

Tensor-network approach to QCD in the strong-coupling expansion



DISSERTATION ZUR ERLANGUNG DES DOKTORGRADES
DER NATURWISSENSCHAFTEN (DR. RER. NAT.)
DER FAKULTÄT FÜR PHYSIK DER UNIVERSITÄT REGENSBURG

vorgelegt von

Thomas Samberger

aus

Vilsbiburg

im Jahr 2026

Promotionsgesuch eingereicht am: 10.04.2026

Die Arbeit wurde angeleitet von: Prof. Dr. Tilo Wettig

Prüfungsausschuss:

Prof. Dr. Sascha Schäfer

Prof. Dr. Tilo Wettig

Prof. Dr. Shinji Takeda

Prof. Dr. Markus Schmitt

Datum Promotionskolloquium: 09.07.2026

Contents

I	Introduction	8
II	Theoretical foundations	13
1	Fundamental concepts of group theory	13
1.1	Basic definitions and properties	13
1.2	Group representations	15
1.3	Reducibility	16
1.4	Characters of representations	17
1.5	Group algebra	18
1.6	The symmetric group	20
1.7	Construction of the irreps of the symmetric group	22
1.8	Murnaghan-Nakayama rule	24
1.9	Lie groups	26
2	Fundamental concepts of quantum field theory	30
2.1	Path integral formulation in quantum mechanics	30
2.2	Path integral formulation for a scalar field theory	32
2.3	Wick rotation	34
2.4	Thermodynamical observables using the path integral	35
2.5	Quantum chromodynamics	36
2.6	Lattice QCD	39
2.6.1	Wilson's plaquette action	40
2.6.2	Discretized fermion action	42
2.6.3	Chemical potential on the lattice	45
2.6.4	Partition function and observables	46
2.7	Chiral symmetry	47
2.8	Sign problem for nonzero chemical potential	49
2.9	Phase diagram of QCD	51
3	Tensor-network procedure	54
3.1	HOSVD	55
3.2	HOTRG	58
3.2.1	Backward-forward symmetry	59
3.2.2	SuperQ and the Xie et al. procedure	62
3.2.3	Stabilized finite difference	63

3.3	GHOTRG	66
3.3.1	Contraction	67
3.3.2	Truncation	69
3.3.3	Iterative procedure	71
3.3.4	Tracing	72
3.3.5	ASAP tracing	72
3.3.6	Truncation of the initial numerical tensor	73
III	Deriving the tensor network	75
4	Dual variables and evaluation of the gauge integral	75
4.1	Taylor expansions of the Boltzmann factor	75
4.2	Disentangling color indices	76
4.3	Gauge-link integral and simplifications	78
4.3.1	Gauge-link integral	78
4.3.2	Range reduction due to Grassmann variables	80
5	Creating local contributions	81
5.1	Grassmann integration	82
5.2	Construction of the tensor network	83
5.2.1	Elimination of color degrees of freedom	83
5.2.2	Tensor formulation	84
5.2.3	Restriction of the edge variables	86
6	Translational invariance for GHOTRG	88
IV	The OS-GHOTRG method	90
7	Order separation in β	90
7.1	Overview	90
7.2	Introduction of edge occupations	91
7.3	Tensor contraction and decomposition	92
7.4	Reduction of the tensor	94
7.5	Truncation of the tensor	95
7.6	Tensor trace	98
8	Translational invariance for OS-GHOTRG	100
8.1	Contributions to the partition function	100
8.2	The X network	101
8.3	The Y network	102
8.4	Matching the coefficients	103
V	Results	104
9	Analytical Calculations	104
9.1	Analytical calculation for small lattices	104
9.2	Pure-gauge theory	106
9.3	Properties of the partition function and observables	109

9.3.1	Asymptotic results for large chemical potential	111
9.3.2	Asymptotic results for large mass	112
9.3.3	Even mass exponents	115
9.3.4	Generalization to arbitrary number of flavors	116
9.4	Interpolation between small and large μ for large L_t	119
9.5	Applicability of the strong-coupling expansion	121
10	Numerical Results	122
10.1	Comparison of SuperQ and the Xie et al. truncation	122
10.2	Pure-gauge theory	123
10.3	One flavor	124
10.3.1	Comparing different expansions	124
10.3.2	Validation of OS-GHOTRG with analytical results and MC data	128
10.3.3	Behavior of the singular values	129
10.3.4	Results for larger lattices near the phase transition	131
10.4	Two flavors	139
VI	Summary and conclusions	142
	Acknowledgements	144
A	Grassmann contribution	145
B	Initial tensor for $n_{\max} = 0$, $N_f = 1$, $N_c = 3$, and $d = 2$	150
C	Proofs for OS-GHOTRG	156
C.1	Proof of tensor decomposition after contraction	156
C.2	Proof of link and site criteria	158
C.3	Proof of tensor decomposition after tracing	159
D	Validating the coefficients of Sec. 8.4	160
E	Expansion coefficients for a 2×2 lattice	160
F	Proof of $\lambda_R(\beta) = \mathcal{O}(\beta^{\bar{q}_1})$	162
G	Taylor expansion of free energy	163
H	The function $p_{V_s}(m, \beta)$	167
I	Fit ansatz using the tanh models	167
I.1	Coefficients of the strong-coupling expansion for the observables	167
I.2	Integrating the model function for the quark-number density	169
J	Explicit fits for various m, L, and $n_{\max}$	169

I. Introduction

It is astonishing that the theory of Quantum Chromodynamics (QCD) exhibits an extremely rich structure, even though its defining Lagrangian density is strikingly simple. It is based on only six different flavors of matter fields - the quarks - and on the fundamental requirement of being invariant under a so-called local $SU(3)$ gauge transformation. This symmetry requirement alone leads to eight types of gluons, which are the mediators of the strong force. One of the resulting key features of QCD is that the strength of the coupling between quarks and gluons decreases with increasing energy scale. This has a number of consequences, two of which are highlighted here. At low energy scales, quarks are always confined within hadrons (e.g., protons and neutrons, which are the constituents of all atomic nuclei). In contrast, for very large energy scales - reached by high temperatures and densities - a state called Quark-Gluon Plasma (QGP) emerges, in which the confinement breaks down as the coupling of quarks and gluons becomes very small. It is expected that the universe was filled with a QGP in the first microseconds after the Big Bang, from which hadrons condensed as the universe cooled [1]. In fact, such a transition is understood very well at the theoretical level: To observe it, one needs to study the temperature dependence of expectation values of observables based on the partition function. This step requires an integration over all existing fields, i.e., gauge and quark fields, which is usually done in the following way: First, the integration over quark fields is performed analytically, yielding a quantity called fermion determinant. Afterwards, the integration over gauge fields is done numerically: To do so, a finite set of gauge configurations is sampled according to a probability distribution, which in turn is given by the fermion determinant and the gauge action. Finally, the observable is computed as the expectation value over this particular sample. Explicit calculation shows that the aforementioned transition to hadron condensation is in fact a crossover, i.e., a smooth transition without a discontinuity in any of the derivatives of the observable with respect to the temperature T [2].

Problems arise when one examines the phases of QCD not only at nonzero temperature, but also at baryon densities other than zero, as is relevant for heavy ion collisions or inside large neutron stars [2]. At the theoretical level, this is achieved by introducing a nonzero chemical potential μ , which creates a particle-antiparticle asymmetry. It turns out that for nonzero chemical potential the fermion determinant is no longer real but complex, and therefore the approach using importance sampling breaks down. In the literature, this is commonly referred to as the sign problem. In order to circumvent this sign problem, a variety of methods have been proposed, including reweighting, Taylor expansion in μ , analytic continuation from imaginary μ , complex Langevin, thimbles and path optimization, see [3, 4] for reviews. Nevertheless, none of the proposed remedies can successfully reach regimes in which $\mu/T > 1$, and therefore even today the full QCD phase diagram in the $\mu - T$ plane cannot be calculated. It is expected that, for sufficiently large chemical potential, the transition to deconfinement is no longer a crossover but a first-order phase transition [2]. Therefore, the existence of a critical endpoint is proposed. Further, for large chemical potential and low temperatures, a state with so-called color superconductivity is expected [5]. However, due to the $SU(N_c)$ nature of the color charge, there are many possible scenarios for how the color superconducting phase may manifest.

There is a promising further method, based on so-called dual variables, that differs crucially from all of the aforementioned approaches. There, the exponentials of the various contributions to the

action are Taylor-expanded, and the order of the resulting summation and the integration with respect to the field variables is exchanged. Also, the integration order of quark and gauge fields is reversed, i.e., the integration over gauge fields is performed first. Altogether, this corresponds to a transformation of the degrees of freedom which are no longer continuous but discrete, referred to as dual variables. This method strongly reduces the sign problem, but until now the dualization was mainly applied to the infinite-coupling limit of QCD [6, 7, 8, 9]. An attempt to go beyond this limit was made using the next-to-leading-order term in the inverse coupling β [10] (see also [11] for an early attempt in mean-field approximation). New strategies to go beyond the infinite-coupling limit in a worldline and worldsheet formulation were proposed using so-called Abelian color cycles [12] and, more recently, by introducing additional dual degrees of freedom called decoupling operator indices [13]. Using the method of Ref. [13], first results for the phase transition in four-dimensional QCD with one flavor of staggered quarks in the chiral limit were obtained up to order β^2 using a vertex model [14, 15]. First attempts to include the full gauge action of two-color QCD in two dimensions were made using numerical integrations of the gauge fields [16, 17].

In most cases, the worm algorithm [18] is used to simulate QCD in its dual formulation. As an alternative to Monte Carlo methods, tensor-network methods have recently been applied successfully to various statistical systems. In this approach, the quantity of interest (which is typically the partition function) is written as a contraction of a large number of local tensors. However, due to the curse of dimensionality, an analytical calculation is still numerically unfeasible in most cases. Therefore, approximations based on singular value decomposition are applied in all of these methods. Ref. [19] provides an overview of different tensor-network methods applied to a range of models. Of special relevance is the higher order tensor renormalization group (HOTRG) approach, see Ref. [20]. This is an algorithm, applicable to systems of arbitrary dimension, in which a tensor is assigned to every site of a closed lattice and a contraction is performed for each link between two neighboring sites. After a pairwise contraction of two neighboring tensors, the resulting coarse tensors are truncated using higher-order singular value decompositions (HOSVD) [21]. This algorithm is repeated until a single tensor remains, from which the trace is taken. An additional difficulty arises when a theory including fermions is considered. In this case, it is in general not possible to integrate out all fermions in terms of local objects, so that the result is suitable for the tensor-network formulation. Instead, all fermions are integrated out, except for auxiliary Grassmann variables, which become part of the considered tensors. The Grassmann structure in each tensor reproduces itself after each contraction step and creates a sign factor that is part of the coarse tensor. Such a procedure is known as Grassmann HOTRG (GHOTRG) [22]. Recently, GHOTRG was tailored to strong-coupling QCD using one flavor of staggered quarks in two dimensions [23]. This enabled the calculation of various observables for arbitrary masses, chemical potentials, and lattice volumes. Moreover, the extension to four dimensions, in order to investigate the phase diagram of infinite-coupling QCD, showed quite promising results [24]. Similar computations were performed very recently to investigate cold and dense QCD in the strong-coupling limit [25]. In both two and four dimensions, it was observed that the sign problem originating from nonzero chemical potential did not cause any numerical instabilities. First results for including the first-order strong-coupling expansion to the GHOTRG approach were obtained in Ref. [26].

In this thesis, we go beyond both the infinite-coupling limit and the single-flavor case. The main goal is to express the partition function of lattice QCD with staggered fermions in the strong-coupling expansion in terms of a tensor network. This is done for an arbitrary number

of dimensions, flavors, and orders of the strong-coupling expansion. Further, the GHOTRG approach shall be extended such that every coefficient Z_n of the partition function $Z = \sum_n Z_n \beta^n$ can be calculated separately. Finally, it is part of this thesis to study the numerical results in two dimensions. The main findings of this thesis have been published in a collaboration with Jacques Bloch, Robert Lohmayer, and Tilo Wettig in Ref. [27] and Ref. [28], with, in particular, the wording of these publications being largely incorporated here. The methods were first developed for two dimensions in this thesis and were generalized to an arbitrary number of dimensions mainly by Jacques Bloch and Tilo Wettig.

This thesis is structured as follows: Part II introduces essential principles on which the work is based and is divided into three sections: In Sec. 1, fundamental concepts of group theory are presented, including matrix representations of group elements and the corresponding characters. Special attention is paid to the symmetric group and the computation of its characters using the Murnaghan-Nakayama rule. Additionally, Lie groups, which are essential for gauge theories such as QCD, are introduced. In Sec. 2, concepts of quantum field theory are introduced, with a focus on the path integral formalism. Further, starting from the continuous QCD Lagrangian density in d -dimensional Minkowski space, we derive the Lagrangian density of lattice QCD in d -dimensional Euclidean space-time with staggered fermions. We study the chiral symmetry of the QCD Lagrangian density, explain the aforementioned sign problem more rigorously, and provide an overview of the expected QCD phase diagram. Finally, Sec. 3 is dedicated to the introduction of tensor-network methods. We introduce HOSVD as well as HOTRG for arbitrary number of dimensions. Two different variants of coarse-tensor truncation are discussed, using either the so-called SuperQ method [29] or the Xie et al. [20] procedure. Moreover, an improvement for calculating observables is presented: Observables are typically obtained through a derivative of the partition function with respect to some parameter. Due to the discrete nature of tensor-network methods, this derivative is generally approximated poorly when computed by a finite difference. In the presented stabilized finite-difference method, the variation of every singular vector for varied parameter is taken into account, which allows a less error-prone calculation of observables. Finally, we discuss the GHOTRG approach in detail, as applied, for example, in Ref. [23] to two-dimensional strong-coupling QCD.

Most of Part III has been published in [27] as it contains one of the main findings of this thesis. The partition function of an $SU(N_c)$ gauge theory in d -dimensional Euclidean space-time with N_f flavors of staggered fermions is expressed in the form of a local tensor network with the Grassmann structure from Ref. [23]. The basic idea is quite straightforward: In addition to the Taylor expansion of the Boltzmann factor involving the fermion action, as used previously in the infinite-coupling limit, one also performs a Taylor expansion of the Boltzmann factor involving the gauge action. The gauge links (i.e., the gauge-field matrices defined on the links of the lattice) contained in the terms that arise from these expansions can still be integrated out exactly using known integrals [30, 31, 32, 13, 33]. The integrals are performed on individual color components of the gauge links and their adjoints. Consequently, the original colored Grassmann variables can be integrated out after introducing new colorless auxiliary Grassmann variables on the links of the lattice. This generates local sign factors that can be incorporated in the numerical tensor. After integration, all color indices can be contracted, and the partition function can be rewritten as a completely contracted tensor network of local numerical and Grassmann tensors. In contrast to the infinite-coupling case, the colorless auxiliary Grassmann variables no longer correspond to baryonic degrees of freedom, but are generic fermionic degrees of freedom. Nevertheless, the GHOTRG blocking procedure remains completely identical to that of the infinite-coupling case

in Ref. [23]. As a consistency check, we also show that the obtained tensor reduces to the one stated in Ref. [23] for the case of one flavor, two dimensions, $N_c = 3$, and $\beta = 0$.

Part IV addresses order-separated GHOTRG (OS-GHOTRG) as described in the corresponding publication [28]. It is a modification of GHOTRG that allows the separate calculation of the coefficients Z_n of the partition function $Z = \sum_n Z_n \beta^n$ up to some order $n_{\max}$. It is based on the observation that plaquettes on the lattice are collapsed onto links, and (after at least two contractions) also onto sites. To distinguish how many plaquettes are already collapsed onto a specific site, we introduce $n_{\max} + 1$ tensors of equal shape on every site. Additionally, the link indices are grouped such that, in each group, both the number of plaquettes collapsed onto the link and the number of uncollapsed adjacent plaquettes in every direction are fixed. In a contraction, all summations are reordered such that this tensor structure is reproduced, albeit on a coarse lattice. Thus there are $n_{\max} + 1$ coarse tensors on every site, whose link indices are grouped according to the number of plaquettes collapsed onto that link and the adjacent uncollapsed plaquettes of that link. When the tensors are truncated, every group of link indices requires a separate reduction to avoid the mixing of different orders in β . As in the usual HOTRG approach, for a given group of link indices, the backward- and forward-truncation matrices need to be identical. After the repeated contraction steps, when the trace of the coarse tensor is taken, all summations are reordered again, so that all terms are sorted with respect to the orders in β . In this Part, we also show how translational invariance in the partition function can be used to improve the efficiency of the OS-GHOTRG method.

The results of this thesis are presented in Part V. We start with several analytical results. We state an algorithm that is used to compute the coefficients Z_n of the partition function $Z = \sum_n Z_n \beta^n$ analytically in an efficient way, based on the initial tensor network derived in Part III. Using this, the analytical coefficients Z_n on a 2×2 lattice are computed up to fourth order for one flavor and up to first order for two flavors. Afterwards, we discuss the simplifications of the tensor-network approach for pure-gauge theory in two dimensions and analytically calculate the coefficients of the strong-coupling expansion for a later comparison with numerical results. Returning to the full theory in arbitrary dimensions, we analytically derive several properties of the partition function and the resulting observables, including their symmetries and asymptotic behavior with respect to the chemical potentials and masses. These analytical predictions are in agreement with the numerical results. Afterwards we propose a model function arising from interpolating between small and large μ that turn out to describe the numerical data very well in many cases for $N_f = 1$. At the end of Sec. 9 we discuss the general applicability of the strong-coupling expansion near a phase transition again for arbitrary N_f .

From then on, we present a number of numerical results for two dimensions with $N_c = 3$. First, we compare SuperQ and the Xie et al. truncation procedure for different lattice volumes, $\beta = 0$, and $N_f = 1$ using the usual GHOTRG procedure. Afterwards, results for the pure-gauge theory using OS-HOTRG (which is the simplification of OS-GHOTRG when no Grassmanns are present) are compared with the exact coefficients from analytical calculations up to order six in β . For one flavor, we first compare different expansions which differ regarding whether the partition function is expanded in β or the observable of interest itself is expanded in β . We validate the implementation of the OS-GHOTRG method by comparing with analytical and Monte Carlo results for small lattices. After a study of the singular values for various quark masses and chemical potentials in several truncation steps, we show results for different observables near the phase transition. To model the behavior of the observables, we introduce a

fit ansatz based on the hyperbolic tangent, where we study the dependence of the fit parameters on the quark mass and the lattice size. It is expected that, by using this model, the range of β in which the strong-coupling expansion is applicable increases significantly. Finally, we also show results for $N_f = 2$ up to $n_{\max} = 1$.

II. Theoretical foundations

In this part, we introduce fundamental definitions and concepts of group theory and quantum field theory that form the basis of this thesis. We then present the QCD Lagrangian density, consider the discretization of space-time and discuss the Wick rotation that is required for the transformation from Minkowski to Euclidean space-time. Furthermore, we examine chiral symmetry and provide an overview of the expected QCD phase diagram. In the third section, we introduce HOTRG, a fundamental tensor-network method, along with its generalization GHOTRG, which allows for Grassmann-valued tensor entries as required for theories involving fermions.

1 Fundamental concepts of group theory

The study of group theory is of great importance as it is the study of symmetry itself, which has consequences in all fields of physics. Using group theory, crucial insights can be extracted from physical systems without necessarily knowing the details of the dynamics. The following brief introduction to group theory is based on [34, 35, 36, 37], where most of the proofs are left out for conciseness.

1.1 Basic definitions and properties

A **group** is a set G assigned with a binary operation called multiplication denoted by “ $\cdot$ ”, which satisfies the following properties:

- **closure:** If $f, g \in G$, then $f \cdot g \in G$.
- **associativity:** For any $f, g, h \in G$, $f \cdot (g \cdot h) = (f \cdot g) \cdot h$.
- **identity element:** There is an identity element $e \in G$ such that for all $f \in G$, $e \cdot f = f \cdot e = f$.
- **inverse element:** Every element $f \in G$ has an inverse element $f^{-1} \in G$ such that $f \cdot f^{-1} = f^{-1} \cdot f = e$.

When the multiplication is additionally commutative, i.e., for $f, g \in G$, $f \cdot g = g \cdot f$, the group is called abelian.

If the number of group elements is finite, the group is called a finite group. In this case, the number of group elements n of a group G is called order of the group, and we write $|G| = n$.

A subset $H \subseteq G$ is called a **subgroup** of the group G , when H together with the multiplication of G is a group. To define the multiplication on H , we naturally restrict the multiplication from the group G to the elements of H . When G is finite, one can show that the order of G is divisible

by the order of H , i.e., $|G|/|H| \in \mathbb{N}$. Every group G has the subgroups $\{e\}$ and G , which are called trivial subgroups. All other subgroups are called nontrivial.

Two group elements $f, g \in G$ are said to be **conjugate** or **equivalent**, if there exists a group element $h \in G$, with $f = h \cdot g \cdot h^{-1}$. Equivalent group elements are denoted by $f \sim g$. This relation is reflexive (i.e., $f \sim f$ for all $f \in G$), symmetric (i.e., $f \sim g \Leftrightarrow g \sim f$ for all $f, g \in G$) and transitive (i.e., $f \sim g$ and $g \sim h \Rightarrow f \sim h$ for all $f, g, h \in G$). Therefore the equivalence relation naturally induces equivalence classes as well: The **equivalence class** of a group element $g \in G$ is the subset of all group elements of G that are equivalent to g , denoted by $[g] \equiv \{f \in G \mid f \sim g\}$. By definition, every element of a group G belongs to one and only one class. Even though, for any class $[g]$, the order of the group $|G|$ is divisible by the number of elements in the class, $[g]$ is, in general, not a group. The class that contains the identity element e is trivial, as this element always forms a class by itself, i.e., $[e] = \{e\}$. If the group is abelian, all the classes become trivial, as each element forms a class by itself.

Two subgroups K and H of G are **conjugate** to each other, if there exists an element $g \in G$ such that $K = gHg^{-1} \equiv \{g \cdot h \cdot g^{-1} \mid h \in H\}$.

H is called an **invariant** subgroup of G , when $ghg^{-1} \in H$ for all $h \in H$ and all $g \in G$. As a consequence, every group G has at least two invariant subgroups, which are the trivial subgroups $\{e\}$ and G .

A group is called **simple**, if it does not have any nontrivial invariant subgroups. Further it is called **semi-simple**, if it does not have any nontrivial abelian invariant subgroups.

For given subgroup $H \subseteq G$ and group element $g \in G$, we define the left **coset** of H with respect to the element g as the set $gH \equiv \{g \cdot h \mid h \in H\}$, and the right **coset** of H with respect to the element g as the set $Hg \equiv \{h \cdot g \mid h \in H\}$. Moreover, two cosets g_1H and g_2H are either identical, i.e., $g_1H = g_2H$ (if $g_1^{-1} \cdot g_2 \in H$) or disjoint, i.e., $g_1H \cap g_2H = \emptyset$ (if $g_1^{-1} \cdot g_2 \notin H$). This can be used to show that every group element $g \in G$ belongs to one and only one of the disjoint cosets of H . From the definition it follows, that the left and right coset coincide when $g \in H$, i.e., $gH = Hg = H$. Also if H is an invariant subgroup, the left and right coset are identical.

The case of an invariant subgroup is of particular interest, as all the cosets build the elements of a new group, where the multiplication of two elements is defined by $(g_1H) \cdot (g_2H) \equiv \{g_1 \cdot h_i \cdot g_2 \cdot h_j \mid h_i \cdot h_j \in H\} = \{g_1 \cdot g_2 \cdot h_k \mid h_k \in H\} = (g_1 \cdot g_2)H$. Here we inserted $h_k = (g_2^{-1} \cdot h_i \cdot g_2) \cdot h_j$, which is again a group element of H , as H is an invariant subgroup. Note that we use the same symbol for the multiplication of two cosets as well as for two elements out of G . The constructed group is called **factor group**, denoted by G/H , with order $|G/H| = |G|/|H|$.

A **group homomorphism** is a mapping Φ of the elements of a group G onto elements of a group H with $\Phi(f \cdot g) = \Phi(f) \cdot \Phi(g)$, for any elements $f, g \in G$, where we use the same symbol for multiplications in G and in H . There always exists at least one such mapping Φ , which is called the trivial homomorphism $\Phi(g) = e'$ for all $g \in G$, where e' is the identity element in H .

From the definition of a homomorphism it follows that $\Phi(e) = e'$, where e is the identity element of the group G and e' is the identity element of the group H . Additionally, inverses are mapped onto each other $\Phi(f^{-1}) = \Phi(f)^{-1}$. Here, $\Phi(f)^{-1}$ denotes the inverse of the group element

$\Phi(f) \in H$.

A crucial property of groups is formulated in the **rearrangement lemma**: If we multiply all elements of a group G by one of the elements of the group, we again obtain the set of group elements, i.e., for given $g \in G$, $G = \{g \cdot h \mid h \in G\} = \{h \cdot g \mid h \in G\}$. Note that when $g \neq e$, we get a new ordering of the elements.

This implies

$$\sum_{g \in G} f(g) = \sum_{g \in G} f(g \cdot h) = \sum_{g \in G} f(h \cdot g) \quad \text{for all } h \in G, \quad (1.1.1)$$

for an arbitrary function f , depending on the group elements.

Finally, we define the **group isomorphism**, which is a group homomorphism that is bijective, i.e., injective and surjective. If there exists such a mapping, the corresponding groups are said to be isomorphic. An example of isomorphic groups are conjugate subgroups K and H of a group G , i.e., $K = gHg^{-1}$, where the isomorphism Φ from H to K is given by $\Phi(h) = g \cdot h \cdot g^{-1}$. Isomorphic groups can be considered identical, up to a relabeling of the group elements.

1.2 Group representations

One of the most powerful concepts in group theory is the so-called group representation, which describes the action of group elements in a linear vector space. This permits to reduce group-theoretical problems to problems in linear algebra.

Let V be a vector space. A **representation** of G is a mapping Γ of the elements of G onto a set of invertible linear operators that map elements of the vector space V again onto elements of the vector space, with the properties

- $\Gamma(e) = \mathbb{1}$, where $\mathbb{1}$ is the identity operator in the vector space V , i.e., $\mathbb{1}v = v$ for all $v \in V$.
- $\Gamma(f \cdot g) = \Gamma(f)\Gamma(g)$, therefore the group multiplication law is mapped onto the composition of the linear operators.

Some comments on terminology: We say the group G acts on the vector space V via the representation Γ . Since the action of $\Gamma(g)$ on any vector in the vector space V is still in the vector space for any group element $g \in G$, we say V is invariant under Γ . The image of Γ , i.e., the set of all linear operators that represent a group element, is called $\Gamma(G)$. V is also referred to as the carrier space of the representation Γ .

From $e = g \cdot g^{-1}$ it follows that $\Gamma(g^{-1}) = \Gamma^{-1}(g)$ for any $g \in G$, i.e., the representation of the inverse group element g^{-1} is given by the inverse linear operator $\Gamma^{-1}(g)$.

For a given vector space V , every group has a trivial representation, given by $\Gamma(g) = \mathbb{1}$ for all $g \in G$. A representation is called faithful, if the mapping of group elements onto linear operators is injective, i.e., $\Gamma(f) = \Gamma(g)$ implies $f = g$. Therefore, the trivial representation is not faithful as long as $|G| \neq 1$.

Let V be finite dimensional, with dimension λ . In this case, we can choose a basis $v = (v_1, \dots, v_\lambda)$, with respect to which the linear operator $\Gamma(g)$ is represented by a matrix $\Gamma^{(v)}(g)$, which is called **matrix representation** of the group element g (with respect to basis v). This uses the isomorphism between linear operators and matrices. Naturally, the dimension of the representation is defined as the dimension λ of the $\lambda \times \lambda$ matrices $\Gamma^{(v)}(g)$. The set of all matrix representations $\Gamma^{(v)}(g)$ for all group elements is denoted by $\Gamma^{(v)}(G)$, and is called matrix representation of the group G . Further, we say the basis v furnishes the matrix representation $\Gamma^{(v)}(G)$ of the group G .

Now let us choose another basis $w = (w_1, \dots, w_\lambda)$. The linear operators $\Gamma(G)$ are therefore now represented by the matrices $\Gamma^{(w)}(G)$. The matrix representations $\Gamma^{(w)}(G)$ and $\Gamma^{(v)}(G)$ are said to be **equivalent matrix representations**, as they describe the same action of group elements in different bases. For any $g \in G$, $\Gamma^{(w)}(g)$ can be calculated by

$$\Gamma^{(w)}(g) = T^{-1}\Gamma^{(v)}(g)T, \quad (1.2.1)$$

where T is the transition matrix between the bases. It can be denoted by $x^{(v)} = Tx^{(w)}$, where $x^{(v)}$ and $x^{(w)}$ are the coordinates of a vector $x \in V$ in the different bases. Note that the matrix T is the same for all group elements. Conversely, Eq. (1.2.1) can be used as the defining equality for equivalent matrix representations.

One can show that every matrix representation $\Gamma^{(v)}(G)$ of a finite group is equivalent to a **unitary matrix representation**, i.e., a matrix representation where

$$\Gamma^{(u)}(g)^{-1} = \Gamma^{(u)}(g)^\dagger \quad \text{for all } g \in G. \quad (1.2.2)$$

The unitary representation $\Gamma^{(u)}(g)$ of the group element g is given by $\Gamma^{(u)}(g) = T^{-1}\Gamma^{(v)}(g)T$, with $T = S^{-1/2}$ and $S = \sum_{g \in G} \Gamma^{(v)}(g)^\dagger \Gamma^{(v)}(g)$, where one can show that S is positive definite, which is why $S^{-1/2}$ exists and is unique.

1.3 Reducibility

As it will turn out, some representations are the “building blocks” of other representations, i.e., the representation can be decomposed into these special representations, which leads us to the concept of **reducibility**: A representation Γ of a group G with vector space V is called reducible if there is a subspace $W \subset V$ that is invariant under Γ , i.e., $\Gamma(g)w \in W$ for all $g \in G$ and $w \in W$. Conversely, the representation is said to be irreducible (or simply an **irrep**), if it is not reducible. We use the notion of reducibility for the carrier space of the representation as well.

Further, a matrix representation $\Gamma^{(v)}(G)$ of a group G is called completely reducible if it is equivalent to a matrix representation $\Gamma^{(w)}(G)$, whose matrices have a block-diagonal form

$$\Gamma^{(w)}(g) = T^{-1}\Gamma^{(v)}(g)T = \begin{pmatrix} \Gamma_1^{(w)}(g) & 0 & \cdots \\ 0 & \Gamma_2^{(w)}(g) & \cdots \\ \vdots & \vdots & \ddots \end{pmatrix}, \quad \text{for all } g \in G, \quad (1.3.1)$$

where $\Gamma_i^{(w)}(g)$, for all i , is an irreducible matrix representation of the group element g , acting on a subspace of V , for all i . Such a matrix representation is said to be the direct sum of subrepresentations $\Gamma_i^{(w)}(g)$, denoted by

$$\Gamma^{(w)}(g) = \sum_{i \oplus} \Gamma_i^{(w)}(g) \equiv \Gamma_1^{(w)}(g) \oplus \Gamma_2^{(w)}(g) \oplus \cdots. \quad (1.3.2)$$

Therefore, the vector space V decomposes into subspaces V_i that are invariant under the irreducible matrix representation $\Gamma_i^{(w)}(G)$ of the group G . Note that in general, several of the occurring irreps in Eq. (1.3.2) can be equivalent. Their respective number of appearances is referred to as the multiplicity of the irrep in the representation $\Gamma(g)$. Also note that there are representations Γ that are reducible, but not completely reducible. In this case, there is an invariant subspace of V under Γ , but its complement is not invariant under Γ .

We now state some important theorems in the context of irreducible matrix representations:

Schur's lemma 1: A matrix M which commutes with all matrices of an irrep $\Gamma^{(v)}(g)$ for all $g \in G$ is proportional to the identity matrix, i.e.,

$$M\Gamma^{(v)}(g) = \Gamma^{(v)}(g)M \Rightarrow M \sim \mathbf{1}. \quad (1.3.3)$$

Schur's lemma 2: Suppose we have two irreps $\Gamma_1^{(v)}(G)$ and $\Gamma_2^{(w)}(G)$ of the group G , acting on vector spaces V and W with dimensions λ_1 and λ_2 and bases v and w , respectively. If there is a $\lambda_2 \times \lambda_1$ matrix M , with $M\Gamma_1^{(v)}(g) = \Gamma_2^{(w)}(g)M$ for all $g \in G$, then either M is zero, or $\Gamma_1^{(v)}(G)$ and $\Gamma_2^{(w)}(G)$ are equivalent. Note that the latter can only happen for $\lambda_1 = \lambda_2$. In this case it also holds that $\det(M) \neq 0$, and therefore $\Gamma_1^{(v)}(g) = M^{-1}\Gamma_2^{(w)}(g)M$.

Orthogonality relations: Let $\Gamma_i^{(v_i)}(G)$ and $\Gamma_k^{(v_k)}(G)$ be two unitary irreps of the group G , which act on the vector spaces with dimensions λ_i and λ_k in the bases v_i and v_k , respectively. Then it holds that

$$\sum_{g \in G} \left(\Gamma_i^{(v_i)}(g)_{\mu\nu} \right)^* \Gamma_k^{(v_k)}(g)_{\mu'\nu'} = \frac{n}{\lambda_i} \delta_{ik} \delta_{\mu\mu'} \delta_{\nu\nu'}, \quad (1.3.4)$$

where the order of the group is denoted by $n \equiv |G|$. For the proof of this relation, one uses Eq. (1.1.1) as well as Schur's lemma 1 and 2. We will use this relation in the following section, where we consider the trace of representations, called characters, yielding the important conclusion that the number of non-equivalent irreps of a group G is given by the number of classes.

1.4 Characters of representations

The **character** of the group element $g \in G$ in representation Γ_i is defined as the trace

$$\chi_i(g) = \text{tr} \left(\Gamma_i^{(v)}(g) \right) = \sum_{\mu} \Gamma_i^{(v)}(g)_{\mu\mu}, \quad (1.4.1)$$

for any basis v of the vector space. What is special about this quantity is that it is the same for equivalent representations $\Gamma_i^{(w)}(g)$ due to the cyclic property of the trace, as

$$\text{tr} \left(\Gamma_i^{(w)}(g) \right) = \text{tr} \left(T^{-1} \Gamma_i^{(v)}(g) T \right) = \text{tr} \left(\Gamma_i^{(v)}(g) T T^{-1} \right) = \text{tr} \left(\Gamma_i^{(v)}(g) \right) \quad \text{for all } g \in G. \quad (1.4.2)$$

Therefore, the characters are independent of the chosen basis.

Furthermore, group elements that belong to the same class have the same character as well, as we have

$$\mathrm{tr} \left(\Gamma(g \cdot p \cdot g^{-1}) \right) = \mathrm{tr} \left(\Gamma(g) \Gamma(p) \Gamma(g^{-1}) \right) = \mathrm{tr} \left(\Gamma(p) \Gamma^{-1}(g) \Gamma(g) \right) = \mathrm{tr} \left(\Gamma(p) \right), \quad (1.4.3)$$

where we wrote $\Gamma \equiv \Gamma_i^{(v)}$ for conciseness. Thus we can write $\chi_i(g) = \chi_i^c$, where c enumerates the classes of the group G and g is in class c . Combining this result with Eq. (1.3.4) yields the orthogonality relation for characters

$$\sum_{c=1}^m n_c (\chi_i^c)^* \chi_k^c = n \delta_{ik}, \quad (1.4.4)$$

which contains a sum over all classes of G , enumerated by c . The total number of classes is m , which is finite for finite groups. The number of group elements in class c is denoted by n_c . The orthogonality relation can be used to show that the number of non-equivalent irreps of a group G is given by the number of classes.

Another crucial property of characters is that for any given completely reducible representation, one can conclude from the characters, whether the representation is reducible. The characters of such a representation $\Gamma(G)$, whose block-diagonal form is given in Eq. (1.3.2), are given by

$$\chi(g) = \sum_i a_i \chi_i(g), \quad (1.4.5)$$

where i labels the non-equivalent irreps of the group G and a_i is the multiplicity of this particular irrep in the representation $\Gamma(G)$. In Eq. (1.3.2) we used basic properties of the trace of block diagonal matrices. Moreover, by using the orthogonality relation for characters, we can conclude

$$\sum_{g \in G} |\chi(g)|^2 = n \sum_i a_i^2. \quad (1.4.6)$$

Now, when $\Gamma(g)$ is irreducible (i.e., exactly one a_i is nonzero, and this particular $a_i = 1$), we get $\sum_g |\chi(g)|^2 = n$, with $n = |G|$. Otherwise, when $\Gamma(g)$ is reducible (i.e., at least two a_i are nonzero, or at least one $a_i \geq 2$), it follows that $\sum_g |\chi(g)|^2 > n$. Thus from the characters alone, we can determine, whether a given representation is reducible.

What is still missing is an algorithm for determining all irreps of a given group, which will be shown later. A fundamental concept of this algorithm is the group algebra, which is introduced in the next section.

1.5 Group algebra

A **group algebra** $\tilde{G}$ is a linear vector space based on a set of elements g_i in which both an addition and a multiplication are defined such that the group properties are satisfied. There is only one exception, as the zero element of the algebra does not have an inverse element.

By definition, the dimension of the vector space is the order of the underlying group $|G|$.

Any element of the group algebra can be written as $\sum_{i=1}^n c_i g_i$ where n is the number of elements g_i . The multiplication in the group algebra is then defined via the multiplication of the group, using:

$$\left(\sum_{i=1}^n c_i g_i \right) \left(\sum_{j=1}^n c_j g_j \right) \equiv \sum_{i,j=1}^n c_i d_j g_i \cdot g_j = \sum_{m=1}^n \bar{c}_m g_m, \quad (1.5.1)$$

which is again an element of the group algebra. In the group algebra, we can write

$$g_i \cdot g_j = \sum_{m=1}^n g_m \Delta_{mj}(g_i) \quad \text{with} \quad \Delta_{mj}(g_i) \equiv \begin{cases} 1 & \text{for } m = k, \\ 0 & \text{else,} \end{cases} \quad (1.5.2)$$

where k is the index of the group element that is obtained by multiplying g_i and g_j , i.e., $g_k = g_i \cdot g_j$. Using this, the coefficients $\bar{c}_m$ in Eq. (1.5.1) are given by

$$\bar{c}_m = \sum_{i,j=1}^n c_i d_j \Delta_{mj}(g_i). \quad (1.5.3)$$

What makes the matrices $\Delta(g)$ special is that they are a (n -dimensional) matrix representation of the group G , as they fulfill both defining properties given in Sec. 1.2:

1. $eg_j = \sum_{m=1}^n g_m \Delta_{mj}(e)$, but also $eg_j = g_j$. As the coefficients must be equal, we find $\Delta_{mj}(e) = \delta_{mj}$ and therefore $\Delta(e) = \mathbb{1}$.
2. Assume $g_a g_b = g_c$. We now calculate the result of applying g_c to a group element g_j . On the one hand, applying Eq. (1.5.2) two times, yields

$$g_a g_b g_j = \sum_{m=1}^n g_a g_m \Delta_{mj}(g_b) = \sum_{k,m} g_k \Delta_{km}(g_a) \Delta_{mj}(g_b) = \sum_k g_k (\Delta(g_a) \Delta(g_b))_{kj}. \quad (1.5.4)$$

On the other hand, we find $g_a g_b g_j = g_c g_j = \sum_k g_k \Delta_{kj}(g_c)$. Again, the coefficients must be equal, from which we get $\Delta(g_a) \Delta(g_b) = \Delta(g_c)$. Thus $\Delta(G)$ is a representation of the group G .

The representation $\Delta(G)$ is called regular representation of the group G . The carrier space of the regular representation is the space of the group algebra $\tilde{G}$.

A crucial property of the regular representation is stated in the following theorem: The regular representation contains all irreps of G , and the multiplicity of an irrep k equals its dimension λ_k . Therefore the representation $\Delta(g_i)$ is equivalent (i.e., equal up to basis transformation) to the matrix representation

$$T^{-1} \Delta(g_i) T = \sum_{k \oplus=1}^m \lambda_k \Gamma^k(g_i) = \lambda_1 \Gamma^1(g_i) \oplus \lambda_2 \Gamma^2(g_i) \oplus \cdots \oplus \lambda_m \Gamma^m(g_i). \quad (1.5.5)$$

Let us simplify this equation for $g_i = e$. In this case, by definition of a representation, we have $\Delta(e) = \mathbb{1}_{n \times n}$ and $\Gamma^k(e) = \mathbb{1}_{\lambda_k \times \lambda_k}$ for all different irreps k , with dimension λ_k .

Further, we take the trace of Eq. (1.5.5), from which we get on the one hand

$$\text{tr} \left(T^{-1} \Delta(e) T \right) = \text{tr} (\Delta(e)) = n \quad (1.5.6)$$

and on the other hand

$$\mathrm{tr} \left(\sum_{k \oplus = 1}^m \lambda_k \Gamma^k(e) \right) = \sum_{k=1}^m \lambda_k \mathrm{tr} \left(\Gamma^k(e) \right) = \sum_{k=1}^m \lambda_k^2, \quad (1.5.7)$$

where we used basic properties of the trace of block diagonal matrices. Therefore we get $n = \sum_{k=1}^m \lambda_k^2$, which is a fundamental identity relating the order of a group to the dimensions of its irreps.

For later use, we also define the following terms:

A subalgebra $\tilde{H} \subset \tilde{G}$ is a vector space that is contained in the group algebra $\tilde{G}$ which is closed under the law of multiplication that is given by the group algebra $\tilde{G}$. Further, if the subalgebra is invariant under left multiplication, i.e., $\tilde{g}\tilde{h} \in \tilde{H}$ for all $\tilde{g} \in \tilde{G}$ and for all $\tilde{h} \in \tilde{H}$, it is called a left ideal. If the subalgebra $\tilde{H}$ in turn does not have any subalgebras which are left ideals, it is called an irreducible left ideal. An irreducible left ideal is especially relevant, as one can extract an irrep from it, by acting with each group element on the basis vectors of the left ideal.

An element $\tilde{g}$ of the group algebra satisfying $\tilde{g}^2 = \tilde{g}$ is called idempotent. If $\tilde{g}^2$ is proportional to $\tilde{g}$, i.e., $\tilde{g}^2 \propto \tilde{g}$, the element is called essentially idempotent.

Additionally, an idempotent is called irreducible or primitive, if it generates an irreducible left ideal, i.e., when $\{\tilde{h}\tilde{g} \mid \tilde{h} \in \tilde{G}\}$ is an irreducible left ideal, which eventually provides an irrep.

It can be shown that all irreps of a group can be constructed, once all primitive idempotents are known. The construction of the primitive idempotents is demonstrated for the so-called symmetric group, which will be introduced in the next section.

The primitive idempotents of S_n are further relevant, as they are also used to construct irreps of some important continuous groups, e.g., $U(m)$ and $SU(m)$ (cf. Sec. 1.9). Further, the study of the group S_n is important due to Cayley's theorem, which states that every group of order n is isomorphic to a subgroup of the symmetric group S_n .

1.6 The symmetric group

The **symmetric group** S_n on a finite set X with n elements is the group whose elements are all bijective functions from X to X and whose group operation is the function composition. In the following, we restrict ourselves without loss of generality to $X = \{1, \dots, n\}$, for which such a bijection π (we call it permutation) is given by $\pi(i) = \pi_i$ for $i \in X$ or equivalently

$$\begin{pmatrix} 1 & 2 & \cdots & n \\ \pi_1 & \pi_2 & \cdots & \pi_n \end{pmatrix}, \quad (1.6.1)$$

where all π_i are distinct and again in X , i.e., $\{\pi_i \mid i \in X\} = X$. Therefore, the symmetric group S_n is the group of permutations of n objects. There are in total $n!$ permutations, i.e., $|S_n| = n!$.

For $\pi, \sigma \in S_n$, the function composition can be denoted by $(\sigma \cdot \pi)(i) \equiv \sigma(\pi(i)) = \sigma(\pi_i)$.

The identity element of the group is given by the permutation $\pi(i) = i$. S_n is non-abelian for $n > 2$.

We can denote a permutation of S_n also with the use of cycles. A cycle of length k (or simply k -cycle) is given by the notation $(a_1 a_2 \cdots a_k)$ with distinct numbers $a_i \in X$. This is a shorthand notation for the permutation π , with $\pi(a_i) = a_{(i+1) \% k}$ for all $i \leq k$ that leaves all other elements unchanged, i.e., $\pi(i) = i$ for all i with $i \neq a_j$ for all $j \in \{1, \dots, k\}$. A permutation that can be represented by a cycle is called cyclic permutation.

Note that the cycle $(a_1 a_2 \cdots a_k)$ represents the same permutation, when the numbers within the cycle are interchanged cyclically. However, in general the order of the numbers is relevant. Cycles of length 2 are called transpositions.

Two cycles $(a_1 a_2 \cdots a_{k_1})$ and $(b_1 b_2 \cdots b_{k_2})$ are said to be disjoint when they do not have any elements in common, i.e., $a_i \neq b_j$ for all $i \leq k_1$ and all $j \leq k_2$. Two cyclic permutations are commutative, when the corresponding cycles are disjoint.

Any permutation $\pi \in S_n$ can be written as a product of pairwise disjoint cyclic permutations. An example of a permutation written as a product of cycles is given by

$$\begin{pmatrix} 1 & 2 & 3 & 4 & 5 & 6 \\ 6 & 5 & 3 & 1 & 2 & 4 \end{pmatrix} = (164) \cdot (25) \cdot (3), \quad (1.6.2)$$

where the order of the cycles is irrelevant, as the cyclic permutations are commutative. The example consists of one 3-cycle, one 2-cycle and one 1-cycle. In general, this information is contained in the so-called cycle structure of a permutation, which lists the lengths of all cycles in descending order. Formally, it is denoted by $(\lambda_1, \lambda_2, \dots, \lambda_r)$ with $\lambda_i \geq \lambda_{i+1}$, where λ_i is the length of a given cycle and r is the number of appearing cycles. In the example (1.6.2), the cycle structure is $(3, 2, 1)$. Every cycle structure is also a partition of n , where a partition $\lambda \equiv (\lambda_1, \lambda_2, \dots, \lambda_r)$ of an integer n is defined as the sequence of positive integers such that $\sum_{i=1}^r \lambda_i = n$ with $\lambda_i \geq \lambda_{i+1}$.

Any cycle of length k with $k \geq 2$ can be written as a product of transpositions, e.g., $(164) = (14) \cdot (16)$. Note that the transpositions are not disjoint, and therefore the permutations are not commutative.

Therefore any permutation can be written as product of s transpositions. If s is even we call π an even permutation, otherwise it is an odd permutation. We denote this by the so-called parity $\text{sgn}(\pi) \equiv (-1)^s$.

For the symmetric group S_n , the subset of all even permutations builds an invariant subgroup of S_n . This group is called alternating group A_n . The factor group S_n/A_n is isomorphic to S_2 .

Two elements of S_n are equivalent (i.e., are in the same class) if they have the same cycle structure. As an example, we consider two permutations with cycle structure $(3, 2, 1)$,

$$(123)(4)(56) = \pi^{-1}(142)(5)(63)\pi, \quad (1.6.3)$$

where π is the permutation which is given by

$$\begin{pmatrix} 1 & 2 & 3 & 4 & 5 & 6 \\ 1 & 4 & 2 & 5 & 6 & 3 \end{pmatrix}. \quad (1.6.4)$$

This is the reason why the analysis of cycle structures - and thus of partitions - is of central importance. To facilitate the study of partitions, we introduce a visualization in the next section.

1.7 Construction of the irreps of the symmetric group

A partition $\lambda \equiv (\lambda_1, \lambda_2, \dots, \lambda_r)$ of an integer n with $\sum_{i=1}^r \lambda_i = n$ and $\lambda_i \geq \lambda_{i+1}$ can be represented graphically by a **Young diagram**. Hereby, n boxes are arranged in r rows such that the i -th row contains λ_i boxes. For example, the cycle structures of the symmetric group S_4 are visualized as

$$\begin{array}{c} \boxed{} \boxed{} \boxed{} \boxed{} \\ \hline \end{array}, \quad \begin{array}{c} \boxed{} \boxed{} \boxed{} \\ \boxed{} \end{array}, \quad \begin{array}{c} \boxed{} \boxed{} \\ \boxed{} \boxed{} \end{array}, \quad \begin{array}{c} \boxed{} \boxed{} \\ \boxed{} \\ \boxed{} \end{array}, \quad \begin{array}{c} \boxed{} \\ \boxed{} \\ \boxed{} \\ \boxed{} \end{array}, \quad (1.7.1)$$

representing the partitions (4) , $(3,1)$, $(2,2)$, $(2,1,1)$, and $(1,1,1,1)$, respectively. Further, a **Young tableau** is a Young diagram containing the numbers 1 to n in the n boxes, where each number appears exactly once, e.g.,

$$\begin{array}{c} \boxed{2} \boxed{3} \boxed{1} \\ \boxed{4} \end{array}, \quad \begin{array}{c} \boxed{4} \boxed{1} \\ \boxed{2} \boxed{3} \end{array}. \quad (1.7.2)$$

There is a one-to-one correspondence between permutations and Young tableaux. For the Young tableaux given in (1.7.2), the corresponding permutations are given by the product of cycles $(231) \cdot (4)$ and $(41) \cdot (23)$, respectively. Additionally, when the numbers from 1 to n appear in ascending order, first from left to right, then from top to bottom, the tableau is called **normal Young tableau**, e.g.,

$$\begin{array}{c} \boxed{1} \boxed{2} \boxed{3} \\ \boxed{4} \end{array}, \quad \begin{array}{c} \boxed{1} \boxed{2} \\ \boxed{3} \boxed{4} \end{array}. \quad (1.7.3)$$

For every Young diagram there is only one normal Young tableau. In a **standard Young tableau** the numbers increase in the rows and columns, but not necessarily in strict order. Two examples for such tableaux are

$$\begin{array}{c} \boxed{1} \boxed{2} \boxed{4} \\ \boxed{3} \end{array}, \quad \begin{array}{c} \boxed{1} \boxed{3} \\ \boxed{2} \boxed{4} \end{array}. \quad (1.7.4)$$

The normal Young tableau of the partition λ is denoted by Θ_λ . Any other Young tableau corresponding to the same Young diagram is obtained from Θ_λ by a permutation π of the n numbers. This is denoted by

$$\Theta_\lambda^\pi = \pi \Theta_\lambda, \quad \text{where we have} \quad \sigma \Theta_\lambda^\pi = \Theta_\lambda^{\sigma \cdot \pi}. \quad (1.7.5)$$

As we will see, for every Young tableau one can define a primitive idempotent that generates an irrep of S_n on the group algebra $\tilde{S}_n$.

For the construction of these primitive idempotents one needs the so-called symmetrizer and anti-symmetrizer. For given Young tableau with shape λ and filling π , the symmetrizer is defined as $s_\lambda^\pi \equiv \sum_{h_\lambda^\pi} h_\lambda^\pi$, where h_λ^π are all permutations, that only permute numbers in the rows of Θ_λ^π . Note that h_λ^π is not restricted to permutations in one single row.

Additionally, we define the antisymmetrizer $a_\lambda^\pi \equiv \sum_{v_\lambda^\pi} (-1)^{\text{sgn}(v_\lambda^\pi)} v_\lambda^\pi$, where v_λ^π are all permutations, that only permute numbers in the columns of Θ_λ^π .

Note that the symmetrizer and the antisymmetrizer are elements of the group algebra $\tilde{S}_n$. From these definitions one can construct the **Young operator** $e_\lambda^\pi \equiv s_\lambda^\pi a_\lambda^\pi$, which is again an element of the group algebra $\tilde{S}_n$.

One can show that e_λ^π are essentially primitive idempotents for all π and all λ . Therefore, every e_λ^π generates an irreducible left ideal L_λ^π . Using the approach explained at the end of Sec. 1.5, we thus can construct all irreducible representations of S_n .

Further, different Young tableaux obtained from the same Young diagram give equivalent irreps. Only representations corresponding to different Young diagrams are inequivalent. Thus the Young operator corresponding to the normal Young tableaux already generates all inequivalent irreps.

As an example, we consider the normal Young tableau

$$\begin{array}{|c|c|} \hline 1 & 2 \\ \hline 3 & \\ \hline \end{array}, \quad (1.7.6)$$

which represents an irrep of S_3 . The symmetrizer of this tableau is given by $s_{(2,1)}^e = e + (12) \cdot (3)$, as the identity element e of S_3 and $(12) \cdot (3)$ are the only permutations that only permute numbers in the rows of the tableau. Likewise, the antisymmetrizer is given by $a_{(2,1)}^e = e - (13) \cdot (2)$, which contains a minus since the permutation $(13) \cdot (2)$ has negative parity. Therefore, the Young operator

$$e_{(2,1)}^e = s_{(2,1)}^e a_{(2,1)}^e = (e + (12) \cdot (3))(e - (13) \cdot (2)) = e + (12) \cdot (3) - (13) \cdot (2) - (321) \quad (1.7.7)$$

generates an irreducible left ideal and thus also an irrep of S_3 . For this example, the dimension of the irreducible left ideal and therefore also of the irrep is two. One can show that in general the dimension of the irrep Γ^λ corresponding to the Young diagram with shape λ equals the number of standard Young tableaux of λ and is given by

$$n_\lambda = \frac{n!}{\prod_{i=1}^r \prod_{k=1}^{\lambda_i} h_{i,k}}, \quad (1.7.8)$$

where $h_{i,k}$ is the hook length of the box in row i and column k . For the box i, k the **hook** consists of all boxes below and to the right of it, plus the box itself. The **hook length** is then given by the number of involved boxes.

From the calculated irreps one can in principle take the trace to obtain the characters for all classes in all irreps of the symmetric group. However, there is an efficient algorithm known as the Murnaghan-Nakayama rule, which allows for a computation of the characters without explicitly calculating all irreps. In the next section, this algorithm is presented and the character for a given class and given irrep is calculated as an example.

1.8 Murnaghan-Nakayama rule

In the following, we go through all necessary steps for calculating the character of the symmetric group for a given class in a given irrep. The proof of this algorithm is omitted here. Before we dive into the details of calculating the character, we need some further definitions related to Young diagrams, which are given next.

The **boundary** of a Young diagram, consists of all boxes with row-index i and column-index j for which the position $i + 1, j + 1$ is already outside the diagram. As an example, we consider the Young diagram corresponding to the partition $(5, 4, 2)$

$$\begin{array}{|c|c|c|c|c|} \hline & & & 2 & 1 \\ \hline & 5 & 4 & 3 & \\ \hline 7 & 6 & & & \\ \hline \end{array}, \tag{1.8.1}$$

where we enumerated all boxes that belong to the boundary from one to seven. Further, a connected part of the boundary that can be removed to leave a proper Young diagram is called **skew hook**. For the Young diagram (1.8.1), there are in total eleven skew hooks. For example the boxes 1 to 3 build a skew hook, since removing them would yield the Young diagram that corresponds to the partition $(3, 3, 2)$. Similarly, the box 1 itself is a skew hook. Its removal corresponds to the partition $(4, 4, 2)$. The remaining skew hooks are 1-4, 1-6, 1-7, 3, 3-4, 3-6, 3-7, 6, and 6-7.

Note that every skew hook corresponds to a hook and vice versa. Hereby, the beginning and ending of hook and corresponding skew hook coincide and lie on the boundary of the Young diagram. As every hook in turn corresponds to one box and vice versa, there are as many skew hooks as boxes in an arbitrary Young diagram. Consider for example the skew hook 1-7. The corresponding hook is that one, that is constructed from the box in row 1 and column 1.

Further, the **length** of a skew hook is the number of involved boxes. Therefore, the skew hook that contains the boxes 1 to 7 has length 7. This length always equals the hook length of the associated hook.

Additionally, we define the **leg length** of a skew hook, which is the number of vertical steps, i.e., the number of rows in the skew hook minus one.

Now, consider a given Young diagram (i.e., an irrep of S_n) corresponding to the partition λ together with a class $c = (c_1, c_2, \dots, c_r)$ of S_n , for which we want to calculate the character. Here, the class c is given in cycle structure, where c_i is the length of the i -th cycle with a total

number of r cycles. Then the corresponding character is given by the recursive formula

$$\chi_c^\lambda = \sum_{\tilde{\lambda}} \pm \chi_{\tilde{c}}^{\tilde{\lambda}}, \quad (1.8.2)$$

which is explained in the following: To apply this equation, we can first freely choose one of the lengths c_i in the cycle structure and determine all skew hooks in the Young diagram that have length c_i . By definition, the removal of such a skew hook yields another (smaller) Young diagram. Using this, the sum in Eq. (1.8.2) runs over all diagrams (i.e., partitions) $\tilde{\lambda}$, that are obtained from the diagram λ , when one of the skew hooks of length c_i is removed. $\tilde{c}$ is the class that is obtained by removing one cycle of length c_i from the class c . Finally, the sign in Eq. (1.8.2) is positive if the skew hook that we removed from λ in order to obtain $\tilde{\lambda}$ has even leg length, and it is negative if the removed skew hook has odd leg length.

By applying this formula, one is left with the task of calculating characters for smaller classes and smaller Young diagrams than we had at the beginning. Therefore we can recursively apply this formula, until the algorithm stops due to one of the following two scenarios: Either there is no skew hook of length c_i , in which case the summation is empty and therefore $\chi_c^\lambda = 0$; or, if the removed skew hook covers the whole diagram. For this case we set $\chi_{(0)}^{(0)} \equiv 1$.

To make this method efficient, the order in which the c_i are removed should be chosen to minimize the total number of skew hooks to be removed. For example, if there is a length c_i for which there is no skew hook with length c_i , this length is always preferred.

In the following, we calculate the character of the class $(7, 2, 2)$ for the partition $(5, 4, 2)$, as an example. The underlying group is S_{11} and the corresponding Young diagram is already drawn in (1.8.1), where also all the skew hooks are stated below this diagram.

For efficiency reasons, we now choose c_1 , which is 7, to begin with. There is only one skew hook of length 7 (it is the skew hook that contains the boxes 1-7). The leg length of this skew hook is 2, which is even, and therefore we get a plus in the formula (1.8.2).

The diagram we get, when we remove this skew hook corresponds to the partition $(3, 1)$, and the class that we get when we remove $c_i = 7$ is $(2, 2)$. Therefore we obtain

$$\lambda_{(7,2,2)}^{(5,4,2)} = +\lambda_{(2,2)}^{(3,1)}. \quad (1.8.3)$$

After a renaming, we repeat the process for given partition $\lambda = (3, 1)$ and given class $c = (c_1, c_2) = (2, 2)$. The corresponding Young diagram with enumerated boundary is given by

$$\begin{array}{|c|c|c|} \hline 3 & 2 & 1 \\ \hline 4 & & \\ \hline \end{array}. \quad (1.8.4)$$

As this diagram contains four boxes, there are four skew hooks as well. They read 1, 1-2, 1-4, and 4.

We now choose the cycle $c_1 = 2$ to continue. In the above diagram, there is only one skew hook of length 2, which is the skew hook that contains the boxes 1-2. The leg length of this skew hook is 0, yielding a plus in the recursive formula.

After removing this skew hook we get the partition $(1, 1)$. The obtained class after removing $c_1 = 2$ is (2) . Therefore we get

$$\lambda_{(2,2)}^{(3,1)} = +\lambda_{(2)}^{(1,1)}. \quad (1.8.5)$$

We repeat the algorithm once more, now for given partition $\lambda = (1, 1)$ and given class $c = (c_1) = (2)$. The relevant Young diagram, together with its boundary, reads

$$\begin{array}{|c|} \hline 1 \\ \hline 2 \\ \hline \end{array}. \quad (1.8.6)$$

Now there are only two skew hooks, containing either the boxes 1-2 or the box 2. The remaining cycle length has $c_1 = 2$. Again, there is only one skew hook of length 2, which is the skew hook that contains the whole diagram. In this case we use the definition $\chi_{(0)}^{(0)} = 1$. Therefore,

$$\lambda_{(2)}^{(1,1)} = -\chi_{(0)}^{(0)} = -1, \quad (1.8.7)$$

where we get a minus, as the leg length of the removed skew hook is 1. Thus, for the class $(7, 2, 2)$, the character in the irrep $(5, 4, 2)$ is given by

$$\lambda_{(7,2,2)}^{(5,4,2)} = -1. \quad (1.8.8)$$

1.9 Lie groups

So far, we mainly discussed finite groups, especially the symmetric group S_n . Now we proceed with continuous groups. These are groups where the group elements $g \in G$ are continuous functions of a finite number of real parameters $(\alpha_1, \alpha_2, \dots, \alpha_n) \equiv \alpha$, i.e.,

$$g = g(\alpha_1, \dots, \alpha_n) = g(\alpha). \quad (1.9.1)$$

Continuous in this sense means that there is some notion of closeness such that two group elements $g(\alpha)$ and $g(\alpha')$ are close together, when the parameters of two group elements are close together, i.e., $\|\alpha - \alpha'\|$ gets very small. Here, $\|\cdot\|$ denotes the euclidean norm. Note that the parameterization of the group elements is not unique. In fact, there is an infinite amount of possible parameterizations. The **order** of a continuous group no longer refers to the number of elements in the group - which is infinite - but rather to the number of parameters n in the group.

The multiplication law for the group now reads

$$g(\alpha) \cdot g(\alpha') = g(\alpha''), \quad (1.9.2)$$

where the α''_i for $i \in \{1, \dots, n\}$ in general depends on all parameters α and α'' . We denote this by n different functions f_i , i.e.,

$$\alpha''_i = f_i(\alpha_1, \dots, \alpha_n, \alpha'_1, \dots, \alpha'_n) \quad \text{for } i \in \{1, \dots, n\}. \quad (1.9.3)$$

Note that by stating these functions f_i the whole structure of the group is determined. Due to the defining properties of a group, there is an identity group element $e \in G$ with parameters $\alpha^0 = (\alpha_1^0, \dots, \alpha_n^0)$ such that

$$\alpha_i = f_i(\alpha_1, \dots, \alpha_n, \alpha_1^0, \dots, \alpha_n^0) = f_i(\alpha_1^0, \dots, \alpha_n^0, \alpha_1, \dots, \alpha_n) \quad \text{for } i \in \{1, \dots, n\}. \quad (1.9.4)$$

for any group element $g \in G$ parameterized by α . Further, due to the existence of an inverse element $g^{-1} \in G$, parameterized with $\bar{\alpha} = (\bar{\alpha}_1, \dots, \bar{\alpha}_n)$, we get the condition

$$\alpha_i^0 = f_i(\alpha_1, \dots, \alpha_n, \bar{\alpha}_1, \dots, \bar{\alpha}_n) = f_i(\bar{\alpha}_1, \dots, \bar{\alpha}_n, \alpha_1, \dots, \alpha_n) \quad \text{for } i \in \{1, \dots, n\}. \quad (1.9.5)$$

When these functions can be inverted, i.e.,

$$\bar{\alpha}_i = h_i(\alpha_1, \dots, \alpha_n, \alpha_1^0, \dots, \alpha_n^0) \quad \text{for } i \in \{1, \dots, n\}, \quad (1.9.6)$$

and when the functions f_i and h_i are analytic for all i , the group is called **Lie group**. Moreover, a Lie group is called **compact** if the intervals of the parameters α_i are finite and closed. Otherwise it is called non-compact.

A typical example of a compact Lie group is $\text{GL}(N, \mathbb{C})$, which is the group of all non-singular, linear transformations in an N -dimensional complex vector space. Its defining (or fundamental) representation is given by the invertible complex $N \times N$ matrices. Therefore, each group element is determined by $2N^2$ real parameters.

An important subgroup of $\text{GL}(N, \mathbb{C})$ is the **unitary group** $\text{U}(N)$, whose elements U are additionally unitary, i.e., they fulfill

$$U^\dagger U = U U^\dagger = \mathbb{1} \quad \text{for all } U \in \text{U}(N), \quad (1.9.7)$$

where $\mathbb{1}$ is the N -dimensional identity matrix. The group elements of the unitary group can be parameterized by N^2 real parameters.

A subgroup of $\text{U}(N)$ in turn is the special unitary group $\text{SU}(N)$, whose elements U have a determinant that is equal to one, i.e., $\det(U) = 1$ for all U in $\text{SU}(N)$. The group elements are parameterized by $N^2 - 1$ parameters.

One crucial change between continuous and finite groups is that the summation over group elements in the relation

$$\sum_{g \in G} f(g) = \sum_{g \in G} f(g \cdot h) = \sum_{g \in G} f(h \cdot g) \quad \text{for all } h \in G, \quad (1.1.1)$$

derived from the rearrangement lemma, must be properly replaced by an integration over the parameters $\alpha_1, \dots, \alpha_n$ of the group G . To obtain an analogue integral, we need an integration measure

$$d\mu(g) \equiv d\mu(\alpha_1, \dots, \alpha_n) \equiv \rho(\alpha_1, \dots, \alpha_n) d\alpha_1 \cdots d\alpha_n \quad (1.9.8)$$

with density ρ such that

$$\int_G d\mu(g) f(g) = \int_G d\mu(g) f(g \cdot h) = \int_G d\mu(g) f(h \cdot g) \quad \text{for all } h \in G. \quad (1.9.9)$$

As this equation again should hold for any function f , we require

$$d\mu(g) = d\mu(g \cdot h) = d\mu(h \cdot g) \quad \text{for all } g \in G. \quad (1.9.10)$$

An integration measure that fulfills this equation is called **Haar measure**. One can show, that for a compact Lie group, there exists such a Haar measure, which is also unique after setting the normalization $\mu(G) = \int_G d\mu(g) = 1$.

After introducing the Haar measure, one can show that a compact Lie group has a countable number of inequivalent irreps, all of which are finite-dimensional. Further, due to the existence of the Haar measure, many crucial results derived for finite groups also hold, in a slightly modified version, for compact Lie groups. For example, for a compact Lie group, every matrix representation is equivalent to a unitary representation, and every representation can be decomposed into a sum of irreps.

For a compact Lie group, there are also orthogonality relations for unitary irreps as in 1.3.4, which are given by

$$\int_G d\mu(g) \left(\Gamma_i^{(v_i)}(g)_{\mu\nu} \right)^* \Gamma_k^{(v_k)}(g)_{\mu'\nu'} = \frac{1}{\lambda_i} \delta_{ik} \delta_{\mu\mu'} \delta_{\nu\nu'} . \quad (1.9.11)$$

From these, we also get orthogonality relations for characters

$$\int_G d\mu(g) (\chi_i(g))^* \chi_k(g) = \delta_{ik} , \quad (1.9.12)$$

where $\chi_i(g)$ is the character of the group element g in irrep i . Additionally, just as in Eq. (1.4.6), from the characters alone we can conclude, whether a given representation of a compact Lie group is reducible. Specifically, the condition

$$\int_G d\mu(g) |\chi_i(g)|^2 = 1 , \quad (1.9.13)$$

is valid if and only if the irrep i is irreducible.

Since compact Lie groups are continuous groups, where additionally the identity element in any representation is given by the identity matrix, we can explicitly parameterize any compact Lie group in the vicinity of the identity. To do so, we choose the parameterization of the group such that $\alpha_i^0 = 0$ for all i , i.e., $g(0, \dots, 0) = e$ is the identity element.

Therefore, by definition, we have for any representation of the compact Lie group

$$\Gamma(g(\alpha))|_{\alpha=0} = \mathbb{1} . \quad (1.9.14)$$

For group elements close to the identity that is parameterized by $\varepsilon = (\varepsilon_1, \dots, \varepsilon_n)$, we then can parameterize any representation by

$$\Gamma(\varepsilon_1, \dots, \varepsilon_n) = \exp \left(i \sum_{j=1}^n \varepsilon_j S_j \right) , \quad (1.9.15)$$

where the linear operators S_j act on the carrier space of Γ and are called **generators** of the Lie group. Note that there is one generator per parameter.

As a direct consequence, one can show that for the special groups (i.e., $\det(\Gamma(g)) = 1$) the generators are traceless, and for unitary groups (i.e., $\Gamma(g) (\Gamma(g))^\dagger = \mathbb{1}$) the generators are hermitian, i.e., $S_j = S_j^\dagger$ for all j . In general, the generators of a given representation Γ can be extracted using the formula

$$S_j = -i \frac{\partial \Gamma(g(\alpha))}{\partial \alpha_j} \Big|_{\alpha=0} . \quad (1.9.16)$$

We now consider the product of two group elements that are both in the vicinity of the identity such that the result is again in the vicinity of the identity, i.e., $\Gamma(\varepsilon)\Gamma(\varepsilon') = \Gamma(\varepsilon'')$. By using the Baker-Campbell-Hausdorff formula, the product becomes

$$\Gamma(\varepsilon)\Gamma(\varepsilon') = \exp \left(i \sum_{j=1}^n \varepsilon_j S_j + i \sum_{k=1}^n \varepsilon'_k S_k - \frac{1}{2} \sum_{j,k=1}^n \varepsilon_j \varepsilon'_k [S_j, S_k] + \dots \right), \quad (1.9.17)$$

which needs to be expressed in the form of Eq. (1.9.15), in order to determine ε'' . However, for arbitrary ε and ε' this is only possible when the commutator of the generators is again a linear combination of the generators [38], i.e.,

$$[S_j, S_k] = i f_{jkl} S_l \quad \text{for all } j, k \in \{1, \dots, n\}, \quad (1.9.18)$$

where the imaginary unit is a matter of definition. The coefficients f_{jkl} are called **structure constants**. By using this commutator relation several times, even the inclusion of higher order terms to an arbitrary order in Eq. (1.9.17) (which yields nested commutators) yields just a linear combination of the n generators, i.e., a representation of a group element in the form of Eq. (1.9.15). Using this, we can calculate ε'' up to arbitrary order in ε and ε' . Therefore, the structure constants contain the whole information of the multiplication law of the group near the identity.

From the properties of the commutator, we can conclude some useful properties of the structure constants. First, from the definition of the commutator, it follows directly that

$$f_{jkl} = -f_{kjl} \quad \text{for all } j, k, l \in \{1, \dots, n\}. \quad (1.9.19)$$

Next, due the Jacobi identity for commutators, the structure constants fulfill a Jacobi identity as well:

$$\sum_{a=1}^n f_{jka} f_{lab} + f_{lja} f_{kab} + f_{kla} f_{jab} = 0 \quad \text{for all } b, j, k, l \in \{1, \dots, n\}. \quad (1.9.20)$$

For a unitary representation, it can be shown that all structure constants f_{jkl} are real.

There is another important consequence of Eq. (1.9.18): As the commutator of two generators is again a linear combination of generators, the vector space with basis vectors S_i for $i \in \{1, \dots, n\}$ is closed with respect to the commutator. Thus we can define an algebra, where the multiplication law (i.e., composition law), is defined as the commutator of two elements. This algebra is called **Lie algebra**. Also the Lie algebra is completely determined by the structure constants.

The structure constants themselves satisfy the commutator relation (1.9.18). For this reason, we say that they form a representation of the Lie algebra. This representation is n -dimensional, and called **adjoint representation**. Formally, it is defined as $\Gamma(S_j)_{lk} = f_{jkl}$.

2 Fundamental concepts of quantum field theory

Quantum field theory (QFT) is one of the most significant advances in physics of the 20th century as it connects both quantum mechanics and special relativity while providing the most remarkable agreement between theoretical predictions and experimental data in the history of science [39].

A crucial conceptual shift in QFT compared to earlier theories is that the fundamental objects of the theory are no longer particles but fields. Particles are then regarded as quantized excitations of these underlying fields. This allows for a natural quantum mechanical description of multi-particle states in which the number of particles is not conserved. In general, this is a crucial property in all scattering and decay processes.

With the exception of gravity, QFT is used to describe all fundamental forces observed in the universe. These are the electromagnetic, strong, and weak interaction between elementary particles. The mediators of these forces are called gauge bosons and together with quarks, leptons, and the Higgs boson they build the standard model of particle physics.

The theory that describes the strong interaction is QCD. As mentioned in the introduction, it is a gauge theory based on the non-abelian group $SU(3)$, whose fundamental particles (called quarks) carry a color charge that occurs in three different types. The eight different mediators of the color charge called gluons describe the strong interactions between the quarks. For each of the three color charges there is an anti-color charge, which are carried by anti-quarks. In contrast to the photon in Quantum Electrodynamics, the gauge bosons of QCD also carry a color charge in the form of color-anticolor combinations, which allows for gluon self-interactions. Further, all observed composite states of the strong interaction, i.e., hadrons, are color singlets of $SU(3)$.

QCD is especially relevant, as it has a rich structure of phenomena, which differ crucially from all other interactions. Among its most remarkable features are the previously mentioned phenomena of confinement and asymptotic freedom. Both are closely connected to the fact that $SU(3)$ is a non-abelian gauge group [40].

In the following, we first introduce the path integral formalism for one-dimensional ordinary quantum mechanics, as well as for a scalar field theory. This approach is a basic formalism that is also used in modern quantum field theory. Next, we formally introduce QCD by defining its Lagrangian density and discuss some crucial properties. After introducing lattice QCD, we will take a look at the phase diagram of QCD. The overview is based on Refs. [2, 41, 42, 43, 44, 45].

2.1 Path integral formulation in quantum mechanics

We start with the path integral formulation for a particle in one dimension (cf. Ref. [46]). A particle at position x is described by the state $|x\rangle$ in position space, which forms an infinite-dimensional Hilbert-space. The states $|x\rangle$ for $x \in \mathbb{R}$ are an orthogonal basis with normalization

$\langle x|y\rangle = \delta(x-y)$, which fulfill the completeness relation

$$\hat{1} = \int_{-\infty}^{\infty} dx |x\rangle \langle x|. \quad (2.1.1)$$

The time evolution of the particle is determined by a Hamilton operator $\hat{H}$, where the probability of measuring a particle with initial position x_i after time T at position x_f is given by $|\langle x_f|e^{-i\hat{H}T}|x_i\rangle|^2$. Throughout the thesis, we use natural units, i.e., $\hbar = c = 1$.

To rewrite this expression in terms of a path integral, we divide the time interval $[0, T]$ into N small segments of length δt such that $T = \sum_{i=1}^N \delta t = N\delta t$.

Since any operator commutes with itself, we can write

$$\begin{aligned} \langle x_f|e^{-i\hat{H}T}|x_i\rangle &= \langle x_f|e^{-i\hat{H}\delta t}e^{-i\hat{H}\delta t}\dots e^{-i\hat{H}\delta t}|x_i\rangle \\ &= \left(\prod_{j=1}^{N-1} \int dx_j\right) \langle x_f|e^{-i\hat{H}\delta t}|x_{N-1}\rangle \langle x_{N-1}|e^{-i\hat{H}\delta t}|x_{N-2}\rangle \dots \langle x_1|e^{-i\hat{H}\delta t}|x_i\rangle, \end{aligned} \quad (2.1.2)$$

where we inserted the completeness relation (2.1.1) between any two exponentials.

We now consider the simple case $\hat{H} = \frac{\hat{p}^2}{2m} + V(\hat{x})$, with eigenvalue equations $\hat{p}|p\rangle = p|p\rangle$ and $\hat{x}|x\rangle = x|x\rangle$. With the use of the approximation

$$\exp(-i\hat{H}\delta t) = \exp\left(-i\frac{\hat{p}^2}{2m}\delta t\right) \exp(-iV(\hat{x})\delta t) + \mathcal{O}(\delta t^2) \quad (2.1.3)$$

and the completeness relation $\hat{1} = \int dp |p\rangle \langle p|$ for the momenta, we can eliminate all operators appearing in Eq. (2.1.2) by applying the eigenvalue equations. After using $\langle x|p\rangle = e^{ipx}$, all the integrals over momenta can be evaluated analytically using Gaussian integrals.

After taking the limit $N \rightarrow \infty$, the transition amplitude is written in the **path integral formulation**

$$\langle x_f|e^{-iT\hat{H}}|x_i\rangle = \int_{x(0)=x_i}^{x(T)=x_f} [\mathcal{D}x(t)] e^{iS[x(t),\dot{x}(t)]}, \quad (2.1.4)$$

where the integration

$$\int [\mathcal{D}x(t)] = \lim_{N \rightarrow \infty} \left(\frac{-im}{2\pi\delta t}\right)^{\frac{N}{2}} \left(\prod_{k=1}^{N-1} \int dx_k\right) \quad (2.1.5)$$

considers all possible paths $x(t)$ with fixed boundaries $x(0) = x_i$ and $x(T) = x_f$. The right hand side of Eq. (2.1.4) contains the action

$$S[x(t), \dot{x}(t)] = \int_0^T dt L[x(t), \dot{x}(t)], \quad (2.1.6)$$

with Lagrange function $L[x(t), \dot{x}(t)] = \frac{1}{2}m\dot{x}^2 - V(x)$. Note that in the path integral formalism, all the operators are replaced by scalars, which makes the formalism predestined for numerical calculations when $\lim_{N \rightarrow \infty}$ is approximated by using large values of N .

The essential property of Eq. (2.1.4) is that it is valid for any Hamilton operator $\hat{H}$ that is a function of products and sums of $\hat{x}$ and $\hat{p}$, as long as it is written in Weyl ordering. In Weyl

ordering, all terms are symmetrized so that the operator sequence reads the same forwards and backwards, e.g., $\hat{x}^2\hat{p}^2 + 2\hat{x}\hat{p}^2\hat{x} + \hat{p}^2\hat{x}^2$. Any Hamiltonian which consists of operators $\hat{x}$ and $\hat{p}$ can be brought to Weyl ordering.

2.2 Path integral formulation for a scalar field theory

The path integral formalism for quantum mechanics can be directly applied also to a field theory, with real scalar field $\phi(t, \vec{x}) \in \mathbb{R}$. In this scenario, the field amplitudes take on the role previously played by the coordinates, while the argument $\vec{x}$ of the field ϕ is just a label. The following formulation is based on Ref. [47].

In the path integral formalism, the transition amplitude for a propagation from given initial field configuration $\phi_i(\vec{x})$ at time t_i to a final field configuration $\phi_f(\vec{x})$ at time t_f is given by

$$\langle \phi_f | e^{-i\hat{H}T} | \phi_i \rangle = \int_{\phi(t_i, \vec{x})=\phi_i(\vec{x})}^{\phi(t_f, \vec{x})=\phi_f(\vec{x})} [\mathcal{D}\phi(t, \vec{x})] \exp(iS[\phi]) , \quad (2.2.1)$$

where the path integral (or functional integral) denotes the integration over all fields ϕ that are constrained at the boundary to $\phi(t_i, \vec{x}) = \phi_i(\vec{x})$ and $\phi(t_f, \vec{x}) = \phi_f(\vec{x})$. In general, the action of a field theory is given in terms of the Lagrange density $\mathcal{L}$, via

$$S[\phi] = \int_{t_i}^{t_f} dt \int d^3x \mathcal{L}[\phi] . \quad (2.2.2)$$

Eq. (2.2.1) is a very fundamental result. One can even reverse the perspective and interpret this relation as the definition of Hamiltonian dynamics. Then, $\mathcal{L}$ can be considered as the most fundamental specification of a quantum field theory.

