

Beyond Procrustes: Balancing-Free Gradient Descent for Asymmetric Low-Rank Matrix Sensing

Cong Ma , Yuanxin Li, and Yuejie Chi , *Senior Member, IEEE*

Abstract—Low-rank matrix estimation plays a central role in various applications across science and engineering. Recently, nonconvex formulations based on matrix factorization are provably solved by simple gradient descent algorithms with strong computational and statistical guarantees. However, when the low-rank matrices are asymmetric, existing approaches rely on adding a regularization term to balance the scale of the two matrix factors which in practice can be removed safely without hurting the performance when initialized via the spectral method. In this paper, we provide a theoretical justification to this for the matrix sensing problem, which aims to recover a low-rank matrix from a small number of linear measurements. As long as the measurement ensemble satisfies the restricted isometry property, gradient descent—in conjunction with spectral initialization—converges linearly without the need of explicitly promoting balancedness of the factors; in fact, the factors stay balanced automatically throughout the execution of the algorithm. Our analysis is based on analyzing the evolution of a new distance metric that directly accounts for the ambiguity due to invertible transforms, and might be of independent interest.

Index Terms—Asymmetric low-rank matrix sensing, nonconvex optimization, gradient descent.

I. INTRODUCTION

LOW-RANK matrix estimation plays a central role in many applications [2]–[4]. Broadly speaking, we are interested in estimating a rank- r matrix $M_\star = X_\star Y_\star^\top \in \mathbb{R}^{n_1 \times n_2}$ by solving a rank-constrained optimization problem:

$$\min_{M \in \mathbb{R}^{n_1 \times n_2}} \mathcal{L}(M) \quad \text{subject to} \quad \text{rank}(M) \leq r, \quad (1)$$

where $\mathcal{L}(\cdot)$ denotes a certain loss function with the rank r typically much smaller than the dimension of the matrix. To reduce computational complexity, a common approach, popularized by the work of Burer and Monteiro [5]–[7], is to factorize $M = XY^\top$ with $X \in \mathbb{R}^{n_1 \times r}$ and $Y \in \mathbb{R}^{n_2 \times r}$, and rewrite the above problem (1) into an unconstrained nonconvex

optimization problem:

$$\min_{X \in \mathbb{R}^{n_1 \times r}, Y \in \mathbb{R}^{n_2 \times r}} f(X, Y) \triangleq \mathcal{L}(XY^\top). \quad (2)$$

Despite nonconvexity, one might be tempted to estimate the low-rank factors (X, Y) via gradient descent, which proceeds via the following update rule

$$\begin{bmatrix} X_{t+1} \\ Y_{t+1} \end{bmatrix} = \begin{bmatrix} X_t \\ Y_t \end{bmatrix} - \eta_t \begin{bmatrix} \nabla_X f(X_t, Y_t) \\ \nabla_Y f(X_t, Y_t) \end{bmatrix} \quad (3)$$

from (X_0, Y_0) some proper initialization. Here, η_t is the step size, $\nabla_X f$ and $\nabla_Y f$ are the gradients of f w.r.t. X and Y , respectively.

Significant progress has been made recently in understanding the performance of gradient descent for nonconvex matrix estimation. Somewhat surprisingly, most of the existing guarantees are not directly applicable to the vanilla gradient descent rule (3). One particular challenge is associated with the identifiability of the factors (X, Y) —they are indistinguishable as long as their product $XY^\top$ is the same. What is worse, if the norms of the factors become highly unbalanced, gradient descent might diverge easily. Consequently, it becomes a routine procedure to insert a regularizer $g(X, Y)$ that balances the two factors [8]–[10]:

$$g(X, Y) \triangleq \lambda \|X^\top X - Y^\top Y\|_F^2, \quad (4)$$

where $\lambda > 0$ is some regularization parameter, and apply gradient descent to the regularized loss function instead:

$$\min_{X \in \mathbb{R}^{n_1 \times r}, Y \in \mathbb{R}^{n_2 \times r}} f_{\text{reg}}(X, Y) \triangleq f(X, Y) + g(X, Y). \quad (5)$$

For a variety of important problems such as low-rank matrix sensing and matrix completion, it has been established that gradient descent over the regularized loss function, when properly initialized, achieves compelling statistical and computational guarantees.

A. Why Balancing is Needed in Prior Work?

Before we investigate the possibility of a balancing-free procedure (i.e. vanilla gradient descent as in (3)), let us first explain using a heuristic argument why balancing is needed in the prior literature.

To handle the asymmetric factorization, it is common to stack the two factors into one augmented factor $Z_\star \triangleq \begin{bmatrix} X_\star \\ Y_\star \end{bmatrix} \in \mathbb{R}^{(n_1+n_2) \times r}$ and then seek to estimate $Z_\star$ directly, by rewriting the loss function with respect to the lifted low-rank matrix: $Z_\star Z_\star^\top = \begin{bmatrix} X_\star X_\star^\top & X_\star Y_\star^\top \\ Y_\star X_\star^\top & Y_\star Y_\star^\top \end{bmatrix} \in \mathbb{R}^{(n_1+n_2) \times (n_1+n_2)}$. It is obvious that the loss function originally with respect to the asymmetric matrix $X_\star Y_\star^\top$ only constrains the off-diagonal blocks

Manuscript received April 26, 2020; revised December 5, 2020; accepted January 7, 2021. Date of publication January 13, 2021; date of current version February 5, 2021. The associate editor coordinating the review of this manuscript and approving it for publication was Prof. Yue M. Lu. This work was supported in part by ONR under the Grants N00014-18-1-2142 and N00014-19-1-2404, in part by ARO under the Grant W911NF-18-1-0303, and in part by NSF under the Grants CAREER ECCS-1818571, CCF-1806154 and CCF-1901199. A preliminary version of this paper was presented at the 2019 Asilomar Conference on Signals, Systems, and Computers [1]. (Corresponding author: Cong Ma.)

Cong Ma is with the Department of Electrical Engineering and Computer Sciences, UC Berkeley, Berkeley, CA 94720 USA (e-mail: congm@berkeley.edu).

Yuanxin Li and Yuejie Chi are with the Department of Electrical and Computer Engineering, Carnegie Mellon University, Pittsburgh, PA 15213 USA (e-mail: li.3822@osu.edu; yuejiechi@cmu.edu).

Digital Object Identifier 10.1109/TSP.2021.3051425

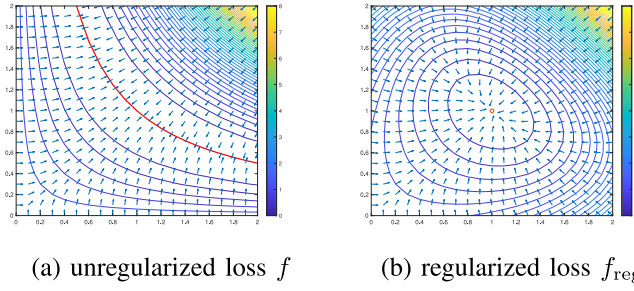


Fig. 1. Geometry for the scalar case $f(x, y) = (xy - 1)^2$ and $g(x, y) = (x^2 - y^2)^2/8$. The regularized loss function is locally strongly convex while the unregularized one is nonconvex; in particular, the Hessian of the unregularized loss function is rank deficient on the ambiguity set $xy = 1$ (colored in red).

of $\mathbf{Z}_* \mathbf{Z}_*^\top$ and not the diagonal ones; correspondingly, the loss function is not (restricted) strongly convex with respect to the augmented factor, unless we appropriately regularize the diagonal blocks. This gives rise to the adoption of the regularization term in (4).

To develop more intuitions regarding why this regularization term (4) may help analysis, consider a toy example of factorizing a rank-one matrix $\mathbf{x}_* \mathbf{y}_*^\top$, where $f(\mathbf{x}, \mathbf{y})$ and $g(\mathbf{x}, \mathbf{y})$ respectively are $f(\mathbf{x}, \mathbf{y}) = \|\mathbf{x} \mathbf{y}^\top - \mathbf{x}_* \mathbf{y}_*^\top\|_F^2/2$ and $g(\mathbf{x}, \mathbf{y}) = (\|\mathbf{x}\|_2^2 - \|\mathbf{y}\|_2^2)^2/8$. Figure 1 illustrates the landscape of the unregularized loss function $f(\mathbf{x}, \mathbf{y})$ and the regularized loss function $f_{\text{reg}}(\mathbf{x}, \mathbf{y})$, respectively, when the arguments are scalar-valued, i.e. $n_1 = n_2 = 1$. One can clearly appreciate the value of the regularizer: $f_{\text{reg}}(\mathbf{x}, \mathbf{y})$ becomes strongly convex in the local neighborhood around the global optimum (1,1). In contrast, the Hessian of the unregularized loss function $f_{\text{reg}}(\mathbf{x}, \mathbf{y})$ remains rank deficient along the ambiguity set $\{(x, y) \mid xy = 1\}$, making the analysis less tractable.

B. This Paper: Balancing-Free Procedure?

The goal of this paper is to understand the effectiveness of vanilla gradient descent (3) when initialized with balanced factors. Indeed, Figure 2 plots the normalized error $\|\mathbf{X}_t \mathbf{Y}_t^\top - \mathbf{M}_*\|_F / \|\mathbf{M}_*\|_F$ for low-rank matrix completion, which aims to recover a low-rank matrix from a subset of its observations [11], with respect to the iteration count, using either a regularized loss function or an unregularized loss function when initialized by the spectral method. The two sequences of iterates converge in almost exactly the same trajectory, suggesting that gradient descent over the unregularized loss function converges almost in the same manner as its regularized counterpart, and perhaps is more natural to use in practice since it eliminates the tuning of the regularization parameters.

This paper justifies formally that even without explicit balancing in asymmetric low-rank matrix sensing, gradient descent converges linearly towards the global optimum, as long as the initialization is (nearly) balanced and close to the optimum. As will be detailed later, our analysis is simple and built on a novel distance metric that directly accounts for the ambiguity due to invertible transformations—in contrast, the ambiguity set reduces to orthonormal transforms when the balancing regularization is present. Our key message is this:

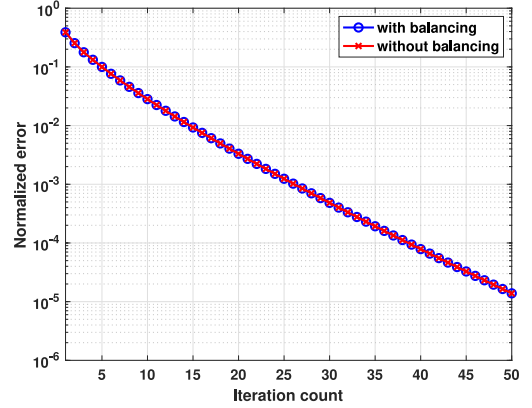


Fig. 2. Normalized reconstruction error $\|\mathbf{X}_t \mathbf{Y}_t^\top - \mathbf{M}_*\|_F / \|\mathbf{M}_*\|_F$ with respect to the iteration count, for completing a 1000×1000 matrix of rank-10 when each entry is observed independently with probability $p = 0.15$. The balancing regularizer is set as $g(\mathbf{X}, \mathbf{Y}) = \frac{1}{64} \|\mathbf{X}^\top \mathbf{X} - \mathbf{Y}^\top \mathbf{Y}\|_F^2$ following the suggestion in [12].

As long as the factors are (nearly) balanced in a basin of attraction at the initialization, they will stay approximately balanced throughout the trajectory of gradient descent, and therefore no additional regularization is necessary.

C. Notation

We use boldface lowercase (resp. uppercase) letters to represent vectors (resp. matrices). We denote by $\|\mathbf{x}\|_2$ the ℓ_2 norm of a vector $\mathbf{x}$, and $\mathbf{X}^\top$, $\mathbf{X}^{-1}$, $\|\mathbf{X}\|$ and $\|\mathbf{X}\|_F$ the transpose, the inverse, the spectral norm and the Frobenius norm of a matrix $\mathbf{X}$, respectively. Furthermore, we denote $\mathbf{X}^{-\top} = (\mathbf{X}^{-1})^\top = (\mathbf{X}^\top)^{-1}$ for an invertible matrix $\mathbf{X}$. The k th largest singular value of a matrix $\mathbf{X}$ is denoted by $\sigma_k(\mathbf{X})$. The inner product between two matrices $\mathbf{X}$ and $\mathbf{Y}$ is defined as $\langle \mathbf{X}, \mathbf{Y} \rangle = \text{Tr}(\mathbf{Y}^\top \mathbf{X})$, where $\text{Tr}(\cdot)$ denotes the trace operator. Denote by $\mathcal{O}^{r \times r}$ the set of $r \times r$ orthonormal matrices. In addition, we use c and C with different subscripts to represent positive numerical constants, whose values may change from line to line.

II. MAIN RESULTS

Let the object of interest $\mathbf{M}_* \in \mathbb{R}^{n_1 \times n_2}$ be a rank- r matrix whose compact Singular Value Decomposition (SVD) is given by

$$\mathbf{M}_* = \mathbf{U}_* \mathbf{\Sigma}_* \mathbf{V}_*^\top,$$

where $\mathbf{U}_* \in \mathbb{R}^{n_1 \times r}$, $\mathbf{V}_* \in \mathbb{R}^{n_2 \times r}$ and $\mathbf{\Sigma}_* \in \mathbb{R}^{r \times r}$ correspond to the left singular vectors, the right singular vectors and the singular values, respectively. Without loss of generality, we denote the ground truth factors as

$$\mathbf{X}_* \triangleq \mathbf{U}_* \mathbf{\Sigma}_*^{1/2}, \quad \text{and} \quad \mathbf{Y}_* \triangleq \mathbf{V}_* \mathbf{\Sigma}_*^{1/2}. \quad (6)$$

Let $\sigma_{\max} \triangleq \sigma_1(\mathbf{M}_*)$ (resp. $\sigma_{\min} \triangleq \sigma_r(\mathbf{M}_*)$) be the largest (resp. smallest) nonzero singular value of $\mathbf{M}_*$. The condition number of $\mathbf{M}_*$ is therefore defined as $\kappa \triangleq \sigma_{\max}/\sigma_{\min}$.

