

Research Reports on Mathematical and Computing Sciences

A trust-region method for box-constrained nonlinear
semidefinite programs

Akihiko Komatsu, Makoto Yamashita

November 2014, B-478

Department of
Mathematical and
Computing Sciences
Tokyo Institute of Technology

SERIES **B:** Operations Research

A trust-region method for box-constrained nonlinear semidefinite programs

Akihiko Komatsu¹ and Makoto Yamashita²

Submitted: November 17, 2014.

Abstract:

We propose a trust-region method for nonlinear semidefinite programs with box-constraints. The penalty barrier method can handle this problem, but the size of variable matrices available in practical time is restricted to be less than 500. We develop a trust-region method based on the approach of Coleman and Li (1996) that utilizes the distance to the boundary of the box-constraints into consideration. To extend this method to the space of positive semidefinite matrices, we devise a new search direction by incorporating the eigenvectors of the variable matrix into the distance to the boundary. In this paper, we establish a global convergence of the proposed method, and preliminary numerical experiments show that our method solves the problems with the size being larger than 5000, and it is faster than the feasible direction for functions with nonlinearity higher than quadratic.

Keywords:

Trust-region method, Nonlinear semidefinite programs

AMS Classification:

90 Operations research, mathematical programming, 90C22 Semidefinite programming, 90C30 Nonlinear programming.

1 Introduction

This paper is concerned with a box-constrained nonlinear semidefinite problem (shortly, box-constrained SDP)

$$\min f(\mathbf{X}) \quad \text{subject to} \quad \mathbf{O} \preceq \mathbf{X} \preceq \mathbf{I}. \quad (1)$$

The variable in this problem is $\mathbf{X} \in \mathbb{S}^n$, and we use $\mathbb{S}^n$ to denote the space of $n \times n$ symmetric matrices. The notation $\mathbf{A} \succeq \mathbf{B}$ for $\mathbf{A}, \mathbf{B} \in \mathbb{S}^n$ means that the matrix $\mathbf{A} - \mathbf{B}$ is positive semidefinite. The matrix $\mathbf{I}$ is the identity matrix of the appropriate dimension. We assume through this paper that the objective function $f : \mathbb{S}^n \rightarrow \mathbb{R}$ is a twice continuously differentiable function on an open set containing the feasible set $\mathcal{F} := \{\mathbf{X} \in \mathbb{S}^n : \mathbf{O} \preceq \mathbf{X} \preceq \mathbf{I}\}$.

The feasible set of (1) can express a more general feasible set, $\{\mathbf{X} \in \mathbb{S}^n : \mathbf{L} \preceq \mathbf{X} \preceq \mathbf{U}\}$ with $\mathbf{L}, \mathbf{U} \in \mathbb{S}^n$ such that $\mathbf{L} \preceq \mathbf{U}$. Since we can assume that $\mathbf{U} - \mathbf{L}$ is positive definite without loss of generality [21], we use a Cholesky factorization matrix $\mathbf{C}$ of $\mathbf{U} - \mathbf{L}$ that satisfies $\mathbf{U} - \mathbf{L} = \mathbf{C}\mathbf{C}^T$ to convert a problem

$$\min f(\mathbf{X}) \quad \text{subject to} \quad \mathbf{L} \preceq \mathbf{X} \preceq \mathbf{U}$$

¹ Equities Department, Tokyo Stock Exchange, INC., 2-1, Nihombashi-kabuto-cho, Chuo-ku, Tokyo 103-8220, Japan.

² (Corresponding Author) Department of Mathematical and Computing Sciences, Tokyo Institute of Technology, 2-12-1-W8-29 Ookayama, Meguro-ku, Tokyo 152-8552, Japan (Makoto.Yamashita@is.titech.ac.jp). The work of the second author was supported by JSPS KAKENHI Grant Numbers 24710151.

into an equivalent problem

$$\min f(\mathbf{C}\bar{\mathbf{X}}\mathbf{C}^T + \mathbf{L}) \quad \text{subject to} \quad \mathbf{O} \preceq \bar{\mathbf{X}} \preceq \mathbf{I}$$

via the relation $\bar{\mathbf{X}} = \mathbf{C}^{-1}(\mathbf{X} - \mathbf{L})\mathbf{C}^{-T}$.

A box-constrained nonlinear optimization problem

$$\min f(\mathbf{x}) \quad \text{subject to} \quad \mathbf{l} \leq \mathbf{x} \leq \mathbf{u}. \quad (2)$$

is an important case of (1). If the variable matrix $\mathbf{X}$ in (1) is a diagonal-block matrix, $\mathbf{X}$ is positive semidefinite if and only if all of the sub matrices are positive semidefinite. In particular, when the dimension of each sub-matrix is one, the box-constrained nonlinear optimization problem (2) emerges. In this problem, the variable is $\mathbf{x} \in \mathbb{R}^n$, and the lower and upper bounds are given by $\mathbf{l} \in \mathbb{R}^n$ and $\mathbf{u} \in \mathbb{R}^n$. The problem (2) is a basic problem in constrained optimization and many methods are proposed. Hei *et. al.* [11] compared the performance of four active-sets methods and two interior-point methods. Trust-region methods are also discussed in [4, 5, 20],

On the other hand, the positive semidefinite condition on a matrix ($\mathbf{X} \succeq \mathbf{O}$) is extensively studied in the context of SDP (semidefinite programs). The range of SDP application is very wide and includes control theory [3], combinatorial optimization [9], polynomial optimization [13] and quantum chemistry [8, 14]. A consider number of studies on SDP can be found in the survey of Todd [19], the handbook edited by Anjos and Lasserre [1] and the references therein.

The penalty barrier method [2, 12] can be applied to solve the box-constrained SDP (1). Though it can handle the problem (1) with/without additional constraints, it requires the full information of the second derivative of the objective function $f(\mathbf{X})$, and it can solve the problems in practical time only when the size of variable matrix is small; $n \leq 500$. To solve larger problems, we should prepare methods specialized for solving (1).

As a method specialized for the box-constrained nonlinear semidefinite problem (1), Xu *et al* [21] proposed a feasible direction method. This method is an iterative method and it tries to find a point which satisfies a first-order optimality condition.

We say that $\mathbf{X}^* \in \mathcal{F}$ satisfies a first-order optimality condition of (1) if

$$\langle \nabla f(\mathbf{X}^*) | \mathbf{X} - \mathbf{X}^* \rangle \geq 0 \quad \text{for} \quad \forall \mathbf{X} \in \mathcal{F}. \quad (3)$$

Here, we use $\langle \mathbf{A} | \mathbf{B} \rangle$ to denote the inner-product between $\mathbf{A} \in \mathbb{S}^n$ and $\mathbf{B} \in \mathbb{S}^n$, and $\nabla f(\mathbf{X}^*) \in \mathbb{S}^n$ is the gradient matrix of f at $\mathbf{X}^*$. In particular, when $f(\mathbf{X})$ is a convex function, $\mathbf{X}^*$ that satisfies (3) is an optimal solution. When $\mathbf{X}^* \in \mathcal{F}$, we can derive an equivalent but more convinient condition;

$$\underline{f}(\mathbf{X}^*) = 0$$

where

$$\underline{f}(\widehat{\mathbf{X}}) := \min \{ \langle \nabla f(\widehat{\mathbf{X}}) | \mathbf{X} - \widehat{\mathbf{X}} \rangle : \mathbf{X} \in \mathcal{F} \}. \quad (4)$$

Xu *et al* [21] proved that the feasible direction method generates an sequence $\{\mathbf{X}^k\} \in \mathcal{F}$ that attains $\lim_{k \rightarrow \infty} \underline{f}(\mathbf{X}^k) = 0$.

From some preliminary experiments, however, we found that the feasible direction methods did not perform well for objective functions with higher nonlinearity. Since the feasible direction computes the search direction by projecting the steepest direction $-\nabla f(\mathbf{X})$ to the feasible set $\mathcal{F}$, it is difficult to capture the effect of the boundary of $\mathcal{F}$.

In this paper, we propose a trust-region method for solving the box-constrained SDP (1). Trust-region methods are iterative methods and they compute the search direction by minimizing

a quadratic approximate function in the trust-region radius. By adjusting the trust-region radius in which the quadratic function approximates the objective function well, trust-region methods are known to attain excellent global convergence properties [6, 10, 17].

We develop the method based on the trust-region method of Coleman and Li [4]. They used the distance from the current point to the boundary of $\mathcal{F}$ to compute the search direction for solving the simple-bound problem (2). However, we can not directly extend their approach to (1), since the positive semidefinite condition $\mathbf{X} \succeq \mathbf{O}$ involves not only eigenvalues but also eigenvectors, and the eigenvectors are not always continuous functions on $\mathbf{X}$. Therefore, we devise a new search direction by taking both the eigenvalues and the eigenvectors into account. This new search direction guarantees non-zero step length, hence it enables us to establish a global convergence of the generated sequence. Through preliminary numerical results which will be reported in this paper, we observe that the proposed trust-region method solve highly nonlinear objective function faster than the feasible direction method, and it can handle large problems than the penalty barrier method implemented in PENLAB [7].

This paper is organized as follows. Section 2 investigates equivalent conditions of the first-optimality conditions. We also introduce the new search direction $\mathbf{D}(\mathbf{X})$, and propose the trust-region method in Algorithm 2.3. Section 3 establishes the global convergence of the proposed method, $\lim_{k \rightarrow \infty} \underline{f}(\mathbf{X}^k) = 0$. Section 4 reports numerical results on the performance comparison of the proposed method, the feasible direction, and the penalty barrier method. Finally, Section 5 gives a conclusion of this paper and discusses future directions.

1.1 Notation and preliminaries

The inner-product between $\mathbf{A} \in \mathbb{R}^{m \times n}$ and $\mathbf{B} \in \mathbb{R}^{m \times n}$ is defined by $\langle \mathbf{A} | \mathbf{B} \rangle := \text{Trace}(\mathbf{A}^T \mathbf{B})$. Here, $\text{Trace}(\mathbf{X})$ for a matrix $\mathbf{X} \in \mathbb{R}^{n \times n}$ is the summation of its diagonal elements, that is, $\text{Trace}(\mathbf{X}) := \sum_{i=1}^n X_{ii}$.

For $\mathbf{A} \in \mathbb{R}^{m \times n}$, we define the Frobenius norm by $\|\mathbf{A}\|_F := \sqrt{\langle \mathbf{A} | \mathbf{A} \rangle}$. From the Cauchy-Schwartz inequality, it holds $|\langle \mathbf{A} | \mathbf{B} \rangle| \leq \|\mathbf{A}\|_F \|\mathbf{B}\|_F$ for $\forall \mathbf{A} \in \mathbb{R}^{m \times n}$ and $\forall \mathbf{B} \in \mathbb{R}^{m \times n}$. We often use the relation $\langle \mathbf{A} | \mathbf{B} \rangle = \langle \mathbf{B} | \mathbf{A} \rangle$. In addition, we use the inequality $\langle \mathbf{A} | \mathbf{B} \rangle \geq 0$ for two positive semidefinite matrices $\mathbf{A} \succeq \mathbf{O}$ and $\mathbf{B} \succeq \mathbf{O}$.

For symmetric matrices, the 2-norm $\|\mathbf{A}\|_2$ is defined by the largest absolute eigenvalue of $\mathbf{A} \in \mathbb{S}^n$. The notation $\text{diag}(\kappa_1, \kappa_2, \dots, \kappa_n)$ stands for the diagonal matrix whose diagonal elements are $\kappa_1, \kappa_2, \dots, \kappa_n$. When $\mathbf{A} = \mathbf{Q}\mathbf{K}\mathbf{Q}^T$ is the eigenvalue decomposition of $\mathbf{A}$ with the diagonal matrix $\mathbf{K} = \text{diag}(\kappa_1, \kappa_2, \dots, \kappa_m)$, the r th power of $\mathbf{A}$ for $r \in \mathbb{R}$ is the symmetric matrix defined by $\mathbf{A}^r := \mathbf{Q}\text{diag}(\kappa_1^r, \kappa_2^r, \dots, \kappa_n^r)\mathbf{Q}^T$.

