

Factor-Based Large Covariance Matrix Estimation: Nonconvex Optimization and Statistical Guarantees

Shanshan Zou, *Student Member, IEEE* and Ziping Zhao, *Member, IEEE*

Abstract—In this paper, we study the problem of large covariance matrix estimation based on factor analysis models. In this setting, the covariance matrix is decomposed into a low-rank component and a sparse component. We formulate the estimation problem as a nonconvex optimization problem and propose an iterative algorithm to solve it. Our algorithm starts with the principal orthogonal complement thresholding (POET) estimator, a widely used optimization-free method for factor-based covariance estimation. We then refine this initial estimator through alternating optimization of the low-rank and sparse components. The resulting nonconvex estimator is referred to as POET with refining iteratively (POETRY). We theoretically prove that the algorithm converges linearly to the ground truth with an improved statistical rate compared to POET, achieving the minimax rate of convergence for factor-based covariance matrix estimation. This matches the statistical rate achieved by existing convex relaxation methods but with significantly reduced per-iteration computational complexity. We further establish the statistical rate of the inverse covariance matrix estimator. Additionally, by introducing the notion of effective rank, we are able to further reduce the sample size requirement and enhance the statistical rate. Numerical experiments validate our theoretical findings and demonstrate the superiority of the proposed algorithm compared to state-of-the-art methods.

Index Terms—Covariance matrix estimation, factor models, DOA estimation, correlated noise, low-rank plus sparse, minimax rate, nonconvex statistical optimization.

I. INTRODUCTION

Covariance matrix estimation is a fundamental problem with applications in various fields related to data analytics, including statistics [2], economics [3], [4], finance [5], [6], biology [7], signal processing [8], [9], and machine learning [10]. The sample covariance matrix (SCM) is widely used in practice due to its simplicity; however, it tends to perform poorly in high-dimensional settings where the number of variables is comparable to or larger than the sample size [11]. Recognizing the criticality of this problem, recent studies have developed diverse regularization techniques for consistent estimation of large covariance matrices. Prominent approaches exploit structural assumptions, including sparsity [4], [12], [13], low-rankness [14], [15], bandedness [16], Toeplitz [17], [18], and Hankel [19] structures.

In this paper, we focus on the estimation of large covariance matrices under the factor model assumption. Covariance matrix estimation based on factor models has wide-ranging applications, including return modeling in finance [20], [21], gene expression studies in genetics [22], spectral estimation

in signal processing [23], [24], and dimensionality reduction in machine learning [10], [25].

A substantial body of literature has emerged on estimating factor-based covariance matrices. A classical assumption in this field is that the covariance matrix follows a strict factor model [26], [27], where it is decomposed into a low-rank component plus a diagonal component [23], [28], [29]. In [30], the estimation problem was formulated as a least squares problem and solved using principal component analysis (PCA) followed by diagonal thresholding, resulting in a direct method that does not require iterative optimization. Furthermore, an alternating projection algorithm was proposed in [31], which iteratively updates both the low-rank and diagonal components of the covariance matrix.

Although simple, the covariance matrix based on the strict factor model can be inefficient because the diagonal component modeling assumes uncorrelated errors. However, this condition is often violated in practical applications, where the noise is weakly or partially correlated. For example, in finance, asset returns are often modeled as a combination of systematic factors and asset-specific residuals. The corresponding return covariance matrix can be expressed as the sum of a low-rank component, which captures the common factors, and a sparse component, which reflects the residual covariances between assets [32], [33]. It is well documented that a small set of common factors, such as macroeconomic shocks, interest rates, or the market factor itself, drives the observed dynamics of a large number of time series. Conditional on these pervasive factors, the idiosyncratic errors are influenced by specific factors, such as sectors, countries, or asset classes, that affect only a smaller number of variables. This structure results in a sparse covariance component that aligns with the frameworks for sparse covariance estimation considered in [12], [16], [34]. Similarly, in array processing, the received signals can be effectively modeled using a factor model that combines a low-dimensional subspace representation with a noise component. This low-dimensional representation arises because the signal component originates from only a few distinct sources. The corresponding covariance matrix, therefore, exhibits a low-rank plus sparse structure, where the sparse noise covariance captures the spatially correlated noise. A sparse noise covariance matrix is a common model for an array of sensors composed of several widely separated subarrays. When the interelement spacing within a subarray is small, the noise among sensors within the subarray tends to be spatially correlated. By contrast, the noise between different subarrays can often be assumed uncorrelated due to the large distances separating them. As a result, the noise covariance matrix of such an

This paper was presented in part at the 49th IEEE International Conference on Acoustics, Speech, and Signal Processing (ICASSP), Seoul, Korea, April 14–19, 2024 [1].

array has a block diagonal structure, which is indeed sparse. Even in a typical nonsparse array, the small spacing between adjacent sensors generally leads to noise coupling between neighboring sensors. Nevertheless, it is often reasonable to assume that nonadjacent sensors have spatially uncorrelated noise, leading to a banded structure that is also sparse [35], [36], [37]. To overcome the limitation of the strict factor model in modeling weakly correlated errors, researchers have turned to the approximate factor model as an alternative [32], [38], [39], [40], [41], [42], [43], [44]. In this case, the covariance matrix can be decomposed into two components: a low-rank matrix and a sparse matrix. Specifically, the true covariance matrix $\Sigma_\star \in \mathbb{R}^{d \times d}$ is expressed as $\Sigma_\star = \mathbf{L}_\star + \Psi_\star$, where $\mathbf{L}_\star$ represents the low-rank component and $\Psi_\star$ denotes the sparse component. These models utilize sparse components to represent the covariance of weakly corrected noise. This broader model subsumes the strict factor model, where the diagonal component is a special case of a sparse structure.

To obtain a covariance estimator under the approximate factor model, one direct approach is the principal orthogonal complement thresholding (POET) estimator proposed by [41], which combines PCA for recovering the low-rank component with a thresholding operation for recovering the sparse component. In addition to POET, researchers have also explored convex optimization techniques based on penalized least squares. These methods leverage the nuclear norm for recovering the low-rank component and the ℓ_1 norm for recovering the sparse component. The error analysis in [45] reveals an estimation error bound in terms of the squared Frobenius norm as $O_p\left(\frac{rd}{n} + \frac{s \log d}{n} + \frac{s}{d^2}\right)$, where n is the number of samples, r corresponds to the rank of $\mathbf{L}_\star$, and s represents the number of nonzero elements of $\Psi_\star$. However, this error bound encompasses an additional error component $\frac{s}{d^2}$, indicating that exact recovery is unattainable even in noiseless conditions due to an additional error term. In [39], [42], [46], a tighter rate was obtained as $O_p\left(\frac{rd}{n} + \frac{s \log d}{n}\right)$, which matches the minimax optimal rate and demonstrates exact recovery. Despite this, the computational complexity of the convex relaxation method is notable, with a spectral decomposition being required to compute the proximal mapping of the nuclear norm at each iteration, resulting in a high per-iteration time complexity of $O(d^3)$. This complexity becomes particularly challenging when dealing with high-dimensional problems where d is large. Additionally, there is another line of work on factor-based covariance matrix estimation that focuses on likelihood-based loss functions [47], [48], [49], [50], [51], [52], although this topic is beyond the scope of this paper.

In this paper, we focus on estimating large covariance matrices under the assumption of the approximate factor model. We formulate the estimation problem as a nonconvex optimization problem and propose an iterative algorithm to compute the estimator. The algorithm begins with the POET estimator and then updates the low-rank and sparse components in an iterative manner. After certain iterations, we obtain an improved estimator of the covariance matrix. We name the finally obtained nonconvex estimator POET with refining iteratively (POETRY). Our theoretical analysis demonstrates that the

sequence generated by the algorithm converges linearly to the true covariance matrix $\Sigma_\star$ up to a statistical rate of $O_p\left(\frac{rd}{n} + \frac{s \log d}{n}\right)$, which matches the minimax rate [42], [45], [46]. Moreover, we establish the statistical rate of the inverse covariance matrix estimate based on the covariance estimator. By introducing the notion of effective rank r_e [53], we reduce the required sample size from $n \geq d$ (as in convex relaxation methods [42], [45], [46]) to $n \gtrsim r_e r \log 2d$ and enhance the statistical rate to $O_p\left(\frac{r_e r \log 2d + s \log d}{n}\right)$. Additionally, our algorithm has a lower per-iteration time complexity of $O(rd^2)$ compared to the $O(d^3)$ per-iteration complexity of convex relaxation-based methods [42], [45], [46].

The remainder of this paper is organized as follows. In Section II, we provide our problem formulation for large covariance matrix estimation based on factor models and discuss relevant existing estimators. The detailed algorithm is elaborated in Section III. Section IV contains our theoretical results, including the statistical convergence rate of the proposed estimator and the analysis of iteration complexity. In Section V, we conduct numerical experiments to compare our proposed algorithm with state-of-the-art methods on both synthetic and real-world data. The numerical experiments demonstrate the effectiveness of our algorithm over existing approaches and validate the obtained theoretical claims regarding the statistical rate of the POETRY estimator. Finally, we conclude our paper in Section VI. The proofs of all theoretical results are provided in the Appendices.

Notation. Standard lower-case and upper-case letters, such as x and X , represent scalars. Boldface lower-case and boldface upper-case letters, such as $\mathbf{x}$ and $\mathbf{X}$, denote vectors and matrices, respectively. $\mathbf{X}_{ij}$ or $(\mathbf{X})_{ij}$ refers to the (i, j) -th entry of $\mathbf{X}$. $|\mathbf{X}|_{ij}$ denotes the absolute value of $\mathbf{X}_{ij}$. $\mathbf{0}$ stands for the all-zero vector or all-zero matrix. $\mathbf{1}$ stands for the all-one vector. $\mathbf{I}$ stands for the identity matrix. $\mathbb{R}^m$ denotes the set of real $m \times 1$ vectors. $\mathbb{R}^{m \times n}$ denotes the set of real $m \times n$ matrices. Calligraphic letters, like $\mathcal{S}$, are used to represent sets.

$\text{diag}(\mathbf{X})$ represents a diagonal matrix containing the diagonal elements of $\mathbf{X}$. $\mathbf{X}^\top$ and $\mathbf{X}^{-1}$ denote the transpose and inverse of $\mathbf{X}$, respectively. We use $\sigma_i(\mathbf{X})$ to denote the i -th largest singular value of $\mathbf{X}$. In addition, $\|\mathbf{X}\|_F$, $\|\mathbf{X}\|_*$, and $\|\mathbf{X}\|_2$ denote the Frobenius norm, the nuclear norm, and the spectral norm of matrix $\mathbf{X}$, respectively. $\|\mathbf{X}\|_1$ and $\|\mathbf{X}\|_\infty$ denote the entrywise ℓ_1 norm and the entrywise ℓ_∞ norm of matrix $\mathbf{X}$, respectively. $\|\mathbf{X}\|_0$ denotes the number of nonzero elements in the matrix $\mathbf{X}$.

$\nabla f(\cdot)$ represents the gradient of a function f . $a \vee b = \max\{a, b\}$. When comparing functionals $f(n)$ and $g(n)$ for large values of n , we denote $f(n) \gtrsim g(n)$ if $f(n) \geq cg(n)$, $f(n) \lesssim g(n)$ if $f(n) \leq Cg(n)$, and $f(n) \asymp g(n)$ if $cg(n) \leq f(n) \leq Cg(n)$ for some positive constants c and C . For sequences of numbers a_m and b_m , we denote $b_m = O(a_m)$ if $\limsup_{m \rightarrow \infty} \frac{|b_m|}{a_m} < \infty$, and $b_m = O_p(a_m)$ if, for any $\varepsilon > 0$, there exists C_ε such that $P(|b_m| \leq C_\varepsilon a_m) > 1 - \varepsilon$ for all m . We use C and c , with appropriate subscripts, to denote universal positive constants whose values may differ from line to line.

II. COVARIANCE MATRIX ESTIMATION BASED ON FACTOR MODELS

The factor model describes a random vector $\mathbf{x} \in \mathbb{R}^d$ as

$$\mathbf{x} = \mathbf{B}_* \mathbf{f} + \boldsymbol{\epsilon}, \quad (1)$$

where $\mathbf{B}_* \in \mathbb{R}^{d \times r}$ with $r \ll d$ is a loading matrix, $\mathbf{f} \in \mathbb{R}^r$ represents the low-dimensional unobservable common factors, and $\boldsymbol{\epsilon} \in \mathbb{R}^d$ are idiosyncratic errors. The factors $\mathbf{f}$ and the errors $\boldsymbol{\epsilon}$ are assumed to be uncorrelated. Additionally, it is commonly assumed that both $\mathbf{f}$ and $\boldsymbol{\epsilon}$ follow normal distributions with zero mean. The factor $\mathbf{f}$ is assumed to have an identity covariance matrix, while the error $\boldsymbol{\epsilon}$ has a covariance matrix $\boldsymbol{\Psi}_*$. The covariance matrix of $\mathbf{x}$, denoted by $\boldsymbol{\Sigma}_*$, is then computed as

$$\boldsymbol{\Sigma}_* = \mathbf{B}_* \mathbf{B}_*^\top + \boldsymbol{\Psi}_* = \mathbf{L}_* + \boldsymbol{\Psi}_*, \quad (2)$$

where $\mathbf{L}_* = \mathbf{B}_* \mathbf{B}_*^\top$ is low-rank, and $\boldsymbol{\Psi}_*$ is sparse due to the weak correlation among errors in the approximate factor structure [32]. In [45], the decomposition in (2) is also referred to as factor analysis with sparse noise.

In practice, the true covariance matrix $\boldsymbol{\Sigma}_*$ is unknown. Instead, we can compute the SCM as $\mathbf{S} = \frac{1}{n} \sum_{i=1}^n \mathbf{x}_i \mathbf{x}_i^\top$, where $\mathbf{x}_i$ for $i = 1, 2, \dots, n$ represents n independent and identically distributed (i.i.d.) samples from the model (1). Due to the limited number of available samples, $\mathbf{S}$ is considered a “noisy” representation of $\boldsymbol{\Sigma}_*$ [45], [54]. This relationship is expressed as

$$\mathbf{S} = \boldsymbol{\Sigma}_* + \mathbf{N}, \quad (3)$$

where $\mathbf{N}$ denotes a re-centered Wishart noise [55] arising from the normal distribution of each $\mathbf{x}_i$. Given $\mathbf{S}$, our goal is to estimate $\boldsymbol{\Sigma}_*$ within the context of factor modeling.