Since the factors are identifiable up to invertible transforms, i.e. $(\mathbf{X}_* \mathbf{P})(\mathbf{Y}_* \mathbf{P}^{-\top})^\top = \mathbf{X}_* \mathbf{Y}_*^\top$ for any invertible matrix $\mathbf{P} \in \mathbb{R}^{r \times r}$, it is natural to measure the distance between two

Algorithm 1: Gradient Descent With Spectral Initialization (Unregularized Procrustes Flow).

Input: Measurements $\mathbf{y} = \{y_i\}_{1 \leq i \leq m}$, and sensing matrices $\{\mathbf{A}_i\}_{1 \leq i \leq m}$.

Parameters: Step size η_t , rank r , and number of iterations T .

Initialization: Initialize $\mathbf{X}_0 = \mathbf{U}\Sigma^{1/2}$ and $\mathbf{Y}_0 = \mathbf{V}\Sigma^{1/2}$, where $\mathbf{U}\Sigma\mathbf{V}^\top$ is the rank- r SVD of the surrogate matrix $\mathbf{K} = \frac{1}{m}\mathcal{A}^*(\mathbf{y}) = \frac{1}{m}\sum_{i=1}^m y_i \mathbf{A}_i$.

Gradient loop: For $t = 0, 1, \dots, T-1$, do

$$\mathbf{X}_{t+1} = \mathbf{X}_t - \frac{\eta_t}{\|\mathbf{Y}_0\|^2} \cdot \left[\sum_{i=1}^m (\langle \mathbf{A}_i, \mathbf{X}_t \mathbf{Y}_t^\top \rangle - y_i) \mathbf{A}_i \mathbf{Y}_t \right]; \quad (9a)$$

$$\mathbf{Y}_{t+1} = \mathbf{Y}_t - \frac{\eta_t}{\|\mathbf{X}_0\|^2} \cdot \left[\sum_{i=1}^m (\langle \mathbf{A}_i, \mathbf{X}_t \mathbf{Y}_t^\top \rangle - y_i) \mathbf{A}_i^\top \mathbf{X}_t \right]. \quad (9b)$$

Output: $\mathbf{X}_T$ and $\mathbf{Y}_T$.

pairs of factors $\mathbf{Z} = \begin{bmatrix} \mathbf{X} \\ \mathbf{Y} \end{bmatrix} \in \mathbb{R}^{(n_1+n_2) \times r}$ and $\mathbf{Z}_\star = \begin{bmatrix} \mathbf{X}_\star \\ \mathbf{Y}_\star \end{bmatrix} \in \mathbb{R}^{(n_1+n_2) \times r}$ via the following function:¹

$$\text{dist}(\mathbf{Z}, \mathbf{Z}_\star) = \min_{\mathbf{P} \in \mathbb{R}^{r \times r} \text{ invertible}} \sqrt{\|\mathbf{X}\mathbf{P} - \mathbf{X}_\star\|_F^2 + \|\mathbf{Y}\mathbf{P}^{-\top} - \mathbf{Y}_\star\|_F^2}.$$

A. Low-Rank Matrix Sensing

Low-rank matrix sensing refers to the problem of recovering a low-rank matrix (i.e. $\mathbf{M}_\star$) from a small number of linear measurements. Specifically, we are given a set of m measurements as follows

$$y_i = \langle \mathbf{A}_i, \mathbf{M}_\star \rangle = \langle \mathbf{A}_i, \mathbf{X}_\star \mathbf{Y}_\star^\top \rangle, \quad i = 1, \dots, m, \quad (7)$$

where $\mathbf{A}_i \in \mathbb{R}^{n_1 \times n_2}$ is the i th sensing matrix. For convenience, we define $\mathcal{A} : \mathbb{R}^{n_1 \times n_2} \rightarrow \mathbb{R}^m$ as an affine transformation from $\mathbb{R}^{n_1 \times n_2}$ to $\mathbb{R}^m$, such that $\mathcal{A}(\mathbf{M}) = \{\langle \mathbf{A}_i, \mathbf{M} \rangle\}_{1 \leq i \leq m}$. Consequently, one can compactly write (7) as $\mathbf{y} = \mathcal{A}(\mathbf{M}_\star)$. The adjoint operator $\mathcal{A}^* : \mathbb{R}^m \rightarrow \mathbb{R}^{n_1 \times n_2}$ is defined as $\mathcal{A}^*(\mathbf{y}) = \sum_{i=1}^m y_i \mathbf{A}_i$.

To recover the low-rank matrix, a natural choice is to minimize the least-squares loss function

$$f(\mathbf{X}, \mathbf{Y}) \triangleq \frac{1}{2} \|\mathbf{y} - \mathcal{A}(\mathbf{X}\mathbf{Y}^\top)\|_2^2. \quad (8)$$

Algorithm 1 describes the gradient descent algorithm initialized by the spectral method [13] for minimizing (8). Compared to the Procrustes Flow (PF) algorithm in [8], which minimizes the regularized loss function in (5), the new algorithm does not include the balancing regularizer $g(\mathbf{X}, \mathbf{Y})$.

¹More rigorously, we should write $\inf$ instead of $\min$ in the definition of $\text{dist}(\cdot, \cdot)$. However, as we will soon see, in the cases we care about, the minimum can always be achieved by some invertible matrix $\mathbf{P}$.

B. Theoretical Guarantee for Local Linear Convergence

To understand the performance of Algorithm 1, we adopt a standard assumption on the sensing operator $\mathcal{A}$, namely the Restricted Isometry Property (RIP).

Definition 1 (RIP): The operator $\mathcal{A}(\cdot)$ is said to satisfy the rank- r RIP with a constant $\delta_r \in [0, 1)$, if

$$(1 - \delta_r) \|\mathbf{M}\|_F^2 \leq \|\mathcal{A}(\mathbf{M})\|_2^2 \leq (1 + \delta_r) \|\mathbf{M}\|_F^2$$

holds for all matrices $\mathbf{M} \in \mathbb{R}^{n_1 \times n_2}$ of rank at most r .

It is well-known that many measurement ensembles satisfy the RIP property [14]. For example, under the Gaussian design where the entries of $\mathbf{A}_i$'s are composed of i.i.d. Gaussian entries $\mathcal{N}(0, 1/m)$, the RIP is satisfied as long as m is on the order of $(n_1 + n_2)r/\delta_r^2$.

Armed with the RIP, we have the following theoretical guarantee for the local convergence of Algorithm 1.

Theorem 1: Suppose that $\mathcal{A}(\cdot)$ satisfies the RIP with $\delta_{2r} \leq c$ for some sufficiently small constant c . Let $\mathbf{Z}_0 \triangleq \begin{bmatrix} \mathbf{X}_0 \\ \mathbf{Y}_0 \end{bmatrix}$ be any initialization point that satisfies

$$\min_{\mathbf{R} \in \mathbb{O}^{r \times r}} \|\mathbf{Z}_0 \mathbf{R} - \mathbf{Z}_\star\|_F \leq c_0 \frac{1}{\kappa^{3/2}} \sigma_{\min}^{1/2}, \quad (10)$$

for some small enough constant $c_0 > 0$. Then there exist some constant $c_1 > 0$ such that at as long as $\eta_t = \eta = c_1$, the iterates of unregularized gradient descent (cf. (9)) satisfy

$$\text{dist}(\mathbf{Z}_t, \mathbf{Z}_\star) \leq \left(1 - \frac{\eta}{50\kappa}\right)^t \text{dist}(\mathbf{Z}_0, \mathbf{Z}_\star).$$

In words, Theorem 1 reveals that if the initialization $\mathbf{Z}_0$ lands in a basin of attraction given by (10), then Algorithm 1 converges linearly with a constant step size. To reach ϵ -accuracy, i.e. $\text{dist}(\mathbf{Z}_t, \mathbf{Z}_\star) \leq \epsilon$, it takes an order of $\kappa \log(1/\epsilon)$ iterations, which is order-wise equivalent to the regularized PF algorithm proposed in [8]. Comparing to [8], which requires $\delta_{6r} \leq c$, Theorem 1 only requires a weaker assumption $\delta_{2r} \leq c$. However, the basin of attraction allowed by Theorem 1 is smaller than that in [8], which is specified by $\min_{\mathbf{R} \in \mathbb{O}^{r \times r}} \|\mathbf{Z}_0 \mathbf{R} - \mathbf{Z}_\star\|_F \leq c_0 \sigma_{\min}^{1/2}$. Compared with prior work that relies on local strong convexity to establish linear convergence, our result suggests the benign behavior of gradient descent even in the absence of local strong convexity.

C. Achieving Global Convergence With a Proper Initialization

We are still in need of finding a good initialization that obeys (10). In general, one could initialize with the balanced factors of the output of projected gradient descent (over the low-rank matrix), i.e.

$$\mathbf{M}_{\tau+1} = \mathcal{P}_r \left(\mathbf{M}_\tau - \frac{1}{m} \sum_{i=1}^m (\langle \mathbf{A}_i, \mathbf{M}_\tau \rangle - y_i) \mathbf{A}_i \right),$$

where $\mathcal{P}_r(\cdot)$ is the Euclidean projection operator to the set of rank- r matrices. The spectral initialization specified in Algorithm 1 can be regarded as the output at the first iteration, initialized at zero $\mathbf{M}_0 = \mathbf{0}$. Based on [8], [15], the balanced factorization of $\mathbf{M}_\tau$, denoted by $\tilde{\mathbf{Z}}_\tau$, satisfy

$$\min_{\mathbf{R} \in \mathbb{O}^{r \times r}} \|\tilde{\mathbf{Z}}_\tau \mathbf{R} - \mathbf{Z}_\star\|_F \leq c_2 (2\delta_{4r})^\tau \frac{\|\mathbf{M}_\star\|_F}{\sigma_{\min}^{1/2}} \quad (11)$$

for some constant c_2 . Thus, to achieve the required initialization condition (10) using the spectral method specified in Algorithm 1 (which corresponds to setting $\mathbf{Z}_0 = \tilde{\mathbf{Z}}_1$ with $\tau = 1$ in (11)), we need

$$\delta_{4r} \leq c_2 \frac{1}{\kappa^{3/2}} \cdot \frac{\sigma_{\min}}{\|\mathbf{M}_\star\|_F}.$$

In particular, under Gaussian design, where each measurement matrix $\mathbf{A}_i$ has i.i.d. $\mathcal{N}(0, 1/m)$ entries, a total of $m \gtrsim nr\kappa^3 \|\mathbf{M}_\star\|_F^2 / \sigma_{\min}^2$ measurements suffice for the above requirement on δ_{4r} . This is worse by a factor of κ^3 compared with the sample complexity guarantee in [8] with the balancing regularizer, which is due to the restriction on the basin of attraction, as we have remarked earlier. Improving the dependence on κ is an interesting future direction.

In order to alleviate the dependency on the condition number κ , we can allow a few iterations of (11) and set the initialization as $\mathbf{Z}_0 = \tilde{\mathbf{Z}}_\tau$, a procedure suggested by [8]. The advantage of this hybrid procedure is that the switch to factored gradient descent allows a smaller per-iteration memory and computation complexity after the iterates of projected gradient descent enter the basin of attraction. Consequently, the algorithm is still guaranteed to succeed when $\delta_{4r} \leq \delta_c$ for a sufficiently small constant δ_c (which implies a near-optimal sample complexity of $m \gtrsim nr$ under Gaussian design), by running at least

$$\tau \geq c_1 \log \left(\kappa^{3/2} \frac{\|\mathbf{M}_\star\|_F}{\sigma_{\min}} \right) / \log(\delta_c^{-1})$$

iterations of projected gradient descent for initialization, which order-wise matches the requirement in [8].

III. RELATED WORK

Low-rank matrix estimation has been extensively studied in recent years [3], [4], due to its broad applicability in collaborative filtering, imaging science, and machine learning, to name a few. Convex relaxation approaches based on nuclear norm minimization are among the first set of algorithms with near-optimal statistical guarantees, e. g. [2], [14], [16]–[22], however, their computational costs are often prohibitive in practice.

To cope with the computational challenges, a popular approach in practice is to invoke low-rank matrix factorization and then apply first-order methods such as gradient descent directly over the factors to recover the underlying low-rank structure. This approach is demonstrated to possess near-optimal statistical and computational guarantees in a variety of low-rank matrix recovery problems, including but not limited to [8], [23]–[31]. The readers are referred to the recent overview [32] for additional references.

To the best of our knowledge, the balancing regularization term (4) was first introduced in [8] to deal with asymmetric matrix factorization, and has become a standard approach to deal with asymmetric low-rank matrix estimation [9], [10], [12], [33]–[35]. A major benefit of adding the regularization term is to reduce the ambiguity set from invertible transforms to orthonormal transforms. For the special rank-one matrix recovery problem, there are some evidence in the prior literature that a balancing regularization is not needed, for example, Ma et al. [27] established that vanilla gradient descent works for blind

deconvolution at a near-optimal sample complexity with spectral initialization. In [36], the trajectory of gradient descent is studied for asymmetric matrix factorization with an infinitesimal and diminishing step size; in contrast, we consider the case when the step size is constant for low-rank matrix estimation with incomplete observations. Finally, very recently, [37] also studied low-rank matrix sensing using a nonsmooth formulation without the balancing regularization via subgradient descent.