The gradient matrix $\nabla f(\mathbf{X}) \in \mathbb{S}^n$ and the Hessian map $\nabla^2 f(\mathbf{X})$ at $\mathbf{X} \in \mathbb{S}^n$ are defined in the way that the Taylor expansion for $\mathbf{D} \in \mathbb{S}^n$ holds by

$$f(\mathbf{X} + \mathbf{D}) = f(\mathbf{X}) + \langle \nabla f(\mathbf{X}) | \mathbf{D} \rangle + \frac{1}{2} \langle \mathbf{D} | \nabla^2 f(\mathbf{X}) | \mathbf{D} \rangle + O(\|\mathbf{D}\|_F^2),$$

where $O(d)$ is of the order of d . For example, for the function $\hat{f}(\mathbf{X}) = \langle \mathbf{X} | \mathbf{X} \rangle$, we have $\nabla \hat{f}(\mathbf{X}) = 2\mathbf{X}$ and $\langle \mathbf{D} | \nabla^2 \hat{f}(\mathbf{X}) | \mathbf{D} \rangle = 2\langle \mathbf{D} | \mathbf{D} \rangle$ from the relation $\langle \mathbf{X} + \mathbf{D} | \mathbf{X} + \mathbf{D} \rangle = \langle \mathbf{X} | \mathbf{X} \rangle + 2\langle \mathbf{X} | \mathbf{D} \rangle + \langle \mathbf{D} | \mathbf{D} \rangle$. In particular, $\nabla f(\mathbf{X})$ corresponds to the Fréchet derivative, and for $\mathbf{A}, \mathbf{B} \in \mathbb{S}^n$ we have $\langle \mathbf{A} | \nabla^2 f(\mathbf{X}) | \mathbf{B} \rangle = \sum_{i,j,k,l=1}^n \frac{\partial^2 f(\mathbf{X})}{\partial X_{kl} \partial X_{ij}} A_{ij} B_{kl}$.

Though this paper, we use the matrices $\mathbf{P}(\mathbf{X})$ and $\Lambda(\mathbf{X})$ to denote the eigenvalue decomposition of $\nabla f(\mathbf{X})$ as $\nabla f(\mathbf{X}) = \mathbf{P}(\mathbf{X})\Lambda(\mathbf{X})\mathbf{P}(\mathbf{X})^T$. The matrix $\Lambda(\mathbf{X})$ is the diagonal matrix whose diagonal elements are the descending-order eigenvalues of $\nabla f(\mathbf{X})$, denoted by $\lambda_1(\mathbf{X}) \geq \lambda_2(\mathbf{X}) \geq \dots \geq \lambda_n(\mathbf{X})$. The j th column of $\mathbf{P}(\mathbf{X})$, denoted by $\mathbf{p}_j(\mathbf{X})$, is the associated eigenvector of $\lambda_j(\mathbf{X})$. We use $n_+(\mathbf{X})$ and $n_-(\mathbf{X})$ to denote the number of positive and non-positive eigenvalues of $\nabla f(\mathbf{X})$,

respectively. We divide $\Lambda(\mathbf{X})$ into the two blocks, $\Lambda_+(\mathbf{X}) := \text{diag}(\lambda_1(\mathbf{X}), \lambda_2(\mathbf{X}), \dots, \lambda_{n_+}(\mathbf{X}))$, $\Lambda_-(\mathbf{X}) := \text{diag}(\lambda_{n_+}(\mathbf{X})+1(\mathbf{X}), \lambda_{n_+}(\mathbf{X})+2(\mathbf{X}), \dots, \lambda_n)$. Note that the sizes of $\Lambda_+(\mathbf{X})$ and $\Lambda_-(\mathbf{X})$ can be zero, but the total is kept $n_+(\mathbf{X}) + n_-(\mathbf{X}) = n$. We also divide $P(\mathbf{X})$ into the two matrices $P_+(\mathbf{X}), P_-(\mathbf{X})$ by collecting the corresponding vectors, so the columns of $P_+(\mathbf{X})$ are $\mathbf{p}_1(\mathbf{X}), \dots, \mathbf{p}_{n_+}(\mathbf{X})$ in this order. We can write the eigenvalue decomposition in another form, $\nabla f(\mathbf{X}) = P_+(\mathbf{X})\Lambda_+(\mathbf{X})P_+(\mathbf{X})^T + P_-(\mathbf{X})\Lambda_-(\mathbf{X})P_-(\mathbf{X})^T$. From properties of eigenvectors, we have $P_+(\mathbf{X})^T P_-(\mathbf{X}) = \mathbf{O}$. We also know that $P_+(\mathbf{X})^T P_+(\mathbf{X})$ is the identity matrix of dimension $n_+(\mathbf{X})$ and $P_-(\mathbf{X})^T P_-(\mathbf{X})$ is the identity matrix of dimension $n_-(\mathbf{X})$. Finally, we define $\lambda_{\max}(\mathbf{X}) := \max\{|\lambda_1(\mathbf{X})|, |\lambda_n(\mathbf{X})|\}$. From a property of the 2-norm, it holds that $\lambda_{\max}(\mathbf{X}) = \|\nabla f(\mathbf{X})\|_2$.

2 Trust-region method for box-constrained SDP

For the simple bound problem (2), Coleman and Li [4] proposed a trust-region method which measures the distance from the current feasible point $\mathbf{x}(\mathbf{l} \leq \mathbf{x} \leq \mathbf{u})$ to the boundary of feasible set by the vector $\mathbf{v}(\mathbf{x}) \in \mathbb{R}^n$ defined as

$$v_i(\mathbf{x}) = \begin{cases} x_i - l_i & \text{if } \frac{\partial f(\mathbf{x})}{\partial x_i} > 0 \\ u_i - x_i & \text{if } \frac{\partial f(\mathbf{x})}{\partial x_i} \leq 0. \end{cases}$$

This vector is used to control the approach to the boundary, and the key property in the discussion of [4] is that $\mathbf{x}^*$ satisfies the first-order optimality condition if and only if $\frac{\partial f(\mathbf{x})}{\partial x_i} v_i(\mathbf{x}) = 0$ for each $i = 1, \dots, n$.

We should emphasize that we can not directly extend the definition of $\mathbf{v}(\mathbf{x})$ to box-constrained SDPs (1) using the conditions on the eigenvalue of $\mathbf{X}$, since the distance to the boundary of $\mathcal{F}$ relates to not only the eigenvalues but also the eigenvectors. If we ignore the eigenvectors, it is difficult to ensure the non-zero step length. To take the effect of eigenvectors into account, we define two positive semidefinite matrices for $\mathbf{X} \in \mathcal{F}$;

$$\mathbf{V}_+(\mathbf{X}) := P_+(\mathbf{X})^T \mathbf{X} P_+(\mathbf{X}), \quad \mathbf{V}_-(\mathbf{X}) := P_-(\mathbf{X})^T (\mathbf{I} - \mathbf{X}) P_-(\mathbf{X}).$$

The definition of these matrices brings us other properties of the first-order optimality condition in Lemma 2.1. In the lemma, we use a matrix $\mathbf{D}(\mathbf{X}) \in \mathbb{S}^n$ and a scalar $N(\mathbf{X})$ defined by

$$\begin{aligned} \mathbf{D}(\mathbf{X}) &:= P(\mathbf{X}) \begin{pmatrix} \mathbf{V}_+(\mathbf{X})^{1/2} \Lambda_+(\mathbf{X}) \mathbf{V}_+(\mathbf{X})^{1/2} & \lambda_{\max} P_+(\mathbf{X})^T \mathbf{X} P_-(\mathbf{X}) \\ \lambda_{\max} P_-(\mathbf{X})^T \mathbf{X} P_+(\mathbf{X}) & \mathbf{V}_-(\mathbf{X})^{1/2} \Lambda_-(\mathbf{X}) \mathbf{V}_-(\mathbf{X})^{1/2} \end{pmatrix} P(\mathbf{X})^T \quad (5) \\ N(\mathbf{X}) &:= \langle \nabla f(\mathbf{X}) | \mathbf{D}(\mathbf{X}) \rangle. \quad (6) \end{aligned}$$

In particular, the definition of the matrix $\mathbf{D}(\mathbf{X})$ enables us to extend the trust-region method of Coleman and Li [4] to the space of symmetric matrices, since we can ensure the positiveness of the step length for the direction $\mathbf{D}(\mathbf{X})$ as shown in Lemma 2.2.

Using the relations $\nabla f(\mathbf{X}) = P_+(\mathbf{X})\Lambda_+(\mathbf{X})P_+(\mathbf{X})^T + P_-(\mathbf{X})\Lambda_-(\mathbf{X})P_-(\mathbf{X})^T$, $\mathbf{V}_+(\mathbf{X}) = \mathbf{V}_+(\mathbf{X})^{1/2} \mathbf{V}_+(\mathbf{X})^{1/2}$, $\mathbf{V}_+(\mathbf{X})^{1/2} = \mathbf{V}_+(\mathbf{X})^{1/4} \mathbf{V}_+(\mathbf{X})^{1/4}$, $\mathbf{V}_-(\mathbf{X}) = \mathbf{V}_-(\mathbf{X})^{1/2} \mathbf{V}_-(\mathbf{X})^{1/2}$, and $\mathbf{V}_-(\mathbf{X})^{1/2} = \mathbf{V}_-(\mathbf{X})^{1/4} \mathbf{V}_-(\mathbf{X})^{1/4}$, we can compute $\|\mathbf{D}(\mathbf{X})\|_F^2$ and $N(\mathbf{X})$ as follow;

$$\begin{aligned} \|\mathbf{D}(\mathbf{X})\|_F^2 &= \|\mathbf{V}_+(\mathbf{X})^{1/2} \Lambda_+(\mathbf{X}) \mathbf{V}_+(\mathbf{X})^{1/2}\|_F^2 + \|\mathbf{V}_-(\mathbf{X})^{1/2} \Lambda_-(\mathbf{X}) \mathbf{V}_-(\mathbf{X})^{1/2}\|_F^2 \\ &\quad + 2\lambda_{\max}^2 \|P_+(\mathbf{X})^T \mathbf{X} P_-(\mathbf{X})\|_F^2, \quad (7) \end{aligned}$$

$$N(\mathbf{X}) = \|\mathbf{V}_+(\mathbf{X})^{1/4} \Lambda_+(\mathbf{X}) \mathbf{V}_+(\mathbf{X})^{1/4}\|_F^2 + \|\mathbf{V}_-(\mathbf{X})^{1/4} \Lambda_-(\mathbf{X}) \mathbf{V}_-(\mathbf{X})^{1/4}\|_F^2. \quad (8)$$

Lemma 2.1. For a matrix $\mathbf{X}^* \in \mathcal{F}$, the following conditions are equivalent.

(a) $\mathbf{X}^*$ satisfies the first-order optimality condition (3).

(b) $\langle \Lambda_+(\mathbf{X}^*) \mid \mathbf{V}_+(\mathbf{X}^*) \rangle = \langle \Lambda_-(\mathbf{X}^*) \mid \mathbf{V}_-(\mathbf{X}^*) \rangle = 0$.

(c) $N(\mathbf{X}^*) = \mathbf{O}$.

(d) $\|\mathbf{D}(\mathbf{X}^*)\|_F = 0$.

Proof: For simplicity, we use $\Lambda_+ := \Lambda_+(\mathbf{X}^*)$, $\Lambda_- := \Lambda_-(\mathbf{X}^*)$, $\mathbf{P}_+ := \mathbf{P}_+(\mathbf{X}^*)$, $\mathbf{P}_- := \mathbf{P}_-(\mathbf{X}^*)$, $\mathbf{V}_+ := \mathbf{P}_+(\mathbf{X}^*)^T \mathbf{X}^* \mathbf{P}_+(\mathbf{X}^*)$, $\mathbf{V}_- := \mathbf{P}_-(\mathbf{X}^*)^T \mathbf{X}^* \mathbf{P}_-(\mathbf{X}^*)$, $\mathbf{D} := \mathbf{D}(\mathbf{X}^*)$ in this proof.

[(a) $\Rightarrow$ (b)] When we set $\widehat{\mathbf{X}} = \mathbf{P}_+ \mathbf{P}_+^T \mathbf{X}^* \mathbf{P}_+ \mathbf{P}_+^T + \mathbf{P}_- \mathbf{P}_-^T \succeq \mathbf{O}$, the property $\mathbf{I} - \widehat{\mathbf{X}} = (\mathbf{P}_+ \mathbf{P}_+^T + \mathbf{P}_- \mathbf{P}_-^T) - (\mathbf{P}_+ \mathbf{P}_+^T \mathbf{X}^* \mathbf{P}_+ \mathbf{P}_+^T + \mathbf{P}_- \mathbf{P}_-^T) = \mathbf{P}_+ \mathbf{P}_+^T (\mathbf{I} - \mathbf{X}^*) \mathbf{P}_+ \mathbf{P}_+^T \succeq \mathbf{O}$ indicates $\widehat{\mathbf{X}} \in \mathcal{F}$. Substituting $\widehat{\mathbf{X}}$ into the inequality $\langle \nabla f(\mathbf{X}^*) \mid \mathbf{X} - \mathbf{X}^* \rangle \geq 0$ leads to

$$\begin{aligned} \langle \nabla f(\mathbf{X}^*) \mid \widehat{\mathbf{X}} - \mathbf{X}^* \rangle &= \langle \mathbf{P}_+ \Lambda_+ \mathbf{P}_+^T + \mathbf{P}_- \Lambda_- \mathbf{P}_-^T \mid \mathbf{P}_+ \mathbf{P}_+^T \mathbf{X}^* \mathbf{P}_+ \mathbf{P}_+^T + \mathbf{P}_- \mathbf{P}_-^T - \mathbf{X}^* \rangle \\ &= \langle \Lambda_- \mid \mathbf{I} \rangle - \langle \Lambda_- \mid \mathbf{P}_-^T \mathbf{X}^* \mathbf{P}_- \rangle = \langle \Lambda_- \mid \mathbf{V}_- \rangle \geq 0. \end{aligned}$$

In the above equalities, we used $\langle \mathbf{A} \mid \mathbf{B} \rangle = \text{Trace}(\mathbf{A}^T \mathbf{B}) = \text{Trace}(\mathbf{B}^T \mathbf{A})$, $\mathbf{P}_+^T \mathbf{P}_+ = \mathbf{I}$ and $\mathbf{P}_+^T \mathbf{P}_- = \mathbf{O}$. Since $-\Lambda_- \succeq \mathbf{O}$ and $\mathbf{V}_- \succeq \mathbf{O}$, we also have $\langle -\Lambda_- \mid \mathbf{V}_- \rangle \geq 0$, so that we obtain $\langle \Lambda_- \mid \mathbf{V}_- \rangle = 0$. Similarly, the usage of $\widehat{\mathbf{X}} = \mathbf{P}_- \mathbf{P}_-^T \mathbf{X}^* \mathbf{P}_- \mathbf{P}_-^T \in \mathcal{F}$ shows $\langle \Lambda_+ \mid \mathbf{V}_+ \rangle = 0$.