For given space-time points $x_1 = (t_1, \vec{x}_1)$ and $x_2 = (t_2, \vec{x}_2)$, we now consider the integral

$$\mathcal{I}(x_1, x_2) = \int_{\phi(-T, \vec{x})=\phi_i(\vec{x})}^{\phi(T, \vec{x})=\phi_f(\vec{x})} [\mathcal{D}\phi(t, \vec{x})] \phi(x_1) \phi(x_2) \exp\left(i \int_{-T}^T d^4x \mathcal{L}[\phi]\right) , \quad (2.2.3)$$

where the field ϕ has boundary conditions $\phi(-T, \vec{x}) = \phi_i(\vec{x})$, $\phi(T, \vec{x}) = \phi_f(\vec{x})$, for some real scalar fields ϕ_i and ϕ_f , depending on the spatial coordinates.

By using Eq. (2.2.1), we want to write this expression as a transition amplitude. To do so, we first rearrange the functional integral

$$\int [\mathcal{D}\phi(t, \vec{x})] = \int [\mathcal{D}\phi_1(\vec{x})] \int [\mathcal{D}\phi_2(\vec{x})] \int_{\substack{\phi(t_1, \vec{x})=\phi_1(\vec{x}) \\ \phi(t_2, \vec{x})=\phi_2(\vec{x})}}^{\phi(t_1, \vec{x})=\phi_1(\vec{x})} [\mathcal{D}\phi(t, \vec{x})] . \quad (2.2.4)$$

In this equation, we first constrained the functional integral at times t_1 and t_2 to the fields $\phi_1(\vec{x})$ and $\phi_2(\vec{x})$, respectively. To restore equality, we additionally introduced two separate path integrals over the fields ϕ_1 and ϕ_2 . Note that these fields solely depend on the spatial coordinates $\vec{x}$.

This rearrangement allows us to replace the extra factors $\phi(x_1)$ and $\phi(x_2)$ in the integral 2.2.3 by $\phi_1(\vec{x}_1)$ and $\phi_2(\vec{x}_2)$, respectively. Thus, they can be pulled out of the main integral, i.e., the

path integral that contains the time variable. Additionally, the main path integral splits into three separate path integrals that are all constrained at its boundaries. For $t_2 > t_1$, the time intervals are $[-T, t_1]$, $[t_1, t_2]$, and $[t_2, T]$, which allows us to write

$$\int_{\phi(t_2, \vec{x})=\phi_2(\vec{x})}^{\phi(t_1, \vec{x})=\phi_1(\vec{x})} [\mathcal{D}\phi(t, \vec{x})] = \int_{\phi(-T, \vec{x})=\phi_i(\vec{x})}^{\phi(t_1, \vec{x})=\phi_1(\vec{x})} [\mathcal{D}\phi(t, \vec{x})] \int_{\phi(t_1, \vec{x})=\phi_1(\vec{x})}^{\phi(t_2, \vec{x})=\phi_2(\vec{x})} [\mathcal{D}\phi(t, \vec{x})] \int_{\phi(t_2, \vec{x})=\phi_2(\vec{x})}^{\phi(T, \vec{x})=\phi_f(\vec{x})} [\mathcal{D}\phi(t, \vec{x})] . \quad (2.2.5)$$

Note that for $t_1 > t_2$ we would need to make the replacements $\phi_1 \leftrightarrow \phi_2$ as well as $t_1 \leftrightarrow t_2$ on the right hand side of Eq. (2.2.5).

Therefore, we can use Eq. (2.2.1) three times to obtain

$$\begin{aligned} \mathcal{I}(x_1, x_2) &= \int [\mathcal{D}\phi_1(\vec{x})] \int [\mathcal{D}\phi_2(\vec{x})] \phi_1(\vec{x}_1) \phi_2(\vec{x}_2) \\ &\quad \times \langle \phi_f | e^{-i\hat{H}(T-t_2)} | \phi_2 \rangle \langle \phi_2 | e^{-i\hat{H}(t_2-t_1)} | \phi_1 \rangle \langle \phi_1 | e^{-i\hat{H}(t_1+T)} | \phi_i \rangle . \end{aligned} \quad (2.2.6)$$

Next, we also want to evaluate the path integrals over the fields ϕ_1 and ϕ_2 by using the completeness relation $\int [\mathcal{D}\phi(\vec{x})] |\phi\rangle \langle \phi| = \hat{1}$. However, at the moment the fields ϕ_1 and ϕ_2 appear additionally as extra factors in the integral. We can eliminate these dependencies by introducing the field operator $\hat{\phi}(\vec{x})$, for which $\hat{\phi}(\vec{x}) |\phi_{1/2}\rangle = \phi_{1/2}(\vec{x}) |\phi_{1/2}\rangle$. This allows us to write

$$\mathcal{I}(x_1, x_2) = \langle \phi_f | e^{-i\hat{H}(T-t_2)} \hat{\phi}(\vec{x}_2) e^{-i\hat{H}(t_2-t_1)} \hat{\phi}(\vec{x}_1) e^{-i\hat{H}(t_1+T)} | \phi_i \rangle , \quad (2.2.7)$$

where the position of the field operators now is relevant. We observe that the exponentials are exactly the time evolution of operators, which generate the Heisenberg operators. Thus we get

$$\mathcal{I}(x_1, x_2) = {}_{\text{H}} \langle \phi_f, T | e^{-i\hat{H}T} T \{ \hat{\phi}_{\text{H}}(x_1) \hat{\phi}_{\text{H}}(x_2) \} e^{-i\hat{H}T} | \phi_i, -T \rangle_{\text{H}} \quad (2.2.8)$$

with $\hat{\phi}_{\text{H}}(x_2) = \hat{\phi}_{\text{H}}(t_2, \vec{x}_2) = e^{i\hat{H}t_2} \hat{\phi}(\vec{x}_2) e^{-i\hat{H}t_2}$ and $|\phi_i, -T\rangle_{\text{H}} = e^{-i\hat{H}T} |\phi_i\rangle$ where the subscript ‘‘H’’ denotes the use of the Heisenberg picture. In this equation, we also introduced the time-ordering operator T , which sorts all fields with respect to their time variable. Therefore, the result is valid for arbitrary times t_1 and t_2 . For two operators and bosonic fields, the time-ordering operator T is formally defined as

$$T \{ \hat{\phi}_{\text{H}}(x_1) \hat{\phi}_{\text{H}}(x_2) \} \equiv \Theta(t_2 - t_1) \hat{\phi}_{\text{H}}(x_2) \hat{\phi}_{\text{H}}(x_1) + \Theta(t_1 - t_2) \hat{\phi}_{\text{H}}(x_1) \hat{\phi}_{\text{H}}(x_2) \quad (2.2.9)$$

with Heaviside step function Θ .¹

Using this result, one can also calculate the two-point correlator

$$G(x_1, x_2) \equiv \langle \Omega | T \{ \hat{\phi}_{\text{H}}(x_1) \hat{\phi}_{\text{H}}(x_2) \} | \Omega \rangle , \quad (2.2.10)$$

where $|\Omega\rangle$ is the ground state of the operator $\hat{H}$, called vacuum state with eigenvalue E_0 . To do so, we need to project the Heisenberg states in Eq. (2.2.9) to the vacuum state $|\Omega\rangle$. In particular this is done in the following way: Let $|n\rangle$ for $n \geq 1$ be the excited eigenstates of

¹For a fermion field, the corresponding field inside the path integral is anticommuting. Therefore the time ordered operator contains an additional minus sign when the number of necessary permutations (starting from the initial operator) is odd.

$\hat{H}$ with eigenvalue E_n .² Together with the ground state $|\Omega\rangle$ we have completeness relation $|\Omega\rangle\langle\Omega| + \sum_{n=1}^{\infty} |n\rangle\langle n| = \hat{1}$. Thus we can write

$$|\phi_i, -T\rangle_H = e^{-i\hat{H}T} |\phi_i\rangle = e^{-iE_0T} |\Omega\rangle\langle\Omega|\phi_i\rangle + \sum_{n=1}^{\infty} e^{-iE_nT} |n\rangle\langle n|\phi_i\rangle, \quad (2.2.11)$$

where we inserted the eigenvalues of $\hat{H}$. We can extract the contribution of the ground state in the limit $T \rightarrow \infty(1 - i\varepsilon)$. In this case, the small rotation into the complex plane introduced by ε ensures that each exponent acquires a nonzero real part. Since $E_n > E_0$ for all $n \geq 1$, all terms in the sum over n are suppressed compared to the term that corresponds to the ground state as long as $|\phi_i\rangle$ has a nonzero overlap with $|\Omega\rangle$, i.e., $\langle\Omega|\phi_i\rangle \neq 0$. Finally, lengthy prefactors can be avoided, when $\mathcal{I}(x_1, x_2)$ is divided by the expression that is obtained, when both terms $\phi(x_1)$ and $\phi(x_2)$ are removed from $\mathcal{I}(x_1, x_2)$, i.e.,

$$\langle\Omega|T\{\hat{\phi}_H(x_1)\hat{\phi}_H(x_2)\}|\Omega\rangle = \lim_{T \rightarrow \infty(1-i\varepsilon)} \frac{\int [\mathcal{D}\phi(t, \vec{x})] \phi(x_1)\phi(x_2) \exp(i \int_{-T}^T d^4x \mathcal{L}[\phi])}{\int [\mathcal{D}\phi(t, \vec{x})] \exp(i \int_{-T}^T d^4x \mathcal{L}[\phi])}. \quad (2.2.12)$$

By complete analogy, higher correlation functions (i.e., those involving a greater number of field insertions) can be computed by simply adding additional factors of ϕ on both sides.

2.3 Wick rotation

The limit $T \rightarrow \infty(1 - i\varepsilon)$ in Eq. (2.2.12) contains a slightly rotated time in the complex plane to ensure convergence of the integral. A more involved procedure, called Wick rotation, utilizes a continuous deformation of the time coordinate to the imaginary axis by a rotation about $\pi/2$ [48]. For a given space-time point $x = (x^0, x^1, x^2, x^3)$ in Minkowski space, the corresponding point in Euclidean spacetime is given by the vector $x^e = (x^1, x^2, x^3, x^4)$, with $x^0 = -ix^4$, where $x^4 \in \mathbb{R}$ is called the Euclidean time.

Consider two space-time points $x = (x^0, x^1, x^2, x^3)$ and $y = (y^0, y^1, y^2, y^3)$ in Minkowski space. The scalar product of these two vectors is defined by

$$x \cdot y \equiv x_\mu y^\mu = g_{\mu\nu} x^\mu y^\nu = x^0 y^0 - x^1 y^1 - x^2 y^2 - x^3 y^3 \quad (2.3.1)$$

with the Minkowski metric $g_{\mu\nu} = \text{diag}(+1, -1, -1, -1)$. For the contravariant vector x^μ , the corresponding covariant vector x_μ is given by $x_\mu = g_{\mu\nu} x^\nu$.

The Wick rotation of these space-time vectors yields the scalar product

$$x \cdot y = -x^4 y^4 - x^1 y^1 - x^2 y^2 - x^3 y^3 = -\delta_{\mu\nu} x^\mu y^\nu, \quad (2.3.2)$$

where we use the convention, that the summation over greek letters for Euclidean quantities runs from 1 to 4. From Eq. (2.3.2) we can read off the corresponding metric in Euclidean space-time, which is given by $-\delta_{\mu\nu}$ with $\mu, \nu \in \{1, 2, 3, 4\}$.³ A covariant vector is then defined via $x_\mu^e = -\delta_{\mu\nu} x^{e,\nu} = -x^{e,\mu}$.

²Here we assume that the ground state is not degenerated, i.e., $E_n > E_0$ for all $n \geq 1$.

³For an Euclidean vector space, the metric $-\delta_{\mu\nu}$ is a rather unconventional choice. However, the usual metric $\delta_{\mu\nu}$ is obtained when the Wick rotation is applied to the Minkowski metric with convention $g_{\mu\nu} = \text{diag}(-1, +1, +1, +1)$. In this case, there is no need to distinguish between upper and lower indices after the Wick rotation.

Applying the Wick transformation to the action, we get

$$S = \int d^4x \mathcal{L}(x) = \int -i(d^4x)^e \mathcal{L}(x^e) = i \int (d^4x)^e \mathcal{L}^e(x^e) = iS^e, \quad (2.3.3)$$

where we used

$$d^4x = dt d^3x = -id\tau d^3x = -i(d^4x)^e \quad (2.3.4)$$

with Euclidean time $\tau \equiv x^4$. We also introduced the Euclidean Lagrangian and action

$$\mathcal{L}^e(x^e) = -\mathcal{L}(x^e) \quad \text{and} \quad S^e = \int (d^4x)^e \mathcal{L}^e(x^e), \quad (2.3.5)$$

respectively. Thus, it holds that $S = iS^e$, where S^e is obtained by $t = -i\tau$ and treating τ as the new time variable. As we focus solely on Euclidean field theory in the remainder, we will omit the superscript “e” whenever the meaning is unambiguous. Note that it is straight-forward to generalize the Wick rotation to a general number of spatial dimensions.

Using this, the transition amplitude given in Eq. (2.2.1) becomes

$$\langle \phi_f | e^{-\hat{H}\tau} | \phi_i \rangle = \int_{\phi(0,x)=\phi_i(\vec{x})}^{\phi(\tau,x)=\phi_f(\vec{x})} [\mathcal{D}\phi(t, \vec{x})] \exp \left(- \int_0^\tau dt \int d^3x \mathcal{L}[\phi] \right) \quad (2.3.6)$$

in Euclidean time, while the Wick rotated version of Eq. (2.2.12) is used as the definition of the expectation value of any observable, i.e.,

$$\langle \phi(x_1) \phi(x_2) \cdots \phi(x_n) \rangle \equiv \frac{\int [\mathcal{D}\phi(x)] \phi(x_1) \phi(x_2) \cdots \phi(x_n) \exp(-S[\phi])}{\int [\mathcal{D}\phi(x)] \exp(-S[\phi])}. \quad (2.3.7)$$

This is an important equation, as the functional integral has now the form of a statistical system with Boltzmann distribution $\exp(-S[\phi])$. This is the starting point for the use of Monte Carlo methods in quantum field theory, where the path integral is estimated by sampling a finite number of paths (i.e., field configurations) whose probability is given by the Boltzmann distribution.⁴

2.4 Thermodynamical observables using the path integral

The correspondence between quantum field theory and statistical mechanics enables the transfer of concepts from one domain to the other. Consider for example the **partition function** Z from statistical mechanics: For a system described by the Hamiltonian $\hat{H}$, the partition function Z is given by

$$Z \equiv \text{tr}(e^{-\beta\hat{H}}) = \int [\mathcal{D}\bar{\phi}(\vec{x})] \langle \bar{\phi} | e^{-\beta\hat{H}} | \bar{\phi} \rangle \quad (2.4.1)$$

with $\beta = \frac{1}{k_B T}$, where T is the temperature of the system.⁵ We set the Boltzmann constant k_B equal to one for convenience. For the evaluation of the trace, we choose the basis that consists of the states $|\bar{\phi}\rangle$. The relevant transition amplitude is a special case of the one calculated in

⁴In practice it is not trivial to relate Euclidean space correlation functions to their Minkowski counterpart. Criteria for the analytic continuation are specified in the Osterwalder-Schrader theorem [49].

⁵To avoid confusion, we stress out that T denoted time before.

Eq. (2.3.6), (obtained by identifying $\tau = \beta$ as well as $\bar{\phi} = \phi_i = \phi_f$). After combining both path integrals via the relation

$$\int [\mathcal{D}\bar{\phi}(\vec{x})] \int_{\phi(0,\vec{x})=\bar{\phi}(\vec{x})}^{\phi(\tau,\vec{x})=\bar{\phi}(\vec{x})} [\mathcal{D}\phi(t,\vec{x})] = \int_{\phi(0,\vec{x})=\phi(\beta,\vec{x})} [\mathcal{D}\phi(t,\vec{x})] \quad (2.4.2)$$

we obtain

$$Z = \int_{\phi(0,\vec{x})=\phi(\beta,\vec{x})} [\mathcal{D}\phi(t,\vec{x})] \exp \left(- \int_0^\beta dt \int d^3x \mathcal{L}[\phi] \right). \quad (2.4.3)$$

with periodic boundary condition $\phi(0,\vec{x}) = \phi(\beta,\vec{x})$.⁶ Note that $\tau = \beta$ makes a connection between the inverse temperature and the considered Euclidean temporal extent. As a result, we can study a field theory at finite temperature by rotating to the Euclidean space and imposing correct boundary conditions for the fields. In the zero temperature limit, we recover the standard Wick rotated quantum field theory on an infinite space-time [46].

Once we have calculated the partition function, we are able to compute the thermodynamical observables. First, the Helmholtz free-energy density f is given by⁷

$$f \equiv \frac{T \ln Z}{V_s} \quad (2.4.4)$$

with spatial volume V_s . Other examples are the energy density and pressure, defined via

$$\varepsilon \equiv -\frac{1}{V_s} \frac{\partial \ln Z}{\partial \beta} \quad \text{and} \quad p \equiv T \frac{\partial \ln Z}{\partial V_s}, \quad (2.4.5)$$

respectively [2]. When we further want to study the grand canonical ensemble, the partition function is extended to

$$Z(T, \mu) = \text{tr} \left(e^{-(\hat{H} - \mu \hat{N})/T} \right), \quad (2.4.6)$$

with the particle number operator $\hat{N}$ depending on the fields and a new parameter μ , called chemical potential. Based on this, the number density corresponding to the operator $\hat{N}$ is given by

$$\rho \equiv \frac{T}{V_s} \frac{\partial \ln Z(T, \mu)}{\partial \mu}. \quad (2.4.7)$$

2.5 Quantum chromodynamics

The Lagrangian density of an $SU(N_c)$ gauge theory in a d -dimensional Minkowski space-time, given by [47]

$$\mathcal{L} = \sum_{f=1}^{N_f} \bar{\psi}^f(x) [i\gamma^\mu D_\mu(x) - m_f] \psi^f(x) - \frac{1}{2} \text{tr}(F_{\mu\nu}(x) F^{\mu\nu}(x)), \quad (2.5.1)$$

has several ingredients, which are discussed in the following. There are N_f different quark flavors, which are distinguished by the index f . For QCD there are in total six different flavors, i.e.,

⁶The periodic boundary condition $\phi(0,\vec{x}) = \phi(\beta,\vec{x})$ results only for bosons. For fermion fields $\psi(t,\vec{x})$, one can show that the boundary condition is anti-periodic, i.e., $\psi(0,\vec{x}) = -\psi(\beta,\vec{x})$.

⁷Strictly speaking, the Helmholtz free-energy density is defined as $f = -\frac{T \ln Z}{V_s}$. However, it is practical to define it with the opposite sign, to avoid overall minus signs in the numerical results.

$N_f = 6$, which are known as “up”, “down”, “charm”, “strange”, “top”, and “bottom”. The corresponding bare quark mass is denoted by m_f . For every flavor f , there is a quark field $\psi^f(x)$, which has a N_s -dimensional Dirac spinor index with $N_s = 2^{\lfloor d/2 \rfloor}$, as well as an N_c -dimensional color index. For QCD $N_c = 3$. Due to readability, those indices are not written explicitly. The quark field $\psi^f(x)$ has a Hermitian-conjugate field $\bar{\psi}^f(x)$, which is defined via $\bar{\psi}^f(x) \equiv \psi^{f,\dagger}(x)\gamma^0$, where the conjugation is taken with respect to the Dirac spinor components. Also the 4×4 dimensional Dirac matrices γ_μ act on these components, which also have an d -dimensional Lorentz index $\mu \in \{0, \dots, d-1\}$. The γ -matrices are defined by the Lorentz-invariant anti-commutator relation

$$\{\gamma_\mu, \gamma_\nu\} \equiv \gamma_\mu \gamma_\nu + \gamma_\nu \gamma_\mu = 2g_{\mu\nu} \mathbb{1} \quad \text{for } \mu, \nu \in \{0, \dots, d-1\}, \quad (2.5.2)$$

which generates a Clifford algebra. Here, $\mathbb{1}$ denotes the identity matrix in spinor space. Further, the QCD Lagrangian contains the covariant derivative

$$D_\mu(x) = \partial_\mu - igA_\mu(x). \quad (2.5.3)$$

While the partial derivative in μ -direction, denoted by ∂_μ , is part of the dynamics of the quark field, the additional term $-igA_\mu(x)$ describes an interaction term of the quark field with the so-called gluon field $A_\mu(x)$. The strength of the interaction is given by the bare coupling g . The gluon field $A_\mu(x)$ lies in the Lie algebra of $SU(N_c)$. Therefore, it can be written as a linear combination of the traceless and hermitian generators t^a of the Lie algebra of $SU(N_c)$, i.e., $A_\mu(x) = A_\mu^a(x)t^a$. Here we use the Einstein summation convention, i.e., there is a sum over the adjoint color index $a \in \{1, \dots, N_c^2 - 1\}$. This is the reason why in QCD ($N_c = 3$) there are 8 different massless gluons, each described by the field $A_\mu^a(x)$. In general, the generators of the Lie algebra of $SU(N_c)$ are defined by the commutation relations

$$[t^a, t^b] = if^{abc}t^c \quad \text{for all } a, b \in \{1, \dots, N_c^2 - 1\} \quad (2.5.4)$$

with structure constants f^{abc} (cf. Sec. 1.9), where again Einstein summation convention is used. For all $N_c \in \mathbb{R}$, the structure constants can be chosen fully antisymmetric. For example, for $N_c = 3$, this is realised by choosing $t^a = \frac{\lambda^a}{2}$, where λ^a are the well-known Gell-Mann matrices. They fulfill $\text{tr}(\lambda^a \lambda^b) = 2\delta_{ab}$.

The gluon field has a dynamics as well, described by the term $\text{tr}(F_{\mu\nu}(x)F^{\mu\nu}(x))$ in 2.5.1, where the trace runs over color indices. It contains the non-abelian field-strength tensor

$$F_{\mu\nu}(x) \equiv \frac{i}{g} [D_\mu(x), D_\nu(x)], \quad (2.5.5)$$

which evaluated reads

$$F_{\mu\nu}(x) = \partial_\mu A_\nu(x) - \partial_\nu A_\mu(x) + gf^{abc}A_\mu^a(x)A_\nu^b(x)t^c. \quad (2.5.6)$$

Remarkably, the term that contains no derivative generates an interaction solely between gluons. From the definition of the gauge field $A_\mu(x)$, it follows immediately that the field-strength tensor is again an element of the Lie algebra of $SU(N_c)$. In particular, it is a $N_c \times N_c$ matrix in color space, whereof the trace is taken in the Lagrangian density (2.5.1).

One can show that the above Lagrangian density is invariant under the local $SU(N_c)$ gauge transformation, where the gluon field $A_\mu(x)$ transforms in the adjoint and the quark field $\psi^f(x)$

transforms in the fundamental representation of $SU(N_c)$, i.e.,

$$A_\mu(x) \rightarrow \Lambda(x)A_\mu(x)\Lambda^{-1}(x) + \frac{i}{g}\Lambda(x)\partial_\mu\Lambda^{-1}(x), \quad (2.5.7a)$$

$$\psi^f(x) \rightarrow \Lambda(x)\psi^f(x), \quad \bar{\psi}^f(x) \rightarrow \bar{\psi}^f(x)\Lambda^{-1}(x), \quad (2.5.7b)$$

for arbitrary $\Lambda(x) \in SU(N_c)$. To do so, one may first calculate the resulting transformation law for the covariant derivative, which yields

$$D_\mu(x) \rightarrow \Lambda(x)D_\mu(x)\Lambda^{-1}(x). \quad (2.5.8)$$

For this derivation, one can use $(\partial_\mu\Lambda(x))\Lambda^{-1}(x) = -\Lambda(x)\partial_\mu\Lambda^{-1}(x)$, which follows directly from $\partial_\mu(\Lambda(x)\Lambda^{-1}(x)) = 0$. From the transformation law (2.5.8) it immediately follows that the covariant derivative acting on a quark spinor transforms again like the quark spinor, i.e.,

$$D_\mu(x)\psi_f(x) \rightarrow \Lambda(x)D_\mu(x)\psi_f(x), \quad (2.5.9)$$

which can be used to show that the fermionic part of the Lagrangian density (2.5.1) is invariant.

Further, by using the definition of the field-strength tensor (2.5.5), together with the transformation law for the covariant derivative it immediately follows that

$$F_{\mu\nu}(x) \rightarrow \Lambda(x)F_{\mu\nu}(x)\Lambda^{-1}(x). \quad (2.5.10)$$

Due to the cyclic property of the trace, $\text{tr}(F_{\mu\nu}F^{\mu\nu})$ is gauge invariant and therefore the whole Lagrangian density is invariant under the local gauge transformation.

The next step is to apply a Wick rotation as described in Sec. 2.3 to the QCD Lagrangian density, in order to obtain the corresponding Lagrangian density in Euclidean space-time. As the γ -matrices as well as the gluon fields carry a Lorentz index, care must be taken in their transformation. The Dirac γ -matrices in Euclidean space-time are defined through the anti-commutation relation

$$\{\gamma_\mu, \gamma_\nu\} = 2\delta_{\mu\nu}\mathbb{1} \quad \text{for all } \mu, \nu \in \{1, \dots, d\}. \quad (2.5.11)$$

This relation is obtained from Eq. (2.5.2) by replacing the Minkowski metric with the Euclidean metric. Consequently, the resulting equation is now invariant under rotation in d dimensions, which takes the place of the Lorentz group in Minkowski space. For given γ -matrices in d -dimensional Minkowski space, we can set $\gamma_{d+1}^e = \gamma^0$ and $\gamma_i^e = -i\gamma^i$, for $i \in \{1, \dots, d-1\}$, since this particular choice fulfills Eq. (2.5.11).

For the term $\gamma^\mu\delta_\mu$, which is part of the QCD Lagrangian, the Wick rotation yields

$$\gamma^\mu\partial_\mu = \gamma^0\partial_0 + \gamma^i\partial_i = i\gamma_4^e\partial_4^e + i\gamma_i^e\partial_i^e = i\gamma_\mu^e\partial_\mu^e, \quad (2.5.12)$$

where we used $\partial_0 = i\partial_4^e$. Note that as before, $\mu \in \{0, \dots, d-1\}$ for space-time vectors in Minkowski space (left) and $\mu \in \{1, \dots, d\}$ for vectors in Euclidean space-time (right). Next, the transformation of the gauge field under Wick rotation is given by $A^0(x) = iA_4^e(x)$ and $A^i(x) = A_i^e(x)$, following the same transformation pattern as the partial derivative. This is motivated by the structure of the gauge transformation of $A_\mu(x)$. There, the gluon field contains

an additive term whose Lorentz structure is given by the partial derivative. Using this, we get $\gamma^\mu A_\mu(x) = i\gamma_\mu^e A_\mu^e(x^e)$, and therefore

$$i\gamma^\mu D_\mu(x) = i\gamma^\mu(\partial_\mu - igA_\mu(x)) = -(\gamma_\mu^e \partial_\mu^e - gi\gamma_\mu^e A_\mu^e(x^e)) = -\gamma_\mu^e D_\mu^e(x^e). \quad (2.5.13)$$

Here, D_μ^e is obtained directly from D_μ by labeling the derivatives and gauge fields with the superscript “e”. Similarly we get,

$$F_{\mu\nu}(x)F^{\mu\nu}(x) = F_{\mu\nu}^e(x^e)F_{\mu\nu}^e(x^e), \quad (2.5.14)$$

where $F_{\mu\nu}^e$ is defined similarly to the Minkowski field-strength tensor $F_{\mu\nu}$, but D_μ is replaced by D_μ^e . Finally, with the use of Eq. (2.3.5), the QCD Lagrangian density in Euclidean space-time takes the form⁸

$$\mathcal{L}(x) = \mathcal{L}_F(x) + \mathcal{L}_G(x), \quad \text{with} \quad (2.5.15)$$

$$\mathcal{L}_F(x) = \sum_{f=1}^{N_f} \bar{\psi}^f(x) [\gamma_\mu D_\mu(x) + m_f] \psi^f(x) \quad \text{and} \quad (2.5.16)$$

$$\mathcal{L}_G(x) = \frac{1}{2} \text{tr}(F_{\mu\nu}(x)F_{\mu\nu}(x)), \quad (2.5.17)$$

where we omitted the superscript “e”, as we will now only be working in Euclidean space-time. Therefore, from this point onward, different directions are denoted by $\mu \in \{1, \dots, d\}$ and all sums or products over μ or ν range from 1 to d unless different limits are stated explicitly.

2.6 Lattice QCD

As the coupling (i.e., the interaction between gauge bosons and leptons or quarks) for the weak and electromagnetic interaction is sufficiently small, one usually applies perturbation theory for computations of matrix elements. This approach breaks down in QCD at low energies, making calculations notoriously difficult. For this situation, a non-perturbative method known as lattice QCD is employed, originally introduced by K. Wilson in 1974 [50]. This approach discretizes space-time and defines all fields on a lattice, making the theory feasible for numerical calculations. The discretization of space-time yields a non-perturbative regularization scheme with ultraviolet cutoff determined by the lattice spacing. Physical observables computed on the lattice must therefore be renormalized to remove the cutoff dependence and to ensure that continuum physics is correctly reproduced [41].

In this thesis, for the discretized formulation, we employ the Wilson plaquette action, originally introduced by K. Wilson [50], to describe the gauge sector. For the fermionic degrees of freedom we will use staggered fermions, first proposed by J. Kogut and L. Susskind [51]. Both approaches are widely used in lattice gauge theory. In the following, we present their construction and demonstrate that they (i) preserve invariance under the discretized form of local $SU(N_c)$ gauge transformations and (ii) are equivalent to the Lagrangian of the continuum theory (2.5.1) in the limit of vanishing lattice spacing. This section is based on [41, 2, 26, 8]. The notation of the lattice QCD action is taken from [27]. From this point onward, we define the Euclidean time direction as $\mu = 1$. This choice is purely conventional.

⁸In Euclidean field theory, it is common to perform the replacement $g \rightarrow -g$, which we adopt here as well.

2.6.1 Wilson's plaquette action

Before we introduce the Wilson plaquette action, it is useful to consider the parallel transporter of a given curve $\mathcal{C} : t \rightarrow c^\mu(t)$ with parametrization $t \in [0, 1]$, defined in continuous space-time, given by

$$U(\mathcal{C}) \equiv \mathcal{P} \exp \left(ig \int_0^1 A_\mu(c(t)) \frac{dc^\mu}{dt} dt \right), \quad (2.6.1)$$

where the path ordering operator $\mathcal{P}$ places matrices $A_\mu(c(t))$ with larger t to the left [47]. It can be shown that $U(\mathcal{C})$ transforms as

$$U(\mathcal{C}) \rightarrow \Lambda(c^\mu(0)) U(\mathcal{C}) \Lambda^\dagger(c^\mu(1)), \quad (2.6.2)$$

under a local $SU(N_c)$ gauge transformation (2.5.7). Thus $\text{tr}(U(\mathcal{C}))$ for closed curves, i.e., curves with $c^\mu(0) = c^\mu(1)$, is gauge invariant due to the cyclic property of the trace. This object is called a Wilson loop and its discretized version builds the basis of Wilson plaquette action.

To see this, we now introduce a finite d -dimensional hypercubic lattice with lattice spacing a as set of discrete points

$$H = \left\{ x \equiv \sum_{\mu=1}^d x_\mu \hat{\mu} \mid x_\mu = 0, \dots, N_\mu \right\}, \quad (2.6.3)$$

where $\hat{\mu} \equiv a\mathbf{e}_\mu$ with unit vector $\mathbf{e}_\mu$ in μ -direction. $(x_1, \dots, x_d)$ are the integer coordinates of the site x . We set $N_1 = N_T$ and $N_j = N$ for $j > 1$, where N_T and N are the number of lattice points in temporal and spatial directions, respectively. Therefore, the lattice has volume $V \equiv a^d |H| = a N_T (aN)^{d-1}$.

For the link (x, μ) which connects the adjacent sites x and $(x + a\hat{\mu})$ the parallel transporter reduces to the gauge-field matrix

$$U_{x,\mu} \equiv e^{igaA_\mu(x)}. \quad (2.6.4)$$

This quantity is called gauge link and is by definition an element of $SU(N_c)$. Formally, this elementary parallel transporter can be derived by introducing the curve $\mathcal{C} : t \rightarrow x + at\hat{\mu}$ in continuous space-time and by using that the discretized field $A_\mu(x)$ is independent of t . The path ordering operator then vanishes, since $A_\mu(x)$ trivially commutes with itself.

In lattice gauge theory, the gauge links $U_{x,\mu}$ are considered the fundamental degrees of freedom associated with the gauge field. From Eq. (2.6.2) we can conclude that it transforms under a local gauge transformation according to

$$U_{x,\mu} \rightarrow \Lambda(x) U_{x,\mu} \Lambda^{-1}(x + \hat{\mu}). \quad (2.6.5)$$

Similarly, the parallel transporter Eq. (2.6.1) for discretized field $A_\mu(x)$ yields for the curve from site $x + \hat{\mu}$ to site x the expression

$$U_{x+\hat{\mu},-\mu} \equiv e^{-igaA_\mu(x)}, \quad (2.6.6)$$

from which it follows that

$$U_{x+\hat{\mu},-\mu} = U_{x,\mu}^{-1} = U_{x,\mu}^\dagger. \quad (2.6.7)$$

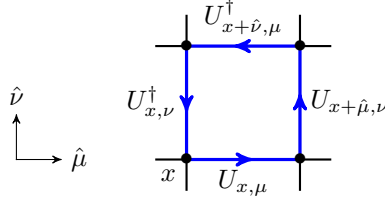


Figure 2.1: Visualization of the gauge links in the plaquette $U_{x,\mu\nu}^P$ defined in Eq. (2.6.10).

For a general discrete path p consisting of the links (x_i, μ_i) for $i \in \{1, \dots, L\}$ the lattice version of the continuum gauge transporter is given by the product

$$\mathcal{U}_p \equiv \prod_{i=1}^L U_{x_i, \mu_i}. \quad (2.6.8)$$

Due to the transformation law for a single gauge link (2.6.5) it transforms as

$$\mathcal{U}_p \rightarrow \Lambda(x_1) \mathcal{U}_p \Lambda^{-1}(x_L + \hat{\mu}_L). \quad (2.6.9)$$

Similarly to the continuum theory, one can build a gauge invariant object by taking the trace of $\mathcal{U}_p$ for any closed path p . The simplest nontrivial closed path on the lattice is the plaquette, which consists of four links forming a square on the lattice. It is specified by a site x as well as two directions $\mu \neq \nu$ ($\mu, \nu \geq 1$) together with an orientation. The corresponding parallel transporter

$$U_{x,\mu\nu}^P \equiv U_{x,\mu} U_{x+\hat{\mu},\nu} U_{x+\hat{\mu}+\hat{\nu},\mu}^\dagger U_{x,\nu}^\dagger, \quad (2.6.10)$$

which will be called plaquette as well, is illustrated in Fig. 2.1. It is the basis of the Wilson plaquette action

$$S^G \equiv \sum_{x, \mu \neq \nu} S_{x,\mu\nu}^G \quad \text{with} \quad S_{x,\mu\nu}^G = \frac{\beta}{2N_c} \text{tr} U_{x,\mu\nu}^P, \quad (2.6.11)$$

with the conventional prefactor including the coupling parameter β , which will be shown to correspond to the continuum coupling g in the subsequent discussion.⁹ Note that due to

$$U_{x,\mu\nu}^{P,\dagger} = U_{x,\nu\mu}^P \quad (2.6.12)$$

the sum over opposite orientations in Eq. (2.6.11) yields Hermitian contributions. Consequently, the Wilson plaquette action S^G is Hermitian as well.

We now demonstrate that, in the limit $a \rightarrow 0$, the Wilson plaquette action differs from the continuum gauge action only by an additive constant, which does not have an influence on physics. To do so, we use the definition of the gauge links (2.6.4) and write

$$S_{x,\mu\nu}^G = \frac{\beta}{2N_c} \text{tr} \left(e^{igaA_\mu(x)} e^{igaA_\nu(x+\hat{\mu})} e^{-igaA_\mu(x+\hat{\mu}+\hat{\nu})} e^{-igaA_\nu(x)} \right). \quad (2.6.13)$$

Next, we apply the Baker–Campbell–Hausdorff formula $e^{aA}e^{aB} = e^{aA+aB+\frac{a^2}{2}[A,B]+\mathcal{O}(a^3)}$ as well as the Taylor expansion of the gauge field $A_\mu(x)$ defined in continuous space-time

$$A_\mu(x + \hat{\nu}) = A_\mu(x) + a\partial_\nu A_\mu(x) + \mathcal{O}(a^2). \quad (2.6.14)$$

⁹This action differs from the one in Ref. [2] not only by an additive constant but also by a sign. However, since we employ the Boltzmann factor e^{S^G} instead of e^{-S^G} (cf. Eq. (2.6.37)), the two formulations are equivalent.

As a result of the expansion, the field-strength tensor appears in the following way

$$S_{x,\mu\nu}^G = \frac{\beta}{2N_c} \text{tr} \left(e^{iga^2 F_{\mu\nu}(x) + \mathcal{O}(a^3)} \right). \quad (2.6.15)$$

Performing a Taylor expansion of the exponential up to second order yields

$$S_{x,\mu\nu}^G = \frac{\beta}{2N_c} \text{tr} \left(\mathbb{1} + \left(iga^2 F_{\mu\nu}(x) + \mathcal{O}(a^3) \right)^2 / 2 \right) = \frac{\beta}{2} - \frac{\beta g^2 a^4}{4N_c} \text{tr} (F_{\mu\nu}(x) F_{\mu\nu}(x)) + \mathcal{O}(a^{d+1}), \quad (2.6.16)$$

where we used that the first order term vanishes due to the trace-less generators appearing in $F_{\mu\nu}(x)$. We now additionally sum over x and $\mu \neq \nu$, and apply the limit $a \rightarrow 0$, where we replace the discrete sum $\sum_x a^d$ by the integral $\int d^d x$. All terms of order a^{d+1} vanish in the limit $a \rightarrow 0$. Thus we get,

$$\sum_{x,\mu \neq \nu} S_{x,\mu\nu}^G \rightarrow c - \frac{1}{2} \int d^d x \text{tr} (F_{\mu\nu}(x) F_{\mu\nu}(x)), \quad (2.6.17)$$

where we identified the coupling parameter as $\beta = \frac{2N_c a^{d-4}}{g^2}$. The constant c is independent of the fields and is irrelevant for observables.

2.6.2 Discretized fermion action

In the following, we introduce a possible discretization of the fermionic section of the QCD action. A fundamental limitation of such a lattice formulation is given by the Nielsen-Ninomiya theorem [52], which states that any discretization preserving chiral symmetry and locality inevitably leads to the appearance of so-called doublers. These lattice artefacts are replicas of the original fermion, as we will see explicitly for the naive discretization of the free theory.

A widely used method to remove all fermion doublers is to introduce a new term in the Lagrangian density that acts like an additional mass term of all replicas. This term ensures that the doublers become heavy and decouple from the theory in the continuum limit [2]. The new action is called Wilson fermion action. However, the added term explicitly breaks chiral symmetry, making this approach less well-suited for the study of chiral symmetry in QCD.

An alternative approach introduces the so-called domain-wall fermions. This formulation employs a five-dimensional lattice to construct fermions with chiral symmetry on a four-dimensional interface of the original lattice.¹⁰ However, as we eventually want to use a tensor-network approach, where additional dimensions can potentially increase computational cost significantly, we choose yet another type of fermions on the lattice, which are called staggered fermions. They sustain a remnant chiral symmetry (cf. Sec. 2.7) and rely on a transformation (called staggered transformation) of the fermion field that decouples components with different Dirac indices. As a result one can consider the theory where the fermion fields have no Dirac structure, which reduces the number of doublers. The remaining doublers can be interpreted as physical flavors in the limit $a \rightarrow 0$. The introduction of staggered fermions is shown explicitly in this section and is based on Ref. [2].

In the lattice discretization of the fermionic section of the QCD action, the quark fields $\psi^f(x)$ $\bar{\psi}^f(x)$ are replaced by the fermion fields ψ_x^f and $\bar{\psi}_x^f$, defined on every site x of the lattice. Similar

¹⁰Domain-wall fermions avoid the Nielsen-Ninomiya theorem by very mildly breaking locality.

to the continuum fields, they have a Dirac and a color index, which is not written explicitly. Each component of ψ_x^f and $\bar{\psi}_x^f$ is Grassmann-valued, which means that all components are totally antisymmetric, i.e.,

$$\psi_{x,i}^f \psi_{y,j}^{f'} = -\psi_{y,j}^{f'} \psi_{x,i}^f \quad (2.6.18)$$

for all lattice sites x, y , for all flavors f, f' , and for all i, j . From this equation it directly follows that $\psi_{x,i}^f \psi_{x,i}^f = 0$ for all x, f , and i . Eq. (2.6.18) also holds for pairs containing (one or two) conjugate Grassmann variables. The reason why $\psi_{x,i}^f$ is Grassmann-valued is that - as a fermion field - it has to obey Fermi statistics, i.e., any observable needs to be antisymmetric under the exchange of the quantum numbers of any two fermions.

In the naive discretization, the derivative with respect to space-time coordinates is replaced using the central finite difference

$$\partial_\mu \psi^f(x) \rightarrow \frac{1}{2a} \left(\psi_{x+\hat{\mu}}^f - \psi_{x-\hat{\mu}}^f \right) \quad (2.6.19)$$

and the covariant derivative acting on $\psi^f(x)$ is replaced by

$$D_\mu \psi^f(x) \rightarrow \frac{1}{2a} \left(U_{x,\mu} \psi_{x+\hat{\mu}}^f - U_{x,-\mu} \psi_{x-\hat{\mu}}^f \right), \quad (2.6.20)$$

using the gauge link defined in Eq. (2.6.4). This discretization is chosen such that both sides have the same transformation behavior under the gauge transformations (2.5.7b) and (2.6.5). This ensures that the obtained discretized action is still gauge-invariant. In the discretization of the action, the integral $\int d^d x$ is again replaced by the finite sum $\sum_x a^d$. Starting from the fermionic Lagrangian density (2.5.16), this leads directly to the so-called naive fermion action

$$S_F = \frac{a^d}{2a} \sum_x \sum_f \bar{\psi}_x^f \left[\sum_\mu \gamma_\mu \left(U_{x,\mu} \psi_{x+\hat{\mu}}^f - U_{x,-\mu} \psi_{x-\hat{\mu}}^f \right) + 2am_f \psi_x^f \right], \quad (2.6.21)$$

where the factor $1/(2a)$ has been extracted and the sum over μ is written explicitly.

In order to see how the discretized theory contains fermion doublers, we rewrite the action (2.6.21) as

$$S_F = a^d \sum_f \sum_{x,y} \bar{\psi}_x^f D_f(x,y) \psi_y^f, \quad (2.6.22)$$

which involves the so-called naive lattice Dirac operator

$$D_f(x,y) = \sum_\mu \gamma_\mu \frac{U_{x,\mu} \delta_{x+\hat{\mu},y} - U_{x,-\mu} \delta_{x-\hat{\mu},y}}{2a} + m_f \delta_{x,y}. \quad (2.6.23)$$

Of particular interest is the inverse of this matrix, referred to as the fermion propagator. One can show that it determines the behavior of n -point correlation functions, according to Wick's theorem [2].

For the free theory (i.e., $U_{x,\mu} = \mathbb{1}$ for all links (x,μ)) the fermion propagator is given by

$$D_f^{-1}(x,y) = \frac{1}{V} \sum_{p \in H} \tilde{D}_f^{-1}(p) e^{ip \cdot (x-y)a}, \quad (2.6.24)$$

with the inverted Fourier transformation of the lattice Dirac operator

$$\tilde{D}_f^{-1}(p) = \frac{m_f - i \sum_\mu \gamma_\mu \sin(p_\mu a)/a}{m_f^2 + \sum_\mu \sin^2(p_\mu a)/a^2} \quad (2.6.25)$$

in discrete momentum space

$$H = \left\{ p \equiv (p_1, \dots, p_d) \mid p_\mu = \frac{2\pi}{aN_\mu} (k_\mu + \theta_\mu), \ k_\mu = -\frac{N_\mu}{2} + 1, \dots, \frac{N_\mu}{2} \right\}. \quad (2.6.26)$$

Here, θ_μ equals zero for periodic and $\frac{1}{2}$ for anti-periodic boundary conditions in the μ -direction. In Eq. (2.6.24), $p \cdot (x - y)$ denotes the usual Euclidean scalar product of the vectors p_μ and $(x - y)_\mu$ and - as before - N_μ denotes the lattice extent.

We now examine the poles of the momentum-space propagator (2.6.25), since they determine the dispersion relation of the described particles. For the present argument it is sufficient to consider the situation $m_f = 0$: for the denominator of the momentum propagator we find $\sum_\mu \sin^2(p_\mu a)/a^2 \rightarrow \sum_\mu p_\mu p_\mu$, in the continuum limit $a \rightarrow 0$. Thus, there is a pole at $p^2 = 0$, which is the correct dispersion relation of a free massless Dirac particle. On the lattice, however, the denominator $\sum_\mu \sin^2(p_\mu a)/a^2$ has zeros at $p = (\delta_1 \pi/a, \dots, \delta_d \pi/a)$ with $\delta_i \in \{0, 1\}$ chosen independently from each other. Thus there are 2^d distinct poles, i.e., 2^d particles, of which only the one with pole at $p = (0, \dots, 0)$ is physical. The remaining modes are unphysical fermions called doublers. Since this argument applies separately to each flavor f , the lattice contains $2^d N_f$ fermions instead of the desired N_f .

The number of doublers is reduced by introducing staggered fermions. To this end, one applies the transformation

$$\psi_x^f = \gamma_1^{x_1} \dots \gamma_d^{x_d} \chi_x^f \quad \text{and} \quad \bar{\psi}_x^f = \bar{\chi}_x^f \gamma_d^{x_d} \dots \gamma_1^{x_1} \quad (2.6.27)$$

with Grassmann variables χ_x^f to the discretized fermion field.¹¹ This transformation mixes space-time and Dirac indices and is referred to as staggered transformation.

Since $\gamma_\mu^2 = \mathbb{1}$ for all μ - a property that follows directly from the anti-commutator relations (2.5.11) - the mass term is invariant under the staggered transformation, i.e., $\bar{\psi}_x^f \psi_x^f = \bar{\chi}_x^f \chi_x^f$.

As the bilinear $\bar{\psi}_x^f \gamma_\mu \psi_{x \pm \hat{\mu}}^f$ transforms as

$$\bar{\psi}_x^f \gamma_\mu \psi_{x \pm \hat{\mu}}^f = (-1)^{\sum_{\nu < \mu} x_\nu} \bar{\chi}_x^f \mathbb{1} \chi_{x \pm \hat{\mu}}^f \quad (2.6.28)$$

under the staggered transformation, the action reads

$$S_F = \frac{a^d}{2a} \sum_x \sum_f \bar{\chi}_x^f \mathbb{1} \left[\sum_\mu \eta_{x,\mu} \left(U_{x,\mu} \chi_{x+\hat{\mu}}^f - U_{x,-\mu} \chi_{x-\hat{\mu}}^f \right) + 2am_f \chi_x^f \right], \quad (2.6.29)$$

where we defined the staggered phases $\eta_{x,\mu} = (-1)^{\sum_{\nu < \mu} x_\nu}$, which replace the gamma matrices γ_μ in the staggered formulation.

¹¹Note that x_μ in $\gamma_\nu^{x_\mu}$ denotes an exponent, i.e., γ_ν raised to the power x_μ .

The obtained action no longer mixes fermions with different Dirac indices and the action has the same form for all Dirac components. Thus we can extract a new action - called staggered fermion action - by keeping only one of the four identical components. To do so, we define

$$\psi_x^f \equiv \sqrt{\frac{a^{d-1}}{2}} \chi_{x,1}^f \quad \bar{\psi}_x^f \equiv \sqrt{\frac{a^{d-1}}{2}} \bar{\chi}_{x,1}^f, \quad (2.6.30)$$

where the subscript “1” refers to the Dirac index. Note that ψ_x^f and $\bar{\psi}_x^f$ are Grassmann-valued vectors with color indices but without Dirac indices, collectively denoted by ψ_x and $\bar{\psi}_x$.

The staggered action takes the form

$$\tilde{S}_F \equiv \sum_{x,\mu} \sum_f \eta_{x,\mu} \left(\bar{\psi}_x^f U_{x,\mu} \psi_{x+\hat{\mu}}^f - \bar{\psi}_{x+\hat{\mu}}^f U_{x,\mu}^\dagger \psi_x^f \right) + 2a \sum_x \sum_f m_f \bar{\psi}_x^f \psi_x^f, \quad (2.6.31)$$

where a shift $x \rightarrow x + \hat{\mu}$ has been applied to the term originally involving $U_{x,-\mu}$. The mass term now only contains the scalar product $\bar{\psi}_x^f \psi_x^f \equiv \sum_{i=1}^{N_c} \bar{\psi}_x^{f,i} \psi_x^{f,i}$. Using this action, the fermion doubling is now reduced to $2^d N_f / N_s$ (N_s is again the number of components of the Dirac spinor index), significantly lowering the number of unphysical degrees of freedom.

2.6.3 Chemical potential on the lattice

While in the vacuum the net number of quarks is zero, it is nonzero in matter. Especially in heavy ion collisions or in dense matter in neutron stars one has to consider also a nonvanishing quark-number density [2]. In this case, one has to study the partition function in the grand canonical ensemble (cf. Sec. 2.4). In the continuum theory of QCD in Euclidean space-time, the quark number corresponding to flavor f is given by $N_q^f = \int d^{d-1}x \bar{\psi}^f(x) \gamma_1 \psi^f(x)$.¹² It is the conserved charge corresponding to the Noether current $j_\mu^f(x) = \bar{\psi}^f(x) \gamma_\mu \psi^f(x)$, where the underlying symmetry the global $U_v(1)$ symmetry in the original Lagrangian density that is preserved even in the quantized theory and also for nonzero quark mass (cf. Sec. 2.7).

Thus we need to add the term $\mu_f N_q^f$ to the continuum Lagrangian density, with chemical potential μ_f . Since a term of the form $\int d^{d-1}x \bar{\psi}^f(x) \gamma_1 \psi^f(x)$ is already contained in the original Lagrangian density, this modification can be implemented simply by replacing the covariant derivative

$$D_\mu \psi^f(x) \rightarrow D_\mu \psi^f(x) + \delta_{\mu,1} \mu_f \psi^f(x). \quad (2.6.32)$$

for every flavor separately. This step introduces the quark chemical potentials μ_f .

The remaining question is how this replacement is incorporated in the discretized version of the covariant derivative Eq. (2.6.20). A widely used choice for this replacement is given by

$$D_\mu \psi^f(x) \rightarrow \frac{1}{2a} \left(e^{a\mu_f \delta_{\mu,1}} U_{x,\mu} \psi_{x+\hat{\mu}}^f - e^{-a\mu_f \delta_{\mu,1}} U_{x,-\mu} \psi_{x-\hat{\mu}}^f \right). \quad (2.6.33)$$

which also preserves time-reversal invariance of the theory. Thus, the staggered fermion action for nonzero chemical can be written as

$$\tilde{S}_F = S_M + \sum_{x,\mu} (S_{x,\mu}^+ + S_{x,\mu}^-), \quad (2.6.34)$$

¹²We used the convention that the time direction is 1.

using the mass term

$$S_M = \sum_{f=1}^{N_f} S_M^f \quad \text{with} \quad S_M^f = 2am_f \sum_x \bar{\psi}_x^f \psi_x^f, \quad (2.6.35)$$

as well as the forward and backward hopping terms

$$S_{x,\mu}^\pm = \sum_{f=1}^{N_f} S_{x,\mu}^{f\pm} \quad \text{with} \quad \begin{cases} S_{x,\mu}^{f+} = \eta_{x,\mu} e^{a\mu_f \delta_{\mu,1}} \bar{\psi}_x^f U_{x,\mu} \psi_{x+\hat{\mu}}^f, \\ S_{x,\mu}^{f-} = -\eta_{x,\mu} e^{-a\mu_f \delta_{\mu,1}} \bar{\psi}_{x+\hat{\mu}}^f U_{x,\mu}^\dagger \psi_x^f. \end{cases} \quad (2.6.36)$$

Note that due to the chemical potentials, the propagation forward in time is enhanced by a factor $\exp(a\mu_f)$, while propagation backward in time is suppressed by $\exp(-a\mu_f)$. Therefore, we introduced a particle–antiparticle asymmetry that gives rise to a nonzero net quark-number density.

2.6.4 Partition function and observables

Using the Wilson plaquette action (2.6.11) as well as the staggered fermion action (2.6.34), we can immediately state the discretized partition function (cf. Sec. 2.4) for the $SU(N_c)$ gauge theory, which is given by[2]

$$Z = \int \left[\prod_x d\psi_x d\bar{\psi}_x \right] \left[\prod_{x,\mu} dU_{x,\mu} \right] e^{S_M + \sum_{x,\mu} (S_{x,\mu}^+ + S_{x,\mu}^-) + \sum_{x,\mu \neq \nu} S_{x,\mu\nu}^G}. \quad (2.6.37)$$

Hereby, $dU_{x,\mu}$ denotes the Haar measure (cf. Sec. 1.9). The Grassmann integration measure is given by

$$d\psi_x d\bar{\psi}_x \equiv \prod_{f=1}^{N_f} \prod_{i=1}^{N_c} d\psi_x^{f,i} d\bar{\psi}_x^{f,i} = \prod_{f=1}^{N_f} \left[\prod_{i=1}^{N_c} d\psi_x^{f,i} \right] \left[\prod_{i=1}^{N_c} d\bar{\psi}_x^{f,i} \right], \quad (2.6.38)$$

where we defined the notation

$$\prod_{k=a}^b \psi^k \equiv \begin{cases} \psi^b \psi^{b-1} \dots \psi^a & \text{for } a \leq b, \\ 1 & \text{for } a > b \end{cases} \quad (2.6.39)$$

for a product of Grassmann variables in reverse order.¹³ Note that in the partition function, the field $\bar{\psi}$ is treated as an independent variable. The boundary conditions are antiperiodic in time for the Grassmann variables. All other boundary conditions are periodic.

On the lattice, the Helmholtz free-energy density (2.4.4), simply referred to as free energy in the rest of this thesis, becomes

$$f = \frac{\ln Z}{V}, \quad (2.6.40)$$

where we used $T = \frac{1}{aN_T}$ (cf. Sec. 2.4), $V_s = (aN)^{d-1}$, and $V = aN_T(aN)^{d-1}$.

For each flavor f , there is a separate quark-number density (2.4.7), which becomes

$$\rho_f = \frac{1}{V} \frac{\partial \ln Z}{\partial (a\mu_f)}. \quad (2.6.41)$$

¹³The symbol $\coprod$ is usually used for the coproduct, which is a very different concept. We use the symbol for notational convenience since there is no danger of confusion with the coproduct.

We used that on the lattice, the dimensionless chemical potential corresponding to flavor f is given by $a\mu_f$.

Finally, the chiral condensate corresponding to flavor f is defined as the expectation value of the product $\bar{\psi}_x^f \psi_x^f$. Using the mass term of the staggered action (2.6.35), it can be expressed via a derivative

$$\langle \bar{\psi}^f \psi^f \rangle \equiv \langle \bar{\psi}_x^f \psi_x^f \rangle = \frac{1}{V} \frac{\partial \ln Z}{\partial (2am_f)}. \quad (2.6.42)$$

The full chiral condensate is obtained when a sum over all flavors f is employed. Note that due to translational invariance, the chiral condensate is independent of the site x .

2.7 Chiral symmetry

One of the fundamental properties of massless QCD is chiral symmetry. As we will see in the following, its study leads to profound insights - most notably - it explains why pions have very small masses compared to the QCD scale. Chirality separates fermion fields in a left-handed and a right-handed sector. This is based on the so-called chirality matrix. For four dimensions in euclidean space-time, the chirality matrix is defined as $\gamma_5 \equiv \gamma_1 \gamma_2 \gamma_3 \gamma_4$, which anti-commutes with all γ -matrices, i.e., $\gamma_5 \gamma_\mu = -\gamma_\mu \gamma_5$ and fulfills $\gamma_5^2 = \mathbb{1}$ as well as $\det(\gamma_5) = 1$.¹⁴

Using the chirality matrix γ_5 , one can define the projection operators

$$P_R = \frac{\mathbb{1} + \gamma_5}{2}, \quad \text{and} \quad P_L = \frac{\mathbb{1} - \gamma_5}{2}. \quad (2.7.1)$$

They are idempotent ($P_{R/L}^2 = P_{R/L}$), orthogonal ($P_R P_L = P_L P_R = 0$), and complete ($P_R + P_L = \mathbb{1}$). Further, they fulfill $P_L \gamma_\mu = \gamma_\mu P_R$ and $\gamma_5 = P_R - P_L$.

With these projectors, a Dirac spinor ψ^f in continuous space-time¹⁵ can be decomposed into left- and right-handed components as

$$\psi_R^f = P_R \psi^f, \quad \text{and} \quad \psi_L^f = P_L \psi^f, \quad (2.7.2)$$

from which the corresponding conjugate spinors follow as

$$\bar{\psi}_R^f = \bar{\psi}^f P_L \quad \text{and} \quad \bar{\psi}_L^f = \bar{\psi}^f P_R. \quad (2.7.3)$$

By inserting $P_R + P_L = \mathbb{1}$ into the Lagrangian density (2.5.16), we obtain

$$\mathcal{L}_F = \bar{\psi}_L \gamma_\mu D_\mu \psi_L + \bar{\psi}_R \gamma_\mu D_\mu \psi_R + [\bar{\psi}_R M \psi_L + \bar{\psi}_L M \psi_R], \quad (2.7.4)$$

where we defined the mass matrix $M \equiv \text{diag}(m_1, \dots, m_{N_f})$. Further, $\psi_{R/L}$ is a vector of size N_f in flavor space that has components $\psi_{R/L}^f$. Remarkably, left- and right-handed Dirac spinors couple to each other only via the mass term and for $M = 0$, the Lagrangian density is invariant under independent global $U(N_f)$ transformations for left- and right-handed spinors, i.e.,

$$\psi_L \rightarrow \Lambda_L \psi_L, \quad \text{and} \quad \psi_R \rightarrow \Lambda_R \psi_R, \quad (2.7.5)$$

¹⁴The generalization of the chirality matrix to arbitrary number of space-time dimensions is straight-forward, as long as the dimensionality is even.

¹⁵Here, we omit the space-time argument x for all spinors as it is not essential in the present context.

for any $\Lambda_L, \Lambda_R \in U(N_f)$. This is called the chiral transformation and the symmetry group is denoted by $U_L(N_f) \times U_R(N_f)$.

As a general property of the unitary group, it can be decomposed as $U(N_f) \cong \frac{U(1) \times SU(N_f)}{Z_{N_f}}$ with cyclic group Z_{N_f} . We apply this relation for $U_L(N_f)$ as well as $U_R(N_f)$, which implies that the above Lagrangian density for $M = 0$ is in particular invariant under the $U_L(1) \times U_R(1)$ transformation

$$\psi_L \rightarrow e^{i\alpha} \psi_L \quad \text{and} \quad \psi_R \rightarrow e^{i\beta} \psi_R \quad (2.7.6)$$

with $\alpha, \beta \in \mathbb{R}$. It is convenient to rewrite this transformations in terms of the spinor ψ . By using $(P_R + P_L) = \mathbb{1}$, we first get

$$\psi \rightarrow (e^{i\beta} P_R + e^{i\alpha} P_L) \psi = e^{i\Theta_v} (e^{-i\Theta_a} P_R + e^{i\Theta_a} P_L) \psi, \quad (2.7.7)$$

where we defined $\Theta_v = \frac{\alpha+\beta}{2}$ and $\Theta_a = \frac{\alpha-\beta}{2}$. Further, one can use the general relation

$$e^{-i\Theta_a} P_R + e^{i\Theta_a} P_L = e^{i\Theta_a \gamma_5} \quad (2.7.8)$$

to obtain the independent vector and axial transformations

$$\psi \rightarrow e^{i\Theta_v} \psi \quad \text{and} \quad \psi \rightarrow e^{i\Theta_a \gamma_5} \psi, \quad (2.7.9)$$

respectively, under which the fermionic Lagrangian density (2.5.16) is invariant in the massless case. We stress out that this invariance only holds due to

$$\{\gamma_\mu D_\mu, \gamma_5\} = 0, \quad (2.7.10)$$

which is the fundamental criterion for chiral symmetry. The corresponding symmetry group of the vector and axial transformations (2.7.9) is denoted by

$$U_v(1) \times U_a(1), \quad (2.7.11)$$

which allows us to write the chiral symmetry of the massless theory in the form

$$SU_L(N_f) \times SU_R(N_f) \times U_v(1) \times U_a(1). \quad (2.7.12)$$

As it turns out, the quantized theory is not invariant under the axial $U_a(1)$ transformation even in the massless case, but only under a discrete $Z_{2N_f}^a$ transformation. The original symmetry is broken explicitly, as the fermion integration measure is no longer invariant under this transformation. This is called axial anomaly and the remaining symmetry is $SU_L(N_f) \times SU_R(N_f) \times U_v(1)$.

For nonzero, but degenerated masses $M = \text{diag}(m, \dots, m)$, the symmetry $SU_L(N_f) \times SU_R(N_f)$ reduces to the subgroup $SU_v(N_f)$, and $SU_v(N_f) \times U_v(1)$ is left. Moreover, for general masses $M = \text{diag}(m_1, \dots, m_{N_f})$, the symmetry is reduced even further to $U_v(1) \times \dots \times U_v(1)$ with N_f factors. The introduction of a chemical potential explicitly breaks all axial symmetries.

In nature, the masses of up and down quark are very small compared to the typical QCD scale. Thus the symmetry $SU_v(N_f) \times SU_a(N_f) \times U_v(1)$ should still be applicable as an approximation. However, the symmetry $SU_a(N_f)$ is broken spontaneously, i.e., even though the theory is invariant under the symmetry, its ground state is not. Relevant for the study of spontaneous chiral symmetry breaking is the chiral condensate

$$\langle \bar{\psi}^f(x) \psi^f(x) \rangle, \quad (2.7.13)$$

where there is a sum over flavors f (and colors). Since it transforms in the same way as the mass term of the Lagrangian density, it is not invariant under the chiral transformation. Thus, a nonzero chiral condensate implies that the chiral symmetry is broken spontaneously. Accordingly, the chiral condensate is an order parameter for the chiral symmetry.¹⁶ As a general consequence of spontaneous symmetry breaking, there are $N_f^2 - 1$ Goldstone bosons. For $N_f = 2$, these are the three pions, which have masses that are much smaller than typical hadronic bound states, like protons and neutrons. This mechanism explains why pions are significantly lighter than protons and neutrons.

In order to study chiral symmetry also for staggered fermions, we first have to identify how γ_5 transforms under the staggered transformation Eq. (2.6.27) for four dimensions.¹⁷ Equivalent to Eq. (2.6.28), we now consider the transformation of the bilinear $\bar{\psi}_x^f \gamma_5 \psi_x^f$ and obtain

$$\bar{\psi}_x^f \gamma_5 \psi_x^f = \eta_x^5 \bar{\chi}_x^f \mathbb{1} \chi_x^f. \quad (2.7.14)$$

Here we used that γ_5 anti-commutes with all γ -matrices and defined

$$\eta_x^5 = (-1)^{x_1+x_2+x_3+x_4}. \quad (2.7.15)$$

The hopping terms (2.6.36) are then invariant under

$$\psi_x^f \rightarrow e^{i\alpha\eta_x^5} \psi_x^f \quad \text{and} \quad \bar{\psi}_x^f \rightarrow \bar{\psi}_x^f e^{i\alpha\eta_x^5}, \quad (2.7.16)$$

which is the discretized version of the $U_a(1)$ symmetry for staggered fermions.¹⁸

There is also a remnant symmetry of the axial $SU_a(N_f)$ symmetry. This is of particular interest for the study of spontaneous symmetry breaking. Due to the fermion doubling, it is actually a remnant of the group $SU_a(N_t)$, where N_t is the number of flavors due to the fermion doubling. Those are also referred to as different “tastes”. To identify this remaining symmetry group for $N_f = 1$, one can use the following procedure, which was developed in [53, 54]: First one makes a transformation of the fermion field such that the new field is defined on hypercubes of the lattice instead of sites. Hereby, a hypercube consists of 2^d adjacent sites, which are arranged in a $2 \times 2 \times \dots \times 2$ configuration. The hypercubes are non-intersecting and even lattice sizes in all directions are assumed. As each hypercube contains 2^d sites, the new field has 2^d additional degrees of freedom. Those can be interpreted as $2^d/N_s$ different tastes of fermions, where each has the usual N_s -dimensional Dirac spinor structure. In this formulation one can determine the nontrivial subgroup of $SU_a(N_t)$ under which the action is invariant. This makes staggered fermions especially well-suited for the study of spontaneous symmetry breaking of chiral symmetry.

2.8 Sign problem for nonzero chemical potential

A standard approach for calculating the quark-number density (2.6.41) or the chiral condensate (2.6.42) first involves the integration over all Grassmann fields.¹⁹ One can show that this results

¹⁶For nonzero quark masses, the chiral condensate is not a strict order parameter, however it can still be used to identify the location of the chiral transition.

¹⁷Again, it is straight-forward to generalize this approach to an arbitrary but even number of dimensions.

¹⁸Note that in contrast to Eq. (2.7.14), the fields ψ_x^f and $\bar{\psi}_x^f$ are the staggered fermions defined in Eq. (2.6.30).

¹⁹The standard approach is in contrast to Part III, where we will first integrate the gauge field.

in the (gauge-link dependent) determinant of the Dirac operator, called fermion determinant. The integration of the gauge field is then approximated by a sum over finite number of gauge field configurations, where

$$e^{-S_{\text{YM}}} \prod_f \det D_f(\mu_f) \quad (2.8.1)$$

is the probability weight that such a configuration of gauge links is generated, based on importance sampling. Hereby, S_{YM} denotes the Yang-Mills action, i.e., the action that corresponds to the pure gauge Lagrangian density (2.5.17) and $D_f(\mu_f) = D_f^0(\mu_f) + m_f$, where D_f^0 is the Dirac operator of flavor f for the massless theory.

Therefore it is crucial that $\det(D_f)$ is real and positive. In the following we will show that this is fulfilled for vanishing chemical potential as done in [55]. The crucial properties to do so are the γ_5 -hermiticity of the Dirac operator, i.e., $(\gamma_5 D_f)^\dagger = \gamma_5 D_f$ (or equivalently $\gamma_5 D_f \gamma_5 = D_f^{\dagger 20}$) as well as the anti-hermiticity of the Dirac operator in the massless case $D_f^{0,\dagger} = -D_f^0$. Let χ_i be an eigenfunction of the Dirac operator with vanishing chemical potential, i.e., $D_f(0)\chi_i = \lambda\chi_i$ with eigenvalue λ_i . Due to the anti-hermiticity of the Dirac operator, λ_i is purely imaginary. Additionally, due to the γ_5 -hermiticity, $\gamma_5\chi_i$ is an independent eigenfunction with eigenvalue $-\lambda_i = \lambda_i^*$. Thus the eigenvalues come in complex-conjugate pairs, which allows us to write

$$\det [D_f(0) + m] = \prod_i (\lambda_i + m_f) (\lambda_i^* + m_f) > 0. \quad (2.8.2)$$

For nonzero chemical potential however, the γ_5 -hermiticity is generalized to

$$\gamma_5 D_f(\mu_f) \gamma_5 = D_f^\dagger(-\mu_f^*), \quad (2.8.3)$$

which yields (due to $\det(\gamma_5) = 1$)

$$[\det D_f(\mu_f)]^* = \det D_f(-\mu_f^*). \quad (2.8.4)$$

Consequently, for nonvanishing real μ_f the determinant of the Dirac operator then is complex and one cannot directly apply importance sampling. This problem is known as the sign problem, or more precise the complex-phase problem. Great effort have been made to circumvent this problem [56, 57]. However, so far none of the remedies allows for precise calculations for general values of μ_f . One known solution is based on the observation that Eq. (2.8.4) yields again a real fermion determinant, when a purely imaginary chemical potential μ_f is used. By an analytic continuation one can also reach the domain of real chemical potentials. Other approaches include a Taylor expansion in μ , reweighting, complex Langevin, thimbles or path optimization, see [3, 4] for reviews.

It is interesting to note yet another scenario, where the Boltzmann weight is real and positive. Consider the situation of $N_f = 2$, where the chemical potentials can be expressed in terms of the baryonic and isospin chemical potential $\mu_B = \frac{\mu_1 + \mu_2}{2}$ and $\mu_I = \frac{\mu_1 - \mu_2}{2}$, respectively. For vanishing baryonic chemical potential $\mu_B = 0$ it follows

$$\prod_f \det D_f(\mu_f) = \det D_1(\mu_I) \det D_2(-\mu_I) = \det D(\mu_I) [\det D(\mu_I)]^* \in \mathbb{R}, \quad (2.8.5)$$

²⁰Note that the same argument can be made when staggered fermions are used, except that the γ_5 -hermiticity is replaced by the corresponding equation $D_f^{\text{st},\dagger}(x, y) = \eta_x^5 D_f^{\text{st}}(x, y) \eta_y^5$.

where we used Eq. (2.8.4) and defined $D(\mu) \equiv D_1(\mu) = D_2(\mu)$, which holds for identical masses. Therefore the fermion determinant is again real and non-negative which allows importance sampling.

The method used in this thesis is based on so-called dual variables. Hereby, in contrast to the conventional approach, the gauge field is integrated out before the Grassmann fields (cf. Part III). This method strongly reduces the sign problem, allowing for a calculation of general values of μ_f .

2.9 Phase diagram of QCD

The phase diagram of QCD has an extremely rich structure, however due to the sign problem, calculations in the regime $\mu/T > 1$ are only feasible by using model scenarios or remedies as explained in Sec. 2.8. Before we discuss the essential properties of the expected QCD phase diagram as done in [2, 8, 55], we briefly introduce the Ehrenfest classifications of phase transitions [58].

A phase transition of n -th order is characterized by a discontinuity of the n -th derivative of the free energy f , while all lower-order derivatives are continuous, i.e.,

$$\frac{\partial^n f_\alpha}{\partial x^n} \neq \frac{\partial^n f_\beta}{\partial x^n} \quad \text{and} \quad \frac{\partial^m f_\alpha}{\partial x^m} = \frac{\partial^m f_\beta}{\partial x^m} \quad \text{for all } m < n, \quad (2.9.1)$$

where α and β distinguishes the different phases and x denotes a parameter of the system like the temperature, the chemical potential, or - specifically in QCD - the quark mass. Of high relevance in practice are the first and second order phase transitions, as both are expected to be present in the $\mu - T$ QCD phase diagram. These transitions are always associated with an order parameter, which is vanishing in one phase and nonzero in the other.

Another relevant phase transition that is also present in the QCD phase diagram is the crossover. Here, all derivatives of the free energy are continuous along the transition, i.e.,

$$\frac{\partial^m f_\alpha}{\partial x^m} = \frac{\partial^m f_\beta}{\partial x^m} \quad \text{for all } m \in \mathbb{N}_0, \quad (2.9.2)$$

and the order parameter still approaches zero in one phase.

There are several order parameters relevant for QCD. Besides the chiral condensate, which is the order parameter for the chiral transition in the limit of small quark masses (cf. Sec. 2.7), there is also the Polyakov loop, defined as

$$P(\vec{x}_s) \equiv \text{tr} \left[\prod_{x_t=1}^{N_T} U_{(x_t, \vec{x}_s), 1} \right]. \quad (2.9.3)$$

It corresponds to a Wilson loop with fixed spatial position $\vec{x}_s$ winding around the compact time direction of the lattice. For pure-gauge theory (i.e., the limit of infinite quark masses), one can show that due to a Z_3 -symmetry of the action, the Polyakov loop is an order parameter for confinement. If it is nonzero the phase is deconfined and free charges can be found.

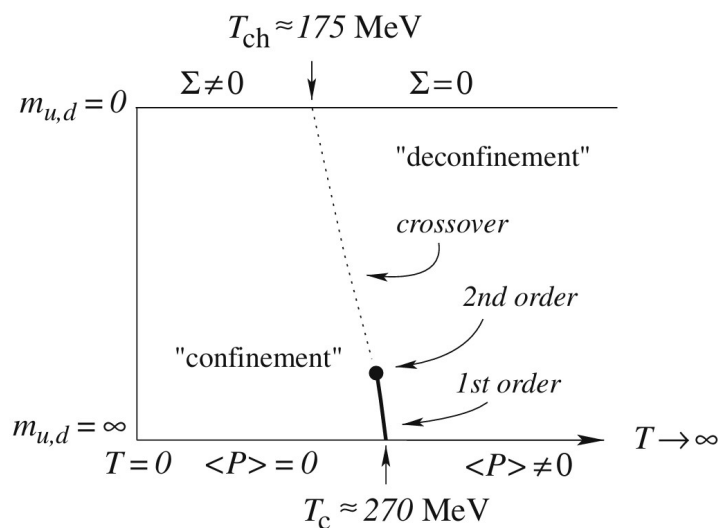


Figure 2.2: QCD phase diagram for two flavors in the $T - m$ plane with $m = m_u = m_d$. The transition from confinement to deconfinement is either first order, second order or a crossover, depending on the quark mass. For small masses, the relevant order parameter is the chiral condensate and for large masses it is the Polyakov loop. The figure is taken from Ref. [2].

The calculation of the Polyakov loop and the chiral condensate results in the QCD phase diagram shown in Fig. 2.2 for two degenerated flavors (denoted by $m_{u,d}$ for up- and down quark) in the $T - m$ plane. As the chemical potentials are set to zero, there is no sign problem and usual methods based on importance sampling are applicable. For vanishing quark masses, the chiral transition is of second order and occurs at a critical temperature of approximately $T_{\text{ch}} \approx 175$ MeV. The chiral condensate is nonzero only below this threshold, which is consistent with the spontaneous breaking of chiral symmetry for zero temperature as discussed in Sec. 2.7. For increasing masses, the chiral condensate is no longer a strict order parameter of the phase transition. However, its value still vanishes only for temperatures above T_{ch} , which is why the chiral condensate is still considered as a valid order parameter in this situation. The transition is then no longer of second order, but a crossover. One can show that the transition between the confined and deconfined phase for small masses also takes place around T_{ch} , which indicates that there is a strong connection between confinement and chiral symmetry [59]. In the limit of infinitely heavy quarks, the Polyakov loop is a strict order parameter, exhibiting a first-order phase transition at $T_c \approx 270$ MeV. The first-order transition weakens for decreasing masses, before it ends in a second-order point.

We now turn to the phase diagram of QCD for nonzero chemical potential. Although the exact locations of phase transitions and the associated critical exponents are not known, there is agreement regarding certain structures of the phase diagram. The expected phase diagram in the $\mu - T$ plane is shown in Fig. 2.3 for vanishing quark masses (left) and for degenerated and small masses close to the physical mass (right) for $N_f = 2$. In both diagrams the isospin chemical potential μ_I is set to zero, and therefore $\mu = \mu_B = \mu_1 = \mu_2$.

For vanishing quark masses, the chiral transition is denoted by the dashed line (second-order

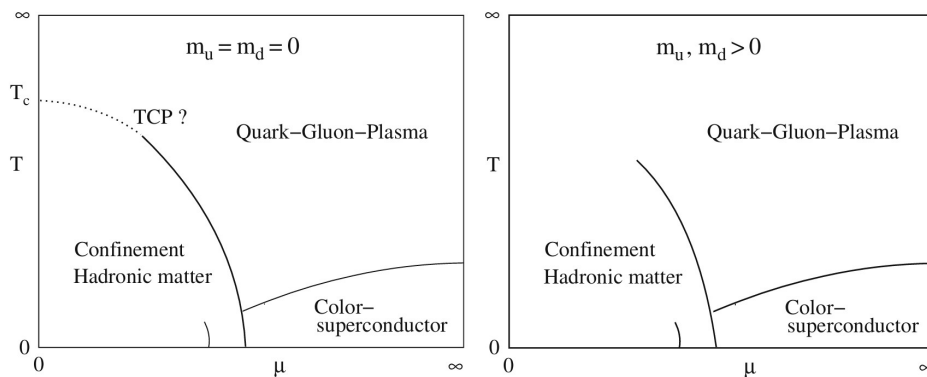


Figure 2.3: Expected QCD phase diagram for massless (left) and degenerated light quarks (right). While for vanishing masses and small chemical potential the chiral transition is of first order, it turns into a crossover, when the masses are increased slightly. The figure is taken from Ref. [2].

transition²¹) that turns for increasing chemical potential at the tricritical point (TCP) into a first order transition (solid line). Inside the quarter-circle the chiral condensate is nonzero. When the quark masses are increased close to the physical mass, the second-order transition becomes a crossover, resulting in a critical endpoint. When the quark masses are increased further, the temperature transition at $\mu = 0$ is again of second- and finally of first order (not shown).

For increased temperature T and chemical potential μ , the energy scale of the theory is increased. Due to asymptotic freedom, quarks and gluons are only weakly coupling leading to a phase generically known as Quark-Gluon-Plasma [60].

Another phase transition in QCD is the liquid-gas phase transition of nuclear matter. The corresponding order parameter is the baryon density n_B . The transition is denoted by a short solid line located in the confined phase for small temperatures. It is expected to be a first order transition which ends in a second-order critical endpoint for increased temperature [55].

From model calculations it follows that the phase diagram of QCD exhibits - analogous to superconductivity for electrons - color superconductivity [5]. Hereby, quarks form cooper pairs resulting in a nonzero diquark condensate, which is an order parameter for color superconductivity. However, there are many possible pairing scenarios of the quarks and the precise behavior of QCD is not clear yet. Only for asymptotically high densities, all models agree that QCD enters a color-flavor locked phase, which, in addition to color superconductivity, shows also a breaking of chiral symmetry. In this limit, many properties are known from first principles [61].

Finally we turn to the discussion of the phase diagram of lattice QCD in the strong-coupling case, i.e., $\beta = 0$, again in the $\mu - T$ plane. The general structure for $N_f = 1$ in four dimensions (thus number of tastes $N_t = 4$) is shown in Fig. 2.4. The phase diagram is calculated by a large dimensional expansion and a mean-field approximation [62]. The chiral transition is similar to the case of general β shown in Fig. 2.3: In the chiral limit, there is a tricritical point that separates the second-order transition (dashed line) from the first-order transition (solid line). For increased masses, the first-order line ends in a critical endpoint, whose movement in the

²¹For $\mu = 0$, this coincides with the second-order chiral transition from Fig. 2.2.

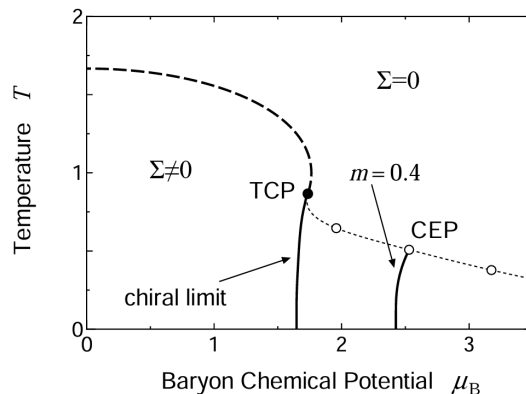


Figure 2.4: Lattice QCD phase diagram at strong coupling for $N_f = 1$ in four dimensions ($N_t = 4$). For vanishing quark mass, the chiral phase transition is of second order (dashed line) and turns into a first-order transition (solid line) at the tricritical point (TCP). For non-vanishing mass, the first-order line is shifted to increased baryon chemical potential and ends in a critical endpoint (CEP). The trajectory of the CEP for increased mass is denoted by a thin dashed line. The figure is taken from Ref. [62] using adapted labels.

phase diagram for changed quark mass is shown by the dotted line. The first-order transition is also shown for quark mass $m = 0.4$. The liquid-gas transition of nuclear matter and the chiral transition coincide for $\beta = 0$.

