



WESTFÄLISCHE
WILHELMS-UNIVERSITÄT
MÜNSTER

MASTER'S THESIS

Evolution of Parton Distribution Functions

FLORIAN HERRMANN

INSTITUTE FOR THEORETICAL PHYSICS

SUPERVISOR: DR. KAROL KOVARIK

SUPERVISOR AND FIRST EXAMINER: PROF. DR. MICHAEL KLASSEN

SECOND EXAMINER: PD DR. JOCHEN HEITGER

OCTOBER 2015

Abstract

This thesis presents preliminary results of a modern C++ package for the evolution of parton distribution functions by solving the DGLAP equations in x -space or in Mellin-space. The x -space solution uses Lagrange interpolation for the x -grid as well as for the convolution, so that the integro-differential evolution equations reduce to sums of parton distribution functions and pre-computable evolution kernels. The actual integration can be performed using an arbitrary ordinary differential equation solver such as the fourth-order Runge–Kutta method, which is employed at present. The solution in Mellin-space is implemented as flexible as possible using functional representations and thus exhibits strong agreement with the theory. It is intended for calculations and fitting procedures in Mellin-space. Nevertheless, an inversion to x -space is also implemented, which especially allows for a comparison between both methods. The evolution methods are accessible through a streamlined user interface, allowing easy change of input parametrizations or evolution methods. Numerical results demonstrate high agreement with already existing tools like HOPPET [42] or QCD-PEGASUS [44]. The comparison between both methods verifies their equivalence, thus allowing them to be used for data fitting in numerous settings in the future.

Contents

1	Introduction	5
2	QCD Summary	7
2.1	General QCD Introduction	7
2.2	QCD Lagrangian	8
2.3	Running Coupling and Asymptotic Freedom	9
2.4	Parton Model	10
2.5	Scaling Violations	14
2.6	Factorization and DGLAP Equations	17
2.7	Relevance of PDFs and Their Evolution	18
3	DGLAP Equations and Related Topics	22
3.1	The Solution in x -Space	22
3.1.1	DGLAP Equation	22
3.1.2	Splitting Functions	23
3.1.3	Evolution Basis	24
3.1.4	Matching Conditions	26
3.1.5	Solution in x -Space	27
3.2	The Solution in Mellin-Space	27
3.2.1	General Mellin Transformation	27
3.2.2	Mellin-Space Solution of the DGLAP Equations	28
3.2.3	Splitting Functions / Anomalous Dimensions	32
3.2.4	Inverse Mellin Transformation	33
3.2.5	Solution in Mellin-Space	33
4	Implementation	35
4.1	The General Code Structure	35
4.2	The x -Space Evolution	35
4.2.1	Discretization and Interpolation of the Parton Momentum Densities	36
4.2.2	Implementation of the Interpolation	37
4.2.3	Discretization and Interpolation of the Convolution	39
4.2.4	Final DGLAP Equation and Evolution Operators	41
4.2.5	Splitting Functions	43
4.2.6	Strong Coupling	45
4.2.7	Matching Conditions	45

4.2.8	General Code Structure	45
4.3	The Mellin-Space Evolution	46
4.3.1	Initial Parametrization	47
4.3.2	Anomalous Dimensions	49
4.3.3	Strong Coupling	49
4.3.4	Matching Conditions	49
4.3.5	Inverse Mellin Transformation	50
4.3.6	Number of Evaluations	50
4.3.7	General Code Structure	51
4.4	Overview and Comparison of Both Methods	52
4.4.1	The x -Space Evolution	52
4.4.2	The Mellin-Space Evolution	53
4.4.3	Comparison	54
5	Numerical Results	55
5.1	Common Settings	55
5.2	The x -Space Evolution	56
5.2.1	Interpolation Order	56
5.2.2	Sampling Size in x	58
5.2.3	Integration Stepsize	60
5.2.4	Comparison to HOPPET	62
5.3	The Mellin-Space Evolution	64
5.3.1	Comparison to QCD-PEGASUS	64
5.3.2	Comparison to Benchmark Tables	67
5.4	Comparison of the Evolution Methods	67
6	Summary and Perspectives	72
Appendix A Additional Numerical Results		73
A.1	Interpolation Order	74
A.2	Comparison to HOPPET	75
A.3	Comparison to QCD-PEGASUS	76

Chapter 1

Introduction

Parton distribution functions (PDFs) are an essential part of collider physics and a necessary ingredient for high energy physics predictions. Moreover, although the proton/neutron PDFs have already been calculated with high-accuracy using data from for example the Drell-Yan process and deep inelastic scattering (DIS) (see e.g. [15]), in nuclear collisions they have to be corrected by cold nuclear matter effects. Furthermore, the gluon distribution is not sufficiently constrained for small Bjorken x . Therefore, due to missing data, the (nuclear) parton distribution functions ([n]PDFs) are still an ongoing and fundamental research topic.

The calculation and accuracy of the PDFs depends – besides the actual experimental data – on two ingredients: on the one hand, the theoretical calculations and predictions have to be carried out up to a required order in perturbation theory, and on the other hand, the PDFs have to be evolved from an initial scale up to the required energy scale. In practice, the PDFs are given at an initial scale through a parametrization, and with the help of the well-known Dokshitzer-Gribov-Lipatov-Altarelli-Parisi (DGLAP) equations these PDFs are evolved up to the scale of measurement. Then the parametrization parameters are adjusted by a fitting procedure.

This master's thesis concentrates on the evolution of the PDFs, i.e. on numerical solutions of the DGLAP equations. The code is intended for future use with nuclear PDFs in the nCTEQ collaboration (see e.g. [32]). In fact, there already exist multiple solutions for this challenge but most of them suffer from some shortcomings for the intended usage:

- The most prominent programs are old, i.e. they were designed in a different decade of computational power and in somewhat outdated programming languages like FORTRAN-77.
- External code is often insufficiently documented so that it is hard for users to gain any insight into the approximations and other sources of deviations in the implementation. Further, it is difficult to adapt third-party code to one's own interests.
- To the authors' best knowledge, there is at present no prominent program which is designed in modular form to be used easily with different input parametriza-

tions or evolution methods. In particular, the presented code will be able to solve the evolution equations in x -space and in Mellin-space. This way, both solution methods can be compared directly within one consistent framework.

At present, the evolution in x -space is able to evolve given parametrizations or PDFs at an initial scale up to a user specified scale using a fixed flavour number (FFN) scheme or a variable flavour number (VFN) scheme at leading order (LO). Also next-to-leading order (NLO) effects have been taken into account, so that no fundamental modifications are required if higher order contributions will be implemented. Likewise, the solution in Mellin-space is able to evolve the PDFs from an initial scale using the FFN or VFN scheme to another scale and afterwards to invert the results back to the x -space. Nevertheless, the evolution currently requires an analytic input parametrization unlike the x -space solution. Again NLO expansions have been considered. Both methods have been embedded in a streamlined user interface.

In future, the presented code is intended to be extended by an interpolation routine, a fitting procedure, and by the calculation of observables such as structure functions and cross sections, in order to be able to fit PDFs.

The structure of this thesis is as follows:

Chapter 2 A short summary of the basics of quantum chromodynamics (QCD), leading to factorization and the derivation of the DGLAP equations, will be presented.

Chapter 3 The theoretical questions concerning the DGLAP equations such as solution methods, evolution basis, flavour thresholds etc. will be discussed in view of the factual solutions.

Chapter 4 The actual implementations, their approximations, advantages and future optimizations are discussed.

Chapter 5 Numerical results of different evolution settings, evolution accuracies, comparisons with existing tools, and of comparisons between the implemented methods are analyzed.

Chapter 6 A summary of the main results and of possible extensions of the code project concludes this thesis.

Chapter 2

QCD Summary

Quantum chromodynamics (QCD) is an extensive field of research. The aim of this chapter is not a complete representation of the current state of research but a rough sketch along the path from the most elementary principles to the DGLAP equations. The description predominantly follows current textbooks like Halzen & Martin [30] (see also Martin [35]), Aitchison & Hey [2,3], Ellis et al. [16] and Muta [37]. Natural units $c = \hbar = 1$ are used throughout.

2.1 General QCD Introduction

One of the most important assumptions of QCD is the concept that hadronic matter consists of quarks. This concept was until the early 70th first only introduced as mathematical model to explain the multiple hadrons which have been observed. In particular Gell-Mann (Nobel Prize 1969) introduced the global $SU(3)_f$ flavour symmetry between three spin one-half constituents, the up, down and strange quark, to explain these observation (see [22]). In this model new particles such as the Ω -Baryons could be predicted and were experimentally proofed. However, at that time, the model was neglected to be of any physical relevance by large parts of the physical community as no free quarks could be observed.

Then, in 1969, it was the detection of Bjorken scaling (see [7]), i.e. the independence of the structure functions from the momentum transfer $Q^2 = -q^2$ at a given Bjorken x value, that lead Richard Feynman to the invention of the (naive) parton model of a nucleon consisting of free pointlike so called “partons” [18]. The measurements of the structure functions showed the partons to be of spin one-half, so that they were soon identified with Gell-Mann’s quarks. Nevertheless, still no free quarks could be observed. The parton model was later improved – and justified theoretically as well as experimentally – to become the current parton model described below.

Further problems occurs with the detection of baryons such as the Δ^{++} ($= uuu$) which seems to violate the Pauli exclusion principle. Therefore a new quantum number “colour” was invented (Greenberg) having three different values: red (r), green (g), and blue (b). However, only the detection of asymptotic freedom and confinement around 1973, for which Gross, Wilczek [28], and Politzer [39] were awarded Nobel Prize 2004, was able to overcome the challenge due to the missing

quark observations. Now, the color symmetry was interpreted in a non-Abelian gauge theory, the local $SU(3)_c$ colour (see [19]). As colour could not be observed, it was concluded that only colour neutral particles, i.e. colour singlets, could be stable. Thus, no free quarks could be observed; they are confined in the nucleon. The $SU(3)_c$ colour symmetry is thought to be exact.

2.2 QCD Lagrangian

Since there is no mixing in the standard model Lagrangian between the electroweak and strong interactions, it is possible to discuss the QCD and its Lagrangian separately. The QCD Lagrangian in general can be separated in the following parts:

$$\mathcal{L}_{\text{QCD}} = \mathcal{L}_{\text{classical}} + \mathcal{L}_{\text{gauge-fixing}} + \mathcal{L}_{\text{ghost}}. \quad (2.1)$$

Putting aside the discussion of the gauge-fixing and ghost parts, the “classical” contribution is given by

$$\mathcal{L}_{\text{classical}} = \bar{q}_a (i \not{D} - m)_{ab} q_b - \frac{1}{4} F_{\mu\nu}^A F_A^{\mu\nu}, \quad (2.2)$$

$$\text{where } F_{\mu\nu}^A = [\partial_\mu A_\nu^A - \partial_\nu A_\mu^A - g f^{ABC} A_\mu^B A_\nu^C] \quad (2.3)$$

and q_a is the triplet representation of the colour group. It is the third non-abelian term in the field strength tensor ($F_{\mu\nu}^A$) which especially distinguishes QCD from quantum electrodynamics (QED) as it gives rise to triplet and quartic gluon self interactions. This, in fact, is the reason for asymptotic freedom (see Section 2.3). The covariant derivative is defined for the quark triplet fields as

$$(D_\mu)_{ab} = \partial_\mu \delta_{ab} + ig(t^C A_\mu^C)_{ab}, \quad (2.4)$$

where a, b, c runs over all colours (r, g, b). The generators t are in the fundamental representation and the generators T belong to the adjoint representation of $SU(3)_c$, obeying the following algebra

$$[t^A, t^B] = if^{ABC} t^C, \quad [T^A, T^B] = if^{ABC} T^C, \quad (T^A)_{BC} = -if^{ABC}. \quad (2.5)$$

In general, the eight Gell-Mann matrices λ^A are chosen as representation for the generators $t^A = \lambda^A/2$. Using the normalization

$$\text{Tr } t^A t^B = T_F \delta^{AB}, \quad T_F = \frac{1}{2}, \quad (2.6)$$

the Casimir operator for the fundamental representation of $SU(N)$ is

$$t_{ab}^A t_{bc}^A = \delta_{ac} \frac{N^2 - 1}{2N} = \delta_{ac} C_F, \quad \text{i.e. for } SU(3): \quad C_F = \frac{4}{3}. \quad (2.7)$$

The Casimir operator of the adjoint representation is given by

$$\text{Tr } T^C T^D = f^{ABC} f^{ABD} = \delta^{CD} N = \delta^{CD} C_A, \quad \text{i.e. for } SU(3): \quad C_A = 3. \quad (2.8)$$

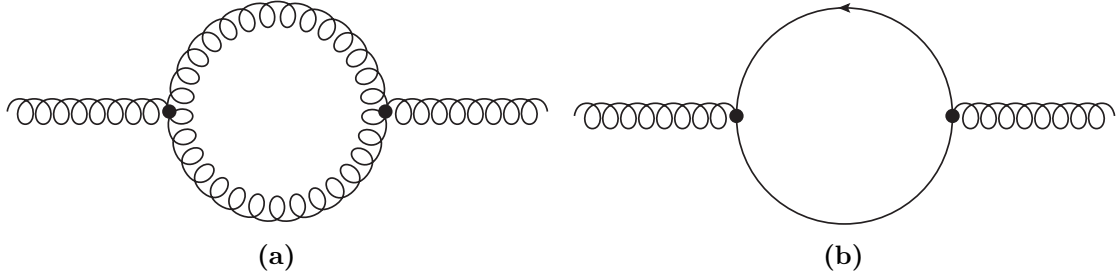


Fig. 2.1: Some contributions to the running of the strong coupling constant at one loop: a) gluon loop (anti-screening), b) quark loop (screening).

2.3 Running Coupling and Asymptotic Freedom

Due to the well-known running of the strong coupling constant, the coupling is large at low energy scales (large distances) and small at higher energies (short distances). This explains why on the one hand quarks and gluons are confined in the nucleon at large distances and on the other hand particles are asymptotic free and perturbation theory is valid at high energies.

This behavior of the coupling constant is in particular related to the self interactions of the exchange bosons, the gluons. It can be determined by calculating higher order perturbative corrections. In Fig. 2.1 some contributions to the one loop corrections are shown. In fact, the QCD coupling shows the behavior described above because the anti-screening of the gluons overcomes the screening due to the quarks.

Applying the renormalization group equation (RGE) yields up to three loops for $a_s = \alpha_s/(2\pi)$

$$Q^2 \frac{\partial a_s}{\partial Q^2} = \frac{\partial a_s}{\partial t} = \beta(a_s(t)), \quad (2.9)$$

$$\beta(a_s(t)) = -a_s(t)(\beta_0 a_s(t) + \beta_1 a_s^2(t) + \beta_2 a_s^3(t)), \quad (2.10)$$

where $t = \ln Q^2$ and the β_i are the contributions from the $i + 1$ loop correction. The β_i are given up to NNLO by

$$\beta_0 = \frac{11C_A - 4T_F n_f}{6} = \frac{11}{2} - \frac{1}{3}n_f, \quad (2.11)$$

$$\beta_1 = \frac{17C_A^2 - (10C_A - 6C_F)T_F n_f}{6} = \frac{51}{2} - \frac{19}{6}n_f, \quad (2.12)$$

$$\begin{aligned} \beta_2 &= \frac{2857C_A^3 + (108C_F^2 - 1230C_F C_A - 2830C_A^2)T_F n_f + (264C_F + 316C_A)T_F^2 n_f^2}{432} \\ &= \frac{2857}{16} - \frac{5033}{144}n_f + \frac{325}{432}n_f^2, \end{aligned} \quad (2.13)$$

where n_f is the number of effective (light) flavours. The first two β functions are independent of the renormalization scheme, and the last (β_3) is defined in the modified minimal subtraction ($\overline{\text{MS}}$) scheme.

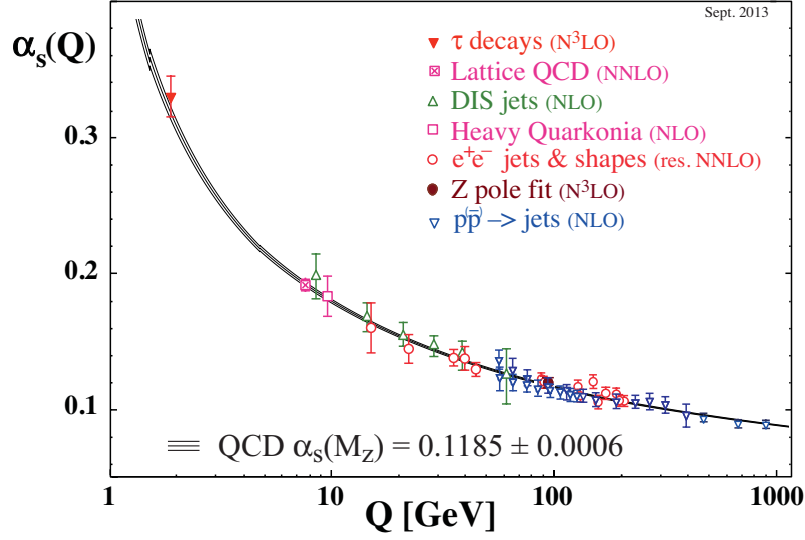


Fig. 2.2: Summary of measurements of α_s as a function of the energy scale Q . The respective degree of QCD perturbation theory used in the extraction of α_s is indicated in brackets (NLO: next-to-leading order; NNLO: next-to-next-to-leading order; res. NNLO: NNLO matched with resummed next-to-leading logs; N³LO: next-to-NNLO) (from [38, p. 134]).

At LO, Eq. (2.9) can be solved analytically. Thus, in the leading-log approximation

$$a_s(Q^2) = \frac{a_s(\mu^2)}{1 + a_s(\mu^2)\beta_0 t}, \quad \text{with } t = \log \frac{Q^2}{\mu_R^2}, \quad (2.14)$$

where μ_R is the renormalization scale.

The current world average value of the coupling according to the Particle Data Group [38, p. 122 ff.] is

$$\alpha_s(M_Z^2) = 0.1185 \pm 0.0006, \quad (2.15)$$

where M_Z is the mass of the Z boson. The behaviour of the coupling as a function of Q is depicted in Fig. 2.2.

2.4 Parton Model

As described above, one of the major advances in QCD was achieved through the introduction of the parton model, i.e. thinking of the proton as consisting of point-like and free constituents: quarks and gluons. The aim of this section is to shortly introduce this model by calculating the main steps in deep inelastic scattering (DIS: $Q^2 = -q^2 \gg M_p^2$, $(p+q)^2 \gg M_p^2$) for the example of electron-proton scattering as shown in Fig. 2.3.

In its most general form the cross section of this lepton-hadron process can be due to photon or Z-boson exchange. The differential cross section including the interference

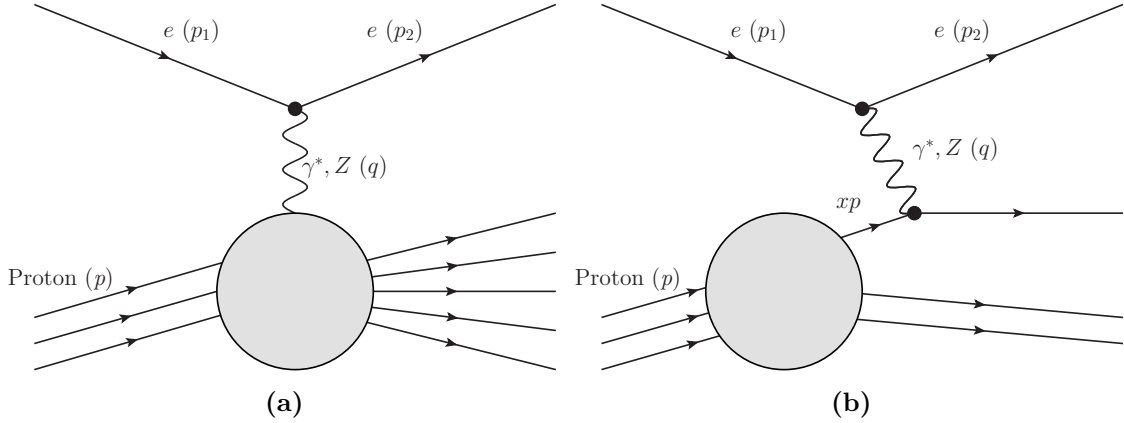


Fig. 2.3: Deep inelastic electron-proton scattering: a) scattering of an electron from a proton, b) scattering of an electron from a quark constituent of the proton in the parton model.

between those processes can be written as

$$\frac{d\sigma}{dE_2 d\cos\theta} = 2\pi \frac{\alpha^2}{q^4} \frac{E_2}{E_1} \sum_i \eta_i L_i^{\mu\nu} W_{\mu\nu}, \quad (2.16)$$

where i runs over $\gamma\gamma$, ZZ and γZ , η_i includes the different contributions from the couplings and propagators, and L_i and W_i are the leptonic and hadronic tensors, respectively.

The leptonic tensors are:

$$L_{\gamma\gamma}^{\mu\nu} = 2 (m^2 g^{\mu\nu} - g^{\mu\nu} p_1 \cdot p_2 + p_1^\nu p_2^\mu + p_1^\mu p_2^\nu), \quad (2.17)$$

$$L_{ZZ}^{\mu\nu} = -2g^{\mu\nu} (m^2 (c_A^2 - c_V^2) + (c_A^2 + c_V^2) p_1 \cdot p_2) + 2 (c_A^2 + c_V^2) (p_1^\nu p_2^\mu + p_1^\mu p_2^\nu) + 4ic_A c_V \epsilon^{\mu\nu\rho\sigma} p_{1\rho} p_{2\sigma}, \quad (2.18)$$

$$L_{\gamma Z}^{\mu\nu} = 2 [2c_V (g^{\mu\nu} (m^2 - p_1 \cdot p_2) + p_1^\nu p_2^\mu + p_1^\mu p_2^\nu) + 2ic_A \epsilon^{\mu\nu\rho\sigma} p_{1\rho} p_{2\sigma}], \quad (2.19)$$

where m is the lepton mass, p_1 and p_2 are the initial state and final state momenta of the lepton, respectively, and

$$c_V^f = c_L^f + c_R^f = t_3^f - 2e^f \sin^2 \theta_W, \quad c_A^f = c_L^f - c_R^f = t_3^f, \quad (2.20)$$

where f denotes the lepton, t_3^f is the third component of the weak isospin, e^f is the charge of the lepton, and θ_W is the Weinberg angle.

Starting from the most general Lorentz structure, the hadronic proton tensor can be parametrized as

$$W^{\mu\nu} = -g^{\mu\nu} W_1 + \frac{p^\mu p^\nu}{M_p^2} W_2 + \frac{ip^\alpha q^\beta \epsilon^{\alpha\beta\mu\nu}}{2M_p^2} W_3 + \frac{q^\mu q^\nu}{M_p^2} W_4 + \frac{(p^\nu q^\mu + p^\mu q^\nu)}{M_p^2} W_5, \quad (2.21)$$

where p is the proton momentum, M_p is the mass of the proton and q is the momentum of the exchange boson. This expression can be simplified considerably by

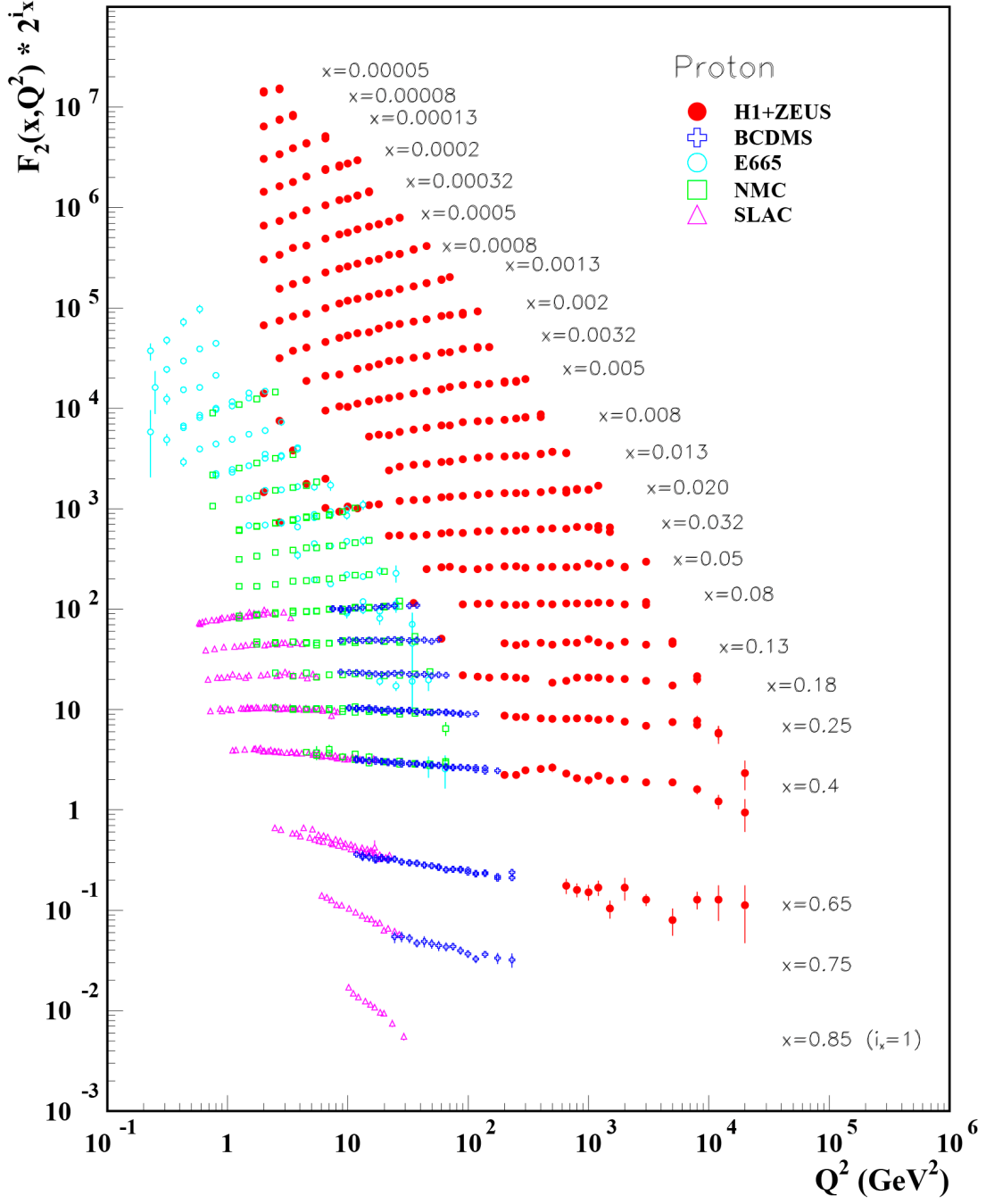


Fig. 2.4: The proton structure function F_2^p measured in electromagnetic scattering of electrons and positrons on protons (collider experiments H1 and ZEUS for $Q^2 \geq 2 \text{ GeV}^2$), in the kinematic domain of the HERA data, and for electrons (SLAC) and muons (BCDMS, E665, NMC) on a fixed target. The data are plotted as a function of Q^2 in bins of fixed x . For the purpose of plotting, F_2^p has been multiplied by 2^{i_x} , where i_x is the number of the x bin (figure from [38, p. 303], where additional information are available).

applying current conservation $q_\mu W^{\mu\nu} = 0$:

$$W^{\mu\nu} = \left(\frac{q^\mu q^\nu}{q^2} - g^{\mu\nu} \right) W_1 + \left(p^\mu - \frac{q^\mu p \cdot q}{q^2} \right) \left(p^\nu - \frac{q^\nu p \cdot q}{q^2} \right) \frac{W_2}{M_p^2} + \frac{i p^\alpha q^\beta \epsilon^{\alpha\beta\mu\nu}}{2M_p^2} W_3. \quad (2.22)$$

For simplicity, the calculations are reduced to the case of massless leptons (high energies) and weak current contributions are neglected. Using the Mandelstam variables

$$s = (p + p_1)^2, \quad u = (p - p_2)^2, \quad t = (p_1 - p_2)^2, \quad (2.23)$$

and the Bjorken variable

$$x = \frac{-q^2}{2p \cdot q} = \frac{Q^2}{2p \cdot q}, \quad (2.24)$$

the result is

$$\left(\frac{d\sigma}{dt du} \right)_{ep \rightarrow eX} = \frac{4\pi\alpha^2}{t^2 s^2} \frac{1}{s+u} [(s+u)^2 x M_p W_1 - u s \nu W_2] \quad (2.25)$$

$$= \frac{4\pi\alpha^2}{t^2 s^2} \frac{1}{s+u} [(s+u)^2 x F_1 - u s F_2], \quad (2.26)$$

where $\nu = E_1 - E_2$ and where the form factors are expressed through the structure functions $F_1 = M_p W_1$ and $F_2 = \nu W_2$. The experimental result of this process was the scaling of the structure functions, i.e. their independence of the scale $Q^2 = -q^2$ at a given value of x . This scaling is – despite logarithmic scaling violations (see below) – visible in Fig. 2.4, especially for $0.1 < x < 0.4$.

