

COMMENTS ON GLOBAL PARTON ANALYSES*

A.D. MARTIN^a, M.G. RYSKIN^{a,b}^aInstitute for Particle Physics Phenomenology, University of Durham
Durham, DH1 3LE, United Kingdom^bPetersburg Nuclear Physics Institute, Kurchatov National Research Centre
Gatchina, St. Petersburg 188300, Russia*(Received January 20, 2020)*

We discuss the causes which can limit the accuracy of the predictions based on the conventional PDFs when including in global parton analyses the data at moderate scales μ . The first is the existence of power corrections $\mathcal{O}(Q_0^2/\mu^2)$ due to the double counting of contributions arising from the region below the input scale Q_0 . The second concerns the possible inclusion of the BFKL re-summation of the $(\alpha_s \ln(1/x))^n$ terms. The third is the treatment of the heavy-quark thresholds. We show how to include the heavy-quark masses (m_h with $h = c, b, t$) in DGLAP evolution which provides the correct *smooth* behaviour through the threshold regions and how to subtract the low parton virtuality $|k^2| < Q_0^2$ contributions from the DIS and Drell–Yan NLO coefficient functions in order to avoid the double counting.

DOI:10.5506/APhysPolB.51.1079

1. Introduction

Recall that the framework for parton analysis is based on the factorization theorem and DGLAP evolution, both formulated and justified for very large scales where the QCD coupling is small and perturbation theory is applicable. Due to the strong k_T ordering of the emitted partons during the evolution, all the contributions from the low k_T , confinement region, can be isolated and factorized into the input parton distribution functions (PDFs) at some boundary scale Q_0 which is not very high, but sufficiently large to justify the applicability of DGLAP evolution, and is smaller than the factorization scale μ_F which separates the ‘hard’ matrix elements describing the subprocess from the partons described by the evolution.

* Funded by SCOAP³ under Creative Commons License, CC-BY 4.0.

As far as we include the NLO, NNLO, ... corrections in the hard matrix element and in the DGLAP splitting functions, there appear loop integrals which contain some contribution from the region with $k_T < Q_0$. Provided $Q_0 \ll \mu_F$, this is not a danger since there are no infrared divergences in the corresponding loop integrals. The contribution from $k_T < Q_0$ may be treated as a power correction of $\mathcal{O}(Q_0^2/\mu_F^2)$ (or even less depending on the particular process).

The situation becomes more complicated when we include in a global parton analysis data with only a moderate scale. In this case, the correction $\mathcal{O}(Q_0^2/\mu_F^2)$ becomes crucial. In the present note, we will discuss three topics relevant when a global analysis includes data of processes at scales comparable to Q_0 .

The first problem is that we have to avoid double counting of the $k_T < Q_0$ contribution, which on the one hand was included in the input PDF, while on the other hand is sampled in the loop integrals in the coefficient functions determining the hard matrix elements. Next, we consider the BFKL re-summation and emphasize the fact that at a not too large scale, the leading order BFKL amplitude is strongly affected by the boundary condition that we have to put at some $k_T \simeq Q_0$ [1]. Here, we also have to exclude the possible contributions from the region with $k_T < Q_0$. Finally, we discuss the treatment of the heavy-quark thresholds. Usually, the contribution of a heavy quark, h , is completely neglected for the scales $Q^2 < m_h^2$, while for a Q above the quark mass, m_h , the heavy-quark evolution is described by the same (massless) expressions as that for the light quarks. This is not a danger when we are interested in parton distributions at a large scale $\mu_F \gg m_h$. It is possible to account for the heavy-quark mass using appropriate ‘matching conditions’ such as ACOT [2] or RT [3]. However, working at a scale comparable with the quark mass, it is better to use the splitting functions (at least at LO) which account for the value of m_h from the beginning. To include the mass m_h in the corresponding Feynman diagrams explicitly can be especially important for the running of the QCD coupling $\alpha_s(\mu)$ (see *e.g.* [4] and Fig. 3 below).

We discuss these three topics in turn in the following three sections.

2. Double counting and the Q_0 subtraction

Recall that the idea of factorization is to separate the small and large virtuality contributions. Formally, the coefficient functions correspond to large virtualities, while all the low virtuality contributions are collected in some phenomenological input. Simultaneously, we have to exclude the low virtuality contributions from the hard matrix element. Otherwise, there will be the *double counting*.

The Q_0 subtraction should therefore be done for every observable fitted in a global analysis. Without doing the Q_0 subtraction, the precision of the PDFs cannot be better than $\mathcal{O}(\alpha_s Q_0^2/\mu_F^2)$, since the contribution from $k_T < Q_0$ is not under control.

2.1. Physical scheme

Strictly speaking, there are two different types of $k_T < Q_0$ contributions. First, there are the contributions from extremely large distances ($k_T \rightarrow 0$) arising from the ϵ -regularization prescription. The point is that in order to regularize the ultraviolet (UV) divergence, the loop integrals were calculated in $4 + 2\epsilon$ dimensional space (with $\epsilon \rightarrow 0$) where the logarithmic divergence results in $1/\epsilon$ terms (which are finally cancelled in the minimal subtraction scheme). However, simultaneously, the $1/\epsilon$ terms come from unphysically large distances, that is from the infrared (IR) ($k_T \rightarrow 0$) region. Together with the parts proportional to ϵ in the splitting and the coefficient functions, these IR $1/\epsilon$ terms give some finite constant ϵ/ϵ contributions.