3 Tensor-network procedure

As explained in Sec. 2.8, the sign problem for nonzero chemical potential makes it impossible to successfully access the regime where $\mu/T > 1$ with any established method. Therefore, it is reasonable to consider alternative approaches that differ crucially from these methods. One promising approach is the use of tensor networks. For a general system, the partition function is expressed as a contraction of tensors, whose total count typically scales with the volume of the system. In order to make the calculation numerically feasible, approximations based on singular value decompositions (SVD) are employed. The cost of computing such a tensor network typically scales only logarithmically with the volume of the system, which is a decisive advantage over usual methods that are based on importance sampling. Tensor-network approaches have already been successfully applied to a wide range of problems in classical and quantum statistical physics, such as spin systems and gauge systems in two, three, and four dimensions. Even some systems with a complex action, i.e., with a sign problem, were successfully studied, for example, the three-dimensional $O(2)$ model with a chemical potential [63]. An overview of different tensor-network approaches for a variety of systems is given in Ref. [19], including a road map towards four-dimensional QCD.

Within this thesis, it is sufficient to consider tensor networks in which a tensor is placed on every site of the space-time lattice and only neighboring tensors are contracted via a shared link index. It is efficient to contract the tensors pairwise, yielding again a tensor network on a coarse grid with half the number of sites. After each such contraction step, the bond dimension

(i.e., the range of the tensor indices) of the coarse-grid tensor rapidly increases. To avoid the curse of dimensionality, an approximation of the coarse-grid tensor based on SVD is employed. The tensor renormalization group (TRG) [64] was the first algorithm with this structure but it is only applicable for two-dimensional systems. An improvement is the higher-order tensor renormalization group (HOTRG) method [20], which is applicable also to higher-dimensional systems and a foundation for the algorithms developed in this thesis. In contrast to TRG, HOTRG is based on the higher-order singular value decomposition (HOSVD) [21] of the coarse grid tensors.

For systems including fermions, Grassmann-valued tensor networks were first discussed in Ref. [65] based on the Hamiltonian formulation. For the Lagrangian formulation, which we will use in this thesis, the Grassmann TRG (GTRG) [66, 67] and the GHOTRG [22] methods were recently developed for cases where the Grassmann variables cannot be integrated out locally.²² This algorithm has already been tailored to strong-coupling QCD in two dimensions using one flavor of staggered quarks in Ref. [23].²³ There, it was shown that the non-local sign factors occurring in the meson-baryon-loop representation of the partition function [6, 7] can be avoided by rewriting the partition function as a tensor network including Grassmann-valued tensors. The method was used to compute the chiral condensate as a function of the quark mass and the volume, and to confirm the absence of dynamical chiral symmetry breaking in the two-dimensional case. Tensor methods are especially well-suited for this investigation, since they can be used for the very large volumes that are required for small masses. Furthermore, the quark-number density at nonzero chemical potential was computed, which hinted at a first-order phase transition. The study of strong-coupling QCD with one flavor of staggered quarks was extended to three and four dimensions in Ref. [69, 24].

In the following, we give a brief introduction to HOSVD, HOTRG, and GHOTRG for a general number of dimensions. Secs. 3.1 and 3.2 are based on Ref. [29], while Sec. 3.3 is based on Refs. [23, 24]. An introduction to HOTRG and GHOTRG in two dimensions is also given in Ref. [26].

3.1 HOSVD

In this subsection, we present the HOSVD approximation, originally introduced in Ref. [21]. We start with the following relevant definitions.

For a real tensor M of order n and dimension $N_1 \times \cdots \times N_n \in \mathbb{N}^n$ and a real matrix A , we define the left matrix-tensor multiplication of the r -th index of M as

$$(A \odot_r M)_{i_1 \dots i_{r-1} j_r i_{r+1} \dots i_n} \equiv \sum_{i_r} A_{j_r i_r} M_{i_1 \dots i_n}. \quad (3.1.1)$$

By definition, we find

$$A \odot_r (B \odot_s M) = B \odot_s (A \odot_r M), \quad (3.1.2)$$

²²In theories with Grassmann variables, the latter can sometimes be integrated out completely during the construction of the tensor network, as for example in the $U(N)$ gauge theory at infinite coupling [6, 68], in which case the usual TRG or HOTRG method can be used.

²³We will refer to GHOTRG for both the procedure introduced in Ref. [22] and its tailored version for lattice QCD in the strong-coupling expansion.

for any $r \neq s$ and for any matrices A and B .

Further, we define the r -unfolding of the tensor M as the matrix

$$M_{i_r, (i_1, \dots, i_{r-1}, i_{r+1}, \dots, i_n)}^{(r)} = M_{i_1 \dots i_n}, \quad (3.1.3)$$

i.e., the row index of $M^{(r)}$ is given by the r -th index of M , while the column index of $M^{(r)}$ is given by a linearization of the remaining indices. We note that the r -unfolding of $(A \odot_r M)$ is given by the matrix product of A and the r -unfolding of M , i.e.,

$$(A \odot_r M)^{(r)} = A \cdot M^{(r)}. \quad (3.1.4)$$

Additionally, the multi-rank $(R_1, \dots, R_n) \in \mathbb{N}^n$ of a tensor M of order n is defined by the r -rank

$$R_r \equiv \text{rank}(M^{(r)}) \quad (3.1.5)$$

for $r = 1, \dots, n$.

We now define the tensor approximation²⁴

$$\hat{M} \equiv P^{(1)} \odot_1 P^{(2)} \odot_2 \dots P^{(n)} \odot_n M, \quad (3.1.6)$$

with $N_r \times N_r$ orthogonal projectors $P^{(r)}$ (i.e., $P^{(r)2} = P^{(r)}$ and $P^{(r)\text{T}} = P^{(r)}$ for all r) of rank $K_r \leq N_r$. Therefore, $\hat{M}$ has the same dimension as M . Due to the property

$$\text{rank}(A \cdot B) \leq \min(\text{rank}(A), \text{rank}(B)) \quad (3.1.7)$$

for general matrices A and B , $\hat{M}$ is a lower-rank approximation of M , i.e., the r -rank of $\hat{M}$ is smaller or equal to the r -rank of M for all $r = 1, \dots, n$.

As a general property of an orthogonal projector of dimension $N_r \times N_r$ with rank K_r , $P^{(r)}$ can be written as

$$P^{(r)} = U^{(r)} \cdot U^{(r)\text{T}} \quad (3.1.8)$$

with $N_r \times K_r$ -dimensional semi-orthogonal matrices $U^{(r)}$, called frames, i.e., the columns of $U^{(r)}$ are orthonormal, but not their rows (unless $K_r = N_r$). Note that the columns of the frame $U^{(r)}$ provide an orthonormal basis of a K_r -dimensional vector space. In particular, we find

$$U^{(r)\text{T}} U^{(r)} = \mathbb{1}_{K_r \times K_r}. \quad (3.1.9)$$

For given frames $U^{(r)}$, we now define the $K_1 \times \dots \times K_n$ dimensional core tensor

$$S = U^{(1)\text{T}} \odot_1 U^{(2)\text{T}} \odot_2 \dots U^{(n)\text{T}} \odot_n M, \quad (3.1.10)$$

such that the tensor approximation $\hat{M}$ can be rewritten as

$$\hat{M} = U^{(1)} \odot_1 U^{(2)} \odot_2 \dots U^{(n)} \odot_n S. \quad (3.1.11)$$

²⁴Due to Eq. (3.1.2), the ordering of the operations is irrelevant.

The frames $U^{(r)}$ together with the core tensor S contain the same information as the tensor approximation $\hat{M}$, but require only $\sum_r K_r N_r + \prod_r K_r$ entries instead of $\prod_r N_r$ entries. Consequently, it is numerically more efficient, not to calculate $\hat{M}$ explicitly, but rather to evaluate operations involving $\hat{M}$ using the frames $U^{(r)}$ and the core tensor S .

In order to quantize the approximation error introduced by the tensor $\hat{M}$ we now introduce the Frobenius norm. For a given tensor M it is defined by

$$\|M\| = \sqrt{\langle M, M \rangle} = \sqrt{\sum_{i_1, \dots, i_n} M_{i_1 \dots i_n}^2}, \quad (3.1.12)$$

where the inner product of two real tensors M and $\tilde{M}$ of same size is defined as

$$\langle M, \tilde{M} \rangle = \sum_{i_1, \dots, i_n} M_{i_1 \dots i_n} \tilde{M}_{i_1 \dots i_n}. \quad (3.1.13)$$

For the inner product of the tensor M and its approximation $\hat{M}$ we find

$$\begin{aligned} \langle M, \hat{M} \rangle &= M_{j_1, \dots, j_n} P_{j_1, i_1}^{(1)} \cdots P_{j_n, i_n}^{(n)} M_{i_1, \dots, i_n} \\ &= M_{j_1, \dots, j_n} P_{j_1, k_1}^{(1)} P_{k_1, i_1}^{(1)} \cdots P_{j_n, k_n}^{(n)} P_{k_n, i_n}^{(n)} M_{i_1, \dots, i_n} = \|\hat{M}\|^2 \\ &= M_{j_1, \dots, j_n} U_{j_1, k_1}^{(1)} U_{k_1, i_1}^{(1)\top} \cdots U_{j_n, k_n}^{(n)} U_{k_n, i_n}^{(n)\top} M_{i_1, \dots, i_n} = \|S\|^2, \end{aligned} \quad (3.1.14)$$

where the second and third line both follow directly from the first line by using the fundamental projection operator property, and by expressing the projectors in terms of the frames, respectively. Using this, the squared approximation error of $\hat{M}$ is given by

$$\|M - \hat{M}\|^2 = \|M\|^2 + \|\hat{M}\|^2 - 2\langle M, \hat{M} \rangle = \|M\|^2 - \|\hat{M}\|^2 = \|M\|^2 - \|S\|^2 \quad (3.1.15)$$

and does not require the explicit calculation of $\hat{M}$.

We now consider the singular value decomposition of the r -unfolding of M ,

$$M^{(r)} = L^{(r)} \Sigma^{(r)} R^{(r)\top} \quad (3.1.16)$$

with $N_r \times N_r$ orthogonal matrix $L^{(r)}$. The columns of $L^{(r)}$ are the left singular vectors of $M^{(r)}$. $\Sigma^{(r)}$ is a $N_r \times N_r^{n-1}$ matrix that is only nonzero at the diagonal. The values at the diagonal are the singular values of $M^{(r)}$, which are real and non-negative. Similarly to $L^{(r)}$, $R^{(r)}$ is a $N_r^{n-1} \times N_r^{n-1}$ orthogonal matrix.

For any matrix $M^{(r)}$, it is well known that the approximation $A^{(r)} \equiv L^{(r)} \bar{\Sigma}^{(r)} R^{(r)\top}$, where $\bar{\Sigma}^{(r)}$ is obtained from $\Sigma^{(r)}$ by keeping only the K_r largest singular values and setting all others to zero, is the best rank- K_r approximation of the matrix $M^{(r)}$. Formally, this means that it minimizes the Frobenius norm $\|M^{(r)} - A^{(r)}\|$ for given rank K_r . The relative truncation error $\varepsilon^{(r)}$ is given by

$$\varepsilon^{(r)} = \frac{\|M^{(r)} - A^{(r)}\|}{\|M^{(r)}\|} = \sqrt{\frac{\sum_{i=K_r+1}^{N_r} \lambda_i^{(r)}}{\sum_{i=1}^{N_r} \lambda_i^{(r)}}}, \quad (3.1.17)$$

where $\lambda_i^{(r)}$ are the squared singular values of the unfolding $M^{(r)}$, ordered such that $\lambda_1^{(r)} \geq \dots \geq \lambda_{N_r}^{(r)}$. Note that $\lambda_i^{(r)}$ are also the eigenvalues of the Gramian $M^{(r)} M^{(r)\top}$.

The HOSVD of M is defined by the Eqs. (3.1.10) and (3.1.11) with $U^{(r)} = L^{(r)}$. As in this case $K_r = N_r$, it holds that $\hat{M} = M$ and there is now information lost by performing the HOSVD.

More relevant is the HOSVD approximation $\hat{M}$ with multi-rank $(K_1, \dots, K_n)$. For given tensor M , it is defined by Eqs. (3.1.10) and (3.1.11), where the columns of $U^{(r)}$ are given by those columns of $L^{(r)}$ (i.e., left singular vectors of $M^{(r)}$) that correspond to the K_r largest singular values of $M^{(r)}$. Note that it holds

$$A^{(r)} = U^{(r)} U^{(r)\text{T}} M^{(r)}. \quad (3.1.18)$$

One can show that the HOSVD tensor approximation $\hat{M}$, in general, is only the best possible approximation for given multi-rank $(K_1, \dots, K_n)$ when $n = 2$, as in this case the Frobenius norm $\|M - \hat{M}\|$ is minimized by definition. But even for general order n the HOSVD approximation is usually very close to the best approximation [21].

There is a known algorithm to construct the best possible approximation, i.e, to minimize the Frobenius norm $\|M - \hat{M}\|$ for given multi-rank $(K_1, \dots, K_n)$ numerically. The algorithm is called Higher Order Orthogonal Iteration (HOOI) [70]. Nevertheless, the HOSVD approximation is particularly attractive due to its computational simplicity while still yielding an almost optimal approximation.

3.2 HOTRG

We now turn to the HOTRG approach [20] for computing a tensor network of the form

$$Z = \sum_{\mathbf{i}} \prod_x (T_x)_{i_{x,-1} i_{x,1} \dots i_{x,-d} i_{x,d}}. \quad (3.2.1)$$

This expression is visualized by a d -dimensional lattice with V sites. On every site $x = (x_1, \dots, x_d)$ there is a tensor T_x . The different directions of the lattice are denoted by $\mu = 1, \dots, d$. Neighboring tensors on sites x and $x + \hat{\mu}$, where $\hat{\mu}$ is the unit vector in μ -direction, share a link index $i_{x,\mu}$ which is summed over, denoted by the sum over all tuples $\mathbf{i} \equiv (i_{x,\mu} | \forall x, \mu)$. For convenience, we introduced the notation $i_{x,-\mu} = i_{x-\hat{\mu},\mu}$. The lattice has a toroidal shape, i.e., it is closed in all directions. For our purposes, it is sufficient to consider the case where the range of the summation over $i_{x,\mu}$ is independent of the link. This range is denoted by D and is referred to as the bond dimension. The result is denoted by Z , as it is typically the partition function of a system that is expressed in the form of a tensor network.

For the moment, we assume that the tensors are independent of the chosen site, i.e., we can write $T = T_x$. As a further assumption, the number of lattice sites in any direction is a power of two. A natural way to calculate the given partition function is to contract the tensors pairwise, yielding again a tensor network in the form of Eq. (3.2.1) albeit on a coarse lattice, i.e., a lattice with half the number of sites. As all tensors are identical, this contraction has to be calculated only once. This procedure is repeated until only a single tensor is left, from which the trace is taken. It is crucial to note that the bond dimensions of the tensors increase rapidly in each contraction. In particular, for any direction perpendicular to the contraction direction, the bond dimension is squared in each step. For a 4×4 lattice, the contraction procedure is visualized in Fig. 3.5.

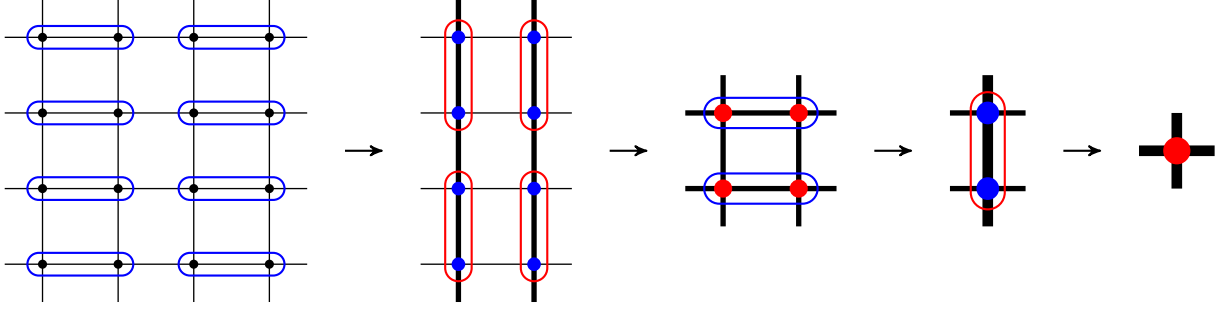


Figure 3.5: Visualization of the contraction procedure for computing a tensor network on a two-dimensional 4×4 lattice. In each contraction step, the indices perpendicular to the contraction direction are combined to a new link index with squared range. This is visualized by progressively thicker links.

Due to the increasing bond dimensions, an exact calculation of the tensor network is only feasible for very small lattices. Therefore, an essential property of HOTRG is to reduce the dimension of the coarse tensor using an HOSVD-based approximation which avoids the curse of dimensionality. This truncation together with the contraction of two fine-grid tensors is referred to as a blocking step. Symbolically, we summarize the HOTRG algorithm for the $(k + 1)$ -th blocking step as

$$T^{[k]} \star_{\mu} T^{[k]} \xrightarrow{\text{truncate}} T^{[k+1]}, \quad (3.2.2)$$

where $T^{[k]}$ is the obtained coarse tensor after k blocking steps and $\star_{\mu}$ denotes the contraction of two tensors $T^{[k]}$ that are neighbored in direction μ . The explicit construction of $T^{[k+1]}$ is shown later on.

Formally, we define the coarse tensor after the first contraction as

$$M_{j_{X,-1}j_{X,1}\cdots j_{X,-d}j_{X,d}} = \sum_{i_{x,1}} T_{i_{x,-1}i_{x,1}\cdots i_{x,-d}i_{x,d}} T_{i_{y,-1}i_{y,1}\cdots i_{y,-d}i_{y,d}} \quad (3.2.3)$$

with $y = x + \hat{1}$ and therefore $i_{y,-1} = i_{x,1}$. The coarse site is denoted by $X \equiv (x, y)$. On the coarse lattice, the link indices are defined as

$$j_{X,-1} = i_{x,-1}, \quad j_{X,1} = i_{y,1}, \quad \text{and} \quad j_{X,\pm\mu} = (i_{x,\pm\mu}, i_{y,\pm\mu}) \text{ for } \mu \neq 1, \quad (3.2.4)$$

where the combined indices are referred to as fat indices. The construction of the link indices associated to the coarse lattice is visualized for three dimensions in Fig. 3.6.

3.2.1 Backward-forward symmetry

A tempting approach would be to replace the coarse tensor (3.2.3) by its HOSVD approximation of multi-rank $(D, \dots, D)$. However, it turns out that this approximation is not optimal in the iterative blocking procedure as shown next.

For this argument it is sufficient to consider the contraction in 1-direction in two dimensions. According to Eq. (3.2.3) the third and fourth index of M are the fat indices, i.e., they have bond dimension D^2 , while the first and second index have bond dimension D .

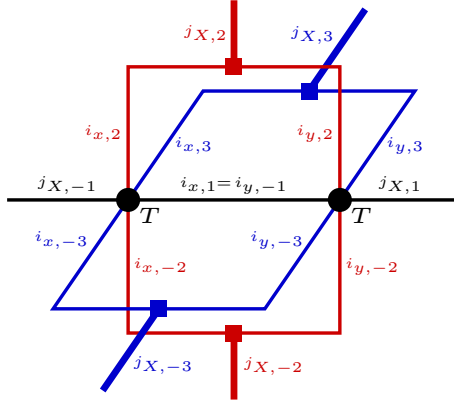


Figure 3.6: Construction of the link indices on a coarse site. Fat indices occur only on links perpendicular to the contraction direction. The figure is taken from Ref. [29].

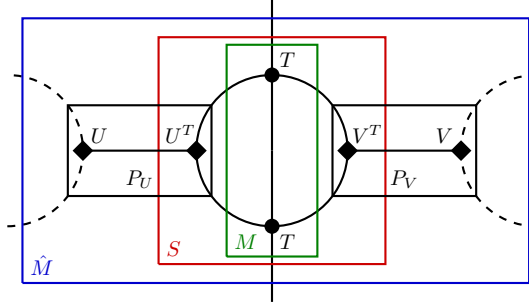


Figure 3.7: Visualization of the constituents of the HOSVD approximation $\hat{M}$. Fat indices are visualized as two separate lines. The figure is taken from Ref. [29].

For the HOSVD approximation with multi-rank (D, D, D, D) , we need to calculate the frames U and V for the unfoldings corresponding to the third and fourth index of M , denoted by $M^{(3)}$ and $M^{(4)}$.²⁵ According to Eq. (3.1.10), the core tensor is then given by

$$S = U^T \odot_3 V^T \odot_4 M, \quad (3.2.5)$$

yielding the coarse tensor approximation

$$\hat{M} = U \odot_3 V \odot_4 S = P_U \odot_3 P_V \odot_4 M \quad (3.2.6)$$

with $D^2 \times D^2$ projectors $P_U = UU^T$ and $P_V = VV^T$ with $U^T U = V^T V = \mathbb{1}_{D \times D}$. The constituents of the tensor approximation are visualized in Fig. 3.7.

Assume that, for the next contraction step, two tensors $\hat{M}$ are contracted as shown in Fig. 3.8, again using the constituents of the tensor approximation $\hat{M}$. We observe that the matrix product $G = U^T V$ occurs. This matrix is called merger, and the whole contraction can be perceived as contracting core tensors S with mergers G in between. It is crucial to note that, in the calculation of the merger G , information may be lost depending on the overlap of the vector spaces spanned by the columns of U and V . Consequently, the rank of the merger is, in general, smaller than

²⁵In principal, the unfoldings $M^{(1)}$ and $M^{(2)}$ corresponding to the first and second index of M are necessary as well. However, as there is no truncation for these links, the frames are orthogonal and no information is lost. Therefore, after a change of basis, these frames are the identity matrix.

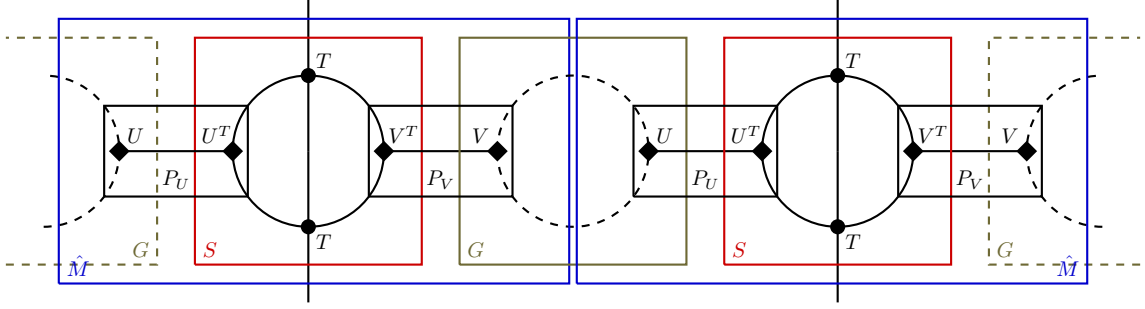


Figure 3.8: Contraction of two HOSVD approximations $\hat{M}$. By using the constituents of $\hat{M}$, we observe that the contraction can be perceived as contracting core tensors and mergers only. The figure is taken from Ref. [29].

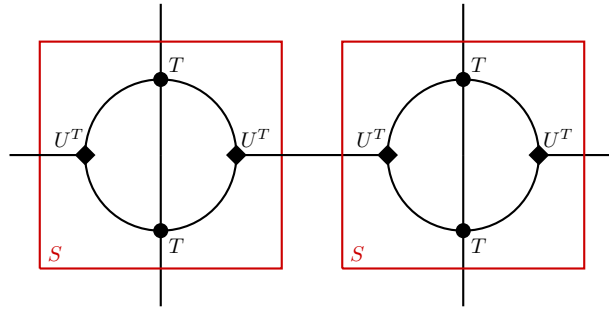


Figure 3.9: Simplification of the contraction procedure in two dimensions after setting $U = V$. The core tensors itself already build a tensor network in the form of Eq. (3.2.1) albeit on the coarse lattice. The figure is taken from Ref. [29].

D. Therefore, even though the HOSVD approximation $\hat{M}$ is a close-to-best approximation of M , $\hat{M} \star_1 \hat{M}$ is not necessarily a good approximation of $M \star_1 M$.

The only possibility in which no information is lost in the merger is when the columns of U and V span the same vector space, which is an additional demand in the HOTRG algorithm. In this case, the merger is orthogonal. For efficiency, we additionally set $U = V$, i.e., the same basis for the vector space. With this condition, it is clear that the usual HOSVD approximation cannot be used to approximate the coarse tensor M and it is unclear a priori how the frames U should be constructed instead. The construction of the frames will be considered in the next section, but for the moment let us discuss the consequences of the condition $U = V$. Due to $G = U^T V = \mathbb{1}_{D \times D}$, the merger drops out and the situation simplifies drastically to the one shown in Fig. 3.9. The contraction of two tensor approximations $\hat{M}$ is replaced by the contraction of two core tensors S only.²⁶

Since the condition of using common frames for forward and backward direction can be directly generalized to arbitrary dimensions, we define the core tensor used in HOTRG, for a contraction in 1-direction, as

$$S = U^{(3)\text{T}} \odot_3 U^{(3)\text{T}} \odot_4 \cdots U^{(d)\text{T}} \odot_{2d-1} U^{(d)\text{T}} \odot_{2d} M. \quad (3.2.7)$$

By writing the tensor network (3.2.1) in terms of the core tensor S , the structure of the tensor

²⁶To be able to construct and eliminate all mergers we use the property that the lattice is closed.

network is preserved albeit on a coarse lattice. From there, the algorithm is applied iteratively. Symbolically, the $(k + 1)$ -th blocking step of the HOTRG approach can be written as

$$T^{[k]} \star_{\mu} T^{[k]} =: M \rightarrow S =: T^{[k+1]}, \quad (3.2.8)$$

when the contraction is applied in μ -direction. When the lattice consists of a single site, one finally needs to calculate the trace of the obtained tensor, i.e., the sum over the link indices in all directions.

The remaining task is to construct the frames corresponding to the truncation, such that forward-backward symmetry is satisfied and the approximation is close to HOSVD. There are two different established methods for constructing such frames. We refer to them as SuperQ [29] and the Xie et al. [20] procedure and introduce them in the following section.

3.2.2 SuperQ and the Xie et al. procedure

For the tensor M , given in Eq. (3.2.3), we denote the forward and backward unfolding corresponding to direction μ with $1 < \mu \leq d$ as

$$M_{\text{b}}^{(\mu)} \equiv M^{(2\mu-1)} \quad \text{and} \quad M_{\text{f}}^{(\mu)} \equiv M^{(2\mu)}, \quad (3.2.9)$$

respectively. Further, the frames $U_{\text{b}}^{(\mu)}$ and $U_{\text{f}}^{(\mu)}$ are the HOSVD truncation frames corresponding to these unfoldings, i.e., the semi-orthogonal matrices, whose columns are given by the left singular vectors corresponding to the D largest singular values of $M_{\text{b}}^{(\mu)}$ and $M_{\text{f}}^{(\mu)}$, respectively. In the originally proposed HOTRG procedure, referred to as Xie et al. truncation procedure [20], the truncation frame $U^{(\mu)}$ for direction μ , is either $U_{\text{b}}^{(\mu)}$ or $U_{\text{f}}^{(\mu)}$, depending on which of both gives the smaller SVD truncation error (3.1.17). Therefore, the chosen frame is a near-to-optimal truncation frame only for one out of the two considered directions, and it is clear that in general, truncation frames should exist that optimize both truncations.

To find such an optimized truncation frame, it is useful to consider the positive semi-definite Gram matrices

$$Q_{\text{b}}^{(\mu)} \equiv M_{\text{b}}^{(\mu)} M_{\text{b}}^{(\mu)\text{T}} \quad \text{and} \quad Q_{\text{f}}^{(\mu)} \equiv M_{\text{f}}^{(\mu)} M_{\text{f}}^{(\mu)\text{T}}. \quad (3.2.10)$$

Note that the columns of $U_{\text{b}}^{(\mu)}$ and $U_{\text{f}}^{(\mu)}$ are the eigenvectors corresponding to the D largest eigenvalues of $Q_{\text{b}}^{(\mu)}$ and $Q_{\text{f}}^{(\mu)}$, respectively.

We now consider a truncation frame $U^{(\mu)}$ for truncating both $M_{\text{b}}^{(\mu)}$ and $M_{\text{f}}^{(\mu)}$, which yields the rank- D approximations (cf. Eq. (3.1.18))

$$A_{\text{b}}^{(\mu)} \equiv U^{(\mu)} U^{(\mu)\text{T}} M_{\text{b}}^{(\mu)} \quad \text{and} \quad A_{\text{f}}^{(\mu)} \equiv U^{(\mu)} U^{(\mu)\text{T}} M_{\text{f}}^{(\mu)}. \quad (3.2.11)$$

The relative truncation errors are given by

$$\varepsilon_b^{(\mu)} \equiv \frac{\|M_b^{(\mu)} - A_b^{(\mu)}\|}{\|M_b^{(\mu)}\|} = \sqrt{1 - \frac{\|A_b^{(\mu)}\|}{\|M_b^{(\mu)}\|}} = \sqrt{1 - \frac{\text{tr}(U^{(\mu)\text{T}} Q_b^{(\mu)} U^{(\mu)})}{\text{tr} Q_b^{(\mu)}}}, \quad (3.2.12)$$

$$\varepsilon_f^{(\mu)} \equiv \frac{\|M_f^{(\mu)} - A_f^{(\mu)}\|}{\|M_f^{(\mu)}\|} = \sqrt{1 - \frac{\|A_f^{(\mu)}\|}{\|M_f^{(\mu)}\|}} = \sqrt{1 - \frac{\text{tr}(U^{(\mu)\text{T}} Q_f^{(\mu)} U^{(\mu)})}{\text{tr} Q_f^{(\mu)}}}, \quad (3.2.13)$$

where in the first step we used Eq. (3.1.15), and in the second step we used the general property $\|\mathcal{M}\| = \text{tr}(\mathcal{M}\mathcal{M}^{\text{T}})$ for real matrices $\mathcal{M}$ as well as the projector property of $U^{(\mu)}U^{(\mu)\text{T}}$.

For an optimized frame $U^{(\mu)}$, the combined truncation error

$$\varepsilon_S^{(\mu)^2} = \frac{1}{2} \left(\varepsilon_b^{(\mu)^2} + \varepsilon_f^{(\mu)^2} \right) = 1 - \frac{1}{2\|M\|^2} \text{tr} \left[U^{(\mu)\text{T}} \left(Q_b^{(\mu)} + Q_f^{(\mu)} \right) U^{(\mu)} \right] \quad (3.2.14)$$

is reduced, where we used $\text{tr} Q_b^{(\mu)} = \text{tr} Q_f^{(\mu)} = \|M\|^2$. This is reached by choosing the columns of $U^{(\mu)}$ to be the eigenvectors corresponding to the D largest eigenvalues of the SuperQ matrix

$$Q_S^{(\mu)} \equiv Q_b^{(\mu)} + Q_f^{(\mu)} \quad (3.2.15)$$

for direction μ . Note that the SuperQ matrix is symmetric and positive semi-definite as it is defined as the sum of two symmetric and positive semi-definite matrices. $U^{(\mu)}$ is also given by the D leading left singular vectors of the extended matrix

$$(M_b^{(\mu)} M_f^{(\mu)}). \quad (3.2.16)$$

The obtained truncation frame is used to calculate the core tensor given in Eq. (3.2.7). This method is referred to as SuperQ.

An advantage of the SuperQ method is that a single singular value decomposition is required for each pair of forward and backward links, i.e., the SuperQ method halves the number of required singular value decompositions compared to the usual HOTRG procedure.

3.2.3 Stabilized finite difference

Thermodynamic observables are typically calculated by taking the derivative of the free energy with respect to a certain control parameter (cf. Sec. 2.6.4). Numerically, this is most easily achieved by using the finite-difference approach. In this method, the partition function is calculated for two slightly different values of the control parameter using two separate tensor networks. It turns out that the evaluated values frequently exhibit discontinuities, prohibiting the calculation of the derivative. These are typically caused by level crossings of nearly degenerate singular values for both terms of the finite difference, resulting in discontinuous changes of the vector spaces that are used for the truncation.

To obtain more stable results for the calculation of the derivative, different approaches were developed. In one method, referred to as the impurity approach, the derivative is taken analytically at the level of the initial tensor network. This results in an impurity system, i.e., a

tensor network that contains an impurity tensor at one specific site of the lattice and identical tensors at all other sites²⁷ [71]. This approach does not suffer from discontinuities originating from the derivative, but it is outperformed by the so-called stabilized finite-difference method [63]. In this approach, the discontinuities are avoided by explicitly matching the bases of the vector spaces used for the truncation. How this is done formally is explained next.

The finite difference of the free energy (2.6.40) with respect to a control parameter p is given as

$$\frac{1}{V\Delta p} \left(\ln Z^+ - \ln Z^- \right), \quad (3.2.17)$$

where both Z^+ and Z^- are calculated using two independent tensor networks. This is no longer the case in the stabilized finite-difference approach, where the truncation frames for calculating Z^+ are modified while taking into account the truncation frames employed in the calculation of Z^- .

As a prerequisite for the stabilized finite-difference approach, one needs to use the same contraction order in both tensor networks and, when the Xie et al. procedure is used, in each blocking step either the forward or backward frame needs to be used in both tensor networks.

For the calculation of Z^- , in a particular blocking step, we denote the obtained left singular vectors by $L_i^{(\mu)-}$ for $i \in \{1, \dots, D^2\}$ (cf. Eq. (3.1.16)), for all directions μ that are perpendicular to the contraction direction. As explained in Sec. 3.1, the truncation frames are then built by choosing the D vectors corresponding to the largest singular values²⁸. We assume that the vectors $L_i^{(\mu)-}$ are ordered such that the corresponding singular values $s_i^{(\mu)-}$ are in descending order, i.e., $s_i^{(\mu)-} \leq s_j^{(\mu)-}$ for $i > j$. Equivalently, the left singular vectors necessary for calculating Z^+ are denoted by $L_i^{(\mu)+}$ and are ordered according to the singular values $s_i^{(\mu)+}$.

Assume that a particular singular value is n -fold nearly degenerate. Thus, the vectors $L_i^{(\mu)-}$, for $i = k_1 \dots k_n$ correspond to nearly degenerate singular values, i.e.,

$$\frac{|s_i^{(\mu)-} - s_{k_1}^{(\mu)-}|}{s_{k_1}^{(\mu)-}} < t \quad \forall i \in \{k_1, \dots, k_n\}. \quad (3.2.18)$$

with threshold t , e.g., $t = 10^{-4}$. For a sufficiently small variation of the control parameter $\Delta p \ll t$, the vectors $L_i^{(\mu)+}$ for $i = k_1 \dots k_n$ correspond to nearly degenerate singular values as well.

Since the vector spaces $\text{span} \{L_i^{(\mu)+} | i \in \{k_1, \dots, k_n\}\}$ and $\text{span} \{L_i^{(\mu)-} | i \in \{k_1, \dots, k_n\}\}$ tend to be identical for vanishing Δp , they will have a large overlap also for a finite variation Δp . However, this is generally not the case for the overlap of the vectors $L_i^{(\mu)+}$ and $L_i^{(\mu)-}$, given by the scalar product $L_i^{(\mu)+} \cdot L_i^{(\mu)-}$. Even if the two vector spaces are identical, the vectors $L_i^{(\mu)+}$ and $L_i^{(\mu)-}$ may form different bases. Therefore, the overlap of the columns of the truncation

²⁷The adaptation of HOTRG in this case is straightforward, except that the truncation matrices are constructed solely by considering the identical tensors. This is necessary for the tensor network to be self-reproducing after each blocking step.

²⁸This applies to general frames, in particular also for SuperQ and the Xie et al. method.

frame, $U_i^{(\mu)+} \cdot U_i^{(\mu)-}$ for $i = 1, \dots, D$, may be small, which can lead to discontinuities in the result.

The optimal solution is to build a new basis of the vector space $\text{span} \{L_i^{(\mu)+} | i \in \{k_1, \dots, k_n\}\}$ with basis vectors

$$\tilde{L}_j^{(\mu)+} = \sum_{i=k_1}^{k_n} c_{ij}^{(\mu)} L_i^{(\mu)+} \text{ with } j \in \{k_1, \dots, k_n\}, \quad (3.2.19)$$

such that the combined overlap of the basis vectors

$$\sum_{i=k_1}^{k_n} \tilde{L}_i^{(\mu)+} \cdot L_i^{(\mu)-} \quad (3.2.20)$$

is maximized. When such a basis is found, we replace $L_j^{(\mu)+}$ with $\tilde{L}_j^{(\mu)+}$ for $j \in \{k_1, \dots, k_n\}$ in the matrix $L^{(\mu)+}$ and continue with the usual HOTRG procedure.

Since the determination of $\tilde{L}_i^{(\mu)+}$ is an expensive task, we use a heuristic approach that is close to the optimal solution. It is given by the following iterative approach.

First, we choose the normalized vector

$$\bar{L}_{k_1}^{(\mu)+} \equiv \frac{1}{\sqrt{\sum_{i=k_1}^{k_n} (L_i^{(\mu)+} \cdot L_{k_1}^{(\mu)-})^2}} \sum_{i=k_1}^{k_n} (L_i^{(\mu)+} \cdot L_{k_1}^{(\mu)-}) L_i^{(\mu)+}, \quad (3.2.21)$$

as this choice maximizes the overlap

$$\bar{L}_{k_1}^{(\mu)+} \cdot L_{k_1}^{(\mu)-}. \quad (3.2.22)$$

One can now use the Gram-Schmidt algorithm to build an orthonormal basis of the subspace $\text{span} \{L_i^{(\mu)+} | i \in \{k_1, \dots, k_n\}\}$ that contains the vector $\bar{L}_{k_1}^{(\mu)+}$ as a basis vector.

As we only want to construct $n - 1$ new vectors using the Gram-Schmidt algorithm, there is one vector $L_j^{(\mu)+}$ with $j \in \{k_1, \dots, k_n\}$ which is not used in the orthogonalization. In principle, it is irrelevant which of the vectors is chosen, as long as it is not orthogonal to $\bar{L}_{k_1}^{(\mu)+}$. For numerical stability, we do not use the vector $L_j^{(\mu)+}$ in the Gram-Schmidt algorithm for which the squared overlap

$$(L_j^{(\mu)+} \cdot L_{k_1}^{(\mu)-})^2 \quad (3.2.23)$$

is maximized. This choice is especially relevant when the squared overlap (3.2.23) is close to one, since otherwise one would orthogonalize almost parallel or antiparallel vectors.

After applying the Gram-Schmidt algorithm, we have determined new basis vectors $\bar{L}_i^{(\mu)+}$ for $i \in \{k_1 + 1, \dots, k_n\}$. We now replace $L_i^{(\mu)+}$ with $\bar{L}_i^{(\mu)+}$ for $i \in \{k_1, \dots, k_n\}$ since at least $\bar{L}_{k_1}^{(\mu)+}$ has an optimized overlap with $L_{k_1}^{(\mu)-}$. After renaming $\bar{L} \rightarrow L$, we repeat the process for the vectors $L_i^{(\mu)+}$ and $L_i^{(\mu)-}$ with $i \in \{k_1 + 1, \dots, k_n\}$. The procedure is applied n times, ensuring that all vectors $L_i^{(\mu)+}$ are adapted to the corresponding vectors $L_i^{(\mu)-}$.

The presented adaptation of the left singular vectors is especially relevant when only some of the vectors with degenerate singular values are included in the truncation frame. However, it is preferable to apply this adaptation to every vector $L_i^{(\mu)+}$ with $i \leq D$, as this further improves numerical stability without significantly increasing the computational cost of the HOTRG algorithm.

3.3 GHOTRG

It turns out that the partition function of lattice QCD in the strong-coupling expansion cannot be expressed in the form of Eq. (3.2.1) in order to apply HOTRG. The integration of the fermion field generates non-local sign factors, i.e., sign factors that cannot be factorized such that each factor only depends on the indices $i_{x,\pm\mu}$ corresponding to site x . However, the partition function can be written as a tensor network, where the local tensors contain, in addition to a numerical part, a specific combination of Grassmann variables (cf. Part III). In the following section, we generalize the HOTRG procedure to make it applicable to tensor networks with this structure, referred to as GHOTRG. In each blocking step, the Grassmann part is self-reproducing, solely yielding a sign factor that is local on the coarse lattice. This sign factor is taken into account in the contraction of two numerical tensors, which are finally truncated using a HOTRG truncation.

Assume that the partition function can be written as

$$Z = \sum_{\mathbf{j}} \int_{\chi} \prod_x (T_x)_{j_{x,-1}j_{x,1}\dots j_{x,-d}j_{x,d}} (K_x)_{F_{x,-1}F_{x,1}\dots F_{x,-d}F_{x,d}} \quad (3.3.1)$$

with $\mathbf{j} \equiv (j_{x,\mu} | \forall x, \mu)$. As for HOTRG, the expression can be visualized by a d -dimensional lattice with tensors T_x and K_x placed on its sites. For every link there is a so-called Grassmann parity $F_{x,\mu} = 0, 1$ that is determined solely by the index $j_{x,\mu}$. T_x is a numerical tensor, while K_x is the Grassmann tensor

$$(K_x)_{F_x} \equiv \prod_{\mu=1}^d (\chi_{x,\mu})^{F_{x,\mu}} \prod_{\mu=1}^d (d\chi_{x,-\mu})^{F_{x,-\mu}}, \quad (3.3.2)$$

depending on the tuple $F_x \equiv (F_{x,-1}F_{x,1}\dots F_{x,-d}F_{x,d})$.²⁹ We also used the notation

$$\prod_{\mu=1}^d \psi_{\mu} \equiv \psi_d \psi_{d-1} \dots \psi_1 \quad (3.3.3)$$

for a product in reverse order. On every link, there is a Grassmann variable $\chi_{x,\mu}$ with corresponding differential $d\chi_{x,\mu}$ which fulfills

$$\int_{\chi_{x,\mu}} d\chi_{x,\mu} \chi_{x,\mu} = 1. \quad (3.3.4)$$

In Eq. (3.3.1), we also introduce the notation

$$\int_{\chi} \equiv \prod_{x,\mu} \left(\int_{\chi_{x,\mu}} \right)^{F_{x,\mu}}. \quad (3.3.5)$$

²⁹For convenience, we again use the notation $i_{x,-\mu} \equiv i_{x-\hat{\mu},\mu}$ for all subsequent quantities that carry a link index, such as $j_{x,\mu}$, $F_{x,\mu}$, and $\chi_{x,\mu}$.

The lattice has a toroidal shape, with boundary conditions for the Grassmann variables chosen to be anti-symmetric in the time direction ($\mu = 1$) and symmetric in all spatial directions. Furthermore, we assume that an entry of T_x is nonzero only if $\sum_{\mu=\pm 1}^{\pm d} F_{x,\mu}$ is even. Under this assumption, K_x can be regarded as commuting.

In a contributing summand in Eq. (3.3.1), there is an integral symbol and a differential associated with every occurring Grassmann variable $\chi_{x,\mu}$ that appears, ensuring that the contribution is real. However, because Grassmann variables are anticommuting, it is not possible to integrate out all of them without producing non-local sign factors. Instead, HOTRG is modified as described below.

3.3.1 Contraction

For a contraction in μ -direction, it is convenient to rewrite

$$(K_x)_{F_x} = \hat{\sigma}_{F_x} (\chi_{x,\mu})^{F_{x,\mu}} \left[\prod_{\nu \neq \mu} (\chi_{x,\nu})^{F_{x,\nu}} \right] \left[\prod_{\nu \neq \mu} (d\chi_{x,-\nu})^{F_{x,-\nu}} \right] (d\chi_{x,-\mu})^{F_{x,-\mu}} \quad (3.3.6)$$

with local sign factor

$$\hat{\sigma}_{F_x} \equiv (-1)^{F_{x,\mu} \sum_{\nu=1}^{\mu-1} F_{x,\nu} + F_{x,-\mu} \sum_{\nu=1}^{\mu-1} F_{x,-\nu}}. \quad (3.3.7)$$

We now define

$$(\mathcal{K}_X)_{F_x, F_y} \equiv \left(\int_{\chi_{x,\mu}} \right)^{F_{x,\mu}} (K_y)_{F_y} (K_x)_{F_x}, \quad (3.3.8)$$

with coarse site $X = (x, y)$. As explained above, the ordering of the Grassmann tensors K irrelevant, but the chosen order is especially convenient, since we directly obtain the factor

$$\left(\int_{\chi_{x,\mu}} \right)^{F_{x,\mu}} (d\chi_{y,-\mu})^{F_{y,-\mu}} (\chi_{x,\mu})^{F_{x,\mu}} = 1 \quad (3.3.9)$$

after inserting Eq. (3.3.6) in Eq. (3.3.8) for both sites. We now reorder the remaining Grassmanns in $(\mathcal{K}_X)_{F_x, F_y}$ in the following way

$$\begin{aligned} (\mathcal{K}_X)_{F_x, F_y} &= \hat{\sigma}_{F_x} \hat{\sigma}_{F_y} \hat{\sigma}_{F_x, F_y} (\chi_{y,\mu})^{F_{y,\mu}} \left[\prod_{\nu \neq \mu} (\chi_{x,\nu})^{F_{x,\nu}} (\chi_{y,\nu})^{F_{y,\nu}} \right] \\ &\times \left[\prod_{\nu \neq \mu} (d\chi_{y,-\nu})^{F_{y,-\nu}} (d\chi_{x,-\nu})^{F_{x,-\nu}} \right] (d\chi_{x,-\mu})^{F_{x,-\mu}}, \end{aligned} \quad (3.3.10)$$

where we defined the sign factor

$$\hat{\sigma}_{F_x, F_y} = (-1)^F, \text{ with } F = \left(\sum_{\substack{\nu=1 \\ \nu \neq \mu}}^d F_{y,-\nu} \right) \left(\sum_{\substack{\nu=1 \\ \nu \neq \mu}}^d F_{x,\nu} \right) + \sum_{\substack{\rho=1 \\ \rho \neq \mu}}^d F_{x,\rho} \sum_{\substack{\nu=\rho \\ \nu \neq \mu}}^d F_{y,\nu} + \sum_{\substack{\rho=1 \\ \rho \neq \mu}}^{d-1} F_{y,-\rho} \sum_{\substack{\nu=\rho+1 \\ \nu \neq \mu}}^d F_{x,-\nu}. \quad (3.3.11)$$

The ordering of the Grassmann variables in Eq. (3.3.10) is particularly useful as the resulting Grassmann structure resembles that in Eq. (3.3.6), where the Grassmann structure for links

perpendicular to the contraction direction is modified. To obtain the initial Grassmann structure albeit on the coarse grid, we insert

$$\left(\int_{\chi_{X,-\nu}} d\chi_{X,-\nu} \chi_{X,-\nu} \right)^{F_{X,-\nu}} = 1 \quad (3.3.12)$$

with the coarse Grassmann parity³⁰

$$F_{X,\nu} \equiv \begin{cases} 0 & \text{if } F_{x,\nu} + F_{y,\nu} \text{ even,} \\ 1 & \text{if } F_{x,\nu} + F_{y,\nu} \text{ odd} \end{cases} \quad (3.3.13)$$

in Eq. (3.3.10) right before the term $(d\chi_{y,-\nu})^{F_{y,-\nu}}$, i.e., inside the product over ν . Consequently, we obtain the term

$$(\chi_{X,-\nu})^{F_{X,-\nu}} (d\chi_{y,-\nu})^{F_{y,-\nu}} (d\chi_{x,-\nu})^{F_{x,-\nu}}, \quad (3.3.14)$$

which, by definition of the coarse Grassmann parity, is commuting. Therefore, we can shift this term in the product over coarse sites from coarse site X to $X - \hat{\nu}$. Since the lattice is closed in ν -direction, at every coarse site X , the term (3.3.14) is replaced by the corresponding term from the coarse site $X + \hat{\nu}$.³¹ Formally, we write³²

$$\int_{\chi} \prod_x (K_x)_{F_x} = \int_{\chi} \prod_X (\mathcal{K}_X)_{F_x, F_y} = \int_{\chi} \prod_X \hat{\sigma}_{F_x} \hat{\sigma}_{F_y} \hat{\sigma}_{F_x, F_y} (\tilde{\mathcal{K}}_X)_{F_X} \quad (3.3.15)$$

with new Grassmann tensor

$$\begin{aligned} (\tilde{\mathcal{K}}_X)_{F_X} &\equiv (\chi_{y,\mu})^{F_{y,\mu}} \left[\prod_{\nu \neq \mu} (\chi_{X,\nu})^{F_{X,\nu}} \left(\int_{\chi_{y,\nu}} d\chi_{y,\nu} \right)^{F_{y,\nu}} \left(\int_{\chi_{x,\nu}} d\chi_{x,\nu} \right)^{F_{x,\nu}} (\chi_{x,\nu})^{F_{x,\nu}} (\chi_{y,\nu})^{F_{y,\nu}} \right] \\ &\quad \times \left[\prod_{\nu \neq \mu} (d\chi_{X,-\nu})^{F_{X,-\nu}} \right] (d\chi_{x,-\mu})^{F_{x,-\mu}} \\ &= (\chi_{y,\mu})^{F_{y,\mu}} \prod_{\nu \neq \mu} (\chi_{X,\nu})^{F_{X,\nu}} \prod_{\nu \neq \mu} (d\chi_{X,-\nu})^{F_{X,-\nu}} (d\chi_{x,-\mu})^{F_{x,-\mu}} \\ &= \hat{\sigma}_{F_X} \prod_{\nu} (\chi_{X,\nu})^{F_{X,\nu}} \prod_{\nu} (d\chi_{X,\nu})^{F_{X,\nu}}, \end{aligned} \quad (3.3.16)$$

where the commuting term originating from the neighbored coarse site is placed in the first bracket. Afterwards, we performed the integration over all Grassmann variables $\chi_{x,\nu}$ and $\chi_{y,\nu}$ for $\nu \neq \mu$, and used Eq. (3.3.6) at the coarse site X , such that the Grassmanns are in canonical ordering again. For convenience, we defined $F_{X,\mu} = F_{y,\mu}$, $F_{X,-\mu} = F_{x,-\mu}$, $\chi_{X,\mu} = \chi_{y,\mu}$, and $\chi_{X,-\mu} = \chi_{x,-\mu}$. It is essential that Eq. (3.3.16) exhibits the same Grassmann structure as Eq. (3.3.2). This shows that the Grassmann structure is self-reproducing after each contraction step.

³⁰We use the same symbol for both the initial and the coarse Grassmann parity, since the meaning is clear from the indices of F . The same applies for the corresponding Grassmann variables.

³¹Note that the boundary conditions must be applied when the shift moves variables over the lattice boundary. For the new Grassmann variables, we impose boundary conditions identical to those of the original variables. As a result, the shift does not generate sign factors, since $F_{X,-\nu} = F_{x,-\nu} + F_{y,-\nu} \bmod 2$.

³²Note that for the integral symbol after the first step, the product over links in Eq. (3.3.5) contains all links of the fine lattice but does not contain contracted links. After the second step, this product runs over all X and μ , i.e., all links of the coarse lattice.

As mentioned before, all signs generated by multiplying two Grassmann tensors are absorbed into the coarse numerical tensor, i.e., we define

$$(T_X)_{j_X} = \sum_{j_{x,\mu}} \sigma_{F_x, F_y} (T_x)_{j_x} (T_y)_{j_y}, \quad (3.3.17)$$

where we introduced the sign³³

$$\sigma_{F_x, F_y} \equiv \hat{\sigma}_{F_x} \hat{\sigma}_{F_y} \hat{\sigma}_{F_x, F_y} \hat{\sigma}_{F_X}, \quad (3.3.18)$$

and used the shorthand notation $j_X \equiv (j_{x,-1}, j_{x,1}, \dots, j_{x,-d}, j_{x,d})$ for fine and coarse site. This yields the partition function

$$Z = \sum_{\mathbf{j}} \int_{\chi} \prod_X (T_X)_{j_X} (K_X)_{F_X} \quad (3.3.19)$$

on the coarse grid. This tensor network has summation indices $j_{X,\mu}$. They are defined as in Eq. (3.2.4), namely as fat indices for directions perpendicular to the contraction direction, and as the original link indices for the remaining links along the contraction direction. j_X is again the tuple containing all indices $j_{X,\mu}$ for $\mu = \pm 1, \dots, \pm d$. All link indices are combined in the tuple $\mathbf{j} \equiv (j_{X,\mu} | \forall X, \mu)$, which we redefined compared to Eq. (3.3.1). Note that by definition, $F_{X,\mu}$ is determined solely by $j_{X,\mu}$.

In order to express the tensor network in the form of Eq. (3.3.19) and Eq. (3.3.17), we used the fact that the new Grassmann tensor

$$(K_X)_{F_X} = \prod_{\nu} (\chi_{X,\nu})^{F_{X,\nu}} \prod_{\nu} (d\chi_{X,\nu})^{F_{X,\nu}} \quad (3.3.20)$$

is independent of the summation index $j_{x,\mu}$. The corresponding integral symbols are collected in the redefined notation

$$\int_{\chi} = \prod_{X,\mu} \left(\int_{\chi_{X,\mu}} \right)^{F_{X,\mu}}. \quad (3.3.21)$$

One can show, $(T_X)_{j_X}$ is nonzero only if

$$\sum_{\rho=\pm 1}^{\pm d} F_{X,\rho} \quad (3.3.22)$$

is even. The proof only requires that the equivalent property on the fine grid holds for T_x and T_y , as assumed at the beginning. Therefore, $(K_X)_{F_X}$ can be regarded as Grassmann even.

3.3.2 Truncation

The remaining task is to truncate $(T_X)_{j_X}$ on all fat links, i.e., on all links perpendicular to the contraction direction μ . As explained in Sec. 3.2.2, for given $\nu = 1, \dots, d$ with $\nu \neq \mu$, the truncation frame $U_{X,\nu}$ ³⁴, for SuperQ and the Xie et al. method, is determined by the unfoldings

$$(M_X)_{j_{X,\nu}, j_X \setminus j_{X,\nu}} = (T_X)_{j_X} \quad \text{and} \quad (M_Y)_{j_{Y,-\nu}, j_Y \setminus j_{Y,-\nu}} = (T_Y)_{j_Y} \quad (3.3.23)$$

³³Due to the structure of σ_{F_x, F_y} , we could assign individual sign factors to the fine-grid tensors T_x and T_y , to the contraction, and to the coarse-grid tensor T_X . The numerical efficiency may depend on whether we apply the sign factors separately or as a whole.

³⁴Not to be confused with the gauge link in lattice QCD.

with coarse site $Y = X + \hat{\nu}$.³⁵ We introduced the tuple $j_X \setminus j_{X,\nu}$, which contains all indices of the tuple j_X except for $j_{X,\nu}$. The entries of the truncation frame are denoted by $(U_{X,\nu})_{j_{X,\nu}, \bar{j}_{X,\nu}}$, where the index $\bar{j}_{X,\nu}$ only has range D , while $j_{X,\nu}$ is the fat link index with range D^2 .

Since the Grassmann tensor depends on the indices $j_{X,\nu}$ through the coarse Grassmann parity, special care must be taken when applying the truncation matrices. We recall that $(T_X)_{j_X}$ is nonzero only if $\sum_{\rho=\pm 1}^{\pm d} F_{X,\rho}$ is even. Therefore, $(M_X)_{j_{X,\nu}, j_X \setminus j_{X,\nu}}$ is, up to permutations in rows and columns, block diagonal. There are two blocks, differentiated by

$$F_{X,\nu} = \left(\sum_{\substack{\rho=\pm 1 \\ \rho \neq \nu}}^{\pm d} F_{X,\rho} \right) \bmod 2 = 1 \quad \text{and} \quad F_{X,\nu} = \left(\sum_{\substack{\rho=\pm 1 \\ \rho \neq \nu}}^{\pm d} F_{X,\rho} \right) \bmod 2 = 0. \quad (3.3.24)$$

Note that $(M_X)_{j_{X,\nu}, j_X \setminus j_{X,\nu}}$ and $(M_Y)_{j_{Y,-\nu}, j_Y \setminus j_{Y,-\nu}}$ have the same vertical block structure³⁶, since $F_{X,\nu} = F_{Y,-\nu}$ holds in the contraction. Therefore, the truncation frame $(U_{X,\nu})_{j_{X,\nu}, \bar{j}_{X,\nu}}$ is, up to permutations in rows and columns, block diagonal as well, regardless of whether one uses SuperQ or the Xie et al. procedure. We can distinguish the blocks of $U_{X,\nu}$ according to $F_{X,\nu} = 1$ or $F_{X,\nu} = 0$.

This result also allows us to define the Grassmann parity of $\bar{j}_{X,\nu}$, denoted by $\bar{F}_{X,\nu}$. Formally, it is defined as follows: for given $\bar{j}_{X,\nu}$, consider all nonzero entries $(U_{X,\nu})_{j_{X,\nu}, \bar{j}_{X,\nu}}$. The indices $j_{X,\nu}$ of these entries share the same Grassmann parity $F_{X,\nu}$. When this parity is zero (one), we define $\bar{F}_{X,\nu}$ to be zero (one), respectively. For the next steps, it is convenient to also define $\bar{j}_{X,\pm\mu} = j_{X,\pm\mu}$ and $\bar{F}_{X,\pm\mu} = F_{X,\pm\mu}$ for links in contraction direction μ .

We now consider the truncation

$$\left(\prod_{\substack{\nu=\pm 1 \\ |\nu| \neq \mu}}^{\pm d} \sum_{j_{X,\nu}} (U_{X,\nu})_{j_{X,\nu}, \bar{j}_{X,\nu}} \right) (T_X)_{j_X} (K_X)_{F_X} = \left(\prod_{\substack{\nu=\pm 1 \\ |\nu| \neq \mu}}^{\pm d} \sum_{j_{X,\nu}} (U_{X,\nu})_{j_{X,\nu}, \bar{j}_{X,\nu}} \right) (T_X)_{j_X} (K_X)_{\bar{F}_X}, \quad (3.3.25)$$

where it holds that $U_{X,-\nu} = U_{X-\hat{\nu},\nu}$, due to the backward-forward symmetry. Since for every nonzero entry $(U_{X,\nu})_{j_{X,\nu}, \bar{j}_{X,\nu}}$, we have $\bar{F}_{X,\nu} = F_{X,\nu}$, we replaced F_X by $\bar{F}_X$ in the Grassmann tensor with $\bar{F}_x = (\bar{F}_{x,-1} \bar{F}_{x,1} \dots \bar{F}_{x,-d} \bar{F}_{x,d})$. This is a crucial result, since it shows that the Grassmann tensor can be pulled out of the sum, obtaining the usual HOTRG truncation for the numerical coarse tensor. Hence, we define the new tensor

$$(\bar{T}_X)_{\bar{j}_X} = \left(\prod_{\substack{\nu=\pm 1 \\ |\nu| \neq \mu}}^{\pm d} \sum_{j_{X,\nu}} (U_{X,\nu})_{j_{X,\nu}, \bar{j}_{X,\nu}} \right) (T_X)_{j_X} \quad (3.3.26)$$

with $\bar{j}_X \equiv (\bar{j}_{X,-1} \bar{j}_{X,1} \dots \bar{j}_{X,-d} \bar{j}_{X,d})$, yielding the partition function³⁷

$$Z = \sum_{\bar{j}} \int_{\chi} \prod_X (\bar{T}_X)_{\bar{j}_X} (K_X)_{\bar{F}_X}. \quad (3.3.27)$$

³⁵Note that this is a slight generalization of Sec. 3.2.2, in which M , in general, depends on the coarse site.

³⁶Note that the number of columns can be different.

³⁷The integral symbol is now expressed in terms of the parities $\bar{F}_{X,\mu}$ as well, so that it can be taken outside of the sums over $j_{x,\mu}$.

This tensor network has new summation indices $\bar{j}_{X,\mu}$, combined in the tuple $\bar{\mathbf{j}} \equiv (\bar{j}_{X,\mu} | \forall X, \mu)$.

Note that by construction, $\bar{F}_{X,\mu}$ is determined solely by $\bar{j}_{X,\mu}$ and $\bar{T}_X$ is nonzero only if $\sum_{\mu=\pm 1}^{\pm d} \bar{F}_{X,\mu}$ is even. Therefore, the obtained tensor network has the same structure as the one in Eq. (3.3.1) albeit on a lattice with half the number of sites.

3.3.3 Iterative procedure

After renaming

$$\bar{T} \rightarrow T, \quad \bar{j} \rightarrow j, \quad X \rightarrow x, \quad \bar{F} \rightarrow F, \quad (3.3.28)$$

the described GHOTRG algorithm can be applied again to block the lattice further. For every blocking step, a new truncation direction is selected. The simplest option is to select the contraction direction $\mu = (k-1) \bmod d+1$ for the k -th blocking step, which was used in this thesis. It is possible to find contraction orderings that reduce the truncation error, as is done in the improved contraction ordering (ICO) [63].

Since the partition function Z grows exponentially with the volume, the tensor entries generally grow exponentially with the number of applied GHOTRG blocking steps. Consequently, we rescale the tensor entries after each contraction step, as explained next. After the first blocking step (including the renaming), we determine the maximal absolute value

$$t_1 = \max_{x,j_x} |(T_x)_{j_x}| \quad (3.3.29)$$

of all tensor entries, and define the rescaled partition function

$$Z^{(1)} = \sum_{\mathbf{j}} \int_{\chi} \prod_x (T_x^{(1)})_{j_x} (K_x)_{F_x} \quad \text{with} \quad (T_x^{(1)})_{j_x} = \frac{1}{t_1} (T_x)_{j_x}. \quad (3.3.30)$$

Consequently, the full partition function is given by

$$Z = t_1^{V_1} Z^{(1)}, \quad (3.3.31)$$

where V_1 is the number of lattice sites after the first blocking step, i.e., $V_1 = V/2$. We can apply the GHOTRG algorithm again, this time to the rescaled partition function $Z^{(1)}$. After this blocking step, we rescale the resulting numerical tensors using their maximum t_2 to obtain $Z^{(2)}$ with $Z^{(1)} = t_2^{V_2} Z^{(2)}$ and $V_2 = V/2^2$. This is repeated for each blocking step. After the k -th blocking step, we rescale the resulting numerical tensors using their maximum t_k to obtain $Z^{(k)}$ with $Z^{(k-1)} = t_k^{V_k} Z^{(k)}$ and $V_k = V/2^k$. $Z^{(k)}$ is connected to the full partition function via

$$Z = Z^{(k)} \prod_{i=1}^k t_i^{V_i}, \quad (3.3.32)$$

resulting in the expression

$$\frac{\ln Z}{V} = \frac{\ln Z^{(k)}}{V} + \sum_{i=1}^k \frac{\ln t_i}{2^i} \quad (3.3.33)$$

for the free energy.

The computational cost of the GHOTRG approach scales logarithmically with the volume when the tensors have no site dependence. In this case, the pairwise contraction of the tensors is the same for all coarse sites, and has to be performed only once per blocking step. Further, there are only $d - 2$ different truncation frames to be calculated in this case. Nevertheless, we considered the more general case of site-dependent tensors, as it will be necessary in a more sophisticated approach later on (cf. Sec. 8). Additionally, the implementation of staggered phases results in site-dependent tensors, however due to their specific structure, the coarse tensor will be site-independent after the blocking has been done at least once in each direction, and there will be no further site dependence caused by the staggered phases.

3.3.4 Tracing

As one possibility, we can repeat the GHOTRG algorithm until the coarse lattice consists of a single site x . Assuming we applied in total k blocking steps, the rescaled partition function reads

$$Z^{(k)} = \sum_{\vec{j}} \int_{\chi} T_{j_{x,1}j_{x,1}\dots j_{x,d}j_{x,d}} K_{F_{x,1}F_{x,1}\dots F_{x,d}F_{x,d}}, \quad (3.3.34)$$

where T and K are the obtained numerical tensor and Grassmann tensor, respectively, after all blocking steps including the rescaling. We can now integrate the Grassmann tensor

$$\int_{\chi} K_{F_{x,1}F_{x,1}\dots F_{x,d}F_{x,d}} = (-1)^{F_{x,1}} \prod_{\mu} \left(\int_{\chi_{x,\mu}} \right)^{F_{x,\mu}} \prod_{\mu} (\chi_{x,\mu})^{F_{x,\mu}} \prod_{\mu} (d\chi_{x,\mu})^{F_{x,\mu}} = (-1)^{\sum_{\nu=2}^d F_{x,\nu}}, \quad (3.3.35)$$

where we used the boundary conditions

$$d\chi_{x,-1} = -d\chi_{x,1}, \quad \text{and} \quad d\chi_{x,-\mu} = d\chi_{x,\mu}, \quad \text{for } \mu \neq 1, \quad (3.3.36)$$

and obtain the result for the rescaled partition function

$$Z^{(k)} = \sum_{j_{x,1}, \dots, j_{x,d}} (-1)^{\sum_{\nu=2}^d F_{x,\nu}} T_{j_{x,1}j_{x,1}\dots j_{x,d}j_{x,d}}. \quad (3.3.37)$$

This step is generally referred to as taking the trace of the tensor T .

3.3.5 ASAP tracing

In an improved version, which is used in this thesis, the tracing is optimized in order to reduce the truncation error. Whenever the lattice extent in a specific direction equals one, this direction can be removed from the network. In particular, the individual sums of the trace (3.3.37) are evaluated as soon as possible (ASAP).

Assume there is a dimension μ on the coarse grid, in which the lattice extent equals one.

Therefore, the Grassmann contribution reads

$$\begin{aligned}
& \left(\int_{\chi_{x,\mu}} \right)^{F_{x,\mu}} (K_x)_{F_x} \\
&= (-1)^{F_{x,\mu} \delta_{\mu,1}} \left(\int_{\chi_{x,\mu}} \right)^{F_{x,\mu}} \prod_{\nu} (\chi_{x,\nu})^{F_{x,\nu}} (d\chi_{x,-d})^{F_{x,-d}} \dots (d\chi_{x,\mu})^{F_{x,\mu}} \dots (d\chi_{x,-1})^{F_{x,-1}} \\
&= \sigma_{\mu} (\bar{K}_x)_{\bar{F}_x},
\end{aligned} \tag{3.3.38}$$

where in the first step we applied the boundary condition for $(d\chi_{x,-\mu})^{F_{x,-\mu}}$, and in the second step we rearranged the order of the Grassmann variables such that $\chi_{x,\mu}$ can be integrated and defined

$$\sigma_{\mu} = (-1)^{F_{x,\mu}(1+\delta_{\mu,\text{time}}+\sum_{\nu=\mu+1}^d (F_{x,\nu}+F_{x,-\nu}))}. \tag{3.3.39}$$

The reduced Grassmann tensor reads

$$(\bar{K}_x)_{\bar{F}_x} \equiv \prod_{\nu \neq \mu} (\chi_{x,\nu})^{F_{x,\nu}} \prod_{\nu \neq \mu} (d\chi_{x,-\nu})^{F_{x,-\nu}}, \tag{3.3.40}$$

where the tuple $\bar{F}_x$ contains all parities $F_{x,\nu}$ for $\nu = \pm 1, \dots, \pm d$, except for $F_{x,\mu}$ and $F_{x,-\mu}$. The resulting sign is absorbed into the lower dimensional tensor

$$(\bar{T}_x)_{\bar{j}_x} \equiv \sum_{j_{x,\mu}} \sigma_{\mu} (T_x)_{j_x}. \tag{3.3.41}$$

Similarly to $\bar{F}_x$, the tuple $\bar{j}_x$ contains all indices $j_{x,\nu}$ for $\nu = \pm 1, \dots, \pm d$, except for $j_{x,\mu}$ and $j_{x,-\mu}$. Assuming k blocking steps including the rescaling are already applied, the rescaled partition function reads

$$Z^{(k)} = \sum_{\bar{j}} \int_{\chi} \prod_x (\bar{T}_x)_{\bar{j}_x} (\bar{K}_x)_{\bar{F}_x}, \tag{3.3.42}$$

where we defined the tuple $\bar{j} \equiv (j_{x,\nu} | \forall x, \nu \text{ with } \nu \neq \mu)$. The product over links inside the integral symbol then only contains directions perpendicular to μ .

After renaming

$$\bar{T} \rightarrow T, \quad \bar{K} \rightarrow K, \quad \bar{j} \rightarrow j, \quad \bar{F} \rightarrow F, \tag{3.3.43}$$

the GHOTRG approach can be applied again, where the directions are now only $1, \dots, \mu-1, \mu+1, \dots, d$, i.e., the lattice only has $d-1$ dimensions. ASAP tracing is used again when the lattice extent in one of the remaining dimensions equals one, to reduce the dimension of the lattice further. This is continued until a 1-dimensional lattice is reduced to a 0-dimensional lattice, i.e., a scalar. Together with Eq. (3.3.33), this yields the final result for the tensor network.

3.3.6 Truncation of the initial numerical tensor

Depending on the size of initial numerical tensor T_x in Eq. (3.3.1), a truncation of the initial tensor may be necessary as well. In this case, the relevant unfoldings are given by

$$(M_x)_{j_{x,\nu}, j_x/j_{x,\nu}} = (T_x)_{j_{x,\nu}}. \tag{3.3.44}$$

Due to the necessary backward-forward symmetry for the truncation frames, the frame $(U_{x,\nu})_{j_{x,\nu},\bar{j}_{x,\nu}}$ with $\nu = 1, \dots, d$ is determined by the unfoldings $(M_x)_{j_{x,\nu},j_x/j_{x,\nu}}$ and $(M_y)_{j_{y,-\nu},j_y/j_{y,-\nu}}$ with $y = x + \hat{\nu}$, regardless of whether one uses SuperQ or the Xie et al. truncation procedure.

Analogue to usual GHOTRG blocking procedure, the truncation matrix $(U_{x,\nu})_{j_{x,\nu},\bar{j}_{x,\nu}}$ is, up to row and column permutations, block diagonal, where both blocks can be distinguished by $F_{x,\nu}$. We can also define the parity $\bar{F}_{x,\mu}$ for given link $\bar{j}_{x,\nu}$ such that $(K_x)_{F_x}$ can be replaced by $(K_x)_{\bar{F}_x}$. The partition function is then given by

$$Z = \sum_{\bar{\mathbf{j}}} \int_{\chi} \prod_x (\bar{T}_x)_{\bar{j}_x} (K_x)_{\bar{F}_x} \quad (3.3.45)$$

with tuple $\bar{\mathbf{j}} \equiv (j_{x,\mu} | \forall x, \mu)$ and truncated tensor

$$(\bar{T}_x)_{\bar{j}_x} = \left(\prod_{\nu=\pm 1}^{\pm d} \sum_{j_{x,\nu}} (U_{x,\nu})_{j_{x,\nu},\bar{j}_{x,\nu}} \right) (T_x)_{j_x}. \quad (3.3.46)$$

After the renaming

$$\bar{T} \rightarrow T, \quad \bar{j} \rightarrow j, \quad \bar{F} \rightarrow F, \quad (3.3.47)$$

we apply the usual GHOTRG algorithm.

III. Deriving the tensor network

In the following, we go through all the necessary steps to rewrite the partition function of lattice QCD in the strong-coupling expansion with staggered fermions in terms of a tensor network that is suitable for the GHOTRG algorithm presented in Sec. 3.3. The derivation is carried out for a general number of flavors N_f , for a general number of space-time dimensions d , for arbitrary order in β , and also for general gauge group $SU(N_c)$. For simplicity, we set the lattice spacing a to one. The derivation of this tensor network is published in Ref. [27], in collaboration with Jacques Bloch, Robert Lohmayer, and Tilo Wettig.

4 Dual variables and evaluation of the gauge integral

In this section, we proceed as follows. First, we Taylor-expand the exponentials of the various contributions to the action, resulting in dual variables (i.e., occupation numbers). In Sec. 4.2 we explicitly introduce all existing color indices and describe their enumeration in detail. In Sec. 4.3 we perform the gauge-link integral for a single link and simplify the result based on the color contraction with the Grassmann variables in the hopping terms.

4.1 Taylor expansions of the Boltzmann factor

At first, we write the partition function (2.6.37) as

$$Z = \int \left[\prod_x d\psi_x d\bar{\psi}_x \right] \left[\prod_{x,\mu} dU_{x,\mu} \right] e^{S_M} \left[\prod_{x,\mu} \prod_f e^{S_{x,\mu}^{f+}} e^{S_{x,\mu}^{f-}} \right] \left[\prod_{x,\mu \neq \nu} e^{S_{x,\mu\nu}^G} \right], \quad (4.1.1)$$

where the exponentials for the hopping terms and the gauge action are factorized over the links and plaquettes, respectively. An intermediate goal is to integrate out the gauge links analytically. To this end, we now perform separate Taylor expansions of the exponentials for the forward and backward hopping terms on each link (x, μ) and for each flavor, with summation indices¹ $k_{x,\mu}^f$ and $\bar{k}_{x,\mu}^f$, respectively. Similarly, we perform a Taylor expansion of the exponential of the gauge action on each oriented plaquette $(x, \mu\nu)$ with summation index $n_{x,\mu\nu}$. All summation indices start at zero. After performing the Taylor expansions we rewrite the products of sums as sums of products,

$$Z = \int \left[\prod_x d\psi_x d\bar{\psi}_x \right] \left[\prod_{x,\mu} dU_{x,\mu} \right] e^{S_M} \sum_{\mathbf{k}, \mathbf{n}} \left[\prod_{x,\mu} \prod_f \frac{(S_{x,\mu}^{f+})^{k_{x,\mu}^f} (S_{x,\mu}^{f-})^{\bar{k}_{x,\mu}^f}}{k_{x,\mu}^f! \bar{k}_{x,\mu}^f!} \right] \left[\prod_{x,\mu \neq \nu} \frac{(S_{x,\mu\nu}^G)^{n_{x,\mu\nu}}}{n_{x,\mu\nu}!} \right] \quad (4.1.2)$$

with dual variables $\mathbf{k} \equiv ((k_{x,\mu}^f, \bar{k}_{x,\mu}^f) | \forall x, \mu, f)$ and $\mathbf{n} \equiv (n_{x,\mu\nu} | \forall x, \mu, \nu; \mu \neq \nu)$ that contain all link and plaquette occupation numbers of the lattice. These occupation numbers can be viewed as dual fields, which also have periodic boundary conditions.

¹We will frequently refer to these and other summation indices as “occupation numbers”.

To be able to express the partition function as the trace of a tensor network, we exchange the order of integration and summation to obtain

$$Z = \sum_{\mathbf{k}, \mathbf{n}} \int \left[\prod_x d\psi_x d\bar{\psi}_x \right] \left[\prod_{x,\mu} dU_{x,\mu} \right] e^{S_M} \left[\prod_{x,\mu} \prod_f \frac{(S_{x,\mu}^{f+})^{k_{x,\mu}^f} (S_{x,\mu}^{f-})^{\bar{k}_{x,\mu}^f}}{k_{x,\mu}^f! \bar{k}_{x,\mu}^f!} \right] \left[\prod_{x,\mu \neq \nu} \frac{(S_{x,\mu\nu}^G)^{n_{x,\mu\nu}}}{n_{x,\mu\nu}!} \right]. \quad (4.1.3)$$

For the summation over $\mathbf{k}$ this exchange is justified since the sums are finite due to the nature of the Grassmann variables and since the gauge integral is convergent. However, because the summation over $\mathbf{n}$ involves infinite sums the exchange is potentially questionable. A formal convergence analysis appears to be nontrivial, and therefore this issue needs to be kept in mind when numerical results are interpreted.

4.2 Disentangling color indices

To do the gauge-link integrals we have to rearrange the gauge links coming from the hopping terms and the plaquette action in such a way that all contributions from a specific gauge link are gathered. This is done by writing all summations over color indices explicitly, which allows us to move components of the gauge links freely. We use the convention that gauge links U always carry color indices i and j , while adjoint gauge links $U^\dagger$ carry color indices ℓ and m . To distinguish all the color indices occurring in Eq. (4.1.3) we need to assign two additional labels to each color index. First, in order to distinguish the color indices of different gauge links, we assign a link label to the color indices. Second, as the same gauge link $U_{x,\mu}$ and its adjoint $U_{x,\mu}^\dagger$ can occur several times in each term contributing to the partition function, we distinguish the color indices corresponding to these different occurrences by a label a . Hence, our color indices are $i_{x,\mu}^a, j_{x,\mu}^a, \ell_{x,\mu}^a$, and $m_{x,\mu}^a$, again with periodic boundary conditions.