[(b) $\Rightarrow$ (a)] Since $\mathbf{P}_+^T \mathbf{X} \mathbf{P}_+ \succeq \mathbf{O}$ and $\mathbf{P}_-^T (\mathbf{I} - \mathbf{X}) \mathbf{P}_- \succeq \mathbf{O}$ for any $\mathbf{X} \in \mathcal{F}$, it holds that

$$\begin{aligned} &\langle \nabla f(\mathbf{X}^*) \mid \mathbf{X} - \mathbf{X}^* \rangle \\ &= \langle \mathbf{P}_+ \Lambda_+ \mathbf{P}_+^T + \mathbf{P}_- \Lambda_- \mathbf{P}_-^T \mid \mathbf{X} - \mathbf{X}^* \rangle \\ &= \langle \Lambda_+ \mid \mathbf{P}_+^T \mathbf{X} \mathbf{P}_+ \rangle - \langle \Lambda_+ \mid \mathbf{V}_+ \rangle - \langle \Lambda_- \mid \mathbf{P}_-^T (\mathbf{I} - \mathbf{X}) \mathbf{P}_- \rangle + \langle \Lambda_- \mid \mathbf{V}_- \rangle \\ &= \langle \Lambda_+ \mid \mathbf{P}_+^T \mathbf{X} \mathbf{P}_+ \rangle + \langle -\Lambda_- \mid \mathbf{P}_-^T (\mathbf{I} - \mathbf{X}) \mathbf{P}_- \rangle \geq 0. \end{aligned}$$

For the last equality, we used $\langle \Lambda_+ \mid \mathbf{V}_+ \rangle = \langle \Lambda_- \mid \mathbf{V}_- \rangle = 0$ from (b).

[(b) $\Rightarrow$ (c)] Since $\langle \Lambda_+ \mid \mathbf{V}_+ \rangle = \text{Trace}(\mathbf{V}_+^{1/2} \Lambda_+ \mathbf{V}_+^{1/2})$ and $\mathbf{V}_+^{1/2} \Lambda_+ \mathbf{V}_+^{1/2} \succeq \mathbf{O}$, $\langle \Lambda_+ \mid \mathbf{V}_+ \rangle = 0$ is equivalent to $\mathbf{V}_+^{1/2} \Lambda_+ \mathbf{V}_+^{1/2} = \mathbf{O}$. From $\mathbf{V}_+^{1/4} \succeq \mathbf{O}$, the condition $\mathbf{V}_+^{1/2} \Lambda_+ \mathbf{V}_+^{1/2} = \mathbf{O}$ is further equivalent to $\mathbf{V}_+^{1/4} \Lambda_+ \mathbf{V}_+^{1/4} = \mathbf{O}$. Similarly, $\langle \Lambda_- \mid \mathbf{V}_- \rangle = 0$ is equivalent to $\mathbf{V}_-^{1/4} \Lambda_- \mathbf{V}_-^{1/4} = \mathbf{O}$. Hence, we obtain (c) by (8).

[(c) $\Rightarrow$ (d)] From (c), we obtain $\mathbf{V}_+^{1/2} \Lambda_+ \mathbf{V}_+^{1/2} = \mathbf{O}$ and $\mathbf{V}_-^{1/2} \Lambda_- \mathbf{V}_-^{1/2} = \mathbf{O}$. Since all the eigenvalues in Λ_+ are positive, the conditions $\langle \Lambda_+ \mid \mathbf{V}_+ \rangle = 0$ and $\mathbf{V}_+ \succeq \mathbf{O}$, indicate $\mathbf{V}_+ = \mathbf{O}$. Furthermore, the decomposition $\mathbf{V}_+ = \mathbf{P}_+^T (\mathbf{X}^*)^{1/2} (\mathbf{X}^*)^{1/2} \mathbf{P}_+ = \mathbf{O}$ implies $\mathbf{P}_+^T (\mathbf{X}^*)^{1/2} = \mathbf{O}$. Therefore, it holds that $\mathbf{P}_+^T \mathbf{X}^* \mathbf{P}_- = \mathbf{P}_+^T (\mathbf{X}^*)^{1/2} (\mathbf{X}^*)^{1/2} \mathbf{P}_- = \mathbf{O}$. Hence, we conclude $\|\mathbf{D}\|_F = 0$ from (7).

[(d) $\Rightarrow$ (b)] From the relation (7), $\|\mathbf{D}\|_F = 0$ indicates $\mathbf{V}_+^{1/2} \Lambda_+ \mathbf{V}_+^{1/2} = \mathbf{O}$ and $\mathbf{V}_-^{1/2} \Lambda_- \mathbf{V}_-^{1/2} = \mathbf{O}$. By taking these traces of these matrices, we obtain (b). $\square$

Lemma 2.1 and (8) indicate that when $\mathbf{X}$ does not satisfy the first-order optimality condition, we can take $-\frac{\mathbf{D}(\mathbf{X})}{\|\mathbf{D}(\mathbf{X})\|_F}$ as a descent direction of $f(\mathbf{X})$, that is, $\langle \nabla f(\mathbf{X}) \mid -\frac{\mathbf{D}(\mathbf{X})}{\|\mathbf{D}(\mathbf{X})\|_F} \rangle < 0$. Hence, we can expect that the decrease of the objective function $f(\mathbf{X} - \alpha \frac{\mathbf{D}(\mathbf{X})}{\|\mathbf{D}(\mathbf{X})\|_F}) < f(\mathbf{X})$ for a certain value $\alpha > 0$. The next lemma gives a range of α to ensure $\mathbf{X} - \alpha \frac{\mathbf{D}(\mathbf{X})}{\|\mathbf{D}(\mathbf{X})\|_F} \in \mathcal{F}$.

Lemma 2.2. If $\mathbf{X} \in \mathcal{F}$ does not satisfy the first-order optimality condition, then $\mathbf{X} - \alpha \frac{\mathbf{D}(\mathbf{X})}{\|\mathbf{D}(\mathbf{X})\|_F} \in \mathcal{F}$ for $\alpha \in [0, \frac{\|\mathbf{D}(\mathbf{X})\|_F}{\lambda_{\max}(\mathbf{X})}]$.

Proof:

In this proof, we drop $(\mathbf{X})$ from $\mathbf{P}(\mathbf{X})$, $\mathbf{D}(\mathbf{X})$, $\mathbf{V}_+(\mathbf{X})$, $\mathbf{V}_-(\mathbf{X})$, $\lambda_{\max}(\mathbf{X})$ for simplicity. From the definition of $\lambda_{\max}$, the matrix $\mathbf{I} - \frac{\Lambda_+}{\lambda_{\max}}$ is a diagonal matrix with nonnegative diagonal elements, hence this matrix is positive semidefinite. Using the property $\mathbf{P}\mathbf{P}^T = \mathbf{I}$, it holds

$$\begin{aligned} \mathbf{X} - \frac{\mathbf{D}}{\lambda_{\max}} &= \mathbf{P}\mathbf{P}^T \left(\mathbf{X} - \frac{\mathbf{D}}{\lambda_{\max}} \right) \mathbf{P}\mathbf{P}^T \\ &= \mathbf{P} \left\{ \mathbf{P}^T \mathbf{X} \mathbf{P} - \begin{pmatrix} \mathbf{V}_+^{1/2} \frac{\Lambda_+}{\lambda_{\max}} \mathbf{V}_+^{1/2} & \mathbf{P}_+^T \mathbf{X} \mathbf{P}_- \\ \mathbf{P}_-^T \mathbf{X} \mathbf{P}_+ & \mathbf{V}_-^{1/2} \frac{\Lambda_-}{\lambda_{\max}} \mathbf{V}_-^{1/2} \end{pmatrix} \right\} \mathbf{P}^T \\ &= \mathbf{P} \left\{ \begin{pmatrix} \mathbf{V}_+^{1/2} \left(\mathbf{I} - \frac{\Lambda_+}{\lambda_{\max}} \right) \mathbf{V}_+^{1/2} & \mathbf{O} \\ \mathbf{O} & \mathbf{V}_-^{1/2} \left(-\frac{\Lambda_-}{\lambda_{\max}} \right) \mathbf{V}_-^{1/2} \end{pmatrix} \right\} \mathbf{P}^T \succeq \mathbf{O}. \end{aligned}$$

Due to the linear combination $\mathbf{X} - \alpha \frac{\mathbf{D}}{\|\mathbf{D}\|_F} = \left(1 - \alpha \frac{\lambda_{\max}}{\|\mathbf{D}\|_F}\right) \mathbf{X} + \alpha \frac{\lambda_{\max}(\mathbf{X})}{\|\mathbf{D}\|_F} \left(\mathbf{X} - \frac{\mathbf{D}}{\lambda_{\max}} \right)$, we obtain $\mathbf{X} - \alpha \frac{\mathbf{D}}{\|\mathbf{D}\|_F} \succeq \mathbf{O}$ for $\alpha \in [0, \frac{\|\mathbf{D}\|_F}{\lambda_{\max}}]$.

In a similar way, we can show that $\mathbf{I} - (\mathbf{X} - \frac{\mathbf{D}}{\lambda_{\max}}) \succeq \mathbf{O}$, and this leads to $\mathbf{X} - \alpha \frac{\mathbf{D}}{\|\mathbf{D}\|_F} \preceq \mathbf{I}$ for $\alpha \in [0, \frac{\|\mathbf{D}\|_F}{\lambda_{\max}}]$. □

We propose a trust-region method for the box-constrained SDP (1) as Algorithm 2.3. Based on the property that $-\frac{\mathbf{D}(\mathbf{X})}{\|\mathbf{D}(\mathbf{X})\|_F}$ is a descent direction of $f(\mathbf{X})$, we can set a matrix $-\mathbf{S}(\mathbf{X})$, where $\mathbf{S}(\mathbf{X}) := \frac{\mathbf{D}(\mathbf{X})}{\|\mathbf{D}(\mathbf{X})\|_F}$, as a normalized search direction to find a local minimizer. In Algorithm 2.3, we use a quadratic approximation of f with the direction $\mathbf{S}(\mathbf{X})$;

$$q(\alpha, \mathbf{X}) := f(\mathbf{X}) - \alpha \langle \nabla f(\mathbf{X}) \mid \mathbf{S}(\mathbf{X}) \rangle + \frac{\alpha^2}{2} \langle \mathbf{S}(\mathbf{X}) \mid \nabla^2 f(\mathbf{X}) \mid \mathbf{S}(\mathbf{X}) \rangle.$$

From Lemma 2.2, the generated sequence by Algorithm 2.3 remains in $\mathcal{F}$, that is, $\{\mathbf{X}^k\} \subset \mathcal{F}$.

3 Convergence properties

As noted in Section 1, $\mathbf{X}^* \in \mathcal{F}$ satisfies the first-order optimality condition (3) if and only if $\underline{f}(\mathbf{X}^*) = 0$. In this section, we show that the sequence $\{\mathbf{X}^k\} \subset \mathcal{F}$ generated by Algorithm 2.3 with the stopping threshold $\epsilon = 0$ attains $\lim_{k \rightarrow \infty} \underline{f}(\mathbf{X}^k) = 0$. We divide the proof into two parts. The first part shows there exists a subsequence of $\{N(\mathbf{X}^k)\}$ that converges to zero. The second part shows $\lim_{k \rightarrow \infty} N(\mathbf{X}^k) = 0$ in Theorem 3.6, and finally the global convergence $\lim_{k \rightarrow \infty} \underline{f}(\mathbf{X}_k) = 0$ in Theorem 3.7.

Though we can employ similar approaches in the proof of the first part as [4] by the usage of $\mathbf{D}(\mathbf{X})$, we can not directly apply the results of [4] to the second part. This is mainly because that the eigenvector matrices $\mathbf{P}_+(\mathbf{X})$ and $\mathbf{P}_-(\mathbf{X})$ are not always continuous functions of $\mathbf{X}$. Instead, our proof relies on the boundedness of $\langle \Lambda_+(\mathbf{X}^k) \mid \mathbf{V}_+(\mathbf{X}^k) \rangle$ and $\langle -\Lambda_-(\mathbf{X}^k) \mid \mathbf{V}_-(\mathbf{X}^k) \rangle$.

