
Geometric Analysis of Noisy Low-rank Matrix Recovery in the Exact Parameterized and the Overparameterized Regimes

¹Ziye Ma*, ²Yingjie Bi*, ²Javad Lavaei, ^{1,3}Somayeh Sojoudi

¹Department of Electrical Engineering and Computer Science

²Department of Industrial Engineering and Operations Research

³Department of Mechanical Engineering

University of California, Berkeley

{ziyema,yingjiebi,lavaei,sojoudi}@berkeley.edu

Abstract

In this paper, we study noisy low-rank matrix recovery problems with linear measurements and arbitrary probability distributions for noise. We investigate the scenario where the search rank r is equal to the true rank r^* of the unknown ground truth (the exact parameterized case), as well as the scenario where r is greater than r^* (the overparameterized case). The objective is to understand under what conditions on the restricted isometry property (RIP) the non-convex factorized formulation of the problem has a benign landscape and thus local search methods can find the ground truth with a small error. First, we develop a global guarantee on the maximum distance between an arbitrary local minimizer of the non-convex problem and the ground truth under the assumption that the RIP constant is smaller than $1/(1 + \sqrt{r^*/r})$. As expected, this distance shrinks to zero as the intensity of the noise reduces, which recovers the state-of-the-art result concerning the noiseless version of this problem. Our new guarantee is sharp in terms of the RIP constant and is much stronger than the existing results. We then present a local guarantee for problems with an arbitrary RIP constant, which states that any local minimizer is either considerably close to the ground truth or far away from it. Next, we prove the strict saddle property under the same RIP assumption as above, which leads to the global convergence of the perturbed gradient descent method. The developed results demonstrate how the noise intensity, the parameterization and the RIP constant affect the landscape of the non-convex optimization problem.

1 Introduction

Low-rank matrix recovery problems arise in various applications, such as matrix completion [1, 2], phase synchronization/retrieval [3–5], robust PCA [6], and several others [7, 8]. In this paper, we study a class of low-rank matrix recovery problems, where the goal is to recover a symmetric and positive semidefinite ground truth matrix M^* with $\text{rank}(M^*) = r^* > 0$ from certain linear measurements corrupted by noise. In many cases, the exact rank r^* is often unknown a priori, which prompts the user to choose a sufficiently large r with $r \geq r^*$ and formulate the above learning problem as the following optimization problem:

$$\begin{aligned} \min_{M \in \mathbb{R}^{n \times n}} \quad & \frac{1}{2} \|\mathcal{A}(M) - b + w\|^2 \\ \text{s. t.} \quad & \text{rank}(M) \leq r, \quad M \succeq 0. \end{aligned} \tag{1}$$

*Equal Contribution

Here, $\mathcal{A} : \mathbb{R}^{n \times n} \rightarrow \mathbb{R}^m$ is a linear operator whose action on a matrix M is given by

$$\mathcal{A}(M) = [\langle A_1, M \rangle, \dots, \langle A_m, M \rangle]^T,$$

where $A_1, \dots, A_m \in \mathbb{R}^{n \times n}$ are called sensing matrices. In addition, $b = \mathcal{A}(M^*)$ represents the perfect measurement on the ground truth M^* and w comes from an arbitrary probability distribution. Note that only the noisy measurement $b - w$ is available to the user, and indeed b is unknown. In other words, from a problem-solving perspective, the random variable w is hidden to the user, and it is explicitly modeled here only for the sake of analysis. In this work, we call r the search rank and r^* the true rank. If $r = r^*$, the problem is said to be exact-parameterized. If $r > r^*$, it is said to be overparameterized.

Although existing literature such as [1, 2, 9] demonstrated that it is possible to solve (1) using convex relaxations under appropriate assumptions, the computationally demanding nature of semidefinite programs makes it difficult to apply such techniques to large-scale problems. A more scalable approach is to use the Burer–Monteiro factorization [10] by expressing M as XX^T with $X \in \mathbb{R}^{n \times r}$, which leads to the following equivalent formulation of the aforementioned problem (1):

$$\min_{X \in \mathbb{R}^{n \times r}} f(X) = \frac{1}{2} \|\mathcal{A}(XX^T) - b + w\|^2. \quad (2)$$

Since (2) is unconstrained, it can be easily solved by local search methods such as gradient descent. However, due to the non-convexity of the objective function $f(X)$, local search algorithms may converge to a local minimizer, leading to a suboptimal or plainly wrong solution. Hence, it is both practically and theoretically important to provide guarantees on the maximum distance between these local minimizers and the ground truth M^* . This will be the main focus of this paper. Moreover, although we focus on the symmetric matrix sensing problem, our results can be also applied to the asymmetric problem in which M^* is allowed to be rectangular. This is due to the fact that any asymmetric problem can be equivalently transformed into a symmetric one [6].

1.1 Related works

The special noiseless case of the problem (2) can be obtained by setting $w = 0$. In this case, any solution Z with $ZZ^T = M^*$ is a global minimizer of the problem (2). Several papers such as [6, 11–21] have shown that the problem has no spurious (non-global) local minimizers under the assumption of restricted isometry property (RIP). Moreover, as demonstrated in [6], the developed techniques under the RIP condition can be adopted to show that other low-rank matrix recovery problems, such as the matrix completion under incoherence condition and the robust PCA problem, also have benign landscape. The RIP condition is equivalent to the restricted strongly convex and smooth property used in [19, 22, 23], and its formal definition is given below.

Definition 1. Given a positive integer g , the linear operator $\mathcal{A}(\cdot) : \mathbb{R}^{n \times n} \rightarrow \mathbb{R}^m$ is said to satisfy the δ -RIP $_g$ property for some constant $\delta \in [0, 1)$ if the inequality

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

holds for all $M \in \mathbb{R}^{n \times n}$ with $\text{rank}(M) \leq g$.

In the recent paper [20], the author developed a sharp bound on the absence of spurious local minima for the noiseless case of problem (2), which says that the problem has no spurious local minima if the measurement operator $\mathcal{A}$ satisfies the δ -RIP $_{r+r^*}$ property with $\delta < 1/(1 + \sqrt{r^*/r})$. In the exact parameterized case, this simplifies to $\delta < 1/2$, which is a sharp bound due to the counterexample given in [24] that has spurious local minima under $\delta = 1/2$.

For the general noisy problem, the relation $X^*X^{*T} = M^*$ is unlikely to be satisfied, where X^* denotes a global minimizer of problem (2), due to the influence of noise. However, X^*X^{*T} should be close to the ground truth M^* if the noise w is small. As a generalization of the above-mentioned results for the noiseless problem, it is natural to study whether all local minimizers, including the global minimizers, are close to the ground truth M^* under the RIP assumption. One such result for the exact parameterized case is presented in [11] and given below.

Theorem 1. ([11], Theorem 3.1) Suppose that $w \sim \mathcal{N}(0, \sigma_w^2 I_m)$, $r = r^*$ and $\mathcal{A}(\cdot)$ has the δ -RIP $_{4r}$ property with $\delta < 1/10$. Then, with probability at least $1 - 10/n^2$, every local minimizer $\hat{X}$ of

problem (2) satisfies the inequality

$$\|\hat{X}\hat{X}^T - M^*\|_F \leq 20\sqrt{\frac{\log(n)}{m}}\sigma_w.$$

Theorem 31 in [6] further improves the above result by replacing the δ -RIP_{4r} property with the δ -RIP_{2r} property. [25] studies a similar noisy low-rank matrix recovery problem with the l_1 norm.

There is an apparent gap between the state-of-the-art results for the noiseless and noisy problems. Even in the exact parameterized regime, the result for the noiseless problem only requires the RIP constant $\delta < 1/2$, but Theorem 1 requires $\delta < 1/10$ regardless of the intensity of the noise. Furthermore, the existing results cannot be applied to the overparameterized noisy problem. This gap will be addressed in this paper by showing that a major generalization of Theorem 1 holds for the noisy problem under the same RIP assumption as the bound for the noiseless problem in both the exact parameterized and the overparameterized regimes.

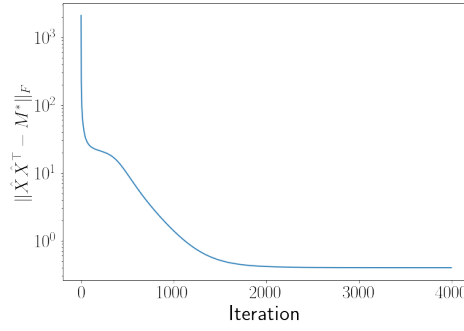


Figure 1: The evolution of the error between the found solution $\hat{X}\hat{X}^T$ and the ground truth M^* during the iterations of the gradient descent method for a noisy problem with the RIP constant $\delta < 1/2$. The error decreases linearly at first, but it cannot be further improved after a certain number of iterations because at that time the remaining error is almost incurred by the measurement noise

Earlier works such as [19, 26, 27] established the strict saddle property for the noiseless and exact parameterized problem, which essentially states that any matrix whose gradient is small and whose Hessian is almost positive semidefinite must be sufficiently close to a global minimizer. This property, together with certain local regularity property near the ground truth, implies the global linear convergence for the perturbed gradient descent method. In other words, the algorithm will return a solution $\hat{X}$ satisfying $\|\hat{X}\hat{X}^T - M^*\|_F \leq e$ after $O(\log(1/e))$ number of iterations. In this paper, we prove a similar strict saddle property for the noisy problem in both the exact parameterized and the overparameterized cases. However, in the noisy problem, even if the local search algorithm finds the global minimum, it cannot recover the ground truth exactly. As such, it is no longer meaningful to discuss the convergence rate because the error between the found solution and the ground truth has two sources: the difference between X^*X^{*T} and the ground truth M^* where X^* denotes an exact global minimizer of the problem, and the difference between $\hat{X}\hat{X}^T$ and X^*X^{*T} where $\hat{X}$ denotes the approximate solution found by the algorithm. Using our strict saddle property, we can characterize the time point when the errors induced by the above two sources are roughly equal. As demonstrated via an example in Fig. 1, it is almost futile to run the algorithm beyond a certain number of iterations since the error will be dominated by the former one after some time.