In the following we discuss the label a in detail. The range of the label a consists of several subranges. For $U_{x,\mu}$, the subrange of a from the forward hopping term for flavor f contains $k_{x,\mu}^f$ elements, and there are N_f such subranges. For the plaquette terms, we note that in d dimensions, a given gauge link $U_{x,\mu}$ is contained in $2(d-1)$ plaquettes, i.e., $U_{x,\mu\nu}^P$ and $U_{x-\hat{\nu},\nu\mu}^P$ with $\nu \neq \mu$. Hence there are $2(d-1)$ subranges of a for the plaquette terms. In summary, the sizes of the $N_f + 2(d-1)$ subranges of a for the various contributions are given by

$$U_{x,\mu} : \quad |\text{subrange}(a)| = \begin{cases} k_{x,\mu}^f & \text{forward fermion hopping for flavor } f, \\ n_{x,\mu\nu} & \text{plaquettes } U_{x,\mu\nu}^P \text{ for } \nu \neq \mu, \\ n_{x-\hat{\nu},\nu\mu} & \text{plaquettes } U_{x-\hat{\nu},\nu\mu}^P \text{ for } \nu \neq \mu, \end{cases} \quad (4.2.1a)$$

where $|\dots|$ denotes the number of elements of the (sub)range. Similarly, for $U_{x,\mu}^\dagger$ the subrange sizes are

$$U_{x,\mu}^\dagger : \quad |\text{subrange}(a)| = \begin{cases} \bar{k}_{x,\mu}^f & \text{backward fermion hopping for flavor } f, \\ n_{x,\nu\mu} & \text{plaquettes } U_{x,\nu\mu}^P \text{ for } \nu \neq \mu, \\ n_{x-\hat{\nu},\mu\nu} & \text{plaquettes } U_{x-\hat{\nu},\mu\nu}^P \text{ for } \nu \neq \mu. \end{cases} \quad (4.2.1b)$$

We now define

$$k_{x,\mu} \equiv \sum_f k_{x,\mu}^f \quad \text{and} \quad \bar{k}_{x,\mu} \equiv \sum_f \bar{k}_{x,\mu}^f \quad (4.2.2)$$

and combine the subranges of a to a full range, whose size is given by

$$U_{x,\mu} : \quad |\text{range}(a)| = g_{x,\mu} \equiv k_{x,\mu} + \sum_{\substack{\nu \\ \nu \neq \mu}} (n_{x,\mu\nu} + n_{x-\hat{\nu},\nu\mu}), \quad (4.2.3a)$$

$$U_{x,\mu}^\dagger : \quad |\text{range}(a)| = \bar{g}_{x,\mu} \equiv \bar{k}_{x,\mu} + \sum_{\substack{\nu \\ \nu \neq \mu}} (n_{x,\nu\mu} + n_{x-\hat{\nu},\mu\nu}), \quad (4.2.3b)$$

respectively. To access the different subranges within the full range, we define the offsets²

$$\kappa_{x,\mu}^f \equiv \sum_{f'=1}^{f-1} k_{x,\mu}^{f'}, \quad \bar{\kappa}_{x,\mu}^f \equiv \sum_{f'=1}^{f-1} \bar{k}_{x,\mu}^{f'}, \quad (4.2.4a)$$

$$\omega_{x,\mu}^{+\nu} \equiv k_{x,\mu} + \sum_{\substack{\sigma < \nu \\ \sigma \neq \mu}} (n_{x,\mu\sigma} + n_{x-\hat{\sigma},\sigma\mu}), \quad \omega_{x,\mu}^{-\nu} \equiv \omega_{x,\mu}^{+\nu} + n_{x,\mu\nu}, \quad (4.2.4b)$$

$$\bar{\omega}_{x,\mu}^{+\nu} \equiv \bar{k}_{x,\mu} + \sum_{\substack{\sigma < \nu \\ \sigma \neq \mu}} (n_{x,\sigma\mu} + n_{x-\hat{\sigma},\mu\sigma}), \quad \bar{\omega}_{x,\mu}^{-\nu} \equiv \bar{\omega}_{x,\mu}^{+\nu} + n_{x,\nu\mu}. \quad (4.2.4c)$$

With these definitions and using the Einstein summation convention, the color contractions in the terms coming from $S_{x,\mu}^{f+}$, $S_{x,\mu}^{f-}$, and $S_{x,\mu\nu}^G$ are written as

$$(\bar{\psi}_x^f U_{x,\mu} \psi_{x+\hat{\mu}}^f)^{k_{x,\mu}^f} = \prod_{a=1}^{k_{x,\mu}^f} \bar{\psi}_x^{f,i_{x,\mu}^{a+\kappa_{x,\mu}^f}} U_{x,\mu}^{i_{x,\mu}^{a+\kappa_{x,\mu}^f}, j_{x,\mu}^{a+\kappa_{x,\mu}^f}} \psi_{x+\hat{\mu}}^{f,j_{x,\mu}^{a+\kappa_{x,\mu}^f}}, \quad (4.2.5a)$$

$$(\bar{\psi}_{x+\hat{\mu}}^f U_{x,\mu}^\dagger \psi_x^f)^{\bar{k}_{x,\mu}^f} = \prod_{a=1}^{\bar{k}_{x,\mu}^f} \bar{\psi}_{x+\hat{\mu}}^{f,\ell_{x,\mu}^{a+\bar{\kappa}_{x,\mu}^f}} U_{x,\mu}^{\ell_{x,\mu}^{a+\bar{\kappa}_{x,\mu}^f}, m_{x,\mu}^{a+\bar{\kappa}_{x,\mu}^f}} \psi_x^{f,m_{x,\mu}^{a+\bar{\kappa}_{x,\mu}^f}}, \quad (4.2.5b)$$

$$\left(\text{tr} U_{x,\mu\nu}^P \right)^{n_{x,\mu\nu}} = \quad (4.2.5c)$$

$$\prod_{a=1}^{n_{x,\mu\nu}} \Delta_{x,\mu\nu}^a U_{x,\mu}^{i_{x,\mu}^{a+\omega_{x,\mu}^{+\nu}}, j_{x,\mu}^{a+\omega_{x,\mu}^{+\nu}}} U_{x+\hat{\mu},\nu}^{i_{x+\hat{\mu},\nu}^{a+\omega_{x+\hat{\mu},\nu}^{-\mu}}, j_{x+\hat{\mu},\nu}^{a+\omega_{x+\hat{\mu},\nu}^{-\mu}}} U_{x+\hat{\nu},\mu}^{\ell_{x+\hat{\nu},\mu}^{a+\bar{\omega}_{x+\hat{\nu},\mu}^{-\nu}}, m_{x+\hat{\nu},\mu}^{a+\bar{\omega}_{x+\hat{\nu},\mu}^{-\nu}}} U_{x,\nu}^{\ell_{x,\nu}^{a+\bar{\omega}_{x,\nu}^{+\mu}}, m_{x,\nu}^{a+\bar{\omega}_{x,\nu}^{+\mu}}}, \quad (4.2.5d)$$

where we have rewritten the matrix multiplications and traces in Eq. (4.2.5d) using the quantity

$$\Delta_{x,\mu\nu}^a \equiv \delta_{j_{x,\mu}^{a+\omega_{x,\mu}^{+\nu}}, i_{x+\hat{\mu},\nu}^{a+\omega_{x+\hat{\mu},\nu}^{-\mu}}} \delta_{j_{x+\hat{\mu},\nu}^{a+\omega_{x+\hat{\mu},\nu}^{-\mu}}, \ell_{x+\hat{\nu},\mu}^{a+\bar{\omega}_{x+\hat{\nu},\mu}^{-\nu}}} \delta_{m_{x+\hat{\nu},\mu}^{a+\bar{\omega}_{x+\hat{\nu},\mu}^{-\nu}}, \ell_{x,\nu}^{a+\bar{\omega}_{x,\nu}^{+\mu}}} \delta_{m_{x,\nu}^{a+\bar{\omega}_{x,\nu}^{+\mu}}, i_{x,\mu}^{a+\omega_{x,\mu}^{+\nu}}}. \quad (4.2.6)$$

The motivation behind this rewriting is that in the construction of the tensor formulation below it is more convenient to retain independent indices on U and $U^\dagger$ to be able to manipulate individual indices. We see from Eq. (4.2.5) that the color indices $i_{x,\mu}^a$ and $m_{x,\mu}^a$ can be associated with site x , while $j_{x,\mu}^a$ and $\ell_{x,\mu}^a$ can be associated with site $x + \hat{\mu}$.

We can now gather all gauge-field contributions corresponding to the same link. After all the dust has settled, the partition function can be written as

$$Z = \sum_{\mathbf{k}, \mathbf{n}} \int \left[\prod_x d\psi_x d\bar{\psi}_x \right] \mathcal{GW} \bar{\Delta} \prod_{x,\mu} \mathcal{I}^{g_{x,\mu} \bar{g}_{x,\mu}}, \quad (4.2.7)$$

²The offsets ω and $\bar{\omega}$ correspond to U and $U^\dagger$, respectively. The superscript $\pm\nu$ indicates that the plaquette of interest extends in the positive/negative ν -direction perpendicular to the link (x, μ) .

where all contributions involving the same gauge-field matrix U are combined in the integral

$$(\mathcal{I}^{g\bar{g}})^{ij,\ell\mathbf{m}} \equiv \int_{\text{SU}(N_c)} dU \left[\prod_{a=1}^g U^{i^a j^a} \right] \left[\prod_{a=1}^{\bar{g}} U^{\dagger \ell^a m^a} \right] \quad (4.2.8)$$

with color indices

$$\mathbf{i} \equiv (i^1, \dots, i^g), \quad \mathbf{j} \equiv (j^1, \dots, j^g), \quad \boldsymbol{\ell} \equiv (\ell^1, \dots, \ell^{\bar{g}}), \quad \mathbf{m} \equiv (m^1, \dots, m^{\bar{g}}). \quad (4.2.9)$$

The gauge-link occupation numbers g and $\bar{g}$ in Eq. (4.2.7) are defined in Eq. (4.2.3). All Grassmann variables on the lattice are combined in the factor

$$\mathcal{G} = e^{S_M} \prod_{x,\mu} \mathcal{G}_{x,\mu} \quad \text{with} \quad \mathcal{G}_{x,\mu} = \prod_f \left[\prod_{a=1}^{k_{x,\mu}^f} \bar{\psi}_x^{f,i_{x,\mu}^{a+\kappa_{x,\mu}^f}} \psi_{x+\hat{\mu}}^{f,j_{x,\mu}^{a+\kappa_{x,\mu}^f}} \right] \left[\prod_{a=1}^{\bar{k}_{x,\mu}^f} \psi_x^{f,m_{x,\mu}^{a+\bar{\kappa}_{x,\mu}^f}} \bar{\psi}_{x+\hat{\mu}}^{f,\ell_{x,\mu}^{a+\bar{\kappa}_{x,\mu}^f}} \right], \quad (4.2.10)$$

where in the last factor we have changed the order of the Grassmann variables to eliminate the minus sign in Eq. (2.6.36). The integration of $\mathcal{G}$ will be discussed in Sec. 5.1. The remaining quantities in Eq. (4.2.7) are

$$\bar{\Delta} \equiv \prod_{x,\mu \neq \nu} \prod_{a=1}^{n_{x,\mu\nu}} \Delta_{x,\mu\nu}^a, \quad (4.2.11)$$

$$\mathcal{W} \equiv \left[\prod_{x,\mu \neq \nu} \frac{(\beta/2N_c)^{n_{x,\mu\nu}}}{n_{x,\mu\nu}!} \right] \left[\prod_{x,\mu} \prod_f \frac{\eta_{x,\mu}^{k_{x,\mu}^f, \mu + \bar{k}_{x,\mu}^f}}{k_{x,\mu}^f! \bar{k}_{x,\mu}^f!} \right] \left[\prod_x \prod_f e^{\mu_f(k_{x,1}^f - \bar{k}_{x,1}^f)} \right]. \quad (4.2.12)$$

4.3 Gauge-link integral and simplifications

4.3.1 Gauge-link integral

Now that we have collected all gauge-field contributions for every link, we are ready to evaluate the resulting $\text{SU}(N_c)$ integrals. The relevant integral, given in Eq. (4.2.8), was first solved by Creutz [30], but we will use a more efficient solution, which was recently presented by Gagliardi and Unger [32, 13] and by Borisenko et al. [33]. These authors gave expressions for the case of $g \geq \bar{g}$ and pointed out that using the defining property of the Haar measure one obtains $(\mathcal{I}^{g\bar{g}})^{ij,\ell\mathbf{m}} = (\mathcal{I}^{\bar{g}g})^{\ell\mathbf{m},ij}$, from which the case of $g < \bar{g}$ follows. After defining $p \equiv \min(g, \bar{g})$ and $q \equiv |g - \bar{g}|/N_c$, both cases can be combined in the result

$$(\mathcal{I}^{g\bar{g}})^{ij,\ell\mathbf{m}} = \left[\prod_{s=1}^{N_c-1} \frac{s!}{(s+q)!} \right] \sum_{\pi, \sigma \in S_p} \widetilde{\text{Wg}}_{N_c}^{q,p}(\pi \cdot \sigma^{-1}) \sum_{\alpha \in A} (I_{\alpha\pi\sigma}^b)^{i\mathbf{m}} (I_{\alpha\pi\sigma}^e)^{j\boldsymbol{\ell}} \quad \text{for } q \in \mathbb{N}_0. \quad (4.3.1)$$

If $q \notin \mathbb{N}_0$ the integral is zero. In the remainder of this section we describe the various ingredients of the result (4.3.1).

The sums over π and σ run over all elements of the symmetric group S_p . Here, we consider the special case that a permutation π acts on the sequence of numbers $(1, \dots, p)$, resulting in the permuted sequence $(\pi_1, \dots, \pi_p)$, and similarly for σ . Note that we use the same symbol π (or σ) for both the permutation and the permuted sequence.

The generalized Weingarten function $\widetilde{\text{Wg}}$ is defined for $\pi \in S_p$ by [32]

$$\widetilde{\text{Wg}}_{N_c}^{q,p}(\pi) \equiv \frac{1}{(p!)^2} \sum_{\substack{\lambda \\ |\lambda| \leq N_c}} \frac{d_\lambda^2}{s_\lambda^{(N_c+q)}} \chi^\lambda(\pi). \quad (4.3.2)$$

Here, λ stands both for an irreducible representation (irrep) of S_p and for the corresponding partition of p .³ The sum over λ is restricted to partitions with length $|\lambda| \leq N_c$. The dimension of the irrep λ is denoted by d_λ , and $\chi^\lambda(\pi)$ is the character of π in irrep λ . The quantity $s_\lambda^{(N_c+q)}$ is defined by

$$s_\lambda^{(N_c+q)} \equiv \prod_{1 \leq i < j \leq N_c+q} \frac{\lambda_i - \lambda_j + j - i}{j - i}, \quad (4.3.3)$$

which is the Schur polynomial evaluated at the $(N_c + q)$ -dimensional point $(1, 1, \dots, 1)$.

The definitions of p and q imply $\max(g, \bar{g}) = qN_c + p$. The sum over $\alpha \in A$ in Eq. (4.3.1) runs over all possible ways to pick q subsets of size N_c each from the set $\{1, \dots, qN_c + p\}$ in such a way that all numbers that have been picked are distinct. There are $\frac{(qN_c + p)!}{p!q!(N_c!)^q}$ ways to do so. The remaining subset has size p . For later use, we order the numbers in each of these $q + 1$ subsets in increasing order, i.e., the subsets become increasing sequences. Therefore, a particular grouping α can be written using integers α_s^c ($1 \leq s \leq q$, $1 \leq c \leq N_c$) as a set of increasing sequences,

$$\alpha \equiv \{(\alpha_1^1, \dots, \alpha_{N_c}^1), \dots, (\alpha_q^1, \dots, \alpha_{N_c}^q)\} \quad \text{with } \alpha_s^c < \alpha_s^{c+1}. \quad (4.3.4a)$$

The remaining increasing sequence is denoted by $\gamma(\alpha)$. It is completely fixed by α and is written using integers γ_s ($1 \leq s \leq p$) as

$$\gamma(\alpha) \equiv (\gamma_1, \dots, \gamma_p) \quad \text{with } \gamma_s < \gamma_{s+1}. \quad (4.3.4b)$$

Recall that the α_s^c and γ_s in Eq. (4.3.4) are all distinct and from $\{1, \dots, qN_c + p\}$.

All color indices in Eq. (4.3.1) are collected in the expressions⁴

$$(I_{\alpha\pi\sigma}^b)^{im} \equiv \begin{cases} \varepsilon_{i_\alpha}^{\otimes q} \delta_{i_\gamma}^{m_\pi} & \text{for } g \geq \bar{g}, \\ \varepsilon_{m_\alpha}^{\otimes q} \delta_{m_\gamma}^{i_\sigma} & \text{for } g < \bar{g}, \end{cases} \quad \text{and} \quad (I_{\alpha\pi\sigma}^e)^{j\ell} \equiv \begin{cases} \varepsilon_{j_\alpha}^{\otimes q} \delta_{j_\gamma}^{\ell_\sigma} & \text{for } g \geq \bar{g}, \\ \varepsilon_{\ell_\alpha}^{\otimes q} \delta_{\ell_\gamma}^{j_\pi} & \text{for } g < \bar{g}, \end{cases} \quad (4.3.5)$$

where b and e stand for the beginning and the end of the link, respectively, and where we omitted the α -dependence of γ for readability. The q -fold Levi-Civita tensor and the Kronecker delta of two tuples are defined as

$$\varepsilon_{i_\alpha}^{\otimes q} \equiv \prod_{s=1}^q \varepsilon_{i_\alpha^1, \dots, i_\alpha^{N_c}} \quad \text{and} \quad \delta_{i_\gamma}^{m_\pi} \equiv \prod_{s=1}^p \delta_{m^{\pi_s}, i^{\gamma_s}}, \quad (4.3.6)$$

respectively. From Eqs. (4.3.5) and (4.3.6) we see that the elements of α , γ , π , and σ correspond to the additional label a on the color indices we defined at the beginning of Sec. 4.2.

³As introduced in Sec. 1.6, a partition λ of the integer p is a sequence $(\lambda_1, \dots, \lambda_k)$ of integers such that $\lambda_i \geq \lambda_{i+1} > 0$ and $\sum_{i=1}^k \lambda_i = p$. The length $|\lambda|$ of the partition is the number of elements in the sequence, i.e., $|\lambda| = k$.

⁴As mentioned after Eq. (4.2.6), the color indices i and m always correspond to the site at the beginning of the link, while j and ℓ correspond to the site at the end of the link.

Note that in Ref. [13] the sums in (4.3.1) are eventually replaced by a sum over so-called decoupling operator indices, which arise from explicitly expressing the character of $\pi \cdot \sigma^{-1} \in S_p$, see (4.3.1) and (4.3.2), in terms of individual matrix elements. The motivation is that the sum over permutations π and σ generates oscillating signs in the Weingarten functions, which leads to a sign problem in the Monte Carlo simulation of their dual formulation [32]. These signs will not be a problem in the tensor formulation, which we develop here, and therefore we keep the sum over permutations.

4.3.2 Range reduction due to Grassmann variables

Recall that the integral (4.3.1) comes from a link between two sites. In the partition function, all color indices in $(I_{\alpha\pi\sigma}^b)^{im}$ and $(I_{\alpha\pi\sigma}^e)^{j\ell}$ that originate from Grassmann hopping terms are contracted with the color indices of the Grassmann variables at these two sites. In this subsection we will show that this observation leads to a simplification of the sum over α , π , and σ . In Ref. [13] a similar reduction was briefly mentioned for the sum over the decoupling operator indices.

We first derive a useful intermediate result. Consider a function f depending on color indices $i^1, \dots, i^k$ that are contracted with Grassmann variables $\chi^{i^1}, \dots, \chi^{i^k}$. Also consider a permutation $\rho \in S_k$. Using the Einstein summation convention, we have

$$f(i^{\rho_1}, \dots, i^{\rho_k}) \prod_{b=1}^k \chi^{i^b} = f(i^{\rho_1}, \dots, i^{\rho_k}) \operatorname{sgn}(\rho) \prod_{b=1}^k \chi^{i^{\rho_b}} = \operatorname{sgn}(\rho) f(i^1, \dots, i^k) \prod_{b=1}^k \chi^{i^b}. \quad (4.3.7)$$

In the first step, we reordered the Grassmann variables, which produces a prefactor given by the sign of the permutation ρ . In the second step, the equality follows from renaming the summation variables.

We will use (4.3.7) with χ replaced by a single color vector of Grassmann variables ψ_x^f or $\bar{\psi}_x^f$, and for color indices in the functions I_b and I_e in (4.3.5) that originate from Grassmann hopping terms. We thus need to identify these color indices in our enumeration scheme. In (4.3.5) we see that for $g \geq \bar{g}$, the entries of α and γ label the indices i and j of U , while the entries of π and σ label the indices ℓ and m of $U^\dagger$. For $g < \bar{g}$ it is the other way around. If we define

$$\varkappa^f \equiv \begin{cases} \kappa^f & \text{for } g \geq \bar{g}, \\ \bar{\kappa}^f & \text{for } g < \bar{g}, \end{cases} \quad \bar{\varkappa}^f \equiv \begin{cases} \bar{\kappa}^f & \text{for } g \geq \bar{g}, \\ \kappa^f & \text{for } g < \bar{g} \end{cases} \quad (4.3.8)$$

with κ^f and $\bar{\kappa}^f$ defined in Eq. (4.2.4a), then for given α , π , and σ , the entries coming from the forward and backward hopping terms of flavor f correspond to

$$\varkappa^f < \alpha_s^c, \gamma_s \leq \varkappa^{f+1} \quad \text{and} \quad \bar{\varkappa}^f < \pi_s, \sigma_s \leq \bar{\varkappa}^{f+1}, \quad (4.3.9)$$

respectively. For given f , there are $\varkappa^{f+1} - \varkappa^f$ entries satisfying the first inequality. The number of entries of γ for which $\varkappa^f < \gamma_s \leq \varkappa^{f+1}$ will be called z^f . This implies that there are $\varkappa^{f+1} - \varkappa^f - z^f$ entries with $\varkappa^f < \alpha_s^c \leq \varkappa^{f+1}$. Note that $\varkappa^{f+1} - \varkappa^f$ equals k^f for $g \geq \bar{g}$ and $\bar{k}^f$ for $g < \bar{g}$, respectively.

It will turn out to be useful to introduce two different equivalence relations, denoted by $\sim$ and $\dot{\sim}$. We define two groupings $\alpha, \alpha' \in A$ to be equivalent (written as $\alpha \sim \alpha'$) if the pairs $\alpha, \gamma(\alpha)$

and $\alpha', \gamma(\alpha')$ can be transformed into each other by permuting, for every flavor f separately, the entries α_s^c and γ_s satisfying $\kappa^f < \alpha_s^c, \gamma_s \leq \kappa^{f+1}$. Furthermore, we define two permutations $\pi, \pi' \in S_p$ to be equivalent (written as $\pi \sim \pi'$) if the corresponding sequences can be transformed into each other by permuting the numbers π_s with $\bar{\kappa}^f < \pi_s \leq \bar{\kappa}^{f+1}$ (for every flavor f separately) and by permuting the numbers s with $\sum_{f'=1}^{f-1} z^{f'} < s \leq \sum_{f'=1}^f z^{f'}$ (for every flavor f separately).⁵ With these two definitions we can use Eq. (4.3.7) to simplify the sum over α, π , and σ . To this end we write

$$\sum_{\alpha \in A} = \sum_{\alpha \in A^\sim} \sum_{\alpha' \in [\alpha]} \quad \text{and} \quad \sum_{\pi \in S_p} = \sum_{\pi \in S_p^\sim} \sum_{\pi' \in [\pi]}, \quad (4.3.10)$$

where $[\cdot]$ denotes an equivalence class and $A^\sim (S_p^\sim)$ denotes a complete set of representatives with respect to the equivalence relation $\sim (\sim)$.⁶

We now insert (4.3.10) into (4.3.1) and use (4.3.7) in the sums over α', π' , and σ' . Taking into account the effect of the Grassmann variables in Eq. (4.2.7), we can replace

$$\mathcal{I}^{g\bar{g}} \rightarrow \sum_{\alpha \in A^\sim} \sum_{\pi, \sigma \in S_p^\sim} \mathcal{L}_{\alpha\pi\sigma}^{g\bar{g}} I_{\alpha\pi\sigma}^b I_{\alpha\pi\sigma}^e. \quad (4.3.11)$$

Equation (4.3.11) becomes an equality after contraction with all Grassmann variables connected to the link in question. The link weight is defined as

$$\mathcal{L}_{\alpha\pi\sigma}^{g\bar{g}} = \left[\prod_{s=1}^{N_c-1} \frac{s!}{(s+q)!} \right] |[\alpha]| \sum_{\pi' \in [\pi]} \sum_{\sigma' \in [\sigma]} \text{sgn}(\pi'^{-1} \cdot \pi) \text{sgn}(\sigma'^{-1} \cdot \sigma) \widetilde{\text{Wg}}_{N_c}^{q,p}(\pi' \cdot \sigma'^{-1}). \quad (4.3.12)$$

Here, the two sign factors are generated by the commutation of Grassmann variables when transforming π' and σ' to their representatives π and σ , respectively. The permutations needed to transform the groupings α' to their representatives α always come in pairs (from a ψ at one end and a $\bar{\psi}$ at the other end of the link), and hence these permutations do not generate sign factors but merely produce a factor of $|[\alpha]|$, which is the cardinality of the equivalence class $[\alpha]$.

What have we gained? The quantity $\mathcal{L}_{\alpha\pi\sigma}^{g\bar{g}}$ is a number that is independent of the color indices. It can easily be precomputed for all desired combinations of $g, \bar{g}, \alpha, \pi$, and σ . The remaining sums over α, π, σ in (4.3.11) are now only over the representatives and thus contain much fewer terms than the sums in (4.3.1). In the following we will use the shorthand notation

$$r \equiv (\alpha, \pi, \sigma) \quad \text{and} \quad \sum_r \equiv \sum_{\alpha \in A^\sim} \sum_{\pi, \sigma \in S_p^\sim}. \quad (4.3.13)$$

5 Creating local contributions

The remaining steps for writing QCD in terms of a tensor network are shown in this section. In Sec. 5.1 we integrate out the original Grassmann variables and introduce integrals over auxiliary

⁵Note that the numbers π_s appear in Eq. (4.3.6) in Kronecker deltas together with γ_s . Since the numbers γ_s are in increasing order, the z^f entries of γ originating from the Grassmann hopping terms of flavor f start at $s = \sum_{f'=1}^{f-1} z^{f'}$. Since we cannot permute the entries of γ , we permute the entries of the sequence π with the same s -indices.

⁶Our concepts of equivalence and equivalence class are closely related, but not identical, to these concepts in group theory. We simply use them to split the sums as in (4.3.10). As in group theory, our classes are disjoint and cover the entire range of the sums.

(color- and flavorless) Grassmann variables. In Sec. 5.2 we eliminate the color degrees of freedom and construct the tensor network.

5.1 Grassmann integration

We now consider the nearest-neighbor Grassmann hopping terms in $\mathcal{G}$, see Eqs. (4.2.7) and (4.2.10), and turn them into purely local contributions, suitable for a tensor network, based on ideas developed for GHOTRG [66, 23]. To this end we introduce a new auxiliary Grassmann variable $\chi_{x,\mu}$ on each link,⁷ with corresponding differential $d\chi_{x,\mu}$, and integrate out the original Grassmann variables. Since these steps are quite technical we perform the corresponding calculations in App. A, where we also Taylor-expand the term $\exp(S_M)$ in the partition function. We obtain

$$\int \left[\prod_x d\psi_x d\bar{\psi}_x \right] \mathcal{G} = \int \prod_{\chi} \prod_x K_x \sigma_x \prod_f \frac{(2m_f)^{w_x^f}}{w_x^f!} E_x^f \Theta(w_x^f) \delta_{h_x^{f,\text{in}}, h_x^{f,\text{out}}}. \quad (5.1.1)$$

Here, the integral on the right-hand side is over all auxiliary Grassmann variables,

$$\int_{\chi} \equiv \prod_{x,\mu} \left(\int \right)^{F_{x,\mu}}, \quad (5.1.2)$$

where the differentials are part of

$$K_x \equiv \prod_{\mu} (\chi_{x,\mu})^{F_{x,\mu}} \prod_{\mu} (d\chi_{x,-\mu})^{F_{x,-\mu}} \quad (5.1.3)$$

with Grassmann parities⁸⁹

$$F_{x,\mu} \equiv \begin{cases} 0 & \text{if } k_{x,\mu} + \bar{k}_{x,\mu} \text{ even,} \\ 1 & \text{if } k_{x,\mu} + \bar{k}_{x,\mu} \text{ odd.} \end{cases} \quad (5.1.4)$$

Furthermore, σ_x is a local sign factor resulting from permutations of Grassmann variables, given in (A.24), (A.9) and (A.20). The quantities

$$h_x^{f,\text{in}} \equiv \sum_{\mu} (k_{x,-\mu}^f + \bar{k}_{x,\mu}^f) \quad \text{and} \quad h_x^{f,\text{out}} \equiv \sum_{\mu} (\bar{k}_{x,-\mu}^f + k_{x,\mu}^f) \quad (5.1.5)$$

are the numbers of incoming and outgoing fermion hopping terms for flavor f at site x . The power w_x^f of m_f in (5.1.1) is called mass occupation number and given by

$$w_x^f = N_c - h_x^{f,\text{in}}. \quad (5.1.6)$$

The Heaviside function and the Kronecker delta in (5.1.1) lead to the constraint $h_x^{f,\text{in}} = h_x^{f,\text{out}} \leq N_c$. For given site and flavor, the integration of the original Grassmann variables produces the factor

$$\begin{aligned} E_x^f = & \varepsilon_{v_{x,1}^{f,1}, \dots, v_{x,w_x^f}^{f,w_x^f}, i_{x,1}^{\kappa_{x,1}^f+1}, \dots, i_{x,1}^{\kappa_{x,1}^f+k_{x,1}^f}, \dots, i_{x,d}^{\kappa_{x,d}^f+1}, \dots, i_{x,d}^{\kappa_{x,d}^f+k_{x,d}^f}, \ell_{x,-1}^{\bar{\kappa}_{x,-1}^f+1}, \dots, \ell_{x,-1}^{\bar{\kappa}_{x,-1}^f+\bar{k}_{x,-1}^f}, \dots, \ell_{x,-d}^{\bar{\kappa}_{x,-d}^f+1}, \dots, \ell_{x,-d}^{\bar{\kappa}_{x,-d}^f+\bar{k}_{x,-d}^f}} \\ & \times \varepsilon_{v_{x,1}^{f,1}, \dots, v_{x,w_x^f}^{f,w_x^f}, m_{x,1}^{\bar{\kappa}_{x,1}^f+1}, \dots, m_{x,1}^{\bar{\kappa}_{x,1}^f+\bar{k}_{x,1}^f}, \dots, m_{x,d}^{\bar{\kappa}_{x,d}^f+1}, \dots, m_{x,d}^{\bar{\kappa}_{x,d}^f+\bar{k}_{x,d}^f}, j_{x,-1}^{\kappa_{x,-1}^f+1}, \dots, j_{x,-1}^{\kappa_{x,-1}^f+k_{x,-1}^f}, \dots, j_{x,-d}^{\kappa_{x,-d}^f+1}, \dots, j_{x,-d}^{\kappa_{x,-d}^f+k_{x,-d}^f}}, \end{aligned} \quad (5.1.7)$$

⁷The boundary conditions for the $\chi_{x,\mu}$ are antiperiodic in time and periodic otherwise.

⁸In Ref. [23] this Grassmann parity is denoted by f instead of F , but we have already used f as a flavor index.

⁹We use the notation $(x, -\mu) \equiv (x - \hat{\mu}, \mu)$, where $\hat{\mu}$ is the unit vector in direction μ .

where all color indices are associated with site x and flavor f . Inside E_x^f , we sum over all color indices $v_x^{f,a}$.

Two remarks are in order. First, the Levi-Civita tensors in E_x^f appear to have many indices. However, because of (5.1.6) and the subsequent constraint, both Levi-Civita tensors have exactly N_c indices each. If $k_{x,\pm\mu}^f$ or $\bar{k}_{x,\pm\mu}^f$ or w_x^f is zero, the corresponding color indices are not present. Second, the constraint $h_x^{f,\text{in}} = h_x^{f,\text{out}}$ can be used to show that $\sum_\mu (F_{x,\mu} + F_{x,-\mu})$ is even. This implies that K_x , although it contains Grassmann variables, is Grassmann even and thus commuting.

5.2 Construction of the tensor network

5.2.1 Elimination of color degrees of freedom

We now consider the full partition function (4.2.7) and insert Eqs. (4.3.11) and (5.1.1) to obtain

$$Z = \sum_{\mathbf{k}, \mathbf{n}} \sum_{\mathbf{r}} \mathcal{W} \bar{\Delta} \int_{\chi} \prod_x K_x \sigma_x \left[\prod_f \frac{(2m_f)^{w_x^f}}{w_x^f!} E_x^f \Theta(w_x^f) \delta_{h_x^{f,\text{in}}, h_x^{f,\text{out}}} \right] \prod_{\mu} I_{x,\mu}^b I_{x,\mu}^e \mathcal{L}_{x,\mu}, \quad (5.2.1)$$

where we defined

$$I_{x,\mu}^b \equiv (I_{r_{x,\mu}}^b)^{i_{x,\mu} \mathbf{m}_{x,\mu}}, \quad I_{x,\mu}^e \equiv (I_{r_{x,\mu}}^e)^{j_{x,\mu} \ell_{x,\mu}}, \quad \mathcal{L}_{x,\mu} \equiv \mathcal{L}_{r_{x,\mu}}^{g_{x,\mu} \bar{g}_{x,\mu}} \quad (5.2.2)$$

for simplicity and introduced the field $\mathbf{r} \equiv (r_{x,\mu} | \forall x, \mu)$, again with periodic boundary conditions.

Let us briefly discuss this intermediate result. First, for every link, the gauge integral results in a sum over r (i.e., over representatives α , π , and σ , see Eq. (4.3.13)) for this link. The range of r for a given link depends on k , $\bar{k}$, and $\vec{n}$ for this link, where $\vec{n}$ contains all occupation numbers of plaquettes which include that link. If the $\text{SU}(N_c)$ condition $|g - \bar{g}|/N_c \in \mathbb{N}_0$ is not satisfied, the range of r is zero and the tuple $(k, \bar{k}, \vec{n})$ does not contribute to the partition function. Second, the integral over the original Grassmann variables has been replaced by an integral over the auxiliary Grassmann variables, which are color- and flavorless. The original Grassmann variables were coupled to neighbors by hopping terms. In contrast, the auxiliary Grassmann variables, which live on the links, are decoupled. This is the main achievement of Sec. 5.1.

A given term in the sum over $\mathbf{k}$, $\mathbf{n}$, and $\mathbf{r}$ in (5.2.1) contains factors from all sites, links, and plaquettes of the lattice. To construct a network of local tensors we need to gather the factors corresponding to a single site and contract their color indices. This can be done for each lattice site separately, as we show in the following.

We first rewrite $\bar{\Delta}$ in (4.2.11) in terms of local objects. To this end, we shift the first factor in (4.2.6) from site x to site $x - \hat{\mu}$, the second factor from site x to site $x - \hat{\mu} - \hat{\nu}$, and the third

factor from site x to site $x - \hat{\nu}$, respectively, to obtain¹⁰

$$\bar{\Delta} = \prod_x \Delta_x \quad \text{with} \quad \Delta_x \equiv \prod_{\mu \neq \nu} \left(\prod_{a=1}^{n_{x-\hat{\mu},\mu\nu}} \delta_{j_{x,-\mu}^{a+\omega_{x,-\mu}^+}, i_{x,\nu}^{a+\omega_{x,\nu}^-}} \right) \left(\prod_{a=1}^{n_{x-\hat{\mu}-\hat{\nu},\mu\nu}} \delta_{j_{x,-\nu}^{a+\omega_{x,-\nu}^-}, \ell_{x,-\mu}^{a+\bar{\omega}_{x,-\mu}^-}} \right) \\ \times \left(\prod_{a=1}^{n_{x-\hat{\nu},\mu\nu}} \delta_{m_{x,\mu}^{a+\bar{\omega}_{x,\mu}^-}, \ell_{x,-\nu}^{a+\bar{\omega}_{x,-\nu}^+}} \right) \left(\prod_{a=1}^{n_{x,\mu\nu}} \delta_{m_{x,\nu}^{a+\bar{\omega}_{x,\nu}^+}, i_{x,\mu}^{a+\omega_{x,\mu}^+}} \right). \quad (5.2.3)$$

This expression depends on all color indices associated with site x that originate from the plaquette action. Analogously, the contributions depending on the color indices in the gauge-link integrals in (5.2.1) can be rewritten locally by shifting the sites in $I_{x,\mu}^e$ from x to $x - \hat{\mu}$,

$$I_x \equiv \prod_{\mu} I_{x,\mu}^b I_{x,-\mu}^e, \quad (5.2.4)$$

so that

$$\prod_x \prod_{\mu} I_{x,\mu}^b I_{x,\mu}^e = \prod_x I_x. \quad (5.2.5)$$

The local expression I_x depends on all color indices (originating from hopping terms and the plaquette action) that are associated with site x .¹¹ We now define the local color factor

$$C_x = \Delta_x I_x \prod_f E_x^f, \quad (5.2.6)$$

which only depends on the variables associated with site x ,

$$\mathbf{k}_x = (k_{x,-\mu}^f, k_{x,\mu}^f, \bar{k}_{x,-\mu}^f, \bar{k}_{x,\mu}^f \mid \forall \mu, f), \quad (5.2.7a)$$

$$\mathbf{n}_x = (n_{x,\mu\nu}, n_{x-\hat{\mu},\mu\nu}, n_{x-\hat{\nu},\mu\nu}, n_{x-\hat{\mu}-\hat{\nu},\mu\nu} \mid \forall \mu \neq \nu), \quad (5.2.7b)$$

$$\mathbf{r}_x = (r_{x,-\mu}, r_{x,\mu} \mid \forall \mu). \quad (5.2.7c)$$

Every color index corresponding to site x appears exactly twice on the right-hand side of Eq. (5.2.6). We computed C_x analytically by explicitly summing over the color indices using `sympy` [72]. The partition function (5.2.1) then becomes

$$Z = \sum_{\mathbf{k}, \mathbf{n}} \sum_{\mathbf{r}} \mathcal{W} \int_{\chi} \prod_x K_x \sigma_x C_x \left[\prod_f \frac{(2m_f)^{w_x^f}}{w_x^f!} \Theta(w_x^f) \delta_{h_x^{f,\text{in}}, h_x^{f,\text{out}}} \right] \prod_{\mu} \mathcal{L}_{x,\mu}. \quad (5.2.8)$$

5.2.2 Tensor formulation

In the partition function we have several summation variables, where $k_{x,\mu}^f$ and $\bar{k}_{x,\mu}^f$ are link occupation numbers, $n_{x,\mu\nu}$ is a plaquette occupation number, and $r_{x,\mu}$ is a link variable whose range depends on k , $\bar{k}$, and n . We rewrite the partition function (5.2.8) as a tensor network

$$Z = \sum_{\mathbf{k}, \mathbf{n}} \sum_{\mathbf{r}} \int_{\chi} \prod_x \bar{T}_x K_x \quad (5.2.9)$$

¹⁰Note that each Kronecker delta just enforces the equality of any two color indices at a plaquette corner.

¹¹A similar decoupling is obtained in Ref. [13] using so-called decoupling operators.

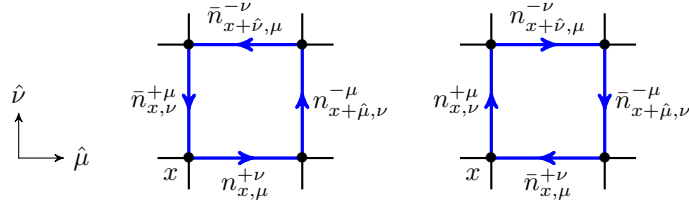


Figure 5.1: Introducing edge variables for plaquettes $U_{x,\mu\nu}^P$ (left) and $U_{x,\nu\mu}^P$ (right).

with the local numerical tensor

$$\bar{T}_x \equiv \sigma_x C_x W_x \left[\prod_f \frac{(2m_f)^{w_x^f}}{w_x^f!} \Theta(w_x^f) \delta_{h_x^{f,\text{in}}, h_x^{f,\text{out}}} \right] \left[\prod_\mu \text{sgn}(\mathcal{L}_{x,\mu}) \right] \prod_{\mu=\pm 1}^{\pm d} \sqrt{|\mathcal{L}_{x,\mu}|} \quad (5.2.10)$$

whose tensor indices are $\mathbf{k}_x$, $\mathbf{n}_x$, and $\mathbf{r}_x$. To obtain this tensor-network structure we also factorized the weight $\mathcal{W}$ of (4.2.12) into local contributions by writing $\mathcal{W} = \prod_x W_x$ with

$$W_x = \left[\prod_{\mu \neq \nu} \frac{\left(\frac{\beta}{2N_c} \right)^{n_{x,\mu\nu} + n_{x-\hat{\mu},\mu\nu} + n_{x-\hat{\mu}-\hat{\nu},\mu\nu} + n_{x-\hat{\nu},\mu\nu}}}{n_{x,\mu\nu}! n_{x-\hat{\mu},\mu\nu}! n_{x-\hat{\mu}-\hat{\nu},\mu\nu}! n_{x-\hat{\nu},\mu\nu}!} \right]^{\frac{1}{4}} \left[\prod_\mu \eta_{x,\mu}^{F_{x,\mu}} \right] \left[\prod_{\mu=\pm 1}^{\pm d} \prod_f \sqrt{\frac{e^{\mu_f(k_{x,\mu}^f - \bar{k}_{x,\mu}^f) \delta_{|\mu|,1}}}{k_{x,\mu}^f! \bar{k}_{x,\mu}^f!}} \right], \quad (5.2.11)$$

where all numerical quantities on a link were evenly distributed via square roots to both endpoints, except for the staggered phases and the sign of $\mathcal{L}_{x,\mu}$ to avoid complex entries, and all numerical quantities on a plaquette were distributed via fourth roots to its four corners. Similarly, the link weight $|\mathcal{L}_{x,\mu}|$ was distributed via square roots to the two tensors at the endpoints of the link. Recall that μ_f in the last factor of (5.2.11) is the chemical potential for flavor f .

In standard tensor-renormalization-group methods, two tensors at neighboring sites are contracted by summing over the shared link index on the link connecting the two sites. However, the sum in (5.2.9) involves plaquette occupation numbers n , which are shared by four tensors. A summation over plaquette occupation numbers therefore does not correspond to a contraction of two local tensors T_x . A similar argument applies to r , whose range is determined by n . To be able to still apply tensor-renormalization-group methods, one option would be to develop a more general blocking procedure for a tensor network containing indices that are shared by more than two tensors. While this would indeed be an interesting option to explore in future work, here we follow a more traditional route. We convert the plaquette occupation numbers $n_{x,\mu\nu}$ to link variables and use the GHOTRG method developed for the infinite-coupling case [23]. This is implemented by introducing a link variable for each edge of a plaquette, called “edge variable” in the following, and requiring that all four edge variables around a plaquette take on the same value for terms that contribute to the partition function. For the plaquette $U_{x,\mu\nu}^P$ we introduce $n_{x,\mu}^{+\nu}$, $n_{x+\hat{\mu},\nu}^{-\mu}$, $\bar{n}_{x+\hat{\nu},\mu}^{-\nu}$, and $\bar{n}_{x,\mu}^{+\nu}$, as illustrated in Fig. 5.1. The superscript $\pm\mu$ indicates that the plaquette extends in the positive/negative μ -direction perpendicular to the link under consideration, while n and $\bar{n}$ correspond to U and $U^\dagger$, respectively.

Formally, we rewrite the sum of a generic function f over a plaquette occupation number as

$$\begin{aligned} \sum_{n_{x,\mu\nu}} f(n_{x,\mu\nu}) &= \sum_{n_{x,\mu\nu}} f(n_{x,\mu\nu}) \sum_{n_{x,\mu}^{+\nu}} \delta_{n_{x,\mu}^{+\nu}, n_{x,\mu\nu}} \sum_{n_{x+\hat{\mu},\nu}^{-\mu}} \delta_{n_{x+\hat{\mu},\nu}^{-\mu}, n_{x,\mu\nu}} \sum_{\bar{n}_{x+\hat{\nu},\mu}^{-\nu}} \delta_{\bar{n}_{x+\hat{\nu},\mu}^{-\nu}, n_{x,\mu\nu}} \sum_{\bar{n}_{x,\nu}^{+\mu}} \delta_{\bar{n}_{x,\nu}^{+\mu}, n_{x,\mu\nu}} \\ &= \sum_{n_{x,\mu}^{+\nu}, n_{x+\hat{\mu},\nu}^{-\mu}, \bar{n}_{x+\hat{\nu},\mu}^{-\nu}, \bar{n}_{x,\nu}^{+\mu}} \delta_{n_{x,\mu}^{+\nu}, n_{x+\hat{\mu},\nu}^{-\mu}} \delta_{n_{x+\hat{\mu},\nu}^{-\mu}, \bar{n}_{x+\hat{\nu},\mu}^{-\nu}} \delta_{\bar{n}_{x+\hat{\nu},\mu}^{-\nu}, \bar{n}_{x,\nu}^{+\mu}} \delta_{\bar{n}_{x,\nu}^{+\mu}, n_{x,\mu}^{+\nu}} f(\text{edge variable}), \end{aligned} \quad (5.2.12)$$

where we have eliminated the original sum over $n_{x,\mu\nu}$ in the second step. Although one of the Kronecker deltas in the second line is redundant, it is useful for the tensor formulation to use this symmetrized form. Effectively, each Kronecker delta imposes the equality of the edge variables around a plaquette corner. Contributions to the partition function that depend on $n_{x,\mu\nu}$ will only be nonzero if all four edge variables are equal, as they represent the same plaquette occupation number. Therefore, every occurrence of $n_{x,\mu\nu}$ can be replaced by any of the edge variables $n_{x,\mu}^{+\nu}$, $n_{x+\hat{\mu},\nu}^{-\mu}$, $\bar{n}_{x+\hat{\nu},\mu}^{-\nu}$, or $\bar{n}_{x,\nu}^{+\mu}$. We gather all these single-plaquette constraints on the lattice in

$$\Delta^P \equiv \prod_{x, \mu \neq \nu} \delta_{n_{x,\mu}^{+\nu}, n_{x+\hat{\mu},\nu}^{-\mu}} \delta_{n_{x+\hat{\mu},\nu}^{-\mu}, \bar{n}_{x+\hat{\nu},\mu}^{-\nu}} \delta_{\bar{n}_{x+\hat{\nu},\mu}^{-\nu}, \bar{n}_{x,\nu}^{+\mu}} \delta_{\bar{n}_{x,\nu}^{+\mu}, n_{x,\mu}^{+\nu}}. \quad (5.2.13)$$

This can be rewritten in terms of local objects by separately shifting the site in each Kronecker delta, which results in

$$\Delta^P = \prod_x \Delta_x^P \quad \text{with} \quad \Delta_x^P = \prod_{\mu \neq \nu} \delta_{n_{x,-\mu}^{+\nu}, n_{x,\nu}^{-\mu}} \delta_{n_{x,-\nu}^{-\mu}, \bar{n}_{x,-\mu}^{-\nu}} \delta_{\bar{n}_{x,\mu}^{-\nu}, \bar{n}_{x,-\nu}^{+\mu}} \delta_{\bar{n}_{x,\nu}^{+\mu}, n_{x,\mu}^{+\nu}}, \quad (5.2.14)$$

where Δ_x^P depends on all edge variables associated with site x . We thus have four new variables $n_{x,\mu}^{+\nu}$, $\bar{n}_{x,\mu}^{+\nu}$, $\bar{n}_{x,\mu}^{-\nu}$, $n_{x,\mu}^{-\nu}$ on every link (x, μ) of the lattice for each plane $\mu\nu$, see Fig. 5.2. The range of $r_{x,\mu}$ depends on the plaquette occupation numbers through $g_{x,\mu}$ and $\bar{g}_{x,\mu}$ defined in (4.2.3). Because of the Kronecker deltas in (5.2.12) the plaquette occupation numbers in (4.2.3) can be replaced by edge variables defined on the link (x, μ) . Thus $r_{x,\mu}$ now only depends on link variables.

Now that all variables are defined on the links we can combine the various link variables in a compound link index $j_{x,\mu} = ((k_{x,\mu}^f, \bar{k}_{x,\mu}^f | \forall f), (n_{x,\mu}^{\pm\nu}, \bar{n}_{x,\mu}^{\pm\nu} | \forall \nu \neq \mu), r_{x,\mu})$ and rewrite the partition function (5.2.9) as

$$Z = \sum_j \int_{\chi} \prod_x T_x K_x, \quad (5.2.15)$$

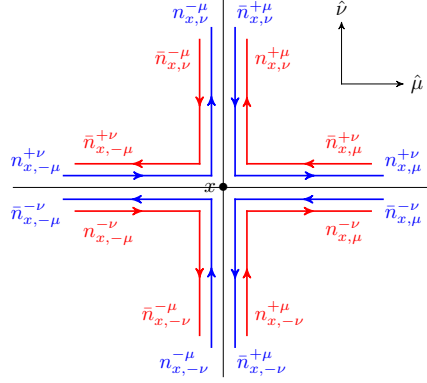
where $j \equiv (j_{x,\mu} | \forall x, \mu)$. The numerical tensor is given by

$$T_x \equiv \Delta_x^P \bar{T}_x(n_x \rightarrow \text{edge variables}) \quad (5.2.16)$$

and depends on $j_{x,\pm\mu}$ ($\mu = 1, \dots, d$). Recall that $\bar{T}_x$, defined in (5.2.10), depends on the plaquette occupation numbers in $\mathbf{n}_x$, see Eq. (5.2.7). We have a choice of how to replace them by edge variables defined on links connected to x , as described after (5.2.12), but the result does not depend on what choice we make.

5.2.3 Restriction of the edge variables

To contract the complete tensor network and compute the partition function and thermodynamical observables using the GHOTRG [23] (or another) method, the bond dimension of the

Figure 5.2: Enumeration of the edge variables associated with site x .

initial tensor has to be finite. While k and $\bar{k}$ are limited by the number of colors N_c , and r is finite as well for a given gauge integral, the plaquette occupation numbers n run from zero to infinity in the original expansion. In the tensor formulation, the plaquette occupation numbers have been turned into edge variables, which also run from zero to infinity.

To make the range of the tensor indices $j_{x,\mu}$ finite, we need to restrict the edge variables. More precisely, we introduce a maximal order $n_{\max}$ and remove, for a given link (x, μ) , all combinations of edge variables that do not satisfy the link criterion

$$\sum_{\substack{\nu=\pm 1 \\ |\nu| \neq \mu}}^{\pm d} (n_{x,\mu}^\nu + \bar{n}_{x,\mu}^\nu) \leq n_{\max}. \quad (5.2.17)$$

With this restriction, all configurations of order β^n with $0 \leq n \leq n_{\max}$ are included in the calculation, while higher orders will be incomplete. Without any TRG truncations, the partition function is then exact up to $\mathcal{O}(\beta^{n_{\max}+1})$ corrections.

Note that we can substantially speed up the tensor construction by setting to zero the entries of the initial tensor for which¹²

$$\frac{1}{2} \sum_{\mu=\pm 1}^{\pm d} \sum_{\substack{\nu=\pm 1 \\ |\nu| \neq |\mu|}}^{\pm d} (n_{x,\mu}^\nu + \bar{n}_{x,\mu}^\nu) > n_{\max}, \quad (5.2.18)$$

which correspond to tensor entries, and hence configurations, of order β^n with $n > n_{\max}$. The condition (5.2.18) is referred to as site criterion. Although this step does not reduce the size of the initial tensor, it also increases the efficiency of the tensor-network method by removing singular values that would correspond to $n > n_{\max}$ and could be large.¹³

In App. B, we explicitly show that the obtained tensor network reduces to the one stated in Ref. [23] in the limit $\beta = 0$ (or equivalently $n_{\max} = 0$), $N_f = 1$, $N_c = 3$, and $d = 2$, which is therefore a consistency check.

¹²The factor of $1/2$ arises since two edge variables correspond to the same plaquette occupation number, see Fig. 5.2.

¹³After choosing $n_{\max}$, the bond dimension of the initial tensor is fixed. However, for an efficient numerical computation, especially for higher orders in β , the initial tensor may additionally be truncated to a particular bond dimension D , using an HOSVD-inspired truncation procedure analogous to that of the GHOTRG method.

6 Translational invariance for GHOTRG

For $n_{\max} \geq 1$ we can increase the numerical efficiency by employing translational invariance on the initial lattice. Configurations that only differ by a translation on the lattice make the same contribution to the partition function. Therefore, we would like to include only a subset of these configurations explicitly and apply a combinatorial factor to also take into account the others.¹⁴

In the following, we present a simple approach that is applicable for arbitrary $n_{\max}$ in two dimensions. The generalization of this method for the OS-(G)HOTRG approach introduced in Sec. 7 is straight-forward, however there is a more sophisticated approach possible in this case (see Sec. 8), which further reduces the number of considered configurations and is typically used instead.

To apply translational invariance for GHOTRG, we first observe that for any valid configuration that contributes to Z_n , there is at most one plaquette at the lattice with an unoriented occupation $n_{x,\mu\nu} + n_{x,\nu\mu} > \lfloor n_{\max}/2 \rfloor$. We want to position this plaquette at the plaquette at the origin (red plaquette in Fig. 8.5), and apply a combinatorial factor V to these configurations, to compensate this. As a consequence, all other plaquettes, which we call restricted plaquettes (black plaquettes), can be restricted to

$$n_{x,\mu\nu} + n_{x,\nu\mu} \leq \lfloor n_{\max}/2 \rfloor. \quad (6.1)$$

Formally, we introduce four initial “special” tensors $S_1, \dots, S_4$ at the corners of the plaquette at the origin. They are built by the corresponding initial tensors T_x by employing the restriction (6.1) to all adjacent restricted plaquettes¹⁵ and multiplying a combinatorial factor $V^{1/4}$ to all tensor entries involving a origin plaquette with unoriented occupation greater than $\lfloor n_{\max}/2 \rfloor$. The tensors on all other lattice sites, which we call “restricted” tensors R , are built by the corresponding initial tensors T_x by employing the restriction (6.1) for all adjacent plaquettes.¹⁶

The blocking procedure now proceeds as shown in Fig. 6.3. Denoting a blocking step in direction μ by $\star_\mu$, the steps can be summarized as

$$\begin{aligned} \text{Step 1:} \quad & R^{(0)} \star_1 R^{(0)} \rightarrow R^{(1)}, \quad S_1^{(0)} \star_1 S_2^{(0)} \rightarrow S_1^{(1)}, \quad S_3^{(0)} \star_1 S_4^{(0)} \rightarrow S_2^{(1)}, \\ \text{Step 2:} \quad & R^{(1)} \star_2 R^{(1)} \rightarrow R^{(2)}, \quad S_1^{(1)} \star_2 S_2^{(1)} \rightarrow S^{(2)}, \\ \text{Step } i+1: \quad & R^{(i)} \star_\mu R^{(i)} \rightarrow R^{(i+1)}, \quad S^{(i)} \star_\mu R^{(i)} \rightarrow S^{(i+1)} \quad \text{for } i \geq 2, \end{aligned} \quad (6.2)$$

where the last line is repeated in all directions until we arrive at a single point.¹⁷

After every blocking step we also apply a rescaling. For simplicity, we do not introduce a new name for the rescaled tensors in the summary (6.2). After the first blocking step, the scaling

¹⁴Translational invariance implies that the contributions of two valid configurations $\mathbf{j}$ and $\mathbf{j}'$ with $j'_x = j_{x-y}$ are equal in (5.2.15), where y is an arbitrary shift on the toroidal lattice. An explicit analysis, taking into account the staggered phases and boundary conditions, confirms that this is indeed the case [26].

¹⁵To do so, the restriction (6.1) is rewritten in terms of the corresponding edge variables.

¹⁶Both the tensors S_i and R have reduced dimensions as some or all of the surrounding plaquettes are restricted.

¹⁷Because of the contractions between S and R we have to revisit the truncation procedure. If we were to apply the truncation procedure of Sec. 3.3.2, the dependence of the tensor on the location x would no longer disappear, which would destroy the $\log V$ scaling of the tensor-network method. We therefore construct the truncation matrices only from the tensors R , except for the truncation of $S_1^{(1)}$ and $S_2^{(1)}$.

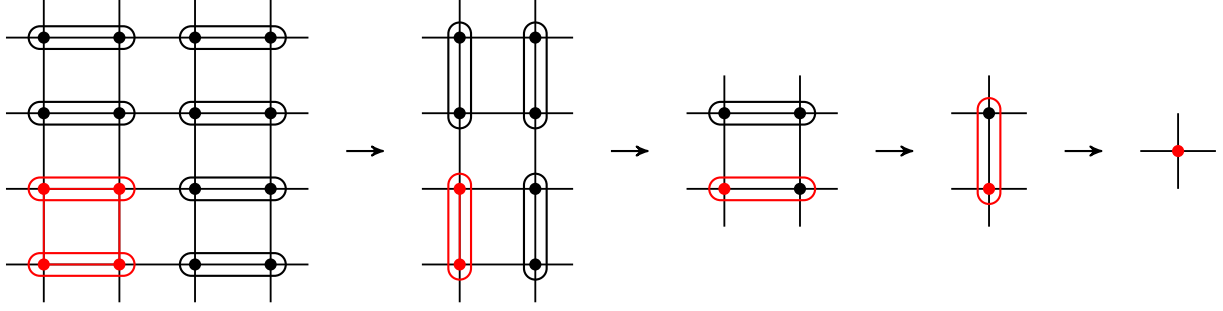


Figure 6.3: Illustration of the blocking procedure using translational invariance on the initial lattice. The plaquette at the origin and the S tensors are indicated in red, while the restricted plaquettes and the tensors R are indicated in black.

factor t_1 is defined as the maximal absolute value all tensors $R^{(1)}$, $S_1^{(1)}$, and $S_2^{(1)}$. After the i -th blocking step, for $i \geq 2$, t_i is the maximal absolute value of the obtained tensors $R^{(i)}$ and $S^{(i)}$. The computation of the partition function using Eq. (3.3.32) remains unchanged.

IV. The OS-GHOTRG method

Using the initial tensor obtained in Part III, one can in principle apply the standard GHOTRG algorithm described in Sec. 3.3. However, even though the initial tensor is truncated to order $n_{\max}$ and we additionally use the restrictions (5.2.17) and (5.2.18), the completely blocked tensor network will contain many incomplete higher-order contributions ($n > n_{\max}$) for larger lattices. Numerical calculations show that this can significantly restrict the range in β for which the result remains close to Monte Carlo data that are not based on a strong-coupling expansion (see Fig. 10.5). By applying the GHOTRG algorithm, we obtain Z (and thus also $\ln Z$ and the thermodynamical observables) as a numerical function of β , and not as an explicit expansion in β . One can try to determine the expansion coefficients Z_n by fitting tensor data to polynomials in β , but this attempt does not work well for larger lattices since the tensor data have limited precision for realistic bond dimensions. Therefore, we developed the OS-GHOTRG algorithm, which is introduced in this chapter and published in [28]. This approach allows for an explicit calculation of the coefficients Z_n using a modified GHOTRG algorithm. The main idea is that each tensor entry is written as a power series in β during the iterative procedure. This allows us to remove unwanted contributions of higher order. Using the expansion of the partition function we can also expand the free energy, i.e., the logarithm of the partition function, and hence the thermodynamical observables, order by order also up to $\beta^{n_{\max}}$. At the end of this part we discuss how the translational invariance on the initial lattice can be used to increase the numerical efficiency of the OS-GHOTRG approach for two dimensions, which is published in [28] as well.

7 Order separation in β

7.1 Overview

In the following, we give a qualitative outline of the OS-GHOTRG approach. As mentioned before, we aim to determine the expansion coefficients of Z in β up to a given order $n_{\max}$. However, the order of a valid configuration $\mathbf{j}^1$ involves the whole lattice and cannot be determined by the local tensor alone. Therefore a crucial property of OS-GHOTRG is that it removes higher-order contributions on the fly. In contrast to the GHOTRG approach, we now divide a blocking step into three substeps: contraction, reduction, and truncation. After the contraction step, we perform a reduction step which removes invalid configurations and contributions of $\mathcal{O}(\beta^{n_{\max}+1})$ from the partition function. This reduces the bond dimension of the fat-link index and sets some tensor entries to zero. Afterwards we perform a truncation step in which the bond dimension of the fat links is cut to a desired value. The truncation step is designed to preserve the structure of the tensor as a power series in β . This completes a blocking step.

¹We use the term “valid configuration” for a configuration $\mathbf{j}$ that makes a nonzero contribution to the partition function.

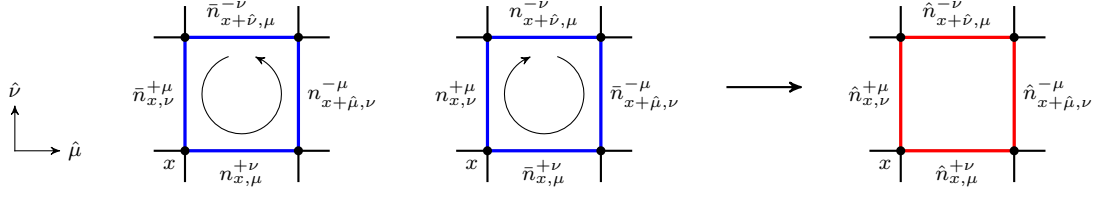


Figure 7.1: Blue: Edge variables n and $\bar{n}$ for the oriented plaquettes starting at site x in the $\mu\nu$ and $\nu\mu$ directions, respectively. The subscript indicates the link, while the superscript indicates the direction (including sign) in which the plaquette extends perpendicular to the link. For a valid configuration the four edge variables around an oriented plaquette must be equal. Red: Edge occupations of the unoriented plaquette as defined in (7.2.1), which must also be equal for a valid configuration.

7.2 Introduction of edge occupations

For the OS-GHOTRG approach we first introduce the edge occupation numbers (or occupations in short)

$$\hat{n}_{x,\mu}^\nu(j_{x,\mu}) \equiv n_{x,\mu}^\nu + \bar{n}_{x,\mu}^\nu \quad (7.2.1)$$

as the sum of edge variables of two superimposed plaquettes with opposite orientations, see Fig. 7.1. The edge occupations therefore refer to unoriented plaquettes. We now reformulate some of the results from Sec. 5.2.2 and Sec. 5.2.3 in terms of the edge occupations. First, the β -dependence of a tensor entry is given by

$$(T_x)_{j_x} \propto \beta^{n_x} \quad \text{with} \quad n_x(j_x) \equiv \frac{1}{8} \sum_{\substack{\mu,\nu=\pm 1 \\ |\mu| \neq |\nu|}}^{\pm d} \hat{n}_{x,\mu}^\nu. \quad (7.2.2)$$

Further, the condition (5.2.14) implies a weaker condition for the edge occupations to obtain nonzero entries of T_x ,

$$(T_x)_{j_x} \propto \Delta_x \quad \text{with} \quad \Delta_x(j_x) \equiv \prod_{\substack{\mu,\nu=\pm 1 \\ |\mu| < |\nu|}}^{\pm d} \delta_{\hat{n}_{x,\mu}^\nu, \hat{n}_{x,\nu}^\mu}, \quad (7.2.3)$$

which is illustrated in Fig. 7.2. We will refer to the Kronecker deltas in (7.2.3) as hook conditions in the following. Finally, the link criterion (5.2.17) now reads

$$\sum_{\substack{\nu=\pm 1 \\ |\nu| \neq |\mu|}}^{\pm d} \hat{n}_{x,\mu}^\nu(j_{x,\mu}) \leq n_{\max}, \quad (7.2.4)$$

while the site criterion (5.2.18) becomes

$$n_x \leq \frac{1}{4} n_{\max}. \quad (7.2.5)$$

A detailed derivation of more general versions of the link and site criteria, which are valid at every step of the blocking procedure, will be given in Sec. 7.4.

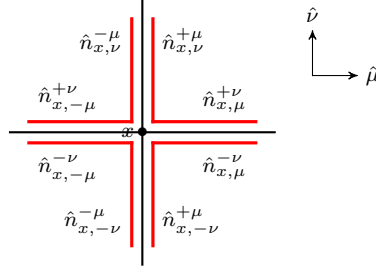


Figure 7.2: For valid configurations we require the two edge occupations of the same unoriented plaquette to be equal and call this requirement the hook condition. The quantity Δ_x in (7.2.3) is a product of hook conditions.

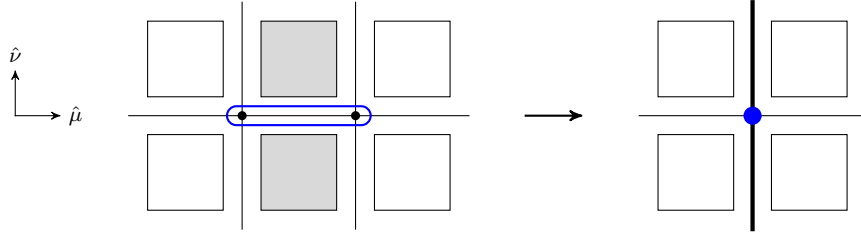


Figure 7.3: Illustration of a contraction step. On the fine grid (left), the tensors at the sites x and $y = x + \hat{\mu}$ are contracted, resulting in a new tensor at the site X on the coarse grid (right). Every plaquette that has the contracted link as an edge is collapsed onto the corresponding fat link perpendicular to $\hat{\mu}$.

7.3 Tensor contraction and decomposition

We now perform a contraction step, in which we contract two adjacent tensors on the link between x and $y = x + \hat{\mu}$. For the initial tensor, the β dependence is given by (7.2.2) and comes exclusively from the edge occupations. To understand what happens in this step we consider all plaquettes that have x or y as a corner and divide them into three groups: (1) plaquettes that have an edge along the contracted link and thus have both x and y as corners (indicated in gray in Fig. 7.3), (2) plaquettes in planes perpendicular to the contraction direction (not shown in Fig. 7.3), and (3) plaquettes that have either x or y as a corner and whose plane includes the contraction direction (indicated in white in Fig. 7.3). Every plaquette in group (1) is collapsed onto the corresponding perpendicular fat link generated in the contraction, see Fig. 7.3. We shall see that the edge occupations of the collapsed plaquettes will be turned into a new link occupation $\ell_{X,\nu} \equiv \frac{1}{2}(\hat{n}_{x,\nu}^{\mu} + \hat{n}_{y,\nu}^{\mu})$ for every fat link (X, ν) perpendicular to $\hat{\mu}$, where positive or negative ν correspond to the upper and lower gray plaquette in Fig. 7.3, respectively. Plaquettes in group (2), say in the plane $(\nu, \lambda) \perp \hat{\mu}$, are merged, and the edge occupations on the coarse grid are given by (7.3.8c) below. For plaquettes in group (3) the edge occupations on the coarse grid are the same as on the fine grid.

In subsequent blocking steps the link occupations ℓ can increase further as more plaquettes in group (1) are collapsed onto fat links. Since links that are merged into fat links (perpendicular to the contraction direction) may already contain collapsed plaquettes, the formula for $\ell_{X,\nu}$ generalizes to (7.3.8b) below.

If the contracted link (x, μ) already contains a collapsed plaquette from a previous blocking step (i.e., $\ell_{x,\mu} \neq 0$), the contraction turns $\ell_{x,\mu}$ into a site occupation $n_X^s \equiv \ell_{x,\mu}$ of the entry of the coarse-grid tensor. If the sites x and y at the ends of the contracted link are already occupied through previous blocking steps, we obtain the more general formula (7.3.8a) below.

In the following (including Secs. 7.4 and 7.5) we will show that, at every blocking step of the OS-GHOTRG method, the numerical tensor T_x can be decomposed as

$$(T_x)_{j_x} = \beta^{n_x(j_x)} \sum_{n_x^s=0}^{n_{\max}} \beta^{n_x^s} (T_x^{n_x^s})_{j_x}, \quad (7.3.1)$$

where the coefficient tensors $T_x^{n_x^s}$ are labeled by site occupations n_x^s , have the same bond dimensions as the full tensor T_x , are independent of β , and satisfy

$$(T_x^{n_x^s})_{j_x} \propto \Delta_x(j_x) \quad \text{with } \Delta_x \text{ given in (7.2.3)}. \quad (7.3.2)$$

The exponent n_x of β in (7.3.1) is given by²

$$n_x(j_x) \equiv \frac{1}{2} \sum_{\nu=\pm 1}^{\pm d} \ell_{x,\nu}(j_{x,\nu}) + \frac{1}{8} \sum_{\substack{\nu,\lambda=\pm 1 \\ |\nu| \neq |\lambda|}}^{\pm d} \hat{n}_{x,\nu}^\lambda(j_{x,\nu}), \quad (7.3.3)$$

which generalizes the definition in (7.2.2) to an arbitrary blocking step. The link and edge occupations $\ell_{x,\mu}$ and $\hat{n}_{x,\mu}^\nu$ are determined by the tensor index $j_{x,\mu}$, and we have $\ell_{x,\mu} \leq n_{\max}$.

The initial tensor (7.2.2) can be decomposed as in (7.3.1) with

$$(T_x^0)_{j_x} \equiv \beta^{-n_x(j_x)} (T_x)_{j_x} \quad \text{and} \quad (T_x^{n_x^s})_{j_x} \equiv 0 \quad \text{for } n_x^s \geq 1, \quad (7.3.4)$$

where (7.3.3) holds with $\ell_{x,\mu} = 0$. From (7.3.4) and (7.2.3) it follows immediately that (7.3.2) holds for the initial tensor. Furthermore, for the initial tensor $\ell_{x,\mu}$ and $\hat{n}_{x,\mu}^\nu$ are determined by $j_{x,\mu}$, the former trivially since it is zero, and the latter since the $\hat{n}_{x,\mu}^\nu$ are defined in terms of the edge variables, which in turn are part of $j_{x,\mu}$.

Now consider the contraction of two adjacent tensors on sites x and $y = x + \hat{\mu}$ in the μ -direction. The contraction of the local Grassmann tensors K_x and K_y into a coarse-grid Grassmann tensor K_X using the GHOTRG procedure is described in detail in Sec. 3.3. The resulting sign factor σ_{F_x, F_y} is given in Eq. (3.3.18) and the numerical tensor on the coarse grid, including the sign factor, is obtained as

$$(T_X)_{j_X} \equiv \sum_{j_{x,\mu}} \sigma_{F_x, F_y} (T_x)_{j_x} (T_y)_{j_y}, \quad (7.3.5)$$

where the contraction is performed by summing over the shared index $j_{x,\mu} = j_{y,-\mu}$. The coarse-grid tensor T_X has fat indices

$$j_{X,\nu} \equiv (j_{x,\nu}, j_{y,\nu}) \quad \text{for } |\nu| \neq \mu \quad (7.3.6)$$

defined on the fat links (perpendicular to the contraction direction), while the external indices in the direction of contraction are $j_{X,-\mu} \equiv j_{x,-\mu}$ and $j_{X,\mu} \equiv j_{y,\mu}$. All these indices are combined in j_X as usual.

²For a valid configuration, which satisfies the hook conditions, n_x may not be an integer, but $\sum_x n_x$, which appears in the exponent of β in Z , is.

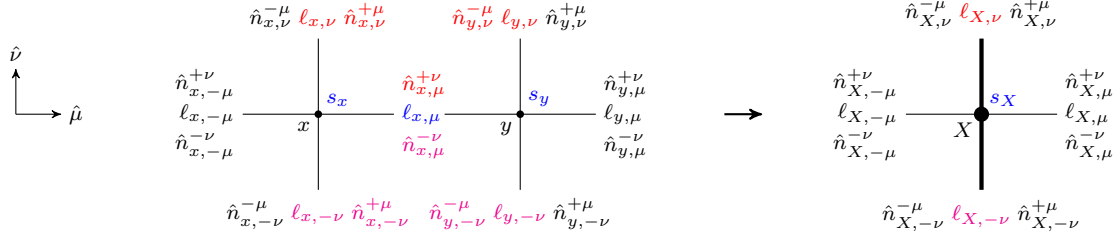


Figure 7.4: Visualization of the construction of the occupation numbers on a coarse site. The occupations on the coarse site (right) are computed from those on the fine sites (left) according to (7.3.8).

Given the occupation numbers on the fine grid and assuming that T_x and T_y both satisfy (7.3.1), we show in Appendix C.1 that (7.3.1) also holds after the contraction for the coarse-grid tensor T_X with

$$(T_X^{n_X^s})_{j_X} = \sum_{n_x^s=0}^{n_X^s} \sum_{n_y^s=0}^{n_X^s-n_x^s} \sum_{\substack{j_{x,\mu} \\ \ell_{x,\mu}=n_X^s-n_x^s-n_y^s}} (T_x^{n_x^s})_{j_x} (T_y^{n_y^s})_{j_y} \sigma_{F_x, F_y} \quad (7.3.7)$$

and occupation numbers given by

$$n_X^s \equiv n_x^s + n_y^s + \ell_{x,\mu}, \quad (7.3.8a)$$

$$\ell_{X,\nu} \equiv \begin{cases} \ell_{x,\nu} & \text{for } \nu = -\mu, \\ \ell_{y,\nu} & \text{for } \nu = \mu, \\ \ell_{x,\nu} + \ell_{y,\nu} + \frac{1}{2}(\hat{n}_{x,\nu}^{\mu} + \hat{n}_{y,\nu}^{-\mu}) & \text{for } |\nu| \neq \mu, \end{cases} \quad (7.3.8b)$$

$$\hat{n}_{X,\nu}^{\lambda} \equiv \begin{cases} \hat{n}_{x,\nu}^{\lambda} & \text{for } \nu = -\mu \text{ or } \lambda = -\mu, \\ \hat{n}_{y,\nu}^{\lambda} & \text{for } \nu = \mu \text{ or } \lambda = \mu, \\ \hat{n}_{x,\nu}^{\lambda} + \hat{n}_{y,\nu}^{\lambda} & \text{for } |\lambda|, |\nu| \neq \mu, \end{cases} \quad (7.3.8c)$$

where in (7.3.8c) we have $|\lambda| \neq |\nu|$. This is visualized in Fig. 7.4. Note that $\ell_{X,\nu}$ and $\hat{n}_{X,\nu}^{\lambda}$ are determined uniquely by $j_{X,\nu}$.

7.4 Reduction of the tensor

After a contraction the numerical tensor can be simplified in three different ways, which we now discuss in turn. The first and second simplification reduce the size of the tensor by removing some of the fat-index values $j_{X,\nu} = (j_{x,\nu}, j_{y,\nu})$. This reduces the range of $j_{X,\nu}$, i.e., the bond dimension of the fat link (X, ν) . The second and third simplification remove information that contributes to $\mathcal{O}(\beta^{n_{\max}+1})$ in the partition function. All three simplifications preserve the structure of (7.3.1).

First, we consider the definition of $\ell_{X,\nu}$ in the third line of (7.3.8b). The edge occupations $\hat{n}_{x,\nu}^{\mu}$ and $\hat{n}_{y,\nu}^{-\mu}$ are uniquely determined by $j_{x,\nu}$ and $j_{y,\nu}$. Although these edge occupations are independent variables for the tensors T_x and T_y , the hook conditions in (7.2.3), together with the link identity $(x, \mu) = (y, -\mu)$, imply that the entries of the coarse-grid tensor T_X are zero if these two edge occupations are not equal. Thus we can simply remove all values of the fat index $j_{X,\nu}$ for which $\hat{n}_{x,\nu}^{\mu} \neq \hat{n}_{y,\nu}^{-\mu}$ without losing information.

Second, we remove all fat-index values $j_{X,\nu}$ for which the link criterion

$$\ell_{X,\nu} + \sum_{\substack{\lambda=\pm 1 \\ |\lambda| \neq |\nu|}}^{\pm d} \hat{n}_{X,\nu}^\lambda \leq n_{\max} \quad (7.4.1)$$

is not satisfied.

Third, we set all entries of the coefficient tensors to zero for which the site criterion³

$$n_X^s + \sum_{\nu=\pm 1}^{\pm d} \ell_{X,\nu} + \sum_{\substack{\nu,\lambda=\pm 1 \\ |\nu| < |\lambda|}}^{\pm d} \hat{n}_{X,\nu}^\lambda \leq n_{\max} \quad (7.4.2)$$

is not satisfied. The link and site criteria, which remove $\mathcal{O}(\beta^{n_{\max}+1})$ contributions to the partition function, are derived in Appendix C.2.

In the actual implementation of the OS-GHOTRG method the reductions described above are applied on the fly during the contraction performed in Sec. 7.3.

7.5 Truncation of the tensor

When computing the partition function (5.2.15) numerically using the GHOTRG method, we contract pairs of adjacent tensors at each blocking step, thus halving the number of tensors on the grid. This is repeated until the tensor network has been reduced to a single tensor, which is then traced over all directions to yield the partition function. If we were to contract and reduce tensors as detailed in Secs. 7.3 and 7.4, the size of the tensor would increase exponentially during the blocking procedure. As explained in Sec. 3.3, a crucial ingredient of the GHOTRG method is the truncation of the bond dimension of the fat links (perpendicular to the contraction direction) after each contraction step to avoid such an exponential blowup. The truncation is based on the HOSVD concept to reduce the multi-rank of a tensor [21], in which an approximation $\hat{T}$ to a tensor T is constructed as

$$\hat{T}_{i_1 i_2 \dots} = U_{i_1 j_1}^{(1)} U_{i_2 j_2}^{(2)} \dots \bar{T}_{j_1 j_2 \dots} \quad \text{with} \quad \bar{T}_{j_1 j_2 \dots} = U_{i_1 j_1}^{(1)} U_{i_2 j_2}^{(2)} \dots T_{i_1 i_2 \dots} \quad (7.5.1)$$

Here, sums over repeated indices are implied, $\bar{T}$ is the core tensor, and the $U^{(k)}$ are tall and skinny semi-orthogonal matrices. The core tensor $\bar{T}$ has smaller multi-rank and smaller bond dimensions than T . The tensor approximation $\hat{T}$ has the same multi-rank as $\bar{T}$ but the same bond dimensions as T .

We now apply this concept to the numerical part T_X of the coarse-grid tensor, while the Grassmann part is discussed at the end of this subsection. We construct the core tensor

$$(\bar{T}_X)_{\bar{j}_X} \equiv \left(\prod_{\substack{\nu=\pm 1 \\ |\nu| \neq \mu}}^{\pm d} \sum_{j_{X,\nu}} \right) \left(\prod_{\substack{\nu=\pm 1 \\ |\nu| \neq \mu}}^{\pm d} (U_{X,\nu})_{j_{X,\nu}, \bar{j}_X, \nu} \right) (T_X)_{j_X}, \quad (7.5.2)$$

³Note that link and site criterion have to be slightly altered when the coarse grid only contains a single site in a particular direction, as the backward and forward links are identical due to the toroidal shape of the grid. This case can be taken into account by including a factor of t_λ (or t_ν) for every summation index λ (or ν) in (7.4.1) and (7.4.2), where $t_\lambda = 1/2$ if the lattice extent in the λ direction is one and $t_\lambda = 1$ otherwise (and similarly for t_ν).

where $\bar{j}_X \equiv (\bar{j}_{X,-1}, \bar{j}_{X,1}, \dots, \bar{j}_{X,-d}, \bar{j}_{X,d})$ with $\bar{j}_{X,\pm\mu} = j_{X,\pm\mu}$ for the contraction direction μ , and $U_{X,\nu}$ is a semi-orthogonal matrix whose columns are typically singular vectors of a matrix constructed using tensor unfoldings of the coarse-grid tensor with respect to its fat links. The first dimension of $U_{X,\nu}$ is the original bond dimension of the fat link, and the second dimension is the desired bond dimension after truncation. We impose the condition

$$U_{X,-\nu} = U_{X-\hat{\nu},\nu}, \quad (7.5.3)$$

such that the core tensor $\bar{T}_X$ builds the new tensor on the blocked lattice and the merger (see Sec. 3.2.1) drops out.

Inserting the decomposition of the coarse-grid tensor T_X (obtained from (7.3.1) with $x \rightarrow X$ and reduced as explained in Sec. 7.4) in (7.5.2) yields

$$(\bar{T}_X)_{\bar{j}_X} = \sum_{n_X^s=0}^{n_{\max}} \beta^{n_X^s} \left(\prod_{\substack{\nu=\pm 1 \\ |\nu| \neq \mu}}^{\pm d} \sum_{j_{X,\nu}} \right) \left(\prod_{\substack{\nu=\pm 1 \\ |\nu| \neq \mu}}^{\pm d} (U_{X,\nu})_{j_{X,\nu}, \bar{j}_{X,\nu}} \right) \beta^{n_X} (T_X^{n_X^s})_{j_X}. \quad (7.5.4)$$

The exponent n_X depends on the $j_{X,\nu}$. Due to the sum over $j_{X,\nu}$ in (7.5.4), a term in the sum over n_X^s contains contributions of different orders in β , i.e., these contributions mix. As we eventually want to obtain a decomposition of the partition function in powers of β we will modify the GHOTRG method so that the structure of (7.3.1) is preserved after truncation, i.e.,

$$(\bar{T}_X)_{\bar{j}_X} = \beta^{n_X} \sum_{n_X^s=0}^{n_{\max}} \beta^{n_X^s} (\bar{T}_X^{n_X^s})_{\bar{j}_X} \quad (7.5.5)$$

with coefficient tensors

$$(\bar{T}_X^{n_X^s})_{\bar{j}_X} = \left(\prod_{\substack{\nu=\pm 1 \\ |\nu| \neq \mu}}^{\pm d} \sum_{j_{X,\nu}} \right) \left(\prod_{\substack{\nu=\pm 1 \\ |\nu| \neq \mu}}^{\pm d} (U_{X,\nu})_{j_{X,\nu}, \bar{j}_{X,\nu}} \right) (T_X^{n_X^s})_{j_X}. \quad (7.5.6)$$

This can be achieved by ensuring that the truncated indices $\bar{j}_{X,\nu}$ again uniquely determine the occupation numbers $\ell_{X,\nu}$ and $\hat{n}_{X,\nu}^\mu$. In the following we explain this central part of the OS-GHOTRG method.