3.1 Convergence of subsequence

To analyze Algorithm 2.3, we define constant values

$$M_1 := \max_{\mathbf{X} \in \mathcal{F}} \|\nabla f(\mathbf{X})\|_2,$$

Algorithm 2.3. Trust-region method for the box-constrained SDP

Step 1: Choose an initial point $\mathbf{X}^0 \in \mathcal{F}$. Set an initial trust-region radius $\Delta_0 > 0$. Choose parameters $0 < \mu_1 < \mu_2 < 1$, $0 < \gamma_1 < 1 < \gamma_2$. Set an iteration count $k = 0$. Set a stopping threshold $\epsilon > 0$.

Step 2: If $N(\mathbf{X}^k) < \epsilon$, output $\mathbf{X}^k$ as a solution and stop.

Step 3: Solve a one dimensional optimization problem with respect to α ;

$$\min \quad q(\alpha, \mathbf{X}^k) \quad \text{subject to} \quad 0 \leq \alpha \leq \min \left\{ \frac{\|\mathbf{D}(\mathbf{X}^k)\|_F}{\lambda_{\max}(\mathbf{X}^k)}, \Delta_k \right\}, \quad (9)$$

and let the step length α_k be the minimizer of (9).

Step 4: Let $\bar{\mathbf{X}}^k := \mathbf{X}^k - \alpha_k \mathbf{S}(\mathbf{X}^k)$. Compute the ratio

$$r_k := \frac{f(\mathbf{X}^k) - f(\bar{\mathbf{X}}^k)}{f(\mathbf{X}^k) - q(\alpha_k, \mathbf{X}^k)}, \quad (10)$$

and set

$$\mathbf{X}^{k+1} = \begin{cases} \bar{\mathbf{X}}^k & \text{if } r_k \geq \mu_1 \\ \mathbf{X}^k & \text{otherwise.} \end{cases}$$

Step 5: Update the trust-region radius Δ_k by

$$\Delta_{k+1} = \begin{cases} \gamma_1 \Delta_k & \text{if } r_k < \mu_1 \\ \Delta_k & \text{if } \mu_1 \leq r_k \leq \mu_2 \\ \gamma_2 \Delta_k & \text{if } r_k > \mu_2. \end{cases}$$

Step 6: Set $k \leftarrow k + 1$ and return to Step 2.

$$M_2 := \max_{\mathbf{X} \in \mathcal{F}, \mathbf{D} \in \mathbb{S}^n, \mathbf{D} \neq \mathbf{O}} \left| \frac{\langle \mathbf{D} | \nabla^2 f(\mathbf{X}) | \nabla \mathbf{D} \rangle}{\langle \mathbf{D} | \mathbf{D} \rangle} \right|.$$

Note that M_1 and M_2 are finite from the assumptions that the feasible set $\mathcal{F}$ is a bounded closed set and the objective function $f(\mathbf{X})$ is a twice continuously differentiable function on an open set containing $\mathcal{F}$. We can assume that $M_1 > 0$ and $M_2 > 0$ without loss of generality. When $M_1 = 0$, $f(\mathbf{X})$ is a constant function; and if $M_2 = 0$, then $\nabla f(\mathbf{X})$ is a constant matrix, so that the global minimizer can be obtained as $\mathbf{X}^* = \mathbf{P}_-(\mathbf{X})\mathbf{P}_-(\mathbf{X})^T$ from the constant matrix $\nabla f(\mathbf{X}) = \mathbf{P}(\mathbf{X})\Lambda(\mathbf{X})\mathbf{P}(\mathbf{X})^T$.

We now evaluate the quadratic approximate function $q(\alpha^k, \mathbf{X}^k)$ compared with $f(\mathbf{X}^k)$.

Lemma 3.1. *The step length α_k in Step 3 satisfies*

$$q(\alpha_k, \mathbf{X}^k) \leq f(\mathbf{X}^k) - \frac{1}{2} \min \left\{ \frac{N(\mathbf{X}^k)^2}{M_2 \|\mathbf{D}(\mathbf{X}^k)\|_F^2}, \frac{N(\mathbf{X}^k)}{\lambda_{\max}(\mathbf{X}^k)}, \frac{\Delta_k N(\mathbf{X}^k)}{\|\mathbf{D}(\mathbf{X}^k)\|_F} \right\}.$$

Proof:

In this proof, we use $\mathbf{D} := \mathbf{D}(\mathbf{X}^k)$, $\mathbf{S} := \mathbf{S}(\mathbf{X}^k)$, $N := N(\mathbf{X}^k)$, $\lambda_{\max} := \lambda_{\max}(\mathbf{X}^k)$.

We define a quadratic function $\phi(\alpha) := -\alpha \frac{N}{\|\mathbf{D}\|_F} + \frac{\alpha^2}{2} M_2$. From the definitions of N and M_2 , we have $q(\alpha, \mathbf{X}^k) \leq f(\mathbf{X}^k) + \phi(\alpha)$, hence,

$$q(\alpha_k, \mathbf{X}^k) \leq f(\mathbf{X}^k) + \min_{\alpha \in \left[0, \min\left\{\frac{\|\mathbf{D}\|_F}{\lambda_{\max}}, \Delta_k\right\}\right]} \phi(\alpha).$$

Since $N = \langle \nabla f(\mathbf{X}^k) \mid \mathbf{D} \rangle > 0$ (otherwise, $\mathbf{X}^k$ already satisfies the first-order optimality condition from Lemma 2.1) and $\phi(\alpha)$ is a quadratic function with respect to α , the minimum of ϕ is attained at one of the three candidates $\frac{\|\mathbf{D}\|_F}{\lambda_{\max}}$, Δ_k or $\hat{\alpha} := \frac{N}{M_2 \|\mathbf{D}\|_F}$. Let $\alpha_{\min}$ be the minimizer of $\phi(\alpha)$ subject to $0 \leq \alpha \leq \min\left\{\frac{\|\mathbf{D}\|_F}{\lambda_{\max}}, \Delta_k\right\}$.

If $\alpha_{\min} = \hat{\alpha}$, we have $\phi(\hat{\alpha}) = -\frac{1}{2} \frac{N^2}{M_2 \|\mathbf{D}\|_F^2}$. For the case when $\alpha_{\min} = \frac{\|\mathbf{D}\|_F}{\lambda_{\max}}$, we have $\frac{\|\mathbf{D}\|_F}{\lambda_{\max}} \leq \hat{\alpha}$, therefore, $\frac{\|\mathbf{D}\|_F^2}{\lambda_{\max}} M_2 \leq N$. Hence, it holds that $\phi\left(\frac{\|\mathbf{D}\|_F}{\lambda_{\max}}\right) = -\frac{N}{\lambda_{\max}} + \frac{1}{2} \frac{\|\mathbf{D}\|_F^2}{\lambda_{\max}^2} M_2 \leq -\frac{1}{2} \frac{N}{\lambda_{\max}}$. Finally, when $\alpha_{\min} = \Delta_k$, the inequality $\Delta_k \leq \hat{\alpha}$ indicates that $\Delta_k \leq \frac{N}{M_2 \|\mathbf{D}\|_F}$. Hence, it holds that $\phi(\Delta_k) = -\Delta_k \frac{N}{\|\mathbf{D}\|_F} + \frac{1}{2} \Delta_k^2 M_2 \leq -\Delta_k \frac{N}{\|\mathbf{D}\|_F} + \frac{1}{2} \Delta_k \frac{N}{\|\mathbf{D}\|_F} \leq -\frac{1}{2} \frac{\Delta_k N}{\|\mathbf{D}\|_F}$.

Taking the maximum of the three cases, we obtain the inequality of this lemma. $\square$

To simplify the inequality of Lemma 3.1, we replace $\lambda_{\max}(\mathbf{X}^k)$ and $\|\mathbf{D}(\mathbf{X}^k)\|_F$ by convenient upper bounds. Since $\lambda_{\max}(\mathbf{X}^k)$ is bounded by M_1 , we will seek an upper bound of $\|\mathbf{D}(\mathbf{X}^k)\|_F$.

Lemma 3.2. For $\mathbf{X} \in \mathcal{F}$, it holds that $\|\mathbf{D}(\mathbf{X})\|_F^2 \leq N(\mathbf{X}) + \frac{1}{2} M_1^2 n^3$.

Proof: In this proof, we use simplified notation $\mathbf{D} := \mathbf{D}(\mathbf{X})$, $\mathbf{V}_+ := \mathbf{V}_+(\mathbf{X})$, $\mathbf{V}_- := \mathbf{V}_-(\mathbf{X})$, $\Lambda_+ := \Lambda_+(\mathbf{X})$, $\Lambda_- := \Lambda_-(\mathbf{X})$, $n_+ := n_+(\mathbf{X})$, $n_- := n_-(\mathbf{X})$, $\mathbf{P}_+ := \mathbf{P}_+(\mathbf{X})$, $\mathbf{P}_- := \mathbf{P}_-(\mathbf{X})$. Let $\mathbf{V}_+ = \mathbf{Q} \mathbf{K} \mathbf{Q}^T$ be the eigenvalue decomposition of $\mathbf{V}_+$ such that $\mathbf{K} = \text{diag}(\kappa_1, \kappa_2, \dots, \kappa_{n_+})$ is the diagonal matrix with the eigenvalues of $\mathbf{V}_+$. Since $\mathbf{O} \preceq \mathbf{X} \preceq \mathbf{I}$, we have $\mathbf{O} \preceq \mathbf{V}_+ \preceq \mathbf{I}$, hence, $0 \leq \kappa_i \leq 1$ for $i = 1, 2, \dots, n_+$. Using a matrix $\mathbf{W} \in \mathbb{S}^{n_+}$ defined by $\mathbf{W} := \mathbf{Q}^T \Lambda_+ \mathbf{Q}$, we compare $\|\mathbf{V}_+^{1/2} \Lambda_+ \mathbf{V}_+^{1/2}\|_F$ and $\|\mathbf{V}_+^{1/4} \Lambda_+ \mathbf{V}_+^{1/4}\|_F$;

$$\begin{aligned} & \|\mathbf{V}_+^{1/4} \Lambda_+ \mathbf{V}_+^{1/4}\|_F^2 - \|\mathbf{V}_+^{1/2} \Lambda_+ \mathbf{V}_+^{1/2}\|_F^2 \\ &= \langle \Lambda_+ \mid \mathbf{V}_+^{1/2} \Lambda_+ \mathbf{V}_+^{1/2} \rangle - \langle \mathbf{V}_+^{1/2} \Lambda_+ \mathbf{V}_+^{1/2} \mid \mathbf{V}_+^{1/2} \Lambda_+ \mathbf{V}_+^{1/2} \rangle \\ &= \langle \Lambda_+ - \mathbf{V}_+^{1/2} \Lambda_+ \mathbf{V}_+^{1/2} \mid \mathbf{V}_+^{1/2} \Lambda_+ \mathbf{V}_+^{1/2} \rangle \\ &= \langle \Lambda_+ - \mathbf{Q} \mathbf{K}^{1/2} \mathbf{Q}^T \Lambda_+ \mathbf{Q} \mathbf{K}^{1/2} \mathbf{Q}^T \mid \mathbf{Q} \mathbf{K}^{1/2} \mathbf{Q}^T \Lambda_+ \mathbf{Q} \mathbf{K}^{1/2} \mathbf{Q}^T \rangle \\ &= \langle \mathbf{W} - \mathbf{K}^{1/2} \mathbf{W} \mathbf{K}^{1/2} \mid \mathbf{K}^{1/2} \mathbf{W} \mathbf{K}^{1/2} \rangle \\ &= \|\mathbf{K}^{1/4} \mathbf{W} \mathbf{K}^{1/4}\|_F^2 - \|\mathbf{K}^{1/2} \mathbf{W} \mathbf{K}^{1/2}\|_F^2 \\ &= \sum_{i=1}^{n_+} \sum_{j=1}^{n_+} (W_{ij} \kappa_i^{1/4} \kappa_j^{1/4})^2 - \sum_{i=1}^{n_+} \sum_{j=1}^{n_+} (W_{ij} \kappa_i^{1/2} \kappa_j^{1/2})^2 \geq 0. \end{aligned}$$

The last inequality comes from $0 \leq \kappa_i \leq \kappa_i^{1/2} \leq \kappa_i^{1/4} \leq 1$. In a similar way, we also have $\|\mathbf{V}_-^{1/2} \Lambda_- \mathbf{V}_-^{1/2}\|_F^2 \leq \|\mathbf{V}_-^{1/4} \Lambda_- \mathbf{V}_-^{1/4}\|_F^2$. We evaluate the last term of (7) by the property of the Frobenius norm;

$$\|\mathbf{P}_+^T \mathbf{X} \mathbf{P}_-\|_F^2 \leq \|\mathbf{P}_+\|_F^2 \cdot \|\mathbf{X}\|_F^2 \cdot \|\mathbf{P}_-\|_F^2 \leq n_+ \cdot n \cdot n_- \leq \frac{n^3}{4}.$$

For the last inequality, we used the relation $n_+ + n_- = n$ to derive $n_+ \cdot n_- \leq \frac{n^2}{4}$.