Complementary to the algorithmic analysis, we remark that a similar regularization term (4) is also adopted when analyzing the optimization landscape of low-rank matrix estimation, e.g. [38]–[42]. It is worth mentioning that when converting a nuclear-norm regularized problem into a nonconvex formulation, [43] demonstrated that the nonconvex problem has benign geometry without adding the balancing regularization, since the nuclear norm regularization induces a term $\frac{1}{2}(\|\mathbf{X}\|_F^2 + \|\mathbf{Y}\|_F^2)$ which ensures both factors have similar sizes. Very recently, [44] showed that the balancing regularizer is unnecessary from the landscape analysis perspective.

After the initial version of the current paper, several other works have further examined the balancing-free low-rank matrix optimization problem. In particular, Tian, Ma and Chi developed a scaled gradient descent algorithm [45] that achieves a faster convergence rate independent of the condition number κ without imposing the balancing regularization for a variety of low-rank matrix estimation problems, which are further extended in [46] to achieve robustness to adversarial outliers.

IV. PROOF OF THEOREM 1

In this section, we provide the proof of Theorem 1. We first discuss some basic properties of aligning two low-rank factors via an invertible transformation. Then we prove a similar result for a warm-up case of low-rank matrix factorization. In the end, viewing matrix sensing as a perturbed version of low-rank matrix factorization helps us finish the proof of Theorem 1.

A. Alignment Via Invertible Transformations

We begin with introducing the alignment matrix will play a key role in the subsequent analysis.

Definition 2: Fix a matrix $\mathbf{Z} = \begin{bmatrix} \mathbf{X} \\ \mathbf{Y} \end{bmatrix} \in \mathbb{R}^{(n_1+n_2) \times r}$. We define the optimal alignment matrix $\mathbf{Q}$ between $\mathbf{Z}$ and $\mathbf{Z}_\star$ as

$$\mathbf{Q} \triangleq \underset{\substack{\mathbf{P} \in \mathbb{R}^{r \times r} \\ \text{invertible}}}{\operatorname{argmin}} \|\mathbf{X}\mathbf{P} - \mathbf{X}_\star\|_F^2 + \|\mathbf{Y}\mathbf{P}^\top - \mathbf{Y}_\star\|_F^2,$$

whenever the minimum is attained.

As we will soon see, for the iterates $\{\mathbf{Z}_t\}_{t \geq 0}$ generated by Algorithm 1, the optimal alignment matrix is always well-defined. Furthermore, we call $\mathbf{Z}$ and $\mathbf{Z}_\star$ aligned if the corresponding optimal alignment matrix is just the identity matrix $\mathbf{I}_r$. Below we provide some basic understandings of this alignment matrix.

The following lemma provides a sufficient condition for the existence of the optimal alignment matrix.

Lemma 1: Fix some matrix $\mathbf{Z} = \begin{bmatrix} \mathbf{X} \\ \mathbf{Y} \end{bmatrix} \in \mathbb{R}^{(n_1+n_2) \times r}$. Suppose that there exists a matrix $\mathbf{P} \in \mathbb{R}^{r \times r}$ with $1/2 \leq \sigma_r(\mathbf{P}) \leq$

$\sigma_1(P) \leq 3/2$ such that

$$\max \{ \|XP - X_\star\|_F, \|YP^{-\top} - Y_\star\|_F \} \leq \delta \leq \frac{\sigma_r(X_\star)}{80}. \quad (12)$$

Then the optimal alignment matrix $Q \in \mathbb{R}^{r \times r}$ between Z and $Z_\star$ exists. In addition, the matrix Q satisfies

$$\|P - Q\| \leq \|P - Q\|_F \leq \frac{5\delta}{\sigma_r(X_\star)}.$$

Next, the lemma below presents a necessary condition for Q to be the optimal alignment matrix between Z and $Z_\star$.

Lemma 2: Let Z and $Z_\star$ be any two matrices. Suppose that the optimal alignment matrix Q between Z and $Z_\star$ exists. Then we have

$$\widetilde{X}^\top (\widetilde{X} - X_\star) = (\widetilde{Y} - Y_\star)^\top \widetilde{Y},$$

where $\widetilde{X} = XQ$ and $\widetilde{Y} = YQ^{-\top}$ are two matrices after the alignment.

Both lemmas provide basic understandings of the solution to the alignment problem with invertible transformations, which can be regarded as a generalization of the classical orthogonal Procrustes problem that only considers orthonormal transformations. Clearly, this generalized problem is more involved and our work provides some basic understandings.

B. A Warm-Up: Low-Rank Matrix Factorization

We consider the following minimization problem for low-rank matrix factorization

$$f_{\text{MF}}(X, Y) = \frac{1}{2} \|XY^\top - M_\star\|_F^2, \quad (13)$$

where $X \in \mathbb{R}^{n_1 \times r}$ and $Y \in \mathbb{R}^{n_2 \times r}$. The gradient descent updates with an initialization (X_0, Y_0) can be written as

$$\begin{aligned} X_{t+1} &= X_t - \frac{\eta}{\sigma_{\max}} \nabla_X f_{\text{MF}}(X_t, Y_t) \\ &= X_t - \frac{\eta}{\sigma_{\max}} (X_t Y_t^\top - M_\star) Y_t; \\ Y_{t+1} &= Y_t - \frac{\eta}{\sigma_{\max}} \nabla_Y f_{\text{MF}}(X_t, Y_t) \\ &= Y_t - \frac{\eta}{\sigma_{\max}} (X_t Y_t^\top - M_\star)^\top X_t. \end{aligned} \quad (14)$$

Here, $\eta > 0$ stands for the step size. We have the following theorem regarding the performance of (14), which parallels Theorem 1.

Theorem 2: Let $Z_0 = \begin{bmatrix} X_0 \\ Y_0 \end{bmatrix} \in \mathbb{R}^{(n_1+n_2) \times r}$ be any initialization point that satisfies

$$\min_{R \in \mathbb{O}^{r \times r}} \|Z_0 R - Z_\star\|_F \leq c_0 \frac{1}{\kappa^{3/2}} \sigma_{\min}^{1/2} \quad (15)$$

for some sufficiently small constant $c_0 > 0$. Then setting the step size $\eta > 0$ to be some sufficiently small constant, the iterates of GD (cf. (14)) satisfy

$$\text{dist}(Z_t, Z_\star) \leq \left(1 - \frac{\eta}{50\kappa}\right)^t \text{dist}(Z_0, Z_\star).$$

To prove Theorem 2, we need the following properties regarding the gradients of $f_{\text{MF}}(X, Y)$; the proofs are deferred to the appendix.

Lemma 3 (Gradient dominance): Suppose that $Z = \begin{bmatrix} X \\ Y \end{bmatrix} \in \mathbb{R}^{(n_1+n_2) \times r}$ is aligned with $Z_\star$, i.e.

$$I_r = \underset{\substack{P \in \mathbb{R}^{r \times r} \\ \text{invertible}}}{\text{argmin}} \|XP - X_\star\|_F^2 + \|YP^{-\top} - Y_\star\|_F^2.$$

Then we have

$$\begin{aligned} &\langle X - X_\star, (XY^\top - M_\star) Y \rangle \\ &\geq \|Y(X - X_\star)^\top\|_F^2 - \frac{1}{4} \|X - X_\star\|_F^4, \end{aligned}$$

and similarly,

$$\begin{aligned} &\langle Y - Y_\star, (XY^\top - M_\star)^\top X \rangle \\ &\geq \|X(Y - Y_\star)^\top\|_F^2 - \frac{1}{4} \|Y - Y_\star\|_F^4. \end{aligned}$$

Lemma 4 (Smoothness): Suppose that $\|Y - Y_\star\| \leq \sigma_1(Y_\star)/4$, then one has

$$\begin{aligned} &\|(XY^\top - M_\star) Y\|_F \\ &\leq \frac{3}{2} \sigma_1(Y_\star) (\|(X - X_\star)Y^\top\|_F + \|X(Y - Y_\star)^\top\|_F \\ &\quad + \|X - X_\star\|_F \|Y - Y_\star\|_F). \end{aligned}$$

Similarly, with the proviso that $\|X - X_\star\| \leq \sigma_1(X_\star)/4$, one has

$$\begin{aligned} &\|(XY^\top - M_\star)^\top X\|_F \\ &\leq \frac{3}{2} \sigma_1(X_\star) (\|(X - X_\star)Y^\top\|_F + \|X(Y - Y_\star)^\top\|_F \\ &\quad + \|X - X_\star\|_F \|Y - Y_\star\|_F). \end{aligned}$$

C. Proof of Theorem 2

With the help of Lemmas 1–4, we are in a position to establish Theorem 2. Denote by $\widehat{R} \in \mathbb{R}^{r \times r}$ the best rotation matrix between Z_0 and $Z_\star$, that is

$$\widehat{R} \triangleq \underset{R \in \mathbb{O}^{r \times r}}{\text{argmin}} \|Z_0 R - Z_\star\|_F.$$

Combine the assumption of initialization (cf. (15)) and Lemma 1 to see that

$$Q_0 \triangleq \underset{\substack{P \in \mathbb{R}^{r \times r} \\ \text{invertible}}}{\text{argmin}} \|X_0 P - X_\star\|_F^2 + \|Y_0 P^{-\top} - Y_\star\|_F^2$$

exists and in addition, one has

$$\|Q_0 - \widehat{R}\| \leq \frac{5c_0}{\kappa^{3/2}} \leq \frac{1}{400\sqrt{\kappa}}$$

as long as $c_0 > 0$ is sufficiently small.

The remaining proof is inductive in nature. In particular, we aim at proving the following induction hypotheses.

- 1) The optimal alignment matrix Q_t between Z_t and $Z_\star$ exists.
- 2) The distance between Z_t and $Z_\star$ obeys

$$\text{dist}(Z_t, Z_\star) \leq \left(1 - \frac{\eta}{50\kappa}\right)^t \text{dist}(Z_0, Z_\star).$$

- 3) The optimal alignment matrix $\mathbf{Q}_t$ is nearly a rotation matrix in the sense that

$$\|\mathbf{Q}_t - \hat{\mathbf{R}}\| \leq \frac{1}{400\sqrt{\kappa}}.$$

It is straightforward to check that these three claims hold for $t = 0$. In what follows, we shall assume that the induction hypotheses hold for all iterations up to the t th iteration and intend to establish that they continue to hold for the $(t + 1)$ th iteration.

a) *Verifying the first induction hypothesis:* We begin with demonstrating the existence of $\mathbf{Q}_{t+1}$. In view of the gradient update rule (14), we have

$$\begin{aligned} \mathbf{X}_{t+1}\mathbf{Q}_t &= \mathbf{X}_t\mathbf{Q}_t - \frac{\eta}{\sigma_{\max}} (\mathbf{X}_t\mathbf{Y}_t^\top - \mathbf{M}_\star) \mathbf{Y}_t\mathbf{Q}_t \\ &= \tilde{\mathbf{X}}_t - \frac{\eta}{\sigma_{\max}} (\tilde{\mathbf{X}}_t\tilde{\mathbf{Y}}_t^\top - \mathbf{M}_\star) \tilde{\mathbf{Y}}_t (\mathbf{Q}_t^\top\mathbf{Q}_t), \\ \mathbf{Y}_{t+1}\mathbf{Q}_t^\top &= \mathbf{Y}_t\mathbf{Q}_t^\top - \frac{\eta}{\sigma_{\max}} (\mathbf{X}_t\mathbf{Y}_t^\top - \mathbf{M}_\star)^\top \mathbf{X}_t\mathbf{Q}_t^\top \\ &= \tilde{\mathbf{Y}}_t - \frac{\eta}{\sigma_{\max}} (\mathbf{X}_t\mathbf{Y}_t^\top - \mathbf{M}_\star)^\top \tilde{\mathbf{X}}_t (\mathbf{Q}_t^\top\mathbf{Q}_t)^{-1}, \end{aligned}$$

where we denote

$$\tilde{\mathbf{X}}_t \triangleq \mathbf{X}_t\mathbf{Q}_t \quad \text{and} \quad \tilde{\mathbf{Y}}_t \triangleq \mathbf{Y}_t\mathbf{Q}_t^\top.$$

As a result, one has the following equality

$$\begin{aligned} \|\mathbf{X}_{t+1}\mathbf{Q}_t - \mathbf{X}_\star\|_F^2 + \|\mathbf{Y}_{t+1}\mathbf{Q}_t^\top - \mathbf{Y}_\star\|_F^2 &= \underbrace{\left\| \tilde{\mathbf{X}}_t - \mathbf{X}_\star - \frac{\eta}{\sigma_{\max}} (\tilde{\mathbf{X}}_t\tilde{\mathbf{Y}}_t^\top - \mathbf{M}_\star) \tilde{\mathbf{Y}}_t\mathbf{\Lambda}_t \right\|_F^2}_{:=\alpha_1} \\ &\quad + \underbrace{\left\| \tilde{\mathbf{Y}}_t - \mathbf{Y}_\star - \frac{\eta}{\sigma_{\max}} (\tilde{\mathbf{X}}_t\tilde{\mathbf{Y}}_t^\top - \mathbf{M}_\star)^\top \tilde{\mathbf{X}}_t\mathbf{\Lambda}_t^{-1} \right\|_F^2}_{:=\alpha_2}, \end{aligned}$$

where we have denoted $\mathbf{\Lambda}_t \triangleq \mathbf{Q}_t^\top\mathbf{Q}_t$. By virtue of the third induction hypothesis, namely $\|\mathbf{Q}_t - \hat{\mathbf{R}}\| \leq 1/(400\sqrt{\kappa})$, it is easy to check that $\|\mathbf{\Lambda}_t - \mathbf{I}_r\| \leq 1/(180\sqrt{\kappa}) \triangleq \zeta$. Let

$$\tilde{\mathbf{E}}_{\mathbf{X}_t} \triangleq \tilde{\mathbf{X}}_t - \mathbf{X}_\star \quad \text{and} \quad \tilde{\mathbf{E}}_{\mathbf{Y}_t} \triangleq \tilde{\mathbf{Y}}_t - \mathbf{Y}_\star.$$