1.2 Notations

In this paper, I_n refers to the identity matrix of size $n \times n$. The notation $M \succeq 0$ means that M is a symmetric and positive semidefinite matrix. $\sigma_i(M)$ denotes the i -th largest singular value of a matrix M , and $\lambda_i(M)$ denotes the i -th largest eigenvalue of M . $\|v\|$ denotes the Euclidean norm of a vector v , while $\|M\|_F$ and $\|M\|_2$ denote the Frobenius norm and induced l_2 norm of a matrix M , respectively. $\langle A, B \rangle$ is defined to be $\text{tr}(A^T B)$ for two matrices A and B of the same size. The Kronecker product between A and B is denoted as $A \otimes B$. For a matrix M , $\text{vec}(M)$ is the

usual vectorization operation by stacking the columns of the matrix M into a vector. For a vector $v \in \mathbb{R}^{n^2}$, $\text{mat}(v)$ converts v to a square matrix and $\text{mat}_S(v)$ converts v to a symmetric matrix, i.e., $\text{mat}(v) = M$ and $\text{mat}_S(v) = (M + M^T)/2$, where $M \in \mathbb{R}^{n \times n}$ is the unique matrix satisfying $v = \text{vec}(M)$. Finally, $\mathcal{N}(\mu, \Sigma)$ refers to the multivariate Gaussian distribution with mean μ and covariance Σ .

2 Main results

2.1 Guarantees on the local minima

We first present the global guarantee on the local minimizers of the problem (2). To simplify the notation, we use a matrix representation of the measurement operator $\mathcal{A}$ as follows:

$$\mathbf{A} = [\text{vec}(A_1), \text{vec}(A_2), \dots, \text{vec}(A_m)]^T \in \mathbb{R}^{m \times n^2}.$$

Then, $\mathbf{A} \text{vec}(M) = \mathcal{A}(M)$ for every matrix $M \in \mathbb{R}^{n \times n}$.

Theorem 2. *Given arbitrary positive integers r and r^* such that $r \geq r^*$, assume that the linear operator $\mathcal{A}$ satisfies the δ -RIP $_{r+r^*}$ property with $\delta < 1/(1 + \sqrt{r^*/r})$. For every $\epsilon > 0$, with probability at least $\mathbb{P}(\|\mathbf{A}^T w\| \leq \epsilon)$, either of the following two inequalities*

$$(1 - \delta) \|\hat{X} \hat{X}^T - M^*\|_F^2 \leq \epsilon \sqrt{r} \|\hat{X} \hat{X}^T - M^*\|_F + 4\epsilon \sqrt{r} \|M^*\|_F, \quad (3a)$$

$$\left(\frac{1 - \delta}{1 + \delta} - \frac{\sqrt{r^*/r}}{2 + \sqrt{r^*/r}} \right) \|\hat{X} \hat{X}^T - M^*\|_F \leq 2\epsilon \sqrt{r^*} + 2\sqrt{2\epsilon(1 + \delta)} (\|\hat{X} \hat{X}^T - M^*\|_F^{1/2} + \|M^*\|_F^{1/2}) \quad (3b)$$

holds for every arbitrary local minimizer $\hat{X} \in \mathbb{R}^{n \times r}$ of problem (2).

Note that two upper bounds on the distance $\|\hat{X} \hat{X}^T - M^*\|_F$ can be obtained for every local minimizer $\hat{X}$ by solving the two quadratic-like inequalities (3a) and (3b), and the larger bound needs to be used because only one of the two inequalities is guaranteed to hold. The reason for the existence of two inequalities in Theorem 2 is the split of its proof into two cases. The first case is associated with the r -th smallest singular value of $\hat{X}$ being small and the second case is the opposite, which are respectively handled by Lemma 2 and Lemma 3.

Theorem 2 is a major extension of the state-of-the-art result stating that the noiseless problem has no spurious local minima under the same assumption of the δ -RIP $_{r+r^*}$ property with $\delta < 1/(1 + \sqrt{r^*/r})$. The reason is that in the case when the noise w is equal to zero, one can choose an arbitrarily small ϵ in Theorem 2 to conclude from the inequalities (3a) and (3b) that $\hat{X} \hat{X}^T = M^*$ for every local minimizer $\hat{X}$. Moreover, when the RIP constant δ further decreases from $1/(1 + \sqrt{r^*/r})$, the upper bound on $\|\hat{X} \hat{X}^T - M^*\|_F$ will also decrease, which means that a local minimizer found by local search methods will be closer to the ground truth M^* . This suggests that the RIP condition is able to not only guarantee the absence of spurious local minima as shown in the previous literature but also mitigate the influence of the noise in the measurements.

Compared with the existing results such as Theorem 1, our new result has two advantages even when specialized to the exact parameterized case $r = r^*$. First, by improving the RIP constant from $1/10$ to $1/2$, one can apply the results on the location of spurious local minima to a much broader class of problems, which can often help reduce the number of measurements. For example, in the case when the measurements are given by random Gaussian matrices, it is proven in [28] that to achieve the δ -RIP $_{2r}$ property the minimum number of measurements needed is on the order of $O(1/\delta^2)$. By improving the RIP constant in the bound, we can significantly reduce the number of measurements while still keeping the benign landscape. In applications such as learning for energy networks, there is a fundamental limit on the number of measurements that can be collected due to the physics of the problem [29]. Finding a better bound on RIP helps with addressing the issues with the number of measurements needed to reliably solve the problem. Second, Theorem 1 is just about the probability of having all spurious solutions in a fixed ball around the ground truth of radius $O(\sigma_w)$ instead of balls of arbitrary radii, and this fixed ball could be a large one depending on whether the noise level

σ_w is fixed or scales with the problem. On the other hand, in Theorem 2, we consider the probability $\mathbb{P}(\|\mathbf{A}^T w\| \leq \epsilon)$ for any arbitrary value of ϵ . By having a flexible ϵ , our work not only improves the RIP constant but also allows computing the probability of having all spurious solutions in any given ball.

In the special case of rank $r = r^* = 1$, the conditions (3a) and (3b) in Theorem 2 can be substituted with a simpler condition as presented below.

Theorem 3. *Consider the case $r = r^* = 1$ and assume that the linear operator $\mathcal{A}$ satisfies the δ -RIP₂ property with $\delta < 1/2$. For every $\epsilon > 0$, with probability at least $\mathbb{P}(\|\mathbf{A}^T w\| \leq \epsilon)$, every arbitrary local minimizer $\hat{X} \in \mathbb{R}^{n \times r}$ of problem (2) satisfies*

$$\|\hat{X} \hat{X}^T - M^*\|_F \leq \frac{3(1 + \sqrt{2})\epsilon(1 + \delta)}{1 - 2\delta}. \quad (4)$$

In the case when the RIP constant δ is not less than $1/(1 + \sqrt{r^*/r})$, it is not possible to achieve a global guarantee similar to Theorem 2 or Theorem 3 since it is known that the problem may have a spurious solution even in the noiseless case. Instead, we turn to local guarantees by showing that every arbitrary local minimizer $\hat{X}$ of problem (2) is either close to the ground truth M^* or far away from it in terms of the distance $\|\hat{X} \hat{X}^T - M^*\|_F$.

Theorem 4. *Assume that the linear operator $\mathcal{A}$ satisfies the δ -RIP _{$r+r^*$} property for some $\delta \in [0, 1)$. Consider arbitrary constants $\epsilon > 0$ and $\tau \in (0, 1)$ such that $\delta < \sqrt{1 - \tau}$. Every arbitrary local minimizer $\hat{X} \in \mathbb{R}^{n \times r}$ of problem (2) satisfying*

$$\|\hat{X} \hat{X}^T - M^*\|_F \leq \tau \lambda_{r^*}(M^*) \quad (5)$$

will satisfy at least one of the following inequalities

$$(1 - \delta)\|\hat{X} \hat{X}^T - M^*\|_F^2 \leq \epsilon \sqrt{r} \|\hat{X} \hat{X}^T - M^*\|_F + 4\epsilon \sqrt{r} \|M^*\|_F, \quad (6a)$$

$$\|\hat{X} \hat{X}^T - M^*\|_F \leq \frac{\sqrt{\epsilon}(1 + \delta)^{3/2} C(\tau, M^*)}{\sqrt{1 - \tau} - \delta} \quad (6b)$$

with probability at least $\mathbb{P}(\|\mathbf{A}^T w\| \leq \epsilon)$, where

$$C(\tau, M^*) = \sqrt{2(\lambda_1(M^*) + \tau \lambda_{r^*}(M^*))}.$$

The upper bounds in (5), (6a) and (6b) define an outer ball and an inner ball centered at the ground truth M^* . Theorem 4 states that there is no local minimizer in the ring between the two balls, which means that bad local minimizers are located outside the outer ball. Note that the problem could be highly non-convex when δ is close to 1, while this theorem shows a benign landscape in a local neighborhood of the solution. Furthermore, similar to Theorem 2 and Theorem 3, as ϵ approaches zero, the inner ball shrinks to the ground truth. Hence, when the problem is noiseless, Theorem 4 shows that every local minimizer $\hat{X}$ satisfying (5) must have $\hat{X} \hat{X}^T = M^*$. In the noiseless and exact parameterized case, this exactly recovers Theorem 5 in [16]. Our theorem significantly generalizes the previous result by showing that the same conclusion also holds in the overparameterized regime.