Consequently, it holds from (7) that

$$\begin{aligned} \|\mathbf{D}\|_F^2 &= \|\mathbf{V}_+^{1/2} \Lambda_+ \mathbf{V}_+^{1/2}\|_F^2 + \|\mathbf{V}_-^{1/2} \Lambda_- \mathbf{V}_-^{1/2}\|_F^2 + 2\lambda_{\max}^2 \|\mathbf{P}_+ \mathbf{X} \mathbf{P}_-\|_F^2 \\ &\leq \|\mathbf{V}_+^{1/4} \Lambda_+ \mathbf{V}_+^{1/4}\|_F^2 + \|\mathbf{V}_-^{1/4} \Lambda_- \mathbf{V}_-^{1/4}\|_F^2 + 2\lambda_{\max}^2 \frac{n^3}{4} \\ &\leq N + \frac{1}{2} M_1^2 n^3. \end{aligned}$$

□

We put Lemma 3.2 into Lemma 3.1 to obtain a new upper bound of $q(\alpha_k, \mathbf{X}^k)$;

$$q(\alpha_k, \mathbf{X}^k) \leq f(\mathbf{X}^k) - \frac{1}{2} \min \left\{ \frac{N(\mathbf{X}^k)^2}{M_2 (N(\mathbf{X}^k) + \frac{1}{2} M_1^2 n^3)}, \frac{N(\mathbf{X}^k)}{M_1}, \frac{\Delta_k N(\mathbf{X}^k)}{\sqrt{N(\mathbf{X}^k) + \frac{1}{2} M_1^2 n^3}} \right\}. \quad (11)$$

In Algorithm 2.3, we call the k th iteration a *successful* iteration if $\mathbf{X}^{k+1}$ is set $\bar{\mathbf{X}}^k$ in Step 4, that is, $r_k \geq \mu_1$. Otherwise, the k th iteration is called an *unsuccessful* iteration. For a successful iteration, we obtain a decrease in the objective function

$$\begin{aligned} f(\mathbf{X}^{k+1}) &\leq f(\mathbf{X}^k) - \mu_1 (f(\mathbf{X}^k) - q(\alpha_k, \mathbf{X}^k)) \\ &\leq f(\mathbf{X}^k) - \frac{\mu_1}{2} \min \left\{ \frac{N(\mathbf{X}^k)^2}{M_2 (N(\mathbf{X}^k) + \frac{1}{2} M_1^2 n^3)}, \frac{N(\mathbf{X}^k)}{M_1}, \frac{\Delta_k N(\mathbf{X}^k)}{\sqrt{N(\mathbf{X}^k) + \frac{1}{2} M_1^2 n^3}} \right\}. \quad (12) \end{aligned}$$

Since $f(\mathbf{X}^{k+1}) = f(\mathbf{X}^k)$ in the unsuccessful iterations, the objective value $f(\mathbf{X}^k)$ is non-increasing in Algorithm 2.3.

We are now prepared to show that there exists a subsequence of $\{N(\mathbf{X}^k)\}$ that converges to zero.

Theorem 3.3. *When $\{\mathbf{X}^k\}$ is the sequence generated by Algorithm 2.3 with the stopping threshold $\epsilon = 0$, it holds that*

$$\liminf_{k \rightarrow \infty} N(\mathbf{X}^k) = 0.$$

Proof: We assume that there exists $\hat{\epsilon} > 0$ and k_0 such that $N(\mathbf{X}^k) \geq \hat{\epsilon}$ for any $k \geq k_0$, and we will derive a contradiction.

Let $\mathcal{K} = \{k_1, k_2, \dots, k_i, \dots\}$ be the successful iterations. If $\mathcal{K}$ is a finite sequence, let k_i be the last iteration of $\mathcal{K}$. Since all of the iterations after k_i are unsuccessful, the update rule of Δ_k (Step 5 of Algorithm 2.3) implies $\Delta_{k_i+j} = \gamma_1^j \Delta_{k_i}$. Hence, we obtain $\lim_{j \rightarrow \infty} \Delta_j = 0$. Next, we consider the case when $\mathcal{K}$ is an infinite sequence. The function $\frac{x^2}{x + \frac{1}{2} M_1^2 n^3}$ is an increasing function for $x > 0$, so that it holds from (12) that for $k_i \in \mathcal{K}$,

$$\begin{aligned} f(\mathbf{X}^{k_i+1}) &\leq f(\mathbf{X}^{k_i}) - \frac{\mu_1}{2} \min \left\{ \frac{N(\mathbf{X}^{k_i})^2}{M_2 (N(\mathbf{X}^{k_i}) + \frac{1}{2} M_1^2 n^3)}, \frac{N(\mathbf{X}^{k_i})}{M_1}, \frac{\Delta_{k_i} N(\mathbf{X}^{k_i})}{\sqrt{N(\mathbf{X}^{k_i}) + \frac{1}{2} M_1^2 n^3}} \right\} \\ &\leq f(\mathbf{X}^{k_i}) - \frac{\mu_1}{2} \min \left\{ \frac{\hat{\epsilon}^2}{M_2 (\hat{\epsilon} + \frac{1}{2} M_1^2 n^3)}, \frac{\hat{\epsilon}}{M_1}, \frac{\Delta_{k_i} \hat{\epsilon}}{\sqrt{\hat{\epsilon} + \frac{1}{2} M_1^2 n^3}} \right\}. \end{aligned}$$

Since f is continuous on a closed set $\mathcal{F}$ and $\mathbf{X}^k \in \mathcal{F}$ for each k , $f(\mathbf{X}^{k_i})$ is bounded below, therefore $\lim_{i \rightarrow \infty} \Delta_{k_i} = 0$. From the update rule of Δ_k (Step 5 of Algorithm 2.3), it holds that $\Delta_j \leq \gamma_2 \Delta_{k_i}$ for the unsuccessful iterations $j = k_i + 1 \dots, k_{i+1} - 1$. Hence, we obtain $\lim_{j \rightarrow \infty} \Delta_j = 0$, regardless of the finiteness of $\mathcal{K}$.

Now, we will take a close look at the ratio r_k . From the Taylor expansion, there exists $\xi \in [0, 1]$ such that

$$f(\mathbf{X}^k - \alpha_k \mathbf{S}(\mathbf{X}^k)) = f(\mathbf{X}^k) - \alpha_k \langle \nabla f(\mathbf{X}^k) \mid \mathbf{S}(\mathbf{X}^k) \rangle + \frac{\alpha_k^2}{2} \langle \mathbf{S}(\mathbf{X}^k) \mid \nabla^2 f(\mathbf{X}^k - \xi \alpha_k \mathbf{S}(\mathbf{X}^k)) \mid \mathbf{S}(\mathbf{X}^k) \rangle.$$

Therefore,

$$\begin{aligned} |f(\overline{\mathbf{X}}^k) - q(\alpha^k, \mathbf{X}^k)| &\leq \frac{\alpha_k^2}{2} \left| \langle \mathbf{S}(\mathbf{X}^k) \mid \nabla^2 f(\mathbf{X}^k - \xi \alpha_k \mathbf{S}(\mathbf{X}^k)) \mid \mathbf{S}(\mathbf{X}^k) \rangle - \langle \mathbf{S}(\mathbf{X}^k) \mid \nabla^2 f(\mathbf{X}^k) \mid \mathbf{S}(\mathbf{X}^k) \rangle \right| \\ &\leq \frac{\Delta_k^2}{2} (M_2 + M_2) = \Delta_k^2 M_2. \end{aligned}$$

On the other hand, from (11), $N(\mathbf{X}^k) > \hat{\epsilon}$ and $\lim_{k \rightarrow \infty} \Delta_k = 0$, it holds for large value k that

$$f(\mathbf{X}^k) - q(\alpha_k, \mathbf{X}^k) \geq \frac{1}{2} \frac{\Delta_k \hat{\epsilon}}{\sqrt{\hat{\epsilon} + \frac{1}{2} M_1^2 n^3}} > 0.$$

Consequently, the ratio r_k can be evaluated by

$$|r_k - 1| = \frac{|f(\overline{\mathbf{X}}^k) - q(\alpha_k, \mathbf{X}^k)|}{|f(\mathbf{X}^k) - q(\alpha_k, \mathbf{X}^k)|} \leq \frac{\Delta_k^2 M_2}{\frac{1}{2} \frac{\Delta_k \hat{\epsilon}}{\sqrt{\hat{\epsilon} + \frac{1}{2} M_1^2 n^3}}} = \Delta_k \frac{2M_2 \sqrt{\hat{\epsilon} + \frac{1}{2} M_1^2 n^3}}{\hat{\epsilon}}.$$

Therefore, $\lim_{k \rightarrow \infty} \Delta_k = 0$ leads to $\lim_{k \rightarrow \infty} r_k = 1$. From the update rule of Δ_k , we have $\Delta_{k+1} \geq \Delta_k$ for large value k . Thus, there exists k_0 such that $\Delta_k \geq \Delta_{k_0}$ for $\forall k \geq k_0$, but this contradicts $\lim_{k \rightarrow \infty} \Delta_k = 0$. Hence, $\liminf_{k \rightarrow \infty} N(\mathbf{X}^k) = 0$. $\square$

3.2 Convergence of the whole sequence

Using the convergence of the subsequence, we will show in Theorem 3.6 that the whole sequence of $\{N(\mathbf{X}^k)\}$ converges to zero. We will use the following two lemmas to prove Theorem 3.6.

Lemma 3.4. For $\mathbf{X} \in \mathcal{F}$ and $\mathbf{A}, \mathbf{B} \in \mathbb{S}^n$, we have

$$\begin{aligned} |\langle \nabla f(\mathbf{X}) \mid \mathbf{A} \rangle| &\leq \sqrt{n} M_1 \|\mathbf{A}\|_F \\ |\langle \mathbf{A} \mid \nabla^2 f(\mathbf{X}) \mid \mathbf{B} \rangle| &\leq 3M_2 \|\mathbf{A}\|_F \|\mathbf{B}\|_F \end{aligned}$$

Proof: The first inequality holds, since $|\langle \nabla f(\mathbf{X}) \mid \mathbf{A} \rangle| \leq \|\nabla f(\mathbf{X})\|_F \|\mathbf{A}\|_F$ for $\forall \mathbf{A} \in \mathbb{S}^n$ and we employ the inequality $\|\mathbf{A}\|_F \leq \sqrt{n} \|\mathbf{A}\|_2$ from [18, (1.2.27)].

For the second inequality, we start with the the following inequality derived from the definition of M_2 :

$$|\langle \mathbf{D} \mid \nabla^2 f(\mathbf{X}) \mid \mathbf{D} \rangle| \leq M_2 \|\mathbf{D}\|_F^2 \quad \text{for } \forall \mathbf{D} \in \mathbb{S}^n.$$

Therefore, we get $|\langle \mathbf{A} \mid \nabla^2 f(\mathbf{X}) \mid \mathbf{A} \rangle| \leq M_2 \|\mathbf{A}\|_F^2$ and $|\langle \mathbf{B} \mid \nabla^2 f(\mathbf{X}) \mid \mathbf{B} \rangle| \leq M_2 \|\mathbf{B}\|_F^2$. Furthermore, we put $\mathbf{A} - t\mathbf{B}$ into $\mathbf{D}$ to obtain the following inequality, which holds for any $t \in \mathbb{R}$;

$$|\langle \mathbf{A} - t\mathbf{B} \mid \nabla^2 f(\mathbf{X}) \mid \mathbf{A} - t\mathbf{B} \rangle| \leq M_2 \|\mathbf{A} - t\mathbf{B}\|_F^2.$$

Therefore, the inequality

$$(M_2\|\mathbf{B}\|_F^2 - \langle \mathbf{B} \mid \nabla^2 f(\mathbf{X}) \mid \mathbf{B} \rangle)t^2 - 2(M_2\langle \mathbf{A} \mid \mathbf{B} \rangle - \langle \mathbf{A} \mid \nabla^2 f(\mathbf{X}) \mid \mathbf{B} \rangle)t + (M_2\|\mathbf{A}\|_F^2 - \langle \mathbf{A} \mid \nabla^2 f(\mathbf{X}) \mid \mathbf{A} \rangle) \geq 0$$

holds for $\forall t \in \mathbb{R}$, and we can derive

$$\begin{aligned} & (M_2\langle \mathbf{A} \mid \mathbf{B} \rangle - \langle \mathbf{A} \mid \nabla^2 f(\mathbf{X}) \mid \mathbf{B} \rangle)^2 \\ & \leq (M_2\|\mathbf{A}\|_F^2 - \langle \mathbf{A} \mid \nabla^2 f(\mathbf{X}) \mid \mathbf{A} \rangle) (M_2\|\mathbf{B}\|_F^2 - \langle \mathbf{B} \mid \nabla^2 f(\mathbf{X}) \mid \mathbf{B} \rangle) \\ & \leq (2M_2\|\mathbf{A}\|_F^2)(2M_2\|\mathbf{B}\|_F^2). \end{aligned}$$

Consequently, it holds that

$$\begin{aligned} \langle \mathbf{A} \mid \nabla^2 f(\mathbf{X}) \mid \mathbf{B} \rangle & \leq M_2\langle \mathbf{A} \mid \mathbf{B} \rangle + \sqrt{(2M_2\|\mathbf{A}\|_F^2)(2M_2\|\mathbf{B}\|_F^2)} \\ & \leq M_2\|\mathbf{A}\|_F\|\mathbf{B}\|_F + 2M_2\|\mathbf{A}\|_F\|\mathbf{B}\|_F = 3M_2\|\mathbf{A}\|_F\|\mathbf{B}\|_F \end{aligned}$$

We replace $\mathbf{A}$ with $-\mathbf{A}$ to obtain

$$\langle -\mathbf{A} \mid \nabla^2 f(\mathbf{X}) \mid \mathbf{B} \rangle \leq 3M_2\|\mathbf{A}\|_F\|\mathbf{B}\|_F.$$

By combining these inequalities, we get $|\langle \mathbf{A} \mid \nabla^2 f(\mathbf{X}) \mid \mathbf{B} \rangle| \leq 3M_2\|\mathbf{A}\|_F\|\mathbf{B}\|_F$. □

Lemma 3.5. For $\mathbf{X}^k \in \mathcal{F}$, it holds that $\underline{f}(\mathbf{X}^k) \geq -n\sqrt{N(\mathbf{X}^k)}$.