Expand α_1 to obtain

$$\begin{aligned} \alpha_1 &= \|\tilde{\mathbf{E}}_{\mathbf{X}_t}\|_F^2 + \left(\frac{\eta}{\sigma_{\max}}\right)^2 \underbrace{\left\| (\tilde{\mathbf{X}}_t\tilde{\mathbf{Y}}_t^\top - \mathbf{M}_\star) \tilde{\mathbf{Y}}_t\mathbf{\Lambda}_t \right\|_F^2}_{:=\beta_1} \\ &\quad - 2\frac{\eta}{\sigma_{\max}} \underbrace{\left\langle \tilde{\mathbf{E}}_{\mathbf{X}_t}, (\tilde{\mathbf{X}}_t\tilde{\mathbf{Y}}_t^\top - \mathbf{M}_\star) \tilde{\mathbf{Y}}_t\mathbf{\Lambda}_t \right\rangle}_{:=\gamma_1}. \end{aligned}$$

Similarly, we can decompose α_2 into

$$\begin{aligned} \alpha_2 &= \|\tilde{\mathbf{E}}_{\mathbf{Y}_t}\|_F^2 + \left(\frac{\eta}{\sigma_{\max}}\right)^2 \underbrace{\left\| (\tilde{\mathbf{X}}_t\tilde{\mathbf{Y}}_t^\top - \mathbf{M}_\star)^\top \tilde{\mathbf{X}}_t\mathbf{\Lambda}_t^{-1} \right\|_F^2}_{:=\beta_2} \\ &\quad - 2\frac{\eta}{\sigma_{\max}} \underbrace{\left\langle \tilde{\mathbf{E}}_{\mathbf{Y}_t}, (\tilde{\mathbf{X}}_t\tilde{\mathbf{Y}}_t^\top - \mathbf{M}_\star)^\top \tilde{\mathbf{X}}_t\mathbf{\Lambda}_t^{-1} \right\rangle}_{:=\gamma_2}. \end{aligned}$$

We intend to apply Lemma 3 to lower bound the terms γ_1 and γ_2 and apply Lemma 4 to upper bound β_1 and β_2 . First, since

$(\tilde{\mathbf{X}}_t, \tilde{\mathbf{Y}}_t)$ is aligned with $(\mathbf{X}_\star, \mathbf{Y}_\star)$, we can invoke Lemma 3 to see that

$$\begin{aligned} \gamma_1 &\geq \left\langle \tilde{\mathbf{E}}_{\mathbf{X}_t}, (\tilde{\mathbf{X}}_t\tilde{\mathbf{Y}}_t^\top - \mathbf{M}_\star) \tilde{\mathbf{Y}}_t \right\rangle \\ &\quad - \left| \left\langle \tilde{\mathbf{E}}_{\mathbf{X}_t}, (\tilde{\mathbf{X}}_t\tilde{\mathbf{Y}}_t^\top - \mathbf{M}_\star) \tilde{\mathbf{Y}}_t (\mathbf{\Lambda}_t - \mathbf{I}_r) \right\rangle \right| \\ &\geq \|\tilde{\mathbf{Y}}_t\tilde{\mathbf{E}}_{\mathbf{X}_t}^\top\|_F^2 - \frac{1}{4}\|\tilde{\mathbf{E}}_{\mathbf{X}_t}\|_F^4 \\ &\quad - \|\mathbf{\Lambda}_t - \mathbf{I}_r\| \cdot \left\| (\tilde{\mathbf{X}}_t\tilde{\mathbf{Y}}_t^\top - \mathbf{M}_\star) \tilde{\mathbf{Y}}_t \right\|_F \|\tilde{\mathbf{E}}_{\mathbf{X}_t}\|_F \\ &\geq \|\tilde{\mathbf{Y}}_t\tilde{\mathbf{E}}_{\mathbf{X}_t}^\top\|_F^2 - \frac{1}{400}\sigma_{\min}\|\tilde{\mathbf{E}}_{\mathbf{X}_t}\|_F^2 \\ &\quad - \zeta \left\| (\tilde{\mathbf{X}}_t\tilde{\mathbf{Y}}_t^\top - \mathbf{M}_\star) \tilde{\mathbf{Y}}_t \right\|_F \|\tilde{\mathbf{E}}_{\mathbf{X}_t}\|_F. \end{aligned} \quad (16)$$

Here the last line follows from the bound $\|\mathbf{\Lambda}_t - \mathbf{I}_r\| \leq \zeta$ and the second induction hypothesis, i.e.

$$\|\tilde{\mathbf{E}}_{\mathbf{X}_t}\|_F^2 \leq \text{dist}^2(\mathbf{Z}_t, \mathbf{Z}_\star) \leq \text{dist}^2(\mathbf{Z}_0, \mathbf{Z}_\star) \leq \frac{1}{100}\sigma_{\min}.$$

The last term in (16) can be further bounded via Lemma 4 as

$$\begin{aligned} \zeta \left\| (\tilde{\mathbf{X}}_t\tilde{\mathbf{Y}}_t^\top - \mathbf{M}_\star) \tilde{\mathbf{Y}}_t \right\|_F \|\tilde{\mathbf{E}}_{\mathbf{X}_t}\|_F &\leq \frac{3\zeta}{2}\sqrt{\sigma_{\max}} \\ &\quad (\|\tilde{\mathbf{E}}_{\mathbf{X}_t}\tilde{\mathbf{Y}}_t^\top\|_F + \|\tilde{\mathbf{X}}_t\tilde{\mathbf{E}}_{\mathbf{Y}_t}^\top\|_F \\ &\quad + \|\tilde{\mathbf{E}}_{\mathbf{X}_t}\|_F\|\tilde{\mathbf{E}}_{\mathbf{Y}_t}\|_F) \|\tilde{\mathbf{E}}_{\mathbf{X}_t}\|_F \\ &= \frac{9\zeta\sqrt{\kappa}}{2}\|\tilde{\mathbf{E}}_{\mathbf{X}_t}\tilde{\mathbf{Y}}_t^\top\|_F \cdot \frac{\sqrt{\sigma_{\min}}}{3}\|\tilde{\mathbf{E}}_{\mathbf{X}_t}\|_F \\ &\quad + \frac{9\zeta\sqrt{\kappa}}{2}\|\tilde{\mathbf{X}}_t\tilde{\mathbf{E}}_{\mathbf{Y}_t}^\top\|_F \cdot \frac{\sqrt{\sigma_{\min}}}{3}\|\tilde{\mathbf{E}}_{\mathbf{X}_t}\|_F \\ &\quad + \frac{3\zeta}{2}\sqrt{\sigma_{\max}}\|\tilde{\mathbf{E}}_{\mathbf{Y}_t}\|_F\|\tilde{\mathbf{E}}_{\mathbf{X}_t}\|_F^2 \\ &\leq \frac{81\zeta^2\kappa}{8}\|\tilde{\mathbf{E}}_{\mathbf{X}_t}\tilde{\mathbf{Y}}_t^\top\|_F^2 + \frac{81\zeta^2\kappa}{8}\|\tilde{\mathbf{X}}_t\tilde{\mathbf{E}}_{\mathbf{Y}_t}^\top\|_F^2 \\ &\quad + \frac{\sigma_{\min}}{8}\|\tilde{\mathbf{E}}_{\mathbf{X}_t}\|_F^2, \end{aligned}$$

where the last inequality arises since $ab \leq (a^2 + b^2)/2$ and

$$\begin{aligned} \frac{3\zeta}{2}\sqrt{\sigma_{\max}}\|\tilde{\mathbf{E}}_{\mathbf{Y}_t}\|_F &\leq \frac{3\zeta}{2}\sqrt{\sigma_{\max}}\text{dist}(\mathbf{Z}_0, \mathbf{Z}_\star) \\ &\leq \frac{3\zeta}{2}\sqrt{\sigma_{\max}}c_0\frac{1}{\kappa^{3/2}}\sqrt{\sigma_{\min}} \leq \frac{\sigma_{\min}}{72} \end{aligned}$$

as long as c_0 is sufficiently small. Combine the above two bounds to reach

$$\begin{aligned} \gamma_1 &\geq \left(1 - \frac{81\zeta^2\kappa}{8}\right) \|\tilde{\mathbf{Y}}_t\tilde{\mathbf{E}}_{\mathbf{X}_t}^\top\|_F^2 \\ &\quad - \frac{81\zeta^2\kappa}{8}\|\tilde{\mathbf{X}}_t\tilde{\mathbf{E}}_{\mathbf{Y}_t}^\top\|_F^2 - \frac{\sigma_{\min}}{7}\|\tilde{\mathbf{E}}_{\mathbf{X}_t}\|_F^2. \end{aligned}$$

Similarly, γ_2 can be lower bounded as

$$\begin{aligned} \gamma_2 &\geq \left(1 - \frac{81\zeta^2\kappa}{8}\right) \|\tilde{\mathbf{X}}_t\tilde{\mathbf{E}}_{\mathbf{Y}_t}^\top\|_F^2 \\ &\quad - \frac{81\zeta^2\kappa}{8}\|\tilde{\mathbf{Y}}_t\tilde{\mathbf{E}}_{\mathbf{X}_t}^\top\|_F^2 - \frac{\sigma_{\min}}{7}\|\tilde{\mathbf{E}}_{\mathbf{Y}_t}\|_F^2, \end{aligned}$$

which together with the bound on γ_1 implies

$$\begin{aligned} & \gamma_1 + \gamma_2 \\ & \geq \left(1 - \frac{81\zeta^2\kappa}{4}\right) \left(\|\tilde{\mathbf{Y}}_t \tilde{\mathbf{E}}_{\mathbf{X}_t}^\top\|_F^2 + \|\tilde{\mathbf{X}}_t \tilde{\mathbf{E}}_{\mathbf{Y}_t}^\top\|_F^2\right) \\ & \quad - \frac{\sigma_{\min}}{7} \text{dist}^2(\mathbf{Z}_t, \mathbf{Z}_\star) \\ & \geq \frac{3}{4} \left(\|\tilde{\mathbf{Y}}_t \tilde{\mathbf{E}}_{\mathbf{X}_t}^\top\|_F^2 + \|\tilde{\mathbf{X}}_t \tilde{\mathbf{E}}_{\mathbf{Y}_t}^\top\|_F^2\right) - \frac{\sigma_{\min}}{7} \text{dist}^2(\mathbf{Z}_t, \mathbf{Z}_\star), \end{aligned}$$

where we plug in the definition of $\zeta = 1/(180\sqrt{\kappa})$.

Now we move on to controlling β_1 and β_2 . Recognizing that $\|\Lambda_t\| \leq 2$, one has

$$\begin{aligned} \beta_1 & \leq 4 \left\| \left(\tilde{\mathbf{X}}_t \tilde{\mathbf{Y}}_t^\top - \mathbf{M}_\star \right) \tilde{\mathbf{Y}}_t \right\|_F^2 \\ & \leq 9\sigma_{\max} \left(\|\tilde{\mathbf{E}}_{\mathbf{X}_t} \tilde{\mathbf{Y}}_t^\top\|_F + \|\tilde{\mathbf{X}}_t \tilde{\mathbf{E}}_{\mathbf{Y}_t}^\top\|_F \right. \\ & \quad \left. + \|\tilde{\mathbf{E}}_{\mathbf{X}_t}\|_F \|\tilde{\mathbf{E}}_{\mathbf{Y}_t}\|_F \right)^2, \end{aligned} \quad (17)$$

where the second line follows from Lemma 4. Apply the elementary inequality $(a + b + c)^2 \leq 3(a^2 + b^2 + c^2)$ to see that

$$\begin{aligned} \beta_1 & \leq 27\sigma_{\max} \left(\|\tilde{\mathbf{E}}_{\mathbf{X}_t} \tilde{\mathbf{Y}}_t^\top\|_F^2 + \|\tilde{\mathbf{X}}_t \tilde{\mathbf{E}}_{\mathbf{Y}_t}^\top\|_F^2 \right. \\ & \quad \left. + \|\tilde{\mathbf{E}}_{\mathbf{X}_t}\|_F^2 \|\tilde{\mathbf{E}}_{\mathbf{Y}_t}\|_F^2 \right) \\ & \leq 27\sigma_{\max} \left(\|\tilde{\mathbf{E}}_{\mathbf{X}_t} \tilde{\mathbf{Y}}_t^\top\|_F^2 + \|\tilde{\mathbf{X}}_t \tilde{\mathbf{E}}_{\mathbf{Y}_t}^\top\|_F^2 \right) \\ & \quad + 27c_0^2 \frac{\sigma_{\max}\sigma_{\min}}{\kappa^3} \|\tilde{\mathbf{E}}_{\mathbf{X}_t}\|_F^2. \end{aligned}$$