We start with the unfolding of the coarse-grid tensor T_X with respect to the fat index $j_{X,\nu}$,⁴

$$(M_{X,\nu})_{j_{X,\nu}, j_X \setminus j_{X,\nu}} \equiv (T_X)_{j_X}, \quad (7.5.7)$$

where $j_X \setminus j_{X,\nu}$ stands for the index tuple j_X without the index $j_{X,\nu}$. The index $j_{X,\nu}$ labels the rows of $M_{X,\nu}$, while the combination of all remaining indices labels its columns.

We first show that the edge occupations are not mixed in (7.5.4). The index $j_{X,\nu}$ uniquely determines all edge occupations $\hat{n}_{X,\nu}^\lambda$ (with $|\lambda| \neq |\nu|$). For nonzero entries the edge occupations satisfy the hook condition $\hat{n}_{X,\nu}^\lambda = \hat{n}_{X,\lambda}^\nu$, and hence the nonzero entries of a row of $M_{X,\nu}$ will only occur for columns for which the hook conditions are satisfied, i.e., all $\hat{n}_{X,\lambda}^\nu$ in the combined

⁴Note that the OS-GHOTRG method explicitly produces the coefficient tensors $T_X^{n_X^s}$. Therefore, to numerically compute the tensor entry $(T_X)_{j_X}$ using (7.3.1) an explicit choice of β is required. This choice could be viewed as an optimization opportunity, but for the numerical results in Part V we always used $\beta = 1$.

column index must match the $\hat{n}_{X,\nu}^\lambda$ of the row index $j_{X,\nu}$. This automatically yields a block-diagonal structure in the unfolding.⁵ An SVD of the unfolding respects this blocking, i.e., the singular vectors⁶ have the same block structure as the unfolding itself, and thus each singular vector, labeled by $\bar{j}_{X,\nu}$, can be uniquely allocated values of $\hat{n}_{X,\nu}^\lambda$ for all λ with $|\lambda| \neq |\nu|$.

We now turn to the link occupations, which are affected by mixing. The fat index $j_{X,\nu}$ also determines the link occupation $\ell_{X,\nu}$. When computing the SVD and applying the truncation matrix, the different ℓ values will mix such that the truncated index $\bar{j}_{X,\nu}$ no longer uniquely determines the link occupation and the decomposition (7.5.5) does not hold. We now introduce a procedure to construct truncation matrices $U_{X,\nu}$ for which the truncated index $\bar{j}_{X,\nu}$ still uniquely determines the link occupation $\ell_{X,\nu}$.

To avoid the mixing of different link occupations $\ell_{X,\nu}$ under truncation, we perform separate SVDs of sub-unfoldings with fixed link occupation. To this end we reorder the rows, labeled by the index $j_{X,\nu}$, in increasing order of link occupation $\ell_{X,\nu}(j_{X,\nu})$ and obtain

$$M_{X,\nu} = \begin{pmatrix} M_{X,\nu}^0 \\ M_{X,\nu}^1 \\ \vdots \\ M_{X,\nu}^{n_{\max}} \end{pmatrix}, \quad (7.5.8)$$

where each block $M_{X,\nu}^m$ corresponds to a fixed value $\ell_{X,\nu} = m$. Provided that (7.5.3) has been satisfied in all preceding blocking steps, $M_{X,\nu}$ and $M_{X+\hat{\nu},-\nu}$ have the same vertical block structure as in (7.5.8), although the number of columns could be different.

To obtain singular vectors that do not mix link occupations and construct the truncation matrix $U_{X,\nu}$, we apply again one of two procedures. Either a symmetric one based on the SuperQ method [29] or an asymmetric one based on the original idea by Xie et al. [20], where symmetric and asymmetric refers to the two tensors at the end of the truncated link.

In SuperQ, we perform SVDs of the $n_{\max} + 1$ extended matrices

$$\begin{pmatrix} M_{X,\nu}^m & M_{X+\hat{\nu},-\nu}^m \end{pmatrix} \quad \text{with} \quad 0 \leq m \leq n_{\max}. \quad (7.5.9)$$

The resulting singular-vector matrices $U_{X,\nu}^m$ are assembled in the block-diagonal square matrix

$$\bar{U}_{X,\nu} = \begin{pmatrix} U_{X,\nu}^0 & 0 & \dots & 0 \\ 0 & U_{X,\nu}^1 & \dots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \dots & U_{X,\nu}^{n_{\max}} \end{pmatrix}. \quad (7.5.10)$$

The semi-orthogonal truncation matrix $U_{X,\nu}$ is then obtained by retaining the columns of $\bar{U}_{X,\nu}$ corresponding to the D largest singular values of the set of matrices in (7.5.9), where D is the desired bond dimension of the truncated coarse-grid tensor. In the Xie et al. procedure, we perform SVDs of the matrices $M_{X,\nu}^m$ and $M_{X+\hat{\nu},-\nu}^m$ separately, again for $0 \leq m \leq n_{\max}$, resulting

⁵A permutation of rows and columns may be necessary to make the block-diagonal structure explicit.

⁶We always consider the left singular vectors, which form a square matrix whose dimension equals the row dimension of the unfolding.

in two matrices $\bar{U}_+$ and $\bar{U}_-$ constructed as in (7.5.10), from which we obtain U_+ and U_- as before. For $U_{X,\nu}$ we then pick either U_+ or U_- depending on which gives the smaller truncation error for $M_{X,\nu}$ or $M_{X+\hat{\nu},-\nu}$, respectively.⁷ SuperQ is preferred since it results in a smaller combined truncation error for $M_{X,\nu}$ and $M_{X+\hat{\nu},-\nu}$ [29]. If it becomes numerically unstable, which can happen close to the phase transition for large volume, we use the Xie et al. procedure instead.

The truncation matrix $U_{X,\nu}$ is block-diagonal, with blocks corresponding to different $\ell_{X,\nu}$. To apply $U_{X,\nu}$ in (7.5.4) its rows are sorted back to the original order of the fat index $j_{X,\nu}$. The truncated index $\bar{j}_{X,\nu}$ corresponds to the columns of $U_{X,\nu}$ and uniquely determines $\ell_{X,\nu}$. Together with the block structure discussed above for $\hat{n}_{X,\nu}^\lambda$ this ensures that the decomposition (7.5.5) holds, with the coefficient tensors given by (7.5.6).

We now turn to the Grassmann part of the coarse-grid tensor. From Sec. 3.3.2 we know that $M_{X,\nu}$ is block-diagonal in the Grassmann parity. This Grassmann-parity-based block structure also holds for the individual $M_{X,\nu}^m$ blocks. This implies that the matrices $U_{X,\nu}^m$ are themselves block-diagonal, where each of the two diagonal blocks corresponds to a different Grassmann parity. Similar to Sec. 3.3.2, if we apply the truncation matrices to the product of T_X and K_X , the “truncated” Grassmann tensor is simply K_X with $F_{X,\nu}(j_{X,\nu})$ replaced by $\bar{F}_{X,\nu}(\bar{j}_{X,\nu})$, which is the Grassmann parity of column $\bar{j}_{X,\nu}$ of $U_{X,\nu}$.

At the end of every blocking step, we rename $X \rightarrow x$ and drop the bar on all quantities. The partition function then has the form of (5.2.15) at every step of the blocking procedure, with half the number of lattice sites and new numerical tensors. As mentioned in Sec. 3.3.3, T_X will be independent of X after a sufficient number of blocking steps (at least $d - 1$ steps). This implies that the complexity of the tensor-network method scales like $\log V$.

7.6 Tensor trace

In the following we generalize the ASAP tracing from Sec. 3.3.5 for OS-GHOTRG. Therefore, when the lattice is reduced to a single site⁸ in a certain direction, say μ , during the iterative blocking procedure, the tensor will be traced in that direction and the dimension of the tensor network is reduced from d to $d - 1$, where d is the current dimension of the network. In this process we again have to separate contributions from different orders in β .

As part of the GHOTRG procedure, we integrate out the remaining Grassmann variable in direction μ and apply the corresponding boundary condition, resulting in a sign factor σ_μ , defined in Eq. (3.3.39). We then compute the trace to obtain

$$(T_X)_{j_X} \equiv \sum_{j_{x,\mu}} (T_x)_{j_x} \sigma_\mu \quad (7.6.1)$$

with summation index $j_{x,\mu} = j_{x,-\mu}$. Here, X is the site on the $(d - 1)$ -dimensional lattice obtained from x by dropping the coordinate x_μ , and $j_X \equiv j_x \setminus (j_{x,-\mu}, j_{x,\mu})$.

Given the occupation numbers on the d -dimensional lattice and assuming that T_x satisfies (7.3.1),

⁷The truncation error is $1 - \sum_{i=1}^D \sigma_i^2 / \sum_{i=1}^N \sigma_i^2$, where the σ_i are the singular values and N is the dimension of the fat index. Usually $N = D^2$.

⁸We assume that the linear extent of the lattice is a power of 2 in every direction.

we show in Appendix C.3 that (7.3.1) also holds after the tracing for T_X on the $(d-1)$ -dimensional lattice, with

$$(T_X^{n_X^s})_{j_X} = \sum_{n_x^s=0}^{n_X^s} \sum_{\substack{j_{x,\mu} \\ \ell_{x,\mu}=n_X^s-n_x^s}} (T_x^{n_x^s})_{j_x} \sigma_\mu, \quad (7.6.2)$$

occupation numbers for $|\nu|, |\lambda| \neq \mu$ and $|\nu| \neq |\lambda|$ given by

$$s_X = s_x + \ell_{x,\mu}, \quad (7.6.3a)$$

$$\ell_{X,\nu} = \ell_{x,\nu} + \frac{1}{2}(\hat{n}_{x,\nu}^\mu + \hat{n}_{x,\nu}^{-\mu}), \quad (7.6.3b)$$

$$\hat{n}_{X,\nu}^\lambda = \hat{n}_{x,\nu}^\lambda, \quad (7.6.3c)$$

and

$$n_X(j_X) = \frac{1}{2} \sum_{\substack{\nu=\pm 1 \\ |\nu| \neq \mu}}^{\pm d} \ell_{X,\nu} + \frac{1}{8} \sum_{\substack{\nu,\lambda=\pm 1 \\ |\nu|,|\lambda| \neq \mu \\ |\nu| \neq |\lambda|}}^{\pm d} \hat{n}_{X,\nu}^\lambda. \quad (7.6.4)$$

Note that $\ell_{X,\nu}$ and $\hat{n}_{X,\nu}^\lambda$ are uniquely determined by $j_{X,\nu}$.

We now rename $X \rightarrow x$ and relabel the $d-1$ remaining directions from 1 to $d-1$ to be able to start the next blocking step using the same notation as before. After each blocking step, we apply a rescaling as in Sec. 3.3.3, except that the scaling factor t_i after the i -th blocking step is determined by the maximal absolute value

$$t_i = \max_{x,j_x,n_x^s} |(T_x^{n_x^s})_{j_x}|, \quad (7.6.5)$$

where $T_x^{n_x^s}$ are the tensors that are obtained after the i -th blocking step (including the renaming).

Note that once we have arrived at a one-dimensional lattice, the blocking no longer involves a truncation since there is no perpendicular direction. Assuming a total number of k blocking steps to obtain a lattice consisting of a single site, the rescaled partition function up to order $n_{\max}$ is given by

$$Z^{(k)} = \sum_{j_{x,1}} (T_x)_{j_{x,-1},j_{x,1}} \sigma_1 \quad (7.6.6)$$

with $j_{x,-1} = j_{x,1}$ and σ_1 from Eq. (3.3.39). Repeating the steps above, now with $d=1$, we obtain the rescaled partition function as a polynomial in β ,

$$Z^{(k)} = \sum_{n=0}^{n_{\max}} \beta^n Z_n^{(k)} \quad \text{with} \quad Z_n^{(k)} = \sum_{n_x^s=0}^n \sum_{\substack{j_{x,1} \\ \ell_{x,1}=n-n_x^s}} (T_x^{n_x^s})_{j_{x,1},j_{x,1}} \sigma_1. \quad (7.6.7)$$

Finally, the partition function (5.2.15) up to order $n_{\max}$ is given by

$$Z = \sum_{n=0}^{n_{\max}} \beta^n Z_n \quad \text{with} \quad Z_n = Z_n^{(k)} \prod_{i=1}^k t_i^{V_i} \quad (7.6.8)$$

with $V_i = V/2^i$.

8 Translational invariance for OS-GHOTRG

Using the OS-GHOTRG algorithm we can again use the translational invariance on the initial lattice to increase the numerical efficiency of the algorithm for $n_{\max} \geq 1$. Similar to Sec. 6, we include only a subset of equivalent configurations and apply a combinatorial factor to also take into account the others. Using the following method, the number of explicitly included configurations is reduced further compared to Sec. 6 for $n_{\max} > 1$. For simplicity, we only discuss this idea for the two-dimensional case. An extension to more than two dimensions is the subject of future work.

8.1 Contributions to the partition function

Let $(q_1, \dots, q_r)$ be a tuple of r positive integers. We define $Z_{(q_1, \dots, q_r)}$ to be the sum of all contributions to the partition function of valid configurations $\mathbf{j}$ with nonzero plaquette occupations⁹ as in the tuple. Furthermore, we define $Z_{(q_1|q_2, \dots, q_r)}$ just like $Z_{(q_1, \dots, q_r)}$, but with the additional constraint that the plaquette at the origin has occupation q_1 .

In $Z_{(q_1, \dots, q_r)}$ the order of the q_i is irrelevant since we sum over all possible locations of the plaquette occupations on the lattice. Hence, to obtain all contributions to Z_n in (7.6.8) it is sufficient to sum $Z_{(n)}$ over all partitions (n) of n , defined by¹⁰ $(n) = (n_1, \dots, n_1, n_2, \dots, n_2, \dots, n_m, \dots, n_m)$, where $n_1 > n_2 > \dots > n_m > 0$ and k_i is the multiplicity of n_i , i.e., $n = \sum_{i=1}^m k_i n_i$. Translational invariance allows us to obtain $Z_{(n)}$ as¹¹

$$Z_{(n)} = \frac{V}{k_i} Z_{(n_i | n \setminus n_i)} \quad \text{for any } i \in \{1, \dots, m\}, \quad (8.1.1)$$

where $(n \setminus n_i)$ means that the first occurrence of n_i in (n) has been removed. The desired coefficients Z_n are thus given by

$$Z_1 = Z_{(1)} = V Z_{(1)}, \quad (8.1.2a)$$

$$Z_2 = Z_{(2)} + Z_{(1,1)} = V Z_{(2)} + \frac{V}{2} Z_{(1|1)}, \quad (8.1.2b)$$

$$Z_3 = Z_{(3)} + Z_{(2,1)} + Z_{(1,1,1)} = V Z_{(3)} + V Z_{(2|1)} + \frac{V}{3} Z_{(1|1,1)}, \quad (8.1.2c)$$

$$Z_4 = Z_{(4)} + Z_{(3,1)} + Z_{(2,2)} + Z_{(2,1,1)} + Z_{(1,1,1,1)}$$

⁹For a valid configuration the hook conditions are satisfied and the plaquette occupation is the same as the four identical edge occupations around this plaquette, see Fig. 5.1. When we speak of plaquette occupations we always mean those of the initial lattice.

¹⁰In the interest of readability we have overloaded the notation, denoting by (n) both a generic partition and the specific partition with $n_1 = n$.

¹¹Consider a valid configuration $\mathbf{j}$ with nonzero plaquette occupations at lattice sites $x_1, \dots, x_r$ given as in (n) , where $r = \sum_{i=1}^m k_i$. There are V shifted replicas $\mathbf{j}'$ of this configuration (including shift $y = 0$) which make identical contributions to $Z_{(n)}$. Out of these V replicas, only those k_i replicas are kept in $Z_{(n_i | n \setminus n_i)}$ for which the shift y is such that the occupation of the plaquette at the origin is n_i . If $(x_1, \dots, x_r)$ has special symmetries such that some shifted configurations are identical to one another on the torus (e.g., if for a lattice of size $L \times L$ and partition $(1, 1)$ with $x_1 = (0, 0)$ and $x_2 = (L/2, L/2)$ we shift by $y = (L/2, L/2)$), the number of replicas is reduced to V/s and k_i/s , respectively, with symmetry factor s . Therefore all configurations considered in $Z_{(n_i | n \setminus n_i)}$ contribute to $Z_{(n)}$ with combinatorial factor V/k_i independent of $x_1, \dots, x_r$.

$$= V Z_{(4)} + V Z_{(3|1)} + \frac{V}{2} Z_{(2|2)} + V Z_{(2|1,1)} + \frac{V}{4} Z_{(1|1,1,1)} \quad (8.1.2d)$$

and so on for larger n . We now construct two tensor networks from which we will determine the Z_n .¹²

8.2 The X network

For the first tensor network, which we call the X network, we require that the occupation of the plaquette at the origin is nonzero. To this end we introduce four initial “special” tensors $S_1, \dots, S_4$ at the corners of the plaquette at the origin. We then restrict the occupations of all other plaquettes, which we call restricted plaquettes, to $n \leq n_R$. Here, n_R is a parameter in the interval $[\lfloor n_{\max}/2 \rfloor, n_{\max} - 1]$,¹³ where the lower bound ensures that all configurations can be included and the upper bound follows from the fact that the occupation of the plaquette at the origin is at least one. The tensors S_i are surrounded by the plaquette at the origin with plaquette occupation $n \in \{1, \dots, n_{\max}\}$ (red plaquette in Fig. 8.5) and three restricted plaquettes with occupation $n \in \{0, \dots, n_R\}$ (black plaquettes). Similar to Sec. 6, the tensors on all other lattice sites, which we call “restricted” tensors R , are surrounded by four restricted plaquettes.

The blocking procedure now proceeds as shown in Fig. 8.5. After one blocking step in each of the two directions, the volume V has been reduced by a factor of 4, and the occupation of the plaquette at the origin has been turned into the site occupation of the resulting S tensor. We decompose this tensor as in (7.3.1) but introduce additional scaling coefficients c_s , resulting in the tensor

$$(\tilde{S})_j \equiv \beta^{n(j)} \sum_{n^s=1}^{n_{\max}} \beta^{n^s} (\tilde{S}^{n^s})_j \quad \text{with} \quad \tilde{S}^{n^s} \equiv c_{n^s} S^{n^s}, \quad (8.2.1)$$

where the c_s are determined below. Note that s corresponds to the occupation of the plaquette at the origin of the initial lattice and that the sum starts at $s = 1$ since here we only consider configurations for which this occupation is at least one.

Denoting a blocking step in direction μ by $\star_\mu$, the steps can be summarized as

$$\begin{aligned} \text{Step 1:} \quad & R^{(0)} \star_1 R^{(0)} \rightarrow R^{(1)}, \quad S_1^{(0)} \star_1 S_2^{(0)} \rightarrow S_1^{(1)}, \quad S_3^{(0)} \star_1 S_4^{(0)} \rightarrow S_2^{(1)}, \\ \text{Step 2:} \quad & R^{(1)} \star_2 R^{(1)} \rightarrow R^{(2)}, \quad S_1^{(1)} \star_2 S_2^{(1)} \rightarrow S^{(2)} \rightarrow \tilde{S}^{(2)} \quad \text{using (8.2.1)}, \\ \text{Step } i+1: \quad & R^{(i)} \star_\mu R^{(i)} \rightarrow R^{(i+1)}, \quad \tilde{S}^{(i)} \star_\mu R^{(i)} \rightarrow \tilde{S}^{(i+1)} \quad \text{for } i \geq 2, \end{aligned} \quad (8.2.2)$$

where the last line is repeated in all directions until we arrive at a single point.¹⁴

With this modified blocking procedure we now compute the full contraction of the tensor network involving one tensor $\tilde{S}$ and $V/4 - 1$ restricted tensors R , written symbolically as

$$X = \text{tr} \left[\tilde{S}^{(2)} \left(R^{(2)} \right)^{V/4-1} \right], \quad (8.2.3)$$

¹²It is conceivable that translational invariance could be exploited in a slightly different way using more than two tensor networks, where in each network the occupation of the origin plaquette is fixed at a particular value. The pros and cons of such an approach have not yet been explored.

¹³ $\lfloor \dots \rfloor$ denotes the floor function.

¹⁴As in Sec. 6, we construct the truncation matrices only from the tensors R , except for the truncation of $S_1^{(1)}$ and $S_2^{(1)}$.

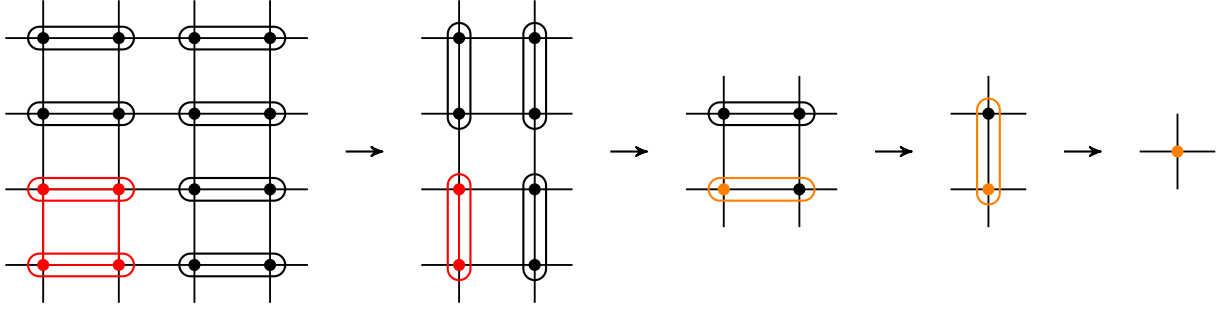


Figure 8.5: Illustration of the blocking procedure using translational invariance on the initial lattice. The plaquette at the origin and the S tensors are indicated in red, the tensors $\tilde{S}$ are indicated in orange, and restricted plaquettes and tensors R are indicated in black, respectively.

using the OS-GHOTRG method of Sec. 7. In the reduction step of the OS-GHOTRG method, see Sec. 7.4, we can replace $n_{\max}$ by $n_{\max} - 1$ when computing $R \star R$ since the occupation of the plaquette at the origin is nonzero. From (7.6.8) we obtain

$$X = \sum_{n=1}^{n_{\max}} X_n \beta^n \quad (8.2.4)$$

with¹⁵

$$X_1 = c_1 Z_{(1|)}, \quad (8.2.5a)$$

$$X_2 = c_2 Z_{(2|)} + c_1 Z_{(1|1)} \quad \text{for } n_R \geq 1, \quad (8.2.5b)$$

$$\begin{aligned} X_3 &= c_3 Z_{(3|)} + c_2 Z_{(2|1)} + c_1 (\Theta(n_R - 2) Z_{(1|2)} + Z_{(1|1,1)}) \quad \text{for } n_R \geq 1 \\ &= c_3 Z_{(3|)} + (c_2 + c_1 \Theta(n_R - 2)) Z_{(2|1)} + c_1 Z_{(1|1,1)}, \end{aligned} \quad (8.2.5c)$$

$$\begin{aligned} X_4 &= c_4 Z_{(4|)} + c_3 Z_{(3|1)} + c_2 (Z_{(2|2)} + Z_{(2|1,1)}) + c_1 (\Theta(n_R - 3) Z_{(1|3)} + Z_{(1|2,1)} + Z_{(1|1,1,1)}) \\ &\quad \text{for } n_R \geq 2 \\ &= c_4 Z_{(4|)} + (c_3 + c_1 \Theta(n_R - 3)) Z_{(3|1)} + c_2 Z_{(2|2)} + (c_2 + 2c_1) Z_{(2|1,1)} + c_1 Z_{(1|1,1,1)}, \end{aligned} \quad (8.2.5d)$$

and so on for larger n . Here, $\Theta(x)$ is the Heaviside function with $\Theta(0) \equiv 1$. To go from the first to the second line in (8.2.5c) and (8.2.5d) we shifted the configurations such that the highest plaquette occupation is at the origin and used

$$Z_{(n_i|n \setminus n_i)} = \frac{k_i}{k_1} Z_{(n_1|n \setminus n_1)}, \quad (8.2.6)$$

which follows from (8.1.1).

8.3 The Y network

We also compute the full contraction of another tensor network only consisting of tensors R , written symbolically as

$$Y = \text{tr} \left[\left(R^{(0)} \right)^V \right] = \text{tr} \left[\left(R^{(2)} \right)^{V/4} \right]. \quad (8.3.1)$$

¹⁵For simplicity, we omit the rescaling in this section.

Note that X and Y can be computed together efficiently, since $R \star R$ is needed for both quantities. Applying the OS-GHOTRG method with restriction to order $n_{\max} - 1$ in each blocking step, we obtain

$$Y = \sum_{n=0}^{n_{\max}-1} Y_n \beta^n, \quad (8.3.2)$$

where the Y_n contain the contributions to Z_n of all configurations for which all plaquettes of the initial lattice have occupation $\leq n_R$. These sums are given by¹⁶

$$Y_0 = Z_0 \quad \text{for } n_{\max} \geq 1, n_R \geq 0, \quad (8.3.3a)$$

$$Y_1 = Z_{(1)} = V Z_{(1|)} \quad \text{for } n_{\max} \geq 2, n_R \geq 1, \quad (8.3.3b)$$

$$Y_2 = \Theta(n_R - 2) Z_{(2)} + Z_{(1,1)} = \Theta(n_R - 2) V Z_{(2|)} + \frac{V}{2} Z_{(1|1)} \quad \text{for } n_{\max} \geq 3, n_R \geq 1, \quad (8.3.3c)$$

$$Y_3 = \Theta(n_R - 3) Z_{(3)} + Z_{(2,1)} + Z_{(1,1,1)} = \Theta(n_R - 3) V Z_{(3|)} + V Z_{(2|1)} + \frac{V}{3} Z_{(1|1,1)} \\ \text{for } n_{\max} \geq 4, n_R \geq 2, \quad (8.3.3d)$$

and so on for larger n . In the second step of each line we have used translational invariance as in (8.1.2).

8.4 Matching the coefficients

To determine the desired Z_n we write them as $Z_n = a_n X_n + b_n Y_n$ and obtain a system of equations to be solved for the coefficients a_n , b_n , and c_n . In App. D we show that a solution for any $n_{\max}$ and $\lfloor n_{\max}/2 \rfloor \leq n_R \leq n_{\max} - 1$ is given by

$$c_n = \begin{cases} n/n_{\max} & \text{for } n \leq n_R, \\ 1 & \text{for } n > n_R, \end{cases} \quad b_n = 1 - \frac{na_n}{Vn_{\max}} \quad \text{with} \quad a_n = \begin{cases} \text{arbitrary} & \text{for } n \leq n_R, \\ V & \text{for } n > n_R. \end{cases} \quad (8.4.1)$$

For $n = 0$ we have $b_0 = 1$, and a_0 is irrelevant since $X_0 = 0$, i.e., $Z_0 = Y_0$. For $1 \leq n \leq n_R$, choosing $a_n = Vn_{\max}/n$ results in $b_n = 0$. If in addition we choose $n_R = n_{\max} - 1$ we have $b_n = 0$ for all $n \geq 1$. Equivalent solutions can be obtained by rescaling $a_n \rightarrow ra_n$ and $c_n \rightarrow c_n/r$ for all n and arbitrary r . All results presented in Part V were obtained with $n_R = \lfloor n_{\max}/2 \rfloor$ and $a_n = Vn_{\max}/n$ for $1 \leq n \leq n_R$, for which the explained approach in this section was developed first. The more general solution (8.4.1) was mainly developed by Robert Lohmayer. We did not perform a systematic study of the dependence of the results on the choice of n_R and a_n , because such a study would be quite expensive numerically.

¹⁶Again, we omit the rescaling for simplicity.

V. Results

In the following chapter, the results of this thesis are presented. The main results are published in [27] and [28]. Before we discuss the numerical results obtained from tensor-network methods in Sec. 10, we first derive several analytical results.

9 Analytical Calculations

The analytical calculations are split into three sections. First, we state an efficient algorithm for calculating the coefficients of the partition function in the strong-coupling expansion exactly. The method is based on the tensor-network formulation, derived in Part III. As a result, the exact coefficients as functions of the masses and chemical potentials up to order 4 (for $N_f = 1$) and order 1 (for $N_f = 2$) for a 2×2 lattice are stated in App. E. In the subsequent section, we consider the analytical simplifications of the tensor network derived in Part III, when pure-gauge theory in two dimensions is considered. In this case, the partition function can be evaluated using HOTRG or OS-HOTRG. In order to be able to compare the corresponding tensor-network results with exact results, as done in Sec. 10, we calculate the coefficients of the strong-coupling expansion in this case and show that the strong-coupling expansion of the partition function on a finite lattice factorizes in the volume up to a certain order. Both results are based on the known exact solution for the pure-gauge theory in two dimensions [73, 74]. Finally, in Sec. 9.3, we consider again arbitrary dimensions and orders of the strong-coupling expansion and derive a variety of analytical properties of the partition function and the observables. They have relevant consequences for the exact expressions in App. E, and, more important, for the numerical results in Sec. 10. Consequently, they serve as a consistency check.

9.1 Analytical calculation for small lattices

In the following, we present an efficient algorithm for calculating the coefficients $Z_n(\mu, m)$ in

$$Z(\mu, m, \beta) = \sum_{n=0}^{\infty} Z_n(\mu, m) \beta^n \quad (9.1.1)$$

exactly. This is done in three steps, whereby the first two steps consist of up-front calculations, and the third step is the actual calculation.

up-front calculation 1:

The aim in this step is to find for given flavor f all tuples $(k_{x,\mu}^f, \bar{k}_{x,\mu}^f \mid \forall x, \mu)$ with $\Theta(w_x^f) \delta_{h_x^{f,\text{in}}, h_x^{f,\text{out}}} = 1$ for all x . The tuples are categorized with respect to $((k_{x,\mu}^f - \bar{k}_{x,\mu}^f) \bmod N_c \mid \forall x, \mu)$, which is called the “fermion gauge key”.

An efficient algorithm to find these tuples is given in the following: First, we iterate through all numbers in base $N_c + 1$, where the number of digits is the number of links dV . Therefore, we

iterate up to the number $(N_c + 1)^{dV}$. Since the link occupation number $k_{x,\mu}^f$ is in $\{0, \dots, N_c\}$, the iteration goes through all possible tuples $(k_{x,\mu}^f | \forall x, \mu)$ with which we begin with. When no site has more than N_c incoming and outgoing hopping terms (i.e., $h_x^{f,\text{in}}, h_x^{f,\text{out}} \leq N_c$ for all x)¹, we add $(k_{x,\mu}^f | \forall x, \mu)$ to a list, creating the so-called “forward hopping fill”.

The “backward hopping fill”, which consists of all tuples $(\bar{k}_{x,\mu}^f | \forall x, \mu)$ with $h_x^{f,\text{in}}, h_x^{f,\text{out}} \leq N_c$ for all x ² is identical to the forward hopping fill, which is why it does not need to be calculated separately.

We can now perform a nested iteration through the forward and backward hopping fill, and whenever the condition $\Theta(w_x^f) \delta_{h_x^{f,\text{in}}, h_x^{f,\text{out}}} = 1$ is fulfilled for all x , we add the tuple $(k_{x,\mu}^f, \bar{k}_{x,\mu}^f | \forall x, \mu)$ to a dictionary, where the key (called “fermion gauge key”) is the tuple $(\tilde{k}_{x,\mu}^f | \forall x, \mu)$ with $\tilde{k}_{x,\mu}^f \equiv (k_{x,\mu}^f - \bar{k}_{x,\mu}^f) \bmod N_c$. The resulting dictionary is called “fermion gauge dictionary”.

Since in this step, different flavors can be treated independently, we do not need to repeat this step for every flavor, but use the same gauge dictionary for every flavor.

up-front calculation 2:

The aim of the second step is to precompute the numerical factor that is already determined by given link occupation numbers $k_{x,\mu}^f$ and $\bar{k}_{x,\mu}^f$ for one flavor f and for all links (x, μ) of the lattice.

To do so, we run through all entries $(k_{x,\mu}^f, \bar{k}_{x,\mu}^f | \forall x, \mu)$ of the fermion gauge dictionary and calculate the expression

$$r_{(k_{x,\mu}^f, \bar{k}_{x,\mu}^f | \forall x, \mu)}(\mu_f, m_f) \equiv \exp \left(\mu_f \sum_x (k_{x,1}^f - \bar{k}_{x,1}^f) \right) \frac{(2m_f)^{\sum_x w_x^f}}{\prod_x w_x^f! \prod_{\mu} k_{x,\mu}^f! \bar{k}_{x,\mu}^f!} . \quad (9.1.2)$$

as a function of μ_f and m_f . By treating these two parameters as variables, different flavors can be treated independently, and we do not need to repeat this step for every flavor.

main calculation:

After we have completed the preliminary work, we can iterate through every order n , for which we want to calculate the coefficient Z_n . For every order n , we iterate through all tuples of oriented plaquette occupation numbers

$$(n_{x,\mu\nu} | \forall x, \mu \neq \nu) \quad \text{with} \quad \sum_{\substack{x,\mu,\nu \\ \mu \neq \nu}} n_{x,\mu\nu} = n . \quad (9.1.3)$$

Hereby we can reduce the number of relevant tuples by using the translational invariance.

For every tuple of oriented plaquette occupation numbers, we now perform $N_f - 1$ nested iterations, each over all fermion gauge keys. The fermion gauge key of the remaining flavor f with

¹At this level, we set $\bar{k}_{x,\mu}^f = 0$ for all x, μ , which occurs inside $h_x^{f,\text{in}}$ and $h_x^{f,\text{out}}$.

²Similarly to the forward hopping fill, we now set $k_{x,\mu}^f = 0$ for all x, μ , which occurs inside $h_x^{f,\text{in}}$ and $h_x^{f,\text{out}}$ as well.

$f = N_f$ is then determined by

$$\left(\left(- \sum_{\substack{\nu=\pm 1 \\ |\nu| \neq \mu}}^{\pm d} (n_{x,\mu\nu} + n_{x,\nu\mu}) - \sum_{f=1}^{N_f-1} \tilde{k}_{x,\mu}^f \right) \bmod N_c \mid \forall x, \mu \right), \quad (9.1.4)$$

which ensures that the $SU(N_c)$ is fulfilled at every link (x, μ) .

For every given order n , tuple of oriented plaquette occupation numbers, and given gauge keys for all N_f flavors, we now perform (additionally) N_f nested iterations. The iteration corresponding to flavor f runs over all entries belonging to the fermion gauge key of flavor f in the fermion gauge dictionary. Thus we have set all link occupation numbers $k_{x,\mu}^f$ and $\bar{k}_{x,\mu}^f$, which allows us to calculate all quantities that do not depend on the auxiliary summation index $r_{x,\mu}$. Thus we explicitly calculate³

$$S_{(k_{x,\mu}^f, \bar{k}_{x,\mu}^f, n_{x,\mu\nu} \mid \forall x, \mu \neq \nu, f)} \equiv \frac{\prod_x \sigma_x \prod_\mu \eta_{x,\mu}^{F_{x,\mu}}}{\prod_{\substack{\mu \neq \nu \\ x, \mu\nu}} n_{x,\mu\nu}!} \left[\prod_f r_{(k_{x,\mu}^f, \bar{k}_{x,\mu}^f \mid \forall x, \mu)}(\mu_f, m_f) \right] \int_\chi \prod_x K_x. \quad (9.1.5)$$

Finally, we additionally iterate through the auxiliary index $r_{x,\mu}$ on every link (x, μ) in nested loops and calculate all terms

$$S_{(k_{x,\mu}^f, \bar{k}_{x,\mu}^f, n_{x,\mu\nu} \mid \forall x, \mu \neq \nu, f)} \prod_x C_x \prod_\mu \mathcal{L}_{x,\mu}. \quad (9.1.6)$$

After applying the factor from translational invariance, as well as $1/(2N_c)^n$, we add all results corresponding to order n in order to obtain Z_n .

In App. E, the coefficients $Z_n(m, \beta)$ obtained with this algorithm are listed for a 2×2 lattice. For $N_f = 1$, the coefficients are given up to $n = 4$, while for $N_f = 2$ they are given up to $n = 1$ for $m_1 = m_2$ and $\mu_1 = \mu_2$.

9.2 Pure-gauge theory

Next, we consider the $SU(N_c)$ pure-gauge theory. Since the derivation of the initial tensor network in Part III holds for an arbitrary number of flavors N_f , the initial tensor network for the pure-gauge theory is obtained directly by setting $N_f = 0$. For clarity, we state the resulting tensor network explicitly. It reads

$$Z^{\text{PG}} = \sum_j \prod_x T_x^{\text{PG}}, \quad (9.2.1)$$

where the indices $j^{\text{PG}} \equiv (j_{x,\mu}^{\text{PG}} \mid \forall x, \mu)$ do not contain link occupation numbers $k_{x,\mu}$ and $\bar{k}_{x,\mu}$, as these originate from Grassmann hopping terms. Thus we obtain

$$j_{x,\mu}^{\text{PG}} = ((n_{x,\mu}^{\pm\nu}, \bar{n}_{x,\mu}^{\pm\nu} \mid \forall \nu \neq \mu), r_{x,\mu}). \quad (9.2.2)$$

³Note that $\int_\chi \prod_x K_x$ only results in a sign factor.

Consequently, the Grassmann part of the local tensor reduces to one, and the tensor network can be evaluated using usual HOTRG and OS-HOTRG.

The numerical part of the initial tensor is given by

$$T_x^{\text{PG}} \equiv \Delta_x^P \bar{T}_x^{\text{PG}} (\mathbf{n}_x \rightarrow \text{edge variables}) \quad (9.2.3)$$

with⁴

$$\bar{T}_x^{\text{PG}} \equiv C_x^{\text{PG}} W_x^{\text{PG}} \left[\prod_{\mu} \text{sgn}(\mathcal{L}_{x,\mu}) \right] \prod_{\mu=\pm 1}^{\pm d} \sqrt{|\mathcal{L}_{x,\mu}|}. \quad (9.2.4)$$

This tensor contains the pure-gauge color factor

$$C_x^{\text{PG}} = \Delta_x I_x \quad (9.2.5)$$

and the pure-gauge weight

$$W_x^{\text{PG}} = \left[\prod_{\mu \neq \nu} \frac{\left(\frac{\beta}{2N_c} \right)^{n_{x,\mu\nu} + n_{x-\hat{\mu},\mu\nu} + n_{x-\hat{\mu}-\hat{\nu},\mu\nu} + n_{x-\hat{\nu},\mu\nu}}}{n_{x,\mu\nu}! n_{x-\hat{\mu},\mu\nu}! n_{x-\hat{\mu}-\hat{\nu},\mu\nu}! n_{x-\hat{\nu},\mu\nu}!} \right]^{\frac{1}{4}}. \quad (9.2.6)$$

Finally, the criterion (5.2.17) ensures that the initial tensor network is finite-dimensional. It is optimized for a calculation of order n_{max} by applying the criterion (5.2.18) for every tensor entry.

As a special property of the pure-gauge theory in two dimensions, there exists yet another criterion that allows us to reduce the size of the initial tensor. It can be used whenever $n_{\text{max}} < V$. To formulate this criterion, we define the “factorization condition” $(n_{x,\mu\nu} - n_{x,\nu\mu})/N_c \in \mathbb{Z}$ at every geometrical plaquette. The $\text{SU}(N_c)$ condition $|g_{x,\mu} - \bar{g}_{x,\mu}|/N_c \in \mathbb{N}_0$ reduces to

$$\begin{aligned} |g_{x,\mu} - \bar{g}_{x,\mu}|/N_c &= \frac{1}{N_c} \left| \sum_{\substack{\nu \\ \nu \neq \mu}} (n_{x,\mu\nu} + n_{x-\hat{\nu},\nu\mu}) - \sum_{\substack{\nu \\ \nu \neq \mu}} (n_{x,\nu\mu} + n_{x-\hat{\nu},\mu\nu}) \right| \\ &= \left| \sum_{\substack{\nu \\ \nu \neq \mu}} (n_{x,\mu\nu} - n_{x,\nu\mu})/N_c - \sum_{\substack{\nu \\ \nu \neq \mu}} (n_{x-\hat{\nu},\mu\nu} - n_{x-\hat{\nu},\nu\mu})/N_c \right| \in \mathbb{N}_0 \end{aligned} \quad (9.2.7)$$

for pure-gauge theory. Note that we rearranged the sums in the second step such that the factorization conditions appear. Assume that for a given link the $\text{SU}(N_c)$ condition is satisfied, but at one adjacent plaquette the factorization condition is not satisfied. In this case, due to Eq. (9.2.7), there needs to be at least one other adjacent plaquette, where the factorization condition is not satisfied. For any valid configuration in two dimensions, it follows that whenever the factorization condition at a plaquette is not satisfied, then it is not satisfied as well at all four neighboring plaquettes. We can repeat this argument for these plaquettes and conclude that the factorization condition is not satisfied at any plaquette of the lattice in this case. We note that such a configuration contributes at least to order β^V . Therefore, whenever $n_{\text{max}} < V$, we can reduce the size of the initial tensor significantly by removing those link-indices from the initial

⁴ $\mathcal{L}_{x,\mu}$ is defined in Eqs. (5.2.2) and (4.3.12). However, as there are no link occupation numbers $k_{x,\mu}^f$ and $\bar{k}_{x,\mu}^f$ for $N_f = 0$, the equivalence classes become trivial, i.e., $[\pi] = \{\pi\}$ for all $\pi \in S_p$ and $[\alpha] = \{\alpha\}$ for all $\alpha \in A$. Thus, the ranges of the sums over α , π , and σ are not reduced.

tensor for which the factorization condition is not satisfied at the adjacent plaquettes. It follows that for $n_{\max} < V$, every configuration consists exclusively of plaquette-antiplaquette contributions and triple (anti-)plaquette contributions, including their multiples, as these contributions satisfy the factorization condition at every plaquette.

We can compare the results obtained from OS-HOTRG with the exact solution (see Sec. 10.2) for the two-dimensional $SU(N_c)$ pure-gauge theory in a finite volume V , given in Ref. [73]. It reads

$$Z^{\text{PG}}(\beta) = \sum_R \lambda_R(\beta)^V, \quad (9.2.8)$$

which contains a sum over irreducible representations $R = (q_1, \dots, q_{N_c-1})$ with Dynkin labels $q_i \in \mathbb{N}_0$. The corresponding Young diagram (cf. Sec. 1.7) is determined by the partition $(\bar{q}_1, \bar{q}_2, \dots, \bar{q}_{N_c-1})$ with $\bar{q}_1 \geq \bar{q}_2 \geq \dots \geq \bar{q}_{N_c-1} \geq 0$ and $\bar{q}_i = \sum_{j=i}^{N_c-1} q_j$. The coefficient $\lambda_R(\beta)$, given by⁵

$$\lambda_R(\beta) \equiv \frac{1}{d_R} \sum_{Q \in \mathbb{Z}} \det \left[I_{\bar{q}_j + i - j + Q}(\beta/N_c) \right]_{i,j=1 \dots N_c}, \quad (9.2.9)$$

has several ingredients, which are discussed in the following. First, the dimension

$$d_R = \Delta(\ell_1, \ell_2, \dots, \ell_{N_c-1}, \ell_{N_c}) / \Delta(N_c - 1, N_c - 2, \dots, 1, 0) \quad (9.2.10)$$

with $\ell_i \equiv \bar{q}_i + N_c - i$ and $\bar{q}_{N_c} = 0$ contains the factor

$$\Delta(x_1, \dots, x_{N_c}) \equiv \prod_{i < j} (x_i - x_j). \quad (9.2.11)$$

Second, $I_n(z)$ are the modified Bessel functions of the first kind, defined by [75]

$$I_n(z) = \frac{1}{\pi} \int_0^\pi e^{z \cos(\theta)} \cos(n\theta) d\theta. \quad (9.2.12)$$

Clearly, the result (9.2.8) does not factorize in the volume (cf. Ref. [74]). However, as shown in App. F, it holds that

$$\lambda_R(\beta) = \mathcal{O}\left((\beta/N_c)^{\bar{q}_1}\right) \quad (9.2.13)$$

and in particular

$$\lambda_R(\beta) = \mathcal{O}(\beta/N_c) \quad \text{for } R \neq (0, \dots, 0). \quad (9.2.14)$$

This allows us to write (cf. Ref. [76])

$$Z^{\text{PG}}(\beta) = \left(\lambda_{(0, \dots, 0)}(\beta) \right)^V + \mathcal{O}\left((\beta/N_c)^V\right), \quad (9.2.15)$$

i.e., the strong-coupling expansion of Z factorizes in the volume whenever $n_{\max} < V$.⁶

Note that this factorization in V for $n_{\max} < V$ corresponds to the simplification which we made in the tensor formulation by removing those links from the initial tensor for which the factorization condition is not satisfied at the adjacent plaquettes.

⁵Note that the expression differs compared to Ref. [73] due to a different definition of the gauge action.

⁶From this, we also observe that the partition function of pure-gauge theory in the strong-coupling expansion factorizes in the limit $V \rightarrow \infty$ to arbitrary order in β .

n	0	1	2	3	4	5	6
$\lambda^{(n)}$	1	0	$\frac{1}{36}$	$\frac{1}{648}$	$\frac{1}{2592}$	$\frac{1}{31104}$	$\frac{13}{3359232}$
f_n^{PG}	0	0	$\frac{1}{36}$	$\frac{1}{648}$	0	$-\frac{1}{93312}$	$-\frac{1}{1119744}$

Table 9.1: Coefficients of the strong-coupling expansion of $\lambda_{(0,\dots,0)}$ and of the free energy f^{PG} for different orders n for the SU(3) pure-gauge theory in two dimensions. The stated values of $\lambda^{(n)}$ and f_n^{PG} are valid as long as $V > n$.

Next, we calculate the strong-coupling expansion of the coefficient

$$\lambda_{(0,\dots,0)}(\beta) = \sum_{n=0}^{n_{\max}} \lambda^{(n)} \beta^n, \quad (9.2.16)$$

which allows us to calculate the coefficients f_n^{PG} of the strong-coupling expansion of the free energy

$$\frac{\ln Z^{\text{PG}}(\beta)}{V} = \sum_{n=0}^{\infty} f_n^{\text{PG}} \beta^n \quad (9.2.17)$$

as well. Table 9.1 shows the resulting coefficients $\lambda^{(n)}$ and f_n^{PG} up to order six for $N_c = 3$. The values of $\lambda^{(n)}$ and f_n^{PG} are valid as long as $V > n$. In this case, the coefficients of the strong-coupling expansion of

$$Z^{\text{PG}}(\beta) = \sum_{n=0}^{\infty} Z_n^{\text{PG}} \beta^n \quad (9.2.18)$$

up to order three for $N_c = 3$ read

$$Z_0^{\text{PG}} = 1, \quad Z_1^{\text{PG}} = 0, \quad Z_2^{\text{PG}} = V/36, \quad Z_3^{\text{PG}} = V/648. \quad (9.2.19)$$

9.3 Properties of the partition function and observables

Even though the full partition function (4.1.3) can not be evaluated analytically for arbitrary volume, one can give its structure with respect to the chemical potentials and the masses, allowing for several analytical results in different limits.

At first, let us consider $N_f = 1$, for an arbitrary number of dimensions including all orders of β . For brevity, we omit the flavor index. A general number of flavors N_f will be considered later in this section. By using time-reversal invariance of the partition function one obtains the general result

$$Z(\mu, m, \beta) = Z(-\mu, m, \beta). \quad (9.3.1)$$

Using the strong-coupling expansions

$$Z(\mu, m, \beta) = \sum_{n=0}^{\infty} Z_n(\mu, m) \beta^n \quad \text{and} \quad \frac{\ln Z(\mu, m, \beta)}{V} = \sum_{n=0}^{\infty} f_n(\mu, m) \beta^n \quad (9.3.2)$$

with coefficients $f_n(\mu, m)$ that are determined explicitly in terms of the coefficients $Z_n(\mu, m)$ in App. G, we note that the coefficients Z_n and f_n are even functions in μ as well.

Consequently, the chiral condensate is an even function in μ as well, i.e.,

$$\Sigma(\mu, m, \beta) = \Sigma(-\mu, m, \beta). \quad (9.3.3)$$

Also the coefficients $\Sigma_n(\mu, m)$ of the strong-coupling expansion

$$\Sigma(\mu, m, \beta) = \sum_{n=0}^{\infty} \Sigma_n(\mu, m) \beta^n \quad \text{with} \quad \Sigma_n(\mu, m) = \frac{\partial f_n(\mu, m)}{\partial m} \quad (9.3.4)$$

are even functions in μ .

Since the quark-number density $\rho(\mu, m, \beta)$ is proportional to the derivative of the free energy with respect to μ , it is an odd function in μ , i.e.,

$$\rho(-\mu, m, \beta) = -\rho(\mu, m, \beta), \quad (9.3.5)$$

and therefore all coefficients $\rho_n(\mu, m)$ of the strong-coupling expansion

$$\rho(\mu, m, \beta) = \sum_{n=0}^{\infty} \rho_n(\mu, m) \beta^n \quad \text{with} \quad \rho_n(\mu, m) = \frac{\partial f_n(\mu, m)}{\partial \mu} \quad (9.3.6)$$

are odd functions in μ . In particular, we stress

$$\rho_n(\mu, m) \rightarrow 0, \quad \text{for } \mu \rightarrow 0, \text{ for all } n \in \mathbb{N}_0. \quad (9.3.7)$$

Next, we want to quantize the general dependence of the chemical potential. To do so, we first note that at infinite coupling, the μ dependence of Z comes from closed baryon and antibaryon loops winding around the torus in the time direction [6, 7]. Every winding contributes a factor of $e^{\pm N_c L_t \mu}$ to a configuration. Since baryon loops are self-avoiding the maximum winding number is V_s . For nonzero β the baryon loops are locally disturbed by plaquettes (more precisely, one or more quark lines in the baryon go around the plaquettes), but the functional form of the μ dependence remains the same. We can formalize this in the equation

$$Z(\mu, m, \beta) = \sum_{k=-V_s}^{V_s} \bar{p}_k(m, \beta) \exp(k N_c L_t \mu), \quad (9.3.8)$$

where L_t is the temporal extent of the d -dimensional space-time lattice⁷ and V_s is the spatial volume.

The functions $\bar{p}_k(m, \beta)$ are polynomials in m and, due to the defining equation (4.1.3), series in β . As consequence of the symmetry (9.3.1), they satisfy

$$\bar{p}_k(m, \beta) = \bar{p}_{-k}(m, \beta). \quad (9.3.9)$$

This property allows us to write

$$Z(\mu, m, \beta) = \bar{p}_0(m, \beta) + \sum_{k=1}^{V_s} \bar{p}_k(m, \beta) (\exp(k N_c L_t \mu) + \exp(-k N_c L_t \mu)), \quad (9.3.10)$$

⁷Note that in Sec. 2, the temporal extent was denoted by N_T .

which can be simplified to

$$Z(\mu, m, \beta) = \sum_{k=0}^{V_s} p_k(m, \beta) \cosh(k N_c L_t \mu), \quad (9.3.11)$$

where we introduced

$$p_k(m, \beta) = (2 - \delta_{k,0}) \bar{p}_k(m, \beta). \quad (9.3.12)$$

We now perform a strong-coupling expansion of the coefficients p_k ,

$$p_k(m, \beta) = \sum_{n=0}^{\infty} p_{k,n}(m) \beta^n, \quad (9.3.13)$$

and obtain the explicit μ dependence of the coefficients Z_n ,

$$Z_n(\mu, m) = \sum_{k=0}^{V_s} p_{k,n}(m) \cosh(k N_c L_t \mu). \quad (9.3.14)$$

9.3.1 Asymptotic results for large chemical potential

In this section, we derive analytical results in the limit of large chemical potential μ . Intermediate results of this calculation are used to verify some of the terms in Eqs. (E.1a) - (E.1e).

In the limit $|\mu| \rightarrow \infty$ the sum (9.3.11) is dominated by the $k = V_s$ term, for which we show in App. H that

$$p_{V_s}(m, \beta) = 2Z^{\text{PG}}(\beta) \quad (9.3.15)$$

independent of m , where Z^{PG} is the partition function of the pure-gauge theory (with the same lattice size and number of colors), whose strong-coupling expansion reads⁸

$$Z^{\text{PG}}(\beta) = \sum_{n=0}^{\infty} Z_n^{\text{PG}} \beta^n. \quad (9.3.16)$$

This leads to

$$p_{V_s,n}(m) = p_{V_s,n} = 2Z_n^{\text{PG}}, \quad (9.3.17)$$

and with $V = L_t V_s$ we obtain in the $|\mu| \rightarrow \infty$ limit

$$Z_n(\mu, m) \rightarrow 2Z_n^{\text{PG}} \cosh(N_c V \mu) \rightarrow Z_n^{\text{PG}} \exp(N_c V |\mu|), \quad (9.3.18)$$

provided that $p_{V_s,n} \neq 0$. This yields

$$\frac{Z_n(\mu, m)}{Z_0(\mu, m)} \rightarrow \frac{Z_n^{\text{PG}}}{Z_0^{\text{PG}}} \quad (9.3.19)$$

⁸For the example of two dimensions and $N_c = 3$, the coefficients Z_n^{PG} are stated explicitly in Eq. (9.2.19) for $V > n$.

for any $n \geq 0$.⁹ The ratio is well-defined since $Z_0^{\text{PG}} = \text{vol}(\text{SU}(N_c))^{V_d} = 1$. Note that for $n \geq 1$ the coefficients f_n are polynomials in $\bar{Z}_n = Z_n/Z_0$, see App. G. We thus obtain

$$f_n(\mu, m) \rightarrow f_n^{\text{PG}} + N_c |\mu| \delta_{n,0} \quad \text{for } |\mu| \rightarrow \infty, \quad (9.3.20)$$

where the Kronecker delta takes into account the different definition of f_0 . From (9.3.20) we can obtain the behavior of the chiral condensate and the quark-number density in this limit. Since there is no dependence on m , all coefficients of the chiral condensate approach zero, i.e.,¹⁰

$$\Sigma_n(\mu, m) \rightarrow 0 \quad \text{for } |\mu| \rightarrow \infty. \quad (9.3.21)$$

For the quark-number density we obtain

$$\rho_n(\mu, m) \rightarrow \pm N_c \delta_{n,0} \quad \text{for } \mu \rightarrow \pm \infty. \quad (9.3.22)$$

As a consistency check, we can compare Eq. (9.3.17) with the analytical expressions of the functions $Z_n(m, \mu)$ given in Eqs. (E.1a) - (E.1e). They are calculated for $d = 2$, $V_s = 2$, $L_t = 2$, and $N_c = 3$. The corresponding values of $p_{V_s, n}$ are shown in Table 9.2. The coefficients Z_n^{PG} are calculated by expanding the exact formula (9.2.8) in β .¹¹ The obtained value $p_{2, n}$ coincides with the prefactor of the term $\cosh(12\mu)$ in $Z_n(m, \mu)$ for all stated values of n .

n	0	1	2	3	4
$p_{2, n}$	2	0	$\frac{2}{9}$	$\frac{1}{81}$	$\frac{325}{26244}$

Table 9.2: Values of $p_{V_s, n}$ for $V_s = 2$, $L_t = 2$, and $N_c = 3$ up to $n = 4$.

9.3.2 Asymptotic results for large mass

In this section, we will derive analytical results for large mass. Intermediate results of this calculation are used to verify some of the terms in Eqs. (E.1a) - (E.1e).

For $N_f = 1$ and large m , the partition function can be approximated, up to $\mathcal{O}(\beta^{L_t})$ corrections, as

$$Z_{\text{factorized}}(\mu, m, \beta) = \left((2m)^{N_c L_t} + 2 \cosh(N_c L_t \mu) \right)^{V_s} Z^{\text{PG}}(\beta). \quad (9.3.23)$$

To derive this result, we start from (9.3.11), and determine, for given k , the leading-order term in $p_k(m, \beta)$ for large m . For powers of β smaller than L_t , this term is obtained by configurations which have k baryon loops (or k antibaryon loops) with all links in the time direction and winding around the torus once, and all other lattice sites occupied by mass terms. For all other configurations, e.g., with mesonic contributions, or baryon lines in a spatial direction,

⁹If $Z_n^{\text{PG}} = 0$, which happens in particular for $n = 1$, we obtain $Z_n(\mu, m) \rightarrow p_{k, n}(m) \cosh(k N_c L_t \mu)$, where $k < V_s$ is the largest value of k for which $p_{k, n}(m) \neq 0$. In this case $\cosh(k N_c L_t \mu) / \cosh(V_s N_c L_t \mu) \rightarrow 0 = Z_n^{\text{PG}}$ for $k < V_s$, and hence the relation (9.3.19) still applies.

¹⁰Strictly speaking, the limit $|\mu| \rightarrow \infty$ and the derivative with respect to m may not commute. However, a more careful analysis in which we first take the derivative and then the limit confirms (9.3.21).

¹¹Since the condition $V > n$ is not satisfied for $p_{2, 4}$ for $L_t = 2$, the partition function does not factorize in the volume at this order (cf. Sec. 9.2).

or combinations of baryon and antibaryon windings, the number of lattice sites that can be occupied by mass terms is reduced. Therefore such configurations do not contribute to p_k at leading order in m . For a configuration which has only mass terms and k baryon (or antibaryon) loops the gauge field only occurs in contributions from the gauge action, see (H.2). For powers of β larger than or equal to L_t , there are also contributions from other configurations when $k \neq 0$ and $k \neq V_s$,¹² but we do not include them explicitly since we are only interested in an expansion of Z up to $\beta^{n_{\max}}$, and typically $n_{\max} < L_t$. We therefore have for large m

$$p_k(m, \beta) \rightarrow (2 - \delta_{k0}) \binom{V_s}{k} (2m)^{N_c L_t (V_s - k)} Z^{\text{PG}}(\beta) + \mathcal{O}(\beta^{L_t}) \quad (9.3.24)$$

since there are $\binom{V_s}{k}$ possibilities to distribute the k parallel baryon loops on the lattice. The factor of $2 - \delta_{k0}$ comes from $e^{k N_c L_t \mu} + e^{-k N_c L_t \mu} = 2 \cosh(k N_c L_t \mu)$, which sums configurations with k baryon loops and configurations with k antibaryon loops. This results in

$$Z(\mu, m, \beta) \rightarrow Z_{\text{asymptotic}}(\mu, m, \beta) \equiv Z^{\text{PG}}(\beta) \sum_{k=0}^{V_s} \binom{V_s}{k} (2m)^{N_c L_t (V_s - k)} (2 - \delta_{k0}) \cosh(k N_c L_t \mu) \quad (9.3.25)$$

up to $\mathcal{O}(\beta^{L_t})$ corrections. Let us now show that the large- m behavior of (9.3.25) is reproduced by (9.3.23). Expanding (9.3.23) gives

$$Z_{\text{factorized}}(\mu, m, \beta) = Z^{\text{PG}}(\beta) \sum_{k=0}^{V_s} \binom{V_s}{k} (2m)^{N_c L_t (V_s - k)} 2^k \cosh^k(N_c L_t \mu). \quad (9.3.26)$$

Defining $x = N_c L_t \mu$ and using

$$\cosh^k(x) = \frac{1}{2^k} (e^x + e^{-x})^k = \frac{1}{2^k} \sum_{j=0}^k \binom{k}{j} e^{(k-2j)x} = \frac{1}{2^k} \sum_{j=0}^k \binom{k}{j} \cosh((k-2j)x), \quad (9.3.27)$$

we can now organize the terms in $Z_{\text{factorized}}$ in complete analogy to (9.3.11). Next, we approximate $Z_{\text{factorized}}$ for large m . For the coefficient of $\cosh(k N_c L_t \mu)$ in $Z_{\text{factorized}}$, we again want to keep only the leading-order term for large m . This term is obtained by setting $j = 0$ or $j = k$ in (9.3.27). Hence, to obtain the large- m behavior of $Z_{\text{factorized}}$, we can replace $2^k \cosh^k(x) \rightarrow (2 - \delta_{k,0}) \cosh(kx)$. This shows that $Z_{\text{factorized}} \rightarrow Z_{\text{asymptotic}}$ for large m .

For the chiral condensate $\Sigma = (\partial_m \ln Z)/V$ in the large- m limit and up to $\mathcal{O}(\beta^{L_t})$ corrections we obtain from (9.3.23)

$$\Sigma(\mu, m) = \bar{b} \frac{2 \sqrt{\left(\frac{1}{2} (2m)^{N_c L_t}\right)^2 - 1}}{\frac{1}{2} (2m)^{N_c L_t} + \cosh(N_c L_t \mu)} = \bar{b} (\tanh(\bar{a}(\mu + \bar{\mu}^c)) - \tanh(\bar{a}(\mu - \bar{\mu}^c))), \quad (9.3.28)$$

where we used the general formula

$$\frac{2 \sinh(2x)}{\cosh(2x) + \cosh(2z)} = \tanh(x+z) + \tanh(x-z). \quad (9.3.29)$$

¹²An example for such a configuration, for $N_c = 3$ and an $L_t \times 2$ lattice, would be two quark loops in the time direction with spatial coordinate 1, one quark loop in the time direction with spatial coordinate 2, and plaquette occupation 1 for all plaquettes $(x, 12)$ starting at sites $x = (t, 1)$ with $t = 1, \dots, L_t$.

and defined

$$\bar{b} \equiv \frac{N_c(2m)^{N_c L_t - 1}}{2\sqrt{\left(\frac{1}{2}(2m)^{N_c L_t}\right)^2 - 1}} \rightarrow \frac{N_c}{2m}, \quad \bar{a} \equiv \frac{N_c L_t}{2}, \quad \bar{\mu}^c \equiv \frac{1}{N_c L_t} \operatorname{arccosh}\left(\frac{1}{2}(2m)^{N_c L_t}\right) \rightarrow \ln(2m). \quad (9.3.30)$$

The parameters $\bar{\mu}^c$ and $\bar{a}$ are referred to as critical chemical potential and transition sharpness, respectively.

Similarly, for the quark-number density $\rho = (\partial_\mu \ln Z)/V$ we obtain in the large- m limit and up to $\mathcal{O}(\beta^{L_t})$ corrections

$$\rho(\mu, m) = \frac{N_c}{2} \frac{2 \sinh(N_c L_t \mu)}{\frac{1}{2}(2m)^{N_c L_t} + \cosh(N_c L_t \mu)} = \frac{N_c}{2} (\tanh(\bar{a}(\mu + \bar{\mu}^c)) + \tanh(\bar{a}(\mu - \bar{\mu}^c))), \quad (9.3.31)$$

where we again used (9.3.29). Note that critical chemical potential and transition sharpness are identical for both observables.

In particular, from Eq. (9.3.23) we conclude $Z \rightarrow (2m)^{N_c V} Z^{\text{PG}}(\beta)$ for large mass¹³. For the coefficients of the strong-coupling expansion we can conclude

$$f_n(\mu, m) \rightarrow f_n^{\text{PG}} + N_c \ln(2m) \delta_{n,0}, \quad \text{for } m \rightarrow \infty, \quad (9.3.32)$$

$$\Sigma_n(\mu, m) \rightarrow \frac{N_c}{m} \delta_{n,0}, \quad \text{for } m \rightarrow \infty, \quad (9.3.33)$$

$$\rho_n(\mu, m, \beta) \rightarrow 0, \quad \text{for } m \rightarrow \infty. \quad (9.3.34)$$

The intermediate result (9.3.24) can be compared with some of the terms in Eqs. (E.1), which serves as a consistency check. To do so, we expand the functions p_k in m and β , i.e.,

$$p_k(m, \beta) = \sum_{n=0}^{\infty} \sum_{j=0}^{w_k} p_{k,n,j} \beta^n m^j, \quad (9.3.35)$$

where we defined $w_k \equiv N_c(V - kL_t)$. w_k is an upper bound for the degree of the polynomial, depending on the volume and the winding number k .

With the use of (9.3.24) we obtain

$$p_{k,n,w_k} = (2 - \delta_{k0}) \binom{V_s}{k} 2^{N_c L_t (V_s - k)} Z_n^{\text{PG}}, \quad (9.3.36)$$

which is valid as long as $n < L_t$. For $k \in \{0, V_s\}$, the formula is valid for arbitrary n . Note that $k = 0$ corresponds to those configurations, where every site contains N_c mass terms. The associated coefficients read

$$p_{0,n,N_c V} = 2^{N_c V} Z_n^{\text{PG}}. \quad (9.3.37)$$

The values of p_{k,n,w_k} for $d = 2$, $V_s = 2$, $L_t = 2$, and $N_c = 3$ are shown in Table 9.3. The coefficients Z_n^{PG} are calculated by expanding the exact formula (9.2.8) in β .¹⁴ All values are

n	0	1	2	3	4
$p_{0,n,12}$	4096	0	$\frac{4096}{9}$	$\frac{2048}{81}$	$\frac{166400}{6561}$
$p_{1,n,6}$	256	0	-	-	-
$p_{2,n,0}$	2	0	$\frac{2}{9}$	$\frac{1}{81}$	$\frac{325}{26244}$

Table 9.3: Values of p_{k,n,w_k} for $V_s = 2$, $L_t = 2$, and $N_c = 3$ up to $n = 4$.

consistent with Eq. (E.1). Note that with the use of Eq. (9.3.36), $p_{1,n,6}$ can only be calculated for $n \leq 1$. Also note that the coefficients $p_{2,n,0}$ are consistent with Table 9.2, since we have $p_{V_s,n} = p_{V_s,n,0}$.

9.3.3 Even mass exponents

Another observation for $N_f = 1$ is that

$$p_{k,n,j} = 0 \quad \text{for all } k, n, \text{ for all } j \text{ odd}, \quad (9.3.38)$$

which is explained by the following argument. We start with a configuration that has N_c mass terms at every site. Such a configuration is proportional to $m^{N_c V}$, i.e., has even exponent.¹⁵ In order to get other configurations, where the Grassmann conditions in (A.23) are still fulfilled at every site¹⁶, we apply the following algorithm: We remove one mass term from a first site and place one hopping term in an arbitrary direction at an adjacent link. Consequently, the Grassmann conditions are no longer satisfied at the two sites adjacent to the chosen link, referred to as first and second site. At the second site, we remove one mass term as well and introduce one hopping term in an arbitrary direction. Therefore, the Grassmann conditions at the second site are satisfied again, but we introduced a third site, where the Grassmann conditions are not satisfied. We can repeat the procedure of removing one mass at the last site and adding a hopping term in arbitrary direction. The set of chosen sites forms a line on the lattice, where the Grassmann conditions are not satisfied only at the first and last site of the line. Only if the first and last site of the line coincide, the Grassmann conditions are satisfied again on all sites of the lattice. In this case the line forms a closed loop. As any closed loop on the lattice contains an even number of links and an even number of sites¹⁷, we have applied this algorithm an even number of times and therefore removed an even number of masses from the lattice. Therefore, the mass exponent is again even. By doing this algorithm several times (starting again at an arbitrary site), we can create all relevant arrangements of hopping terms. All of the corresponding valid configurations have an even mass exponent. Note that for a valid configuration, we also need to fulfill the gauge condition at every link. This is also influenced by placing plaquettes on the

¹³We assume that μ does not go to infinity

¹⁴Since the condition $V > n$ is not satisfied for $p_{0,4,12}$ and $p_{2,4,0}$ for $L_t = 2$, the partition function does not factorize in the volume at this order (cf. Sec. 9.2).

¹⁵Recall that for staggered quarks the number of lattice sites needs to be even.

¹⁶The $SU(N_c)$ gauge condition may not be satisfied for the moment.

¹⁷Loops that wind around the lattice contain an even number of links and sites, since spatial and temporal extent of the lattice are always even when staggered quarks are used.

lattice, but this does not change the mass exponent. Due to Eq. (9.3.38) all mass exponents in Eqs. (E.1) are even.

9.3.4 Generalization to arbitrary number of flavors

We now generalize some of the main statements for a general number of flavors N_f . For a given valid configuration, let k_f be the number of hopping terms of flavor f in forward time direction minus the number of hopping terms of flavor f in backward time direction, i.e., the net number of hopping terms in forward time direction. For a compact notation, we define the net-number vector $\vec{k} = (k_1, k_2, \dots, k_{N_f})$, the mass vector $\vec{m} = (m_1, m_2, \dots, m_{N_f})$, and the chemical potential vector $\vec{\mu} = (\mu_1, \mu_2, \dots, \mu_{N_f})$.

Using this, the partition function for an arbitrary number of flavors N_f , in arbitrary dimensions, for general L_t and V_s , and including all orders of β , takes the general form

$$Z(\vec{\mu}, \vec{m}, \beta) = \sum_{k_1, \dots, k_{N_f} = -N_c V_s}^{N_c V_s} \bar{p}_{\vec{k}}(\vec{m}, \beta) \exp(L_t \vec{k} \cdot \vec{\mu}), \quad (9.3.39)$$

where the functions $\bar{p}_{\vec{k}}(\vec{m}, \beta)$ are polynomials in m_f and series β . From the $SU(N_c)$ condition it follows that $\bar{p}_{\vec{k}}$ is only nonzero when $\sum_{f=1}^{N_f} k_f = 0 \pmod{N_c}$. Therefore, k_1 is a multiple of N_c only for $N_f = 1$, in general.

From time-reversal invariance we get the condition

$$Z(\vec{\mu}, \vec{m}, \beta) = Z(-\vec{\mu}, \vec{m}, \beta) \quad (9.3.40)$$

and therefore

$$\bar{p}_{\vec{k}}(\vec{m}, \beta) = \bar{p}_{-\vec{k}}(\vec{m}, \beta). \quad (9.3.41)$$

Consequently, the chiral condensate is symmetric under $\vec{\mu} \rightarrow -\vec{\mu}$, i.e.,

$$\Sigma^{(f)}(-\vec{\mu}, \vec{m}, \beta) = \Sigma^{(f)}(\vec{\mu}, \vec{m}, \beta) \quad (9.3.42)$$

while the quark-number density is antisymmetric, i.e.,

$$\rho^{(f)}(-\vec{\mu}, \vec{m}, \beta) = -\rho^{(f)}(\vec{\mu}, \vec{m}, \beta). \quad (9.3.43)$$

As a generalization of Eq. (9.3.15), we find

$$\bar{p}_{(k_1, \dots, \pm N_c V_s, \dots, k_{N_f})}(\vec{m}, \beta) = \bar{p}_{\vec{k} \setminus k_f}(\vec{m} \setminus m_f, \beta) \quad (9.3.44)$$

with

$$\vec{v} \setminus v_f \equiv (v_1, \dots, v_{f-1}, v_{f+1}, \dots, v_{N_f}), \quad (9.3.45)$$

which is explained as follows. Since $k_f = \pm N_c V_s$, there are V_s baryon loops of flavor f that wind around the lattice in time direction. This results in a gauge-field-independent background, and we obtain the corresponding polynomial of the theory with $N_f - 1$ flavors.

Consequently, for $k_f = \pm N_c V_s$ for all f ¹⁸, we obtain

$$\bar{p}_{N_c V_s(\pm 1, \dots, \pm 1)}(\vec{m}, \beta) = Z^{\text{PG}}(\beta), \quad (9.3.46)$$

by applying Eq. (9.3.44) for every flavor separately and using Eq. (9.3.15)¹⁹. In this case, every flavor results in a separate constant background.

The property (9.3.44) is relevant for the limit $\mu_f \rightarrow \pm\infty$ ²⁰, as in this case any relevant configuration contains V_s baryon lines of flavor f that wind around the lattice and extend in time direction only. We get²¹

$$Z(\vec{\mu}, \vec{m}, \beta) \rightarrow \exp(N_c V |\mu_f|) Z^{(N_f-1)}(\vec{\mu} \setminus \mu_f, \vec{m} \setminus m_f, \beta), \quad (9.3.47)$$

where $Z^{(N_f-1)}$ is the partition function of the theory involving $N_f - 1$ flavors. It follows²²

$$f_n(\vec{\mu}, \vec{m}) \rightarrow f_n^{(N_f-1)}(\vec{\mu} \setminus \mu_f, \vec{m} \setminus m_f) + N_c |\mu_f| \delta_{n,0} \quad \text{for } \mu_f \rightarrow \pm\infty, \quad (9.3.48)$$

$$\Sigma_n^{(f)}(\vec{\mu}, \vec{m}) \rightarrow 0 \quad \text{for } \mu_f \rightarrow \pm\infty, \quad (9.3.49)$$

$$\rho_n^{(f)}(\vec{\mu}, \vec{m}) \rightarrow \pm N_c \delta_{n,0} \quad \text{for } \mu_f \rightarrow \pm\infty. \quad (9.3.50)$$

In the limit $\mu_1, \dots, \mu_{N_f} \rightarrow \infty$ we get

$$Z(\vec{\mu}, \vec{m}, \beta) \rightarrow \exp\left(N_c V \sum_{f=1}^{N_f} |\mu_f|\right) Z^{\text{PG}}(\beta), \quad (9.3.51)$$

where we applied Eq. (9.3.47) for every flavor separately. For the next steps, it is convenient to expand the functions $p_{\vec{k}}$ in β and in the masses m_f , i.e.,

$$\bar{p}_{\vec{k}}(\vec{m}, \beta) = \sum_{n=0}^{\infty} \sum_{j_1=0}^{w_{k_1}} \cdots \sum_{j_{N_f}=0}^{w_{k_{N_f}}} \bar{p}_{\vec{k}, n, \vec{j}} \beta^n \prod_{f=1}^{N_f} m_f^{j_f}, \quad (9.3.52)$$

where we used the upper bound $w_{k_f} = N_c(V - k_f L_t)$.

The following observation will be relevant for large masses. For $k_f = 0$ and $j_f = N_c V$ we obtain

$$\bar{p}_{\vec{k}, n, \vec{j}} = 2^{N_c V} \bar{p}_{\vec{k}/k_f, n, \vec{j}/j_f}. \quad (9.3.53)$$

In this case, there are N_c mass terms of flavor f at every site of the lattice, which results in a constant background. If $k_f = 0$ and $j_f = N_c V$ is satisfied for all flavors f , we find

$$\bar{p}_{(0, \dots, 0), n, N_c V(1, \dots, 1)} = 2^{N_f N_c V} Z_n^{\text{PG}}, \quad (9.3.54)$$

¹⁸The sign can differ for different flavors.

¹⁹Note that $\bar{p}_{(N_c V_s)} = \bar{p}_{V_s} = \frac{1}{2} p_{V_s}$, where we used Eq. (9.3.12) for $V_s > 0$.

²⁰We assume that m_f does not go to infinity.

²¹Like before, the condition $Z_n(\vec{\mu}, \vec{m}, \beta) \rightarrow \exp(N_c V |\mu_f|) Z_n^{(N_f-1)}(\vec{\mu} \setminus \mu_f, \vec{m} \setminus m_f, \beta)$ for the n -th coefficient of the strong-coupling expansion in β is correct only for $Z_n^{(N_f-1)}(\vec{\mu} \setminus \mu_f, \vec{m} \setminus m_f, \beta) \neq 0$. Otherwise the expression is dominated by a smaller exponent. This effect is again irrelevant for the limit of the ratio $Z_n(\vec{\mu}, \vec{m}, \beta)/Z_0(\vec{\mu}, \vec{m}, \beta)$. Therefore it is irrelevant for the observables.

²²The results (9.3.49) and (9.3.50) are confirmed by a more careful analysis in which we first take the derivative and then the limit.

by applying Eq. (9.3.53) for every flavor separately and using Eq. (9.3.37).

The property (9.3.53) is relevant for the limit $m_f \rightarrow \infty$ ²³, as in this case any relevant configuration contains N_c mass terms of flavor f at every site. We get²⁴

$$Z(\vec{\mu}, \vec{m}, \beta) \rightarrow (2m_f)^{N_c V} Z^{(N_f-1)}(\vec{\mu}/\mu_f, \vec{m}/m_f, \beta), \quad (9.3.55)$$

where $Z^{(N_f-1)}$ is again the partition function of the theory involving $N_f - 1$ flavors. For the coefficients of the strong-coupling expansion we conclude²⁵

$$f_n(\vec{\mu}, \vec{m}) \rightarrow f_n^{(N_f-1)}(\vec{\mu}/\mu_f, \vec{m}/m_f) + N_c \ln(2m_f) \delta_{n,0}, \quad \text{for } m_f \rightarrow \infty, \quad (9.3.56)$$

$$\Sigma_n^{(f)}(\vec{\mu}, \vec{m}) \rightarrow \frac{N_c}{m_f} \delta_{n,0}, \quad \text{for } m_f \rightarrow \infty, \quad (9.3.57)$$

$$\rho_n^{(f)}(\vec{\mu}, \vec{m}, \beta) \rightarrow 0, \quad \text{for } m_f \rightarrow \infty. \quad (9.3.58)$$

For the limit $m_1, \dots, m_{N_f} \rightarrow \infty$ we can apply Eq. (9.3.55) for every flavor separately to obtain

$$Z(\vec{\mu}, \vec{m}, \beta) \rightarrow Z^{\text{PG}}(\beta) \prod_{f=1}^{N_f} (2m_f)^{N_c V}. \quad (9.3.59)$$

Next, we observe $\bar{p}_{\vec{k}, n, \vec{j}} = 0$ for all $\vec{k}$, n , when j_f is odd for at least one flavor f . This statement follows from the algorithm described below Eq. (9.3.38), which can be applied for every flavor separately. For the Eqs. (E.2), we conclude that there are no odd mass exponents, which serves as a consistency check.

By using the symmetry (9.3.40), we can express the partition function in terms of the hyperbolic cosines also for an arbitrary number of flavors. To do so, we write

$$\begin{aligned} Z(\vec{\mu}, \vec{m}, \beta) &= \frac{1}{2} (Z(\vec{\mu}, \vec{m}, \beta) + Z(-\vec{\mu}, \vec{m}, \beta)) \\ &= \sum_{k_1, \dots, k_{N_f} = -N_c V_s}^{N_c V_s} \bar{p}_{\vec{k}}(\vec{m}, \beta) \frac{1}{2} \left(\exp(L_t \vec{k} \cdot \vec{\mu}) + \exp(-L_t \vec{k} \cdot \vec{\mu}) \right), \end{aligned} \quad (9.3.60)$$

resulting in

$$Z(\vec{\mu}, \vec{m}, \beta) = \sum_{k_1, \dots, k_{N_f} = -N_c V_s}^{N_c V_s} \bar{p}_{\vec{k}}(\vec{m}, \beta) \cosh(L_t \vec{k} \cdot \vec{\mu}). \quad (9.3.61)$$

Finally, the symmetry of permuting flavors is formalized in the condition²⁶

$$Z(\vec{\mu}, \vec{m}, \beta) = Z(\vec{\mu}_\pi, \vec{m}_\pi, \beta) \quad (9.3.62)$$

²³we assume that μ_f does not go to infinity.

²⁴Again, $Z_n(\vec{m}, \vec{\mu}, \beta)$ is dominated by a smaller exponent than $N_c V$, when $Z_n^{(N_f-1)}(\vec{m}/m_f, \vec{\mu}/\mu_f, \beta) = 0$.

²⁵Again, a more careful analysis in which we first take the derivative and then the limit confirms this result.