A. Existing Estimators

1) *Direct Method*: We present the principal orthogonal complement thresholding (POET) estimator [41], a well-established direct method in the field of factor-based covariance matrix estimation. POET involves extracting the top r principal components of $\mathbf{S}$ and subsequently applying thresholding to the remaining components. The thresholding methods that can be employed include soft thresholding, hard thresholding, smoothly clipped absolute deviations, and adaptive thresholding [41]. Specifically, let $\phi_1 \geq \phi_2 \geq \dots \geq \phi_d$ denote the eigenvalues of $\mathbf{S}$ in descending order, and let $\{\boldsymbol{\xi}_i\}_{i=1}^d$ denote the corresponding eigenvectors. The spectral decomposition of $\mathbf{S}$ can be written as

$$\mathbf{S} = \sum_{i=1}^r \phi_i \boldsymbol{\xi}_i \boldsymbol{\xi}_i^\top + \mathbf{R},$$

where the residual component $\mathbf{R} = \sum_{i=r+1}^d \phi_i \boldsymbol{\xi}_i \boldsymbol{\xi}_i^\top$ is referred to as the “principal orthogonal complement.” We define the hard-thresholding operator, which retains only the τ largest entries in magnitude of the input matrix, as follows:

$$(\mathbf{H}_\tau(\mathbf{M}))_{ij} = \begin{cases} \mathbf{M}_{ij} & \text{if } |\mathbf{M}_{ij}| \text{ is among the largest } \tau \\ & \text{entries in magnitude of } \mathbf{M} \\ 0 & \text{otherwise.} \end{cases}$$

Algorithm 1: Algorithm for POET

Input: $\mathbf{S}$, r , τ ;
1 let $\boldsymbol{\Xi} \boldsymbol{\Phi} \boldsymbol{\Xi}^\top$ be the top- r spectral decomposition of $\mathbf{S}$;
2 $\mathbf{B}_{\text{POET}} = \boldsymbol{\Xi} \boldsymbol{\Phi}^{1/2}$;
3 $\boldsymbol{\Psi}_{\text{POET}} = \mathbf{H}_\tau(\mathbf{S} - \mathbf{B}_{\text{POET}} \mathbf{B}_{\text{POET}}^\top)$;
Output: $\boldsymbol{\Sigma}_{\text{POET}} = \mathbf{B}_{\text{POET}} \mathbf{B}_{\text{POET}}^\top + \boldsymbol{\Psi}_{\text{POET}}$.

The POET estimator is then given by

$$\boldsymbol{\Sigma}_{\text{POET}} = \sum_{i=1}^r \phi_i \boldsymbol{\xi}_i \boldsymbol{\xi}_i^\top + \mathbf{H}_\tau(\mathbf{R}). \quad (4)$$

This method is optimization-free, making it computationally efficient. A detailed algorithm for computing POET is provided in Algorithm 1.

2) *Convex Optimization Method*: Convex relaxation-based methods have been proposed for factor-based covariance matrix estimation [39], [42], [45]. This approach formulates the optimization problem as a least squares problem that incorporates the nuclear norm and the ℓ_1 norm to simultaneously promote a low-rank component and a sparse component, respectively. The corresponding problem is given by

$$\underset{\mathbf{L}, \boldsymbol{\Psi}}{\text{minimize}} \quad \frac{1}{2} \|\mathbf{L} + \boldsymbol{\Psi} - \mathbf{S}\|_F^2 + \kappa_1 \|\mathbf{L}\|_* + \kappa_2 \|\boldsymbol{\Psi}\|_1, \quad (5)$$

where $\kappa_1, \kappa_2 > 0$ are tuning parameters, and the resulting estimator is referred to as the low rank and sparse covariance (LOREC) estimator [45], [42]. The problem is convex and can be solved to global optimality. Many standard algorithms, such as the accelerated proximal gradient method [42] and alternating minimization, can be used to solve it. However, it is important to note that one limitation of the LOREC estimator is that the low-rank component $\mathbf{L}$ does not necessarily guarantee positive semidefiniteness.

B. Comparison with Robust PCA

In this section, we compare the factor-based covariance matrix estimation problem with the well-established robust PCA problem [56], [57], [58], [59]. We begin by briefly reviewing the robust PCA formulation. Given an observation matrix $\mathbf{X} \in \mathbb{R}^{d \times n}$, the robust PCA model assumes the following decomposition:

$$\mathbf{X} = \mathbf{U} \mathbf{Y} + \mathbf{E} + \mathbf{O},$$

where $\mathbf{U} \in \mathbb{R}^{d \times r}$ is a basis matrix with $r \ll d$, $\mathbf{Y} \in \mathbb{R}^{r \times n}$ contains the low-dimensional representations of $\mathbf{X}$, $\mathbf{E} \in \mathbb{R}^{d \times n}$ is a noise matrix, and $\mathbf{O} \in \mathbb{R}^{d \times n}$ is a sparse outlier matrix. The goal of robust PCA is to obtain a low-rank representation $\mathbf{U} \mathbf{Y}$ of the possibly noisy measurement $\mathbf{X}$ by effectively separating and removing the sparse outliers $\mathbf{O}$. Although both robust PCA and factor-based covariance estimation involve a low-rank plus sparse decomposition, existing methods and theoretical results for robust PCA cannot be directly applied to our setting due to several key differences, which we summarize below:

- (i) Robust PCA typically focuses on decomposing an observed data matrix directly into low-rank and sparse

components, often without a probabilistic interpretation. In contrast, the factor-based covariance estimation problem is rooted in the statistical factor model, where the decomposition is applied to the covariance matrix, motivated by a generative latent variable model.

- (ii) The low-rank component in robust PCA is generally asymmetric, while the low-rank component in covariance estimation is symmetric and positive semidefinite due to the factor modeling.
- (iii) The sparse component in robust PCA reflects a small number of corrupted entries (i.e., sparse outliers) at the data level. Two types of sparsity are typically considered: element-wise sparsity and column-wise sparsity, corresponding to different corruption models. In contrast, factor-based covariance estimation assumes sparsity in the residual noise covariance matrix, where the sparse component captures weak correlations among residual errors. This sparse component is required to be symmetric and may exhibit additional structure, such as a block-diagonal pattern.
- (iv) In robust PCA, the noise term represents the reconstruction or representation error. It can also be absent, leading to the so-called noiseless robust PCA formulation. By contrast, in factor-based covariance estimation, the noise arises from sample variability under the finite-sample regime and is always present.
- (v) In covariance estimation, it is often desirable for the combined low-rank and sparse estimator to be positive semidefinite. Robust PCA, on the other hand, does not impose such a requirement.

These distinctions lead to fundamentally different algorithmic designs, assumptions, and statistical guarantees between the two problems.

III. THE PROPOSED ESTIMATOR AND ALGORITHM

A. Proposed Estimator

In this paper, we introduce a nonconvex optimization method for estimating covariance matrices under the factor model assumption (2). In contrast to the regularization technique used for estimating the low-rank part of the covariance matrix, $\mathbf{L}$ [45], [42], we choose a nonconvex matrix factorization technique to address the computational challenges associated with rank regularization. We represent $\mathbf{L}$ as the product of smaller matrices, denoted as $\mathbf{L} = \mathbf{B}\mathbf{B}^\top$, where $\mathbf{B} \in \mathbb{R}^{d \times r}$. Furthermore, we apply a cardinality constraint to the sparse component of the covariance matrix, Ψ .

The basic idea behind our approach is to find the "best fit" between the estimated covariance $\Sigma = \mathbf{B}\mathbf{B}^\top + \Psi$ and the sample covariance $\mathbf{S}$. To achieve this, we consider the following nonconvex optimization problem:

$$\begin{aligned} & \underset{\mathbf{B}, \Psi}{\text{minimize}} && \frac{1}{2} \left\| \mathbf{B}\mathbf{B}^\top + \Psi - \mathbf{S} \right\|_F^2 \\ & \text{subject to} && \|\Psi\|_0 \leq \tau, \end{aligned} \quad (6)$$

where $\tau > 0$ is a tuning parameter that controls the number of nonzero elements in Ψ . When $\tau = 0$, the estimator simplifies to the sparse covariance matrix estimator [12]. When $\tau =$

Algorithm 2: Algorithm for POETRY

Input: $\mathbf{S}$, r , τ , η , T ;
1 let $\Xi\Phi\Xi^\top$ be the top- r spectral decomposition of $\mathbf{S}$;
2 $\mathbf{B}_0 = \Xi\Phi^{1/2}$;
3 $\Psi_0 = H_\tau(\mathbf{S} - \mathbf{B}_0\mathbf{B}_0^\top)$;
4 **for** $t = 0, 1, \dots, T-1$ **do**
5 $\mathbf{B}_{t+1} = \mathbf{B}_t - \eta \nabla_{\mathbf{B}} L(\mathbf{B}_t, \Psi_t) (\mathbf{B}_t^\top \mathbf{B}_t)^{-1}$;
6 $\Psi_{t+1} = H_\tau(\mathbf{S} - \mathbf{B}_{t+1}\mathbf{B}_{t+1}^\top)$;
7 **end**
Output: $\Sigma_T = \mathbf{B}_T\mathbf{B}_T^\top + \Psi_T$.

0, the estimator reduces to the low-rank covariance matrix estimator [60]. When either r or τ is sufficiently large, the estimator becomes the SCM. Additionally, the POET estimator presented in (4) can be viewed as a heuristic solution to (6).

B. Proposed Algorithm

We propose to solve problem (6) through an iterative optimization procedure, which is outlined in Algorithm 2. In the initialization phase, we utilize POET as the initial estimate, where we set $\mathbf{B}_0 = \mathbf{B}_{\text{POET}}$, $\Psi_0 = \Psi_{\text{POET}}$, and $\Sigma_0 = \Sigma_{\text{POET}}$ based on Algorithm 1. Subsequently, we iteratively optimize the variables $\mathbf{B}$ and Ψ over sufficient iterations, resulting in an estimator which we name POET with refining iteratively (POETRY). This acronym highlights the iterative refinement nature of the algorithm. During the iterative refining stage, we update $\mathbf{B}$ using one step of preconditioned gradient descent with a step size of η [61]. The preconditioner $(\mathbf{B}_t^\top \mathbf{B}_t)^{-1}$ is right-multiplied to the gradient

$$\nabla_{\mathbf{B}} L(\mathbf{B}, \Psi) = (\mathbf{B}\mathbf{B}^\top + \Psi - \mathbf{S}) \mathbf{B}, \quad (7)$$

where L denotes the objective function in (6). The preconditioner is employed to alleviate the effect of the potentially large condition number of the low-rank component of the covariance matrix. This adaptive preconditioner varies with each iteration and has a lower computational complexity than computing the gradient. Therefore, it does not introduce additional high-order complexity compared to standard gradient descent. To update Ψ , we employ one step of hard thresholding.

Algorithm 2 is notably efficient, as it eliminates the need for spectral decomposition in each iteration, making it substantially faster than algorithms depending on convex relaxations [45], [42], [46]. The per-iteration computational overhead of Algorithm 2 mainly arises from computing the gradient, which has a time complexity of $O(rd^2)$. This time complexity provides a significant advantage over the per-iteration complexity of $O(d^3)$ associated with the convex methods [45], [42], [46].

IV. MAIN THEORY

In this section, we present our assumptions, definitions, and main theoretical results. We begin with the necessary assumptions and definitions essential for our subsequent analysis.

A. Preliminaries

Assumption 1. $0 < \sigma_d(\Sigma_\star) \leq \sigma_1(\Sigma_\star) \leq \mu < \infty$.

This assumption requires that the eigenvalues of the true covariance matrix $\Sigma_\star$ are finite and bounded below by a positive number. It is commonly adopted in the context of factor-based covariance matrix estimation [41].

The covariance matrix estimation problem with the decomposition $\Sigma_\star = \mathbf{L}_\star + \Psi_\star$ can be ill-posed, with an identifiability issue arising when the low-rank matrix $\mathbf{L}_\star$ is also sparse [62]. Therefore, certain conditions are required to ensure that the decomposition is well-posed. For example, the well-known incoherence condition can be adopted to prevent the low-rank matrix from being excessively sparse by restricting the degree of coherence between singular vectors and the standard basis [63]. However, this condition is only applicable in noiseless models [45], making it unsuitable for our problem, which is inherently “noisy by nature.” In this paper, we introduce the spikiness condition as an alternative, which is milder than the incoherence condition and can be applied to noisy and noiseless models [45]. The spikiness condition involves constraining the spikiness ratio, defined as $\alpha := d \|\mathbf{L}_\star\|_\infty / \|\mathbf{L}_\star\|_F$ for the matrix $\mathbf{L}_\star \in \mathbb{R}^{d \times d}$ [64].

Assumption 2 (Spikiness condition [45]). *For a low-rank matrix $\mathbf{L}_\star$, we assume that there exists a constant k such that*

$$\|\mathbf{L}_\star\|_\infty = \frac{\alpha \|\mathbf{L}_\star\|_F}{d} \leq \frac{k}{d}.$$

Next, we introduce the concept of the effective rank of a matrix, which is instrumental in deriving the error bounds. The covariance matrices based on factor models have effectively reduced rank, thus we make use of the notation of the effective rank, which is first suggested by [53].

Definition 3 ([53]). The effective rank of the covariance matrix $\Sigma_\star$ is defined as

$$r_e = \frac{\text{tr}(\Sigma_\star)}{\|\Sigma_\star\|_2}.$$

The effective rank can be considered as a measure of the concentration level of the spectrum of $\Sigma_\star$. As we will show in Section IV-B, when the effective rank of the covariance matrix is much smaller than d , our theoretical results provide a tight squared Frobenius norm estimation error bound.

B. Theoretical Analysis

The following theorems characterize the theoretical properties of the sequence $\{\Sigma_t\}_{t \geq 0}$, where the composite estimator $\Sigma_t = \mathbf{B}_t \mathbf{B}_t^\top + \Psi_t$, generated by Algorithm 2. We first derive the statistical error of the POET estimator, i.e., Σ_0 , in Theorem 4, and then we derive the error bound for Σ_t when $t = 1, 2, \dots$, which encompasses both the statistical error and the optimization error, in Theorem 5. By taking into account the effective rank, we obtain an improved statistical error characterization in Theorem 7.