On the one hand, such an ϵ/ϵ contribution is unphysical. It comes from an infinitely large distances which are forbidden by confinement. However, it turns out easier not to fight with it, but to retain it via a re-definition of the factorization scheme. Indeed, since the ϵ/ϵ contribution, $\Delta C_a(z)$, in the NLO coefficient function calculated within the $\overline{\text{MS}}$ approach originates from very small $k_T \rightarrow 0$, it can be written as the convolution

$$\Delta C_a \equiv C_a^{\text{NLO}}(\overline{\text{MS}}) - C_a^{\text{NLO}}(\text{phys}) = \frac{\alpha_s}{2\pi} \sum_b C_b^{\text{LO}} \otimes \delta P_{ab}(z), \quad (1)$$

where $\delta P_{ab}(z)$ is the proportional to ϵ part of the LO $\overline{\text{MS}}$ splitting

$$P_{ab}^{\overline{\text{MS}}}(z) = P_{ab}^{\text{LO}}(z) + \epsilon \delta P_{ab}(z), \quad (2)$$

and $a, b = g, q$ denote the type of partons, while $\otimes$ denotes the convolution in z distribution. The difference in coefficient functions, ΔC_a , can be compensated by a redefinition of the parton distributions

$$\begin{aligned} a^{\overline{\text{MS}}}(x) &= a^{\text{phys}}(x) - \frac{\alpha_s}{2\pi} \int \frac{dz}{z} \sum_b \delta P_{ab}(z) b^{\text{phys}}(x/z) \\ &\equiv a^{\text{phys}} - \frac{\alpha_s}{2\pi} \sum_b \delta P_{ab} \otimes b^{\text{phys}}. \end{aligned} \quad (3)$$

Correspondingly, if we re-define the splitting function, then we reproduce at NLO level the original DGLAP evolution (see [5] for details).

That is working at NLO, in the $\overline{\text{MS}}$ scheme we do not deal with the original (physical) quarks and gluons, which are pictured in Feynman diagrams but, instead, with a slightly ‘rotated’ partons where a quark/gluon with momentum fraction x has an $O(\alpha_s)$ admixture of other partons which can be of another type and may carry another momentum fraction. This is not a danger but one has to clearly understand what was calculated. In this respect, see the comment at the end of the introduction to Section 4.

2.2. Subtraction of the contribution from finite $k_T < Q_0$

Unfortunately, we cannot replace the subtraction of the contribution from finite $k_T < Q_0$ just by the choice of a new factorization scheme¹. The problem is that this contribution depends on the particular ‘hard’ matrix element (say, on the transverse jet energy, E_T , in the case of the coefficient function for dijet production) and on the factorization scale. Therefore, it is impossible to re-define the splitting functions in such a way as to ensure the same DGLAP evolution of universal PDFs which can be used at different factorization scales and be convoluted with different ‘hard’ matrix elements². The corrections are large for $\mu_F \sim Q_0$, while in the limit of $\mu_F \gg Q_0$, the corrections become negligibly small and we come back to the ‘physical’ (or the $\overline{\text{MS}}$) scheme. The only way to avoid double counting is to exclude the $k_T < Q_0$ contribution (analogous to that which occurs in DGLAP evolution) from the perturbative NLO (and higher α_s order) calculations moving it to some phenomenological input at Q_0 .

We emphasize that these $k_T < Q_0$ contributions are not admixtures of higher twist terms. Recall that twist is defined as the dimension of the operator minus its spin. When we calculate the $k_T < Q_0$ contribution, we deal with the same operator (of the same spin and dimension). That is we are concerned with the same twist. So it is just a *power* correction to the contribution of the old leading twist operator (of conventional DGLAP).

Numerically, these power corrections are most important at relatively low scales. In such a case, we practically have no place for the logarithmic DGLAP evolution. Thus first of all we have to consider the corrections (caused by the subtraction of $k_T < Q_0$ contributions) to the coefficient

¹ Recall that factorization theorems are proven within the logarithmic approximation; that is assuming a strong $k_{i-1} \ll k_i$ ordering (where $k_i \equiv k_{Ti}$). In this limit, we can consider the ϵ/ϵ contributions coming from very large distances (which satisfy $k \ll k_i$) as that corresponding to another factorization scheme. However, the Q_0^2/k_i^2 power correction which (a) is not negligibly small and (b) cannot be written in terms of the (one or a few powers of) $\ln(Q_0/k_i)$ do depend on a particular process and so cannot be accounted for by choosing an appropriate scheme.

² In other words, if we replace the Q_0 subtraction by another factorization scheme, then we are unable to justify the factorization.

functions, process by process. As examples we present in Appendix the results for the NLO coefficient functions in DIS and for the Drell–Yan lepton pair production.

Note also that the Q_0^2/μ^2 power correction in the splitting function destroys the logarithmic structure of the evolution in $\ln(\mu^2)$. The major part of this correction corresponding to the lower limit of the integral is absorbed into the phenomenological input PDF while the upper limit of the integral contains an additional QCD coupling α_s *without* the $\ln(\mu^2)$ and should be considered as the power correction to the next (now NNLO, since we are talking about the NLO contribution), *i.e.* a higher order α_s term.

3. BFKL re-summation

To improve the accuracy of the PDF determinations in the low- x region, the calculations of the splitting and the coefficient functions are often supplemented by the re-summation of the $(\alpha_s \ln(1/x))^n$ terms generated by the BFKL equation (see, for example, [6, 7] and [8] for a short review). For a large scales, this is a good procedure. However, there may be a danger using the BFKL re-summation at scales comparable with Q_0 .