Proof: In this proof, we use $\mathbf{P}_+ := \mathbf{P}_+(\mathbf{X}^k)$, $\mathbf{P}_- := \mathbf{P}_-(\mathbf{X}^k)$, $\Lambda_+ := \Lambda_+(\mathbf{X}^k)$, $\Lambda_- := \Lambda_-(\mathbf{X}^k)$, $\mathbf{V}_+ := \mathbf{V}_+(\mathbf{X}^k)$, $\mathbf{V}_- := \mathbf{V}_-(\mathbf{X}^k)$, $n_+ := n_+(\mathbf{X}^k)$, $n_- := n_-(\mathbf{X}^k)$. Since all of the eigenvalues of $\mathbf{V}_+$ and $\mathbf{V}_-$ are between 0 and 1, we have $\|\mathbf{V}_+^{1/4}\|_F \leq \sqrt{n_+}$ and $\|\mathbf{V}_-^{1/4}\|_F \leq \sqrt{n_-}$.

The objective function of (4) at $\mathbf{X} \in \mathcal{F}$ can be evaluated from below by

$$\begin{aligned} \langle \nabla f(\mathbf{X}^k) \mid \mathbf{X} - \mathbf{X}^k \rangle & = \langle \mathbf{P}_+ \Lambda_+ \mathbf{P}_+^T + \mathbf{P}_- \Lambda_- \mathbf{P}_-^T \mid \mathbf{X} - \mathbf{X}^k \rangle \\ & = \langle \Lambda_+ \mid \mathbf{P}_+^T \mathbf{X} \mathbf{P}_+ \rangle - \langle \Lambda_- \mid \mathbf{P}_-^T (\mathbf{I} - \mathbf{X}) \mathbf{P}_- \rangle - \langle \Lambda_+ \mid \mathbf{V}_+ \rangle + \langle \Lambda_- \mid \mathbf{V}_- \rangle \\ & \geq -\langle \Lambda_+ \mid \mathbf{V}_+ \rangle + \langle \Lambda_- \mid \mathbf{V}_- \rangle. \end{aligned}$$

We can obtain an upper bound on $\langle \Lambda_+ \mid \mathbf{V}_+ \rangle$;

$$\begin{aligned} \langle \Lambda_+ \mid \mathbf{V}_+ \rangle & = \text{Trace}(\mathbf{V}_+^{1/4} \Lambda_+ \mathbf{V}_+^{1/4}) \\ & \leq \|\mathbf{V}_+^{1/4}\|_F \|\mathbf{V}_+^{1/4} \Lambda_+ \mathbf{V}_+^{1/4}\|_F \leq n_+ \|\mathbf{V}_+^{1/4} \Lambda_+ \mathbf{V}_+^{1/4}\|_F. \end{aligned}$$

In a similar way, it also holds $\langle -\Lambda_- \mid \mathbf{V}_- \rangle \leq n_- \|\mathbf{V}_-^{1/4} \Lambda_- \mathbf{V}_-^{1/4}\|_F$.

Consequently, we obtain

$$\begin{aligned} \underline{f}(\mathbf{X}^k) & \geq -n_+ \|\mathbf{V}_+^{1/4} \Lambda_+ \mathbf{V}_+^{1/4}\|_F - n_- \|\mathbf{V}_-^{1/4} \Lambda_- \mathbf{V}_-^{1/4}\|_F \\ & \geq -(n_+ + n_-) \sqrt{\|\mathbf{V}_+^{1/4} \Lambda_+ \mathbf{V}_+^{1/4}\|_F^2 + \|\mathbf{V}_-^{1/4} \Lambda_- \mathbf{V}_-^{1/4}\|_F^2} \\ & = -n\sqrt{N(\mathbf{X}^k)}. \end{aligned}$$

□

We are now ready to prove the convergence of the whole sequence.

Theorem 3.6. When $\{\mathbf{X}^k\}$ is the sequence generated by Algorithm 2.3 with $\epsilon = 0$, it holds that

$$\lim_{k \rightarrow \infty} N(\mathbf{X}^k) = 0.$$

Proof:

We take a small positive number ϵ_1 such that $0 < \epsilon_1 \leq 16n^2M_1^2$. We assume that there is an infinite subsequence $\mathcal{K} := \{k_1, k_2, \dots, k_i, \dots\} \subset \{1, 2, \dots\}$ such that $N(\mathbf{X}^{k_i}) \geq \epsilon_1$ for $\forall k_i \in \mathcal{K}$, and we will derive a contradiction.

From Theorem 3.3, we can take a subsequence $\mathcal{L} := \{l_1, l_2, \dots, l_i, \dots\} \subset \{1, 2, \dots\}$ such that

$$\begin{cases} N(\mathbf{X}^k) & \geq & \epsilon_2^2 \\ N(\mathbf{X}^{l_i}) & < & \epsilon_2^2. \end{cases} \quad \text{for } k = k_i, k_i + 1, \dots, l_i - 1$$

where $\epsilon_2 := \frac{\epsilon_1}{4nM_1}$. Note that this is consistent with $N(\mathbf{X}^{k_i}) \geq \epsilon_1$, since we took $0 < \epsilon_1 \leq 16M_1^2$.

If the k th iteration is a successful iteration and $k_i \leq k < l_i$, we put $N(\mathbf{X}^k) \geq \epsilon_2^2$ into (12) and obtain

$$f(\mathbf{X}^{k+1}) \leq f(\mathbf{X}^k) - \frac{\mu_1}{2} \min \left\{ \frac{\epsilon_2^4}{M_2(\epsilon_2^2 + \frac{1}{2}M_1^2n^3)}, \frac{\epsilon_2^2}{M_1}, \frac{\Delta_k \epsilon_2^2}{\sqrt{N\epsilon_2^2 + \frac{1}{2}M_1^2n^3}} \right\}.$$

Since f is bounded below, when the k th iteration ($k_i \leq k < l_i$) is a successful iteration and k is large enough, it holds that

$$f(\mathbf{X}^{k+1}) \leq f(\mathbf{X}^k) - \Delta_k \epsilon_3$$

where $\epsilon_3 := \frac{\mu_1}{2} \frac{\epsilon_2^2}{\sqrt{N\epsilon_2^2 + \frac{1}{2}M_1^2n^3}}$. Since we update the matrix by $\mathbf{X}^{k+1} = \mathbf{X}^k - \alpha_k \mathbf{S}(\mathbf{X}^k)$ in a successful iteration, we use $\alpha_k \leq \Delta_k$ and $\|\mathbf{S}(\mathbf{X}^k)\|_F = 1$ to derive

$$\|\mathbf{X}^k - \mathbf{X}^{k+1}\|_F \leq \Delta_k \leq \frac{f(\mathbf{X}^k) - f(\mathbf{X}^{k+1})}{\epsilon_3}.$$

This inequality is also valid when the k th iteration is an unsuccessful iteration, since the matrix is kept by $\mathbf{X}^{k+1} = \mathbf{X}^k$. Hence, it holds that

$$\begin{aligned} & \|\mathbf{X}^{k_i} - \mathbf{X}^{l_i}\|_F \\ & \leq \|\mathbf{X}^{k_i} - \mathbf{X}^{k_i+1}\|_F + \|\mathbf{X}^{k_i+1} - \mathbf{X}^{k_i+2}\|_F \dots + \|\mathbf{X}^{l_i-1} - \mathbf{X}^{l_i}\|_F \\ & \leq \frac{1}{\epsilon_3} \left((f(\mathbf{X}^{k_i}) - f(\mathbf{X}^{k_i+1})) + (f(\mathbf{X}^{k_i+1}) - f(\mathbf{X}^{k_i+2})) + \dots + (f(\mathbf{X}^{l_i-1}) - f(\mathbf{X}^{l_i})) \right) \\ & = \frac{f(\mathbf{X}^{k_i}) - f(\mathbf{X}^{l_i})}{\epsilon_3}. \end{aligned}$$

Since the objective function $f(\mathbf{X}^k)$ is non-increasing and bounded below, this implies that $\lim_{i \rightarrow \infty} \|\mathbf{X}^{k_i} - \mathbf{X}^{l_i}\|_F = 0$. Therefore, for $\epsilon_4 := \frac{\sqrt{n}\epsilon_2}{M_1+3M_2} > 0$, there exists i_0 such that $\|\mathbf{X}^{k_i} - \mathbf{X}^{l_i}\|_F < \epsilon_4$ for $\forall i \geq i_0$.

Since $-\mathbf{I} \preceq \mathbf{X} - \mathbf{X}^{k_i} \preceq \mathbf{I}$ for $\mathbf{X} \in \mathcal{F}$, we have an inequality $\|\mathbf{X} - \mathbf{X}^{k_i}\|_F \leq \sqrt{n}$. From Lemma 3.4, it holds for $\mathbf{X} \in \mathcal{F}$ and $i \geq i_0$ that

$$\left| \langle \nabla f(\mathbf{X}^{k_i}) \mid \mathbf{X} - \mathbf{X}^{k_i} \rangle - \langle \nabla f(\mathbf{X}^{l_i}) \mid \mathbf{X} - \mathbf{X}^{l_i} \rangle \right|$$

$$\begin{aligned}
&= \left| \langle \nabla f(\mathbf{X}^{l_i} + (\mathbf{X}^{k_i} - \mathbf{X}^{l_i})) \mid \mathbf{X} - \mathbf{X}^{k_i} \rangle - \langle \nabla f(\mathbf{X}^{l_i}) \mid \mathbf{X} - \mathbf{X}^{l_i} \rangle \right| \\
&= \left| \langle \nabla f(\mathbf{X}^{l_i}) \mid \mathbf{X} - \mathbf{X}^{k_i} \rangle + \int_0^1 \langle \mathbf{X}^{k_i} - \mathbf{X}^{l_i} \mid \nabla^2 f(\mathbf{X}^{l_i} + \xi(\mathbf{X}^{k_i} - \mathbf{X}^{l_i})) \mid \mathbf{X} - \mathbf{X}^{k_i} \rangle d\xi \right. \\
&\quad \left. - \langle \nabla f(\mathbf{X}^{l_i}) \mid \mathbf{X} - \mathbf{X}^{l_i} \rangle \right| \\
&= \left| \int_0^1 \langle \mathbf{X}^{k_i} - \mathbf{X}^{l_i} \mid \nabla^2 f(\mathbf{X}^{l_i} + \xi(\mathbf{X}^{k_i} - \mathbf{X}^{l_i})) \mid \mathbf{X} - \mathbf{X}^{k_i} \rangle d\xi - \langle \nabla f(\mathbf{X}^{l_i}) \mid \mathbf{X}^{k_i} - \mathbf{X}^{l_i} \rangle \right| \\
&\leq \int_0^1 \left| \langle \mathbf{X}^{k_i} - \mathbf{X}^{l_i} \mid \nabla^2 f(\mathbf{X}^{l_i} + \xi(\mathbf{X}^{k_i} - \mathbf{X}^{l_i})) \mid \mathbf{X} - \mathbf{X}^{k_i} \rangle \right| d\xi + \left| \langle \nabla f(\mathbf{X}^{l_i}) \mid \mathbf{X}^{k_i} - \mathbf{X}^{l_i} \rangle \right| \\
&\leq 3M_2 \|\mathbf{X}^{k_i} - \mathbf{X}^{l_i}\|_F \|\mathbf{X} - \mathbf{X}^{k_i}\|_F + \sqrt{n}M_1 \|\mathbf{X}^{k_i} - \mathbf{X}^{l_i}\|_F \\
&\leq 3M_2 \epsilon_4 \sqrt{n} + \sqrt{n}M_1 \epsilon_4 = \sqrt{n}(3M_2 + M_1)\epsilon_4 = n\epsilon_2.
\end{aligned}$$