As a remark, all the theorems in this section are applicable to arbitrary noise models since they make no explicit use of the probability distribution of the noise. The only required information is the probability $\mathbb{P}(\|\mathbf{A}^T w\| \leq \epsilon)$, which can be computed or bounded when the probability distribution of the noise is given, as illustrated in Section 4.

2.2 Strict saddle property and global convergence

The results presented above are all about the locations of the local minimizers. They do not automatically imply the global convergence of local search methods with a fast convergence rate. To provide performance guarantees for local search methods, the next theorem establishes a stronger property for the landscape of the noisy problem that is similar to the strict saddle property in the literature, which essentially states that all approximate second-order critical points are close to the ground truth. For notational convenience, in the following $\nabla^2 f(\hat{X})$ refers to the Hessian of the objective function f in the matrix form, i.e., $\nabla^2 f(\hat{X})$ is the matrix satisfying the equation

$$(\text{vec}(U))^T \nabla^2 f(\hat{X}) \text{vec}(V) = \sum_{i,j,k,l} \frac{\partial^2 f}{\partial X_{ij} \partial X_{kl}}(\hat{X}) U_{ij} V_{kl},$$

for all $U, V \in \mathbb{R}^{n \times r}$.

Theorem 5. *Given arbitrary positive integers r and r^* such that $r \geq r^*$, assume that the linear operator $\mathcal{A}$ satisfies the δ -RIP $_{r+r^*}$ property with $\delta < 1/(1 + \sqrt{r^*/r})$. For every $\epsilon > 0$ and $\kappa \geq 0$, with probability at least $\mathbb{P}(\|\mathbf{A}^T w\| \leq \epsilon)$, either of the following two inequalities*

$$(1 - \delta)\|\hat{X}\hat{X}^T - M^*\|_F^2 \leq (\epsilon + \kappa/2)\sqrt{r}\|\hat{X}\hat{X}^T - M^*\|_F + (4\epsilon + 2\kappa)\sqrt{r}\|M^*\|_F, \quad (7a)$$

$$\begin{aligned} \left(\frac{1 - \delta}{1 + \delta} - \frac{\sqrt{r^*/r}}{2 + \sqrt{r^*/r}} \right) \|\hat{X}\hat{X}^T - M^*\|_F &\leq (2\epsilon + \kappa)\sqrt{r^*} \\ &+ 2\sqrt{(2\epsilon + \kappa)(1 + \delta)}(\|\hat{X}\hat{X}^T - M^*\|_F^{1/2} + \|M^*\|_F^{1/2}) \end{aligned} \quad (7b)$$

holds for every matrix $\hat{X} \in \mathbb{R}^{n \times r}$ satisfying

$$\|\nabla f(\hat{X})\| \leq \kappa\|\hat{X}\|_2, \quad \nabla^2 f(\hat{X}) \succeq -\kappa I_{nr}. \quad (8)$$

By Theorem 5, the error $\|\hat{X}\hat{X}^T - M^*\|_F$ in both (7a) and (7b) is induced by the measurement noise characterized by ϵ , together with the inaccuracy of the local search algorithm captured by κ . $\hat{X}\hat{X}^T$ will be close to the ground truth if ϵ and κ are relatively small, and the contribution from κ to the bounds is exactly half of that from ϵ . Since ϵ is a constant which cannot be decreased during the iterations, it is reasonable to design an algorithm to find an approximate solution $\hat{X}$ satisfying (8) with $\epsilon = \kappa/2$ to strike a balance between the probabilistic lower bound and the required number of iterations.

To simplify the analysis, our strict saddle property in Theorem 5 is different from the traditional ones which are usually stated as that $\|\hat{X}\hat{X}^T - M^*\|_F$ is small if $\hat{X}$ satisfies

$$\|\nabla f(\hat{X})\| \leq \tilde{\kappa}, \quad \nabla^2 f(\hat{X}) \succeq -\tilde{\kappa} I_{nr}, \quad (9)$$

for a sufficiently small $\tilde{\kappa} > 0$. In [26], it is proven that the perturbed gradient descent method with an arbitrary initialization will find a solution $\hat{X}$ satisfying (9) with a high probability in $O(\text{poly}(1/\tilde{\kappa}))$ iterations. Using the assumption that $r^* > 0$ and thus $0_{n \times n}$ is not the ground truth, in the proof of the next theorem, we will see that the conditions in (9) will imply the ones in (8) if $\tilde{\kappa}$ is chosen appropriately. This establishes the global convergence for the noisy low-rank matrix recovery problems in both the exact parameterized regime and the overparameterized regime.

Theorem 6. *Let $D \in (0, 1]$ and $\lambda_0 > 0$ be constants such that*

$$\lambda_{\min}(\nabla^2 f_0(\hat{X})) < -\lambda_0$$

holds for every matrix $\hat{X} \in \mathbb{R}^{n \times r}$ with $\|\hat{X}\|_2 < D$, where f_0 is the noiseless objective function, i.e. the function f in (2) satisfying $w = 0$. Assume that the linear operator $\mathcal{A}$ satisfies the δ -RIP $_{r+r^}$ property with $\delta < 1/(1 + \sqrt{r^*/r})$. For every $\epsilon \in (0, \lambda_0)$, the perturbed gradient descent method will find a solution $\hat{X}$ satisfying either of the two inequalities (3a) and (3b) with probability at least $\mathbb{P}(\|\mathbf{A}^T w\| \leq \epsilon/2)$ in $O(\text{poly}(1/\epsilon))$ number of iterations.*

Note that the constants D and λ_0 stated above always exist. As we mentioned before Theorem 6, $0_{n \times n}$ is not the ground truth. Since the RIP assumption implies the unique recovery, $0_{n \times r}$ is not a global minimizer of the function f_0 , and Theorem 2 (or the previous equivalent result in the noiseless case) implies that $0_{n \times r}$ is also not a local minimizer of f_0 . As $\nabla f_0(0_{n \times r}) = 0$, $\nabla^2 f_0(0_{n \times r})$ cannot be positive semidefinite, which implies the existence of D and λ_0 by a smoothness argument. Moreover, the two constants can be directly calculated out after the measurement operator $\mathcal{A}$ is explicitly given.

3 Proofs of main results

Before presenting the proofs, we first compute the gradient and the Hessian of the objective function $f(\hat{X})$ of the problem (2):

$$\begin{aligned} \nabla f(\hat{X}) &= \hat{\mathbf{X}}^T \mathbf{A}^T (\mathbf{A} \mathbf{e} + w), \\ \nabla^2 f(\hat{X}) &= 2I_r \otimes \text{mat}_S(\mathbf{A}^T (\mathbf{A} \mathbf{e} + w)) + \hat{\mathbf{X}}^T \mathbf{A}^T \mathbf{A} \hat{\mathbf{X}}, \end{aligned}$$

where

$$\mathbf{e} = \text{vec}(\hat{X}\hat{X}^T - M^*),$$

and $\hat{\mathbf{X}} \in \mathbb{R}^{n^2 \times nr}$ is the matrix satisfying

$$\hat{\mathbf{X}} \text{vec}(U) = \text{vec}(\hat{X}U^T + U\hat{X}^T), \quad \forall U \in \mathbb{R}^{n \times r}.$$

3.1 Proofs of Theorem 2, Theorem 5 and Theorem 6

The first step in the proofs of Theorem 2 and Theorem 5 is to derive necessary conditions for a matrix $\hat{X} \in \mathbb{R}^{n \times r}$ to be an approximate second-order critical point, which depend on the linear operator $\mathcal{A}$, the noise $w \in \mathbb{R}^m$, the solution $\hat{X}$, and the parameter κ characterizing how close $\hat{X}$ is to a true second-order critical point.

Lemma 1. *Given $\kappa \geq 0$, assume that $\hat{X} \in \mathbb{R}^{n \times r}$ satisfies*

$$\|\nabla f(\hat{X})\| \leq \kappa \|\hat{X}\|_2, \quad \nabla^2 f(\hat{X}) \succeq -\kappa I_{nr}.$$

Then, it must also satisfy the following inequalities:

$$\|\hat{\mathbf{X}}^T \mathbf{H} \mathbf{e}\| \leq (2\|\mathbf{A}^T w\| + \kappa) \|\hat{X}\|_2, \quad (10a)$$

$$2I_r \otimes \text{mat}_S(\mathbf{H} \mathbf{e}) + \hat{\mathbf{X}}^T \mathbf{H} \hat{\mathbf{X}} \succeq -(2\|\mathbf{A}^T w\| + \kappa) I_{nr}, \quad (10b)$$

where $\mathbf{H} = \mathbf{A}^T \mathbf{A}$.

Proof. To obtain condition (10a), notice that $\|\nabla f(\hat{X})\| \leq \kappa \|\hat{X}\|_2$ implies that

$$\|\hat{\mathbf{X}}^T \mathbf{H} \mathbf{e}\| \leq \|\hat{\mathbf{X}}^T \mathbf{A}^T w\| + \|\nabla f(\hat{X})\| \leq \|\hat{\mathbf{X}}\|_2 \|\mathbf{A}^T w\| + \kappa \|\hat{X}\|_2 \leq (2\|\mathbf{A}^T w\| + \kappa) \|\hat{X}\|_2,$$

in which the last inequality is due to

$$\|\hat{\mathbf{X}} \text{vec}(U)\| = \|\hat{X}U^T + U\hat{X}^T\|_F \leq 2\|\hat{X}\|_2 \|U\|_F,$$

for every $U \in \mathbb{R}^{n \times r}$. Similarly, $\nabla^2 f(\hat{X}) \succeq -\kappa I_{nr}$ implies that

$$2I_r \otimes \text{mat}_S(\mathbf{H} \mathbf{e}) + \hat{\mathbf{X}}^T \mathbf{H} \hat{\mathbf{X}} \succeq -2I_r \otimes \text{mat}_S(\mathbf{A}^T w) - \kappa I_{nr}.$$

On the other hand, the eigenvalues of $I_r \otimes \text{mat}_S(\mathbf{A}^T w)$ are the same as those of $\text{mat}_S(\mathbf{A}^T w)$, and each eigenvalue $\lambda_i(\text{mat}_S(\mathbf{A}^T w))$ of the latter matrix further satisfies

$$|\lambda_i(\text{mat}_S(\mathbf{A}^T w))| \leq \|\text{mat}_S(\mathbf{A}^T w)\|_F \leq \|\mathbf{A}^T w\|,$$

which proves condition (10b). $\square$

If $\hat{X}$ is a local minimizer of the problem (2), Lemma 1 shows that $\hat{X}$ satisfies the inequalities (10a) and (10b) with $\kappa = 0$. Similarly, Theorem 2 can also be regarded as a special case of Theorem 5 with $\kappa = 0$. The proofs of these two theorems consist of inspecting two cases. The following lemma deals with the first case in which $\hat{X}$ is an approximate second-order critical point with $\sigma_r(\hat{X})$ being close to zero.

