilon} \frac{1}{t} 2P_{qq}(z, \varepsilon) |\overline{\mathcal{M}}_0|^2 , \quad (4.8)$$

onde $P_{qq}(z, \varepsilon)$ é a função de Altarelli-Parisi que descreve o splitting $q \rightarrow q + g$ no regime colinear

$$P_{qq}(z, \varepsilon) = \frac{4}{3} \left[\frac{1+z^2}{1-z} - \varepsilon(1-z) \right] . \quad (4.9)$$

Expressão semelhante é obtida para $\bar{q} \rightarrow \bar{q} + g$, com t trocado por s em (4.8).

Essas expressões são agora integradas nas regiões soft e colinear do espaço de fase aproximado de 3 partículas. A região soft é definida por

$$0 < E_{gluon} < \delta_s \frac{Q}{2} \quad (4.10)$$

e o espaço de fase nessa aproximação pode ser encontrado no capítulo 6. A região colinear, por sua vez, é definida por

$$0 < t, s < \delta_c Q^2 \quad (4.11)$$

e o espaço de fase aproximado é o de (C.36).

Após integração, obtemos os seguintes resultados

$$\sigma_{soft} = \frac{\alpha_s}{2\pi} C_F \left(\frac{4\pi\mu^2}{Q^2} \right)^\varepsilon \frac{\Gamma(1-\varepsilon)}{\Gamma(1-2\varepsilon)} \left(\frac{2}{\varepsilon^2} - \frac{4}{\varepsilon} \ln \delta_s + 4 \ln^2(\delta_s) \right) \times \sigma_0, \quad (4.12)$$

$$\begin{aligned} \sigma_{HC} = & \frac{\alpha_s}{2\pi} C_F \left(\frac{4\pi\mu^2}{Q^2} \right)^\varepsilon \frac{1}{\Gamma(1-\varepsilon)} \left(\frac{1}{\varepsilon} (3 + 4 \ln \delta_s) - 3 \ln \delta_c - 4 \ln \delta_c \ln \delta_s - 2 \ln^2(\delta_s) \right. \\ & \left. + 7 - \frac{2\pi^2}{3} \right) \times \sigma_0, \end{aligned} \quad (4.13)$$

onde nossas expressões para σ_V , σ_{soft} e σ_{HC} concordam com aquelas dadas em [38].

4.2.4 Seção de Choque de 2 Corpos

A seção de choque de 2 corpos agora é dada pela soma

$$\sigma_2 = \sigma_V + \sigma_{soft} + \sigma_{HC} \quad (4.14)$$

$$= \sigma_0 \left[1 + \frac{\alpha_s}{2\pi} C_F (2 \ln^2(\delta_s) - 4 \ln \delta_s \ln \delta_c - 3 \ln \delta_c - 1) \right]. \quad (4.15)$$

Veja que os pólos em ε se cancelam nesse estágio, o que faz de σ_2 um observável *infrared safe*, o que significa que, se os mesmos cortes que definimos aqui para delimitar as regiões soft e colinear forem usados em laboratório, podemos medir e comparar a seção de choque de 2 jatos predita pela QCD. Nesse contexto, σ_2 seria um observável físico mesmo com a dependência nos cut-offs.

A seção de choque de 2 jatos foi calculada por Sterman e Weinberg na década de 70, de uma forma muito parecida com a que fizemos aqui, mas no esquema de glúon massivo [39].

4.2.5 Seção de Choque de 3 Corpos

Calculado σ_2 , está terminada a fase de cálculos analíticos. Agora basta obter o quadrado do elemento de matriz do processo $\gamma^* \rightarrow q\bar{q}g$ e implementar os resultados em um programa numérico de integração em espaços de fase em 4 dimensões. Posteriormente será preciso checar a invariância da seção de choque total inclusiva $\sigma_2 + \sigma_3$ com relação à escolha dos valores de δ_s e δ_c , respeitando sempre o fato de que devem assumir valores bem menores do que 1, de maneira que sejam válidas as aproximações. Isso será feito quando estudarmos o processo de Drell-Yan. Por hora, queremos mostrar de uma maneira formal, que a utilização do formalismo nos leva a um resultado correto, a saber, a seção de choque em NLO do processo $e^+e^- \rightarrow$ hádrons dada por (3.66).

Para tanto será necessário integrar analiticamente no espaço de fase de 3 partículas, dado por (C.43), na região complementar $\overline{HC}$ definida por

$$\delta_s \frac{Q}{2} < E_{gluon} < \frac{Q}{2} , \quad (4.16)$$

$$\delta_c Q^2 < s, t < Q^2 . \quad (4.17)$$

Ao invés de x_1 e x_2 , escolhemos t e u como variáveis de integração, através das relações (3.16) e (3.17), de modo que integrando a expressão

$$|\overline{\mathcal{M}}_R|^2 = 32\pi\alpha Q_q^2 \cdot 4\pi\alpha_s \times 12 \left[\frac{s}{t} + \frac{t}{s} + 2 \frac{uQ^2}{st} \right] \quad (4.18)$$

obtida de (3.24) fazendo $\varepsilon = 0$ e somando sobre as cores do estado final, obtemos

$$\sigma_3 = \frac{\alpha_s}{2\pi} C_F \left[-2 \ln^2(\delta_s) + 4 \ln \delta_s \ln \delta_c + 3 \ln \delta_c + \frac{5}{2} \right] \times \sigma_0 . \quad (4.19)$$

Esse resultado foi obtido com o auxílio do pacote Mathematica e desprezados todos os termos que anulavam-se no limite $\delta_s, \delta_c \rightarrow 0$.

Somando σ_2 e σ_3 , podemos ver que todos os termos dependentes dos cut-offs se cancelam restando um termo finito que corresponde exatamente à seção de choque em NLO dada por (3.65), de maneira que trocando $\sigma_0(\gamma^* \rightarrow q\bar{q})$ por $\sigma_0(e^+e^- \rightarrow q\bar{q})$ obtemos (3.66)

$$\sigma_2 + \sigma_3 = \left(1 + \frac{\alpha_s}{\pi} \right) \times \sigma_0(e^+e^- \rightarrow q\bar{q}) . \quad (4.20)$$

A seguir iniciamos os cálculos em NLO de QCD para o processo de Drell-Yan usando o formalismo aqui apresentado.

Capítulo 5

Correções em NLO de QCD para as Funções de Distribuição de Pártons

5.1 O Modelo a Pártons

Sob a luz da realidade experimental do *scaling de Bjorken*, o modelo a pártons surge no final da década de 60 como uma poderosa e eficiente ferramenta teórica no estudo da estrutura interna dos hádrons. Sua principal virtude é a simplicidade com que descreve a estrutura interna de um hádron no regime de altas energias, descrevendo de forma precisa o scaling de Bjorken e permitindo o cálculo de seções de choque inclusivas. Vamos rever quais são os princípios básicos do modelo a pártons.

Nesse modelo, um hádron é descrito como um objeto extenso, constituído de partículas pontuais que não interagem entre si no regime de colisões altamente energéticas. Tais partículas são o que chamamos genericamente de *pártons* e que serão, dentro do contexto da QCD , identificados com quarks e glúons.

Sendo mais específicos, considere o processo de espalhamento inelástico profundo (DIS) entre um elétron e um hádron pela troca de um fóton de grande virtualidade. No referencial do centro de massa do sistema ($e^- \oplus$ hádron) dois fenômenos importantes ocorrem com o hádron: o fenômeno de dilatação do tempo que diminui drasticamente a razão na qual os pártons interagem entre si, de modo que no curto espaço de tempo em que um fóton virtual interage com um párton carregado, ele se comporta como uma partícula livre, e a contração espacial na direção da colisão. Dessa forma, conforme a energia do centro de massa aumenta, o tempo de vida dos estados partônicos virtuais é aumentada, enquanto a distância a ser percorrida pelo fóton através do hádron é diminuída. O hádron portanto pode ser caracterizado como um soma de estados partônicos virtuais estáveis não-interagentes durante todo o tempo de interação elétron-hádron.

Como os pártons não interagem durante esse período, podemos atribuir a cada párton uma fração $0 < x < 1$ de energia e momento do hádron do qual ele faz

parte, de modo que a soma dos momentos e energias de todos os pártons resulta no quadrimomento do hádron. Faz sentido, dessa maneira, dizer que o elétron interage agora com pártons de momento definido, ao invés de interagir com o hádron como um todo. Além disso, se o momento transferido é muito alto, o fóton virtual não pode viajar muito longe, logo, se a densidade de pártons não for muito grande, o elétron será capaz de interagir com apenas um párton.

O modelo assume que a distribuição dos momentos entre os pártons é independente, ou seja, a probabilidade de que um dado párton i carregue uma fração x_i do momento total é independente da probabilidade de que um outro párton j tenha fração de momento x_j do momento total, de uma forma mais precisa, se $\phi_i(x)$ representa uma tal probabilidade e x é a fração de momento que o párton i carrega então

$$\sum_i \int_0^1 dx x \phi_i(x) = 1 . \quad (5.1)$$

A partir dessas hipóteses, o espalhamento de alta energia torna-se essencialmente incoerente e clássico. Ou seja, as interações entre pártons ocorrem em uma escala de tempo muito maior do que a interação elétron-párton, de forma que não há interferência desses dois fenômenos. A seção de choque pode, então, ser calculada combinando-se probabilidades (ϕ_i) ao invés de amplitudes.

A cinemática do modelo pode ser resumida da seguinte forma

Hádron

$$\text{Energia} \rightarrow E$$

$$\text{Momento} \rightarrow P_L , P_T = 0$$

Párton

$$\text{Energia} \rightarrow xE$$

$$\text{Momento} \rightarrow xP_L , p_T = 0 ,$$

onde assumimos que o hádron e os seus pártons movem-se ao longo do eixo z com momento longitudinal P_L e xP_L , respectivamente, sem qualquer momento transversal $P_T = 0$, de modo que os seus 4-momentos podem ser escritos de acordo com

$$P = (E, P_L, 0, 0) \quad (5.2)$$

$$p = (xE, xP_L, 0, 0) . \quad (5.3)$$

À primeira vista, essas relações cinemáticas parecem estar incorretas, ou pelo menos contraditórias. Se o parton carrega momento xP_L , então sua energia será dada por

$$\begin{aligned} E_p^2 &= m^2 + (xP_L)^2 \\ E_p^2 &= m^2 + x^2(E^2 - M^2) \\ E_p &= \sqrt{m^2 - x^2M^2 + x^2E^2} , \end{aligned} \quad (5.4)$$

ou seja, a energia do parton só é igual a xE se a sua massa e a do hádron são iguais a zero. No cálculo do espalhamento $e^-e^+ \rightarrow q\bar{q}$ vimos que se o quark tivesse massa nula, assim como o glúon, havia o problema da divergência colinear, pois nesse caso, o glúon podia ser emitido na mesma direção do quark. Logo, se hádron e pártons viajam na mesma direção, ou seja, colinearmente, então ambos tem de ter massa nula. O impasse, contudo, se resolve se trabalhamos em um referencial onde $E, |\vec{P}| \gg m, M$ de modo que todas as massas possam ser ignoradas. No referencial onde o momento é infinito, a cinemática se torna exata justificando a associação

$$p_\mu = xP_\mu . \quad (5.5)$$

Note também que nesse referencial $m = xM$, o que sugere que não há energia de interação.

Por último, devido ao fenômeno da liberdade assintótica da QCD, no regime de altas energias é permitido o uso da QCD perturbativa para a análise da influência da presença de glúons no interior dos hádrons, já antecipando a natureza dos pártons.

Dessa maneira, a seção de choque total de uma colisão hádron-hádron em léptons mais estados hadrônicos

$$H_1(P_1) + H_2(P_2) \rightarrow \ell_1 + \ell_2 + X , \quad (5.6)$$

que é o tipo que estamos interessados em estudar, é dada, pelo *teorema de fatorização* [40], inspirado no modelo a pártons, por

$$\sigma(H_1 H_2 \rightarrow \ell_1 \ell_2 X) = \sum_{i,j} \int dx_a dx_b [f_i^{H_1}(x_a, Q^2) f_j^{H_2}(x_b, Q^2) + a \leftrightarrow b] \times \hat{\sigma}_{ij}(x_a x_b S, Q^2) , \quad (5.7)$$

onde $S = (P_1 + P_2)^2$, sendo P_1 e P_2 os momentos dos hádrons 1 e 2 respectivamente e $f_i^H(x, Q^2)$ representam as *funções de distribuição de pártons* no interior dos hádrons com $\hat{\sigma}_{ij}$ a seção de choque partônica em d dimensões. Para o processo de Drell-Yan no Tevatron temos, $H_1 \mapsto$ próton, $H_2 \mapsto$ antipróton e $\ell_{1,2} \mapsto e^\pm$, por exemplo.

5.2 As Funções de Distribuição de Pártons

A simplicidade do modelo a pártons deve-se em boa parte à suposição de que as partículas que constiuem os hádrons estão essencialmente livres durante o intervalo

de tempo em que a colisão dura acontece. Ainda assim, sabemos que tais partículas nunca podem ser isoladas de seus hádrons devido aos fenômenos não-perturbativos das interações fortes, de modo que se faz absolutamente essencial o conhecimento da exata proporção de cada tipo de parton na constituição de um dado hádron quando estamos interessados em determinar as propriedades dessas partículas e as interações a que estão sujeitas.

Tais informações estão contidas nas *funções de distribuição de pártons* (PDF's) de um determinado hádron: $f_i^H(x, Q^2)$, e dependem da natureza dos fenômenos em escala de massa hadrônica, portanto, do regime não-perturbativo da QCD.

Aqui $f_i^H(x, \mu)dx$ é interpretado como a probabilidade de encontrar um parton do tipo i ($=$ glúon, $u, d, s, c, \bar{u}, \bar{d}, \dots$) em um hádron do tipo H carregando uma fração entre x e $x + dx$ do momento total do hádron.

Podemos entender agora o profundo significado do *teorema de fatorização* contido na expressão (5.7). Dado que a seção de choque partônica

$$d\hat{\sigma}_{ij}(\hat{s}, Q^2) = \frac{1}{2\hat{s}} |\overline{M}_{ij}|^2 d\Phi_2^d, \quad (5.8)$$

onde Φ_2^d é o espaço de fase de 2 partículas em d dimensões, pode ser expandida em uma série perturbativa em α_s , e portanto contém toda a informação a respeito do regime duro da colisão, então (5.7) nos permite separar a contribuição dos regimes perturbativo e não-perturbativo no cômputo de uma seção de choque hadrônica.

Contudo, ao contrário de $\hat{\sigma}_{i,j}$, as PDFs não podem ser calculadas em teoria de perturbação. Não há outra alternativa senão extraí-las da experiência, assim como a constante de acoplamento forte α_s . Hoje em dia alguns grupos ao redor do mundo trabalham com ajustes globais para as funções de distribuição. A idéia do ajuste global é ajustar as PDFs de modo que teoria e experimento concordem para uma vasta gama de reações.

Conforme a energia e a luminosidade dos aceleradores de partículas aumentar, melhores ajustes serão necessários, dada a precisão com que os dados experimentais serão obtidos. Os grandes desvios Leading-Order/Next-to-Leading-Order típicos da QCD poderão ser cada vez mais testados, exigindo esforços teóricos no sentido de se calcular $\hat{\sigma}_{i,j}$ em ordens superiores e de se eliminar as incertezas nas PDFs.

A dependência de f_i^H com a escala de massa Q pode ser entendida quando definimos as funções de distribuição de forma rigorosa em termos de produto de operadores, onde as f_i^H são definidas como elementos de matriz hadrônico do operador que conta o número de quarks ou glúons (operadores de criação e de aniquilação) [40]. A escala Q é introduzida quando renormaliza-se os operadores.

Na próxima seção veremos; em linhas gerais, em que condições é válido o *teorema de fatorização*.

5.2.1 Universalidade das PDFs

As correções de ordem α_s ao processo de *Drell-Yan* são de duas naturezas:

- Correções ao subprocesso partônico $q\bar{q} \rightarrow \gamma, Z \rightarrow e^+e^-$. Tais correções são provenientes dos seguintes subprocessos
 1. $q + \bar{q} \rightarrow e^+e^-$: Glúons Virtuais.
 2. $q + \bar{q} \rightarrow g + e^+e^-$.
 3. $q(\bar{q}) + g \rightarrow q(\bar{q}) + e^+e^-$.
- Correções às funções de distribuição de pártons (PDFs). Cujas correções de ordem α_s podem ser calculadas a partir dos processos partônicos do *Deep Inelastic Scattering - DIS*
 1. $e^- + q(\bar{q}) \rightarrow e^- + q(\bar{q})$: Glúons Virtuais.
 2. $e^- + q(\bar{q}) \rightarrow g + q(\bar{q})e^-$.
 3. $e^- + g \rightarrow q\bar{q} + e^-$.

Essas correções sempre devem ser incluídas em conjunto, pois a seção de choque hadrônica é dada por um produto de PDFs e seções de choque partônicas.

Além daqueles subprocessos elementares que corrigem o processo de Drell-Yan em NLO de QCD, existe, ainda em ordem α_s , uma classe bastante abrangente de processos envolvendo os pártons que participam do subprocesso partônico duro e os quarks e glúons espectadores. Essa situação é ilustrada na figura abaixo.

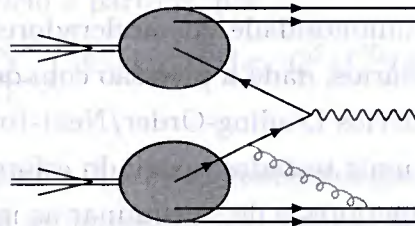


Figura 5.1: Contribuição de pártons espectadores no processo Drell-Yan.

Esse tipo de processo é suprimido por potências da razão Λ^2/Q^2 , onde $\Lambda \sim 200$ MeV é o parâmetro de escala da QCD e Q é virtualidade do fóton produzido no processo partônico, sendo que no limite de altas energias $\Lambda \ll Q$. Tais correções são denominadas “*Higher Twists*” e são completamente desprezíveis em nosso processo. Esse fato é muito importante para a validade do modelo a pártons, pois essas

contribuições poderiam impedir qualquer tentativa de fatorização para a seção de choque hadrônica como em (5.7).

Correções análogas surgem no DIS e também são suprimidas por potências de Λ^2/Q^2 .

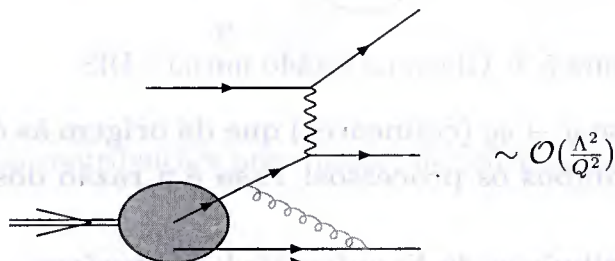


Figura 5.2: Contribuição de pártons espectadores no DIS.

Existe uma outra fonte de problemas ao processo de fatorização. As correções de QCD ao modelo a pártons dão origem a termos logarítmicos de ordem $\mathcal{O}(\alpha_s \ln \frac{Q^2}{M^2})$. Embora α_s seja pequena para $M \ll Q$, $\ln Q^2/M^2$ com certeza não é, o que invalida a teoria de perturbação, implicando em contradições com o modelo a pártons, onde as partículas que estão no interior dos hádrons são tratadas como partículas livres para grandes energias e ordem zero em QCD, que é a ordem dominante.

Veremos mais adiante que tais logaritmos surgem através de divergências colineares, contudo, felizmente elas se cancelam no cálculo da seção de choque total hadrônica por estarem adequadamente contidas tanto em $\hat{\sigma}$ quanto nas PDFs.

Podemos começar a entender isso de forma intuitiva, considerando os subprocessos representados nas figuras (5.3) e (5.4). Suponha que um glúon incidente decaia em um par $q\bar{q}$ colineares à direção do glúon no processo de Drell-Yan

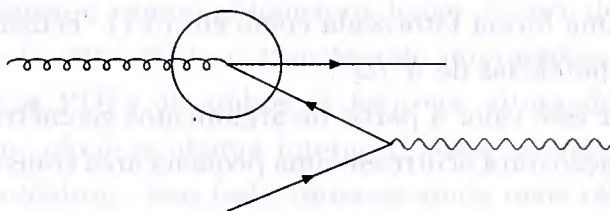


Figura 5.3: Glúon no estado inicial : Drell-Yan.

o antiquark então se aniquila com o quark incidente que vem do outro hádron dando origem ao fóton.

No DIS, por outro lado, o glúon também pode decair em um par $q\bar{q}$ colineares, com a diferença que o quark é atingido posteriormente pelo γ criado pela interação com o lépton do estado inicial.

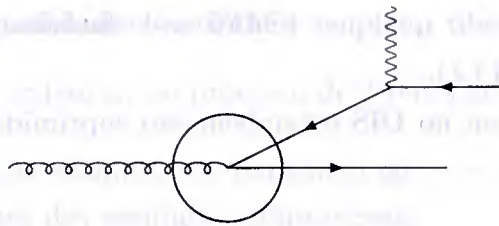


Figura 5.4: Glúon no estado inicial : DIS.

O processo $g \rightarrow q\bar{q}$ (colineares) que dá origem às divergências colineares é comum a ambos os processos! Essa é a razão dos cancelamentos.

5.2.2 Contribuição de Eventos Multipartônicos

Como sabemos, no modelo a pártons, hádrons energéticos podem ser considerados como uma superposição de estados partônicos confinados em um disco de largura longitudinal L/γ , onde γ é o fator de Lorentz, e largura transversal $L \sim 1fm$, que é a escala de comprimento típica de um nucleon. Cada párton (quark ou glúon) carrega uma fração finita de momento x do momento do hádron. Além disso, existe também uma nuvem de “soft quanta” associado ao hádron (soft glúons, pares virtuais $q\bar{q}$, etc).

Com base nesse modelo, o processo de *Drell-Yan* é um processo no qual dois hádrons relativísticos aproximam-se rapidamente um do outro, ao longo da direção z por exemplo, a colisão em si ocorre entre um párton de um dos hádrons e um párton do outro, que se aniquilam dando origem a um bóson vetorial de momento q^μ onde $|q^2| \gg \Lambda^2$.

Processos multipartônicos, do tipo $n \rightarrow m$, $n \geq 3$, onde mais de dois pártons participam de um único subprocesso duro, fazem parte daquela classe de processos que poderiam invalidar qualquer tentativa de expressar a seção de choque hadrônica inclusiva total em uma forma fatorizada como em (5.7). Felizmente, tais processos são suprimidos por potências de Λ^2/Q^2 .

Podemos estimar esse valor à partir de argumentos geométricos . Pelo *princípio da incerteza*, a interação dura ocorre em uma pequena área transversal cuja dimensão é da ordem

$$\Delta x \cdot Q \sim \hbar = 1 \Rightarrow \Delta x \sim \frac{1}{Q} , \quad (5.9)$$

logo $S_I \sim \pi(\Delta x)^2 = \pi/Q^2$ é a área em que ocorre a interação dura.

A probabilidade de um outro párton de um dos hádrons ser encontrado nessa região e tomar parte do processo partônico duro é dada por

$$\mathcal{P} = \frac{\pi(\Delta x)^2}{\pi(R_0)^2} , \quad (5.10)$$

onde $R_0 =$ raio do nucleon $\sim 1 fm \sim \frac{1}{\Lambda}$ e Λ é o parâmetro de escala da QCD. Dessa forma, a probabilidade de que um processo multipartônico ocorra na pequena região de interação dura é da ordem

$$\mathcal{P} \sim \frac{\Lambda^2}{Q^2} . \quad (5.11)$$

A seguir, analisamos as contribuições provenientes de soft glúons.

5.2.3 Correlação entre as funções de onda dos hádrons: soft glúons

Como já havíamos comentado, as correções de QCD ao processo de *Drell-Yan* modificam tanto a seção de choque do subprocesso partônico ($\hat{\sigma}(q\bar{q} \rightarrow \gamma, Z \rightarrow e^+e^-)$) quanto as distribuições de pártons desses mesmos hádrons (PDFs).

Nesse último caso, as PDFs podem ser extraídas de experiências de espalhamento inelástico profundo, e usadas em outros processos hadrônicos, como o DY. Essa propriedade é chamada de *universalidade*.

Para que isso seja válido, é preciso que não haja correlação entre as PDFs dos hádrons antes que o subprocesso hadrônico ocorra, ou seja, os hádrons não podem trocar informações antes que seus pártons colidam.

Isso poderia muito bem acontecer já que cada um dos hádrons é acompanhado por uma nuvem de soft glúons (e eventualmente, pares $q\bar{q}$ virtuais), quer dizer, campos de gauge de longo alcance. Dessa forma, poderia acontecer que muito antes da colisão partônica, o campo gluônico de longo alcance de um hádron interagisse com os pártons do outro hádron transferindo informações entre as cargas de cor, correlacionando as PDFs de ambos os hádrons, alterando-as em relação àquelas do DIS, que não sofrem nenhuma interação dessa natureza por se tratar de uma interação lépton-hádron. Isso tudo torna-se ainda mais relevante quando levamos em conta que a forma das PDFs é governada pelo regime não-perturbativo da QCD.

Em termos um pouco mais quantitativos, considere um glúon soft de momento $|q^\mu| \ll Q$, de modo que a componente z do seu momento (supondo que o glúon viaje na direção $+z$) é da ordem $q^z \sim \Lambda(1 fm) \sim 1/R_0$, podemos então estimar que o campo “soft” do hádron se estende a uma distância R_0 à frente dele, que por sua vez, tem sua largura na direção z (ao longo do qual ele se move) contraída pelo fator γ de Lorentz, pictoricamente teríamos algo parecido com a figura (5.5).

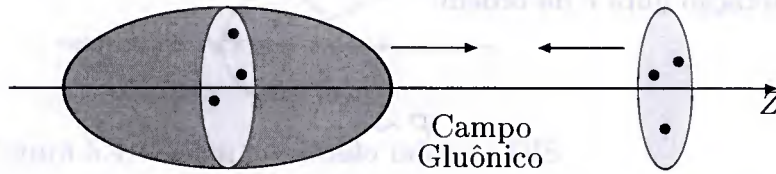


Figura 5.5: Campo clássico de um hádron devido a cargas de cor.

Dessa forma, interações entre o campo de longo alcance do hádron A começariam a interagir

$$\Delta t \sim \frac{R_0}{c} = \frac{1}{\Lambda c} \quad (5.12)$$

com os pártons do hádron B antes do processo duro partônico ocorrer, processo esse que ocorre em um intervalo de tempo muito menor, da ordem de

$$\frac{1}{Qc} \ll \Delta t, \quad (5.13)$$

permitindo amplo tempo para a correlação das PDFs, não permitindo mais que usemos as funções de distribuição de pártons obtidas do DIS nas fórmulas da seção de choque de Drell-Yan.

No limite clássico, no entanto, isso de fato não ocorre, por um motivo que nos é muito familiar e análogo ao eletromagnetismo clássico: as linhas de campo de uma partícula carregada movimentando-se relativisticamente concentram-se na região transversal ao movimento da partícula, quanto maior a sua velocidade maior é a concentração de linhas nessa região. Esse efeito é uma consequência direta da contração do espaço.

Efeito análogo ocorre às cargas de cor dos pártons, fazendo com que as linhas de campo se concentrem na direção perpendicular ao eixo z . O eixo de colisão estaria, então, “livre” dos campos de longo alcance contanto que os hádrons tivessem grandes momentos.

Dessa forma, os pártons de um hádron apenas sentiriam a presença dos pártons do outro hádron quando eles estivessem perto o bastante um dos outros para interagirem fortemente, sendo suas PDFs governadas por efeitos não-perturbativos dentro dos hádrons que os contêm e não influenciados pelos efeitos que ocorrem no outro. Dessa forma as PDFs permaneceriam não-correlacionadas até o momento da colisão dura de um par de pártons.

Na realidade, é possível mostrar, a partir de um modelo simples baseado no eletromagnetismo clássico [41, 42] que a força experimentada por um quark em

um dado hádron H decresce de acordo com m^2/s^2 , onde m é massa da partícula carregada e $s = \gamma m^2$ onde γ é o fator de Lorentz. Uma quebra da fatorização de ordem $1/s^2$ é, portanto, esperada em teoria de perturbação tendo esse fato já sido explicitamente provado [43].

5.2.4 Breve Análise dos Higher Twists Através da Expansão em Produtos de Operadores

O aniquilamento $e^+e^- \rightarrow$ hádrons e o espalhamento inelástico profundo são processos cujas seções de choque são dominadas por um produto de correntes eletromagnéticas próximas à região do cone de luz $x^2 \sim 0$ [11]. No caso do DIS, isso nos permite separar a seção de choque em uma forma conveniente usando funções de estrutura que parametrizam a estrutura interna do hádron, dado que não conhecemos o seu estado no momento da colisão.

O método de escrever um dado processo em termos de produtos de correntes analisando o seu comportamento no regime de altas energias transferidas, ou seja, para pequenas separações espaço-temporais $x^2 \sim 0$, na verdade, constitui uma poderosa ferramenta de investigação teórica em teoria quântica de campo e é conhecido como *Expansão de Produtos de Operadores* (OPE).

A OPE, proposta originalmente por Wilson em 1964 [44], nos dá uma definição plausível e coerente de um operador composto em termos de uma expansão em série de operadores locais bem-definidos e cujas singularidades do operador composto estão inseridos totalmente nas funções-coeficiente da expansão.

Vamos nos concentrar aqui no DIS, que é o exemplo não trivial mais simples em colisões hadrônicas [45]. Nossa meta é justificar, de forma mais rigorosa, o modelo a pártons como uma ótima aproximação no limite de altas energias, para um processo extremamente complicado onde vários efeitos diferentes ocorrem no interior do hádron. Entre esses efeitos estão aqueles em que um quark participante do processo elementar partônico interage com um quark espectador antes do processo duro (veja a figura (5.2)). Tais contribuições são denominadas *Higher Twists* pois provêm de operadores de dimensão superior da OPE do produto de correntes em termos da qual a seção de choque depende.

Graças ao *Teorema Óptico*, é possível relacionar a seção de choque do espalhamento inelástico profundo à amplitude do espalhamento Compton. Considere a amplitude para o espalhamento Compton de um hádron alvo como o representado na figura (5.6)

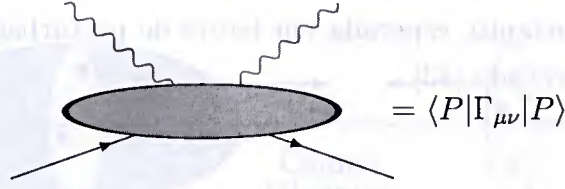


Figura 5.6: Amplitude do espalhamento Compton.

onde

$$\Gamma_{\mu\nu} = i \int d^4x e^{iq \cdot x} T(j_\mu(x) j_\nu(0)) , \quad (5.14)$$

com $j_\mu(x)$ sendo a corrente eletromagnética do quark, T representando o produto ordenado temporalmente e q o momento do fóton incidente.

O elemento de matriz de $\Gamma_{\mu\nu}$ entre estados de nucleons

$$S_{\mu\nu} = \langle P | \Gamma_{\mu\nu} | P \rangle \quad (5.15)$$

contraído com o tensor leptônico $L_{\mu\nu}$

$$L^{\mu\nu} = \sum_{spins} \langle k' | j_e^\nu(0) | k \rangle \langle k | j_e^\mu(0) | k' \rangle , \quad (5.16)$$

onde j_e^μ é corrente eletromagnética leptônica, nos dá a amplitude do espalhamento Compton. O tensor $S_{\mu\nu}$ deve ter as mesmas propriedades de $W_{\mu\nu}$ do DIS de modo que deve ser expandido em termos dos mesmos tensores de $W_{\mu\nu}$, que é dado por

$$W_{\mu\nu} = \frac{1}{2} \frac{1}{4\pi M} \int dX \langle P | J_\mu | X \rangle \langle X | J_\nu | P \rangle (2\pi)^4 \delta^4(P + q - X) , \quad (5.17)$$

onde M é a massa e J_μ a corrente eletromagnética do hádron respectivamente.

Por outro lado, a seção de choque do DIS é dada pela expressão correspondente ao diagrama abaixo

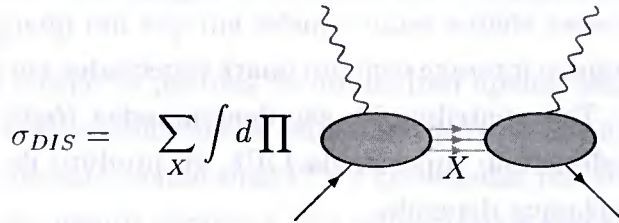


Figura 5.7: Seção de choque do espalhamento inelástico profundo.

O Teorema Óptico, por sua vez, implica a relação

$$2 \operatorname{Im} L_{\mu\nu} S^{\mu\nu} = \sigma_{DIS} , \quad (5.18)$$

e dessa forma é possível estudar o DIS com a amplitude do espalhamento Compton.

A idéia chave que permite o cálculo de $S_{\mu\nu}$ em QCD é a expansão em produto de operadores. Considere o produto de dois operadores locais

$$O_a(x)O_b(0) , \quad (5.19)$$

onde no limite em que $x \rightarrow 0$, os operadores estão praticamente no mesmo ponto. Neste limite, o produto pode ser escrito como uma expansão em operadores locais O_k

$$\lim_{x \rightarrow 0} O_a(x)O_b(0) = \sum_k C_{abk}(x) O_k(0) . \quad (5.20)$$

Note que as funções coeficiente dependem da separação x entre os dois operadores. No espaço de momentos, essa expressão é da forma

$$\lim_{q \rightarrow \infty} \int d^4x e^{iq \cdot x} O_a(x)O_b(0) = \sum_k \tilde{C}_{abk}(q) O_k(0) . \quad (5.21)$$

Dessa maneira, para $q^2 \gg \Lambda^2$, podemos expandir o produto de correntes $\Gamma_{\mu\nu}$ em um produto de operadores locais da QCD, invariantes de gauge por construção, criados a partir de espinores de quarks Ψ , tensores de campo gluônico $F_{\mu\nu}$ e derivadas covariantes D_μ .

Relacionando as expressões (5.14) e (5.15), obtemos

$$S_{\mu\nu} = \langle P | \Gamma_{\mu\nu} | P \rangle = i \int d^4x e^{iq \cdot x} \langle P | T(j_\mu(x)j_\nu(0)) | P \rangle . \quad (5.22)$$

Agora, utilizando a OPE, podemos expressar o elemento de matriz do produto de correntes na expressão acima em termos de operadores locais, sempre mantendo em mente o fato de que $q^2 \equiv Q^2 \gg \Lambda^2$,

$$S_{\mu\nu} = \sum_k \tilde{C}_{\mu\nu\mu_1 \dots \mu_n}^k(q) \langle P | O_{k[d,n]}^{\mu_1 \dots \mu_n}(0) | P \rangle , \quad (5.23)$$

onde d e n denotam a dimensão de massa e o spin total do operador respectivamente. Cada componente de spin contribui com um fator P^μ para o elemento de matriz, de modo que

$$\langle P | O_{k[d,n]}^{\mu_1 \dots \mu_n}(0) | P \rangle = M^{d-n-2} A P^{\mu_1} \dots P^{\mu_n} , \quad (5.24)$$

onde A é uma constante adimensional e M é a massa do próton, por exemplo. Note que o elemento de matriz não pode depender explicitamente do momento q^μ , pois sua dependência está isolada completamente nos coeficientes, de modo que apenas P^μ está disponível para carregar os índices do elemento de matriz do lado esquerdo de (5.24). O fato do elemento de matriz ser proporcional a M^{-2} é devido à dimensão do estado do próton ser -1 levando em conta a normalização relativística convencional. A dimensão do elemento de matriz é portanto M^{d-2} .

Como $S_{\mu\nu}$ é adimensional, pois é da forma de $W_{\mu\nu}$, concluímos que $\Gamma_{\mu\nu}$ tem dimensão de massa 2, logo

$$[\Gamma_{\mu\nu}] = 2 = [C_k(q)][O^k(0)] , \quad (5.25)$$

de modo que

$$[C_k(q)] = 2 - d . \quad (5.26)$$

Agora, como os coeficientes só dependem dos momentos q^μ e têm dimensão $2 - d$, então, podemos escrevê-los da seguinte forma

$$\tilde{C}_k^{\mu\nu\mu_1\cdots\mu_n} = a_k q^\mu q^\nu Q^{2-d} \frac{q^{\mu_1}}{Q} \cdots \frac{q^{\mu_n}}{Q} . \quad (5.27)$$

A seção de choque do DIS é dada por (5.18) usando (5.16), (5.23) e (5.24), logo

$$\sigma_{DIS} = 2Im L_{\mu\nu} S^{\mu\nu} = \sum_k \tilde{C}_{k\mu_1\cdots\mu_n}^{\mu\nu} \langle P | O_{k[d,n]}^{\mu_1\cdots\mu_n}(0) | P \rangle L_{\mu\nu} , \quad (5.28)$$

portanto

$$\begin{aligned} \sigma_{DIS} &= \sum_k \hat{a}_k \frac{q^{\mu_1}}{Q} \cdots \frac{q^{\mu_n}}{Q} Q^{2-d} M^{d-n-2} P_{\mu_1} \cdots P_{\mu_n} \\ &= \sum_k \hat{a}_k \frac{(P \cdot q)^n}{Q^n} Q^{2-d} M^{d-n-2} \\ &= \sum_k \hat{a}_k \frac{(P \cdot q)^n}{Q^{2n}} Q^{2-d+n} M^{d-n-2} \\ &= \sum_k \hat{a}_k \omega^n (Q/M)^{2-t} , \end{aligned} \quad (5.29)$$

onde definimos $\hat{a}_k = a_k q^\mu q^\nu L_{\mu\nu}$. O *twist* t é definido como

$$\text{twist} = t \equiv d - n = \text{dimensão} - \text{spin} . \quad (5.30)$$

Os operadores mais relevantes na expansão em produto de operadores, e portanto para a seção de choque são aqueles com o menor *twist* possível. Operadores de *twist* 2 são os que darão a contribuição líder, os outros entrarão com contribuições que são suprimidas por potências de

$$(M/Q)^k$$

onde $k = 1, 2, 3, \dots$.

Os campos fundamentais da QCD são os campos dos quarks Ψ e glúons G_μ , de modo que os operadores que formam a base da expansão em produtos de operadores podem ser escritos em termos de Ψ , o tensor de campo gluônico $G_{\mu\nu}$ e a derivada covariante D_μ . A seguir apresentamos uma tabela listando os objetos básicos, com sua dimensão, spin e twist

	Ψ	$G_{\mu\nu}$	D_μ
d	3/2	2	1
s	1/2	1	1
t	1	1	0

Tabela 5.1: Dimensão, spin e twist de operadores básicos da QCD.

Qualquer operador invariante de gauge deve conter, no mínimo, dois campos de quarks $[\Psi\bar{\Psi}]$ ou dois tensores de campo do glúon contraídos $[F_{\mu\nu}F^{\mu\nu}]$, e, pela tabela, vemos que esses operadores possuem twist 2.

É possível mostrar através de uma análise mais detalhada do DIS em termos da OPE [13], que os operadores de twist igual a 2 dados por

$$\mathcal{O}_i^{(n)\mu_1\cdots\mu_n} = \bar{\Psi}_i \gamma^{\{\mu_1} (iD^{\mu_2}) \cdots (iD^{\mu_n}) \Psi_i, \quad (5.31)$$

onde o símbolo $\{\mu_1 \cdots \mu_n\}$ denota simetrização dos índices e D^μ é derivada covariante da QCD (que possui $t = 0$, pela tabela (5.1)), dão as contribuições dominantes ao tensor $W_{\mu\nu}$, quando contraídos com potências de q . O operador (5.31), por sua vez, corresponde à corrente eletromagnética de um quark de sabor i quando retiramos todas as derivadas covariantes, que são responsáveis pela interação entre os quarks que compõem a corrente e glúons. A partir de (5.31), pode-se mostrar que o tensor $W^{\mu\nu}$ tem exatamente a mesma forma predita pelo modelo a pártons, sendo possível, inclusive, decompô-lo em termos de operadores cujos coeficientes, os fatores de forma, obedecem à relação de Callan-Gross. Todos os operadores com twist maior do que 2 são suprimidos por potências de $(M/Q)^k$, onde $k = 1, 2, 3, \dots$

5.3 As Correções da QCD Perturbativa para o DIS

Nosso principal objetivo nesta seção é entender como definir as funções de distribuição de pártons experimentalmente. Para isso, o espalhamento inelástico profundo é o ambiente ideal para tal definição. Antes de mais nada vamos definir a notação e vários conceitos que serão usados nos próximos passos do cálculo.

O *espalhamento inelástico profundo* (DIS) é um processo onde um lépton incidente de momento k^μ emite um fóton virtual altamente energético de momento

q^μ e retrocede com momento k'^μ , o qual é medido na experiência. O fóton atinge um próton de momento P^μ , dissociando-o em um “chuveiro” de hádrons X . Tal processo está ilustrado na figura (5.8).

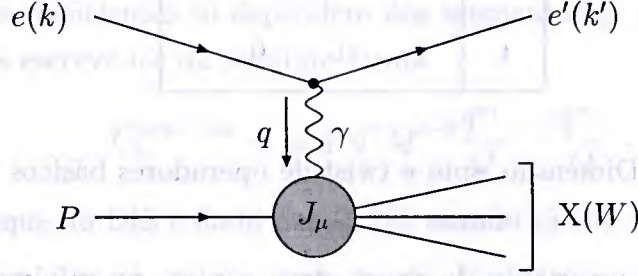


Figura 5.8: Diagrama de Feynman para o espalhamento inelástico profundo.

As variáveis cinemáticas do DIS podem ser obtidas a partir dos 4-momentos indicados na figura acima. Em primeiro lugar, temos a equação de conservação de 4-momento

$$k + P = k' + W, \quad (5.32)$$

os invariantes relevantes são dados por

$$S = (k + P)^2, \quad (5.33)$$

$$Q^2 \equiv -q^2 = 2k \cdot k', \quad (5.34)$$

$$\nu = \frac{q \cdot P}{M}, \quad (5.35)$$

onde M é a massa do próton. As massas dos léptons foram desprezadas frente as energias envolvidas no processo e

$$W^2 = (P + k - k')^2 = 2M\nu + M^2 - Q^2 \quad (5.36)$$

é o quadrado da massa invariante do estado hadrônico X produzido após a dissociação do próton. No referencial de laboratório, onde o próton está parado, os 4-momentos das partículas são dadas explicitamente por

$$\begin{aligned} k^\mu &= (E, 0, 0, E) \\ k'^\mu &= (E', E' \sin \theta, 0, E' \cos \theta) \\ P^\mu &= (M, 0, 0, 0) \end{aligned} \quad (5.37)$$

e em termos dos quais temos

$$Q^2 = 2k \cdot k' = 2EE'(1 - \cos \theta) = 4EE' \sin^2(\theta/2) , \quad (5.38)$$

onde θ é o ângulo de espalhamento do lépton no referencial de repouso do próton, e

$$\nu = \frac{P \cdot q}{M} = \frac{M(E - E')}{M} = E - E' . \quad (5.39)$$

Note que, para o espalhamento elástico $e^-P \rightarrow e^-P$, teríamos $W^2 = M^2$, o que corresponderia a $Q^2 = 2M\nu$.

Nosso ponto de partida é a amplitude do processo ilustrado na figura (5.8), dada por

$$i\mathcal{M} = (ie)(-ie)\frac{i}{Q^2}\bar{u}(k')\gamma^\mu u(k)\langle P|J_\mu|X\rangle \quad (5.40)$$

A quantidade $\langle P|J_\mu|X\rangle$ é denominada corrente hadrônica e descreve a dissociação do próton e, ao contrário da corrente leptônica, não pode ser calculada perturbativamente por conter informações a respeito da função de onda do próton, que é determinada por efeitos complicados de interações fortes de longa distância.

A única coisa que sabemos, a priori, é que essa corrente deve ser conservada por interações eletromagnéticas, ou seja

$$q^\mu \langle P|J_\mu|X\rangle = 0 . \quad (5.41)$$

O quadrado da amplitude, após soma sobre os spins finais e média sobre os spins iniciais é dada por

$$|\overline{\mathcal{M}}|^2 = \frac{e^4}{Q^4} \frac{1}{4} \text{Tr}[k'\gamma^\mu k\gamma^\nu] \langle P|J_\mu|X\rangle \langle X|J_\nu^\dagger|P\rangle . \quad (5.42)$$

A seção de choque diferencial inclusiva é calculada a partir da expressão

$$d\sigma = \frac{1}{2S} \frac{e^4}{Q^4} (4\pi M) L^{\mu\nu} W_{\mu\nu} \frac{d^3k'}{(2\pi)^3 2E'} , \quad (5.43)$$

onde os tensores leptônico e hadrônico são definidos por

$$L^{\mu\nu} = \frac{1}{2} \text{Tr}[k'\gamma^\mu k\gamma^\nu] \quad (5.44)$$

$$W_{\mu\nu} = \frac{1}{2} \frac{1}{4\pi M} \int dX \langle P|J_\mu|X\rangle \langle X|J_\nu^\dagger|P\rangle (2\pi)^4 \delta^4(P + q - X) . \quad (5.45)$$

Aqui dX representa o espaço de fase dos hádrons do estado final decorrentes da dissociação do próton.

O tensor $W_{\mu\nu}$ carrega, agora, todas as informações a respeito da interação γ - próton, mas não é possível calculá-lo perturbativamente. As únicas informações

que temos sobre $W_{\mu\nu}$ vêm de argumentos de simetria: primeiro, $W_{\mu\nu}$ é invariante de Lorentz por construção, logo deve ser formado por uma combinação linear dos tensores: $P_\mu P_\nu$, $P_\mu q_\nu$, $P_\nu q_\mu$, $q_\mu q_\nu$ e $g_{\mu\nu}$

$$\begin{aligned} W_{\mu\nu} = & V_1 g_{\mu\nu} + V_2 P_\mu P_\nu + V_3 (P_\mu q_\nu + P_\nu q_\mu) \\ & + V_4 (P_\mu q_\nu - P_\nu q_\mu) + V_5 q_\mu q_\nu + V_6 \varepsilon_{\mu\nu\alpha\beta} P^\alpha q^\beta . \end{aligned} \quad (5.46)$$

Como $L^{\mu\nu}$ é um tensor simétrico nos índices μ e ν , os termos antisimétricos de $W_{\mu\nu}$, proporcionais a V_4 e V_6 não contribuem para a seção de choque total. Se por outro lado, a corrente hadrônica se acopla a uma corrente leptônica de estrutura $V - A$, como no caso do espalhamento próton-neutrino, o termo em V_6 não pode ser desprezado.