Richards Feynman's approach and subsequently further elaborated parton model is able to explain this fact. In this model, the proton consists of point-like quarks. At high energies, the proton/hadron can be described using the so called infinite momentum frame, so that the partons can be thought of as having a momentum collinearly to the proton and carrying a longitudinal momentum fraction x of it. Because of the time dilation and Lorentz contraction, the constituents can be handled as independent, and hence cross sections can be expressed through the incoherent sum of its partonic contributions. Therefore, parton distribution functions (PDFs) $f_q(x)$ can be introduced, reflecting the probability of finding a quark of type q with momentum fraction x of the proton in the proton. This way, the inelastic electron-proton cross section is explained by elastic electron-parton scattering:

$$\left(\frac{d\sigma}{dt du} \right)_{ep \rightarrow eX} = \sum_q \int dx f_q(x) \left(\frac{d\sigma}{dt du} \right)_{eq \rightarrow eq}, \quad (2.27)$$

where $q = u, \bar{u}, d, \dots$. The electron-quark elastic scattering cross section equals apart from a colour factor the well known electron-muon scattering, so that

$$\left(\frac{d\sigma}{dt du} \right)_{ep \rightarrow eX} = \sum_q \int dx f_q(x) x \frac{2\pi\alpha^2 e_q^2}{t^2 s^2} [(s+u)^2 - 2su] \delta(t + x(s+u)), \quad (2.28)$$

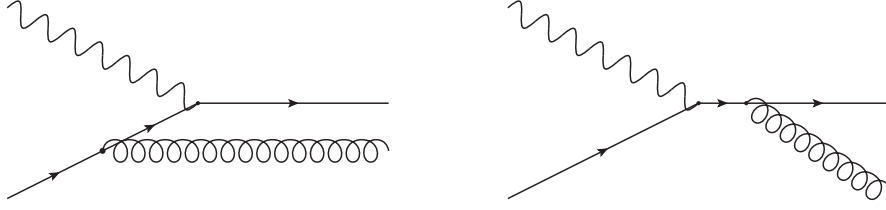


Fig. 2.5: $\mathcal{O}(\alpha_s)$ corrections to the naive parton model through real gluon emission.

where the delta distribution fixes the parton momentum fraction of the proton to the Bjorken variable x . Comparing both solutions, Eqs. (2.26) and (2.28), yields the Callan-Gross relation

$$2xF_1(x) = F_2(x) = \sum_q e_q^2 x f_q(x), \quad (2.29)$$

and explains Bjorken scaling, since the structure functions only depend on the Bjorken scaling variable x .

Instead of the structure functions F_1 and F_2 , alternate structure functions for the absorption of a transversal and longitudinal polarized virtual photon can be introduced. For large Q^2 ,

$$F_T \rightarrow F_1, \quad F_L \rightarrow F_2/2x - F_1. \quad (2.30)$$

The measured vanishing of F_L proofs the quarks to be of spin one-half, because spin one-half quarks cannot absorb longitudinal polarized vector bosons. This is of course in agreement with the spin one-half calculation above; especially with the Callan-Gross relation $2xF_1(x) = F_2(x)$ (see Eq. 2.29).

2.5 Scaling Violations

In contrast to the “naive” parton model above, scaling violations, proportional to $\alpha_s \ln(Q^2)$, can be observed. These logarithmic effects can already be seen in Fig. 2.4. For example these $\mathcal{O}(\alpha_s)$ corrections are due to gluon emissions of the quarks struck by the photon (see Fig. 2.5). A straightforward calculation of this cross section $\gamma^* + q \rightarrow q + g$ for the real gluon emission yields the squared matrix element

$$|\mathcal{M}|^2 = 32\pi^2(e_q^2\alpha_s)C_F \left(\frac{-\hat{t}}{\hat{s}} - \frac{\hat{s}}{\hat{t}} + \frac{2\hat{u}Q^2}{\hat{s}\hat{t}} \right), \quad (2.31)$$

where hats ($\hat{}$) indicate partonic variables, Mandelstam variables are used, and running coupling effects have to be considered. Since the virtual photon still scatters from a spin one-half quark, the Callan-Gross relation (see Eq. 2.29) is also valid for the structure function of each parton separately. That is, in the DIS limit

$$\frac{\hat{F}_2}{x} = 2\hat{F}_1 = \hat{\sigma}_T \frac{\hat{s}}{4\pi^2\alpha} \quad (2.32)$$

$$= \frac{\hat{s}}{4\pi^2\alpha} \int_{-1}^1 d\cos\theta C_F \frac{\pi e_q^2 \alpha_s}{\hat{s}} \left(\frac{-\hat{t}}{\hat{s}} - \frac{\hat{s}}{\hat{t}} + \frac{2\hat{u}Q^2}{\hat{s}\hat{t}} \right). \quad (2.33)$$

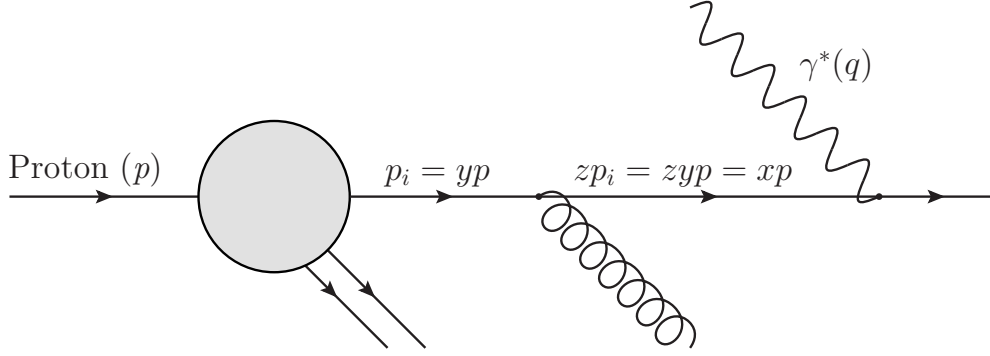


Fig. 2.6: Momentum fractions for the real gluon emission parton process.

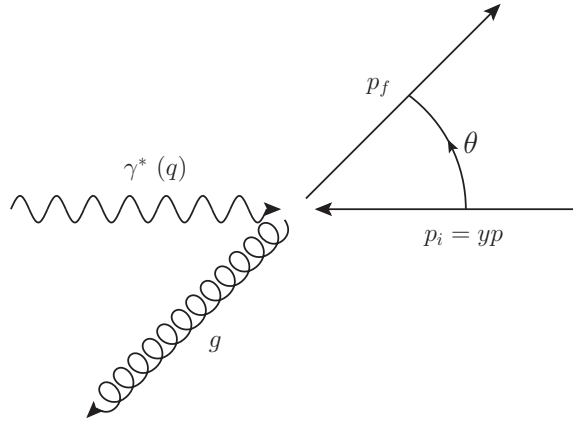


Fig. 2.7: Kinematics for the real gluon emission of Fig. 2.6.

Introducing the momentum fractions: $p \rightarrow p_i = yp \rightarrow p_i' = zp_i = zyp = xp$ as described in Fig. 2.6, a partonic version of x and can be expressed through

$$z = \frac{Q^2}{2p_i \cdot q} = \frac{Q^2}{(p_i + q)^2 - q^2} = \frac{Q^2}{\hat{s} + Q^2}. \quad (2.34)$$

Using $c = \cos \theta$ and the kinematics described in Fig. 2.7, the partonic Mandelstam variables are

$$\hat{s} = (q + p_i)^2 = (p_2 + g)^2 = \frac{Q^2(1 - z)}{z}, \quad (2.35)$$

$$\hat{t} = (q - p_f)^2 = (g - p_i)^2 = \frac{-Q^2(1 - c)}{2z}, \quad (2.36)$$

$$\hat{u} = (p_i - p_f)^2 = (q - g)^2 = \frac{-Q^2(1 + c)}{2z}. \quad (2.37)$$

Thus,

$$2\hat{F}_1 = C_F \frac{e_q^2 \alpha_s}{4\pi} \int_{-1}^1 dc \left(\frac{1-c}{2(1-z)} + \frac{2(1-z)}{1-c} + \frac{2z(1+c)}{(1-c)(1-z)} \right). \quad (2.38)$$

In this last expression are multiple sources of infra-red divergences. First, the integral diverge for $c \rightarrow 1$, in which case the gluon is emitted collinear to the initial quark p_i . It is exactly this long-range contribution which will be later factorized in the PDFs. In addition, the integral is also divergent for $z \rightarrow 1$, which occurs if the gluon is soft, i.e. has vanishing momentum, or is collinear to the outgoing quark. An accurate handling of this infra-red divergences would imply regularization techniques like dimensional regularization. Instead, in this short summary, the $c \rightarrow 1$ collinear divergence will be regularized using a lower cut-off μ in transverse momentum

$$p_T^2 = \frac{1}{4} \frac{Q^2(1-z)}{z} (1-c^2). \quad (2.39)$$

In the collinear limit, $c \approx 1$, therefore

$$2\hat{F}_1 \approx e_q^2 \int_{\mu^2}^{\hat{s}^{1/4}} dp_T^2 \left[\frac{1}{p_T^2} \left(\frac{\alpha_s}{2\pi} C_F \frac{1+z^2}{1-z} \right) + \mathcal{O}\left(\frac{p_T^2}{Q^2}\right) \right], \quad (2.40)$$

where the lower p_T scale μ regularizes the collinear singularity and gives rise to logarithmic corrections (dimensional regularization yields a pole in $1/\epsilon$). Thus,

$$2\hat{F}_1 = P'_{qq}(z) \log \frac{Q^2}{\mu^2} + C(z), \quad (2.41)$$

with the incomplete and unregularized splitting function

$$P'_{qq}(z) = C_F \left(\frac{1+z^2}{1-z} \right). \quad (2.42)$$

The remainder term $C(z)$ contains no initial state collinear singularities. The additional soft infra-red divergences for $z \rightarrow 1$ will be canceled by the virtual diagrams of Fig. 2.8. Instead of explicitly calculating the contributions from these diagrams, an easier way to consider them is proposed by Altarelli and Parisi [5]. The virtual diagrams contributes only for $z = 1$ and are proportional to $\delta(1-z)$. Moreover, the splitting function, regularized by the virtual contributions, can be expressed by introducing the plus distribution:

$$\int_0^1 dz \frac{f(z)}{(1-z)_+} = \int_0^1 dz \frac{f(z) - f(1)}{(1-z)}, \text{ with } (1-z)_+ = (1-z) \text{ for } z < 1, \quad (2.43)$$

and requiring that $f(z)$ is sufficiently regular at the end points. Furthermore, due to valence quark number conservation the integral $\int_0^1 dz P_{qq}(z)$ has to vanish, which can be used to fix the proportionality constant of the delta distribution. Thus,

$$P_{qq}(z) = C_F \frac{1+z^2}{(1-z)_+} + 2\delta(1-z). \quad (2.44)$$

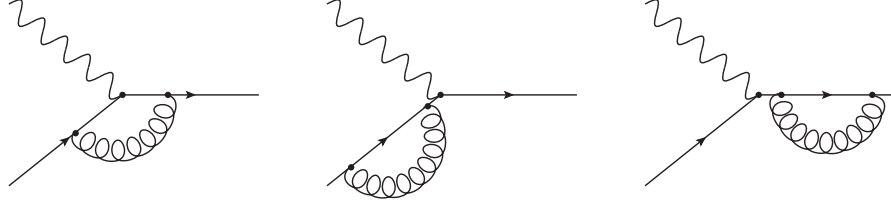


Fig. 2.8: Virtual diagrams contributing to the order $\mathcal{O}(\alpha_s)$ through interference with the leading order parton diagram.

Finally, the hadronic structure function F_1 is given as the sum of all its partonic contributions

$$\frac{F_2}{x} = 2F_1 = \sum_q \int_0^1 dy f_q(y) \int_0^1 dz 2\hat{F}_1^q(Q^2, \mu^2, z) \delta(x - yz) \quad (2.45)$$

$$= \sum_q \frac{e_q^2 \alpha_s}{2\pi} \int_x^1 \frac{dy}{y} f_q(y) \left[P_{qq}\left(\frac{x}{y}\right) \log \frac{Q^2}{\mu^2} + C\left(\frac{x}{y}\right) \right]. \quad (2.46)$$

Including the $\mathcal{O}(\alpha)$ contributions, the total structure function up to $\mathcal{O}(\alpha\alpha_s)$ is

$$\frac{F_2}{x} = \sum_q e_q^2 \int_x^1 \frac{dy}{y} f_q(y) \left(\delta\left(1 - \frac{x}{y}\right) + \frac{\alpha_s}{2\pi} \left[P_{qq}\left(\frac{x}{y}\right) \log \frac{Q^2}{\mu^2} + C\left(\frac{x}{y}\right) \right] \right), \quad (2.47)$$

where the term proportional to $\log Q^2$ violates Bjorken scaling.

2.6 Factorization and DGLAP Equations

The result of the last section still includes the collinear singularity for $\mu \rightarrow 0$. One way to handle this long-range contribution is to define – similar to the renormalization of ultra-violet divergences – dressed parton densities $f_q(x, \mu_F^2)$ using the factorization scale μ_F

$$f_q(x, \mu_F^2) = f_q(x) + \frac{\alpha_s}{2\pi} \int_x^1 \frac{dy}{y} f_q(y) P_{qq}\left(\frac{x}{y}\right) \left(\log \frac{\mu_F^2}{\mu^2} + C\left(\frac{x}{y}\right) \right), \quad (2.48)$$

so that

$$\frac{F_2}{x} = \sum_q e_q^2 \int_x^1 \frac{dy}{y} f_q(y, \mu_F^2) \left(\delta\left(1 - \frac{x}{y}\right) + \frac{\alpha_s}{2\pi} \left[P_{qq}\left(\frac{x}{y}\right) \log \frac{Q^2}{\mu_F^2} \right] \right). \quad (2.49)$$

The treatment of the term $C(x/y)$ depends on the choice of the factorization scheme. The collinear singularity is factorized in the well defined “running” parton density $f_q(x, \mu_F^2)$, which therefore includes non-perturbative long-distance physics. This density has to be obtained experimentally, but its evolution can be calculated since

the structure functions should not depend on an arbitrarily chosen factorization scale. Using renormalization group techniques yields

$$\frac{\partial f_q(x, \mu_F^2)}{\partial \log \mu_F^2} = \frac{\alpha_s}{2\pi} \int_x^1 \frac{dy}{y} f_q(y, \mu_F^2) P_{qq}\left(\frac{x}{y}\right), \quad (2.50)$$

and hence allows to calculate the μ_F^2 dependence of the PDFs. However, until now gluon splitting into two quarks P_{qg} has been neglected, which couples the quark distribution to the gluon distribution. Including this contribution, the well-known Dokshitzer-Gribov-Lipatov-Altarelli-Parisi (DGLAP) equation [5, 14, 27, 34] for a quark density reads

$$\frac{\partial f_q(x, \mu_F^2)}{\partial \log \mu_F^2} = \frac{\alpha_s}{2\pi} \int_x^1 \frac{dy}{y} f_q(y, \mu_F^2) P_{qq}\left(\frac{x}{y}\right) + f_g(y, \mu_F^2) P_{qg}\left(\frac{x}{y}\right). \quad (2.51)$$

Then, using equal factorization and renormalization scales $\mu_F^2 = \mu_R^2 = Q^2$, and denoting the convolution by $\otimes$, the evolution equation for PDFs is:

$$\frac{\partial f_q(x, Q^2)}{\partial \log Q^2} = \frac{\alpha_s(Q^2)}{2\pi} \left[[f_q \otimes P_{qq}](x, Q^2) + [f_g \otimes P_{qg}](x, Q^2) \right], \quad (2.52)$$

and considering the triple gluon vertex, the equation for the gluon density is:

$$\frac{\partial f_g(x, Q^2)}{\partial \log Q^2} = \frac{\alpha_s(Q^2)}{2\pi} \left[[f_g \otimes P_{gg}](x, Q^2) + \sum_q [f_q \otimes P_{gq}](x, Q^2) \right]. \quad (2.53)$$

This derivation is, of course, only a sketch. A more formal and accurate derivation would include operator product expansion techniques etc. [23, 29, 37], which is beyond the scope of this thesis.

2.7 Relevance of PDFs and Their Evolution

The above sketch of the derivation of the DGLAP equations has similarities in other QCD calculations. First, the factorization, used above to handle initial state collinear singularities, can be used to calculate hard parton processes. This way, they can be split into a product of the unperturbative PDFs and perturbative short-distance cross sections,

$$\sigma_{pp \rightarrow X} = \sum_{i,j} \int dx_1 dx_2 f_i(x_1, \mu_F^2) f_j(x_2, \mu_F^2) \hat{\sigma}_{ij}(p_1, p_2, \alpha_s(\mu_R^2 = \mu_F^2), Q^2/\mu_F^2), \quad (2.54)$$

which is illustrated in Fig. 2.9. Despite this result, the PDFs have still to be determined and cannot be predicted *a priori* at any scale Q_0 . But once they are determined, they are universal and so perturbative QCD can be predictive.

The procedure for determining PDFs is, in principle, to use existing experimental data to fit PDFs, and to use these PDFs afterwards as input for predictive calculations. In concrete:

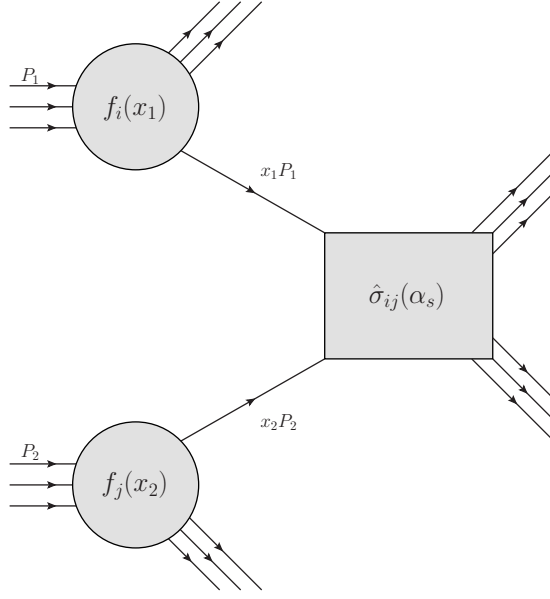


Fig. 2.9: Factorization in the parton model hard scattering process.

- First a parametrization for the PDFs at an initial scale is proposed. This parametrization can be inspired using spectator counting rules and Regge behaviour to be proportional to $x^\alpha(1-x)^\beta$, but details are beyond the scope of this work. This central part can be extended for example by polynomials. Typical parametrizations are [41, 44]

$$xf_q(x, Q_0^2) = N_q x^{p_{q,1}} (1-x)^{p_{q,2}} e^{p_{q,3}} (1 + p_{q,4}x)^{p_{q,5}} \quad \text{or} \quad (2.55)$$

$$xf_q(x, Q_0^2) = N_q p_{q,1} x^{p_{q,2}} (1-x)^{p_{q,3}} [1 + p_{q,5}x^{p_{q,4}} + p_{q,6}x], \quad (2.56)$$

where $q = g, u, \bar{u}, \dots$, N_q are the normalizations, and $p_{q,j}$ are fitting parameters.

- Then the PDFs will be evolved using the DGLAP equations up to the experimental scale and used to calculate observables such as structure functions and cross sections. Data are typically obtained from deep inelastic scattering (DIS) and the Drell-Yan process.
- Using a fitting procedure the parameters of the parametrization will be adjusted to the experimental data.

Finally, the PDFs can be evolved to other scales and hence used for predictions. Furthermore, besides the experimental data so called “sum rules” put theoretical restrictions on the PDFs. For instance, for the proton there are the valence quark

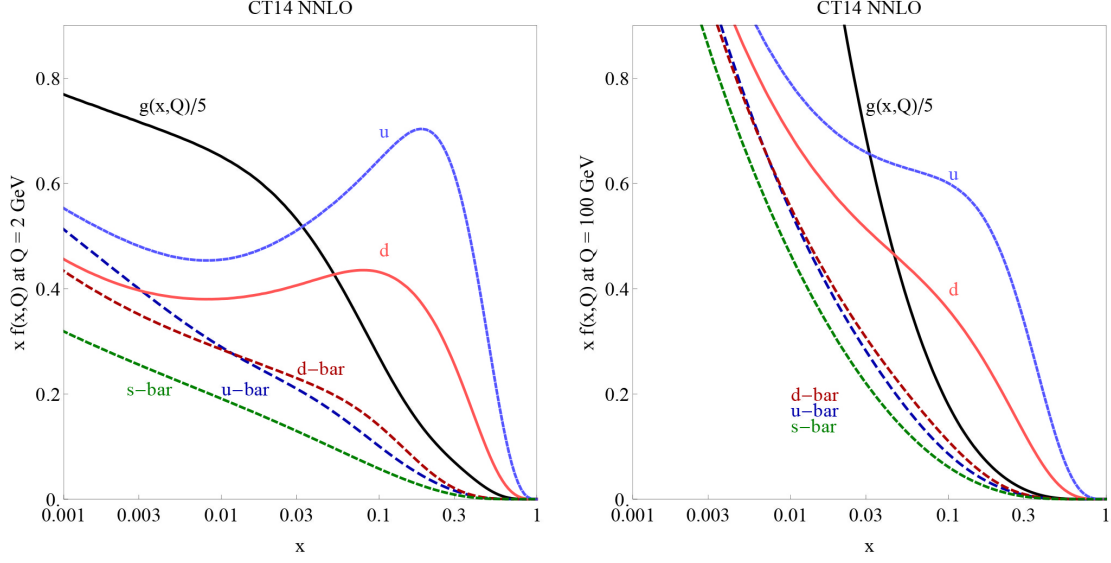


Fig. 2.10: The recent CTEQ CT14 parton distribution functions at $Q = 2 \text{ GeV}$ and $Q = 100 \text{ GeV}$ for u , $\bar{u}$, d , $\bar{d}$, $s = \bar{s}$ and g (from [15]).

numbers:

$$\int_0^1 dx u_v(x) = \int_0^1 dx (f_u(x) - f_{\bar{u}}(x)) = 2, \quad \text{i.e. two } u\text{-valence quarks}, \quad (2.57)$$

$$\int_0^1 dx d_v(x) = \int_0^1 dx (f_d(x) - f_{\bar{d}}(x)) = 1, \quad \text{i.e. one } d\text{-valence quarks}, \quad (2.58)$$

$$\int_0^1 dx [s(x) - \bar{s}(x)] = 0, \quad \text{i.e. no } s\text{-valence quark}. \quad (2.59)$$

And, of course, momentum has to be conserved

$$\int_0^1 dx x \left[f_g(x) + \sum_q (f_q(x) + f_{\bar{q}}(x)) \right] = 1, \quad (2.60)$$

which experimentally requires that the gluon carries – depending on the scale – about 50 % of the proton’s momentum.

In Fig. 2.10, a recent proton PDF plot from the CTEQ collaboration is shown (as usual, parton momenta xf_q (xPDFs) are plotted). The valence distributions of the up and down quark peaks at $x \approx 0.2$ with the zero limit for $x \rightarrow 1$ and $x \rightarrow 0$, whereas the sea-quarks increase strongly for $x \rightarrow 0$ and the gluon distribution dominates for $x < 0.2$. Of course, this picture is not complete and especially for small $x < 10^{-3}$ the PDFs are only weakly restricted by the current data. At these small x -scales, saturation effects are present and other mathematical descriptions beyond DGLAP are required such as the BFKL equation. A depiction of the kinematical coverage of current measurements is given in Fig. 2.11.

This whole fitting procedure, of course, requires a great deal of PDF evaluations and evolutions. Thus, an exact and fast solver of the DGLAP evolution equations is of fundamental importance.

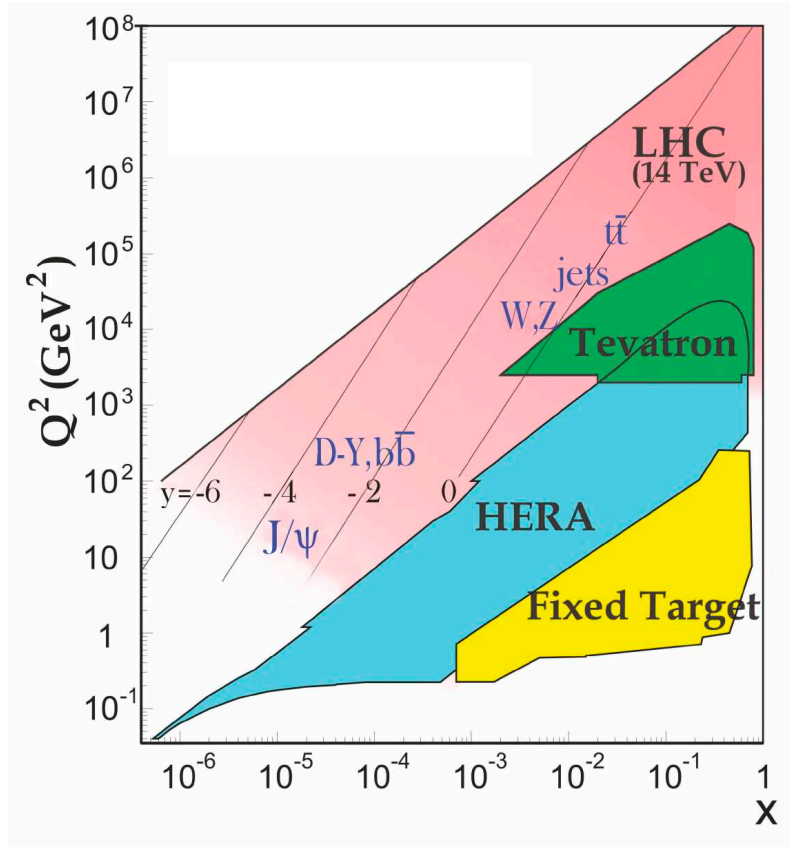


Fig. 2.11: Kinematic domains in x and Q^2 probed by fixed-target and collider experiments. Some of the final states accessible at the LHC are indicated in the appropriate regions, where y is the rapidity (from [38, p. 299]).

Chapter 3

DGLAP Equations and Related Topics

In this chapter, further properties of the DGLAP equations in view of their numerical solution will be discussed. First, the solution in x -space is dealt with, and afterwards its modifications in Mellin-space are presented if differences are significant. A general overview can be found in [16], especially on x -space solutions see [42] and for the solutions in Mellin-space see [21, 44].

3.1 The Solution in x -Space

There is no analytic solution of the DGLAP equations in x -space, but the complete equations can be simplified by basis transformations. In addition, challenges due to heavy flavour thresholds will be discussed.

3.1.1 DGLAP Equation

As described in Chapter 2, the DGLAP equations are renormalization group equations for the parton densities. They are coupled integro-differential equation with the splitting functions as kernel elements. These splitting functions $P_{ab}(z)$ describe, for $z < 1$, the probability of finding a quark of type a from a quark of type b carrying a fraction z of the longitudinal momentum of b . From now on, the PDFs, resolved at a scale Q^2 , are denoted by $q_i(x, Q^2)$ for quarks and $g(x, Q^2)$ for gluons, where x is the fraction of the longitudinal momentum of the hadron and $i = u, \bar{u}, d, \dots$.

Introducing the new scale variable $t = \ln Q^2$, using the convention for $a_s = \alpha_s/2\pi$ (see Section 2.3) and considering the most general form of the DGLAP equations, Eqs. (2.52) and (2.53) can be written as

$$\frac{\partial}{\partial t} \begin{pmatrix} q_i(x, t) \\ g(x, t) \end{pmatrix} = a_s(t) \sum_{j=1}^{2n_f} \int_x^1 \frac{dz}{z} \left[\begin{pmatrix} P_{q_i q_j} \left(\frac{x}{z}, a_s(t) \right) & P_{q_i g} \left(\frac{x}{z}, a_s(t) \right) \\ P_{g q_j} \left(\frac{x}{z}, a_s(t) \right) & P_{g g} \left(\frac{x}{z}, a_s(t) \right) \end{pmatrix} \begin{pmatrix} q_j(z, t) \\ g(z, t) \end{pmatrix} \right], \quad (3.1)$$

where the i and j runs over all flavours: $u, \bar{u}, d, \bar{d}, \dots$

3.1.2 Splitting Functions

Moreover, the splitting functions themselves have a perturbative expansion in the strong coupling constant. The above calculation of P_{qq} is only the first order term of

$$P_{ij}(x, t) = P_{ij}^{(0)}(x) + a_s^1(t)P_{ij}^{(1)}(x) + a_s^2(t)P_{ij}^{(2)}. \quad (3.2)$$

In the last chapter, one splitting function was already calculated explicitly: $P_{q_i q_i}$. However, the general formulation of the DGLAP equation (Eq. 3.1) includes multiple splitting functions. Due to charge conjugation invariance and flavour symmetry their number can be reduced significantly:

$$P_{q_i q_j} = P_{\bar{q}_i \bar{q}_j} \quad (3.3)$$

$$P_{q_i \bar{q}_j} = P_{\bar{q}_i q_j} \quad (3.4)$$

$$P_{q_i g} = P_{\bar{q}_i g} = P_{qg} \quad (3.5)$$

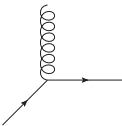
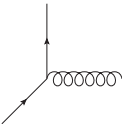
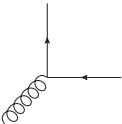
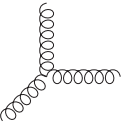
$$P_{gq_i} = P_{g\bar{q}_i} = P_{gg}. \quad (3.6)$$

Furthermore, at LO, $P_{q_i q_j}^{(0)}$ is proportional to δ_{ij} , and the probability of a gluon coming from a quark with momentum fraction z should be equal to the probability of quark with the momentum fraction $(1 - z)$ coming from the same quark:

$$P_{qq}(z) = P_{qq}(1 - z). \quad (3.7)$$

Thus, at LO, four splitting functions remain [5], which are listed in Tab. 3.1. The LO splitting functions are independent of the factorization scheme. Higher order contributions to the splitting functions have already been calculated (e.g. NLO: [13, 20]; NNLO: [36, 45]).