First, in this region the solution of BFKL equation is strongly affected by the boundary condition at $k_T = Q_0$ [1]. While at very large scales, the x behaviour is controlled by the position of the vacuum (BFKL) singularity in complex j -plane, at k_T close to the confinement region the x behaviour is driven by an unknown boundary condition. This fact is usually not accounted for (when implementing in BFKL re-summation) by keeping only the BFKL results justified for large scales $\mu \gg Q_0$. On the other hand, we have no such a problem in a pure DGLAP approach where this input x behaviour at scale equal to Q_0 is considered as the phenomenological function fitted from the experiment.

Next, we have to recall that the BFKL equation includes not only the leading twist contributions but also higher twists as well. These higher twists are hidden in the gluon reggeization terms which cannot be neglected since without these terms we are unable to eliminate the IR singularity of the BFKL kernel. Thus after the BFKL re-summation is included, we cannot claim that the resulting predictions correspond to a leading twist contribution.

4. Heavy-quark thresholds

The correct treatment of heavy quarks in an analysis of parton distributions is essential for precision measurements at hadron colliders. The up, down and strange quarks, with $m^2 \ll Q_0^2$, can be treated as massless partons.

However, for charm, bottom or top quarks, we must allow for the effects of their mass, m_h with $h = c, b$ or t . The problem is that we require a consistent description of the evolution of parton distribution functions (PDFs) over regions which include *both* the $Q^2 \sim m_h^2$ domain and the region $Q^2 \gg m_h^2$ where the heavy quark, h , can be treated as an additional massless quark.

During the logarithmic DGLAP evolution in $\ln Q^2$, the quark mass affects the splitting function only within a *finite* interval of $\ln Q^2$ (that is at $Q^2 \sim m_h^2$). Thus, the mass correction to the leading order splitting function enters at the same level as the NLO correction. In other words, it is sufficient to include the mass corrections to the LO splitting function to provide the NLO accuracy (and so on — the mass corrections to the NLO functions provide NNLO accuracy). We give more detail how this arises below.

We have to emphasize that these mass corrections should be implemented in the *physical* factorization scheme where the heavy-quark PDF has *no* admixture of gluons or light quarks.

4.1. NLO heavy-quark mass effects already included at LO

We have just mentioned that as the heavy-quark mass effects come only from a finite interval of the $\ln Q^2$ evolution to reach the NLO accuracy, it is sufficient to account for m_h only in the LO diagrams. Moreover, if we keep the mass in the NLO (two-loop) graphs, then it leads to an NNLO correction. It is informative to describe in more detail how this happens.

As usual, we use the axial gauge, where only the ladder (real emission) and the self-energy (virtual-loop contribution) diagrams give Leading Logarithms. Actually, for real emission, we need to consider only the ‘gluon-to-heavy quark’ splitting function. Indeed, the heavy-quark mass effects can be identified in the following subset of integrations:

$$\dots \int \frac{dk_{i-1}^2}{k_{i-1}^2} \int \frac{dk_i^2 k_i^2}{(k_i^2 + m_h^2)^2} \int \frac{dk_{i+1}^2}{k_{i+1}^2} \dots \quad (4)$$

corresponding to the part of the parton chain containing the $g \rightarrow h\bar{h}$ transition, as shown in Fig. 1. The k^2 s are the virtualities of the t -channel partons, and the heavy-quark mass effects enter in the k_i^2 integration that results from the $g \rightarrow h\bar{h}$ transition. The kinematics responsible for the LO result are when the virtualities are strongly ordered ($\dots k_{i-1}^2 \ll k_i^2 \ll k_{i+1}^2 \dots$). If two of the partons have comparable virtuality, $k_j^2 \sim k_{j+1}^2$, then we lose a $\ln Q^2$ and obtain an NLO contribution of the form of $\alpha_s(\alpha_s \ln Q^2)^{n-1}$ for n emitted partons.

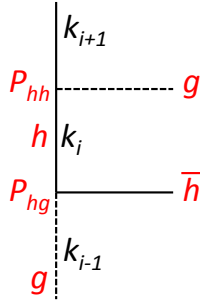


Fig. 1. Part of the parton evolution chain which contains the $g \rightarrow h\bar{h}$ transition.

At first sight, it appears that m_h^2 should also have been retained in the integration over the heavy-quark line with virtuality k_{i+1} . However, the heavy quark was produced at $Q^2 \sim m_h^2$ via the $g \rightarrow h$ splitting. Due to the strong ordering $k_{i+1}^2 \gg k_i^2$ in the evolution chain, we have $k_{i+1}^2 \gg m_h^2$, and so we may neglect m_h^2 in the k_{i+1}^2 integration; otherwise, this would be a NNLO effect.

Note that in our NLO calculations, described below, we use a fixed number $m_h(m_h)$ for the heavy-quark mass³. All the effects of the running quark mass should be regarded as part of the NNLO corrections.

4.2. Smooth evolution of α_s across a heavy-quark threshold

Here, we demonstrate the role of the effect of the heavy-quark mass in the running QCD coupling [4]. At NLO, the Q^2 evolution of $\alpha_s(Q^2)$ is described by the equation

$$\frac{d}{d \ln Q^2} \left(\frac{\alpha_s}{4\pi} \right) = -\beta_0 \left(\frac{\alpha_s}{4\pi} \right)^2 - \beta_1 \left(\frac{\alpha_s}{4\pi} \right)^3, \quad (5)$$

where the β -function coefficients are

$$\beta_0(n_f) = 11 - \frac{2}{3}n_f, \quad \beta_1(n_f) = 102 - \frac{38}{3}n_f. \quad (6)$$

The fermion loop insertion is responsible for the $-(2/3)n_f$ term in the LO β -function. Including the mass m_h we find that, instead of changing n_f from 3 to 4 (at $Q^2 = m_c^2$), and from 4 to 5 (at $Q^2 = m_b^2$), we must include in n_f a term

$$\kappa(r) = \left[1 - 6r + 12 \frac{r^2}{\sqrt{1+4r}} \ln \frac{\sqrt{1+4r} + 1}{\sqrt{1+4r} - 1} \right] \quad (7)$$

for each heavy quark, where $r \equiv m_h^2/Q^2$. In Fig. 2, we plot κ as a function of Q^2/m_h^2 .