Lemma 2. *Assume that the linear operator $\mathcal{A}$ satisfies the δ -RIP $_{r+r^*}$ property. Given $\hat{X} \in \mathbb{R}^{n \times r}$ and arbitrary constants $\epsilon > 0$ and $\kappa \geq 0$, the inequalities*

$$\sigma_r(\hat{X}) \leq \sqrt{\frac{\epsilon + \kappa/2}{1 + \delta}}, \quad \|\nabla f(\hat{X})\| \leq \kappa \|\hat{X}\|_2, \quad \nabla^2 f(\hat{X}) \succeq -\kappa I_{nr}$$

and $\|\mathbf{A}^T w\| \leq \epsilon$ will together imply the inequality (7a).

Proof. Let $G = \text{mat}_S(\mathbf{H} \mathbf{e})$ and $u \in \mathbb{R}^n$ be a unit eigenvector of G corresponding to its minimum eigenvalue, i.e.,

$$\|u\| = 1, \quad Gu = \lambda_{\min}(G)u.$$

In addition, let $v \in \mathbb{R}^r$ be a singular vector of $\hat{X}$ such that

$$\|v\| = 1, \quad \|\hat{X}v\| = \sigma_r(\hat{X}).$$

Let $\mathbf{U} = \text{vec}(uv^T)$. Then, $\|\mathbf{U}\| \leq 1$ and (10b) imply that

$$\begin{aligned}
-2\epsilon - \kappa &\leq 2\mathbf{U}^T(I_r \otimes \text{mat}_S(\mathbf{He}))\mathbf{U} + \mathbf{U}^T \hat{\mathbf{X}}^T \mathbf{H} \hat{\mathbf{X}} \mathbf{U} \\
&\leq 2\text{tr}(vu^T G uv^T) + (1 + \delta) \|\hat{\mathbf{X}} vu^T + uv^T \hat{\mathbf{X}}^T\|_F^2 \\
&\leq 2\lambda_{\min}(G) + 4(1 + \delta) \sigma_r(\hat{\mathbf{X}})^2 \\
&\leq 2\lambda_{\min}(G) + 4\epsilon + 2\kappa.
\end{aligned} \tag{11}$$

On the other hand,

$$\begin{aligned}
(1 - \delta) \|\hat{\mathbf{X}} \hat{\mathbf{X}}^T - M^*\|_F^2 &\leq \mathbf{e}^T \mathbf{He} = \text{vec}(\hat{\mathbf{X}} \hat{\mathbf{X}}^T)^T \mathbf{He} - \text{vec}(M^*)^T \mathbf{He} \\
&= \frac{1}{2} \text{vec}(\hat{\mathbf{X}})^T \hat{\mathbf{X}}^T \mathbf{He} - \langle M^*, \text{mat}_S(\mathbf{He}) \rangle \\
&\leq \frac{1}{2} \|\hat{\mathbf{X}}\|_F \|\hat{\mathbf{X}}^T \mathbf{He}\| + \left(3\epsilon + \frac{3\kappa}{2}\right) \text{tr}(M^*) \\
&\leq \left(\epsilon + \frac{\kappa}{2}\right) \|\hat{\mathbf{X}}\|_F^2 + \left(3\epsilon + \frac{3\kappa}{2}\right) \text{tr}(M^*),
\end{aligned}$$

in which the second last inequality is due to (11) and the last inequality is due to (10a). Furthermore, the right-hand side of the above inequality can be relaxed as

$$\begin{aligned}
&\left(\epsilon + \frac{\kappa}{2}\right) \|\hat{\mathbf{X}}\|_F^2 + \left(3\epsilon + \frac{3\kappa}{2}\right) \text{tr}(M^*) \\
&\leq \left(\epsilon + \frac{\kappa}{2}\right) \sqrt{r} \|\hat{\mathbf{X}} \hat{\mathbf{X}}^T\|_F + \left(3\epsilon + \frac{3\kappa}{2}\right) \sqrt{r^*} \|M^*\|_F \\
&\leq \left(\epsilon + \frac{\kappa}{2}\right) \sqrt{r} \|\hat{\mathbf{X}} \hat{\mathbf{X}}^T - M^*\|_F + \left(\epsilon + \frac{\kappa}{2}\right) \sqrt{r} \|M^*\|_F + \left(3\epsilon + \frac{3\kappa}{2}\right) \sqrt{r} \|M^*\|_F \\
&= \left(\epsilon + \frac{\kappa}{2}\right) \sqrt{r} \|\hat{\mathbf{X}} \hat{\mathbf{X}}^T - M^*\|_F + (4\epsilon + 2\kappa) \sqrt{r} \|M^*\|_F,
\end{aligned}$$

which leads to the inequality (7a). $\square$

The remaining case with

$$\sigma_r(\hat{\mathbf{X}}) > \sqrt{\frac{\epsilon + \kappa/2}{1 + \delta}}$$

will be handled in the following lemma using a different method.

Lemma 3. Assume that the linear operator $\mathcal{A}$ satisfies the δ -RIP $_{r+r^*}$ property with $\delta < 1/(1 + \sqrt{r^*/r})$. Given $\hat{\mathbf{X}} \in \mathbb{R}^{n \times r}$ and arbitrary constants $\epsilon > 0$ and $\kappa \geq 0$, the inequalities

$$\sigma_r(\hat{\mathbf{X}}) > \sqrt{\frac{\epsilon + \kappa/2}{1 + \delta}}, \quad \|\nabla f(\hat{\mathbf{X}})\| \leq \kappa \|\hat{\mathbf{X}}\|_2, \quad \nabla^2 f(\hat{\mathbf{X}}) \succeq -\kappa I_{nr}$$

and $\|\mathbf{A}^T w\| \leq \epsilon$ will together imply the inequality (7b).

The proofs of both Lemma 3 and the local guarantee in Theorem 4 later generalize the proof of the absence of spurious local minima for the noiseless problem in [15]. Our innovation here is to develop new techniques to analyze approximate optimality conditions for the solutions because unlike the noiseless problem the local minimizers of the noisy one are only approximate second-order critical points of the distance function $\|\mathcal{A}(X X^T) - b\|^2$. For a fixed solution $\hat{\mathbf{X}}$ and noise w , one can find an operator $\hat{\mathcal{A}}$ satisfying the δ -RIP $_{r+r^*}$ property with the smallest possible δ such that $\hat{\mathbf{X}}$ and $\hat{\mathcal{A}}$ satisfy the necessary conditions stated in Lemma 1. Let $\delta^*(\hat{\mathbf{X}})$ be the RIP constant of the found measurement operator $\hat{\mathcal{A}}$ in the worst-case scenario. Then, if $\hat{\mathbf{X}}$ in Lemma 3 is a solution of the current problem with the linear operator $\mathcal{A}$ satisfying the δ -RIP $_{r+r^*}$ property, it holds that $\delta \geq \delta^*(\hat{\mathbf{X}})$, which can further lead to an upper bound on the distance $\|\hat{\mathbf{X}} \hat{\mathbf{X}}^T - M^*\|_F$.

Proof of Lemma 3. The $\delta^*(\hat{X})$ defined above is the optimal value of the following optimization problem:

$$\begin{aligned}
& \min_{\delta, \hat{\mathbf{H}}} \quad \delta \\
& \text{s.t.} \quad \|\hat{\mathbf{X}}^T \hat{\mathbf{H}} \mathbf{e}\| \leq (2\epsilon + \kappa) \|\hat{X}\|_2, \\
& \quad 2I_r \otimes \text{mat}_S(\hat{\mathbf{H}} \mathbf{e}) + \hat{\mathbf{X}}^T \hat{\mathbf{H}} \hat{\mathbf{X}} \succeq -(2\epsilon + \kappa) I_{nr}, \\
& \quad \hat{\mathbf{H}} \text{ is symmetric and satisfies the } \delta\text{-RIP}_{r+r^*} \text{ property.}
\end{aligned} \tag{12}$$

Here, a matrix $\hat{\mathbf{H}} \in \mathbb{R}^{n^2 \times n^2}$ is said to satisfy the δ -RIP $_{r+r^*}$ property if

$$(1 - \delta) \|\mathbf{U}\|^2 \leq \mathbf{U}^T \hat{\mathbf{H}} \mathbf{U} \leq (1 + \delta) \|\mathbf{U}\|^2$$

holds for every matrix $U \in \mathbb{R}^{n \times n}$ with $\text{rank}(U) \leq r + r^*$ and $\mathbf{U} = \text{vec}(U)$. Obviously, for a linear operator $\hat{\mathcal{A}}$, $\hat{\mathbf{H}} = \hat{\mathcal{A}}^T \hat{\mathcal{A}}$ satisfies the δ -RIP $_{r+r^*}$ property if and only if $\hat{\mathcal{A}}$ satisfies the δ -RIP $_{r+r^*}$ property. By the discussion above, we have $\delta \geq \delta^*(\hat{X})$.