Table 3.1: Leading order splitting functions and associated diagrams.

Diagram	Splitting Function
	$P_{qq}^{(0)}(x) = C_F \left[\frac{1+x^2}{(1-x)_+} + \frac{3}{2}\delta(1-x) \right]$
	$P_{gq}^{(0)}(x) = C_F \left[\frac{1+(1-x)^2}{x} \right]$
	$P_{qg}^{(0)}(x) = T_F [x^2 + (1-x)^2]$
	$P_{gg}^{(0)}(x) = 2C_A \left[\frac{1-x}{x} + x(1-x) + \frac{x}{1-x_+} \right] + \delta(1-x) \frac{11C_A - 4T_F n_f}{6}$

3.1.3 Evolution Basis

The physical basis of the PDFs used thus far can be modified to a more appropriate basis, which decouples non-singlet distributions from the gluon solution [9, 21]. This means that a new basis is introduced, which will also depend on the number of active flavours n_f :

$$\text{Singlet: } \Sigma = \sum_{i=1}^{n_f} (q_i + \bar{q}_i), \quad (3.8)$$

$$\text{Non-singlet: } q_{ij}^{\pm} = (q_i \pm \bar{q}_i) - (q_j \pm \bar{q}_j), \quad (3.9)$$

$$\text{Valence (non-singlet): } q^V = \sum_{i=1}^{n_f} (q_i - \bar{q}_i), \quad (3.10)$$

$$\text{Gluon: } g = g. \quad (3.11)$$

In general there are $3 + 2n_f(n_f - 1)$ possible nontrivial distribution functions but not all of them are independent, so that the number can be reduced to $3 + 2(n_f - 1)$ which is equivalent to the size of the physical basis. Different basis schemes are possible and in use. In this thesis the following basis is chosen:

$$q_{i1}^{\pm} = (q_i \pm \bar{q}_i) - (q_1 \pm \bar{q}_1) \quad \text{for } i > 1, \quad (3.12)$$

since $q_{21}^{\pm} - q_{31}^{\pm} = q_{32}^{\pm} = -q_{23}^{\pm}$ etc. The numbering of the flavours follows an increasing order: $u (= 1)$, $d (= 2)$, $s (= 3)$, $c (= 4)$, $b (= 5)$, $t (= 6)$.

Basis Transformation

As in the calculation of observables the PDFs are required in the physical basis the PDFs have to be transformed from one basis to the other. The transformation procedure adopted in this thesis is:

- First $u_p = (u + \bar{u})$ and $u_m = (u - \bar{u})$ are calculated using the singlet (valence) basis and subtracting the non-singlet q_{i1}^+ (q_{i1}^-) contributions. Thus, for u_p :

$$\begin{aligned} u_p &= \frac{1}{n_f} (\Sigma - q_{21}^+ - q_{31}^+ \cdots) \\ &= \frac{1}{n_f} ((u + \bar{u}) + (d + \bar{d}) + (s + \bar{s}) + \cdots \\ &\quad - [(d + \bar{d}) - (u + \bar{u})] - [(s + \bar{s}) - (u + \bar{u})] - \cdots) \\ &= u_p. \end{aligned} \quad (3.13)$$

- Then the other quark combinations $q_{i,p}$ and $q_{i,m}$ are calculated by adding u_p and u_m to the non-singlet q_{i1}^+ (q_{i1}^-) basis. For example d_p :

$$d_p = q_{21}^+ + u_p = [(d + \bar{d}) - (u + \bar{u})] + (u + \bar{u}) = (d + \bar{d}). \quad (3.14)$$

- Finally, the physical quark distributions read

$$s = \frac{1}{2}(s_p + s_m), \quad \bar{s} = \frac{1}{2}(s_p - s_m). \quad (3.15)$$

Splitting Functions in the Evolution Basis

The basis change also alters the DGLAP equations and introduces new combinations of splitting functions [9, 21]. Separating P_{qq} and $P_{q\bar{q}}$ in valence and singlet contributions leads to

$$P_{q_i q_k} = P_{\bar{q}_i \bar{q}_k} = \delta_{ik} P_{qq}^V + P_{qq}^S, \quad (3.16)$$

$$P_{q_i \bar{q}_k} = P_{\bar{q}_i q_k} = \delta_{ik} P_{q\bar{q}}^V + P_{q\bar{q}}^S. \quad (3.17)$$

Charge conjugation and flavour symmetry up to two loops yields

$$P_{qq}^V = P_{\bar{q}\bar{q}}^V, \quad P_{q\bar{q}}^V = P_{\bar{q}q}^V, \quad (3.18)$$

$$P_{qq}^S = P_{\bar{q}\bar{q}}^S = P_{q\bar{q}}^S = P_{\bar{q}q}^S, \quad (3.19)$$

$$P_{q_i g} = P_{\bar{q}_i g} = P_{gq}, \quad (3.20)$$

$$P_{gq_i} = P_{g\bar{q}_i} = P_{gq}. \quad (3.21)$$

Therefore, for the non-singlet evolution holds

$$P^\pm = P_{qq}^V \pm P_{q\bar{q}}^V \quad \text{and} \quad (3.22)$$

$$P^V = P_{qq}^V - P_{q\bar{q}}^V + n_f(P_{qq}^S - P_{q\bar{q}}^S), \quad (3.23)$$

but while P_{qq}^V starts at first order and $P_{q\bar{q}}^V$ and $P_{q\bar{q}}^S$ unequal zero at second order, $(P_{qq}^S - P_{q\bar{q}}^S)$ vanishes up to second order. Thus,

$$P^\pm = P_{qq}^{(0)} + a_s P^{(1)\pm} + a_s^2 P^{(2)\pm}, \quad (3.24)$$

$$P^V = P_{qq}^{(0)} + a_s P^{(1)-} + a_s^2 P^{(2)V}. \quad (3.25)$$

In the singlet case the new defined splitting functions are

$$P_{qg} = 2n_f P_{q_i, g} = 2n_f P_{\bar{q}_i, g} = 2n_f P_{q, g}, \quad (3.26)$$

$$P_{gq} = P_{gq_i} = P_{g\bar{q}_i}, \quad (3.27)$$

$$P_{qq} = P^+ + P^{\text{PS}}, \quad \text{with} \quad P^{\text{PS}} = n_f(P_{qq}^S + P_{\bar{q}\bar{q}}^S), \quad (3.28)$$

$$P_{gg} = P_{gg}, \quad (3.29)$$

where PS stands for purely singlet.

DGLAP Equations in the Evolution Basis

In summary, the evolution equations for the non-singlets are

$$\frac{\partial q_{i1}^\pm}{\partial t} = a_s(t) \int_x^1 \frac{dz}{z} P^\pm\left(\frac{x}{z}, a_s(t)\right) q_{i1}^\pm(z, t), \quad (3.30)$$

$$\frac{\partial q^V}{\partial t} = a_s(t) \int_x^1 \frac{dz}{z} P^V\left(\frac{x}{z}, a_s(t)\right) q^V(z, t). \quad (3.31)$$

And the coupled equations for the singlet and gluon read

$$\frac{\partial}{\partial t} \begin{pmatrix} \Sigma(x, t) \\ g(x, t) \end{pmatrix} = a_s(t) \int_x^1 \frac{dz}{z} \begin{bmatrix} \left(P_{qq}\left(\frac{x}{z}, a_s(t)\right) & P_{qg}\left(\frac{x}{z}, a_s(t)\right) \right) \\ \left(P_{gq}\left(\frac{x}{z}, a_s(t)\right) & P_{gg}\left(\frac{x}{z}, a_s(t)\right) \right) \end{bmatrix} \begin{pmatrix} \Sigma(z, t) \\ g(z, t) \end{pmatrix}. \quad (3.32)$$

3.1.4 Matching Conditions

There is still a problem left concerning the number of active flavours. While it is a valid approximation to exclude flavours with pole masses beyond the virtuality $M_{hf}^2 \gg Q^2$, it can happen that the evolution crosses such a flavour threshold and new flavour contributions have to be considered. This is relevant for both the strong coupling constant and for the PDFs. Nevertheless, if the evolution is performed in fixed flavour number (FVN) scheme, these contributions are neglected. In contrast, in variable flavour number (VFN) scheme the evolution is carried out up to the threshold with the number of light quarks n_f , and from the thresholds upwards using $n_f + 1$ flavours. Thus, a matching between both effective theories, above and below the threshold, has to be carefully defined.

Matching Conditions for a_s

One way of formulating a matching condition for the strong coupling a_s is to require continuity at pole masses $Q^2 = M_{hf}^2$. However, since both effective theories are only valid for $Q^2 \gg M_{hf}^2$ and $Q^2 \ll M_{hf}^2$ continuity is not a requirement of the theory. The required consistency of the full theory leads to discontinuous thresholds at least at NNLO. Since the number of active flavours change at a heavy flavour threshold the slope will be discontinuous anyway.

The threshold conditions have been calculated and can be found in [9, 11, 33, 44].

Using $\kappa = \frac{\mu_R^2}{\mu_F^2}$ at N^lLO holds:

$$a_s^{(n_f+1)}(\kappa t_{hf}) = a_s^{(n_f)}(\kappa t_{hf}) + \sum_{n=1}^l \left\{ [a_s^{(n_f)}(\kappa t_{hf})]^{n+1} \sum_{j=0}^n C_{n,j} \ln^j \kappa \right\}, \quad (3.33)$$

where l is the order of perturbation theory, $a_s = a_s/2\pi$, and t_{hf} is the heavy flavour threshold, which is the heavy quark pole mass. The relevant factors up to NNLO for equal renormalization and factorization scale are

$$\begin{aligned} C_{1,0} &= 0, \text{ i.e. continuous at NLO for } \mu_F^2 = \mu_R^2 \text{ and} \\ C_{2,0} &= \frac{7}{6}, \text{ i.e. always discontinuous at NNLO.} \end{aligned}$$

Matching Conditions for the PDFs

The matching conditions for the PDFs have been calculated in [10] and can be found also in [9, 44]. The general structures of these matching conditions up to NNLO are:

$$q_{i,n_f+1}^\pm(x, t_{hf}) = q_{i,n_f}^\pm(x, t_{hf}) + a_s^2 [A_{qq}^{(2)} \otimes q_{i,n_f}^\pm](x, t_{hf}), \quad (3.34)$$

$$g_{n_f+1}(x, t_{hf}) = g_{n_f}(x, t_{hf}) + a_s^2 \{ [A_{gq}^{(2)} \otimes \Sigma_{n_f}](x, t_{hf}) + [A_{gg}^{(2)} \otimes g_{n_f}](x, t_{hf}) \}, \quad (3.35)$$

$$h_{n_f+1}^+(x, t_{hf}) = a_s^2 \left\{ [A_{hq}^{(2)} \otimes \Sigma_{n_f}] + [A_{hg}^{(2)} \otimes g_{n_f}](x, t_{hf}) \right\}, \quad (3.36)$$

$$h_{n_f+1}^-(x, t_{hf}) = 0, \quad (3.37)$$

where $h^\pm$ denotes the new heavy flavour basis elements, t_{hf} is the heavy flavour threshold, and the A 's are renormalized operator-matrix elements (OME's) calculated in [10]. At present, the main conclusion for this thesis is that the PDFs are continuous at LO and NLO.

3.1.5 Solution in x -Space

Despite all the modifications described above, the DGLAP evolution equations still cannot be solved analytically in x -space. Therefore the derived and partially decoupled Eqs. (3.30), (3.31), and (3.32) have to be solved numerically with an ordinary differential equation solver such as the Runge-Kutta method. The numerical solution will include additional modifications and optimizations, which are described in detail in Chapter 4.2.

3.2 The Solution in Mellin-Space

In the solution in Mellin-space, the basis transformation described above in x -space can be equally applied, so that the problem reduces to a coupled equation for the singlet and gluon and non-singlet equations. The matching conditions can likewise be transformed to the Mellin-space, but as everything is continuous up to NLO (for equal renormalization and factorization scale $\mu_F = \mu_R$) an explicit consideration is skipped in the following discussion.

3.2.1 General Mellin Transformation

The Mellin transformation [6, 12, 31, 44] for a function $f(x)$ is defined

$$f(n) := \int_0^1 dx x^{n-1} f(x), \quad (3.38)$$

with $n \in \mathbb{C}$. If $f(n)$ is holomorphic in a strip $a < \text{Re}(n) < b$ and if $f(c \pm i\omega)$ tends uniformly to zero for $\omega \rightarrow \infty$, then the inverse transformation for $a < c < b$ ($a, c, b \in \mathbb{R}$) is given by

$$f(x) = \frac{1}{2\pi i} \int_{c-i\infty}^{c+i\infty} dn x^{-n} f(n). \quad (3.39)$$

The main advantage of this transformation is that in Mellin-space the Mellin-convolution reduces to a product

$$\begin{aligned} \int_0^1 dx x^{n-1} \int_x^1 \frac{dy}{y} h(y) g(x/y) &= \int_0^1 dx x^{n-1} \left[\int_0^1 dy \int_0^1 dz h(y) g(z) \delta(x - yz) \right] \\ &= \int_0^1 dy y^{n-1} h(y) \int_0^1 dz z^{n-1} g(z) \\ &= h(n) \cdot g(n). \end{aligned} \quad (3.40)$$

3.2.2 Mellin-Space Solution of the DGLAP Equations

Solutions of the DGLAP equations in Mellin-space are discussed in [17,21,44]. Starting from the general DGLAP equation

$$\frac{\partial \mathbf{q}(x, t)}{\partial t} = \frac{a_s(t)}{2\pi} \int_x^1 \frac{dz}{z} \mathbf{P}\left(\frac{x}{z}, a_s(t)\right) \mathbf{q}(z, t) = \frac{a_s(t)}{2\pi} [\mathbf{P} \otimes \mathbf{q}](x, t), \quad (3.41)$$

using $a_s = \alpha_s/(2\pi)$ and the Mellin transformation, the problem reduces to

$$\frac{\partial \mathbf{q}(n, t)}{\partial t} = a_s(t) \mathbf{P}(n, a_s(t)) \mathbf{q}(n, t), \quad (3.42)$$

where the bold print indicates vectors/matrices in the singlet and gluon case and scalars in the non-singlet cases. Using $\mathbf{P}' = a_s \mathbf{P}$ and a_s as independent integration variable results in

$$\frac{\partial \mathbf{q}(n, a_s)}{\partial a_s} = \frac{1}{\beta_{N^m LO}(a_s)} \mathbf{P}'(n, a_s) \mathbf{q}(n, a_s). \quad (3.43)$$

Considering the series expansion for the splitting functions Eq. (3.2) and for the strong coupling Eq. (2.10), and sorting everything in a powerseries in a_s , yields

$$\begin{aligned} \frac{\partial \mathbf{q}(n, a_s)}{\partial a_s} = & -\frac{1}{\beta_0 a_s} \left[\mathbf{P}^{(0)}(n) + a_s (\mathbf{P}^{(1)}(n) - \beta_1 \mathbf{P}^{(0)}(n)) \right. \\ & \left. + a_s^2 (\mathbf{P}^{(2)}(n) - \beta_1 \mathbf{P}^{(1)}(n) + (\beta_1^2 - \beta_2) \mathbf{P}^{(0)}(n)) + \dots \right] \mathbf{q}(n, a_s). \end{aligned} \quad (3.44)$$

Introducing further the abbreviations

$$\mathbf{R}_0 = \frac{1}{\beta_0} \mathbf{P}^{(0)}, \quad \mathbf{R}_k = \frac{1}{\beta_0} \mathbf{P}^{(k)} - \sum_{i=1}^k \beta_i \mathbf{R}_{k-i} \quad (3.45)$$

leads to a compact form of the DGLAP equation

$$\frac{\partial \mathbf{q}(n, a_s)}{\partial a_s} = -\frac{1}{a_s} \left[\mathbf{R}_0(n) + \sum_{k=1}^{\infty} a_s^k \mathbf{R}_k(n) \right] \mathbf{q}(n, a_s). \quad (3.46)$$

At LO, the exponential solution of this equation reads

$$\begin{aligned} \mathbf{q}_{LO}(n, a_s) &= \exp \left[-\ln \left(\frac{a_s}{a_0} \right) \mathbf{R}_0(n) \right] \mathbf{q}(n, a_0) \\ &= \left(\frac{a_s}{a_0} \right)^{-\mathbf{R}_0(n)} \mathbf{q}(n, a_0) = \mathbf{L}(n, a_s, a_0) \mathbf{q}(n, a_0), \end{aligned} \quad (3.47)$$

where a_s denotes the coupling at the final scale and a_0 the coupling at the initial scale. Beyond leading order, the solution is more complicated, since in the singlet and gluon case the matrices $\mathbf{R}_k$ do not commute and therefore cannot be written in an exponential form.

The Singlet and Gluon Case

One strategy to solve the equation in the coupled singlet and gluon case is to expand the solution around the LO solution (Eq. 3.47)

$$\mathbf{q}(n, a_s) = \mathbf{U}(n, a_s) \mathbf{L}(n, a_s, a_0) \mathbf{U}^{-1}(n, a_0) \mathbf{q}(n, a_0) \quad (3.48)$$

$$= \left[1 + \sum_{k=1}^{\infty} a_s^k \mathbf{U}_k(n) \right] \mathbf{L}(N, a_s, a_0) \left[1 + \sum_{k=1}^{\infty} a_0^k \mathbf{U}_k(n) \right]^{-1} \mathbf{q}(N, a_0), \quad (3.49)$$

where $\mathbf{U}$ is expressed through a powerseries in a_s and matrices $\mathbf{U}_k(n)$, which only dependent on n . The calculation of this matrices is described below. Moreover, the third factor normalizes the evolution operator to the unit matrix at $a_s = a_0$.

Calculation of the $\mathbf{U}$ -Matrices For all higher order solutions the already introduced $\mathbf{U}$ matrices have to be defined. Inserting the general solution Eq. (3.49) into the evolution Eq. (3.46) and sorting powers of a_s results in a series of commutation relations

$$\begin{aligned} [\mathbf{U}_1, \mathbf{R}_0] &= \mathbf{R}_1 + \mathbf{U}_1, \\ [\mathbf{U}_2, \mathbf{R}_0] &= \mathbf{R}_2 + \mathbf{R}_1 \mathbf{U}_1 + 2\mathbf{U}_2, \\ &\vdots \end{aligned} \quad (3.50)$$

$$[\mathbf{U}_k, \mathbf{R}_0] = \mathbf{R}_k + \sum_{i=1}^{k-1} \mathbf{R}_{k-i} \mathbf{U}_i + k\mathbf{U}_k = \tilde{\mathbf{R}}_k + k\mathbf{U}_k.$$

Then, performing an eigenvalue decomposition of Eq. (3.45)

$$\mathbf{R}_0 = \frac{1}{\beta_0} \begin{pmatrix} P_{qq}^{(0)} & P_{qg}^{(0)} \\ P_{gq}^{(0)} & P_{gg}^{(0)} \end{pmatrix} \quad (3.51)$$

with the eigenvalues

$$\begin{aligned} r_{\pm} &= \frac{1}{2\beta_0} \left[P_{qq}^{(0)} + P_{gg}^{(0)} \pm \sqrt{\left(P_{qq}^{(0)} - P_{gg}^{(0)} \right)^2 - 4P_{qg}^{(0)} P_{gq}^{(0)}} \right] \\ &= \frac{1}{2\beta_0} \left[P_{qq}^{(0)} + P_{gg}^{(0)} \pm \sqrt{\left(P_{qq}^{(0)} - P_{gg}^{(0)} \right)^2 + 4P_{qg}^{(0)} P_{gq}^{(0)}} \right] \end{aligned} \quad (3.52)$$

leads to

$$\mathbf{R}_0 = r_+ \mathbf{e}_+ + r_- \mathbf{e}_- \quad (3.53)$$

$$\text{with } \mathbf{e}_{\pm} = \frac{1}{r_+ - r_-} [\mathbf{R}_0 - r_{\mp} \mathbb{1}], \quad (3.54)$$

where $\mathbb{1}$ is the identity matrix. The matrices $\mathbf{e}_-$ and $\mathbf{e}_+$ obey the following relations

$$\mathbf{e}_+ \mathbf{e}_- = \mathbf{e}_- \mathbf{e}_+ = 0, \quad (3.55)$$

$$\mathbf{e}_+ \mathbf{e}_+ = \mathbf{e}_+, \quad (3.56)$$

$$\mathbf{e}_- \mathbf{e}_- = \mathbf{e}_-, \quad (3.57)$$

$$\mathbf{e}_+ + \mathbf{e}_- = \mathbb{1}. \quad (3.58)$$

Inserting all that into the exponential representation of the LO solution (Eq. 3.47) yields

$$\mathbf{L}(a_s, a_0, n) = \mathbf{e}_-(n) \left(\frac{a_s}{a_0} \right)^{-r_-(n)} + \mathbf{e}_+(n) \left(\frac{a_s}{a_0} \right)^{-r_+(n)}. \quad (3.59)$$

Moreover, Eq. (3.58) can be used to decompose the $\mathbf{U}_k$ matrices of Eq. (3.49)

$$\begin{aligned} \mathbf{U}_k &= (\mathbf{e}_+ + \mathbf{e}_-) U_k (\mathbf{e}_+ + \mathbf{e}_-) \\ &= \mathbf{e}_- U_k \mathbf{e}_- + \mathbf{e}_- U_k \mathbf{e}_+ + \mathbf{e}_+ U_k \mathbf{e}_- + \mathbf{e}_+ U_k \mathbf{e}_+. \end{aligned} \quad (3.60)$$

This expression can be used in the commutation relations Eqs. (3.50) so that

$$\mathbf{U}_k = -\frac{1}{k} \left[\mathbf{e}_- \tilde{\mathbf{R}}_k \mathbf{e}_- + \mathbf{e}_+ \tilde{\mathbf{R}}_k \mathbf{e}_+ \right] + \frac{\mathbf{e}_+ \tilde{\mathbf{R}}_k \mathbf{e}_-}{r_- - r_+ - k} + \frac{\mathbf{e}_- \tilde{\mathbf{R}}_k \mathbf{e}_+}{r_+ - r_- - k}. \quad (3.61)$$

The poles in the denominator are cancelled by the $\mathbf{U}^{-1}$ terms. Now, there are at least three possibilities to define N^mLO solutions.

Solution 1 (x-Space Solution) In the first truncation scheme, the terms originating from $\beta_{\text{N}^m\text{LO}}$ (Eq. 2.10) and $P_{\text{N}^m\text{LO}}$ (Eq. 3.2) are kept in accordance with N^mLO solution in the x -space, described above. This means that the sums in the evolution matrices $\mathbf{U}$ run up to infinity (or up to contributions that are sufficiently high numerically). At NLO the differential equation reads

$$\begin{aligned} \frac{\partial \mathbf{q}(n, a_s)}{\partial a_s} &= -\frac{1}{\beta_0 a_s} \left[\mathbf{P}^{(0)}(n) + a_s (\mathbf{P}^{(1)}(n) - \beta_1 \mathbf{P}^{(0)}(n)) \right. \\ &\quad \left. + a_s^2 (-\beta_1 \mathbf{P}^{(1)}(N) + \beta_1^2 \mathbf{P}^{(0)}(n)) \right] \mathbf{q}(n, a_s), \end{aligned} \quad (3.62)$$

with

$$\mathbf{R}_0 = \frac{1}{\beta_0} \mathbf{P}^{(0)}, \quad \mathbf{R}_1 = \frac{1}{\beta_0} \mathbf{P}^{(1)} - \beta_1 \mathbf{R}_0, \quad \mathbf{R}_k = -\beta_1 \mathbf{R}_{k-1} \quad \text{for } k \geq 2.$$

Solution 2 Another choice is to keep terms in the matrices $\mathbf{R}$ up to the specified order but still keep them up to infinity in $\mathbf{U}$. Therefore at NLO remains

$$\frac{\partial \mathbf{q}(n, a_s)}{\partial a_s} = -\frac{1}{\beta_0 a_s} \left[\mathbf{P}^{(0)}(n) + a_s (\mathbf{P}^{(1)}(n) - \beta_1 \mathbf{P}^{(0)}(n)) \right] \mathbf{q}(n, a_s), \quad (3.63)$$

with

$$\mathbf{R}_0 = \frac{1}{\beta_0} \mathbf{P}^{(0)}, \quad \mathbf{R}_1 = \frac{1}{\beta_0} \mathbf{P}^{(1)} - \beta_1 \mathbf{R}_0, \quad \mathbf{R}_k = 0 \quad \text{for } k \geq 2.$$

Solution 3 (Truncated Solution) Alternatively, one could keep terms in $\mathbf{R}$ and $\mathbf{U}$ only up to the specified order. At NLO this leads to the differential equation

$$\frac{\partial \mathbf{q}(n, a_s)}{\partial a_s} = -\frac{1}{\beta_0 a_s} \left[\mathbf{P}^{(0)}(n) + a_s (\mathbf{P}^{(1)}(n) - \beta_1 \mathbf{P}^{(0)}(n)) \right] \mathbf{q}(n, a_s), \quad (3.64)$$

with

$$\mathbf{R}_0 = \frac{1}{\beta_0} \mathbf{P}^{(0)}, \quad \mathbf{R}_1 = \frac{1}{\beta_0} \mathbf{P}^{(1)} - \beta_1 \mathbf{R}_0, \quad \mathbf{R}_k = 0 \quad \text{for } k \geq 2.$$

Thus, the solution reads

$$\mathbf{q}(n, a_s) = [1 + a_s \mathbf{U}_1(n)] \mathbf{L}(n, a_s, a_0) [1 + a_0 \mathbf{U}_1(n)]^{-1} \mathbf{q}(n, a_0). \quad (3.65)$$

This can be further simplified with the help of a series expansion of the third term $\mathbf{U}^{-1}$ and by neglecting all terms of a_s^2

$$\mathbf{q}(n, a_s) = [\mathbf{L}(N, a_s, a_0) + a_s \mathbf{U}_1(n) \mathbf{L}(n, a_s, a_0) - \mathbf{L}(N, a_s, a_0) a_0 \mathbf{U}_1(n)] \mathbf{q}(n, a_0). \quad (3.66)$$

Finally, using the explicit LO solution, the Solution 3 at NLO, where $\tilde{\mathbf{R}}_1 = \mathbf{R}_1$ holds, is

$$\begin{aligned} \mathbf{q}(n, a_s) = & \left[\left(\frac{a_s}{a_0} \right)^{-r_-} \left\{ \mathbf{e}_- + (a_0 - a_s) \mathbf{e}_- \mathbf{R}_1 \mathbf{e}_- \right. \right. \\ & \left. \left. - \left(a_0 - a_s \left(\frac{a_s}{a_0} \right)^{r_- - r_+} \right) \frac{\mathbf{e}_- \mathbf{R}_1 \mathbf{e}_+}{r_+ - r_- - 1} \right\} + (+ \leftrightarrow -) \right] \mathbf{q}(n, a_0), \end{aligned} \quad (3.67)$$

where for $(r_+ - r_-) = 1$ not only the denominator but also the coupling prefactor vanishes. In contrast to Solution 3, the Solutions 1 and 2 include an infinite sum over $\mathbf{U}_k$, which will require in each numerical solution an abort at sufficient high k (e.g. for $k \gtrsim 15$ as used in QCD-PEGASUS [44]).

All three solutions differ only in contributions beyond the order under consideration. In particular, the first solution should reproduce the solution in x -space. While the first two solutions fulfill Eqs. (3.62) and (3.63), respectively, exact (if the whole infinite sums in $\mathbf{U}$ are considered), the third solution fulfills Eq. (3.64) only up to order m . The differences between these solutions can be used as an estimate for a lower limit of the uncertainties due to higher-order corrections [44].

The Non-Singlet Case

In the non-singlet case, the evolution equation is a scalar equation so that the commutation relations vanish. The different N^mLO solutions from above are defined as follows.