Here the second line relies on the fact that $\|\tilde{\mathbf{E}}_{\mathbf{Y}_t}\|_F^2 \leq \text{dist}^2(\mathbf{Z}_0, \mathbf{Z}_\star) \leq c_0^2 \sigma_{\min}/\kappa^3$. Similarly, one can bound β_2 as

$$\begin{aligned} \beta_2 & \leq 27\sigma_{\max} \left(\|\tilde{\mathbf{E}}_{\mathbf{X}_t} \tilde{\mathbf{Y}}_t^\top\|_F^2 + \|\tilde{\mathbf{X}}_t \tilde{\mathbf{E}}_{\mathbf{Y}_t}^\top\|_F^2 \right) \\ & \quad + 27c_0^2 \frac{\sigma_{\max}\sigma_{\min}}{\kappa^3} \|\tilde{\mathbf{E}}_{\mathbf{Y}_t}\|_F^2, \end{aligned}$$

which in conjunction with the bound on β_1 yields

$$\begin{aligned} \beta_1 + \beta_2 & \leq 54\sigma_{\max} \left(\|\tilde{\mathbf{E}}_{\mathbf{X}_t} \tilde{\mathbf{Y}}_t^\top\|_F^2 + \|\tilde{\mathbf{X}}_t \tilde{\mathbf{E}}_{\mathbf{Y}_t}^\top\|_F^2 \right) \\ & \quad + 27c_0^2 \frac{\sigma_{\max}\sigma_{\min}}{\kappa^3} \text{dist}^2(\mathbf{Z}_t, \mathbf{Z}_\star). \end{aligned}$$

Collect all the bounds on α_1 and α_2 to arrive at

$$\begin{aligned} & \|\mathbf{X}_{t+1} \mathbf{Q}_t - \mathbf{X}_\star\|_F^2 + \|\mathbf{Y}_{t+1} \mathbf{Q}_t^\top - \mathbf{Y}_\star\|_F^2 \\ & \leq \left(1 + \frac{27c_0^2\eta^2}{\kappa^4}\right) \text{dist}^2(\mathbf{Z}_t, \mathbf{Z}_\star) \\ & \quad + \left(\frac{54\eta^2}{\sigma_{\max}}\right) \left(\|\tilde{\mathbf{E}}_{\mathbf{X}_t} \tilde{\mathbf{Y}}_t^\top\|_F^2 + \|\tilde{\mathbf{X}}_t \tilde{\mathbf{E}}_{\mathbf{Y}_t}^\top\|_F^2\right) \\ & \quad - 2\frac{\eta}{\sigma_{\max}} \left[\frac{3}{4} \left(\|\tilde{\mathbf{E}}_{\mathbf{X}_t} \tilde{\mathbf{Y}}_t^\top\|_F^2 + \|\tilde{\mathbf{X}}_t \tilde{\mathbf{E}}_{\mathbf{Y}_t}^\top\|_F^2\right) \right. \\ & \quad \left. - \frac{\sigma_{\min}}{7} \text{dist}^2(\mathbf{Z}_t, \mathbf{Z}_\star)\right] \\ & = \left(1 + \frac{27c_0^2\eta^2}{\kappa^4} + \frac{\eta}{3.5\kappa}\right) \text{dist}^2(\mathbf{Z}_t, \mathbf{Z}_\star) \end{aligned}$$

$$\begin{aligned} & + \left(\frac{54\eta^2}{\sigma_{\max}} - \frac{3\eta}{2\sigma_{\max}}\right) \cdot \left(\|\tilde{\mathbf{E}}_{\mathbf{X}_t} \tilde{\mathbf{Y}}_t^\top\|_F^2 + \|\tilde{\mathbf{X}}_t \tilde{\mathbf{E}}_{\mathbf{Y}_t}^\top\|_F^2\right) \\ & \leq \left(1 + \frac{\eta}{3\kappa}\right) \text{dist}^2(\mathbf{Z}_t, \mathbf{Z}_\star) \\ & \quad - \frac{3\eta}{4\sigma_{\max}} \left(\sigma_r^2(\tilde{\mathbf{Y}}_t) \|\tilde{\mathbf{E}}_{\mathbf{X}_t}\|_F^2 + \sigma_r^2(\tilde{\mathbf{X}}_t) \|\tilde{\mathbf{E}}_{\mathbf{Y}_t}\|_F^2\right), \end{aligned}$$

where the last line follows as long as $\eta \leq 1/24$. Furthermore, since $\sigma_r^2(\tilde{\mathbf{Y}}_t) \geq \sigma_{\min}/2$ and $\sigma_r^2(\tilde{\mathbf{X}}_t) \geq \sigma_{\min}/2$, we have

$$\begin{aligned} & \|\mathbf{X}_{t+1} \mathbf{Q}_t - \mathbf{X}_\star\|_F^2 + \|\mathbf{Y}_{t+1} \mathbf{Q}_t^\top - \mathbf{Y}_\star\|_F^2 \\ & \leq \left(1 - \frac{\eta}{24\kappa}\right) \text{dist}^2(\mathbf{Z}_t, \mathbf{Z}_\star). \end{aligned} \quad (18)$$

Lemma 1 then ensures the existence of $\mathbf{Q}_{t+1}$.

b) Verifying the second induction hypothesis. The second induction hypothesis for the $(t+1)$ th iteration follows immediately from the above proof. Since $\mathbf{Q}_{t+1}$ exists, by definition, one has

$$\begin{aligned} & \text{dist}(\mathbf{Z}_{t+1}, \mathbf{Z}_\star) \\ & = \sqrt{\|\mathbf{X}_{t+1} \mathbf{Q}_{t+1} - \mathbf{X}_\star\|_F^2 + \|\mathbf{Y}_{t+1} \mathbf{Q}_{t+1}^\top - \mathbf{Y}_\star\|_F^2} \\ & \leq \sqrt{\|\mathbf{X}_{t+1} \mathbf{Q}_t - \mathbf{X}_\star\|_F^2 + \|\mathbf{Y}_{t+1} \mathbf{Q}_t^\top - \mathbf{Y}_\star\|_F^2} \\ & \leq \left(1 - \frac{\eta}{50\kappa}\right) \text{dist}(\mathbf{Z}_t, \mathbf{Z}_\star). \end{aligned}$$

c) Verifying the third induction hypothesis. It remains to show the last induction hypothesis, namely $\|\mathbf{Q}_{t+1} - \hat{\mathbf{R}}\| \leq 1/(400\sqrt{\kappa})$. In view of (18), one has $\max\{\|\mathbf{X}_{t+1} \mathbf{Q}_t - \mathbf{X}_\star\|_F, \|\mathbf{Y}_{t+1} \mathbf{Q}_t^\top - \mathbf{Y}_\star\|_F\} \leq \text{dist}(\mathbf{Z}_t, \mathbf{Z}_\star)$. Invoke Lemma 1 again to arrive at

$$\begin{aligned} \|\mathbf{Q}_{t+1} - \mathbf{Q}_t\| & \leq \frac{5}{\sigma_r(\mathbf{X}_\star)} \text{dist}(\mathbf{Z}_t, \mathbf{Z}_\star) \\ & \leq \frac{5}{\sigma_r(\mathbf{X}_\star)} \left(1 - \frac{\eta}{50\kappa}\right)^t c_0 \frac{1}{\kappa^{3/2}} \sigma_r(\mathbf{X}_\star) \\ & \leq 5c_0 \left(1 - \frac{\eta}{50\kappa}\right)^t \frac{1}{\kappa^{3/2}}. \end{aligned}$$

Hence, by the triangle inequality and the telescoping sum, we obtain

$$\begin{aligned} \|\mathbf{Q}_{t+1} - \hat{\mathbf{R}}\| & \leq \sum_{s=0}^t \|\mathbf{Q}_{s+1} - \mathbf{Q}_s\| + \|\mathbf{Q}_0 - \hat{\mathbf{R}}\| \\ & \leq 5c_0 \sum_{s=0}^t \left(1 - \frac{\eta}{50\kappa}\right)^s \frac{1}{\kappa^{3/2}} + 5c_0 \frac{1}{\kappa^{3/2}} \\ & < 5c_0 \sum_{s=0}^{\infty} \left(1 - \frac{\eta}{50\kappa}\right)^s \frac{1}{\kappa^{3/2}} + 5c_0 \frac{1}{\kappa^{3/2}} \\ & = 5c_0 \frac{50\kappa}{\eta} \frac{1}{\kappa^{3/2}} + 5c_0 \frac{1}{\kappa^{3/2}} \\ & \leq \frac{1}{400\sqrt{\kappa}}, \end{aligned}$$

as long as c_0 is small enough and η is some constant.

Putting everything together, we finish the induction step and the proof is then completed.

D. Analysis for Matrix Sensing

We now extend the techniques used in the proof of Theorem 2 to the matrix sensing case by leveraging the RIP. Suppose that the initialization $\mathbf{Z}_0$ satisfies the condition (10). By a standard argument as in [8], [9], [33],² it is sufficient to consider the following update rule:

$$\begin{aligned}\mathbf{X}_{t+1} &= \mathbf{X}_t - \frac{\eta}{\sigma_{\max}} [\mathcal{A}^* \mathcal{A}(\mathbf{X}_t \mathbf{Y}_t^\top - \mathbf{M}_*)] \mathbf{Y}_t; \\ \mathbf{Y}_{t+1} &= \mathbf{Y}_t - \frac{\eta}{\sigma_{\max}} [\mathcal{A}^* \mathcal{A}(\mathbf{X}_t \mathbf{Y}_t^\top - \mathbf{M}_*)]^\top \mathbf{X}_t.\end{aligned}\quad (19)$$

Compared with the update rule (14) for low-rank matrix factorization, the update rule for matrix sensing differs by the operation of $\mathcal{A}^* \mathcal{A}$ when forming the gradient. Therefore, we expect that GD has similar behaviors as earlier as long as the operator $\mathcal{A}^* \mathcal{A}$ behaves as a near isometry on low-rank matrices. This can be supplied by the following consequence of the RIP.

Lemma 5: Suppose that $\mathcal{A}$ satisfies $2r$ -RIP with a constant δ_{2r} . Then, for all matrices $\mathbf{M}_1$ and $\mathbf{M}_2$ of rank at most r , we have

$$|\langle \mathcal{A}(\mathbf{M}_1), \mathcal{A}(\mathbf{M}_2) \rangle - \langle \mathbf{M}_1, \mathbf{M}_2 \rangle| \leq \delta_{2r} \|\mathbf{M}_1\|_F \|\mathbf{M}_2\|_F.$$

Equivalently, we can write this as

$$|\text{Tr}[(\mathcal{A}^* \mathcal{A} - \mathcal{I})(\mathbf{M}_1) \mathbf{M}_2^\top]| \leq \delta_{2r} \|\mathbf{M}_1\|_F \|\mathbf{M}_2\|_F.$$

A simple consequence is that for any $\mathbf{A} \in \mathbb{R}^{n_2 \times r}$

$$\|(\mathcal{A}^* \mathcal{A} - \mathcal{I})(\mathbf{M}_1) \mathbf{A}\|_F \leq \delta_{2r} \|\mathbf{M}_1\|_F \|\mathbf{A}\|.$$

Similar to before, we denote $\tilde{\mathbf{X}}_t = \mathbf{X}_t \mathbf{Q}_t$ and $\tilde{\mathbf{Y}}_t = \mathbf{Y}_t \mathbf{Q}_t^{-\top}$, which are aligned with $(\mathbf{X}_*, \mathbf{Y}_*)$. With this notation in place, we can rewrite the update rule as

$$\begin{aligned}\mathbf{X}_{t+1} \mathbf{Q}_t &= \tilde{\mathbf{X}}_t - \frac{\eta}{\sigma_{\max}} [\mathcal{A}^* \mathcal{A}(\mathbf{X}_t \mathbf{Y}_t^\top - \mathbf{M}_*)] \tilde{\mathbf{Y}}_t \mathbf{\Lambda}_t, \\ \mathbf{Y}_{t+1} \mathbf{Q}_t^{-\top} &= \tilde{\mathbf{Y}}_t - \frac{\eta}{\sigma_{\max}} [\mathcal{A}^* \mathcal{A}(\mathbf{X}_t \mathbf{Y}_t^\top - \mathbf{M}_*)]^\top \tilde{\mathbf{X}}_t \mathbf{\Lambda}_t^{-1}.\end{aligned}$$

where we recall $\mathbf{\Lambda}_t = \mathbf{Q}_t^\top \mathbf{Q}_t$. By the definition of the distance function, we further obtain