However, since problem (12) is non-convex due to the RIP constraint, we instead solve the following convex reformulation:

$$\begin{aligned}
& \min_{\delta, \hat{\mathbf{H}}} \quad \delta \\
& \text{s.t.} \quad \|\hat{\mathbf{X}}^T \hat{\mathbf{H}} \mathbf{e}\| \leq (2\epsilon + \kappa) \|\hat{X}\|_2, \\
& \quad 2I_r \otimes \text{mat}_S(\hat{\mathbf{H}} \mathbf{e}) + \hat{\mathbf{X}}^T \hat{\mathbf{H}} \hat{\mathbf{X}} \succeq -(2\epsilon + \kappa) I_{nr}, \\
& \quad (1 - \delta) I_{n^2} \preceq \hat{\mathbf{H}} \preceq (1 + \delta) I_{n^2}.
\end{aligned} \tag{13}$$

Lemma 14 in [17] proves that problem (12) and problem (13) have the same optimal value. The remaining step is to solve the optimization problem (13) for given $\hat{X}$, ϵ and κ . To further simplify (13), one can replace its decision variable δ with η and introduce the following optimization problem:

$$\begin{aligned}
& \max_{\eta, \hat{\mathbf{H}}} \quad \eta \\
& \text{s.t.} \quad \|\hat{\mathbf{X}}^T \hat{\mathbf{H}} \mathbf{e}\| \leq (2\epsilon + \kappa) \|\hat{X}\|_2, \\
& \quad 2I_r \otimes \text{mat}_S(\hat{\mathbf{H}} \mathbf{e}) + \hat{\mathbf{X}}^T \hat{\mathbf{H}} \hat{\mathbf{X}} \succeq -(2\epsilon + \kappa) I_{nr}, \\
& \quad \eta I_{n^2} \preceq \hat{\mathbf{H}} \preceq I_{n^2}.
\end{aligned} \tag{14}$$

Given any feasible solution $(\delta, \hat{\mathbf{H}})$ to (13), the tuple

$$\left(\frac{1 - \delta}{1 + \delta}, \frac{1}{1 + \delta} \hat{\mathbf{H}} \right)$$

is a feasible solution to problem (14). Therefore, if the optimal value of (14) is denoted as $\eta^*(\hat{X})$, then it holds that

$$\eta^*(\hat{X}) \geq \frac{1 - \delta^*(\hat{X})}{1 + \delta^*(\hat{X})} \geq \frac{1 - \delta}{1 + \delta}. \tag{15}$$

In the remaining part, we will prove the following upper bound on $\eta^*(\hat{X})$:

$$\eta^*(\hat{X}) \leq \frac{\sqrt{r^*/r}}{2 + \sqrt{r^*/r}} + \frac{(2\epsilon + \kappa)\sqrt{r^*} + 2\sqrt{(2\epsilon + \kappa)(1 + \delta)} \|\hat{X}\|_2}{\|\mathbf{e}\|}. \tag{16}$$

The inequality (7b) is a consequence of (15), (16) and the inequality

$$\|\hat{X}\|_2 \leq \|\hat{X} \hat{X}^T\|_F^{1/2} \leq \|\hat{X} \hat{X}^T - M^*\|_F^{1/2} + \|M^*\|_F^{1/2}.$$

The proof of the upper bound (16) can be completed by finding a feasible solution to the dual problem of (14):

$$\begin{aligned}
& \min_{U_1, U_2, W, G, \lambda, y} \quad \text{tr}(U_2) + \langle \hat{\mathbf{X}}^T \hat{\mathbf{X}}, W \rangle + (2\epsilon + \kappa) \text{tr}(W) + (2\epsilon + \kappa)^2 \|\hat{X}\|_2^2 \lambda + \text{tr}(G) \\
& \text{s. t.} \quad \text{tr}(U_1) = 1, \\
& \quad (\hat{\mathbf{X}}y - w)\mathbf{e}^T + \mathbf{e}(\hat{\mathbf{X}}y - w)^T = U_1 - U_2, \\
& \quad \begin{bmatrix} G & -y \\ -y^T & \lambda \end{bmatrix} \succeq 0, \\
& \quad U_1 \succeq 0, \quad U_2 \succeq 0, \quad W = \begin{bmatrix} W_{1,1} & \cdots & W_{r,1}^T \\ \vdots & \ddots & \vdots \\ W_{r,1} & \cdots & W_{r,r} \end{bmatrix} \succeq 0, \\
& \quad w = \sum_{i=1}^r \text{vec}(W_{i,i}).
\end{aligned} \tag{17}$$

Before describing the choice of the dual feasible solution, we need to represent the error vector $\mathbf{e}$ in a different form. Let $\mathcal{P} \in \mathbb{R}^{n \times n}$ be the orthogonal projection matrix onto the range of $\hat{X}$, and $\mathcal{P}_\perp \in \mathbb{R}^{n \times n}$ be the orthogonal projection matrix onto the orthogonal complement of the range of $\hat{X}$. Furthermore, let $Z \in \mathbb{R}^{n \times r}$ be a matrix satisfying $ZZ^T = M^*$. Then, $\text{rank}(Z) = r^*$, and Z can be decomposed as $Z = \mathcal{P}Z + \mathcal{P}_\perp Z$, so there exists a matrix $R \in \mathbb{R}^{r \times r}$ such that $\mathcal{P}Z = \hat{X}R$. Note that

$$ZZ^T = \mathcal{P}ZZ^T\mathcal{P} + \mathcal{P}ZZ^T\mathcal{P}_\perp + \mathcal{P}_\perp ZZ^T\mathcal{P} + \mathcal{P}_\perp ZZ^T\mathcal{P}_\perp.$$

Thus, if we choose

$$\hat{Y} = \frac{1}{2}\hat{X} - \frac{1}{2}\hat{X}RR^T - \mathcal{P}_\perp ZR^T, \quad \hat{y} = \text{vec}(\hat{Y}), \tag{18}$$

then it can be verified that

$$\begin{aligned}
& \hat{X}\hat{Y}^T + \hat{Y}\hat{X}^T - \mathcal{P}_\perp ZZ^T\mathcal{P}_\perp = \hat{X}\hat{X}^T - ZZ^T, \\
& \langle \hat{X}\hat{Y}^T + \hat{Y}\hat{X}^T, \mathcal{P}_\perp ZZ^T\mathcal{P}_\perp \rangle = 0.
\end{aligned} \tag{19}$$

Moreover, we have

$$\begin{aligned}
\|\hat{X}\hat{Y}^T + \hat{Y}\hat{X}^T\|_F^2 &= 2\text{tr}(\hat{X}^T\hat{X}\hat{Y}^T\hat{Y}) + \text{tr}(\hat{X}^T\hat{Y}\hat{X}^T\hat{Y}) + \text{tr}(\hat{Y}^T\hat{X}\hat{Y}^T\hat{X}) \\
&\geq 2\text{tr}(\hat{X}^T\hat{X}\hat{Y}^T\hat{Y}) \geq 2\sigma_r(\hat{X})^2\|\hat{Y}\|_F^2,
\end{aligned} \tag{20}$$

in which the first inequality is due to

$$\text{tr}(\hat{X}^T\hat{Y}\hat{X}^T\hat{Y}) = \frac{1}{4}\text{tr}((\hat{X}^T\hat{X}(I_r - RR^T))^2) = \frac{1}{4}\text{tr}((\hat{X}(I_r - RR^T)\hat{X}^T)^2) \geq 0.$$

Assume first that $Z_\perp = \mathcal{P}_\perp Z \neq 0$. The other case will be handled at the end of this proof. In the case when $Z_\perp \neq 0$, we also have $\hat{X}\hat{Y}^T + \hat{Y}\hat{X}^T \neq 0$. Otherwise, the inequality (20) and the assumption $\sigma_r(\hat{X}) > 0$ imply that $\hat{Y} = 0$. The orthogonality and the definition of $\hat{Y}$ in (18) then give rise to

$$\hat{X} - \hat{X}RR^T = 0, \quad \mathcal{P}_\perp ZR^T = 0.$$

The first equation above implies that R is invertible since $\hat{X}$ has full column rank, which contradicts $Z_\perp \neq 0$. Now, define the unit vectors

$$\hat{u}_1 = \frac{\hat{\mathbf{X}}\hat{y}}{\|\hat{\mathbf{X}}\hat{y}\|}, \quad \hat{u}_2 = \frac{\text{vec}(Z_\perp Z_\perp^T)}{\|Z_\perp Z_\perp^T\|_F}.$$

Then, $\hat{u}_1 \perp \hat{u}_2$ and

$$\mathbf{e} = \|\mathbf{e}\|(\sqrt{1 - \alpha^2}\hat{u}_1 - \alpha\hat{u}_2) \tag{21}$$

with

$$\alpha = \frac{\|Z_\perp Z_\perp^T\|_F}{\|\hat{X}\hat{X}^T - ZZ^T\|_F}. \tag{22}$$

We first describe our choices of the dual variables W and y (which will be rescaled later). Let

$$\hat{X}^T \hat{X} = QSQ^T, \quad Z_\perp Z_\perp^T = PGP^T,$$

with orthogonal matrices Q, P and diagonal matrices S, G such that $S_{11} = \sigma_r(\hat{X})^2$. Fix a constant $\gamma \in [0, 1]$ that is to be determined and define

$$V_i = k^{1/2} G_{ii}^{1/2} P E_{i1} Q^T, \quad \forall i \in \{1, \dots, r\},$$

$$W = \sum_{i=1}^r \text{vec}(V_i) \text{vec}(V_i)^T, \quad y = l \hat{y},$$

with $\hat{y}$ defined in (18) and

$$k = \frac{\gamma}{\|\mathbf{e}\| \|Z_\perp Z_\perp^T\|_F}, \quad l = \frac{\sqrt{1-\gamma^2}}{\|\mathbf{e}\| \|\hat{\mathbf{X}} \hat{y}\|}.$$