Solution 1 (x -Space Solution) Since no commutation relations have to be considered, Eq. (3.46) can be directly integrated with a truncation similar to Solution 1 in the singlet and gluon case, i.e. keeping the terms from $\beta_{\text{N}^{\text{mLO}}}$ (Eq. 2.10) and $P_{\text{N}^{\text{mLO}}}$ (Eq. 3.2) up to the specified order. In NLO using $U_1 = -R_1$:

$$q^{\pm, \text{V}}(n, a_s) = \exp \left[\frac{U_1^{\pm, \text{V}}}{\beta_1} \ln \left(\frac{1 + \beta_1 a_s}{1 + \beta_1 a_0} \right) \right] \left(\frac{a_s}{a_0} \right)^{-R_0^{\pm, \text{V}}} q^{\pm, \text{V}}(n, a_0). \quad (3.68)$$

Solution 2 Keeping only the terms in R up to the chosen order yields, in NLO with $U_1 = -R_1$,

$$q^{\pm, \text{V}}(n, a_s) = \exp \left[U_1^{\pm, \text{V}}(a_s - a_0) \right] \left(\frac{a_s}{a_0} \right)^{-R_0^{\pm, \text{V}}} q^{\pm, \text{V}}(n, a_0). \quad (3.69)$$

Solution 3 (Truncated Solution) The full untruncated solution is

$$q^{\pm, \text{V}}(n, a_s) = \left[1 + \sum_{k=1}^m a_s^k U_k^{\pm, \text{V}} \right] \left[1 + \sum_{k=1}^m a_0^k U_k^{\pm, \text{V}} \right]^{-1} \left(\frac{a_s}{a_0} \right)^{-R_0^{\pm, \text{V}}} q^{\pm, \text{V}}(n, a_0). \quad (3.70)$$

Truncation of the U_k at NLO yields

$$q^{\pm, \text{V}}(n, a_s) = \left[\frac{1 + a_s U_1^{\pm, \text{V}}}{1 + a_0 U_1^{\pm, \text{V}}} \right] \left(\frac{a_s}{a_0} \right)^{-R_0^{\pm, \text{V}}} q^{\pm, \text{V}}(n, a_0), \quad (3.71)$$

and fully truncated

$$q^{\pm, \text{V}}(n, a_s) = \left[1 + a_s U_1^{\pm, \text{V}} - a_0 U_1^{\pm, \text{V}} \right] \left(\frac{a_s}{a_0} \right)^{-R_0^{\pm, \text{V}}} q^{\pm, \text{V}}(n, a_0). \quad (3.72)$$

Again, using this solution in Eq. (3.46) returns $U_1 = -R_1$.

3.2.3 Splitting Functions / Anomalous Dimensions

In the solutions described above already the splitting functions in Mellin-space have been used. The Mellin transformations of the LO x -space splitting functions Tab. 3.1

are

$$\gamma_{qq}^{(0)}(n) = C_F \left[-\frac{1}{2} + \frac{1}{n(n+1)} - 2 \sum_{k=2}^n \frac{1}{k} \right] \quad (3.73)$$

$$= C_F \left[\frac{3}{2} + \frac{1}{n(n+1)} - 2S_1(n) \right], \quad (3.74)$$

$$\gamma_{gq}^{(0)}(n) = C_F \left[\frac{(2+n+n^2)}{n(n^2-1)} \right], \quad (3.75)$$

$$\gamma_{gg}^{(0)}(n) = T_F \left[\frac{(2+n+n^2)}{n(n+1)(n+2)} \right], \quad (3.76)$$

$$\gamma_{gg}^{(0)}(n) = 2C_A \left[-\frac{1}{12} + \frac{1}{n(n-1)} + \frac{1}{(j+1)(j+2)} - \sum_{k=2}^n \frac{1}{k} \right] - \frac{2}{3}n_f T_F \quad (3.77)$$

$$= 2C_A \left[\frac{11}{12} + \frac{1}{n(n-1)} + \frac{1}{(j+1)(j+2)} - S_1(n) \right] - \frac{2}{3}n_f T_F, \quad (3.78)$$

where S_1 is a harmonic sum of depths one, i.e. a simple harmonic sum. The plus distributions

$$\int_0^1 dx \frac{x^{n-1} f(x)}{(1-x)_+} = \int_0^1 dx \frac{x^{n-1} f(x) - f(1)}{(1-x)} \quad (3.79)$$

are transformed, as in [43], using the relation

$$\frac{x^m}{1-x} = \sum_{k=0}^{\infty} x^k. \quad (3.80)$$

3.2.4 Inverse Mellin Transformation

In a final step perhaps a back-transformation to the x -space is required. In general, the inverse Mellin transformation is given by

$$f(x) \equiv \frac{1}{2\pi i} \int_{c-i\infty}^{c+i\infty} dn x^{-n} f(n), \quad (3.81)$$

with $c \in \mathbb{R}$, which requires c to lie right of all singularities of $f(n)$ (see Section 3.2.1). In the above solution two sources of singularities have to be considered: the anomalous dimensions as well as the initial parametrization in Mellin-space. The LO anomalous dimensions include the rightmost pole for $n = 1$, so that depending on the initial conditions c has to be chosen at least greater than one. In addition, for the inversion the anomalous dimensions have to be continued in the complex plane (see Section 4.3.2).

3.2.5 Solution in Mellin-Space

In this section the solutions of the DGLAP equations up to NLO in Mellin-space have been derived explicitly. Whereas the solutions in Mellin-space are analytic, the

inversion is not, and hence an numerical procedure is required – if x -space PDFs are demanded. The analytic solutions can be implemented almost as described above. The details of the implementation are described in Chapter 4.3.

Chapter 4

Implementation of x -Space and Mellin-Space PDF Evolutions

In this chapter a more detailed description of the actual implementation of the DGLAP solutions is given. First, the general code structure and user interface is sketched, and then the details of the numerical implementations first in x -space and afterwards in Mellin-space are described.

4.1 The General Code Structure

The code consists of a streamlined user interface, which allows easy switching between parametrizations and evolution methods. Its general structure uses two inheritance levels. In short, one central abstract class *CORE* has been implemented as generally as possible. It controls the x - Q -Grid and the PDFs and contains the two main virtual and abstract functions *fillGridatInputScale* and *fillGrid*. These virtual functions are implemented by new classes inheriting virtually from *CORE*: *IntialParametrization<Name>*, which implements *fillGridatInputScale*, and *Evolution<Name>*, which implements *fillGrid*. Finally, the user has to declare a new class just inheriting from a chosen parametrization and evolution class. Through this class both mentioned methods and the grid data are accessible.

There is, of course, a larger amount of classes controlling the strong coupling beta function, splitting functions, interpolation weights etc., but normally there is no need for the user to access them. The evolution settings can be controlled with an input file using YAML for data serialization.

4.2 The x -Space Evolution

The x -space evolution is oriented on the well established HOPPET toolkit [42]. Nevertheless, beyond the object-orientated implementation in C++ also some modifications and optimizations have been done.

The general idea is to use a grid in x -space and interpolate it piecewise polynomial with the help of Lagrange interpolation. The splitting functions can be represented

in terms of their convolution with these basis functions. Afterwards, the evolution in Q is carried out, currently using fourth-order Runge–Kutta integration.

4.2.1 Discretization and Interpolation of the Parton Momentum Densities

First, the evolution is designed to operate on parton momentum densities $xq_i(x, Q^2)$ instead of parton densities $q_i(x, Q^2)$, as they are smoother and so numerical deviations are reduced. Using $\tilde{q}_i(x, Q^2) = xq_i(x, Q^2)$ the general DGLAP equation reads:

$$\frac{\partial \tilde{q}(x, t)}{\partial t} = a_s \int_x^1 dz \frac{x}{z^2} P\left(\frac{x}{z}, a_s(t)\right) \tilde{q}(z, t) = a_s \int_x^1 dz P(z, a_s(t)) \tilde{q}\left(\frac{x}{z}, t\right). \quad (4.1)$$

The introduced momentum densities are represented by an equidistant grid of N_x points in a variable $y = \ln(1/x)$, so that $y_\alpha = \alpha \delta y = \ln(1/x_\alpha)$ with $\alpha = 1 \dots N_x$. Thus, the moments at an arbitrary y are given with the help of Lagrange interpolation by

$$xq(y = \ln(1/x), t) = \sum_{\alpha} w_{\alpha}(y) q_{\alpha}(t) \quad \text{with} \quad q_{\alpha}(t) = x_{\alpha} q(y_{\alpha}, t), \quad (4.2)$$

whereby α runs over all $n+1$ points used for the interpolation of order n . The $w_{\alpha}(y)$ represent the Lagrange interpolation weights

$$w_{\alpha}(y) = \prod_{i \neq \alpha} \frac{y - y_i}{y_{\alpha} - y_i}, \quad (4.3)$$

whereby i runs over all $n+1$ points except $i = \alpha$.

General Properties of Lagrange-Like Interpolations

To take care of all possible sources of errors, the interpolation is the first approximation to be considered. Note, therefore, that it is indeed continuous at the grid points, but the first derivatives may already be discontinuous.

Furthermore, the accuracy of the interpolated values will increase with the number of grid points (N_x), i.e. smaller grid distances, but that requires more computation time. Moreover, the interpolation order n could be increased, but – besides higher computational efforts – the accuracy is not automatically increased, because of possible oscillations in higher order polynomials. Lower order interpolations such as linear interpolation ($n = 1$) as well as high interpolation orders like $n > 10$ tend to misfit the data, and an appropriate value in between has to be chosen [40, p. 110 ff.]. Numerical tests indicate $n = 6$ to be a reasonable choice (see Chapter 5).

Chosen Interpolation Scheme and Boundary Problems

The interpolation only requires choosing points in the vicinity of the interpolation point. However, instead of the obvious choice of the $n+1$ points centered around

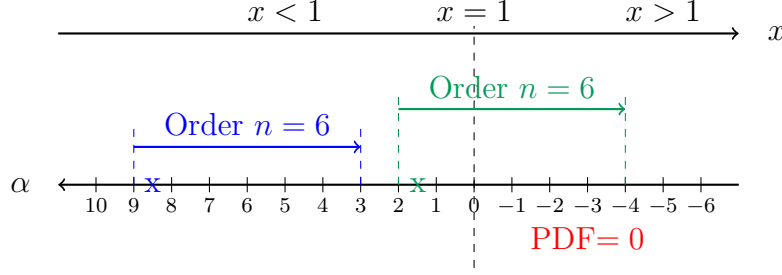


Fig. 4.1: Implemented interpolation scheme.

the evaluation point, another scheme was chosen: as shown in Fig. 4.1 a point is interpolated by the next lower (higher) grid point in y (x) and the n higher (lower) grid points in y (x). The reason for this choice is an optimization of calculation time and a simplification of the boundary value problems. Since this scheme is shift-invariant, it reduces the number of terms to be calculated roughly by a factor of the interpolation order as will be shown in Sections 4.2.3-4.2.3. Moreover, for small x all supporting points are in the grid. In contrast, for large x the supporting points lie outside the grid, i.e. correspond to $x > 1$, which is also shown in Fig. 4.1. That is of course not physical, but as all PDFs should vanish for $x \rightarrow 1$, all these values can be set to zero whilst only producing small deviations (depending on the sampling size in x). This way, the first n sections are interpolated with lower-order polynomials. In sum, the following calculations are done with an interpolation using the $n + 1$ points $\{y_i, \dots, y_{i-n}\}$ for interpolating a value in the section $[y_i, y_{i-1}]$ (see Fig. 4.1).

4.2.2 Implementation of the Interpolation

As Eq. (4.2) shows, the $n + 1$ interpolation weights have to be used each time a point is interpolated. To reduce computation time, their calculation can be modified with regard to the direct implementation of Eq. (4.3).

Each value of the introduced variable y can be represented through

$$y = y_\beta + x\delta y, \quad \text{with } x \in [0, -1] \quad \text{or} \quad y \in [y_\beta, y_{\beta-1}], \quad (4.4)$$

in the equidistant (δy) grid. The $n + 1$ weights w_α are given by

$$w_\alpha(y) = \prod_{i=\beta-n \neq \alpha}^{\beta} \frac{\beta + x - i}{\alpha - i}. \quad (4.5)$$

Using an index shift $i \rightarrow i - \beta$ and setting $x' = -x \geq 0$ results in

$$w_\alpha(y) = \prod_{i=0 \neq (\beta-\alpha)}^n \frac{x' - i}{(\beta - \alpha) - i}. \quad (4.6)$$

One important feature of this representation is that it only depends on the difference $\beta - \alpha$, and not on α and β separately. This will be used later. Note also that if

$x' = 0$ and $x' = 1$, then the weights are evaluated at grid points. At these points holds:

$$\text{if } x' = 0, \text{ then } w_\alpha(y) = \begin{cases} 1, & \text{for } \beta - \alpha = 0, \\ 0, & \text{else,} \end{cases} \quad (4.7)$$

$$\text{if } x' = 1, \text{ then } w_\alpha(y) = \begin{cases} 1, & \text{for } \beta - \alpha = 1, \\ 0, & \text{else.} \end{cases} \quad (4.8)$$

Furthermore, to reduce computation time, the derived expression can be split in weight normalizations, which have to be calculated only once because of the equidistant grid, and point dependent contributions:

$$w_\alpha(y) = \frac{1}{\prod_{i=0 \neq (\beta-\alpha)}^n (\beta - \alpha) - i} \prod_{i=0 \neq (\beta-\alpha)}^n x' - i. \quad (4.9)$$

Since there are only $n+1$ possible values of $\beta \in \{\alpha, \alpha+n\}$ for each α , a new variable can be introduced $v = \beta - \alpha \in \{0, n\}$ so that

$$w(v, x) = \frac{1}{\prod_{i=0 \neq v}^n v - i} \prod_{i=0 \neq v}^n x' - i. \quad (4.10)$$

Splitting this function into the normalization $wn(v)$ and the point dependent factor $wx(v, x')$ reads

$$w(v, x') = wn(v) \cdot wx(v, x'). \quad (4.11)$$

Weight Normalization

As the normalization does not depend on α and β separately, but only on their difference, it can be pre-calculated. First $wn(-n+1)$ is calculated and afterwards a loop is performed, which uses the following recursive relation:

$$1/wn(v+1) = \prod_{i=0 \neq v+1}^n v+1-i = \prod_{i'=-1 \neq v}^{n-1} v-i' = \frac{1}{wn(v)} \frac{v+1}{v-n}. \quad (4.12)$$

Point Dependent Weight Factor

Instead of considering the exception $i \neq v$ in the product of the point dependent weight factors explicitly (see Eq. 4.9), the product is performed over all $i = 0$ to n and is divided afterwards by the term belonging to $i = v$. This way, divisions by zero are introduced for the special cases Eqs. (4.7) and (4.8), which have to be excluded. In the notation of Eq. (4.11) these cases read

$$w(v, x' = 0) = \begin{cases} 1, & \text{for } v = 0, \\ 0, & \text{for all other cases,} \end{cases} \quad (4.13)$$

$$w(v, x' = 1) = \begin{cases} 1, & \text{for } v = 1, \\ 0, & \text{for all other cases.} \end{cases} \quad (4.14)$$

4.2.3 Discretization and Interpolation of the Convolution

In the same way as the momentum distributions were described on the grid, the whole convolution can be expressed by

$$(P \otimes q)(y, t) = \int_{e^{-y}}^1 dz P(z, a_s(t)) q\left(\frac{e^{-y}}{z}, t\right) = \sum_{\alpha} \omega_{\alpha}(y) (P \otimes q)_{\alpha}(t), \quad (4.15)$$

where q denotes momentum densities and $\otimes$ now denotes the convolution with the momentum densities. The evaluation of the convolution at a grid point x_{α} is given by

$$(P \otimes q)_{\alpha}(t) = \int_{x_{\alpha}}^1 dz P(z, a_s(t)) q\left(\frac{e^{-y_{\alpha}}}{z}, t\right), \quad (4.16)$$

where q is again the momentum density, which itself can be expressed by the interpolation. Introducing the integer division symbol “ $\setminus$ ”, returning the integer quotient of a division, this can be written as

$$\begin{aligned} (P \otimes q)_{\alpha}(t) &= \int_{e^{-y_{\alpha}}}^1 dz P(z, t) \sum_{\beta=\alpha-(\ln(z)\setminus\delta_y)-n}^{\alpha-(\ln(z)\setminus\delta_y)-n} w_{\beta}(\ln(z/x_{\alpha})) q_{\beta}(t) \\ &= \int_{e^{-y_{\alpha}}}^1 dz P(z, t) \sum_{\beta=\alpha-(\ln(z)\setminus\delta_y)}^{\alpha-(\ln(z)\setminus\delta_y)-n} w_{\beta}(y_{\alpha} + \ln(z)) q_{\beta}(t). \end{aligned} \quad (4.17)$$

The expression is rather complicated because the support points depend on the argument of the weights w , i.e. on z . If the integral is decomposed in each grid section, the expression can be simplified. Therefore, when integrating over the interval $[e^{-y_i}, e^{-y_{i-1}}]$, then

$$\ln(z) \in [-y_i, -y_{i-1}] \quad \text{or} \quad \ln(z) = -(y_{i-1} + x\delta_y) = -y_{i-1} + x'\delta_y, \quad (4.18)$$

with $x' = -\frac{\ln(z)\setminus\delta_y}{\delta_y} \in [0, -1]$. Thus,

$$\begin{aligned} (P \otimes q)_{\alpha}(t) &= \sum_{i=1}^{\alpha} \int_{e^{-y_i}}^{e^{-y_{i-1}}} dz P(z, t) \sum_{\beta=\alpha-i-(n-1)}^{\alpha-i+1} w_{\beta}(y_{\alpha} - y_{i-1} + x'\delta_y) q_{\beta}(t) \\ &= \sum_{i=1}^{\alpha} \sum_{\beta=\alpha-i-(n-1)}^{\alpha-i+1} q_{\beta}(t) \int_{e^{-y_i}}^{e^{-y_{i-1}}} dz P(z, t) w_{\beta}(y_{\alpha} - y_{i-1} + x'\delta_y). \end{aligned} \quad (4.19)$$

This is the central representation which is used in the numerical calculations and includes $\alpha \cdot n$ terms. To gain more insights into the properties of the solution, it can be further modified. Since each q_{β} contributes only in n consecutive intervals,

i.e for $i \in [\alpha - \beta - n + 1, \alpha - \beta + 1]$, the sum may be rearranged to

$$\begin{aligned}
 (P \otimes q)_\alpha(t) &= \sum_{\beta=-n+1}^{\alpha} \left[\sum_{i=0}^n \int_{e^{-y_{\alpha-\beta+1-i}}}^{e^{-y_{\alpha-\beta-i}}} dz P(z, t) w_\beta(y_\alpha - y_{\alpha-\beta-i} + x' \delta_y) \right] q_\beta(t) \\
 &= \sum_{\beta=-n+1}^{\alpha} \left[\sum_{i=0}^n \int_{e^{-y_{\alpha-\beta+1-i}}}^{e^{-y_{\alpha-\beta-i}}} dz P(z, t) w_\beta((\beta + i) \delta_y + x' \delta_y) \right] q_\beta(t)
 \end{aligned} \tag{4.20}$$

if the additional condition

$$-\beta + 1 - i < 0 \quad \forall i \text{ and } \beta, \tag{4.21}$$

which avoids over counting, is fulfilled. Moreover, using the above described property of the interpolation weights, $w_\alpha(y) = w(\beta - \alpha, x')$, the problem reduces to

$$(P \otimes q)_\alpha(t) = \sum_{\beta=-n+1}^{\alpha} \left[\sum_{i=0}^n \int_{e^{-y_{\alpha-\beta+1-i}}}^{e^{-y_{\alpha-\beta-i}}} dz P(z, t) w(i \delta_y + x' \delta_y) \right] q_\beta(t) \tag{4.22}$$

$$= \sum_{\beta=-n+1}^{\alpha} P_{\alpha\beta}(t) q_\beta(t), \tag{4.23}$$

with

$$P_{\alpha\beta} = \sum_{i=0}^n \int_{e^{-y_{\alpha-\beta+1-i}}}^{e^{-y_{\alpha-\beta-i}}} dz P(z, t) w(i \delta_y + x' \delta_y), \tag{4.24}$$

whereby $P_{\alpha\beta}$ now only depends on the difference of α and β : $P_{\alpha-\beta}$. Therefore, it is enough to calculate $P_{N_x-\beta}$ for all $\beta \in [-n+1, N_x]$ ($= \{P_0, P_{N_x+n-1}\}$). Instead of directly computing these terms, Eq. (4.19) is used and thereby the terms for each q_β are collected, since this procedure avoids the additional condition Eq. (4.21).

Irrespective of this last remark, the convolution at grid point y_α under the assumption $q_\beta = 0$ for $\beta < 0$ (see Section 4.2.2) is finally

$$(P \otimes q)_\alpha(t) = \sum_{\beta=0}^{\alpha} P_{\alpha-\beta} q_\beta. \tag{4.25}$$

Number of Terms

Due to the equidistant grid and chosen shift invariant interpolation scheme the number of calculations has been significantly reduced. Whereas in general for each y_α (N_x) one interpolation (n) has to be performed, in total $\mathcal{O}(N_x \cdot n)$ terms, now, in the above formulation, only $\mathcal{O}(N_x + n)$ terms are sufficient, since the kernels P depend only on the difference of α and β .

4.2.4 Final DGLAP Equation and Evolution Operators

Finally, the DGLAP equations can be reformulated on the grid. Since the interpolation weights for the momentum densities and the convolution terms are identical, the result is simply

$$\frac{\partial q_\alpha(t)}{\partial t} = a_s(t) \sum_{\beta=0}^{\alpha} P_{\alpha\beta}(t) q_\beta(t), \quad (4.26)$$

taking into account that $a_s = \alpha_s/(2\pi)$. In the evolution basis described in Section 3.1.3 they read

$$\frac{\partial q_\alpha^{\pm,V}(t)}{\partial t} = a_s(t) \sum_{\beta=0}^{\alpha} P_{\alpha\beta}^{\pm,V}(t) q_\beta^{\pm,V}(t), \quad (4.27)$$

and

$$\frac{\partial}{\partial t} \begin{pmatrix} \Sigma_\alpha(t) \\ g_\alpha(t) \end{pmatrix} = a_s(t) \sum_{\beta=0}^{\alpha} \begin{pmatrix} P_{\alpha\beta}^{qq}(t) & P_{\alpha\beta}^{qg}(t) \\ P_{\alpha\beta}^{gq}(t) & P_{\alpha\beta}^{gg}(t) \end{pmatrix} \begin{pmatrix} \Sigma_\beta(t) \\ g_\beta(t) \end{pmatrix}. \quad (4.28)$$

Now, if the PDFs are used in a fitting procedure, a repeated evaluation of these equations is highly time consuming. As for each basis element the derived DGLAP equation is a matrix in the x -grid space (α) of homogeneous ordinary differential equations with non-constant coefficients, it has a fundamental matrix solution (fundamental system). Therefore the problem can be further simplified by splitting the evolution in evolution operators M and the initial conditions.

The Non-Singlet Case

This way, each grid point in the uncoupled case can be described by

$$q_\alpha(t) = \sum_{\beta=0}^{\alpha} M_{\alpha\beta}(t) q_\beta(t_0) \quad \text{with} \quad M_{\alpha\beta}(t_0) = \delta_{\alpha\beta}. \quad (4.29)$$

Inserting Eq. (4.29) into the derived DGLAP equation for the uncoupled case, Eq. (4.26), yields

$$\begin{aligned} \sum_{\beta=0}^{\alpha} q_\beta(t_0) \frac{\partial M_{\alpha\beta}(t)}{\partial t} &= a_s(t) \sum_{\beta=0}^{\alpha} P_{\alpha\beta}(t) \left[\sum_{\gamma=0}^{\beta} M_{\beta\gamma}(t) q_\gamma(t_0) \right] \\ &= a_s \sum_{\beta=0}^{\alpha} \sum_{\gamma=0}^{\beta} P_{\alpha\beta}(t) M_{\beta\gamma}(t) q_\gamma(t_0). \end{aligned} \quad (4.30)$$

A careful investigation of these summations

$$\begin{aligned} \beta = 0 & \quad P_{\alpha 0} M_{00} q_0 \\ \beta = 1 & \quad P_{\alpha 1} M_{10} q_0 \quad + \quad P_{\alpha 1} M_{11} q_1 \\ \beta = 2 & \quad P_{\alpha 2} M_{20} q_0 \quad + \quad P_{\alpha 2} M_{21} q_1 \quad \dots \end{aligned}$$

proofs that Eq. (4.30) can be rewritten as

$$\begin{aligned} \sum_{\beta=0}^{\alpha} q_{\beta}(t_0) \frac{\partial M_{\alpha\beta}(t)}{\partial t} &= a_s(t) \sum_{\beta=0}^{\alpha} q_{\beta}(t_0) \left[\sum_{\gamma=\beta}^{\alpha} P_{\alpha\gamma}(t) M_{\gamma\beta}(t) \right] \\ \Rightarrow \frac{\partial M_{\alpha\beta}(t)}{\partial t} &= a_s(t) \sum_{\gamma=\beta}^{\alpha} P_{\alpha\gamma}(t) M_{\gamma\beta}(t). \end{aligned} \quad (4.31)$$

Since $P_{\alpha\beta} = P_{\alpha-\beta}$ and $M_{\alpha\beta}(t_0) = M_{\alpha-\beta}(t_0)$, it also follows that $M_{\alpha\beta}(t) = M_{\alpha-\beta}(t)$. Thus, again only one row, e.g. $\beta = 0$, has to be evolved:

$$\frac{\partial M_{\alpha}(t)}{\partial t} = a_s(t) \sum_{\gamma=0}^{\alpha} P_{\alpha-\gamma}(t) M_{\gamma}(t). \quad (4.32)$$

The Singlet/Gluon Case

The gluon and singlet are coupled to each other so that they are described by

$$\Sigma_{\alpha}(t) = \sum_{\beta=0}^{\alpha} [M_{\alpha\beta}^{qq}(t) \Sigma_{\beta}(t_0) + M_{\alpha\beta}^{qg}(t) g_{\beta}(t_0)] \quad (4.33)$$

$$\text{with } M_{\alpha\beta}^{qq}(t_0) = \delta_{\alpha\beta}, \quad M_{\alpha\beta}^{qg}(t_0) = 0,$$

$$g_{\alpha}(t) = \sum_{\beta=0}^{\alpha} [M_{\alpha\beta}^{gq}(t) \Sigma_{\beta}(t_0) + M_{\alpha\beta}^{gg}(t) g_{\beta}(t_0)] \quad (4.34)$$

$$\text{with } M_{\alpha\beta}^{gq}(t_0) = 0, \quad M_{\alpha\beta}^{gg}(t_0) = \delta_{\alpha\beta}.$$

Inserting these expressions into the coupled DGLAP equations, without marking all the explicit t dependencies, yields

$$\sum_{\beta=0}^{\alpha} \frac{\partial}{\partial t} \begin{pmatrix} M_{\alpha\beta}^{qq} & M_{\alpha\beta}^{qg} \\ M_{\alpha\beta}^{gq} & M_{\alpha\beta}^{gg} \end{pmatrix} \begin{pmatrix} \Sigma_{\beta}(t_0) \\ g_{\beta}(t_0) \end{pmatrix} = a_s \sum_{\beta=0}^{\alpha} \sum_{\gamma=0}^{\beta} \begin{pmatrix} P_{\alpha\beta}^{qq} & P_{\alpha\beta}^{qg} \\ P_{\alpha\beta}^{gq} & P_{\alpha\beta}^{gg} \end{pmatrix} \begin{pmatrix} M_{\beta\gamma}^{qq} & M_{\beta\gamma}^{qg} \\ M_{\beta\gamma}^{gq} & M_{\beta\gamma}^{gg} \end{pmatrix} \begin{pmatrix} \Sigma_{\gamma}(t_0) \\ g_{\gamma}(t_0) \end{pmatrix}. \quad (4.35)$$

The result of the same rearrangement of the sums as in the uncoupled case is

$$\begin{aligned} \sum_{\beta=0}^{\alpha} \frac{\partial}{\partial t} \begin{pmatrix} M_{\alpha\beta}^{qq} & M_{\alpha\beta}^{qg} \\ M_{\alpha\beta}^{gq} & M_{\alpha\beta}^{gg} \end{pmatrix} \begin{pmatrix} \Sigma_{\beta}(t_0) \\ g_{\beta}(t_0) \end{pmatrix} &= a_s \sum_{\beta=0}^{\alpha} \sum_{\gamma=\beta}^{\alpha} \begin{pmatrix} P_{\alpha\gamma}^{qq} & P_{\alpha\gamma}^{qg} \\ P_{\alpha\gamma}^{gq} & P_{\alpha\gamma}^{gg} \end{pmatrix} \begin{pmatrix} M_{\gamma\beta}^{qq} & M_{\gamma\beta}^{qg} \\ M_{\gamma\beta}^{gq} & M_{\gamma\beta}^{gg} \end{pmatrix} \begin{pmatrix} \Sigma_{\beta}(t_0) \\ g_{\beta}(t_0) \end{pmatrix} \\ \Rightarrow \frac{\partial}{\partial t} \begin{pmatrix} M_{\alpha\beta}^{qq} & M_{\alpha\beta}^{qg} \\ M_{\alpha\beta}^{gq} & M_{\alpha\beta}^{gg} \end{pmatrix} &= a_s \sum_{\gamma=\beta}^{\alpha} \begin{pmatrix} P_{\alpha\gamma}^{qq} & P_{\alpha\gamma}^{qg} \\ P_{\alpha\gamma}^{gq} & P_{\alpha\gamma}^{gg} \end{pmatrix} \begin{pmatrix} M_{\gamma\beta}^{qq} & M_{\gamma\beta}^{qg} \\ M_{\gamma\beta}^{gq} & M_{\gamma\beta}^{gg} \end{pmatrix}. \end{aligned} \quad (4.36)$$

Then again, as in the uncoupled case, for $P_{\alpha\beta} = P_{\alpha-\beta}$ and $M_{\alpha\beta}(t_0) = M_{\alpha-\beta}(t_0)$ this can be reduced to:

$$\frac{\partial}{\partial t} \begin{pmatrix} M_{\alpha}^{qq} & M_{\alpha}^{qg} \\ M_{\alpha}^{gq} & M_{\alpha}^{gg} \end{pmatrix} = a_s \sum_{\gamma=0}^{\alpha} \begin{pmatrix} P_{\alpha-\gamma}^{qq} & P_{\alpha-\gamma}^{qg} \\ P_{\alpha-\gamma}^{gq} & P_{\alpha-\gamma}^{gg} \end{pmatrix} \begin{pmatrix} M_{\gamma}^{qq} & M_{\gamma}^{qg} \\ M_{\gamma}^{gq} & M_{\gamma}^{gg} \end{pmatrix} \quad (4.37)$$

$$= a_s \sum_{\gamma=0}^{\alpha} \begin{pmatrix} P_{\alpha-\gamma}^{qq} M_{\gamma}^{qq} + P_{\alpha-\gamma}^{qg} M_{\gamma}^{gq} & P_{\alpha-\gamma}^{qg} M_{\gamma}^{qg} + P_{\alpha-\gamma}^{gg} M_{\gamma}^{gg} \\ P_{\alpha-\gamma}^{gq} M_{\gamma}^{qq} + P_{\alpha-\gamma}^{gg} M_{\gamma}^{gq} & P_{\alpha-\gamma}^{gg} M_{\gamma}^{qg} + P_{\alpha-\gamma}^{gg} M_{\gamma}^{gg} \end{pmatrix}. \quad (4.38)$$

This rather long derivation using two evolution matrices for both the singlet and the gluon is in fact the same as in the non-singlet case, constructing a basis using N_x gluon and N_x singlet values with one $2N_x \times 2N_x$ evolution matrix. In this thesis the presented solution is used, because it avoids different array sizes for the evolution matrices in the code.