A conservação da corrente hadrônica, por outro lado, garante que

$$q^\mu W_{\mu\nu} = q^\nu W_{\mu\nu} = 0 ,$$

o que implica

$$V_1 q_\nu + V_2 (q \cdot P) P_\nu + V_3 (q \cdot P q_\nu + q^2 P_\nu) + V_5 q^2 q_\nu = 0 , \quad (5.47)$$

de forma que os coeficientes dos 4-vetores P_ν e q_ν devem se anular separadamente, o que nos leva ao seguinte sistema de equações

$$\begin{aligned} V_1 + (q \cdot P) V_3 + q^2 V_5 &= 0 \\ (q \cdot P) V_2 + q^2 V_3 &= 0 , \end{aligned} \quad (5.48)$$

a partir do qual obtemos V_3 e V_5 em função de V_1 e V_2 . O tensor hadrônico pode, então, ser expresso em termos dessas duas variáveis, e que, por convenção, redefinimos como $V_1 = W_1$ e $V_2 = W_2/M^2$. Os objetos adimensionais W_1 e W_2 são denominados de *funções de estrutura* para o espalhamento inelástico elétron-próton, e parametrizam todas as informações a respeito da estrutura dos hádrons na colisão. A princípio, elas dependem de P e q , ou em linguagem invariante de Lorentz, de ν e Q^2 , logo $W_{\mu\nu}$ pode ser escrito como

$$W_{\mu\nu} = - \left(g_{\mu\nu} - \frac{q_\mu q_\nu}{q^2} \right) W_1(\nu, Q^2) + \frac{1}{M^2} \left(P_\mu - q_\mu \frac{P \cdot q}{q^2} \right) \left(P_\nu - q_\nu \frac{P \cdot q}{q^2} \right) W_2(\nu, Q^2) . \quad (5.49)$$

A seção de choque diferencial do espalhamento inelástico profundo, dada em (5.43), em termos das quantidades medidas em laboratório (as do elétron espalhado com momento k') é da forma

$$\frac{d^2\sigma}{dE' d\Omega'} = \frac{4\alpha E'^2}{Q^4} \left(2W_1(\nu, Q^2) \sin^2 \frac{\theta}{2} + W_2(\nu, Q^2) \cos^2 \frac{\theta}{2} \right) \quad (5.50)$$

após o cálculo da contração dos tensores leprônico e hadrônico, e onde $d\Omega' = d(\cos \theta)d\phi$.

No modelo a pártons assume-se a existência de uma função de distribuição de pártons $\phi_i(y)$, tal que $\phi_i(y)dy$ representa a probabilidade de se encontrar um párton do tipo i = (quark, antiquark ou glúon) no interior de um hádron levando uma fração de momento compreendida entre y e $y + dy$. Se a seção de choque de um dado evento partônico é dada por $d\hat{\sigma}$, então, a seção de choque diferencial hadrônica será

$$d\sigma = \sum_i \int_{y_{\min}}^{y_{\max}} dy \phi_i(y) d\hat{\sigma}_i(y) . \quad (5.51)$$

A variável y na verdade varia de $x = Q^2/2P \cdot q$, a variável de Bjorken, até 1, pois

$$(p + q)^2 = 2yP \cdot q - Q^2 \geq 0 \\ \implies y \geq x$$

onde $p = yP$.

Analogamente a (5.43), podemos escrever $d\hat{\sigma}_i$ como

$$d\hat{\sigma}_i = \frac{1}{2s} \frac{e^4}{Q^4} (4\pi M) L_{\mu\nu} \hat{W}_i^{\mu\nu} \frac{d^3 k'}{(2\pi)^3 2E'} , \quad (5.52)$$

onde $s = (p + q)^2 = 2p \cdot q = yS$.

O tensor $\hat{W}_i^{\mu\nu}$, por sua vez é o análogo de $W_{\mu\nu}$ para o caso do espalhamento de um párton do tipo i , cujo estado quântico corresponde a $|yP\rangle$, através da corrente eletromagnética J_μ^i

$$\hat{W}_i^{\mu\nu} = \frac{1}{2} \int \langle yP | J_\mu^i | X' \rangle \langle X' | J_\nu^i | yP \rangle (2\pi)^4 \delta^4(p + q - X') dX' , \quad (5.53)$$

enquanto isso, $|X'\rangle$ corresponde ao quark do estado final. Em outras palavras

$$\langle p | J_\mu^i | X' \rangle = -ie_i \bar{u}(p') \gamma_\mu u(p) \quad (5.54)$$

nos dá exatamente a amplitude do processo $\gamma q \rightarrow q'$ ao nível de árvore. A integral sobre dX' indica que estamos integrando sobre o espaço de fase do quark espalhado q' e somando sobre spins e cores finais. O fator 1/2 indica a média sobre os spins iniciais do quark q e e_i é carga elétrica do quark de sabor i . Logo, substituindo (5.54) em (5.53), obtemos

$$\hat{W}_{\mu\nu}^i = \frac{1}{2} e_i^2 \int \text{Tr}[\not{p}' \gamma_\mu \not{p} \gamma_\nu] \frac{d^3 p'}{(2\pi)^3 2E'} (2\pi)^4 \delta^4(p + q - p') , \quad (5.55)$$

onde o fator de espaço de fase de 1 partícula pode ser escrito como

$$\begin{aligned} \frac{d^3 p'}{(2\pi)^3 2E'} \delta^4(p + q - p') &= d^4 p' \delta(p'^2) \delta^4(p + q - p') \\ &= \delta(2p \cdot q - Q^2) . \end{aligned} \quad (5.56)$$

Calculando o traço, (5.55) é dada agora por

$$\hat{W}_{\mu\nu}^i = 2e_i^2 [2p_\mu p_\nu + q_\mu p_\nu + q_\nu p_\mu - p \cdot q g_{\mu\nu}] 2\pi \delta(2p \cdot q - Q^2) , \quad (5.57)$$

e lembrando que $p = yP$ e que $x = Q^2/2P \cdot q$, então

$$\hat{W}_{\mu\nu}^i = 2e_i^2 [2y^2 P_\mu P_\nu + yq_\mu P_\nu + yq_\nu P_\mu - yP \cdot q g_{\mu\nu}] 2\pi \frac{\delta(y - x)}{2P \cdot q} . \quad (5.58)$$

Para relacionar $W^{\mu\nu}$ a $\hat{W}_i^{\mu\nu}$ basta substituir (5.52) em (5.51), lembrando que $s = yS$, de modo que, comparando com (5.43) obtém-se

$$W^{\mu\nu} = \frac{1}{4\pi M} \sum_i \int_x^1 \frac{dy}{y} \phi_i(y) \hat{W}_i^{\mu\nu} , \quad (5.59)$$

contudo, por invariância de Lorentz, $W_{\mu\nu}$ deve ter a forma dada em (5.49), consequentemente, $\hat{W}_{\mu\nu}^i$ também deve ser daquela forma, dada sua relação com $W_{\mu\nu}$ através de (5.59). Portanto, para isolar W_1 e W_2 , precisamos apenas calcular os coeficientes de $P_\mu P_\nu$ e $g_{\mu\nu}$ quando inserimos (5.58) em (5.59), ou seja

$$W_{\mu\nu} = \frac{1}{4\pi M} \frac{2\pi}{P \cdot q} \sum_i e_i^2 \int_x^1 \frac{dy}{y} \phi_i(y) [2y^2 P_\mu P_\nu + \dots - yP \cdot q g_{\mu\nu}] \delta(x - y) , \quad (5.60)$$

de modo que quando usamos que $\nu = P \cdot q/M$, obtemos

$$W_{\mu\nu} = \sum_i \frac{e_i^2}{\nu M^2} x \phi_i(x) P_\mu P_\nu + \dots - \sum_i \frac{e_i^2}{2M} \phi_i(x) g_{\mu\nu} . \quad (5.61)$$

Comparando com a nossa expressão geral de $W_{\mu\nu}$ dada em (5.49), vemos que

$$W_1 = \sum_i \frac{e_i^2}{2M} \phi_i(x) \quad (5.62)$$

$$W_2 = \sum_i \frac{e_i^2}{\nu} x \phi_i(x) . \quad (5.63)$$

Os fatores de forma F_1 e F_2 são definidos a partir de W_1 e W_2 por

$$F_1(x, Q^2) = MW_1(\nu, Q^2) = \sum_i e_i^2 \phi_i(x) , \quad (5.64)$$

$$F_2(x, Q^2) = \nu W_2(\nu, Q^2) = \sum_i e_i^2 x \phi_i(x) . \quad (5.65)$$

A seguir vamos calcular as correções de ordem α_s às funções de distribuição de quarks usando o fator de forma $F_2(x, Q^2)$ obtido no modelo a pártons para definir as PDFs em qualquer ordem de teoria de perturbação, ou seja, nosso objetivo é calcular a correção de ordem α_s ao fator de forma F_2 dado por

$$F_2(x, Q^2) = \sum_q Q_q^2 x [q(x, Q^2) + \bar{q}(x, Q^2)] , \quad (5.66)$$

onde estão computadas as contribuições de quarks e antiquarks*. Para extrair F_2 do tensor $W_{\mu\nu}$, é conveniente definir as funções de estrutura transversal e longitudinal em $d = 4 - 2\varepsilon$ dimensões

$$W_T \equiv -g^{\mu\nu}W_{\mu\nu} = (3 - 2\varepsilon)W_1 - \left(1 + \frac{\nu^2}{Q^2}\right)W_2, \quad (5.67)$$

$$W_L \equiv P^\mu P^\nu W_{\mu\nu} = -M^2 \left(1 + \frac{\nu^2}{Q^2}\right)W_1 + M^2 \left(1 + \frac{\nu^2}{Q^2}\right)^2 W_2. \quad (5.68)$$

No regime de espalhamento inelástico profundo, onde $\nu^2/Q^2 \gg 1$, podemos realizar as seguintes aproximações para W_T e W_L

$$W_T = (3 - 2\varepsilon)W_1 - \frac{\nu^2}{Q^2}W_2 \quad (5.69)$$

$$W_L = -\frac{M^2\nu^2}{Q^2}W_1 + \frac{M^2\nu^4}{Q^4}W_2 \quad (5.70)$$

Usando a definição (5.65), podemos escrever W_2 em termos de W_T e W_L , obtendo

$$(1 - \varepsilon)\frac{F_2}{M} = xW_T + \frac{4x^3}{Q^2}(3 - 2\varepsilon)W_L. \quad (5.71)$$

A partir de (5.59) W_T e W_L são dadas por

$$W_T = -g_{\mu\nu}W^{\mu\nu} = \frac{1}{4\pi M} \sum_i \int_x^1 \frac{dy}{y} \phi_i(y) \hat{W}_T^i \quad (5.72)$$

$$W_L = P_\mu P_\nu W^{\mu\nu} = \frac{1}{4\pi M} \sum_i \int_x^1 \frac{dy}{y} \phi_i(y) \hat{W}_L^i \quad (5.73)$$

onde $\hat{W}_T^i = -g_{\mu\nu}\hat{W}_i^{\mu\nu}$ e $\hat{W}_L^i = P_\mu P_\nu \hat{W}_i^{\mu\nu}$. A função de estrutura W_T^i agora representa o quadrado da amplitude do espalhamento $\gamma q \rightarrow q'$ cujos índices do fóton são contraídos pelos tensor $-g_{\mu\nu}$ e integrado no espaço de fase do quark do estado final, enquanto que W_L^i tem a mesma estrutura, mas com os índices do fóton contraídos pelo tensor $P_\mu P_\nu$.

O próximo passo consiste em determinar as correções de ordem α_s à W_T e W_L , dado que W_T^i e W_L^i podem ser calculados perturbativamente, e, dessa maneira, usando a equação (5.71), a correção para F_2 . Com a definição (5.66) das PDFs em mãos, obtemos $q(x, Q^2)$ e $\bar{q}(x, Q^2)$. Seguimos de perto, aqui, o roteiro da ref. [46], onde é calculada a correção de ordem α_s a $q(x)$. O caso da distribuição de antiquarks é absolutamente análogo.

*Em ordem zero $q(x, Q^2) = q_0(x)$. A partir da agora denotamos a carga do quark q por Q_q e a função de distribuição de quarks por $q(x)$.

5.3.1 Glúons Virtuais

Assim como a seção de choque elementar, as PDFs também contém uma parcela devida a glúons em loops. Ao nível de 1-loop, essa contribuição é calculada com base nos seguintes diagramas de Feynman

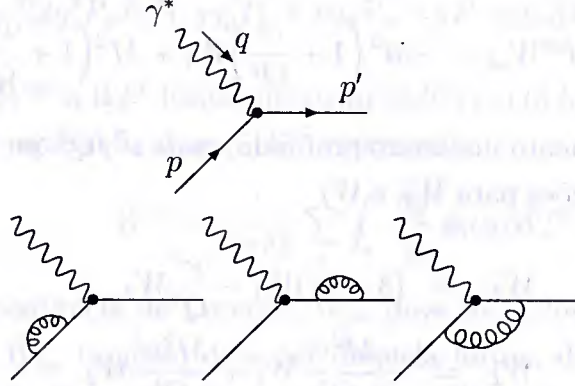


Figura 5.9: Deep Inelastic Scattering ao nível de 1-loop em QCD.

Exatamente como no caso $e^+e^- \rightarrow q\bar{q}$, adotando o gauge de Landau, apenas o vértice contribui. Ao nível de 1-loop o vértice é da forma

$$\begin{aligned} i\Gamma^\mu &= -iQ_q\gamma^\mu \left\{ 1 + \frac{\alpha_s}{4\pi} C_F \left(\frac{4\pi\mu^2}{Q^2} \right)^\epsilon \frac{\Gamma(1+\epsilon)\Gamma^2(1-\epsilon)}{\Gamma(1-2\epsilon)} \left(-\frac{2}{\epsilon^2} - \frac{3}{\epsilon} - 8 \right) \right\} \\ &\equiv -iQ_q\gamma^\mu (1 + \alpha_s \langle V \rangle) . \end{aligned} \quad (5.74)$$

Logo, a amplitude corresponde a

$$\mathcal{M}_{virt} = \bar{u}(p') i\Gamma^\mu u(p) \epsilon_\mu(q) \quad (5.75)$$

e o seu quadrado, após soma sobre os spins e cores do quark final, média sobre os spins e cores do estado inicial e média sobre as polarizações do fóton é

$$|\overline{\mathcal{M}}|_{virt}^2 = 2Q_q^2(1-\epsilon)Q^2(1+2\alpha_s \langle V \rangle) , \quad (5.76)$$

onde mantivemos termos até ordem α_s apenas.

A contribuição virtual à F_2 é inserida por meio da expressão

$$\frac{F_2}{x} = \sum_q Q_q^2 \int_x^1 \frac{dy}{y} q_0(y) \delta(1-x/y) (1 + 2\alpha_s \langle V \rangle) , \quad (5.77)$$

onde vamos definir $z = x/y$. Finalmente obtemos, então

$$\begin{aligned} \frac{F_2}{x} \Big|_{virt} &= \sum_q Q_q^2 \int_x^1 \frac{dy}{y} q_0(y) \delta(1-z) \times \\ &\quad \left\{ 1 + \frac{\alpha_s}{2\pi} C_F \left(\frac{4\pi\mu^2}{Q^2} \right)^\epsilon \frac{\Gamma(1-\epsilon)}{\Gamma(1-2\epsilon)} \left(-\frac{2}{\epsilon^2} - \frac{3}{\epsilon} - 8 - \frac{\pi^2}{3} \right) \right\} . \end{aligned} \quad (5.78)$$

O pólo simples $1/\varepsilon$ corresponde a uma divergência colinear e o pólo duplo $1/\varepsilon^2$, a uma superposição composta de divergências IR.

5.3.2 Glúon no Estado Inicial

Calcularemos, agora, a contribuição do subprocesso $\gamma g \rightarrow q\bar{q}$ a W_T e W_L . A correção de ordem α_s devida a glúons iniciais no subprocesso $\gamma g \rightarrow q\bar{q}$ às funções de estrutura W_T e W_L são obtidas pela contração de $-g_{\mu\nu}$ e $P_\mu P_\nu$, respectivamente, aos índices do fóton na amplitude ao quadrado do processo $\gamma g \rightarrow q\bar{q}$ e integrada no espaço de fase do par $q\bar{q}$.

Os diagramas de Feynman são os da figura (5.10)

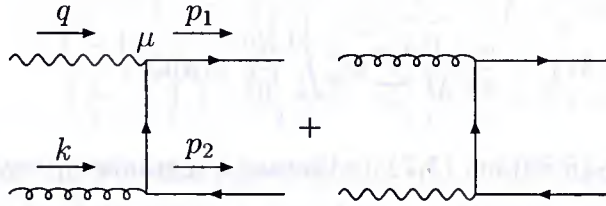


Figura 5.10: Diagramas de Feynman para o processo $\gamma g \rightarrow q\bar{q}$.

A amplitude do processo, com os índices do fóton livres, ou seja sem os vetores de polarização do fóton, é dado pela expressão

$$\mathcal{M}_i^\mu = -iQ_q g_s T^a \bar{u}(p_1) \left[\gamma^\mu \frac{(\not{p}_1 - \not{q})}{(p_1 - q)^2} \not{\epsilon}(k) + \not{\epsilon}(k) \frac{(\not{p}_1 - \not{k})}{(p_1 - k)^2} \gamma^\mu \right] v(p_2) . \quad (5.79)$$

Os invariantes de Mandelstam são alterados pela virtualidade do fóton $q^2 \equiv -Q^2$

$$s = 2p_1 \cdot p_2 = -Q^2 + 2q \cdot k , \quad (5.80)$$

$$t = -2k \cdot p_2 = -Q^2 - 2p_1 \cdot q , \quad (5.81)$$

$$u = -2k \cdot p_1 = -Q^2 - 2p_2 \cdot q , \quad (5.82)$$

$$s + t + u = -Q^2 . \quad (5.83)$$

Estes podem ser escritos em termos de $v = \frac{1}{2}(1 + \cos \theta)$, onde θ é ângulo formado pelos vetores $\vec{p}_1$ e $\vec{q}$ e da varável z definida por $z = x/y$, onde x é variável de Bjorken e $k = yP$. Dessa forma, temos o seguinte

$$s = \frac{Q^2}{z}(1 - z) , \quad (5.84)$$

$$t = -\frac{Q^2}{z}(1 - v) , \quad (5.85)$$

$$u = -\frac{Q^2}{z}v . \quad (5.86)$$

Após a soma sobre os números quânticos do estado final e média sobre os iniciais[†], integrando $-g_{\mu\nu}(M_i^\mu)^\dagger M_i^\nu$ no espaço de fase de 2 partículas dado por

$$\Phi_2^d = \frac{1}{8\pi} \left(\frac{4\pi}{s} \right)^\varepsilon \frac{1}{\Gamma(1-\varepsilon)} \int_0^1 dv v^{-\varepsilon} (1-v)^{-\varepsilon}, \quad (5.87)$$

obtemos o seguinte resultado em d dimensões

$$W_T = \frac{\alpha_s}{2\pi} \frac{1}{M} \sum_q Q_q^2 \int_x^1 \frac{dy}{y} g(y) [z^2 + (1-z)^2] \left[-\frac{1}{\varepsilon} \frac{\Gamma(1-\varepsilon)}{\Gamma(1-2\varepsilon)} + \ln \left(\frac{Q^2}{4\pi} \frac{1-z}{z} \right) \right], \quad (5.88)$$

onde já identificamos $\phi(y) = g(y)$ como a função de distribuição de glúons no interior do próton. O pólo $1/\varepsilon$ origina-se de uma divergência colinear.

O cálculo de W_L não apresenta nenhuma complicação adicional, sendo, aliás, finito no limite $\varepsilon \rightarrow 0$

$$W_L = \frac{\alpha_s}{4\pi} \frac{1}{M} \sum_q Q_q^2 \int_x^1 \frac{dy}{y^3} g(y) Q^2 \frac{1-z}{z}. \quad (5.89)$$

Substituindo (5.88) e (5.89) em (5.71), obtemos a seguinte correção ao fator de forma F_2

$$\left. \frac{F_2}{x} \right|_g = \frac{\alpha_s}{2\pi} \frac{1}{1-\varepsilon} \sum_q Q_q^2 \int_x^1 dy g(y) z \times \left\{ [z^2 + (1-z)^2] \left(-\frac{1}{\varepsilon} \frac{\Gamma(1-\varepsilon)}{\Gamma(1-2\varepsilon)} + \ln \frac{Q^2}{4\pi\mu^2} \right) + 6z(1-z) \right\}. \quad (5.90)$$

5.3.3 Emissão de Glúon Real

Vamos agora considerar o processo de emissão de glúon $\gamma q(\bar{q}) \rightarrow q(\bar{q})g$. Os diagramas de Feynman são

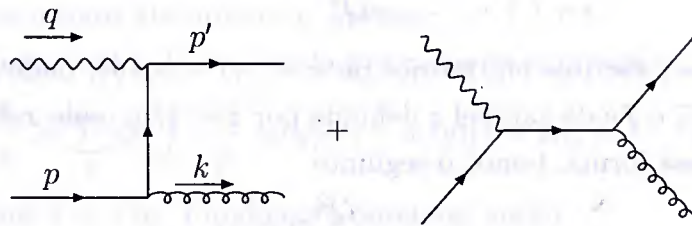


Figura 5.11: Diagramas de Feynman para emissão de glúon.

[†]Em d dimensões, se temos 1 dimensão espacial longitudinal, então restam $d-2$ dimensões espaciais transversais que corresponde ao número de estados de polarização do glúon. Contudo isso é uma convenção, dado que poderíamos ter escolhido 2 dimensões transversais e $d-1$ longitudinais.

A amplitude desse processo pode ser obtida a partir da (5.79) pelas substituições $q \rightarrow q$, $p_1 \rightarrow p'$, $p_2 \rightarrow -p$ e $k \rightarrow -k$ o que implica $s \leftrightarrow u$ em relação às variáveis de Mandelstam. Além disso, como o glúon agora está no estado final e o quark no estado inicial é preciso incluir o seguinte fator adicional

$$8 \times 2(1 - \varepsilon) \times \frac{1}{2} \times \frac{1}{3} ,$$

onde o fator 8 compensa a média sobre as cores do glúon inicial no processo $\gamma g \rightarrow q\bar{q}$ e $2(1 - \varepsilon)$ compensa a média sobre os spins do glúon. Finalmente, os fatores $1/2$ e $1/3$ indicam a média sobre cores e spins do quark no processo em questão.

No mais, excetuando a álgebra, todo o procedimento de cálculo é análogo ao caso dos glúons iniciais. Vamos, portanto, nos restringir a apresentar o resultado para o fator de forma

$$\begin{aligned} \frac{F_2}{x} \Big|_q &= \frac{\alpha_s}{2\pi} C_F \left(\frac{4\pi\mu^2}{Q^2} \right) \frac{\Gamma(1 - \varepsilon)}{\Gamma(1 - 2\varepsilon)} \sum_q Q_q^2 \int_x^1 \frac{dy}{y} q_0(y) \\ &\quad \left(\frac{z}{1 - z} \right)^\varepsilon \left\{ -\frac{1}{\varepsilon} \frac{1 + z^2}{1 - z} - \frac{3}{2} \frac{1}{1 - z} + 3 + 2z - \varepsilon \frac{1}{1 - z} \right\} . \end{aligned} \quad (5.91)$$

A experiência nos diz que essa expressão deveria conter uma divergência quadrática em ε , a qual, no final das contas, cancelaria o pólo correspondente da contribuição virtual. Há, no entanto, uma divergência associada ao limite $z \rightarrow 1$, que é o limite onde o glúon emitido se torna soft. Precisamos explicitar essa divergência em termos de um pólo em ε .

Isso é realizado introduzindo as chamadas funções “+” seguindo o roteiro original da ref. [47]. Tais funções são, estritamente falando, distribuições bem comportadas quando convoluídas com funções suaves que vão a zero suficientemente rápido quando $x \rightarrow 1$. A definição mais usada é

$$\int_0^1 dx g(x) [F(x)]_+ = \int_0^1 dx [g(x) - g(1)] F(x) , \quad (5.92)$$

onde $[F(x)]_+ = F(x)$ no intervalo $0 \leq x < 1$.

Se $g(x) = 1$, segue a importante propriedade das funções “+”

$$\int_0^1 dx [F(x)]_+ = 0 . \quad (5.93)$$

Nosso interesse é explicitar a divergência soft contida na função

$$F(z) \equiv \left(\frac{z}{1 - z} \right)^\varepsilon \frac{1}{1 - z} , \quad (5.94)$$

que no limite $\varepsilon \rightarrow 0$ admite a expansão

$$\left(\frac{z}{1 - z} \right)^\varepsilon \frac{1}{1 - z} \approx \frac{1}{1 - z} + \varepsilon \frac{\ln(z)}{1 - z} - \varepsilon \frac{\ln(1 - z)}{1 - z} .$$

Usando a prescrição positiva, por outro lado, temos

$$\int_0^1 dz \left[\left(\frac{z}{1-z} \right)^\varepsilon \frac{1}{1-z} \right]_+ \approx \int_0^1 dz \left[\frac{1}{(1-z)_+} + \varepsilon \frac{\ln(z)}{(1-z)_+} - \varepsilon \left(\frac{\ln(1-z)}{1-z} \right)_+ \right] \quad (5.95)$$

A partir de (5.92) com $g(z) = \ln z$, podemos calcular o segundo termo do lado esquerdo de (5.95)

$$\begin{aligned} \int_0^1 dz \ln(z) \frac{1}{(1-z)_+} &= \int_0^1 dz \frac{\ln(z) - \ln(1)}{1-z} \\ &= \int_0^1 dz \frac{\ln(z)}{1-z} = -Li_2(1) = -\frac{\pi^2}{6} \end{aligned} \quad (5.96)$$

A expressão (5.93), então nos permite escrever

$$\int_0^1 dz \left[\left(\frac{z}{1-z} \right)^\varepsilon \frac{1}{1-z} \right]_+ \approx - \int_0^1 dz \varepsilon \frac{\pi^2}{6} \delta(1-z) \quad (5.97)$$

até ordem ε .

A integral em $F(z)$, por outro lado, pode ser calculada exatamente

$$\begin{aligned} \int_0^1 dz z^\varepsilon (1-z)^{-1-\varepsilon} &= \frac{\Gamma(1+\varepsilon)\Gamma(-\varepsilon)}{\Gamma(1)} \\ &\approx -\frac{1}{\varepsilon} - \frac{\pi^2}{6}\varepsilon \end{aligned} \quad (5.98)$$

Portanto, até ordem ε , sabemos que

$$\int_0^1 dz \left\{ z^\varepsilon (1-z)^{-1-\varepsilon} + \left(\frac{1}{\varepsilon} + \frac{\pi^2}{6}\varepsilon \right) \delta(1-z) \right\} = 0, \quad (5.99)$$

onde, substituindo o termo em π^2 dado em (5.97), temos

$$\int_0^1 dz \left\{ \frac{1}{\varepsilon} \delta(1-z) + z^\varepsilon (1-z)^{-1-\varepsilon} - \left[\left(\frac{z}{1-z} \right)^\varepsilon \frac{1}{1-z} \right]_+ \right\} = 0. \quad (5.100)$$

De (5.95) e (5.96) obtemos a expansão desejada da função $F(z)$

$$\left(\frac{z}{1-z} \right)^\varepsilon \frac{1}{1-z} \approx -\frac{1}{\varepsilon} \delta(1-z) + \frac{1}{(1-z)_+} - \varepsilon \left(\frac{\ln(1-z)}{1-z} \right)_+ + \varepsilon \frac{\ln(z)}{1-z}. \quad (5.101)$$

Finalmente, inserindo essa expressão em (5.91), obtemos a contribuição a F_2 pela emissão de glúon

$$\begin{aligned} \frac{F_2}{x} \Big|_q &= \frac{\alpha_s}{2\pi} C_F \left(\frac{4\pi\mu^2}{Q^2} \right)^\varepsilon \frac{\Gamma(1-\varepsilon)}{\Gamma(1-2\varepsilon)} \\ &\sum_q Q_q^2 \int_x^1 \frac{dy}{y} q_0(y) \left\{ \frac{2}{\varepsilon^2} \delta(1-z) + \frac{3}{2\varepsilon} \delta(1-z) - \frac{1}{\varepsilon} \frac{1+z^2}{(1-z)_+} \right. \\ &\quad \left. (1+z^2) \left(\frac{\ln(1-z)}{1-z} \right)_+ - \frac{3}{2} \frac{1}{(1-z)_+} - \frac{1+z^2}{1-z} \ln(z) \right. \\ &\quad \left. + 3 + 2z + \frac{7}{2} \delta(1-z) \right\}. \end{aligned} \quad (5.102)$$

Note que o pólo soft associado ao limite $z \rightarrow 1$, multiplicado pelo pólo colinear de (5.91), produz um pólo duplo como era esperado, e que irá cancelar o pólo duplo da amplitude virtual.

5.4 PDFs em Ordem α_s

Calculadas todas as contribuições ao fator de forma F_2 , somando-as de acordo com a expressão

$$\frac{F_2}{x} = \frac{F_2}{x} \Big|_{virt} + \frac{F_2}{x} \Big|_q + \frac{1}{2} \frac{F_2}{x} \Big|_g, \quad (5.103)$$

onde o fator $1/2$ é necessário se desejamos calcular as correções apenas para $q_0(x)$, já que o processo $\gamma g \rightarrow q\bar{q}$ contribui igualmente para $q_0(x)$ e $\bar{q}_0(x)$, e inserindo o resultado em (5.66), verifica-se que os pólos quadráticos se cancelam restando, contudo, pólos simples originados de divergências colineares. A função de distribuição de quarks em NLO de QCD é dada então por

$$\begin{aligned} q(x, Q_F^2) &= q_0(x) \\ &+ \frac{\alpha_s}{2\pi} \int_x^1 \frac{dy}{y} q_0(y) \left[\left(-\frac{1}{\bar{\epsilon}} + \ln \frac{Q_F^2}{\mu^2} \right) P_{qq}(z) + F_{qq}(z) \right] \\ &+ \frac{\alpha_s}{2\pi} \int_x^1 \frac{dy}{y} g(y) \left[\left(-\frac{1}{\bar{\epsilon}} + \ln \frac{Q_F^2}{\mu^2} \right) P_{qg}(z) + F_{qg}(z) \right], \end{aligned} \quad (5.104)$$

onde $1/\bar{\epsilon}$ está definida em (2.92). A função $P_{qq}(z)$ dada por

$$P_{qq}(z) = C_F \frac{1+z^2}{1-z} \quad (5.105)$$

corresponde à função de splitting de Altarelli-Parisi [48] que representa a probabilidade de um quark emitir um glúon e tornar-se, dessa forma, um quark de momento reduzido por uma fração z . Note que $P_{qq}(z)$ é singular no limite $z \rightarrow 1$, o que corresponde à emissão de um glúon soft. Da mesma forma

$$P_{qg}(z) = T_R [z^2 + (1-z)^2] \quad (5.106)$$

corresponde a uma função de splitting de Altarelli-Parisi que representa a probabilidade de um glúon decair em um par $q\bar{q}$ onde o quark (ou o antiquark) carrega uma fração z do seu momento.[†] As funções $F_{qq}(z)$ e $F_{qg}(z)$ são finitas e dadas logo abaixo

$$F_{qq}(z) = C_F \left[(1+z^2) \left(\frac{\ln(1-z)}{1-z} \right)_+ - \frac{3}{2} \frac{1}{(1-z)_+} - \frac{1+z^2}{1-z} \ln(z) \right]$$

[†]Os fatores de cor $T_R = \frac{1}{2}$ e $C_F = \frac{4}{3}$ são deduzidos no apêndice B.

$$+3 + 2z - \left(\frac{9}{2} + \frac{\pi^2}{3}\right)\delta(1-z) \Big] \quad (5.107)$$

$$F_{qg}(z) = T_R \left[[z^2 + (1-z)^2] \ln\left(\frac{1-z}{z}\right) - 8z^2 + 8z - 1 \right] . \quad (5.108)$$

Como sabemos, não é possível calcular as funções de distribuição de pártons através da QCD perturbativa, contudo o regime perturbativo da QCD nos permite obter a evolução das PDFs com a escala de energia através da chamada equação *DGLAP*[§], que é a análoga da equação da função β que descreve a variação de α_s com a escala de energia. A equação DGLAP é dada por

$$\frac{d q(x, Q^2)}{dQ^2} = \frac{\alpha_s}{2\pi} \int_x^1 \frac{d\xi}{\xi} \left[q(\xi, Q^2) P_{qq}\left(\frac{x}{\xi}\right) + g(\xi, Q^2) P_{qg}\left(\frac{x}{\xi}\right) \right] \quad (5.109)$$

para um dado sabor de quark q . Veja que as funções de splitting de Altarelli-Parisi têm um papel fundamental na equação de evolução e, por isso, também são chamadas de *evolution kernels*.

Cabe salientar que a função de distribuição de quarks $q(x, Q_F^2)$ pode ser definida de formas diferentes, diferindo de (5.104) por termos finitos. Tudo depende de quantos e quais termos finitos são absorvidos em $q_0(x)$, a PDF “nua” junto com o pólo colinear $(1/\varepsilon)P_{q,qq}(z)$. A maneira com que escolhemos os termos finitos de $q(x, Q_F^2)$ define um *esquema de fatorização*.

As seções de choque partônicas que contém divergências colineares, geradas pela radiação de glúons ou quarks não-massivos em um dado processo partônico, têm uma estrutura universal, fatorizada para todas as ordens de teoria de perturbação, expressa por

$$\hat{\sigma}_{ij} = \int_0^1 dx_a dx_b \sum_{l,m} \Gamma_{li}(x_a, Q_F^2, \mu^2, \varepsilon) \Gamma_{mj}(x_b, Q_F^2, \mu^2, \varepsilon) \times \sigma_{lm}^{red}(Q_F^2) . \quad (5.110)$$

Os índices $i-m$ caracterizam os pártons do estado inicial e $\sigma_{lm}^{red}(Q_F^2)$ representa uma seção de choque reduzida, livre de singularidades. As funções de splitting universais Γ_{ij} , que contém divergências colineares, representam a probabilidade de encontrar, dentro de um párton inicial j , uma partícula i com uma fração de momento x do momento longitudinal total a uma escala de energia Q_F . Tais divergências podem ser absorvidas em uma redefinição das densidades de pártons em NLO no cômputo de uma seção de choque inclusiva hadrônica [40]. A esse procedimento dá-se o nome de *fatorização de massa*. Até ordem α_s as funções de splitting universais são dadas por

$$\Gamma_{ij}(x, Q_F^2, \mu^2, \varepsilon) = \delta_{ij} \delta(1-x) + \frac{\alpha_s}{2\pi} \left[\left(-\frac{1}{\varepsilon} + \ln\left(\frac{Q_F^2}{\mu^2}\right) \right) P_{ij}(x) + f_{ij}(x, Q_F^2, \mu^2) \right] , \quad (5.111)$$

[§]Dokshitzer–Gribov–Lipatov–Altarelli–Parisi.

onde $P_{ij}(x)$ são as funções de Altarelli-Parisi. As funções f_{ij} são arbitrárias e dependem do esquema de fatorização. No esquema $\overline{MS}$, por exemplo, $f_{ij} = 0$ simplesmente, o que nos permite reescrever Γ_{ij} de uma maneira mais conveniente aos nossos cálculos

$$\Gamma_{ij}(x, Q_F^2, \mu^2, \epsilon) = \Gamma_{ij}^{\overline{MS}}(x, Q_F^2, \mu^2, \epsilon) + \frac{\alpha_s}{2\pi} f_{ij}(x, Q_F^2, \mu^2) . \quad (5.112)$$

A função de distribuição de quarks (5.104), por exemplo, pode ser escrita em termos de (5.112) como

$$q(x, Q_F^2) = \int_x^1 \frac{dy}{y} \left[q_0(y) \left(\Gamma_{qq}^{\overline{MS}} + \frac{\alpha_s}{2\pi} F_{qq} \right) + g_0(y) \left(\Gamma_{qg}^{\overline{MS}} + \frac{\alpha_s}{2\pi} F_{qg} \right) \right] \quad (5.113)$$

onde F_{qq} e F_{qg} podem ser encontradas em (5.107) e (5.108), respectivamente.

As funções de estrutura, que parametrizam a estrutura de um próton da maneira “vista” por um fóton, são também dependentes do esquema escolhido, por exemplo, $F_2(x)$, que é dada por

$$F_2(x) = x \sum_q e_q^2 q_0(x) \quad (5.114)$$

em ordem zero de QCD, no esquema $\overline{MS}$, é dada por

$$\begin{aligned} F_2(x, Q_F^2) &= x \sum_q e_q^2 \int_x^1 \frac{d\xi}{\xi} \left\{ q(\xi, Q_F^2) \left[\delta\left(1 - \frac{x}{\xi}\right) + \frac{\alpha_s}{2\pi} C_q^{\overline{MS}}\left(\frac{x}{\xi}\right) + \dots \right] \right. \\ &\quad \left. + g(\xi, Q_F^2) \left[\frac{\alpha_s}{2\pi} C_g^{\overline{MS}} + \dots \right] \right\} \end{aligned} \quad (5.115)$$

em ordem α_s . As funções $C_q^{\overline{MS}}$ e $C_g^{\overline{MS}}$, chamadas de *funções coeficiente*, são dadas exatamente pelas expressões de F_{qq} e F_{qg} , respectivamente.

Com essa definição de $F_2(x, Q_F^2)$, temos, neste esquema

$$q(x, Q_F^2) = \int_x^1 \frac{dy}{y} [q_0(y) \Gamma_{qq}^{\overline{MS}} + g_0(y) \Gamma_{qg}^{\overline{MS}}] . \quad (5.116)$$

Outra escolha usual é manter a forma original de F_2 em ordem α_s , ou seja,

$$F_2(x, Q_F^2) = x \sum_q e_q^2 q(x, Q_F^2) \quad (5.117)$$

absorvendo todos os termos finitos em $q(x, Q_F^2)$. Esse é o esquema de fatorização *Deep Inelastic Scattering-DIS*, também chamado de esquema físico, e que é mais útil quando se quer extrair as PDFs de experiências de espalhamento inelástico profundo lépton-hádron. No esquema DIS, as densidades de quarks renormalizadas são dadas por (5.113), ou alternativamente, por (5.104), visto que usamos (5.117) como definição de F_2 para todas as ordens de teoria de perturbação.

Esses dois esquemas podem ser definidos por uma constante λ_{FC} de maneira que

$$\int_x^1 \frac{dy}{y} [q_0(y) \Gamma_{qq}^{\overline{MS}} + g_0(y) \Gamma_{qg}^{\overline{MS}}] = q(x, Q_F^2) - \frac{\alpha_s}{2\pi} \lambda_{FC} \int_{x_a}^1 \frac{dy}{y} [q(y, Q_F^2) F_{qq}(z) + g(y, Q_F^2) F_{qg}(z)] , \quad (5.118)$$

onde $\lambda_{FC} = 0$ define o esquema $\overline{MS}$ enquanto $\lambda_{FC} = 1$ define o esquema DIS.

Capítulo 6

Correções de Ordem α_s ao Processo de Drell-Yan

$$p\bar{p} \rightarrow e^+e^- X$$

6.1 Drell-Yan em Leading-Log Approximation

Em ordem zero, a produção de pares elétron-pósitron em colisões hádron-hádron é dada pela reação

$$q(p_a) + \bar{q}(p_b) \rightarrow e^-(k_1) + e^+(k_2) . \quad (6.1)$$

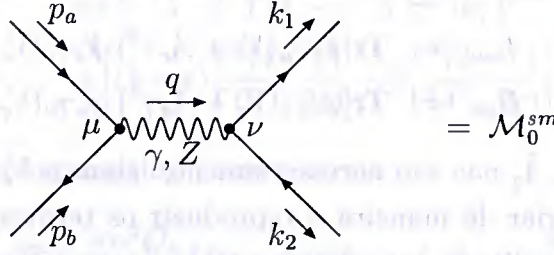


Figura 6.1: Processo de Drell-Yan em nível de árvore.

A amplitude do processo de aniquilação é obtida a partir do diagrama de Feynman mostrado na figura (6.1)

$$\mathcal{M}_0^{sm} = \frac{i}{k_V^2} \bar{u}_e(k_1) \gamma^\mu (V_e + A_e \gamma^5) v_e(k_2) \bar{v}_q(p_b) \gamma_\mu (V_q + A_q \gamma^5) u_q(p_a) \quad (6.2)$$

onde $V = \gamma, Z$.

Na expressão acima V_e, V_q e A_e, A_q correspondem aos acoplamentos vetorial e axial dos léptons e quarks, respectivamente, ao fóton e ao bóson Z. No caso do acoplamento com o fóton, $A_e = A_q = 0$, e quando a interação envolve o bóson Z, temos

$$V_f = \frac{1}{\sqrt{2}} \left(\frac{G_F M_Z^2}{\sqrt{2}} \right)^{1/2} (R_f + L_f) ,$$

$$A_f = \frac{1}{\sqrt{2}} \left(\frac{G_F M_Z^2}{\sqrt{2}} \right)^{1/2} (R_f - L_f) ,$$

onde as constantes R_e , R_q , L_e e L_q são dadas no apêndice A.

O quadrado do elemento de matriz é dado, após soma sobre os números quânticos finais e média dos iniciais, por

$$|\overline{\mathcal{M}_0}|_V^2 = \frac{1}{4} \frac{3}{9} \frac{1}{|k_V^2|^2} L_{\mu\nu} H^{\mu\nu} , \quad (6.3)$$

sendo que k_γ^2 corresponde ao denominador do propagador do fóton, k_Z^2 ao denominador do propagador do bóson Z, onde o pólo do Z é regularizado pela Breit-Wigner fazendo $k_Z^2 = k^2 - m_Z^2 + i\Gamma_Z m_Z$ e $|k_{\gamma Z}^2|^2 = k^2[k^2 - m_Z^2 + i\Gamma_Z m_Z]$ para o termo de interferência.

A amplitude completa é obtida pela soma

$$|\overline{\mathcal{M}_0}|^2 = |\overline{\mathcal{M}_\gamma}|^2 + |\overline{\mathcal{M}_Z}|^2 + 2\Re(\mathcal{M}_\gamma^\dagger \mathcal{M}_Z) , \quad (6.4)$$

$$\equiv \sum_V |\overline{\mathcal{M}_0}|_V^2 , \quad (6.5)$$

aqui V denota as contribuições de γ , Z e a interferência, dada pelo último termo de (6.4).

Os tensores leptônico e hadrônico correspondem aos seguintes traços

$$L_{\mu\nu} = \text{Tr}[k_1 \gamma_\mu (V_e + A_e \gamma^5) k_2 \gamma_\nu (\tilde{V}_e + \tilde{A}_e \gamma^5)] \quad (6.6)$$

$$H_{\mu\nu} = \text{Tr}[\not{p}_b \gamma_\mu (V_q + A_q \gamma^5) \not{p}_a \gamma_\nu (\tilde{V}_q + \tilde{A}_q \gamma^5)] , \quad (6.7)$$

onde $\tilde{V}_e, \tilde{A}_e, \tilde{V}_q, \tilde{A}_q$ não são necessariamente iguais a V_e, A_e, V_q, A_q , respectivamente, mas podem variar de maneira a reproduzir os termos correspondentes ao módulo quadrado da amplitude que contém a troca de um fóton ou um bóson Z, e os termos de interferência γ/Z .

Expandindo os traços, obtemos a contração dos tensores leptônico e hadrônico

$$L_{\mu\nu} H^{\mu\nu} = 8[\hat{u}^2 + \hat{t}^2 - \varepsilon \hat{s}^2] d_e d_q - 4(1 - 2\varepsilon)(2 - 2\varepsilon)(\hat{u}^2 - \hat{t}^2) h_e h_q ,$$

onde

$$d_f \equiv V_f \tilde{V}_f + A_f \tilde{A}_f , \quad (6.8)$$

$$h_f \equiv V_f \tilde{A}_f + \tilde{V}_f A_f . \quad (6.9)$$

As variáveis de Mandelstam do processo $2 \rightarrow 2$, $\hat{s}$, $\hat{t}$ e $\hat{u}$ são definidas por

$$\hat{s} = (p_a + p_b)^2 , \quad (6.10)$$

$$\hat{t} = (p_a - k_1)^2 , \quad (6.11)$$

$$\hat{u} = (p_b - k_1)^2 . \quad (6.12)$$

Definindo $z = \cos \theta$, onde θ é o ângulo formado pelos vetores $\vec{p}_a$ e $\vec{k}_1$ no sistema de centro de massa (CM) do par $q\bar{q}$, obtemos

$$L_{\mu\nu}H^{\mu\nu} = 8 \left[\frac{\hat{s}^2}{2}(1+z^2) - \varepsilon \hat{s}^2 \right] d_e d_q - 4(1-2\varepsilon)(2-2\varepsilon)z\hat{s}^2 h_e h_q . \quad (6.13)$$

O espaço de fase de 2 partículas não massivas em $d = 4 - 2\varepsilon$ dimensões espaço-temporais é dado por (C.12)

$$\Phi_2^d = \frac{2^{2\varepsilon}}{16\pi} \left(\frac{4\pi}{\hat{s}} \right)^\varepsilon \frac{1}{\Gamma(1-\varepsilon)} \int_{-1}^1 dz (1-z^2)^{-\varepsilon} . \quad (6.14)$$

Note que o termo proporcional a $h_e h_q$ na equação (6.13) não contribue quando integrado em Φ_2^d , pois a função multiplicada por $h_e h_q$ é ímpar e está sendo integrada no intervalo simétrico $[-1, +1]$.

Integrar o termo em $d_e d_q$ envolve realizar a mudança de variável $z = 2v - 1$, de modo que o resultado é expresso em termos de funções Γ . Usando as identidades e expansões da função Γ do apêndice D, a seção de choque de Born em d dimensões é dada por

$$\begin{aligned} \hat{\sigma}_0^{sm}(\hat{s}) &= \frac{1}{8\pi} \left(\frac{4\pi\mu^2}{\hat{s}} \right)^\varepsilon \frac{1}{\Gamma(1-\varepsilon)} \frac{1}{3} \sum_V \frac{\hat{s}}{|k_V^2|^2} d_e d_q \\ &\times \left[(1-\varepsilon) \frac{\Gamma^2(1-\varepsilon)}{\Gamma(2-2\varepsilon)} - 2 \frac{\Gamma^2(2-\varepsilon)}{\Gamma(4-2\varepsilon)} \right] . \end{aligned} \quad (6.15)$$

Em 4 dimensões espaço-temporais $\hat{\sigma}_0^{sm}$ tem a forma

$$\begin{aligned} \hat{\sigma}_0^{sm}(\hat{s}) &= \frac{4\pi\alpha^2}{9\hat{s}} Q_q^2 - \frac{\pi\alpha^2 Q_q}{18s_w^2 c_w^2} (R_e + L_e)(R_q + L_q) \frac{\hat{s} - m_Z^2}{(\hat{s} - m_Z^2)^2 + \Gamma_Z^2 m_Z^2} \\ &+ \frac{\pi\alpha^2}{144s_w^4 c_w^4} (R_e^2 + L_e^2)(R_q^2 + L_q^2) \frac{\hat{s}^2}{(\hat{s} - m_Z^2)^2 + \Gamma_Z^2 m_Z^2} . \end{aligned} \quad (6.16)$$

A seção de choque hadrônica de Drell-Yan em Leading-Logarithm Approximation é obtida pela convolução da seção de choque partônica $\hat{\sigma}_0^{sm}$ dada acima com as funções de distribuição de quarks obtidas pelas experiências pelo ajuste em next-to-leading-order, ou seja, pelas densidades dependentes de escala (Q), e finalmente, somando sobre todos os pártons que participam do processo

$$\sigma^{LL}(p\bar{p} \rightarrow e^+e^-X) = \sum_{q,\bar{q}} \int dx_a dx_b [q(x_a, Q^2)\bar{q}(x_b, Q^2) + \bar{q}(x_a, Q^2)q(x_b, Q^2)] \times \hat{\sigma}_0^{sm}(\hat{s}) . \quad (6.17)$$

A seguir iniciaremos o cálculo do processo de Drell-Yan em ordem α_s em QCD utilizando o método NLO Monte Carlo descrito no capítulo 4.