Here, E_{ij} is the elementary matrix of size $n \times r$ with the (i, j) -entry being 1. By our construction, $\hat{X}^T V_i = 0$, which implies that

$$\begin{aligned} \langle \hat{\mathbf{X}}^T \hat{\mathbf{X}}, W \rangle &= \sum_{i=1}^r \|\hat{X} V_i^T + V_i \hat{X}^T\|_F^2 = 2 \sum_{i=1}^r \text{tr}(\hat{X}^T \hat{X} V_i^T V_i) \\ &= 2k \sigma_r(\hat{X})^2 \sum_{i=1}^r G_{ii} = 2\beta\gamma, \end{aligned} \tag{23}$$

with

$$\beta = \frac{\sigma_r(\hat{X})^2 \text{tr}(Z_\perp Z_\perp^T)}{\|\hat{X} \hat{X}^T - Z Z^T\|_F \|Z_\perp Z_\perp^T\|_F}. \tag{24}$$

In addition,

$$\text{tr}(W) = \sum_{i=1}^r \|V_i\|_F^2 = k \sum_{i=1}^r G_{ii} = k \text{tr}(Z_\perp Z_\perp^T) \leq \frac{\sqrt{r^*}}{\|\mathbf{e}\|}, \tag{25}$$

and

$$w = \sum_{i=1}^r \text{vec}(W_{i,i}) = \sum_{i=1}^r V_i V_i^T = k Z_\perp Z_\perp^T.$$

Therefore,

$$\hat{\mathbf{X}} y - w = \frac{1}{\|\mathbf{e}\|} (\sqrt{1-\gamma^2} \hat{u}_1 - \gamma \hat{u}_2),$$

which together with (21) implies that

$$\|\mathbf{e}\| \|\hat{\mathbf{X}} y - w\| = 1, \quad \langle \mathbf{e}, \hat{\mathbf{X}} y - w \rangle = \gamma \alpha + \sqrt{1-\gamma^2} \sqrt{1-\alpha^2} = \psi(\gamma). \tag{26}$$

Next, the inequality (20) and the assumption on $\sigma_r(\hat{X})$ imply that

$$(4\epsilon + 2\kappa) \|y\| \leq \frac{\sqrt{1-\gamma^2} (4\epsilon + 2\kappa)}{\sqrt{2} \sigma_r(\hat{X}) \|\mathbf{e}\|} \leq \frac{2\sqrt{(2\epsilon + \kappa)(1+\delta)}}{\|\mathbf{e}\|}. \tag{27}$$

Define

$$M = (\hat{\mathbf{X}} y - w) \mathbf{e}^T + \mathbf{e} (\hat{\mathbf{X}} y - w)^T,$$

and decompose M as $M = [M]_+ - [M]_-$ in which both $[M]_+ \succeq 0$ and $[M]_- \succeq 0$. Let θ be the angle between $\mathbf{e}$ and $\hat{\mathbf{X}} y - w$. By Lemma 14 in [15], we have

$$\text{tr}([M]_+) = \|\mathbf{e}\| \|\hat{\mathbf{X}} y - w\| (1 + \cos \theta), \quad \text{tr}([M]_-) = \|\mathbf{e}\| \|\hat{\mathbf{X}} y - w\| (1 - \cos \theta).$$

Now, one can verify that $(U_1^*, U_2^*, W^*, G^*, \lambda^*, y^*)$ defined as

$$\begin{aligned} U_1^* &= \frac{[M]_+}{\text{tr}([M]_+)}, \quad U_2^* = \frac{[M]_-}{\text{tr}([M]_-)}, \quad y^* = \frac{y}{\text{tr}([M]_+)}, \\ W^* &= \frac{W}{\text{tr}([M]_+)}, \quad \lambda^* = \frac{\|y^*\|}{(2\epsilon + \kappa) \|\hat{X}\|_2}, \quad G^* = \frac{1}{\lambda^*} y^* y^{*T} \end{aligned}$$

forms a feasible solution to the dual problem (17) whose objective value is equal to

$$\frac{\text{tr}([M]_-) + \langle \hat{\mathbf{X}}^T \hat{\mathbf{X}}, W \rangle + (2\epsilon + \kappa) \text{tr}(W) + (4\epsilon + 2\kappa) \|\hat{X}\|_2 \|y\|}{\text{tr}([M]_+)}$$

Substituting (23), (25), (26), and (27) into the above equation, we obtain

$$\begin{aligned} \eta^*(\hat{X}) &\leq \frac{2\beta\gamma + 1 - \psi(\gamma) + ((2\epsilon + \kappa)\sqrt{r^*} + 2\sqrt{(2\epsilon + \kappa)(1 + \delta)}\|\hat{X}\|_2)/\|\mathbf{e}\|}{1 + \psi(\gamma)} \\ &\leq \frac{2\beta\gamma + 1 - \psi(\gamma)}{1 + \psi(\gamma)} + \frac{(2\epsilon + \kappa)\sqrt{r^*} + 2\sqrt{(2\epsilon + \kappa)(1 + \delta)}\|\hat{X}\|_2}{\|\mathbf{e}\|}. \end{aligned}$$

Choosing the best value of the parameter $\gamma \in [0, 1]$ to minimize the far right-side of the above inequality leads to

$$\frac{2\beta\gamma + 1 - \psi(\gamma)}{1 + \psi(\gamma)} \leq \eta_0(\hat{X}),$$

with

$$\eta_0(\hat{X}) := \begin{cases} \frac{1 - \sqrt{1 - \alpha^2}}{1 + \sqrt{1 - \alpha^2}}, & \text{if } \beta \geq \frac{\alpha}{1 + \sqrt{1 - \alpha^2}}, \\ \frac{\beta(\alpha - \beta)}{1 - \beta\alpha}, & \text{if } \beta \leq \frac{\alpha}{1 + \sqrt{1 - \alpha^2}}. \end{cases}$$

Here, α and β are defined in (22) and (24), respectively. In the proof of Theorem 1.2 in [20], it is shown that

$$\frac{1 - \eta_0(\hat{X})}{1 + \eta_0(\hat{X})} \geq \frac{1}{1 + \sqrt{r^*/r}}$$

for every $\hat{X}$ with $\hat{X}\hat{X}^T \neq ZZ^T$ and $\text{rank}(Z) = r^*$. Therefore,

$$\eta_0(\hat{X}) \leq \frac{\sqrt{r^*/r}}{2 + \sqrt{r^*/r}},$$

which gives the upper bound (16).

Finally, we still need to deal with the case when $\mathcal{P}_\perp Z = 0$. In this case, we know that $\hat{\mathbf{X}}\hat{y} = \mathbf{e}$ with $\hat{y}$ defined in (18). Then, it is easy to check that $(U_1^*, U_2^*, W^*, G^*, \lambda^*, y^*)$ defined as

$$\begin{aligned} U_1^* &= \frac{\mathbf{e}\mathbf{e}^T}{\|\mathbf{e}\|^2}, \quad U_2^* = 0, \quad y^* = \frac{\hat{y}}{2\|\mathbf{e}\|^2}, \\ W^* &= 0, \quad \lambda^* = \frac{\|y^*\|}{(2\epsilon + \kappa)\|\hat{X}\|_2}, \quad G^* = \frac{1}{\lambda^*} y^* y^{*T} \end{aligned}$$

forms a feasible solution to the dual problem (17) whose objective value is $(4\epsilon + 2\kappa)\|\hat{X}\|_2\|y^*\|$. By the inequality (20), we have

$$\eta^*(\hat{X}) \leq (4\epsilon + 2\kappa)\|\hat{X}\|_2\|y^*\| \leq \frac{(2\epsilon + \kappa)\|\hat{X}\|_2}{\sqrt{2}\sigma_r(\hat{X})\|\mathbf{e}\|} \leq \frac{\sqrt{(2\epsilon + \kappa)(1 + \delta)}\|\hat{X}\|_2}{\|\mathbf{e}\|}$$

Hence, the upper bound (16) still holds in this case. $\square$

Finally, Theorem 5 is a direct consequence of Lemma 2 and Lemma 3. Theorem 2 is a special case of Theorem 5 with $\kappa = 0$, and the global convergence in Theorem 6 is also a corollary of Theorem 5.

Proof of Theorem 6. Assume that $\|\mathbf{A}^T w\| \leq \epsilon/2$ holds. Since

$$\nabla^2 f(\hat{X}) = 2I_r \otimes \text{mat}_S(\mathbf{A}^T(\mathbf{A}\mathbf{e} + w)) + \hat{\mathbf{X}}^T \mathbf{A}^T \mathbf{A} \hat{\mathbf{X}} = \nabla^2 f_0(\hat{X}) + 2I_r \otimes \text{mat}_S(\mathbf{A}^T w),$$

and

$$|\lambda_{\max}(I_r \otimes \text{mat}_S(\mathbf{A}^T w))| \leq \|\text{mat}_S(\mathbf{A}^T w)\|_F \leq \|\mathbf{A}^T w\| \leq \frac{\epsilon}{2}$$

as shown in the proof of Lemma 1, by the assumption it holds that

$$\begin{aligned} -\lambda_0 &> \lambda_{\min}(\nabla^2 f_0(\hat{X})) \geq \lambda_{\min}(\nabla^2 f(\hat{X})) - 2\lambda_{\max}(I_r \otimes \text{mat}_S(\mathbf{A}^T w)) \\ &\geq \lambda_{\min}(\nabla^2 f(\hat{X})) - \epsilon \end{aligned} \quad (28)$$

for every matrix $\hat{X} \in \mathbb{R}^{n \times r}$ with $\|\hat{X}\|_2 < D$. The perturbed gradient descent method in [26] will find a solution $\hat{X}$ satisfying (9) with

$$\tilde{\kappa} = \min\{\lambda_0 - \epsilon, D\epsilon\}$$

in $O(\text{poly}(1/\epsilon))$ number of iterations. The inequality (28) and the second condition in (9) together imply that $\|\hat{X}\|_2 \geq D$, and thus the conditions in (8) are automatically satisfied for $\hat{X}$ with $\kappa = \epsilon$, which gives the desired result after we apply Theorem 5 with the original ϵ in Theorem 5 replaced with $\epsilon/2$ and κ replaced with ϵ . $\square$

3.2 Proof of Theorem 3

The proof of Theorem 3 is similar to the above proof of Lemma 3 in the situation with $\kappa = 0$, and we will only emphasize the difference here.