$$\begin{aligned}\text{dist}^2(\mathbf{Z}_{t+1}, \mathbf{Z}_*) &\leq \|\mathbf{X}_{t+1} \mathbf{Q}_t - \mathbf{X}_*\|_F^2 + \|\mathbf{Y}_{t+1} \mathbf{Q}_t^{-\top} - \mathbf{Y}_*\|_F^2 \\ &= \left\| \tilde{\mathbf{X}}_t - \frac{\eta}{\sigma_{\max}} [\mathcal{A}^* \mathcal{A}(\mathbf{X}_t \mathbf{Y}_t^\top - \mathbf{M}_*)] \tilde{\mathbf{Y}}_t \mathbf{\Lambda}_t - \mathbf{X}_* \right\|_F^2 \\ &\quad + \left\| \tilde{\mathbf{Y}}_t - \frac{\eta}{\sigma_{\max}} [\mathcal{A}^* \mathcal{A}(\mathbf{X}_t \mathbf{Y}_t^\top - \mathbf{M}_*)]^\top \tilde{\mathbf{X}}_t \mathbf{\Lambda}_t^{-1} - \mathbf{Y}_* \right\|_F^2 \\ &= \|\tilde{\mathbf{E}}_{\mathbf{X}_t}\|_F^2 + \|\tilde{\mathbf{E}}_{\mathbf{Y}_t}\|_F^2 \\ &\quad + \left(\frac{\eta}{\sigma_{\max}} \right)^2 \left(\underbrace{\left\| [\mathcal{A}^* \mathcal{A}(\mathbf{X}_t \mathbf{Y}_t^\top - \mathbf{M}_*)] \tilde{\mathbf{Y}}_t \mathbf{\Lambda}_t \right\|_F^2}_{:=\tilde{\beta}_1} \right. \\ &\quad \left. + \underbrace{\left\| [\mathcal{A}^* \mathcal{A}(\mathbf{X}_t \mathbf{Y}_t^\top - \mathbf{M}_*)]^\top \tilde{\mathbf{X}}_t \mathbf{\Lambda}_t^{-1} \right\|_F^2}_{:=\tilde{\beta}_2} \right. \\ &\quad \left. - \frac{2\eta}{\sigma_{\max}} \left(\underbrace{\left\langle \tilde{\mathbf{E}}_{\mathbf{X}_t}, [\mathcal{A}^* \mathcal{A}(\mathbf{X}_t \mathbf{Y}_t^\top - \mathbf{M}_*)] \tilde{\mathbf{Y}}_t \mathbf{\Lambda}_t \right\rangle}_{:=\tilde{\gamma}_1} \right. \right. \\ &\quad \left. \left. + \underbrace{\left\langle \tilde{\mathbf{E}}_{\mathbf{Y}_t}, [\mathcal{A}^* \mathcal{A}(\mathbf{X}_t \mathbf{Y}_t^\top - \mathbf{M}_*)]^\top \tilde{\mathbf{X}}_t \mathbf{\Lambda}_t^{-1} \right\rangle}_{:=\tilde{\gamma}_2} \right) \right), \quad (20)\end{aligned}$$

where $\tilde{\mathbf{E}}_{\mathbf{X}_t} \triangleq \tilde{\mathbf{X}}_t - \mathbf{X}_*$ and $\tilde{\mathbf{E}}_{\mathbf{Y}_t} \triangleq \tilde{\mathbf{Y}}_t - \mathbf{Y}_*$. From the high level, the four terms $\tilde{\beta}_1, \tilde{\beta}_2, \tilde{\gamma}_1$ and $\tilde{\gamma}_2$ are the perturbed versions of $\beta_1, \beta_2, \gamma_1$ and γ_2 in Section IV-C, respectively.

For the first term, we have

$$\begin{aligned}&\sqrt{\tilde{\beta}_1} - \sqrt{\beta_1} \\ &\stackrel{(i)}{\leq} \left\| [\mathcal{A}^* \mathcal{A}(\mathbf{X}_t \mathbf{Y}_t^\top - \mathbf{M}_*)] \tilde{\mathbf{Y}}_t \mathbf{\Lambda}_t \right. \\ &\quad \left. - (\mathbf{X}_t \mathbf{Y}_t^\top - \mathbf{M}_*) \tilde{\mathbf{Y}}_t \mathbf{\Lambda}_t \right\|_F \\ &= \left\| (\mathcal{A}^* \mathcal{A} - \mathcal{I})(\mathbf{X}_t \mathbf{Y}_t^\top - \mathbf{M}_*) \tilde{\mathbf{Y}}_t \mathbf{\Lambda}_t \right\|_F \\ &\stackrel{(ii)}{\leq} \left\| (\mathcal{A}^* \mathcal{A} - \mathcal{I}) \tilde{\mathbf{E}}_{\mathbf{X}_t} \tilde{\mathbf{Y}}_t^\top \right\|_F \|\tilde{\mathbf{Y}}_t \mathbf{\Lambda}_t\|_F \\ &\quad + \left\| (\mathcal{A}^* \mathcal{A} - \mathcal{I}) \tilde{\mathbf{X}}_t \tilde{\mathbf{E}}_{\mathbf{Y}_t}^\top \right\|_F \|\tilde{\mathbf{Y}}_t \mathbf{\Lambda}_t\|_F \\ &\quad + \left\| (\mathcal{A}^* \mathcal{A} - \mathcal{I}) \tilde{\mathbf{E}}_{\mathbf{X}_t} \tilde{\mathbf{E}}_{\mathbf{Y}_t}^\top \right\|_F \|\tilde{\mathbf{Y}}_t \mathbf{\Lambda}_t\|_F \\ &\stackrel{(iii)}{\leq} \delta_{2r} \left(\|\tilde{\mathbf{E}}_{\mathbf{X}_t} \tilde{\mathbf{Y}}_t^\top\|_F + \|\tilde{\mathbf{X}}_t \tilde{\mathbf{E}}_{\mathbf{Y}_t}^\top\|_F + \|\tilde{\mathbf{E}}_{\mathbf{X}_t} \tilde{\mathbf{E}}_{\mathbf{Y}_t}^\top\|_F \right) \|\tilde{\mathbf{Y}}_t \mathbf{\Lambda}_t\|_F \\ &\stackrel{(iv)}{\leq} 4\delta_{2r} \sqrt{\sigma_{\max}} \left(\|\tilde{\mathbf{E}}_{\mathbf{X}_t} \tilde{\mathbf{Y}}_t^\top\|_F + \|\tilde{\mathbf{X}}_t \tilde{\mathbf{E}}_{\mathbf{Y}_t}^\top\|_F + \|\tilde{\mathbf{E}}_{\mathbf{X}_t} \tilde{\mathbf{E}}_{\mathbf{Y}_t}^\top\|_F \right). \quad (21)\end{aligned}$$

Here, the first (i) and second (ii) inequalities follow from the triangle inequality. The third one (iii) uses Lemma 5 and the last relation (iv) depends on $\|\mathbf{\Lambda}_t\| \leq 2$ and $\|\tilde{\mathbf{Y}}_t\| \leq 2\sqrt{\sigma_{\max}}$. Comparing (21) with (17) reveals that $\tilde{\beta}_1 - \beta_1$ constitutes a small perturbation to β_1 when δ_{2r} is small. Similar bounds hold for $\sqrt{\tilde{\beta}_2} - \sqrt{\beta_2}$. As a result, when δ_{2r} is sufficiently small, we have

$$\begin{aligned}\tilde{\beta}_1 + \tilde{\beta}_2 &\leq 108\sigma_{\max} \left(\|\tilde{\mathbf{E}}_{\mathbf{X}_t} \tilde{\mathbf{Y}}_t^\top\|_F^2 + \|\tilde{\mathbf{X}}_t \tilde{\mathbf{E}}_{\mathbf{Y}_t}^\top\|_F^2 \right) \\ &\quad + 54c_0^2 \frac{\sigma_{\max} \sigma_{\min}}{\kappa^3} \text{dist}^2(\mathbf{Z}_t, \mathbf{Z}_*).\end{aligned}$$

We now proceed to $\tilde{\gamma}_1$, for which we have

$$\begin{aligned}&|\tilde{\gamma}_1 - \gamma_1| \\ &= \left| \left\langle \tilde{\mathbf{E}}_{\mathbf{X}_t}, [\mathcal{A}^* \mathcal{A}(\mathbf{X}_t \mathbf{Y}_t^\top - \mathbf{M}_*)] \tilde{\mathbf{Y}}_t \mathbf{\Lambda}_t \right\rangle \right.\end{aligned}$$

²Since (i) the initialization $\mathbf{Z}_0$ is close to the ground truth $\mathbf{Z}_*$, (ii) $\mathbf{X}_0$ and $\mathbf{Y}_0$ are balanced, it is obvious that the operator norm $\|\mathbf{X}_0\|^2 = \|\mathbf{Y}_0\|^2$ is orderwise equivalent to $\sigma_{\max}$. Therefore, all the convergence claims on using $\sigma_{\max}$ can be translated to those on using $\|\mathbf{X}_0\|^2$ and $\|\mathbf{Y}_0\|^2$ by adjusting η up to some absolute constant.

$$\begin{aligned}
& - \left\langle \tilde{\mathbf{E}}_{\mathbf{X}_t}, (\mathbf{X}_t \mathbf{Y}_t^\top - \mathbf{M}_*) \tilde{\mathbf{Y}}_t \mathbf{\Lambda}_t \right\rangle \Big| \\
& = \left| \left\langle \tilde{\mathbf{E}}_{\mathbf{X}_t}, [(\mathcal{A}^* \mathcal{A} - \mathcal{I})(\mathbf{X}_t \mathbf{Y}_t^\top - \mathbf{M}_*)] \tilde{\mathbf{Y}}_t \mathbf{\Lambda}_t \right\rangle \right| \\
& = \left| \text{Tr} \left([(\mathcal{A}^* \mathcal{A} - \mathcal{I})(\mathbf{X}_t \mathbf{Y}_t^\top - \mathbf{M}_*)] \tilde{\mathbf{Y}}_t \mathbf{\Lambda}_t \tilde{\mathbf{E}}_{\mathbf{X}_t}^\top \right) \right| \\
& \leq \left| \text{Tr} \left([(\mathcal{A}^* \mathcal{A} - \mathcal{I})(\tilde{\mathbf{E}}_{\mathbf{X}_t} \tilde{\mathbf{Y}}_t^\top)] \tilde{\mathbf{Y}}_t \mathbf{\Lambda}_t \tilde{\mathbf{E}}_{\mathbf{X}_t}^\top \right) \right| \\
& \quad + \left| \text{Tr} \left([(\mathcal{A}^* \mathcal{A} - \mathcal{I})(\tilde{\mathbf{X}}_t \tilde{\mathbf{E}}_{\mathbf{Y}_t}^\top)] \tilde{\mathbf{Y}}_t \mathbf{\Lambda}_t \tilde{\mathbf{E}}_{\mathbf{X}_t}^\top \right) \right| \\
& \quad + \left| \text{Tr} \left([(\mathcal{A}^* \mathcal{A} - \mathcal{I})(\tilde{\mathbf{E}}_{\mathbf{X}_t} \tilde{\mathbf{E}}_{\mathbf{Y}_t}^\top)] \tilde{\mathbf{Y}}_t \mathbf{\Lambda}_t \tilde{\mathbf{E}}_{\mathbf{X}_t}^\top \right) \right| \\
& \leq \delta_{2r} \left(\|\tilde{\mathbf{E}}_{\mathbf{X}_t} \tilde{\mathbf{Y}}_t^\top\|_F + \|\tilde{\mathbf{X}}_t \tilde{\mathbf{E}}_{\mathbf{Y}_t}^\top\|_F + \|\tilde{\mathbf{E}}_{\mathbf{X}_t} \tilde{\mathbf{E}}_{\mathbf{Y}_t}^\top\|_F \right) \cdot \\
& \quad \|\tilde{\mathbf{Y}}_t \mathbf{\Lambda}_t \tilde{\mathbf{E}}_{\mathbf{X}_t}^\top\|_F.
\end{aligned}$$

Here once again, we utilize the triangle inequality and Lemma 5. Noticing that $\|\mathbf{\Lambda}_t - \mathbf{I}\|$ is small, we further have

$$\begin{aligned}
\|\tilde{\mathbf{Y}}_t \mathbf{\Lambda}_t \tilde{\mathbf{E}}_{\mathbf{X}_t}^\top\|_F & \leq \|\tilde{\mathbf{Y}}_t \tilde{\mathbf{E}}_{\mathbf{X}_t}^\top\|_F + \|\tilde{\mathbf{Y}}_t (\mathbf{\Lambda}_t - \mathbf{I}_r) \tilde{\mathbf{E}}_{\mathbf{X}_t}^\top\|_F \\
& \leq \|\tilde{\mathbf{E}}_{\mathbf{X}_t} \tilde{\mathbf{Y}}_t^\top\|_F + 2\sigma_{\max}^{1/2} \zeta \|\tilde{\mathbf{E}}_{\mathbf{X}_t}\|_F,
\end{aligned}$$

where we use $\|\mathbf{\Lambda}_t - \mathbf{I}_r\| \leq \zeta$ and $\|\tilde{\mathbf{Y}}_t\| \leq 2\sqrt{\sigma_{\max}}$. Combine the previous two bounds and apply the basic inequality $2ab \leq a^2 + b^2$ to see