6.2 Contribuições Virtuais ao Nível de 1-Loop

Nesta seção calculamos as contribuições virtuais de ordem α_s ao processo de Drell-Yan o que inclui a contribuição das autoenergias do quarks e da correção ao vértice $(\gamma, Z) q\bar{q}$.

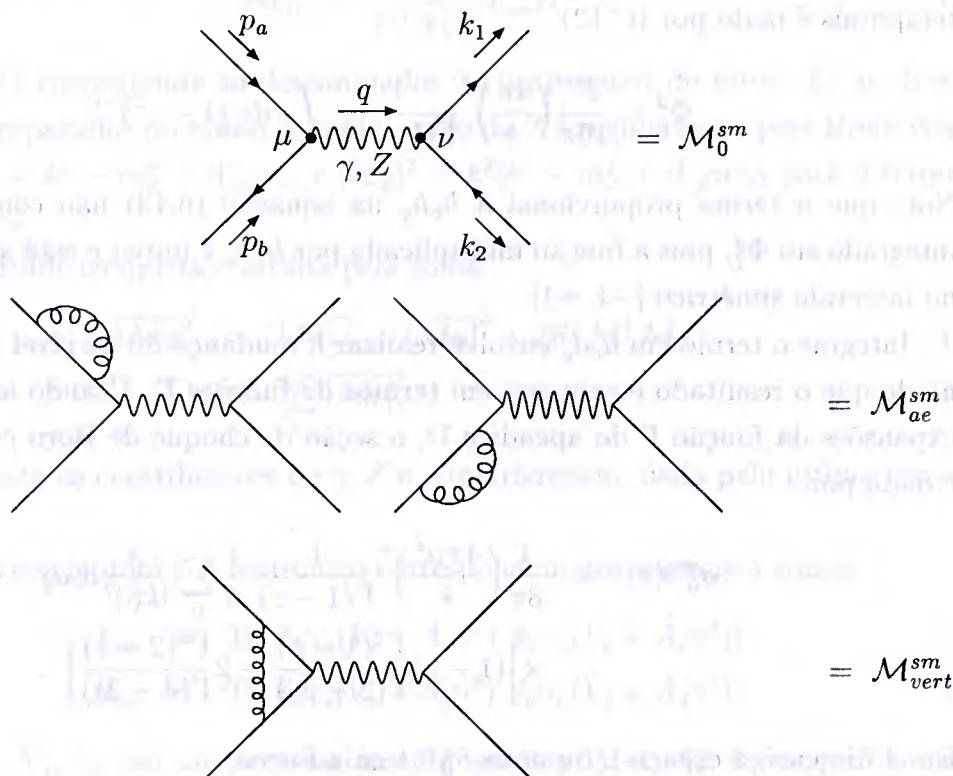


Figura 6.2: Amplitudes virtuais ao nível de 1-loop no SM.

As divergências das correções virtuais são regularizadas efetuando os cálculos em $d = 4 - 2\epsilon$ dimensões tanto para as integrais dos loops quanto para os traços de matrizes de Dirac. Tais divergências são de natureza ultravioleta (UV), soft e colinear (também chamada de singularidade de massa), essas duas últimas são denominadas coletivamente de divergências infravermelhas (IR). Todas elas se apresentam como pólos da forma ϵ^{-i} ($i = 1, 2$). Pólos duplos são gerados em duas situações distintas: quando um glúon virtual é trocado entre duas linhas de quarks não-massivas, como no caso da correção do vértice e quando um glúon real é emitido no regime soft e colinear ao mesmo tempo. Em qualquer outra situação somente pólos simples ocorrem.

As amplitudes virtuais foram obtidas em termos de integrais escalares de Passarino-Veltman [49], com o auxílio do pacote FeynCalc para Mathematica, de modo a facilitar o procedimento de renormalização das amplitudes divergentes no UV. Ao

contrário do processo de aniquilamento e^+e^- em pares $q\bar{q}$, utilizamos o gauge de Feynman $\eta = 0$ nos cálculos e o esquema *on-shell* para realizar as renormalizações. As correções de autoenergia das partículas externas *on-shell* (quarks não-massivos) são multiplicadas pelo fator 1/2 de modo a levar em conta de maneira apropriada a transição das funções de Green para a matriz-S.

A seguir iniciamos os cálculos em NLO de QCD para o processo de Drell-Yan.

6.2.1 Autoenergia dos Quarks

A autoenergia de um quark não-massivo foi calculada na subseção (3.3.2) em um gauge covariante arbitrário η e sua expressão é dada por (3.39). Verificamos naquela ocasião a conveniência do uso do gauge de Landau para o cálculo das amplitudes virtuais dado que a constante de renormalização da função de onda do quark Z_2 é finita e vale 1, simplesmente. Então, pela identidade de Ward-Takahashi, podíamos relacionar Z_1 e Z_2 de uma forma trivial simplificando enormemente o cálculo. O mesmo poderia ser feito agora, no entanto, no próximo capítulo, vamos somar a contribuição de leptoquarks escalares às amplitudes de Drell-Yan até ordem α_s , de modo que o uso do gauge de Feynman será mais conveniente devido à forma simples com que o propagador do glúon, e consequentemente todas as outras expressões, tem nesse gauge. No entanto, no gauge de Feynman, Z_2 é divergente, contendo pólos em ε , sendo preciso utilizar o procedimento de renormalização para separar o pólo UV do pólo IR.

Para tanto, introduzimos o contratermo $i\delta_2 \not{p}$ para absorver a divergência UV

$$-i\Sigma_q(p) = i \not{p} \frac{\alpha_s}{4\pi} C_F \left(\frac{-p^2}{4\pi\mu^2} \right)^{-\varepsilon} (1 + \varepsilon) \Gamma(\varepsilon) + i\delta_2 \not{p} , \quad (6.18)$$

onde $\delta_2 = Z_2 - 1$, como em (2.80). A seguir vamos considerar tal renormalização em dois esquemas diferentes.

1. Esquema Off-Shell

No esquema off-shell, a subtração do pólo UV é feita em um ponto euclidiano $p^2 = -\lambda^2$ de modo a regular a divergência IR associada ao limite $p^2 \rightarrow 0$, dado que os quarks estão *on-shell*. Logo, aplicamos a seguinte condição de renormalização

$$\frac{d\Sigma_q}{d\not{p}}(p^2 = -\lambda^2) = 0 , \quad (6.19)$$

o que resulta no seguinte contratermo

$$\delta_2 = -\frac{\alpha_s}{4\pi} C_F \left(\frac{\lambda^2}{\mu^2} \right)^{-\varepsilon} \left(\frac{1}{\varepsilon} + 1 \right) . \quad (6.20)$$

A autoenergia renormalizada é, portanto, dada por

$$-i\Sigma_q^{ren}(p) = i\frac{\alpha_s}{4\pi}C_F \not{p} \left[\left(\frac{-p^2}{\mu^2} \right)^{-\varepsilon} - \left(\frac{-\lambda^2}{\mu^2} \right)^{-\varepsilon} \right] \left(\frac{1}{\varepsilon} + 1 \right). \quad (6.21)$$

Colocando o quark na camada de massa, e expandindo em ε obtemos

$$\Sigma_q^{ren}(p^2 = 0) = \frac{\alpha_s}{4\pi}C_F \not{p} \left(\frac{1}{\varepsilon} + 1 - \ln \frac{\lambda^2}{\mu^2} \right). \quad (6.22)$$

Note que fazendo $\lambda = 0$ em (6.21), para $\varepsilon < 0$, o que corresponde ao esquema on-shell, a expressão resultante seria da forma

$$\frac{1}{\varepsilon} - \ln \left(\frac{-p^2}{\mu^2} \right),$$

que é divergente para $p^2 = 0$, exibindo uma singularidade de massa.

2. Esquema On-Shell

Vamos primeiramente definir a função $B(p^2)$ dada por

$$B(p^2) = -i\frac{\Sigma_q^{ren}(p)}{\not{p}}, \quad (6.23)$$

onde $B(p^2)$ pode escrita na seguinte forma integral

$$B(p^2) = i\frac{\alpha_s}{4\pi}C_F(1+\varepsilon)\Gamma(1+\varepsilon) \int_{-p^2}^{\infty} \frac{dq^2}{q^2} \left(\frac{q^2}{4\pi\mu^2} \right)^{-\varepsilon} + i\delta_2, \quad (6.24)$$

que pode ser dividida em duas partes, uma contendo a divergência UV (limite superior igual a ∞) e a outra contendo a divergência de massa (limite inferior igual a $-p^2$)

$$B(p^2) = i\frac{\alpha_s}{4\pi}C_F(1+\varepsilon)\Gamma(1+\varepsilon) \left(\int_{\mu^2}^{\infty} dq^2 (q^2)^{-1-\varepsilon} + \int_{-p^2}^{\mu^2} dq^2 (q^2)^{-1-\varepsilon} \right) (4\pi\mu^2)^\varepsilon + i\delta_2. \quad (6.25)$$

A subtração agora é feita no ponto $p^2 = 0$ usando o contratermo

$$\delta_2 = -\frac{\alpha_s}{4\pi}C_F(1+\varepsilon) \int_{\mu^2}^{\infty} dq^2 (q^2)^{-1-\varepsilon} (4\pi\mu^2)^\varepsilon \Gamma(1+\varepsilon), \quad (6.26)$$

de modo que, calculando as integrais e subtraindo o pólo UV, obtemos

$$-i\Sigma_q^{Ren}(p) = -i\frac{\alpha_s}{4\pi}C_F N_\varepsilon \mu^{2\varepsilon} \frac{1}{\varepsilon} \not{p}, \quad (6.27)$$

onde $N_\varepsilon = (4\pi)^\varepsilon \Gamma(1+\varepsilon)$. Todas as integrais divergentes no IR são reguladas com $\varepsilon < 0$.

6.2.2 Correção do Vértice $(\gamma, Z) q\bar{q}$

A correção de ordem α_s ao vértice $Vq\bar{q}$, onde $V = \gamma, Z$, é dada pela amplitude correspondente ao seguinte diagrama de Feynman

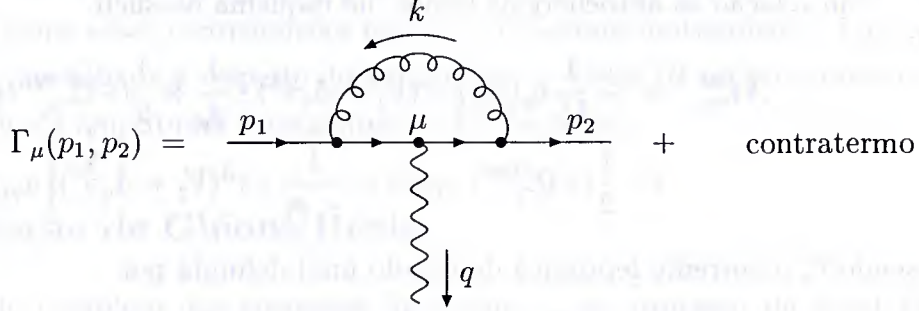


Figura 6.3: Correção do vértice $(\gamma, Z)q\bar{q}$ ao nível de 1-loop em QCD.

A expressão do loop é dada por

$$\Gamma_\mu(p_1, p_2) = -ig_s^2 C_F \int D^d k \frac{N}{k^2(k+p_1)^2(k-p_2)^2} \gamma_\mu \quad (6.28)$$

$$N = -4p_1 \cdot p_2 + 4(p_1 - p_2) \cdot k + 2(1 - \varepsilon)k^2, \quad (6.29)$$

onde

$$D^d k \equiv \frac{d^d k}{(2\pi)^d}. \quad (6.30)$$

Renormalização no Esquema On-Shell

Renormalizamos o vértice subtraindo de $\Gamma_\mu(p_1, p_2)$ seu valor no ponto $p_1^2 = p_2^2 = 0$, $q = p_1 - p_2 = 0$, ou seja, com as partículas externas on-shell e para diferença de momentos nula, como no caso da QED. Dessa forma, substituindo o valor do contratermo, obtemos o resultado

$$\Gamma_\mu^{Ren}(p_1, p_2) = \Gamma_\mu(p_1, p_2) - \Gamma_\mu(p_1^2 = p_2^2 = 0, q = 0) \quad (6.31)$$

$$= -ig_s^2 C_F [2q^2 C_0(0, 0, -q^2, 0, 0) - 4B_0(-q^2, 0, 0)] \gamma_\mu. \quad (6.32)$$

Calculando as funções escalares e expandindo em ε temos

$$\Gamma_\mu^{Ren}(p_1, p_2) = \frac{\alpha_s}{4\pi} C_F N_\varepsilon \left(\frac{\mu^2}{-q^2} \right)^\varepsilon \left(-\frac{2}{\varepsilon^2} - \frac{4}{\varepsilon} + F_{vert}^{sm} \right) \gamma_\mu \quad (6.33)$$

a função F_{vert}^{sm} representa as partes finitas da expressão do vértice e pode ser obtida a partir das funções escalares apresentadas no Apêndice D.

6.2.3 Contribuição Virtual Total

A contribuição virtual total é obtida pela inserção das autoenergias e da correção do vértice nas amplitudes obtidas pelos diagramas de Feynman da figura (6.2).

Em relação às autoenergias temos, no esquema on-shell,

$$\begin{aligned} \mathcal{M}_{ae}^{sm} = & -\frac{1}{k_V^2} \bar{v}_q(p_b) \left[\gamma^\mu (V_q + A_q \gamma^5) \cdot \frac{1}{\not{p}_a} \times \frac{1}{2} (-i\Sigma_q^{Ren}(p_a)) \right. \\ & \left. + \frac{1}{2} (-i\Sigma_q^{Ren}(-p_b)) \times \frac{1}{-\not{p}_b} \cdot \gamma^\mu (V_q + A_q \gamma^5) \right] u_q(p_a) \cdot \ell_\mu \end{aligned} \quad (6.34)$$

sendo ℓ_μ a corrente leptônica do estado final definida por

$$\ell_\mu = \bar{u}_e(k_1) \gamma_\mu (V_e + A_e \gamma^5) v_e(k_2) . \quad (6.35)$$

Após soma e média sobre números quânticos obtemos o termo de interferência

$$2\Re(\mathcal{M}_0^{sm*} \mathcal{M}_{ae}^{sm}) = \frac{\alpha_s}{2\pi} C_F N_\epsilon \left(\frac{\mu^2}{\hat{s}} \right)^\epsilon \frac{1}{\epsilon} |\overline{\mathcal{M}_0^{sm}}|^2 . \quad (6.36)$$

A contribuição do vértice, por sua vez, é dada por

$$\mathcal{M}_{vert}^{sm} = \frac{-i}{k_V^2} \bar{v}_q(p_b) [-i\Gamma_{Ren}^\mu (V_q + A_q \gamma^5)] u_q(p_a) \cdot \ell_\mu \quad (6.37)$$

e a sua interferência com a amplitude de Born é da forma

$$2\Re(\mathcal{M}_0^{sm*} \mathcal{M}_{vert}^{sm}) = \frac{\alpha_s}{2\pi} C_F N_\epsilon \left(\frac{\mu^2}{\hat{s}} \right)^\epsilon \left(-\frac{2}{\epsilon^2} - \frac{4}{\epsilon} - 8 + \pi^2 \right) |\overline{\mathcal{M}_0^{sm}}|^2 . \quad (6.38)$$

Finalmente, somando todas as contribuições virtuais temos que

$$\begin{aligned} 2\Re(\mathcal{M}_0^{sm*} \mathcal{M}_V^{sm}) &= \frac{\alpha_s}{2\pi} C_F N_\epsilon \left(\frac{\mu^2}{\hat{s}} \right)^\epsilon \left[\frac{1}{\epsilon} - \frac{2}{\epsilon^2} - \frac{4}{\epsilon} + F_V^{sm} \right] |\overline{\mathcal{M}_0^{sm}}|^2 \\ &= \frac{\alpha_s}{2\pi} C_F N_\epsilon \left(\frac{\mu^2}{\hat{s}} \right)^\epsilon \left[-\frac{2}{\epsilon^2} - \frac{3}{\epsilon} \right] |\overline{\mathcal{M}_0^{sm}}|^2 + \frac{\alpha_s}{2\pi} C_F |\overline{\mathcal{M}_0^{sm}}|^2 \times F_V^{sm} \end{aligned} \quad (6.39)$$

onde F_V^{sm} representa as partes finitas provenientes das amplitudes virtuais e é obtida através das funções escalares do Apêndice D.

Integrando no espaço de fase de 2 partículas obtemos a contribuição virtual à seção de choque de 2 corpos

$$\hat{\sigma}_{virt}(\hat{s}) = \frac{\alpha_s}{2\pi} C_F N_\epsilon \left(\frac{\mu^2}{\hat{s}} \right)^\epsilon \left[-\frac{2}{\epsilon^2} - \frac{3}{\epsilon} + F_V^{sm} \right] \times \hat{\sigma}_0^{sm}(\hat{s}) . \quad (6.40)$$

Note que a estrutura de divergências surge de uma forma diferente neste caso, em relação ao processo $e^+e^- \rightarrow q\bar{q}$. Naquela ocasião, onde todos os cálculos foram

feitos no gauge de Landau, todos os pólos IR eram provenientes do vértice, dado que as autoenergias não contribuíam. No caso do processo de Drell-Yan no gauge de Feynman, contudo, a autoenergia dos quarks contribuem com um pólo colinear extra, enquanto o vértice também possui um pólo colinear adicional mas de sinal oposto, de modo que na soma essas contribuições extras se cancelam mutuamente. Para uma discussão mais detalhada a respeito do surgimento de pólos IR no procedimento de renormalização de amplitudes divergentes no UV veja [50].

6.3 Radiação de Glúons Reais

Para o cálculo completo das correções de ordem α_s ao processo de Drell-Yan é preciso levar em conta, além das já mencionadas correções virtuais, as contribuições provenientes da emissão de glúons reais e de quarks não-massivos no estado final. A radiação de glúons reais é obtida pela adição de um glúon no estado final a partir da reação partônica em LO

$$q(p_a) + \bar{q}(p_b) \rightarrow e^+(k_1) + e^-(k_2) + g(p) . \quad (6.41)$$

Os diagramas de Feynman, a partir dos quais obtemos as amplitudes em ordem α_s são os dados na figura (6.4)

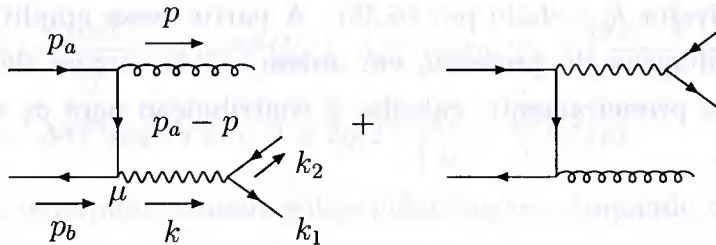


Figura 6.4: Diagramas de Feynman para emissão de glúon real.

Os 4-momentos envolvidos no processo obedecem às seguintes identidades, lembrando que todas as partículas não possuem massa

$$\begin{aligned} p_a + p_b &= p + k_1 + k_2 , \\ p_a^2 = p_b^2 &= p^2 = k_1^2 = k_2^2 = 0 , \\ k &\equiv k_1 + k_2 . \end{aligned}$$

Os invariantes de Mandelstam linearmente independentes, podem ser escolhidos da seguinte forma

$$\hat{s} = (p_a + p_b)^2 , \quad (6.42)$$

$$\hat{t} = (p_a - p)^2, \quad (6.43)$$

$$\hat{u} = (p_b - p)^2, \quad (6.44)$$

$$\hat{t}_1 = (p_a - k_1)^2, \quad (6.45)$$

$$\hat{t}_2 = (p_b - k_1)^2. \quad (6.46)$$

Além disso, existem cinco outros invariantes linearmente dependentes

$$k^2 = (k_1 + k_2)^2 = \hat{s} + \hat{t} + \hat{u}, \quad (6.47)$$

$$\hat{s}_{12} = (p + k_1)^2 = -k^2 - \hat{t}_1 - \hat{t}_2, \quad (6.48)$$

$$\hat{s}_{13} = (p + k_2)^2 = \hat{s} + \hat{t}_1 + \hat{t}_2, \quad (6.49)$$

$$\hat{t}_{b3} = (p_b - k_2)^2 = -k^2 - \hat{t}_2 + \hat{t}, \quad (6.50)$$

$$\hat{u}_{a3} = (p_a - k_2)^2 = -k^2 - \hat{t}_1 + \hat{u}. \quad (6.51)$$

A cinemática do processo está ilustrada na figura (6.5) abaixo. A amplitude correspondente aos diagramas é dada por

$$\mathcal{M}_R^{sm} = \frac{-ig_s T^a}{k_V^2} \bar{v}_q(p_b) \left[\gamma^\mu (V_q + A_q \gamma^5) \frac{(\not{p}_a - \not{p})}{\hat{t}} \not{\epsilon}^*(p) + \not{\epsilon}^*(p) \frac{(\not{p} - \not{p}_b)}{\hat{u}} \gamma^\mu (V_q + A_q \gamma^5) \right] u_q(p_a) \times \ell_\mu, \quad (6.52)$$

onde, o quadrivetor ℓ_μ é dado por (6.35). A partir dessa amplitude podemos calcular as contribuições do processo, em ordem α_s , às secções de choque de 2 e 3 corpos. Vamos primeiramente, calcular a contribuição para σ_2 no regime soft do glúon emitido.

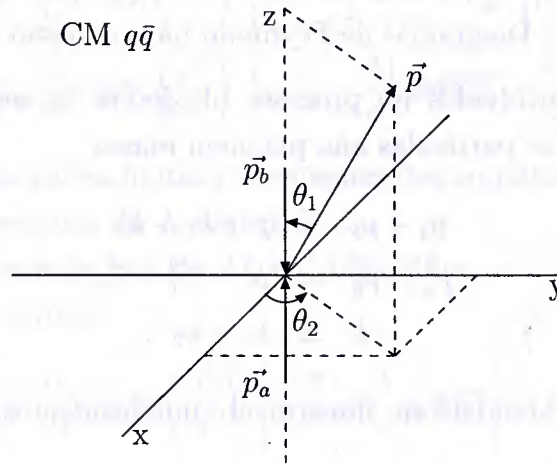


Figura 6.5: Cinemática do processo $q\bar{q} \rightarrow e^+e^-g$.

6.3.1 Emissão de Glúon Soft

No referencial do centro de massa do par $q\bar{q}$, os 4-momentos p_a , p_b e p podem ser escritos como

$$p_a = \frac{\sqrt{\hat{s}}}{2}(1, 0, 0, 1) , \quad (6.53)$$

$$p_b = \frac{\sqrt{\hat{s}}}{2}(1, 0, 0, -1) , \quad (6.54)$$

$$p = E_g(1, \sin\theta_1 \cos\theta_2, \sin\theta_1 \sin\theta_2, \cos\theta_1) , \quad (6.55)$$

de modo que os invariantes $\hat{t}$ e $\hat{u}$ são dados por

$$\hat{t} = -E_g\sqrt{\hat{s}}(1 - \cos\theta_1) \quad (6.56)$$

$$\hat{u} = -E_g\sqrt{\hat{s}}(1 + \cos\theta_1) . \quad (6.57)$$

Note agora que a amplitude de emissão de glúon real é divergente nos limites em que $\hat{t}, \hat{u} \rightarrow 0$, ou seja, quando $E_g \rightarrow 0$, que corresponde à divergência soft, ou quando $\cos\theta_1 \rightarrow \pm 1$, que corresponde à divergências colineares.

Para o cálculo do limite soft usaremos a chamada *aproximação eikonal* que consiste na identificação $\not{p} \rightarrow 0$ no numerador de (6.52). A expressão resultante pode ser ainda mais simplificada com o uso das equações de Dirac para os espiniores u e v , de modo que a amplitude resultante na aproximação soft é dada por

$$\begin{aligned} \mathcal{M}_{R,soft}^{sm} &= \frac{-ig_s T^a}{k_V^2} \bar{v}_q(p_b) \gamma^\mu (V_q + A_q \gamma^5) u_q(p_a) \ell_\mu \times \left[\frac{2p_a \cdot \epsilon^*(p)}{\hat{t}} - \frac{2p_b \cdot \epsilon^*(p)}{\hat{u}} \right] \\ &= \mathcal{M}_0^{sm}(q\bar{q} \rightarrow e^+e^-) \times 2g_s T^a \left[\frac{p_b^\mu}{\hat{u}} - \frac{p_a^\mu}{\hat{t}} \right] \epsilon_\mu^*(p) , \end{aligned} \quad (6.58)$$

e cujo módulo quadrado, somado sobre polarizações e tomando médias é da forma

$$|\overline{\mathcal{M}_{R,soft}^{sm}}|^2 = |\overline{\mathcal{M}_0^{sm}}|^2 \frac{4\hat{s}}{\hat{u}\hat{t}} g_s^2 C_F . \quad (6.59)$$

O próximo passo, agora, consiste em calcular o espaço de fase de 3 partículas na aproximação de glúon soft.

Espaço de Fase de 3 Partículas na Aproximação Soft

O espaço de fase de 3 partículas é dado por

$$d^d\Phi_3 = \frac{d^{d-1}p}{(2\pi)^{d-1}2E_p} \frac{d^{d-1}k_1}{(2\pi)^{d-1}2E_1} \frac{d^{d-1}k_2}{(2\pi)^{d-1}2E_2} \times (2\pi)^d \delta^d(p_a + p_b - p - k_1 - k_2) . \quad (6.60)$$

Aqui, a aproximação é realizada na função δ

$$\delta^d(p_a + p_b - p - k_1 - k_2) \approx \delta^d(p_a + p_b - k_1 - k_2) , \quad (6.61)$$

de modo que

$$d^d\Phi_3 \simeq \frac{d^{d-1}p}{(2\pi)^{d-1}2E_p} d^d\Phi_2 . \quad (6.62)$$

A seção de choque diferencial é obtida através da expressão

$$\begin{aligned} d\hat{\sigma}_{soft} &= \frac{1}{4\sqrt{(p_a \cdot p_b)^2}} d^d\Phi_3 |\overline{\mathcal{M}}_R^{sm}|_{soft}^2 \\ &= \frac{1}{2\hat{s}} d^d\Phi_2 |\overline{\mathcal{M}}_0^{sm}|^2 4\pi\alpha_s\mu^{2\varepsilon} C_F \int \frac{d^{d-1}p}{(2\pi)^{d-1}2E_p} \frac{4\hat{s}}{\hat{u}\hat{t}} , \end{aligned} \quad (6.63)$$

onde $\hat{t}$ e $\hat{u}$ são dados por (6.43) e (6.44), respectivamente, enquanto o seu produto é da forma

$$\hat{u}\hat{t} = \hat{s}E_p^2 \sin^2 \theta_1 . \quad (6.64)$$

Substituindo

$$\begin{aligned} \frac{d^{d-1}p}{(2\pi)^{d-1}2E_p} &= \frac{1}{2(2\pi)^{d-1}} dE_p E_p^{d-2} \times \\ &\times \sin^{d-3} \theta_1 d\theta_1 \sin^{d-4} \theta_2 d\theta_2 \times \cdots \times d\theta_{d-2} \end{aligned} \quad (6.65)$$

na equação (6.63), integrando sobre as variáveis angulares triviais, usando as identidades da função Γ do apêndice D, e, finalmente integrando sobre a energia do glúon na região soft

$$0 < E_p < \delta_s \frac{\sqrt{\hat{s}}}{2} \quad (6.66)$$

obtemos

$$\begin{aligned} d\hat{\sigma}_{soft} &= d\hat{\sigma}_0^{sm} C_F \frac{\alpha_s}{2\pi} \left(\frac{4\pi\mu^2}{\hat{s}} \right)^\varepsilon \frac{\Gamma(1-\varepsilon)}{\Gamma(1-2\varepsilon)} \frac{2}{\varepsilon^2} e^{-2\varepsilon \ln(\delta_s)} \\ &\simeq d\hat{\sigma}_0^{sm} C_F \frac{\alpha_s}{2\pi} \left(\frac{4\pi\mu^2}{\hat{s}} \right)^\varepsilon \frac{\Gamma(1-\varepsilon)}{\Gamma(1-2\varepsilon)} \left[\frac{2}{\varepsilon^2} - \frac{4}{\varepsilon} \ln(\delta_s) + 4 \ln^2(\delta_s) \right] . \end{aligned} \quad (6.67)$$

A integral sobre a região soft da energia do glúon é responsável por um pólo em ε e pelos termos logarítmicos em δ_s , após a expansão em ε da integral

$$\int_0^{\delta_s \frac{\sqrt{\hat{s}}}{2}} dE_p E_p^{1-2\varepsilon} = \frac{1}{-2\varepsilon} \frac{\hat{s}^{-\varepsilon}}{2^{-2\varepsilon}} e^{-2\varepsilon \ln(\delta_s)} .$$

Integrando sobre o espaço de fase de 2 partículas $d^d\Phi_2$ dada por (C.12), finalmente obtemos a seção de choque do subprocesso de emissão de glúon na aproximação soft

$$\hat{\sigma}_{soft} = \frac{\alpha_s}{2\pi} C_F \left(\frac{4\pi\mu^2}{\hat{s}} \right)^\varepsilon \frac{\Gamma(1-\varepsilon)}{\Gamma(1-2\varepsilon)} \left[\frac{2}{\varepsilon^2} - \frac{4}{\varepsilon} \ln(\delta_s) + 4 \ln^2(\delta_s) \right] \times \hat{\sigma}_0^{sm}(\hat{s}) . \quad (6.68)$$

Somando as equações (6.40) e (6.68), vemos que os pólos quadráticos em ε se cancelam

$$\hat{\sigma}_{virt} + \hat{\sigma}_{soft} = C_F \frac{\alpha_s}{2\pi} N_c \left(\frac{\mu^2}{\hat{s}} \right)^\varepsilon \left[-\frac{1}{\varepsilon} (3 + 4 \ln(\delta_s)) - 8 + \pi^2 + 4 \ln^2(\delta_s) \right] , \quad (6.69)$$

e, como, já havíamos antecipado, os termos finitos dependem sensivelmente do cut-off soft δ_s .

6.3.2 Emissão de Glúon Colinear

O elemento de matriz para esse processo é dado por (6.52), vamos, no entanto definir as quantidades

$$\Gamma_{q,e}^\mu = \gamma^\mu (V_{q,e} + A_{q,e} \gamma^5) , \quad (6.70)$$

$$C_V = \frac{-ig_s T^a}{k_V^2(z)} , \quad (6.71)$$

de maneira que a amplitude toma a seguinte forma

$$\mathcal{M}_R^{sm} = C_V \bar{v}_q(p_b) \left[\Gamma^\mu \frac{(\not{p}_a - \not{p})}{\hat{t}} \not{\epsilon}^*(p) + \not{\epsilon}^*(p) \frac{(\not{p} - \not{p}_b)}{\hat{u}} \Gamma^\mu \right] u_q(p_a) \times \ell_\mu . \quad (6.72)$$

Note que as divergências ocorrem quando:

- O glúon é emitido na mesma direção e sentido do quark , ou seja, quando $\theta_1 \rightarrow 0$ de forma que

$$\begin{aligned} \hat{t} &= -\sqrt{\hat{s}} E_p (1 - \cos \theta_1) \rightarrow 0 \\ \hat{u} &= -\sqrt{\hat{s}} E_p (1 + \cos \theta_1) \rightarrow -2\sqrt{\hat{s}} E_p , \end{aligned}$$

onde $E_p > \delta_s \frac{\sqrt{\hat{s}}}{2}$.

- O glúon é emitido na mesma direção e sentido do antiquark, ou seja, quando $\theta_1 \rightarrow \pi$ de forma que

$$\begin{aligned} \hat{t} &= -\sqrt{\hat{s}} E_p (1 - \cos \theta_1) \rightarrow -2\sqrt{\hat{s}} E_p \\ \hat{u} &= -\sqrt{\hat{s}} E_p (1 + \cos \theta_1) \rightarrow 0 \end{aligned}$$

onde $E_p > \delta_s \frac{\sqrt{\hat{s}}}{2}$.

Seguindo o procedimento sugerido na ref. [51], vamos calcular o caso em que o glúon é emitido colinearmente ao quark , o caso do antiquark será exatamente igual. Em primeiro lugar decompomos a amplitude em duas partes

$$\mathcal{M}_R^{sm} = \mathcal{M}_T^{sm} + \mathcal{M}_U^{sm} , \quad (6.73)$$

onde

$$\mathcal{M}_T^{sm} = C_V \bar{v}_q(p_b) \Gamma^\mu \frac{(\not{p}_a - \not{p})}{\hat{t}} \not{\epsilon}^*(p) u_q(p_a) \ell_\mu \quad (6.74)$$

$$\mathcal{M}_U^{sm} = C_V \bar{v}_q(p_b) \not{\epsilon}^*(p) \frac{(\not{p} - \not{p}_b)}{\hat{u}} \Gamma^\mu u_q(p_a) \ell_\mu . \quad (6.75)$$

Quando $\hat{t} \rightarrow 0$, os termos singulares em $|\mathcal{M}_R^{sm}|^2$ são $|\mathcal{M}_T|^2$ e $\mathcal{M}_T^\dagger \mathcal{M}_U + \mathcal{M}_T \mathcal{M}_U^\dagger = 2\Re(\mathcal{M}_T^\dagger \mathcal{M}_U)$. A aproximação a que chamamos *Leading Pole Approximation* (LPA) consiste em manter apenas esses termos em $|\mathcal{M}_R^{sm}|^2$ usando cinemática colinear. Para facilitar a álgebra, reescrevemos $\mathcal{M}_T$ e $\mathcal{M}_U$ da seguinte forma

$$\mathcal{M}_T^{sm} = \epsilon_\alpha^*(p) \mathcal{M} \frac{(\not{p}_a - \not{p})}{\hat{t}} \gamma^\alpha u_q(p_a) , \quad (6.76)$$

onde $\mathcal{M} = C_V \ell_\mu \bar{v}_q(p_b) \Gamma^\mu$ e ℓ_μ está definida em (6.35), e

$$\mathcal{M}_U^{sm} = \epsilon_\alpha^*(p) \mathcal{N}^\alpha u_q(p_a) , \quad (6.77)$$

onde

$$\mathcal{N}^\alpha = C_V \ell_\mu \bar{v}_q(p_b) \Gamma^\alpha \frac{(\not{p} - \not{p}_b)}{\hat{u}} \Gamma^\mu . \quad (6.78)$$

No limite colinear, o 4-momento do glúon pode ser escrito como

$$p = (1 - z)p_a , \quad (6.79)$$

onde desprezamos as componentes transversais do momento e z representa a quantidade de energia e momento longitudinal que o quark leva para a interação dura após emitir um glúon de momento $(1 - z)p_a$.

Usando essa aproximação e que $\gamma_\alpha \not{p}_a \gamma^\alpha = -2(1 - \varepsilon) \not{p}_a$, temos

$$|\mathcal{M}_T^{sm}|^2 = -2(1 - \varepsilon) \frac{1}{\hat{t}} (1 - z) \text{Tr}[\not{p}_a \mathcal{M}^\dagger \mathcal{M}] , \quad (6.80)$$

onde a soma sobre os spins da partículas já foi realizada.

O traço $\text{Tr}[\not{p}_a \mathcal{M}^\dagger \mathcal{M}]$ é dado por

$$\text{Tr}[\not{p}_a \mathcal{M}^\dagger \mathcal{M}] = g_s^2 C_F \frac{1}{k_V^4} \text{Tr}[\not{p}_a \Gamma_{q,\mu} \not{p}_b \Gamma_{q,\nu}] \text{Tr}[k_1 \Gamma_e^\mu k_2 \Gamma_e^\nu] , \quad (6.81)$$

que corresponde exatamente ao traço envolvido na amplitude de Born.

Para o termo de interferência obtemos

$$2\mathcal{M}_T^{sm\dagger} \mathcal{M}_U^{sm} = -\frac{4z}{\hat{t}} \text{Tr}[\not{p}_a \mathcal{M}^\dagger \mathcal{N} \cdot p_a] . \quad (6.82)$$

Agora é preciso calcular o produto escalar $\mathcal{N} \cdot p_a$. Para isso vamos explorar a invariância de gauge da amplitude $\mathcal{M}_R^{sm}$ sabendo que

$$p_\alpha \mathcal{M}_R^\alpha = 0 , \quad (6.83)$$

onde

$$\mathcal{M}_R^\alpha = \mathcal{M}_T^\alpha + \mathcal{M}_U^\alpha = \left[\mathcal{M} \frac{\not{p}_a - \not{p}}{\hat{t}} \gamma^\alpha + \mathcal{N}^\alpha \right] \cdot u_q(p_a) , \quad (6.84)$$

logo

$$\begin{aligned} p \cdot \mathcal{M}_R &= \left[\frac{1}{\hat{t}} \mathcal{M} \not{p}_a \not{p} + p \cdot \mathcal{N} \right] \cdot u_q(p_a) \\ &= -\frac{\mathcal{M}}{\hat{t}} \not{p} \cdot \not{p}_a u_q(p_a) + \left[\frac{2p_a \cdot p}{\hat{t}} \mathcal{M} + p \cdot \mathcal{N} \right] \cdot u_q(p_a) \\ &= [-\mathcal{M} + (1-z)\mathcal{N} \cdot p_a] \cdot u_q(p_a) = 0 , \end{aligned} \quad (6.85)$$

de onde obtemos

$$\mathcal{N} \cdot p_a = \frac{\mathcal{M}}{1-z} \quad (6.86)$$

e que substituído em (6.82), nos dá para o termo de interferência o seguinte resultado

$$2\mathcal{M}_T^{sm\dagger} \mathcal{M}_U^{sm} = \frac{4z}{(1-z)(-\hat{t})} \text{Tr}[\not{p}_a \mathcal{M}^\dagger \mathcal{M}] . \quad (6.87)$$

Dessa forma, obtemos ao somar (6.80) e (6.82)

$$|\overline{\mathcal{M}_R^{sm}}|_{col}^2 = 4\pi\alpha_s \mu^{2\varepsilon} \frac{1}{|\hat{t}|} 2P_{qq}(z, \varepsilon) |\overline{\mathcal{M}_0^{sm}}|^2 , \quad (6.88)$$

onde

$$P_{qq}(z, \varepsilon) = C_F \left[\frac{1+z^2}{1-z} - \varepsilon(1-z) \right] . \quad (6.89)$$

é uma das funções de splitting de Altarelli-Parisi em $d = 4 - 2\varepsilon$ dimensões.

Quando o glúon sai colinear ao antiquark, substituímos $|\hat{t}|$ por $|\hat{u}|$ na expressão acima.

6.4 Glúon no Estado Inicial

Os diagramas de Feynman para o subprocesso

$$g(p_a) + \bar{q}(p_b) \rightarrow \bar{q}(p) + e^-(k_1) + e^+(k_2) \quad (6.90)$$

são os ilustrados abaixo

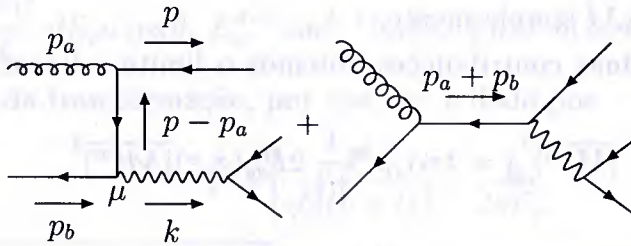


Figura 6.6: Diagramas de Feynman para glúon inicial.

Para o caso $g(p_a) + q(p_b) \rightarrow q(p) + e^-(k_1) + e^+(k_2)$, os diagramas de Feynman são os mesmos apenas trocando o sentido da linha fermiônica e todos os cálculos desenvolvidos a seguir são exatamente análogos.

A amplitude correspondente aos diagramas de Feynman da figura (6.6) é dada por

$$\mathcal{M}_I^{sm} = \frac{-ig_s T^a}{k_V^2} \bar{v}_q(p_b) \left[\Gamma^\mu \frac{(\not{p} - \not{p}_a)}{\hat{t}} \not{\epsilon}(p_a) + \not{\epsilon}(p_a) \frac{(\not{p}_a + \not{p}_b)}{\hat{s}} \Gamma^\mu \right] v_q(p) \times \ell_\mu . \quad (6.91)$$

Como no caso anterior, vamos desmembrar $\mathcal{M}_I^{sm}$

$$\mathcal{M}_I^{sm} = \mathcal{M}_T^i + \mathcal{M}_S^i , \quad (6.92)$$

onde

$$\mathcal{M}_T^i = \epsilon_\alpha(p_a) \mathcal{M} \frac{(\not{p} - \not{p}_a)}{\hat{t}} \gamma^\alpha v_q(p) \quad (6.93)$$

$$\mathcal{M}_S^i = \epsilon_\alpha(p_a) \mathcal{N}^\alpha v_q(p) . \quad (6.94)$$

Nesse processo, a divergência está associada ao limite $\hat{t} \rightarrow 0$ quando o antiquark é espalhado colinearmente ao glúon inicial (que participa do processo duro) e ao limite $\hat{u} \rightarrow 0$ quando o quark é espalhado colinearmente ao glúon inicial.

No caso em que o antiquark sai colinear ao glúon, usando a *leading pole approximation* temos

$$|\mathcal{M}_I^{sm}|^2 \rightarrow |\mathcal{M}_I^{sm}|_{col}^2 = |\mathcal{M}_T^i|^2 + 2\Re(\mathcal{M}_T^{i\dagger} \mathcal{M}_S^i) . \quad (6.95)$$

Seguindo exatamente os mesmos passos do caso da emissão de glúon colinear obtemos, após soma sobre polarizações, spins e cores e tomando a média sobre os números quânticos das partículas iniciais, o resultado

$$\overline{|\mathcal{M}_T^i|^2} = \frac{1}{2(1-\varepsilon)} \frac{1}{2} \frac{(1-\varepsilon)}{|\hat{t}|} 4\pi\alpha_s \mu^{2\varepsilon} |\overline{\mathcal{M}_0}|^2 , \quad (6.96)$$

$$2\overline{\mathcal{M}_T^{i\dagger} \mathcal{M}_S^i} = \frac{2}{2(1-\varepsilon)} \frac{1}{2} \frac{z(1-z)}{-|\hat{t}|} |\overline{\mathcal{M}_0}|^2 ; \quad (6.97)$$

nesse caso, $\mathcal{N} \cdot p_a = \mathcal{M}$ simplesmente.

Somando essas duas contribuições obtemos o limite colinear da amplitude do processo

$$|\overline{\mathcal{M}_I^{sm}}|_{col}^2 = 4\pi\alpha_s \mu^{2\varepsilon} \frac{1}{|\hat{t}|} 2P_{qg}(z, \varepsilon) |\overline{\mathcal{M}_0^{sm}}|^2 , \quad (6.98)$$

onde

$$P_{qg}(z, \varepsilon) = \frac{1}{2(1-\varepsilon)} [z^2 + (1-z)^2 - \varepsilon] \quad (6.99)$$

é uma função de splitting de Altarelli-Parisi em $d = 4 - 2\varepsilon$ dimensões.

Quando o quark sai colinear ao glúon, basta trocar $|\hat{t}|$ por $|\hat{u}|$ na expressão acima.

Tendo em mãos as amplitudes no limite colinear, é preciso agora obter o espaço de fase de 3 partículas nesse limite.

Espaço de fase de 3 partículas na aproximação colinear

Partindo da expressão (C.15) e usando a identidade

$$1 = \int dk^2 \int \frac{d^{d-1}k}{2E_k} \delta^d(k - k_1 - k_2) \quad (6.100)$$

é possível introduzir um sistema intermediário transformando $d^d\Phi_3$ em um produto de 2 espaços de fase de 2 partículas, um descrevendo a produção conjunta de um párton e um bóson vetorial ($V = \gamma, Z$) e outro correspondendo ao processo de decaimento $V \rightarrow$ léptons.*

Dessa maneira é fácil mostrar que

$$d^d\Phi_3 = (2\pi)^{3-3d} dM^2 \frac{d^{d-1}p}{2E_p} \delta[(p_a + p_b - p)^2 - M^2] d^d\Phi_2, \quad (6.101)$$

onde $k^2 = M^2$ e p é o momento do párton no estado final. Inserindo (6.65) na equação (6.101) temos

$$d^d\Phi_3 = \frac{(2\pi)^d}{2} dM^2 dE_p E_p^{d-3} d\Omega_{d-1} \delta(\hat{s} + \hat{t} + \hat{u} - M^2) d^d\Phi_2. \quad (6.102)$$

Como as divergências colineares manifestam-se em $\hat{t}$ e $\hat{u}$ no limite em que se anulam, é coerente mudar as variáveis de integração para $(|\hat{t}|, |\hat{u}|)$, que são as variáveis contidas nas expressões das amplitudes. É mais simples efetuar essas mudanças a partir das variáveis $(E_p, \cos\theta_1)$ dado que

$$|\hat{t}| = E_p \sqrt{\hat{s}} (1 - \cos\theta_1) \quad (6.103)$$

$$|\hat{u}| = E_p \sqrt{\hat{s}} (1 + \cos\theta_1). \quad (6.104)$$

O espaço de fase que queremos transformar é dado por

$$d^d\Phi_3 = -\frac{(2\pi)^d}{2} dE_p d\cos\theta_1 E_p^{d-3} \sin^{d-4}\theta_1 d\Omega_{d-2} dM^2 \delta(\hat{s} + \hat{t} + \hat{u} - M^2) d^d\Phi_2, \quad (6.105)$$

e o jacobiano da transformação, por sua vez, é dado por

$$J = \left| \frac{1}{\sqrt{\hat{s}}(\hat{u} + \hat{t})} \right| = \frac{1}{2\hat{s}E_p}. \quad (6.106)$$

*Ver apêndice C para maiores detalhes a respeito desse procedimento, o qual será amplamente utilizado na implementação numérica das seções de choque de 3 corpos.

Logo, integrando em dM^2 , e usando as equações (6.103) até (6.106) obtemos

$$d^d\Phi_3 = (2\pi)^{3-3d} \frac{d\hat{t}d\hat{u}}{\hat{s}} \left(\frac{\hat{t}\hat{u}}{\hat{s}} \right)^{-\varepsilon} \frac{\pi}{2} d\Omega_{d-3} d^d\Phi_2, \quad (6.107)$$

onde já integramos sobre $d\theta_{d-2} \equiv d\phi$ no intervalo $[0, 2\pi]$.

A aproximação colinear consiste na introdução da variável z através das expressões

$$p = (1-z)p_a \quad (6.108)$$

$$\hat{u} = -2p_b \cdot p = -(1-z)\hat{s}, \quad (6.109)$$

antecipando $\hat{t} \rightarrow 0$

$$p = (1-z)p_b \quad (6.110)$$

$$\hat{t} = -2p_a \cdot p = -(1-z)\hat{s}, \quad (6.111)$$

antecipando $\hat{u} \rightarrow 0$.