Hence, we have

$$\langle \nabla f(\mathbf{X}^{k_i}) \mid \mathbf{X} - \mathbf{X}^{k_i} \rangle \geq \langle \nabla f(\mathbf{X}^{l_i}) \mid \mathbf{X} - \mathbf{X}^{l_i} \rangle - n\epsilon_2. \quad (13)$$

From the assumption $N(\mathbf{X}^{k_i}) \geq \epsilon_2^2$ we know that $\lambda_{\max}(\mathbf{X}^{k_i}) > 0$ (If $\lambda_{\max}(\mathbf{X}^{k_i}) = 0$, then $N(\mathbf{X}^{k_i}) = 0$ from Lemma 2.1, and this means $\mathbf{X}^{k_i}$ already satisfies the first-order optimality condition). Since $\mathbf{X}^{k_i} - \frac{\mathbf{D}(\mathbf{X}^{k_i})}{\lambda_{\max}(\mathbf{X}^{k_i})} \in \mathcal{F}$ from Lemma 2.2, we can put $\mathbf{X}^{k_i} - \frac{\mathbf{D}(\mathbf{X}^{k_i})}{\lambda_{\max}(\mathbf{X}^{k_i})}$ into (13) to get

$$\langle \nabla f(\mathbf{X}^{k_i}) \mid -\frac{\mathbf{D}(\mathbf{X}^{k_i})}{\lambda_{\max}(\mathbf{X}^{k_i})} \rangle \geq \langle \nabla f(\mathbf{X}^{l_i}) \mid \left(\mathbf{X}^{k_i} - \frac{\mathbf{D}(\mathbf{X}^{k_i})}{\lambda_{\max}(\mathbf{X}^{k_i})} \right) - \mathbf{X}^{l_i} \rangle - n\epsilon_2 \geq \underline{f}(\mathbf{X}^{l_i}) - n\epsilon_2.$$

With Lemma 3.5 and $N(\mathbf{X}^{l_i}) < \epsilon_2^2$, we have an upper bound on $N(\mathbf{X}^{k_i})$;

$$\begin{aligned}
N(\mathbf{X}^{k_i}) &= \langle \nabla f(\mathbf{X}^{k_i}) \mid \mathbf{D}(\mathbf{X}^{k_i}) \rangle \leq \lambda_{\max}(\mathbf{X}^{k_i})(-\underline{f}(\mathbf{X}^{l_i}) + n\epsilon_2) \\
&\leq \lambda_{\max}(\mathbf{X}^{k_i})(n\sqrt{N(\mathbf{X}^{l_i})} + n\epsilon_2) \leq M_1(n\epsilon_2 + n\epsilon_2) = 2M_1n\epsilon_2.
\end{aligned}$$

Now, we reach the contradiction;

$$\epsilon_1 \leq N(\mathbf{X}^{k_i}) \leq 2M_1n\epsilon_2 = \frac{1}{2}\epsilon_1 < \epsilon_1.$$

Hence, $\lim_{k \rightarrow \infty} N(\mathbf{X}^k) = 0$. □

Combining Lemma 3.5 and Theorem 3.6, we derive a property for the first-order optimality condition.

Theorem 3.7. *When $\{\mathbf{X}^k\}$ is the sequence generated by Algorithm 2.3 with $\epsilon = 0$, it holds that*

$$\lim_{k \rightarrow \infty} \underline{f}(\mathbf{X}^k) = 0.$$

Proof: From Lemma 3.5 and the definition of $\underline{f}(\mathbf{X})$, we know that $-n\sqrt{N(\mathbf{X}^k)} \leq \underline{f}(\mathbf{X}^k) \leq 0$. Hence, Theorem 3.6 indicates $\lim_{k \rightarrow \infty} \underline{f}(\mathbf{X}^k) = 0$. □

To make the generated sequence $\{\mathbf{X}^k\}$ itself converge, we need a stronger assumption on the objective function, for example, strong convexity.

Corollary 3.8. *If the objective function f is strongly convex, that is, there exists $\nu > 0$ such that*

$$f(\mathbf{Y}) \geq f(\mathbf{X}) + \langle \nabla f(\mathbf{X}) \mid \mathbf{Y} - \mathbf{X} \rangle + \frac{\nu}{2} \|\mathbf{Y} - \mathbf{X}\|_F^2 \quad \text{for } \forall \mathbf{X}, \forall \mathbf{Y} \in \mathcal{F},$$

then the sequence $\{\mathbf{X}^k\}$ generated by Algorithm 2.3 with $\epsilon = 0$ converges. Furthermore, the accumulation point $\mathbf{X}^ := \lim_{k \rightarrow \infty} \mathbf{X}^k$ satisfies the first-order optimality condition (3).*

Proof:

From $\mathbf{X}^k \in \mathcal{F}$ and the definition of $\underline{f}(\mathbf{X}^j)$, we have an inequality $\underline{f}(\mathbf{X}^j) \leq \langle \nabla f(\mathbf{X}^j) \mid \mathbf{X}^k - \mathbf{X}^j \rangle$. By swapping $\mathbf{X}^k$ and $\mathbf{X}^j$, we also obtain the inequality $\underline{f}(\mathbf{X}^k) \leq \langle \nabla f(\mathbf{X}^k) \mid \mathbf{X}^j - \mathbf{X}^k \rangle$. The addition of these two inequalities results in

$$\langle \nabla f(\mathbf{X}^k) - \nabla f(\mathbf{X}^j) \mid \mathbf{X}^k - \mathbf{X}^j \rangle \leq -\underline{f}(\mathbf{X}^k) - \underline{f}(\mathbf{X}^j).$$

Theorem 2.1.9 of [16] gives equivalent conditions of strongly-convexity, and one of them is

$$\langle \nabla f(\mathbf{Y}) - \nabla f(\mathbf{X}) \mid \mathbf{Y} - \mathbf{X} \rangle \geq \nu \|\mathbf{Y} - \mathbf{X}\|_F^2 \quad \forall \mathbf{X}, \forall \mathbf{Y} \in \mathcal{F}.$$

Due to this inequality, we get

$$\|\mathbf{X}^k - \mathbf{X}^j\|_F \leq \frac{1}{\nu} \sqrt{-\underline{f}(\mathbf{X}^k) - \underline{f}(\mathbf{X}^j)}.$$

Theorem 3.7 implies that the sequence $\{\mathbf{X}^k\}$ is a Cauchy sequence. Since $\{\mathbf{X}^k\}$ is generated in the closed and bounded set $\mathcal{F}$, it converges to a point of $\mathcal{F}$. Consequently, the accumulation point $\mathbf{X}^* = \lim_{k \rightarrow \infty} \mathbf{X}^k$ satisfies the first-order optimality condition. $\square$

4 Numerical Results

To evaluate the performance of the proposed trust-region method, we conducted preliminary numerical experiments. We used Matlab R2014a and the computing environment was Debian Linux run on AMD Opteron Processor 4386 (3 GHz) and 128 GB of memory space.

The test functions used in the experiments are listed below and they are classified into the two groups. The functions of Group I were selected from the test functions in [21]. We added new functions as Group II. In particular, Function 5 and 6 are an extension of Generalized Rosenbrock function [15] and its variant with cosine functions.

Group I: Function 1. $f(\mathbf{X}) = -2\langle \mathbf{C}_1, \mathbf{X} \rangle + \langle \mathbf{X}, \mathbf{X} \rangle$;

Function 2. $f(\mathbf{X}) = 3\cos(\langle \mathbf{X}, \mathbf{X} \rangle) + \sin(\langle \mathbf{X} + \mathbf{C}_1, \mathbf{X} + \mathbf{C}_1 \rangle)$;

Function 3. $f(\mathbf{X}) = \log(\langle \mathbf{X}, \mathbf{X} \rangle + 1) + 5\langle \mathbf{C}_1, \mathbf{X} \rangle$;

Group II: Function 4. $f(\mathbf{X}) = \frac{\langle \mathbf{X}, \mathbf{X} \rangle^3}{n^3}$;

Function 5. $f(\mathbf{X}) = 1 + \sum_{i=1}^n \sum_{j=i}^n (A_{ij} - X_{ij})^2$
 $+ 100 \sum_{i=1}^{n-1} \sum_{j=i}^{n-1} \left(\frac{A_{ij}^2}{A_{i,j+1}} X_{i,j+1} - X_{ij}^2 \right)^2$
 $+ 100 \sum_{i=1}^{n-1} \left(\frac{A_{in}^2}{A_{i+1,i+1}} X_{i+1,i+1} - X_{in}^2 \right)^2$;

Function 6. $f(\mathbf{X}) = \frac{1}{n^2} \sum_{i=1}^n \left(\sum_{j=1, j \neq i}^n \frac{X_{ij}}{A_{ij}} - (n-1) \frac{X_{ii}^2}{A_{ii}^2} \right)^2$
 $- \frac{1}{n^2} \sum_{i=1}^n \sum_{j=1}^n \cos((X_{ij} - A_{ij})^2)$;

Function 7. $f(\mathbf{X}) = \langle \mathbf{C}_1, \mathbf{X} \rangle - \log \det(\mathbf{X} + \bar{\epsilon} \mathbf{I}) - \log \det((1 + \bar{\epsilon}) \mathbf{I} - \mathbf{X})$;

To generate the matrix $\mathbf{C}_1$, we chose the eigenvalues $\kappa_1, \dots, \kappa_n$ randomly from the interval $[-1, 2]$ and multiply a randomly-generated orthogonal matrix $\mathbf{Q}$, that is, $\mathbf{C}_1 := \mathbf{Q} \text{diag}(\kappa_1, \dots, \kappa_n) \mathbf{Q}^T$. The matrix $\mathbf{A}$ appeared in Functions 5 and 6 was set by $A_{ii} = \frac{1}{2}$ for $i = 1, \dots, n$ and $A_{ij} = \frac{1}{2(n-1)}$ if $i \neq j$. The parameter $\bar{\epsilon}$ of Function 7 determines the effect of the boundary of the feasible set $\mathcal{F}$, and we used $\bar{\epsilon} = 0.02$.

We compared the performance of three methods, TR (Algorithm 2.3), FEAS (the feasible direction method of Xu *et al.* [21]), and PEN (the penalty barrier method [2, 12] implemented in PENLAB [7]). We started TR and FEAS with the initial point $\mathbf{X}^0 = \frac{1}{2} \mathbf{I}$, while PEN automatically chose its initial point. We used the following condition as the stopping criterion;

$$\begin{aligned} \text{TR} \quad & N(\mathbf{X}^k) < 10^{-7} \text{ or } \frac{|f(\mathbf{X}^k) - f(\mathbf{X}^{k-1})|}{\max\{|f(\mathbf{X}^k)|, 1\}} < 10^{-6} \\ \text{FEAS} \quad & |\text{Trace}(\Lambda_-(\mathbf{X}^k)) - \langle f(\mathbf{X}^k) | \mathbf{X}^k \rangle| < 10^{-6} \text{ or } \frac{|f(\mathbf{X}^k) - f(\mathbf{X}^{k-1})|}{\max\{|f(\mathbf{X}^k)|, 1\}} < 10^{-6} \\ \text{PEN} \quad & \text{the default parameter of PENLAB.} \end{aligned}$$

For details of the stopping conditions on FEAS and PENLAB, refer to [21] and [7], respectively.

Tables 1 and 2 report the numerical results of Group I and Group II, respectively. The first column is the function type, and the second column n is the size of the matrix $\mathbf{X}$. The third column indicates the method we applied. The fourth column is the objective value. The fifth column is the number of main iterations, and the six column is the computation time in seconds. The last three columns correspond to the number of the evaluation of the function value $f(\mathbf{X})$, the gradient matrix $\nabla f(\mathbf{X})$, and the Hessian map $\nabla^2 f(\mathbf{X})$.

From these tables, PEN was much slow compared to TR and FEAS, and it is difficult for PEN to solve large problems with $n > 100$ in 24 hours. PENLAB [7] handled the symmetric matrix $\mathbf{X}$ as $n(n+1)/2$ independent variables, $X_{11}, X_{12}, \dots, X_{1n}, X_{22}, \dots, X_{2n}, \dots, X_{nn}$, and it stored all the elements of the Hessian map $\nabla^2 f(\mathbf{X})$, so that the computation cost estimated from [12] was $\mathcal{O}(n^4)$. This heavy cost restricted PENLAB to the small sizes. TR also used the information of the Hessian map, but TR stored only the scalar value $\langle \mathbf{S} | \nabla^2 f(\mathbf{X}) | \mathbf{S} \rangle$. hence, the memory space required by TR was only $\mathcal{O}(n^2)$.