The useful result of the above derivations is that the evolution has been split in evolution operators and initial conditions. Thus, in a fitting procedure, the evolution matrices have to be calculated only once (!) and can afterwards be applied to the modified initial conditions.

4.2.5 Splitting Functions

Nevertheless, in any case the kernels Eq. (4.24) have to be calculated, but the splitting functions are singular for $x \rightarrow 1$ (plus distribution) or for $x \rightarrow 0$. The implementation has to take care of these special cases. In the code the decomposition of the splitting functions used in QCDNUM [9] has been adapted for reasons which will become obvious below:

$$P(x) = A(x) + [B(x)]_+ + R(x)[S(x)]_+ + K(x)\delta(1-x), \quad (4.39)$$

where $+$ indicates plus distributions.

All splitting functions will be evaluated in the integral $\int_{e^{-y_a}}^1 dz$ so that the singularities for $x \rightarrow 0$ can be neglected. Because each grid section will be integrated separately as described by Eq. (4.19) the main part $\int_{e^{-y_a}}^{e^{-y_1}}$ also involves none of these special cases. Therefore

$$I = \int_{e^{-y_a}}^{e^{-y_1}} P(z, t) w_\beta(y_\alpha + \ln z) dz \quad (4.40)$$

$$= \int_{e^{-y_a}}^{e^{-y_1}} [A(z) + B(z) + R(z)S(z)] w_\beta(y_\alpha + \ln z) dz. \quad (4.41)$$

The remaining part can be evaluated using the plus distribution

$$\begin{aligned} \int_{e^{-y_1}}^1 P(z, t) w_\beta(y_a + \ln z) dz &= \int_{e^{-y_1}}^1 [A(x) w_\beta(y_a + \ln z)] dz \\ &\quad \left[+ \int_{e^{-y_1}}^1 [w_\beta(y_a + \ln z) - w_\beta(y_a)] B(z) dz - w_\beta(y_a) \int_0^{e^{-y_1}} B(z) dz \right] \\ &\quad + \int_{e^{-y_1}}^1 [R(z) w_\beta(y_a + \ln z) - R(1) w_\beta(y_a)] S(z) dz - R(1) w_\beta(y_a) \int_0^{e^{-y_1}} S(z) dz \\ &\quad + K(1) w_\beta(y_a). \end{aligned} \quad (4.42)$$

In LO, only plus distributions of the type $S = \frac{1}{(1-x)_+}$ occur, so that the expression can be simplified to

$$\begin{aligned} \int_{e^{-y_1}}^1 P(z, t) w_\beta(y_a + \ln z) dz &= \int_{e^{-y_1}}^1 [A(x) w_\beta(y_a + \ln z)] dz \\ &+ \int_{e^{-y_1}}^1 \frac{R(z) w_\beta(y_a + \ln z) - R(1) w_\beta(y_a)}{1 - z} dz + R(1) w_\beta(y_a) \ln(1 - e^{-y_1}) \\ &+ K(1) w_\beta(y_a). \end{aligned} \quad (4.43)$$

This again can be rearranged to

$$\begin{aligned} \int_{e^{-y_1}}^1 P(z, t) w_\beta(y_a + \ln z) dz &= \int_{e^{-y_1}}^1 A(x) w_\beta(y_a + \ln z) \\ &+ \frac{R(z) w_\beta(y_a + \ln z) - R(1) w_\beta(y_a)}{1 - z} dz + [R(1) \ln(1 - e^{-y_1}) + K(1)] w_\beta(y_a). \end{aligned} \quad (4.44)$$

There is still a singularity in the fraction of the integrand for $z \rightarrow 1$, but since both the numerator and the denominator vanish, this case can be handled by using L'Hôpital's rule

$$\lim_{z \rightarrow 1} \frac{R(z) w_\beta(y_a + \ln z) - R(1) w_\beta(y_a)}{1 - z} = - \left[R'(1) w_\beta(y_a) + R(1) \frac{d}{dz} w_\beta(y_a + \ln z) \right]. \quad (4.45)$$

Using the separation of the weights in normalization and x dependent part, the derivative of the weight is

$$\frac{d}{dz} w_\beta(y_a + \ln z) \stackrel{z \rightarrow 1}{\equiv} w n(\beta') \frac{(-1)}{\delta y} \begin{cases} \prod_{i=1}^n (-i), & \text{for } \beta' = 0, \\ 0, & \text{else,} \end{cases} \quad (4.46)$$

where $\beta' = \beta - \alpha$ and it is implied that the interpolation order is greater than one.

Leading Order Splitting Functions

The LO splitting functions and their representations in the above scheme are

$$P_{qq}^{(0)}(x) = \frac{4}{3} \left[\frac{1+x^2}{(1-x)_+} + \frac{3}{2} \delta(1-x) \right], \quad (4.47)$$

$$R_{qq}^{(0)} = \frac{4}{3}(1+x^2), \quad S_{qq}^{(0)} = \frac{1}{1-x}, \quad K_{qq}^{(0)} = 2, \quad (4.48)$$

$$P_{qg}^{(0)}(x) = \frac{1}{2}(x^2 + (1-x)^2), \quad (4.49)$$

$$A_{qg}^{(0)} = \frac{1}{2}(x^2 + (1-x)^2), \quad (4.50)$$

$$P_{gq}^{(0)}(x) = \frac{4}{3} \frac{1 + (1-x)^2}{x}, \quad (4.51)$$

$$A_{gq}^{(0)} = \frac{4}{3} \frac{1 + (1-x)^2}{x}, \quad (4.52)$$

$$P_{gg}^{(0)}(x) = 6 \left[\frac{1-x}{x} + x(1-x) + \frac{x}{1-x_+} \right] + \delta(1-x) \frac{33-2n_f}{6}, \quad (4.53)$$

$$A_{qq}^{(0)} = 6 \left[\frac{1-x}{x} + x(1-x) \right], \quad R_{gg}^{(0)} = 6x, \quad S_{qq}^{(0)} = \frac{1}{1-x}, \quad K_{qq}^{(0)} = \frac{33-2n_f}{6}. \quad (4.54)$$

And in the used evolution basis (see Section 3.1.3) they are given by

$$P^{\pm(0)} = P_{qq}^{(0)}, \quad P^{V(0)} = P_{qq}^{(0)}, \quad P_{qq}^{(0)} = P_{qq}^{(0)}, \quad P_{gg}^{(0)} = P_{gg}^{(0)}, \quad P_{qg}^{(0)} = 2n_f P_{qig}^{(0)}, \quad P_{gq}^{(0)} = P_{gqi}^{(0)}. \quad (4.55)$$

4.2.6 Strong Coupling

Furthermore, the value of the strong coupling at the scale t is required by DGLAP equations. It is primarily implemented in its differential form of Eq. (2.9). Indeed, at LO, an analytic expression exists, too, but having higher orders in mind, the differential form is the default implementation. In this form, the strong coupling is first evolved from its value at a given scale, e.g. M_Z , down (up) to the initial scale of the PDFs using a precise integration procedure, i.e. adaptive Burlisch–Stoer integration [40, Chapter 17]. For congruence with the PDFs, the strong coupling is subsequently evolved in parallel with the PDFs using their integration procedure, which, at the moment, is fourth-order Runge–Kutta integration.

4.2.7 Matching Conditions

If the evolution is performed using a VFN scheme, both the strong coupling and the PDFs matching conditions have to be evaluated. Neglecting higher-order contributions, they are all continuous at the heavy quark thresholds, i.e. the pole masses (see Section 3.1.4). But since this concerns only the physical PDFs and not the evolution basis, still the new basis contributions have to be determined. At the threshold t_{hf} they are, at LO and at NLO, of type $hf^+(t_{\text{hf}}) = -u_p = -(u + \bar{u})$ and $hf^-(t_{\text{hf}}) = -u_m = -(u - \bar{u})$. The quantities u_p and u_m are calculated from the evolution basis using the same procedure as for the basis transformation (see Section 3.1.3). Actually, u_m is a non-singlet basis so that it can simply be evolved from the beginning, but since that is not possible for u_p , both are treated in the same way, at the moment.

4.2.8 General Code Structure

In sum, the general code contains the following steps.

1. The grid variables, flavour thresholds etc. are calculated (C++ object: class *CORE*).
2. The grid at the initial scale (stored in class *CORE*) is filled using a user-chosen initial parametrization (C++ object: class *IntialParametrization*<Name>; method: *fillGridatInputScale*).
3. The main evolution procedure is called (C++ object: class *EvolutionInX*; method: *fillGrid*)
 - 3.1 The initial PDFs are transformed to the evolution basis.
 - 3.2 The weight normalizations are calculated.
 - 3.2 The strong coupling at initial scale is calculated using adaptive Burlisch–Stoer integration [40, Chapter 17].
 - 3.3 The evolution kernels are calculated using an adaptive thirteen-point Gauss–Lobatto method with a Kronrod extension [40, Webnote 4].
 - 3.4 The PDFs or the evolution matrices are evolved.
 - 3.4.1 For direct PDF evolution: the PDFs and the strong coupling are jointly evolved up to the final scale or next threshold using fourth-order Runge–Kutta integration. If there is a threshold, the matching conditions are evaluated and the number of active flavours is increased by one for both, i.e. in the splitting functions and in the strong coupling beta function. Afterwards the evolution is continued using the extended basis.
 - 3.4.2 For matrix evolution: the evolution matrices and the strong coupling are jointly evolved up to the final scale or next threshold using fourth-order Runge–Kutta integration. If there is a threshold, the matching conditions are evaluated and the number of active flavours is increased by one for both, i.e. in the splitting functions and in the strong coupling beta function. Afterwards the evolution is continued using the extended basis. Finally, the evolution matrices are applied to the PDFs at initial scale.
 - 3.5 The evolved PDFs are transformed back to the physical basis.
4. Finally an output procedure can be implemented.

4.3 The Mellin-Space Evolution

The Mellin-space evolution is based on the well established QCD-PEGASUS program [44]. Nevertheless, as in the case of the x -space evolution, some modifications and optimizations have been implemented along with the object-orientated implementation in C++.

In principal, the implementation in Mellin-space is straightforward. Since all expressions are given analytically, numerical challenges occurs especially in the infinite

sums, in initial parametrizations, and in the inverse Mellin transformation. Since the Mellin PDFs are generally intended to be used in Mellin-space applications, normally the PDFs need not be inverted. Nevertheless, as the aim is to develop a flexible toolkit, the program is kept as generally as possible, and an inversion is therefore also implemented.

Furthermore, there are two main approaches to a numerical implementation. On the one hand, all the calculations can be performed with fixed (or user specified) grid *values* as in x -space. For instance, this method is faster if a whole grid in x -space is required, since the inversion can be optimized in view of the chosen grid points and the evaluations of functions are reduced to a minimum. QCD-PEGASUS has adopted this procedure.

On the other hand, C++ offers the possibility to directly implement the analytic *functions* and products of functions using functors. This way, the implementation almost directly reflects the theory. Besides this, the advantages of this approach are a fast evaluation of single PDF values in Mellin-space and great flexibility. The drawback of this procedure is that functions will unnecessarily be evaluated multiple times in a full PDF inversion to the x -space, expanding computation time significantly. Nevertheless, with little additional effort, a method using *values* can be derived from this approach.

Therefore, at present, the second way is implemented because of its greater flexibility. In future, it is desirable to implement both options depending on the requirements. Below, the main numerical challenges are discussed.

4.3.1 Initial Parametrization

In principal, the Mellin-space method applies to all kinds of initial PDFs such as functions or grid values in x -space. However, if the initial PDFs cannot be transformed analytically to Mellin-space, great numerical efforts are required to perform a numerical Mellin transformation. In the future, this option should be implemented too, but until now, the code depends on an analytic initial parametrization in Mellin-space. This is of course available for the x -space initial parametrization Eq. (2.56), but not for the CTEQ6 basis Eq. (2.55). Therefore, to compare both x -space and Mellin-space evolutions, the first one is preferable. At the moment, this and an extended version of it are implemented [44]. The first also equals the PDF benchmarks from [24] and is therefore particularly appropriate for testing and comparison:

$$x f_q(x, Q_0^2) = N_q p_{q,1} x^{p_{q,2}} (1-x)^{p_{q,3}} [1 + p_{q,5} x^{p_{q,4}} + p_{q,6} x], \quad (4.56)$$

where q indicates the basis and is often used for: $u-\bar{u}$, $d-\bar{d}$, $2(\bar{u}+\bar{d})$, $\bar{d}-\bar{u}$, g , $s\pm\bar{s}$. The moments of these functions are given in terms of Euler's beta functions, which itself can be represented as gamma functions for $\text{Re } m, n > 0$:

$$B(m, n) = \frac{\Gamma(m)\Gamma(n)}{\Gamma(m+n)} = \int_0^1 du u^{m-1} (1-u)^{n-1}. \quad (4.57)$$

The gamma function can be analytically continued to the complex numbers except for non-positive integers. Its arguments can be shifted to the real positive numbers

with the functional equation

$$\Gamma(m+1) = m\Gamma(m). \quad (4.58)$$

Therefore, in Mellin-space, the first parametrization Eq. (4.56) is given by

$$\begin{aligned} f_q(Q_0^2, n) &= N_g p_{q,1} \left[\int_0^1 dx x^{p_{q,2}+n-2} (1-x)^{p_{q,3}} + p_{q,5} \int_0^1 dx x^{p_{q,2}+p_{q,4}+n-2} (1-x)^{p_{q,3}} \right. \\ &\quad \left. + p_{q,6} \int_0^1 dx x^{p_{q,2}+n-1} (1-x)^{p_{q,3}} \right] \\ &= N_q p_{q,1} \left[B(p_{q,2}+n-1, p_{q,3}+1) + p_{q,5} B(p_{q,2}+p_{q,4}+n-1, p_{q,3}+1) \right. \\ &\quad \left. + p_{q,6} B(p_{q,2}+n, p_{q,3}+1) \right] \end{aligned} \quad (4.59)$$

$$\begin{aligned} &= N_q p_{q,1} \left[B(p_{q,2}+n-1, p_{q,3}+1) \left[1 + \frac{p_{q,6}(p_{q,2}+n-1)}{p_{q,2}+p_{q,3}+n} \right] \right. \\ &\quad \left. + p_{q,5} B(p_{q,2}+p_{q,4}+n-1, p_{q,3}+1) \right], \end{aligned} \quad (4.60)$$

so that only two beta functions have to be evaluated. The normalization factors N_q have to be calculated for a given initial distribution using the sum rules described in Chapter 2.7.

Basis transformations in Mellin-space are generally performed via a linear combination of the functions of the initial PDFs that are for example in the physical basis. Since *functions/functors* are used in the current implementation, this leads to superfluous computations (see below). In the future, it might be more efficient to implement the PDFs directly in the evolution basis or switching at least at this point to the *value* approach, fixing the PDFs at initial scale by a fixed grid and interpolation scheme.

Euler's Beta Function and Gamma Function

The beta function is implemented with the help of the gamma functions as described in Eq. (4.57). The gamma functions are calculated using the Lanczos expansion and algorithm described in Numerical Recipes [40, Chapter 6.1]:

$$\Gamma(z+1) = \left(z + \gamma + \frac{1}{2} \right)^{z+\frac{1}{2}} e^{-(z+\gamma+\frac{1}{2})} \sqrt{2\pi} \left[c_0 + \frac{c_1}{z+1} + \cdots + \frac{c_n}{z+N} + \epsilon \right] \quad (4.61)$$

with specific calculated γ and c 's, and an error smaller than $|\epsilon| < 10^{-15}$ if terms up to $N = 14$ are considered. This approximation applies for the complex gamma function everywhere in the complex plane with $\text{Re } z > 0$ and can be shifted to negative values using Eq. (4.58).

4.3.2 Anomalous Dimensions

Moreover, the anomalous dimensions in the complex Mellin-space are required. They have been easily derived by transforming the splitting functions as described in Section 3.2.3. However, at LO, this includes harmonic sums $S_1(n)$, which have to be continued in the complex plane for a Mellin inversion. Their analytic continuation is given by

$$S_1(n) = \psi(n) + \gamma_E, \quad (4.62)$$

where ψ is the digamma function and γ_E is the Euler–Mascheroni constant [4, 8, 25].

Digamma Function

The digamma function is implemented using the asymptotic series expansion [1], valid for $|z| \rightarrow \infty$ and $\arg(z) < \pi$:

$$\psi(z) = \ln(z) - \frac{1}{2z} - \frac{1}{12z^2} + \frac{1}{120z^4} - \frac{1}{252z^6} + \frac{1}{250z^8} + \mathcal{O}\left(\frac{1}{z^{10}}\right), \quad (4.63)$$

which can easily be expanded to higher order terms. The second condition concerning the argument will be automatically fulfilled because of the chosen contour of the inverse Mellin transformation (see Section 4.3.5). The first condition, $|z| \rightarrow \infty$, can be handled using the recurrence relation for the digamma function

$$\psi(z+1) = \psi(z) + \frac{1}{z}, \quad (4.64)$$

and requiring for example that $\text{Re } z \geq 10$ (see [25, Appendix A]).

4.3.3 Strong Coupling

In addition, like in x -space, multiple evaluations of the strong coupling are required by the solution (see e.g. Eq. 3.47). Unlike in x -space, in this approach the analytic implementation of the coupling – if possible – is preferable. Therefore, at LO, the analytic expression for the coupling is used, but with little additional effort, Burlisch–Stoer integration of the differential definition can be implemented.

4.3.4 Matching Conditions

Again, if the program is used in VFN scheme, matching conditions have to be evaluated. Their implementation equals the implementation in the x -space solution and follows Section 3.1.3 with the difference that now functions and not values are combined. And since at threshold, $t_{\text{hf}}, hf^-(t_{\text{hf}}) = -u_{\text{m}} = -(u - \bar{u})$ is a non-singlet, it can be directly evolved from the beginning. This option is also available in the code to speed up a PDF inversion.

4.3.5 Inverse Mellin Transformation

The final inverse Mellin transformation is at present just a test option for the comparison of results with respect to x -space PDFs, and is also only designed for PDF inversion. The method used is adopted from QCD-PEGASUS [26, 44] and modifies the general inversion contour of Eq. (3.81) by changing the argument from $\pi/2$ to $(3/4)\pi$ (see Fig. 4.2).

For a general contour the inversion is given by

$$xf(x) = \frac{1}{2\pi i} \left[\int_0^\infty dz \left[e^{i\phi} x^{1-c-z \exp(i\phi)} f(n = c + ze^{i\phi}) \right] + \int_\infty^0 dz \left[e^{-i\phi} x^{1-c-z \exp(-i\phi)} f(n = c + ze^{-i\phi}) \right] \right]. \quad (4.65)$$

Further, as the PDFs and splitting functions are real in x -space, it follows that $f^*(n) = f(n^*)$. And because of $v - v^* = 2i \cdot \text{Im}(v) \forall v \in \mathbb{C}$,

$$xf(x) = \frac{1}{\pi} \int_0^\infty dz \text{Im} \left[e^{i\phi} x^{1-c-z \exp(i\phi)} f(n = c + ze^{i\phi}) \right], \quad (4.66)$$

with c larger than the the rightmost singularity and $\phi > \pi$. As already mentioned, the anomalous dimensions are singular for $n = 1$. Using the above initial conditions with typical parameter values as described in [24] (see Eq. 5.2) yields no further singularities for $n > 1$. Therefore $c > 1$ has to be chosen, and the choice $c = 1.9$ from QCD-PEGASUS [44] was adopted. In principle, numerical tests of these values should be performed to obtain an inversion for small and for high x as fast and as accurate as possible.

If $\phi > \pi/2$, there is exponential damping because of the term proportional $\exp\{z \log(1/x) \cos(\phi)\}$ which allows for a smaller upper limit in the integration depending on the x value. Again, in general, numerical tests have to be performed to optimize this cut. Currently, the following values from QCD-PEGASUS are taken:

$$\begin{aligned} x < 0.01 : \quad z &< 5, \\ 0.01 \leq x < 0.3 : \quad z &< 14, \\ 0.3 \leq x < 0.7 : \quad z &< 32, \\ x \geq 0.7 : \quad z &< 80. \end{aligned} \quad (4.67)$$

In contrast to fixed grid *values*, the implementation of *functions* allows for the use of an adaptive integration procedure. Currently an adaptive thirteen-point Gauss-Lobatto method with Kronrod extension [40, Webnote 4] is used.

4.3.6 Number of Evaluations

A short sketch of the evaluations for a full x -grid shows the differences between both previously described approaches using *fixed values* or *functions*. In the inversion for each $a_s(Q)$, each flavour, and for each in the adaptive procedure chosen value of n the LO solutions, eigenvalue decompositions, basis transformations, and matching

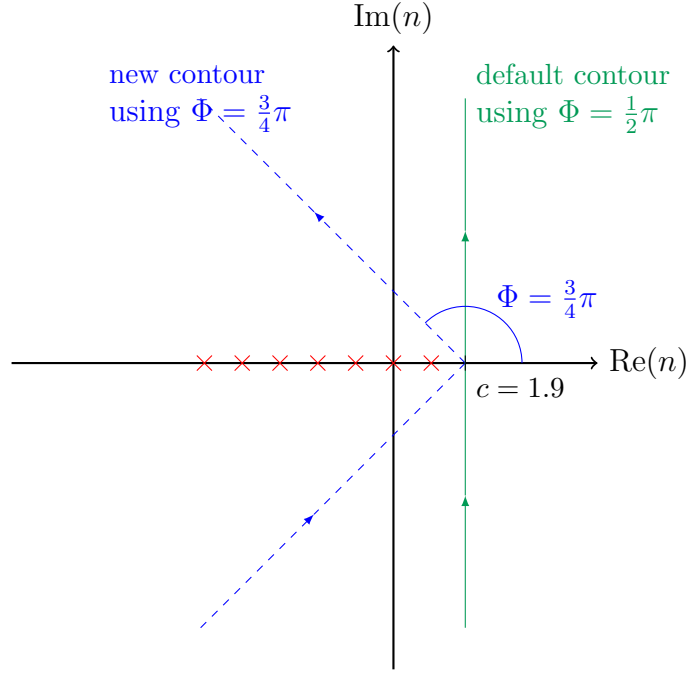


Fig. 4.2: Chosen and default integration contour for the inverse Mellin transformation described in Section 4.3.5. Crosses indicate the singularities of the Mellin-space function.

conditions have to be evaluated. Thus, there are a lot of redundant calculations. For example, the LO solutions $\mathbf{L}$ for all non-singlet PDFs at a given n and a_s are identical, but since the adaptive procedure can choose (slightly) different values of n in the integration, each of them has to be calculated anew. Fixing instead the supporting points of the inversion-integration, it is enough to evaluate the PDFs and (matrix-)solutions only once at the expense of adaptive integration and flexibility. Thus, an implementation of both solutions – if x -grids by Mellin methods are required – is a future objective.

4.3.7 General Code Structure

At present, the Mellin code is not designed for a specific application but besides single Mellin moments the provided methods can be used to create either an x -grid or a grid of Mellin moments (see the numerical analysis in Chapter 5). The x -grid solution proceeds as follows:

1. The grid variables, flavour thresholds etc. are calculated (C++ object: class *CORE*).
2. The x - Q -grid at initial scale (stored in class *CORE*) is filled using a user chosen initial parametrization (C++ object: class *IntialParametrization<Name>*; method: *fillGridatInputScale*). Both the analytic forms of the x -space and the Mellin-space parametrizations are implemented, so that no transformation is necessary at this step. Furthermore, the parametrizations in both spaces

are handed over to the evolution method using function pointers in the class *CORE*.

3. The main evolution procedure is called (C++ object: class *EvolutionInX*; method: *fillGrid*)
 - 3.1 The new basis functions are constructed using the initial parametrization functions.
 - 3.2 The strong coupling at initial scale and at thresholds (recall that a_s is the integration variable) is calculated using the analytic LO expression or Burlisch–Stoer integration [40, Chapter 17].
 - 3.3 Functions returning the PDF Mellin moments are constructed from the LO solutions and the initial conditions.
 - 3.4 The inversion is called for each flavour, for each x , and for each a_s (or Q^2), calling itself for each Mellin n the LO solution $\mathbf{L}$ (its eigenvalue decomposition etc.) and using adaptive thirteen-point Gauss–Lobatto integration with a Kronrod extension [40, Webnote 4].
 - 3.5 The evolved PDFs are transformed back to the physical basis.
4. Finally an output procedure can be implemented.

At alternative implementation to produce a grid in Mellin-space can be easily obtained. The Mellin moments can be directly returned from the defined functions under point 3.3. At least, just a basis transformation has to be performed.

4.4 Overview and Comparison of Both Methods

At this stage, a summary of the approximations used in each method and of the possible future optimizations is possible. In addition, a first comparison between both methods is reasonable.

4.4.1 The x -Space Evolution

The following approximations and sources of uncertainty occur in the LO code of the x -space solution (i.e. besides numerical accuracy – all real numbers are implemented as *double*):

- The general discretizations of the PDFs and the convolutions, using Lagrange interpolation, are sources of deviations. Their accuracy depends on the sampling size in x and on the order of interpolation (see Chapter 5).
- The distribution of grid points equidistant in $y = \ln(1/x)$ implies large distances in the x -grid for large x . A possible optimization will be the implementation of subgrids.

- The interpolation scheme contributes to further deviations because of the unphysical interpolation points for $x > 1$ described in Section 4.2.2.
- The integration procedures: while the adaptive thirteen-point Gauss–Lobatto–Konrod integration for the splitting functions and the adaptive Burlisch–Stoer integration for the strong coupling constant at initial scale only contribute marginal deviations, the fourth-order Runge–Kutta integration includes a total accumulated error of $\mathcal{O}(\text{stepsize}^4)$ (see Chapter 5). Changing this last integration procedure to an adaptive method or increasing the stepsize will increase the precision. The higher computational efforts can be neglected in a fitting procedure, because the calculation of the evolution matrices is performed only once.
- The choice of equal factorization and renormalization scale is only negligible if all orders in perturbation theory are considered. Their variation can be used to estimate the error due to higher order contributions.

Beyond accuracy, some optimization of computation time and the development of additional tools are short-term objectives. In sum these are:

- Implementing subgrids in x , i.e. $y = \ln(1/x)$, for large x values.
- Implementing an interpolation scheme that is *centered* around the interpolation point and does not use unphysical PDF values for $x > 1$. Thus, the error in the above implementation can be estimated.
- Transforming the evolved evolution matrices instead of the PDFs to the physical basis. That way, a lot of computation time can be saved in a fitting procedure.