A região onde p se torna colinear a p_a ou p_b é definida por

$$|\hat{t}| < \delta_c \hat{s} \quad (6.112)$$

$$|\hat{u}| < \delta_c \hat{s}, \quad (6.113)$$

sendo que $\delta_c \ll 1$.

Para o caso em que $\hat{t} \rightarrow 0$, temos, usando as equações (6.107) até (6.113), e já integrando nas variáveis angulares triviais [52]

$$d^d\Phi_{3,col} = \frac{(4\pi)^\varepsilon}{16\pi^2\Gamma(1-\varepsilon)} d|\hat{t}|dz [(1-z)|\hat{t}|]^{-\varepsilon} d^d\Phi_2. \quad (6.114)$$

A seção de choque diferencial colinear é dada então por

$$d\hat{\sigma}_{col} = \frac{1}{\mathcal{F}} d^d\Phi_{3,col} |\overline{\mathcal{M}}|_{col}^2 \quad (6.115)$$

$$= \frac{(4\pi)^\varepsilon}{16\pi^2\Gamma(1-\varepsilon)} d|\hat{t}|dz z [(1-z)|\hat{t}|]^{-\varepsilon} d\hat{\sigma}_0(z\hat{s}). \quad (6.116)$$

Note que o fator de fluxo que efetivamente está envolvido no termo $d\hat{\sigma}_0$ é

$$\begin{aligned} \mathcal{F}' &= 4\sqrt{(zp_a \cdot p_b)^2} \\ &= 2z(2p_a \cdot p_b) = 2z\hat{s}. \end{aligned}$$

6.5 Seções de Choque Colineares

Obter as seções de choque colineares se reduz, agora, a integrar as amplitudes colineares obtidas na subseção (6.3.2) e na seção (6.4) usando o espaço de fase colinear que acabamos de calcular. No entanto ainda é preciso deduzir os extremos de integração da integral em z , pois na região colinear já sabemos que $[0, \delta_c \hat{s}]$ é o intervalo em que $|\hat{t}|, |\hat{u}|$ variam. Veremos que a variação em z está intimamente ligada às variáveis do modelo a párons. Vamos, portanto, relembrar algumas definições:

$$p_a = x_a P_a , \quad (6.117)$$

$$p_b = x_b P_b , \quad (6.118)$$

onde P_a e P_b são os 4-momentos do próton e do antipróton ou vice-versa.

Além disso introduzimos a variável adimensional

$$\tau \equiv \frac{\hat{s}}{S} , \quad (6.119)$$

onde $\hat{s} = 2p_a \cdot p_b$ e $S = 2P_a \cdot P_b$. Com isso deduzimos que $\tau = x_a x_b$.

O valor mínimo de $\hat{s}$ é M^2 , pois essa é a energia mínima necessária para a produção do par leptônico e^+e^- , logo, o valor mínimo de τ é

$$\tau_0 = \frac{M^2}{S} . \quad (6.120)$$

Vamos definir também a variável z do modelo a párons

$$z \equiv \frac{M^2}{\hat{s}} = \frac{\tau_0}{x_a x_b} . \quad (6.121)$$

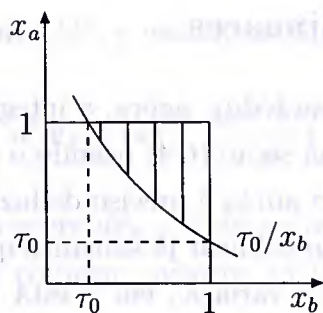
As variáveis x_a e x_b só podem assumir valores reais de 0 até 1, já que representam a fração de energia que um quark ou antiquark carrega do seu hádron de origem para a colisão elementar dura. Dessa forma, como $M^2 \leq \hat{s} \leq S$, então obtemos para τ os limites $\tau_0 \leq \tau \leq 1$, e, por outro lado, como ainda existe a igualdade $\tau = x_a x_b$, temos

$$\tau_0 \leq x_a x_b \leq 1 . \quad (6.122)$$

Apesar de x_a e x_b poderem, a princípio, assumir o valor 0, existe um limite inferior para o produto $x_a x_b$ para que a reação ocorra. Dessa maneira, para um dado x_b fixo, temos para x_a , por exemplo, os limites

$$\frac{\tau_0}{x_b} \leq x_a \leq 1 . \quad (6.123)$$

Graficamente temos a seguinte situação


 Figura 6.7: Região de integração no plano $x_a \otimes x_b$.

ou seja, $\tau_0 \leq x_b \leq 1$.

A variação em z é obtida a partir da variação de x_a para um dado x_b fixo, de forma que obtemos

$$\frac{\tau_0}{x_b} \leq z \leq 1. \quad (6.124)$$

Por último veja que, no regime colinear do processo $q\bar{q} \rightarrow ge^+e^-$, por exemplo, temos

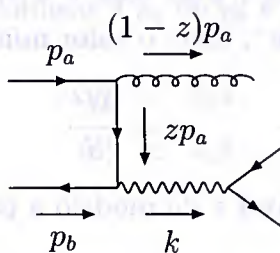


Figura 6.8: Emissão de glúon colinear.

Como há conservação de momento nos vértices, podemos escrever

$$\begin{aligned} zp_a + p_b &= k = k_1 + k_2 \\ (zp_a + p_b)^2 &= k^2 = M^2 \\ z\hat{s} &= M^2 \longrightarrow z = \frac{M^2}{\hat{s}}, \end{aligned}$$

ou seja, o z que usamos como parâmetro para definir a quantidade de momento que o quark ou o antiquark levam para a colisão dura após emitir um glúon colinear é o mesmo z definido no modelo a pártons!

Visto que já subtraímos a região soft do espaço de fase, a variável z , na realidade, varia de acordo com

$$\frac{\tau_0}{x} \leq z \leq 1 - \delta_s. \quad (6.125)$$

Após integração na região colinear (6.112), obtemos, finalmente as seções de choque colineares

$$\hat{\sigma}_{q \rightarrow q}(x) = \int_{\frac{\tau_0}{x}}^{1-\delta_s} dz \left\{ \delta(1-z) + \frac{\alpha_s}{2\pi} \left[\left(-\frac{1}{\varepsilon} + \ln\left(\frac{Q_F^2}{\mu^2}\right) \right) P_{qq}(z) + \Delta_{qq}(z) \right] \right\} \times \hat{\sigma}_0(z\hat{s}) \quad (6.126)$$

para os subprocessos de emissão de glúon, e

$$\hat{\sigma}_{q \rightarrow g}(x) = \frac{\alpha_s}{2\pi} \int_{\frac{\tau_0}{x}}^1 dz \left[\left(-\frac{1}{\varepsilon} + \ln\left(\frac{Q_F^2}{\mu^2}\right) \right) P_{qg}(z) + \Delta_{qg}(z) \right] \hat{\sigma}_0(z\hat{s}) \quad (6.127)$$

para os subprocessos que contém glúon no estado inicial, e onde Q_F define a escala onde será efetuada a fatorização das divergências colineares, antecipando a discussão da seção (6.6). Lembrando que os pártons do estado inicial carregavam uma fração x do momento do seu hádron de origem.

As funções finitas Δ são da forma

$$\Delta_{qq} = \ln\left(\frac{1-z}{z} \delta_c \frac{\hat{s}}{Q_F^2}\right) P_{qq}(z) - P'_{qq}(z) , \quad (6.128)$$

$$\Delta_{qg}(z) = \ln\left(\frac{1-z}{z} \delta_c \frac{\hat{s}}{Q_F^2}\right) P_{qg}(z) - P'_{qg}(z) , \quad (6.129)$$

onde as funções $P'_{qq}(z)$ e $P'_{qg}(z)$ são definidas por

$$P_{qq}(z, \varepsilon) = P_{qq}(z) + \varepsilon P'_{qq}(z) , \quad (6.130)$$

$$P_{qg}(z, \varepsilon) = P_{qg}(z) + \varepsilon P'_{qg}(z) . \quad (6.131)$$

Nas expressões acima x assume, hora o valor x_a , hora o valor x_b .

Veja que em $\hat{\sigma}_{q \rightarrow g}$ a integral sobre a função de splitting tem seu limite superior indo até 1, pois esses processos não apresentam divergências soft. A presença da função $\delta(1-z)$ sob o sinal da integral de $\hat{\sigma}_{q \rightarrow q}(x)$ é inócua, dado que sua contribuição é identicamente nula, no entanto, veremos que sua introdução nos auxiliará no procedimento de fatorização de massa que realizamos a seguir.

6.6 Fatorização de Massa

Em termos de (5.112), as seções de choque colineares são escritas como

$$\hat{\sigma}_{q \rightarrow q} = \int_{\frac{\tau_0}{x}}^{1-\delta_s} dz \left[\Gamma_{qq}^{\overline{MS}}(z, Q_F^2, \mu^2, \varepsilon) + \frac{\alpha_s}{2\pi} \Delta_{qq} \right] \times \hat{\sigma}_0^{sm}(z\hat{s}) \quad (6.132)$$

$$\hat{\sigma}_{q \rightarrow g} = \int_{\frac{\tau_0}{x}}^1 dz \left[\Gamma_{qg}^{\overline{MS}}(z, Q_F^2, \mu^2, \varepsilon) + \frac{\alpha_s}{2\pi} \Delta_{qg} \right] \times \hat{\sigma}_0^{sm}(z\hat{s}) , \quad (6.133)$$

onde Δ_{qq} e Δ_{qg} , dadas por (6.128) e (6.129), respectivamente, representam as funções $f_{ij}(x, Q_F^2, \mu^2)$ de (5.112).

A seção de choque colinear hadrônica total é dada pela convolução das PDFs e seções de choque partônicas colineares

$$\begin{aligned} \sigma_{col} = & \sum_q \int dx_a dx_b \{ q_0(x_a) \hat{\sigma}_{q \rightarrow q}(x_a) \bar{q}_0(x_b) + \bar{q}_0(x_a) \hat{\sigma}_{q \rightarrow q}(x_a) q_0(x_b) \\ & + q_0(x_b) \hat{\sigma}_{q \rightarrow q}(x_b) \bar{q}_0(x_a) + \bar{q}_0(x_b) \hat{\sigma}_{q \rightarrow q}(x_b) q_0(x_a) \} \\ & + \sum_q \int dx_a dx_b \{ [q_0(x_a) + \bar{q}_0(x_a)] g(x_b) \hat{\sigma}_{q \rightarrow g}(x_b) \\ & + [q_0(x_b) + \bar{q}_0(x_b)] g(x_a) \hat{\sigma}_{q \rightarrow g}(x_a) \} . \end{aligned} \quad (6.134)$$

Vamos elucidar o procedimento de fatorização de massa tomando um exemplo específico. Considere a seção de choque hadrônica σ_q definida por

$$\begin{aligned} \sigma_q = & \int dx_a dx_b [q_0(x_a) \hat{\sigma}_{q \rightarrow q}(x_a) \bar{q}_0(x_b) + g_0(x_a) \hat{\sigma}_{q \rightarrow g}(x_a) \bar{q}_0(x_b)] \\ = & \int dx_b \bar{q}_0(x_b) \int dx_a [q_0(x_a) \hat{\sigma}_{qq}(x_a) + g_0(x_a) \hat{\sigma}_{qg}(x_a)] . \end{aligned} \quad (6.135)$$

Agora substituímos $\hat{\sigma}_{q \rightarrow q}(x_a)$ e $\hat{\sigma}_{q \rightarrow g}(x_a)$ por suas expressões dadas por (6.126) e (6.127), respectivamente, de forma que

$$\begin{aligned} \sigma_q = & \int dx_b \bar{q}_0(x_b) \int dx_a \left[q_0(x_a) \int_{\frac{\tau_0}{x_a}}^{1-\delta_s} dz \left(\Gamma_{qq}^{\overline{MS}}(z, Q_F^2, \varepsilon) + \frac{\alpha_s}{2\pi} \Delta_{qq}(z) \right) \right. \\ & \left. + g_0(x_a) \int_{\frac{\tau_0}{x_a}}^1 dz \left(\Gamma_{qg}^{\overline{MS}}(z, Q_F^2, \varepsilon) + \frac{\alpha_s}{2\pi} \Delta_{qg}(z) \right) \right] \times \hat{\sigma}_0^{sm}(z\hat{s}) . \end{aligned} \quad (6.136)$$

Antes de mais nada, precisamos reescrever a integral em $\Gamma_{qq}^{\overline{MS}}(z, Q_F^2, \varepsilon)$ cujo limite superior vale $1 - \delta_s$ a partir da identidade

$$\int_{\frac{\tau_0}{x_a}}^{1-\delta_s} dz F(z) = \int_{\frac{\tau_0}{x_a}}^1 dz F(z) - \int_{1-\delta_s}^1 dz F(z) ,$$

de modo a combinar os termos das seções de choque elementares e os termos provenientes das *PDFs*. Em seguida é preciso mostrar que as seguintes integrais duplas são equivalentes:

$$\int_{\tau_0}^1 dx f(x) \int_{\frac{\tau_0}{x}}^1 dz P(z) , \quad (6.137)$$

proveniente das seções de choque elementares - o limite superior é $1 - \delta_s$ para os termos de emissão de glúon, e

$$\int_{\tau_0}^1 dx \int_x^1 \frac{dy}{y} f(y) P(x/y) , \quad (6.138)$$

proveniente das *PDFs*. Isso é realizado simplesmente escrevendo essa segunda integral na forma

$$\int_{\tau_0}^1 dx \int_0^1 dx_1 \int_0^1 dx_2 f(x_1) P(x_2) \delta(x - x_1 x_2) . \quad (6.139)$$

A função $\delta(x - x_1x_2)$ força os extremos de integração a adquirirem a mesma forma daqueles da integral (6.137). Depois, basta integrar em x para mostrar a equivalência. O mesmo pode ser mostrado quando o extremo superior é $1 - \delta_s$.

Com isso, σ_q assume a seguinte forma

$$\begin{aligned} \sigma_q = & \int dx_b \bar{q}_0(x_b) \int_{\tau_0}^1 dx_a \left\{ \int_{x_a}^1 \frac{dy}{y} [q_0(y) \Gamma_{qq}^{\overline{MS}}(z, Q_F^2, \varepsilon) + g_0(y) \Gamma_{qg}^{\overline{MS}}(z, Q_F^2, \varepsilon)] \right. \\ & + \frac{\alpha_s}{2\pi} \left[\int_{x_a}^{1-\delta_s} \frac{dy}{y} q_0(y) \Delta_{qq}(z) + \int_{x_a}^1 \frac{dy}{y} g_0(y) \Delta_{qg}(z) \right] \\ & \left. - \int_{1-\delta_s}^1 dz q_0(x_a) \Gamma_{qq}^{\overline{MS}}(z, Q_F^2, \varepsilon) \right\} \times \hat{\sigma}_0^{sm}(z\hat{s}) . \end{aligned} \quad (6.140)$$

Note, agora, que a primeira linha de (6.140) contém precisamente a expressão da densidade de quarks “renormalizada” $q(x, Q_F^2)$ dada por (5.118), a menos de termos finitos! Basta, portanto redefinirmos a primeira linha de (6.140) de maneira a absorver a divergência colinear, contida nas funções Γ_{ij} , em $q(x, Q_F^2)$. Isso, contudo, pode ser feito de uma forma arbitrária, de modo que podemos escolher quais termos finitos serão absorvidos junto com a divergência na escala de fatorização Q_F como discutimos na seção (5.4).

Após a fatorização de massa σ_q tem a seguinte forma

$$\begin{aligned} \sigma_q = & \int dx_b \bar{q}_0(x_b) \int_{\tau_0}^1 dx_a \left\{ q(x_a, Q_F^2) + \frac{\alpha_s}{2\pi} \left[\int_{x_a}^{1-\delta_s} \frac{dy}{y} q(y, Q_F^2) \Delta_{qq}(z) \right. \right. \\ & + \left. \int_{x_a}^1 \frac{dy}{y} g(y, Q_F^2) \Delta_{qg}(z) - \lambda_{FC} \int_{x_a}^1 \frac{dy}{y} [q(y, Q_F^2) F_q(z) + g(y, Q_F^2) F_g(z)] \right] \\ & \left. - \int_{1-\delta_s}^1 dz \Gamma_{qq}^{\overline{MS}}(z, Q_F^2, \varepsilon) \cdot q(x_a, Q_F^2) \right\} \times \hat{\sigma}_0^{sm}(z\hat{s}) . \end{aligned} \quad (6.141)$$

Usando novamente a identidade

$$\int_{\frac{\tau_0}{x_a}}^{1-\delta_s} dz F(z) = \int_{\frac{\tau_0}{x_a}}^1 dz F(z) - \int_{1-\delta_s}^1 dz F(z)$$

para reescrever a integral

$$\lambda_{FC} \int_{x_a}^1 \frac{dy}{y} q(y, Q_F^2) F_q(z)$$

e substituindo a forma explícita de $\Gamma_{qq}^{\overline{MS}}(z, Q_F^2, \varepsilon)$ em (6.141) acima, obtemos, finalmente

$$\sigma_q = \sigma_{HC}^q + \sigma_{SC}^q , \quad (6.142)$$

onde

$$\begin{aligned} \sigma_{HC}^q = & \frac{\alpha_s}{2\pi} \int dx_b \bar{q}(x_b, Q_F^2) \int_{\tau_0}^1 dx_a \left[\int_{x_a}^{1-\delta_s} \frac{dy}{y} q(y, Q_F^2) \tilde{P}_{qq}(z) \right. \\ & \left. + \int_{x_a}^1 \frac{dy}{y} g(y, Q_F^2) \tilde{P}_{qg}(z) \right] \times \hat{\sigma}_0^{sm}(z\hat{s}) \end{aligned} \quad (6.143)$$

$$\sigma_{SC}^q = -\frac{\alpha_s}{2\pi} \int dx_b \bar{q}(x_b, Q_F^2) \int_{\tau_0}^1 dx_a \int_{1-\delta_s}^1 \frac{dy}{y} q(y, Q_F^2) \left[\left(-\frac{1}{\bar{\epsilon}} + \ln\left(\frac{Q_F^2}{\mu^2}\right) \right) P_{qq}(z) + \lambda_{FC} \cdot F_{qq}(z) \right] \times \hat{\sigma}_0^{sm}(z\hat{s}) \quad (6.144)$$

onde σ_{HC}^q contribui à seção de choque hard-colinear que será apresentada na próxima, e σ_{SC}^q contribui para a seção de choque soft-colinear que será calculada logo em seguida.

6.6.1 Termo Hard-Colinear

Após a fatorização de massa, as partes finitas provenientes da soma de todas as seções de choque σ_{HC}^q que contribuem para (6.134) são absorvidas na seção de choque Hard-Colinear dada por

$$\begin{aligned} \sigma_{HC} = & \frac{\alpha_s}{2\pi} \sum_q \int dx_a dx_b \times \\ & \left\{ q(x_a, Q_F^2) \int_{x_b}^{1-\delta_s} \frac{dz}{z} \bar{q}\left(\frac{x_b}{z}, Q_F^2\right) \tilde{P}_{qg \leftarrow q}(z) + q(x_a, Q_F^2) \int_{x_b}^1 \frac{dz}{z} g\left(\frac{x_b}{z}, Q_F^2\right) \tilde{P}_{q\bar{q} \leftarrow g}(z) + \right. \\ & \bar{q}(x_b, Q_F^2) \int_{x_a}^{1-\delta_s} \frac{dz}{z} q\left(\frac{x_a}{z}, Q_F^2\right) \tilde{P}_{qg \leftarrow q}(z) + \bar{q}(x_b, Q_F^2) \int_{x_a}^1 \frac{dz}{z} g\left(\frac{x_a}{z}, Q_F^2\right) \tilde{P}_{q\bar{q} \leftarrow g}(z) \\ & \left. + a \leftrightarrow b \right\} \times \hat{\sigma}_0^{sm}(z\hat{s}) \quad , \quad (6.145) \end{aligned}$$

onde

$$\tilde{P}_{qq}(z) = P_{qq}(z) \ln\left(\frac{1-z}{z} \delta_c \frac{\hat{s}}{Q_F^2}\right) - P'_{qq}(z) - \lambda_{FC} F_{qq}(z) \quad , \quad (6.146)$$

$$\tilde{P}_{qg}(z) = P_{qg}(z) \ln\left(\frac{1-z}{z} \delta_c \frac{\hat{s}}{Q_F^2}\right) - P'_{qg}(z) - \lambda_{FC} F_{qg}(z) \quad , \quad (6.147)$$

sendo que

$$P_{qq}(z) = C_F \frac{1+z^2}{1-z} \quad , \quad (6.148)$$

$$P'_{qq}(z) = -C_F(1-z) \quad , \quad (6.149)$$

$$F_q(z) = C_F \left\{ \frac{1+z^2}{1-z} \ln\left(\frac{1-z}{z}\right) - \frac{3}{2} \frac{1}{1-z} + 3 + 2z \right\} \quad , \quad (6.150)$$

$$P_{qg}(z) = T_R [z^2 + (1-z)^2] \quad , \quad (6.151)$$

$$P'_{qg}(z) = T_R [z^2 + (1-z)^2 - 1] \quad , \quad (6.152)$$

$$F_g(z) = T_R \left\{ [z^2 + (1-z)^2] \ln\left(\frac{1-z}{z}\right) + 8z(1-z) - 1 \right\} \quad . \quad (6.153)$$

Essas expressões mostram o caráter universal das correções em NLO de QCD no contexto do método de cálculo utilizado e concorda com os resultados das refs. [31, 34, 35]. A seguir tratamos do termo soft-colinear.

6.6.2 Termo Soft-Colinear

Pelo cálculo das seções de choque colineares σ_q fomos capazes, através do procedimento de fatorização, de eliminar as divergências colineares da seção de choque total hadrônica.

Contudo, desse procedimento restou um termo que contém o limite $z \rightarrow 1$, correspondente à superposição das região soft e colinear do espaço de fase. Vamos denominar essa contribuição à seção de choque de σ_{SC} e verificar que a divergência colinear remanescente se cancela com o pólo remanescente na expressão de $\hat{\sigma}_{virt} + \hat{\sigma}_{soft}$, equação (6.69).

O ponto de partida para o cálculo de σ_{SC} é soma de todas as contribuições σ_{SC}^q

$$\begin{aligned} \hat{\sigma}_{SC} = & -\frac{\alpha_s}{2\pi} \left\{ q(x_a, Q_F^2) \int_{1-\delta_s}^1 \frac{dz}{z} \bar{q}\left(\frac{x_b}{z}, Q_F^2\right) \left(-\frac{1}{\bar{\epsilon}} P_{qq}(z) + P_{qq}(z) \ln \frac{Q_F^2}{\mu^2} + \lambda_{FC} F_{qq}(z) \right) \right. \\ & + \bar{q}(x_b, Q_F^2) \int_{1-\delta_s}^1 \frac{dz}{z} q\left(\frac{x_a}{z}, Q_F^2\right) \left(-\frac{1}{\bar{\epsilon}} P_{qq}(z) + P_{qq}(z) \ln \frac{Q_F^2}{\mu^2} + \lambda_{FC} F_{qq}(z) \right) \\ & \left. + a \leftrightarrow b \right\} \times \hat{\sigma}_0^{sm}(z\hat{s}) , \end{aligned} \quad (6.154)$$

onde a função de splitting, nesse caso, pode ser escrita como

$$P_{qq}(z) = C_F \left(\frac{1+z^2}{(1-z)_+} + \frac{3}{2} \delta(1-z) \right) . \quad (6.155)$$

Para $\delta_s \ll 1$ podemos aproximar (6.154) por

$$\begin{aligned} \hat{\sigma}_{SC} \approx & -\frac{\alpha_s}{2\pi} [q(x_a, Q_F^2) \bar{q}(x_b, Q_F^2) + q(x_b, Q_F^2) \bar{q}(x_a, Q_F^2)] \\ & \times 2 \int_{1-\delta_s}^1 dz \left(-\frac{1}{\bar{\epsilon}} P_{qq}(z) + P_{qq}(z) \ln \frac{Q_F^2}{\mu^2} + \lambda_{FC} F_{qq}(z) \right) \times \hat{\sigma}_0^{sm}(\hat{s}) , \end{aligned} \quad (6.156)$$

onde fizemos $z = 1$ em todos os lugares onde era possível.

Usando a seguinte definição equivalente das funções “+”

$$[F(x)]_+ \equiv \lim_{\beta \rightarrow 0} [F(x) \theta(1-x-\beta) - \delta(1-x-\beta) \int_0^{1-\beta} F(y) dy] , \quad (6.157)$$

podemos mostrar as identidades

$$\int_{\xi}^1 dz \frac{g(z)}{(1-z)_+} = \int_{\xi}^1 dz \frac{g(z) - g(1)}{1-z} + g(1) \ln(1-\xi) \quad (6.158)$$

$$\int_{\xi}^1 dz g(z) \left(\frac{\ln(1-z)}{(1-z)_+} \right) = \int_{\xi}^1 dz \frac{g(z) - g(1)}{1-z} \ln(1-z) + \frac{1}{2} g(1) \ln^2(1-\xi) , \quad (6.159)$$

que nos permitem calcular as integrais envolvidas na expressão (6.156), desde que $g(z)$ seja uma função bem comportada no limite $z \rightarrow 1$. O resultado final, integrado

nas frações de momento x_a e x_b , é dado por

$$\sigma_{SC} = \sum_q \int dx_a dx_b [q(x_a, Q^2) \bar{q}(x_b, Q^2) + a \leftrightarrow b] \times \hat{\sigma}_0^{sm}(\hat{s}) \times \quad (6.160)$$

$$\frac{\alpha_s}{2\pi} C_F N_\epsilon \left(\frac{\mu^2}{Q_F^2} \right)^\epsilon \left\{ \frac{1}{\epsilon} (3 + 4 \ln(\delta_s)) + 9 + \frac{2\pi^2}{3} + 3 \ln(\delta_s) - 2 \ln^2(\delta_s) \right\}.$$

Todos os termos dependentes de potências do cut-off soft δ_s foram desprezados no cômputo desse resultado.

6.7 Contribuição de 2 Corpos à Seção de Choque Hadrônica

Tendo já calculado todos os termos que contribuem à seção de choque de 2 corpos, temos, finalmente

$$\sigma_2 = \sigma^{LL} + \sigma_{HC} + \sigma_{virt} + \sigma_{soft} + \sigma_{SC}, \quad (6.161)$$

cujas expressões completas são dadas por

$$\begin{aligned} \sigma_2 = & \sum_q \int dx_a dx_b [q(x_a, Q_F^2) \bar{q}(x_b, Q_F^2) + a \leftrightarrow b] \times \hat{\sigma}_0(\hat{s}) \times \\ & \left\{ 1 + \frac{\alpha_s}{2\pi} C_F \left[(3 + 4 \ln(\delta_s)) \ln\left(\frac{\hat{s}}{Q_F^2}\right) + 1 + \frac{5\pi^2}{3} + 3 \ln(\delta_s) + 2 \ln^2(\delta_s) \right] \right\} \\ & + \frac{\alpha_s}{2\pi} \sum_q \int dx_a dx_b \times \hat{\sigma}_0(z\hat{s}) \times \\ & \left\{ q(x_a, Q_F^2) \left(\int_{x_b}^{1-\delta_s} \frac{dz}{z} \bar{q}\left(\frac{x_b}{z}, Q_F^2\right) \tilde{P}_{qq}(z) + \int_{x_b}^1 \frac{dz}{z} g\left(\frac{x_b}{z}, Q_F^2\right) \tilde{P}_{qg}(z) \right) \right. \\ & + \bar{q}(x_b, Q_F^2) \left(\int_{x_a}^{1-\delta_s} \frac{dz}{z} q\left(\frac{x_a}{z}, Q_F^2\right) \tilde{P}_{qq}(z) + \int_{x_a}^1 \frac{dz}{z} g\left(\frac{x_a}{z}, Q_F^2\right) \tilde{P}_{qg}(z) \right) \\ & \left. + a \leftrightarrow b \right\}. \quad (6.162) \end{aligned}$$

Como esperávamos, o pólo remanescente na soma $\hat{\sigma}_{soft} + \hat{\sigma}_{virt}$ cancelou-se com o pólo da contribuição soft-colinear deixando o resultado livre de quaisquer divergências. Contudo, as partes finitas são sensivelmente dependentes dos cut-offs δ_s e δ_c . Essa expressão concorda com a da referência [31], onde correções em NLO de QCD são calculadas à produção hadrônica de $W^\pm$ pelo mesmo método que estamos utilizando. Tendo em mãos todas as contribuições à seção de choque de 2 de corpos que contém pólos IR, a estrutura de divergências pode ser resumida pela tabela (6.1) abaixo

δ	Gauge de Feynman ($\eta = 0$)	Gauge de Landau ($\eta = -1$)
δ_{AE}^q	$\frac{1}{\epsilon}$	0
δ_{vert}^{sm}	$-\frac{2}{\epsilon^2} - \frac{4}{\epsilon}$	$-\frac{2}{\epsilon^2} - \frac{3}{\epsilon}$
δ_{Soft}	$\frac{2}{\epsilon^2} - \frac{4}{\epsilon} \ln \delta_s$	$\frac{2}{\epsilon^2} - \frac{4}{\epsilon} \ln \delta_s$
δ_{SC}	$\frac{1}{\epsilon} (3 + 4 \ln \delta_s)$	$\frac{1}{\epsilon} (3 + 4 \ln \delta_s)$

Tabela 6.1: Estrutura de divergências IR no processo de Drell-Yan $p\bar{p} \rightarrow e^+e^-X$ em NLO de QCD.

As contribuições δ_{AE}^q , δ_{vert} , δ_{soft} e δ_{SC} , correspondem aos pólos originados pelas amplitudes virtuais das autoenergias dos quarks, correção virtual do vértice (γ, Z) $q\bar{q}$, emissão de glúon soft e termo soft-colinear, respectivamente. A soma de todas as contribuições, dada por

$$\Delta^{sm} = \delta_{AE}^q + \delta_{vert} + \delta_{soft} + \delta_{SC} \quad (6.163)$$

é nula, como podemos ver pela tabela (6.1), em ambos os gauges.

6.8 Seção de Choque de 3 Corpos

Nessa seção, apresentamos a contribuição à seção de choque total que provém dos subprocessos $q\bar{q} \rightarrow ge^+e^-$ e $q(\bar{q})g \rightarrow q(\bar{q})e^+e^-$ na região $\overline{HC}$ do espaço de fase, definida por

$$\delta_s \frac{\sqrt{\hat{s}}}{2} < E_g < E_{g,max} , \quad (6.164)$$

$$\delta_e \hat{s} < |\hat{t}|, |\hat{u}| < |\hat{t}_{max}|, |\hat{u}_{max}| , \quad (6.165)$$

onde E_g é energia do glúon no processo de emissão de glúon. Essas contribuições são calculadas em 4 dimensões espaço-temporais, e já não contém divergências a serem reguladas. As amplitudes foram calculadas com o auxílio do pacote FeynCalc e, concordam com os resultados da referência [53].

6.8.1 Glúons Iniciais

O elemento de matriz para esse processo é dado pela expressão (6.91). O quadrado de $\mathcal{M}_V^I$, após somas e médias dos números quânticos relevantes é dado por

$$\begin{aligned} |\overline{\mathcal{M}_I^{sm}}|^2 = & -\frac{8\pi}{3}\alpha_s \sum_V \frac{k^2}{|k_V^2|^2} \left\{ \left[\frac{\hat{t} - 2(\hat{t}_1 + k^2)}{\hat{s}} + \frac{\hat{s} + 2(\hat{t}_1 + \hat{t}_2)}{\hat{t}} \right. \right. \\ & + \frac{2}{\hat{s}\hat{t}} [(\hat{t}_1 + \hat{t}_2 + k^2)^2 + \hat{t}_1^2 - \hat{t}_2 k^2] \Big] d_e d_q \\ & \left. + \left[\frac{2(\hat{t}_1 + k^2) - \hat{t}}{\hat{s}} + \frac{\hat{s} + 2(\hat{t}_1 + \hat{t}_2)}{\hat{t}} - \frac{2(2\hat{t}_1 + \hat{t}_2 + k^2)k^2}{\hat{s}\hat{t}} \right] h_e h_q \right\} . \end{aligned} \quad (6.166)$$

6.8.2 Emissão de Glúon

O elemento de matriz é o da expressão (6.52), de modo que

$$|\overline{\mathcal{M}_R^{sm}}|^2 = -\frac{64\pi}{9}\alpha_s \sum_V \frac{k^2}{|k_V^2|^2} \left\{ \left[1 + \frac{\hat{s} - 2\hat{t}_1 - k^2}{\hat{t}} - \frac{\hat{t}_1^2 + \hat{t}_2^2 + \hat{s}(\hat{t}_1 + \hat{t}_2 + k^2)}{\hat{t}\hat{u}} \right] \right\} \quad (6.167)$$

$$+[\hat{t} \leftrightarrow \hat{u}, \hat{t}_1 \leftrightarrow \hat{t}_2] \left[d_e d_q + \left[\frac{\hat{s} + 2\hat{t}_2 + k^2}{\hat{t}} + \frac{\hat{t}_1 - \hat{t}_2}{\hat{t}\hat{u}} - [\hat{t} \leftrightarrow \hat{u}, \hat{t}_1 \leftrightarrow \hat{t}_2] \right] h_e h_q \right\} .$$

Os invariantes cinemáticos são dados nas equações (6.42) até (6.46) e as constantes d_e, d_q, h_e, h_q são as definidas na seção (6.1).

A amplitude completa é obtida simplesmente somando sobre as contribuições do fóton, do Z e do termo de interferência como fizemos com a amplitude de Born na seção (6.1).

No apêndice C mostramos o espaço de fase de 3 partículas que usamos na implementação numérica dos resultados.

Capítulo 7

Introdução de Leptoquarks Escalares

7.1 Extensões do Modelo Padrão com Acoplamentos Quark-Lépton

7.1.1 Motivação de Estudo

Como sabemos, o modelo padrão das interações fortes (QCD) e eletrofracas (modelo Glashow–Salam–Weinberg), é uma teoria quântica de campos invariante sob transformações do grupo de gauge $SU(3)_C \otimes SU(2)_L \otimes U(1)_Y$ cuja simetria do setor eletrofraco é espontaneamente quebrada abaixo de uma determinada escala de energia, sendo, portanto, renormalizável.

Correções radiativas envolvendo loops, no entanto, podem violar leis de conservação clássicas, derivadas à partir da invariância de gauge via teorema de Noether. A esse fenômeno dá-se o nome de *anomalia*. Suas consequências são desastrosas, invalidando a renormalizabilidade da teoria, ao quebrar as identidades de Ward. O cancelamento de anomalias é, nesse contexto, um ingrediente fundamental para o sucesso do programa de renormalização do modelo padrão.

Uma teoria é livre de anomalias se

$$F^{abc} = F_+^{abc} - F_-^{abc} = 0 \quad , \quad (7.1)$$

onde $F_{\pm}^{abc}$ é dado pelo seguinte traço de geradores

$$F_{\pm}^{abc} \equiv \text{Tr}[\{T_{\pm}^a, T_{\pm}^b\}T_{\pm}^c] \quad , \quad (7.2)$$

onde $T_{\pm}^a$ são os geradores na representação de mão direita (+) e mão esquerda (−) dos campos de matéria.

Em uma teoria do tipo $V - A$, como o SM, as únicas anomalias possíveis vêm de loops triangulares do tipo VVA , ou seja, de loops com dois vértices vetoriais e um vértice axial, como ocorre no decaimento pseudoescalar $\pi^0 \rightarrow \gamma\gamma$ ilustrado na figura

abaixo, onde o vértice em forma de cruz denota a corrente axial

$$J_A^\mu = (\bar{u}\gamma^\mu\gamma^5 u - \bar{d}\gamma^\mu\gamma^5 d) .$$

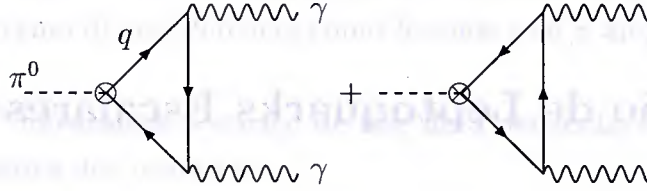


Figura 7.1: Loops triangulares de quarks que geram anomalias.

Nesse caso existem duas possibilidades

$$[SU(2) \otimes SU(2)] \otimes U(1) : \text{Tr}[\{\tau^a, \tau^b\}Y] = \text{Tr}[\{\tau^a, \tau^b\}]Tr[Y] \propto \sum_{\text{doublets}} Y$$

$$U(1) \otimes U(1) \otimes U(1) : \text{Tr}[Y^3] \propto \sum_{\text{férmions}} Y^3 .$$

Para os léptons, os valores das hipercargas Y são -1 para os de mão esquerda e -2 para os de mão direita, enquanto que para os quarks *up* e *down* valem $1/3$ para os de mão esquerda, e $4/3$ para o *up* e $-2/3$ para o *down* de mão direita, respectivamente. Assim sendo, temos para $[SU(2) \otimes SU(2)] \otimes U(1)$

$$F^{abc} \propto -F_-^{abc} = - \sum_{\text{doublets}} Y = - \left[-1 + 3 \left(\frac{1}{3} \right) \right] = 0 ,$$

enquanto que para o caso de $U(1) \otimes U(1) \otimes U(1)$, temos

$$F^{abc} \propto \sum_{\text{férmions}} [Y_+^3 - Y_-^3] = \left\{ (-2)^3 + 3 \left[\left(\frac{4}{3} \right)^3 + \left(\frac{-2}{3} \right)^3 \right] \right\}$$

$$- \left\{ 2 \times (-1)^3 + 3 \times 2 \times \left(\frac{1}{3} \right)^3 \right\} = 0 ,$$

onde as 3 cores dos quarks já foram levadas em conta. Isso mostra que o modelo padrão é livre de anomalias se os férmions aparecem em multipletos completos, de acordo com a estrutura geral*

$$\left\{ \begin{pmatrix} \nu_e \\ e \end{pmatrix}_L, e_R, \begin{pmatrix} u \\ d \end{pmatrix}_L, u_R, d_R \right\} ,$$

*Neutrinos sem massa.

que deve se repetir, sempre respeitando a mesma estrutura

$$\left\{ \left(\begin{pmatrix} \nu_\mu \\ \mu \end{pmatrix}_L, \mu_R, \begin{pmatrix} c \\ s \end{pmatrix}_L, c_R, s_R \right) \right\}.$$

A descoberta do lépton τ em 1975, e do quinto sabor de quark, o *bottom* b , dois anos mais tarde, evidenciou a existência de uma terceira família

$$\left\{ \left(\begin{pmatrix} \nu_\tau \\ \tau \end{pmatrix}_L, \tau_R, \begin{pmatrix} t \\ b \end{pmatrix}_L, t_R, b_R \right) \right\}.$$

A existência de gerações completas de férmions é absolutamente essencial para o cancelamento de anomalias. Isso foi um argumento teórico importante em favor da existência do quark *top* antes mesmo de sua descoberta em 1995 no Fermilab.

Por outro lado, essa aparente simetria entre famílias de quarks e léptons continua sendo um mistério, o que tem motivado a construção de modelos além do SM, visando relacionar aquelas partículas em um nível mais fundamental. Como resultado, muitos destes modelos naturalmente contém partículas que se acoplam diretamente a pares quark-lépton.

Vários têm sido os modelos propostos que relacionam quarks e léptons de maneira mais fundamental, entre eles: os modelos de grande unificação baseados em grupos de gauge mais amplos como o $SU(5)$ [54], o $SU(4)$ do modelo de Pati-Salam [55], modelos de technicolor [56], modelos onde léptons e quarks possuem uma subestrutura [57] e modelos inspirados em *superstrings* [58]. A extensão supersimétrica do modelo padrão com a chamada *violação de paridade- R* , também dá origem a acoplamentos com pares quark-lépton + partícula supersimétrica, e que pode ter sinais experimentais parecidos com o de um leptoquark.

Em todos aqueles modelos, novos campos são introduzidos e têm a propriedade de mediar transições quark-lépton: são os *leptoquarks*. Podem ter spin 0 (escalares) ou 1 (vetoriais).

A maioria dos modelos, contudo, têm sido fortemente vinculados em relação ao que vem sendo estabelecido nos laboratórios e que tem corroborado cada vez mais o modelo padrão gerando vínculos bastante restritivos à existência de leptoquarks.

Em 1997, as colaborações H1 e ZEUS comunicaram ter encontrado um excesso de eventos para grandes valores de Q^2 em seus dados das seções de choque do processo de espalhamento inelástico profundo (DIS) e^+p [59], em relação ao número previsto pelo modelo padrão. Tal excesso, se confirmado, seria um claro sinal de um nível de composição de quarks ou léptons se os dados não apresentassem uma estrutura cinemática específica, ou de uma ressonância, caso contrário. Nesse último caso poderiam ser interpretados como a produção de um leptoquark de massa

aproximadamente igual a 200 GeV. Contudo, devido à baixíssima estatística, não foi possível confirmar tal excesso de eventos, e posteriormente, constatou-se que o excesso era apenas uma flutuação estatística.

Apesar do DIS poder revelar ressonâncias no canal s da interação, existem outros processos que podem igualmente revelar a existência dessas partículas. Várias outras colaborações têm procurado por sinais de produção de leptoquarks em aceleradores de altas energias. Os processos estudados são do tipo

$$q + \bar{q} \rightarrow S + \bar{S} , \quad (7.3)$$

$$g + g \rightarrow S + \bar{S} , \quad (7.4)$$

$$q + g \rightarrow \ell + S , \quad (7.5)$$

$$e^+ + e^- \rightarrow S + \bar{S} , \quad (7.6)$$

Os processos (7.3)–(7.4) e (7.5) são canais de produção dupla e simples no Tevatron respectivamente e, futuramente, no LHC, enquanto (7.6) é um dos canais de produção no LEP e, futuramente, no NLC. Os dados combinados das experiências colocam um limite inferior para a massa de leptoquarks escalares de 1ª geração maior do que 237 GeV [60].

Contudo, análises indiretas são possíveis, como por exemplo, a partir da influência de leptoquarks sobre os parâmetros oblíquos eletrofracos [61], violação de paridade atômica [62] e efeitos virtuais de leptoquarks nos canais de espalhamento t e u em processos como $p\bar{p}, pp \rightarrow \ell\bar{\ell} + X$ e $e^+e^- \rightarrow$ hádrons [63].

Nesse projeto, em especial, investigamos os efeitos virtuais de leptoquarks escalares que se acoplam a pares férmion–antiférmion da primeira geração no processo de Drell-Yan no Tevatron

$$p\bar{p} \rightarrow e^-e^+ + X ,$$

incluindo correções de QCD em ordem α_s . A seguir apresentamos um modelo efetivo para interações de leptoquarks.

7.1.2 Lagrangiana de Interação de Leptoquarks

Se leptoquarks realmente existem, devem estar inseridos em um modelo efetivo que respeita as simetrias e predições do modelo padrão até o presente momento. Dessa forma é preferível que:

- Leptoquarks não se acoplem a pares de quarks e léptons para evitar violação de número bariônico (B) e leptônico (L), caso contrário, suas massas deveriam ser muito grandes. Em modelos de *grande unificação*, $SU(5)$ por exemplo, tais processos são permitidos, ocasionando o decaimento do próton. No $SU(5)$

existem 12 tipos diferentes de leptoquarks, que denominamos bósons X e Y . A lagrangiana de interação que envolve X e Y e a primeira geração de quarks e léptons é dada por

$$\begin{aligned}\mathcal{L}_{X,Y} = & i\frac{g_X}{\sqrt{2}}X_\mu^i[\varepsilon_{ijk}\bar{u}_{kL}^c\gamma^\mu u_{jL} + \bar{d}_{iL}\gamma^\mu e_L^+] \\ & + i\frac{g_Y}{\sqrt{2}}Y_\mu^i[\varepsilon_{ijk}\bar{u}_{kL}^c\gamma^\mu d_{jL} + \bar{d}_{iR}\gamma^\mu \mu_R^c - \bar{u}_{iL}\gamma^\mu e_L^+]\end{aligned}\quad (7.7)$$

e dessa forma, os seguintes processo são possíveis

$$\begin{aligned}\bar{X} & \rightarrow uu \quad (B = 2/3) \\ & \rightarrow \bar{d}e^+ \quad (B = -1/3) \quad \text{e} \\ \bar{Y} & \rightarrow ud \quad (B = 2/3) \\ & \rightarrow \bar{u}e^+ \quad (B = -1/3) .\end{aligned}$$

O que leva ao decaimento do próton através da reação $p \rightarrow e^+\pi^0$ com uma razão de decaimento dada por

$$\Gamma(p \rightarrow e^+\pi^0) \approx \frac{g_{X,Y}^2}{m_{X,Y}^2} \quad (7.8)$$

levando a uma estimativa para o tempo de vida médio do próton de aproximadamente

$$\tau_p \approx \frac{m_X^4}{g_X^4 m_p^5} , \quad (7.9)$$

o que nos levaria a concluir que as massas dos leptoquarks X e Y estariam próximas à escala de Planck levando-se em conta os atuais limites sobre a vida média do próton!

- Leptoquarks acoplem-se a uma única geração de quarks e léptons. Caso contrário podem dar a origem a processos de corrente neutra com mudança de sabor, que são processos altamente suprimidos no modelo padrão.
- Leptoquarks possuam acoplamentos quirais, ou seja, não se acoplem simultaneamente a férmions de mão direita e de mão esquerda para evitar os fortes vínculos provenientes das experiências de decaimento de mésons pseudoescalares como píons.
- Devem posuir acoplamentos de gauge invariantes sob o grupo $SU(3)_c \otimes SU(2)_L \otimes U(1)_Y$ pois espera-se que suas massas sejam maiores do que a escala de quebra espontânea da simetria eletrofraca.

Buchmüller *et.al.* [64] postularam uma lagrangiana efetiva de interações de leptoquarks escalares e vetoriais com as partículas do modelo padrão, que levam em conta os itens apresentados acima. O setor escalar de tal Lagrangiana é dado por

$$\mathcal{L}_{eff} = \mathcal{L}_{F=2} + \mathcal{L}_{F=0} + h.c. \quad (7.10)$$

$$\begin{aligned} \mathcal{L}_{F=2} = & g_{1L} \bar{q}_L^c i\tau_2 \ell_L S_{1L} + g_{1R} \bar{u}_R^c e_R S_{1R} \\ & + \tilde{g}_{1R} \bar{d}_R^c e_R \tilde{S}_1 + g_{3L} \bar{q}_L^c i\tau_2 \vec{\ell}_L \cdot \vec{S}_3 \end{aligned} \quad (7.11)$$

$$\begin{aligned} \mathcal{L}_{F=0} = & h_{2L} R_{2L}^T \bar{u}_R i\tau_2 \ell_L + h_{2R} \bar{q}_L e_R R_{2R} \\ & + \tilde{h}_{2L} \tilde{R}_2^T \bar{d}_R i\tau_2 \ell_L, \end{aligned} \quad (7.12)$$

onde $F = 3B + L$ é o número fermiônico dos leptoquarks. $S_{1R(L)}$ e $\tilde{S}_1$ são singletos, $R_{2R(L)}$ e $\tilde{R}_2$ dubletos e S_3 tripleto do grupo $SU(2)_L$. Os campos ℓ_L e q_L são dubletos de léptons e quarks, enquanto e_R , u_R e d_R são singletos de $SU(2)$; h e g são acoplamentos de Yukawa. Podemos ver das interações acima, que os principais modos de decaimento de leptoquarks são em pares $e^{(\pm)}q$ e/ou $\nu_e q'$. O setor vetorial não é renormalizável e possui vínculos experimentais muito mais fortes.