³ Strictly speaking, we may choose any *reasonable* fixed value for m_h , say m_c (1.4 GeV), so that the NNLO correction is not large.

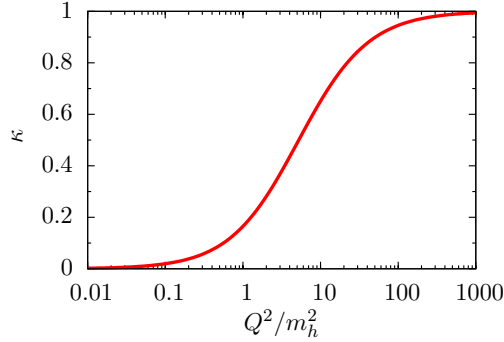


Fig. 2. The contribution of a heavy quark to the running of α_s showing a smooth behaviour across the heavy-quark threshold. If $\kappa = 1$, the heavy quark acts as if it were massless.

Next, in Fig. 3, we compare the evolution of α_s in which the effects of the heavy-quark masses are included, with an evolution assuming all quarks are massless. In the latter case, a prescription has been used to ensure that α_s is continuous across the heavy-quark thresholds. Different prescriptions are possible, but it is not possible to make the derivative also continuous, as can be seen from Fig. 3(b). Indeed, with massless evolution, different reasonable prescriptions can lead to a difference of more than 0.5% in going from $Q^2 \sim 20 \text{ GeV}^2$ up to $Q^2 = M_Z^2$. However, when the heavy-quark masses are properly accounted for, we see that the difference over this interval is about 4%, and in fact up to 14% starting from $Q^2 = 1 \text{ GeV}^2$. The fact that

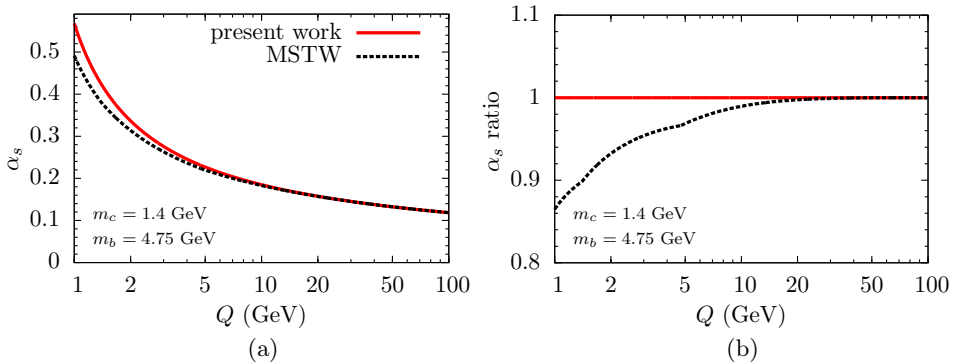


Fig. 3. (a) The running of α_s at NLO: the continuous curve is obtained with the effects of the heavy-quark masses m_c , m_b included, and the dashed curve is that used, for example, by the MSTW global parton analysis [9]. Both evolutions are normalised to $\alpha_s(M_Z^2) = 0.12$. (b) The ratio of the above two evolutions of α_s . The figure is taken from [4].

the α_s curve, obtained with mass effects included, lies consistently above that for massless evolution in Fig. 3(a) follows from the behaviour of κ in Fig. 2 and that we have required both curves to have $\alpha_s(M_Z^2) = 0.12$.

4.3. Heavy-quark mass effects in the LO splitting functions

We may summarize the LO evolution equations in the symbolic form

$$\begin{aligned}\dot{g} &= P_{gg} \otimes g + \sum_q P_{gq} \otimes q + \sum_h P_{gh} \otimes h, \\ \dot{q} &= P_{qg} \otimes g + P_{qq} \otimes q, \\ \dot{h} &= P_{hg} \otimes g + P_{hh} \otimes h,\end{aligned}\tag{8}$$

where $q = u, d, s$ denotes the light-quark density functions and $h = c, b, t$ are the heavy-quark densities. We have abbreviated P^{LO} by P , and $\dot{a} = (2\pi/\alpha_s)\partial a/\partial \ln Q^2$. The formulae for the individual splitting functions P_{ij} including the m_h effects can be found in [4]. In that paper, the splitting functions given in Eqs. (12) and (13) are in error. They should be replaced by

$$\begin{aligned}P_{hh}^{\text{real}}(z, Q^2) &= C_F \left(\frac{1+z^2}{1-z} \frac{Q^2}{(1-z)m_h^2 + Q^2} + z(z-3) \frac{Q^2 m_h^2}{(Q^2 + (1-z)m_h^2)^2} \right), \\ P_{gh}(z, Q^2) &= C_F \left(\frac{1+(1-z)^2}{z} \frac{Q^2}{zm_h^2 + Q^2} + (z^2 + z - 2) \frac{Q^2 m_h^2}{(Q^2 + zm_h^2)^2} \right),\end{aligned}$$

respectively⁴. The $z \leftrightarrow (1-z)$ symmetry between these two equations enables overall momentum conservation to be satisfied during the evolution.