Theorem 4. *Suppose Assumptions 1 and 2 hold. Assume that the sparsity parameter satisfies $\tau \asymp s$. Then, with probability at least $\min\{1 - 2\exp(-C_1 d), 1 - 2\exp(-C_2 \log d)\}$, the*

initialization of Algorithm 2, which is also the POET estimator, satisfies

$$\|\Sigma_0 - \Sigma_\star\|_F^2 \leq c_{1,4} \frac{rd}{n} + c_{2,4} \frac{s \log d}{n} + c_3 \frac{rsd}{n} + c_4 rs + c_5 r.$$

Theorem 4 demonstrates that the POET estimator achieves a statistical rate of $O_p\left(\frac{rd}{n} + \frac{s \log d}{n} + \frac{rsd}{n}\right)$, in which we neglect the last two terms in the error bound of Theorem 4 since we focus on the high-dimensional case, i.e., $n \ll d$. This rate includes an additional error term, $\frac{rsd}{n}$, compared to the minimax optimal rate $O_p\left(\frac{rd}{n} + \frac{s \log d}{n}\right)$ [42], [46]. This discrepancy suggests that POET can hardly achieve a good rate of convergence under the squared Frobenius norm, as also noted by [41].

Theorem 5. *Suppose Assumptions 1 and 2 hold. Assume that the sample size satisfies $n > 48rd/\sigma_r^2(\mathbf{L}_\star)$ and the sparsity parameter satisfies $\tau \asymp s$. Suppose the step size obeys $\frac{2-\sqrt{2}}{20} < \eta < \frac{1}{2}$. Then, with probability at least $\min\{1 - 2\exp(-C_1 d), 1 - 2\exp(-C_2 \log d)\}$, the output of Algorithm 2 for $t = 1, 2, \dots$ satisfies*

$$\|\Sigma_t - \Sigma_\star\|_F^2 \leq \underbrace{c_{1,5} \frac{rd}{n} + c_{2,5} \frac{s \log d}{n}}_{\text{statistical error}} + \underbrace{c_6 \rho^T \sigma_r^2(\mathbf{L}_\star)}_{\text{optimization error}},$$

where $\rho = 2(1 - 10\eta)^2$.

Theorem 5 suggests that the estimation error of the output from Algorithm 2 comprises two components: the statistical error and the optimization error. The statistical error scales as $O_p\left(\frac{rd}{n} + \frac{s \log d}{n}\right)$, which matches the minimax rate [42], [46]. The term $\frac{s \log d}{n}$ corresponds to the statistical error of the sparse component $\Psi_\star$, while the term $\frac{rd}{n}$ represents the statistical error of the low-rank component $\mathbf{L}_\star$. Regarding the optimization error, the contraction parameter ρ depends only on the step size η . It is possible to choose an appropriate step size such that ρ remains strictly between 0 and 1, which ensures a linear convergence rate for Algorithm 2. The assumption $n > 48rd/\sigma_r^2(\mathbf{L}_\star)$ is required to ensure that the initial estimator is sufficiently close to the true covariance matrix, and similar conditions have been adopted in previous works, e.g., [45], [42], [65], [66]. We conjecture that this requirement primarily arises from employing the sample covariance matrix as the input to our estimation procedure. If a more accurate initial estimator were used (as suggested in [42]), the sample size requirement could potentially be further relaxed.

It can be shown that by choosing $T \geq O\left(\log\left(\frac{n}{rd+s \log d}\right)\right)$, the total estimation error achieves the same order as the statistical error. From Theorem 5, we also have that Σ_T satisfies the following spectral norm bound:

$$\|\Sigma_T - \Sigma_\star\|_2^2 \leq c_1 \frac{rd}{n} + c_2 \frac{s \log d}{n}. \quad (8)$$

Similar to POET in [45] and LOREC in [42], the POETRY estimator Σ_T from Algorithm 2 is not necessarily positive definite. However, the above spectral norm bound implies that Σ_T will be positive definite if $\sigma_d(\Sigma_\star)$ is larger than $\sqrt{c_1 \frac{rd}{n} + c_2 \frac{s \log d}{n}}$. Furthermore, this result asymptotically ensures the positive definiteness of Σ_T under the regularity

condition that $\sigma_d(\Sigma_*)$ is bounded away from zero. Consequently, we conclude that Σ_T will be positive definite with an overwhelming probability tending to one.

C. Statistical Error Bound for the Inverse Covariance Matrix

In this section, we provide the statistical estimation result for the inverse covariance matrix, also known as the precision matrix. This error bound is particularly useful in scenarios where the inverse covariance serves as an input to the statistical methods, including linear and quadratic discriminant analysis, Mahalanobis distance computation, weighted least squares, and Markowitz portfolio selection.

Corollary 6. *Under the conditions in Theorem 5 and assuming that $\sigma_d(\Sigma_*) \geq c_0 \sqrt{c_1 \frac{rd}{n} + c_2 \frac{s \log d}{n}}$ with $c_0 > 1$, then, with probability at least $\min\{1 - 2 \exp(-C_2 \log d), 1 - 1/(2d)\}$, the output of Algorithm 2 with $T \geq O\left(\log\left(\frac{n}{rd + s \log d}\right)\right)$ satisfies*

$$\|\Sigma_T^{-1} - \Sigma_*^{-1}\|_F^2 \leq c_{1,6} \frac{rd}{n} + c_{2,6} \frac{s \log d}{n}.$$

This result shows that the statistical rate for covariance estimation also holds for the inverse covariance estimation. The statistical error matches the minimax rate for the low-rank plus sparse inverse covariance matrix estimation problem [45], [65], [67].

D. Improved Statistical Error Bound

In real-world scenarios, the effective rank r_e of the true covariance matrix Σ_* is often much smaller than d . Then, leveraging recent progress on the asymptotic properties of the sample covariance matrix [68], we can derive a much tighter error bound which only depends on d logarithmically, as stated in the subsequent theorem.

Theorem 7. *Suppose Assumptions 1 and 2 hold. Assume that the sample size satisfies $n > 12r_e r \log 2d / \sigma_r^2(\mathbf{L}_*)$ and the sparsity parameter satisfies $\tau \asymp s$. Suppose the step size obeys $\frac{2-\sqrt{2}}{20} < \eta < \frac{1}{2}$. Then with probability at least $\min\{1 - 2 \exp(-C_2 \log d), 1 - 1/(2d)\}$, the output of Algorithm 2 for $t = 1, 2, \dots$ satisfies*

$$\|\Sigma_t - \Sigma_*\|_F^2 \leq \underbrace{c_{1,7} \frac{r_e r \log 2d}{n}}_{\text{statistical error}} + \underbrace{c_{2,7} \frac{s \log d}{n} + c_{6,\rho}^T \sigma_r^2(\mathbf{L}_*)}_{\text{optimization error}},$$

where $\rho = 2(1 - 10\eta)^2$.

The statistical error in Theorem 7 scales as $O_p\left(\frac{r_e r \log 2d + s \log d}{n}\right)$. When $r_e \ll d$, it is significantly tighter compared to the error $O_p\left(\frac{rd}{n} + \frac{s \log d}{n}\right)$ in Theorem 5. Moreover, the condition for the sample size, $n \gtrsim r_e r \log 2d$, is much less stringent. This result implies the efficiency of factor-based covariance matrix estimation when the true covariance model has a low effective rank.

V. NUMERICAL EXPERIMENTS

In this section, we present numerical experiments on both synthetic and real datasets to validate the properties of our proposed estimator, POETRY. We compare its performance with several methods, including the sample covariance matrix (SCM) estimator and three estimators for approximate factor models: POET (principal orthogonal complement thresholding) [41], LOREC (low rank and sparse covariance estimator) [42], which is based on convex relaxation, and the alternating projection method, AltProj [59], which employs spectral decomposition for low-rank matrix projection and hard thresholding for sparse matrix projection. Additionally, we compare POETRY with two estimators for strict factor models: POET-D (principal orthogonal complement thresholding diagonally) [30], [41], which thresholds the diagonal of the residual components of the SCM after extracting the first r principal components, and AltProj-D [31], [69], which uses spectral decomposition to obtain the low-rank matrix component followed by diagonal thresholding. The simulations are performed in MATLAB with a 2.9 GHz Intel Core i7 CPU.

A. Synthetic Data

We first generate a loading matrix $\mathbf{B}_* = \mathbf{V}_* \mathbf{\Lambda}_*^{1/2}$, where $\mathbf{V}_* \in \mathbb{R}^{d \times r}$ and $\mathbf{\Lambda}_* \in \mathbb{R}^{r \times r}$ is a diagonal matrix. To obtain $\mathbf{V}_*$, we generate a $d \times r$ matrix with i.i.d. random signs and then take its r left singular vectors. The diagonal of $\mathbf{\Lambda}_*$ are set to values that are linearly distributed from $\frac{1}{10}$ to 1. Next, we sample $\mathbf{f}_1, \dots, \mathbf{f}_n$ from a multivariate normal distribution with zero mean and identity covariance matrix. We sample $\boldsymbol{\epsilon}_1, \dots, \boldsymbol{\epsilon}_n$ from a multivariate normal distribution with zero mean and covariance matrix $\boldsymbol{\Psi}_*$, which is a $d \times d$ positive definite matrix with a sparsity level $s = 0.02 \times d^2$. Finally, we generate the samples $\mathbf{x}_1, \dots, \mathbf{x}_n$ based on model (1).

1) *Estimation Performance:* First, we compare the estimation performance of different methods based on the estimation errors in Frobenius norm for the low-rank matrix (i.e., $\|\mathbf{L}_T - \mathbf{L}_*\|_F$), the sparse matrix (i.e., $\|\boldsymbol{\Psi}_T - \boldsymbol{\Psi}_*\|_F$), and the covariance matrix (i.e., $\|\Sigma_T - \Sigma_*\|_F$). Furthermore, we assess the selection performance of the sparse matrix by analyzing both the false positive rate (FPR) and the true positive rate (TPR). FPR refers to the proportion of non-significant elements that are incorrectly identified as significant in the estimated sparse matrix. In contrast, TPR measures the proportion of actual significant elements that are correctly identified in the estimated sparse matrix. These metrics together provide a comprehensive evaluation of both estimation precision and the effectiveness of identifying significant elements within the sparse matrices. The definitions of FPR and TPR are given as follows:

$$\text{FPR} = \frac{\#\{(i, j) : \boldsymbol{\Psi}_{ij} \neq 0 \& (\boldsymbol{\Psi}_*)_{ij} = 0\}}{\#\{(i, j) : (\boldsymbol{\Psi}_*)_{ij} = 0\}},$$

$$\text{TPR} = \frac{\#\{(i, j) : \boldsymbol{\Psi}_{ij} \neq 0 \& (\boldsymbol{\Psi}_*)_{ij} \neq 0\}}{\#\{(i, j) : (\boldsymbol{\Psi}_*)_{ij} \neq 0\}}.$$

TABLE I
QUANTITATIVE COMPARISON AMONG SEVEN DIFFERENT METHODS

		Strict factor models		Approximate factor models			
		POET-D	AltProj-D	POET	LOREC	AltProj	POETRY (prop.)
SCM							
$n = 50, d = 100, r = 2, s = 200$							
$\ \mathbf{L}_T - \mathbf{L}_\star\ _F$	—	1.6242±0.1307	1.3632±0.1096	1.6242±0.1307	1.4283±0.0869	1.3653±0.1314	1.3584±0.1281
$\ \boldsymbol{\Psi}_T - \boldsymbol{\Psi}_\star\ _F$	—	2.4714±0.0409	2.3842±0.0455	2.3845±0.0262	2.6151±0.0298	2.2983±0.0267	2.2921±0.0289
$\ \boldsymbol{\Sigma}_T - \boldsymbol{\Sigma}_\star\ _F$	2.7091±0.0423	2.6586±0.0502	2.5374±0.0502	2.5418±0.0479	2.5447±0.0354	2.4397±0.0453	2.4378±0.0462
FPR	—	—	—	0.0220±0.0011	0.0369±0.0150	0.0343±0.0010	0.0200±0.0011
TPR	—	0.8981±0.3156	0.9434±0.0001	0.9472±0.0080	0.9377±0.0150	0.7953±0.0278	0.9491±0.0091
Time	—	—	0.2281	—	1.1719	0.3266	0.1516
$n = 80, d = 200, r = 5, s = 800$							
$\ \mathbf{L}_T - \mathbf{L}_\star\ _F$	—	2.2082±0.0553	1.9106±0.0487	2.2082±0.0553	2.1194±0.0362	1.9247±0.0841	1.9072±0.0814
$\ \boldsymbol{\Psi}_T - \boldsymbol{\Psi}_\star\ _F$	—	3.8206±0.0400	3.7043±0.0552	3.7159±0.0319	4.1742±0.0347	3.6637±0.0393	3.6526±0.0426
$\ \boldsymbol{\Sigma}_T - \boldsymbol{\Sigma}_\star\ _F$	4.0526±0.0264	4.0300±0.0381	3.9011±0.0424	3.9144±0.0296	3.9279±0.0301	3.8254±0.0362	3.8249±0.0361
FPR	—	—	—	0.0110±0.0004	0.0141±0.0027	0.0175±0.0008	0.0100±0.0004
TPR	—	0.8971±0.5127	0.9524±0.0002	0.9571±0.0081	0.9543±0.0149	0.8505±0.0265	0.9581±0.0080
Time	—	—	0.6578	—	3.2141	1.4719	0.5828
$n = 100, d = 500, r = 10, s = 5000$							
$\ \mathbf{L}_T - \mathbf{L}_\star\ _F$	—	2.6486±0.0284	2.5134±0.1010	2.6486±0.0284	2.6562±0.0246	2.4847±0.0346	2.4800±0.0333
$\ \boldsymbol{\Psi}_T - \boldsymbol{\Psi}_\star\ _F$	—	7.0320±0.0205	6.9668±0.0311	6.9375±0.0107	7.5082±0.0141	6.8706±0.0134	6.8668±0.0131
$\ \boldsymbol{\Sigma}_T - \boldsymbol{\Sigma}_\star\ _F$	7.2616±0.0113	7.2427±0.0235	7.1926±0.0246	7.1378±0.0147	7.2592±0.0159	7.0933±0.0145	7.0924±0.0149
FPR	—	—	—	0.0041±0.0006	0.0045±0.0002	0.0044±0.0006	0.0040±0.0003
TPR	—	0.9315±0.0057	0.9545±0.0003	0.9630±0.0057	0.9363±0.0058	0.8795±0.0085	0.9637±0.0060
Time	—	—	9.6219	—	16.5750	8.1828	4.5719
$n = 400, d = 1000, r = 20, s = 20000$							
$\ \mathbf{L}_T - \mathbf{L}_\star\ _F$	—	4.5096±0.0245	3.8785±0.0396	4.5096±0.0245	5.6019±0.0293	3.7858±0.0374	3.7697±0.0360
$\ \boldsymbol{\Psi}_T - \boldsymbol{\Psi}_\star\ _F$	—	8.3495±0.0300	8.1280±0.0346	7.9495±0.0220	9.8759±0.0230	7.6904±0.0254	7.6815±0.0242
$\ \boldsymbol{\Sigma}_T - \boldsymbol{\Sigma}_\star\ _F$	8.8933±0.0138	8.7787±0.0284	8.5077±0.0325	8.3367±0.0208	8.6932±0.0164	8.0873±0.0243	8.0869±0.0243
FPR	—	—	—	0.0022±0.0003	0.0033±0.0002	0.0036±0.0001	0.0021±0.0002
TPR	—	0.9015±0.0057	0.9557±0.0008	0.9714±0.0060	0.9708±0.0036	0.9315±0.0069	0.9715±0.0059
Time	—	—	38.8047	—	207.6719	35.4063	27.0766

We conduct a series of experiments by varying parameters n , d , r , and s . For each parameter setting, we generate ten datasets to get ten replication results and calculate the mean and standard error of the evaluation metrics. In implementing our algorithm, we set the thresholding parameter τ as $2s$, and choose a step size of $\eta = 0.25$. In both our algorithm and the benchmarks, where the rank needs to be specified, we use the ground truth rank r . For the other parameters in the benchmarks, we select them using 4-fold cross-validation.