Leptoquarks de número fermiônico $F = 2$ acoplam-se a pares ℓq , enquanto os de $F = 0$, os quais estamos interessados em estudar, têm acoplamentos $\ell \bar{q}$. O acoplamento de gauge dos leptoquarks escalares com os bósons de gauge eletrofracos podem ser encontrados em [61], por exemplo; o acoplamento com glúons é dado pela QCD escalar.

A lagrangiana efetiva que utilizamos para descrever a propagação e a interação dos leptoquarks, a partir da qual obtivemos as regras de Feynman do apêndice A, é dada por

$$\mathcal{L}_S = \mathcal{L}_g + \mathcal{L}_{Yukawa} \quad (7.13)$$

$$\mathcal{L}_g = D_\mu \phi_0 (D^\mu \phi_0)^* - m_0^2 \phi_0 \phi_0^* \quad (7.14)$$

$$\mathcal{L}_{Yukawa} = \lambda_0 \bar{q}_0 (a + b\gamma^5) \ell_0 \phi_0 + \lambda_0 \bar{\ell}_0 (a - b\gamma^5) q_0 \phi_0^*, \quad (7.15)$$

onde $D_\mu = \partial_\mu - ig_s^0 T^a A_{\mu 0}^a$ é derivada covariante que introduz o acoplamento de gauge glúon-leptoquark. Os campos ϕ_0 escalar, q_0 e ℓ_0 espinoriais e $A_{\mu 0}^a$ vetorial descrevem o leptoquark S , o quark de sabor q e o elétron[†], e o glúon, respectivamente. O subscrito 0 indica que esses são os “campos nus” não renormalizados. Essa lagrangiana tem a mesma forma de (7.12), e pode reproduzi-la ao especificarmos os campos escalares ϕ e as constantes λ , a e b . A renormalização de (7.15) é feita através das seguintes redefinições

$$\phi_0 = \sqrt{Z_\phi} \phi_R \quad (7.16)$$

[†]Ou o pósitron dependendo do tipo de leptoquark.

$$A_{\mu 0}^a = \sqrt{Z_g} A_{\mu R}^a \quad (7.17)$$

$$q_0 = \sqrt{Z_q} q_R \quad (7.18)$$

$$\ell_0 = \sqrt{Z_l} \ell_R . \quad (7.19)$$

As constantes de renormalização, por outro lado, podem ser escritas em função de *contratermos*

$$Z_\phi = 1 + \delta\phi \quad (7.20)$$

$$Z_\phi \sqrt{Z_g} g_s^0 \equiv \tilde{Z}_g g_s^0 = g_s^R + \delta g \quad (7.21)$$

$$Z_\phi m_0^2 = m_R^2 + \delta m^2 \quad (7.22)$$

$$\sqrt{Z_\phi Z_q Z_l} \lambda_0 \equiv \tilde{Z}_\lambda \lambda_0 = \lambda_R + \delta\lambda . \quad (7.23)$$

A lagrangiana dos campos renormalizados e dos contratermos agora pode ser obtida pela substituição dessas definições na lagrangiana de campos não renormalizados original, o que nos dá

$$\begin{aligned} \mathcal{L}_g &= \partial_\mu \phi_R \partial^\mu \phi_R^* - m_R^2 \phi_R \phi_R^* \\ &+ \delta\phi \partial_\mu \phi_R \partial^\mu \phi_R^* - \delta m^2 \phi_R \phi_R^* + \dots \end{aligned} \quad (7.24)$$

$$\begin{aligned} \mathcal{L}_{Yukawa} &= \lambda_R \bar{q}_R (a + b\gamma^5) \ell_R \phi_R + c.c. \\ &+ \delta\lambda \bar{q}_R (a + b\gamma^5) \ell_R \phi_R + c.c. , \end{aligned} \quad (7.25)$$

onde as reticências indicam termos que contém os vértices tríplice SSg e o vértice quártico $SSgg$ que não precisam ser renormalizados em ordem α_s considerando o processo de Drell-Yan.

A partir destas lagrangianas são obtidas as regras de Feynman que contém leptoquarks escalares e os respectivos contratermos, e que são dadas no Apêndice A. É possível mostrar que a contribuição dos leptoquarks à função β da QCD, que calculamos na seção (2.4.2), é dada por

$$\beta_S(g_s) = -\frac{g_s^3}{16\pi^2} \left(\frac{11}{3} C_A - \frac{1}{3} T_R \right) < 0 . \quad (7.26)$$

A liberdade assintótica da QCD é portanto mantida após a introdução dos graus de liberdade escalares. Na próxima seção iniciamos os cálculos das amplitudes virtuais com a inclusão de leptoquarks.

7.2 Efeitos Virtuais de Leptoquarks Escalares no Processo Drell-Yan ao Nível de Árvore

O elemento de matriz de ordem zero incluindo a troca de um leptoquark no canal t da reação é obtido a partir dos seguintes diagramas de Feynman

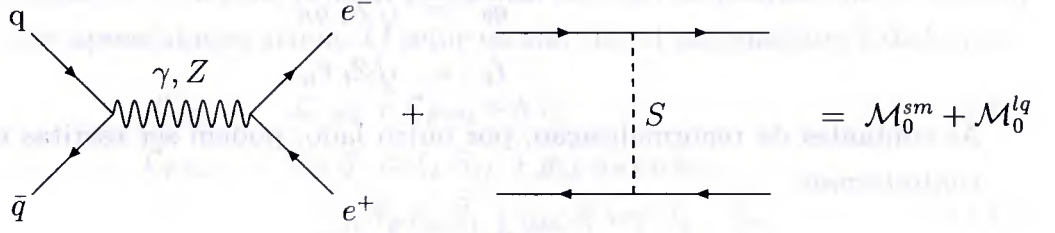


Figura 7.2: Contribuição de leptoquarks escalares ao processo $q\bar{q} \rightarrow e^+e^-$.

Utilizando as regras de Feynman do apêndice A e adotando as mesmas convenções do capítulo 6, obtemos

$$\mathcal{M}_0^{sm} = \frac{i}{k_V^2} \bar{u}_e(k_1) \gamma^\mu (V_e + A_e \gamma^5) v_e(k_2) \bar{v}_q(p_b) \gamma_\mu (V_q + A_q \gamma^5) u_q(p_a) \quad (7.27)$$

$$\mathcal{M}_0^{lq} = \frac{-i\lambda^2}{\hat{t} - m^2} \bar{u}_e(k_1) \gamma_+ u_q(p_a) \bar{v}_q(p_b) \gamma_- v_e(k_2) , \quad (7.28)$$

onde as matrizes $\gamma_\pm$ contidas nos acoplamentos de Yukawa são dadas por $\gamma_\pm = a \pm b\gamma^5$ e m é a massa do leptoquark escalar S .

O quadrado da amplitude total é dado por

$$|\mathcal{M}_0|^2 = |\mathcal{M}_0^{sm}|^2 + 2\Re(\mathcal{M}_0^{sm*} \mathcal{M}_0^{lq}) + |\mathcal{M}_0^{lq}|^2 , \quad (7.29)$$

onde o cálculo dos traços foi realizado em $d = 4 - 2\epsilon$ dimensões com o auxílio do FeynCalc e o resultado vem a seguir

$$|\mathcal{M}_0^{sm}|^2 = \frac{1}{4} \frac{3}{9} \frac{8}{|k_V^2|^2} \{ (\hat{u}^2 + \hat{t}^2) d_e d_q - (\hat{u}^2 - \hat{t}^2) h_e h_q + \epsilon [-\hat{s}^2 d_e d_q + 3(\hat{u}^2 - \hat{t}^2) h_e h_q] - 2\epsilon^2 (\hat{u}^2 - \hat{t}^2) h_e h_q \} \quad (7.30)$$

$$2\Re(\mathcal{M}_0^{sm*} \mathcal{M}_0^{lq}) = -2\Re\left(\frac{1}{k_V^2}\right) \frac{\lambda^2}{\hat{t} - m^2} \frac{1}{4} \frac{3}{9} \times 4\hat{t}(\hat{t} + \epsilon\hat{s}) [(V_e V_q - A_e A_q)(a^2 + b^2) + (V_e A_q - A_e V_q) 2ab] \quad (7.31)$$

$$|\mathcal{M}_0^{lq}|^2 = \frac{1}{4} \frac{3}{9} \frac{\lambda^4}{(\hat{t} - m^2)^2} \times 4(a^2 + b^2)^2 \hat{t}^2 . \quad (7.32)$$

As constantes de acoplamento estão definidas no apêndice A, assim como k_V^2 . A seção de choque total partônica ao nível de árvore é denotada por $\hat{\sigma}_0^{sm+lq}$. A seguir apresentamos as correções em NLO de QCD aos processos citados acima.

7.3 Correções Virtuais em Ordem α_s

As correções virtuais de ordem α_s incluem, além daquelas amplitudes já calculadas no setor do modelo padrão, as amplitudes correspondentes à autoenergia do leptoquark escalar, à correção do vértice de Yukawa $S\ell q$ e à caixa mostradas na figura abaixo. Os esquemas e convenções são os mesmos utilizados no capítulo 6.

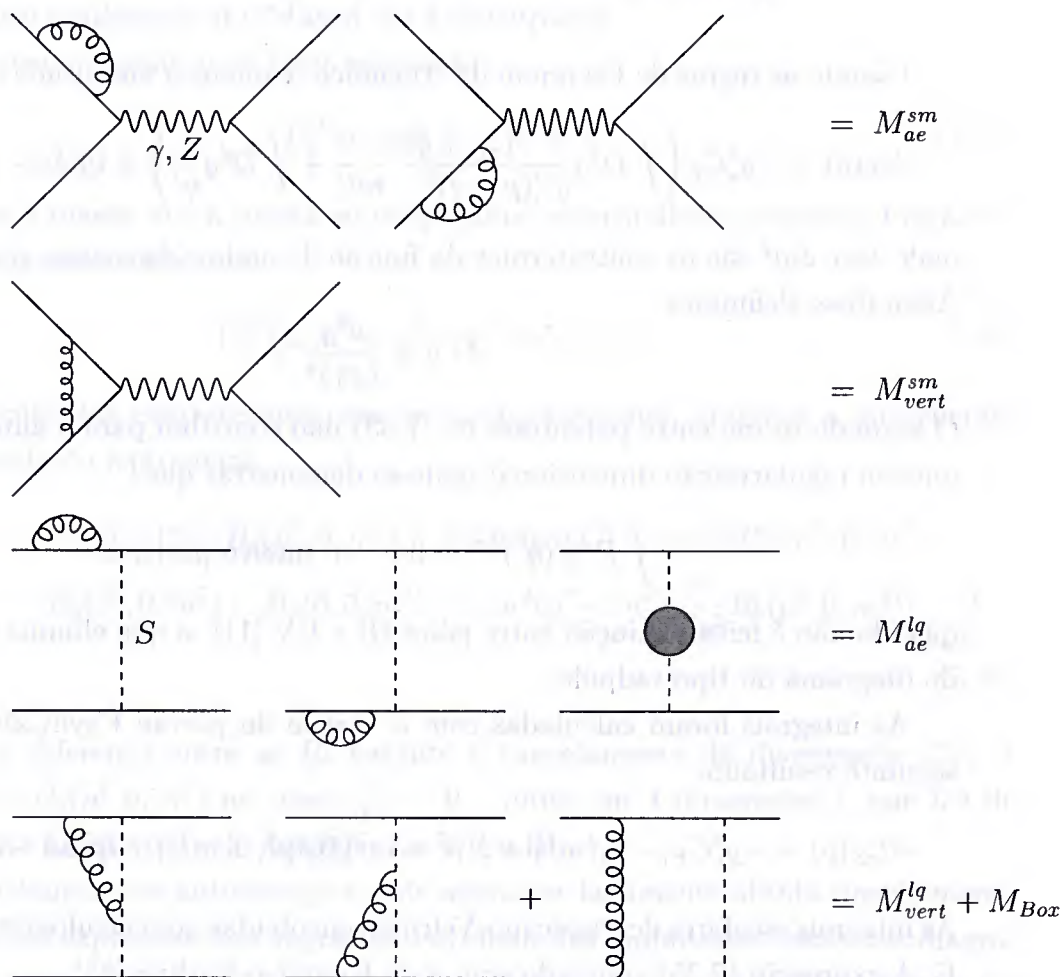


Figura 7.3: Amplitudes virtuais ao nível de 1-loop.

A seguir passamos ao cálculo das amplitudes virtuais contendo leptoquarks escalares. As amplitudes que só contêm partículas do modelo padrão podem ser encontradas no capítulo 6.

7.3.1 Autoenergia do Leptoquark

A correção de 1-loop à autoenergia do leptoquark escalar é obtida pela soma dos diagramas de Feynman da figura (7.4).

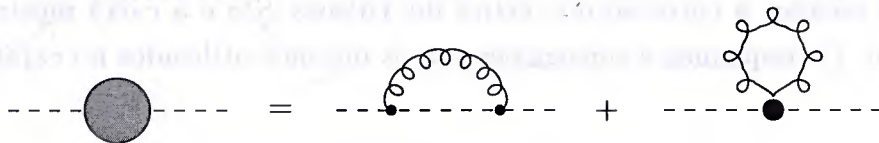


Figura 7.4: Autoenergia do leptoquark ao nível de 1-loop.

Usando as regras de Feynman do Apêndice A temos a amplitude da autoenergia

$$-i\Sigma_S(p) = -g_s^2 C_F \left(\int D^d q \frac{(2p+q)^2}{q^2[(p+q)^2 - m^2]} + \int D^d q \frac{1}{q^2} \right) + ip^2 \delta\phi - i\delta m^2, \quad (7.33)$$

onde $\delta\phi$ e δm^2 são os contratermos da função de onda e de massa, respectivamente. Além disso definimos

$$D^d q \equiv \frac{d^d q}{(2\pi)^d}.$$

O segundo termo entre parênteses de (7.33) não contribui para a autoenergia, dado que em regularização dimensional pode-se demonstrar que

$$\int D^d q (q^2)^{-n} = 0, \quad n \text{ inteiro positivo} \quad (7.34)$$

quando não é feita distinção entre pólos IR e UV [11], o que elimina a contribuição do diagrama do tipo tadpole.

As integrais foram calculadas com o auxílio do pacote FeynCalc fornecendo o seguinte resultado

$$-i\Sigma_S(p) = -g_s^2 C_F [-A_0(m^2) + 2(p^2 + m^2)B_0(p^2, 0, m^2)] + ip^2 \delta\phi - i\delta m^2. \quad (7.35)$$

As integrais escalares de Passarino-Veltman envolvidas nos cálculos estão no apêndice E. A expressão (7.35) concorda com a de Kunszt e Stirling [65].

Condições de Renormalização

Vamos aqui seguir o esquema proposto por *Nason et.al.* na ref. [66]. Nesse esquema as partículas pesadas se desacoplam a baixas energias e os pártons leves continuam a obedecer as mesmas equações do grupo de renormalização que obedeciam na ausência das partículas pesadas [67]. As seguintes condições de renormalização

são usadas para a função de Green de 2-pontos de uma partícula pesada de massa m :

1. Renormalização da Função de onda do Leptoquark

A condição

$$\left. \frac{d}{dp^2} \Gamma_S^{(2)}(p, m) \right|_{p^2=0} = 1 , \quad (7.36)$$

fixa a renormalização da função de onda do leptoquark.

2. Renormalização da Massa do Leptoquark

Nesse caso, a condição de renormalização

$$\left. \Gamma_S^{(2)}(p, m) \right|_{p^2=m^2} = 0 , \quad (7.37)$$

implica que a massa m é a massa no propagador renormalizado. Sendo a função de Green de 2-pontos dada por

$$\Gamma_S^{(2)}(p, m) = p^2 - m^2 + \Sigma_S(p) , \quad (7.38)$$

após o cálculo dos contratermos com as condições acima, obtemos a autoenergia renormalizada do leptoquark

$$\begin{aligned} \Sigma_S^{Ren}(p) = & -ig_s^2 C_F \{ 2p^2 [B_0(p^2, 0, m^2) - B_0(0, 0, m^2)] - 2m^2 [2B_0(m^2, 0, m^2) \\ & - B_0(p^2, 0, m^2) - B_0(0, 0, m^2)] - 2m^2 (p^2 - m^2) \left. \frac{d}{dp^2} B_0(p^2, 0, m^2) \right|_{p^2=0} \} . \end{aligned} \quad (7.39)$$

Note que a diferença entre as B_0 garante o cancelamento da divergência UV. A derivada de $B_0(p^2, 0, m^2)$ no ponto $p^2 = 0$ é finita em 4 dimensões, o que faz de $\Sigma_S^{Ren}(p)$ uma função livre de divergências UV e IR.

A contribuição das autoenergias pode agora ser facilmente obtida simplesmente inserindo suas expressões nas regras de Feynman das amplitudes totais dos diagramas de Feynman da figura (7.3). As correções de autoenergia das partículas externas (quarks não-massivos) são multiplicadas pelo fator 1/2 de modo a levar em conta de maneira apropriada a transição das funções de Green para a matriz-S.

Dessa forma obtemos os seguintes resultados[†]

$$2\Re(\mathcal{M}_0^{sm*} \mathcal{M}_{ae}^{sm}) = \frac{\alpha_s}{2\pi} C_F N_\epsilon \left(\frac{\mu^2}{\hat{s}} \right)^\epsilon \left(\frac{1}{\epsilon} + F_1^V \right) |\mathcal{M}_0^{sm}|^2 \quad (7.40)$$

[†]Lembrando $N_\epsilon = (4\pi)^\epsilon \Gamma(1 + \epsilon)$.

$$2\Re(\mathcal{M}_0^{lq*} \mathcal{M}_{ae}^{sm}) = \frac{\alpha_s}{4\pi} C_F N_\epsilon \left(\frac{\mu^2}{\hat{s}} \right)^\epsilon \left(\frac{1}{\epsilon} + F_1^V \right) \times 2\Re(\mathcal{M}_0^{sm*} \mathcal{M}_0^{lq}) \quad (7.41)$$

$$2\Re(\mathcal{M}_0^{lq*} \mathcal{M}_{ae}^{lq}) = \frac{\alpha_s}{2\pi} C_F N_\epsilon \left(\frac{\mu^2}{\hat{s}} \right)^\epsilon \left(\frac{1}{\epsilon} + F_2^V \right) |\mathcal{M}_0^{lq}|^2 \quad (7.42)$$

$$2\Re(\mathcal{M}_0^{sm*} \mathcal{M}_{ae}^{lq}) = \frac{\alpha_s}{4\pi} C_F N_\epsilon \left(\frac{\mu^2}{\hat{s}} \right)^\epsilon \left(\frac{1}{\epsilon} + F_2^V \right) \times 2\Re(\mathcal{M}_0^{sm*} \mathcal{M}_0^{lq}) , \quad (7.43)$$

onde F_1^V e F_2^V representam as partes finitas das autoenergias. Apesar da autoenergia renormalizada do leptoquark ser finita em 4 dimensões, $\mathcal{M}_{ae}^{lq}$, definida na figura (7.3), possui um pólo IR proveniente da soma das amplitudes que contém a correção de autoenergia do quark.

7.3.2 Correção do Vértice leptoquark–quark–lépton

A correção de QCD, ao nível de 1-loop, do vértice *leptoquark–quark–lépton* é obtida a partir do seguinte diagrama de Feynman

$$\Gamma_+(p_a, k_1) = \text{diagrama} + \text{contratermo}$$

Figura 7.5: Correção do vértice $S\ell q$ ao nível de 1-loop em QCD.

A amplitude correspondente é da forma

$$\Gamma_+(p_a, k_1) = -ig_s^2 C_F \int D^d k \frac{N_+}{k^2(k+p_1)^2[(k+p_a-k_1)^2-m^2]} \gamma_+ , \quad (7.44)$$

onde

$$N_+ = k^2 + 2k \cdot (p_a - k_1) - 4k_1 \cdot p_a . \quad (7.45)$$

Renormalização no Esquema On-Shell

A exemplo do vértice $(\gamma, Z) q\bar{q}$, o contatermo é fixado subtraindo-se o valor do vértice no ponto $p_a^2 = k_1^2 = 0$, $p_a - k_1 = 0$. Logo

$$\begin{aligned} \Gamma_+^{Ren} &= \Gamma_+(p_a, k_1) - \Gamma_+(p_a^2 = k_1^2 = 0 , p_a - k_1 = 0) \\ &= -ig_s^2 C_F \{ B_0(0, 0, m^2) - B_0(\hat{t}, 0, m^2) \\ &\quad + 2m^2 [C_0(0, 0, \hat{t}, m^2, 0, 0) - C_0(0, 0, 0, m^2, 0, 0)] \} . \end{aligned} \quad (7.46)$$

Note que a diferença entre as B_0 's deixa claro o cancelamento da divergência UV.

Usando as funções do Apêndice E, obtemos

$$\Gamma_{\pm}^{Ren} = \frac{\alpha_s}{4\pi} C_F N_{\epsilon} \left(\frac{\mu^2}{\hat{s}} \right)^{\epsilon} \left[-\frac{2}{\epsilon} - \frac{2m^2}{\hat{t}} \frac{1}{\epsilon} \ln \left(1 - \frac{\hat{t}}{m^2} \right) + F_4^V \right], \quad (7.47)$$

onde a função finita F_4^V pode ser obtida pelas funções escalares do Apêndice E. O vértice *leptoquark-antiquark-antilépton* tem as mesmas correções do vértice Γ_+ .

Analogamente ao caso da correção das autoenergias, as amplitudes dos vértices geram as seguintes correções ao quadrado do elemento de Born

$$2\Re(\mathcal{M}_0^{sm*} \mathcal{M}_{vert}^{sm}) = \frac{\alpha_s}{2\pi} C_F N_{\epsilon} \left(\frac{\mu^2}{\hat{s}} \right)^{\epsilon} \left(-\frac{2}{\epsilon^2} - \frac{4}{\epsilon} + F_{vert}^{sm} \right) |\mathcal{M}_0^{sm}|^2 \quad (7.48)$$

$$2\Re(\mathcal{M}_0^{sm*} \mathcal{M}_{vert}^{lq}) = \frac{\alpha_s}{4\pi} C_F N_{\epsilon} \left(\frac{\mu^2}{\hat{s}} \right)^{\epsilon} \left[-\frac{4}{\epsilon} - \frac{4m^2}{\epsilon \hat{t}} \ln \left(1 - \frac{\hat{t}}{m^2} \right) + F_{vert}^{lq} \right] \times 2\Re(\mathcal{M}_0^{sm*} \mathcal{M}_0^{lq}) \quad (7.49)$$

$$2\Re(\mathcal{M}_0^{lq*} \mathcal{M}_{vert}^{sm}) = \frac{\alpha_s}{4\pi} C_F N_{\epsilon} \left(\frac{\mu^2}{\hat{s}} \right)^{\epsilon} \left(-\frac{2}{\epsilon^2} - \frac{4}{\epsilon} + F_{vert}^{sm} \right) \times 2\Re(\mathcal{M}_0^{sm*} \mathcal{M}_0^{lq}) \quad (7.50)$$

$$2\Re(\mathcal{M}_0^{lq*} \mathcal{M}_{vert}^{lq}) = \frac{\alpha_s}{2\pi} C_F N_{\epsilon} \left(\frac{\mu^2}{\hat{s}} \right)^{\epsilon} \left[-\frac{4}{\epsilon} - \frac{4m^2}{\epsilon \hat{t}} \ln \left(1 - \frac{\hat{t}}{m^2} \right) + F_{vert}^{lq} \right] \times |\mathcal{M}_0^{lq}|^2. \quad (7.51)$$

As funções F_{vert}^{sm} e F_{vert}^{lq} representam as partes finitas das expressões do vértice $(\gamma, Z)q\bar{q}$ e $S q(\bar{q})\ell(\bar{\ell})$ após a renormalização das divergências UV.

7.3.3 Contribuição Virtual da Caixa

Sem dúvida, a contribuição mais trabalhosa e difícil de tratar é a da “caixa”, que corresponde ao diagrama de Feynman da figura abaixo.

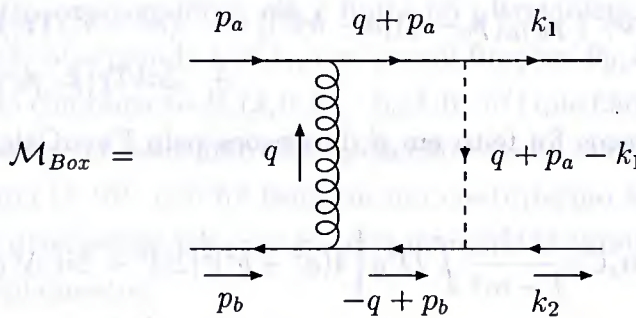


Figura 7.6: Contribuição virtual da caixa.

A amplitude correspondente é dada por

$$\begin{aligned} \mathcal{M}_{Box} = & \bar{u}_e(k_1) \cdot i\lambda(a + b\gamma^5) \int D^d q \frac{i(\not{q} + \not{p}_a)}{(q + p_a)^2} \cdot ig_s T^a \gamma^\mu \cdot u_q(p_a) \times \frac{-ig_{\mu\nu}\delta^{ab}}{q^2} \\ & \times \bar{v}_q(p_b) \cdot ig_s T^b \gamma^\nu \cdot \frac{i(\not{q} - \not{p}_b)}{(q - p_b)^2} \cdot i\lambda(a - b\gamma^5) \cdot v_e(k_2) \\ & \times \frac{i}{(q + p_a - k_1)^2 - m^2} . \end{aligned} \quad (7.52)$$

Após várias abordagens, concluímos que a forma mais fácil e direta de se calcular essa contribuição é primeiramente calcular os traços da interferência entre $\mathcal{M}_{Box}$ e $\mathcal{M}_0$ e depois integrar no momento do loop.

1. Interferência $\langle M_{Box} \otimes M_0^{lq} \rangle$

Nesse caso obtemos o seguinte

$$\begin{aligned} |\mathcal{M}_{int}^{lq}|^2 &= 2\Re(\mathcal{M}_0^{lq*} \mathcal{M}_{Box}) \\ &= 2C_F g_s^2 \frac{\lambda^4}{\hat{t} - m^2} \frac{1}{4} \times \int D^d q \text{Tr}[\not{p}_a(a - b\gamma^5) \not{k}_1(a + b\gamma^5)(\not{q} + \not{p}_a)\gamma^\mu] \\ &\times \text{Tr}[\not{k}_2(a + b\gamma^5) \not{p}_b\gamma_\mu(\not{p}_b - \not{q})(a - b\gamma^5)] \times \frac{1}{D_4} , \end{aligned} \quad (7.53)$$

onde

$$D_4 = q^2(q + p_a)^2(q - p_b)^2[(q + p_a - k_1)^2 - m^2] . \quad (7.54)$$

Os traços que contêm γ^5 foram calculados em um esquema onde

$$\{\gamma^5, \gamma^\mu\} = 0 \quad \mu = 0, 1, 2, 3 , \quad (7.55)$$

$$\gamma_5^2 = 1 , \quad (7.56)$$

$$\text{Tr}[\gamma^\mu \gamma^\nu \gamma^\sigma \gamma^\rho \gamma^5] = -4i\epsilon_D^{\mu\nu\rho\sigma} \quad (7.57)$$

e podem ser decompostos da seguinte forma utilizando as regras acima

$$\text{Tr}[\not{p}_a(a - b\gamma^5) \not{k}_1(a + b\gamma^5)(\not{q} + \not{p}_a)\gamma^\mu] = (a^2 + b^2)\text{Tr}[\not{p}_a \not{k}_1(\not{q} + \not{p}_a)\gamma^\mu] \quad (7.58)$$

$$+ 2ab\text{Tr}[\not{p}_a \not{k}_1\gamma^5(\not{q} + \not{p}_a)\gamma^\mu] \quad (7.59)$$

$$\text{Tr}[\not{k}_2(a + b\gamma^5) \not{p}_b\gamma_\mu(\not{p}_b - \not{q})(a - b\gamma^5)] = (a^2 + b^2)\text{Tr}[\not{k}_2 \not{p}_b\gamma_\mu(\not{p}_b - \not{q})] \quad (7.60)$$

$$+ 2ab\text{Tr}[\not{k}_2 \not{p}_b\gamma_\mu\gamma^5(\not{p}_b - \not{q})] . \quad (7.61)$$

O cálculo dos traços foi feito em d dimensões pelo FeynCalc e o resultado vem a seguir

$$\begin{aligned} |\mathcal{M}_{int}^{lq}|^2 &= \pi\alpha_s C_F \frac{\lambda^4}{\hat{t} - m^2} \int D^d q \left\{ 4(a^2 + b^2)^2 [2\hat{s}\hat{t}^2 + 2\hat{s}\hat{t} \Lambda^\mu q_\mu + \Gamma^{\mu\nu} q_\mu q_\nu - \hat{t}^2 q^2] \right. \\ &\quad \left. + 16a^2 b^2 (d - 3) [\Gamma^{\mu\nu} q_\mu q_\nu - (\hat{s}^2 - \hat{u}^2) q^2] \right\} \frac{1}{D_4} , \end{aligned} \quad (7.62)$$

onde

$$\Lambda^\mu = (p_a - p_b + k_2 - k_1)^\mu, \quad (7.63)$$

$$\Gamma^{\mu\nu} = -2s(p_a^\mu p_b^\nu + k_1^\mu k_2^\nu) - 2u(k_1^\mu p_b^\nu + k_2^\mu p_a^\nu). \quad (7.64)$$

O próximo passo consistiu em integrar no momento do loop q , de onde obtivemos o resultado das integrais abaixo relacionadas

$$\mu^{2\epsilon} \int D^d q \frac{1}{D_4} \doteq D_0(0, \hat{s}, \hat{s} + \hat{t} + \hat{u}, \hat{t}, 0, 0, 0, 0, m^2) \equiv I_1 \quad (7.65)$$

$$\mu^{2\epsilon} \int D^d q \frac{q_\mu}{D_4} \Lambda^\mu = C_0(0, 0, \hat{s}, 0, 0, 0) + (m^2 - \hat{t})D_0 - C_0(0, \hat{s}, \hat{s} + \hat{t} + \hat{u}, m^2, 0, 0) \equiv I_2 \quad (7.66)$$

$$\begin{aligned} \mu^{2\epsilon} \int D^d q \frac{q_\mu q_\nu}{D_4} \Gamma^{\mu\nu} &= \hat{t}[B_0(\hat{s}, 0, 0) - B_0(\hat{t}, 0, m^2)] + (m^2 - \hat{t}) \left\{ \hat{s}[C_0(0, 0, \hat{s}, 0, 0, 0) \right. \\ &+ (m^2 - \hat{t})D_0] + \hat{t}[C_0(0, 0, \hat{t}, m^2, 0, 0) + C_0(0, \hat{t}, \hat{s} + \hat{t} + \hat{u}, 0, 0, m^2)] \\ &\left. - \frac{1}{2}F(\hat{s}, \hat{t}, \hat{u}, m^2)C_0(0, \hat{s}, \hat{s} + \hat{t} + \hat{u}, m^2, 0, 0) \right\} \equiv I_3 \end{aligned} \quad (7.67)$$

$$\mu^{2\epsilon} \int D^d q \frac{q^2}{D_4} = C_0(0, \hat{s}, \hat{s} + \hat{t} + \hat{u}, m^2, 0, 0) \equiv I_4, \quad (7.68)$$

onde a função $F(\hat{s}, \hat{t}, \hat{u}, m^2)$ é um polinômio de seus argumentos e é dada explicitamente no Apêndice E. A função D_0 contida em I_2 e I_3 é a mesma definida por I_1 .

Finalmente podemos obter a contribuição virtual em uma forma compacta

$$\begin{aligned} |\mathcal{M}_{int}^{lq}|^2 &= \pi \alpha_s C_F \frac{\lambda^4}{\hat{t} - m^2} \{ 4(a^2 + b^2)^2 [\hat{s}\hat{t}^2 I_1 + 2\hat{s}\hat{t} I_2 + I_3 - \hat{t}^2 I_4] + 16a^2 b^2 (1 - 2\epsilon) [I_3 \\ &+ (\hat{s}^2 - \hat{u}^2) I_4] \}. \end{aligned} \quad (7.69)$$

Divergências UV

Como a amplitude da caixa possui, no máximo, duas potências de momento no numerador e quatro propagadores, ela é finita no ultravioleta. Isso pode ser visto em nossa amplitude observando que I_3 , que possui funções B_0 , contém essas funções somente através da combinação $B_0(\hat{s}, 0, 0) - B_0(\hat{t}, 0, m^2)$ que não possui pólo UV. Isso reflete o fato da teoria ser renormalizável, pois não havendo termos quárticos do tipo $q\bar{q}\ell\bar{\ell}$ na lagrangiana (7.15), não há também um contratermo associado. Logo, se a caixa tivesse uma divergência UV, não haveria maneira de eliminá-la pela redefinição dos campos e acoplamentos.

Limite Soft da Caixa: Pólos Duplos

No limite em que o momento q do glúon se torna soft é válida a aproximação

$$\mu^{2\varepsilon} \int D^d q \frac{1}{D_4} \approx \frac{1}{\hat{t} - m^2} \mu^{2\varepsilon} \int D^d q \frac{1}{q^2(q + p_a)^2(q - p_b)^2} \approx \frac{1}{\hat{t} - m^2} C_0(0, 0, \hat{s}, 0, 0, 0) ,$$

de maneira que a combinação

$$C_0(0, 0, \hat{s}, 0, 0, 0) + (m^2 - \hat{t})D_0(0, \hat{s}, \hat{s} + \hat{t} + \hat{u}, \hat{t}, 0, 0, 0, 0, m^2)$$

não contém pólos duplos soft. O único pólo duplo de (7.69), proveniente de I_1 , é proporcional ao elemento de Born.

Termos não-proporcionais a $\mathcal{M}_0^{lq}$ - Pólos Duplos e Simples

Pela expressão (7.69) vemos que as integrais I_3 e I_4 estão contidas em termos não proporcionais a $\mathcal{M}_0^{lq}$. No entanto, como veremos adiante, os pólos IR das contribuições de glúons reais continuam sendo proporcionais ao elemento de Born quando incluímos a contribuição dos leptoquarks escalares. Se quisermos ter esperança de que os pólos IR se cancelem em uma seção de choque total inclusiva, devemos mostrar que I_3 e I_4 são finitas em 4 dimensões.

A integral I_4 , na verdade, nada mais é do que a função escalar de Passarino-Veltman $C_0(0, \hat{s}, \hat{s} + \hat{t} + \hat{u}, m^2, 0, 0)$. O cálculo explícito dessa função mostrou que ela é de fato finita. Essa função está reproduzida de forma explícita no apêndice E.

No que diz respeito a I_3 , podemos encontrar dois argumentos distintos para mostrar que ela é finita. O primeiro é devido a *Nason et.al.* [18, 68] e consiste na análise da seguinte combinação de integrais escalares que ocorre em I_3

$$\begin{aligned} & \hat{s}[C_0(0, 0, \hat{s}, 0, 0, 0) + (m^2 - \hat{t})D_0] + \hat{t}[C_0(0, 0, \hat{t}, m^2, 0, 0) + C_0(0, \hat{t}, \hat{s} + \hat{t} + \hat{u}, 0, 0, m^2)] \\ &= \mu^{2\varepsilon} \int D^d q \left\{ \frac{\hat{s}}{q^2(q - p_a)^2(q + p_b)^2} + \frac{\hat{s}(m^2 - \hat{t})}{q^2(q - p_a)^2(q + p_b)^2[(q + k_1 - p_a)^2 - m^2]} \right. \\ & \quad \left. + \frac{\hat{t}}{q^2(q - p_a)^2[(q + k_1 - p_a)^2 - m^2]} + \frac{\hat{t}}{q^2(q + p_b)^2[(q + k_1 - p_a)^2 - m^2]} \right\} \\ &= \mu^{2\varepsilon} \int D^d q \frac{\hat{s}[(q + k_1 - p_a)^2 - m^2] + \hat{s}(m^2 - \hat{t}) + \hat{t}(q + p_b)^2 + \hat{t}(q - p_a)^2}{D_4} \\ &= \mu^{2\varepsilon} \int D^d q \frac{(\hat{s} + 2\hat{t})q^2 + 2q \cdot [\hat{s}(k_1 - p_a) + \hat{t}(p_b - p_a)]}{D_4} . \end{aligned} \quad (7.70)$$

Em primeiro lugar, o numerador não contém termo em q^0 e, portanto, não há pólo duplo soft. O termo em q^2 é proporcional a $C_0(0, \hat{s}, \hat{s} + \hat{t} + \hat{u}, m^2, 0, 0)$ que é finita em 4 dimensões. Além disso, o numerador vai a zero no limite soft do glúon ($q \rightarrow 0$). O termo de ordem q , por sua vez, envolve uma projeção de q em uma direção ortogonal

a p_a e p_b , veja

$$\begin{aligned} 2p_a \cdot [\hat{s}(k_1 - p_a) + \hat{t}(p_b - p_a)] &= \hat{s}(-\hat{t}) + \hat{t}(\hat{s}) = 0 \\ 2p_b \cdot [\hat{s}(k_1 - p_a) + \hat{t}(p_b - p_a)] &= \hat{s}(-\hat{u} - \hat{s}) + \hat{t}(-\hat{s}) \\ &= \hat{s}(\hat{t}) + \hat{t}(-\hat{s}) = 0 , \end{aligned}$$

de maneira que não ocorrem pólos simples colineares em (7.70), tornando I_3 finita em 4 dimensões!

O segundo argumento na verdade é uma constatação. No Apêndice E, mostramos que (7.70) pode ser escrita em termos de $C_0(0, \hat{s}, \hat{s} + \hat{t} + \hat{u}, m^2, 0, 0)$ e além disso mostramos também que a função D_0 pode ser escrita em termos de funções escalares de 3-pontos, que são muito mais fáceis de serem calculadas explicitamente.

Com isso, usando o resultado do cálculo das integrais escalares do Apêndice E, obtemos a seguinte contribuição

$$\begin{aligned} 2\Re(\mathcal{M}_0^{lq*} \mathcal{M}_{box}) &= \frac{\alpha_s}{2\pi} C_F N_\epsilon \left(\frac{\mu^2}{\hat{s}} \right)^\epsilon \left[-\frac{2}{\epsilon^2} + \frac{4}{\epsilon} \frac{m^2}{\hat{t}} \ln \left(1 - \frac{\hat{t}}{m^2} \right) + F_{Box}^1 \right] \times |\mathcal{M}_0^{lq}|^2 \\ &+ \frac{\alpha_s}{2\pi} C_F \times F_1^{nonBorn} , \end{aligned} \quad (7.71)$$

a função $F_1^{nonBorn}$ contém todas as partes finitas que não são proporcionais ao elemento de Born. Termos finitos não proporcionais ao elemento de Born é um fenômeno teórico comum em processos cujas correções de ordem superior envolvem caixas.

2. Interferência $< \mathcal{M}_{Box} \otimes \mathcal{M}_0^{sm} >$

O procedimento de cálculo de $\mathcal{M}_{Box} \otimes \mathcal{M}_0^{sm}$ foi análogo ao da interferência $\mathcal{M}_{Box} \otimes \mathcal{M}_0^{lq}$. A estrutura é, na verdade bem parecida, com termos proporcionais às integrais I_1, I_2, I_3 e I_4 . No entanto, nesse caso, não ocorrem partes finitas não proporcionais ao elemento de Born. A seguir apresentamos o resultado

$$2\Re(\mathcal{M}_0^{sm*} \mathcal{M}_{box}) = \frac{\alpha_s}{4\pi} C_F N_\epsilon \left(\frac{\mu^2}{\hat{s}} \right)^\epsilon \left[-\frac{2}{\epsilon^2} + \frac{4}{\epsilon} \frac{m^2}{\hat{t}} \ln \left(1 - \frac{\hat{t}}{m^2} \right) + F_{Box}^2 \right] \times 2\Re(\mathcal{M}_0^{sm*} \mathcal{M}_0^{lq}) , \quad (7.72)$$

onde F_{Box}^2 representa as partes finitas.

Somando agora todas as contribuições virtuais em ordem α_s , temos o seguinte

$$\begin{aligned} |\mathcal{M}_0^{sm} + \mathcal{M}_0^{lq} + \mathcal{M}_V^{sm} + \mathcal{M}_V^{lq}|^2 &= |\mathcal{M}_0^{sm} + \mathcal{M}_0^{lq}|^2 + 2\Re(\mathcal{M}_0^{sm*} \mathcal{M}_V^{sm}) + 2\Re(\mathcal{M}_0^{sm*} \mathcal{M}_V^{lq}) \\ &+ 2\Re(\mathcal{M}_0^{lq*} \mathcal{M}_V^{sm}) + 2\Re(\mathcal{M}_0^{lq*} \mathcal{M}_V^{lq}) , \end{aligned} \quad (7.73)$$

onde

$$\mathcal{M}_V^{sm} = \mathcal{M}_{ae}^{sm} + \mathcal{M}_{vert}^{sm} \quad (7.74)$$

$$\mathcal{M}_V^{lq} = \mathcal{M}_{ae}^{lq} + \mathcal{M}_{vert}^{lq} + \mathcal{M}_{box} . \quad (7.75)$$

Substituindo essas amplitudes em (7.73) obtemos a correção em *next-to-leading-order* de QCD devida a troca de glúons virtuais ao processo de Drell-Yan com inclusão de leptoquarks escalares

$$\begin{aligned} \Delta_V &= 2\Re(\mathcal{M}_0^{sm*} \mathcal{M}_V^{sm}) + 2\Re(\mathcal{M}_0^{sm*} \mathcal{M}_V^{lq}) + 2\Re(\mathcal{M}_0^{lq*} \mathcal{M}_V^{sm}) + 2\Re(\mathcal{M}_0^{lq*} \mathcal{M}_V^{lq}) \\ &= \frac{\alpha_s}{2\pi} C_F N_\epsilon \left(\frac{\mu^2}{\hat{s}} \right)^\epsilon \left[\frac{1}{\epsilon} - \frac{2}{\epsilon^2} - \frac{4}{\epsilon} \right] |\mathcal{M}_0^{sm}|^2 \\ &\quad + \frac{\alpha_s}{4\pi} C_F N_\epsilon \left(\frac{\mu^2}{\hat{s}} \right)^\epsilon \left[\frac{1}{\epsilon} - \frac{4}{\epsilon} - \frac{4m^2}{\epsilon \hat{t}} \ln \left(1 - \frac{\hat{t}}{m^2} \right) - \frac{2}{\epsilon^2} + \frac{4m^2}{\epsilon \hat{t}} \ln \left(1 - \frac{\hat{t}}{m^2} \right) \right] \\ &\quad \times 2\Re(\mathcal{M}_0^{sm*} \mathcal{M}_0^{lq}) \\ &\quad + \frac{\alpha_s}{4\pi} C_F N_\epsilon \left(\frac{\mu^2}{\hat{s}} \right)^\epsilon \left[\frac{1}{\epsilon} - \frac{4}{\epsilon} - \frac{2}{\epsilon^2} \right] \times 2\Re(\mathcal{M}_0^{sm} \mathcal{M}_0^{lq*}) \\ &\quad + \frac{\alpha_s}{2\pi} C_F N_\epsilon \left(\frac{\mu^2}{\hat{s}} \right)^\epsilon \left[\frac{1}{\epsilon} - \frac{4}{\epsilon} - \frac{4m^2}{\epsilon \hat{t}} \ln \left(1 - \frac{\hat{t}}{m^2} \right) - \frac{2}{\epsilon^2} + \frac{4m^2}{\epsilon \hat{t}} \ln \left(1 - \frac{\hat{t}}{m^2} \right) \right] \\ &\quad \times |\mathcal{M}_0^{lq}|^2 + \text{partes finitas} . \end{aligned} \quad (7.76)$$

Note que os pólos simples multiplicados por $\ln(1 - \hat{t}/m^2)$ se cancelam entre as contribuições do vértice $S\ell q$ e da caixa. Somando todas elas, temos, finalmente

$$\Delta_V = \frac{\alpha_s}{2\pi} C_F N_\epsilon \left(\frac{\mu^2}{\hat{s}} \right)^\epsilon \left[-\frac{2}{\epsilon^2} - \frac{3}{\epsilon} \right] \times |\mathcal{M}_0|^2 + \Delta_V^{\text{finito}} , \quad (7.77)$$

onde Δ_V^{finito} contém todas as partes finitas do cálculo.

Mais uma vez, a exemplo de vários outros processos aos quais são calculadas correções de QCD em ordem α_s , constatamos que as divergências IR devidas a glúons virtuais são proporcionais ao elemento de Born e estão contidas no fator global

$$\frac{\alpha_s}{2\pi} C_F N_\epsilon \left(\frac{\mu^2}{\hat{s}} \right)^\epsilon \left[-\frac{2}{\epsilon^2} - \frac{3}{\epsilon} \right] . \quad (7.78)$$

O cancelamento das divergências IR foi também checado no gauge de Landau. O termo adicional do propagador do glúon mantém, como sabemos, o setor do modelo padrão inalterado. Com a inclusão dos leptoquarks escalares, as amplitudes nesse setor tiveram de ser recalculadas, na verdade, os cálculos se resumiram na soma do novo termo do propagador aos resultados do gauge de Feynman. Verificamos

que o vértice $S\ell q$ permaneceu inalterado em relação à estrutura de pólos, enquanto a autoenergia do leptoquark e a caixa receberam contribuições adicionais mas de sinais opostos o que garantiu, no final do dia, o cancelamento das divergências IR, a exemplo do gauge de Feynman.

7.4 Seção de Choque Virtual Total

A contribuição das amplitudes de glúons virtuais é, portanto, dada pela expressão

$$\hat{\sigma}_V = \frac{\alpha_s}{2\pi} C_F N_\epsilon \left(\frac{\mu^2}{\hat{s}} \right)^\epsilon \left[-\frac{2}{\epsilon^2} - \frac{3}{\epsilon} + F_V \right] \times \hat{\sigma}_0^{sm+lq} + \frac{\alpha_s}{2\pi} C_F \hat{\sigma}^{nonBorn} \quad (7.79)$$

nos gauges de Feynman e Landau. A função F_V representa as partes finitas em 4 dimensões e foi calculada apenas no gauge de Feynman, onde a sua forma é mais simples, enquanto que $\frac{\alpha_s}{2\pi} C_F \hat{\sigma}^{nonBorn}$ representa os termos não proporcionais ao elemento de Born. A seguir passaremos ao cálculo das amplitudes reais com a inclusão dos leptoquarks escalares.