Note that there are evolution equations, (8), for *all* type of partons (including heavy quarks) just starting from Q_0 . The input heavy-quark distribution $h(x, Q_0^2)$ should be treated as an ‘intrinsic’ PDF introduced in [10]. Of course, at low $Q^2 \ll m_h^2$, the corresponding splitting functions are strongly suppressed by the small value of the ratio Q^2/m_h^2 . So, actually the evolution of the heavy quark will start somewhere in the region $Q^2 \simeq m_h^2$.

4.4. Quark mass effects in NLO diagrams

It turns out that to include heavy-quark mass effects in NLO evolution, we do not need to modify the usual NLO splitting functions. In the absence of an intrinsic heavy quark, we only have to take m_h into account in P_{hg} and then only in the LO part $P_{hg}^{(0)}$. (Of course, as a consequence, we must adjust the virtual corrections to P_{gg} .) The argument is as follows.

⁴ We thank Valerio Bertone for drawing our attention to this error.

The k_i^2 integral of (4) written with NLO accuracy, has the form of

$$\int \frac{dk_i^2 A(k_i^2, k_{i+1}^2, m_h^2, z)}{(k_i^2 + m_h^2)^2} = \int A_1(z) \frac{d(k_i^2 + m_h^2)}{(k_i^2 + m_h^2)} + \int A_2(z) \frac{m_h^2 dk_i^2}{(k_i^2 + m_h^2)^2} + \int A_3(z) \frac{dk_i^2}{k_{i+1}^2}. \quad (9)$$

The first term gives the leading logarithm contribution. To be specific, we have

$$\int_{k_{i-1}^2}^{Q^2} \frac{dk^2}{(k^2 + m_h^2)} = \ln \frac{Q^2 + m_h^2}{m_h^2} \quad (10)$$

for $k_{i-1}^2 \ll m_h^2$. Both the second term in (9), which is concentrated in the region $k_i^2 \sim m_h^2$, and the third term, which is concentrated near the upper limit, at $k_i^2 \sim k_{i+1}^2$, give non-logarithmic contributions.

In the axial gauge, the two first terms on the right-hand side of (9) come only from the pure ladder (and the corresponding self-energy) diagrams, from the region of $k_i^2 \ll k_{i+1}^2$. That is, these two terms are exactly the same as those generated by $\text{LO} \otimes \text{LO}$ evolution, in which we have already accounted for the m_h effects. To avoid double counting, we have to subtract these contributions from (9). Thus, the true NLO contribution is given by the third term only, in which we can omit the m_h dependence since: (a) $k_{i+1}^2 \gg m_h^2$, and, (b) these order of $\mathcal{O}(m_h^2/k_{i+1}^2)$ terms kill the large logarithm in the further $\int dk_{i+1}^2/k_{i+1}^2$ integration. That is, at NLO accuracy, we can use the old, well-known, NLO splitting functions $P_{ik}^{(1)}(z)$. If we were to account for the mass effect in $P_{ik}^{(1)}(z)$, then we would be calculating an NNLO correction⁵.

In summary, to reach NLO accuracy, one may neglect the heavy-quark mass effects in the NLO splitting functions (where the quark mass results in a NNLO correction). Moreover, in the absence of an intrinsic heavy quark only the LO $P_{hg}^{(0)}$ needs to be modified.

5. Conclusion

We consider the role of low $k_T < Q_0$ contributions which can limit the accuracy of the parton distributions at moderate scales. We recall that:

⁵ Before proceeding to NNLO, a phenomenological way to provide very smooth behaviour of the NLO contribution would be to multiply the ‘heavy-quark’ NLO terms (that is, those NLO terms which contains the heavy quark) simply by the factor $Q^2/(Q^2 + m_h^2)$.

- In conventional DGLAP evolution, *all* the low virtuality contributions are collected in the input PDFs at a scale equal to Q_0 . Therefore, to avoid the double counting, we have to exclude the $|k^2| < Q_0^2$ loop integration from the NLO (and the higher α_s order) coefficient and splitting functions. Without doing this, the $|k^2| < Q_0^2$ contributions result in an order of $\alpha_s Q_0^2/\mu^2$ power corrections which are not under control. These corrections limit the accuracy of the pQCD predictions at moderate scales μ .
In Appendix, we present the formulae which allow the subtraction of the $|k^2| < Q_0^2$ terms from the NLO DIS and the Drell–Yan coefficient functions.
- An analogous subtraction is needed for the $(\alpha_s \ln(1/x))^n$ terms in the case of the BFKL re-summation. Moreover, note that at moderate scales, the behaviour of BFKL amplitude is strongly affected by the phenomenological boundary condition at Q_0 (which is not well-known at the moment).
- Finally, we consider the role of the heavy-quark mass effects and present the formulae which provide a smooth transition of the LO splitting functions over the heavy-quark threshold. We show that using these formulae, one can reach NLO accuracy while replacing an explicit mass effect by an appropriate matching of the massless expressions we already observe about a 4% correction in α_s value at $Q^2 = 20 \text{ GeV}^2$.

We thank Valerio Bertone, Robert Thorne and Stefano Forte for valuable discussions. M.G.R. thanks the IPPP at Durham University for hospitality.

Appendix

Power corrections to coefficient functions

Here, we describe in detail the power corrections which arise from the double counting of the $k_T < Q_0$ contribution using as examples the NLO coefficient functions for DIS and for the Drell–Yan production.