In the comparison between TR and FEAS, the computation time of FEAS was shorter than TR in Table 1, but longer in Table 2. The functions in Group I involved the variable matrix $\mathbf{X}$ in the linear form $\langle \mathbf{C}_1 | \mathbf{X} \rangle$ or the quadratic form $\langle \mathbf{X} | \mathbf{X} \rangle$, and this simple structure was favorable for the feasible direction that was based on a steepest descent direction. In contrast, the functions in Group II have higher nonlinearity than Group I. The operation count with respect to the function value (co.f) shows that this higher nonlinearity demanded that FEAS have a large number of back-step loop. In particular, FEAS needed many iterations for Ronsenbrock-type functions (Functions 5 and 6). TR reduced the number of iterations by the properties of the search direction $\mathbf{D}(\mathbf{X})$ and the quadratic approximation with the Hessian map. In particular, $\mathbf{D}(\mathbf{X})$ encompassed the information of the distance to the boundary to the box-constraints as $\Lambda_+(\mathbf{X})$ and $\Lambda_-(\mathbf{X})$. Consequently, TR was faster than FEAS for the functions of Group II.

5 Conclusions and Future Directions

In this paper, we proposed a trust-region method for a box-constrained nonlinear semidefinite programs. The search direction $\mathbf{D}(\mathbf{X})$ studied in Section 2 enabled us to devise a trust-region method based on Coleman and Li [4] in the space of positive semidefinite matrices. We presented the global convergence property of the generated sequence for the first-order optimality condition.

Table 1: Numerical results on Group I.

type	n	method	obj	iter	cpu	co. f	co. ∇f	co. $\nabla^2 f$
1	50	TR	-3.631×10^1	48	0.08	95	48	48
1	50	FEAS	-3.633×10^1	36	0.04	215	36	0
1	50	PEN	-3.633×10^1	22	323.70	62	31	22
1	100	TR	-7.930×10^1	67	0.30	133	67	67
1	100	FEAS	-7.932×10^1	36	0.11	239	36	0
1	100	PEN	-7.932×10^1	23	5554.30	64	32	23
1	500	TR	-3.572×10^2	81	7.70	161	81	81
1	500	FEAS	-3.574×10^2	37	2.24	250	37	0
1	1000	TR	-8.648×10^2	64	30.16	127	64	64
1	1000	FEAS	-8.651×10^2	32	9.62	204	32	0
1	5000	TR	-3.861×10^3	80	3497.86	159	80	80
1	5000	FEAS	-3.862×10^3	36	1111.89	232	36	0
1	10000	TR	-7.731×10^3	73	24730.04	145	73	73
1	10000	FEAS	-7.734×10^3	34	7782.18	213	34	0
2	50	TR	-4.000	23	0.04	45	23	23
2	50	FEAS	-4.000	31	0.04	293	31	0
2	50	PEN	-4.000	115	1808.54	1857	124	116
2	100	TR	-4.000	40	0.19	79	40	40
2	100	FEAS	-4.000	13	0.05	122	13	0
2	100	PEN	-3.985	15	4581.35	114	21	18
2	500	TR	-4.000	26	2.40	51	26	26
2	500	FEAS	-4.000	17	1.31	183	17	0
2	1000	TR	-4.000	17	6.32	33	17	17
2	1000	FEAS	-4.000	10	4.30	134	10	0
2	5000	TR	-4.000	28	1205.24	55	28	28
2	5000	FEAS	-4.000	13	449.59	161	13	0
2	10000	TR	-4.000	27	8461.01	53	27	27
2	10000	FEAS	-3.951	8	2009.73	73	8	0
3	50	TR	-3.756×10^1	201	0.35	401	201	201
3	50	FEAS	-3.756×10^1	2	0.01	3	2	0
3	50	PEN	-3.756×10^1	28	418.41	76	36	28
3	100	TR	-7.418×10^1	208	0.93	415	208	208
3	100	FEAS	-7.419×10^1	7	0.02	24	7	0
3	100	PEN	-7.419×10^1	30	7316.99	81	37	30
3	500	TR	-3.625×10^2	257	25.62	513	257	257
3	500	FEAS	-3.625×10^2	2	0.13	3	2	0
3	1000	TR	-7.739×10^2	269	128.23	537	269	269
3	1000	FEAS	-7.741×10^2	2	0.65	3	2	0
3	5000	TR	-4.129×10^3	257	11996.45	513	257	257
3	5000	FEAS	-4.129×10^3	2	74.63	3	2	0
3	10000	TR	-8.294×10^3	256	92901.29	511	256	256
3	10000	FEAS	-8.295×10^3	2	575.84	3	2	0

Table 2: Numerical results on Group II.

type	n	method	obj	iter	cpu	co. f	co. ∇f	co. $\nabla^2 f$
4	50	TR	4.036×10^{-3}	12	0.02	23	12	12
4	50	FEAS	3.882×10^{-3}	18	0.02	81	18	0
4	50	PEN	3.879×10^{-3}	117	1747.06	420	149	117
4	100	TR	1.426×10^{-2}	23	0.11	45	23	23
4	100	FEAS	1.412×10^{-2}	19	0.054	105	19	0
4	100	PEN	1.411×10^{-2}	42	10210.67	139	61	42
4	500	TR	1.086×10^{-2}	10	0.98	19	10	10
4	500	FEAS	1.039×10^{-2}	22	1.28	115	22	0
4	1000	TR	1.072×10^{-2}	9	4.24	17	9	9
4	1000	FEAS	9.943×10^{-3}	21	6.12	114	21	0
4	5000	TR	1.240×10^{-2}	8	367.77	15	8	8
4	5000	FEAS	1.039×10^{-2}	20	625.73	107	20	0
4	10000	TR	1.438×10^{-2}	8	2798.12	15	8	8
4	10000	FEAS	1.142×10^{-2}	23	5329.24	126	23	0
5	50	TR	1.122	4	0.01	7	4	4
5	50	FEAS	1.126	19	0.05	252	19	0
5	50	PEN	1.000	20	294.58	61	30	20
5	100	TR	1.117	6	0.06	11	6	6
5	100	FEAS	1.125	16	0.15	226	16	0
5	100	PEN	1.000	20	4814.74	61	30	20
5	500	TR	1.004	4	0.82	7	4	4
5	500	FEAS	1.125	16	4.28	286	16	0
5	1000	TR	1.008	4	4.02	7	4	4
5	1000	FEAS	1.125	18	26.96	352	18	0
5	5000	TR	1.002	4	192.25	7	4	4
5	5000	FEAS	1.125	90	6345.17	2279	90	0
5	10000	TR	1.013	4	1332.14	7	4	4
5	10000	FEAS	1.124	122	51611.04	3285	122	0
6	50	TR	-1.000	20	0.11	39	20	20
6	50	FEAS	-1.000	12	0.10	92	12	0
6	50	PEN	-1.000	300	4577.01	915	1218	300
6	100	TR	-1.000	20	0.358	39	20	20
6	100	FEAS	-1.000	16	0.564	150	16	0
6	100	PEN	-9.997×10^{-1}	300	73262.02	1005	1308	300
6	500	TR	-1.000	18	10.00	35	18	18
6	500	FEAS	-1.000	12	9.36	110	12	0
6	1000	TR	-1.000	4	9.42	7	4	4
6	1000	FEAS	-1.000	12	56.33	110	12	0
6	5000	TR	-1.000	4	406.01	7	4	4
6	5000	FEAS	-1.000	13	2046.55	130	13	0
6	10000	TR	-1.000	3	2076.17	5	3	3
6	10000	FEAS	-1.000	14	10416.77	130	14	0
7	50	TR	7.817×10^1	10	0.03	19	10	10
7	50	FEAS	7.817×10^1	15	0.06	108	15	0
7	50	PEN	7.817×10^1	13	195.41	38	19	13
7	100	TR	1.583×10^2	10	0.11	19	10	10
7	100	FEAS	1.583×10^2	17	0.302	13	17	0
7	100	PEN	1.583×10^2	14	3427.86	40	20	14
7	500	TR	7.825×10^2	10	2.73	19	10	10
7	500	FEAS	7.825×10^2	16	6.22	116	16	0
7	1000	TR	1.556×10^3	10	12.02	19	10	10
7	1000	FEAS	1.556×10^3	10	15.79	60	10	0
7	5000	TR	7.707×10^3	11	1708.11	21	11	11
7	5000	FEAS	7.707×10^3	16	4931.70	115	16	0
7	10000	TR	1.533×10^4	11	13379.96	21	11	11
7	10000	FEAS	1.533×10^4	14	32643.49	94	14	0

The preliminary numerical experiments in Section 4 showed that the proposed trust-region method was more favorable than the feasible direction method for functions with high nonlinearity, mainly due to the properties derived from the search direction $\mathbf{D}(\mathbf{X})$. Since our method did not hold the Hessian map in memory space, it handled the larger problems than the penalty barrier method.

The combination of the feasible direction and trust-region method will be a next step of the research direction, since the feasible direction method fits simple functions. Such a combination, however, would require a more complicated scheme to ensure the positive semidefinite condition as Lemma 2.2. Coleman and Li [4] proved the convergence for a second-order optimality condition. Since the proof in [4] required further stronger assumptions than this paper, we remain it as a matter to be discussed further.

References

- [1] M. F. Anjos and J. B. Lasserre, editors. *Handbook on Semidefinite, Cone and Polynomial Optimization: Theory, Algorithms, Software and Applications*. Springer, NY, USA, 2012.
- [2] A. Ben-tal and M. Zibulevsky. Penalty/barrier multiplier methods for convex programming problems. *SIAM J. Optim.*, 7:347–366, 1997.
- [3] S. Boyd, L. El Ghaoui, E. Feron, and V. Balakrishnan. *Linear Matrix Inequalities in System and Control Theory*. SIAM, PA, USA, 2004.
- [4] T. F. Coleman and Y. Li. An interior trust region approach for nonlinear minimization subject to bounds. *SIAM J. Optim.*, 6(2):418–445, 1996.
- [5] A. R. Conn, N. I. M. Gould, and P. L. Toint. Global convergence of a class of trust region algorithms for optimization with simple bounds. *SIAM J. Numeric. Anal.*, 25:433–460, 1988.
- [6] A. R. Conn, N. I. M. Gould, and P. L. Toint. *Trust-region methods*. SIAM, PA, USA, 2000.
- [7] J. Fiala, M. Kočvara, and M. Stingl. PENLAB - a solver for nonlinear semidefinite programming. <http://web.mat.bham.ac.uk/kocvara/penlab/>, October 2013.
- [8] M. Fukuda, B. J. Braams, M. Nakata, M. L. Overton, J. K. Percus, M. Yamashita, and Z. Zhao. Large-scale semidefinite programs in electronic structure calculation. *Math. Program. Series B*, 109(2-3):553–580, 2007.
- [9] M. X. Goemans and D. P. Williamson. Improved approximation algorithms for maximum cut and satisfiability problems using semidefinite programming. *J. Assoc. Comput. Mach.*, 42(6):1115–1145, 1995.
- [10] I. Griva, S. G. Nash, and A. Sofer. *Linear and Nonlinear Optimization*. SIAM, PA, USA, 2009.
- [11] L. Hei, J. Nocedal, and R. A. Waltz. A numerical study of active-set and interior-point methods for bound constrained optimization. In H. G. Bock, E. Kostina, H. X. Phu, and R. Rannacher, editors, *Modeling, Simulation and Optimization of Complex Processes*, pages 273–292. Springer, Berlin, Germany, 2008.
- [12] M. Kočvara and M. Stingl. PENNON: A code for convex nonlinear and semidefinite programming. *Optim. Methods Softw.*, 18(3):317–333, 2003.

- [13] J. B. Lasserre. Global optimization with polynomials and the problems of moments. *SIAM J. Optim.*, 11(3):796–817, 2001.
- [14] K. Nakata, K. Fujisawa, M. Fukuda, M. Kojima, and K. Murota. Exploiting sparsity in semidefinite programming via matrix completion II: implementation and numerical results. *Math. Program., Ser B*, 95:303–327, 2003.
- [15] S. Nash. Newton-type minimization via the lanczos method. *SIAM J. Numer. Anal.*, 21(4):770–788, 1984.
- [16] Y. E. Nesterov. *Introductory Lectures on Convex Optimization: A Basic Course*. Kluwer Academic Publishers, Dordrecht, The Netherlands, 2004.
- [17] J. Nocedal and S. J. Wright. *Numerical Optimization*. Springer, NY, USA, 2006.
- [18] W. Sun and Y. X. Yuan. *Optimization Theory and Methods Nonlinear Programming*. Springer, NY, USA, 2006.
- [19] M. J. Todd. Semidefinite optimization. *Acta Numerica*, 10:515–560, 2001.
- [20] X. Wang and Y. X. Yuan. A trust region method based on a new affine scaling technique for simple bounded optimization. *Optim. Methods and Softw.*, 28(4):871–888, 2013.
- [21] Y. Xu, W. Sun, and L. Qi. A feasible direction method for the semidefinite program with box constraints. *Appl. Math. Lett.*, 24:1874–1881, 2011.