7.5 Radiação de Glúons Reais

Em ordem α_s , os diagramas de Feynman relevantes ao cálculo das amplitudes de glúons reais são os da figura abaixo:

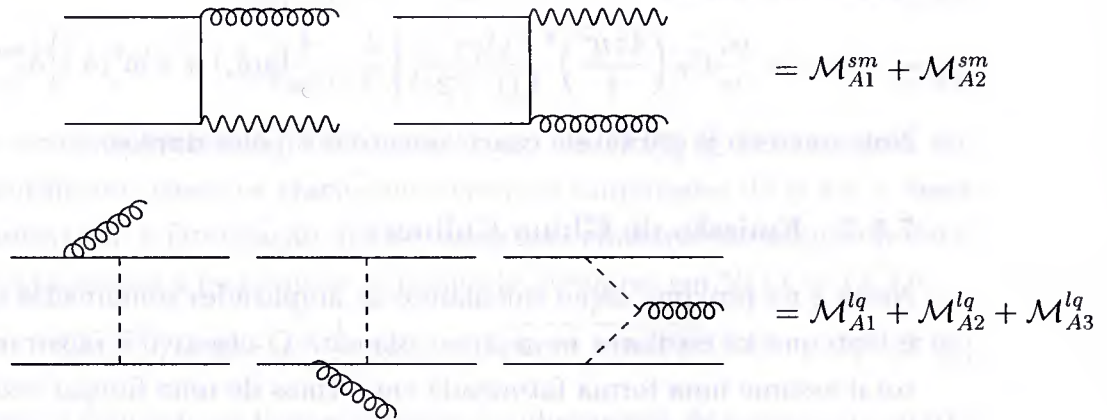


Figura 7.7: Radiação de glúons reais com a inclusão de leptoquarks escalares.

As amplitudes correspondentes a estes diagramas de Feynman são dadas por

$$\mathcal{M}_{A1}^{sm} = -ig_s T^a \frac{1}{\hat{t} \cdot k_V^2} \cdot \bar{v}_q(p_b) \gamma_\nu (V_q + A_q \gamma^5) (\not{p}_a - \not{p}_1) \not{\epsilon}^* u_q(p_a) \cdot \ell^\nu \quad (7.80)$$

$$\mathcal{M}_{A2}^{sm} = -ig_s T^a \frac{1}{\hat{u} \cdot k_V^2} \cdot \bar{v}_q(p_b) \not{\epsilon}^* (\not{p} - \not{p}_b) \gamma_\nu (V_q + A_q \gamma^5) u_q(p_a) \cdot \ell^\nu \quad (7.81)$$

$$\mathcal{M}_{A1}^{lq} = i\lambda^2 g_s T^a \frac{1}{\hat{t}(\hat{t}_{b3} - m^2)} \cdot \bar{v}_q(p_b) \gamma_- v_e(k_2) \bar{u}_e(k_1) \gamma_+ (\not{p}_a - \not{p}) \not{\epsilon}^* u_q(p_a) \quad (7.82)$$

$$\mathcal{M}_{A2}^{lq} = i\lambda^2 g_s T^a \frac{1}{\hat{u}(\hat{t}_1 - m^2)} \cdot \bar{v}_q(p_b) \not{\epsilon}^* (\not{p} - \not{p}_b) \gamma_- v_e(k_2) \bar{u}_e(k_1) \gamma_+ u_q(p_a) \quad (7.83)$$

$$\begin{aligned} \mathcal{M}_{A3}^{lq} &= i\lambda^2 g_s T^a \frac{1}{(\hat{t}_{b3} - m^2)(\hat{t}_1 - m^2)} \cdot \bar{v}_q(p_b) \gamma_- v_e(p_3) \bar{u}_e(k_1) \gamma_+ u_q(p_a) \\ &\quad \times (2p_a - 2p_2 - p_1) \cdot \epsilon^* , \end{aligned} \quad (7.84)$$

onde ℓ^ν é a corrente leptônica definida em (6.35) e os invariantes de Mandelstam são os mesmos da seção (6.3).

A amplitude correspondente ao diagrama onde o glúon é emitido a partir do leptoquark virtual, $\mathcal{M}_{A3}^{lq}$ é finita tanto regime soft $p \rightarrow 0$ quanto no regime colinear onde $\hat{t}, \hat{u} \rightarrow 0$ e pode ser desprezada na aproximação de glúon soft e no *leading pole approximation*, contribuindo de forma efetiva para σ_3 somente.

7.5.1 Emissão de Glúon Soft

O cálculo é exatamente igual ao do correspondente do modelo padrão e resulta na forma fatorizada já deduzida anteriormente

$$\mathcal{M}_{soft}^R \simeq 2g_s T^a \left[\frac{p_b^\mu}{\hat{u}} - \frac{p_a^\mu}{\hat{t}} \right] \epsilon_\mu^*(p) \times (\mathcal{M}_0^{sm} + \mathcal{M}_0^{lq}) . \quad (7.85)$$

Integrando o módulo-quadrado dessa expressão no espaço de fase $2 \rightarrow 3$ na aproximação soft, temos

$$\hat{\sigma}_{soft} = \frac{\alpha_s}{2\pi} C_F \left(\frac{4\pi\mu^2}{\hat{s}} \right)^\epsilon \frac{\Gamma(1-\epsilon)}{\Gamma(1-2\epsilon)} \left[\frac{2}{\epsilon^2} - \frac{4}{\epsilon} \ln(\delta_s) + 4 \ln^2(\delta_s) \right] \hat{\sigma}_0^{sm+lq}(\hat{s}) . \quad (7.86)$$

Note que isso já garante o cancelamento dos pólos duplos.

7.5.2 Emissão de Glúon Colinear

Nesta e na próxima seção calculamos as amplitudes combinadas do modelo padrão e leptoquarks escalares no regime colinear. O objetivo é mostrar que a amplitude total assume uma forma fatorizada em termos de uma função vezes a amplitude de Born, a exemplo do que ocorreu no setor do modelo padrão.

Para tanto, vamos reescrever as amplitudes de emissão de glúon (7.80)–(7.84) de uma maneira mais conveniente, a exemplo do que fizemos na seção (6.3.2), usando exatamente as mesmas aproximações

$$\mathcal{M}_{A1}^{sm} = g_s T^a \epsilon_\mu^* \cdot \mathcal{M}_{sm}^c \times \frac{(\not{p}_a - \not{p})}{\hat{t}} \gamma^\mu u_q(p_a) \quad (7.87)$$

$$\mathcal{M}_{A2}^{sm} = g_s T^a \epsilon_\mu^* \cdot \mathcal{N}_{sm}^\mu \times u_q(p_a) \quad (7.88)$$

$$\mathcal{M}_{A1}^{lq} = g_s T^a \epsilon_\mu^* \cdot \mathcal{M}_{lq}^c \times \frac{(\not{p}_a - \not{p})}{\hat{t}} \gamma^\mu u_q(p_a) \quad (7.89)$$

$$\mathcal{M}_{A2}^{lq} + \mathcal{M}_{A3}^{lq} = g_s T^a \epsilon_\mu^* \cdot \mathcal{N}_{lq}^\mu \times u_q(p_a) . \quad (7.90)$$

A amplitude total, dada pela soma dessas amplitudes, pode agora ser escrita de uma maneira mais útil ao cálculo do limite colinear, no caso, associada a $\hat{t} \rightarrow 0$

$$\mathcal{M}_F = g_s T^a \epsilon_\mu^* \cdot \left\{ [\mathcal{M}_{sm}^c + \mathcal{M}_{lq}^c] \frac{(\not{p}_a - \not{p})}{\hat{t}} \gamma^\mu + [\mathcal{N}_{sm}^\mu + \mathcal{N}_{lq}^\mu] \right\} , \quad (7.91)$$

onde a divergência está, agora, completamente isolada no primeiro termo da amplitude total (7.91).

O módulo quadrado, em *Leading Pole Approximation*, é da forma

$$\begin{aligned} \overline{|\mathcal{M}_F|}_{LPA}^2 &= -2(1 - \varepsilon) g_s^2 (T^a T^a) \frac{1}{\hat{t}} (1 - z) \left\{ \text{Tr}\{\not{p}_a \cdot |\mathcal{M}_{sm}^c|^2\} + \text{Tr}\{\not{p}_a \cdot 2\Re[\mathcal{M}_{sm}^c \mathcal{M}_{lq}^{c*}] \right. \\ &\quad \left. + \text{Tr}\{\not{p}_a \cdot |\mathcal{M}_{lq}^c|^2\} \right\} \\ &\quad - 2z g_s^2 (T^a T^a) \frac{1}{\hat{t}} \left\{ \text{Tr}\{\not{p}_a \cdot 2\Re[\mathcal{M}_{sm}^c \mathcal{N}_{sm}^* \cdot p_a]\} + \text{Tr}\{\not{p}_a \cdot [2\Re(\mathcal{M}_{sm}^c \mathcal{N}_{lq}^* \cdot p_a) \right. \\ &\quad \left. + 2\Re(\mathcal{M}_{lq}^c \mathcal{N}_{sm}^* \cdot p_a)] + \text{Tr}\{\not{p}_a \cdot 2\Re[\mathcal{M}_{lq}^c \mathcal{N}_{lq}^* \cdot p_a]\} \right\} . \end{aligned} \quad (7.92)$$

A invariância de gauge da amplitude total é usada para relacionar os produtos escalares $\mathcal{N}_{sm} \cdot p_a$ e $\mathcal{N}_{lq} \cdot p_a$ com $\mathcal{M}_{sm}^c$ e $\mathcal{M}_{lq}^c$, respectivamente. Outra vez, a exemplo do que ocorreu na seção (6.3.2), temos

$$\mathcal{N}_{sm} \cdot p_a = \frac{\mathcal{M}_{sm}^c}{1 - z} \quad (7.93)$$

$$\mathcal{N}_{lq} \cdot p_a = \frac{\mathcal{M}_{lq}^c}{1 - z} . \quad (7.94)$$

Os traços remanescentes em (7.92) após o uso de (7.93) e (7.94) acima são facilmente identificados como os traços que geram as amplitudes de Born, e dessa maneira, demonstra-se a fatorização das divergências colineares da amplitude total de emissão de glúon com a inclusão de leptoquarks escalares em NLO de QCD[§]

$$\overline{|\mathcal{M}_F|}_{col}^2 = 4\pi\alpha_s \mu^{2\varepsilon} \frac{1}{|\hat{t}|} 2P_{qq}(z, \varepsilon) \overline{|\mathcal{M}_0^{sm} + \mathcal{M}_0^{lq}|}^2 . \quad (7.95)$$

A próxima seção é dedicada ao limite colinear do subprocesso de emissão de quarks não-massivos em ordem α_s , o que nos permitirá a obtenção das contribuições Hard-Colinear e Soft-Colinear à seção de choque de 2 corpos, finalizando dessa maneira os cálculos analíticos do método de *NLO Monte Carlo* incluindo efeitos virtuais de leptoquarks.

[§]Quando o glúon é emitido colinearmente ao antiquark, devemos trocar $\hat{t}$ por $\hat{u}$ em (7.95).

7.6 Estados Finais com Um Quark Não-Massivo Adicional

Nesta seção apresentamos as reações partônicas que envolvem glúon no estado inicial e um (anti)quark não-massivo no estado final incluindo troca de leptoquarks escalares. A cinemática é mesma do setor do modelo padrão. Em ordem α_s tais reações são dadas por

$$g(p_a) + q(\bar{q})(p_b) \longrightarrow q(\bar{q})(p) + \gamma, Z(k) \longrightarrow q(\bar{q})(p) + e^-(k_1) + e^+(k_2) \quad (7.96)$$

$$g(p_a) + q(\bar{q})(p_b) \longrightarrow q(\bar{q})(p) + e^-(k_1) + e^+(k_2) \quad (7.97)$$

$$g(p_a) + q(\bar{q})(p_b) \longrightarrow e^-(k_1) + S(k) \longrightarrow e^-(k_1) + e^+(k_2) + q(\bar{q})(p) . \quad (7.98)$$

Os diagramas de Feynman correspondentes a essas reações estão ilustrados logo abaixo

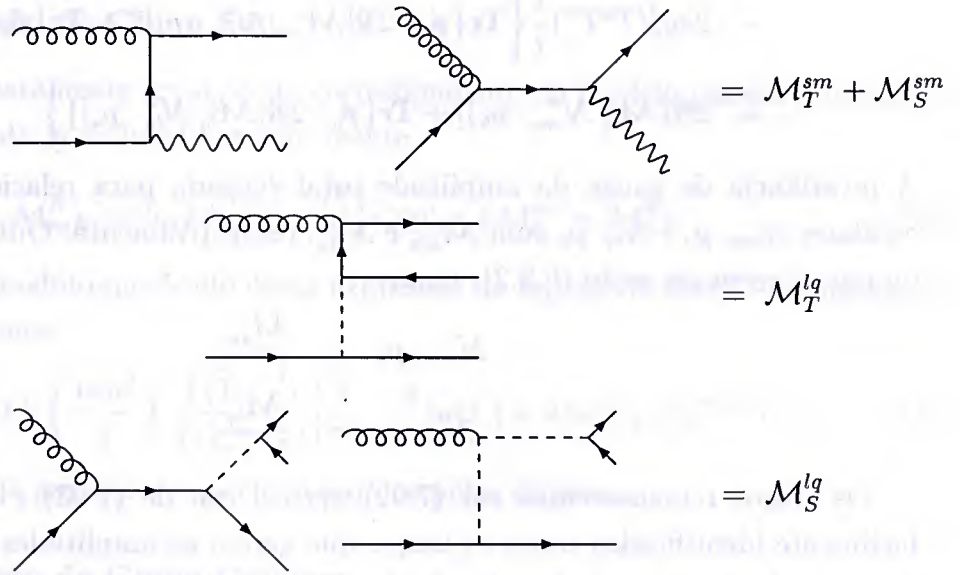


Figura 7.8: Glúon no estado inicial incluindo contribuições de leptoquarks escalares.

As reações do modelo padrão (7.96) correspondem às amplitudes $\mathcal{M}_T^{sm} + \mathcal{M}_S^{sm}$ com o subsequente decaimento dos bósons vetoriais. A reação (7.97) corresponde à amplitude $\mathcal{M}_T^{lq}$ onde o leptoquark é trocado no canal t da reação, enquanto (7.98) são reações onde o leptoquark é produzido no canal $\hat{s}_{13}$, definido em (6.49). Se $E_{cm} > m$ e $\hat{s}_{13} = m^2$, então podemos ter leptoquarks escalares sendo produzidos em sua camada de massa e, subsequentemente decaindo em pares $e^\pm q(\bar{q})$ nesta ordem de teoria de perturbação. Como a produção simples de leptoquarks pode ser distinguida em

laboratório é mais consistente estudar esse processo separadamente, apesar do estado final gerado ser o mesmo do processo de Drell-Yan. Na próxima seção vamos tratar desse problema detalhadamente. Por hora não vamos nos preocupar com isso e prosseguir com os cálculos.

As amplitudes são dadas explicitamente pelas seguintes expressões

$$\mathcal{M}_T^{sm} = -ig_s T^a \frac{1}{\hat{t} \cdot k_V^2} \cdot \bar{u}_q(p) \not{\epsilon}(\not{p} - \not{p}_a) \gamma_\mu (V_q + A_q \gamma^5) u_q(p_b) \cdot \ell^\mu, \quad (7.99)$$

$$\mathcal{M}_S^{sm} = -ig_s T^a \frac{1}{\hat{s} \cdot k_V^2} \cdot \bar{u}_q(p) \gamma_\mu (V_q + A_q \gamma^5) (\not{p}_a + \not{p}_b) \not{\epsilon} u_q(p_b) \cdot \ell^\mu, \quad (7.100)$$

$$\mathcal{M}_T^{lq} = i\lambda^2 g_s T^a \frac{1}{\hat{t}(\hat{t}_2 - m^2)} \times \bar{u}_q(p) \not{\epsilon}(\not{p} - \not{p}_a) \gamma_{-v_e}(k_1) \cdot \bar{u}_e(k_2) \gamma_{+u_q}(p_b) \quad (7.101)$$

$$\mathcal{M}_S^{sm} = i\lambda^2 g_s T^a \frac{1}{(\hat{s}_{13} - m^2)} \times \bar{u}_q(p) \gamma_{-v_e}(k_1) \cdot \bar{u}_e(k_2) \gamma_{+}(\not{P} \cdot \epsilon) u_q(p_b), \quad (7.102)$$

onde

$$P^\mu = \frac{(\not{p}_a + \not{p}_b)}{\hat{s}} \gamma^\mu + \frac{Q^\mu}{\hat{t}_2 - m^2}, \quad (7.103)$$

$$Q^\mu = (2p_a + p_b - 2k_1)^\mu. \quad (7.104)$$

Vamos agora redefinir essas amplitudes da seguinte maneira

$$\mathcal{M}_T^{sm} = g_s T^a \epsilon_\mu \cdot \bar{u}_q(p) \gamma^\mu \frac{(\not{p} - \not{p}_a)}{\hat{t}} \times \mathcal{M}_{sm}^c, \quad (7.105)$$

$$\mathcal{M}_S^{sm} = g_s T^a \epsilon_\mu \cdot \bar{u}_q(p) \times \mathcal{N}_{sm}^\mu, \quad (7.106)$$

$$\mathcal{M}_T^{lq} = g_s T^a \epsilon_\mu \cdot \bar{u}_q(p) \gamma^\mu \frac{(\not{p} - \not{p}_a)}{\hat{t}} \times \mathcal{M}_{lq}^c, \quad (7.107)$$

$$\mathcal{M}_S^{lq} = g_s T^a \epsilon_\mu \cdot \bar{u}_q(p) \times \mathcal{N}_{lq}^\mu. \quad (7.108)$$

Somando todas amplitudes, temos uma expressão muito conveniente ao cálculo do limite colinear

$$\mathcal{M}_I = g_s T^a \epsilon_\mu \cdot \bar{u}_q(p) \left\{ \gamma^\mu \frac{(\not{p} - \not{p}_a)}{\hat{t}} [\mathcal{M}_{sm}^c + \mathcal{M}_{lq}^c] + [\mathcal{N}_{sm}^\mu + \mathcal{N}_{lq}^\mu] \right\}, \quad (7.109)$$

onde, agora, a divergência associada ao limite $\hat{t} \rightarrow 0$ está completamente contida no primeiro termo de $\mathcal{M}_I$, a exemplo de $\mathcal{M}_F$

Tomando o módulo quadrado dessa amplitude em *Leading Pole Approximation*, temos

$$\begin{aligned} |\overline{\mathcal{M}}_I|_{LPA}^2 &= -2(1 - \varepsilon) g_s^2 (T^a T^a) \frac{1}{\hat{t}} \left\{ \text{Tr}\{\not{p}_a \cdot |\mathcal{M}_{sm}^c|^2\} + \text{Tr}\{\not{p}_a \cdot 2\Re[\mathcal{M}_{sm}^c \mathcal{M}_{lq}^{c*}] \right. \\ &\quad \left. + \text{Tr}\{\not{p}_a \cdot |\mathcal{M}_{lq}^c|^2\} \right\} \end{aligned}$$

$$\begin{aligned}
& + 2z(1-z)g_s^2(T^a T^a) \frac{1}{\hat{t}} \left\{ \text{Tr}\{\not{p}_a \cdot 2\Re[\mathcal{M}_{sm}^c \mathcal{N}_{sm}^* \cdot p_a]\} + \text{Tr}\{\not{p}_a \cdot [2\Re(\mathcal{M}_{sm}^c \mathcal{N}_{lq}^* \cdot p_a) \right. \\
& \left. + 2\Re(\mathcal{M}_{lq}^c \mathcal{N}_{sm}^* \cdot p_a)] + \text{Tr}\{\not{p}_a \cdot 2\Re[\mathcal{M}_{lq}^c \mathcal{N}_{lq}^* \cdot p_a]\} \right\} . \quad (7.110)
\end{aligned}$$

Pela propriedade de invariância de gauge das amplitudes, podemos demonstrar, a exemplo das seções (6.3.2) e (6.4) que

$$\mathcal{N}_{sm} \cdot p_a = \mathcal{M}_{sm}^c \quad (7.111)$$

$$\mathcal{N}_{lq} \cdot p_a = \mathcal{M}_{lq}^c . \quad (7.112)$$

Os traços remanescentes após a substituição de (7.111) e (7.112) em (7.110) são facilmente identificados como os traços que geram as amplitudes de Born, e dessa forma, demonstra-se a fatorização das divergências colineares com a introdução de leptoquarks escalares em ordem α_s

$$\overline{|\mathcal{M}_I|}_{col}^2 = 4\pi\alpha_s\mu^{2\epsilon} \frac{1}{|\hat{t}|} 2P_{qg}(z, \epsilon) \overline{|\mathcal{M}_0^{sm} + \mathcal{M}_0^{lq}|}^2 . \quad (7.113)$$

A abordagem que utilizamos aqui pode ser usada para a demonstração de teoremas de fatorização de divergências soft e colineares para uma vasta gama de processos [51]. Para processos que envolvem múltiplos pártons no estado inicial ou final, e apenas um pártão em regime soft ou colinear existem também teoremas de fatorização de divergências, veja, por exemplo [38] e [26].

7.7 Termos Hard-Colinear

Demonstrados a fatorização de divergências colineares, as seções de choque colineares podem ser integradas em (6.107) exatamente da mesma maneira que fizemos no setor do modelo padrão. As divergências colineares são, então, absorvidas na redefinição das funções de distribuição de quarks pelo procedimento de fatorização de massa explicado na seção (6.6). A seção de choque Hard-Colinear obtida nesse caso é dada por (6.145) com $\hat{\sigma}_0^{sm}$ trocada por $\hat{\sigma}_0^{sm+lq}$

$$\begin{aligned}
\sigma_{HC} = & \frac{\alpha_s}{2\pi} \sum_q \int dx_a dx_b \times \\
& \left\{ q(x_a, Q^2) \int_{x_b}^{1-\delta_s} \frac{dz}{z} \bar{q}\left(\frac{x_b}{z}, Q^2\right) \tilde{P}_{qg \leftarrow q}(z) + q(x_a, Q^2) \int_{x_b}^1 \frac{dz}{z} g\left(\frac{x_b}{z}, Q^2\right) \tilde{P}_{q\bar{q} \leftarrow g}(z) \right. \\
& + \bar{q}(x_b, Q^2) \int_{x_a}^{1-\delta_s} \frac{dz}{z} q\left(\frac{x_a}{z}, Q^2\right) \tilde{P}_{qg \leftarrow q}(z) + \bar{q}(x_b, Q^2) \int_{x_a}^1 \frac{dz}{z} g\left(\frac{x_a}{z}, Q^2\right) \tilde{P}_{q\bar{q} \leftarrow g}(z) \\
& \left. + a \leftrightarrow b \right\} \times \hat{\sigma}_0^{sm+lq}(z\hat{s}) , \quad (7.114)
\end{aligned}$$

onde as funções $\tilde{P}_{qg}$ e $\tilde{P}_{q\bar{q}}$ são exatamente as mesmas dadas por (6.146)–(6.147).

7.8 Termo Soft-Colinear

Da mesma forma que σ_{HC} , a seção de choque Soft-Colinear não sofre nenhuma alteração além da troca de $\hat{\sigma}_0^{sm}$ por $\hat{\sigma}_0^{sm+lq}$, graças ao fato de que os teoremas de fatorização de divergências colineares continuam válidos com a inclusão de leptoquarks escalares. Logo

$$\sigma_{SC} = \sum_q \int dx_a dx_b [q(x_a, Q^2) \bar{q}(x_b, Q^2) + a \leftrightarrow b] \times \hat{\sigma}_0^{sm+lq}(\hat{s}) \times \frac{\alpha_s}{2\pi} C_F N_\epsilon \left(\frac{\mu^2}{Q_F^2} \right)^\epsilon \left\{ \frac{1}{\epsilon} (3 + 4 \ln(\delta_s)) + 9 + \frac{2\pi^2}{3} + 3 \ln(\delta_s) - 2 \ln^2(\delta_s) \right\} . \quad (7.115)$$

7.9 Seção de Choque Total Hadrônica de 2 Corpos

Tendo já calculado todos os termos que contribuem à seção de choque de 2 corpos, temos, finalmente

$$\sigma_2 = \sigma^{LL} + \sigma_{HC} + \sigma_{virt} + \sigma_{soft} + \sigma_{SC} , \quad (7.116)$$

cujas expressões completas são dadas por

$$\begin{aligned} \sigma_2 = & \sum_q \int dx_a dx_b [q(x_a, Q_F^2) \bar{q}(x_b, Q_F^2) + a \leftrightarrow b] \times \hat{\sigma}_0^{sm+lq}(\hat{s}) \times \\ & \left\{ 1 + \frac{\alpha_s}{2\pi} C_F \left[(3 + 4 \ln(\delta_s)) \ln\left(\frac{\hat{s}}{Q_F^2}\right) + 9 + \frac{2\pi^2}{3} + F_V + 3 \ln(\delta_s) + 2 \ln^2(\delta_s) \right] \right\} \\ & + \frac{\alpha_s}{2\pi} \sum_q \int dx_a dx_b \times \\ & \left\{ q(x_a, Q_F^2) \left(\int_{x_b}^{1-\delta_s} \frac{dz}{z} \bar{q}\left(\frac{x_b}{z}, Q_F^2\right) \tilde{P}_{qq}(z) + \int_{x_b}^1 \frac{dz}{z} g\left(\frac{x_b}{z}, Q_F^2\right) \tilde{P}_{qg}(z) \right) \times \hat{\sigma}_0^{sm+lq}(z\hat{s}) \right. \\ & + \bar{q}(x_b, Q_F^2) \left(\int_{x_a}^{1-\delta_s} \frac{dz}{z} \bar{q}\left(\frac{x_a}{z}, Q_F^2\right) \tilde{P}_{qq}(z) + \int_{x_a}^1 \frac{dz}{z} g\left(\frac{x_a}{z}, Q_F^2\right) \tilde{P}_{qg}(z) \right) \times \hat{\sigma}_0^{sm+lq}(z\hat{s}) \\ & \left. + a \leftrightarrow b \right\} + \frac{\alpha_s}{2\pi} C_F \sigma_{nonBorn} . \quad (7.117) \end{aligned}$$

Mais uma vez verificamos que a seção de choque total hadrônica inclusiva é livre de divergências infravermelhas em 4 dimensões espaço-temporais. Contudo, as partes finitas são sensivelmente dependentes dos cut-offs δ_s e δ_c . A função F_V inclui todas as partes finitas das amplitudes virtuais, e das expansões de funções Γ e outros termos dependentes de ϵ . O termo $\frac{\alpha_s}{2\pi} C_F \sigma^{nonBorn}$ representa todos os termos finitos não proporcionais ao elemento de Born, provenientes da amplitude virtual da caixa, integrados no espaço de fase de 2 partículas e nas frações de momento dos párons. A exemplo do que fizemos no final do capítulo 6, vamos resumir a estrutura de pólos IR, dessa vez com a inclusão de leptoquarks escalares. Na tabela (7.1) temos

as contribuições divergentes no IR para cada uma das amplitudes singulares no gauge de Feynman e no gauge de Landau, são elas: δ_{AE}^q —autoenergia dos quarks, $\delta_{AE}^{\ell q}$ —autoenergia do leptoquark, δ_{vert}^{sm} —correção do vértice $(\gamma, Z) q\bar{q}$, $\delta_{vert}^{\ell q}$ —correção do vértice Slq , δ_{Box} —correção virtual da caixa, δ_{Soft} —emissão de glúon soft e, finalmente, δ_{SC} —termo soft-colinear.

δ	Gauge de Feynman ($\eta = 0$)	Gauge de Landau ($\eta = -1$)
δ_{AE}^q	$\frac{1}{\epsilon}$	0
$\delta_{AE}^{\ell q}$	0	0
δ_{vert}^{sm}	$-\frac{2}{\epsilon^2} - \frac{4}{\epsilon}$	$-\frac{2}{\epsilon^2} - \frac{3}{\epsilon}$
$\delta_{vert}^{\ell q}$	$-\frac{4}{\epsilon} - \frac{4}{\epsilon} \frac{m^2}{\hat{t}} \ln(1 - \hat{t}/m^2)$	$-\frac{4}{\epsilon} - \frac{4}{\epsilon} \frac{m^2}{\hat{t}} \ln(1 - \hat{t}/m^2)$
δ_{Box}	$-\frac{2}{\epsilon^2} + \frac{4}{\epsilon} \frac{m^2}{\hat{t}} \ln(1 - \hat{t}/m^2)$	$-\frac{2}{\epsilon^2} + \frac{1}{\epsilon} + \frac{4}{\epsilon} \frac{m^2}{\hat{t}} \ln(1 - \hat{t}/m^2)$
δ_{Soft}	$\frac{2}{\epsilon^2} - \frac{4}{\epsilon} \ln \delta_s$	$\frac{2}{\epsilon^2} - \frac{4}{\epsilon} \ln \delta_s$
δ_{SC}	$\frac{1}{\epsilon} (3 + 4 \ln \delta_s)$	$\frac{1}{\epsilon} (3 + 4 \ln \delta_s)$

Tabela 7.1: Estrutura de divergências IR no processo de Drell-Yan $p\bar{p} \rightarrow e^+e^-X$ em NLO de QCD com inclusão de leptoquarks escalares.

Novamente, a exemplo do que fizemos na seção (6.7), podemos mostrar explicitamente o cancelamento das divergências IR resumidas na tabela (7.1). Para o setor do modelo padrão somente, temos, a partir de (6.163), que $\Delta^{sm} = 0$, de forma que todos os termos divergentes proporcionais a $|\overline{\mathcal{M}}_0^{sm}|^2$ se cancelam na soma. Para os termos divergentes proporcionais a $|\overline{\mathcal{M}}_0^{lq}|^2$, a soma dos termos divergentes é dada por

$$\Delta^{lq} = \delta_{AE}^q + \delta_{AE}^{\ell q} + \delta_{vert}^{\ell q} + \delta_{Box} + \delta_{Soft} + \delta_{SC} , \quad (7.118)$$

e pela tabela (7.1) podemos mostrar que $\Delta^{lq} = 0$, demonstrando de forma explícita o cancelamento dos pólos IR. Os pólos proporcionais ao termo de interferência $2\Re(\mathcal{M}_0^{sm\dagger} \mathcal{M}_0^{lq})$ têm a sua soma dada por $\Delta^{sm} + \Delta^{lq} = 0$, provando de forma explícita a finitude da seção de choque total inclusiva na região do infravermelho nos gauges de Feynman e de Landau, após a inclusão de leptoquarks escalares.

7.10 Seção de Choque Total Hadrônica de 3 Corpos

A contribuição à seção de choque total de 3 corpos provém dos subprocessos $q\bar{q} \rightarrow ge^+e^-$ dada pela soma de amplitudes

$$\mathcal{M}_A = \mathcal{M}_{A1}^{sm} + \mathcal{M}_{A2}^{sm} + \mathcal{M}_{A1}^{lq} + \mathcal{M}_{A2}^{lq} + \mathcal{M}_{A3}^{lq} \quad (7.119)$$

e $q(\bar{q})g \rightarrow q(\bar{q})e^+e^-$, cuja amplitude correspondente é da forma

$$\mathcal{M}_C = \mathcal{M}_T^{sm} + \mathcal{M}_S^{sm} + \mathcal{M}_T^{lq} + \mathcal{M}_S^{lq} . \quad (7.120)$$

O módulo quadrado dessas amplitudes em 4 dimensões espaço-temporais, calculadas com o auxílio do pacote FeynCalc para Mathematica, foram geradas diretamente em formato FORTRAN de modo a otimizar sua implementação numérica nos Monte Carlos na região $\overline{HC}$ do espaço de fase definida por

$$\delta_s \frac{\sqrt{\hat{s}}}{2} < E_g < E_{g,max} \quad (7.121)$$

$$\delta_c \hat{s} < |\hat{t}|, |\hat{u}| < |\hat{t}|_{max}, |\hat{u}|_{max} , \quad (7.122)$$

onde E_g é energia do glúon no processo de emissão de glúon.

Ainda assim uma última divergência, de caráter cinemático, resta após eliminarmos as divergências UV e IR. Veremos como isso ocorre na próxima seção.

7.10.1 Subtração dos Estados Intermediários On-Shell

Além das singularidades UV e IR, um outro tipo de singularidade, de natureza cinemática, ocorre em *next-to-leading-order* pela ocasião da inclusão dos leptoquarks e corresponde à produção simples dessas partículas no canal de decaimento e^-q onde $\hat{s}_{13} = m^2$, desde que $\sqrt{\hat{s}} > m$. Tal processo tem o mesmo estado final do processo de Drell-Yan, mas não interfere com as amplitudes que estamos calculando pois corresponde a um processo de produção e, como tal, pode ser calculada consistentemente à parte de nossas discussões. Ainda assim, leptoquarks *off-shell* podem contribuir em ordem α_s . Fenômeno semelhante ocorre nas correções de NLO em QCD supersimétrica aos processos de produção de squarks e gluinos [27, 28]. A análise desenvolvida nessa seção está baseada nessas referências.

Tais singularidades surgem quando somamos a contribuição dos diagramas de Feynman da figura (7.8) que dão origem à amplitude $\mathcal{M}_S^{lq}$. A produção simples de leptoquarks, no entanto, pode ser calculada consistentemente em leading e next-to-leading-order à parte do processo de Drell-Yan que estamos estudando, como já dissemos. Estados intermediários de leptoquarks *off shell*, no entanto, são uma contribuição legítima em NLO de QCD. Porém, quando integrarmos as contribuições de ordem α_s dos processos de emissão de quarks não-massivos sobre o espaço de fase de 3 partículas, o pólo correspondente à configuração onde $\hat{s}_{13} = m^2$ surgirá. É preciso, portanto, subtrair essa contribuição.

Para tanto, é preciso, em primeiro lugar, regularizar o pólo pela introdução de uma largura de decaimento para o leptoquark Γ_S . Isso é feito pela substituição do

propagador do leptoquark pela Breit-Wigner correspondente, ou seja

$$\frac{1}{\hat{s}_{13} - m^2} \rightarrow \frac{1}{\hat{s}_{13} - m^2 + im\Gamma_S} . \quad (7.123)$$

A princípio, a introdução de uma Breit-Wigner da forma com que fizemos acima, não é aconselhável, pois isso destrói a invariância de gauge das amplitudes. Para ilustrar esse fato, considere o processo de espalhamento Compton de um fóton por um leptoquark escalar de carga elétrica Q_S

$$\gamma(k_1) + S(p_1) \rightarrow \gamma(k_2) + S(p_2) , \quad (7.124)$$

cujos diagramas de Feynman são os da figura abaixo.

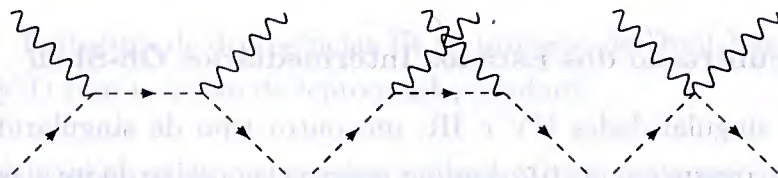


Figura 7.9: Espalhamento Compton $\gamma + S \rightarrow \gamma + S$.

As regras de Feynman são as da eletrodinâmica escalar que podem ser obtidas a partir das regras de Feynman do apêndice A fazendo $T^a = 1$, dado que o grupo de simetria é o $U(1)_{em}$. Dessa forma podemos mostrar que, introduzindo a forma da Breit-Wigner na expressão do propagador do leptoquark, a invariância de gauge da amplitude é quebrada da seguinte maneira[¶]

$$\mathcal{M}^\mu k_{1\mu} = -Q_S^2 \frac{m\Gamma_S}{s - m^2 + i\Gamma_S m} \times (2p_2 + k_2) \cdot \epsilon^*(k_2) . \quad (7.125)$$

Se a largura de decaimento do leptoquark for suficientemente fina, o que é plausível para uma partícula tão pesada quanto o leptoquark, então, fora da região ressonante, o lado direito de (7.125) é suprimido cinematicamente de forma que invariância de gauge é fracamente violada.

[¶]A procura de uma forma analítica do propagador de Breit-Wigner consistente com a invariância de gauge tem sido objeto de vários estudos no Modelo Padrão [69].

Denotando a contribuição ressonante ao elemento de matriz, definido por $p^2 = m^2$, por $\mathcal{M}_{res}$, e a soma das contribuições off-shell e dos outros diagramas por $\mathcal{M}_{DY}$, o quadrado do elemento de matriz pode ser decomposto da seguinte forma

$$|\mathcal{M}|^2 = |\mathcal{M}_{res}|^2 + 2\Re[\mathcal{M}_{res}\mathcal{M}_{DY}^*] + |\mathcal{M}_{DY}|^2 . \quad (7.126)$$

Integrando sobre todo o espaço de fase do estado final e^+e^-q , a contribuição ressonante representa a seção de choque de produção de $S + e^-$ em nível de árvore, com o posterior decaimento do leptoquark $S \rightarrow e^+ + q$, enquanto o restante deve ser atribuído às correções radiativas de ordem α_s ao processo de Drell-Yan com efeitos virtuais de leptoquarks escalares no canal t da reação.

A implementação numérica da radiação de quarks não-massivos em NLO procede, especificamente, da seguinte forma. Em primeiro lugar, sendo $\Gamma_S \ll m$ podemos utilizar a *narrow width approximation*, onde

$$-i \frac{\Gamma m}{(\hat{s}_{13} - m^2)^2 + \Gamma^2 m^2} \simeq -i\pi\delta(\hat{s}_{13} - m^2) \quad (7.127)$$

para calcular as contribuições onde o leptoquark aparece no canal de produção $\hat{s}_{13}$.

O quadrado do elemento de matriz completo pode ser escrito como

$$|\mathcal{M}_I|^2 = \frac{F(\hat{s}_{13})}{(\hat{s}_{13} - m^2)^2 + \Gamma^2 m^2} , \quad (7.128)$$

de forma que a subtração dos estados intermediários on-shell é definida por

$$\frac{F(\hat{s}_{13})}{(\hat{s}_{13} - m^2)^2 + \Gamma^2 m^2} - \frac{F(m^2)}{(\hat{s}_{13} - m^2)^2 + \Gamma^2 m^2} \Theta(\hat{s} - m^2) . \quad (7.129)$$

Como o fator global $\delta(\hat{s}_{13} - m^2)$ está ausente no termo de subtração, o propagador de Breit-Wigner tem de ser integrado sobre todo o espaço de fase da variável $\hat{s}_{13}$.

Nossas simulações numéricas demonstraram que a contribuição da amplitude de produção de leptoquarks à seção de choque de 3 corpos σ_3 , de fato, é muito pequena.

A seguir apresentamos nossas análises numéricas para o processo de Drell-Yan $p\bar{p} \rightarrow e^+e^-X$ em NLO de QCD incluindo os efeitos virtuais de ordem λ^2 e λ^4 do acoplamento de Yukawa no Tevatron.

Capítulo 8

Análises Numéricas

Discutimos neste capítulo as implicações fenomenológicas das correções de QCD em NLO para o processo de Drell-Yan $p\bar{p} \rightarrow e^+e^-X$ com inclusão de efeitos virtuais de leptoquarks escalares no Tevatron ($\sqrt{S} = 1.8$ TeV).

Os parâmetros eletrofracos utilizados foram: $M_Z = 91.187$ GeV, $\Gamma_Z = 2.490$ GeV, $\sin^2 \theta_W = 0.23124$, $\alpha = 7.3 \times 10^{-3}$ e $G_F = 1.633 \times 10^{-5}$ GeV⁻². A constante de estrutura fina da QCD foi obtida diretamente da CERNLIB. Para as PDF's utilizamos o conjunto CTEQ4D [70] com $\Lambda_4 = 296$ GeV no esquema *DIS* em NLO, portanto na implementação numérica dos resultados assumimos $\lambda_{FC} = 1$. A não ser quando dito em contrário, as escalas de renormalização e fatorização são dadas por uma escala única $Q^2 = Q_R^2 = Q_F^2 = M_Z^2$. Os programas numéricos foram escritos em FORTRAN 77 utilizando o pacote VEGAS [71] para o cálculo das integrais multidimensionais.

8.1 Variação da Seção de Choque Total com os Cut-Offs Teóricos δ_s e δ_c

Como discutimos no capítulo 4, a seção de choque total inclusiva obtida pela soma das seções de 2 e 3 corpos deve ser insensível à variação dos cut-offs soft e colinear, δ_s e δ_c , respectivamente. Nessa seção, estudamos a variação de $\sigma_{total} = \sigma_2 + \sigma_3$, com δ_s e δ_c , no esquema *DIS*, utilizando um corte para a massa invariante do par e^+e^- dado por $M(e^+e^-) \geq 120$ GeV. Os resultados são apresentados nas figuras (8.1) e (8.2), e mostram que a seção de choque total é extremamente estável por uma variação de quase 3 ordens de magnitude em δ_c e de 2 ordens de magnitude em δ_s . É possível demonstrar por argumentos cinemáticos que δ_s deve ser maior do que δ_c sempre, o que foi levado em consideração nas simulações numéricas e na obtenção das figuras (8.1) e (8.2).

Note que a seção de choque de 2 corpos é negativa para quase todos os valores

dos cut-off's devido à dependência de σ_2 com logaritmos de δ_s e δ_c , como pode ser visto nas expressões (6.162) e (7.117). A seção de choque de 3 corpos σ_3 , por outro lado é sempre positiva e maior em magnitude do que σ_2 , o que resulta em uma seção de choque total positiva.

8.2 Efeitos das Correções de QCD Sobre as Seções de Choque do Processo de Drell-Yan

O efeito das correções em NLO sobre as seções de choque em Leading-Logarithm (LL) pode ser avaliado pelo cálculo dos chamados *K-factors*, definidos por

$$K = \frac{\sigma^{NLO}}{\sigma^{LL}} . \quad (8.1)$$

Para o processo de Drell-Yan no Tevatron, com $M(e^+e^-) \geq 120$ GeV, o K-factor obtido pela razão das seções de choque totais a partir de (8.1) é de aproximadamente $K = 1.41$, considerando apenas o espectro do modelo padrão, o que mostra que as correções de QCD são positivas e aumentam as seções de choque em LL substancialmente, o que é típico em correções de QCD. Esse valor é um pouco maior do que o K-factor de Drell-Yan encontrado na literatura, que vale aproximadamente $K = 1.33$ [46], no entanto esse valor é obtido no esquema $\overline{MS}$, enquanto estamos trabalhando no esquema DIS, além disso diferentes cortes na massa invariante do par e^+e^- podem produzir K-factors distintos. Na ref. [72], onde correções de QCD ao processo de Drell-Yan no Tevatron e no LHC são calculadas em ordem α_s^2 , mostra-se que a seção de choque de Born é um pouco maior no esquema $\overline{MS}$ em relação ao DIS, enquanto as seções de choque corrigidas são menores tanto em ordem α_s quanto em α_s^2 , o que faz do K-factor calculado em $\overline{MS}$ de fato menor do que o calculado no DIS. O K-factor obtido de seções de choque incluindo leptosquarks escalares é essencialmente o mesmo.

O K-factor, na realidade, não é uma constante, dependendo da região do espaço de fase e das escalas de renormalização/fatorização. Nas figuras (8.3) e (8.4) mostramos os K-factors em função da massa invariante do par e^+e^- $M(e^+e^-)$, e do momento transversal do elétron do estado final $P_T(e^-)$. Para $M(e^+e^-)$ crescente, o K-factor tem uma tendência de queda, por outro lado K em função de $P_T(e^-)$ tem um comportamento aproximadamente constante.

8.3 Estabilidade das Seções de Choque em NLO em Relação à Escolha das Escalas de Renormalização e Fatorização

Uma das principais motivações para o cálculo de correções radiativas de ordem superior em QCD é a necessidade de reduzir a dependência das seções de choque em Leading-Log em relação à escolha das escalas de renormalização Q_R e fatorização Q_F . Em teoria de perturbação, quando truncamos a série em uma dada ordem $\mathcal{O}(\alpha_s^n)$, espera-se uma dependência explícita dos observáveis físicos nas escalas e esquemas adotados em ordem $\mathcal{O}(\alpha_s^{n+1})$. Portanto quando desprezamos essas correções de ordem superior, uma dependência residual ocorre na ordem de teoria de perturbação que estamos trabalhando. Em relação ao processo de Drell-Yan, a seção de choque de Born calculada em Leading-Log só depende da escala de fatorização, inserida via PDF's. No Tevatron, o processo de Drell-Yan é dominado pelo aniquilamento de quarks de valência cujas funções de distribuição têm pouca dependência em Q_F^2 para os valores de x acessados naquele acelerador de partículas. Logo, as seções de choque em LL apresentam pouca dependência com a escolha das escalas. Já em NLO, uma dependência adicional é introduzida pela constante de estrutura fina forte, que depende da escala de renormalização, e que decresce com o aumento de Q_R . Os resultados do Tevatron são, portanto, exemplos de observáveis cuja dependência na escolha das escalas aumenta em NLO. Outros exemplos onde isso ocorre podem ser encontrados na produção hadrônica de $W^\pm\gamma$ e $Z\gamma$ [32] e produção direta de fótons [35]. Constatamos que o mesmo ocorre no processo de Drell-Yan com correções em NLO no esquema DIS, sem inclusão de leptarquarks. Enquanto a seção de choque em LL decresce cerca de 11% quando variamos $Q = Q_R = Q_f$ de $0.5M_Z$ até $5.5M_Z$, a seção de choque em NLO varia cerca de 14% no mesmo intervalo. Isso está ilustrado na figura (8.5) onde comparamos as seções de choque em LL e NLO.

Note que para grandes valores de Q^2 as seções de choque variam cada vez menos, tendendo a um comportamento constante com a variação da escala única.

Apesar de ser frequentemente usada na literatura, a escolha $Q = \sqrt{s}$ não é sempre consistente. Na ref. [40] mostra-se que a escolha de qualquer escala dependente das frações de energia dos pártons implica um erro de ordem α_s nas seções de choque não importando o quão exato é o cálculo perturbativo. Isso se deve ao fato de que a escolha $Q = \sqrt{s}$ quebra a invariância da seção de choque em relação às transformações do grupo de renormalização. Logo, quando trabalhamos em um nível superior ao nível de Born, é preferível escolher uma escala típica fixa.

As seções de choque também dependem da escolha das PDF's utilizadas na convolução com as seções de choque partônicas. Na tabela (8.1) apresentamos os

valores das seções de choque totais em LL e NLO, com o corte $M(e^+e^-) \geq 120$ GeV, em relação a três diferentes escolhas de parametrizações de PDF's em NLO no esquema DIS.

σ (pb)	CTEQ4D	GRV94 (HO)	MRS(H)
σ^{LL}	1.809	1.893	1.820
σ^{NLO}	2.553	2.675	2.585
$K = \frac{\sigma^{NLO}}{\sigma^{LL}}$	1.41	1.41	1.42

Tabela 8.1: Seções de choque totais em LL e NLO approximation e K-factors em função de diferentes parametrizações de PDF's: CTEQ4D [70], GRV94 (HO) [73] e MRS(H) [74] todas obtidas no esquema DIS.

Em LL e NLO, a variação do maior valor (GRV94) para o menor (CTEQ4D) valor é de aproximadamente 5% apenas. Note que os K-factors praticamente não variam de um caso para o outro. A seguir apresentamos algumas distribuições com a inclusão de efeitos virtuais de leptoquarks em NLO de QCD.