A. Deep inelastic scattering

A.1 Coefficient functions in the ‘physical’ renormalization scheme

Recall that to avoid double counting, we need to subtract the terms generated by the convolution of the LO splitting and the LO coefficient functions, $P^{\text{LO}} \otimes C^{\text{LO}}$. As a result, the NLO coefficient functions do not have an infrared divergency. Thus, we may perform an explicit calculation of the

corresponding Feynman diagrams; we have no problem with infrared regularization (and we automatically obtain a result in the ‘physical scheme’). However, the absence of infrared divergences does not exclude non-divergent contributions from a quark or gluon of low virtuality $|k^2|$; that is, from the region of $|k^2| < Q_0^2$. Moreover, in many cases (and, in particular, in the case of the NLO gluon contribution to F_2), we deal with exactly the same diagrams as those which occur in DGLAP evolution. Thus to be consistent, we have to exclude the soft, $|k^2| < Q_0^2$, contributions to the coefficient functions as well. This will result in power corrections of the order of Q_0^2/Q^2 for DIS where the value of $\mu_F^2 = Q^2$ is conventionally used.

The new DIS NLO coefficient functions, which account for the $Q^2 > Q_0^2$ cutoff are for the longitudinal structure function F_L , given by

$$C_{Lg}(z) = 4T_R z(1-z) \left(1 - z \frac{Q_0^2}{Q^2}\right), \quad (11)$$

$$C_{Lq}(z) = C_F 2z \left(1 - \left(z \frac{Q_0^2}{Q^2}\right)^2\right). \quad (12)$$

The expressions for C_{Lq} and C_{Lg} (where from the beginning there are no infrared divergences) are, in the limit of $Q_0 \rightarrow 0$, scheme-independent.

The situation is more complicated for the structure function F_2 . Here, taking into account the cutoff Q_0 , we find the NLO coefficient functions are

$$C_{2g}(z) = T_R \left\{ [(1-z)^2 + z^2] \ln \frac{1}{z} + [6z(1-z) - 1] \left(1 - z \frac{Q_0^2}{Q^2}\right) \right\}, \quad (13)$$

and *accounting for the Adler sum rule*

$$\begin{aligned} C_{2q}(z) = C_F & \left\{ \left(\frac{1+z^2}{1-z} \right) \ln \frac{1}{z} + 3z \left(1 - \left(z \frac{Q_0^2}{Q^2} \right)^2 \right) + \right. \\ & -\delta(1-z) \left[\frac{5}{2} - \frac{\pi^2}{3} - 3 \frac{Q_0^2}{Q^2} - \frac{3}{4} \frac{Q_0^4}{Q^4} \right] + \\ & \left. + \left[2 - 2 \left(\frac{1}{1-z} \right)_+ \right] \left(1 - z \frac{Q_0^2}{Q^2} \right) + \left(\frac{1/2}{1-z} \right)_+ \right\}. \quad (14) \end{aligned}$$

Here, the notation and normalization of [11] are used.

As emphasized above, since there are no infrared divergences, the calculation of the contribution caused by the production of a new real parton can be performed in the normal $D = 4$ space. Thus the F_2 coefficient functions of (13) and (14) coincide, in the limit $Q_0 \ll Q$, with those calculated in the ‘physical scheme’ [12]; but, as described in the next section, differ from those in the $\overline{\text{MS}}$ scheme.

A.2 Scheme dependence of F_2 coefficient functions

As mentioned in Section 5, after the $P^{\text{LO}} \otimes C^{\text{LO}}$ contribution was subtracted, there are no infrared divergences in the NLO coefficient functions. Thus, the F_2 coefficient functions of (13) and (14) coincide, in the limit $Q_0 \ll Q$, with those calculated in the ‘physical scheme’ [5, 12], but differ from those in the $\overline{\text{MS}}$ scheme. The differences are the ϵ/ϵ and ϵ^2/ϵ^2 terms arising from infinitely large distances in the $\overline{\text{MS}}$ scheme. To be more precise, these terms are of the form of $(\epsilon/\epsilon)P_{qa}(z)\ln(1-z)$ (with $a = q, g$) and $(\epsilon/\epsilon)2z(1-z)$ or $(\epsilon/\epsilon)(1-z)$ entering the C_{2g} and C_{2q} functions, respectively; and a term $(\epsilon^2/\epsilon^2)(\pi^2/3)\delta(1-z)$ in the C_{2q} function.

To calculate the power corrections, we must trace the origin of each term. We demonstrate this based on the formulae of the well-known paper Ref. [13], which works in the $\overline{\text{MS}}$ scheme. As an example, we consider the C_{2q} NLO coefficient function. Its ‘real’ contribution is due to the emission of an additional s -channel real gluon. It is given by Eq. (50) of [13], which is written in the $\gamma^*q \rightarrow qg$ centre-of-mass frame. We reproduce the relevant factor of this equation

$$F_2^{\text{real}} = \dots \left\{ 3z + z^\epsilon(1-z)^{-\epsilon} \int_0^1 dy (y(1-y))^{-\epsilon} \left[\left(\frac{1-z}{1-y} + \frac{1-y}{1-z} \right) (1-\epsilon) + \frac{2zy}{(1-z)(1-y)} \right] \right\}, \quad (15)$$

where the variable of angular integration had been changed to $y = \frac{1}{2}(1 + \cos\theta)$. The integral $\int_0^1 dy$ is actually an integration over the t -channel quark virtuality

$$k^2 = t = -\frac{Q^2}{z}(1-y) \quad (16)$$

from 0 to $-\hat{s} = -Q^2/z$. Note that $y = 0$ corresponds to $t = -\hat{s} = -Q^2/z$. Now, however, from this integral, we have to keep only the part from Q_0^2 up to $\hat{s}$. In other words, the upper limit $y = 1$ should be replaced by $y_0 = 1 - zQ_0^2/Q^2$.