²⁶For $N_f = 2$ and $m_1 = m_2$ this implies that the partition function is invariant under $\mu_1 \rightarrow -\mu_1$, where $\mu_1 \equiv (\mu_1 - \mu_2)/2$.

with

$$\vec{v}_\pi \equiv (v_{\pi_1}, v_{\pi_2}, \dots, v_{\pi_{N_f}}), \quad (9.3.63)$$

for a general permutation $\pi \in S_{N_f}$. From this we conclude by comparison of coefficients

$$\bar{p}_k(\vec{m}, \beta) = \bar{p}_{\vec{k}_\pi}(\vec{m}_\pi, \beta) \quad (9.3.64)$$

For the chiral condensate with respect to the mass m_f ,

$$\Sigma^{(f)}(\vec{\mu}, \vec{m}, \beta) \equiv \frac{1}{V} \frac{\partial \ln Z(\vec{\mu}, \vec{m}, \beta)}{\partial m_f}, \quad (9.3.65)$$

we conclude

$$\Sigma^{(f)}(\vec{\mu}, \vec{m}, \beta) = \Sigma^{(\pi_f)}(\vec{\mu}_\pi, \vec{m}_\pi, \beta) \quad (9.3.66)$$

and similarly, for the quark-number density with respect to the chemical potential μ_f ,

$$\rho^{(f)}(\vec{\mu}, \vec{m}, \beta) \equiv \frac{1}{V} \frac{\partial \ln Z(\vec{\mu}, \vec{m}, \beta)}{\partial \mu_f}, \quad (9.3.67)$$

we obtain

$$\rho^{(f)}(\vec{\mu}, \vec{m}, \beta) = \rho^{(\pi_f)}(\vec{\mu}_\pi, \vec{m}_\pi, \beta). \quad (9.3.68)$$

9.4 Interpolation between small and large μ for large L_t

In this section we derive interpolating formulas for $N_f = 1$ that turn out to describe the data very well in many cases, with limitations discussed below. From (9.3.11) and (9.3.15) we see that for $|\mu| \rightarrow \infty$ we have $Z = 2Z^{\text{PG}}(\beta) \cosh(N_c V \mu (1 + \mathcal{O}(e^{-N_c L_t |\mu|})))$. On the other hand, in the vicinity of $\mu = 0$ and for large L_t we expect Z to remain approximately constant, i.e., $Z(\mu, m, \beta) \approx Z(\mu = 0, m, \beta) \approx p_0(m, \beta)$ for the following reason. The coefficients p_k in (9.3.11) have to decrease exponentially with k since p_k contains powers of m and combinatorial factors for distributing mass terms and mesons on up to $(V_s - k)L_t$ lattice sites, i.e., $\ln p_k \sim (V_s - k)L_t$. Hence, $Z(\mu)$ will deviate from its value at $\mu = 0$ only when $|\mu|$ is sufficiently large for $\cosh(k N_c L_t \mu)$ to compensate the exponential suppression of p_k . We expect this to happen in the vicinity of the phase transition.

A simple interpolation between the two regimes is provided by²⁷

$$Z_{\text{IP}}(\mu, m, \beta) = c(m, \beta) + 2Z^{\text{PG}}(\beta) \cosh(N_c V \mu) \quad \text{with} \quad c(m, \beta) = Z(\mu = 0, m, \beta) - 2Z^{\text{PG}}(\beta). \quad (9.4.1)$$

For the chiral condensate and the quark-number density this yields the symmetrized (see Footnote 33) and antisymmetrized (see Footnote 34) tanh model, respectively,

$$\Sigma_{\text{IP}}(\mu, m, \beta) = b_{\text{IP}} (\tanh(a_{\text{IP}}(\mu + \mu_{\text{IP}}^c)) - \tanh(a_{\text{IP}}(\mu - \mu_{\text{IP}}^c))), \quad (9.4.2)$$

$$\rho_{\text{IP}}(\mu, m, \beta) = \frac{N_c}{2} (\tanh(a_{\text{IP}}(\mu + \mu_{\text{IP}}^c)) + \tanh(a_{\text{IP}}(\mu - \mu_{\text{IP}}^c))), \quad (9.4.3)$$

²⁷This interpolation formula is rooted in an idea originally proposed by Robert Lohmayer.

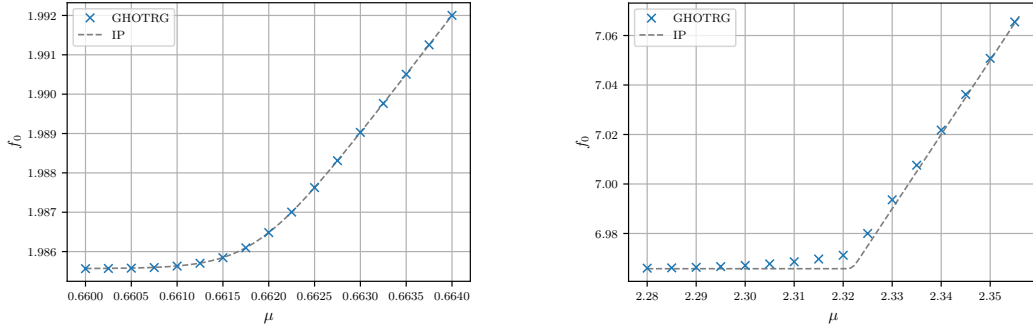


Figure 9.1: Numerical results for $f_0 = (\ln Z_0)/V$ obtained from GHOTRG on a 32×32 lattice with bond dimension $D = 60$ and interpolating formula (9.4.1) for $m = 0.5$ (left) and $m = 5$ (right).

where we used (9.3.29) and defined

$$b_{\text{IP}} \equiv \frac{\partial_m c(m, \beta)}{4V Z^{\text{PG}}(\beta) \sqrt{\left(\frac{c(m, \beta)}{2Z^{\text{PG}}(\beta)}\right)^2 - 1}} \approx \frac{\partial_m c(m, \beta)}{2V c(m, \beta)} \approx \frac{1}{2} \Sigma(\mu = 0, m, \beta), \quad (9.4.4)$$

$$a_{\text{IP}} \equiv \frac{N_c V}{2}, \quad (9.4.5)$$

$$\mu_{\text{IP}}^c \equiv \frac{1}{N_c V} \text{arccosh} \frac{c(m, \beta)}{2Z^{\text{PG}}(\beta)} \approx \frac{1}{N_c V} \ln \frac{c(m, \beta)}{Z^{\text{PG}}(\beta)}. \quad (9.4.6)$$

For the approximations we used $c \gg 2Z^{\text{PG}}$, which follows from $Z(\mu = 0) \approx p_0 \gg p_{V_s} = 2Z^{\text{PG}}$. We compute $Z(\mu = 0, m, \beta)$ numerically using OS-GHOTRG and use $Z^{\text{PG}}(\beta) = 1 + \frac{1}{2}d(d-1)V(1 + \delta_{N_c, 2})(\beta/2N_c)^2 + \mathcal{O}(\beta^3)$, which is obtained from the definition of the gauge action [27] using the results of [30]. As usual, we expand b_{IP} and μ_{IP}^c in β , with coefficients $b_{\text{IP}, n}$ and $\mu_{\text{IP}, n}^c$, respectively.

In Fig. 9.1 we compare numerical results for Z obtained in the infinite-coupling limit from the GHOTRG method with the interpolating formula (9.4.1). We see that the interpolation works very well away from the phase transition, but an interesting question is whether it also provides a reasonable approximation to Z in the vicinity of the transition. In Fig. 9.1 we see that the answer to this question depends on the parameters of the system such as the mass. We can distinguish two scenarios for the transition:

- (a) For any $\mu < \mu^c$, the $k = 0$ term is the dominating term in (9.3.11), while for any $\mu > \mu^c$ the $k = V_s$ term dominates. In this case we expect the interpolation to provide a reasonable approximation to the true Z also across the phase transition.
- (b) In the vicinity of μ^c , terms other than $k = 0$ or $k = V_s$ make the largest contribution to (9.3.11). In this case, keeping only the $k = 0$ or $k = V_s$ term will not provide a reasonable approximation to Z when $\mu \approx \mu^c$. We then expect the phase transition to be less sharp with $a < a_{\text{IP}}$.

Which scenario occurs in practice depends on m , L , and β through the precise values of $p_k(m, \beta)$. For large m , we have scenario (b) since in this case $e^{\mu^c} \approx 2m$ (see Sec. 9.3.2) and

$p_k \sim (2m)^{N_c L_t (V_s - k)}$, which implies that in (9.3.11) all terms are of the same order of magnitude. (At $\mu = \mu^c$, the dominating terms are those with $k \approx V_s/2$ because of the binomial coefficient in (9.3.25).) This is reflected in $a = N_c L_t/2$ for large m (see Sec. 9.3.2). Hence, the actual transition is less sharp than that of Σ_{IP} and ρ_{IP} , where $a_{\text{IP}} = N_c V/2$. This explains why the interpolation deviates from the data for the larger mass in Fig. 9.1 in the vicinity of the phase transition. Nevertheless, we will see in Fig. 10.15 (left) that μ_{IP}^c still provides an excellent approximation to μ^c for all m (for lattice size 16×16). Note that for large m we have $c(m, \beta) \rightarrow (2m)^{N_c V}$, resulting in $\mu_{\text{IP}}^c \rightarrow \ln(2m)$, which is consistent with μ^c of Sec. 9.3.2.

9.5 Applicability of the strong-coupling expansion

In this subsection we discuss the applicability of the strong-coupling expansion, both in the vicinity of and away from a phase transition. The arguments are generic, but to be specific we consider the phase transition as a function of μ . This section is based on an idea by Robert Lohmayer.

Away from a phase transition we use the expansions (9.3.4) and (9.3.6) and expect Σ_n/Σ_0 and ρ_n/ρ_0 to be independent of V for large V . Therefore in this case the range in β for which the expansions provide a reasonable approximation is independent of V .

We now analyze the behavior of a generic observable Ω in the vicinity of a phase transition. To this end we consider a function $\omega(x)$ with a smooth crossover at $x = 0$, e.g., $\omega(x) = \tanh(x)$ or $\arctan(x)$, or $\text{erf}(x)$, and derivatives $\omega^{(n)} = \partial^n \omega / \partial x^n$. Near the phase transition we assume

$$\Omega(\mu) \sim \omega(a(\mu - \mu^c)), \quad (9.5.1)$$

where μ^c is the critical chemical potential and we have omitted a possible constant offset. We assume $a \sim V^\gamma$ ($\gamma > 0$), which implies that the transition becomes steeper with increasing volume. For the time being we assume for simplicity that the width of the transition, which is proportional to $1/a$, is independent of β and that the critical chemical potential is linear in β , i.e., $\mu^c(\beta) = \mu_0^c + \beta \mu_1^c$. The strong-coupling expansion of Ω at fixed μ then reads

$$\Omega(\mu, \beta) = \sum_{n=0}^{n_{\max}} \Omega_n(\mu) \beta^n + \mathcal{O}(\beta^{n_{\max}+1}) \quad \text{with} \quad \Omega_n(\mu) \sim \frac{1}{n!} (-a \mu_1^c)^n \omega^{(n)}(a(\mu - \mu_0^c)) \sim a^n. \quad (9.5.2)$$

This implies that the range of β for which (9.5.2) gives reasonable approximations is limited by $\beta_{\max} \sim 1/a \sim V^{-\gamma}$.

Let us now drop the assumptions on the β dependence of a and μ^c and assume polynomial approximations in β with coefficients a_n and μ_n^c , respectively. We expect the a_n/a_0 and μ_n^c to have finite limits for large V . Therefore the polynomial approximations for $a(\beta)$ and $\mu_c(\beta)$ should give reasonable results in a range of β that is independent of V . In the same range of β , we expect (9.5.1) to give a reasonable approximation to the observable.

The difference in the ranges of β ($\sim 1/V^\gamma$ or independent of V) can be understood more generally. Near the phase transition, the function $\Omega(\mu)$ is very steep (with slope $\sim a$), while the inverse function $\mu(\Omega)$ is nearly flat (with slope $\sim 1/a$). As a function of β , the steep function is shifted horizontally since μ^c is a function of β , see Fig. 10.12 below. At fixed μ close to μ^c ,

the observable therefore changes drastically as a function of β . Therefore it is not a good idea to expand $\Omega(\mu, \beta)$ in β at fixed μ close to μ^c , since the coefficients of this expansion would need to be large, see Fig. 10.11 below. In contrast, expanding μ^c in β corresponds to an expansion of the inverse function $\mu(\Omega, \beta)$ in β . Since $\mu(\Omega)$ is nearly flat we expect only small corrections.

The arguments made in this section will be confirmed by the numerical results for the observables Σ and ρ presented in the following. However, we will see that a_n/a_0 for $n \geq 1$ and μ_n^c for $n \geq 2$ become increasingly difficult to determine for increasing V . Since we expect a_n/a_0 and μ_n^c to remain finite for large V , one possible way to deal with this issue is to combine a_0 determined from large V with a_n/a_0 and μ_n^c determined from small V to still obtain reasonable approximations for the observable even for large V , again in a range of β that is independent of V .

10 Numerical Results

We now turn to the discussion of the numerical results. The tensor-network method was implemented in C++, in particular using code elements written by Jacques Bloch for Ref. [23], as well as tensor libraries developed by him. All of the results presented here are for the two-dimensional case. The study of numerical results in higher dimensions remains a task for future investigations. In all of the results, the spatial extent of the lattice equals the temporal extent, denoted by L and we used $N_c = 3$. We use the SuperQ truncation method of Sec. 7.5 unless stated otherwise.

We start with a short comparison of SuperQ and the Xie et al. truncation procedure (cf. Sec. 3.2.2) near a phase transition for different volumes in strong-coupling QCD. Afterwards, we show the results for the pure-gauge theory obtained with OS-HOTRG up to order six in the strong-coupling expansion. The results are compared with the analytical solution, derived in Sec. 9.2. In Sec. 10.3.1, we compare two natural expansions for the free energy and the observables, when the coefficients of the strong-coupling expansion are calculated up to order $n_{\max}$ for $N_f = 1$. The coefficients corresponding to a higher order than $n_{\max}$ differ in both expansions, which has an influence on the applicable range in β of the expansion. We motivate the OS-GHOTRG method by comparing the results with GHOTRG results as well as Monte Carlo results for 2×2 and 8×8 lattices. Afterwards we show detailed results obtained from the OS-GHOTRG method for various quark masses, chemical potentials, lattice sizes, and bond dimensions for $N_c = 3$ and $N_f = 1, 2$ up to $n_{\max} = 3$.

10.1 Comparison of SuperQ and the Xie et al. truncation

A comparison of SuperQ and the Xie et al. truncation procedure has already been presented in the literature [29]. SuperQ improves the local truncation error and reduces the computational cost, since the number of necessary singular value decompositions is halved compared to the Xie et al. truncation method (see Sec. 3.2.2).

However, for calculations near a phase transition and for larger lattices, the total truncation

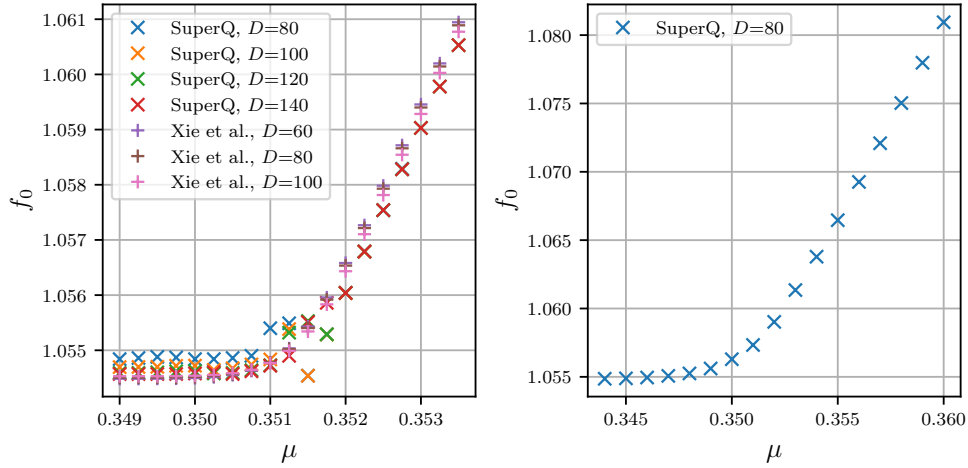


Figure 10.2: Free energy of strong-coupling QCD as a function of μ for $m = 0.1$ on a 32×32 (left) and a 16×16 (right) lattice. In the left plot, we compare the tensor-network results obtained using the SuperQ and the Xie et al. procedure. Near the phase transition, the results using the Xie et al. procedure are more stable and already converged for smaller bond dimension. In the right plot we observe that the results using SuperQ are stable even for bond dimension $D = 80$. For $D \geq 80$, the result is independent of the bond dimension and the chosen truncation method (not shown).

error may not be reduced by using SuperQ. In the left plot of Fig. 10.2 we compare SuperQ with the Xie et al. truncation procedure using the free energy of strong-coupling QCD for one flavor on a 32×32 lattice versus the chemical potential μ . For both methods, the results are shown for several values of the bond dimension D . We observe that near the phase transition the Xie et al. procedure is preferred as this yields more stable results that converge faster with respect to the bond dimension. Discontinuous results emerge for small masses and large lattices near the phase transition only when SuperQ is used. As a comparison, we show the result for the free energy on a 16×16 , i.e., smaller, lattice in the right plot of Fig. 10.2. In this case, also SuperQ yields stable results for bond dimension $D = 80$. For the 16×16 lattice with $m = 0.1$, the result is independent of the bond dimension and the chosen truncation method (not shown) for $D \geq 80$. When SuperQ is used, the final trace of the tensor network may involve the subtraction of two dominating, nearly equal tensor entries which can lead to cancellation effects and numerical instability.

10.2 Pure-gauge theory

We now consider pure-gauge theory in two dimensions with $N_c = 3$ and compare the results obtained from OS-HOTRG with the exact results that are calculated in Sec. 9.2. Figure 10.3 shows a convergence study of the coefficients f_2 , f_3 , f_5 and f_6 of the free energy $f = \sum_n f_n \beta^n$, obtained from OS-HOTRG with respect to the bond dimension D for $n_{\max} = 6$. Additionally, the exact values as given in Table 9.1 are shown in red. For the OS-HOTRG approach we use the translational invariance as explained in Sec. 6, i.e., all plaquettes except the origin plaquette

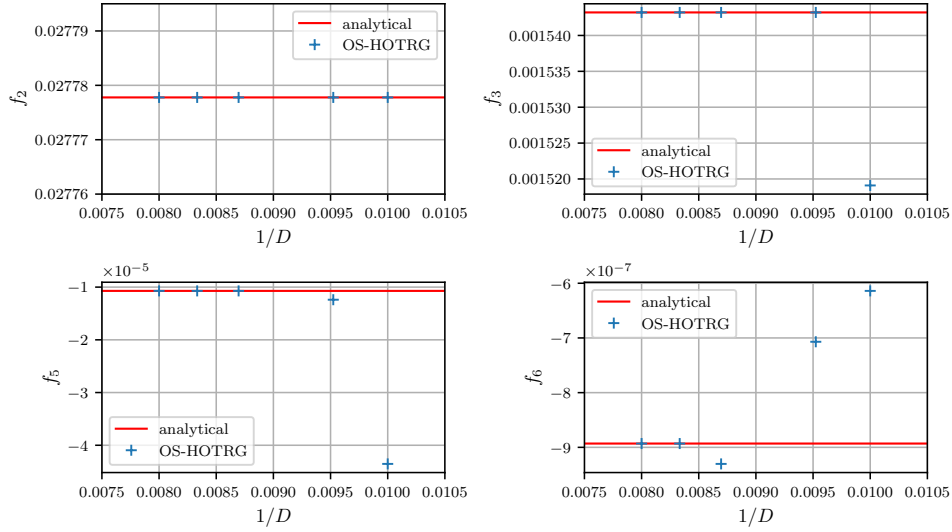


Figure 10.3: Nonzero coefficients f_i for pure-gauge theory with lattice size $L = 8$. We compare the results gained from OS-HOTRG with the exact results given in Table 9.1. Note that the axes for f_5 and f_6 are scaled by the factors 10^{-5} and 10^{-7} , respectively.

are restricted to occupation 3 on the initial lattice.²⁸ We observe that for increasing bond dimension D the OS-HOTRG results approach the exact coefficients to very high accuracy. For conciseness, we omit the plots showing the convergence of the vanishing coefficients f_0 , f_1 , and f_4 . We find very precise agreement for these as well when $D \geq 100$.

10.3 One flavor

10.3.1 Comparing different expansions

In the following we present tensor-network results for the strong-coupling expansion of two-dimensional QCD with three colors and one flavor, i.e., $N_c = 3$ and $N_f = 1$. We omit the flavor index in the following.

As we eventually want to use the initial tensor, constructed in Part III, as input to tensor-network methods on large lattices, it is enlightening to first investigate the exact results that can be obtained with this tensor on a 2×2 lattice. Starting from the initial tensor $T_x K_x$ in Eq. (5.2.15) and using the tensor (5.2.16) to order $n_{\max}$ by imposing the constraints (5.2.17) and (5.2.18), the partition function on a 2×2 lattice can be computed exactly to order $n_{\max}$, by exhaustively enumerating all terms in the partition function Eq. (5.2.15) in an efficient way (see Sec. 9.1). Each such term corresponds to a specific order n in β , and all nonzero contributions satisfy $n \leq n_{\max}$.²⁹ After gathering the terms order by order we can compute all the coefficients

²⁸Since we are using the OS-HOTRG procedure, we could also apply the more sophisticated method for translational invariance explained in Sec. 8 for $n_{\max} = 6$. However, as Fig. 10.3 only serves as a proof of concept, we use the simpler version.

²⁹On a 2×2 lattice no terms with order $n > n_{\max}$ appear in the exact computation of Z . Indeed, for this lattice size the edge occupation numbers of the four tensors are shared in such a way that if the initial tensors have

Z_n separately and write the partition function as

$$Z^{(n_{\max})} = \sum_{n=0}^{n_{\max}} Z_n \beta^n. \quad (10.3.1)$$

Analytical expressions for Z_n as functions of m and μ are given in App. E up to $n = 4$.

To compare tensor results with those of Monte Carlo simulations we will compute the chiral condensate,

$$\Sigma = 2 \langle \bar{\psi} \psi \rangle = \frac{1}{V} \frac{\partial \ln Z}{\partial m}, \quad (10.3.2)$$

where the factor of 2 is due to the rescaling of the Grassmann variables in Eq. (2.6.30). A natural way to perform a tensor-network computation of the chiral condensate is to rewrite Eq. (10.3.2) as

$$\Sigma_A^{(n_{\max})} = \frac{1}{V} \frac{\partial \ln Z^{(n_{\max})}}{\partial m} = \frac{1}{V Z^{(n_{\max})}} \frac{\partial Z^{(n_{\max})}}{\partial m}, \quad (10.3.3)$$

where we approximated Z by the polynomial $Z^{(n_{\max})}$. The derivative can be computed analytically for the 2×2 lattice as the coefficients Z_n are known as explicit functions of m and μ , and Eq. (10.3.3) becomes a ratio of polynomials in β .

Alternatively, we can expand Eq. (10.3.3) or, equivalently, $\ln Z$ in β and truncate that expansion to order $n_{\max}$ when computing the observable. We therefore expand the free energy, i.e.,

$$\frac{\ln Z(\beta)}{V} = \sum_{n=0}^{n_{\max}} f_n \beta^n + \mathcal{O}(\beta^{n_{\max}+1}). \quad (10.3.4)$$

As shown in App. G, up to order three we obtain with $\bar{Z}_n \equiv Z_n/Z_0$

$$\frac{\ln Z(\beta)}{V} = \frac{\ln Z_0}{V} + \frac{\beta}{V} \bar{Z}_1 + \frac{\beta^2}{V} \left(\bar{Z}_2 - \frac{1}{2} \bar{Z}_1^2 \right) + \frac{\beta^3}{V} \left(\bar{Z}_3 - \bar{Z}_1 \bar{Z}_2 + \frac{1}{3} \bar{Z}_1^3 \right) + \mathcal{O}(\beta^4). \quad (10.3.5)$$

Taking the derivative with respect to m yields

$$\Sigma_B^{(n_{\max})} = \sum_{n=0}^{n_{\max}} \Sigma_n \beta^n \quad \text{with} \quad \Sigma_n = \frac{\partial f_n}{\partial m}. \quad (10.3.6)$$

In the following we will refer to the alternatives (10.3.3) and (10.3.6) as “expansion A” and “expansion B”, respectively.

In Fig. 10.4 we compare both expansions for the chiral condensate, with $n_{\max} = 1, 2, 3, 4$, to the full Monte Carlo results for $m = 0.1$. In the left plot, in which $\mu = 0$, we observe that expansion A only agrees with the Monte Carlo results for small values of β and deviates quite substantially as β increases.³⁰ For expansion B, we observe a better agreement with the Monte Carlo results over a larger range in β . In the right plot we see that expansion B outperforms expansion A also for $\mu \neq 0$. It thus seems preferable to expand the observable itself in β and truncate it to the same order $n_{\max}$ as the partition function in Eq. (10.3.1).

maximal order $n_{\max}$, Eq. (5.2.18) ensures that all nonzero contributions to the fully contracted tensor network have at most order $n_{\max}$ as well. On larger lattices higher order contributions with $n > n_{\max}$ also contribute to Z .

³⁰These results are in agreement with those of Ref. [13], where $\Sigma/2$ is plotted.

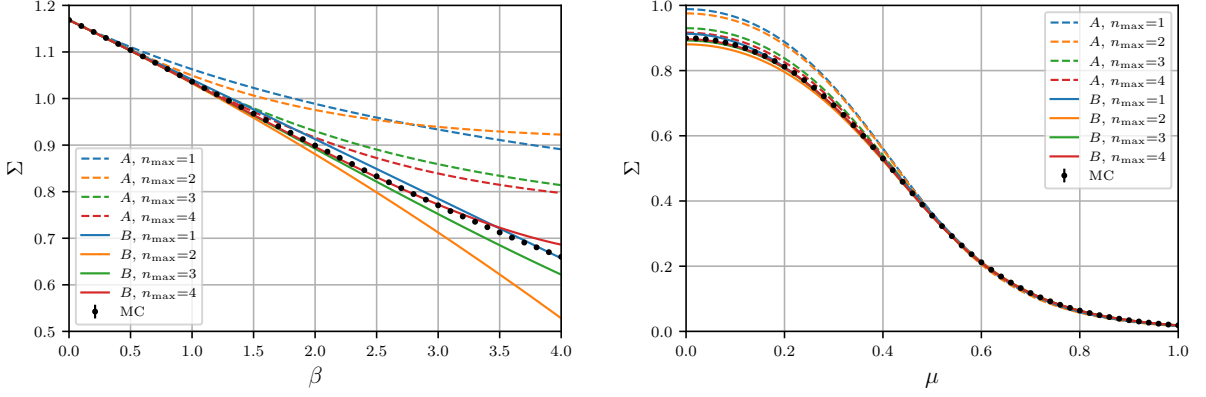


Figure 10.4: Chiral condensate as a function of β (left, for $\mu = 0$) and μ (right, for $\beta = 2$) for $m = 0.1$ on a 2×2 lattice. We compare Monte Carlo (MC) results with the two tensor-network expansions $\Sigma_A^{(n_{\max})}$ and $\Sigma_B^{(n_{\max})}$ of Eqs. (10.3.3) and (10.3.6), respectively. The MC results for $\mu \neq 0$ were obtained through phase-quenched reweighting, with an average phase factor between 0.81 and 1.

Let us compare expansions A and B on theoretical grounds for arbitrary volume. For the free-energy density to be finite in the infinite-volume limit, the coefficients f_n in (10.3.4) must remain finite as $V \rightarrow \infty$. Looking at Eq. (10.3.5) term by term, this implies that the ratios $\bar{Z}_n$ must behave like V^n for large V , and that cancellations must take place in every order $n \geq 2$ to ensure that the f_n are finite for large V . Let us now consider expansion A , which uses $\ln Z^{(n_{\max})}$ in Eq. (10.3.3). If we were to expand $\ln Z^{(n_{\max})}$ in β (which is not actually done in expansion A) we would also obtain terms of order β^n with $n > n_{\max}$, but these higher-order contributions would be incomplete and would have the wrong volume dependence because the cancellations just discussed cannot take place. Moreover, because $\bar{Z}_n \propto V^n$ for large V , Eq. (10.3.1) with Z_0 factored out is then effectively an expansion in powers of βV . For large βV , $Z^{(n_{\max})}$ is dominated by the term with $n = n_{\max}$. In this case, Eq. (10.3.3) with $Z^{(n_{\max})} \approx Z_0 \bar{Z}_{n_{\max}} \beta^{n_{\max}}$ results in $\Sigma_A^{(n_{\max})} \approx \Sigma_0 + \frac{1}{V} \partial_m \ln \bar{Z}_{n_{\max}}$, which is independent of β . In summary, we expect expansion B to perform better than expansion A when compared to Monte Carlo results for large βV (and β not too large), which is consistent with our results for $L = 2$ (see Fig. 10.4) and $L = 8$ (see Fig. 10.5).

In order to use Eq. (10.3.6) for larger lattices we need to be able to compute the coefficients Z_n of the partition function obtained from the tensor network for such lattices. This motivates the usage of OS-GHOTRG, since when applying the standard GHOTRG method to contract the tensor network numerically, we obtain Z (and thus also $\ln Z$ and the thermodynamical observables) as a numerical function of β , and not as an explicit expansion in β . We also tried to determine the expansion coefficients Z_n by fitting tensor data to polynomials in β , but this attempt did not work well on larger lattices since the tensor data have limited precision for realistic bond dimensions, and since on larger lattices the completely contracted tensor network will contain many incomplete higher-order contributions ($n > n_{\max}$) even though the initial tensor is truncated to order $n_{\max}$.³¹

³¹Similar to expansion A , the incomplete higher-order terms in the standard GHOTRG method also have a wrong volume dependence, and Σ tends to a plateau, see Fig. 10.5.

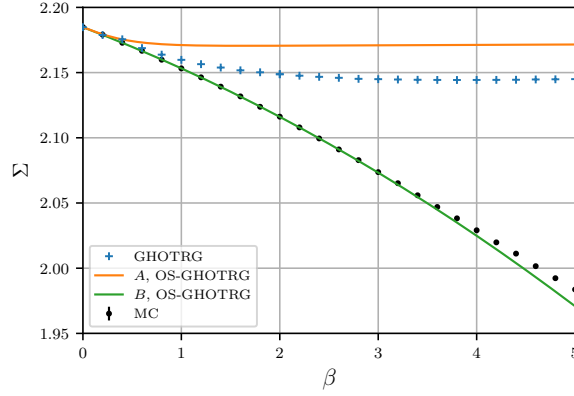


Figure 10.5: Chiral condensate versus β for $L = 8$, $m = 0.5$, $\mu = 0$ using the standard GHOTRG method with initial tensor truncated to $n_{\max} = 2$ and the OS-GHOTRG method with expansions A and B up to order $n_{\max} = 2$, compared to Monte Carlo data. For expansion A we observe the plateau for large βV derived earlier. As expected, expansion B agrees well with the Monte Carlo data over a larger range in β .

Once the coefficients Z_n are known, we can compute expansion coefficients for all thermodynamical observables. In Fig. 10.5 we show that the results for $L = 8$ obtained with expansion B using the OS-GHOTRG method agree much better with Monte Carlo data than those from the standard GHOTRG method and from expansion A using the OS-GHOTRG method (the latter two are identical for $L = 2$).

More detailed data for the comparison of GHOTRG and OS-GHOTRG (using expansion B) with Monte Carlo are shown in Fig. 10.6 for $m = 0.5$, $\mu = 0$, and $L = 8$. In the left plot, we show the GHOTRG results for $n_{\max} = 1$ (blue) and $n_{\max} = 2$ (orange). The β -range over which the expansion agrees with the Monte Carlo data increases for $n_{\max} = 2$ only marginally compared to $n_{\max} = 1$, which is in fact not visible on this scale. The higher-order coefficients are changed significantly in this case, causing a deviation from the Monte Carlo data that is larger for $n_{\max} = 2$ for $\beta \geq 0.25$. One possible modification that has a significant influence on the values of the higher order coefficients is the use of the translational invariance as explained in Sec. 6. Using this for $n_{\max} = 1$, there are no configurations considered whose order in β is greater than one and therefore the higher-order coefficients are introduced solely by taking the logarithm. The GHOTRG results using translational invariance are shown also in the left plot of Fig. 10.6 for $n_{\max} = 1$ (green) and $n_{\max} = 2$ (red). Even though many of the higher-order configurations are removed by this approach, there is no systematic improvement by using the translational invariance for GHOTRG. The deviation from the Monte Carlo data is larger (smaller) when translational invariance is used for $n_{\max} = 1$ ($n_{\max} = 2$). All GHOTRG results (with and without translational invariance) are shown for two different bond dimensions, and all of the data points are identical for both bond dimensions on the considered scale. If the values were still dependent on the bond dimension, one could rule out that a convergence with respect to the bond dimension is already reached. The right plot of Fig. 10.6 shows again the Monte Carlo data, but in comparison with expansion B , obtained from OS-GHOTRG (including the translational invariance as explained in Sec. 8), which is preferred to expansion A . In this case, all higher-order coefficients are set to zero. We observe that the β -range over which the expansion agrees with the Monte Carlo data increases significantly for $n_{\max} = 2$ compared to $n_{\max} = 1$. While the expansion yields accurate results up to $\beta = 1$ for $n_{\max} = 1$, the expansion is accurate

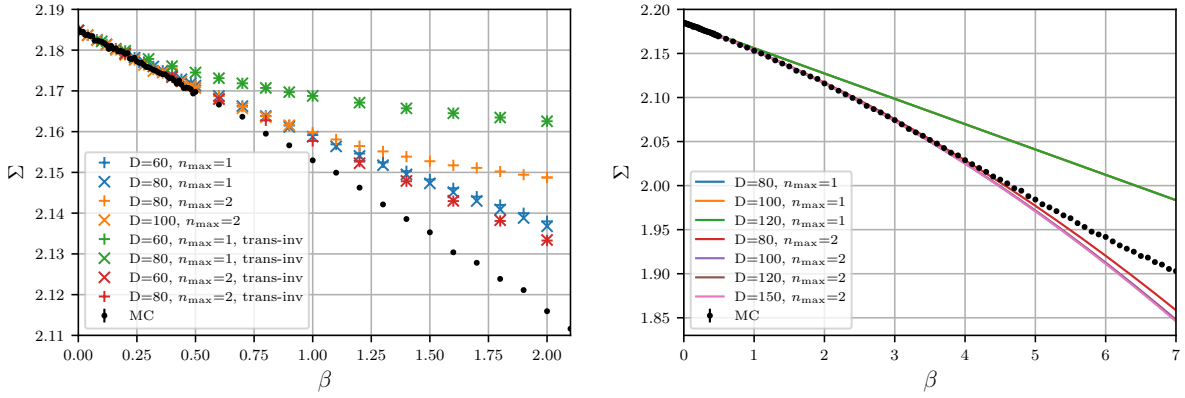


Figure 10.6: Chiral condensate as a function of β for $m = 0.5$ with $\mu = 0$ on an 8×8 lattice. We compare the Monte Carlo results with the tensor-network results obtained from GHOTRG (left) and OS-GHOTRG using expansion B (right). In the left plot we also compare the results with the ones, where additional translational invariance is used. All results are shown for different bond dimensions D . In the right plot, all lines corresponding to $n_{\max} = 1$ lie on top of each other. Note that the range in β is significantly increased compared to the left plot.

up to $\beta = 4.5$ for $n_{\max} = 2$. To study the convergence with respect to the bond dimension, the OS-GHOTRG results are shown for various bond dimensions. All lines corresponding to $n_{\max} = 1$ lie on top of each other, which is a notable sign for convergence. The situation is similar for $n_{\max} = 2$, where a deviation is visible only for the smallest bond dimension, i.e., $D = 80$.

Additionally to Fig. 10.5, Fig. 10.6 strongly motivates the usage of OS-GHOTRG. We stress out that while for GHOTRG a separate tensor network has to be evaluated for every single data point, i.e., β -value, an analytic function of β is obtained when OS-GHOTRG is used. This reduces the computational cost for many applications significantly and supports the usage of OS-GHOTRG further.

10.3.2 Validation of OS-GHOTRG with analytical results and Monte Carlo data

As a consistency check for the correct implementation of the OS-GHOTRG method we first compare tensor data with analytical results for $L = 2$, see (E.1), and with Monte Carlo data for $L = 8$.

Fig. 10.7 shows the coefficients f_n , Σ_n , and ρ_n as a function of the chemical potential for $n_{\max} = 2$ and for different masses. The tensor data are obtained for bond dimension $D = 230$. The data agree nicely with the analytical results (solid lines) which were calculated using the exact coefficients of the partition function, given in (E.1). For $n_{\max} = 3$ and $L = 2$, bond dimensions for which all coefficients are converged were not feasible for small masses and small chemical potentials. The convergence of the coefficients of the free energy with respect to the bond dimension is shown in Fig. 10.8 for $m = 0.1$ (left) and $m = 0.5$ (right). While for $n_{\max} = 2$, all coefficients f_n are converged for $D \geq 120$, for $n_{\max} = 3$ the tensor results are not converged for $D \leq 220$. We observe that the error of the coefficients is generally smaller for the larger mass,

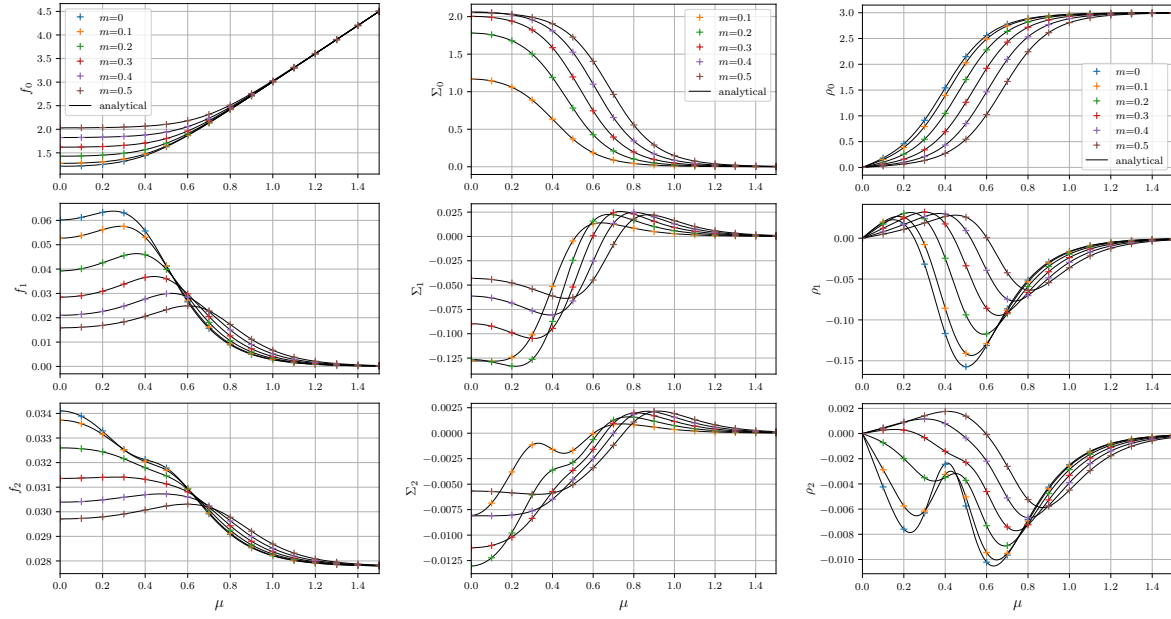


Figure 10.7: Coefficients of the free energy (left), chiral condensate (middle) and the quark-number density (right) up to order 2 as a function of μ for different values of m on a 2×2 lattice. We compare the analytical results obtained from App. E to the OS-GHOTRG results with bond dimension $D = 230$. For Σ we do not show the $m = 0$ result, which is trivially zero on a finite lattice. The large- μ limit of all coefficients is consistent with the analytical results (9.3.20), (9.3.21), and (9.3.22). As stated in Eq. (9.3.7), the quark-number density vanishes for $\mu = 0$ to all orders.

especially for the coefficient f_3 . For large m at fixed μ or large μ at fixed m , the convergence with respect to D can be reached even for $n_{\max} = 3$ (not shown in the figure).

In Fig. 10.9 we show the quark-number density as a function of μ near the phase transition for $L = 8$ and various values of β and compare results obtained from the OS-GHOTRG method using the expansion (9.3.6) and from Monte Carlo simulations using reweighting. For larger lattices the reweighted Monte Carlo method can no longer be used as the sign problem becomes unmanageable. We observe that the OS-GHOTRG results agree very well with the Monte Carlo data for small β . For larger β we observe deviations as expected, since the validity of the expansion is limited as explained in Sec. 9.5.

10.3.3 Behavior of the singular values

We now analyze the fall-off of the squared singular values of $R \star R$, see (8.2.2), in different blocking steps for $n_{\max} = 2$ and various values of m and μ , still for $N_f = 1$. In Fig. 10.10 (left) we show all singular values obtained in the first blocking step for $\mu = 0$ and different values of m . We observe that the singular values decrease faster for larger mass. Similarly, in Fig. 10.10 (middle) we show the singular values for $m = 0.1$ and different values of μ . The singular values decrease faster for larger μ . We do not observe a change in the behavior of the singular values close to the phase transition, which is located at about $\mu_c = 0.4$ for $L = 2$ and $m = 0.1$, as

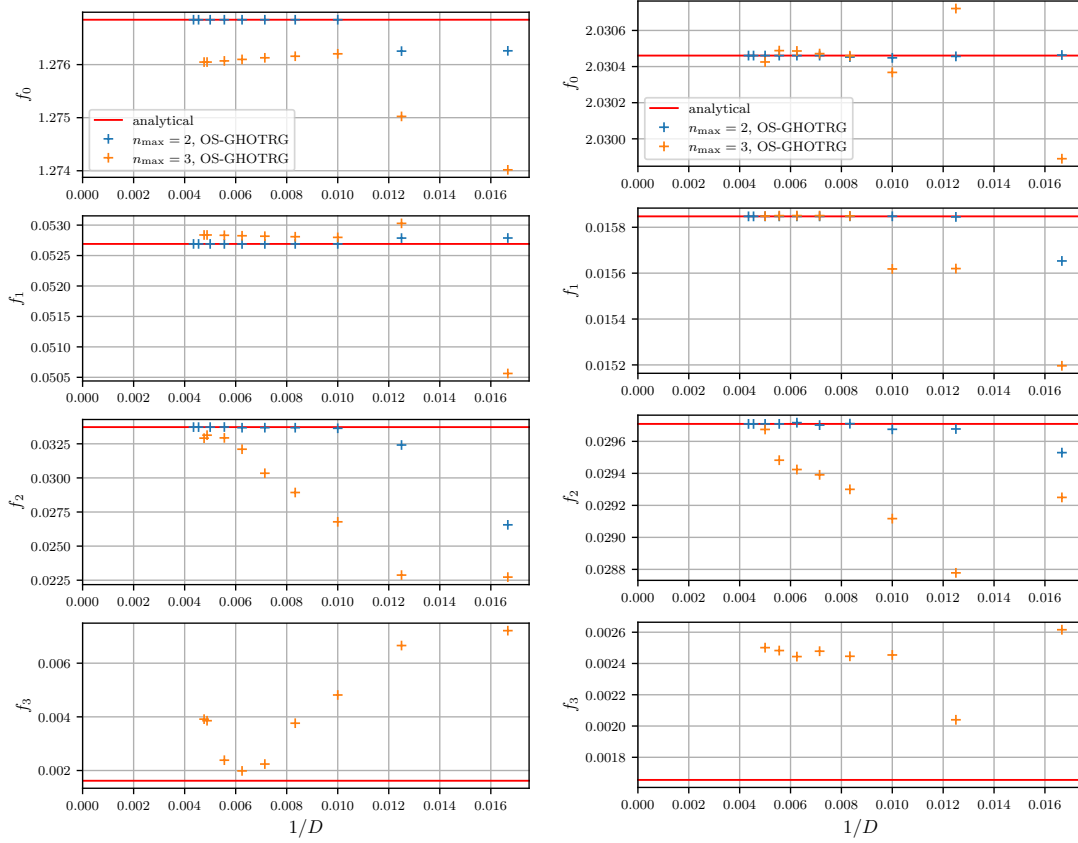


Figure 10.8: Coefficients of the free energy in the strong-coupling expansion versus the inverse bond dimension for $L = 2$ and $\mu = 0$. For $n_{\max} = 2$ the exact values are obtained for $D \geq 120$. For $n_{\max} = 3$, neither the data for $m = 0.1$ (left) nor the ones for $m = 0.5$ (right) are converged with respect to the bond dimension for $D \leq 220$. The discrepancy between exact results and tensor data is generally smaller for $m = 0.5$.

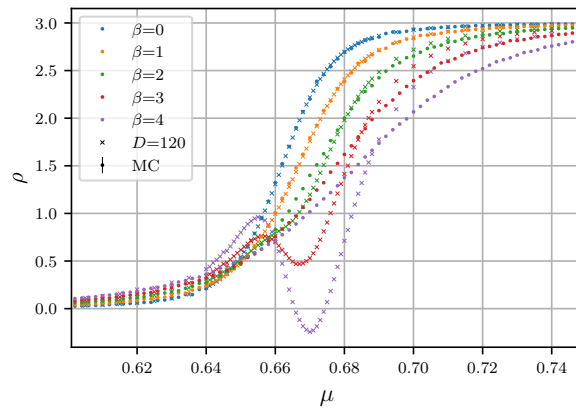


Figure 10.9: Quark-number density ρ as a function of μ for $L = 8$, $m = 0.5$, and various values of β obtained from Monte Carlo simulations (circles) and from the OS-GHOTRG method with bond dimension $D = 120$ (crosses). The systematic deviations for larger β are expected, see Sec. 9.5.

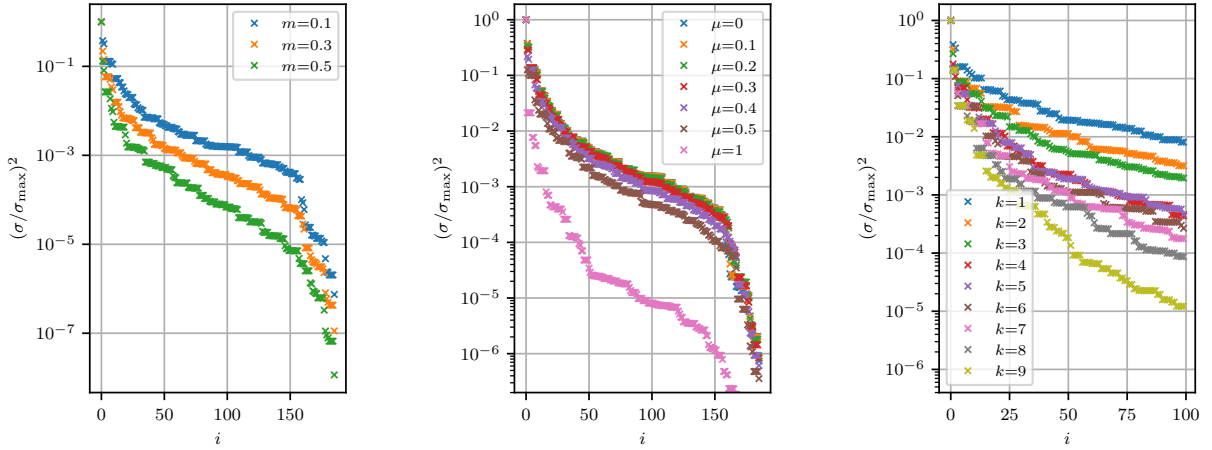


Figure 10.10: Relative size of the squared singular values of $R \star R$, enumerated by i . Left: First blocking step for various values of m at $\mu = 0$. Middle: First blocking step for various values of μ at $m = 0.1$. Right: Successive blocking steps k at $m = 0.1$ and $\mu = 0$. All three plots are for $n_{\max} = 2$. We observe many plateaus of nearly degenerated singular values.

can be seen in Fig. 10.7. Finally, Fig. 10.10 (right) shows the singular values for nine successive blocking steps for $m = 0.1$, $\mu = 0$, and $D = 100$. We observe that the singular values decrease faster in later blocking steps.

10.3.4 Results for larger lattices near the phase transition

In this subsection we present results for larger lattices, still for $N_f = 1$, and introduce a fit ansatz based on Sec. 9.5 that describes the data very nicely in the vicinity of the phase transition. A summary of this is also given in Ref. [77]. In Figs. 10.11 to 10.14 we present OS-GHOTRG results obtained from a particular example, i.e., $L = 32$, $m = 0.5$, $n_{\max} = 3$, and several bond dimensions D . In Figs. 10.15 and 10.17 we then analyze the dependence of the OS-GHOTRG results on m and L .

Figure 10.11 shows the coefficients of the chiral condensate and the quark-number density as a function of μ near the phase transition for three different bond dimensions. The small discrepancies between the data for different values of D imply that the results have not yet fully converged. As a consistency check, we also show the coefficients Σ_n and ρ_n obtained with $n_{\max} = n$, which converge for smaller values of D .

In order to model the behavior of these coefficients in the vicinity of the phase transition we introduce a fit ansatz for each observable, motivated by Fermi-Dirac statistics and the represen-

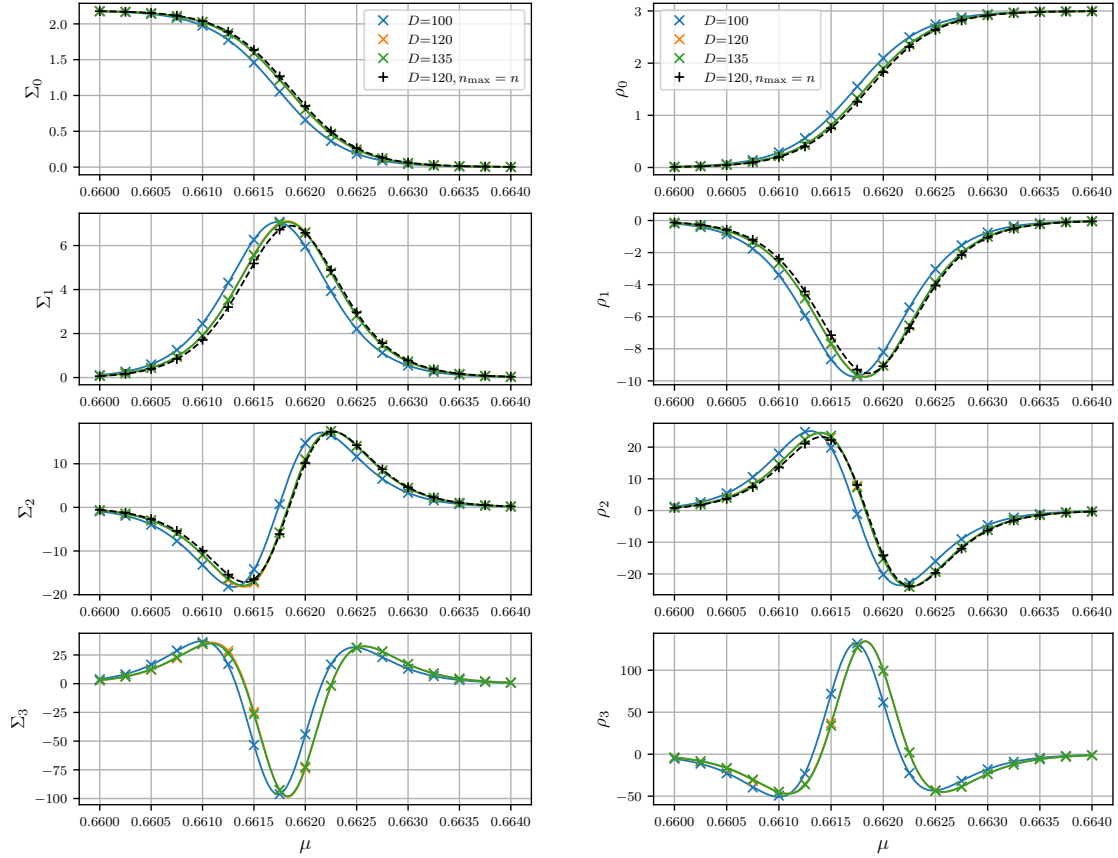


Figure 10.11: Coefficients of the chiral condensate (left) and the quark-number density (right) as a function of μ near the phase transition for $n_{\max} = 3$ with $L = 32$ and $m = 0.5$ using the Xie et al. truncation method for various bond dimensions D . We observe excellent agreement between the tensor data (crosses) and the model functions (solid lines) given in Appendix I.1. The large- μ limits of the coefficients Σ_n and ρ_n , see Sec. 9.3, are consistent with the analytical results.

tation of Z in (9.3.11).³² For the chiral condensate we use³³

$$\Sigma^{\text{model}}(\mu, \beta) = b_{\Sigma}(\beta) (1 - \tanh[a_{\Sigma}(\beta) (\mu - \mu_{\Sigma}^c(\beta))]) \quad (10.3.7)$$

with parameters that also depend on m and the lattice size. Using polynomial approximations

$$\mu_{\Sigma}^c(\beta) = \sum_{n=0}^{n_{\max}} \mu_{\Sigma,n}^c \beta^n, \quad a_{\Sigma}(\beta) = \sum_{n=0}^{n_{\max}} a_{\Sigma,n} \beta^n, \quad b_{\Sigma}(\beta) = \sum_{n=0}^{n_{\max}} b_{\Sigma,n} \beta^n \quad (10.3.8)$$

for the parameters and expanding the hyperbolic tangent in β , we obtain model functions for all coefficients Σ_n , which are given in Appendix I.1 up to $n = 3$. For every bond dimension shown in Fig. 10.11 (left), the tensor data for Σ_n are fitted with the model functions $\Sigma_n(\mu)$ simultaneously for all n . To take into account the different ranges of the Σ_n , given by $r_n = \max_{\mu} \Sigma_n(\mu) - \min_{\mu} \Sigma_n(\mu)$, we assign a weight $1/r_n^2$ to every data point.

For the quark-number density we use the fit ansatz^{34,35}

$$\rho^{\text{model}}(\mu, \beta) = \frac{N_c}{2} (1 + \tanh[a_{\rho}(\beta) (\mu - \mu_{\rho}^c(\beta))]), \quad (10.3.9)$$

which contains the critical chemical potential and transition sharpness corresponding to ρ . Both parameters also depend on m and the lattice size. Using polynomial approximations

$$\mu_{\rho}^c(\beta) = \sum_{n=0}^{n_{\max}} \mu_{\rho,n}^c \beta^n \quad \text{and} \quad a_{\rho}(\beta) = \sum_{n=0}^{n_{\max}} a_{\rho,n} \beta^n \quad (10.3.10)$$

and expanding the hyperbolic tangent, we obtain model functions for all coefficients ρ_n , also given in Appendix I.1 up to $n = 3$. For every bond dimension, we again perform simultaneous fits of the model functions to the tensor data with weights assigned in the same way as for the chiral condensate. As shown in Fig. 10.11, we find excellent agreement between the model functions and the tensor data for both observables, with 12 and 8 fit parameters for Σ and ρ , respectively. Furthermore, we observe that the coefficients for $n \geq 1$ become large near the phase transition, as anticipated in Sec. 9.5.

In Fig. 10.12 we plot the results for $\Sigma(\mu, \beta)$ and $\rho(\mu, \beta)$ as a function of μ , close to the phase transition, for $n_{\max} = 3$, $m = 0.5$, $L = 32$, and several values of β . For small values of β (e.g., $\beta \leq 0.1$ in Fig. 10.12) the expansions in (10.3.6) coincide with the tanh model of (10.3.7) and (10.3.9), respectively. However, for larger values of β , we observe that the expansions in (10.3.6) exhibit increasingly unphysical behavior (not shown in the figure). This is to be expected since the coefficients Σ_n and ρ_n must become large (for $n \geq 1$) near the critical chemical potential as

³²This ansatz will be substantiated further in Secs. 9.3.2 and 9.4.

³³An improved fit ansatz would be $\Sigma^{\text{model}}(\mu, \beta) + \Sigma^{\text{model}}(-\mu, \beta) - 2b_{\Sigma}(\beta)$, which is even in μ as required. However, for $\mu \geq 0$ the additional term is negligible if $a_{\Sigma} \mu_{\Sigma}^c \gg 1$, which is the case for the quark masses and lattice sizes we use. Note that the improved fit ansatz has the structure of the result (9.3.28), which corresponds to the large- m limit.

³⁴Again, an improved fit ansatz would be $\rho^{\text{model}}(\mu, \beta) - \rho^{\text{model}}(-\mu, \beta)$, which is odd in μ . This fit ansatz has the structure of the result (9.3.31), which corresponds to the large- m limit.

³⁵Since this model function includes the analytical dependence on the chemical potential μ , we can integrate this expression with respect to μ in order to obtain a model function for the free energy instead. Consequently, we can determine $\mu_{\rho,n}^c$ and $a_{\rho,n}$ from numerical data for the coefficients of the free energy. This is preferred since this eliminates the numerical derivative with respect to μ . The derivation of the model function for the free energy is shown in Appendix I.2.

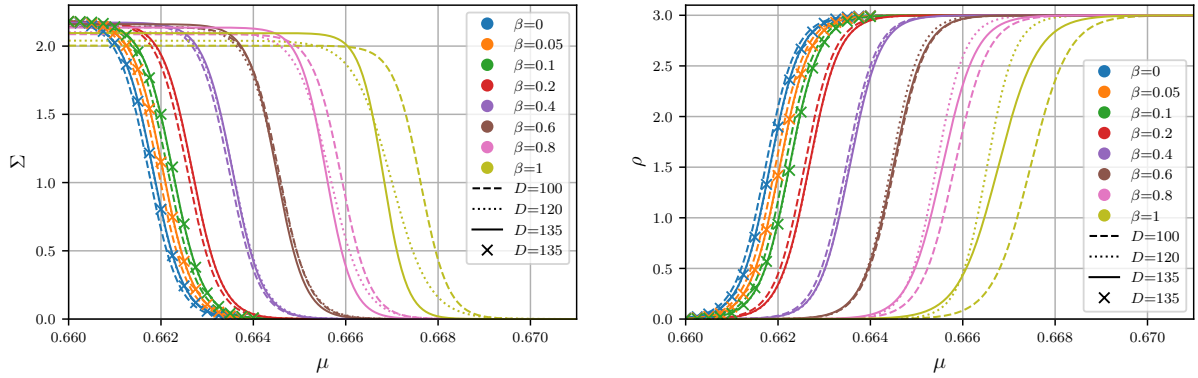


Figure 10.12: Chiral condensate (left) and quark-number density (right) as a function of μ for $m = 0.5$, $L = 32$, different bond dimensions, and several values of β for $n_{\max} = 3$. The data points (crosses) were obtained from (10.3.6), while the lines correspond to the tanh model of (10.3.7) and (10.3.9), respectively. In both plots, we observe an increasing critical chemical potential with increasing coupling parameter β .

explained in Sec. 9.5, see also Fig. 10.11. In contrast, in the tanh model of (10.3.7) and (10.3.9) with the expansions (10.3.8) and (10.3.10), the coefficients μ_n^c , a_n/a_0 , and b_n do not have to become large, which is also what we find, see Fig. 10.13. Therefore we expect the tanh model to yield reliable results also for larger values of β (e.g., for $\beta > 0.1$ in Fig. 10.12). Note that, as expected, the region of β where the expansion (10.3.6) works is much smaller for $L = 32$ than for $L = 8$, see Sec. 9.5. Moreover, we also note a qualitative change in the phase transition. For $L = 8$ the transition sharpness decreases with β , whereas for $L = 32$ it is rather insensitive to β , and it is merely the position μ^c of the phase transition that is shifted to larger values when β increases. The convergence (or lack thereof) of the fit parameters as a function of the bond dimension is shown in Fig. 10.13.

In Fig. 10.14 we show the β dependence of critical chemical potential $\mu^c(\beta)$ and transition sharpness $a(\beta)$ for both observables, obtained from (10.3.8) and (10.3.10), using $n_{\max} = 3$, $L = 32$, and the coefficients shown in Fig. 10.13 for bond dimensions $D = 100, 120, 135$. We observe that the results obtained from ρ and Σ coincide within the systematic errors of the tensor-network method.

We now turn to the dependence of the OS-GHOTRG results on m and L . In Fig. 10.15 we show the coefficients of (10.3.8) and (10.3.10) as a function of m for $n_{\max} = 2$. Again we observe that for both observables the critical chemical potential is very similar, and so is the transition sharpness. As expected, the critical chemical potential increases with the quark mass. Also shown in the figure are asymptotic limits for large m (see Sec. 9.3.2), which corresponds to the quenched theory. Finally, the dashed lines (marked IP) correspond to formulas obtained in Sec. 9.4 from interpolating Z between $\mu = 0$ and $\mu = \infty$. The value of $a_0 = 3V/2$ observed for small m (for both observables) is consistent with this interpolation, see (9.4.4). Figure 10.16 shows the same as Fig. 10.15 except that $n_{\max} = 3$ and we focus on masses between $m = 0.1$ and $m = 0.5$. The interpolation is consistent with the tensor data for $n_{\max} = 3$ as well. The used tensor data in Fig. 10.15 and Fig. 10.16 are shown explicitly in App. J.

In Fig. 10.17 we show the dependence of the coefficients of (10.3.10) on L for $m = 0.5$. We

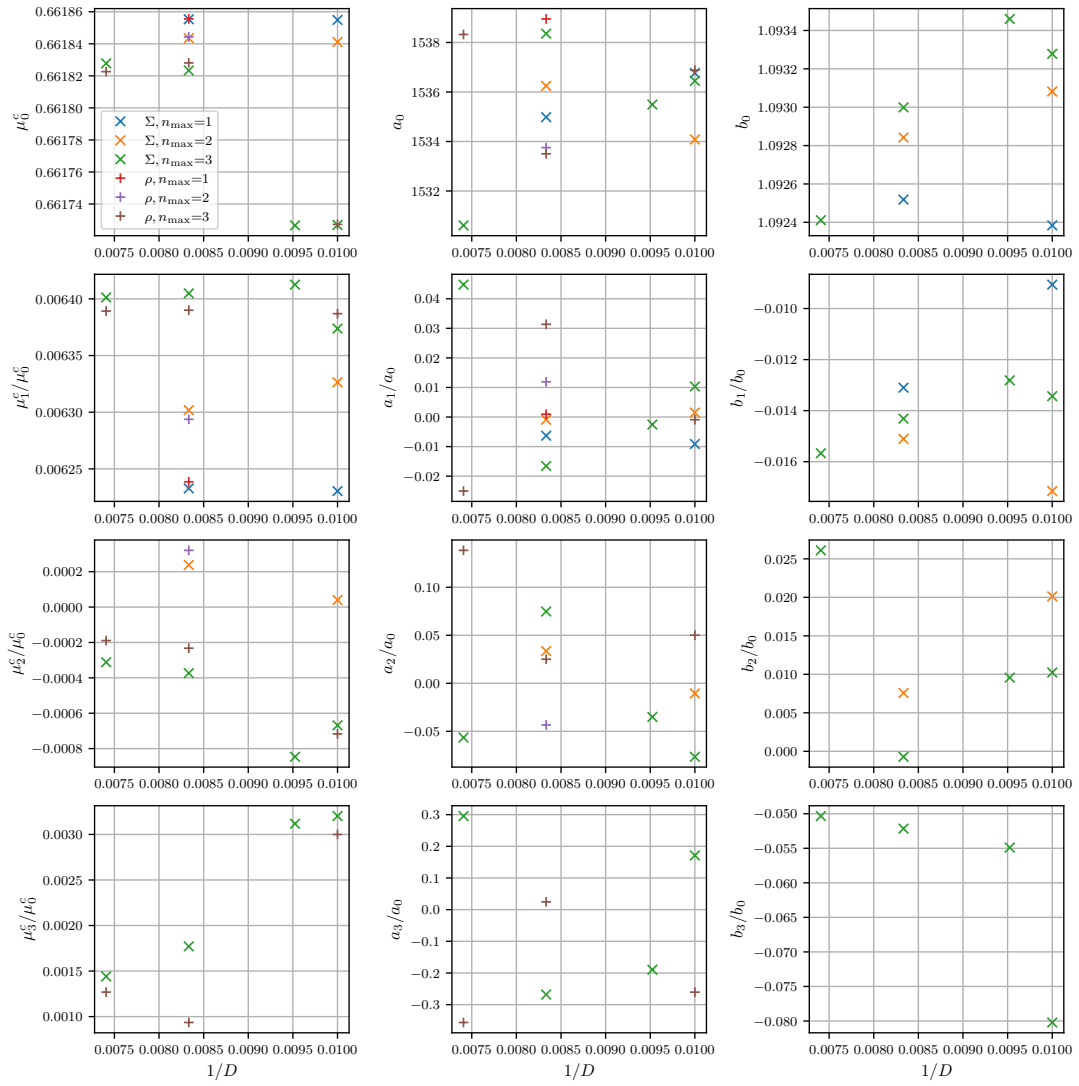


Figure 10.13: Dependence of the coefficients of the tanh model, see (10.3.8) and (10.3.10), on the bond dimension D for $n_{\max} = 1, 2, 3$, $m = 0.5$, and $L = 32$. To remove the overall scale we plot the coefficients for $n \geq 1$ relative to the coefficient for $n = 0$.

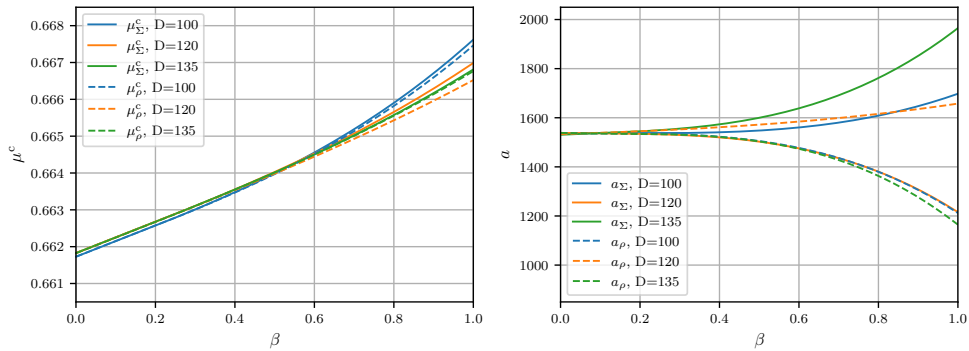


Figure 10.14: Critical chemical potential and transition sharpness as a function of β for the quark-number density and the chiral condensate for $n_{\max} = 3$, $m = 0.5$, and $L = 32$.

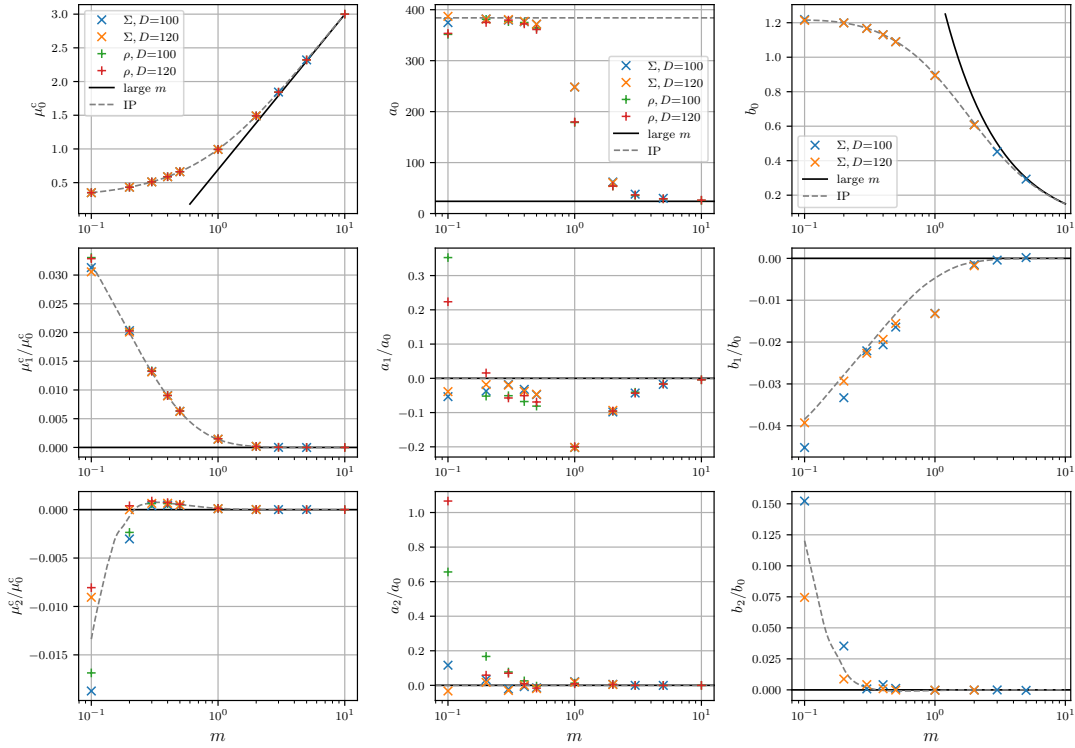


Figure 10.15: Fitted coefficients of the tanh models as functions of the quark mass m . In contrast to the previous figures, we set $L = 16$ in order to be able to reach small values of m with manageable numerical effort (determined by D). Smaller m requires higher bond dimensions to obtain convergence, see Fig. 10.10. In the OS-GHOTRG method, larger L also requires higher bond dimensions due to volume-dependent cancellations in (10.3.5) for $n \geq 2$, see Sec. 10.3.1. The large- m and interpolation (IP) results are given in Sec. 9.3.2 and Sec. 9.4, respectively.

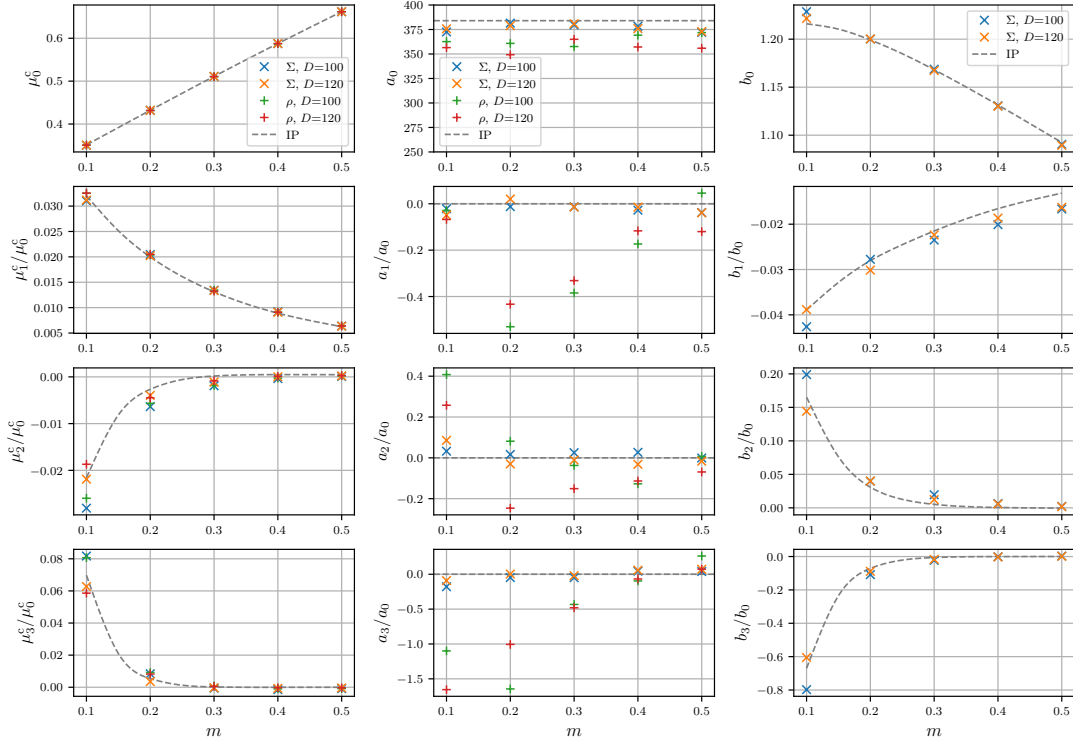


Figure 10.16: Fitted coefficients of the tanh models as functions of the quark mass m for $n_{\max} = 3$ and small masses. The interpolation (IP) results are given in Sec. 9.4.

chose $n_{\max} = 1$ because for larger volumes larger bond dimensions are required to determine a_n/a_0 (for $n \geq 1$) and μ_n^c (for $n \geq 2$) reliably. This is expected in the vicinity of the phase transition, whose sharpness increases with the lattice size (with exponent γ depending on m , see Fig. 10.15), see Sec. 9.5 and Appendix I.1. This issue already occurs for the largest values of L in Fig. 10.17 and becomes more severe for even larger L . Nevertheless, within the accuracy of the OS-GHOTRG results, a_1/a_0 is independent of L for large L , as assumed earlier. The coefficients for μ^c are nearly constant in L so that the location μ^c of the phase transition seems to converge rapidly to its infinite-volume limit. Similarly, in Fig. 10.18, we show the dependence of the coefficients of Eq. (10.3.8) and (10.3.10) on L for $n_{\max} = 2$, for $m = 0.5$ and $L \leq 64$. The values of μ_1^c vary more significantly with respect to the bond dimension for $n_{\max} = 2$, i.e., larger D is required to obtain reliable results. As expected, especially for large lattices, larger D is required to determine a_2/a_0 and μ_2^c reliably. The used tensor data in Fig. 10.17 and Fig. 10.18 are shown explicitly in App. J.

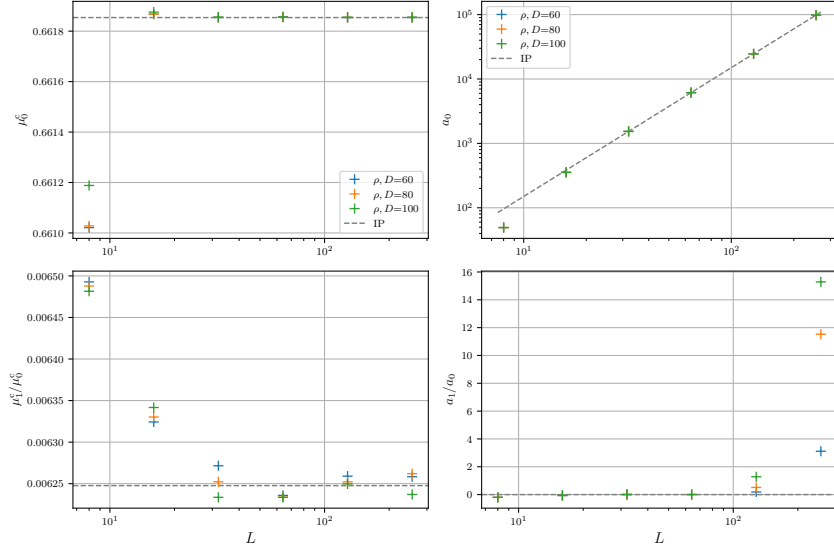


Figure 10.17: Fitted coefficients of (10.3.10) as functions of L for $n_{\max} = 1$ and $m = 0.5$. Note that μ_0^c and μ_1^c vary only very slightly with L . The interpolation (IP) results are given in Sec. 9.4.

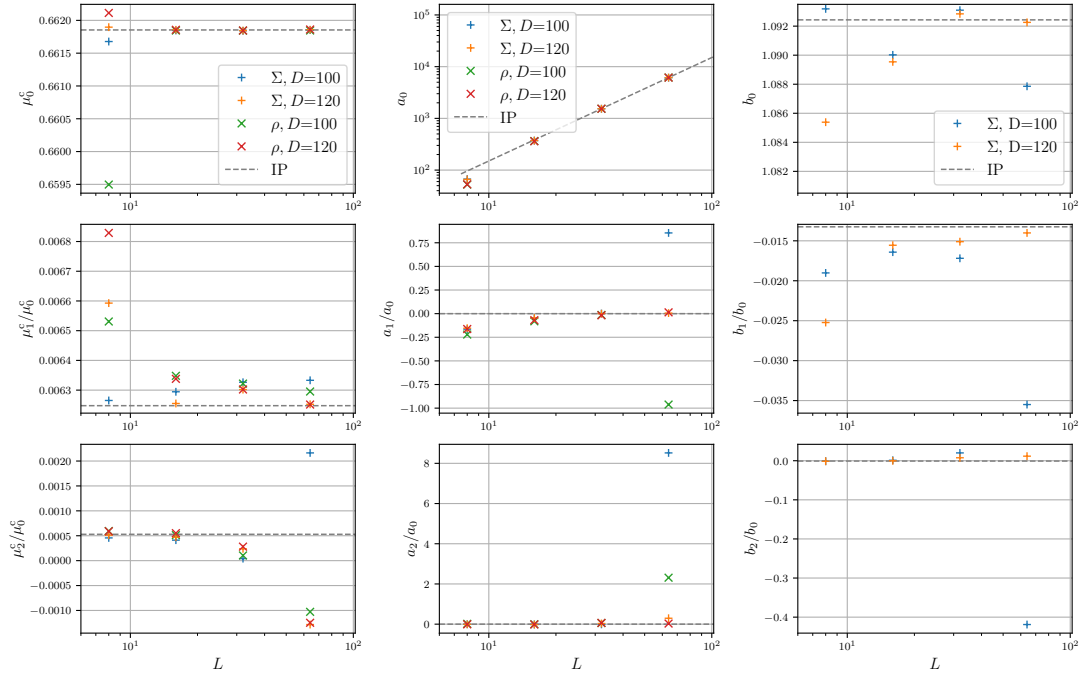


Figure 10.18: Fitted coefficients of Eq. (10.3.8) and (10.3.10) as functions of L for $n_{\max} = 2$ and $m = 0.5$. The interpolation (IP) results are given in Sec. 9.4.

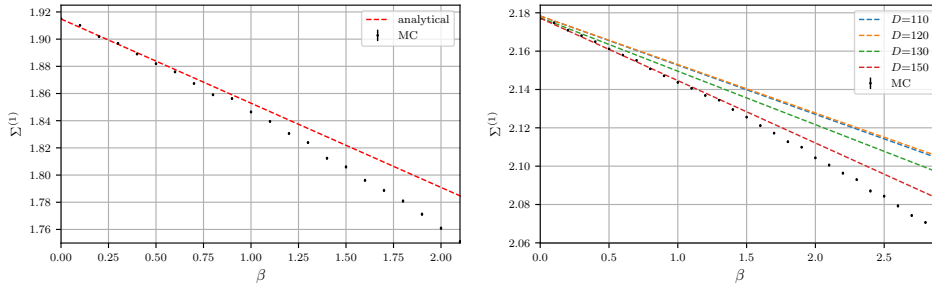


Figure 10.19: Chiral condensate $\Sigma^{(1)} = -\partial f / \partial m_1$ for $N_f = 2$ as a function of β for $m_1 = m_2 = 0.5$ with $\mu_1 = \mu_2 = 0$. Left: Comparison of analytical results obtained from the initial tensor for $n_{\max} = 1$ with Monte Carlo (MC) data on a 2×2 lattice. Right: Comparison of the OS-GHOTRG results for $n_{\max} = 1$ with Monte Carlo data on an 8×8 lattice.

10.4 Two flavors

We now present results for the case of two flavors. Note that for $L = 2$ the partition function and its derivatives can be computed analytically, as explained in Sec. 9. They are given explicitly in Eq. (E.2). We will use these analytical results to validate the numerical results. We will also compare to Monte Carlo data obtained at $\mu = 0$.