Table I summarizes the results. In terms of Frobenius norm estimation errors, POETRY consistently outperforms the other methods. In terms of selection performance, POETRY achieves a higher TPR and a lower FPR than the benchmark methods, which highlights its effectiveness in accurately identifying significant elements. We do not report the FPR for POET-D and AltProj-D, as both methods are designed for the low-rank plus diagonal structure, which naturally yields a zero FPR. Additionally, to underscore the efficiency of our algorithm, we have included the average CPU time in Table I. Since the SCM, POET-D, and POET estimators are direct methods, we do not provide their average CPU time. We observe that POETRY offers considerable computational advantages, being nearly three times faster than the convex method. When the dimension d gets larger, the average CPU time for estimators rises, with our estimator increasing at the slowest rate. Particularly, as d scales up to several hundred, the

repetitive computation of spectral decomposition in LOREC becomes extremely time-consuming, leading to a dramatic increase in computational time.

2) *Statistical Convergence Rate*: In this section, we verify the theoretical results presented in Theorems 5 and 7. We first analyze the squared estimation error, $\|\Sigma_T - \Sigma_\star\|_F^2$, with the sample size n set to the magnitude of rd , as prescribed by Theorem 5. We plot this estimation error against the theoretical statistical error $s \log d/n$ for $r = 2, 4, 6, 8$ in Fig. 1(a). This figure shows that $\|\Sigma_T - \Sigma_\star\|_F^2$ increases linearly with the statistical error $s \log d/n$. Similarly, Fig. 1(b) illustrates the estimation error against the statistical error rd/n for $s = 200, 300, 400, 500$. We observe that $\|\Sigma_T - \Sigma_\star\|_F^2$ grows linearly with the statistical error rd/n . These results corroborate the theoretical statistical rate stated in Theorem 5, validating our guarantee of achieving the minimax optimal statistical rate. Next, we calculate the squared estimation error $\|\Sigma_T - \Sigma_\star\|_F^2$ with the sample size n set to the magnitude of $r_e r \log 2d$, as required in Theorem 7. The estimation error is plotted against the statistical error $s \log d/n$ for $r = 2, 4, 6, 8$ in Fig. 2(a). We observe that $\|\Sigma_T - \Sigma_\star\|_F^2$ grows linearly with the statistical error $s \log d/n$. Furthermore, we also plot the estimation error against the statistical error $r \log(2d)/n$ for $s = 200, 300, 400, 500$ in Fig. 2(b). This plot shows that the estimation error grows linearly with the statistical error $r \log(2d)/n$. These numerical results align with the

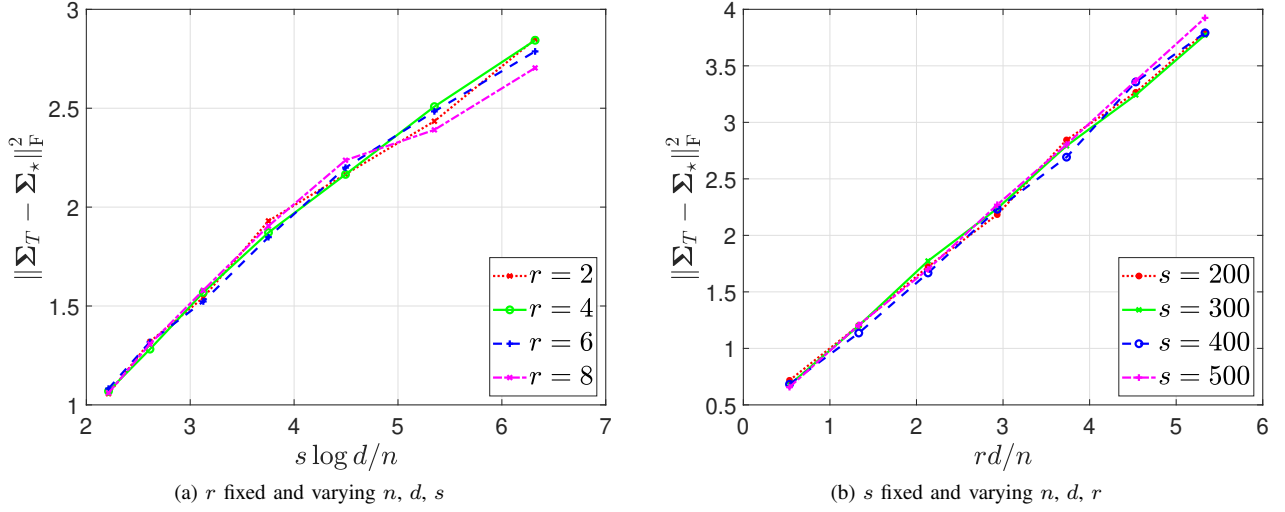


Fig. 1. Estimation errors $\|\Sigma_T - \Sigma_\star\|_F^2$ versus theoretical statistical errors.

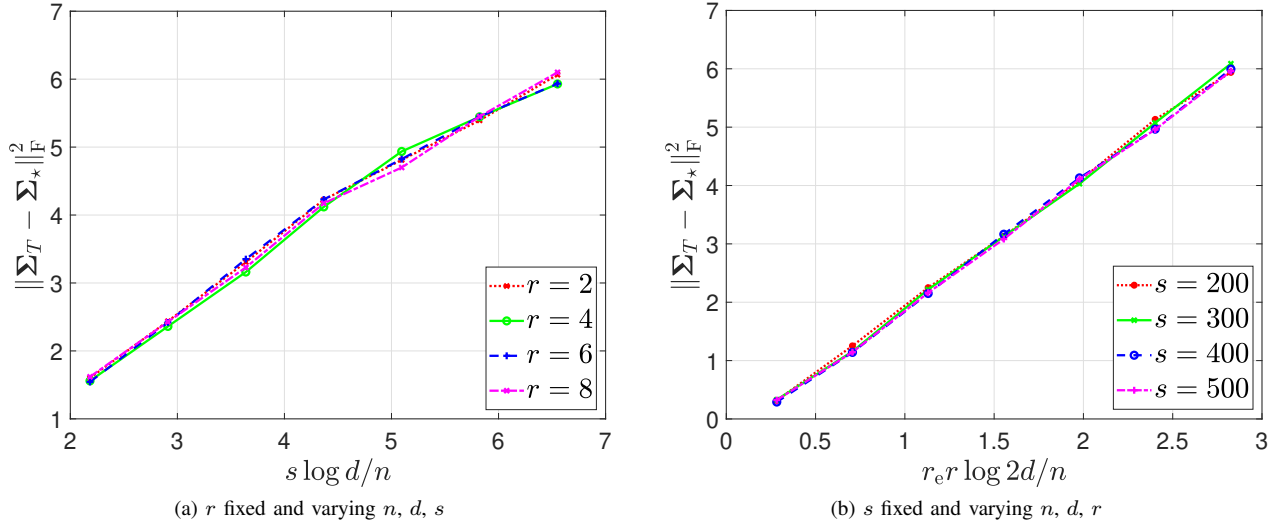


Fig. 2. Estimation errors $\|\Sigma_T - \Sigma_\star\|_F^2$ versus theoretical statistical errors.

theoretical findings of Theorem 7.

B. Real Datasets

1) *Macroeconomic Factors*: The federal reserve economic data-monthly database (FRED-MD)¹ is an open-source dataset containing monthly U.S. macroeconomic indicators. It is maintained by the Federal Reserve Bank of St. Louis [70] and serves as a common benchmark for evaluating model performance and enhancing research reproducibility. FRED-MD covers a wide range of macroeconomic time series that cover key aspects of the U.S. economy, including employment, production, income, and prices. The dataset includes 127 standard U.S. macroeconomic indicators, with data available from January 1, 1959. These 127 indicators are categorized into seven groups based on their characteristics. To prepare

the data for statistical modeling, various transformations are applied to the time series to make them stationary. The transformations include: (1) no transformation; (2) first-order difference; (3) second-order difference; (4) logarithm; (5) first-order logarithmic difference; (6) second-order logarithmic difference; and (7) percentage change. In [70], eight factors from the FRED-MD dataset were estimated using the POET method [41]. In this section, we conduct an analysis based on the proposed method, POETRY. Our analysis begins with a sample starting on January 1, 1960. After accounting for two lost observations due to data transformation, the dataset used for analysis spans from March 1, 1960, to April 1, 2015, resulting in a total of 658 observations.

First, we will verify the number of factors identified by [70] and then analyze the estimated sparse component. We define the correlation matrix estimator as follows:

$$\hat{\Gamma}_T = \left[\text{diag} \left(\hat{\Sigma}_T \right) \right]^{-1/2} \hat{\Sigma}_T \left[\text{diag} \left(\hat{\Sigma}_T \right) \right]^{-1/2}.$$

¹The dataset is publicly available and can be downloaded from the federal reserve economic data website at <https://fred.stlouisfed.org/>.

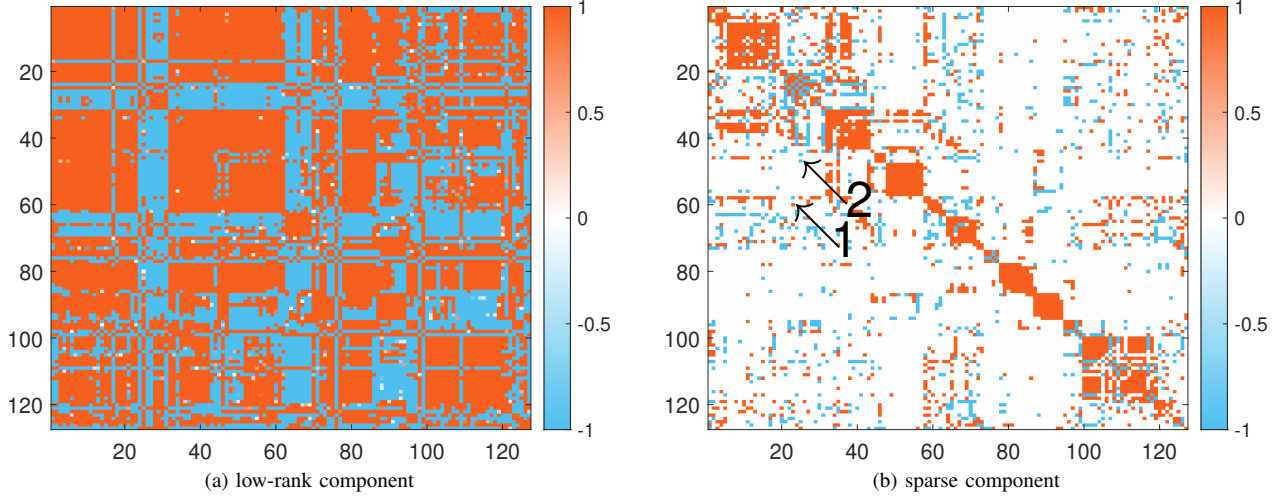


Fig. 3. Heatmaps of the low-rank component and sparse component.

According to [71], the number of eigenvalues greater than 1 in the correlation matrix is often used to determine the number of factors present. When we apply the POETRY method and set $r \geq 8$, we consistently observe 8 eigenvalues greater than 1. This suggests that the rank of the low-rank component should be chosen as 8, aligning with the results presented in [70]. Setting $r = 8$, we obtain eight factors from the correlation matrix, with eigenvalues greater than 1 being 18.70, 9.58, 8.59, 6.66, 5.65, 4.43, 3.94, and 2.90. For the parameter τ , we select its value from $0.05 \times d^2$, $0.25 \times d^2$, $0.45 \times d^2$, and $0.65 \times d^2$ by choosing the value that yields the lowest objective function value. We visualize both the low-rank component and the sparse component of the covariance matrix using heatmaps in Fig. 3. Both components are normalized by the largest element in magnitude within their respective matrices. This visualization helps us understand the underlying structure and correlations within the dataset. Next, we examine the estimated sparse component, which is illustrated in Fig. 3(b). The sparse component provides intriguing insights. For example, the element indicated by the first arrow shows a positive correlation between the civilian labor force and orders for consumer goods. An expanding labor force typically corresponds to increased consumer goods orders, reflecting overall economic growth and consumer confidence. In addition, the element pointed to by the second arrow reveals a negative correlation between the average duration of unemployment and average weekly hours worked. Generally, these two variables are inversely related: a longer duration of unemployment is associated with fewer average weekly hours worked, indicating a weaker labor market. We also examine other significant nonzero correlations that seem reasonable within this context.

2) *Portfolio Design*: We further apply our covariance estimation methods for portfolio design in financial stock markets². We collect the historical daily stock prices of the components listed in the S&P 500 Index from January 1st, 2010, to December 31st, 2020. After removing stocks with missing values, we obtain daily returns of 415 stocks ($d = 415$). We

construct 100 daily return time series datasets with random starting dates, each containing daily returns over 252 consecutive trading days. We first estimate the covariance matrix of each dataset. For each method, we select the parameters that yield the lowest portfolio risk. We then compute the risk of the portfolio return designed based on these estimated covariance matrices. In this context, the quality of the estimated covariance matrix is directly related to portfolio risk; a higher quality estimation leads to lower portfolio risk. This allows us to assess and compare the effectiveness of different estimation methods in minimizing portfolio risk. In this paper, we adopt the global minimum variance portfolio (GMVP) [72]. Let Σ be a generic estimator of the return covariance matrix, and let w be the allocation vector of a portfolio consisting of the corresponding d stocks. The GMVP is defined as the solution to the following convex optimization problem:

$$\begin{aligned} & \underset{w \in \mathbb{R}^d}{\text{minimize}} && w^\top \Sigma w \\ & \text{subject to} && w^\top \mathbf{1} = 1, \quad w \geq 0, \end{aligned}$$

where the objective represents the risk of a portfolio. The solution to this optimization problem is obtained using CVX, a MATLAB-based modeling system for convex optimization [73].