After the subtraction of the $P_{qa}^{\text{LO}} \otimes C^{\text{LO}}$ contribution (to avoid double counting), the logarithmic, $1/(1-y)$, terms are cancelled exactly for all $|k^2| < \mu_F^2 = Q^2$ and $Q_0^2 < Q^2$. Therefore, there are no power corrections to the logarithmic part. The non-logarithmic terms result *either* from an integral of the form of

$$\int_{1-y=1-y_0}^{1-y=1} 2(1-y)d(1-y) = 1 - \left(z \frac{Q_0^2}{Q^2} \right)^2 \quad (17)$$

as in the second term in [...] on the r.h.s. of (15), or from

$$\int_{1-y_0}^1 dy = 1 - z \frac{Q_0^2}{Q^2} \quad (18)$$

as in the third term of (15), or, finally, from

$$\int_0^1 (1-y)dy = 1/2 \quad (19)$$

as in the last term in (15).

The final contribution arises from the $(1-y)/(1-z)$ term in (15). In terms of the cross section, it comes from the quark–gluon cut of quark self-energy diagram where the virtuality of each off-mass-shell quark is large ($k^2 = (1/z-1)Q^2$) and the value of $t = (p_q - p_g)^2$ reflects just the kinematics of the $q + \gamma \rightarrow q^* \rightarrow g + q$ subprocess (with a heavy virtual s -channel quark q^*), rather than the parton virtuality. Here, p_q and p_g denote the momenta of the incoming quark and the final gluon, respectively.

Besides this, in the case of C_{2q} function, the cutoff $|k^2| > Q_0^2$ should be included into the calculation of virtual loop contribution. This results in the Q_0^2/Q^2 correction to the $\delta(1-z)$ term.

A.3 Coefficient functions in the $\overline{\text{MS}}$ scheme

As discussed above, we note that the coefficient functions in the $\overline{\text{MS}}$ scheme are different to those in the physical scheme due to ϵ/ϵ and ϵ^2/ϵ^2 terms coming from the integration over infinitely large distances. Strictly speaking, these terms are not physical. Confinement will kill such contributions. On the other hand, these terms are not *power* corrections. Nevertheless, when working with $\overline{\text{MS}}$ PDFs (and $\overline{\text{MS}}$ evolution), we must keep such terms. These terms must be retained to compensate for the analogous ϵ/ϵ contributions in the definition of NLO PDFs used in the $\overline{\text{MS}}$ scheme.

Therefore, we calculate the expressions for the F_2 coefficient functions C_{2g} and C_{2q} with power corrections in the $\overline{\text{MS}}$ scheme⁶. Indeed, if we keep in (15) (and in the corresponding virtual contributions) all the ϵ/ϵ and ϵ^2/ϵ^2 terms, then we find the following NLO coefficient functions for F_2 in the $\overline{\text{MS}}$

⁶ The expressions for the longitudinal coefficient functions, C_{Lg} and C_{Lq} , are the same as before: namely (11) and (12).

scheme

$$C_{2g}^{\overline{\text{MS}}}(z) = T_{\text{R}} \left\{ [(1-z)^2 + z^2] \ln \frac{1-z}{z} + 2z(1-z) + [6z(1-z) - 1] \left(1 - z \frac{Q_0^2}{Q^2} \right) \right\}, \quad (20)$$

$$C_{2q}^{\overline{\text{MS}}}(z) = C_{\text{F}} \left\{ 2 \left(\frac{\ln(1-z)}{1-z} \right)_+ - (1+z) \ln(1-z) - \frac{1+z^2}{1-z} \ln z + 3z \left(1 - \left(z \frac{Q_0^2}{Q^2} \right)^2 \right) + \delta(1-z) \frac{3}{4} \frac{Q_0^4}{Q^4} - \delta(1-z) \left[\frac{\pi^2}{3} + \frac{9}{2} - 3 \frac{Q_0^2}{Q^2} \right] + \left[2 - 2 \left(\frac{1}{1-z} \right)_+ \right] \left(1 - z \frac{Q_0^2}{Q^2} \right) + (1-z) + \left(\frac{1/2}{1-z} \right)_+ \right\}. \quad (21)$$

Finally, for NLO correction to F_3 structure function, we get

$$C_{3q}^{\overline{\text{MS}}}(z) = C_{\text{F}} \left\{ 2 \left(\frac{\ln(1-z)}{1-z} \right)_+ - (1+z) \ln(1-z) - \frac{1+z^2}{1-z} \ln z + (2z-1) \left(1 - \left(z \frac{Q_0^2}{Q^2} \right)^2 \right) + \delta(1-z) \frac{3}{4} \frac{Q_0^4}{Q^4} - \delta(1-z) \left[\frac{\pi^2}{3} + \frac{9}{2} - 3 \frac{Q_0^2}{Q^2} \right] + \left[2 - 2 \left(\frac{1}{1-z} \right)_+ \right] \left(1 - z \frac{Q_0^2}{Q^2} \right) + (1-z) + \left(\frac{1/2}{1-z} \right)_+ \right\}. \quad (22)$$

These expressions reduce to the usual $\overline{\text{MS}}$ coefficient functions (given, for example, by Eq. (4.85) in [11]) in the absence of power corrections, that is, in limit $Q_0 \rightarrow 0$.