$$\begin{aligned}
& |\tilde{\gamma}_1 - \gamma_1| \\
& \leq \delta_{2r} \|\tilde{\mathbf{E}}_{\mathbf{X}_t} \tilde{\mathbf{Y}}_t^\top\|_F^2 + 2\sigma_{\max}^{1/2} \zeta \delta_{2r} \|\tilde{\mathbf{E}}_{\mathbf{X}_t} \tilde{\mathbf{Y}}_t^\top\|_F \|\tilde{\mathbf{E}}_{\mathbf{X}_t}\|_F \\
& \quad + \delta_{2r} \|\tilde{\mathbf{X}}_t \tilde{\mathbf{E}}_{\mathbf{Y}_t}^\top\|_F \|\tilde{\mathbf{E}}_{\mathbf{X}_t} \tilde{\mathbf{Y}}_t^\top\|_F \\
& \quad + 2\sigma_{\max}^{1/2} \zeta \delta_{2r} \|\tilde{\mathbf{X}}_t \tilde{\mathbf{E}}_{\mathbf{Y}_t}^\top\|_F \|\tilde{\mathbf{E}}_{\mathbf{X}_t}\|_F \\
& \quad + \delta_{2r} \|\tilde{\mathbf{E}}_{\mathbf{X}_t} \tilde{\mathbf{E}}_{\mathbf{Y}_t}^\top\|_F \|\tilde{\mathbf{E}}_{\mathbf{X}_t} \tilde{\mathbf{Y}}_t^\top\|_F \\
& \quad + 2\sigma_{\max}^{1/2} \zeta \delta_{2r} \|\tilde{\mathbf{E}}_{\mathbf{X}_t} \tilde{\mathbf{E}}_{\mathbf{Y}_t}^\top\|_F \|\tilde{\mathbf{E}}_{\mathbf{X}_t}\|_F \\
& \lesssim \delta_{2r} \left(\|\tilde{\mathbf{E}}_{\mathbf{X}_t} \tilde{\mathbf{Y}}_t^\top\|_F^2 + \|\tilde{\mathbf{X}}_t \tilde{\mathbf{E}}_{\mathbf{Y}_t}^\top\|_F^2 + \sigma_{\min} \text{dist}(\mathbf{Z}_t, \mathbf{Z}_*) \right), \\
& \ll \|\tilde{\mathbf{E}}_{\mathbf{X}_t} \tilde{\mathbf{Y}}_t^\top\|_F^2 + \|\tilde{\mathbf{X}}_t \tilde{\mathbf{E}}_{\mathbf{Y}_t}^\top\|_F^2 + \sigma_{\min} \text{dist}(\mathbf{Z}_t, \mathbf{Z}_*)
\end{aligned}$$

as long as δ_{2r} is sufficiently small. The same bound applies to $|\tilde{\gamma}_2 - \gamma_2|$. As a result, as long as δ_{2r} is small enough, $\tilde{\gamma}_1 + \tilde{\gamma}_2$ is lower bounded on the same order as $\gamma_1 + \gamma_2$, say

$$\begin{aligned}
\tilde{\gamma}_1 + \tilde{\gamma}_2 & \geq \frac{1}{2} \left(\|\tilde{\mathbf{Y}}_t \tilde{\mathbf{E}}_{\mathbf{X}_t}^\top\|_F^2 + \|\tilde{\mathbf{X}}_t \tilde{\mathbf{E}}_{\mathbf{Y}_t}^\top\|_F^2 \right) \\
& \quad - \frac{\sigma_{\min}}{6} \text{dist}^2(\mathbf{Z}_t, \mathbf{Z}_*).
\end{aligned}$$

One can then repeat the same arguments for the matrix factorization case to obtain the linear convergence. For the sake of space, we omit it.

V. CONCLUSION

This paper establishes the local linear convergence of gradient descent for asymmetric low-rank matrix sensing without explicit regularization of factor balancedness under the standard RIP assumption, as long as a balanced initialization is provided in the basin of attraction. Coupled with the standard spectral

initialization, this leads to the global convergence guarantee of the balancing-free gradient descent algorithm for asymmetric low-rank matrix sensing. Different from previous work, we analyzed a new error metric that takes into account the ambiguity due to invertible transforms, and showed that it contracts linearly even without local restricted strong convexity. We believe that our technique can be used for other low-rank matrix estimation problems. To conclude, we outline a few future research directions.

- *Low-rank matrix completion:* We believe it is possible to extend our analysis to study rectangular matrix completion without regularization, by combining the leave-one-out technique in [27], [35] to carefully bound the incoherence of the iterates for both factors even without explicit balancing.
- *Improving dependence on κ and r :* The current paper does not try to optimize the dependence with respect to κ and r in terms of sample complexity and the size of the basin of attraction, which are slightly worse than their regularized counterparts. A finer analysis will likely lead to better dependencies, which we leave to the future work.

APPENDIX A

A. Proof of Lemma 1

For notational convenience, we define the following function

$$g(\mathbf{Q}) \triangleq \|\mathbf{X}\mathbf{Q} - \mathbf{X}_*\|_F^2 + \|\mathbf{Y}\mathbf{Q}^{-\top} - \mathbf{Y}_*\|_F^2. \quad (22)$$

Clearly, the optimal alignment matrix, if exists, must be $\text{argmin } g(\mathbf{P})$. With this notation in place, we consider the following constrained minimization problem:

$$\begin{aligned}
& \min_{\mathbf{Q} \in \mathbb{R}^{r \times r}: \mathbf{Q} \text{ is invertible}} g(\mathbf{Q}) \\
& \text{subject to} \quad \|\mathbf{Q} - \mathbf{P}\|_F \leq \frac{5\delta}{\sigma_{\min}(\mathbf{X}_*)}.
\end{aligned}$$

In view of Weyl's inequality, we obtain that for any feasible $\mathbf{Q}$,

$$\sigma_{\min}(\mathbf{Q}) \geq \sigma_{\min}(\mathbf{P}) - \frac{5\delta}{\sigma_{\min}(\mathbf{X}_*)} \geq \frac{1}{2} - \frac{1}{4} = \frac{1}{4}$$

as long as $\delta \leq \sigma_{\min}(\mathbf{X}_*)/80$. As a result, one sees that $g(\mathbf{Q})$ is a continuous function over $\{\mathbf{Q} : \|\mathbf{Q} - \mathbf{P}\|_F \leq 5\delta/\sigma_{\min}(\mathbf{X}_*)\}$, which is a compact set over invertible matrices. Applying the Weierstrass extreme value theorem yields the claim that the minimizer of the constrained problem exists. Denote this minimizer by $\mathbf{Q}_1$. In what follows, we intend to show that $\mathbf{Q}_1$ is also the minimizer of the unconstrained problem. Letting $\mathbf{Q}$ be an arbitrary matrix with $g(\mathbf{Q}) \leq 2\delta^2$ (the existence is assured since $g(\mathbf{Q}_1) \leq g(\mathbf{P}) \leq 2\delta^2$), we have

$$\sqrt{2}\delta \geq \|\mathbf{X}\mathbf{Q} - \mathbf{X}_*\|_F \geq \|\mathbf{X}\mathbf{Q} - \mathbf{X}\mathbf{P}\|_F - \|\mathbf{X}\mathbf{P} - \mathbf{X}_*\|_F,$$

which in conjunction with (12) implies

$$(1 + \sqrt{2})\delta \geq \|\mathbf{X}(\mathbf{Q} - \mathbf{P})\|_F \geq \sigma_{\min}(\mathbf{X}) \|\mathbf{Q}_1 - \mathbf{P}\|_F. \quad (23)$$

We now turn to investigating $\sigma_{\min}(\mathbf{X})$. Weyl's inequality tells us that

$$\begin{aligned}
|\sigma_{\min}(\mathbf{X}\mathbf{P}) - \sigma_{\min}(\mathbf{X}_*)| & \leq \|\mathbf{X}\mathbf{P} - \mathbf{X}_*\|_F \\
& \leq \delta \leq \frac{1}{4} \sigma_{\min}(\mathbf{X}_*),
\end{aligned}$$

which further implies

$$\begin{aligned} \frac{3}{4}\sigma_{\min}(\mathbf{X}_*) &\leq \sigma_{\min}(\mathbf{X}\mathbf{P}) \\ &\leq \sigma_{\min}(\mathbf{X})\sigma_{\max}(\mathbf{P}) \leq \frac{3}{2}\sigma_{\min}(\mathbf{X}). \end{aligned}$$

Therefore we arrive at $\sigma_{\min}(\mathbf{X}) \geq \sigma_{\min}(\mathbf{X}_*)/2$. Putting this back to (23) yields which finally gives

$$\|\mathbf{Q} - \mathbf{P}\|_F \leq 2(1 + \sqrt{2})\frac{\delta}{\sigma_{\min}(\mathbf{X}_*)} < \frac{5\delta}{\sigma_{\min}(\mathbf{X}_*)}.$$

In all, the above arguments reveal that any matrix $\mathbf{Q}$ such that $g(\mathbf{Q}) \leq 2\delta^2$ must obey the above bound. Therefore the minimizer of the constrained problem and that of the unconstrained one coincide with each other. This finished the proof.

B. Proof of Lemma 2

Recall the function $g(\mathbf{P})$ defined in (22) as

$$\begin{aligned} g(\mathbf{P}) &= \|\mathbf{X}\mathbf{P} - \mathbf{X}_*\|_F^2 + \|\mathbf{Y}\mathbf{P}^\top - \mathbf{Y}_*\|_F^2 \\ &= \text{Tr}(\mathbf{X}\mathbf{P}\mathbf{P}^\top\mathbf{X}^\top) - 2\text{Tr}(\mathbf{P}^\top\mathbf{X}^\top\mathbf{X}_*) \\ &\quad + \text{Tr}(\mathbf{X}_*^\top\mathbf{X}_*) + \text{Tr}(\mathbf{P}^{-1}\mathbf{Y}^\top\mathbf{Y}\mathbf{P}^{-\top}) \\ &\quad - 2\text{Tr}(\mathbf{P}^{-1}\mathbf{Y}^\top\mathbf{Y}_*) + \text{Tr}(\mathbf{Y}_*^\top\mathbf{Y}_*). \end{aligned}$$

The gradient is given by

$$\begin{aligned} \nabla g(\mathbf{P}) &= 2\mathbf{X}^\top\mathbf{X}\mathbf{P} - 2\mathbf{X}^\top\mathbf{X}_* \\ &\quad - 2(\mathbf{P}\mathbf{P}^\top)^{-1}\mathbf{Y}^\top\mathbf{Y}(\mathbf{P}\mathbf{P}^\top)^{-1}\mathbf{P} + 2\mathbf{P}^{-\top}\mathbf{Y}_*^\top\mathbf{Y}\mathbf{P}^{-\top}. \end{aligned}$$

Since $\mathbf{Q}$ minimizes $g(\mathbf{P})$, it must satisfy the first-order optimality condition, i.e.

$$\nabla g(\mathbf{Q}) = \mathbf{0}.$$

Identify $\widetilde{\mathbf{X}} = \mathbf{X}\mathbf{Q}$ and $\widetilde{\mathbf{Y}} = \mathbf{Y}\mathbf{Q}^\top$ to yield the condition

$$\widetilde{\mathbf{X}}^\top\widetilde{\mathbf{X}} - \widetilde{\mathbf{X}}^\top\mathbf{X}_* = \widetilde{\mathbf{Y}}^\top\widetilde{\mathbf{Y}} - \mathbf{Y}_*^\top\widetilde{\mathbf{Y}}.$$

C. Proof of Lemma 3

We prove the first part and the second part follows by symmetry. Denote $\mathbf{E}_x = \mathbf{X} - \mathbf{X}_*$ and $\mathbf{E}_y = \mathbf{Y} - \mathbf{Y}_*$. We have

$$\mathbf{X}\mathbf{Y}^\top - \mathbf{M}_* = \mathbf{E}_x\mathbf{Y}^\top + \mathbf{X}_*\mathbf{E}_y^\top.$$

Since $\mathbf{Z}$ is aligned with $\mathbf{Z}_*$, Lemma 2 tells us that

$$\mathbf{X}^\top\mathbf{E}_x = \mathbf{E}_y^\top\mathbf{Y}.$$

As a result, one has

$$\begin{aligned} &\langle \mathbf{X} - \mathbf{X}_*, (\mathbf{X}\mathbf{Y}^\top - \mathbf{M}_*) \mathbf{Y} \rangle \\ &= \text{Tr}(\mathbf{E}_x^\top(\mathbf{E}_x\mathbf{Y}^\top + \mathbf{X}_*\mathbf{E}_y^\top)\mathbf{Y}) \\ &= \text{Tr}(\mathbf{E}_x^\top\mathbf{E}_x\mathbf{Y}^\top\mathbf{Y}) + \text{Tr}(\mathbf{E}_x^\top\mathbf{X}_*\mathbf{E}_y^\top\mathbf{Y}) \\ &= \|\mathbf{Y}\mathbf{E}_x^\top\|_F^2 + \text{Tr}(\mathbf{E}_y^\top\mathbf{Y}\mathbf{E}_x^\top\mathbf{X}) - \text{Tr}(\mathbf{E}_y^\top\mathbf{Y}\mathbf{E}_x^\top\mathbf{E}_x) \\ &= \|\mathbf{Y}\mathbf{E}_x^\top\|_F^2 + \|\mathbf{X}^\top\mathbf{E}_x\|_F^2 - \text{Tr}(\mathbf{X}^\top\mathbf{E}_x\mathbf{E}_x^\top\mathbf{E}_x). \quad (24) \end{aligned}$$

Complete the squares to see that

$$\begin{aligned} &\|\mathbf{X}^\top\mathbf{E}_x\|_F^2 - \text{Tr}(\mathbf{X}^\top\mathbf{E}_x\mathbf{E}_x^\top\mathbf{E}_x) \\ &= \|\mathbf{E}_x^\top\mathbf{X} - \frac{1}{2}\mathbf{E}_x^\top\mathbf{E}_x\|_F^2 - \frac{1}{4}\|\mathbf{E}_x\|_F^4. \end{aligned}$$

Combine the previous two bounds to yield the desired result.