4.4.2 The Mellin-Space Evolution

In Mellin-space the following sources of deviation have to be used:

- (Neglecting higher-order terms in the series expansions of the beta/gamma function of the initial parametrization.)
- Neglecting of higher-order terms in the implementation of the digamma function in the anomalous dimensions. In NLO further approximations will be required (e.g. in the implementation of the dilogarithm).
- In higher-order solutions, uncertainties due to truncations of series expansions occur. In fact, irrespective of the distinction in three solutions all solutions have to truncate the infinite sum in the $\mathbf{U}$ matrices at a specified or sufficient high order.
- (The chosen Mellin inversion; especially the abort criterion for the integration contour.)

The items in brackets denote sources of deviations which are not due to the Mellin evolution itself. Here, the short-term objective is

- Implement *value*-orientated scheme to create fast x -grids.

4.4.3 Comparison

A detailed comparison of both code projects is not possible. Numerical comparisons would include computation time and accuracy. Particularly the former is only reasonable if a fixed goal for the use of both methods is formulated, which is not presently the case. However, some general remarks are possible: for a given initial parametrization in Mellin-space the Mellin evolution is more flexible, e.g. it does not depend on a any grid, and can return the PDF value at a given Mellin n at marginal computation time. In contrast, the x -space solution can be easily applied to multiple initial conditions but requires a sufficiently large sampling in x to evolve the PDFs. In general, the x -space solution is more computationally intensive, and the above list of approximations suggests more uncertainties in the case of the x -space evolution. Numerical results will be discussed in the next chapter.

Chapter 5

Numerical Results

In this chapter the numerical properties of the derived evolution methods will be analyzed. This includes primarily comparisons with respect to the well established tools HOPPET [42] and QCD-PEGASUS [44]. First, both methods will be tested separately and afterwards they will be compared.

5.1 Common Settings

The tests described in this chapter were all carried out using the settings mentioned below, which corresponds to the setup used for the reference results from the Les Houches meeting, 2001 [24].

The initial scale Q_0^2 of the input parametrization is

$$Q_0^2 = 2 \text{ GeV}^2, \quad (5.1)$$

and the parametrization of the PDF moments reads

$$\begin{aligned} xu_v(x, Q_0^2) &= x(u - \bar{u}) = 5.107200x^{0.8}(1-x)^3, \\ xd_v(x, Q_0^2) &= x(d - \bar{d}) = 3.064320x^{0.8}(1-x)^4, \\ xg(x, Q_0^2) &= 1.700000x^{-0.1}(1-x)^5, \\ x\bar{d}(x, Q_0^2) &= 0.1939875x^{-0.1}(1-x)^6, \\ x\bar{u}(x, Q_0^2) &= (1-x)x\bar{d}(x, Q_0^2), \\ xs(x, Q_0^2) &= x\bar{s}(x, Q_0^2) = 0.2x(\bar{u} + \bar{d})(x, Q_0^2). \end{aligned} \quad (5.2)$$

Furthermore, equality of the factorization and renormalization scale was assumed, the strong coupling constant was initialized by

$$\alpha_s(Q^2 = 2 \text{ GeV}^2) = 0.35, \quad (5.3)$$

and the heavy quark thresholds are defined by the pole masses

$$m_c = \sqrt{2} \text{ GeV}, \quad (5.4)$$

$$m_b = 4.5 \text{ GeV}, \quad (5.5)$$

$$m_t = 175 \text{ GeV}. \quad (5.6)$$

If not explicitly mentioned otherwise, all evolutions were run in LO mode and in VFN scheme.

5.2 The x -Space Evolution

The x -space solution depends in particular on the interpolation, the grid sampling size in x , and the Runge–Kutta integration stepsize. These parameters are tested first, in order to be able to estimate their influence, and to optimize the configuration. Then the results are compared to HOPPET.

5.2.1 Interpolation Order

The influence of the interpolation order has already been mentioned in the implementation. Especially for orders higher than eight, oscillations can be observed, as expected (see Section 4.2.1 or [40, p. 110 ff.]). In Fig. 5.1 a reference plot of the momentum PDF (xPDF) at order $n = 7$ for integration stepsize $h = 0.5$ (see below) is shown, and below the absolute values of the relative deviations from another calculation using one order less, i.e. $n = 6$. Significant deviations are present for large x because of the already mentioned low grid density for these values. Moreover, because of the convolution of the DGLAP equations, these large x deviations contribute to the xPDF for all smaller x . However, they vanish at smaller x due to the high absolute xPDF values for these x . Thus, for small x , i.e. a fine grid, despite of this large x influence, the interpolation scheme is of vanishing influence. Therefore, a higher sampling size N_x (a plot for a coarser grid is shown in Appendix A.1) or a subgrid for large x will improve the numerical stability. The latter is preferable as it requires less computation time.

A comparison of different interpolation orders for the xPDFs with the strongest deviations (i.e. for g , u , d) is shown in Fig. 5.2. The sum of squared deviations was chosen as indicator for the accuracy, summing over the least number of grid points ($N_x = 161$). From this figure it can be concluded that interpolation orders for $n \geq 4$ and $n \leq 8$ are of marginal difference in view of the *total* xPDF (e.g. for large x this fact holds not necessarily). In the currently implemented interpolation scheme using $\mathcal{O}(N_x + n)$ terms instead of a more general scheme with $\mathcal{O}(N_x \cdot n)$ (see Section 4.2.3), this choice is of little importance with regard to computation time. By default, $n = 6$ is chosen.

Moreover, the in Fig. 5.2 also depicted comparison between different sampling sizes in x for the same interpolation orders demonstrates again the improved accuracy for higher sampling sizes. It further suggests that the influence of the interpolation scheme depends only weakly on the sampling size.

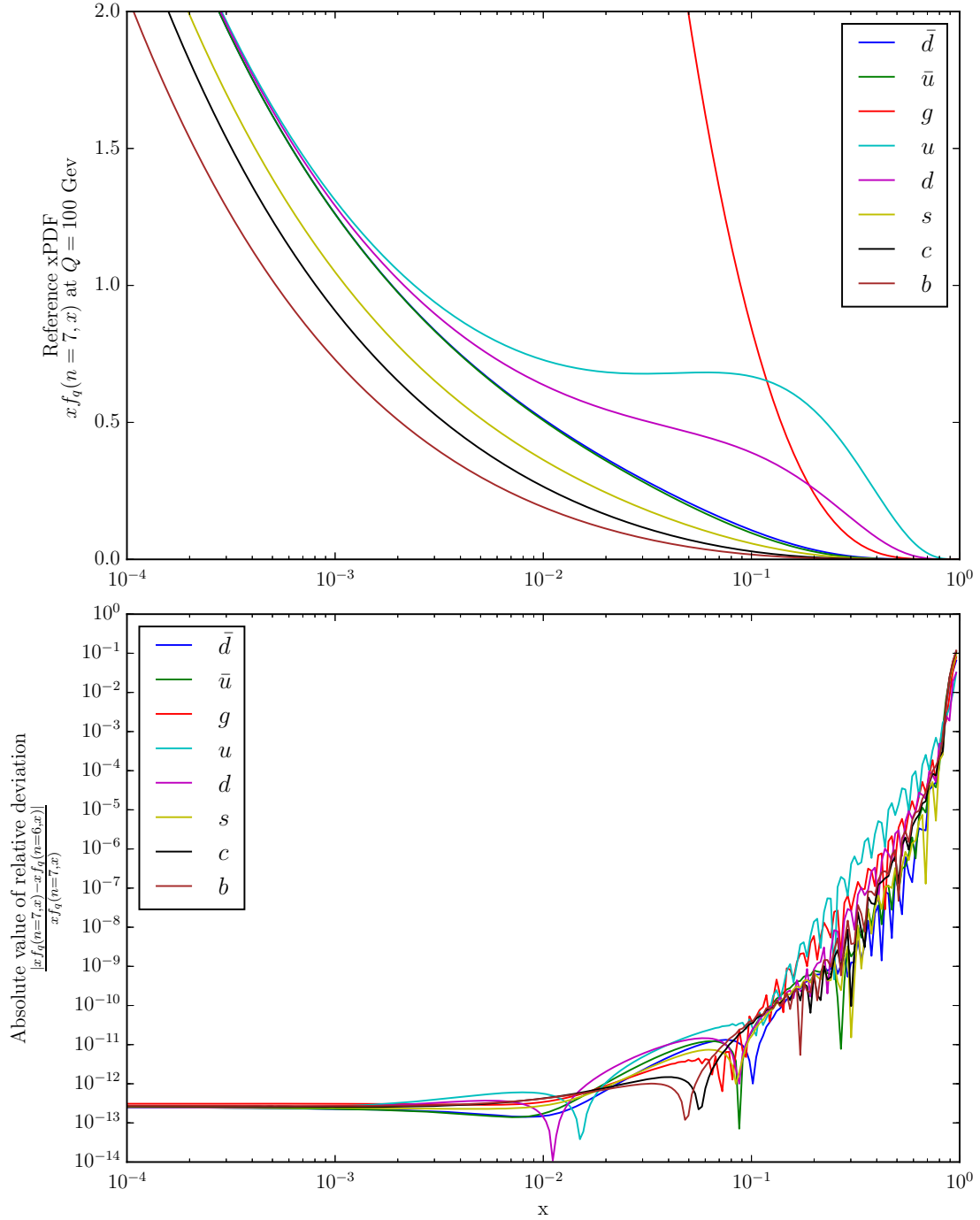


Fig. 5.1: Upper panel: reference xPDF using a sampling size in x of $N_x = 321$ and an interpolation order of $n = 7$ at scale $Q = 100$ GeV. Lower panel: absolute values of the relative deviations of a xPDF with interpolation order $n = 6$ from the reference xPDF.

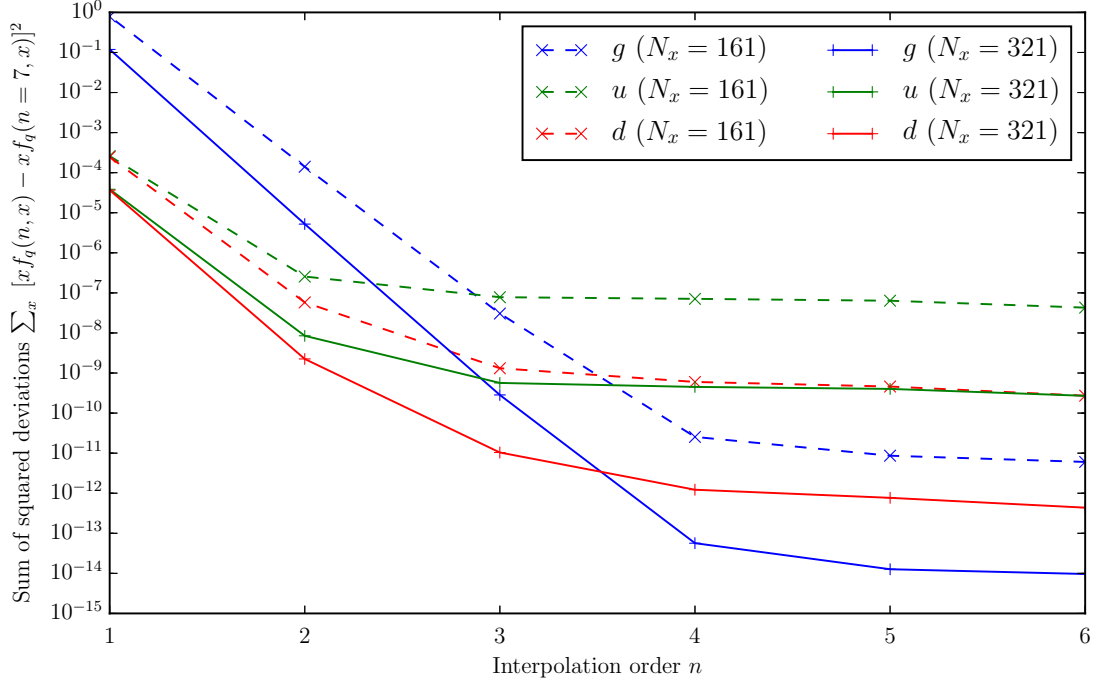


Fig. 5.2: Deviations of xPDFs with different interpolation orders from the reference xPDF with an interpolation order of $n = 7$ at scale $Q = 100$ GeV for sampling sizes in x of $N_x = 321$ and $N_x = 161$.

5.2.2 Sampling Size in x

Besides the interpolation order, the sampling size is of great influence as was shown in the last example. Since the computation time depends approximately linearly on the size of the sampling, a careful choice is required. For interpolation order $n = 6$ and Runge–Kutta integration stepsize $h = 0.05$ (see below) different sizes have been tested and their deviations with respect to the highest sampling size of $N_x = 1921$ points have been determined. The deviations for a low sampling size $N_x = 161$ are shown in Fig. 5.3. Again, for large x (> 0.1) deviations become relevant. This is not only due to the sampling size itself, but also due to the unphysical interpolation, using $x > 1$, that becomes more relevant in this region.

Fig. 5.4 shows for different sampling sizes the sum of the squared deviations, summing over the grid with the least sampling size of $N_x = 61$, from the reference xPDF using $N_x = 1921$. The linear relationship in the double logarithmic scale implies a power law behavior of the deviations depending on the sampling size. In general, a sampling size of $N_x = 161$ was chosen if not mentioned otherwise because of the computation time relevance of the sampling size (this is also true in the solution using evolution matrices!).

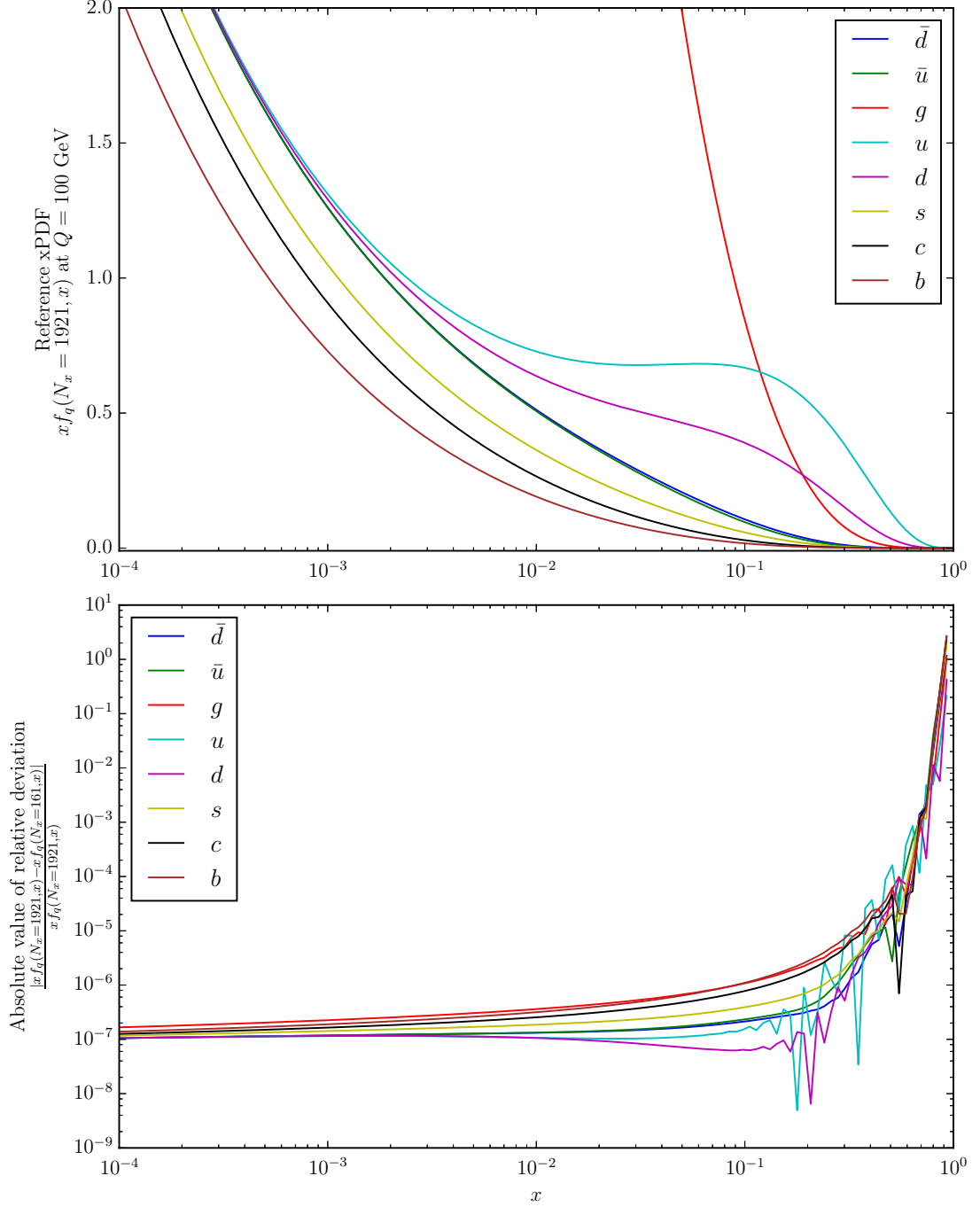


Fig. 5.3: Upper panel: reference xPDF using a sampling size in x of $N_x = 1921$ and a interpolation order $n = 6$ at scale $Q = 100$ GeV. Lower panel: absolute values of the relative deviations of a xPDF with a sampling size of $N_x = 161$ from the reference xPDF.

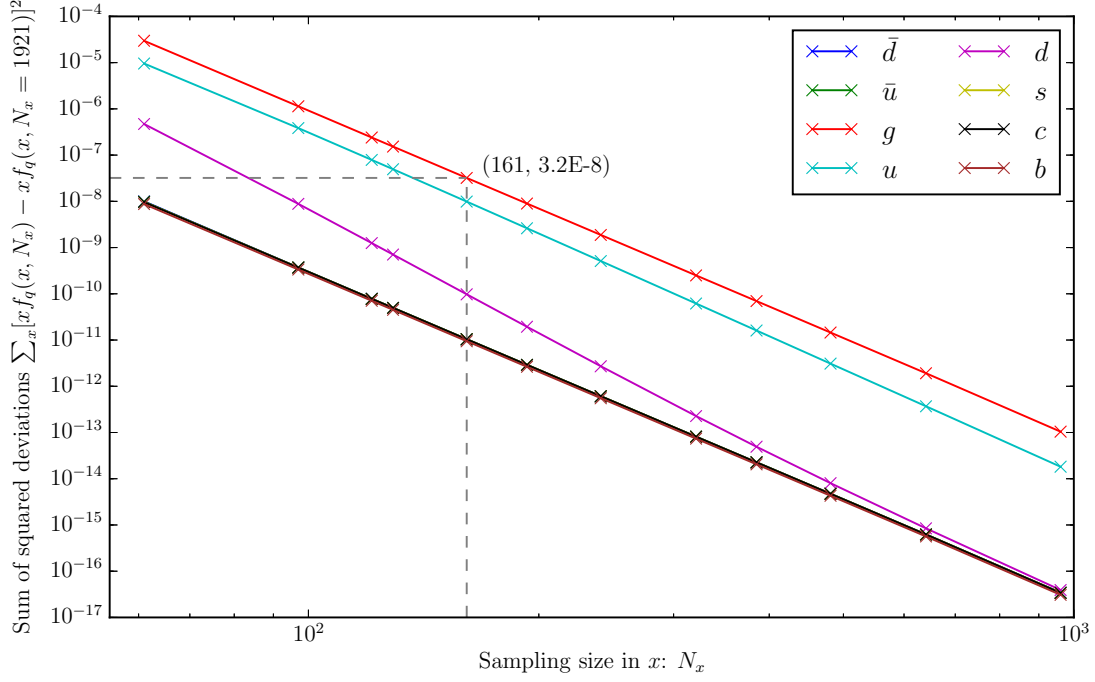


Fig. 5.4: Double logarithmic plot of the sum of squared deviations with respect to the reference xPDF with a sampling size in x of $N_x = 1921$, and an interpolation order $n = 7$ at scale $Q = 100$ GeV.

5.2.3 Integration Stepsize

The uncertainty of fourth-order Runge–Kutta integration can be well described by a global $\mathcal{O}(h^4)$ behavior, where h is the integration stepsize. Nevertheless, if the stepsize h is too small, round-off errors become relevant. In Fig. 5.6 the sum of the squared deviations shows clearly the expected $\mathcal{O}(h^8)$ behavior. The offset of the gluon deviations is due to its higher values. The deviation below 10^{-2} from the Runge–Kutta behavior is due to the limited precision of the data type *double*.

Fig. 5.5 depicts the absolute values of the relative deviations of a xPDF with a stepsize $h = 0.05$ from an evolution using $h = 10^{-4}$. Even at this large stepsize, deviations are already of order $\mathcal{O}(10^{-8})$ and thus beyond any significance. Therefore, a stepsize $h = 0.05$ is taken to be the default choice. Since the stepsize is only relevant for the one-time calculation of the evolution matrices, this choice is not of high relevance concerning computation time.

Furthermore, the Q -grid is related to the stepsize: if the integration stepsize is chosen larger than the Q -grid spacing, then the actual stepsize will be reduced and the choice of the Q -grid will become (again only one-time) computation time relevant.

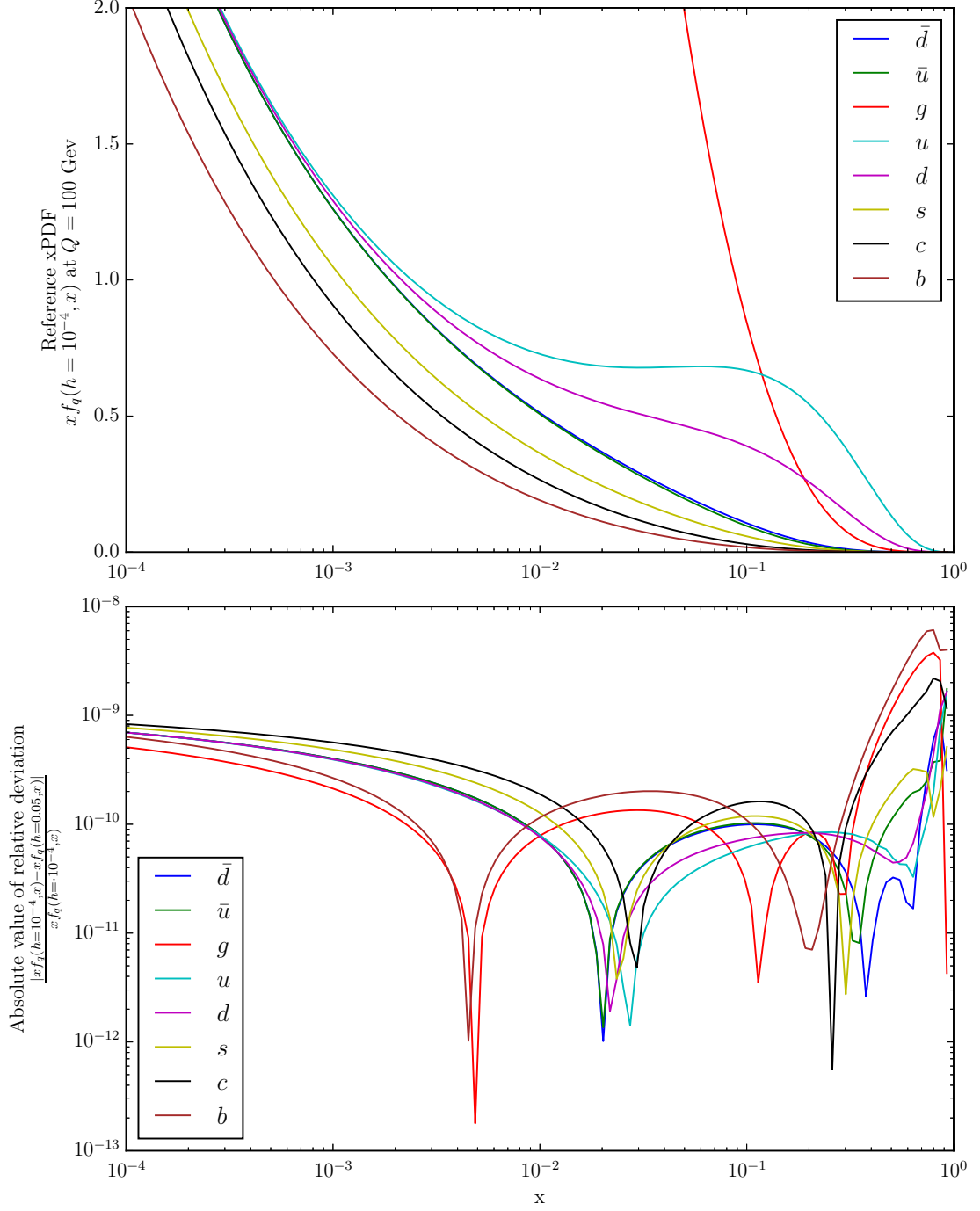


Fig. 5.5: Upper panel: reference xPDF using a stepsize $h = 10^{-4}$, an interpolation order $n = 6$, and a sampling size $N_x = 161$ at scale $Q = 100$ GeV. Lower panel: absolute values of the relative deviations of a xPDF with stepsize $h = 0.5$ from the reference xPDF.

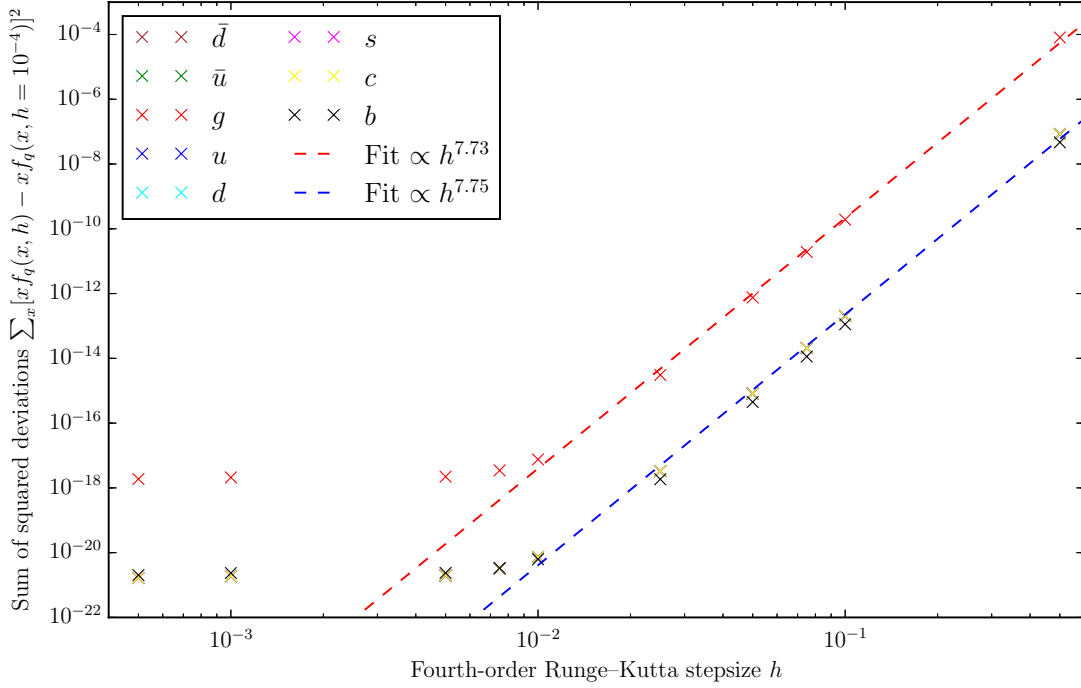


Fig. 5.6: Double logarithmic plot of the sum of squared deviations with respect to the reference xPDF using a stepsize $h = 10^{-4}$, an interpolation order $n = 6$, and a sampling size $N_x = 161$ at scale $Q = 100$ GeV.

5.2.4 Comparison to HOPPET

As already mentioned, the x -space solution is based on the HOPPET [42] toolkit. Therefore a direct comparison of both methods can be used to test the accuracy of the developed code. Despite the wide-ranging similarities, there are still differences in both codes to be considered. Whereas HOPPET evolves the PDFs using a variable $u = (\ln \ln Q^2 / \Lambda_{\text{QCD}}) / \beta_0$ (so that $du/d \ln Q^2 \simeq \alpha_s(Q^2)$), in the presented solution the integration variable is $t = \ln Q^2$. Further differences in the implementation are for example integration methods for the splitting kernels or a different implementation of the splitting functions. For the comparison, the default setting $du = 0.1$ of HOPPET has been modified to a higher precision $du = 0.01$, and the x -code stepsize was fixed to the default value of $h = 0.05$. Both programs used a sampling size $N_x = 161$ with equidistant gridpoints in the variable $y = \ln(1/x)$ and an interpolation order $n = 6$. The results are presented in Fig. 5.7. The deviations are again of order $\mathcal{O}(10^{-8})$ and therefore of the order expected from the Runge–Kutta methods. Thus, the developed code reproduces – with high accuracy – the output of the HOPPET program.

Nevertheless, there are still deviations for large x and equivalently for very small xPDF values. The main reason for these deviations are the different integration variables mention above, which lead to non-equivalent integration stepsizes. Therefore, increasing the integration stepsizes to $du = 0.001$ and $h = 0.01$ reduces the deviations further, as can be seen in the appendix (see Appendix A.2). Other possible

sources of deviations are the implementation of the kernels (integration procedure, splitting functions) or differences in the computation precision because of the small absolute values.