8.4 Distribuições

Nesta seção apresentamos algumas distribuições em NLO de QCD incluindo efeitos virtuais de leptoquarks escalares de 1ª geração no processo de Drell-Yan $p\bar{p} \rightarrow e^+e^-X$ no Tevatron onde $\sqrt{S} = 1.8$ TeV. Analisamos aqui os efeitos da inclusão de apenas um tipo de leptoquark de número fermiônico $F = 0$, a saber, aquele que tem acoplamentos do tipo Sue^+ e $S^*\bar{u}e^-$ envolvendo o quark up . Esse tipo de leptoquark pode ser identificado com os leptoquarks R_2 da lagrangiana efetiva (7.12) e possui uma carga elétrica fracionária $5/3$ em unidades da carga do elétron.

Em todos os casos, o único corte usado foi o da massa invariante do par e^+e^- , dado por

$$M(e^+e^-) \geq 120 \text{ GeV} .$$

As PDF's utilizadas foram as do conjunto CTEQ4D no esquema DIS. Os parâmetros eletrofracos usados foram os definidos no início do capítulo. O valor do cut-off soft foi fixado em $\delta_s = 10^{-2}$, enquanto o cut-off colinear foi fixado em $\delta_c = 10^{-4}$.

Uma das maiores vantagens do método de NLO Monte Carlo para correções de QCD é que podemos obter qualquer número de distribuições diferenciais simultaneamente, bastando incluir o histograma desejado no programa FORTRAN que calcula as correções de QCD.

Na figura (8.6) temos a distribuição diferencial em termos da massa invariante do par e^+e^- para um leptoquark de massa $m = 250$ GeV, respeitando o limite experimental de 237 (GeV) obtido pela procura de produção simples de leptoquarks em colisões $p\bar{p}$ no Tevatron, e para acoplamentos variando de $\lambda = 0$ (SM em LL e NLO), $\lambda = 0.3$ e 0.6. Os parâmetros a e b contidos no vértice $S\ell q$ foram fixados em $a = b = 1/2$ em todos os cálculos. Para λ da ordem do acoplamento eletromagnético ($\lambda = 0.3$) podemos ver que a diferença em relação à curva do SM é muito pequena, aumentando para $\lambda = 0.6$. Cabe salientar também, que nossas simulações mostram que para $\lambda = 0.6$, leptoquarks do tipo R_2 com massas maiores do que 250 GeV devem produzir sinais muito pequenos. Isso pode ser visto na figura (8.7), onde calculamos a seção de choque em NLO incluindo leptoquarks em função da massa dos mesmos em comparação a σ^{NLO} dado pelo modelo padrão, onde pode-se ver que os leptoquarks se desacoplam rapidamente com o aumento de suas massas. Além disso, verificamos que para massas extremamente altas ou $\lambda = 0$, nossos programas numéricos reproduzem as seções de choque do modelo padrão exatamente.

Para quantificar o efeito causado pelos leptoquarks, vamos definir a quantidade Δ dada por

$$\Delta \equiv \frac{|\sigma^{sm+lq} - \sigma^{sm}|}{\sigma^{sm}}. \quad (8.2)$$

Para $\lambda = 0.3$ e $m = 250$ GeV, $\Delta \simeq 1.4\%$ com um K-factor de 1.40, enquanto para $\lambda = 0.6$, $\Delta \simeq 4.0\%$ com um K-factor de 1.38.

Um corte maior na massa invariante aumenta a diferença entre as seções de choque do modelo padrão e aquelas com inclusão de leptoquarks, no entanto, em termos práticos, a baixa estatística experimental na cauda da distribuição pode esconder o efeito causado pela nova física.

Nas figuras (8.8) e (8.9) temos as distribuições diferenciais normalizadas por σ^{NLO} do modelo padrão para os momentos transversos do elétron ($P_T(e)$) e do par e^+e^- ($P_T(e^+e^-)$) com $\lambda = 0$ (SM LL e NLO), $\lambda = 0.3$ e 0.6. No caso da distribuição em $P_T(e)$, diferenças significativas aparecem para $\lambda = 0.6$ para grandes valores daquela variável, enquanto que para $P_T(e^+e^-)$ as diferenças em relação à distribuição do SM são muito pequenas em qualquer caso e em todo o intervalo de variação de $P_T(e^+e^-)$.

Como argumentamos na seção (7.10.1), a produção simples de leptoquarks escalares é um processo que tem o mesmo estado final do processo de Drell-Yan, mas que deve ser calculado em separado quando o leptoquark está on-shell, contudo, no regime off-shell, os diagramas de Feynman com leptoquarks sendo produzidos no canal $\hat{s}_{13}$ dão uma contribuição legítima ao processo de Drell-Yan com inclusão de efeitos virtuais de leptoquarks trocados no canal t da reação em ordem α_s . Dessa

maneira, era preciso subtrair a contribuição de leptquarks on-shell produzidos no canal $\hat{s}_{13}$ após integração sobre o todo o espaço de fase disponível para a reação $q(\bar{q}) + g \rightarrow q(\bar{q}) + e^+ + e^-$. Como podemos constatar pela figura (8.10), onde temos a distribuição diferencial de $\hat{s}_{13}$, não há a ocorrência de ressonâncias, indicando que a subtração foi bem sucedida. Além disso, existe muito pouca diferença entre as distribuições geradas com e sem a inclusão dos leptquarks.

8.5 Região de Exclusão no Espaço de Parâmetros $\lambda \times M_{\ell q}$

Além das distribuições diferenciais com inclusão de leptquarks, obtivemos também as curvas de exclusão no espaço de parâmetros λ (acoplamento) versus $M_{\ell q}$ (massa do leptquark) em LL e NLO approximation. Para essa análise usamos um corte na massa invariante maior, $M(e^+e^-) \geq 180$ GeV, de modo a realçar o efeito causado pela inclusão do escalar. Denotando o número de eventos esperados no SM, e em uma teoria com a inclusão de leptquarks escalares por n_{SM} e $n_{\ell q}$ respectivamente, exigimos que

$$|n_{\ell q} - n_{SM}| \geq 3\sqrt{n_{SM}} \quad (8.3)$$

para que o efeito seja considerado visível. Isso nos permite obter curvas de significância estatística no espaço de parâmetros $M_{\ell q} - \lambda$ para o leptquark R_2 considerado na seção anterior. A região acima das respectivas curvas na figura (8.11) podem, então, ser excluías em um nível de significância de 3σ . A curva foi obtida com a luminosidade do Run I do Tevatron, $\mathcal{L} \sim 100 \text{ pb}^{-1}$. Para leptquarks de massa $m = 250$ GeV, acoplamentos maiores do que 0.68 podem ser excluídos nesse nível de significância.

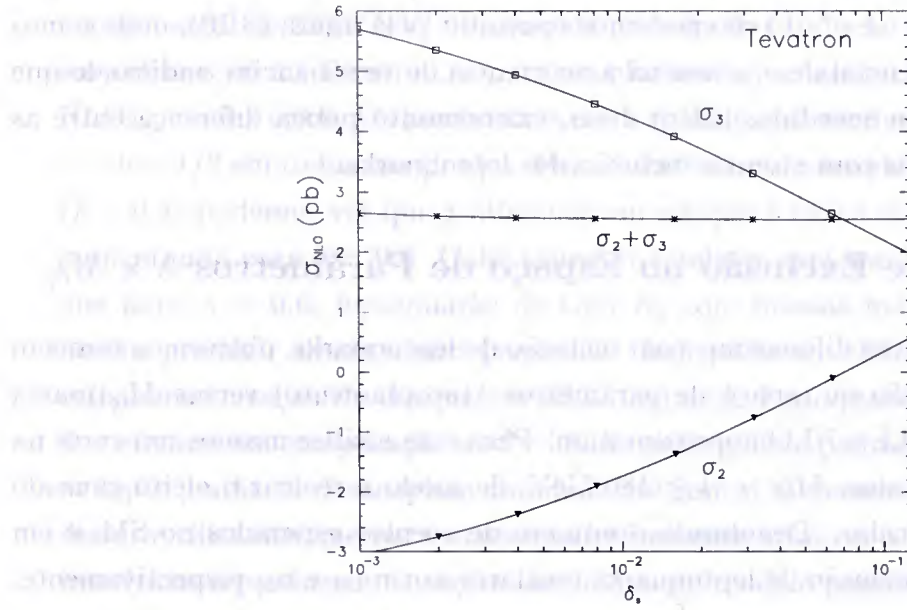


Figura 8.1: Variação da seção de choque em NLO de QCD com o parâmetro de corte soft δ_s , mantendo $\delta_c = 10^{-4}$ constante.

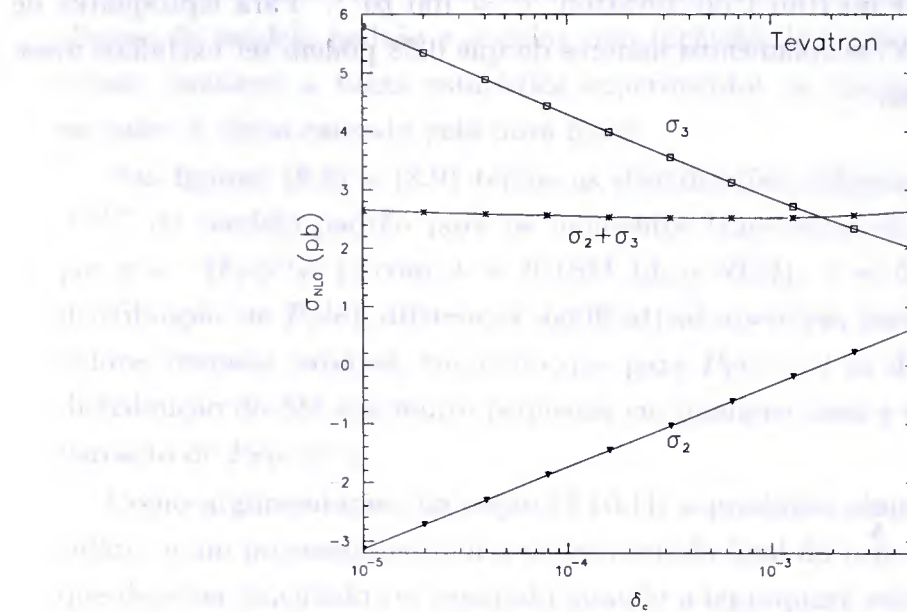


Figura 8.2: Variação da seção de choque em NLO de QCD com o parâmetro de corte colinear δ_c , mantendo $\delta_s = 10^{-2}$ constante.

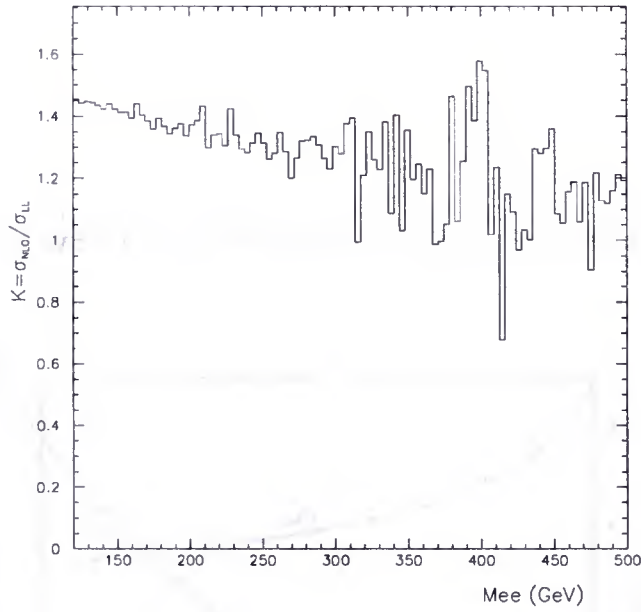


Figura 8.3: K-factor em função da massa invariante do par e^+e^- para $M(e^+e^-) \geq 120$ GeV com $Q = Q_R = Q_F = M_Z$.

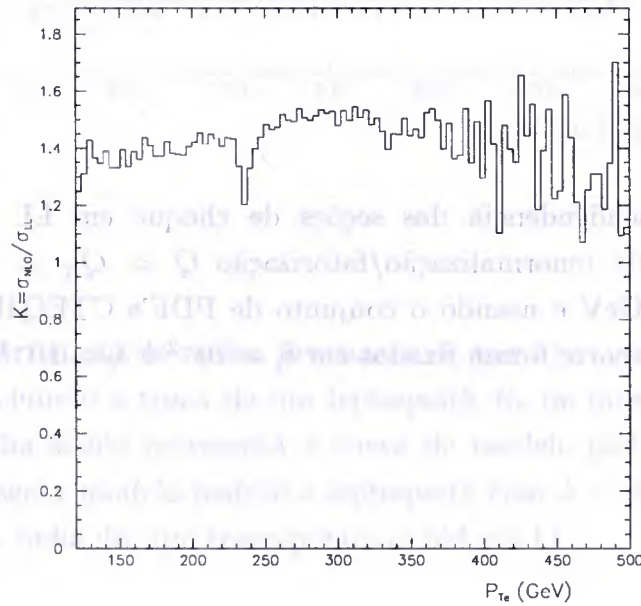


Figura 8.4: K-factor em função do momento transversal do elétron do estado final para $M(e^+e^-) \geq 120$ GeV com $Q = Q_R = Q_F = M_Z$.

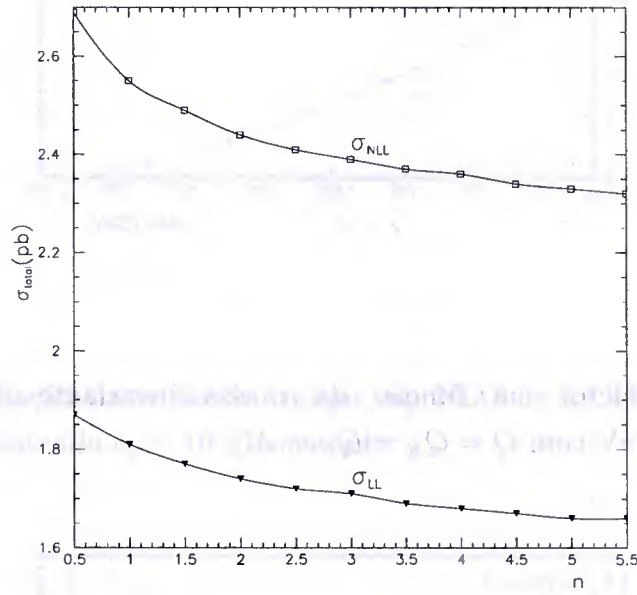


Figura 8.5: Dependência das seções de choque em LL e NLO com a escolha da escala de renormalização/fatorização $Q = Q_R = Q_F = n.M_Z$ para $M(e^+e^-) \geq 120$ GeV e usando o conjunto de PDF's CTEQ4D no esquema DIS. Os parâmetros de corte foram fixados em $\delta_s = 10^{-2}$ e $\delta_c = 10^{-4}$.

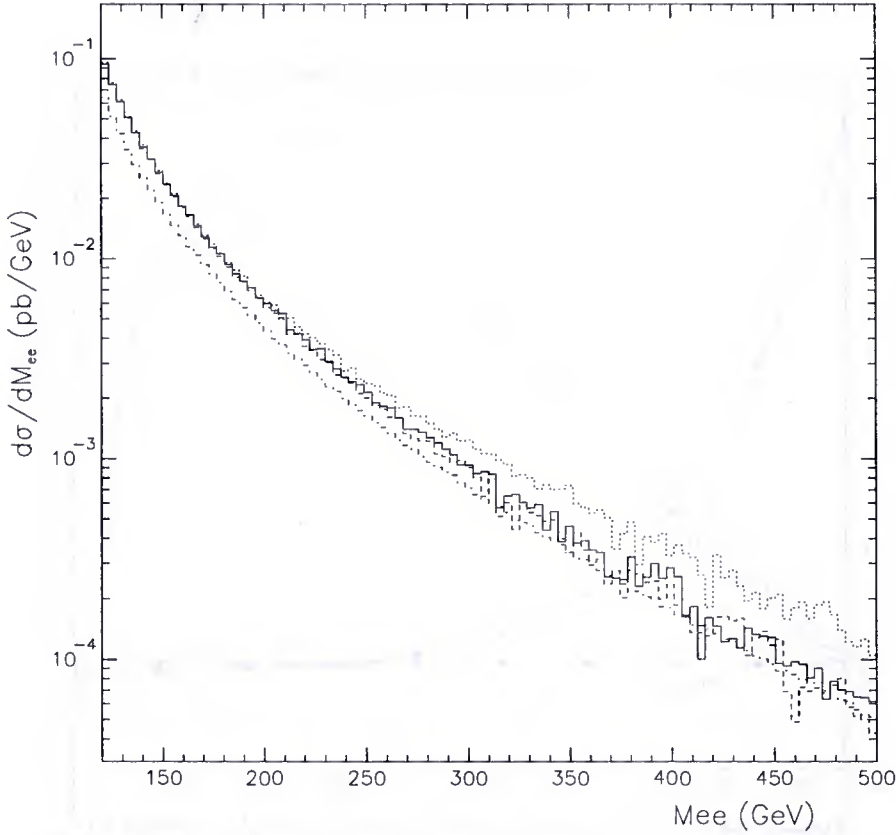


Figura 8.6: Distribuição de massa invariante do par e^+e^- na reação $p\bar{p} \rightarrow e^+e^-X$ no Tevatron incluindo a troca de um leptoquark R_2 de massa 250 GeV em NLO de QCD. A linha sólida representa a curva do modelo padrão em NLO, a linha tracejada representa modelo padrão e leptoquark com $\lambda = 0.3$, a linha pontilhada com $\lambda = 0.6$ e a linha do tipo traço-ponto ao SM em LL.

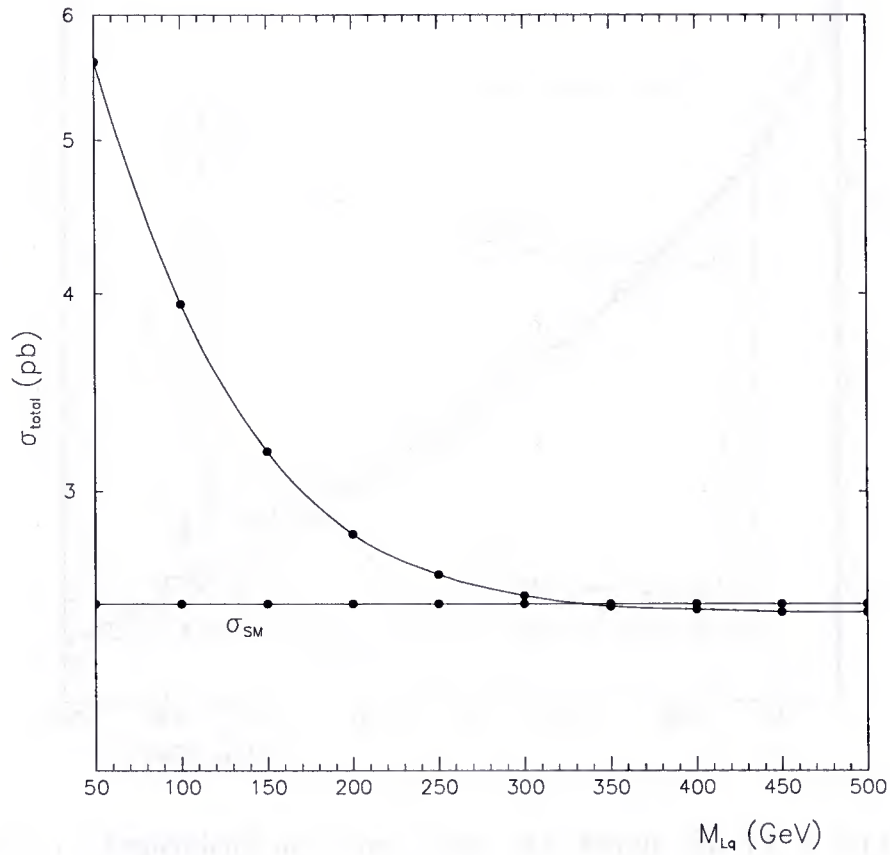


Figura 8.7: Seções de choque totais em NLO em função da massa do leptoquark. A curva constante indica a seção de choque total do modelo padrão em NLO.

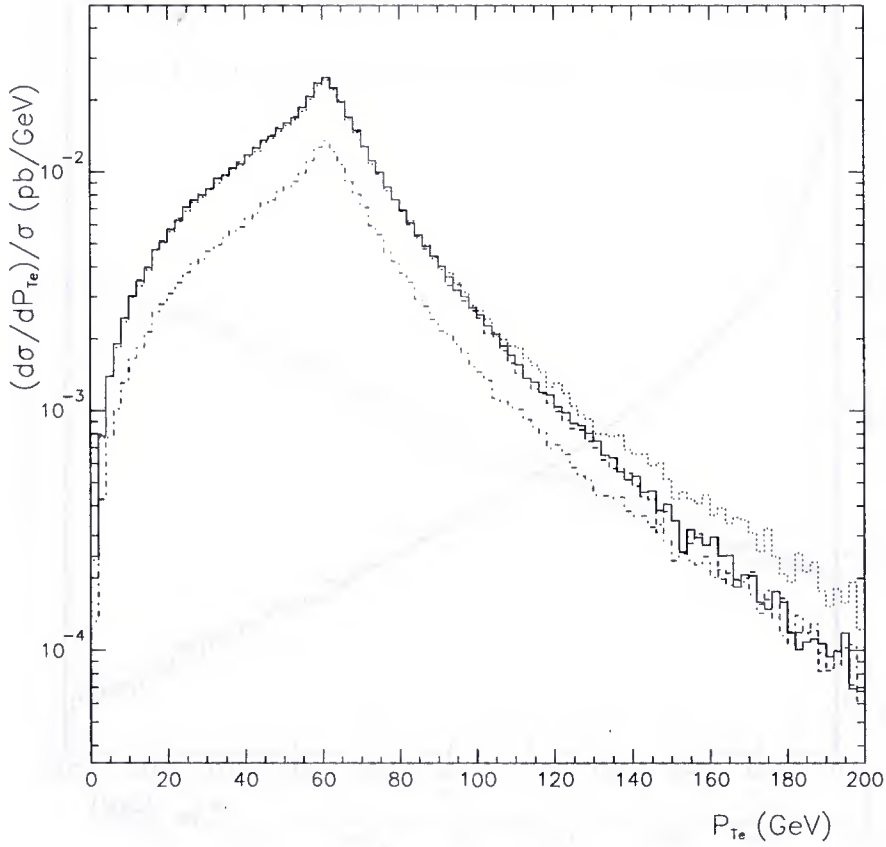


Figura 8.8: Distribuição do momento transversal do e^- do estado final na reação $p\bar{p} \rightarrow e^+e^-X$ no Tevatron incluindo a troca de um leptoquark R_2 de massa 250 GeV em NLO de QCD. A linha sólida representa a curva do modelo padrão, a linha tracejada representa modelo padrão e leptoquark com $\lambda = 0.3$, a linha pontilhada com $\lambda = 0.6$ e a linha do tipo traço-ponto ao SM em LL.

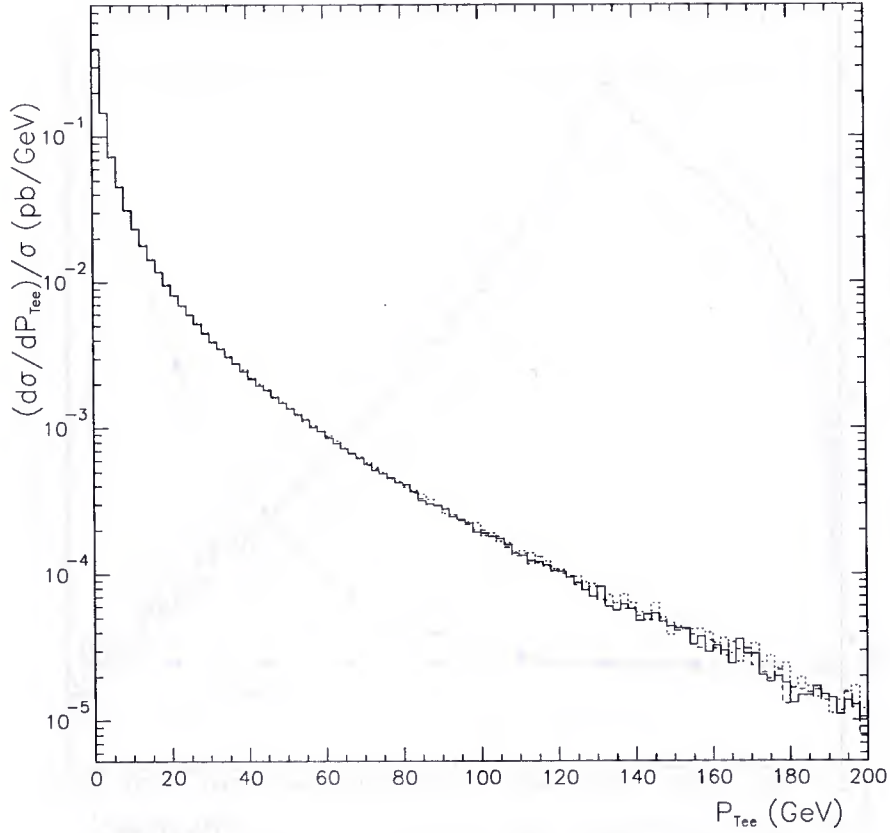


Figura 8.9: Distribuição do momento transverso do par e^+e^- do estado final na reação $p\bar{p} \rightarrow e^+e^-X$ no Tevatron incluindo a troca de um leptoquark do tipo R_2 de massa 250 GeV em NLO de QCD. A linha sólida representa a curva do modelo padrão, a linha tracejada representa modelo padrão e leptoquark com $\lambda = 0.3$ e a linha pontilhada com $\lambda = 0.6$.

2000/05/16 10.36

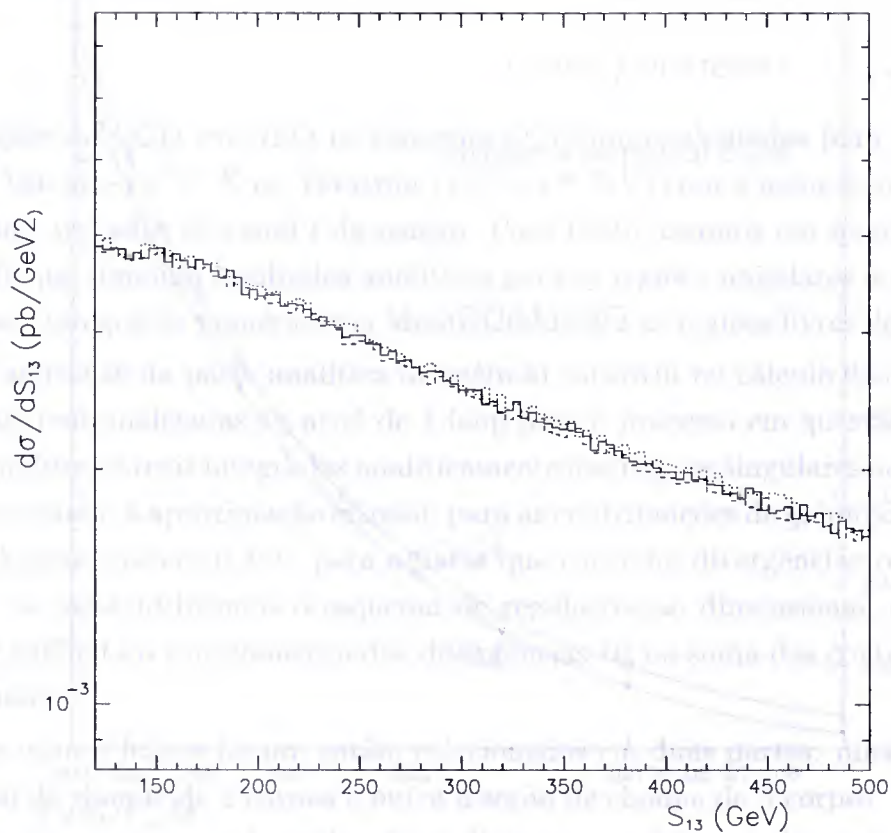


Figura 8.10: Distribuição diferencial para $\hat{s}_{13}$ incluindo apenas a produção simples de leptoquarks off-shell.

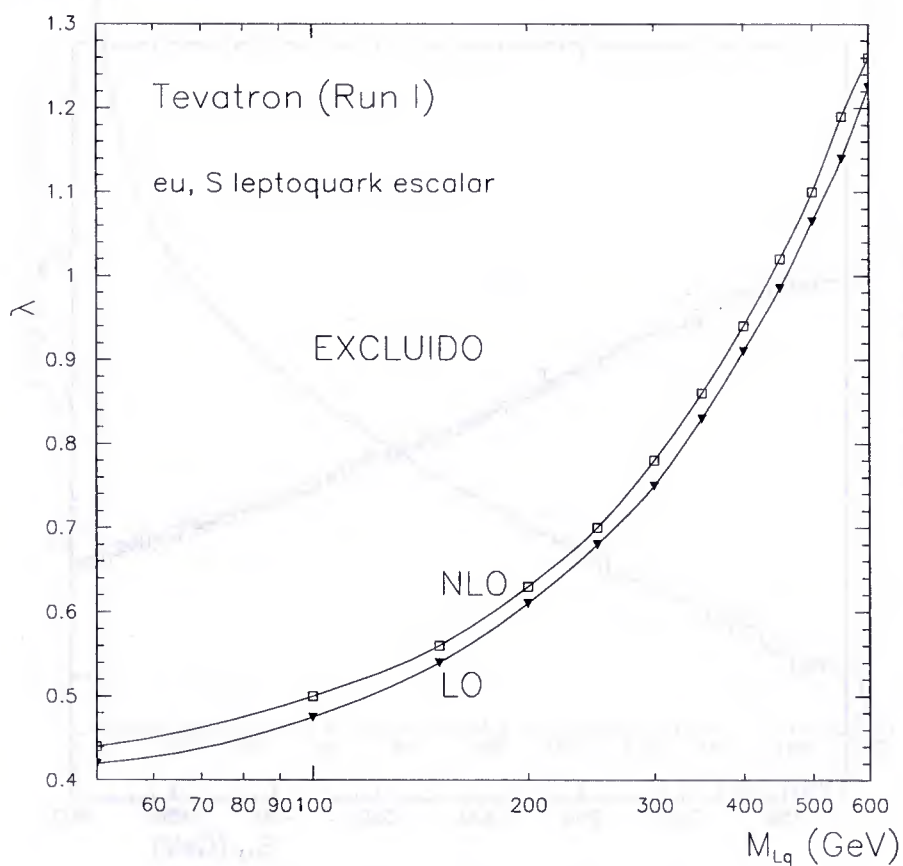


Figura 8.11: Curvas de exclusão no plano $M_{\ell q} - \lambda$ para o leptouark escalar de primeira geração R_2 no Tevatron (Run I). A região acima das respectivas curvas podem ser excluídas a um nível de 3σ .

Capítulo 9

Conclusão

Correções de QCD em NLO no esquema DIS foram calculadas para o processo de Drell-Yan $p\bar{p} \rightarrow e^+e^-X$ no Tevatron ($\sqrt{S} = 1.8$ TeV) com a inclusão de leptquarks escalares trocados no canal t da reação. Para tanto, usamos um método de cálculo híbrido que combina resultados analíticos para as regiões singulares no IR do espaço de fase e integração numérica via Monte Carlo para as regiões livres de divergências.

A aplicação da parte analítica do método consistiu no cálculo das contribuições virtuais renormalizadas ao nível de 1-loop para o processo em questão e no cálculo das amplitudes reais integradas analiticamente nas regiões singulares no IR do espaço de fase usando a aproximação eikonal, para as contribuições de glúon soft e a Leading Pole Approximation (LPA) para a parte que continha divergências colineares. Em todos os casos utilizamos o esquema de regularização dimensional, mostrando de forma explícita o cancelamento das divergências IR na soma das contribuições reais e virtuais.

Os termos finitos foram, então, colecionados em duas partes: uma contribuindo à seção de choque de 2 corpos e outra à seção de choque de 3 corpos. Em ambos os casos encontra-se uma dependência explícita nos parâmetros de corte δ_s e δ_c usados para delimitar as regiões divergentes e não divergentes do espaço de fase. A seção de choque total inclusiva, dada pela soma das contribuições de 2 e 3 corpos, foi então integrada numericamente no espaço de fase de 4 dimensões, onde demonstramos que as dependências nos parâmetros de corte se cancelam mutuamente entre as duas contribuições para uma vasta gama de valores de δ_s e δ_c , tornando a seção de choque total inclusiva livre de parâmetros não-físicos.

Como resultado da aplicação da parte numérica do método, constatamos que as correções de QCD aumentam as seções de choque de Drell-Yan em +41% no modelo padrão. As seções de choque com inclusão de leptquarks escalares que se acoplam a pares ue^- também recebem grandes contribuições em ordem α_s , chegando até +40% dependendo da intensidade do acoplamento de Yukawa e para uma massa de 250 GeV

para os leptoquarks.

A inclusão das correções de QCD não diminuíram a dependência das seções de choque de Born nas escalas de renormalização/fatorização, ao contrário a dependência aumentou de 11% para aproximadamente 14%, a exemplo de outros processos estudados no Tevatron. Por outro lado, constatamos que a dependência na escolha da parametrização das PDF é pequena.

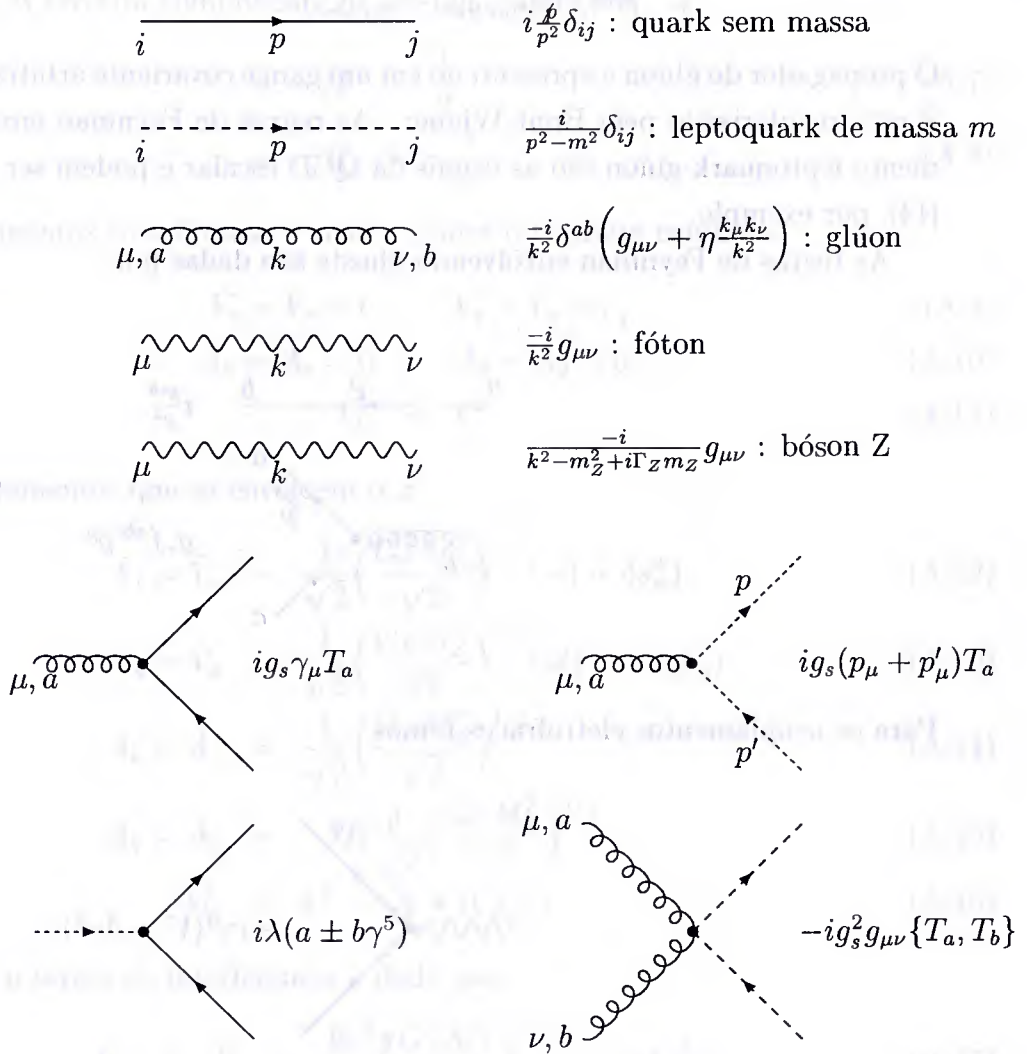
Distribuições diferenciais foram geradas para (I) a massa invariante M_{ee} e momento transversal $P_T(ee)$ do par e^+e^- , (II) para o momento transversal do elétron $P_T(e)$ e (III) para a massa invariante do par $e^- + jato$, $\hat{s}_{12}$, de maneira simultânea, demonstrando a versatilidade do método. Para uma massa de 250 GeV e acoplamento de $\lambda = 0.6$ diferenças significativas podem ser notadas para grandes valores de M_{ee} e $P_T(ee)$ no caso (I), enquanto (II) e (III) mostraram-se casos menos interessantes.

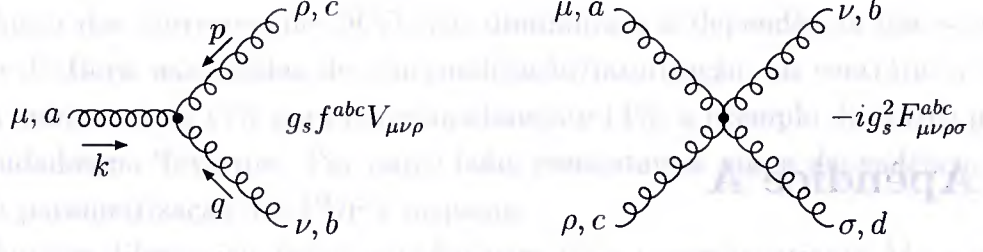
A região de exclusão ao nível de 3σ no espaço de parâmetros $\lambda \times M_{\ell q}$ foi obtida para o leptoquark em questão usando um corte de 180 GeV para M_{ee} . Para $m = 250$ GeV, encontramos que os valores do acoplamento tais que $\lambda > 0.68$ podem ser excluídos nesse nível de significância.

Apêndice A

Regras de Feynman

As regras de Feynman utilizadas no cálculo das amplitudes são dadas logo abaixo. A maioria delas está deduzida nas refs. [11, 13, 14].





o vértice $i \lambda(a - b\gamma^5)$ corresponde ao acoplamento antiquark-antilépton.

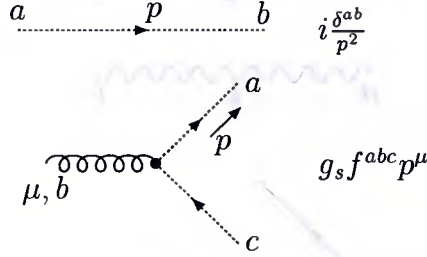
Os tensores envolvidos na expressão dos vértices tríplice e quártico envolvendo apenas glúons são dados por

$$V_{\mu\nu\rho} = g_{\mu\nu}(k-p)_\rho + g_{\nu\rho}(p-q)_\mu + g_{\rho\mu}(q-k)_\nu \quad (\text{A.1})$$

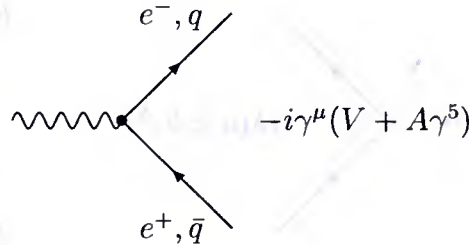
$$\begin{aligned} F_{\mu\nu\rho\sigma}^{abc} = & f^{abe} f^{cde} (g_{\mu\rho} g_{\nu\sigma} - g_{\mu\sigma} g_{\nu\rho}) + f^{ace} f^{bde} (g_{\mu\nu} g_{\rho\sigma} - g_{\mu\sigma} g_{\nu\rho}) \\ & + f^{ade} f^{bce} (g_{\mu\nu} g_{\rho\sigma} - g_{\mu\rho} g_{\nu\sigma}) . \end{aligned} \quad (\text{A.2})$$

O propagador do glúon é apresentado em um gauge covariante arbitrário η e do bóson Z está regularizado pela Breit-Wigner. As regras de Feynman envolvendo acoplamento leptoquark-glúon são as usuais da QCD escalar e podem ser encontradas em [14], por exemplo.

As regras de Feynman envolvendo ghosts são dadas por



Para os acoplamentos eletrofracos temos



onde $V = V_q, A = A_q$ para acoplamento com quarks e $V = V_e, A = A_e$ para acoplamento com léptons.

As constantes de acoplamento dos elétrons e quarks ao bóson Z são obtidos a partir de

$$R_e = 2s_w^2, \quad (\text{A.3})$$

$$L_e = -1 + 2s_w^2, \quad (\text{A.4})$$

para os elétrons, onde $s_w \equiv \sin \theta_W$ e θ_W é o ângulo de Weinberg. Para os quarks, temos

$$R_q = -2Q_q s_w^2, \quad (\text{A.5})$$

$$L_q = 2I_3^q - 2Q_q s_w^2, \quad (\text{A.6})$$

onde Q_q é a carga do quark em unidades da carga do elétron, ou seja, $Q_u = 2/3$ e $Q_d = -1/3$. A terceira componente de isospin é dada por

$$I_3^u = \frac{1}{2}, \quad (\text{A.7})$$

$$I_3^d = -\frac{1}{2}. \quad (\text{A.8})$$

Para acoplamentos com fótons, somente, temos o seguinte conjunto

$$V_e = \tilde{V}_e = e; \quad V_q = \tilde{V}_q = e_q \quad (\text{A.9})$$

$$A_e = \tilde{A}_e = 0; \quad A_q = \tilde{A}_q = 0 \quad (\text{A.10})$$

$$k_V^2 = k^2, \quad (\text{A.11})$$

e para acoplamentos que só envolvem o Z

$$V_e = \tilde{V}_e = \frac{1}{\sqrt{2}} \left(\frac{G_F M_Z^2}{\sqrt{2}} \right)^{1/2} (-1 + 4s_w^2) \quad (\text{A.12})$$

$$V_q = \tilde{V}_q = \frac{1}{\sqrt{2}} \left(\frac{G_F M_Z^2}{\sqrt{2}} \right)^{1/2} (2I_3^q - 4Q_q s_w^2) \quad (\text{A.13})$$

$$A_e = \tilde{A}_e = \frac{1}{\sqrt{2}} \left(\frac{G_F M_Z^2}{\sqrt{2}} \right)^{1/2} \quad (\text{A.14})$$

$$A_q = \tilde{A}_q = -2I_3^q \frac{1}{\sqrt{2}} \left(\frac{G_F M_Z^2}{\sqrt{2}} \right)^{1/2} \quad (\text{A.15})$$

$$k_V^2 = k^2 - m_Z^2 + i\Gamma_Z m_Z. \quad (\text{A.16})$$

Por último, o termo de interferência é dado por

$$V_e = e, \tilde{V}_e = \frac{1}{\sqrt{2}} \left(\frac{G_F M_Z^2}{\sqrt{2}} \right)^{1/2} (-1 + 4s_w^2) \quad (\text{A.17})$$

$$V_q = e_q, \quad \tilde{V}_q = \frac{1}{\sqrt{2}} \left(\frac{G_F M_Z^2}{\sqrt{2}} \right)^{1/2} (2I_3^q - 4Q_q s_w^2) \quad (\text{A.18})$$

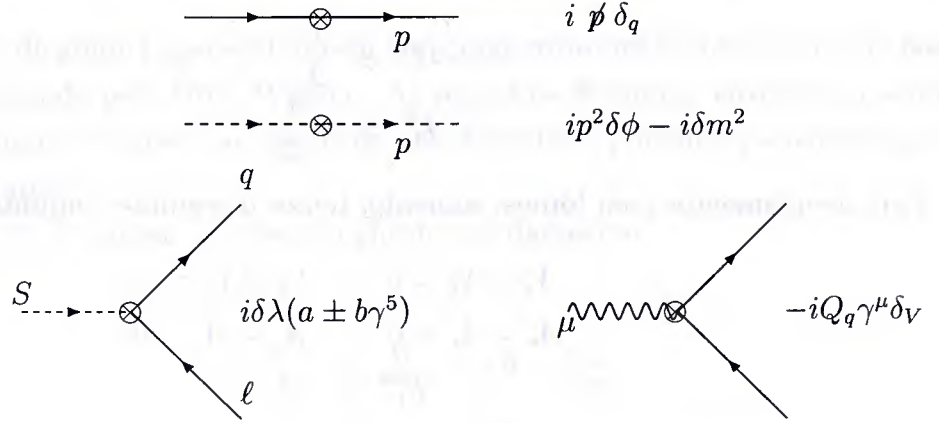
$$A_e = 0, \quad \tilde{A}_e = \frac{1}{\sqrt{2}} \left(\frac{G_F M_Z^2}{\sqrt{2}} \right)^{1/2} \quad (\text{A.19})$$

$$A_q = 0, \quad \tilde{A}_q = -2I_3^q \frac{1}{\sqrt{2}} \left(\frac{G_F M_Z^2}{\sqrt{2}} \right)^{1/2} \quad (\text{A.20})$$

$$|k_V^2|^2 = k^2 [k^2 - m_Z^2 + i\Gamma_Z m_Z] . \quad (\text{A.21})$$

Regras de Feynman para os Contratermos

Relacionamos abaixo as regras de Feynman para os contratermos do propagador do quark e do propagador do leptoquark de massa m , e também para os vértices $(\gamma, Z)q\bar{q}$ e $S\ell q$. Esses são os únicos contratermos necessários nos cálculos.



Apêndice B

Álgebra do Grupo $SU(N)$

O $SU(N)$ é o grupo das matrizes unitárias U $N \times N$, e que satisfazem

$$UU^\dagger = U^\dagger U = I_N \quad (\text{B.1})$$

$$\det U = 1 \quad , \quad (\text{B.2})$$

onde I_N é a matriz unitária N -dimensional.

Os geradores da álgebra do $SU(N)$, denotados por T^a ($a = 1, 2, \dots, N^2 - 1$), são matrizes hermitianas, de traço nulo e satisfazem às seguintes relações de comutação

$$[T^a, T^b] = i f^{abc} T^c \quad , \quad (\text{B.3})$$

onde as constantes de estrutura f^{abc} são reais e totalmente antisimétricas, e obedecem à chamada *identidade de Jacobi*

$$f^{abe} f^{cde} + f^{ace} f^{dbe} + f^{ade} f^{bce} = 0 \quad . \quad (\text{B.4})$$

A representação fundamental $T^a = \lambda^a/2$ é N -dimensional. O caso $N = 2$ corresponde às matrizes de Pauli usuais, enquanto para $N = 3$, temos as oito matrizes de Gell-Mann

$$\lambda^1 = \begin{pmatrix} 0 & 1 & 0 \\ 1 & 0 & 0 \\ 0 & 0 & 0 \end{pmatrix}, \lambda^2 = \begin{pmatrix} 0 & -i & 0 \\ i & 0 & 0 \\ 0 & 0 & 0 \end{pmatrix}, \lambda^3 = \begin{pmatrix} 1 & 0 & 0 \\ 0 & -1 & 0 \\ 0 & 0 & 0 \end{pmatrix}, \lambda^4 = \begin{pmatrix} 0 & 0 & 1 \\ 0 & 0 & 0 \\ 1 & 0 & 0 \end{pmatrix}, \quad (\text{B.5})$$

$$\lambda^5 = \begin{pmatrix} 0 & 0 & -i \\ 0 & 0 & 0 \\ i & 0 & 0 \end{pmatrix}, \lambda^6 = \begin{pmatrix} 0 & 0 & 0 \\ 0 & 0 & 1 \\ 0 & 1 & 0 \end{pmatrix}, \lambda^7 = \begin{pmatrix} 0 & 0 & 0 \\ 0 & 0 & -i \\ 0 & i & 0 \end{pmatrix}, \lambda^8 = \frac{1}{\sqrt{3}} \begin{pmatrix} 1 & 0 & 0 \\ 0 & 1 & 0 \\ 0 & 0 & -2 \end{pmatrix}.$$

Estas matrizes satisfazem as relações de anticomutação

$$\{\lambda^a, \lambda^b\} = \frac{4}{N} \delta^{ab} I_N + 2 d^{abc} \lambda^c \quad , \quad (\text{B.6})$$

onde as constantes d^{abc} são totalmente simétricas nos índices a, b e c .