B. Q_0 -cut correction for the NLO Drell–Yan cross section

Here, we have used the normalization of the [13] paper where the LO cross section for Drell–Yan $q\bar{q} \rightarrow \gamma^*$ subprocess is written as

$$\frac{d\sigma_{q\bar{q}}(z, Q^2)}{dQ^2} = \delta(1-z). \quad (23)$$

Recall also that the incoming parton–parton energy square $s = Q^2(1-z)/z$. Accounting only for the contributions with the virtualities $|t|, |u| > Q_0^2$, we get the following corrections:

- I. All the ‘real’ NLO contributions caused by the $q\bar{q} \rightarrow g + \gamma^*$ or the $qg \rightarrow g + \gamma^*$ ($\bar{q}g \rightarrow g + \gamma^*$) subprocesses, that is all the dependent on z terms except of the terms proportional to $\delta(1 - z)$, should be multiplied by the $\Theta(Q^2(1 - z)/z - Q_0^2)$ function. This provides the possibility to satisfy the condition $|t|, |u| > Q_0^2$.
- II. The Q_0 -cut correction to cross section is denoted as $\Delta d\sigma(Q_0, z, Q^2)$; that is the final result reads

$$\frac{d\sigma(Q_0, z, Q^2)}{dQ^2} = \frac{d\sigma(Q_0 = 0, z, Q^2)}{dQ^2} + \frac{\Delta d\sigma(Q_0, z, Q^2)}{dQ^2}, \quad (24)$$

where the first term is the usual $d\sigma(z, Q^2)/dQ^2$ cross section given in [13] while the corrections are:

$$\frac{\Delta d\sigma_{q\bar{q}}(z, Q^2)}{dQ^2} = \frac{\alpha_s}{2\pi} \frac{8}{3} \frac{zQ_0^2}{Q^2}, \quad (25)$$

$$\frac{\Delta d\sigma_{qg}(z, Q^2)}{dQ^2} = -\frac{\alpha_s}{2\pi} \frac{1}{4} \left[\frac{(zQ_0^2)^2}{Q^4} + 4 \frac{z^2 Q_0^2}{Q^2} \right]. \quad (26)$$

REFERENCES

- [1] H. Kowalski, L.N. Lipatov, D.A. Ross, «The behaviour of the Green function for the BFKL pomeron with running coupling», *Eur. Phys. J. C* **76**, 23 (2016), [arXiv:1508.05744 \[hep-ph\]](#); H. Kowalski, L.N. Lipatov, D.A. Ross, O. Schulz, «Decoupling of the leading contribution in the discrete BFKL analysis of high-precision HERA data», *Eur. Phys. J. C* **77**, 777 (2017), [arXiv:1707.01460 \[hep-ph\]](#).
- [2] M. Aivazis, J.C. Collins, F. Olness, W.K. Tung, «Leptoproduction of heavy quarks. II. A unified QCD formulation of charged and neutral current processes from fixed-target to collider energies», *Phys. Rev. D* **50**, 3102 (1994); W.K. Tung, S. Kretzer, C. Schmidt, «Open heavy flavour production: conceptual framework and implementation issues», *J. Phys. G: Nucl. Part. Phys.* **28**, 983 (2002); S. Kretzer *et al.*, «CTEQ6 parton distributions with heavy quark mass effects», *Phys. Rev. D* **69**, 114005 (2004).
- [3] R.S. Thorne, R.G. Roberts, «A practical procedure for evolving heavy flavour structure functions», *Phys. Lett. B* **421**, 303 (1998); «Ordered analysis of heavy flavor production in deep-inelastic scattering», *Phys. Rev. D* **57**, 6871 (1998); R.S. Thorne, «Variable-flavor number scheme for next-to-next-to-leading order», *Phys. Rev. D* **73**, 054019 (2006); «Effect of changes of variable flavor number scheme on parton distribution functions and predicted cross sections», *Phys. Rev. D* **86**, 074017 (2012).
- [4] E.G. Oliviera, A.D. Martin, M.G. Ryskin, «Treatment of heavy quarks in QCD», *Eur. Phys. J. C* **73**, 2616 (2013), [arXiv:1307.3508 \[hep-ph\]](#).

- [5] E.G. Oliveira, A.D. Martin, M.G. Ryskin, «Physical factorisation scheme for PDFs for non-inclusive applications», *J. High Energy Phys.* **1311**, 156 (2013), [arXiv:1310.8289 \[hep-ph\]](#).
- [6] S. Forte, G. Altarelli, R.D. Ball, «Can we trust small x resummation?», *Nucl. Phys. Proc. Suppl.* **191**, 64 (2009), [arXiv:0901.1294 \[hep-ph\]](#).
- [7] R.D. Ball *et al.*, «Parton distributions with small- x resummation: evidence for BFKL dynamics in HERA data», *Eur. Phys. J. C* **78**, 321 (2018), [arXiv:1710.05935 \[hep-ph\]](#).
- [8] M. Bonvini, «Recent Developments in Small- x Resummation» *Acta Phys. Pol. B Proc. Suppl.* **12**, 873 (2019), [arXiv:1812.01958 \[hep-ph\]](#).
- [9] A.D. Martin, W.J. Stirling, R.S. Thorne, G. Watt, «Parton distributions for the LHC», *Eur. Phys. J. C* **63**, 189 (2009), [arXiv:0901.0002 \[hep-ph\]](#).
- [10] S.J. Brodsky, P. Hoyer, C. Peterson, N. Sakai, «The intrinsic charm of the proton», *Phys. Lett. B* **93**, 451 (1980).
- [11] R.K. Ellis, W.J. Stirling, B.R. Webber, in: «QCD and collider physics», *Cambridge University Press*, 1996.
- [12] E.G. Oliveira, A.D. Martin, M.G. Ryskin, «The treatment of the infrared region in perturbative QCD», *J. High Energy Phys.* **1302**, 060 (2013), [arXiv:1206.2223 \[hep-ph\]](#).
- [13] G. Altarelli, R.K. Ellis, G. Martinelli, «Large perturbative corrections to the Drell–Yan process in QCD», *Nucl. Phys. B* **157**, 461 (1979).