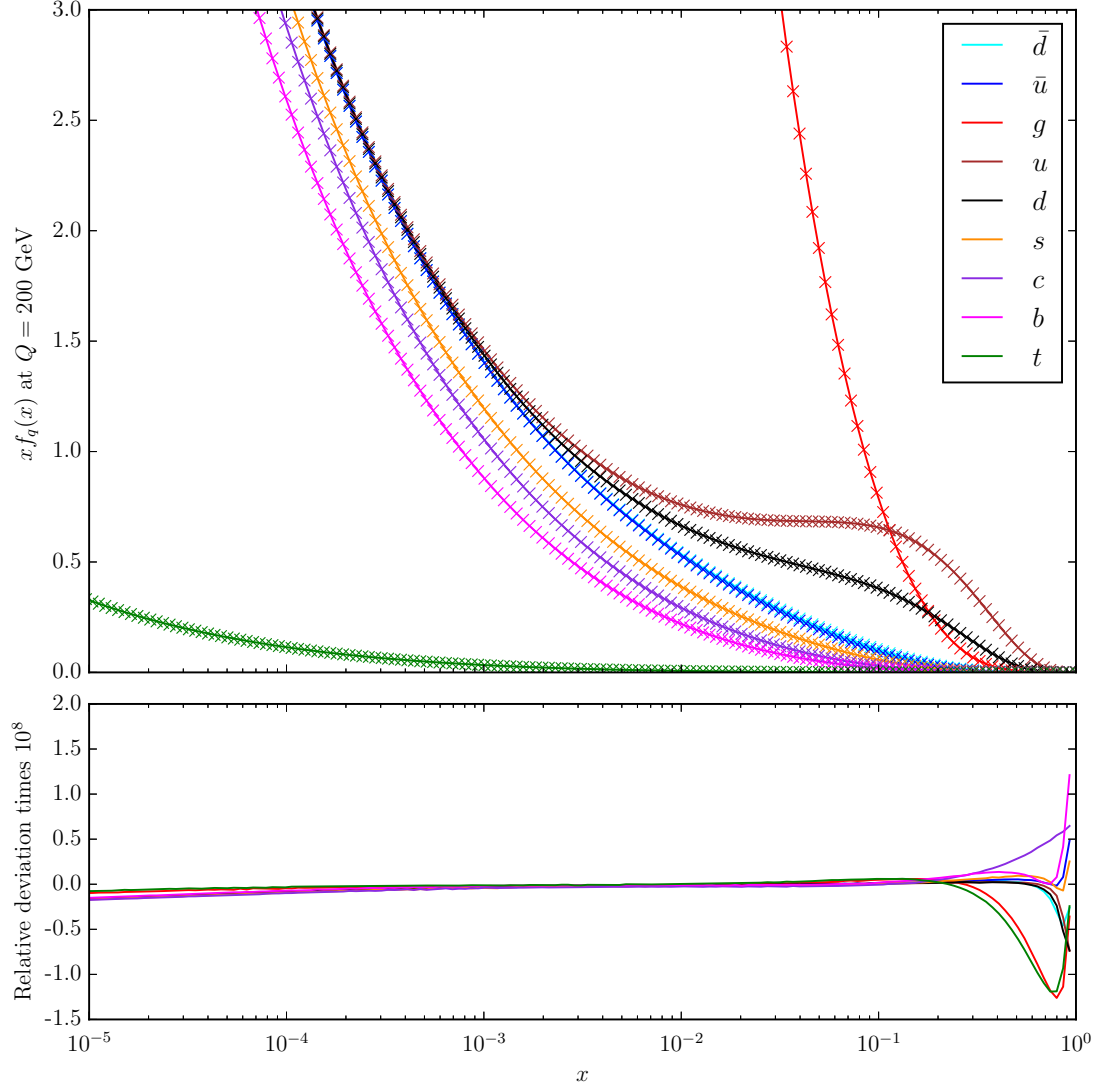


Fig. 5.7: Upper panel: comparison between the x -space code (crosses) and the HOPPET (lines) output at $Q = 200$ GeV using a sampling size $N_x = 161$, an interpolation order $n = 6$, and integration stepsizes $du = 0.01$ (HOPPET) and $h = 0.05$. Lower panel: relative deviations from the HOPPET results.

5.3 The Mellin-Space Evolution

5.3.1 Comparison to QCD-PEGASUS

The Mellin code can be tested in multiple ways. Leaving aside the inversion, the solutions are nearly analytic, so that no initial tests like in x -space have to be done. Therefore a comparison of Mellin moments should be almost exact. In QCD-PEGASUS the normalization of the initial conditions is calculated which is responsible for small deviations from the values given in Eq. (5.2). Therefore, an analogous routine for the QCD-PEGASUS normalization options $IMOMIN = 1$ and $ISSIMP = 0$ has been implemented in the Mellin code. Fig. 5.8 shows the comparison for the second Mellin moment at different scales Q (see Appendix A.3 for a comparison of the fifth Mellin moment $n = 5$). Again the deviations are of $\mathcal{O}(10^{-8})$. An investigation of the deviations revealed that in QCD-PEGASUS the value of the strong coupling at initial scale is only exactly stored up to this precision, so that the agreement is as exact as possible.

Beyond this, a comparison in x -space is interesting. The inversion depends on multiple ingredients: the chosen values for the angle ϕ , the crossing of the real axis c , the abort of the integration at large z , and the integration procedure. Each of these aspects has to be tested and carefully optimized. As the inversion is only of lesser interest at this stage of development, the specifications from QCD-PEGASUS are adopted. Thus, besides differences in the integration, both methods should create comparable results. Fig. 5.9 depicts that for $x \lesssim 0.6$ the relative deviations are of $\mathcal{O}(10^{-6})$, which is in rough agreement with the asserted accuracy of QCD-PEGASUS of five digits in $10^{-7} < x < 0.9$ [44, p. 13]. For higher x , significant deviations are observable. Besides the less precise strong coupling value in QCD-PEGASUS described above, the difference between the developed code and QCD-PEGASUS is probably also due to the integration procedure, which is fixed in the latter tool and adaptive in this tool (see below). Moreover, at large x values, the damping $\propto \ln(1/x)$ is small. Hence the truncation at small z in the inversion integration leads to deviations in both methods with respect to the exact/real solution.

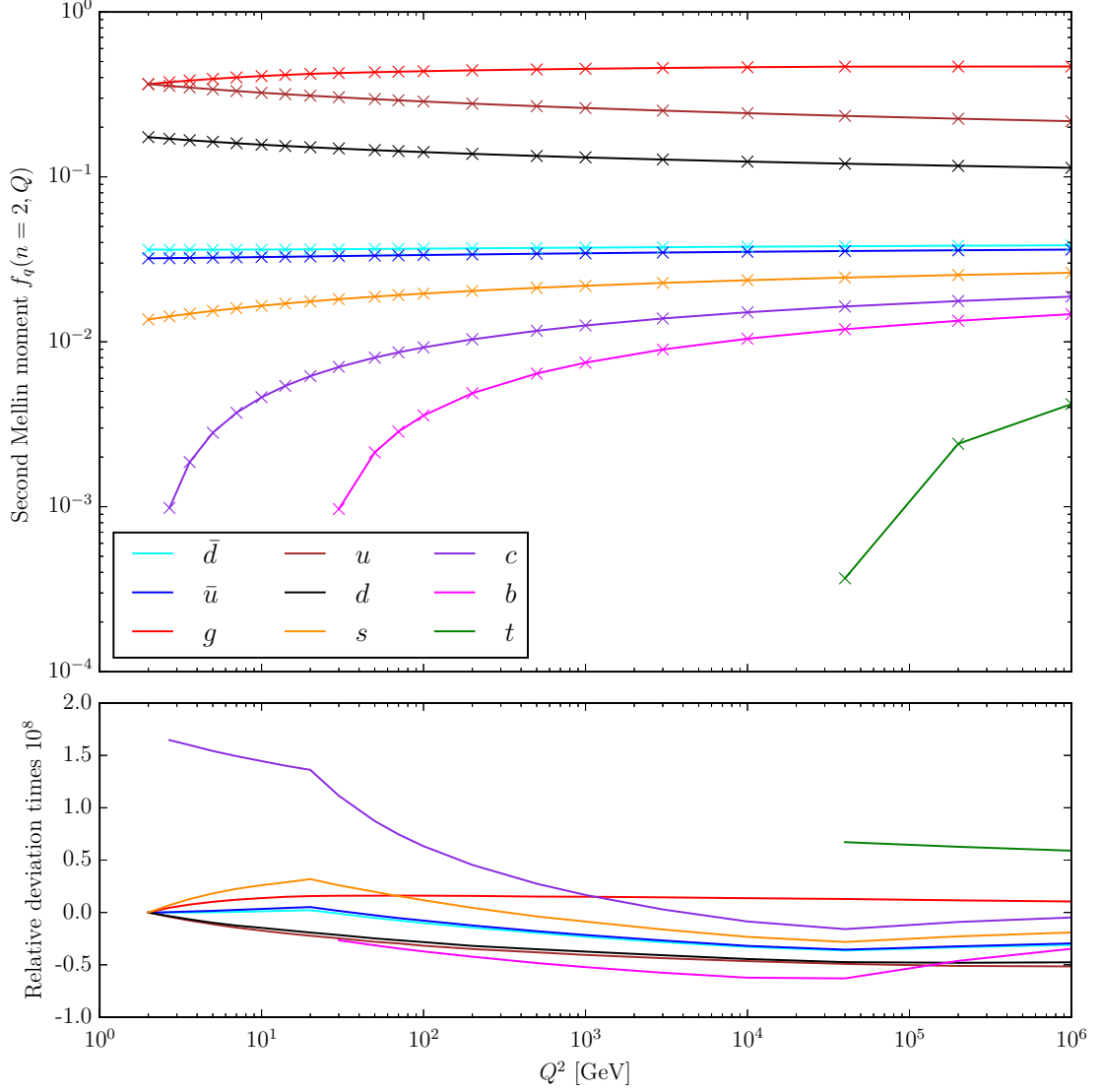


Fig. 5.8: Upper panel: comparison between the Mellin code (crosses) and the QCD-PEGASUS (lines) output of the second Mellin moment for different scales Q^2 using the QCD-PEGASUS normalization ($IMOMIN = 1$, $ISSIMP = 0$) of the initial conditions in both programs. Lower panel: relative deviations of the Mellin code results from the QCD-PEGASUS results.

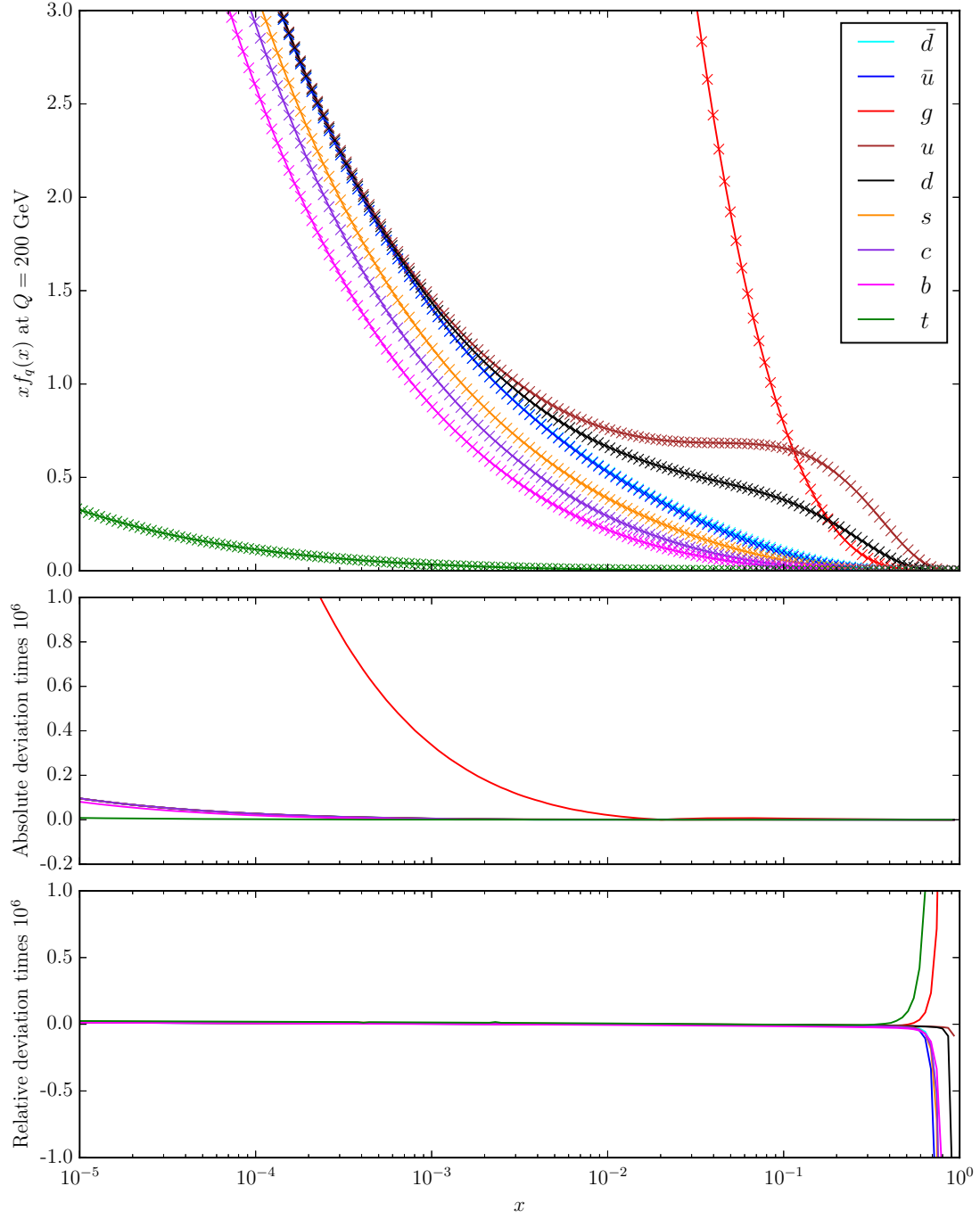


Fig. 5.9: Upper panel: comparison between the Mellin code (crosses) and the QCD-PEGASUS (lines) output at $Q = 200$ GeV using the QCD-PEGASUS normalization ($IMOMIN = 1$, $ISSIMP = 0$) for the initial conditions in both programs and using in the adaptive integration of the Mellin code a precision of 10^{-10} . Lower panels: absolute and relative deviations of the Mellin code results from the QCD-PEGASUS results.

5.3.2 Comparison to Benchmark Tables

The x -grid produced by the Mellin evolution can also be compared to the above mentioned benchmark results from Les Houches, 2001 [24]. Table 5.1 reproduces exactly the benchmarks results up to the stated precision. In contrast, the results of QCD-PEGASUS deviate for large x and are less accurate in this region as table 5.2 demonstrates. This supports the hypothesis that in QCD-PEGASUS the error for large x is underestimated and that the different integration procedures – adaptive or fixed values – are of non-negligible influence.

Table 5.1: Results of the Mellin code with the default options to reproduce the Les Houches [24] benchmark results for VFN at scale $Q = 100$ GeV. The notation is: $L_- = \bar{d} - \bar{u}$, $L_+ = \bar{d} + \bar{u}$, $q_+ = q + \bar{q}$. To keep the table concise the superscripted values in the PDF columns denote decimal powers.

x	xu_v	xd_v	xL_-	xL_+	xs_+	xc_+	xb_+	xg
10^{-7}	5.8771^{-5}	3.4963^{-5}	7.8233^{-7}	1.0181^{+2}	4.9815^{+1}	4.9088^{+1}	4.6071^{+1}	1.3272^{+3}
10^{-6}	3.3933^{-4}	2.0129^{-4}	5.1142^{-6}	5.1182^{+1}	2.4725^{+1}	2.4148^{+1}	2.2239^{+1}	6.0117^{+2}
10^{-5}	1.9006^{-3}	1.1229^{-3}	3.2249^{-5}	2.4693^{+1}	1.1659^{+1}	1.1201^{+1}	1.0037^{+1}	2.5282^{+2}
10^{-4}	1.0186^{-2}	5.9819^{-3}	1.9345^{-4}	1.1406^{+1}	5.1583^{+0}	4.7953^{+0}	4.1222^{+0}	9.6048^{+1}
10^{-3}	5.0893^{-2}	2.9576^{-2}	1.0730^{-3}	5.0424^{+0}	2.0973^{+0}	1.8147^{+0}	1.4582^{+0}	3.1333^{+1}
10^{-2}	2.2080^{-1}	1.2497^{-1}	4.9985^{-3}	2.0381^{+0}	7.2625^{-1}	5.3107^{-1}	3.8106^{-1}	7.7728^{+0}
10^{-1}	5.7166^{-1}	2.8334^{-1}	1.0428^{-2}	4.0496^{-1}	1.1596^{-1}	5.8288^{-2}	3.5056^{-2}	8.4358^{-1}
0.3	3.7597^{-1}	1.4044^{-1}	3.2629^{-3}	3.9592^{-2}	1.0363^{-2}	4.0740^{-3}	2.2039^{-3}	7.8026^{-2}
0.5	1.3284^{-1}	3.4802^{-2}	4.2031^{-4}	2.8066^{-3}	7.1707^{-4}	2.5958^{-4}	1.3522^{-4}	7.4719^{-3}
0.7	2.2643^{-2}	3.5134^{-3}	1.5468^{-5}	6.7201^{-5}	1.7278^{-5}	6.3958^{-6}	3.3996^{-6}	3.5241^{-4}
0.9	4.2047^{-4}	2.1529^{-5}	1.0635^{-8}	3.4998^{-8}	9.8394^{-9}	4.7330^{-9}	2.8903^{-9}	1.0307^{-6}

Table 5.2: Last row of QCD-PEGASUS results for the default options including deviations with respect to the Les Houches [24] benchmark results for VFN at scale $Q = 100$ GeV. The notation is: $L_- = \bar{d} - \bar{u}$, $L_+ = \bar{d} + \bar{u}$, $q_+ = q + \bar{q}$. To keep the table concise the superscripted values in the PDF columns denote decimal powers.

x	xu_v	xd_v	xL_-	xL_+	xs_+	xc_+	xb_+	xg
0.9	4.2047^{-4}	2.1529^{-5}	1.0635^{-8}	3.5020^{-8}	9.8449^{-9}	4.7349^{-9}	2.8916^{-9}	1.0306^{-6}

5.4 Comparison of the Evolution Methods

At the beginning the intend was formulated, to create a tool that is able to solve the DGLAP equations in both spaces. That way, for instance, they could be used in fitting procedures at the same time. Therefore, it is required that both evolution engines are equivalent in sufficiently high numerical limits. Until now, such a comparison can only be done by transforming the Mellin solutions back to the x -space. Thus, the comparison will strongly depend on the Mellin inversion.

A first comparison has been performed using the default options for the inversion and the x -space settings: $n = 6$, $N_x = 161$, $h = 0.05$. In Fig. 5.10 a good agreement

at small x but large deviations at large x are apparent. This is not surprising since both methods have shortcomings exactly in this region. Therefore multiple runs with different options have been performed. Optimizing all options, i.e. $N_x = 641$, $h = 0.001$, and $z_{\max} = 200$ for all $x > 0.3$, reduces the deviations approximately by two orders of magnitude, as Fig. 5.11 demonstrates.

As depicted in Fig. 5.12, the sampling size in x is especially relevant for the deviations. Since also the influence of the chosen interpolation scheme decreases with a larger sampling size, it is probably also not negligible. Recall that although the interpolation only directly affects n grid points, all other points are indirectly affected by large x deviations through the convolution. However, due to the small absolute values for $x \rightarrow 1$, these deviations become irrelevant for smaller x and larger xPDF values.

Finally, it can be concluded that despite the deviations for large x , which can be reduced by subgrids or a higher sampling size and of course also larger inversion limits $z_{\max} \rightarrow \infty$, both solution methods are in very good agreement and satisfy the PDF benchmark tables of [24], since the Mellin code agrees with them. A direct test of the x -space solutions with these benchmark tables requires an interpolation of the final grid, which is still unfinished.

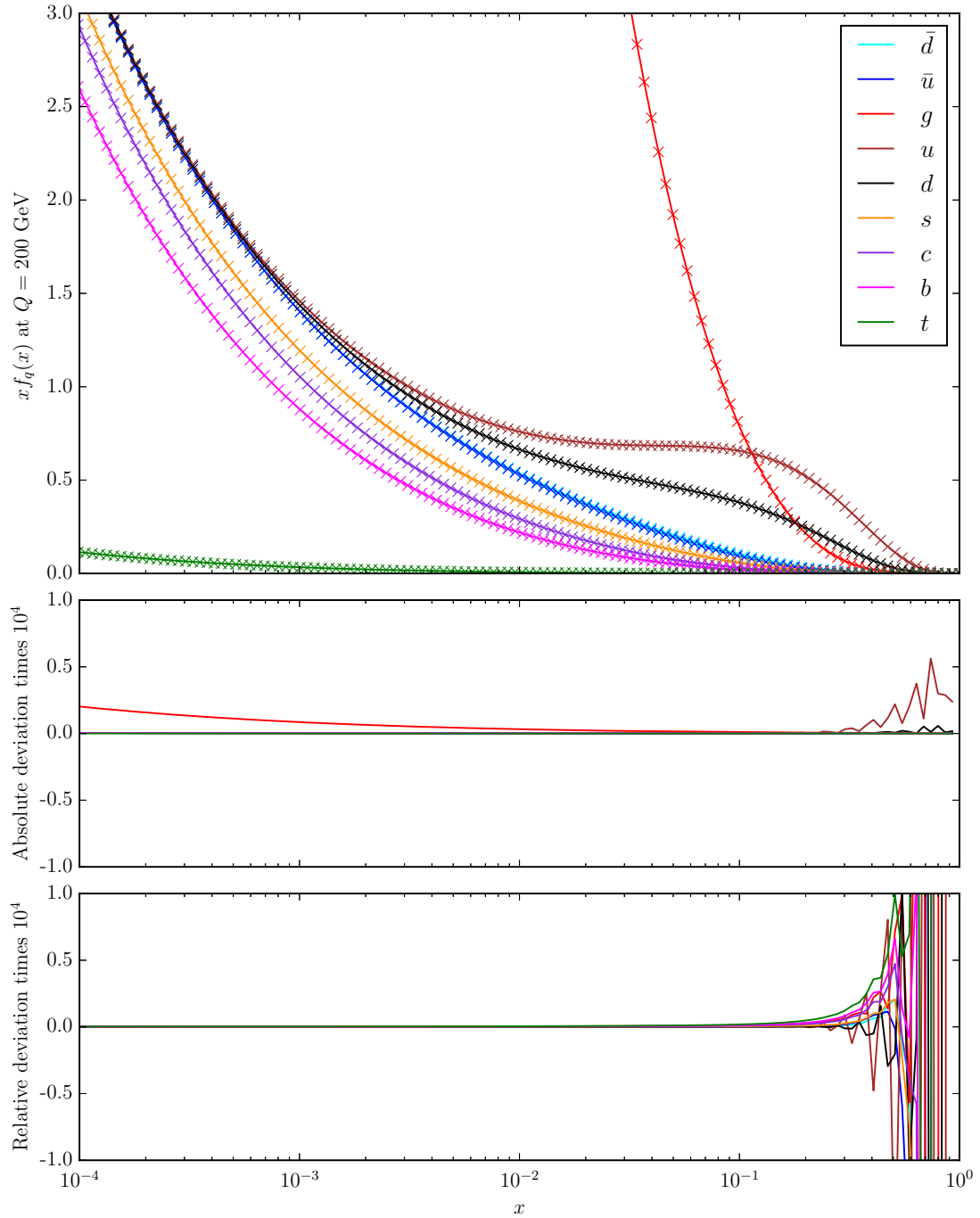


Fig. 5.10: Comparison between the results of the Mellin-space (lines) and x -space (crosses) evolution at $Q = 200$ GeV using the default inversion and the settings $n = 6$, $N_x = 161$, $h = 0.05$. Upper panel: xPDF values. Lower panels: absolute and relative deviations of the x -space results from the Mellin code results.

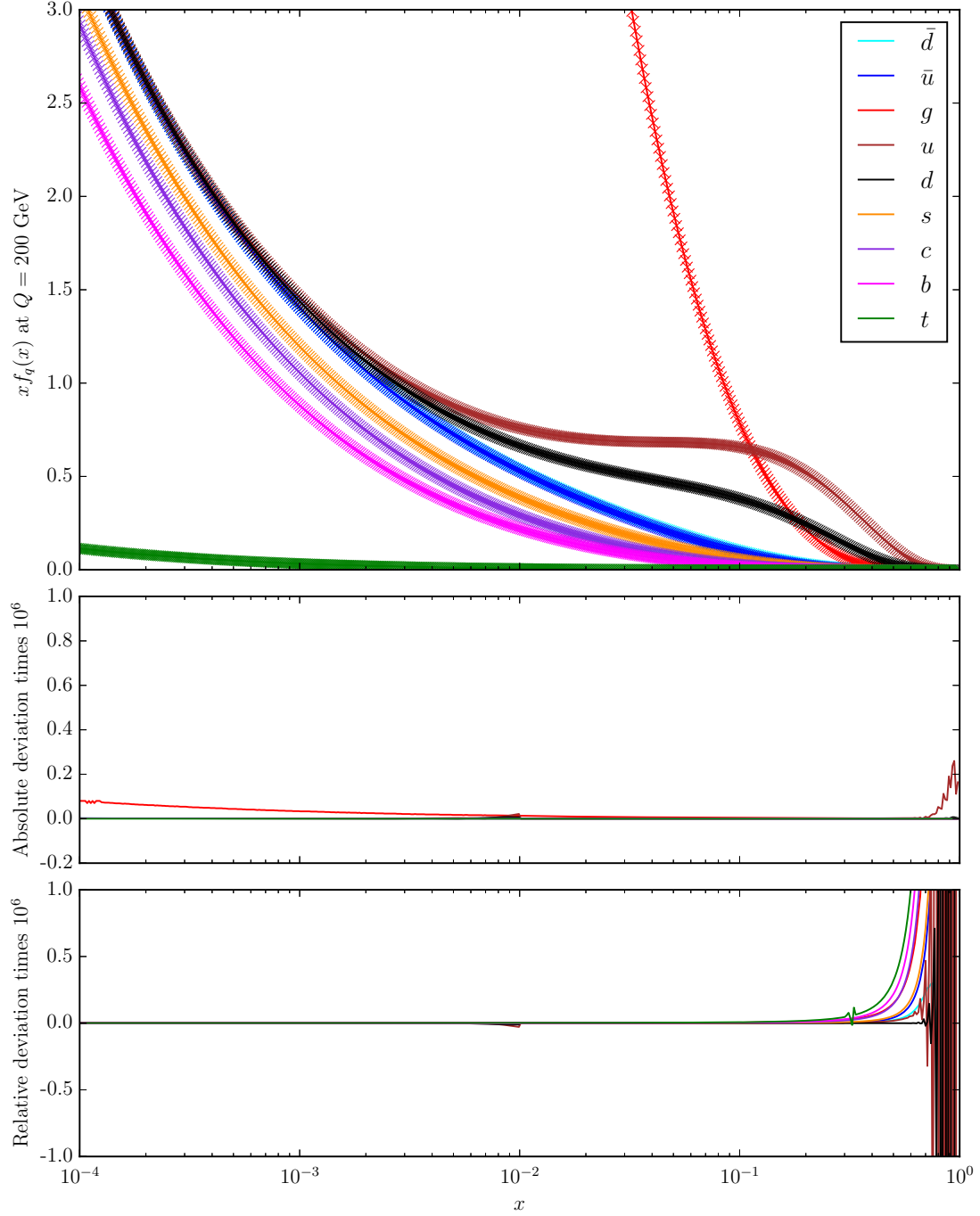


Fig. 5.11: Comparison between the results of the Mellin-space (lines) and x -space (crosses) evolution at $Q = 200$ GeV using the default inversion except for $z_{\max} = 200$ and the settings $n = 6$, $N_x = 641$, $h = 0.001$. Upper panel: xPDF values. Lower panels: absolute and relative deviations of the x -space results from the Mellin code results.