D. Proof of Lemma 4

Again, we demonstrate the claim on $\mathbf{X}$ and the claim on $\mathbf{Y}$ follows by symmetry. Given the decomposition

$$\begin{aligned} \mathbf{X}\mathbf{Y}^\top - \mathbf{M}_* &= (\mathbf{X} - \mathbf{X}_*)\mathbf{Y}^\top + \mathbf{X}(\mathbf{Y} - \mathbf{Y}_*)^\top \\ &\quad + (\mathbf{X}_* - \mathbf{X})(\mathbf{Y} - \mathbf{Y}_*)^\top, \end{aligned}$$

we obtain

$$\begin{aligned} \|\mathbf{X}\mathbf{Y}^\top - \mathbf{M}_*\|_F &\leq \sigma_1(\mathbf{Y})\|\mathbf{X}\mathbf{Y}^\top - \mathbf{M}_*\|_F \\ &\leq \frac{3}{2}\sigma_1(\mathbf{Y}_*)\left(\|(\mathbf{X} - \mathbf{X}_*)\mathbf{Y}^\top\|_F \right. \\ &\quad \left. + \|\mathbf{X}(\mathbf{Y} - \mathbf{Y}_*)^\top\|_F + \|(\mathbf{X}_* - \mathbf{X})(\mathbf{Y} - \mathbf{Y}_*)^\top\|_F\right), \end{aligned}$$

where the last line combines the triangle inequality and Weyl's inequality

$$\sigma_1(\mathbf{Y}) \leq \sigma_1(\mathbf{Y}_*) + \|\mathbf{Y} - \mathbf{Y}_*\| \leq \frac{3}{2}\sigma_1(\mathbf{Y}_*).$$

The proof is then finished.

REFERENCES

- [1] C. Ma, Y. Li, and Y. Chi, "Beyond procrustes: Balancing-free gradient descent for asymmetric low-rank matrix sensing," in *Proc. 53rd Asilomar Conf. Signals, Syst., Comput.*, 2019, pp. 721–725.
- [2] E. Candès and T. Tao, "The power of convex relaxation: Near-optimal matrix completion," *IEEE Trans. Inf. Theory*, vol. 56, no. 5, pp. 2053–2080, May 2010.
- [3] Y. Chen and Y. Chi, "Harnessing structures in big data via guaranteed low-rank matrix estimation: Recent theory and fast algorithms via convex and nonconvex optimization," *IEEE Signal Process. Mag.*, vol. 35, no. 4, pp. 14–31, Jul. 2018.
- [4] M. A. Davenport and J. Romberg, "An overview of low-rank matrix recovery from incomplete observations," *IEEE J. Sel. Topics Signal Process.*, vol. 10, no. 4, pp. 608–622, Jun. 2016.
- [5] S. Burer and R. Monteiro, "A nonlinear programming algorithm for solving semidefinite programs via low-rank factorization," *Math. Program.*, vol. 95, no. 2, pp. 329–357, 2003.
- [6] S. Bhojanapalli, A. Kyrillidis, and S. Sanghavi, "Dropping convexity for faster semi-definite optimization," in *Proc. Conf. Learn. Theory*, 2016, pp. 530–582.
- [7] N. Boumal, V. Voroninski, and A. Bandeira, "The non-convex Burer-Monteiro approach works on smooth semidefinite programs," in *Proc. Adv. Neural Inf. Process. Syst.*, 2016, pp. 2757–2765.
- [8] 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.*, 2016, pp. 964–973.
- [9] Q. Zheng and J. Lafferty, "Convergence analysis for rectangular matrix completion using Burer-Monteiro factorization and gradient descent," 2016, *arXiv:1605.07051*.
- [10] D. Park, A. Kyrillidis, C. Caramanis, and S. Sanghavi, "Finding low-rank solutions via nonconvex matrix factorization, efficiently and provably," *SIAM J. Imag. Sci.*, vol. 11, no. 4, pp. 2165–2204, 2018.
- [11] Y. Chi, "Low-rank matrix completion," *IEEE Signal Process. Mag.*, vol. 35, no. 5, pp. 178–181, 2018.
- [12] X. Yi, D. Park, Y. Chen, and C. Caramanis, "Fast algorithms for robust PCA via gradient descent," in *Proc. Adv. Neural Inf. Process. Syst.*, 2016, pp. 4152–4160.
- [13] Y. Chen, Y. Chi, J. Fan, and C. Ma, "Spectral methods for data science: A statistical perspective," 2020, *arXiv:2012.08496*.
- [14] B. Recht, M. Fazel, and P. A. Parrilo, "Guaranteed minimum-rank solutions of linear matrix equations via nuclear norm minimization," *SIAM Rev.*, vol. 52, no. 3, pp. 471–501, 2010.
- [15] S. Oymak, B. Recht, and M. Soltanolkotabi, "Sharp time-data tradeoffs for linear inverse problems," *IEEE Trans. Inf. Theory*, vol. 64, no. 6, pp. 4129–4158, Jun. 2018.
- [16] E. J. Candès and B. Recht, "Exact matrix completion via convex optimization," *Foundations Comput. Math.*, vol. 9, no. 6, pp. 717–772, 2009.

- [17] E. J. Candès and Y. Plan, "Matrix completion with noise," *Proc. IEEE*, vol. 98, no. 6, pp. 925–936, Jun. 2010.
- [18] S. Negahban and M. Wainwright, "Restricted strong convexity and weighted matrix completion: Optimal bounds with noise," *J. Mach. Learn. Res.*, vol. 13, no. 53, pp. 1665–1697, May 2012.
- [19] D. Gross, "Recovering low-rank matrices from few coefficients in any basis," *IEEE Trans. Inf. Theory*, vol. 57, no. 3, pp. 1548–1566, Mar. 2011.
- [20] B. Recht, "A simpler approach to matrix completion," *J. Mach. Learn. Res.*, vol. 12, no. 104, pp. 3413–3430, Feb. 2011.
- [21] Y. Chen and Y. Chi, "Robust spectral compressed sensing via structured matrix completion," *IEEE Trans. Inf. Theory*, vol. 60, no. 10, pp. 6576–6601, Oct. 2014.
- [22] S. Negahban and M. J. Wainwright, "Estimation of (near) low-rank matrices with noise and high-dimensional scaling," *Ann. Statist.*, vol. 39, no. 2, pp. 1069–1097, 2011.
- [23] Q. Zheng and J. Lafferty, "A convergent gradient descent algorithm for rank minimization and semidefinite programming from random linear measurements," in *Proc. Adv. Neural Inf. Process. Syst.*, 2015, pp. 109–117.
- [24] R. Keshavan, A. Montanari, and S. Oh, "Matrix completion from a few entries," *IEEE Trans. Inf. Theory*, vol. 56, no. 6, pp. 2980–2998, Jun. 2010.
- [25] R. Sun and Z.-Q. Luo, "Guaranteed matrix completion via nonconvex factorization," in *Proc. Symp. Foundations Comput. Sci.*, 2015, pp. 270–289.
- [26] P. Jain, P. Netrapalli, and S. Sanghavi, "Low-rank matrix completion using alternating minimization," in *Proc. 45th Annu. ACM Symp. Theory Comput.*, 2013, pp. 665–674.
- [27] C. Ma, K. Wang, Y. Chi, and Y. Chen, "Implicit regularization in non-convex statistical estimation: Gradient descent converges linearly for phase retrieval, matrix completion, and blind deconvolution," *Foundations Comput. Math.*, vol. 20, no. 3, pp. 1–182, 2019.
- [28] Y. Chen and M. J. Wainwright, "Fast low-rank estimation by projected gradient descent: General statistical and algorithmic guarantees," 2015, *arXiv:1509.03025*.
- [29] Y. Li, C. Ma, Y. Chen, and Y. Chi, "Nonconvex matrix factorization from rank-one measurements," in *Proc. 22nd Int. Conf. Artif. Intell. Statist.*, 2019, pp. 1496–1505.
- [30] X. Li, S. Ling, T. Strohmer, and K. Wei, "Rapid, robust, and reliable blind deconvolution via nonconvex optimization," *Appl. Comput. Harmon. Anal.*, vol. 47, no. 3, pp. 893–934, 2019.
- [31] Y. Chen, Y. Chi, J. Fan, and C. Ma, "Gradient descent with random initialization: Fast global convergence for nonconvex phase retrieval," *Math. Program.*, vol. 176, no. 1, pp. 5–37, 2019.
- [32] Y. Chi, Y. M. Lu, and Y. Chen, "Nonconvex optimization meets low-rank matrix factorization: An overview," *IEEE Trans. Signal Process.*, vol. 67, no. 20, pp. 5239–5269, Oct. 2019.
- [33] Y. Li, Y. Chi, H. Zhang, and Y. Liang, "Non-convex low-rank matrix recovery with arbitrary outliers via median-truncated gradient descent," *Inf. Inference: J. IMA*, vol. 9, no. 2, pp. 289–325, 2020.
- [34] X. Zhang, S. Du, and Q. Gu, "Fast and sample efficient inductive matrix completion via multi-phase procrustes flow," in *Proc. Int. Conf. Mach. Learn.*, 2018, pp. 5751–5760.
- [35] J. Chen, D. Liu, and X. Li, "Nonconvex rectangular matrix completion via gradient descent without $\ell_{2,\infty}$ regularization," *IEEE Trans. Inf. Theory*, vol. 66, no. 9, pp. 5806–5841, 2020.
- [36] S. S. Du, W. Hu, and J. D. Lee, "Algorithmic regularization in learning deep homogeneous models: Layers are automatically balanced," in *Proc. Adv. Neural Inf. Process. Syst.*, 2018, pp. 384–395.
- [37] V. Charisopoulos, Y. Chen, D. Davis, M. Diaz, L. Ding, and D. Drusvyatskiy, "Low-rank matrix recovery with composite optimization: Good conditioning and rapid convergence," 2019, *arXiv:1904.10020*.
- [38] R. Ge, J. D. Lee, and T. Ma, "Matrix completion has no spurious local minimum," in *Proc. Adv. Neural Inf. Process. Syst.*, 2016, pp. 2973–2981.
- [39] 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.*, 2017, pp. 1233–1242.
- [40] Z. Zhu, Q. Li, G. Tang, and M. B. Wakin, "Global optimality in low-rank matrix optimization," *IEEE Trans. Signal Process.*, vol. 66, no. 13, pp. 3614–3628, Jul. 2018.
- [41] Z. Zhu, D. Soudry, Y. C. Eldar, and M. B. Wakin, "The global optimization geometry of shallow linear neural networks," *J. Math. Imag. Vis.*, vol. 62, no. 3, pp. 1–14, 2019.
- [42] X. Li et al., "Symmetry, saddle points, and global optimization landscape of nonconvex matrix factorization," *IEEE Trans. Inf. Theory*, vol. 65, no. 6, pp. 3489–3514, Jun. 2019.
- [43] Q. Li, Z. Zhu, and G. Tang, "The non-convex geometry of low-rank matrix optimization," *Inf. Inference: J. IMA*, vol. 8, no. 1, pp. 51–96, 2019.
- [44] S. Li, Q. Li, Z. Zhu, G. Tang, and M. B. Wakin, "The global geometry of centralized and distributed low-rank matrix recovery without regularization," *IEEE Signal Process. Lett.*, vol. 27, pp. 1400–1404, 2020.
- [45] T. Tong, C. Ma, and Y. Chi, "Accelerating ill-conditioned low-rank matrix estimation via scaled gradient descent," 2020, *arXiv:2005.08898*.
- [46] T. Tong, C. Ma, and Y. Chi, "Low-rank matrix recovery with scaled subgradient methods: Fast and robust convergence without the condition number," 2020, *arXiv:2010.13364*.

Cong Ma received the B.Eng. degree from Tsinghua University, Beijing, China, in 2015, and the Ph.D. degree from Princeton University, NJ, USA, in 2020. He is currently a Postdoctoral Researcher with the Department of Electrical Engineering and Computer Sciences, UC Berkeley, and will be joining the Department of Statistics, the University of Chicago as an Assistant Professor in July 2021. His research interests include mathematics of data science, machine learning, high-dimensional statistics, convex and nonconvex optimization as well as their applications to neuroscience. Dr. Ma was the recipient of the School of Engineering and Applied Science Award for Excellence from Princeton University in 2019, the AI Labs Fellowship from Hudson River Trading in 2019, and the Student Paper Award from International Chinese Statistical Association in 2017.

Yuanxin Li received the B.Eng. degree from the Nanjing University of Posts and Telecommunications, Nanjing, China, in 2010, the M.Eng. degree from Tsinghua University, Beijing, China, in 2013, the M.S. degree as well as the Ph.D. degree from The Ohio State University, Columbus, OH, USA, in 2016 and 2018, respectively. He was a Visiting Student and subsequently a Postdoctoral Researcher with the Department of Electrical and Computer Engineering, Carnegie Mellon University in 2018. He is currently with Samsung Semiconductor, Inc. as a Senior Engineer in San Diego, CA. His research interests include statistical signal processing, convex optimization, computer vision and data analysis.

Yuejie Chi (Senior Member, IEEE) received the Ph.D. and M.A. degrees in electrical engineering from Princeton University in 2012 and 2009, and the B.E. (Hons.) degree in electrical engineering from Tsinghua University, Beijing, China, in 2007. She was with The Ohio State University from 2012 to 2017. Since 2018, she is an Associate Professor with the Department of Electrical and Computer Engineering, Carnegie Mellon University, where she holds the Robert E. Doherty Early Career Development Professorship. Her research interests include the theoretical and algorithmic foundations of data science, signal processing, machine learning and inverse problems, with applications in sensing systems, broadly defined. Among others, she was the recipient of Presidential Early Career Award for Scientists and Engineers (PECASE), the inaugural IEEE Signal Processing Society Early Career Technical Achievement Award for contributions to high-dimensional structured signal processing, and named the 2021 Goldsmith Lecturer by IEEE Information Theory Society. She currently is an Associate Editor for IEEE TRANSACTIONS ON SIGNAL PROCESSING and IEEE TRANSACTIONS ON PATTERN RECOGNITION AND MACHINE INTELLIGENCE.