The boxplot results presented in Fig. 4 illustrate the risk of the portfolios generated by different covariance estimators. This figure indicates that the portfolio risk associated with SCM is higher than other methods, suggesting its inferior performance. Additionally, we find that POET-D and AltProj-D, which are based on strict factor models for covariance estimation, perform less effectively compared to those utilizing approximate factor models. This outcome highlights the advantages of the approximate factor model assumption. Most notably, the GMVP generated by POETRY exhibits the lowest risk among the methods analyzed. This result underscores the effectiveness of POETRY in constructing a robust and low-risk portfolio, thereby validating the superiority of our method in practical financial applications.

²The dataset can be downloaded from <https://finance.yahoo.com/>.

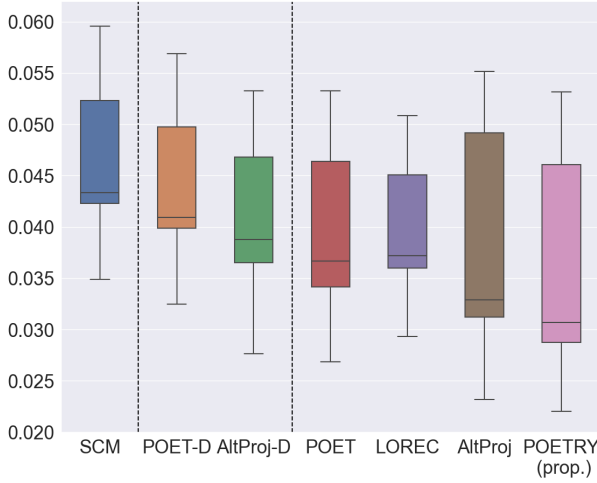


Fig. 4. Boxplots of daily volatility under different estimators on stock market data.

VI. CONCLUSIONS

In this paper, we have studied the problem of estimating large covariance matrices based on factor models, where the covariance matrix is decomposed into a low-rank matrix and a sparse matrix. We have proposed a nonconvex optimization model, together with an efficient numerical algorithm. Furthermore, we have established that the sequence generated by our algorithm converges linearly to the ground truth at a minimax statistical rate. We have also derived the statistical rate of the inverse covariance matrix estimator. By introducing the concept of effective rank, we have reduced the required sample size and further improved the statistical rate. Numerical experiments support our theoretical findings and demonstrate the superiority of our algorithm.

APPENDIX

A. Auxiliary Lemmata

Lemma 8 ([45]). *For the noise matrix $N = \frac{1}{n} \sum_{i=1}^n \mathbf{y}_i \mathbf{y}_i^\top - \Sigma_\star$, where each $\mathbf{y}_i \in \mathbb{R}^d$ is a Gaussian random variable with mean zero and covariance $\Sigma_\star$, with probability at least $1 - 2 \exp(-C_1 d)$, it holds that*

$$\|N\|_2 \leq \alpha \sqrt{\frac{d}{n}},$$

where $\alpha = 4 \|\Sigma_\star\|_2$. Moreover, with probability at least $1 - 2 \exp(-C_2 \log d)$, we have

$$\|N\|_\infty \leq \beta \sqrt{\frac{\log d}{n}},$$

and $\beta = 8 \times \max_j (\Sigma_\star)_{jj}$.

Lemma 9 ([68]). *For the noise matrix $N = \frac{1}{n} \sum_{i=1}^n \mathbf{y}_i \mathbf{y}_i^\top - \Sigma_\star$, where each $\mathbf{y}_i \in \mathbb{R}^d$ is a Gaussian random variable with mean zero and covariance $\Sigma_\star$, with probability at least $1 - 1/(2d)$, it holds that*

$$\|N\|_2 \leq \gamma \left(\sqrt{\frac{2r_e \log(2d)}{n}} \vee \frac{2r_e \log(2d)(3/8 + \log n)}{n} \right),$$

where γ is a positive constant related to $\Sigma_\star$.

Given the spectral decomposition $B_\star B_\star^\top = V_\star \Lambda_\star V_\star^\top$, we define $B_\star = V_\star \Lambda_\star^{1/2}$.

Definition 10. Define the distance between B and $B_\star$ as

$$\text{dist}(B, B_\star) := \min_{Q \in \mathbb{O}^r} \|(BQ - B_\star) \Lambda_\star^{1/2}\|_F, \quad (9)$$

where $\mathbb{O}^r = \{Q \in \mathbb{R}^{r \times r} \mid QQ^\top = I\}$.

The factorization of the low-rank component $L_\star = B_\star B_\star^\top$ is not unique, as $L_\star = (B_\star Q)(B_\star Q)^\top$ holds for any orthogonal matrix Q . Therefore, in Definition 10, we define the distance error metric within the ambiguity class considering orthonormal transformations. Additionally, we take into account diagonal scaling in the metric to properly address the impact of preconditioning [63]. We denote the optimal alignment matrix Q between B and $B_\star$ as

$$Q := \arg \min_{Q \in \mathbb{O}^r} \|(BQ - B_\star) \Lambda_\star^{1/2}\|_F.$$

This implies that Q satisfies the following first-order necessary optimality condition:

$$(BQ)^\top (BQ - B_\star) \Lambda_\star = 0. \quad (10)$$

Based on Definition 10, we can obtain the following results.

Lemma 11 ([74], [75]). *For any loading matrix B , the distance between B and $B_\star$ satisfies*

$$\text{dist}(B, B_\star) \leq \left(\frac{\sqrt{2} + 1}{2} \right)^{1/2} \|BB^\top - L_\star\|_F.$$

Lemma 12. *For any loading matrix B that satisfies the condition $\|(B - B_\star) \Lambda_\star^{-1/2}\|_2 < 1$, the following holds:*

$$\|B(B^\top B)^{-1} \Lambda_\star^{1/2}\|_2 \leq \frac{1}{1 - \|(B - B_\star) \Lambda_\star^{-1/2}\|_2},$$

Proof: Notice that

$$\|B(B^\top B)^{-1} \Lambda_\star^{1/2}\|_2 = \frac{1}{\sigma_r(B \Lambda_\star^{-1/2})}.$$

Based on Weyl's inequality, we obtain

$$\begin{aligned} \sigma_r(B \Lambda_\star^{-1/2}) &\geq \sigma_r(B_\star \Lambda_\star^{-1/2}) - \|(B - B_\star) \Lambda_\star^{-1/2}\|_2 \\ &= 1 - \|(B - B_\star) \Lambda_\star^{-1/2}\|_2, \end{aligned}$$

where we have used $\sigma_r(B_\star \Lambda_\star^{-1/2}) = \sigma_r(V_\star) = 1$. The result is proved by combining the preceding two relations. ■

Lemma 13. *Let $\Psi = H_\tau(S - BB^\top)$, where $s < \tau \leq \pi s$ with $\pi \geq 1$, we have*

$$\|\Psi - \Psi_\star\|_\infty \leq 2 \|BB^\top - L_\star\|_\infty + 2 \|N\|_\infty, \quad (11)$$

and

$$\begin{aligned} & \left| \text{tr}((\Psi - \Psi_\star) Z^\top) \right| \\ & \leq \left(2 \|BB^\top - L_\star\|_F + (1 + \sqrt{\pi} + 1) \sqrt{s} \|N\|_\infty \right) \|Z\|_F, \end{aligned} \quad (12)$$

for any matrix $\mathbf{Z}$.

Proof: Let $\mathcal{S}$ and $\mathcal{S}_*$ represent the support set of Ψ and Ψ_* , respectively. Consequently, $\Psi - \Psi_*$ is supported on $\mathcal{S} \cup \mathcal{S}_*$. Define $\Delta_L = \mathbf{B}\mathbf{B}^\top - \mathbf{L}_*$. For any $(i, j) \in \mathcal{S}$, based on the definition of $\mathbf{H}_\tau(\cdot)$, we have $(\Psi - \Psi_*)_{ij} = \mathbf{N}_{ij} - (\Delta_L)_{ij}$. For any $(i, j) \in \mathcal{S}_* \setminus \mathcal{S}$, one necessarily has $\Psi_{ij} = 0$, thus leading to $(\Psi - \Psi_*)_{ij} = -(\Psi_*)_{ij}$. Again, by the definition of $\mathbf{H}_\tau(\cdot)$, we know that $|\Psi_* + \mathbf{N} - \Delta_L|_{ij} \leq |\Psi_* + \mathbf{N} - \Delta_L|_\tau$, where $|\mathbf{A}|_\tau$ is the τ -th largest absolute entry of matrix $\mathbf{A}$. Furthermore, we know that Ψ_* contains at most s nonzero entries and $\tau > s$. Consequently, one has $|\Psi_* + \mathbf{N} - \Delta_L|_{ij} \leq |\mathbf{N} - \Delta_L|_\tau$. Combining the two cases above, we conclude that

$$\begin{aligned} & |\Psi - \Psi_*|_{ij} \\ & \leq \begin{cases} |\mathbf{N}|_{ij} + |\Delta_L|_{ij}, & (i, j) \in \mathcal{S}, \\ |\mathbf{N}|_{ij} + |\Delta_L|_{ij} + |\mathbf{N}|_\tau + |\Delta_L|_\tau, & (i, j) \in \mathcal{S}_* \setminus \mathcal{S}, \end{cases} \end{aligned} \quad (13)$$

which immediately implies the result in (11).

Recall that $\Psi - \Psi_*$ is supported on $\mathcal{S} \cup \mathcal{S}_*$. Therefore, we have

$$\begin{aligned} & \left| \text{tr}((\Psi - \Psi_*) \mathbf{Z}^\top) \right| \\ & \leq \sum_{(i,j) \in \mathcal{S}} |\Psi - \Psi_*|_{ij} |\mathbf{Z}|_{ij} + \sum_{(i,j) \in \mathcal{S}_* \setminus \mathcal{S}} |\Psi - \Psi_*|_{ij} |\mathbf{Z}|_{ij} \\ & \leq \underbrace{\sum_{(i,j) \in \mathcal{S} \cup \mathcal{S}_*} |\mathbf{N}|_{ij} |\mathbf{Z}|_{ij}}_{I_1} + \underbrace{\sum_{(i,j) \in \mathcal{S} \cup \mathcal{S}_*} |\Delta_L|_{ij} |\mathbf{Z}|_{ij}}_{I_2} \\ & \quad + \underbrace{\sum_{(i,j) \in \mathcal{S}_* \setminus \mathcal{S}} |\mathbf{N}|_\tau |\mathbf{Z}|_{ij}}_{I_3} + \underbrace{\sum_{(i,j) \in \mathcal{S}_* \setminus \mathcal{S}} |\Delta_L|_\tau |\mathbf{Z}|_{ij}}_{I_4}, \end{aligned}$$

where the second inequality follows from (13). We can bound I_1 as follows:

$$I_1 \leq \sqrt{\sum_{(i,j) \in \mathcal{S} \cup \mathcal{S}_*} |\mathbf{N}|_{ij}^2} \|\mathbf{Z}\|_F \leq \sqrt{(\pi + 1)s} \|\mathbf{N}\|_\infty \|\mathbf{Z}\|_F,$$

where the first inequality follows from the Cauchy inequality, and the second inequality follows from the fact that $\mathcal{S} \cup \mathcal{S}_*$ contains at most $\tau + s$ nonzero entries with $s < \tau \leq \pi s$. For I_2 , we invoke the Cauchy inequality to obtain

$$I_2 \leq \|\Delta_L\|_F \|\mathbf{Z}\|_F.$$

We can further upperbound I_3 as

$$I_3 \leq \|\mathbf{N}\|_\infty \sqrt{s \sum_{(i,j) \in \mathcal{S}_*} |\mathbf{Z}|_{ij}^2} \leq \sqrt{s} \|\mathbf{N}\|_\infty \|\mathbf{Z}\|_F,$$

where the first inequality follows from the Jensen's inequality $(\frac{1}{n} \sum_{i=1}^n x_i)^2 \leq \frac{1}{n} \sum_{i=1}^n x_i^2$.

We can also further bound I_4 as

$$I_4 \leq |\Delta_L|_\tau \sum_{(i,j) \in \mathcal{S}_*} |\mathbf{Z}|_{ij} \leq \sqrt{s} |\Delta_L|_\tau \|\mathbf{Z}\|_F \leq \|\Delta_L\|_F \|\mathbf{Z}\|_F,$$

where the last inequality is due to $\sqrt{s} |\Delta_L|_\tau \leq \|\Delta_L\|_F$. Combine I_1 , I_2 , I_3 , and I_4 together to obtain the result in (12). $\blacksquare$

B. Propositions

Proposition 14. Suppose Assumptions 1 and 2 hold. Assume that the sample size satisfies $n \gtrsim rd$. Choose the thresholding parameter as $s < \tau \leq \pi s$ with $\pi \geq 1$. Then, with probability at least $1 - 2 \exp(-C_1 d)$, the initialization $\mathbf{B}_0$ of Algorithm 2 satisfies

$$\text{dist}(\mathbf{B}_0, \mathbf{B}_*) < \sigma_r(\mathbf{L}_*),$$

where C_1 is defined in Lemma 8.