Proof of Theorem 3. In the case when $\hat{X} \neq 0$, after constructing the feasible solution to the dual problem (17), we have

$$\frac{1 - \delta}{1 + \delta} \leq \eta^*(\hat{X}) \leq \frac{\text{tr}([M]_-) + \langle \hat{\mathbf{X}}^T \hat{\mathbf{X}}, W \rangle + 2\epsilon \text{tr}(W) + 4\|\hat{X}\|_2 \epsilon \|y\|}{\text{tr}([M]_+)}. \quad (29)$$

Note that in the rank-1 case, one can write $\sigma_r(\hat{X}) = \|\hat{X}\|_2$ and

$$\|y\| \leq \frac{\|\hat{y}\|}{\|\mathbf{e}\| \|\hat{\mathbf{X}} \hat{y}\|} \leq \frac{1}{\sqrt{2} \|\hat{X}\|_2 \|\mathbf{e}\|},$$

in which the last inequality is due to (20). Substituting (23), (25), (26) and the above inequality into (29) and choosing an appropriate γ as shown in the proof of Lemma 3, we obtain

$$\begin{aligned} \frac{1 - \delta}{1 + \delta} \leq \eta^*(\hat{X}) &\leq \frac{2\beta\gamma + 1 - \psi(\gamma) + (2\epsilon + 2\sqrt{2}\epsilon)/\|\mathbf{e}\|}{1 + \psi(\gamma)} \\ &\leq \frac{1}{3} + \frac{2\epsilon + 2\sqrt{2}\epsilon}{\|\mathbf{e}\|}, \end{aligned}$$

which implies inequality (4) under the probabilistic event that $\|\mathbf{A}^T w\| \leq \epsilon$.

In the case when $\hat{X} = 0$, $(U_1^*, U_2^*, W^*, G^*, \lambda^*, y^*)$ with

$$\begin{aligned} U_1^* &= \frac{\mathbf{e}\mathbf{e}^T}{\|\mathbf{e}\|^2}, \quad U_2^* = 0, \quad y^* = 0, \\ W^* &= \frac{ZZ^T}{2\|\mathbf{e}\|^2}, \quad \lambda^* = 0, \quad G^* = 0 \end{aligned}$$

forms a feasible solution to the dual problem (17), which shows that

$$\frac{1 - \delta}{1 + \delta} \leq \eta^*(\hat{X}) \leq \frac{\epsilon}{\|\mathbf{e}\|}.$$

The above inequality also implies inequality (4) under the probabilistic event that $\|\mathbf{A}^T w\| \leq \epsilon$. $\square$

3.3 Proof of Theorem 4

Now, we turn to the proof of the local guarantee in Theorem 4.

Proof of Theorem 4. Similar to the proof of Theorem 2, we assume that the probabilistic event $\|\mathbf{A}^T w\| \leq \epsilon$ occurs and also break down the analysis into two cases. Consider the case when

$$\sigma_r(\hat{X}) > \sqrt{\frac{\epsilon}{1+\delta}},$$

otherwise it is already handled by Lemma 2 that leads to the inequality (6a). Here, we further relax the optimization problem (14) in Lemma 3 with $\kappa = 0$ by removing the constraint corresponding to the second-order optimality condition, which gives rise to the optimization problem

$$\begin{aligned} \max_{\eta, \hat{\mathbf{H}}} \quad & \eta \\ \text{s. t.} \quad & \|\hat{\mathbf{X}}^T \hat{\mathbf{H}} \mathbf{e}\| \leq 2\epsilon \|\hat{X}\|_2, \\ & \eta I_{n^2} \preceq \hat{\mathbf{H}} \preceq I_{n^2}. \end{aligned} \quad (30)$$

By denoting the optimal value of (30) as $\eta_f^*(\hat{X})$, it holds that

$$\eta_f^*(\hat{X}) \geq \eta^*(\hat{X}) \geq \frac{1-\delta}{1+\delta}. \quad (31)$$

Without loss of generality, we can assume that $\hat{X}$ is in the block form

$$\begin{bmatrix} X_1 \\ 0 \end{bmatrix}$$

with $X_1 \in \mathbb{R}^{r \times r}$ being invertible. Otherwise, there exists an orthogonal matrix $Q \in \mathbb{R}^{n \times n}$ such that $Q^T \hat{X}$ satisfies this requirement. We can then replace $\hat{X}$ and M^* with $Q^T \hat{X}$ and $Q^T M^* Q$ respectively due to the invariance of (30) under the transformation. Moreover, we select a matrix $Z \in \mathbb{R}^{n \times r}$ such that $Z Z^T = M^*$ and Z is in the form

$$\begin{bmatrix} Z_1^* & 0 \\ Z_2^* & 0 \end{bmatrix}$$

with $Z_1^* \in \mathbb{R}^{r \times r^*}$, $Z_2^* \in \mathbb{R}^{(n-r) \times r^*}$. Then, $\|Z_1^*(Z_2^*)^T\|_F^2 \geq \lambda_{r^*}((Z_1^*)^T Z_1^*) \|Z_2^*\|_F^2$, and

$$\begin{aligned} \lambda_{r^*}((Z_1^*)^T Z_1^*) &\geq \lambda_{r^*}((Z_1^*)^T Z_1^* + (Z_2^*)^T Z_2^*) - \lambda_1((Z_2^*)^T Z_2^*) \\ &\geq \lambda_{r^*}(Z^T Z) - \|Z_2^*(Z_2^*)^T\|_F \\ &\geq (1-\tau) \lambda_{r^*}(M^*) > 0, \end{aligned} \quad (32)$$

in which the last inequality is due to the assumption $0 < \tau < 1$ and the second last inequality is due to

$$\begin{aligned} \|Z_2^*(Z_2^*)^T\|_F &\leq (\|Z_1^*(Z_1^*)^T - X_1 X_1^T\|_F^2 + 2\|Z_1^*(Z_2^*)^T\|_F^2 + \|Z_2^*(Z_2^*)^T\|_F^2)^{1/2} \\ &= \|\hat{X} \hat{X}^T - Z Z^T\|_F \leq \tau \lambda_{r^*}(M^*). \end{aligned} \quad (33)$$

To prove the inequality (6b), we need to bound $\eta_f^*(\hat{X})$ from above, which can be achieved by finding a feasible solution to the dual problem of (30) given below:

$$\begin{aligned} \min_{U_1, U_2, G, \lambda, y} \quad & \text{tr}(U_2) + 4\epsilon^2 \|\hat{X}\|_2^2 \lambda + \text{tr}(G) \\ \text{s. t.} \quad & \text{tr}(U_1) = 1, \\ & (\hat{\mathbf{X}} y) \mathbf{e}^T + \mathbf{e} (\hat{\mathbf{X}} y)^T = U_1 - U_2, \\ & \begin{bmatrix} G & -y \\ -y^T & \lambda \end{bmatrix} \succeq 0, \\ & U_1 \succeq 0, \quad U_2 \succeq 0. \end{aligned} \quad (34)$$

If we choose $\hat{Y}$ and $\hat{y} = \text{vec}(\hat{Y})$ as (18) in the proof of Lemma 3, and let θ be the angle between $\hat{\mathbf{X}} \hat{y}$ and $\mathbf{e}$, then (19) implies that

$$\begin{aligned} \sin^2 \theta &= \frac{\|\hat{\mathbf{X}} \hat{y} - \mathbf{e}\|^2}{\|\mathbf{e}\|^2} = \frac{\|\mathcal{P}_\perp Z Z^T \mathcal{P}_\perp\|_F^2}{\|\hat{X} \hat{X}^T - Z Z^T\|_F^2} \\ &= \frac{\|Z_2^*(Z_2^*)^T\|_F^2}{\|Z_1^*(Z_1^*)^T - X_1 X_1^T\|_F^2 + 2\|Z_1^*(Z_2^*)^T\|_F^2 + \|Z_2^*(Z_2^*)^T\|_F^2}. \end{aligned}$$

Following an argument similar to the one at the end of the proof of Lemma 7 in [16] and using (32) and (33), we can obtain

$$\sin^2 \theta \leq \frac{\tau}{2 - \tau} \leq \tau. \quad (35)$$

Define

$$M = (\hat{\mathbf{X}}\hat{\mathbf{y}})\mathbf{e}^T + \mathbf{e}(\hat{\mathbf{X}}\hat{\mathbf{y}})^T,$$

and then decompose M as $M = [M]_+ - [M]_-$ with $[M]_+ \succeq 0$ and $[M]_- \succeq 0$. Then, it is easy to verify that $(U_1^*, U_2^*, G^*, \lambda^*, y^*)$ defined as

$$y^* = \frac{\hat{\mathbf{y}}}{\text{tr}([M]_+)}, \quad U_1^* = \frac{[M]_+}{\text{tr}([M]_+)}, \quad U_2^* = \frac{[M]_-}{\text{tr}([M]_+)},$$

$$G^* = \frac{y^*(y^*)^T}{\lambda^*}, \quad \lambda^* = \frac{\|y^*\|}{2\epsilon\|\hat{\mathbf{X}}\|_2}$$

forms a feasible solution to the dual problem (34) with the objective value

$$\frac{\text{tr}([M]_-) + 4\epsilon\|\hat{\mathbf{X}}\|_2\|\hat{\mathbf{y}}\|}{\text{tr}([M]_+)}. \quad (36)$$