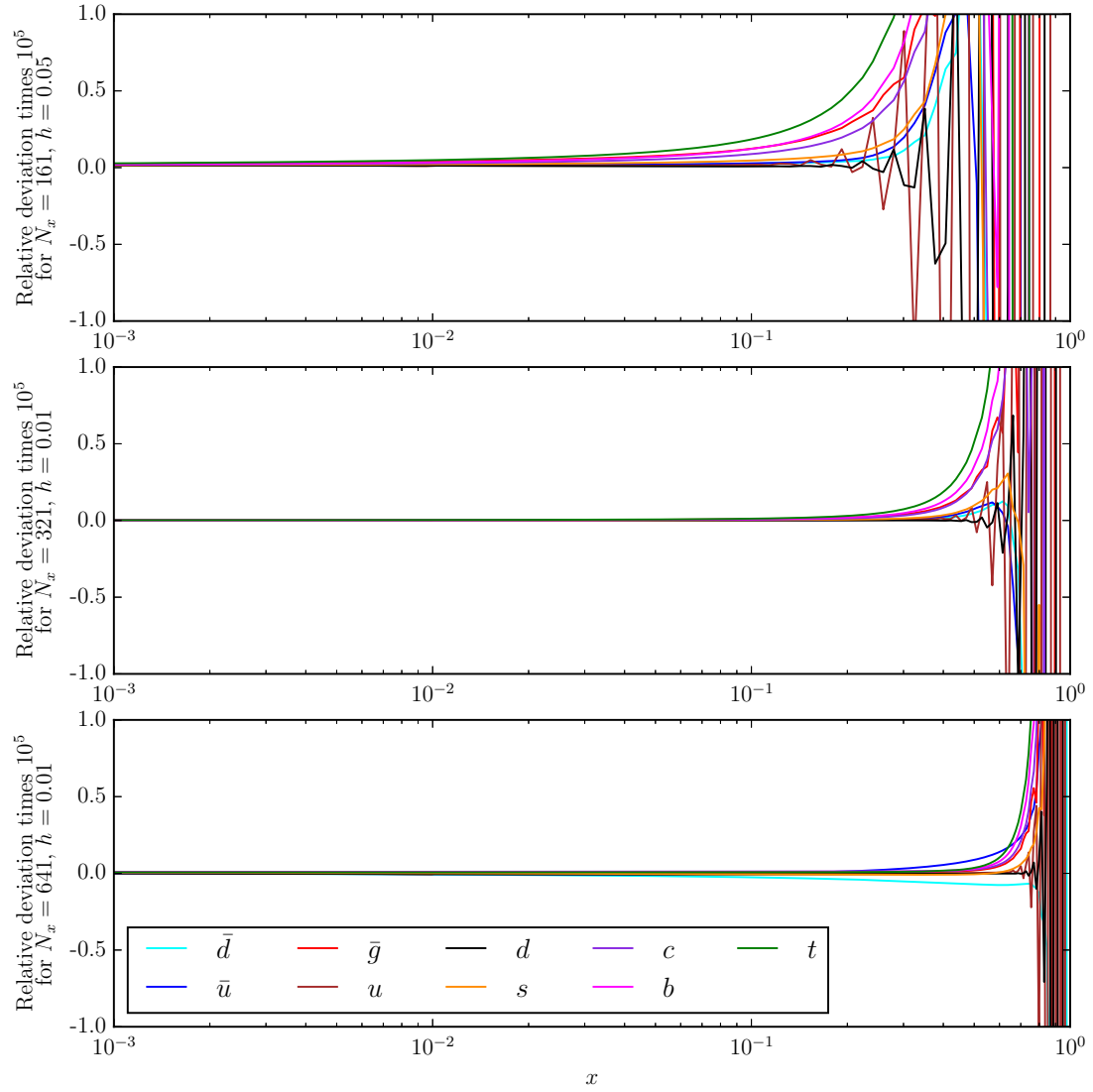


Fig. 5.12: Relative deviations of the x -space results from the Mellin-space results for different x evolution settings using the default options in Mellin-space.

Chapter 6

Summary and Perspectives

In this thesis, a computational tool has been developed which is able to evolve parton densities given at an initial scale up to a higher scale by solving the DGLAP equations. At present, the evolution can be performed at leading order using a fixed flavour number scheme or a variable flavour number scheme in x -space or in Mellin-space. The developed tool has been tested with the established tools QCD-PEGASUS and HOPPET and shows very good agreement. A comparison of the Mellin-space results with the Les Houches, 2001, benchmark results was also successful. Only the direct comparison between the solution in Mellin-space and x -space is not yet as accurate as desired. This is due to the inaccuracies at large x ($x > 0.3$), which are due to the interpolation scheme, the sampling size in the x -space solution, and the inversion in the Mellin code. Possible ways of improving this situation, e.g. by using subgrids in the x -space evolution, have been proposed. Further short-term objectives have been mentioned in Chapter 4.

In the near future, the described code will be merged with an interpolation routine for the final x - Q -grid and theoretical calculations of observables in a fitting procedure, and hence will be able to fit parton distribution functions to experimental data.

Regarding long-term objectives for the evolution methods: these can, first of all, be extended to next-to-leading order. This primarily involves implementing the NLO splitting functions and anomalous dimensions. Furthermore, in Mellin-space the evolution matrices U , described in Chapter 4, have to be implemented. The code is already designed with such an extension in mind, so that this will require little additional effort. However, in NNLO, also the matching conditions have to be implemented since the PDFs are not continuous anymore. Other long-term objectives are a variable factorization and renormalization scale, polarized PDFs, a Mellin transformation of initial PDFs, and an optimized Mellin inversion.

In summary, an extensive computation tool for PDF evolutions has been developed in the modern object-orientated programming language C++ with a simple user interface. Despite the preliminary status of the project, highly accurate results can already be produced in an efficient manner. Thus, the evolution methods form a substantial basis for extensions outlined above.

Appendix A

Additional Numerical Results

A.1 Interpolation Order

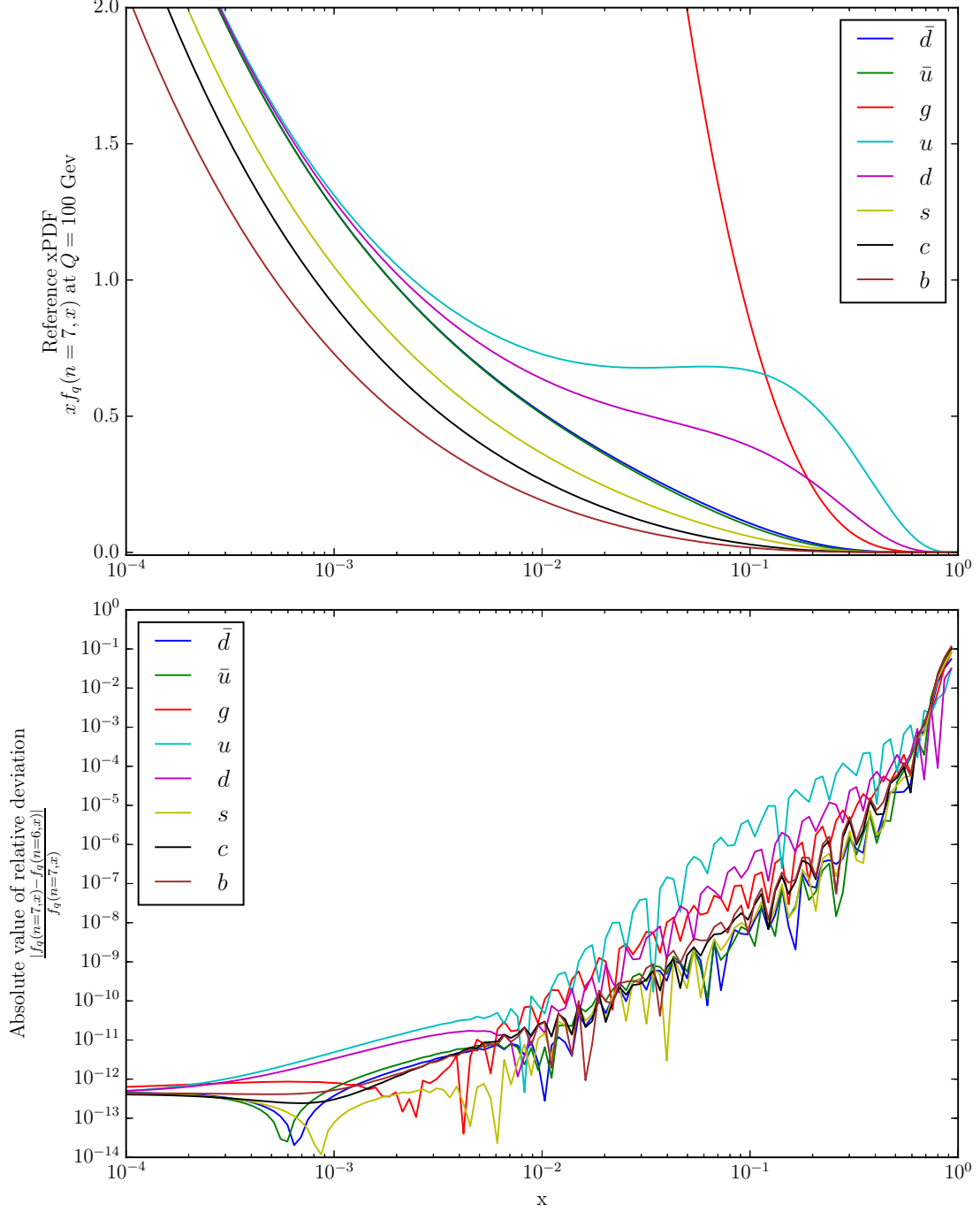


Fig. A.1: Upper panel: reference xPDF using a sampling size in x of $N_x = 161$ and a interpolation order of $n = 7$ at scale $Q = 100$ GeV. Lower panel: deviations of a xPDF with interpolation order $n = 6$ from the reference xPDF.

A.2 Comparison to HOPPET

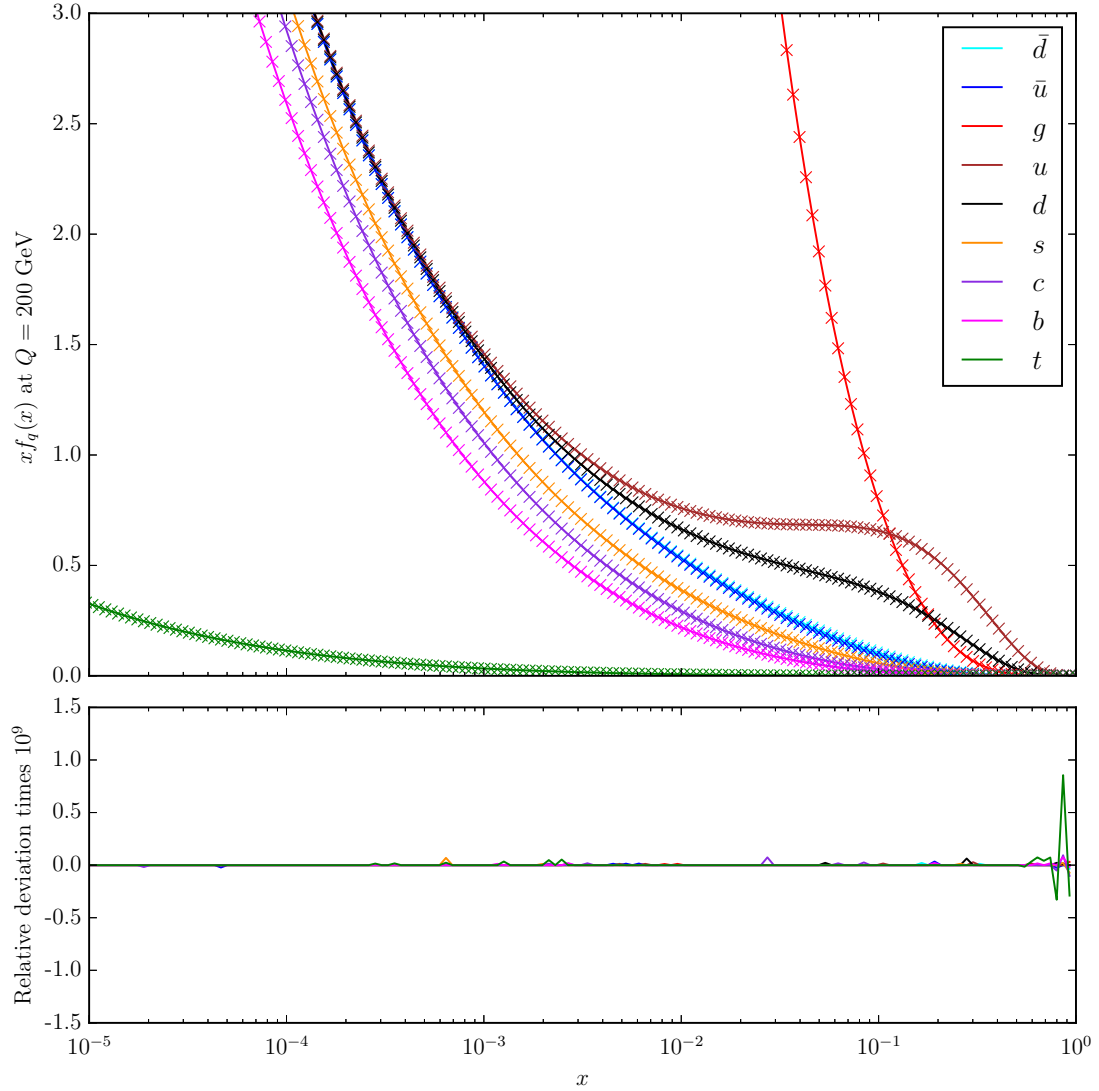


Fig. A.2: Upper panel: comparison between the x -space code (crosses) and the HOPPET (lines) output at $Q = 200 \text{ GeV}$ using a sampling $N_x = 161$, an interpolation order $n = 6$ and integration stepsizes $du = 0.001$ (HOPPET), and $h = 0.01$. Lower panel: relative deviations from the HOPPET results.

A.3 Comparison to QCD-PEGASUS

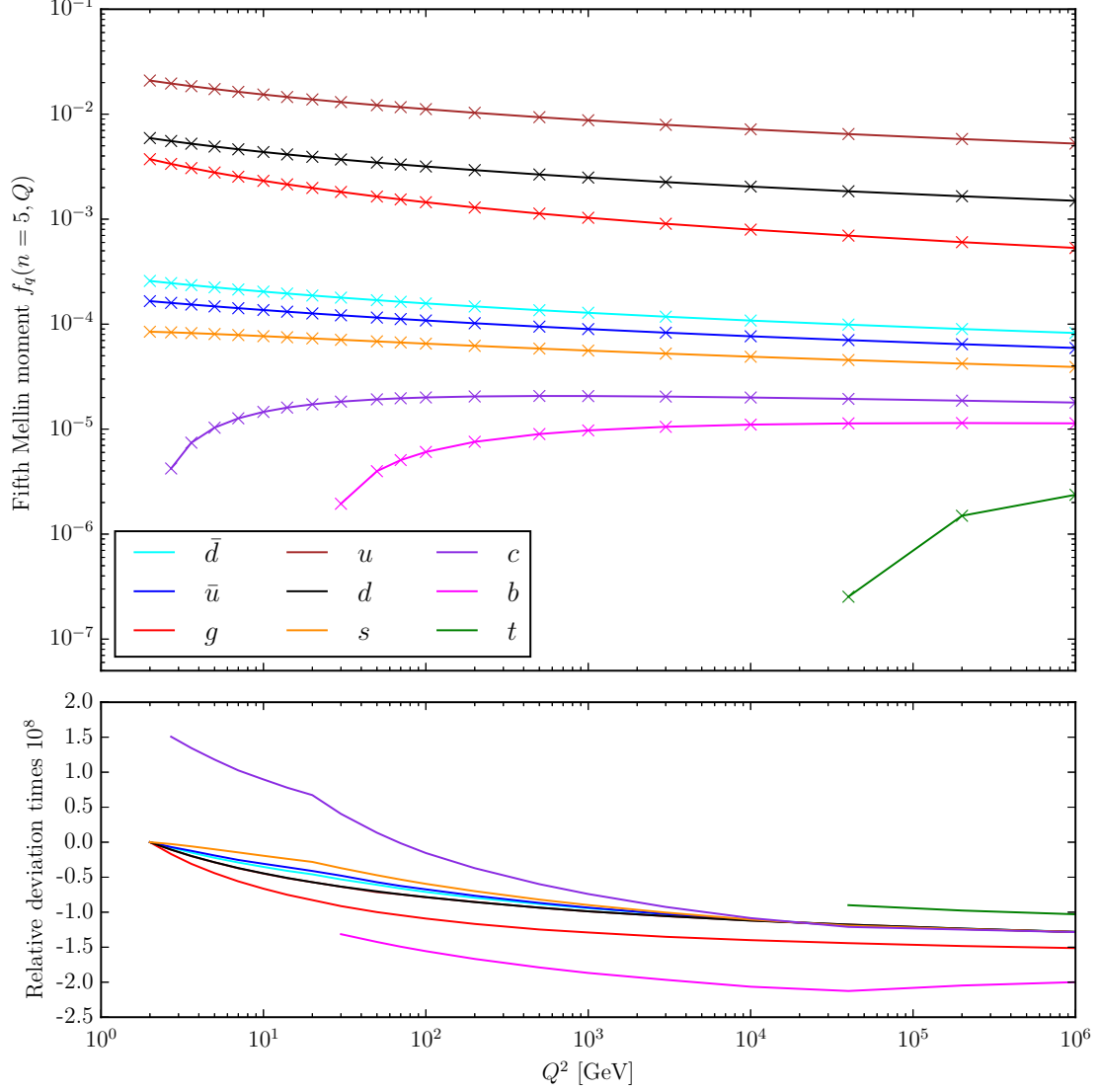


Fig. A.3: Upper panel: comparison between the Mellin code (x) and the QCD-PEGASUS (lines) output for the fifth Mellin moment and for different Q^2 using the QCD-PEGASUS normalization ($IMOMIN = 1$, $ISSIMP = 0$) of the initial conditions in both programs. Lower panel: relative deviations of the Mellin code results from the QCD-PEGASUS results.

Bibliography

- [1] ABRAMOWITZ, M. and I. A. STEGUN: *Handbook of Mathematical Functions with Formulas, Graphs, and Mathematical Tables*. Dover Publishing, New York, 9. ed., 1972.
- [2] AITCHISON, I. J. R. and A. J. G. HEY: *Gauge Theories in Particle Physics: A Practical Introduction. Vol. 1: From Relativistic Quantum Mechanics to QED*. CRC Press, Boca Raton et al., 4. ed., 2013.
- [3] AITCHISON, I. J. R. and A. J. G. HEY: *Gauge Theories in Particle Physics: A Practical Introduction. Vol. 2: Non-Abelian Gauge Theories: QCD and the Electroweak Theory*. CRC Press, Boca Raton et al., 4. ed., 2013.
- [4] ALBINO, S.: *Analytic Continuation of Harmonic Sums*. Phys. Lett., B674:41–48, 2009 ([arXiv:0902.2148 \[hep-ph\]](#)).
- [5] ALTARELLI, G. and G. PARISI: *Asymptotic Freedom in Parton Language*. Nucl. Phys., B126:298–318, 1977.
- [6] BERTRAND, J., P. BERTRAND and J. OVARLEZ: "The Mellin Transform". In POULARIKAS, A. D. (ed.): *The Transforms and Applications Handbook*. CRC Press LLC, Boca Raton et al., 2. ed., 2000.
- [7] BJORKEN, J. D.: *Asymptotic Sum Rules at Infinite Momentum*. Phys. Rev., 179:1547–1553, 1969.
- [8] BLUMLEIN, J.: *Analytic Continuation of Mellin Transforms up to Two Loop Order*. Comput. Phys. Commun., 133:76–104, 2000 ([arXiv:hep-ph/0003100](#)).
- [9] BOTJE, M.: *QCDNUM: Fast QCD Evolution and Convolution*. Comput. Phys. Commun., 182:490–532, 2011 ([arXiv:1005.1481 \[hep-ph\]](#)).
- [10] BUZA, M., Y. MATIOUNINE, J. SMITH and W. L. VAN NEERVEN: *Charm Electroproduction Viewed in the Variable Flavor Number Scheme Versus Fixed Order Perturbation Theory*. Eur. Phys. J., C1:301–320, 1998 ([arXiv:hep-ph/9612398](#)).
- [11] CHETYRKIN, K. G., B. A. KNIEHL and M. STEINHAUSER: *Strong Coupling Constant with Flavor Thresholds at Four Loops in the $\overline{MS}$ Scheme*. Phys. Rev. Lett., 79:2184–2187, 1997 ([arXiv:hep-ph/9706430](#)).

- [12] COURAND, R. and D. HILBERT: *Methoden der Mathematischen Physik I*. Julius Springer Verlag, Berlin, 1924.
- [13] CURCI, G., W. FURMANSKI and R. PETRONZIO: *Evolution of Parton Densities Beyond Leading Order: The Nonsinglet Case*. Nucl. Phys., B175:27, 1980.
- [14] DOKSHITZER, Y. L.: *Calculation of the Structure Functions for Deep Inelastic Scattering and e^+e^- Annihilation by Perturbation Theory in Quantum Chromodynamics*. Sov. Phys. JETP, 46:641–653, 1977. [Zh. Eksp. Teor. Fiz., 73:1216, 1977].
- [15] DULAT, S. et al.: *The CT14 Global Analysis of Quantum Chromodynamics*. 2015 (arXiv:1506.07443 [hep-ph]).
- [16] ELLIS, R., W. STERLING and B. WEBBER: *QCD and Collider Physics*. Cambridge University Press, Cambridge, 1. paperback ed., 2003.
- [17] ELLIS, R. K., Z. KUNSZT and E. M. LEVIN: *The Evolution of Parton Distributions at Small x* . Nucl. Phys., B420:517–549, 1994. [Erratum: Nucl. Phys., B433: 498, 1995].
- [18] FEYNMAN, R. P.: *The Behavior of Hadron Collisions at Extreme Energies*. Conf. Proc., C690905:237–258, 1969.
- [19] FRITZSCH, H., M. GELL-MANN and H. LEUTWYLER: *Advantages of the Color Octet Gluon Picture*. Phys. Lett., B47:365–368, 1973.
- [20] FURMANSKI, W. and R. PETRONZIO: *Singlet Parton Densities Beyond Leading Order*. Phys. Lett., B97:437, 1980.
- [21] FURMANSKI, W. and R. PETRONZIO: *Lepton – Hadron Processes Beyond Leading Order in Quantum Chromodynamics*. Z. Phys., C11:293, 1982.
- [22] GELL-MANN, M.: *A Schematic Model of Baryons and Mesons*. Phys. Lett., 8:214–215, 1964.
- [23] GEORGI, H. and H. D. POLITZER: *Electroproduction Scaling in an Asymptotically Free Theory of Strong Interactions*. Phys. Rev., D9:416–420, 1974.
- [24] GIELE, W. et al.: *"The QCD / SM Working Group: Summary Report"*. In *Physics at TeV Colliders. Proceedings, Euro Summer School, Les Houches, France, May 21-June 1, 2001*, pp. 275–426, 2002 (arXiv:hep-ph/0204316).
- [25] GLUCK, M., E. REYA and A. VOGT: *Radiatively Generated Parton Distributions for High-Energy Collisions*. Z. Phys., C48:471–482, 1990.
- [26] GRAUDENZ, D., M. HAMPEL, A. VOGT and C. BERGER: *The Mellin Transform Technique for the Extraction of the Gluon Density*. Z. Phys., C70:77–82, 1996 (arXiv:hep-ph/9506333).

- [27] GRIBOV, V. N. and L. N. LIPATOV: *Deep Inelastic $e p$ Scattering in Perturbation Theory*. Sov. J. Nucl. Phys., 15:438–450, 1972. [Yad. Fiz., 15:781, 1972].
- [28] GROSS, D. J. and F. WILCZEK: *Ultraviolet Behavior of Nonabelian Gauge Theories*. Phys. Rev. Lett., 30:1343–1346, 1973.
- [29] GROSS, D. J. and F. WILCZEK: *Asymptotically Free Gauge Theories. 2*. Phys. Rev., D9:980–993, 1974.
- [30] HALZEN, F. and A. D. MARTIN: *Quarks and Leptons: An Introductory Course in Modern Particle Physics*. John Wiley & Sons, New York et al., 1. paperback ed., 1984.
- [31] KOECHER, M. and A. KRIEG: *Elliptische Funktionen und Modulformen*. Springer-Verlag, Berlin et al., 1998.
- [32] KOVARIK, K. et al.: *nCTEQ15 - Global Analysis of Nuclear Parton Distributions with Uncertainties in the CTEQ Framework*. arXiv:1509.00792 [hep-ph], 2015.
- [33] LARIN, S. A., T. VAN RITBERGEN and J. A. M. VERMASEREN: *The Large Quark Mass Expansion of $\Gamma(Z^0 \rightarrow \text{hadrons})$ and $\Gamma(\tau^- \rightarrow \nu_\tau + \text{hadrons})$ in the order α_s^3* . Nucl. Phys., B438:278–306, 1995 (arXiv:hep-ph/9411260).
- [34] LIPATOV, L. N.: *The Parton Model and Perturbation Theory*. Sov. J. Nucl. Phys., 20:94–102, 1975. [Yad. Fiz., 20:181, 1974].
- [35] MARTIN, A. D.: *Proton Structure, Partons, QCD, DGLAP and Beyond*. Acta Phys. Polon., B39:2025–2062, 2008 (arXiv:0802.0161 [hep-ph]).
- [36] MOCH, S., J. A. M. VERMASEREN and A. VOGT: *The Three Loop Splitting Functions in QCD: The Nonsinglet Case*. Nucl. Phys., B688:101–134, 2004 (arXiv:hep-ph/0403192).
- [37] MUTA, T.: *Foundations of Quantum Chromodynamics: An Introduction to Perturbative Methods in Gauge Theories*, vol. 78 of *World scientific Lecture Notes in Physics*. World Scientific, Hackensack, 3. ed., 2010.
- [38] OLIVE, K. A. et al.: *Review of Particle Physics*. Chin. Phys., C38:090001, 2014.
- [39] POLITZER, H. D.: *Reliable Perturbative Results for Strong Interactions?*. Phys. Rev. Lett., 30:1346–1349, 1973.
- [40] PRESS, W. H., S. A. TEUKOLSKY et al.: *Numerical Recipes. The Art of Scientific Computing*. Cambridge University Press, Cambridge, 3. ed., 2007.
- [41] PUMPLIN, J., D. R. STUMP, J. HUSTON, H. L. LAI, P. M. NADOLSKY and W. K. TUNG: *New Generation of Parton Distributions with Uncertainties from Global QCD Analysis*. JHEP, 07:012, 2002 (arXiv:hep-ph/0201195).

- [42] SALAM, G. P. and J. ROJO: *A Higher Order Pertubative Parton Evolution Toolkit (HOPPET). Version 1.1.5.* Comput. Phys. Commun., 180:120–156, 2009 ([arXiv:0804.3755v1 \[hep-ph\]](#)).
- [43] VERMASEREN, J. A. M.: *Harmonic Sums, Mellin Transforms and Integrals.* Int. J. Mod. Phys., A14:2037–2076, 1999 ([arXiv:hep-ph/9806280](#)).
- [44] VOGT, A.: *Efficient Evolution of Unpolarized and Polarized Parton Distributions with QCD-PEGASUS.* Comput. Phys. Commun., 170:65–92, 2005 ([arXiv:hep-ph/0408244](#)).
- [45] VOGT, A., S. MOCH and J. A. M. VERMASEREN: *The Three-Loop Splitting Functions in QCD: The Singlet Case.* Nucl. Phys., B691:129–181, 2004.

Danksagungen

Ich möchte mich an dieser Stelle bei all denjenigen bedanken, die mich während der Anfertigung dieser Masterarbeit unterstützt und begleitet haben:

Ich danke Prof. Dr. Michael Klasen für die freundliche Aufnahme in seine Arbeitsgruppe und die Unterstützung auch über den Horizont dieser Arbeit hinaus.

Ganz besonders danke ich Dr. Karol Kovarik, der die Arbeit vorrangig betreut hat, bei Fragen zur Verfügung stand und durch seine umfangreiche Kritik zu dieser Arbeit beigetragen hat.

Ebenfalls danke ich Dr. Thomas Jezo, welcher in einer intensiven Arbeitswoche wichtige Impulse zur Gestaltung der Benutzeroberfläche beigesteuert hat.

Auch meinem langjährigen Freund Dr. Axel Finke gilt meine Dankbarkeit für das umfassende und gewissenhafte Korrekturlesen und die vielen hilfreichen Anmerkungen zur L^AT_EX Typographie.

Schließlich gilt mein unbegrenzter Dank meiner Frau Lisa und meinem Sohn Anton sowie meinen Eltern, die mir allezeit mit viel Geduld und Unterstützung zur Seite standen.

Plagiatserklärung¹

Hiermit versichere ich, dass die vorliegende Arbeit

„Evolution of Parton Distribution Functions“

selbstständig verfasst worden ist, dass keine anderen Quellen und Hilfsmittel als die angegebenen benutzt worden sind und dass die Stellen der Arbeit, die anderen Werken – auch elektronischen Medien – dem Wortlaut oder Sinn nach entnommen wurden, auf jeden Fall unter Angabe der Quelle als Entlehnung kenntlich gemacht worden sind.

Münster, 30. Oktober 2015

Florian Herrmann

Ich erkläre mich mit einem Abgleich der Arbeit mit anderen Texten zwecks Auffindung von Übereinstimmungen sowie mit einer zu diesem Zweck vorzunehmenden Speicherung der Arbeit in einer Datenbank einverstanden.

Münster, 30. Oktober 2015

Florian Herrmann

¹Nach dem Vorbild des Prüfungsamtes Physik der WWU Münster.