For $N_f = 2$ the chiral condensate for particle i is given by $\Sigma^{(i)} = \partial_{m_i} \ln Z / V$. In Fig. 10.19 (left) we show the chiral condensate as a function of β for $L = 2$, $m_1 = m_2 = 0.5$, and $\mu_1 = \mu_2 = 0$ and compare the analytical results obtained from the initial tensor for $n_{\max} = 1$ with Monte Carlo data. Both the infinite-coupling limit and the slope are in good agreement. In Fig. 10.19 (right) we show a similar plot for $L = 8$, where we compare OS-GHOTRG results obtained using the expansion (9.3.4) with various bond dimensions to Monte Carlo data. For the largest bond dimensions the slope of the tensor data agrees well with that of the Monte Carlo calculation.

In Fig. 10.20 we turn on the chemical potentials and show the infinite-coupling limit, i.e., $n_{\max} = 0$, of the free energy f , the chiral condensate $\Sigma^{(1)}$, and the baryon number density $\rho_B = -\partial f / \partial \mu_B$ as a function of the baryon chemical potential $\mu_B = 3(\mu_1 + \mu_2)/2$ for various mass values and zero isospin chemical potential $\mu_I = (\mu_1 - \mu_2)/2$, still for $L = 2$. The OS-GHOTRG results again agree very well with the analytical predictions. Note that the smaller masses required a larger bond dimension to obtain an acceptable accuracy. This is related to the behavior of the singular values, see also Fig. 10.10.

Next, we turn on the isospin chemical potential, i.e., we set $\mu_1 \neq \mu_2$. In Fig. 10.21 we show analytical results, obtained with the $n_{\max} = 1$ initial tensor, for the coefficients $\rho_{B,0}$ and $\rho_{B,1}$ of the baryon density as a function of the baryon and isospin chemical potential for $L = 2$ and $m_1 = m_2 = 0.5$. In the heatmap of $\rho_{B,0}$ we clearly recognize the pattern of the phase transition with plateaus at $\rho_B = 1$ and 2, to which $\rho_{B,1}$ yields small corrections as a function of β .

We now also consider larger lattices for which the OS-GHOTRG was developed. In Fig. 10.22 we show $\rho_{B,0}$ and $\rho_{B,1}$ as a function of μ_B for $\mu_I = 0$ (left) and $\mu_I = 0.5$ (right) for $L = 32$. When turning on the isospin chemical potential we see that there are two phase transitions at different values of μ_B , which correspond to the excitation of the quarks with different chemical potentials

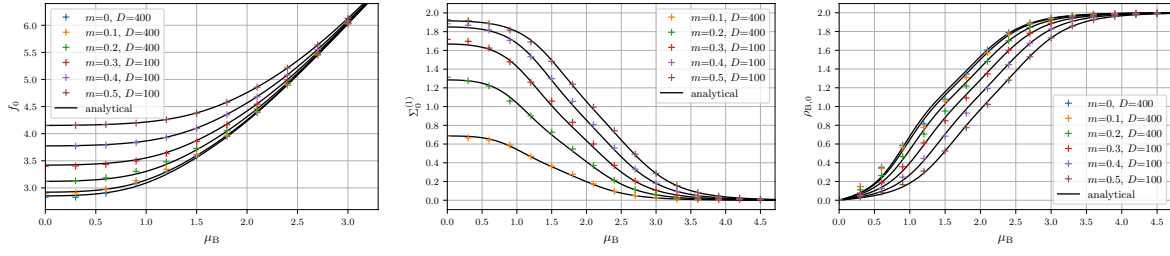


Figure 10.20: Plot of $-f = (\ln Z)/V$ (left), $\Sigma^{(1)}$ (middle), and ρ_B (right) as a function of μ_B for $N_f = 2$, $\mu_I = 0$, $n_{\max} = 0$, i.e., in the infinite coupling limit, and $m_1 = m_2 = m$ on a 2×2 lattice. The tensor-network results are compared with analytical tensor results.

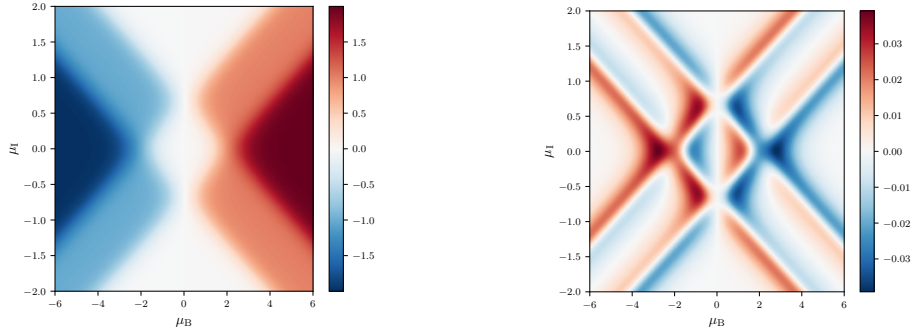


Figure 10.21: Analytical results for $\rho_{B,0}$ (left) and $\rho_{B,1}$ (right) as a function of μ_B and μ_I for $L = 2$, $n_{\max} = 1$, and $m_1 = m_2 = 0.5$.

$\mu_1 \neq \mu_2$. We observe that for the smaller values of μ_B , the results obtained for $n_{\max} = 1$ have not yet converged for the bond dimensions that we could reach on our computers.

A straight-forward model function for the free energy for a general number of flavors is given by

$$f(\vec{\mu}, \beta) = f^{\text{PG}}(\beta) + \sum_{f=1}^{N_f} \frac{N_c}{2a_f(\beta)} \ln(\cosh[a_f(\beta)(\mu_f - \mu_f^c(\beta))] \cosh[a_f(\beta)(\mu_f + \mu_f^c(\beta))]), \quad (10.4.1)$$

which is consistent with the condition (9.3.51). When all masses are equal, one could constrain a_f and μ_f^c to be independent of f , which simplifies the ansatz. In principle, one could expand this expression in β and derive model functions for the coefficients $\rho_{B,n}$ that include both transitions shown in Fig. 10.22 (right). However, reliable results are only expected when the tensor data are rather converged with respect to D , which is not the case for both transitions.

Finally, Fig. 10.23 shows the coefficients $\Sigma_n^{(1)}$ versus μ_B for $m_1 = m_2 = 0.5$, $L = 32$, $n_{\max} = 1$, $\mu_I = 0$ (left), $\mu_I = 0.1$ (middle), and $\mu_I = 0.5$ (right). Additionally, we show the fit that is obtained by simply using the model function (10.3.7), where μ is replaced by μ_B . Even though, the data is not yet converged with respect to the bond dimension, we find good agreement between the model functions and the tensor data. As expected, $\Sigma_n^{(1)}$ is not sensitive to the excitation of quarks corresponding to flavor 2, i.e., we did not observe a second phase transition in another region of μ_B (not shown). For increasing μ_B , $\Sigma_n^{(1)}$ goes to zero, as expected from Eq. (9.3.49).³⁶

³⁶Equation (9.3.49) is applicable, since m_1 and μ_I are fixed and therefore $\mu_1 \rightarrow \infty$ for $\mu_B \rightarrow \infty$.

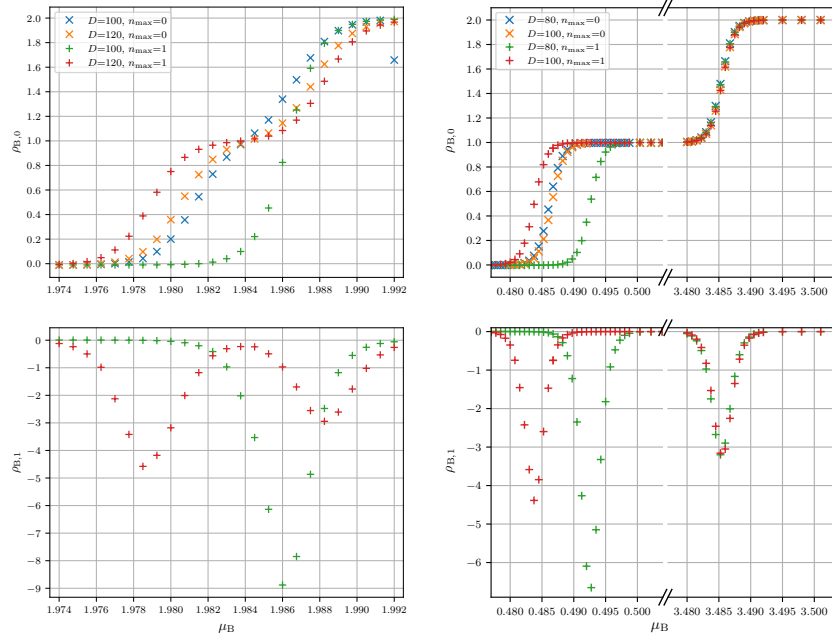


Figure 10.22: $\rho_{B,0}$ (top) and $\rho_{B,1}$ (bottom) versus μ_B for $\mu_I = 0$ (left) and $\mu_I = 0.5$ (right), $N_f = 2$, $m_1 = m_2 = 0.5$, $L = 32$, and using the Xie et al. truncation method.

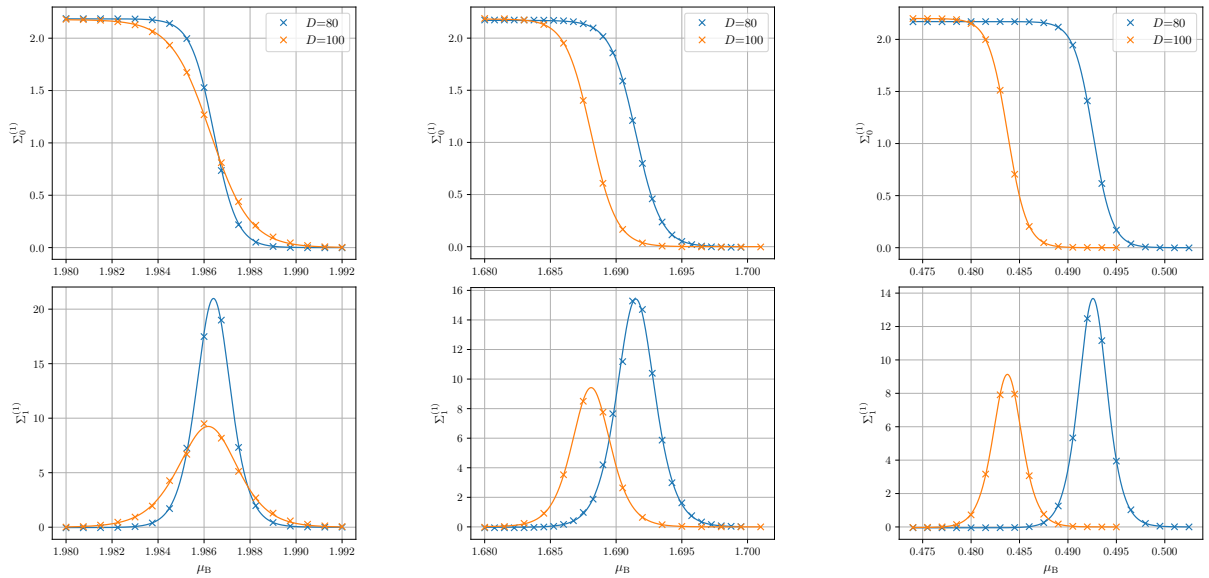


Figure 10.23: $\Sigma_n^{(1)}$ versus μ_B for $m_1 = m_2 = 0.5$, $L = 32$, $n_{\max} = 1$, $\mu_I = 0$ (left), $\mu_I = 0.1$ (middle), and $\mu_I = 0.5$ (right), and using the Xie et al. truncation method.

VI. Summary and conclusions

In this thesis, we derived a tensor-network formulation to compute the partition function of QCD in the strong-coupling expansion for arbitrary space-time dimension, number of colors, number of staggered fermion flavors, and order in β . To do so, we first performed Taylor expansions of the exponentials involving the gauge and fermion actions, which generate fermion and plaquette occupation numbers. We then integrated out the gauge field and observed that, due to the presence of Grassmann variables at both ends of every link, the subsequent calculation can be simplified, i.e., the sum over permutations and groupings of color indices in the gauge integral can be reduced significantly. In the next step we introduced colorless auxiliary Grassmann variables on the links of the lattice, which enabled us to integrate out the original Grassmann variables. Finally, we summed over all color indices to obtain a partition function given by the full contraction of a tensor network of mixed numerical and Grassmann tensors.

We observed that the results of the tensor-network method strongly depend on how the expansion of Z is used to compute observables. The analytical results for $L = 2$, numerical results for $L = 8$, and a theoretical analysis for arbitrary volume show that reliable results in a significant range of β require an expansion of $\ln Z$ or the observable itself, and not only an expansion of Z . This in turn requires knowledge of the expansion coefficients Z_n , also for larger lattices where they are not available through standard tensor-network methods.

In order to be able to calculate the coefficients Z_n , we have presented the order-separated Grassmann higher-order tensor renormalization group (OS-GHOTRG) method that is applicable still for arbitrary space-time dimension, number of colors, number of staggered fermion flavors, and order in β . The crucial ingredient of the method is that at every blocking step, the entries of the coarse-grid tensor are given by a power series in β . If the initial tensor is exact up to order $\beta^{n_{\max}}$ and contains no higher-order contributions, then the partition function computed with the OS-GHOTRG method is also exact up to this order, up to HOSVD truncation errors, and does not contain any higher-order contributions. The free energy, and hence also the thermodynamical observables, are computed up to $\beta^{n_{\max}}$ by a truncated Taylor expansion of $(\ln Z)/V$. For the reliability of the results it is crucial that this expansion contains no incomplete higher-order contributions. As shown in Fig. 10.5, this explains the essential improvement achieved by OS-GHOTRG compared to GHOTRG. The performance of the OS-GHOTRG method is substantially improved by making use of the translational invariance of the system.

We have validated the OS-GHOTRG method up to $n_{\max} = 2$ on a 2×2 lattice by comparing tensor results order by order with analytical calculations for $N_f = 1$, and on an 8×8 lattice by comparing tensor results for the quark-number density with Monte Carlo data. For larger lattices ($L > 8$) and with $n_{\max} = 3$ we argued, and confirmed numerically, that if we expand the observables near the phase transition in β , the higher-order coefficients become large and the expansion has a convergence radius that decreases with the lattice volume. However, we found that the chiral condensate and the quark-number density are fitted very well by tanh-based models, where the fit parameters are the critical chemical potential μ^c , the sharpness a of the phase transition, and the magnitude b of the chiral condensate. Since the fit parameters are expected to vary smoothly with β we expressed them as polynomials in β . The fits then allow for an extrapolation of the observables to much larger β values, and the range of validity appears to

be nearly independent of the lattice volume. We also discussed how the fit parameters depend on the bond dimension, the quark mass, and the lattice volume. Note that we presented results for moderate values of the mass, as the computational effort rises very quickly as the mass is taken to be small.

We also presented $N_f = 2$ results, only up to $n_{\max} = 1$ as larger bond dimensions are required to obtain reliable results. We compared observables computed with the OS-GHOTRG method with analytical results for a 2×2 lattice and with Monte Carlo data on an 8×8 lattice and found good agreement up to higher-order corrections in β . We also showed the evolution of the observables as a function of the baryon and isospin chemical potentials. For a 32×32 lattice we showed the chiral condensate and the baryon-number density versus the baryon chemical potential for different values of the isospin chemical potential.

In future work, it would be interesting to apply the OS-GHOTRG method in three and four dimensions, where the computational complexity and the memory requirements become much larger.

Acknowledgements

An dieser Stelle möchte ich allen Personen danken, die zur Entstehung dieser Arbeit in verschiedenster Weise beigetragen haben:

Tilo Wettig dafür, dass er mir es ermöglichte, an einem so spannenden Thema zu arbeiten. Seine wissenschaftliche Kompetenz und seine Expertise haben in hohem Maße zum Gelingen dieser Arbeit beigetragen. Ich bin dankbar, dass er mir die Teilnahme an der Gitter Konferenz 2024 in Liverpool ermöglichte.

Jacques Bloch für die außerordentlich gute Betreuung in allen Phasen der Promotion. Dafür, dass er zu jeder Zeit - auch spontan - ein offenes Ohr für mich hatte und mich in jederlei Hinsicht mit Rat und Tat unterstützte.

Robert Lohmayer für die engagierte Unterstützung und Mitarbeit, für unzählige Impulse und hilfreiche Diskussionen, die wesentlich zu den Fortschritten beigetragen haben.

Patrick Gotzler für die vielen Diskussionen zu allen möglichen Themen der theoretischen Physik während meiner Zeit an der Uni Regensburg. Ihm danke ich auch, neben Julian Parrino, für das sorgfältige Korrekturlesen dieser Arbeit.

Meinen Eltern und meinem Bruder Georg, für ihre Hilfe und ihren Beistand in allen Belangen. Zu guter Letzt möchte ich meiner Freundin Katharina meinen herzlichen Dank aussprechen, dafür, dass sie mich unermüdlich unterstützt und mir immer Rückhalt gibt.

A Grassmann contribution

In this appendix we integrate out all the original Grassmann variables in $\mathcal{G}$ of Eq. (4.2.10), i.e., we compute

$$\int \left[\prod_x d\psi_x d\bar{\psi}_x \right] \mathcal{G}, \quad (\text{A.1})$$

to arrive at an expression that is suitable for the tensor-network formulation. A basic ingredient will be the relation

$$\prod_{i=1}^k \alpha^i \beta^i = \prod_{i=1}^k \alpha^i \prod_{i=1}^k \beta^i = \prod_{i=1}^k \alpha^i \prod_{i=1}^k \beta^i \quad (\text{A.2})$$

which is valid if, for every i , the product $\alpha^i \beta^i$ is commuting. The relation is most easily shown by moving suitable pairs of Grassmann variables. The symbol $\prod$ denotes the reverse product defined in Eq. (2.6.39).

The expression for $\mathcal{G}_{x,\mu}$ in Eq. (4.2.10) can be manipulated as follows,

$$\begin{aligned} \mathcal{G}_{x,\mu} &= \left[\prod_f \prod_{a=1}^{k_{x,\mu}^f} \bar{\psi}_x^{f,i_{x,\mu}^{a+\kappa_{x,\mu}^f}} \psi_{x+\hat{\mu}}^{f,j_{x,\mu}^{a+\kappa_{x,\mu}^f}} \right] \left[\prod_f \prod_{a=1}^{\bar{k}_{x,\mu}^f} \psi_x^{f,m_{x,\mu}^{a+\bar{\kappa}_{x,\mu}^f}} \bar{\psi}_{x+\hat{\mu}}^{f,\ell_{x,\mu}^{a+\bar{\kappa}_{x,\mu}^f}} \right] \\ &= \left[\prod_f \prod_{a=1}^{k_{x,\mu}^f} \bar{\psi}_x^{f,i_{x,\mu}^{a+\kappa_{x,\mu}^f}} \prod_{a=1}^{k_{x,\mu}^f} \psi_{x+\hat{\mu}}^{f,j_{x,\mu}^{a+\kappa_{x,\mu}^f}} \right] \left[\prod_f \prod_{a=1}^{\bar{k}_{x,\mu}^f} \psi_x^{f,m_{x,\mu}^{a+\bar{\kappa}_{x,\mu}^f}} \prod_{a=1}^{\bar{k}_{x,\mu}^f} \bar{\psi}_{x+\hat{\mu}}^{f,\ell_{x,\mu}^{a+\bar{\kappa}_{x,\mu}^f}} \right] \\ &= \left[\prod_f \prod_{a=1}^{k_{x,\mu}^f} \bar{\psi}_x^{f,i_{x,\mu}^{a+\kappa_{x,\mu}^f}} \right] \left[\prod_f \prod_{a=1}^{k_{x,\mu}^f} \psi_{x+\hat{\mu}}^{f,j_{x,\mu}^{a+\kappa_{x,\mu}^f}} \right] \left[\prod_f \prod_{a=1}^{\bar{k}_{x,\mu}^f} \psi_x^{f,m_{x,\mu}^{a+\bar{\kappa}_{x,\mu}^f}} \right] \left[\prod_f \prod_{a=1}^{\bar{k}_{x,\mu}^f} \bar{\psi}_{x+\hat{\mu}}^{f,\ell_{x,\mu}^{a+\bar{\kappa}_{x,\mu}^f}} \right] \\ &= \left[\prod_f \prod_{a=1}^{k_{x,\mu}^f} \bar{\psi}_x^{f,i_{x,\mu}^{a+\kappa_{x,\mu}^f}} \right] \left[\prod_f \prod_{a=1}^{\bar{k}_{x,\mu}^f} \psi_x^{f,m_{x,\mu}^{a+\bar{\kappa}_{x,\mu}^f}} \right] \left[\prod_f \prod_{a=1}^{\bar{k}_{x,\mu}^f} \bar{\psi}_{x+\hat{\mu}}^{f,\ell_{x,\mu}^{a+\bar{\kappa}_{x,\mu}^f}} \right] \left[\prod_f \prod_{a=1}^{k_{x,\mu}^f} \psi_{x+\hat{\mu}}^{f,j_{x,\mu}^{a+\kappa_{x,\mu}^f}} \right]. \quad (\text{A.3}) \end{aligned}$$

In the first step, we have redistributed the product over flavors in Eq. (4.2.10) by reordering the commuting brackets.¹ In the second step, we applied Eq. (A.2) to the products over a .² In the third step, we applied Eq. (A.2) to the products over f , which is justified since both factors in each product are either commuting or anticommuting depending on the value of $k_{x,\mu}^f$ or $\bar{k}_{x,\mu}^f$. In the last step, we moved the product of the last two brackets, which is commuting, between the first and second bracket, so that the Grassmann variables are now sorted by site.

As the Grassmann variables in the last two brackets live on a different site, we need to reorder the brackets in the product over (x, μ) in Eq. (4.2.10) in order to gather all Grassmann variables living on a single site. However, the brackets are only commuting when $k_{x,\mu} + \bar{k}_{x,\mu}$, see Eq. (4.2.2), is even. In order to construct commuting expressions we insert an auxiliary Grassmann variable $\chi_{x,\mu}$ on each link using

$$\left(\int d\chi_{x,\mu} \chi_{x,\mu} \right)^{F_{x,\mu}} = 1 \quad (\text{A.4})$$

¹For simplicity, we say “bracket” instead of “bracketed term”.

²We use the reverse product for $\bar{\psi}$ and the ordinary product for ψ throughout.

with the Grassmann parity

$$F_{x,\mu} \equiv \begin{cases} 0 & \text{if } k_{x,\mu} + \bar{k}_{x,\mu} \text{ even,} \\ 1 & \text{if } k_{x,\mu} + \bar{k}_{x,\mu} \text{ odd.} \end{cases} \quad (5.1.4)$$

For a given link, we insert Eq. (A.4) in front of Eq. (A.3) to obtain

$$\begin{aligned} \mathcal{G}_{x,\mu} &= \left(\int d\chi_{x,\mu} \chi_{x,\mu} \right)^{F_{x,\mu}} \prod_f \prod_{a=1}^{k_{x,\mu}^f} \bar{\psi}_x^{f,\ell_{x,\mu}^{a+\kappa_{x,\mu}^f}} \prod_f \prod_{a=1}^{\bar{k}_{x,\mu}^f} \psi_x^{f,m_{x,\mu}^{a+\bar{\kappa}_{x,\mu}^f}} \prod_f \prod_{a=1}^{\bar{k}_{x,\mu}^f} \bar{\psi}_{x+\hat{\mu}}^{f,\ell_{x,\mu}^{a+\bar{\kappa}_{x,\mu}^f}} \prod_f \prod_{a=1}^{k_{x,\mu}^f} \psi_{x+\hat{\mu}}^{f,j_{x,\mu}^{a+\kappa_{x,\mu}^f}} \\ &= \left(\int \right)^{F_{x,\mu}} \left[(\chi_{x,\mu})^{F_{x,\mu}} \prod_f \prod_{a=1}^{k_{x,\mu}^f} \bar{\psi}_x^{f,\ell_{x,\mu}^{a+\kappa_{x,\mu}^f}} \prod_f \prod_{a=1}^{\bar{k}_{x,\mu}^f} \psi_x^{f,m_{x,\mu}^{a+\bar{\kappa}_{x,\mu}^f}} \right] \\ &\quad \times \left[(d\chi_{x,\mu})^{F_{x,\mu}} \prod_f \prod_{a=1}^{\bar{k}_{x,\mu}^f} \bar{\psi}_{x+\hat{\mu}}^{f,\ell_{x,\mu}^{a+\bar{\kappa}_{x,\mu}^f}} \prod_f \prod_{a=1}^{k_{x,\mu}^f} \psi_{x+\hat{\mu}}^{f,j_{x,\mu}^{a+\kappa_{x,\mu}^f}} \right], \end{aligned} \quad (A.5)$$

where in the second step we have moved the differential of the auxiliary Grassmann variable in order to obtain commuting expressions. Owing to Eq. (5.1.4), no sign factors are produced in this step.

Next, we perform the product over links in Eq. (4.2.10). Since we now have commuting expressions, we can shift $(x, \mu) \rightarrow (x - \hat{\mu}, \mu)$ in the last bracket of Eq. (A.5) to sort all original Grassmann variables according to their site,

$$\begin{aligned} \mathcal{G} &= e^{S_M} \int_{\chi} \prod_x \mathcal{G}_x \quad \text{with} \quad \mathcal{G}_x = \prod_{\mu} \left[(\chi_{x,\mu})^{F_{x,\mu}} \prod_f \prod_{a=1}^{k_{x,\mu}^f} \bar{\psi}_x^{f,\ell_{x,\mu}^{a+\kappa_{x,\mu}^f}} \right] \left[\prod_f \prod_{a=1}^{\bar{k}_{x,\mu}^f} \psi_x^{f,m_{x,\mu}^{a+\bar{\kappa}_{x,\mu}^f}} \right] \\ &\quad \times \left[(d\chi_{x,-\mu})^{F_{x,-\mu}} \prod_f \prod_{a=1}^{\bar{k}_{x,-\mu}^f} \bar{\psi}_x^{f,\ell_{x,-\mu}^{a+\bar{\kappa}_{x,-\mu}^f}} \right] \left[\prod_f \prod_{a=1}^{k_{x,-\mu}^f} \psi_x^{f,j_{x,-\mu}^{a+\kappa_{x,-\mu}^f}} \right], \end{aligned} \quad (A.6)$$

where we have defined

$$\int_{\chi} \equiv \prod_{x,\mu} \left(\int \right)^{F_{x,\mu}} \quad (5.1.2)$$

to indicate the integration over the auxiliary Grassmann variables and where $(x, -\mu)$ again is a shorthand notation for $(x - \hat{\mu}, \mu)$.³

We define the boundary conditions for the $\chi_{x,\mu}$ to be antiperiodic in time and periodic otherwise. The advantage of this definition is that the signs arising from the antiperiodic boundary conditions of ψ and $\bar{\psi}$ in the step from Eq. (A.5) to Eq. (A.6) are absorbed through the antiperiodic boundary conditions of χ (or, more precisely, $d\chi$).

In Eq. (A.6), ψ_x and $\bar{\psi}_x$ are not yet ordered in the integrand. In the next step we wish to gather all ψ_x and $\bar{\psi}_x$ in $\mathcal{G}_x$. To do so, we first note that $\mathcal{G}_x$ can be written symbolically as $\prod_{\mu} A_{\mu} B_{\mu} C_{\mu} D_{\mu}$, where A_{μ} and C_{μ} include $(\chi_{x,\mu})^{F_{x,\mu}}$ and $(d\chi_{x,-\mu})^{F_{x,-\mu}}$, respectively. Note that

³To avoid potential confusion we point out that $\mathcal{G}_x \neq \prod_{\mu} \mathcal{G}_{x,\mu}$.

A_μ and B_μ have the same Grassmann parity, and so do C_μ and D_μ . Thus both $A_\mu B_\mu$ and $C_\mu D_\mu$ are commuting. This allows us to write

$$\begin{aligned} \prod_\mu (A_\mu B_\mu)(C_\mu D_\mu) &= \left[\prod_\mu A_\mu B_\mu \right] \left[\prod_\mu C_\mu D_\mu \right] = \left[\prod_\mu A_\mu \prod_\mu B_\mu \right] \left[\prod_\mu C_\mu \prod_\mu D_\mu \right] \\ &= \prod_\mu C_\mu \prod_\mu A_\mu \prod_\mu B_\mu \prod_\mu D_\mu, \end{aligned} \quad (\text{A.7})$$

where in the first step we have distributed the product over μ , in the second step we have applied Eq. (A.2) twice, and in the final step we have moved the first commuting bracket inside the second one. Using Eq. (A.7) in $\mathcal{G}_x$ yields

$$\begin{aligned} \mathcal{G}_x &= \left[\prod_\mu (d\chi_{x,-\mu})^{F_{x,-\mu}} \prod_f \prod_{a=1}^{\bar{k}_{x,-\mu}^f} \bar{\psi}_x^{f,\ell_{x,-\mu}^{a+\bar{\kappa}_{x,-\mu}^f}} \right] \left[\prod_\mu (\chi_{x,\mu})^{F_{x,\mu}} \prod_f \prod_{a=1}^{k_{x,\mu}^f} \bar{\psi}_x^{f,\ell_{x,\mu}^{a+\kappa_{x,\mu}^f}} \right] \\ &\quad \times \left[\prod_\mu \prod_f \prod_{a=1}^{\bar{k}_{x,\mu}^f} \psi_x^{f,m_{x,\mu}^{a+\bar{\kappa}_{x,\mu}^f}} \right] \left[\prod_\mu \prod_f \prod_{a=1}^{k_{x,-\mu}^f} \psi_x^{f,j_{x,-\mu}^{a+\kappa_{x,-\mu}^f}} \right] \\ &= \sigma_x^{(1)} K_x \left[\prod_\mu \prod_f \prod_{a=1}^{\bar{k}_{x,-\mu}^f} \bar{\psi}_x^{f,\ell_{x,-\mu}^{a+\bar{\kappa}_{x,-\mu}^f}} \right] \left[\prod_\mu \prod_f \prod_{a=1}^{k_{x,\mu}^f} \bar{\psi}_x^{f,\ell_{x,\mu}^{a+\kappa_{x,\mu}^f}} \right] \\ &\quad \times \left[\prod_\mu \prod_f \prod_{a=1}^{\bar{k}_{x,\mu}^f} \psi_x^{f,m_{x,\mu}^{a+\bar{\kappa}_{x,\mu}^f}} \right] \left[\prod_\mu \prod_f \prod_{a=1}^{k_{x,-\mu}^f} \psi_x^{f,j_{x,-\mu}^{a+\kappa_{x,-\mu}^f}} \right] \end{aligned} \quad (\text{A.8})$$

with

$$K_x = \prod_\mu (\chi_{x,\mu})^{F_{x,\mu}} \prod_\mu (d\chi_{x,-\mu})^{F_{x,-\mu}}. \quad (5.1.3)$$

After the first step in Eq. (A.8) all ψ and all $\bar{\psi}$ are now gathered. In the second step we have moved the auxiliary Grassmann variables to the front and reordered them according to the prescription of [23] by commuting the Grassmann variables, which produces the local sign factor

$$\begin{aligned} \sigma_x^{(1)} &\equiv (-)^{p_1+p_2+p_3} \quad \text{with} \quad p_1 \equiv \sum_{\mu=1}^{d-1} F_{x,-\mu} \sum_{\nu=\mu+1}^d \bar{k}_{x,-\nu}, \\ p_2 &\equiv \sum_{\mu=1}^{d-1} F_{x,\mu} \sum_{\nu=\mu+1}^d (F_{x,\nu} + k_{x,\nu}), \\ p_3 &\equiv \sum_{\mu=1}^d F_{x,\mu} \sum_{\nu=1}^d (F_{x,-\nu} + \bar{k}_{x,-\nu}), \end{aligned} \quad (\text{A.9})$$

where we used the sum (4.2.2) over flavors for k and $\bar{k}$.

Equation (A.6) still contains the mass term, which we now also Taylor-expand. Starting from Eq. (2.6.35) we write

$$e^{S_M^f} = \prod_x \sum_{w_x^f} \frac{(2m_f)^{w_x^f}}{w_x^f!} (\bar{\psi}_x^f \psi_x^f)^{w_x^f}. \quad (\text{A.10})$$

In principle, the sum over w_x^f goes from zero to infinity. However, we will now show that the sum collapses to a single term once we integrate over all Grassmann variables. The point is that, due to the Grassmann integral, the only nonzero contributions to the partition function come from terms in which every Grassmann variable $\psi_x^{f,i}$ (and $\bar{\psi}_x^{f,i}$) occurs once. Let us fix the site x and the flavor f for the time being and define

$$h_x^{f,\text{in}} \equiv \sum_{\mu} (k_{x,-\mu}^f + \bar{k}_{x,\mu}^f) \quad \text{and} \quad h_x^{f,\text{out}} \equiv \sum_{\mu} (\bar{k}_{x,-\mu}^f + k_{x,\mu}^f), \quad (5.1.5)$$

which is the number of incoming and outgoing hopping terms, respectively, at site x for flavor f , see, for example, Eq. (4.2.10), where “in” and “out” correspond to ψ and $\bar{\psi}$, respectively. Since the total number of ψ and $\bar{\psi}$ variables for fixed x and f equals N_c each, we obtain the constraints

$$w_x^f + h_x^{f,\text{in}} = N_c = w_x^f + h_x^{f,\text{out}} \quad \text{and} \quad w_x^f \geq 0. \quad (A.11)$$

Hence, in the following we can, for each flavor, omit the sum over w_x^f in Eq. (A.10), replace w_x^f by $N_c - h_x^{f,\text{in}}$, and include a Heaviside step function $\Theta(w_x^f)$ as well as a Kronecker delta $\delta_{h_x^{f,\text{in}}, h_x^{f,\text{out}}}$.

We now use Eq. (A.2) to sort the ψ and $\bar{\psi}$ in Eq. (A.10),

$$(\bar{\psi}_x^f \psi_x^f)^{w_x^f} = \prod_{a=1}^{w_x^f} \bar{\psi}_x^{f,v_x^{f,a}} \psi_x^{f,v_x^{f,a}} = \prod_{a=1}^{w_x^f} \bar{\psi}_x^{f,v_x^{f,a}} \prod_{a=1}^{w_x^f} \psi_x^{f,v_x^{f,a}}, \quad (A.12)$$

where the Einstein summation convention is used for the color indices $v_x^{f,a}$. We then perform the product over flavors and again use Eq. (A.2) to obtain

$$\prod_f \prod_{a=1}^{w_x^f} \bar{\psi}_x^{f,v_x^{f,a}} \prod_{a=1}^{w_x^f} \psi_x^{f,v_x^{f,a}} = \prod_f \prod_{a=1}^{w_x^f} \bar{\psi}_x^{f,v_x^{f,a}} \prod_f \prod_{a=1}^{w_x^f} \psi_x^{f,v_x^{f,a}}. \quad (A.13)$$

We now substitute Eq. (A.10), multiplied over flavors, into Eq. (A.6) (using Eq. (A.8) for $\mathcal{G}_x$ and Eq. (A.13)) in such a way that ψ and $\bar{\psi}$ are not mixed. We then split off the part that only contains the original Grassmann variables, which is

$$\begin{aligned} \Psi \equiv & \prod_x \left[\prod_{\mu} \prod_f \prod_{a=1}^{\bar{k}_{x,-\mu}^f} \bar{\psi}_x^{f,\ell_{x,-\mu}^{a+\bar{\kappa}_x^f}} \right] \left[\prod_{\mu} \prod_f \prod_{a=1}^{k_{x,\mu}^f} \bar{\psi}_x^{f,\ell_{x,\mu}^{a+\kappa_x^f}} \right] \left[\prod_f \prod_{a=1}^{w_x^f} \bar{\psi}_x^{f,v_x^{f,a}} \right] \\ & \times \left[\prod_f \prod_{a=1}^{w_x^f} \psi_x^{f,v_x^{f,a}} \right] \left[\prod_{\mu} \prod_f \prod_{a=1}^{\bar{k}_{x,\mu}^f} \psi_x^{f,m_{x,\mu}^{a+\bar{\kappa}_x^f}} \right] \left[\prod_{\mu} \prod_f \prod_{a=1}^{k_{x,-\mu}^f} \psi_x^{f,j_{x,-\mu}^{a+\kappa_x^f}} \right]. \end{aligned} \quad (A.14)$$

To prepare for the integration over the original Grassmann variables it is useful to change the order of the products over μ and f in each of the factors in Eq. (A.14) so that the Grassmann variables are sorted by flavor. We obtain

$$\begin{aligned} \Psi = & \prod_x \left[\prod_f \prod_{\mu} \prod_{a=1}^{\bar{k}_{x,-\mu}^f} \bar{\psi}_x^{f,\ell_{x,-\mu}^{a+\bar{\kappa}_x^f}} \right] \left[\prod_f \prod_{\mu} \prod_{a=1}^{k_{x,\mu}^f} \bar{\psi}_x^{f,\ell_{x,\mu}^{a+\kappa_x^f}} \right] \left[\prod_f \prod_{a=1}^{w_x^f} \bar{\psi}_x^{f,v_x^{f,a}} \right] \\ & \times \left[\prod_f \prod_{a=1}^{w_x^f} \psi_x^{f,v_x^{f,a}} \right] \left[\prod_f \prod_{\mu} \prod_{a=1}^{\bar{k}_{x,\mu}^f} \psi_x^{f,m_{x,\mu}^{a+\bar{\kappa}_x^f}} \right] \left[\prod_f \prod_{\mu} \prod_{a=1}^{k_{x,-\mu}^f} \psi_x^{f,j_{x,-\mu}^{a+\kappa_x^f}} \right]. \end{aligned} \quad (A.15)$$

All sign factors that are generated by changing the order of the products cancel after the product over sites is performed. To show this, we first consider the two brackets involving $k_{x,\mu}^f$ in Eq. (A.14), which can symbolically be written as

$$\prod_{\mu} \prod_f (\alpha_{\mu}^f)^{k_{\mu}^f} \quad \text{and} \quad \prod_{\mu} \prod_f (\beta_{\mu}^f)^{k_{\mu}^f}, \quad (\text{A.16})$$

where we use $(\alpha)^k$ for a coproduct and $(\beta)^k$ for a product of k different Grassmann variables (the Grassmann variables inside these products do not get reordered). When reversing the order of the products over μ and f , the commutations of Grassmann variables generate a sign factor σ that is identical for both terms since the same commutations are performed but in opposite directions,

$$\begin{aligned} \prod_{\mu} \prod_f (\alpha_{\mu}^f)^{k_{\mu}^f} &= (\alpha_d^{N_f})^{k_d^{N_f}} \dots (\alpha_d^1)^{k_d^1} \dots (\alpha_1^{N_f})^{k_1^{N_f}} \dots (\alpha_1^1)^{k_1^1} \\ &= \sigma \times (\alpha_d^{N_f})^{k_d^{N_f}} \dots (\alpha_1^{N_f})^{k_1^{N_f}} \dots (\alpha_d^1)^{k_d^1} \dots (\alpha_1^1)^{k_1^1} = \sigma \prod_f \prod_{\mu} (\alpha_{\mu}^f)^{k_{\mu}^f}, \end{aligned} \quad (\text{A.17a})$$

$$\begin{aligned} \prod_{\mu} \prod_f (\beta_{\mu}^f)^{k_{\mu}^f} &= (\beta_1^1)^{k_1^1} \dots (\beta_1^{N_f})^{k_1^{N_f}} \dots (\beta_d^1)^{k_d^1} \dots (\beta_d^{N_f})^{k_d^{N_f}} \\ &= \sigma \times (\beta_1^1)^{k_1^1} \dots (\beta_d^1)^{k_d^1} \dots (\beta_1^{N_f})^{k_1^{N_f}} \dots (\beta_d^{N_f})^{k_d^{N_f}} = \sigma \prod_f \prod_{\mu} (\beta_{\mu}^f)^{k_{\mu}^f}. \end{aligned} \quad (\text{A.17b})$$

Hence, since $\sigma^2 = 1$, no net sign factor is generated when considering the product of both terms. A similar argument applies to the two brackets involving $\bar{k}_{x,\mu}^f$ in Eq. (A.14).

In Eq. (A.15) we now merge three products over f into one, for both $\bar{\psi}$ and ψ , resulting in

$$\Psi = \prod_x \sigma_x^{(2)} \prod_f \bar{\Psi}_x^f \prod_f \Psi_x^f \quad (\text{A.18})$$

with

$$\begin{aligned} \bar{\Psi}_x^f &= \left[\prod_{\mu} \prod_{a=1}^{\bar{k}_{x,-\mu}^f} \bar{\psi}_x^{f,\ell_{x,-\mu}^{a+\bar{\kappa}_{x,-\mu}^f}} \right] \left[\prod_{\mu} \prod_{a=1}^{k_{x,\mu}^f} \bar{\psi}_x^{f,\ell_{x,\mu}^{a+\kappa_{x,\mu}^f}} \right] \left[\prod_{a=1}^{w_x^f} \bar{\psi}_x^{f,v_{x,a}^{f,a}} \right], \\ \Psi_x^f &= \left[\prod_{a=1}^{w_x^f} \psi_x^{f,v_{x,a}^{f,a}} \right] \left[\prod_{\mu} \prod_{a=1}^{\bar{k}_{x,\mu}^f} \psi_x^{f,m_{x,\mu}^{a+\bar{\kappa}_{x,\mu}^f}} \right] \left[\prod_{\mu} \prod_{a=1}^{k_{x,-\mu}^f} \psi_x^{f,j_{x,-\mu}^{a+\kappa_{x,-\mu}^f}} \right] \end{aligned} \quad (\text{A.19})$$

and a new sign factor

$$\begin{aligned} \sigma_x^{(2)} &\equiv (-)^{p_4+p_5} \quad \text{with} \quad p_4 \equiv \sum_{f=1}^{N_f-1} k_x^{f+} \sum_{f'=f+1}^{N_f} w_x^{f'} + \sum_{f=1}^{N_f-1} \bar{k}_x^{f-} \sum_{f'=f+1}^{N_f} (w_x^{f'} + k_x^{f'+}) \\ &\quad \text{and} \quad p_5 \equiv p_4(k \leftrightarrow \bar{k}), \end{aligned} \quad (\text{A.20})$$

where

$$k_x^{f\pm} \equiv \sum_{\mu} k_{x,\pm\mu}^f \quad \text{and} \quad \bar{k}_x^{f\pm} \equiv \sum_{\mu} \bar{k}_{x,\pm\mu}^f. \quad (\text{A.21})$$

Since $\bar{\Psi}_x^f$ and Ψ_x^f have the same Grassmann parity due to Eq. (A.11), we can use Eq. (A.2) to rewrite Eq. (A.18) as

$$\Psi = \prod_x \sigma_x^{(2)} \prod_f \bar{\Psi}_x^f \Psi_x^f. \quad (\text{A.22})$$

We can now integrate out the original Grassmann variables on each site flavor by flavor. For a single flavor on a single site we have

$$\int d\psi d\bar{\psi} \prod_{a=1}^A \bar{\psi}^{i_a} \prod_{b=1}^B \psi^{j_b} = \int \prod_{k=1}^{N_c} d\psi^k \prod_{\ell=1}^{N_c} d\bar{\psi}^\ell \prod_{a=1}^A \bar{\psi}^{i_a} \prod_{b=1}^B \psi^{j_b} = \delta_{A,N_c} \varepsilon_{i_1 \dots i_A} \cdot \delta_{B,N_c} \varepsilon_{j_1 \dots j_B}, \quad (\text{A.23})$$

where we have omitted the site and flavor indices for simplicity. The Kronecker deltas δ_{A,N_c} and δ_{B,N_c} are referred to as Grassmann conditions. To apply the integral Eq. (A.23) to Eq. (A.22) we move each commuting pair of differentials $d\bar{\psi}_x^f d\psi_x^f$ in Eq. (A.1) in front of the corresponding $\bar{\Psi}_x^f \Psi_x^f$. After integrating $\bar{\psi}_x^f$ and ψ_x^f for every flavor f and every site x we obtain

$$\int \left[\prod_x d\psi_x d\bar{\psi}_x \right] \mathcal{G} = \int_{\chi} \prod_x K_x \sigma_x \prod_f \frac{(2m_f)^{w_x^f}}{w_x^f!} E_x^f \Theta(w_x^f) \delta_{h_x^{f,\text{in}}, h_x^{f,\text{out}}}, \quad (\text{5.1.1})$$

where

$$\sigma_x = \sigma_x^{(1)} \sigma_x^{(2)} \quad (\text{A.24})$$

is the combined sign factor on site x and

$$\begin{aligned} E_x^f = & \varepsilon_{v_{x,1}^{f,1}, \dots, v_{x,w_x^f}^{f,w_x^f}, i_{x,1}^{f,1}+1, \dots, i_{x,1}^{f,1}+k_{x,1}^f, \dots, i_{x,d}^{f,d}+1, \dots, i_{x,d}^{f,d}+k_{x,d}^f, \ell_{x,-1}^{f,-1}+1, \dots, \ell_{x,-1}^{f,-1}+\bar{k}_{x,-1}^f, \dots, \ell_{x,-d}^{f,-d}+1, \dots, \ell_{x,-d}^{f,-d}+\bar{k}_{x,-d}^f} \\ & \times \varepsilon_{v_{x,1}^{f,1}, \dots, v_{x,w_x^f}^{f,w_x^f}, m_{x,1}^{f,1}+1, \dots, m_{x,1}^{f,1}+\bar{k}_{x,1}^f, \dots, m_{x,d}^{f,d}+1, \dots, m_{x,d}^{f,d}+\bar{k}_{x,d}^f, j_{x,-1}^{f,-1}+1, \dots, j_{x,-1}^{f,-1}+k_{x,-1}^f, \dots, j_{x,-d}^{f,-d}+1, \dots, j_{x,-d}^{f,-d}+k_{x,-d}^f}, \end{aligned} \quad (\text{5.1.7})$$

is the result of the integration over the original Grassmann variables for flavor f .

B Initial tensor for $n_{\text{max}} = 0$, $N_f = 1$, $N_c = 3$, and $d = 2$

In the following appendix, we show that the initial tensor, derived in Part III, reduces to the tensor given in [23] when $n_{\text{max}} = 0$, $N_f = 1$, $N_c = 3$, and $d = 2$. Therefore, this serves as a consistency check.

In the full tensor network

$$Z = \sum_j \int_{\chi} \prod_x T_x K_x, \quad (\text{5.2.15})$$

we first consider the simplification of the tuple $\mathbf{j} \equiv (j_{x,\mu} | \forall x, \mu)$: Using the criterion (5.2.17) for $n_{\text{max}} = 0$, it immediately follows that all edge variables are zero, i.e., $n_{x,\mu}^{\pm\nu} = \bar{n}_{x,\mu}^{\pm\nu} = 0$. Therefore, the compound link index reduces to $j_{x,\mu} = ((k_{x,\mu}, \bar{k}_{x,\mu}), (0, 0, 0, 0), r_{x,\mu})$. Here, we also used $k_{x,\mu} = k_{x,\mu}^1$ and $\bar{k}_{x,\mu} = \bar{k}_{x,\mu}^1$, which follow from Eq. (4.2.2) for $N_f = 1$. Furthermore, the replacement of edge variables inside the numerical tensor

$$T_x \equiv \Delta_x^{\text{P}} \bar{T}_x(\mathbf{n}_x \rightarrow \text{edge variables}) \quad (\text{5.2.16})$$

is trivial and we obtain $\Delta_x^P = 1$. For one flavor, the tensor T_x from Eq. (5.2.10) reduces to

$$T_x = \sigma_x C_x W_x \left[\frac{(2m_1)^{w_x^1}}{w_x^1!} \Theta(w_x^1) \delta_{h_x^{1,\text{in}}, h_x^{1,\text{out}}} \right] \left[\prod_{\mu} \text{sgn}(\mathcal{L}_{x,\mu}) \right] \prod_{\mu=\pm 1}^{\pm 2} \sqrt{|\mathcal{L}_{x,\mu}|} \quad (\text{B.1})$$

with weight

$$W_x = \left[\prod_{\mu} \eta_{x,\mu}^{F_{x,\mu}} \right] \left[\prod_{\mu=\pm 1}^{\pm 2} \sqrt{\frac{e^{\mu_1(k_{x,\mu} - \bar{k}_{x,\mu})\delta_{|\mu|,1}}}{k_{x,\mu}! \bar{k}_{x,\mu}!}} \right]. \quad (\text{B.2})$$

The sign $\sigma_x = \sigma_x^{(1)} \sigma_x^{(2)}$ is given in Eq. (A.24), with $\sigma_x^{(1)}$ and $\sigma_x^{(2)}$ defined in Eq. (A.9) and Eq. (A.20), respectively. In two dimensions and for one flavor they reduce to

$$\begin{aligned} \sigma_x^{(1)} &= (-)^{p_1+p_2+p_3} \quad \text{with} \quad p_1 = F_{x,-1} \bar{k}_{x,-2}, \\ p_2 &= F_{x,1} (F_{x,2} + k_{x,2}), \\ p_3 &= (F_{x,1} + F_{x,2}) (F_{x,-1} + \bar{k}_{x,-1} + F_{x,-2} + \bar{k}_{x,-2}), \end{aligned} \quad (\text{B.3})$$

and $\sigma_x^{(2)} = 1$.

The Grassmann contribution K_x in two dimensions reduces to

$$K_x = (\chi_{x,1})^{F_{x,1}} (\chi_{x,2})^{F_{x,2}} (d\chi_{x,-2})^{F_{x,-2}} (d\chi_{x,-1})^{F_{x,-1}}, \quad (\text{B.4})$$

with Grassmann parity

$$F_{x,\mu} = \begin{cases} 0 & \text{if } k_{x,\mu} + \bar{k}_{x,\mu} \text{ even,} \\ 1 & \text{if } k_{x,\mu} + \bar{k}_{x,\mu} \text{ odd,} \end{cases} \quad (\text{B.5})$$

which is precisely the Grassmann contribution given in [23].

Next, we will show that the auxiliary summation index $r_{x,\mu} = (\alpha_{x,\mu}, \pi_{x,\mu}, \sigma_{x,\mu})$ has range one. To do so, we first observe that the gauge-link occupation numbers $g_{x,\mu}$ and $\bar{g}_{x,\mu}$, defined in Eq. (4.2.3), simplify in the case of strong coupling to $g_{x,\mu} = k_{x,\mu}$ and $\bar{g}_{x,\mu} = \bar{k}_{x,\mu}$, from which we follow

$$p_{x,\mu} = \min(g_{x,\mu}, \bar{g}_{x,\mu}) = \min(k_{x,\mu}, \bar{k}_{x,\mu}) \quad \text{and} \quad q_{x,\mu} = |k_{x,\mu} - \bar{k}_{x,\mu}|/N_c. \quad (\text{B.6})$$

Furthermore, using the definition (4.2.4a), we find $\kappa_{x,\mu}^1 = \bar{\kappa}_{x,\mu}^1 = 0$ and $\kappa_{x,\mu}^2 = k_{x,\mu}$, $\bar{\kappa}_{x,\mu}^2 = \bar{k}_{x,\mu}$, which yields for $\varkappa$, defined in Eq. (4.3.8), $\varkappa_{x,\mu}^1 = \bar{\varkappa}_{x,\mu}^1 = 0$ and

$$\varkappa_{x,\mu}^2 \equiv \begin{cases} k_{x,\mu} & \text{for } k_{x,\mu} \geq \bar{k}_{x,\mu}, \\ \bar{k}_{x,\mu} & \text{for } k_{x,\mu} < \bar{k}_{x,\mu}, \end{cases} \quad \bar{\varkappa}_{x,\mu}^2 \equiv \begin{cases} \bar{k}_{x,\mu} & \text{for } k_{x,\mu} \geq \bar{k}_{x,\mu}, \\ k_{x,\mu} & \text{for } k_{x,\mu} < \bar{k}_{x,\mu}. \end{cases} \quad (\text{B.7})$$

This expression can be abbreviated as

$$\varkappa_{x,\mu}^2 = \max(k_{x,\mu}, \bar{k}_{x,\mu}), \quad \bar{\varkappa}_{x,\mu}^2 = \min(k_{x,\mu}, \bar{k}_{x,\mu}), \quad (\text{B.8})$$

which is equivalent to

$$\varkappa_{x,\mu}^2 = q_{x,\mu} N_c + p_{x,\mu}, \quad \bar{\varkappa}_{x,\mu}^2 = p_{x,\mu}, \quad (\text{B.9})$$

$\bar{j}_{x,\mu}$	0	1	2	3	4	5
$k_{x,\mu}$	0	1	2	3	3	0
$\bar{k}_{x,\mu}$	0	1	2	3	0	3
$\mathcal{L}_{x,\mu}$	1	1/3	1/3	1	1/6	1/6

Table B.1: Major index $\bar{j}_{x,\mu}$ and corresponding link weight $\mathcal{L}_{x,\mu}$ for $N_f = 1$, $N_c = 3$, and $n_{\max} = 0$.

where we used $\max(g_{x,\mu}, \bar{g}_{x,\mu}) = q_{x,\mu}N_c + p_{x,\mu}$.

To determine the range of the sum¹

$$\sum_r \equiv \sum_{\alpha \in A^\sim} \sum_{\pi, \sigma \in S_p^\sim}, \quad (4.3.13)$$

we recall from Sec. 4.3.2 that for $N_f = 1$ two groupings $\alpha, \alpha' \in A$ are equivalent if the pairs $\alpha, \gamma(\alpha)$ and $\alpha', \gamma(\alpha')$, defined in Eq. (4.3.4), can be transformed into each other by permuting the entries α_s^c and γ_s satisfying $\varkappa^1 < \alpha_s^c, \gamma_s \leq \varkappa^2$, i.e., permuting the entries satisfying $0 < \alpha_s^c, \gamma_s \leq qN_c + p$. This is satisfied for all numbers. Therefore α and α' are equivalent for all $\alpha, \alpha' \in A$ and the complete set of representatives $A^\sim$ contains only one element, which can be chosen arbitrarily from A . We choose canonical ordering $\alpha_s^c \equiv (s-1)N_c + c$ ($1 \leq s \leq q, 1 \leq c \leq N_c$) and therefore $\gamma_s \equiv qN_c + s$ ($1 \leq s \leq p$).

Likewise, we recall from Sec. 4.3.2 that for one flavor two permutations π and π' out of S_p are equivalent if the corresponding sequences can be transformed into each other by permuting the numbers π_s with $\bar{\varkappa}^1 < \pi_s \leq \bar{\varkappa}^2$ and by permuting the numbers s with $0 < s \leq z^1$. Here, z^1 is the number of entries of γ for which $\varkappa^1 < \gamma_s \leq \varkappa^2$, which simplifies to $0 < \gamma_s \leq qN_c + p$. This is fulfilled for all numbers γ_s . Since s runs from 1 to p , we find $z^1 = p$.

When we additionally use the simplified expressions, for $\bar{\varkappa}^1$ and $\bar{\varkappa}^2$, we find that π and π' are equivalent if the corresponding sequences can be transformed into each other by permuting the numbers π_s with $0 < \pi_s \leq p$ and by permuting the numbers s with $0 < s \leq p$. Consequently all pairs of permutations π and π' out of S_p are equivalent and a complete set of representatives $S_p^\sim$ contains only one element, which can be chosen arbitrarily from S_p . For convenience, we choose the identity permutation $\pi_s \equiv s$ (and also $\sigma_s \equiv s$). This shows, that the auxiliary sum index $r_{x,\mu}$ has range 1.

With this result, we can explicitly state all values of the major index $j_{x,\mu}$ with one further observation: Due to $g_{x,\mu} = k_{x,\mu}$ and $\bar{g}_{x,\mu} = \bar{k}_{x,\mu}$ for $N_c = 3$, the SU(3) condition reduces to $(k_{x,\mu} - \bar{k}_{x,\mu})/3 \in \mathbb{Z}$. As $k_{x,\mu}$ and $\bar{k}_{x,\mu}$ now both run from 0 to $N_c = 3$, there are six different combinations - labeled by $\bar{j}_{x,\mu}$ - that fulfill this condition, which are given in Table B.1. This is consistent with the initial index given in [23].² For later usage, we call a link with $\bar{j}_{x,\mu} = 0$ neutral link, a link with $\bar{j}_{x,\mu} \in \{1, 2, 3\}$ mesonic link, and a link with $\bar{j}_{x,\mu} = 4$ (5) forward (backward) baryonic link, respectively.

¹As before, the link index (x, μ) is neglected for readability.

²Note that for $\bar{j}_{x,\mu} = 4$ and $\bar{j}_{x,\mu} = 5$ the definition of $k_{x,\mu}$ used in this thesis differs from [23].

Next, we calculate the link weight $\mathcal{L}$, defined in Eq. (4.3.12). Since all elements of A are equivalent, we find $[\alpha] = A$ and therefore $||[\alpha]|| = \frac{(qN_c+p)!}{p!q!(N_c!)^q}$. Similarly, as all elements π (and σ) are equivalent we find $[\pi] = [\sigma] = S_p$. Therefore, the link weight reduces to

$$\mathcal{L}_{x,\mu} = \left[\prod_{s=1}^2 \frac{s!}{(s+q)!} \right] \frac{(3q+p)!}{p!q!(3!)^q} \sum_{\pi', \sigma' \in S_p} \text{sgn}(\pi') \text{sgn}(\sigma') \widetilde{\text{Wg}}_{N_c}^{q,p}(\pi' \cdot \sigma'^{-1}), \quad (\text{B.10})$$

where we also used³

$$\text{sgn}(\pi'^{-1} \cdot \pi) = \text{sgn}(\pi'^{-1}) \text{sgn}(\pi) = \text{sgn}(\pi'^{-1}) = \text{sgn}(\pi'), \quad (\text{B.11})$$

as well as the analogue equation for σ and σ' . The prefactor of $\mathcal{L}_{x,\mu}$ for relevant values of q and p reads

$$\left[\prod_{s=1}^2 \frac{s!}{(s+q)!} \right] \frac{(3q+p)!}{p!q!(3!)^q} = \begin{cases} 1 & \text{for } q = 0 \\ \frac{1}{3!} & \text{for } q = 1 \text{ and } p = 0, \end{cases} \quad (\text{B.12})$$

and the result for $\mathcal{L}_{x\mu}$ is given for all six values of the major index $\bar{j}_{x,\mu}$ in Table B.1.

Another ingredient in $\bar{T}_x$ from Eq. (B.1) that needs to be calculated, is the color factor C_x . For $N_f = 1$ it reduces to $C_x = I_x E_x^1$, with $I_x = \prod_{\mu} I_{x,\mu}^b I_{x,-\mu}^e$ and

$$\begin{aligned} E_x^1 = & \varepsilon_{v_x^{1,1}, \dots, v_x^{1,w_x^1}, i_{x,1}^{1,1}, \dots, i_{x,1}^{k_{x,1},1}, i_{x,2}^{1,1}, \dots, i_{x,2}^{k_{x,2},2}, \ell_{x,-1}^{1,1}, \dots, \ell_{x,-1}^{\bar{k}_{x,-1},-1}, \ell_{x,-2}^{1,1}, \dots, \ell_{x,-2}^{\bar{k}_{x,-2},-2}} \\ & \times \varepsilon_{v_x^{1,1}, \dots, v_x^{1,w_x^1}, m_{x,1}^1, \dots, m_{x,1}^{\bar{k}_{x,1},1}, m_{x,2}^1, \dots, m_{x,2}^{\bar{k}_{x,2},2}, j_{x,-1}^1, \dots, j_{x,-1}^{k_{x,-1},-1}, j_{x,-2}^1, \dots, j_{x,-2}^{k_{x,-2},-2}}, \end{aligned} \quad (\text{B.13})$$

where we used $d = 2$ and $\kappa_{x,\mu}^1 = \bar{\kappa}_{x,\mu}^1 = 0$.

It can be shown by explicit calculation, that the condition $\Theta(w_x^1) \delta_{h_x^{1,\text{in}}, h_x^{1,\text{out}}}$ (with $w_x^1 = 3 - h_x^{1,\text{in}}$) inside the initial tensor, together with the six initial link occupations, only results in nonzero tensor entries, when either all four links are neutral or mesonic links (we call such a site a “mesonic site” and denote it by $\delta_{x \in \mathcal{M}}$), or when there is exactly one incoming and one outgoing baryonic link as well as two neutral links (in this case, we call the site a “baryonic site” and denote it by $\delta_{x \in \mathcal{B}}$).⁴ Therefore, the tensor network can be written as

$$Z = \sum_{\bar{j}} \int_{\chi} \prod_x \left(\delta_{x \in \mathcal{M}} \bar{T}_x^{\text{M}} + \delta_{x \in \mathcal{B}} \bar{T}_x^{\text{B}} \right) K_x, \quad (\text{B.14})$$

with $\bar{j} \equiv (\bar{j}_{x,\mu} | \forall x, \mu)$.

In the following we calculate the contributions $\bar{T}_x^{\text{M}}$ and $\bar{T}_x^{\text{B}}$.

mesonic site: When the site x is mesonic it holds that $q_{x,\mu} = 0$ for all four links adjacent to this site. Therefore, I^b and I^e defined in Eq. (5.2.2) and Eq. (4.3.5) do not contain a Levi-Civita tensor and can be written, using $g_{x,\mu} = \bar{g}_{x,\mu}$, as

$$I_{x,\mu}^b = \prod_{s=1}^{p_{x,\mu}} \delta_{m_{x,\mu}^s, i_{x,\mu}^s} \quad \text{and} \quad I_{x,\mu}^e = \prod_{s=1}^{p_{x,\mu}} \delta_{\ell_{x,\mu}^s, j_{x,\mu}^s}, \quad (\text{B.15})$$

³Recall that π is the identity permutation.

⁴For a given site x , all backward baryonic links that are positioned in forward direction are said to be “incoming” baryonic links and all forward baryonic links that are positioned in backward direction are said to be “incoming” baryonic links. For “outgoing” baryonic links it is the other way around.

where we used the canonical form of the representatives: $\gamma_s = qN_c + s = s$ (α is absent) and $\pi_s = \sigma_s = s$ at this link.

After we have inserted the expression

$$\prod_{\mu} I_{x,\mu}^b I_{x,-\mu}^e = \prod_{s=1}^{p_{x,1}} \delta_{m_{x,1}, i_{x,1}^s} \prod_{s=1}^{p_{x,-1}} \delta_{\ell_{x,-1}, j_{x,-1}^s} \prod_{s=1}^{p_{x,2}} \delta_{m_{x,2}, i_{x,2}^s} \prod_{s=1}^{p_{x,-2}} \delta_{\ell_{x,-2}, j_{x,-2}^s} \quad (\text{B.16})$$

in the color factor, we can, for example, contract all indices m and ℓ , yielding

$$C_x = \varepsilon_{v_x^{1,1}, \dots, v_x^{1,w_x^1}, i_{x,1}^1, \dots, i_{x,1}^{k_{x,1}}, i_{x,2}^1, \dots, i_{x,2}^{k_{x,2}}, j_{x,-1}^1, \dots, j_{x,-1}^{\bar{k}_{x,-1}}, j_{x,-2}^1, \dots, j_{x,-2}^{\bar{k}_{x,-2}}} \times \varepsilon_{v_x^{1,1}, \dots, v_x^{1,w_x^1}, i_{x,1}^1, \dots, i_{x,1}^{k_{x,1}}, i_{x,2}^1, \dots, i_{x,2}^{k_{x,2}}, j_{x,-1}^1, \dots, j_{x,-1}^{\bar{k}_{x,-1}}, j_{x,-2}^1, \dots, j_{x,-2}^{\bar{k}_{x,-2}}} = N_c! = 6, \quad (\text{B.17})$$

where we used $p_{x,\mu} = \bar{k}_{x,\mu} = k_{x,\mu}$ several times and $\varepsilon_{i_1, \dots, i_n} \varepsilon_{i_1, \dots, i_n} = n!$ for $N_c = 3$. The number of indices of the Levi-Civita tensor is guaranteed to be 3, since $w_x^1 + h_x^{1,\text{out}} = 3$, with $h_x^{1,\text{out}} = k_{x,1} + k_{x,2} + \bar{k}_{x,-1} + \bar{k}_{x,-2}$.

We now observe, using the values of $\mathcal{L}$ given in Table B.1 that

$$\frac{1}{k_{x,\nu}! \bar{k}_{x,\nu}!} |\mathcal{L}_{x,\mu}| = \frac{(3 - k_{x,\nu})!}{3! k_{x,\nu}!} \quad \text{for } k_{x,\nu} = \bar{k}_{x,\nu} \in \{0, 1, 2, 3\}, \quad (\text{B.18})$$

which allows us to write the contribution of a mesonic site as

$$\bar{T}_x^M \equiv 6 \left[\frac{(2m_1)^{w_x^1}}{w_x^1!} \right] \prod_{\mu=\pm 1}^{\pm 2} \sqrt{\frac{(3 - k_{x,\nu})!}{3! k_{x,\nu}!}}. \quad (\text{B.19})$$

Here we also used $\text{sgn}(\mathcal{L}_{x,\mu}) = 1$ for all values of $\bar{j}_{x,\mu}$. There are no signs arising from the staggered phases or $\sigma_x^{(1)}$, since $F_{x,\mu} = 0$ for all four links adjacent to site x . This is precisely the tensor entry for a mesonic site given in Ref. [23].

The remaining task is to calculate the weight of a baryonic site and to compare it with the weight stated in Ref. [23] as well.

baryonic site: For a baryonic site it holds that $p_{x,\mu} = 0$ for all four links adjacent to this site. Therefore, I^b and I^e defined in Eq. (5.2.2) and Eq. (4.3.5) do not contain a Kronecker delta and simplify to

$$(I_{\alpha\pi\sigma}^b)^{im} = \begin{cases} \varepsilon_{i_\alpha}^{\otimes q} & \text{for } g \geq \bar{g}, \\ \varepsilon_{m_\alpha}^{\otimes q} & \text{for } g < \bar{g}, \end{cases} \quad \text{and} \quad (I_{\alpha\pi\sigma}^e)^{j\ell} = \begin{cases} \varepsilon_{j_\alpha}^{\otimes q} & \text{for } g \geq \bar{g}, \\ \varepsilon_{\ell_\alpha}^{\otimes q} & \text{for } g < \bar{g}, \end{cases} \quad (\text{B.20})$$

where we again omitted the link index due to readability. For a baryonic link index $q = 1$ and there is exactly one Levi-Civita tensor in $I^{b/e}$. Due to the chosen canonical grouping α , the corresponding enumeration of different indices is increasing. For a baryonic site, there is exactly one incoming and one outgoing baryonic link, for which we get such a Levi-Civita tensor. Without loss of generality, we call the indices of these Levi-Civita tensors r^i and t^i with $i \in \{1, 2, 3\}$. Therefore we get

$$\prod_{\mu} I_{x,\mu}^b I_{x,-\mu}^e = \varepsilon_{r^1 r^2 r^3} \varepsilon_{t^1 t^2 t^3}. \quad (\text{B.21})$$

The same substitution can be made for

$$E_x^1 = \varepsilon_{i_{x,1}, \dots, i_{x,1}^{k_{x,1}}, i_{x,2}, \dots, i_{x,2}^{k_{x,2}}, \ell_{x,-1}, \dots, \ell_{x,-1}^{\bar{k}_{x,-1}}, \ell_{x,-2}, \dots, \ell_{x,-2}^{\bar{k}_{x,-2}}} \times \varepsilon_{m_{x,1}, \dots, m_{x,1}^{\bar{k}_{x,1}}, m_{x,2}, \dots, m_{x,2}^{\bar{k}_{x,2}}, j_{x,-1}, \dots, j_{x,-1}^{k_{x,-1}}, j_{x,-2}, \dots, j_{x,-2}^{k_{x,-2}}}, \quad (\text{B.22})$$

where we used $w_x^1 = 0^5$. Since the first Levi-Civita tensor refers to the outgoing baryonic link and the second refers to the incoming baryonic link, both Levi-Civita tensors have exactly three indices, which both have ascending enumeration. Therefore we get $E_x^1 = \varepsilon_{r^1 r^2 r^3} \varepsilon_{t^1 t^2 t^3}$, which yields the color factor

$$C_x = \varepsilon_{r^1 r^2 r^3} \varepsilon_{t^1 t^2 t^3} \varepsilon_{r^1 r^2 r^3} \varepsilon_{t^1 t^2 t^3} = 6^2. \quad (\text{B.23})$$

The remaining link weights are given by

$$\prod_{\mu=\pm 1}^{\pm 2} \sqrt{\frac{1}{k_{x,\mu}! \bar{k}_{x,\mu}!}} = \sqrt{\frac{1}{6^2}} = \frac{1}{6} \quad \text{and} \quad \prod_{\mu=\pm 1}^{\pm 2} \sqrt{|\mathcal{L}_{x,\mu}|} = \sqrt{\left(\frac{1}{6}\right)^2} = \frac{1}{6}, \quad (\text{B.24})$$

where again we used that there are two neutral links and two baryonic links on a baryonic site. This allows us write the full contribution of a baryonic site as

$$\bar{T}_x^B \equiv \sigma_x^{(1)} \left[\prod_{\mu} \eta_{x,\mu}^{F_{x,\mu}} \right] \left[\prod_{\mu=\pm 1}^{\pm 2} \sqrt{e^{\mu_1(k_{x,\mu} - \bar{k}_{x,\mu}) \delta_{|\mu|,1}}} \right], \quad (\text{B.25})$$

where we used $w_x^1 = 0$ once more.

As a very last step, we need to evaluate the sign factor $\sigma_x^{(1)}$, given in Eq. (A.9). To do so, we first notice that p_3 can be simplified, using

$$p_3 = (F_{x,1} + F_{x,2})(F_{x,-1} + F_{x,-2}) + (F_{x,1} + F_{x,2})(\bar{k}_{x,-1} + \bar{k}_{x,-2}). \quad (\text{B.26})$$

Since every contributing site is Grassmann even, it holds that $(F_{x,1} + F_{x,2})(F_{x,-1} + F_{x,-2}) \bmod 2 = F_{x,1} + F_{x,2} \bmod 2$ and therefore the replacement

$$p_3 \rightarrow F_{x,1} + F_{x,2} + (F_{x,1} + F_{x,2})(\bar{k}_{x,-1} + \bar{k}_{x,-2}) \quad (\text{B.27})$$

results in an unchanged sign factor $\sigma_x^{(1)}$.

Next, we define $l_{x,\mu}^{\pm}$, which is equal to one if and only if the link is a forward (backward) baryonic link. For a baryonic or neutral link, we can express $F_{x,\mu}$, $k_{x,\mu}$, and $\bar{k}_{x,\mu}$ via $l_{x,\mu}^{\pm}$ as

$$\begin{aligned} F_{x,\mu} &= (l_{x,\mu}^+ + l_{x,\mu}^-), \\ k_{x,\mu} &= 3l_{x,\mu}^+, \quad \text{and} \\ \bar{k}_{x,\mu} &= 3l_{x,\mu}^-. \end{aligned} \quad (\text{B.28})$$

⁵Since there is exactly one incoming and one outgoing baryonic link on a baryonic site, we have $h_x^{1,\text{in}} = h_x^{1,\text{out}} = 3$ and therefore $w_x^1 = 0$.

Since the sign $\sigma_x^{(1)}$ is only influenced by the parity of these quantities, we can also replace $k_{x,\mu} \rightarrow l_{x,\mu}^+$ and $\bar{k}_{x,\mu} \rightarrow l_{x,\mu}^-$ ⁶. Therefore we can replace the sign exponents using

$$\begin{aligned} p_1 &\rightarrow (l_{x,-1}^+ + l_{x,-1}^-)l_{x,-2}^-, \\ p_2 &\rightarrow (l_{x,1}^+ + l_{x,1}^-)(l_{x,2}^+ + l_{x,2}^- + l_{x,2}^+), \\ p_3 &\rightarrow (l_{x,1}^+ + l_{x,1}^- + l_{x,2}^+ + l_{x,2}^-)(1 + l_{x,-1}^- + l_{x,-2}^-). \end{aligned} \quad (\text{B.29})$$

We now gather parts of these exponents inside the factor

$$\omega_x = (-1)^{l_{x,1}^+ + l_{x,2}^+ + l_{x,2}^- + l_{x,1}^- + l_{x,-2}^- + l_{x,-1}^+}, \quad (\text{B.30})$$

which is part of the baryonic initial tensor as given in Ref. [23] and show that the remaining exponent

$$\begin{aligned} l_{x,-1}^- l_{x,-2}^- + l_{x,1}^+ (l_{x,2}^+ + l_{x,2}^- + l_{x,2}^+) + l_{x,1}^- (l_{x,2}^- + l_{x,2}^+) \\ + l_{x,1}^- + l_{x,2}^- + (l_{x,1}^+ + l_{x,1}^- + l_{x,2}^+ + l_{x,2}^-)(l_{x,-1}^- + l_{x,-2}^-) \end{aligned} \quad (\text{B.31})$$

is even for any baryonic site. To do so, we use that for any baryonic site there is exactly one incoming and exactly one outgoing baryonic link, i.e., only one variable out of $l_{x,1}^-$, $l_{x,2}^-$, $l_{x,-1}^+$ and $l_{x,-2}^+$ is one, and only one variable out of $l_{x,1}^+$, $l_{x,2}^+$, $l_{x,-1}^-$ and $l_{x,-2}^-$ is one, while all other variables are zero. Therefore, any product of two variables $l_{x,\pm\mu}^\pm$ in B.31 where both factors belong to the same type of baryonic link (i.e., incoming or outgoing) is zero. The remaining exponential reads

$$l_{x,1}^+ l_{x,2}^- + l_{x,1}^- l_{x,2}^+ + l_{x,1}^- + l_{x,2}^- + (l_{x,1}^- + l_{x,2}^-)(l_{x,-1}^- + l_{x,-2}^-). \quad (\text{B.32})$$

We now observe that all terms contain either $l_{x,1}^-$ or $l_{x,2}^-$, which allows us to factorize the exponent as

$$l_{x,1}^- (1 + l_{x,2}^+ + l_{x,-1}^- + l_{x,-2}^-) + l_{x,2}^- (1 + l_{x,1}^+ + l_{x,-1}^- + l_{x,-2}^-) = l_{x,1}^- (1 + 1 - l_{x,1}^+) + l_{x,2}^- (1 + 1 - l_{x,2}^+), \quad (\text{B.33})$$

where we again used that for any baryonic site there is exactly one outgoing baryonic link, i.e., $l_{x,1}^+ + l_{x,2}^+ + l_{x,-1}^- + l_{x,-2}^- = 1$. The resulting expression is even, since $l_{x,1}^- l_{x,1}^+ = l_{x,2}^- l_{x,2}^+ = 0$ for any baryonic site. Thus we can write

$$\bar{T}_x^B \equiv \omega_x \left[\prod_{\mu} \eta_{x,\mu}^{F_{x,\mu}} \right] \left[\prod_{\mu=\pm 1}^{\pm 2} \sqrt{e^{\mu_1(k_{x,\mu} - \bar{k}_{x,\mu})\delta_{|\mu|,1}}} \right]. \quad (\text{B.34})$$

C Proofs for OS-GHOTRG

C.1 Proof of tensor decomposition after contraction

In this appendix we prove that (7.3.1) is reproduced after one contraction step, before the OS-GHOTRG truncation is applied. This also yields the occupation numbers defined in (7.3.8). We

⁶Formally, we use $(-1)^{ak_{x,\mu}} = (-1)^{al_{x,\mu}^+} = (-1)^{al_{x,\mu}^-}$ for any a and the equivalent equation involving $\bar{k}_{x,\mu}$ and $l_{x,\mu}^-$.

assume that (7.3.1) holds prior to a contraction step and substitute it in the contraction (7.3.5),

$$(T_X)_{j_X} = \sum_{n_x^s=0}^{n_{\max}} \sum_{n_y^s=0}^{n_{\max}} \sum_{j_{x,\mu}} \beta^{n_x^s+n_x+n_y^s+n_y} (T_x^{n_x^s})_{j_x} (T_y^{n_y^s})_{j_y} \sigma_{F_x, F_y}, \quad (\text{C.1})$$

where we dropped the arguments of n_x and n_y for simplicity. Assuming that (7.3.3) holds for the current fine-grid tensors, we find

$$\begin{aligned} & n_x^s + n_x + n_y^s + n_y \\ &= n_x^s + n_y^s + \frac{1}{2} \sum_{\nu=\pm 1}^{\pm d} (\ell_{x,\nu} + \ell_{y,\nu}) + \frac{1}{8} \sum_{\substack{\nu, \lambda=\pm 1 \\ |\nu| \neq |\lambda|}}^{\pm d} (\hat{n}_{x,\nu}^\lambda + \hat{n}_{y,\nu}^\lambda) \\ &= \underbrace{n_x^s + n_y^s + \frac{1}{2}(\ell_{x,\mu} + \ell_{y,-\mu})}_{n_X^s} + \frac{1}{2} \underbrace{(\ell_{x,-\mu} + \ell_{y,\mu})}_{\ell_{X,\mu}} + \frac{1}{2} \sum_{\substack{\nu=\pm 1 \\ |\nu| \neq \mu}}^{\pm d} \underbrace{\left[\ell_{x,\nu} + \ell_{y,\nu} + \frac{1}{4}(\hat{n}_{x,\nu}^\mu + \hat{n}_{y,\nu}^{-\mu} + \hat{n}_{x,\mu}^\nu + \hat{n}_{y,-\mu}^\nu) \right]}_{\ell_{X,\nu}} \\ &\quad + \frac{1}{8} \sum_{\substack{\nu=\pm 1 \\ |\nu| \neq \mu}}^{\pm d} \underbrace{(\hat{n}_{x,-\mu}^\nu + \hat{n}_{y,\mu}^\nu + \hat{n}_{x,\nu}^{-\mu} + \hat{n}_{y,\nu}^\mu)}_{\hat{n}_{X,\mu}^\nu} + \frac{1}{8} \sum_{\substack{\nu, \lambda=\pm 1 \\ |\nu|, |\lambda| \neq \mu \\ |\nu| \neq |\lambda|}}^{\pm d} \underbrace{(\hat{n}_{x,\nu}^\lambda + \hat{n}_{y,\nu}^\lambda)}_{\hat{n}_{X,\nu}^\lambda}, \end{aligned} \quad (\text{C.2})$$

where the underbraces are identified with the site, link, and edge occupation numbers on the coarse grid. Note that $(x, \mu) = (y, -\mu)$. When the hook conditions in (7.2.3) are satisfied, i.e., $\Delta_x \neq 0$ and $\Delta_y \neq 0$, we define the coarse-grid occupations as given in Eq. (7.3.8), where the link occupation $\ell_{X,\nu}$ in the third line of (7.3.8b) was redefined using the hook conditions. We also use these definitions when the hook conditions are not satisfied and the tensor entries are zero. Note that $\ell_{X,\nu}$ and $\hat{n}_{X,\nu}^\lambda$ are determined uniquely by $j_{X,\nu}$. The exponent of β in (C.1) has the form $n_X^s + n_X$ if we define

$$n_X(j_X) \equiv \frac{1}{2} \sum_{\nu=\pm 1}^{\pm d} \ell_{X,\nu} + \frac{1}{8} \sum_{\substack{\nu, \lambda=\pm 1 \\ |\nu| \neq |\lambda|}}^{\pm d} \hat{n}_{X,\nu}^\lambda, \quad (\text{C.3})$$

which is consistent with (7.3.3). Next, we split the sum over the contracted link index $j_{x,\mu}$ in (C.1) into partial sums which run over the values of $j_{x,\mu}$ with the same link occupation $\ell_{x,\mu}(j_{x,\mu})$,

$$(T_X)_{j_X} = \beta^{n_X} \sum_{n_x^s=0}^{n_{\max}} \sum_{n_y^s=0}^{n_{\max}} \sum_{m=0}^{n_{\max}} \beta^{n_X^s} \sum_{\substack{j_{x,\mu} \\ \ell_{x,\mu}=m}} (T_x^{n_x^s})_{j_x} (T_y^{n_y^s})_{j_y} \sigma_{F_x, F_y}, \quad (\text{C.4})$$

where we could pull out β^{n_X} since n_X does not depend on any of the summation variables and where the sum over m terminates at $n_{\max}$ since this is the maximum link occupation by construction. We replace the sum over m by a sum over

$$n_X^s = n_x^s + n_y^s + m, \quad (\text{C.5})$$

see (7.3.8a) with $\ell_{x,\mu} = m$, and then change the order of the sums and adapt the ranges of the sums over n_x^s and n_y^s accordingly. Since we consider the strong-coupling expansion only up to $\beta^{n_{\max}}$, we truncate the resulting sum over n_X^s at $n_{\max}$ to obtain

$$(T_X)_{j_X} = \beta^{n_X} \sum_{n_X^s=0}^{n_{\max}} \beta^{n_X^s} \sum_{n_x^s=0}^{n_X^s} \sum_{n_y^s=0}^{n_X^s - n_x^s} \sum_{\substack{j_{x,\mu} \\ \ell_{x,\mu}=n_X^s - n_x^s - n_y^s}} (T_x^{n_x^s})_{j_x} (T_y^{n_y^s})_{j_y} \sigma_{F_x, F_y}, \quad (\text{C.6})$$

which shows that the tensor on the coarse site X can be brought into the form (7.3.1), with $(T_X^{n_X^s})_{j_X}$ given by (7.3.7), the occupation numbers on the coarse grid defined in (7.3.8), and n_X defined in (C.3).