Proof: From Lemma 11, we have

$$\begin{aligned} \text{dist}(\mathbf{B}_0, \mathbf{B}_*) & \leq \left(\frac{\sqrt{2} + 1}{2} \right)^{1/2} \left\| \mathbf{B}_0 \mathbf{B}_0^\top - \mathbf{L}_* \right\|_F \\ & \leq \sqrt{(\sqrt{2} + 1)r} \left\| \mathbf{B}_0 \mathbf{B}_0^\top - \mathbf{L}_* \right\|_2, \end{aligned}$$

where the last inequality is because $\mathbf{B}_0 \mathbf{B}_0^\top - \mathbf{L}_*$ has a rank of at most $2r$. By applying the triangle inequality, we can derive

$$\begin{aligned} \left\| \mathbf{B}_0 \mathbf{B}_0^\top - \mathbf{L}_* \right\|_2 & \leq \left\| \mathbf{S} - \mathbf{B}_0 \mathbf{B}_0^\top \right\|_2 + \left\| \mathbf{S} - \mathbf{L}_* \right\|_2 \\ & \leq 2 \left\| \mathbf{S} - \mathbf{L}_* \right\|_2 \\ & = 2 \left\| \mathbf{N} + \Sigma_* - \mathbf{L}_* \right\|_2 \\ & \leq 2 (\|\mathbf{N}\|_2 + \|\Sigma_*\|_2 + \|\mathbf{L}_*\|_2), \end{aligned} \quad (14)$$

where the second inequality is based on the fact that $\mathbf{B}_0 \mathbf{B}_0^\top$ is the best rank- r approximation of $\mathbf{S}$. Hence, we have

$$\text{dist}(\mathbf{B}_0, \mathbf{B}_*) \leq 2 \sqrt{(\sqrt{2} + 1)r} (\|\mathbf{N}\|_2 + \mu + k), \quad (15)$$

where the inequality is due to $\|\Sigma_*\|_2 \leq \mu$ from Assumption 1 and $\|\mathbf{L}_*\|_2 \leq \|\mathbf{L}_*\|_F \leq d \|\mathbf{L}_*\|_\infty \leq k$ from Assumption 2, we can apply Lemma 8 to obtain the following result

$$\text{dist}(\mathbf{B}_0, \mathbf{B}_*) \leq 2 \sqrt{(\sqrt{2} + 1)r} \left(\beta \sqrt{\frac{d}{n}} + \mu + k \right),$$

which holds with probability at least $1 - 2 \exp(-C_1 d)$. With $n > 48rd/\sigma_r^2(\mathbf{L}_*)$, we can then conclude that

$$\text{dist}(\mathbf{B}_0, \mathbf{B}_*) < \sigma_r(\mathbf{L}_*). \quad \blacksquare$$

Proposition 15. Suppose Assumptions 1 and 2 hold. Assume that the sample size satisfies $n \gtrsim r_e r \log 2d$. Choose the thresholding parameter as $s < \tau \leq \pi s$ with $\pi \geq 1$. Then, with probability at least $1 - 1/(2d)$, the initialization $\mathbf{B}_0$ of Algorithm 2 satisfies

$$\text{dist}(\mathbf{B}_0, \mathbf{B}_*) < \sigma_r(\mathbf{L}_*).$$

Proof: From (15), we have

$$\text{dist}(\mathbf{B}_0, \mathbf{B}_*) \leq 2 \sqrt{(\sqrt{2} + 1)r} (\|\mathbf{N}\|_2 + \mu + k).$$

Based on Lemma 9 and $r_e \ll d$, we have that with probability at least $1 - 1/(2d)$,

$$\text{dist}(\mathbf{B}_0, \mathbf{B}_*) \leq 2 \sqrt{(\sqrt{2} + 1)r} \left(\gamma \sqrt{\frac{2r_e \log(2d)}{n}} + \mu + k \right).$$

Then, with $n > 12r_e r \log 2d/\sigma_r^2(\mathbf{L}_*)$, we have

$$\text{dist}(\mathbf{B}_0, \mathbf{B}_*) < \sigma_r(\mathbf{L}_*).$$

Proposition 16. *Given the update of $\mathbf{B}_{t+1}$ in Algorithm 2, if the step size satisfies $\frac{2-\sqrt{2}}{20} < \eta < \frac{1}{2}$, then for any $t = 0, 1, \dots$, the following inequality holds*

$$\text{dist}^2(\mathbf{B}_{t+1}, \mathbf{B}_*) \leq \rho^{t+1} \sigma_r^2(\mathbf{L}_*) + \frac{\omega}{1-\rho},$$

where $\rho = 2(1-10\eta)^2$ and $\omega = \frac{4\eta^2 s(1+\sqrt{\pi+1})^2}{(1-\delta)^2} \|\mathbf{N}\|_\infty^2 + \frac{2\eta^2 r}{(1-\delta)^2} \|\mathbf{N}\|_2^2$.

Proof: By the definition of $\text{dist}(\mathbf{B}_{t+1}, \mathbf{B}_*)$, it can be shown that

$$\text{dist}^2(\mathbf{B}_{t+1}, \mathbf{B}_*) \leq \left\| (\mathbf{B}_{t+1} \mathbf{Q}_t - \mathbf{B}_*) \mathbf{A}_*^{1/2} \right\|_F^2,$$

where $\mathbf{Q}_t$ is the optimal alignment matrix between $\mathbf{B}_t$ and $\mathbf{B}_*$. Denote $\mathbf{B} = \mathbf{B}_t \mathbf{Q}_t$, $\Psi = \Psi_t$, $\Delta_B = \mathbf{B} - \mathbf{B}_*$, and $\Delta_L = \mathbf{B} \mathbf{B}^\top - \mathbf{L}_*$. Given the update rule $\mathbf{B}_{t+1} = \mathbf{B}_t - \eta(\mathbf{B}_t \mathbf{B}_t^\top + \Psi_t - \mathbf{S}) \mathbf{B}_t (\mathbf{B}_t^\top \mathbf{B}_t)^{-1}$, one has

$$\begin{aligned} & (\mathbf{B}_{t+1} \mathbf{Q}_t - \mathbf{B}_*) \mathbf{A}_*^{1/2} \\ &= \left(\mathbf{B} - \eta(\mathbf{B} \mathbf{B}^\top + \Psi - \mathbf{S}) \mathbf{B} (\mathbf{B}^\top \mathbf{B})^{-1} - \mathbf{B}_* \right) \mathbf{A}_*^{1/2} \\ &= \Delta_B \mathbf{A}_*^{1/2} - \eta(\Psi - \Psi_*) \mathbf{B} (\mathbf{B}^\top \mathbf{B})^{-1} \mathbf{A}_*^{1/2} \\ &\quad - \eta \Delta_L \mathbf{B} (\mathbf{B}^\top \mathbf{B})^{-1} \mathbf{A}_*^{1/2} + \eta \mathbf{N} \mathbf{B} (\mathbf{B}^\top \mathbf{B})^{-1} \mathbf{A}_*^{1/2} \\ &= (1-\eta) \Delta_B \mathbf{A}_*^{1/2} - \eta(\Psi - \Psi_*) \mathbf{B} (\mathbf{B}^\top \mathbf{B})^{-1} \mathbf{A}_*^{1/2} \\ &\quad - \eta \mathbf{B}_* \Delta_B^\top \mathbf{B} (\mathbf{B}^\top \mathbf{B})^{-1} \mathbf{A}_*^{1/2} + \eta \mathbf{N} \mathbf{B} (\mathbf{B}^\top \mathbf{B})^{-1} \mathbf{A}_*^{1/2}, \end{aligned} \quad (16)$$

where the second equality is due to $\mathbf{S} = \mathbf{L}_* + \Psi_* + \mathbf{N}$ and the last equality is due to the decomposition $\Delta_L = \Delta_B \mathbf{B}^\top + \mathbf{B}_* \Delta_B^\top$. Take the Frobenius norm of both sides of (16) and use $(a+b)^2 \leq 2a^2 + 2b^2$ to obtain (17), where $\mathbf{Z} = \mathbf{B}_* \Delta_B^\top \mathbf{B} (\mathbf{B}^\top \mathbf{B})^{-1} \mathbf{A}_* (\mathbf{B}^\top \mathbf{B})^{-1} \mathbf{B}^\top$. For the first term of (17), it is easy to verify that $\text{tr}(\Delta_B \mathbf{A}_* \Delta_B^\top) = \text{dist}^2(\mathbf{B}_t, \mathbf{B}_*)$. In the following proof, we will focus on controlling J_1, J_2, J_3, J_4, J_5 , and J_6 separately.

We start with J_1 . Since $\mathbf{B}_* = \mathbf{B} - \Delta_B$ and $\mathbf{A}_* \Delta_B^\top \mathbf{B} = \mathbf{0}$ from (10), we have

$$\begin{aligned} J_1 &= \text{tr} \left(\mathbf{B} (\mathbf{B}^\top \mathbf{B})^{-1} \mathbf{A}_* \Delta_B^\top \mathbf{B}_* \Delta_B^\top \right) \\ &= \text{tr} \left(\mathbf{B} (\mathbf{B}^\top \mathbf{B})^{-1} \mathbf{A}_* \Delta_B^\top \mathbf{B} \Delta_B^\top \right) \\ &\quad - \text{tr} \left(\mathbf{B} (\mathbf{B}^\top \mathbf{B})^{-1} \mathbf{A}_* \Delta_B^\top \Delta_B \Delta_B^\top \right) \\ &= -\text{tr} \left(\mathbf{B} (\mathbf{B}^\top \mathbf{B})^{-1} \mathbf{A}_* \Delta_B^\top \Delta_B \Delta_B^\top \right). \end{aligned}$$

We intend to show that for any $t = 0, 1, \dots$, it holds that

$$\text{dist}(\mathbf{B}_t, \mathbf{B}_*) \leq \delta \sigma_r(\mathbf{L}_*), \quad (18)$$

where $0 < \delta < 1$. We prove this result by induction. For $t = 0$, the induction hypothesis (18) holds trivially, as established by Propositions 14 and 15. We therefore concentrate on the induction step from t to $t+1$, for $t \geq 0$. Specifically, suppose

$$\text{dist}(\mathbf{B}_t, \mathbf{B}_*) = \left\| \Delta_B \mathbf{A}_*^{-1/2} \mathbf{A}_* \right\|_F \leq \delta \sigma_r(\mathbf{L}_*).$$

Since $\left\| \Delta_B \mathbf{A}_*^{-1/2} \mathbf{A}_* \right\|_F \geq \sigma_r(\mathbf{L}_*) \left\| \Delta_B \mathbf{A}_*^{-1/2} \right\|_F \geq \sigma_r(\mathbf{L}_*) \left\| \Delta_B \mathbf{A}_*^{-1/2} \right\|_2$, we have

$$\left\| \Delta_B \mathbf{A}_*^{-1/2} \right\|_2 \leq \delta. \quad (19)$$

Combining (19) with Lemma 12 leads to

$$\left\| \mathbf{B} (\mathbf{B}^\top \mathbf{B})^{-1} \mathbf{A}_*^{1/2} \right\|_2 \leq \frac{1}{1-\delta}. \quad (20)$$

Then, we have

$$\begin{aligned} |J_1| &= \left| \text{tr} \left(\mathbf{A}_*^{-1/2} \Delta_B^\top \mathbf{B} (\mathbf{B}^\top \mathbf{B})^{-1} \mathbf{A}_* \Delta_B^\top \Delta_B \mathbf{A}_*^{1/2} \right) \right| \\ &\leq \left\| \mathbf{A}_*^{-1/2} \Delta_B^\top \mathbf{B} (\mathbf{B}^\top \mathbf{B})^{-1} \mathbf{A}_*^{1/2} \right\|_2 \text{tr} \left(\mathbf{A}_*^{1/2} \Delta_B^\top \Delta_B \mathbf{A}_*^{1/2} \right) \\ &\leq \left\| \Delta_B \mathbf{A}_*^{-1/2} \right\|_2 \left\| \mathbf{B} (\mathbf{B}^\top \mathbf{B})^{-1} \mathbf{A}_*^{1/2} \right\|_2 \text{tr} \left(\Delta_B \mathbf{A}_* \Delta_B^\top \right) \\ &\leq \frac{\delta}{1-\delta} \text{tr} \left(\Delta_B \mathbf{A}_* \Delta_B^\top \right), \end{aligned}$$

where the last inequality is due to (19) and (20).

Moving on to bound J_2 , we can utilize the fact $\mathbf{B}^\top \Delta_B \mathbf{A}_* = \mathbf{0}$ from (10) to obtain

$$\begin{aligned} J_2 &= \text{tr} \left(\mathbf{B}_* \Delta_B^\top \mathbf{B} (\mathbf{B}^\top \mathbf{B})^{-1} \mathbf{A}_* (\mathbf{B}^\top \mathbf{B})^{-1} \mathbf{B}^\top \Delta_B \mathbf{B}_*^\top \right) \\ &= \text{tr} \left(\mathbf{B} (\mathbf{B}^\top \mathbf{B})^{-1} \mathbf{A}_* (\mathbf{B}^\top \mathbf{B})^{-1} \mathbf{B}^\top \Delta_B \mathbf{A}_* \Delta_B^\top \right) = 0. \end{aligned}$$

For J_3 , we have

$$\begin{aligned} |J_3| &\leq (2 \|\Delta_L\|_F + (1 + \sqrt{\pi+1}) \sqrt{s} \|\mathbf{N}\|_\infty) \\ &\quad \cdot \left\| \Delta_B \mathbf{A}_* (\mathbf{B}^\top \mathbf{B})^{-1} \mathbf{B}^\top \right\|_F \\ &\leq (2(2+\delta) \left\| \Delta_B \mathbf{A}_*^{1/2} \right\|_F + (1 + \sqrt{\pi+1}) \sqrt{s} \|\mathbf{N}\|_\infty) \\ &\quad \cdot \left\| \Delta_B \mathbf{A}_*^{1/2} \right\|_F \left\| \mathbf{B} (\mathbf{B}^\top \mathbf{B})^{-1} \mathbf{A}_*^{1/2} \right\|_2 \\ &\leq \frac{2(2+\delta)}{1-\delta} \left\| \Delta_B \mathbf{A}_*^{1/2} \right\|_F^2 \\ &\quad + \frac{(1 + \sqrt{\pi+1}) \sqrt{s}}{1-\delta} \|\mathbf{N}\|_\infty \left\| \Delta_B \mathbf{A}_*^{1/2} \right\|_F, \end{aligned}$$

where the first inequality is due to (12) in Lemma 13, and in the second inequality, we have used the relation $\|\mathbf{A} \mathbf{X}\|_F \leq$

$$\begin{aligned}
\text{dist}^2(\mathbf{B}_{t+1}, \mathbf{B}_\star) &\leq 2 \left\| (1-\eta) \mathbf{\Delta}_B \mathbf{\Lambda}_\star^{1/2} - \eta (\mathbf{\Psi} - \mathbf{\Psi}_\star) \mathbf{B} (\mathbf{B}^\top \mathbf{B})^{-1} \mathbf{\Lambda}_\star^{1/2} - \eta \mathbf{B}_\star \mathbf{\Delta}_B^\top \mathbf{B} (\mathbf{B}^\top \mathbf{B})^{-1} \mathbf{\Lambda}_\star^{1/2} \right\|_F^2 \\
&\quad + 2 \left\| \eta \mathbf{N} \mathbf{B} (\mathbf{B}^\top \mathbf{B})^{-1} \mathbf{\Lambda}_\star^{1/2} \right\|_F^2 \\
&\leq 2(1-\eta)^2 \text{tr} \left(\mathbf{\Delta}_B \mathbf{\Lambda}_\star \mathbf{\Delta}_B^\top \right) - 4\eta(1-\eta) \underbrace{\text{tr} \left(\mathbf{B}_\star \mathbf{\Delta}_B^\top \mathbf{B} (\mathbf{B}^\top \mathbf{B})^{-1} \mathbf{\Lambda}_\star \mathbf{\Delta}_B^\top \right)}_{J_1} \\
&\quad + 2\eta^2 \underbrace{\left\| \mathbf{B}_\star \mathbf{\Delta}_B^\top \mathbf{B} (\mathbf{B}^\top \mathbf{B})^{-1} \mathbf{\Lambda}_\star^{1/2} \right\|_F^2}_{J_2} - 4\eta(1-\eta) \underbrace{\text{tr} \left((\mathbf{\Psi} - \mathbf{\Psi}_\star) \mathbf{B} (\mathbf{B}^\top \mathbf{B})^{-1} \mathbf{\Lambda}_\star \mathbf{\Delta}_B^\top \right)}_{J_3} \\
&\quad + 4\eta^2 \underbrace{\text{tr} \left((\mathbf{\Psi} - \mathbf{\Psi}_\star) \mathbf{Z}^\top \right)}_{J_4} + 2\eta^2 \underbrace{\left\| (\mathbf{\Psi} - \mathbf{\Psi}_\star) \mathbf{B} (\mathbf{B}^\top \mathbf{B})^{-1} \mathbf{\Lambda}_\star^{1/2} \right\|_F^2}_{J_5} + 2\eta^2 \underbrace{\left\| \mathbf{N} \mathbf{B} (\mathbf{B}^\top \mathbf{B})^{-1} \mathbf{\Lambda}_\star^{1/2} \right\|_F^2}_{J_6}
\end{aligned} \tag{17}$$