Para o $SU(3)$, as únicas constantes f^{abc} e d^{abc} diferentes de zero (a menos de permutações) são

$$\begin{aligned} \frac{1}{2}f^{123} &= f^{147} = -f^{156} = f^{246} = f^{257} = f^{345} = -f^{367} \\ &= \frac{1}{\sqrt{3}}f^{458} = \frac{1}{\sqrt{3}}f^{678} = \frac{1}{2}, \\ d^{146} &= d^{157} = -d^{247} = d^{256} = d^{344} = d^{355} = -d^{366} = -d^{377} = \frac{1}{2}, \\ d^{118} &= d^{228} = d^{338} = -2d^{448} = -2d^{558} = -2d^{688} = -2d^{788} = -d^{888} = \frac{1}{\sqrt{3}}. \end{aligned} \quad (\text{B.7})$$

A representação adjunta do $SU(N)$ é dada pelas matrizes $(N^2 - 1) \times (N^2 - 1)$

$$(T_A^a)_{bc} \equiv -if^{abc}. \quad (\text{B.8})$$

As relações

$$\begin{aligned} \text{Tr}(\lambda^a \lambda^b) &= 4T_R \delta_{ab}, \quad T_R = \frac{1}{2}, \\ (\lambda^a \lambda^b)_{\alpha\beta} &= 4C_F \delta_{\alpha\beta}, \quad C_F = \frac{N^2 - 1}{2N}, \end{aligned} \quad (\text{B.9})$$

$$\text{Tr}(T_A^a T_A^b) = f^{acd} f^{bcd} = C_A \delta_{ab}, \quad C_A = N, \quad (\text{B.10})$$

definem os invariantes do $SU(N)$: T_R , C_F e C_A . Para o $SU(3)$, esses invariantes são dados por $1/2$, $4/3$ e 3 , respectivamente.

Apêndice C

Espaços de Fase em d Dimensões

C.1 Definição do Espaço de Fase de N Partículas

O espaço de fase de N partículas no estado final em processos do tipo $1, 2 \rightarrow N$, é o espaço das possíveis configurações de energia e momento das partículas do estado final que obedecem às leis de conservação cinemáticas

$$p_a + \{p_b\} = \sum_{i=1}^N p_i, \quad (\text{C.1})$$

onde p_a e p_b são os 4-momentos das partículas do estado inicial, em um processo de colisão, ou p_a simplesmente para processos de decaimento $1 \rightarrow N$; e p_i ($i = 1, \dots, N$), os 4-momentos das partículas finais. O espalhamento $2 \rightarrow N$ é ilustrado abaixo

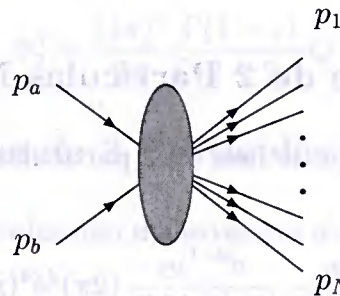


Figura C.1: Espalhamento $a + b \rightarrow 1 + 2 + \dots + N$.

onde a elipse representa um processo de espalhamento arbitrário.

O elemento de espaço de fase diferencial em d dimensões espaço-temporais é dado por

$$d^d \Phi_N = (2\pi)^d \prod_{i=1}^N \frac{d^{d-1} p_i}{(2\pi)^{d-1} 2E_i} \delta^d(p_a + p_b - \sum_{i=1}^N p_i), \quad (\text{C.2})$$

onde E_i representa a energia da partícula de 4-momento p_i . O elemento de integração $d^{d-1} p_i / ((2\pi)^{d-1} 2E_i)$ é dado por

$$\begin{aligned} \frac{d^{d-1}p_i}{(2\pi)^{d-1}2E_i} &= \frac{1}{2(2\pi)^{d-1}} dE_i E_i^{d-2} \times \\ &\times \sin^{d-3}\theta_1 d\theta_1 \sin^{d-4}\theta_2 d\theta_2 \times \cdots \times d\theta_{d-2} \equiv \frac{1}{2(2\pi)^{d-1}} dE_i E_i^{d-2} d\Omega_{d-1} . \end{aligned} \quad (C.3)$$

A partir dessas expressões obtivemos todos os espaços de fase utilizados nos cálculos. A seguir vamos especificar (C.2) para os casos de 1, 2 e 3 partículas no estado final.

C.2 Espaço de Fase de 1 Partícula

O espaço de fase de 1 partícula em d dimensões é dado por (C.2) com $N = 1$

$$d^d\Phi_1 = \frac{d^d p_1}{(2\pi)^{d-1}2E_1} (2\pi)^d \delta^d(p_a + p_b - p_1) , \quad (C.4)$$

o qual pode ser facilmente integrado usando a identidade

$$\int \frac{d^{d-1}p_1}{2E_1} = \int d^d p_1 \delta(p_1^2 - m_1^2) \quad (C.5)$$

e que inserida em (C.4) nos dá

$$\Phi_1 = 2\pi \delta[s - m_1^2] ; \quad (C.6)$$

onde definimos $s = (p_a + p_b)^2$.

C.3 Espaço de Fase de 2 Partículas Não Massivas

Consideremos agora o espaço de fase de 2 partículas não-massivas em d dimensões espaço-temporais

$$\Phi_2^d = \int \frac{d^{d-1}p_1}{(2\pi)^{d-1}2E_1} \frac{d^{d-1}p_2}{(2\pi)^{d-1}2E_2} (2\pi)^d \delta^d(p_a + p_b - p_1 - p_2) . \quad (C.7)$$

Usando a identidade (C.5) para p_1 , podemos escrever (C.7) como

$$\Phi_2^d = (2\pi)^{2-d} \int \frac{d^{d-1}p_2}{2E_2} \delta[(p_a + p_b - p_2)^2] . \quad (C.8)$$

No referencial do centro de massa (CMS) das partículas iniciais, a função δ pode ser escrita explicitamente na forma

$$\delta[(p_a + p_b - p_2)^2] = \delta(s - 2\sqrt{s}E_2) = \frac{1}{2\sqrt{s}} \delta\left(\frac{\sqrt{s}}{2} - E_2\right) . \quad (C.9)$$

Dessa maneira, temos que

$$\begin{aligned}\Phi_2^d &= (2\pi)^{2-d} \int dE_2 \frac{E_2^{d-3}}{4\sqrt{s}} \times \delta\left(\frac{\sqrt{s}}{2} - E_2\right) \\ &\times \int_0^{2\pi} d\theta_{d-2} \int_0^\pi \sin^{d-3} \theta_1 d\theta_1 \times \cdots \times \int_0^\pi \sin^{d-3} \theta_{d-3} d\theta_{d-3} \\ &= (2\pi)^{-1+2\varepsilon} \frac{4^\varepsilon s^{-\varepsilon}}{8} \prod_{m=1}^{d-4} \frac{\Gamma(\frac{m}{2} + \frac{1}{2})}{\Gamma(\frac{m}{2} + 1)} \sqrt{\pi} \int_0^\pi \sin^{1-2\varepsilon} \theta_1 d\theta_1, \quad (C.10)\end{aligned}$$

onde usamos que

$$\int_0^\pi \sin^m \theta d\theta = \sqrt{\pi} \frac{\Gamma(\frac{m}{2} + \frac{1}{2})}{\Gamma(\frac{m}{2} + 1)}, \quad (C.11)$$

e que $d = 4 - 2\varepsilon$. Expandindo a produtória e definindo $z \equiv \cos \theta_1$, temos, finalmente, que

$$\Phi_2^d = \frac{2^{2\varepsilon}}{16\pi} \left(\frac{4\pi}{s}\right) \frac{1}{\Gamma(1-\varepsilon)} \int_{-1}^1 (1-z^2)^{-\varepsilon} dz, \quad (C.12)$$

onde $E_1 = E_2 = \sqrt{s}/2$. Em 4 dimensões, (C.12) se reduz simplesmente a

$$\Phi_2^4 = \frac{1}{8\pi}. \quad (C.13)$$

Em um espalhamento $2 \rightarrow 2$, o elemento de matriz de um dado processo pode depender de θ_1 . Por exemplo, no processo de Drell-Yan, esse ângulo é identificado como o ângulo formado entre o 3-momento do elétron do estado final e o 3-momento do quark inicial. Por outro lado, em um processo do tipo $1 \rightarrow 2$, como o decaimento de um fóton de virtualidade Q^* , podemos integrar em z , tornando o espaço de fase um mero fator a ser multiplicado pelo elemento de matriz, nesse caso, (C.12) se reduz a

$$\Phi_2^d = \frac{(4\pi)^\varepsilon}{8\pi} \frac{\Gamma(1-\varepsilon)}{\Gamma(2-2\varepsilon)} Q^{-2\varepsilon}. \quad (C.14)$$

C.4 Espaço de Fase de 3 Partículas Não Massivas

O espaço de fase de 3 partículas não massivas em d dimensões é dado explicitamente por

$$\Phi_3^d = \int \frac{d^{d-1}p_1}{(2\pi)^{d-1}(2E_1)} \frac{d^{d-1}p_2}{(2\pi)^{d-1}(2E_2)} \frac{d^{d-1}p_3}{(2\pi)^{d-1}(2E_3)} (2\pi)^d \delta^d(p_a + p_b - p_1 - p_2 - p_3). \quad (C.15)$$

Ao contrário de Φ_2^d , o espaço de fase de 3 partículas em d dimensões completo não foi utilizado em nenhum cálculo de processos do tipo $2 \rightarrow 3$. Em todos os casos onde (C.15) foi colocada em uso, aproximações foram usadas para escrevê-la na forma

$$\Phi_3^d \approx \Phi_{\text{singular}} \times \Phi_2^d, \quad (C.16)$$

*Nesse caso identificamos $p_a + p_b \rightarrow q$, onde $q^2 = Q^2$

onde Φ_{singular} denota a região do espaço de fase que gera singularidades IR.

Contudo, Φ_3^4 , o espaço de fase de 3 partículas em 4 dimensões, é usada em sua forma exata para a implementação numérica das integrais envolvidas em processos $2 \rightarrow 3$ nas regiões não singulares do espaço de fase. Essa implementação é feita dividindo-se Φ_3^4 em dois espaços de fase de 2 partículas, através da identidade

$$1 = \int dM_{23}^2 d^4q \delta(q^2 - M_{23}^2) \delta^4(q - p_2 - p_3) , \quad (\text{C.17})$$

onde M_{23} é massa invariante do par 23.

Inserindo (C.17) em (C.15), obtemos

$$\Phi_3^d = \int dM_{23}^2 \Phi_2^4(p_a + p_b - p_1 - q) \times \Phi_2^4(q - p_2 - p_3) , \quad (\text{C.18})$$

onde $\Phi_2^4(p_a + p_b - p_1 - q)$ é calculado no CMS do par $a + b$, enquanto $\Phi_2^4(q - p_2 - p_3)$, é calculado no CMS do par $2 + 3$, e são dados por

$$\Phi_2^4(p_a + p_b - p_1 - q) = \frac{\lambda^{1/2}(s, 0, M_{23}^2)}{8s} \int d^2\Omega_1^*(a + b \rightarrow 1 + q(23)) \quad (\text{C.19})$$

$$\Phi_2^4(q - p_2 - p_3) = \frac{\lambda^{1/2}(M_{23}^2, 0, 0)}{8M_{23}^2} \int d^2\Omega_2^*(q(23) \rightarrow 2 + 3) , \quad (\text{C.20})$$

onde a função $\lambda(x, y, z)$ é dada por

$$\lambda(x, y, z) = (x - y - z)^2 - 4yz , \quad (\text{C.21})$$

e os elementos de ângulo sólido por

$$d^2\Omega_1^*(a + b \rightarrow 1 + q(23)) = d\varphi_1^* d\cos\theta_1^* , \quad (\text{C.22})$$

$$d^2\Omega_2^*(q(23) \rightarrow 2 + 3) = d\varphi_2^{(23)} d\cos\theta_2^{(23)} . \quad (\text{C.23})$$

No CMS de $a + b$, os momentos p_a , p_b e p_1 são dados por

$$p_a = \frac{\sqrt{s}}{2}(1, 0, 0, 1) , \quad (\text{C.24})$$

$$p_b = \frac{\sqrt{s}}{2}(1, 0, 0, -1) , \quad (\text{C.25})$$

$$p_1 = E_1^*(1, \sin\theta_1^*, 0, \cos\theta_1^*) , \quad (\text{C.26})$$

onde escolhemos $\varphi_1^* = 0$ por simplicidade dado que a integração nessa variável produz um fator 2π somente. A energia da partícula 1 é dada por

$$E_1^* = \frac{s - M_{23}^2}{2\sqrt{2}} . \quad (\text{C.27})$$

No CMS do par $2 + 3$ temos que

$$p_2^{(23)} = E_2^{(23)}(1, \sin \theta_2^{(23)} \cos \varphi_2^{(23)}, \sin \theta_2^{(23)} \sin \varphi_2^{(23)}, \cos \theta_2^{(23)}) , \quad (\text{C.28})$$

$$p_3^{(23)} = E_3^{(23)}(1, -\sin \theta_2^{(23)} \cos \varphi_2^{(23)}, -\sin \theta_2^{(23)} \sin \varphi_2^{(23)}, -\cos \theta_2^{(23)}) , \quad (\text{C.29})$$

onde as energias são dadas por

$$E_2^{(23)} = E_3^{(23)} = \frac{M_{23}}{2} . \quad (\text{C.30})$$

O espaço de fase pode agora ser escrito em termos de variáveis de integração que variam de 0 até 1. Denotamos essas variáveis por x_1, x_2, x_3 e x_4 , já integrando sobre φ_1^* , as quais são definidas por

$$M_{23} = \sqrt{s} \cdot x_1 , \quad (\text{C.31})$$

$$\cos \theta_1^* = -1 + 2x_2 , \quad (\text{C.32})$$

$$\varphi_2^{(23)} = 2\pi x_3 , \quad (\text{C.33})$$

$$\cos \theta_2^{(23)} = -1 + 2x_4 . \quad (\text{C.34})$$

Logo, calculando todos os jacobianos, temos que

$$\Phi_3^4 = \frac{\lambda^{1/2}(s, 0, M_{23}^2) \lambda^{1/2}(M_{23}^2, 0, 0)}{64\pi^3 M_{23} \sqrt{s}} \int_0^1 dx_1 dx_2 dx_3 dx_4 . \quad (\text{C.35})$$

Essa expressão foi a base para a implementação numérica dos resultados no espaço de fase 3 partículas.

Por último, apresentamos o espaço de fase de 3 partículas na aproximação colinear usado na integração do elemento de matriz (4.8) correspondente à emissão de gluons colineares no processo $e^+e^- \rightarrow q\bar{q}g$

$$d\Phi_3^{col} \approx \frac{(4\pi)^\epsilon}{16\pi^2 \Gamma(1-\epsilon)} dt dz [tz(1-z)]^{-\epsilon} \Theta(t_{min} - t) . \quad (\text{C.36})$$

O cálculo de $d\Phi_3^{col}$ é análogo ao realizado na seção (6.4) para o processo de Drell-Yan e pode ser encontrado em detalhes em [38].

C.4.1 Espaço de Fase do Decaimento $\gamma^* \rightarrow q\bar{q}g$

Na subseção (3.3.1), as frações de energia das partículas, definidas por

$$x_i = \frac{2E_i}{Q} \quad i = 1, 2, 3 \quad (\text{C.37})$$

são usadas para expressar o módulo quadrado da amplitude e o espaço de fase de 3 partículas que derivamos agora.

Partindo de (C.15) para o caso do decaimento de uma partícula de 4-momento q

$$\Phi_3^d = \int \frac{d^{d-1}p_1}{(2\pi)^{d-1}(2E_1)} \frac{d^{d-1}p_2}{(2\pi)^{d-1}(2E_2)} \frac{d^{d-1}p_3}{(2\pi)^{d-1}(2E_3)} (2\pi)^d \delta^d(q - p_1 - p_2 - p_3) , \quad (\text{C.38})$$

integraremos em p_3 usando a função delta de Dirac, de modo que

$$\int d^{d-1}p_3 \delta^d(q - p_1 - p_2 - p_3) = \delta(Q - E_1 - E_2 - E_3) . \quad (\text{C.39})$$

Por outro lado, usando (C.3), as integrais angulares de p_1 e p_2 podem ser calculadas, de modo que

$$\frac{d^2\Phi_3^d}{dE_1 dE_2} = \frac{2^{d-2}\pi^{d-2}}{(2\pi)^{2d-3}\Gamma(d-2)} (E_1 E_2)^{d-3} \int_{-1}^1 dz (1-z^2)^{\frac{d}{2}-2} \frac{\delta(Q - E_1 - E_2 - E_3)}{2E_3} , \quad (\text{C.40})$$

onde $z \equiv \cos \theta_{d-3} = \cos \theta_{12}$. Usando (C.37), a função (C.39) pode ser reescrita como

$$\frac{\delta(Q - E_1 - E_2 - E_3)}{2E_3} = \delta \left[Q^2(1 - x_1 - x_2 + \frac{1}{2}x_1 x_2(1 - z)) \right] , \quad (\text{C.41})$$

de maneira que

$$\int_{-1}^1 dz (1-z^2)^{\frac{d}{2}-2} \frac{\delta(Q - E_1 - E_2 - E_3)}{2E_3} = \frac{2(1-z^2)^{\frac{d}{2}-2}}{x_1 x_2 Q^2} . \quad (\text{C.42})$$

Usando (C.42) em (C.40) temos finalmente

$$d^2\Phi_3^d = \frac{Q^2}{16(2\pi)^3} \left(\frac{Q^2}{4\pi} \right)^{-2\varepsilon} \frac{1}{\Gamma(2-2\varepsilon)} \left(\frac{1-z^2}{4} \right)^{-\varepsilon} x_1^{-2\varepsilon} dx_1 x_2^{-2\varepsilon} dx_2 , \quad (\text{C.43})$$

onde usamos que $dE_i = (Q/2) dx_i$.

A variável z é, na verdade, dependente das frações de energia. Isso pode ser visto calculando o produto escalar entre p_1 e p_2

$$\begin{aligned} 2p_1 \cdot p_2 &= 2E_1 E_2 (1 - z) = Q^2(-1 + x_1 + x_2) \\ 1 - z &= -\frac{2}{x_1 x_2} (1 - x_1 - x_2) \\ z &= 1 + \frac{2}{x_1 x_2} (1 - x_1 - x_2) , \end{aligned} \quad (\text{C.44})$$

onde usamos que $E_1 = E_2 = Q/2$ no referencial de repouso do fóton virtual.

Apêndice D

Fórmulas e Integrais Úteis

Neste apêndice constam algumas propriedades e expansões da função Γ utilizadas nos cálculos e algumas integrais úteis no cálculo dos loops.

A função $\Gamma(x)$ é definida como

$$\Gamma(x) = \int_0^1 dy e^{-y} y^{x-1} , \quad (D.1)$$

com $x > 0$ e tem as seguintes propriedades

$$\Gamma(n+1) = n! , \quad (D.2)$$

$$\Gamma(x+1) = x\Gamma(x) , \quad (D.3)$$

$$\Gamma(2x) = \frac{2^{2x-1}}{\sqrt{2}} \Gamma(x) \Gamma\left(x + \frac{1}{2}\right) . \quad (D.4)$$

A primeira e segunda derivadas de $\Gamma(x)$ calculadas no ponto $x = 1$ são dadas por

$$\Gamma'(1) = -\gamma , \quad (D.5)$$

$$\Gamma''(1) = \gamma^2 + \frac{1}{6}\pi^2 , \quad (D.6)$$

onde $\gamma = 0.5772157\dots$ é a constante de Euler-Mascheroni.

A função B (beta) é definida por

$$B(\alpha, \beta) = \int_0^1 x^\alpha (1-x)^\beta = \frac{\Gamma(\alpha+1)\Gamma(\beta+1)}{\Gamma(\alpha+\beta+2)} . \quad (D.7)$$

Usando as derivadas de $\Gamma(x)$ mostra-se que a expansão em série de Taylor de $\Gamma(1+x)$ é da forma

$$\Gamma(1+x) = 1 - \gamma x + \frac{1}{2} \left(\gamma^2 + \frac{\pi^2}{6} \right) x^2 + \dots , \quad (D.8)$$

e também que

$$\Gamma(-1)(1+x) = 1 + \gamma x + \frac{1}{2} \left(\gamma^2 - \frac{\pi^2}{6} \right) x^2 + \dots . \quad (D.9)$$

Integrais Paramétricas

$$\frac{1}{ab} = \int_0^1 dy \frac{1}{[ay + b(1-y)]^2} , \quad (\text{D.10})$$

$$\frac{1}{c^2 d} = \int_0^1 dy \frac{2y}{[cy + d(1-y)]^3} , \quad (\text{D.11})$$

$$\frac{1}{c^2 d^2} = \int_0^1 dy \frac{6y(1-y)}{[cy + d(1-y)]^4} . \quad (\text{D.12})$$

Integral de Momento Em d Dimensões

$$\int \frac{d^d l}{(2\pi)^d} \frac{(l^2)^\alpha}{[l^2 - \Delta]^\beta} = \frac{i(-1)^{\alpha-\beta}}{(4\pi)^{n/2}} \Delta^{(\alpha-\beta+n/2)} \frac{\Gamma(\alpha + \frac{n}{2})\Gamma(\beta - \alpha - \frac{n}{2})}{\Gamma(\frac{n}{2})\Gamma(\beta)} . \quad (\text{D.13})$$

Apêndice E

Funções Escalares de Passarino-Veltman

Neste apêndice estão colecionadas todas as funções escalares de 2,3 e 4 pontos provenientes da integração no momento dos loops das amplitudes virtuais calculadas na seção 2.

As integrais foram calculadas em $d = 4 - 2\varepsilon$ dimensões, no esquema de regularização dimensional, e usando o método de parametrização de Feynman. Todas as integrais vetoriais e tensoriais foram reduzidas a escalares usando o pacote FeynCalc para Mathematica.

1. Função Escalar de 1-Ponto

$$A_0(m^2) = \mu^{2\varepsilon} \int D^d k \frac{1}{k^2 - m^2} = \frac{i}{16\pi^2} N_\varepsilon \left(\frac{\mu^2}{m^2} \right)^\varepsilon m^2 \left[\frac{1}{\varepsilon} + 1 \right] , \quad (\text{E.1})$$

onde

$$\begin{aligned} D^d k &= \frac{d^d k}{(2\pi)^d} , \\ N_\varepsilon &= (4\pi)^\varepsilon \Gamma(1 + \varepsilon) . \end{aligned}$$

2. Funções Escalares de 2-Pontos

A função escalar de 2-pontos é definida por

$$B_0(p^2, m_0^2, m_1^2) = \mu^{2\varepsilon} \int D^d k \frac{1}{[k^2 - m_0^2][(k+p)^2 - m_1^2]} . \quad (\text{E.2})$$

Os seguintes tipos de funções B_0 aparecem nas amplitudes virtuais

$$B_0(p^2, 0, m^2) = \frac{i}{16\pi^2} N_\varepsilon \left(\frac{\mu^2}{m^2} \right)^\varepsilon \left[\frac{1}{\varepsilon} + 2 - \left(1 - \frac{m^2}{p^2} \right) \ln \left(1 - \frac{p^2}{m^2} \right) \right] , \quad (\text{E.3})$$

$$B_0(m^2, 0, m^2) = \frac{i}{16\pi^2} N_\epsilon \left(\frac{\mu^2}{m^2} \right)^\epsilon \left[\frac{1}{\epsilon} + 2 \right], \quad (\text{E.4})$$

$$B_0(0, 0, m^2) = \frac{i}{16\pi^2} N_\epsilon \left(\frac{\mu^2}{m^2} \right)^\epsilon \left[\frac{1}{\epsilon} + 1 \right], \quad (\text{E.5})$$

$$B_0(p^2, 0, 0) = \frac{i}{16\pi^2} N_\epsilon \left(\frac{\mu^2}{p^2} \right)^\epsilon \left[\frac{1}{\epsilon} + 2 \right]. \quad (\text{E.6})$$

A derivada de $B_0(p^2, 0, m^2)$ em $p^2 = 0$ é dada por

$$\frac{d}{dp^2} B_0(p^2, 0, m^2)(p^2 = 0) = \frac{i}{16\pi} \frac{1}{2m^2} + \mathcal{O}(\epsilon). \quad (\text{E.7})$$

3. Funções Escalares de 3-Pontos

A função escalar de 3-pontos é definida por

$$C_0(p_1^2, (p_1 - p_2)^2, p_2^2, m_0^2, m_1^2, m_2^2) = \frac{1}{\mu^{2\epsilon} \int D^d k} \frac{1}{[k^2 - m_0^2][(k + p_1)^2 - m_1^2][(k + p_2)^2 - m_2^2]}. \quad (\text{E.8})$$

Sete tipos de funções C_0 foram encontradas, a saber

$$C_0(0, 0, \hat{s}, 0, 0, 0) = \frac{i}{16\pi^2} N_\epsilon \left(\frac{\mu^2}{\hat{s}} \right)^\epsilon \frac{1}{\hat{s}} \left[\frac{1}{\epsilon^2} - \frac{\pi^2}{2} \right], \quad (\text{E.9})$$

$$\begin{aligned} C_0(0, 0, \hat{t}, m^2, 0, 0) &= \frac{-i}{16\pi^2} N_\epsilon \left(\frac{\mu^2}{\hat{s}} \right)^\epsilon \frac{1}{\hat{t}} \left[\frac{1}{\epsilon} \ln \left(1 - \frac{\hat{t}}{m^2} \right) \right. \\ &\quad - \ln \left(\frac{-\hat{t}}{\hat{s}} \right) \ln \left(1 - \frac{\hat{t}}{m^2} \right) + \frac{1}{2} \ln^2 \left(\frac{-\hat{t}}{m^2} \right) \\ &\quad \left. - \frac{1}{2} \ln^2 \left(1 - \frac{m^2}{\hat{t}} \right) - Li_2 \left(\frac{-\hat{t}}{m^2} \right) + Li_2 \left(\frac{\hat{t}}{\hat{t} - m^2} \right) \right] \end{aligned} \quad (\text{E.10})$$

$$C_0(0, \hat{t}, \hat{s} + \hat{t} + \hat{u}, 0, 0, m^2) = C_0(0, 0, \hat{t}, m^2, 0, 0), \quad (\text{E.11})$$

$$C_0(0, s, s + t + u, m^2, 0, 0) = \frac{i}{16\pi^2} \frac{1}{2\hat{s}} \ln^2 \left(1 + \frac{\hat{s}}{m^2} \right), \quad (\text{E.12})$$

$$C_0(0, 0, 0, m^2, 0, 0) = \frac{i}{16\pi^2} N_\epsilon \left(\frac{\mu^2}{m^2} \right)^\epsilon \frac{1}{m^2} \frac{1}{\epsilon}, \quad (\text{E.13})$$

$$C_0(0, 0, 0, 0, 0, m^2) = C_0(0, 0, 0, m^2, 0, 0), \quad (\text{E.14})$$

$$C_0(0, p^2, p^2, 0, 0, m^2) = \frac{-i}{16\pi^2} N_\epsilon \frac{1}{\epsilon} + \text{termos finitos}. \quad (\text{E.15})$$

4. Função Escalar de 4-Pontos

A integral escalar de 4-pontos é dada por

$$D_0(p_1^2, (p_1 - p_2)^2, (p_2 - p_3)^2, p_3^2, p_2^2, (p_1 - p_3)^2, m_0^2, m_1^2, m_2^2, m_3^2) = \frac{1}{\mu^{2\epsilon} \int D^d k} \frac{1}{[k^2 - m_0^2][(k + p_1)^2 - m_1^2][(k + p_2)^2 - m_2^2][(k + p_3)^2 - m_3^2]} \quad (\text{E.16})$$

O cálculo da integral D_0 que surge na amplitude da caixa foi baseado nos métodos de [75].

A seguir apresentamos o resultado apenas das partes divergentes.

$$D_0(0, \hat{s}, \hat{s} + \hat{t} + \hat{u}, \hat{t}, 0, 0, 0, 0, 0, m^2) = \frac{i}{16\pi^2} N_\epsilon \left(\frac{\mu^2}{\hat{s}} \right)^\epsilon \frac{1}{\hat{s}(m^2 - \hat{t})} \left[-\frac{1}{\epsilon^2} + \frac{2}{\epsilon} \ln \left(1 - \frac{\hat{t}}{m^2} \right) \right] + \text{partes finitas} \quad (\text{E.17})$$

O cálculo da partes finitas da D_0 é extremamente complicado, em função disso apresentamos a seguir uma maneira alternativa de obter a expressão da D_0 .

A função escalar D_0 de Passarino-Veltman.

Mesmo com todos os argumentos a favor da finitude da função I_3 discutido no capítulo 7, é preciso calculá-la explicitamente de maneira a extrair as suas partes finitas. A maior dificuldade dessa tarefa é o cálculo explícito da função escalar D_0 .

Em primeiro lugar, não pudemos encontrar na literatura uma tal função que, a exemplo da nossa, possuisse 3 propagadores não-massivos e 1 massivo calculada em regularização dimensional, portanto, mesmo com o cálculo explícito, não poderíamos comparar nossos resultados de modo a obter confiança a respeito das partes finitas principalmente, dado que a estrutura de pólos está fixada pelo fato de que I_3 é finita. Ainda assim, uma tentativa foi feita seguindo os métodos e sugestões da ref. [75]. A estrutura de pólos foi de fato confirmada, mas as partes finitas ainda deixavam dúvidas.

Felizmente fomos capazes de encontrar uma solução simples e segura para o cálculo da D_0 .

Como argumentamos, a combinação de integrais escalares

$$\mathcal{F} = \hat{s}[C_0(0, 0, 0, \hat{s}, 0, 0) + (m^2 - \hat{t}) D_0] + 2\hat{t} C_0(0, 0, \hat{t}, m^2, 0, 0) \quad (\text{E.18})$$

é finita em 4 dimensões. A questão é: sendo $\mathcal{F}$ finita, é possível que possamos identificá-la com alguma função escalar finita? Caso isso fosse possível, estaríamos ao mesmo tempo confirmando a finitude de I_3 e expressando a D_0 em função de outras funções escalares mais simples de serem calculadas através da expressão (E.18) acima. Veremos que isso de fato é possível!

Primeiramente vamos reescrever (E.18) de uma forma mais conveniente

$$\mathcal{F} = \mu^{2\varepsilon} \int D^d q \frac{q^2}{D_4} (\hat{s} + 2\hat{t}) + \mu^{2\varepsilon} \int D^d q \frac{2q \cdot Q}{D_4} , \quad (\text{E.19})$$

onde o 4-vetor

$$Q^\mu = \hat{s}(k_1^\mu - p_a^\mu) + \hat{t}(p_b^\mu - p_a^\mu)$$

goza das seguintes propriedades

$$p_a \cdot Q = 0$$

$$p_b \cdot Q = 0 .$$

Além disso, o cálculo explícito da integral

$$\mu^{2\varepsilon} \int D^d q \frac{q^2}{D_4} \equiv \tilde{C}_0 \quad (\text{E.20})$$

mostra que ela é finita no limite $\varepsilon \rightarrow 0$ e dada por (E.12).

Dessa maneira, $\mathcal{F}$ é dada por

$$\mathcal{F} = (\hat{s} + 2\hat{t})\tilde{C}_0 + \mu^{2\varepsilon} \int D^d q \frac{2q \cdot Q}{D_4} . \quad (\text{E.21})$$

Se pudermos mostrar que o segundo termo pode ser escrito em termos de uma combinação finita de integrais escalares, então nosso objetivo terá sido alcançado.

Vamos, portanto, nos concentrar nesta integral.

Escrevendo-a explicitamente, temos o seguinte

$$\mu^{2\varepsilon} \int D^d q \frac{2q \cdot Q}{D_4} = \hat{s} \mu^{2\varepsilon} \int D^d q \frac{2q \cdot (k_1 - p_a)}{D_4} + \hat{t} \mu^{2\varepsilon} \int D^d q \frac{2q \cdot (p_b - p_a)}{D_4} , \quad (\text{E.22})$$

veja agora que o numerador da primeira integral pode ser escrito como

$$\begin{aligned} (q + p_a - k_1)^2 - m^2 &= q^2 - m^2 + 2q \cdot (p_a - k_1) + (p_a - k_1)^2 \\ &= q^2 + \hat{t} - m^2 + 2q \cdot (p_a - k_1) \\ 2q \cdot (p_a - k_1) &= q^2 + (\hat{t} - m^2) - [(q + p_a - k_1)^2 - m^2] . \end{aligned} \quad (\text{E.23})$$

O numerador da segunda integral, por sua vez, pode ser reescrito usando as identidades

$$2q \cdot p_b = q^2 - (q - p_b)^2 \quad (\text{E.24})$$

$$-2q \cdot p_a = q^2 - (q + p_a)^2 , \quad (\text{E.25})$$

de forma que

$$\begin{aligned} \mu^{2\epsilon} \int D^d q \frac{2q \cdot Q}{D_4} &= \hat{s} \left[\mu^{2\epsilon} \int D^d q \frac{q^2}{D_4} + (\hat{t} - m^2) \mu^{2\epsilon} \int D^d q \frac{1}{D_4} \right. \\ &\quad \left. - \mu^{2\epsilon} \int D^d q \frac{(q + p_a - k_1)^2 - m^2}{D_4} \right] \\ &\quad + \hat{t} \left[\mu^{2\epsilon} \int D^d q \frac{2q^2}{D_4} - \mu^{2\epsilon} \int D^d q \frac{(q + p_a)^2 + (q - p_b)^2}{D_4} \right]. \end{aligned} \quad (\text{E.26})$$

Agora é fácil identificar as integrais escalares

$$\mu^{2\epsilon} \int D^d q \frac{2q \cdot Q}{D_4} = \hat{s} [\tilde{C}_0 - (m^2 - \hat{t}) D_0 - C_0(0, 0, 0, \hat{s}, 0, 0)] + 2\hat{t} \tilde{C}_0 - 2\hat{t} C_0(0, 0, \hat{t}, m^2, 0, 0). \quad (\text{E.27})$$

Substituindo na expressão de $\mathcal{F}$ temos, finalmente

$$\begin{aligned} \mathcal{F} &= (\hat{s} + 2\hat{t}) \tilde{C}_0 + (\hat{s} + 2\hat{t}) \tilde{C}_0 - \{ \hat{s} [C_0(0, 0, 0, \hat{s}, 0, 0) + (m^2 - \hat{t}) D_0] + 2\hat{t} C_0(0, 0, \hat{t}, m^2, 0, 0) \} \\ &= 2(\hat{s} + 2\hat{t}) \tilde{C}_0 - \mathcal{F}, \\ \mathcal{F} &= (\hat{s} + 2\hat{t}) \tilde{C}_0. \end{aligned} \quad (\text{E.28})$$

A função D_0 pode agora ser escrita em termos de funções escalares de 3-pontos, muito mais simples de serem calculadas

$$D_0 = \frac{1}{\hat{s}(m^2 - \hat{t})} [(\hat{s} + 2\hat{t}) \tilde{C}_0 - \hat{s} C_0(0, 0, 0, \hat{s}, 0, 0) - 2\hat{t} C_0(0, 0, \hat{t}, m^2, 0, 0)]. \quad (\text{E.29})$$

Essa expressão para a D_0 foi a usada na implementação numérica dos resultados em NLO de QCD.

A função $F(\hat{s}, \hat{t}, \hat{u})$ citada após a equação (7.68) é dada por

$$F(\hat{s}, \hat{t}, \hat{u}) = \frac{1}{2} \hat{s} m^2 + \frac{1}{4\hat{s}\hat{u}} [\hat{t}(\hat{t} - 2\hat{u})(\hat{s}^2 - \hat{t}^2) - 3\hat{u}^2(\hat{s}^2 - \hat{u}^2) - 2\hat{t}\hat{u}^2(\hat{t} - \hat{u})]. \quad (\text{E.30})$$

Referências

- [1] M. Gell-Mann, Phys. Lett. **8**, 214 (1964); G. Zweig, Erice Lecture 1964, em *Symmetries in Elementary Particle Physics*, A. Zichichi, Ed. Academic Press, New York, 1965.
- [2] O.W. Greenberg, Phys. Rev. Lett. **13**, 598 (1964).
- [3] B.R. Martin e G. Shaw, *Particle Physics* (John Wiley & Sons, New York, 1992).
- [4] R. P. Feynman, *Photon Hadron Interactions*, W. A. Benjamin, New York.
- [5] H. Fritzsch, M. Gell-Mann e H. Leutwyler, Phys. Lett. **B47**, 365 (1973).
- [6] D.J. Gross e F. Wilczek, Phys. Rev. **D10**, 3633 (1973).
- [7] C.N. Yang e R.L. Mills, Phys. Rev. **96**, 191 (1954).
- [8] I. Hinchliffe e A. Manohar, hep-ph/0004186.
- [9] W. Greiner e J. Reinhardt, *Field Quantization* (Springer-Verlag, Berlin, 1996).
- [10] L.D. Fadeev e Y.N. Popov, Phys. Lett. **B25**, 29 (1967).
- [11] T. Muta, *Foundations on Quantum Chromodynamics* (World Scientific Lectures Notes in Physics - Vol. 5, Singapore, 1987).
- [12] P. Pascual e R. Tarrach, Nucl. Phys. **B174**, 123 (1980).
- [13] M.E. Peskin e C.V. Schroeder, *An Introduction to Quantum Field Theory* (Addison-Wesley Publishing Company, 1995).
- [14] C. Itzykson e J.B. Zuber, *Quantum Field Theory*, (McGraw-Hill, New York, 1980).
- [15] C. G. Bollini e J. J. Giambiagi, Phys. Lett. **B40**, 566 (1972); G. 't Hooft e M. Veltman, Nucl. Phys. **B44**, 189 (1972).

- [16] G. 't Hooft, Nucl. Phys. **B190**, 455 (1981); K. Wilson, Phys. Rev. **D10**, 2445 (1974); R.W. Haymaker, Phys.Rep. **315**, 153 (1999).
- [17] S.G. Gorishny, A.L. Kataev e S.A. Larin, Phys. Lett., **B259**, 144 (1991); L.R. Surguladze e M.A. Samuel, Phys. Rev. Lett. **66**, 560 (1991); K.G. Chetyrkin, A.L. Kataev e F.V. Tkachov, Phys. Lett. **B85**, 277 (1979); M.Dine e J.Sapirstein, Phys. Rev. Lett. **43**, 668 (1979); W.Celmaster e R.Gonsalves, Phys. Rev. Lett. **44**, 560 (1980).
- [18] P. Nason e C. Oleari, Nucl. Phys. **B521**, 237 (1998).
- [19] R.K. Ellis, D.A. Ross e E. Terrano, Nucl.Phys **B178**, 421 (1981); G. Rodrigo, *Quark mass effects in QCD jets.*, hep-ph/9703359 (Tese de Doutorado).
- [20] R.D. Field, *Applications of Perturbative QCD* (Addison-Wesley, Frontiers in Physics, Lecture Notes Series - Vol. 77, 1989).
- [21] W.J. Marciano, Phys. Rev. **D12**, 3861 (1975).
- [22] R. Mertig, Comput. Phys. Commun. **64**, 345 (1991).
- [23] G. 't Hooft, Nucl. Phys. **B33**, 173 (1971); G. 't Hooft, Nucl. Phys. **B35**, 167 (1971).
- [24] T. Kinoshita, J. Math. Phys. **3**, 650 (1962); T.D. Lee e M. Nauenberg, Phys. Rev. **133**, B1549 (1964).
- [25] F. Bloch e A. Nordsieck, Phys. Rev. **52**, 54 (1937).
- [26] S. Cattani e M.H. Seymour, Nucl. Phys. **B485**, 291 (1997) [hep-ph/9605323].
- [27] T. Plhen, *Production of supersymmetric particles at high-energy colliders.*, hep-ph/9809319 (Tese de Doutorado).
- [28] W. Beenakker, R. Höpker, M. Spira e P.M. Zerwas, Nucl. Phys. **B492**, 51 (1997) [hep-ph/9610490].
- [29] W. Beenakker, H. Kuijf e W.L. van Neerven, Phys. Rev. **D40**, 54 (1989).
- [30] M. Krämer, Nucl. Phys. **B459**, 3 (1996) [hep-ph/9508409].
- [31] M.H. Reno e H. Baer, Phys. Rev. **D43**, 2892 (1991); I. Avaliani, hep-ph/9505390.
- [32] J. Ohnemus, Phys.Rev **D47**, 940 (1993).

- [33] U. Baur, T. Han e J. Ohnemus, Phys. Rev. **D48**, 5140 (1993); U. Baur, T. Han e J. Ohnemus, Phys. Rev. **D51**, 3381 (1995).
- [34] J. Ohnemus e J.F. Owens, Phys. Rev. **D43**, 3626 (1991).
- [35] J. Ohnemus, J.F. Owens e H. Baer, Phys. Rev. **D42**, 61 (1990); J. Ohnemus e J.F. Owens, Phys. Lett. **B234**, 127 (1990).
- [36] L. Bergmann, FSU-HEP-890215 (Tese de Doutorado).
- [37] K. Fabricius, G. Kramer, G. Schierholz e I. Schmitt, Z. Phys. **C11**, 315 (1981); G. Kramer e B. Lampe, Fortschr. Phys. **37**, 161 (1989).
- [38] W.T. Giele e E.W.N. Glover, Phys. Rev. **D46**, 1980 (1992).
- [39] G. Sterman e S. Weinberg, Phys. Rev. Lett. **39**, 1436 (1977).
- [40] J.C. Collins, D.E. Soper e G. Sterman, *Factorization of hard processes in QCD*. em *Perturbative Quantum Chromodynamics*, A.H. Mueller, Advanced Series on Directions in High-Energy Physics - Vol. 5, Ed. World Scientific; D.E. Soper, hep-lat/9609018.
- [41] R.K. Ellis, W.J. Stirling e B.R. Webber, *QCD and Collider Physics* (Cambridge University Press, 1996).
- [42] R. Basu *et al.*, Nucl. Phys. **B244**, 221 (1984).
- [43] R. Doria, J. Frenkel e J.C. Taylor, Nucl. Phys. **B168**, 93 (1980).
- [44] K.G. Wilson, Phys. Rev. **179**, 1499 (1969).
- [45] A. Manohar, hep-ph/9204208.
- [46] S.S. Willenbrock, *QCD Corrections to $p\bar{p} \rightarrow W^+ + X$: A Case Study.*, Lecture Notes, TASI-89.
- [47] G. Altarelli, R.K. Ellis e G. Martinelli, Nucl. Phys. **B157**, 461 (1979).
- [48] G. Altarelli e G. Parisi, Nucl. Phys. **B126**, 298 (1977).
- [49] G. Passarino e M. Veltman, Nucl. Phys. **B160**, 151 (1979).
- [50] R. Gastmans e J. Verwaest, Nucl. Phys. **B105**, 454 (1976).
- [51] T. Gottschalk, E. Monsay e D. Sivers, Phys. Rev. **D21**, 1799 (1980).

- [52] W.T. Giele, E.W.N. Glover e D.A. Kosower, Nucl. Phys. **B403**, 633 (1993).
- [53] P. Aurenche e J. Lindfors, Nucl. Phys. **B185**, 274 (1981).
- [54] H. Georgi e S.L. Glashow, Phys. Rev. Lett. **32**, 438 (1974).
- [55] J.C. Pati e A. Salam, Phys. Rev. **D10**, 275 (1974); G. Senjanovic e A. Sokorac, Z.Phys. **C20**, 255 (1983).
- [56] K. Lane, hep-ph/9401324.
- [57] B. Schremp e F. Schremp, Phys. Lett. **B153**, 101 (1985).
- [58] E. Witten, Nucl. Phys. **B258**, 75 (1985); M. Dine *et al.*, Nucl. Phys. **B259**, 519 (1985); J.L. Hewett e T.G. Rizzo, Phys. Rep. **183**, 193 (1989).
- [59] H1 Collab., C. Adloff *et al.*, Z. Phys. **C74**, 191 (1997); ZEUS Collab., J. Breitweg *et al.*, Z. Phys. **C74**, 207 (1997).
- [60] *Particle Data Group*, Eur. Phys. J. **C3**, 1 (1998).
- [61] J.K. Mizukoshi, O.J.P. Éboli e M.C. Gonzalez-Garcia, Nucl. Phys. **B443**, 20 (1995) [hep-ph/9411392].
- [62] V. Barger e K. Cheung, hep-ph/0002259; A.F. Zarnecki, hep-ph/0003271.
- [63] R. Rückl e P.M. Zerwas, CERN-87-07 (1987) em *La Thuile/CERN Workshop on Physics at Future Accelerators*, ed. J.M. Mulvey.
- [64] W. Buchmüller, R. Rückl e D. Wyler, Phys. Lett. **B191**, 442 (1987).
- [65]] Z. Kunszt e W.J. Stirling, Z. Phys. **C75**, 453 (1997) [hep-ph/9703427].
- [66] P. Nason, S. Dawson e R.K. Ellis, Nucl. Phys. **B303**, 607 (1988).
- [67] J. Collins, F. Wilczek e A. Zee, Phys. Rev. **D18**, 242 (1978); W.J. Marciano, Phys. Rev. **D29**, 580 (1984).
- [68] B. Mele, P. Nason e G. Ridolfi, Nucl. Phys. **B357**, 409 (1991).
- [69] R.G. Stuart, Phys. Lett. **B262**, 113 (1991); S. Willenbrock e G. Valencia, Phys. Lett. **B259**, 373 (1991); M. Nowakowski e A. Pilaftsis, hep-ph/9305321.
- [70] H.L. Lai *et al.*, CTEQ Group, Phys. Rev. **D55**, 1280 (1997) [hep-ph/9606399].

- [71] G.P. Lepage, *VEGAS: An Adaptative Multidimensional Integration Program*, CLNS-80/447 (1980).
- [72] W.L. van Neerven e E.B. Zijlstra, Nucl. Phys. **B382**, 11 (1992).
- [73] E. Gluck, E. Reya e A. Vogt, Z. Phys. **C67**, 433 (1995).
- [74] A.D. Martin, R.G. Roberts e W.J. Stirling, RAL-93-077 (1993).
- [75] G. Rodrigo e A. Santamaria, hep-ph/9703360.