Finally, it is straightforward to show that (7.3.2) is satisfied after the contraction step (by using Δ_x , Δ_y , and (7.3.8), or by using Fig. 7.3 and invoking the hook conditions).

C.2 Proof of link and site criteria

After a contraction step, the partition function has the form of (5.2.15) with x replaced by X . Consider the product

$$\prod_X (T_X^{n_X^s})_{j_X} \quad (\text{C.7})$$

for a single valid configuration $\mathbf{j}$. This configuration will contribute to the term of order β^N in the partition function. Using (7.3.1) and (7.3.3) we have

$$N = \sum_X (n_X^s + n_X) = \sum_X \left(n_X^s + \frac{1}{2} \sum_{\mu=\pm 1}^{\pm d} \ell_{X,\mu} + \frac{1}{8} \sum_{\substack{\mu,\nu=\pm 1 \\ |\mu| \neq |\nu|}}^{\pm d} \hat{n}_{X,\mu}^\nu \right) = \sum_X \left(n_X^s + \sum_{\mu=1}^d \ell_{X,\mu} + \frac{1}{2} \sum_{\substack{\mu,\nu=1 \\ \mu \neq \nu}}^d \hat{n}_{X,\mu}^\nu \right), \quad (\text{C.8})$$

where in the last step we shifted some sites and used the link identity $(X, -\mu) = (X - \hat{\mu}, \mu)$ as well as the hook condition (7.2.3).

The link criterion (7.4.1) is derived from (C.8) for $\mu > 0$ as follows,

$$\begin{aligned} N &\geq \ell_{X,\mu} + \frac{1}{2} \sum_{\substack{\nu=1 \\ \nu \neq \mu}}^d (\hat{n}_{X,\mu}^\nu + \hat{n}_{X,\nu}^\mu) + \frac{1}{2} \sum_{\substack{\nu=1 \\ \nu \neq \mu}}^d (\hat{n}_{X-\hat{\nu},\mu}^\nu + \hat{n}_{X-\hat{\nu},\nu}^\mu) \\ &= \ell_{X,\mu} + \sum_{\substack{\nu=1 \\ \nu \neq \mu}}^d \hat{n}_{X,\mu}^\nu + \sum_{\substack{\nu=1 \\ \nu \neq \mu}}^d \hat{n}_{X,\mu}^{-\nu} = \ell_{X,\mu} + \sum_{\substack{\nu=\pm 1 \\ |\nu| \neq |\mu|}}^{\pm d} \hat{n}_{X,\mu}^\nu, \end{aligned} \quad (\text{C.9})$$

where in the first line we picked a subset of the terms in (C.8), and to obtain the second line we applied the link identity and the hook condition. Note that we derived (C.9) for $\mu > 0$, but using the link identity it is clear that it also holds for $\mu < 0$. To have $N \leq n_{\max}$, the link criterion (7.4.1) must be satisfied. Note that the link criterion is identical for all $T_X^{n_X^s}$.

The site criterion (7.4.2) is derived from (C.8) in a similar way,

$$\begin{aligned} N &\geq n_X^s + \sum_{\mu=1}^d (\ell_{X,\mu} + \ell_{X-\hat{\mu},\mu}) + \frac{1}{2} \sum_{\substack{\mu,\nu=1 \\ \mu \neq \nu}}^d (\hat{n}_{X,\mu}^\nu + \hat{n}_{X-\hat{\mu},\mu}^\nu + \hat{n}_{X-\hat{\nu},\mu}^\nu + \hat{n}_{X-\hat{\mu}-\hat{\nu},\mu}^\nu) \\ &= n_X^s + \sum_{\mu=1}^d (\ell_{X,\mu} + \ell_{X,-\mu}) + \frac{1}{2} \sum_{\substack{\mu,\nu=1 \\ \mu \neq \nu}}^d (\hat{n}_{X,\mu}^\nu + \hat{n}_{X,-\mu}^\nu + \hat{n}_{X,\mu}^{-\nu} + \hat{n}_{X,-\mu}^{-\nu}) \\ &= n_X^s + \sum_{\mu=\pm 1}^{\pm d} \ell_{X,\mu} + \sum_{\substack{\mu,\nu=\pm 1 \\ |\mu| < |\nu|}}^{\pm d} \hat{n}_{X,\mu}^\nu. \end{aligned} \quad (\text{C.10})$$

The arguments for the first and the second line are the same as in the derivation of the link criterion. In the last we have also ordered $|\mu| < |\nu|$. To have $N \leq n_{\max}$, the site criterion (7.4.2) must be satisfied. Note that the site criterion is applied to each $T_X^{n_X}$ individually.

C.3 Proof of tensor decomposition after tracing

The manipulations in this section closely parallel those in Appendix C.1. Inserting (7.3.1) in (7.6.1) yields

$$(T_X)_{j_X} = \sum_{n_x^s=0}^{n_{\max}} \sum_{j_{x,\mu}} \beta^{n_x^s+n_x} (T_x^{n_x^s})_{j_x} \sigma_\mu, \quad (\text{C.11})$$

where here and below we have $j_{x,-\mu} = j_{x,\mu}$. Using (7.3.3) and separating the occupation numbers related to the tracing direction μ we obtain

$$n_x^s + n_x = \underbrace{s_x + \frac{1}{2}(\ell_{x,\mu} + \ell_{x,-\mu})}_{s_X} + \frac{1}{2} \sum_{\substack{\nu=\pm 1 \\ |\nu| \neq \mu}}^{\pm d} \underbrace{\left[\ell_{x,\nu} + \frac{1}{4}(\hat{n}_{x,\nu}^\mu + \hat{n}_{x,\nu}^{-\mu} + \hat{n}_{x,\mu}^\nu + \hat{n}_{x,-\mu}^\nu) \right]}_{\ell_{X,\nu}} + \frac{1}{8} \sum_{\substack{\nu,\lambda=\pm 1 \\ |\nu|,|\lambda| \neq \mu \\ |\nu| \neq |\lambda|}}^{\pm d} \underbrace{\hat{n}_{x,\nu}^\lambda}_{\hat{n}_{X,\nu}^\lambda}, \quad (\text{C.12})$$

where the underbraces are identified with the site, link, and edge occupation numbers on the $(d-1)$ -dimensional lattice. Using the link and hook conditions we define these occupations in Eq. (7.6.3). As before, we also use these definitions when the hook conditions are not satisfied and the tensor entries are zero. Note that $\ell_{X,\nu}$ and $\hat{n}_{X,\nu}^\lambda$ are determined uniquely by $j_{X,\nu}$.

The exponent of β in (C.11) has the form $s_X + n_X$ if we define n_X as in (7.6.4). Next, we split the sum over the link index $j_{x,\mu}$ in (C.11) into partial sums over the values of $j_{x,\mu}$ with the same link occupation $\ell_{x,\mu}$,

$$(T_X)_{j_X} = \beta^{n_X} \sum_{n_x^s=0}^{n_{\max}} \sum_{m=0}^{n_{\max}-n_x^s} \sum_{\substack{j_{x,\mu} \\ \ell_{x,\mu}=m}} \beta^{n_X} (T_x^{n_x^s})_{j_x} \sigma_\mu, \quad (\text{C.13})$$

where β^{n_X} could be pulled out since n_X does not depend on any of the summation variables and where the upper limit on the sum over m is reduced since we apply the site criterion (7.4.2) in every blocking step. We now eliminate m in favor of $n_X^s = n_x^s + m$,

$$(T_X)_{j_X} = \beta^{n_X} \sum_{n_x^s=0}^{n_{\max}} \sum_{n_X^s=n_x^s}^{n_{\max}} \sum_{\substack{j_{x,\mu} \\ \ell_{x,\mu}=n_X^s-n_x^s}} \beta^{n_X} (T_x^{n_x^s})_{j_x} \sigma_\mu = \beta^{n_X} \sum_{n_X^s=0}^{n_{\max}} \sum_{n_x^s=0}^{n_X^s} \sum_{\substack{j_{x,\mu} \\ \ell_{x,\mu}=n_X^s-n_x^s}} \beta^{n_X} (T_x^{n_x^s})_{j_x} \sigma_\mu, \quad (\text{C.14})$$

where we switched the order of the sums over n_x^s and n_X^s in the last step. This shows that $(T_X)_{j_X}$ can be written in the form of (7.3.1) with $(T_X^{n_X^s})_{j_X}$ given in (7.6.2). Therefore the tensor decomposition is preserved after the tracing.

D Validating the coefficients of Sec. 8.4

Since $n_R \geq \lfloor n_{\max}/2 \rfloor$, all partitions that appear in Z also appear in X . We now show that the coefficients of terms in Z_n match those of $a_n X_n + b_n Y_n$ if a_n , b_n , and c_n are given by (8.4.1).

Consider a specific partition $(n) = (n_1, \dots, n_1, n_2, \dots, n_2, \dots, n_m, \dots, n_m)$ of n as defined in Sec. 8.1, where k_i is the multiplicity of n_i and $n = \sum_{i=1}^m k_i n_i$. The contribution of this partition to Z_n is $Z_{(n)}$. To compute the contribution of this partition to $a_n X_n + b_n Y_n$ we distinguish two cases.

The first case is $n_1 \leq n_R$. Since the n_i are ordered, this implies $n_i \leq n_R$ for all i . Then all terms contained in X_n that are related to (n) by permutation of the n_i are given by

$$\sum_{i=1}^m c_{n_i} Z_{(n_i | n \setminus n_i)} \stackrel{(8.1.1)}{=} \frac{Z_{(n)}}{V} \sum_{i=1}^m c_{n_i} k_i \stackrel{(8.4.1)}{=} \frac{Z_{(n)}}{V} \sum_{i=1}^m \frac{n_i k_i}{n_{\max}} = \frac{n}{V n_{\max}} Z_{(n)}. \quad (\text{D.1})$$

Since $n_i \leq n_R$ for all i , the term in Y_n related to (n) is given by $Z_{(n)}$. Therefore, in $a_n X_n + b_n Y_n$, all terms related to (n) are given by

$$a_n \frac{n}{V n_{\max}} Z_{(n)} + b_n Z_{(n)} \stackrel{(8.4.1)}{=} Z_{(n)}. \quad (\text{D.2})$$

The second case is $n_1 > n_R$. Due to $n_R \geq \lfloor n_{\max}/2 \rfloor$, this implies $k_1 = 1$ and $n_i \leq n_R$ for $i \geq 2$. Then the only term in X_n related to (n) is given by

$$c_{n_1} Z_{(n_1 | n \setminus n_1)} \stackrel{(8.4.1)}{=} Z_{(n_1 | n \setminus n_1)} \stackrel{(8.1.1) \text{ for } k_1=1}{=} \frac{1}{V} Z_{(n)}. \quad (\text{D.3})$$

Any tuple obtained from (n) by a non-trivial permutation cannot appear in X_n , since $n_1 > n_R$. For the same reason, Y_n does not contain terms related to (n) . Therefore, in $a_n X_n + b_n Y_n$ with $n \geq n_1 > n_R$, all terms related to (n) are given by

$$\frac{a_n}{V} Z_{(n)} \stackrel{(8.4.1) \text{ for } n > n_R}{=} Z_{(n)}. \quad (\text{D.4})$$

In both cases, the sum of all terms in $a_n X_n + b_n Y_n$ that are related to (n) is given by $Z_{(n)}$. Since the arguments above apply for any partition (n) of n , this concludes the proof.

E Expansion coefficients for a 2×2 lattice

The expansion coefficients Z_n in Eq. (9.1.1) can be computed analytically for a 2×2 lattice using the procedure described in Sec. 9.1. For $N_f = 1$, they are given, up to $n = 4$, by

$$\begin{aligned} Z_0(m, \mu) = & 4096m^{12} + 24576m^{10} + 53248m^8 + 50944m^6 + 21248m^4 + \frac{9760m^2}{3} + \frac{998}{9} \\ & + (256m^6 + 384m^4 + 160m^2 + 16) \cosh(6\mu) + 2 \cosh(12\mu), \end{aligned} \quad (\text{E.1a})$$

$$Z_1(m, \mu) = \frac{1024m^8}{3} + \frac{10240m^6}{9} + \frac{29632m^4}{27} + \frac{26080m^2}{81} + \frac{5380}{243} + \left(\frac{128m^4}{3} + \frac{416m^2}{9} + \frac{80}{9} \right) \cosh(6\mu), \quad (\text{E.1b})$$

$$Z_2(m, \mu) = \frac{4096m^{12}}{9} + \frac{8192m^{10}}{3} + \frac{482048m^8}{81} + \frac{1405312m^6}{243} + \frac{1820704m^4}{729} + \frac{100624m^2}{243} + \frac{13034}{729} + \left(\frac{2432m^6}{81} + \frac{11936m^4}{243} + \frac{17648m^2}{729} + \frac{2344}{729} \right) \cosh(6\mu) + \frac{2}{9} \cosh(12\mu), \quad (\text{E.1c})$$

$$Z_3(m, \mu) = \frac{2048m^{12}}{81} + \frac{4096m^{10}}{27} + \frac{804352m^8}{2187} + \frac{325312m^6}{729} + \frac{21232m^4}{81} + \frac{134452m^2}{2187} + \frac{8887}{2187} + \left(\frac{448m^6}{243} + \frac{16768m^4}{2187} + \frac{14548m^2}{2187} + \frac{2714}{2187} \right) \cosh(6\mu) + \frac{1}{81} \cosh(12\mu), \quad (\text{E.1d})$$

$$Z_4(m, \mu) = \frac{166400m^{12}}{6561} + \frac{332800m^{10}}{2187} + \frac{2201008m^8}{6561} + \frac{2201512m^6}{6561} + \frac{1788793m^4}{11664} + \frac{2968507m^2}{104976} + \frac{646769}{419904} + \left(\frac{3880m^6}{2187} + \frac{22054m^4}{6561} + \frac{13085m^2}{6561} + \frac{938}{2916} \right) \cosh(6\mu) + \frac{325}{26244} \cosh(12\mu). \quad (\text{E.1e})$$

Next, we state the expansion coefficients Z_n in Eq. (10.3.1) for $N_f = 2$ up to $n = 1$. Although we derived the analytical results for arbitrary m_1 , m_2 , μ_B , and μ_I (with $\mu_B = 3(\mu_1 + \mu_2)/2$ and $\mu_I = (\mu_1 - \mu_2)/2$), we only give the formulas for $m = m_1 = m_2$ and $\mu_I = 0$ for conciseness. They read

$$\begin{aligned} Z_0(m, \mu_B) = & \left(16777216m^{24} + 201326592m^{22} + 1048576000m^{20} + 3107979264m^{18} + 5785255936m^{16} \right. \\ & + \frac{21158297600m^{14}}{3} + \frac{51398721536m^{12}}{9} + 3059761152m^{10} + \frac{9588832768m^8}{9} + \frac{6309162496m^6}{27} \\ & + \frac{2498746240m^4}{81} + \frac{181396288m^2}{81} + \frac{2306680}{27} \Big) + \left(4194304m^{18} + 31457280m^{16} + 95551488m^{14} \right. \\ & + \frac{456589312m^{12}}{3} + \frac{1239973888m^{10}}{9} + \frac{1936113664m^8}{27} + \frac{186022400m^6}{9} + \frac{242845952m^4}{81} \\ & + \frac{14802272m^2}{81} + \frac{464}{27} \Big) \cosh(2\mu_B) + \left(147456m^{12} + 638976m^{10} + 1116160m^8 + \frac{8640512m^6}{9} \right. \\ & + \frac{11054528m^4}{27} + \frac{229376m^2}{3} + \frac{42968}{9} \Big) \cosh(4\mu_B) + \left(1024m^6 + 1536m^4 + 672m^2 \right. \\ & \left. + 80 \right) \cosh(6\mu_B) + 2 \cosh(8\mu_B), \end{aligned} \quad (\text{E.2a})$$

$$\begin{aligned} Z_1(m, \mu_B) = & \left(\frac{8388608m^{20}}{3} + \frac{243269632m^{18}}{9} + \frac{2950430720m^{16}}{27} + \frac{19506135040m^{14}}{81} + \frac{76959875072m^{12}}{243} \right. \\ & + \frac{20798980096m^{10}}{81} + \frac{3472271360m^8}{27} + \frac{9499891456m^6}{243} + \frac{186985124m^4}{27} + \frac{53274248m^2}{81} \end{aligned}$$

$$\begin{aligned}
& + \frac{13825109}{486} \Big) + \left(\frac{2097152m^{16}}{3} + \frac{49283072m^{14}}{9} + \frac{447938560m^{12}}{27} + \frac{2056437760m^{10}}{81} \right. \\
& + \frac{1742405632m^8}{81} + \frac{836475904m^6}{81} + \frac{222070144m^4}{81} + \frac{10028344m^2}{27} + \frac{1587950}{81} \Big) \cosh(2\mu_B) \\
& + \left(\frac{81920m^{10}}{3} + \frac{2385920m^8}{27} + \frac{9005056m^6}{81} + \frac{5262208m^4}{81} + \frac{448096m^2}{27} + \frac{37864}{27} \right) \cosh(4\mu_B) \\
& + \left(\frac{512m^4}{3} + 192m^2 + \frac{128}{3} \right) \cosh(6\mu_B). \tag{E.2b}
\end{aligned}$$

F Proof of $\lambda_R(\beta) = \mathcal{O}(\beta^{\bar{q}_1})$

In this appendix, we prove the relation

$$\lambda_R(\beta) = \mathcal{O}\left((\beta/N_c)^{\bar{q}_1}\right) \tag{9.2.13}$$

with

$$\lambda_R(\beta) \equiv \frac{1}{d_R} \sum_{Q \in \mathbb{Z}} \det \left[I_{\bar{q}_j + i - j + Q}(\beta/N_c) \right]_{i,j=1 \dots N_c}. \tag{9.2.9}$$

To do so, we expand the exponential in the modified Bessel function, given in 9.2.12, and obtain

$$I_n(z) = \frac{1}{\pi} \int_0^\pi \sum_{k=0}^\infty \frac{z^k}{k!} \cos^k(\theta) \cos(n\theta) d\theta = \frac{1}{\pi} \sum_{k=0}^\infty \int_0^\pi \frac{z^k}{k!} \cos^k(\theta) \cos(n\theta) d\theta = \sum_{k=0}^\infty I_n^k z^k, \tag{F.1}$$

where we interchanged the order of the integral and the series in the second step.¹ In the last step, we defined the coefficient

$$I_n^k \equiv \frac{1}{\pi k!} \int_0^\pi \cos^k(\theta) \cos(n\theta) d\theta, \tag{F.2}$$

which we evaluate explicitly in the following. To this end, we write

$$\cos^k(\theta) = \left(\frac{e^{i\theta} + e^{-i\theta}}{2} \right)^k = \frac{1}{2^k} \sum_{m=0}^k \binom{k}{m} e^{im\theta} e^{-i(k-m)\theta} = \frac{1}{2^k} \sum_{m=0}^k \binom{k}{m} \cos((2m-k)\theta), \tag{F.3}$$

where we expressed the cosine in terms of complex exponentials and applied the binomial theorem. In the last step, we reintroduced the cosine, which is valid since the expression on the far left is real. Thus we can write

$$\cos^k(\theta) = \sum_{m=0}^k a_m^k \cos((k-2m)\theta) \quad \text{with} \quad a_m^k = \frac{1}{2^k} \binom{k}{m}, \tag{F.4}$$

where we used the symmetry of the cosine. We now use the orthogonality of the cosine functions

$$\int_0^\pi \cos(a\theta) \cos(b\theta) d\theta = \frac{\pi}{2} \delta_{a,b} (1 + \delta_{a,0}) \quad \text{for all } a, b \in \mathbb{Z}, \tag{F.5}$$

¹The validity of this step can be proven using the inequality $|\frac{z^k}{k!} \cos^k(\theta) \cos(n\theta)| \leq \frac{|z|^k}{k!}$.

which results in

$$I_n^k = \frac{1}{\pi k!} \sum_{m=0}^k a_m^k \frac{\pi}{2} \delta_{k-2m,n} (1 + \delta_{n,0}) = \begin{cases} \frac{1}{2k!} a_{(k-n)/2}^k (1 + \delta_{n,0}) & \text{for } |n| \leq k \text{ and } |n| + k \text{ even,} \\ 0 & \text{else.} \end{cases} \quad (\text{F.6})$$

Therefore, the modified Bessel functions of the first kind have the structure

$$I_n(\beta/N_c) = \mathcal{O}\left((\beta/N_c)^{|n|}\right). \quad (\text{F.7})$$

Consequently, for the determinant in Eq. (9.2.9), it follows that

$$\begin{aligned} \det \left[I_{\bar{q}_j + i - j + Q}(\beta/N_c) \right]_{i,j=1\dots N_c} &= \det \left[\mathcal{O}\left((\beta/N_c)^{|\bar{q}_i + j - i + Q|}\right) \right]_{i,j=1\dots N_c} \\ &= \varepsilon_{j_1, \dots, j_{N_c}} \mathcal{O}\left((\beta/N_c)^{|\bar{q}_1 + j_1 - 1 + Q|}\right) \dots \mathcal{O}\left((\beta/N_c)^{|\bar{q}_{N_c} + j_{N_c} - N_c + Q|}\right) \\ &= \varepsilon_{j_1, \dots, j_{N_c}} \mathcal{O}\left((\beta/N_c)^{\sum_{i=1}^{N_c} |\bar{q}_i + j_i - i + Q|}\right), \end{aligned} \quad (\text{F.8})$$

where we introduced the Levi-Civita tensor $\varepsilon_{j_1, \dots, j_{N_c}}$ with summation indices $j_1, \dots, j_{N_c} \in \{1, \dots, N_c\}$.

The remaining task is to prove the relation

$$\sum_{i=1}^{N_c} |\bar{q}_i + j_i - i + Q| \geq f_1 \quad (\text{F.9})$$

for all values of $\bar{q}_i$, Q and permutations j_i . To do so, we define $a_i \equiv \bar{q}_i + j_i - i$ and obtain

$$\sum_{i=1}^{N_c} |\bar{q}_i + j_i - i + Q| = \sum_{i=1}^{N_c} |a_i + Q| \geq \max(a_i) - \min(a_i), \quad (\text{F.10})$$

where the inequality holds for any values of Q and tuples $(a_1, \dots, a_{N_c})$. We now observe that from

$$a_1 = \bar{q}_1 + j_1 - 1 \geq \bar{q}_1 \quad \text{it follows that} \quad \max(a_i) \geq \bar{q}_1 (\geq 0), \quad (\text{F.11})$$

and from

$$a_{N_c} = j_{N_c} - N_c \leq 0 \quad \text{it follows that} \quad \min(a_i) \leq 0. \quad (\text{F.12})$$

This allows us to write

$$\max(a_i) - \min(a_i) \geq \bar{q}_1, \quad (\text{F.13})$$

which concludes the proof.

G Taylor expansion of free energy

Calculation for general order

In this appendix we expand the expression

$$\mathcal{A} \equiv \ln \left(\sum_{i=0}^{\infty} \beta^i \frac{Z_i}{Z_0} \right), \quad (\text{G.1})$$

which is relevant for the expansion of $\frac{\ln Z}{V} = \frac{\ln Z_0}{V} + \frac{\mathcal{A}}{V}$ (cf. Eqs. (9.3.2) and (10.3.4)). To do so, we write

$$\mathcal{A} = \ln \left(1 + \sum_{i=1}^{n_{\max}} \beta^i \bar{Z}_i \right) = \sum_{n=0}^{\infty} \frac{(-1)^n}{n+1} \left(\sum_{i=1}^{n_{\max}} \beta^i \bar{Z}_i \right)^{n+1}, \quad (\text{G.2})$$

where in the first step we defined $\bar{Z}_i \equiv Z_i/Z_0$, followed by the expansion of the logarithm with convergence radius $|\sum_{i=1}^{n_{\max}} \beta^i \bar{Z}_i| \leq 1$. In the current form, terms belonging to the same order are not sorted. Therefore, we additionally use the multinomial theorem for the expression $(\sum_{i=1}^{n_{\max}} \beta^i \bar{Z}_i)^{n+1}$. The theorem reads

$$\left(\sum_{i=1}^m x_i \right)^n = \sum_{\substack{k_1, k_2, \dots, k_m \geq 0 \\ k_1 + k_2 + \dots + k_m = n}} \binom{n}{k_1, k_2, \dots, k_m} \prod_{t=1}^m x_t^{k_t}, \quad (\text{G.3})$$

with multinomial coefficient given by

$$\binom{n}{k_1, k_2, \dots, k_m} \equiv \frac{n!}{k_1! k_2! \dots k_m!}. \quad (\text{G.4})$$

Using this, we get the expression

$$\mathcal{A} = \sum_{n=0}^{\infty} \frac{(-1)^n}{n+1} \sum_{\substack{k_1, k_2, \dots, k_{n_{\max}} \geq 0 \\ k_1 + k_2 + \dots + k_{n_{\max}} = n+1}} \binom{n+1}{k_1, k_2, \dots, k_{n_{\max}}} \prod_{t=1}^{n_{\max}} (\beta^t \bar{Z}_t)^{k_t}, \quad (\text{G.5})$$

where we can collect all appearances of β by using

$$\prod_{t=1}^{n_{\max}} (\beta^t \bar{Z}_t)^{k_t} = \beta^{\sum_{t=1}^{n_{\max}} t k_t} \prod_{t=1}^{n_{\max}} \bar{Z}_t^{k_t}. \quad (\text{G.6})$$

After this step, we are left with sorting the terms that belong to the same order in β . The order of each contributing term is given by the exponent $\sum_{t=1}^{n_{\max}} t k_t$ and the sorting thereafter can be done by introducing an auxiliary sum over a summation index i together with the additional constraint $i = \sum_{t=1}^{n_{\max}} t k_t$, yielding

$$\mathcal{A} = \sum_{n=0}^{\infty} \frac{(-1)^n}{n+1} \sum_{\substack{k_1, k_2, \dots, k_{n_{\max}} \geq 0 \\ k_1 + k_2 + \dots + k_{n_{\max}} = n+1 \\ 1 \cdot k_1 + 2 \cdot k_2 + \dots + n_{\max} \cdot k_{n_{\max}} = i}} \binom{n+1}{k_1, k_2, \dots, k_{n_{\max}}} \beta^i \prod_{t=1}^{n_{\max}} \bar{Z}_t^{k_t} + \mathcal{O}(\beta^{n_{\max}+1}), \quad (\text{G.7})$$

where the constraint is used to reduce the range of the sum over tuples $(k_1, k_2, \dots, k_{n_{\max}})$. Additionally, we used that we are only interested in terms up to order $\beta^{n_{\max}+1}$, which is why the sum over i can be limited by $n_{\max}$. For the same reason, the sum over the index n can be limited as well: We observe that

$$k_1 + k_2 + \dots + k_{n_{\max}} \leq 1 \cdot k_1 + 2 \cdot k_2 + \dots + n_{\max} \cdot k_{n_{\max}} \quad \forall k_1, k_2, \dots, k_{n_{\max}} \geq 0 \quad (\text{G.8})$$

Therefore, both conditions can only be fulfilled simultaneously for a valid tuple $(k_1, k_2, \dots, k_{n_{\max}})$, whenever $n+1 \leq i$ or equivalently $n \leq i-1$, which is the chosen upper bound for the summation

index n , which we use in

$$\mathcal{A} = \sum_{n=0}^{i-1} \frac{(-1)^n}{n+1} \sum_{i=1}^{n_{\max}} \sum_{\substack{k_1, k_2, \dots, k_i \geq 0 \\ k_1 + k_2 + \dots + k_i = n+1 \\ 1 \cdot k_1 + 2 \cdot k_2 + \dots + i \cdot k_i = i}} \binom{n+1}{k_1, k_2, \dots, k_i} \beta^i \prod_{t=1}^i \bar{Z}_t^{k_t} + \mathcal{O}(\beta^{n_{\max}+1}). \quad (\text{G.9})$$

Here, we also reduced the complexity of the summation over the tuples $(k_1, k_2, \dots, k_{n_{\max}})$ as well: All indices k_j with $j > i$ have to be zero, otherwise the relation $j k_j > i$ makes it impossible for any tuple to fulfill the last constraint of the sum. Note that all indices k_j that are zero result in a trivial factor 1, as these indices appear either in a factor $0! = 1$ inside the multinomial coefficient, yielding a multinomial coefficient with less indices, or in a zero as exponent in the terms q_t which results in a factor one as well.

Now we can read off the coefficients of the strong-coupling expansion of the free energy for $i \geq 1$, by writing

$$\mathcal{A} = V \sum_{i=1}^{n_{\max}} \beta^i f_i + \mathcal{O}(\beta^{n_{\max}+1}), \quad (\text{G.10})$$

which are

$$f_i = \frac{1}{V} \sum_{n=0}^{i-1} \frac{(-1)^n}{n+1} \sum_{\substack{k_1, k_2, \dots, k_i \geq 0 \\ k_1 + k_2 + \dots + k_i = n+1 \\ 1 \cdot k_1 + 2 \cdot k_2 + \dots + i \cdot k_i = i}} \binom{n+1}{k_1, k_2, \dots, k_i} \prod_{t=1}^i \bar{Z}_t^{k_t} \quad \text{with } i \geq 1. \quad (\text{G.11})$$

Explicit calculation up to order three

Using the formular for the coefficients of the strong-coupling expansion for general order i , given in (G.11), we now calculate the explicit results up to third order.

For $i = 1$ we find

$$f_1 = \frac{1}{V} \sum_{\substack{k_1 \geq 0 \\ k_1=1 \\ k_1=1}} \binom{1}{k_1} \bar{Z}_1^{k_1} = \frac{1}{V} \bar{Z}_1, \quad (\text{G.12})$$

where we wrote down the resulting equation for both constraints separately, but in fact, they result both in the equation $k_1 = 1$.

Next we calculate the second-order coefficient, i.e., $i = 2$. The resulting expression reads

$$f_2 = \frac{1}{V} \sum_{n=0}^1 \frac{(-1)^n}{n+1} \sum_{\substack{k_1, k_2 \geq 0 \\ k_1 + k_2 = n+1 \\ 1 \cdot k_1 + 2 \cdot k_2 = 2}} \binom{n+1}{k_1, k_2} \bar{Z}_1^{k_1} \bar{Z}_2^{k_2}. \quad (\text{G.13})$$

For given n , there is only one tuple (k_1, k_2) fulfilling both constraint, which can be determined by solving the corresponding system of linear equations.

$$\left(\begin{array}{cc|c} 1 & 1 & n+1 \\ 1 & 2 & 2 \end{array} \right) \rightarrow \left(\begin{array}{cc|c} 1 & 1 & n+1 \\ 0 & 1 & 1-n \end{array} \right). \quad (\text{G.14})$$

In the second line, we subtracted the first line once and we can read off the solutions $k_1 = 2n$ and $k_2 = 1 - n$. Therefore, the second-order coefficient reads

$$\begin{aligned} f_2 &= \frac{1}{V} \sum_{n=0}^1 \frac{(-1)^n}{n+1} \binom{n+1}{2n, 1-n} \bar{Z}_1^{2n} \bar{Z}_2^{1-n} \\ &= \frac{1}{V} \left[\binom{1}{0, 1} \bar{Z}_2 - \frac{1}{2} \binom{2}{2, 0} \bar{Z}_1^2 \right] = \frac{1}{V} \left(\bar{Z}_2 - \frac{1}{2} \bar{Z}_1^2 \right). \end{aligned} \quad (\text{G.15})$$

Finally, we calculate the coefficient for $i = 3$ as well. We find

$$f_3 = \frac{1}{V} \sum_{n=0}^2 \frac{(-1)^n}{n+1} \sum_{k_1=0}^{n+1} \sum_{\substack{k_2, k_3 \geq 0 \\ k_1+k_2+k_3=n+1 \\ 1 \cdot k_1 + 2 \cdot k_2 + 3 \cdot k_3 = 3}} \binom{n+1}{k_1, k_2, k_3} \bar{Z}_1^{k_1} \bar{Z}_2^{k_2} \bar{Z}_3^{k_3}, \quad (\text{G.16})$$

where we wrote the sum over k_1 separately. Due to the constraints for the indices k_j , the sum over the indices k_2 and k_3 can be eliminated by solving a system of linear equations. We get

$$\left(\begin{array}{cc|c} 1 & 1 & n+1-k_1 \\ 2 & 3 & 3-k_1 \end{array} \right) \rightarrow \left(\begin{array}{cc|c} 1 & 1 & n+1-k_1 \\ 0 & 1 & k_1+1-2n \end{array} \right), \quad (\text{G.17})$$

where we subtracted the first line twice from the second line. The resulting indices are

$$0 \leq k_3 = k_1 + 1 - 2n \quad \text{and} \quad 0 \leq k_2 = 3n - 2k_1.$$

This is also used to limit the range of the summation over the index k_1 as

$$2n - 1 \leq k_1 \leq \frac{3n}{2} \quad \text{and therefore} \quad \max(0, 2n - 1) \leq k_1 \leq \lfloor \frac{3n}{2} \rfloor,$$

where we used that k_1 is a non-negative integer. Thus we obtain

$$f_3 = \frac{1}{V} \sum_{n=0}^2 \frac{(-1)^n}{n+1} \sum_{k_1=\max(0, 2n-1)}^{\lfloor \frac{3n}{2} \rfloor} \binom{n+1}{k_1, (3n-2k_1), (k_1+1-2n)} \bar{Z}_1^{k_1} \bar{Z}_2^{3n-2k_1} \bar{Z}_3^{k_1+1-2n}. \quad (\text{G.18})$$

We observe that the sum over k_1 collapses to a single term with

$$n = 0 \Rightarrow k_1 = 0, \quad n = 1 \Rightarrow k_1 = 1, \quad n = 2 \Rightarrow k_1 = 3, \quad (\text{G.19})$$

and we get

$$f_3 = \frac{1}{V} \left[\binom{1}{0, 0, 1} \bar{Z}_3 - \frac{1}{2} \binom{2}{1, 1, 0} \bar{Z}_1 \bar{Z}_2 + \frac{1}{3} \binom{3}{3, 0, 0} \bar{Z}_1^3 \right] = \frac{1}{V} \left(\bar{Z}_3 - \bar{Z}_1 \bar{Z}_2 + \frac{1}{3} \bar{Z}_1^3 \right). \quad (\text{G.20})$$

H The function $p_{V_s}(m, \beta)$

The purpose of this appendix is to prove the relation (9.3.15). We can read off the expression for p_{V_s} from Eq. (4.1.3) and Eq. (2.6.36) for $N_f = 1$ by identifying the coefficient of $\exp(N_c V \mu)$ and obtain

$$\frac{1}{2} p_{V_s}(m, \beta) = \sum_{\mathbf{n}} \int \left[\prod_{x, \mu} dU_{x, \mu} \right] \left[\prod_x \mathcal{B}_x \right] \left[\prod_{x, \mu \neq \nu} \frac{(S_{x, \mu \nu}^G)^{n_{x, \mu \nu}}}{n_{x, \mu \nu}!} \right] \quad (\text{H.1})$$

with¹

$$\begin{aligned} \mathcal{B}_x &\equiv \frac{1}{N_c!} \int \left[\prod_{i=1}^{N_c} d\psi_{x+\hat{1}}^i \right] \left[\prod_{i=1}^{N_c} d\bar{\psi}_x^i \right] \left(\bar{\psi}_x U_{x,1} \psi_{x+\hat{1}} \right)^{N_c} \\ &= \frac{1}{N_c!} \int \left[\prod_{i=1}^{N_c} d\psi_{x+\hat{1}}^i \right] \left[\prod_{i=1}^{N_c} d\bar{\psi}_x^i \right] \left[\prod_{k=1}^{N_c} \bar{\psi}_{x,1}^{i_k} \right] \left[\prod_{k=1}^{N_c} U_{x,1}^{i_k j_{x,1}^k} \right] \left[\prod_{k=1}^{N_c} \psi_{x+\hat{1}}^{j_{x,1}^k} \right] \\ &= \frac{1}{N_c!} \varepsilon_{i_{x,1}^1 \dots i_{x,1}^{N_c}} \left[\prod_{k=1}^{N_c} U_{x,1}^{i_{x,1}^k j_{x,1}^k} \right] \varepsilon_{j_{x,1}^1 \dots j_{x,1}^{N_c}} = \frac{1}{N_c!} \det(U_{x,1}) \varepsilon_{j_{x,1}^1 \dots j_{x,1}^{N_c}} \varepsilon_{i_{x,1}^1 \dots i_{x,1}^{N_c}} = 1, \end{aligned} \quad (\text{H.2})$$

where sums over repeated color indices are implied. We integrated out the Grassmann variables using Eq. (A.23) and used

$$\varepsilon_{i_{x,1}^1 \dots i_{x,1}^{N_c}} \left[\prod_{k=1}^{N_c} U_{x,1}^{i_{x,1}^k j_{x,1}^k} \right] = \det(U_{x,1}) \varepsilon_{j_{x,1}^1 \dots j_{x,1}^{N_c}}, \quad (\text{H.3})$$

which follows directly from the Leibniz formula. In the last step of (H.2) we used $\det(U_{x,1}) = 1$ and contracted the Levi-Civita tensors, which results in a factor of $N_c!$.

Equation (H.2) shows that the V_s parallel baryon loops in the time direction result in a constant background, independent of the gauge field. Therefore (H.1) is simply the partition function of the pure-gauge theory with the same lattice size and number of colors as the original theory. This concludes the proof.

I Fit ansatz using the tanh models

I.1 Coefficients of the strong-coupling expansion for the observables

For both model functions (10.3.7) and (10.3.9) we first expand the argument of the tanh in β ,

$$a(\beta)(\mu - \mu^c(\beta)) = \sum_{n=0}^{n_{\max}} \bar{\mu}_n \beta^n + \mathcal{O}(\beta^{n_{\max}+1}) \quad (\text{I.1})$$

¹The Grassmann differentials have been reordered compared to [27], but this does not result in a sign change since V is even.

with

$$\bar{\mu}_0 = a_0(\mu - \mu_0^c), \quad (\text{I.2a})$$

$$\bar{\mu}_1 = a_1(\mu - \mu_0^c) - a_0\mu_1^c, \quad (\text{I.2b})$$

$$\bar{\mu}_2 = a_2(\mu - \mu_0^c) - a_1\mu_1^c - a_0\mu_2^c, \quad (\text{I.2c})$$

$$\bar{\mu}_3 = a_3(\mu - \mu_0^c) - a_2\mu_1^c - a_1\mu_2^c - a_0\mu_3^c. \quad (\text{I.2d})$$

Here and below we omit the labels Σ and ρ for the parameters a_n and μ_n^c . Using $t \equiv \tanh(\bar{\mu}_0)$ and $s \equiv \text{sech}(\bar{\mu}_0)$ we obtain

$$\Sigma_0^{\text{model}}(\mu) = b_0(1 - t), \quad (\text{I.3a})$$

$$\Sigma_1^{\text{model}}(\mu) = b_1(1 - t) - b_0\bar{\mu}_1s^2, \quad (\text{I.3b})$$

$$\Sigma_2^{\text{model}}(\mu) = b_2(1 - t) - b_1\bar{\mu}_1s^2 - b_0(\bar{\mu}_2 - \bar{\mu}_1^2t)s^2, \quad (\text{I.3c})$$

$$\Sigma_3^{\text{model}}(\mu) = b_3(1 - t) - b_2\bar{\mu}_1s^2 - b_1(\bar{\mu}_2 - \bar{\mu}_1^2t)s^2 - b_0\left(\bar{\mu}_3 - 2\bar{\mu}_2\bar{\mu}_1t - \bar{\mu}_1^3\left(s^2 - \frac{2}{3}\right)\right)s^2. \quad (\text{I.3d})$$

The results for ρ_n can be obtained by replacing $b_0 \rightarrow \frac{3}{2}$, $b_n \rightarrow 0$ for $n \geq 1$, and $\bar{\mu}_n \rightarrow -\bar{\mu}_n$ for all n (resulting also in $t \rightarrow -t$) in the expressions for Σ_n . This yields

$$\rho_0^{\text{model}}(\mu) = \frac{3}{2}(1 + t), \quad (\text{I.4a})$$

$$\rho_1^{\text{model}}(\mu) = \frac{3}{2}\bar{\mu}_1s^2, \quad (\text{I.4b})$$

$$\rho_2^{\text{model}}(\mu) = \frac{3}{2}(\bar{\mu}_2 - \bar{\mu}_1^2t)s^2, \quad (\text{I.4c})$$

$$\rho_3^{\text{model}}(\mu) = \frac{3}{2}\left(\bar{\mu}_3 - 2\bar{\mu}_2\bar{\mu}_1t - \bar{\mu}_1^3\left(s^2 - \frac{2}{3}\right)\right)s^2. \quad (\text{I.4d})$$

Using these results we can explain the low sensitivity of the model coefficients in (I.3) and (I.4) on a_n/a_0 for $n \geq 1$ and on μ_n^c for $n \geq 2$ when a_0 becomes large, which is the case for large volume. For simplicity we focus on ρ , but similar arguments can be made for Σ . Let us first consider ρ_1 (suppressing the superscript “model” for brevity). Because of $\text{sech}(\bar{\mu}_0)$, ρ_1 becomes exponentially suppressed when $|\bar{\mu}_0| \gg 1$. We therefore only need to consider $|\bar{\mu}_0| \leq \mathcal{O}(1)$, i.e., the vicinity of the phase transition, in (I.2b), which can be written as

$$\bar{\mu}_1 = \frac{a_1}{a_0}\bar{\mu}_0 - a_0\mu_1^c. \quad (\text{I.5})$$

We see that the first term becomes less and less relevant when a_0 becomes large at fixed a_1/a_0 , provided that $\mu_1^c \neq 0$. Therefore ρ_1 in (I.4b) becomes increasingly insensitive to a_1/a_0 in this scenario.

Similarly, for large a_0 we see that ρ_2 is dominated by the $(a_0\mu_1^c)^2$ term originating from $\bar{\mu}_1^2$ and that ρ_3 is dominated by the $(a_0\mu_1^c)^3$ term originating from $\bar{\mu}_1^3$. This pattern continues to larger values of n , and therefore all ρ_n for $n \geq 1$ are dominated by a_0^n and insensitive to $a_1/a_0, \dots, a_n/a_0$ when a_0 becomes large for large V .

The same argument also applies to μ_n^c for $n \geq 2$. For example, the subleading terms in ρ_2/s^2 are $a_1\mu_1^c$ and $a_0\mu_2^c$. Both terms are suppressed by $1/a_0$ compared to the leading term $(a_0\mu_1^c)^2$. Hence

determining μ_2^c or a_1 from ρ_2 should be equally hard. A more careful analysis of the zeros of ρ_2 (which we do not present here) does not change this expectation. To conclude, in a combined fit for all ρ_n , it appears that a_n/a_0 and μ_{n+1}^c should be equally difficult to determine for large V , provided that all a_n/a_0 and μ_n^c have finite large- V limits which are of similar magnitude.

I.2 Integrating the model function for the quark-number density

In this appendix we determine a model function for the free energy starting from the symmetrized model function for the quark-number density given in Footnote 34.

Integration of the model function given in Footnote 34 with respect to μ yields

$$f_{\text{sym}}^{\text{model}}(\mu, \beta) = \frac{N_c}{2a_\rho(\beta)} \ln \left(\cosh[a_\rho(\beta)(\mu - \mu_\rho^c(\beta))] \cosh[a_\rho(\beta)(\mu + \mu_\rho^c(\beta))] \right) + \tilde{c}(\beta), \quad (\text{I.6})$$

with integration constant $\tilde{c}(\beta)$. In the limit $\mu \rightarrow \pm\infty$, we obtain

$$f_{\text{sym}}^{\text{model}}(\mu, \beta) \rightarrow \frac{N_c}{2a_\rho(\beta)} \ln \left(\frac{1}{4} \exp[2a_\rho(\beta)|\mu|] \right) + \tilde{c}(\beta) = \frac{-N_c \ln 4}{2a_\rho(\beta)} + N_c |\mu| + \tilde{c}(\beta). \quad (\text{I.7})$$

Using

$$f(\mu, \beta) \rightarrow f^{\text{PG}}(\beta) + N_c |\mu|, \text{ for } \mu \rightarrow \pm\infty, \quad (\text{I.8})$$

(cf. Eq. (9.3.20)), we can determine the integration constant

$$\tilde{c}(\beta) = f^{\text{PG}}(\beta) + \frac{N_c \ln 4}{2a_\rho(\beta)}. \quad (\text{I.9})$$

For large $a_\rho(\beta)\mu_\rho^c(\beta)$ and $\mu > 0$, we can replace $\cosh[a_\rho(\beta)(\mu + \mu_\rho^c(\beta))] \rightarrow \frac{1}{2} \exp[a_\rho(\beta)(\mu + \mu_\rho^c(\beta))]$ and obtain the simplified fit ansatz

$$f^{\text{model}}(\mu, \beta) = \frac{N_c}{2a_\rho(\beta)} \ln \left(2 \cosh[a_\rho(\beta)(\mu - \mu_\rho^c(\beta))] \right) + \frac{N_c}{2} (\mu + \mu_\rho^c(\beta)) + f^{\text{PG}}(\beta). \quad (\text{I.10})$$

By expanding this simplified fit ansatz in β , we obtain the expansion

$$f^{\text{model}}(\mu, \beta) = \sum_{n=0}^{\infty} f_n^{\text{model}}(\mu) \beta^n \quad (\text{I.11})$$

with coefficients $f_n^{\text{model}}(\mu)$.

J Explicit fits for various m , L , and $n_{\max}$

In this appendix we show the coefficients of the chiral condensate and the free energy versus μ for various m , L , and $n_{\max}$. Additionally, we show the corresponding fits using the model functions Σ_n^{model} (given in Eq. (I.3)) and f_n^{model} (obtained from expanding Eq. (I.10) in β). Using these fits one can determine the dependence of the fit parameters μ_n^c , a_n , and b_n on m and L .

Figure J.1 and Fig. J.2 show the coefficients of the chiral condensate and the free energy, respectively, versus μ obtained from OS-GHOTRG for various masses, for $L = 16$ and $n_{\max} = 2$ together with the corresponding fits. The m -dependence of the fit parameters is shown in Fig. 10.15. For $m = 1$, we observe systematic deviations between fit and tensor data.

Figure J.3 and Fig. J.4 are similar to Fig. J.1 and Fig. J.2, respectively, except that $n_{\max} = 3$. The resulting m -dependence of the fit parameters is shown in Fig. 10.16.

Figure J.5 shows the coefficients of the free energy versus μ obtained from OS-GHOTRG for various lattice sizes, for $m = 0.5$ and $n_{\max} = 1$ together with the corresponding fits. The resulting values of the fit parameters are shown in Fig. 10.17.

Similarly, Fig. J.6 and Fig. J.7 show the coefficients of the chiral condensate and the free energy, respectively, versus μ obtained from OS-GHOTRG for various lattice sizes, for $m = 0.5$ and $n_{\max} = 2$ together with the corresponding fits. The resulting values of the fit parameters are shown in Fig. 10.18. We observe that the largest systematic discrepancy between fit and tensor data occurs for the small lattice, i.e., $L = 8$.

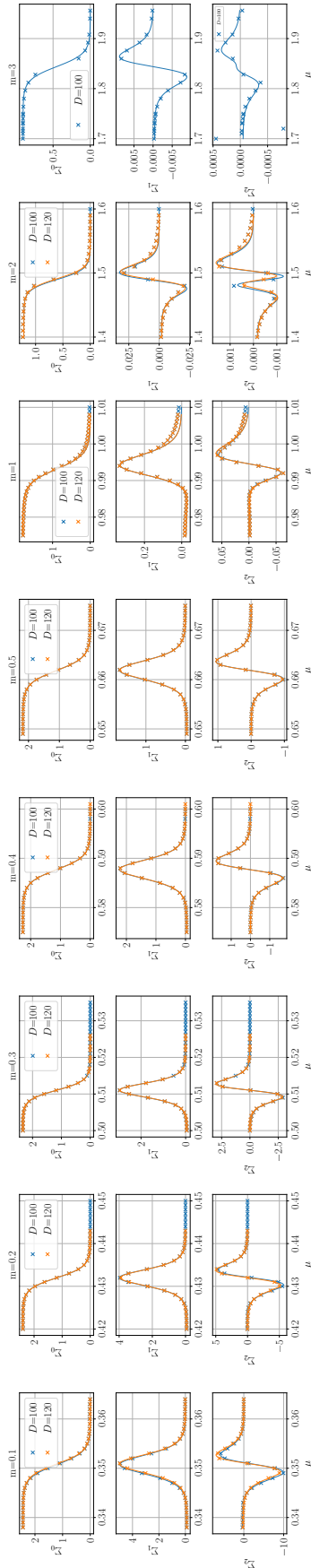


Figure J.1: Coefficients of the chiral condensate versus μ for different masses for $L = 16$ and $n_{\max} = 2$. The solid lines show the fits using the model functions given in Eq. (I.3).

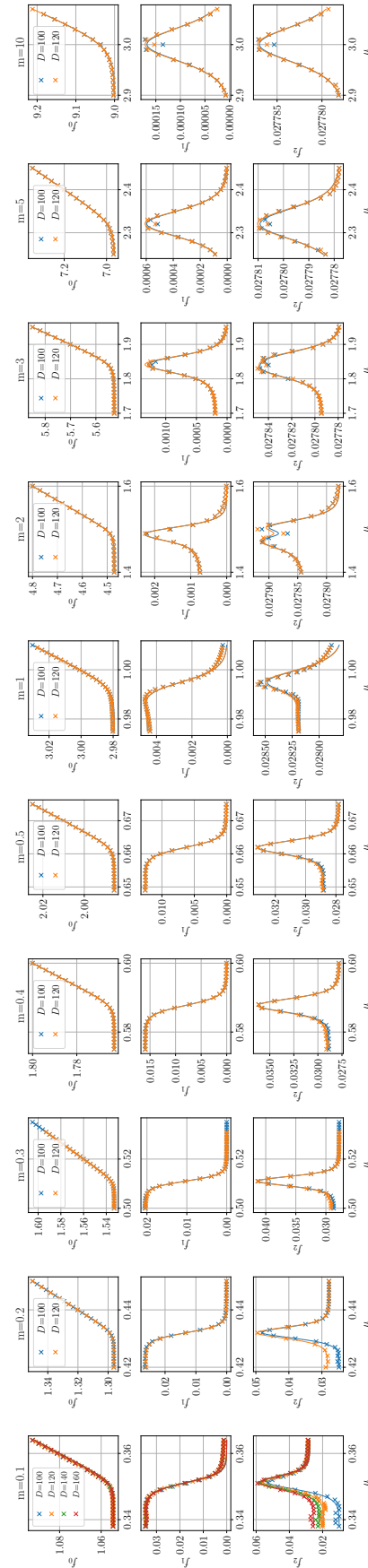


Figure J.2: Coefficients of the free energy versus μ for different masses for $L = 16$ and $n_{\max} = 2$. The solid lines show the fits using the model functions f_n^{model} .

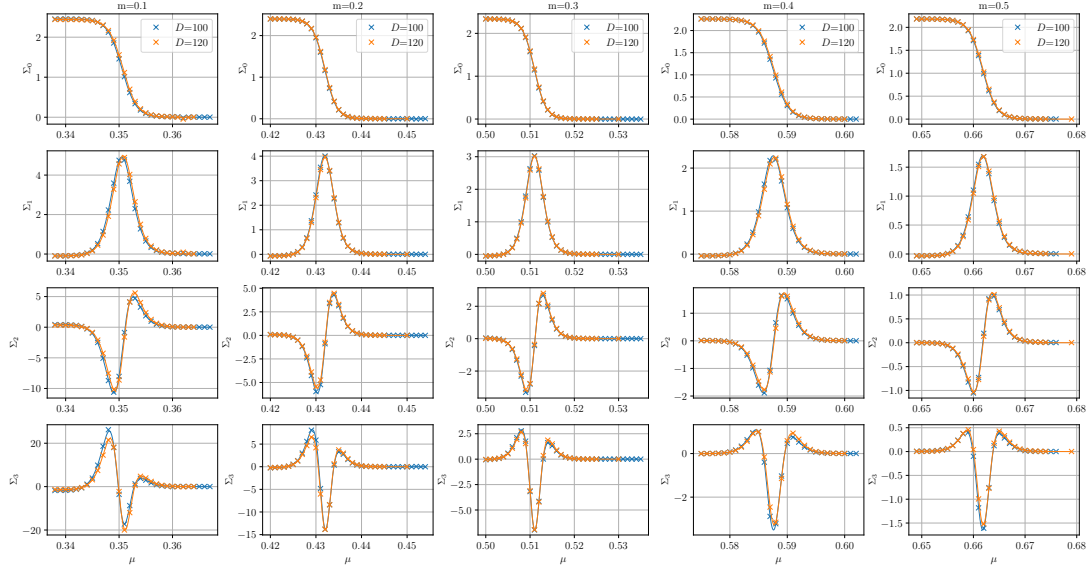


Figure J.3: Coefficients of the chiral condensate versus μ for different masses for $L = 16$ and $n_{\max} = 3$. The solid lines show the fits using the model functions given in Eq. (I.3).

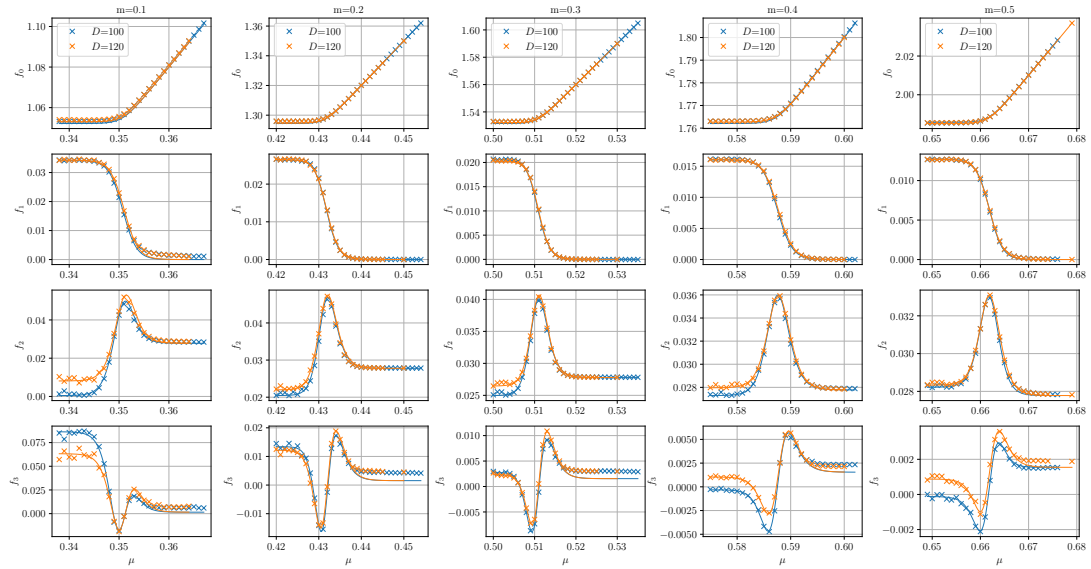


Figure J.4: Coefficients of the free energy versus μ for different masses for $L = 16$ and $n_{\max} = 3$. The solid lines show the fits using the model functions f_n^{model} .

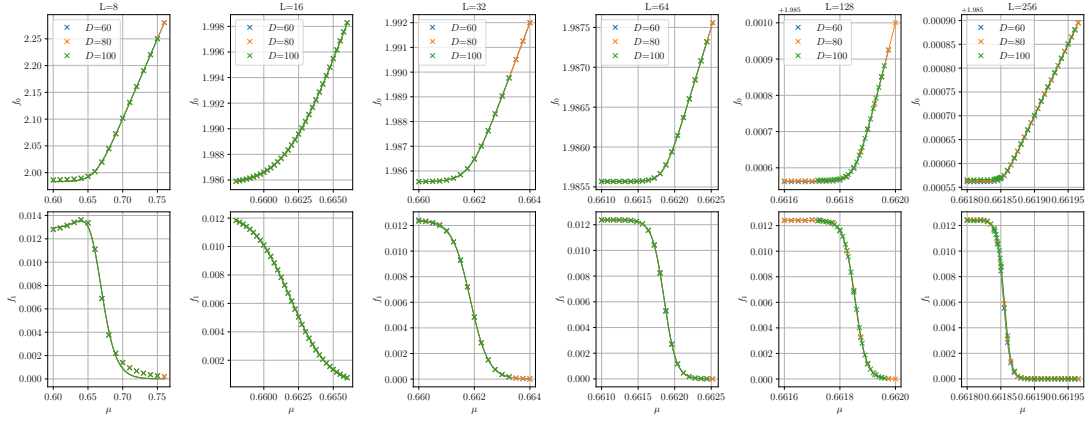


Figure J.5: Coefficients of the free energy versus μ for different lattices sizes for $m = 0.5$ and $n_{\max} = 1$. The solid lines show the fits using the model functions f_n^{model} .

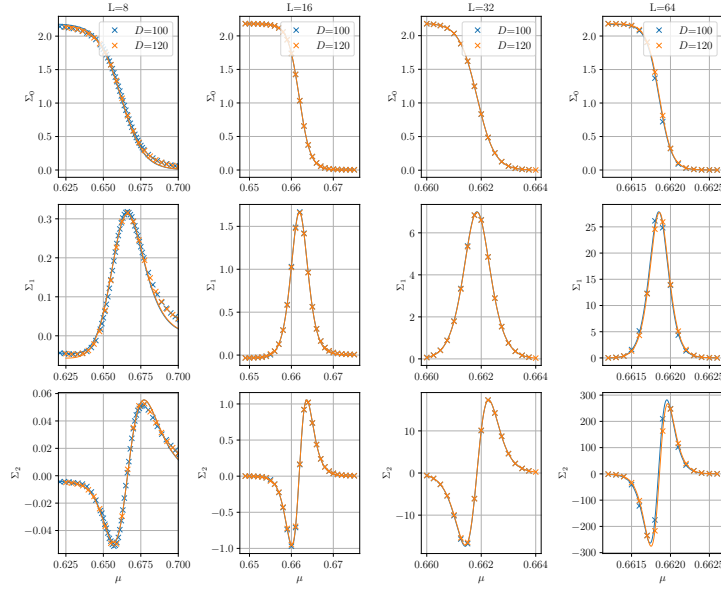


Figure J.6: Coefficients of the chiral condensate versus μ for different lattices sizes for $m = 0.5$ and $n_{\max} = 2$. The solid lines show the fits using the model functions given in Eq. (I.3).

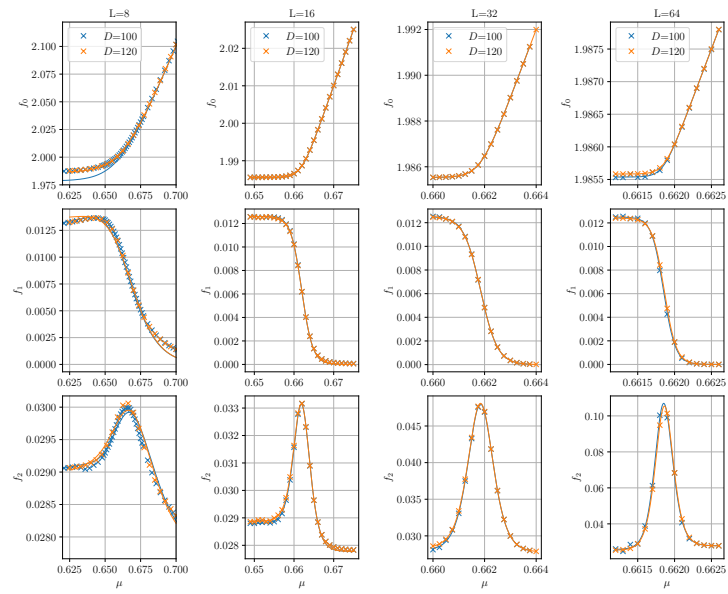


Figure J.7: Coefficients of the free energy versus μ for different lattices sizes for $m = 0.5$ and $n_{\max} = 2$. The solid lines show the fits using the model functions f_n^{model} .

References

- [1] P. Rosnet, *Quark-Gluon Plasma: from accelerator experiments to early Universe*. 2015. arXiv: 1510.04200 [hep-ph].
- [2] C. Gattringer and C. B. Lang, *Quantum Chromodynamics on the Lattice: An Introductory Presentation*. Lecture Notes in Physics. Springer Berlin Heidelberg, 2010. DOI: 10.1007/978-3-642-01850-3.
- [3] P. de Forcrand, *Simulating QCD at finite density*. In: *PoS LAT2009* (2009), p. 010. DOI: 10.22323/1.091.0010. arXiv: 1005.0539 [hep-lat].
- [4] G. Aarts, *Introductory lectures on lattice QCD at nonzero baryon number*. In: *J. Phys. Conf. Ser.* 706.2 (2016), p. 022004. DOI: 10.1088/1742-6596/706/2/022004. arXiv: 1512.05145 [hep-lat].
- [5] M. G. Alford, A. Schmitt, K. Rajagopal, and T. Schäfer, *Color superconductivity in dense quark matter*. In: *Reviews of Modern Physics* 80 (2008), pp. 1455–1515. ISSN: 1539-0756. DOI: 10.1103/revmodphys.80.1455.
- [6] P. Rossi and U. Wolff, *Lattice QCD With Fermions at Strong Coupling: A Dimer System*. In: *Nucl. Phys. B* 248 (1984), p. 105. DOI: 10.1016/0550-3213(84)90589-3.
- [7] F. Karsch and K.-H. Mütter, *Strong coupling QCD at finite baryon number density*. In: *Nucl. Phys. B* 313 (1989), p. 541. DOI: 10.1016/0550-3213(89)90396-9.
- [8] M. Fromm, *Lattice QCD at strong coupling: thermodynamics and nuclear physics*. Ph.D. thesis, ETH Zürich, 2010. DOI: 10.3929/ETHZ-A-006414247.
- [9] P. de Forcrand and M. Fromm, *Nuclear Physics from lattice QCD at strong coupling*. In: *Phys. Rev. Lett.* 104 (2010), p. 112005. DOI: 10.1103/PhysRevLett.104.112005. arXiv: 0907.1915 [hep-lat].
- [10] P. de Forcrand, J. Langelage, O. Philipsen, and W. Unger, *Lattice QCD Phase Diagram In and Away from the Strong Coupling Limit*. In: *Phys. Rev. Lett.* 113 (2014), p. 152002. DOI: 10.1103/PhysRevLett.113.152002. arXiv: 1406.4397 [hep-lat].
- [11] N. Bilic, K. Demeterfi, and B. Petersson, *Strong coupling analysis of the chiral phase transition at finite chemical potential and finite temperature*. In: *Nucl. Phys. B* 377 (1992), pp. 651–665. DOI: 10.1016/0550-3213(92)90305-U.
- [12] C. Marchis and C. Gattringer, *Dual representation of lattice QCD with worldlines and worldsheets of Abelian color fluxes*. In: *Phys. Rev. D* 97 (2018), p. 034508. DOI: 10.1103/PhysRevD.97.034508. arXiv: 1712.07546 [hep-lat].
- [13] G. Gagliardi and W. Unger, *New dual representation for staggered lattice QCD*. In: *Phys. Rev. D* 101 (2020), p. 034509. DOI: 10.1103/PhysRevD.101.034509. arXiv: 1911.08389 [hep-lat].
- [14] W. Unger, J. Kim, and P. Pattanaik, *The chiral critical point from the strong coupling expansion*. In: *PoS LATTICE2024* (2025), p. 166. DOI: 10.22323/1.466.0166. arXiv: 2502.06679 [hep-lat].
- [15] J. Kim, P. Pattanaik, and W. Unger, *Nuclear liquid-gas transition in the strong coupling regime of lattice QCD*. In: *Phys. Rev. D* 107 (2023), p. 094514. DOI: 10.1103/PhysRevD.107.094514. arXiv: 2303.01467 [hep-lat].

- [16] M. Asaduzzaman, S. Catterall, Y. Meurice, R. Sakai, and G. C. Toga, *Tensor network representation of non-abelian gauge theory coupled to reduced staggered fermions*. In: *JHEP* 05 (2024), p. 195. DOI: 10.1007/JHEP05(2024)195. arXiv: 2312.16167 [hep-lat].
- [17] K. H. Pai, S. Akiyama, and S. Todo, *Grassmann tensor renormalization group approach to (1+1)-dimensional two-color lattice QCD at finite density*. In: *JHEP* 03 (2025), p. 027. DOI: 10.1007/JHEP03(2025)027. arXiv: 2410.09485 [hep-lat].
- [18] N. Prokof'ev and B. Svistunov, *Worm Algorithms for Classical Statistical Models*. In: *Phys. Rev. Lett.* 87 (2001), p. 160601. DOI: 10.1103/PhysRevLett.87.160601.
- [19] Y. Meurice, R. Sakai, and J. Unmuth-Yockey, *Tensor lattice field theory for renormalization and quantum computing*. In: *Rev. Mod. Phys.* 94 (2022), p. 025005. DOI: 10.1103/RevModPhys.94.025005. arXiv: 2010.06539 [hep-lat].
- [20] Z. Y. Xie et al., *Coarse-graining renormalization by higher-order singular value decomposition*. In: *Phys. Rev. B* 86 (2012), p. 045139. DOI: 10.1103/physrevb.86.045139.
- [21] L. De Lathauwer, B. De Moor, and J. Vandewalle, *A Multilinear Singular Value Decomposition*. In: *SIAM Journal on Matrix Analysis and Applications* 21 (2000), p. 1253. DOI: 10.1137/S0895479896305696.
- [22] R. Sakai, S. Takeda, and Y. Yoshimura, *Higher order tensor renormalization group for relativistic fermion systems*. In: *PTEP* 2017 (2017), 063B07. DOI: 10.1093/ptep/ptx080. arXiv: 1705.07764 [hep-lat].
- [23] J. Bloch and R. Lohmayer, *Grassmann higher-order tensor renormalization group approach for two-dimensional strong-coupling QCD*. In: *Nucl. Phys. B* 986 (2023), p. 116032. DOI: 10.1016/j.nuclphysb.2022.116032. arXiv: 2206.00545 [hep-lat].
- [24] P. Milde, *Tensor networks for four-dimensional strong-coupling QCD*. Master's thesis, University of Regensburg, 2023.
- [25] Y. Sugimoto, S. Akiyama, and Y. Kuramashi, *Tensor renormalization group study of cold and dense QCD in the strong coupling limit*. 2026. DOI: 10.48550/arXiv.2601.20690. arXiv: 2601.20690 [hep-lat].
- [26] T. Samberger, *Grassmann higher-order tensor renormalization group approach for two-dimensional QCD in next-to-leading order strong-coupling expansion*. Master's thesis, University of Regensburg, 2022.
- [27] T. Samberger, J. Bloch, R. Lohmayer, and T. Wettig, *Tensor-network formulation of QCD in the strong-coupling expansion*. In: *Nucl. Phys. B* 1024 (2025), p. 117267. DOI: 10.1016/j.nuclphysb.2025.117267.
- [28] T. Samberger, J. Bloch, R. Lohmayer, and T. Wettig, *Order-separated tensor-network method for QCD in the strong-coupling expansion*. 2026. arXiv: 2603.24487 [hep-lat].
- [29] J. Bloch, R. Lohmayer, M. Meister, and M. Nunhofer, *Improved local truncation schemes for the higher-order tensor renormalization group method*. In: *Nucl. Phys. B* 987 (2023), p. 116107. DOI: 10.1016/j.nuclphysb.2023.116107. arXiv: 2210.02266 [hep-lat].
- [30] M. Creutz, *On invariant integration over $SU(N)$* . In: *J. Math. Phys.* 19 (1978), p. 2043. DOI: 10.1063/1.523581.
- [31] J. Carlsson, *Integrals over $SU(N)$* . 2008. arXiv: 0802.3409 [hep-lat].
- [32] G. Gagliardi and W. Unger, *Towards a Dual Representation of Lattice QCD*. In: *PoS LATTICE2018* (2018), p. 224. DOI: 10.22323/1.334.0224. arXiv: 1811.02817 [hep-lat].

- [33] O. Borisenko, S. Voloshyn, and V. Chelnokov, *SU(N) polynomial integrals and some applications*. In: *Rept. Math. Phys.* 85 (2020), pp. 129–145. DOI: 10.1016/S0034-4877(20)30015-X. arXiv: 1812.06069 [hep-lat].
- [34] H. Georgi, *Lie algebras in particle physics*. Westview Press, Boulder, Co., 1999.
- [35] T. Wettig, *Group theory for physicists*. Lecture notes, Regensburg, 2020.
- [36] R. Lohmayer, *Eigenvalue distributions of Wilson loops*. Ph.D. thesis, University of Regensburg, 2010. DOI: 10.5283/epub.16031.
- [37] P. Etingof et al., *Introduction to representation theory*. 2011. arXiv: 0901.0827 [math.RT]. URL: <https://arxiv.org/abs/0901.0827>.
- [38] A. Kirillov, *An introduction to Lie groups and Lie algebras*. Cambridge studies in advanced mathematics. Cambridge University Press, 2008. DOI: 10.1017/CB09780511755156.
- [39] M. D. Schwartz, *Quantum Field Theory and the Standard Model*. Cambridge University Press, 2014.
- [40] H. J. Rothe, *Lattice Gauge Theories*. World Scientific, 2012. DOI: 10.1142/8229.
- [41] I. Montvay and G. Münster, *Quantum fields on a lattice*. Cambridge University Press, 1994. DOI: 10.1017/CB09780511470783.
- [42] F. Halzen and A. D. Martin, *Quarks and leptons: An introductory course in modern particle physics*. John Wiley and Sons, 1984.
- [43] F. Gross et al., *50 Years of Quantum Chromodynamics*. In: *Eur. Phys. J. C* 83 (2023), p. 1125. DOI: 10.1140/epjc/s10052-023-11949-2. arXiv: 2212.11107 [hep-ph].
- [44] A. Schäfer, *Quantum Chromodynamics*. Lecture notes, Regensburg, 2022.
- [45] D. J. Griffiths, *Introduction to elementary particles*. Wiley-VCH, 2010.
- [46] A. Zee, *Quantum field theory in a nutshell*. Princeton University Press, 2010.
- [47] M. E. Peskin and D. Schröder, *An introduction to quantum field theory*. Westview Press, Boulder, Co., 1995.
- [48] M. Maggiore, *A modern introduction to quantum field theory*. Oxford University Press, 2005.
- [49] K. Osterwalder and R. Schrader, *Axioms for Euclidean Green's functions II*. In: *Communications in Mathematical Physics* 42 (1975), pp. 281–305. DOI: 10.1007/BF01608978.
- [50] K. G. Wilson, *Confinement of Quarks*. In: *Phys. Rev. D* 10 (1974), pp. 2445–2459. DOI: 10.1103/PhysRevD.10.2445.
- [51] J. Kogut and L. Susskind, *Hamiltonian formulation of Wilson's lattice gauge theories*. In: *Phys. Rev. D* 11 (1975), pp. 395–408. DOI: 10.1103/PhysRevD.11.395.
- [52] H. B. Nielsen and M. Ninomiya, *No Go Theorem for Regularizing Chiral Fermions*. In: *Phys. Lett. B* 105 (1981), pp. 219–223. DOI: 10.1016/0370-2693(81)91026-1.
- [53] F. Gliozzi, *Spinor algebra of the one-component lattice fermions*. In: *Nucl. Phys. B* 204 (1982), pp. 419–428. DOI: [https://doi.org/10.1016/0550-3213\(82\)90199-7](https://doi.org/10.1016/0550-3213(82)90199-7).
- [54] H. Kluberg-Stern, A. Morel, O. Napoly, and B. Petersson, *Flavours of lagrangian Susskind fermions*. In: *Nucl. Phys. B* 220 (1983), pp. 447–470. DOI: [https://doi.org/10.1016/0550-3213\(83\)90501-1](https://doi.org/10.1016/0550-3213(83)90501-1).
- [55] K. Fukushima and T. Hatsuda, *The phase diagram of dense QCD*. In: *Reports on Progress in Physics* 74 (2010), p. 014001. DOI: 10.1088/0034-4885/74/1/014001.

- [56] G. Aarts, *Introductory lectures on lattice QCD at nonzero baryon number*. In: *Journal of Physics: Conference Series* 706 (2016), p. 022004. DOI: 10.1088/1742-6596/706/2/022004.
- [57] C. Gattringer and K. Langfeld, *Approaches to the sign problem in lattice field theory*. In: *International Journal of Modern Physics A* 31 (2016), p. 1643007. DOI: 10.1142/s0217751x16430077.
- [58] W. Nolting, *Grundkurs Theoretische Physik 6: Statistische Physik*. Springer-Lehrbuch. Springer Berlin Heidelberg, 2006.
- [59] J. Kogut et al., *Deconfinement and Chiral Symmetry Restoration at Finite Temperatures in $SU(2)$ and $SU(3)$ Gauge Theories*. In: *Phys. Rev. Lett.* 50 (1983), pp. 393–396. DOI: 10.1103/PhysRevLett.50.393.
- [60] E. V. Shuryak, *Quantum Chromodynamics and the Theory of Superdense Matter*. In: *Phys. Rept.* 61 (1980), pp. 71–158. DOI: 10.1016/0370-1573(80)90105-2.
- [61] M. Alford, K. Rajagopal, and F. Wilczek, *Color-flavor locking and chiral symmetry breaking in high density QCD*. In: *Nucl. Phys. B* 537 (1999), pp. 443–458. DOI: 10.1016/s0550-3213(98)00668-3.
- [62] Y. Nishida, *Phase structures of strong coupling lattice QCD with finite baryon and isospin density*. In: *Physical Review D* 69 (2004). DOI: 10.1103/PhysRevD.69.094501.
- [63] J. Bloch, R. G. Jha, R. Lohmayer, and M. Meister, *Tensor renormalization group study of the three-dimensional $O(2)$ model*. In: *Phys. Rev. D* 104 (2021), p. 094517. DOI: 10.1103/PhysRevD.104.094517. arXiv: 2105.08066 [hep-lat].
- [64] M. Levin and C. P. Nave, *Tensor Renormalization Group Approach to Two-Dimensional Classical Lattice Models*. In: *Phys. Rev. Lett.* 99 (2007), p. 120601. DOI: 10.1103/PhysRevLett.99.120601. arXiv: cond-mat/0611687 [cond-mat.stat-mech].
- [65] Z.-C. Gu, F. Verstraete, and X.-G. Wen, *Grassmann tensor network states and its renormalization for strongly correlated fermionic and bosonic states*. 2010. DOI: 10.48550/arXiv.1004.2563. arXiv: 1004.2563 [cond-mat.str-el].
- [66] Y. Shimizu and Y. Kuramashi, *Grassmann tensor renormalization group approach to one-flavor lattice Schwinger model*. In: *Phys. Rev. D* 90 (2014), p. 014508. DOI: 10.1103/PhysRevD.90.014508. arXiv: 1403.0642 [hep-lat].
- [67] S. Takeda and Y. Yoshimura, *Grassmann tensor renormalization group for the one-flavor lattice Gross-Neveu model with finite chemical potential*. In: *PTEP* 2015 (2015), 043B01. DOI: 10.1093/ptep/ptv022. arXiv: 1412.7855 [hep-lat].
- [68] P. Milde, J. Bloch, and R. Lohmayer, *Tensor-network simulation of the strong-coupling $U(N)$ model*. In: *PoS LATTICE2021* (2021), p. 462. DOI: 10.22323/1.396.0462. arXiv: 2112.01906 [hep-lat].
- [69] J. Bloch and R. Lohmayer, *Grassmann tensor-network method for strong-coupling QCD*. In: *PoS LATTICE2022* (2022), p. 004. DOI: 10.22323/1.430.0004. arXiv: 2210.08935 [hep-lat].
- [70] L. De Lathauwer, B. De Moor, and J. Vandewalle, *On the Best Rank-1 and Rank- $(R_1, R_2, \dots, R_N)$ Approximation of Higher-Order Tensors*. In: *SIAM Journal on Matrix Analysis and Applications* 21 (2000), pp. 1324–1342. DOI: 10.1137/S0895479898346995.

- [71] H.-H. Zhao, Z.-Y. Xie, T. Xiang, and M. Imada, *Tensor network algorithm by coarse-graining tensor renormalization on finite periodic lattices*. In: *Phys. Rev. B* 93 (2016), p. 125115. DOI: 10.1103/physrevb.93.125115.
- [72] A. Meurer et al., *SymPy: symbolic computing in Python*. In: *PeerJ Computer Science* 3 (2017), e103. DOI: 10.7717/peerj-cs.103.
- [73] M. Fukuma, D. Kadoh, and N. Matsumoto, *Tensor network approach to two-dimensional Yang–Mills theories*. In: *Progress of Theoretical and Experimental Physics* 2021 (2021). DOI: 10.1093/ptep/ptab143.
- [74] E. Witten, *On quantum gauge theories in two dimensions*. In: *Communications in Mathematical Physics* 141 (1991), pp. 153–209. DOI: 10.1007/BF02100009.
- [75] M. Abramowitz and I. A. Stegun, eds., *Handbook of mathematical functions: With formulas, graphs and mathematical tables*. Dover Publications, 1965.
- [76] R. G. Jha, *Finite N unitary matrix model*. 2020. arXiv: 2003.00341 [hep-lat].
- [77] T. Samberger, J. Bloch, and R. Lohmayer, *Grassmann tensor approach for two-dimensional QCD in the strong-coupling expansion*. In: *PoS LATTICE2024* (2025), p. 156. DOI: 10.22323/1.466.0156. arXiv: 2501.19192 [hep-lat].