$\|\mathbf{A}\|_2 \|\mathbf{X}\|_F$ and the following result

$$\begin{aligned}
\|\mathbf{\Delta}_L\|_F &= \left\| \mathbf{\Delta}_B \mathbf{B}_\star^\top + \mathbf{B}_\star \mathbf{\Delta}_B^\top + \mathbf{\Delta}_B \mathbf{\Delta}_B^\top \right\|_F \\
&\leq \left\| \mathbf{\Delta}_B \mathbf{B}_\star^\top \right\|_F + \left\| \mathbf{B}_\star \mathbf{\Delta}_B^\top \right\|_F + \left\| \mathbf{\Delta}_B \mathbf{\Delta}_B^\top \right\|_F \\
&= 2 \left\| \mathbf{\Delta}_B \mathbf{\Lambda}_\star^{1/2} \right\|_F + \left\| \mathbf{\Delta}_B \mathbf{\Lambda}_\star^{1/2} \left(\mathbf{\Delta}_B \mathbf{\Lambda}_\star^{-1/2} \right)^\top \right\|_F \\
&\leq 2 \left\| \mathbf{\Delta}_B \mathbf{\Lambda}_\star^{1/2} \right\|_F + \left\| \mathbf{\Delta}_B \mathbf{\Lambda}_\star^{-1/2} \right\|_2 \left\| \mathbf{\Delta}_B \mathbf{\Lambda}_\star^{1/2} \right\|_F \\
&\leq (2+\delta) \left\| \mathbf{\Delta}_B \mathbf{\Lambda}_\star^{1/2} \right\|_F,
\end{aligned} \tag{21}$$

and the last inequality is due to (20).

The term J_4 can be bounded similarly to J_3 . We invoke (12) in Lemma 13 to obtain

$$|J_4| \leq (2 \|\mathbf{\Delta}_L\|_F + (1 + \sqrt{\pi+1}) \sqrt{s} \|\mathbf{N}\|_\infty) \|\mathbf{Z}\|_F.$$

Since

$$\begin{aligned}
\|\mathbf{Z}\|_F &= \left\| \mathbf{B}_\star \mathbf{\Delta}_B^\top \mathbf{B} (\mathbf{B}^\top \mathbf{B})^{-1} \mathbf{\Lambda}_\star (\mathbf{B}^\top \mathbf{B})^{-1} \mathbf{B}^\top \right\|_F \\
&\leq \left\| \mathbf{\Delta}_B \mathbf{\Lambda}_\star^{1/2} \right\|_F \left\| \mathbf{B} (\mathbf{B}^\top \mathbf{B})^{-1} \mathbf{\Lambda}_\star^{1/2} \right\|_2^2,
\end{aligned}$$

we have

$$\begin{aligned}
|J_4| &\leq (2 \|\mathbf{\Delta}_L\|_F + (1 + \sqrt{\pi+1}) \sqrt{s} \|\mathbf{N}\|_\infty) \\
&\quad \cdot \left\| \mathbf{\Delta}_B \mathbf{\Lambda}_\star^{1/2} \right\|_F \left\| \mathbf{B} (\mathbf{B}^\top \mathbf{B})^{-1} \mathbf{\Lambda}_\star^{1/2} \right\|_2^2 \\
&\leq \frac{2(2+\delta)}{(1-\delta)^2} \left\| \mathbf{\Delta}_B \mathbf{\Lambda}_\star^{1/2} \right\|_F^2 \\
&\quad + \frac{(1 + \sqrt{\pi+1}) \sqrt{s}}{(1-\delta)^2} \|\mathbf{N}\|_\infty \left\| \mathbf{\Delta}_B \mathbf{\Lambda}_\star^{1/2} \right\|_F,
\end{aligned}$$

where the last inequality is due to (20) and (21).

For the last term J_5 , utilize the variational representation of the Frobenius norm to see

$$\sqrt{J_5} = \text{tr} \left((\mathbf{\Psi} - \mathbf{\Psi}_\star) \mathbf{B} (\mathbf{B}^\top \mathbf{B})^{-1} \mathbf{\Lambda}_\star^{1/2} \tilde{\mathbf{B}}^\top \right)$$

for some $\tilde{\mathbf{B}} \in \mathbb{R}^{d \times r}$ obeying $\|\tilde{\mathbf{B}}\|_F = 1$. Set $\mathbf{Z} = \tilde{\mathbf{B}} \mathbf{\Lambda}_\star^{1/2} (\mathbf{B}^\top \mathbf{B})^{-1} \mathbf{B}^\top$ and invoke (12) in Lemma 13 to arrive at

$$\begin{aligned}
\sqrt{J_5} &\leq (2 \|\mathbf{\Delta}_L\|_F + (1 + \sqrt{\pi+1}) \sqrt{s} \|\mathbf{N}\|_\infty) \\
&\quad \cdot \left\| \tilde{\mathbf{B}} \mathbf{\Lambda}_\star^{1/2} (\mathbf{B}^\top \mathbf{B})^{-1} \mathbf{B}^\top \right\|_F \\
&\leq (2 \|\mathbf{\Delta}_L\|_F + (1 + \sqrt{\pi+1}) \sqrt{s} \|\mathbf{N}\|_\infty) \\
&\quad \cdot \left\| \mathbf{B} (\mathbf{B}^\top \mathbf{B})^{-1} \mathbf{\Lambda}_\star^{1/2} \right\|_2 \\
&\leq \frac{2(2+\delta)}{1-\delta} \left\| \mathbf{\Delta}_B \mathbf{\Lambda}_\star^{1/2} \right\|_F + \frac{(1 + \sqrt{\pi+1}) \sqrt{s}}{1-\delta} \|\mathbf{N}\|_\infty,
\end{aligned} \tag{22}$$

where the last inequality is due to (20) and (21). Square both sides of (22) and apply the elementary inequality $(a+b)^2 \leq 2a^2 + 2b^2$ to reach

$$J_5 \leq \frac{8(2+\delta)^2}{(1-\delta)^2} \left\| \mathbf{\Delta}_B \mathbf{\Lambda}_\star^{1/2} \right\|_F^2 + \frac{2s(1 + \sqrt{\pi+1})^2}{(1-\delta)^2} \|\mathbf{N}\|_\infty^2.$$

For J_6 , we have

$$\begin{aligned}
J_6 &\leq \|\mathbf{N}\|_2^2 \left\| \mathbf{B} (\mathbf{B}^\top \mathbf{B})^{-1} \mathbf{\Lambda}_\star^{1/2} \right\|_F^2 \\
&\leq \|\mathbf{N}\|_2^2 r \left\| \mathbf{B} (\mathbf{B}^\top \mathbf{B})^{-1} \mathbf{\Lambda}_\star^{1/2} \right\|_2^2 \leq \frac{r}{(1-\delta)^2} \|\mathbf{N}\|_2^2,
\end{aligned}$$

where the last inequality is due to (20). Taking collectively the bounds for J_1 , J_2 , J_3 , J_4 , J_5 , and J_6 , yields (23). Then identify $\text{dist}^2(\mathbf{B}_t, \mathbf{B}_\star) = \left\| \mathbf{\Delta}_B \mathbf{\Lambda}_\star^{1/2} \right\|_F^2$ and set $0 < \eta \leq \frac{1-\delta}{2-\delta}$ to reach

$$\text{dist}^2(\mathbf{B}_{t+1}, \mathbf{B}_\star) \leq \rho \text{dist}^2(\mathbf{B}_t, \mathbf{B}_\star) + \omega, \tag{24}$$

$$\begin{aligned}
\text{dist}^2(\mathbf{B}_{t+1}, \mathbf{B}_\star) &\leq 2(1-\eta)^2 \text{tr}(\mathbf{\Delta}_B \mathbf{A}_\star \mathbf{\Delta}_B^\top) - 4\eta(1-\eta) \frac{\delta}{1-\delta} \text{tr}(\mathbf{\Delta}_B \mathbf{A}_\star \mathbf{\Delta}_B^\top) - 8\eta(1-\eta) \frac{(2+\delta)}{1-\delta} \|\mathbf{\Delta}_B \mathbf{A}_\star^{1/2}\|_F^2 \\
&\quad - 4\eta(1-\eta) \frac{(1+\sqrt{\pi+1})\sqrt{s}}{1-\delta} \|\mathbf{N}\|_\infty \|\mathbf{\Delta}_B \mathbf{A}_\star^{1/2}\|_F + 8\eta^2 \frac{(2+\delta)}{(1-\delta)^2} \|\mathbf{\Delta}_B \mathbf{A}_\star^{1/2}\|_F^2 \\
&\quad + 4\eta^2 \frac{(1+\sqrt{\pi+1})\sqrt{s}}{(1-\delta)^2} \|\mathbf{N}\|_\infty \|\mathbf{\Delta}_B \mathbf{A}_\star^{1/2}\|_F + \frac{16\eta^2(2+\delta)^2}{(1-\delta)^2} \|\mathbf{\Delta}_B \mathbf{A}_\star^{1/2}\|_F^2 \\
&\quad + \frac{4\eta^2 s (1+\sqrt{\pi+1})^2}{(1-\delta)^2} \|\mathbf{N}\|_\infty^2 + \frac{2\eta^2 r}{(1-\delta)^2} \|\mathbf{N}\|_2^2 \\
&\leq \left(2(1-\eta)^2 - 4\eta(1-\eta) \frac{\delta}{1-\delta} - 8\eta(1-\eta) \frac{(2+\delta)}{1-\delta} + 8\eta^2 \frac{(2+\delta)}{(1-\delta)^2} + \frac{16\eta^2(2+\delta)^2}{(1-\delta)^2} \right) \|\mathbf{\Delta}_B \mathbf{A}_\star^{1/2}\|_F^2 \\
&\quad + \left(\frac{4\eta^2}{1-\delta} - 4\eta(1-\eta) \right) \frac{(1+\sqrt{\pi+1})\sqrt{s}}{1-\delta} \|\mathbf{N}\|_\infty \|\mathbf{\Delta}_B \mathbf{A}_\star^{1/2}\|_F \\
&\quad + \frac{4\eta^2 s (1+\sqrt{\pi+1})^2}{(1-\delta)^2} \|\mathbf{N}\|_\infty^2 + \frac{2\eta^2 r}{(1-\delta)^2} \|\mathbf{N}\|_2^2
\end{aligned} \tag{23}$$

where $\omega = \frac{4\eta^2 s (1+\sqrt{\pi+1})^2}{(1-\delta)^2} \|\mathbf{N}\|_\infty^2 + \frac{2\eta^2 r}{(1-\delta)^2} \|\mathbf{N}\|_2^2$, and the contraction rate ρ is given by

$$\begin{aligned}
\rho &= 2(1-\eta)^2 - 4\eta(1-\eta) \frac{\delta}{1-\delta} \\
&\quad - 8\eta(1-\eta) \frac{(2+\delta)}{1-\delta} + 8\eta^2 \frac{(2+\delta)}{(1-\delta)^2} + \frac{16\eta^2(2+\delta)^2}{(1-\delta)^2} \\
&= 2(1-\eta)^2 - 4\eta(1-\eta) \frac{\delta + 2(2+\delta)}{1-\delta} \\
&\quad + 8\eta^2 \frac{(2+\delta) + 2(2+\delta)^2}{(1-\delta)^2} \\
&\leq 2(1-10\eta)^2.
\end{aligned}$$

With $\frac{2-\sqrt{2}}{20} < \eta < \frac{1}{2}$, one has $\rho = 2(1-10\eta)^2 < 1$.

Suppose $\text{dist}(\mathbf{B}_t, \mathbf{B}_\star) < \sigma_r(\mathbf{L}_\star)$, then we have

$$\begin{aligned}
\text{dist}^2(\mathbf{B}_{t+1}, \mathbf{B}_\star) &\leq \rho \delta^2 \sigma_r^2(\mathbf{L}_\star) + (1-\rho) \delta^2 \sigma_r^2(\mathbf{L}_\star) \\
&= \delta^2 \sigma_r^2(\mathbf{L}_\star),
\end{aligned}$$

where in the first inequality we have used the fact that when the sample size n is sufficiently large, we can always ensure $\omega \leq (1-\rho) \delta^2 \sigma_r^2(\mathbf{L}_\star)$. Therefore, we have

$$\text{dist}(\mathbf{B}_{t+1}, \mathbf{B}_\star) \leq \delta \sigma_r(\mathbf{L}_\star). \tag{25}$$

By mathematical induction, we then have $\text{dist}(\mathbf{B}_t, \mathbf{B}_\star) \leq \delta \sigma_r(\mathbf{L}_\star)$ in (25), for any $t = 0, 1, \dots$.

Since (24) holds uniformly for $t = 0, 1, \dots$, we further obtain that

$$\begin{aligned}
\text{dist}^2(\mathbf{B}_{t+1}, \mathbf{B}_\star) &\leq \rho \text{dist}^2(\mathbf{B}_t, \mathbf{B}_\star) + \omega \\
&\leq \rho^{t+1} \sigma_r^2(\mathbf{L}_\star) + \frac{\omega}{1-\rho}.
\end{aligned}$$

C. Proof of Theorem 4

Proof: We have

$$\begin{aligned}
\|\mathbf{\Sigma}_0 - \mathbf{\Sigma}_\star\|_F &= \|\mathbf{B}_0 \mathbf{B}_0^\top + \mathbf{\Psi}_0 - \mathbf{L}_\star - \mathbf{\Psi}_\star\|_F \\
&\leq \|\mathbf{B}_0 \mathbf{B}_0^\top - \mathbf{L}_\star\|_F + \|\mathbf{\Psi}_0 - \mathbf{\Psi}_\star\|_F \\
&\leq \sqrt{2r} \|\mathbf{B}_0 \mathbf{B}_0^\top - \mathbf{L}_\star\|_2 + \sqrt{(1+\pi)s} \|\mathbf{\Psi}_0 - \mathbf{\Psi}_\star\|_\infty,
\end{aligned}$$

where the last inequality uses the fact that $\mathbf{B}_0 \mathbf{B}_0^\top - \mathbf{L}_\star$ has rank at most $2r$, and $\mathbf{\Psi}_0 - \mathbf{\Psi}_\star$ has nonzero elements at most $(1+\pi)s$. Based on (11) in Lemma 13, we have

$$\|\mathbf{\Psi}_0 - \mathbf{\Psi}_\star\|_\infty \leq 2 \|\mathbf{B}_0 \mathbf{B}_0^\top - \mathbf{L}_\star\|_\infty + 2 \|\mathbf{N}\|_\infty.$$

Therefore, it follows that

$$\begin{aligned}
\|\mathbf{\Sigma}_0 - \mathbf{\Sigma}_\star\|_F &\leq \sqrt{2r} \|\mathbf{B}_0 \mathbf{B}_0^\top - \mathbf{L}_\star\|_2 + 2\sqrt{(1+\pi)s} \|\mathbf{B}_0 \mathbf{B}_0^\top - \mathbf{L}_\star\|_\infty \\
&\quad + 2\sqrt{(1+\pi)s} \|\mathbf{N}\|_\infty \\
&\leq 2 \left(\sqrt{2r} + 2\sqrt{2(1+\pi)rs} \right) (\|\mathbf{N}\|_2 + \|\mathbf{\Sigma}_\star\|_2 + \|\mathbf{L}_\star\|_2) \\
&\quad + 2\sqrt{(1+\pi)s} \|\mathbf{N}\|_\infty \\
&\leq 2 \left(\sqrt{2r} + 2\sqrt{2(1+\pi)rs} \right) \left(\alpha \sqrt{\frac{d}{n}} + \mu + k \right) \\
&\quad + 2\beta \sqrt{(1+\pi)s} \sqrt{\frac{\log d}{n}},
\end{aligned}$$

where the second inequality is due to (14), and the last inequality is due to Lemma 8. Using the inequality $(a+b)^2 \leq 2a^2 + 2b^2$, we can conclude that with probability at least $\min\{1 - 2\exp(-C_1 d), 1 - 2\exp(-C_2 \log d)\}$,

$$\|\mathbf{\Sigma}_0 - \mathbf{\Sigma}_\star\|_F^2 \leq c_{1,4} \frac{rd}{n} + c_{2,4} \frac{s \log d}{n} + c_3 \frac{rsd}{n} + c_4 rs + c_5 r.$$

- [20] Y. Feng, D. P. Palomar *et al.*, "A signal processing perspective on financial engineering," *Found. Trends® Signal Process.*, vol. 9, no. 1–2, pp. 1–231, 2016.
- [21] S. C. Ludvigson and S. Ng, "The empirical risk–return relation: A factor analysis approach," *J. Financ. Econ.*, vol. 83, no. 1, pp. 171–222, 2007.
- [22] J. T. Leek and J. D. Storey, "Capturing heterogeneity in gene expression studies by surrogate variable analysis," *PLoS Genetics*, vol. 3, no. 9, p. e161, 2007.
- [23] R. Schmidt, "Multiple emitter location and signal parameter estimation," *IEEE Trans. Antennas Propag.*, vol. 34, no. 3, pp. 276–280, 1986.
- [24] R. Roy and T. Kailath, "ESPRIT-Estimation of signal parameters via rotational invariance techniques," *IEEE Trans. Acoust., Speech, Signal Process.*, vol. 37, no. 7, pp. 984–995, 1989.
- [25] L. Van Der Maaten, E. O. Postma, H. J. Van Den Herik *et al.*, "Dimensionality reduction: A comparative review," *J. Mach. Learn. Res.*, vol. 10, no. 66–71, p. 13, 2009.
- [26] S. A. Ross, "The arbitrage theory of capital asset pricing," *J. Econ. Theory*, vol. 13, no. 3, pp. 341–360, 1976.
- [27] R. Roll and S. A. Ross, "An empirical investigation of the arbitrage pricing theory," *The Journal of Finance*, vol. 35, no. 5, pp. 1073–1103, 1980.
- [28] J. Bai, "Inferential theory for factor models of large dimensions," *Econometrica*, vol. 71, no. 1, pp. 135–171, 2003.
- [29] D. Passemier and J.-F. Yao, "Variance estimation and goodness-of-fit test in a high-dimensional strict factor model," *Submitted to Statist. Sinica, arXiv*, vol. 1308, 2013.
- [30] J. Fan, Y. Fan, and J. Lv, "High dimensional covariance matrix estimation using a factor model," *J. Econom.*, vol. 147, no. 1, pp. 186–197, 2008.
- [31] B. Liao, S.-C. Chan, L. Huang, and C. Guo, "Iterative methods for subspace and DOA estimation in nonuniform noise," *IEEE Trans. Signal Process.*, vol. 64, no. 12, pp. 3008–3020, 2016.
- [32] G. Chamberlain and M. Rothschild, "Arbitrage, factor structure, and mean-variance analysis on large asset markets," *Econometrica: J. Econom.*, pp. 1281–1304, 1983.
- [33] J. Boivin and S. Ng, "Are more data always better for factor analysis?" *J. Econom.*, vol. 132, no. 1, pp. 169–194, 2006.
- [34] A. J. Rothman, E. Levina, and J. Zhu, "Generalized thresholding of large covariance matrices," *J. Am. Stat. Assoc.*, vol. 104, no. 485, pp. 177–186, 2009.
- [35] M. Agrawal and S. Prasad, "A modified likelihood function approach to DOA estimation in the presence of unknown spatially correlated Gaussian noise using a uniform linear array," *IEEE Trans. Signal Process.*, vol. 48, no. 10, pp. 2743–2749, 2002.
- [36] M. Malek-Mohammadi, M. Jansson, A. Owrang, A. Koochakzadeh, and M. Babaie-Zadeh, "DOA estimation in partially correlated noise using low-rank/sparse matrix decomposition," in *Proc. IEEE Workshop Sensor Array Multichannel Signal Process.*, 2014, pp. 373–376.
- [37] S. A. Vorobyov, A. B. Gershman, and K. M. Wong, "Maximum likelihood direction-of-arrival estimation in unknown noise fields using sparse sensor arrays," *IEEE Trans. Signal Process.*, vol. 53, no. 1, pp. 34–43, 2004.
- [38] J. Bai and S. Ng, "Determining the number of factors in approximate factor models," *Econometrica*, vol. 70, no. 1, pp. 191–221, 2002.
- [39] X. Luo, "High dimensional low rank and sparse covariance matrix estimation via convex minimization," *arXiv preprint arXiv:1111.1133*, vol. 199, 2011.
- [40] J. Fan, Y. Liao, and M. Mincheva, "High dimensional covariance matrix estimation in approximate factor models," *Annals of statistics*, vol. 39, no. 6, p. 3320, 2011.
- [41] —, "Large covariance estimation by thresholding principal orthogonal complements," *J. R. Stat. Soc. B (Stat. Methodol.)*, vol. 75, no. 4, 2013.
- [42] X. Luo, "Recovering model structures from large low rank and sparse covariance matrix estimation," *arXiv preprint arXiv:1111.1133*, 2013.
- [43] M. Farnè and A. Montanari, "A large covariance matrix estimator under intermediate spikiness regimes," *J. Multivariate Anal.*, vol. 176, p. 104577, 2020.
- [44] —, "Large factor model estimation by nuclear norm plus ℓ_1 norm penalization," *J. Multivariate Anal.*, vol. 199, p. 105244, 2024.
- [45] A. Agarwal, S. Negahban, and M. J. Wainwright, "Noisy matrix decomposition via convex relaxation: Optimal rates in high dimensions," *Ann. Statist.*, vol. 40, no. 2, 2012.
- [46] E. Yang and P. K. Ravikumar, "Dirty statistical models," in *Proc. Adv. Neural Inf. Process. Syst.*, vol. 26, 2013.
- [47] K. G. Jöreskog, "Some contributions to maximum likelihood factor analysis," *Psychometrika*, vol. 32, no. 4, pp. 443–482, 1967.
- [48] D. B. Rubin and D. T. Thayer, "EM algorithms for ML factor analysis," *Psychometrika*, vol. 47, pp. 69–76, 1982.
- [49] J. Bai and K. Li, "Statistical analysis of factor models of high dimension," *Ann. Statist.*, vol. 40, no. 1, p. 436, 2012.
- [50] P. Stoica and A. Nehorai, "Performance study of conditional and unconditional direction-of-arrival estimation," *IEEE Trans. Acoust., Speech, Signal Process.*, vol. 38, no. 10, pp. 1783–1795, 1990.
- [51] —, "MUSIC, maximum likelihood, and Cramer-Rao bound," *IEEE Trans. Acoust., Speech, Signal Process.*, vol. 37, no. 5, pp. 720–741, 2002.
- [52] P. Stoica and P. Babu, "Low-rank covariance matrix estimation for factor analysis in anisotropic noise: Application to array processing and portfolio selection," *IEEE Trans. Signal Process.*, vol. 71, pp. 1699–1711, 2023.
- [53] R. Vershynin, "Introduction to the non-asymptotic analysis of random matrices," *arXiv preprint arXiv:1011.3027*, 2010.
- [54] A. A. Amini and M. J. Wainwright, "High-dimensional analysis of semidefinite relaxations for sparse principal components," *Ann. Statist.*, vol. 37, no. 5B, 2009.
- [55] T. W. Anderson, T. W. Anderson, T. W. Anderson, T. W. Anderson, and E.-U. Mathématicien, *An introduction to multivariate statistical analysis*. Wiley New York, 1958.
- [56] E. J. Candès, X. Li, Y. Ma, and J. Wright, "Robust principal component analysis," *J. ACM*, vol. 58, no. 3, pp. 1–37, 2011.
- [57] H. Xu, C. Caramanis, and S. Sanghavi, "Robust PCA via outlier pursuit," in *Adv. Neural Inf. Process. Syst.*, vol. 23, 2010.
- [58] Z. Lin, M. Chen, and Y. Ma, "The augmented Lagrange multiplier method for exact recovery of corrupted low-rank matrices," *J. Struct. Biol.*, vol. 181, no. 2, pp. 116–127, 2013.
- [59] P. Netrapalli, N. U N, S. Sanghavi, A. Anandkumar, and P. Jain, "Non-convex robust PCA," in *Adv. Neural Inf. Process. Syst.*, vol. 27, 2014.
- [60] A. Hippert-Ferrer, M. N. El Korso, A. Breloy, and G. Ginolhac, "Robust low-rank covariance matrix estimation with a general pattern of missing values," *Signal Process.*, vol. 195, p. 108460, 2022.
- [61] J. Tanner and K. Wei, "Low rank matrix completion by alternating steepest descent methods," *Appl. Comput. Harmon. A.*, vol. 40, no. 2, pp. 417–429, 2016.
- [62] V. Chandrasekaran, S. Sanghavi, P. A. Parrilo, and A. S. Willsky, "Rank-sparsity incoherence for matrix decomposition," *SIAM Journal on Optimization*, vol. 21, no. 2, pp. 572–596, 2011.
- [63] T. Tong, C. Ma, and Y. Chi, "Accelerating ill-conditioned low-rank matrix estimation via scaled gradient descent," *J. Mach. Learn. Res.*, vol. 22, pp. 150–1, 2021.
- [64] S. Negahban and M. J. Wainwright, "Restricted strong convexity and weighted matrix completion: Optimal bounds with noise," *J. Mach. Learn. Res.*, vol. 13, no. 1, pp. 1665–1697, 2012.
- [65] V. Chandrasekaran, P. A. Parrilo, and A. S. Willsky, "Latent variable graphical model selection via convex optimization," *Ann. Stat.*, pp. 1935–1967, 2012.
- [66] P. Xu, J. Ma, and Q. Gu, "Speeding up latent variable gaussian graphical model estimation via nonconvex optimization," in *Proc. Adv. Neural Inf. Process. Syst.*, vol. 30, 2017.
- [67] Z. Meng, B. Eriksson, and A. Hero, "Learning latent variable gaussian graphical models," in *Proc. Int. Conf. Mach. Learn.* PMLR, 2014, pp. 1269–1277.
- [68] K. Lounici, "High-dimensional covariance matrix estimation with missing observations," *Bernoulli*, pp. 1029–1058, 2014.
- [69] A. R. Zhang, T. T. Cai, and Y. Wu, "Heteroskedastic PCA: Algorithm, optimality, and applications," *Ann. Statist.*, vol. 50, no. 1, pp. 53–80, 2022.
- [70] M. W. McCracken and S. Ng, "FRED-MD: A monthly database for macroeconomic research," *J. Bus. Econ. Stat.*, vol. 34, no. 4, pp. 574–589, 2016.
- [71] J. Fan, J. Guo, and S. Zheng, "Estimating number of factors by adjusted eigenvalues thresholding," *J. Am. Stat. Assoc.*, vol. 117, no. 538, pp. 852–861, 2022.
- [72] H. Markowitz, "Portfolio selection," *J. Finance.*, vol. 7, no. 1, pp. 77–91, 1952.
- [73] M. Grant and S. Boyd, "CVX: Matlab software for disciplined convex programming," 2009.
- [74] S. Tu, R. Boczar, M. Simchowitz, M. Soltanolkotabi, and B. Recht, "Low-rank solutions of linear matrix equations via Procrustes flow," in *Proc. Int. Conf. Mach. Learn.* PMLR, 2016, pp. 964–973.
- [75] R. Ge, C. Jin, and Y. Zheng, "No spurious local minima in nonconvex low rank problems: A unified geometric analysis," in *Proc. Int. Conf. Mach. Learn.* PMLR, 2017, pp. 1233–1242.