Furthermore, it follows from the Wielandt–Hoffman theorem that

$$|\lambda_1(\hat{\mathbf{X}}\hat{\mathbf{X}}^T) - \lambda_1(M^*)| \leq \|\hat{\mathbf{X}}\hat{\mathbf{X}}^T - M^*\|_F \leq \tau\lambda_{r^*}(M^*).$$

Thus, using the above inequality, the inequality (20) and the assumption on $\sigma_r(\hat{\mathbf{X}})$, we have

$$\frac{2\|\hat{\mathbf{X}}\|_2\|\hat{\mathbf{y}}\|}{\|\hat{\mathbf{X}}\hat{\mathbf{y}}\|} \leq \frac{2\|\hat{\mathbf{X}}\|_2}{\sqrt{2}\sigma_r(\hat{\mathbf{X}})} \leq \sqrt{\frac{2(1+\delta)(\lambda_1(M^*) + \tau\lambda_{r^*}(M^*))}{\epsilon}} = C(\tau, M^*)\sqrt{\frac{1+\delta}{\epsilon}}. \quad (37)$$

Next, according to Lemma 14 of [15], one can write

$$\text{tr}([M]_+) = \|\hat{\mathbf{X}}\hat{\mathbf{y}}\|\|\mathbf{e}\|(1 + \cos \theta),$$

$$\text{tr}([M]_-) = \|\hat{\mathbf{X}}\hat{\mathbf{y}}\|\|\mathbf{e}\|(1 - \cos \theta).$$

Substituting the above two equations and (37) into the dual objective value (36), one can obtain

$$\eta_f^*(\hat{\mathbf{X}}) \leq \frac{1 - \cos \theta + 2\sqrt{\epsilon(1+\delta)}C(\tau, M^*)/\|\mathbf{e}\|}{1 + \cos \theta},$$

which together with (31) implies that

$$\|\mathbf{e}\| \leq \sqrt{\epsilon}(1+\delta)^{3/2}C(\tau, M^*)(\cos \theta - \delta)^{-1}.$$

The inequality (6b) can then be proved by combining the above inequality and (35). $\square$

4 Numerical illustration

In the next section, we will empirically study the developed probabilistic guarantees and demonstrate the distance $\|\hat{\mathbf{X}}\hat{\mathbf{X}}^T - M^*\|_F$ between any local minimizer $\hat{\mathbf{X}}$ and the ground truth M^* as well as the value of the RIP constant δ required to be satisfied by the linear operator $\mathcal{A}$, in both the exact parameterized regime and the overparameterized regime.

Before delving into the numerical illustration, note that the probability $\mathbb{P}(\|\mathbf{A}^T w\| \leq \epsilon)$ used in our theorems can be exactly calculated as long as the distribution of the noise w is explicitly given. On the other hand, if we only have partial information for the distribution of w , a lower bound for the probability $\mathbb{P}(\|\mathbf{A}^T w\| \leq \epsilon)$ can still be obtained using certain tail bounds. For example, if w is a σ -sub-Gaussian vector, then applying Lemma 1 in [30] leads to

$$1 - 2e^{-\frac{w_0^2}{16m\sigma^2}} \leq \mathbb{P}(\|w\| \leq w_0) \leq \mathbb{P}(\|\mathbf{A}^T w\| \leq \epsilon),$$

where $w_0 = \epsilon/\|\mathbf{A}\|_2$.

For numerical illustration, assume that $n = 50$, $m = 10$ and $\|\mathbf{A}\|_2 \leq 2$, while the noise w is a 0.05-sub-Gaussian vector. We also assume that the search rank r is 10, $\|M^*\|_F = 3.3$, the largest eigenvalue of M^* is 1.5 and its smallest nonzero eigenvalue is 1.

First, we explore the two inequalities (3a) and (3b) in Theorem 2 to obtain two upper bounds on $\|\hat{X}\hat{X}^T - M^*\|_F$, where $\hat{X}$ denotes any arbitrary (worst) local minimizer. For both the exact parameterized case with $r = r^* = 10$ and the overparameterized case with $r = 10$ and $r^* = 2$, Fig. 2 gives the contour plots of the two upper bounds on $\|\hat{X}\hat{X}^T - M^*\|_F$, which hold with the probability given on the y -axis and the RIP constant δ from 0 to $1/(1 + \sqrt{r^*/r})$ given on the x -axis. Regardless of the parameterization type, when δ is close to the maximum allowable value $1/(1 + \sqrt{r^*/r})$, (3a) usually dominates the bound, and as δ decreases to 0, (3b) dominates instead. Furthermore, in the overparameterized regime, (3b) leads to a tighter bound, while (3a) remains the same. A similar visualization of the upper bounds given by Theorem 4 for the distance $\|\hat{X}\hat{X}^T - M^*\|_F$ is also presented in Fig. 3. We only show the exact parameterized case here because the result is the same for the overparameterized one. It can be observed that (6b) dominates the bound when δ is closer to 1 and (6a) dominates when δ is closer to 0.

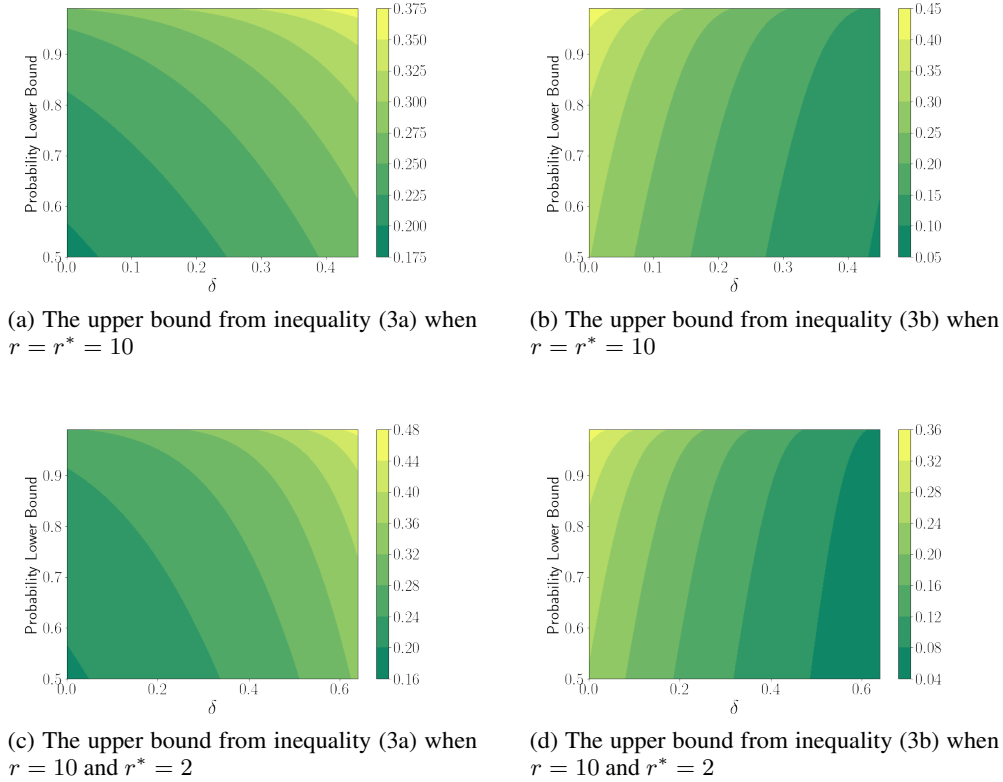


Figure 2: Comparison of the upper bounds given by Theorem 2 for the distance $\|\hat{X}\hat{X}^T - M^*\|_F$ with $\hat{X}$ being an arbitrary local minimizer

Next, we compare the bounds given by Theorem 2 and Theorem 4. Fig. 4 shows the contour plots of the maximum RIP constant δ that is necessary to guarantee that each local minimizer $\hat{X}$ (satisfying the inequality (5) when Theorem 4 is applied) lies within a certain neighborhood of the ground truth (measured via the distance $\|\hat{X}\hat{X}^T - M^*\|_F$ on the x -axis) with a given probability on the y -axis, as implied by the respective global and local guarantees. Fig. 4 clearly shows how a smaller RIP constant δ leads to a tighter bound on the distance $\|\hat{X}\hat{X}^T - M^*\|_F$ with a higher probability. In addition, the local guarantee generally requires a looser RIP assumption as it still holds even when $\delta > 1/2$. However, as the parameter τ in Theorem 4 increases, the local bound also degrades quickly.

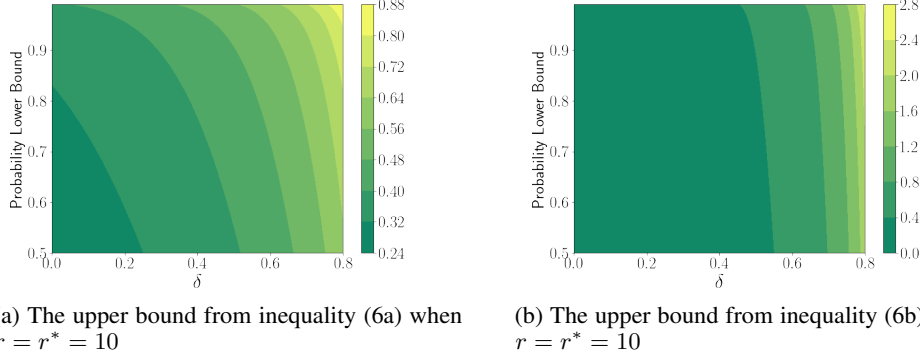


Figure 3: Comparison of the upper bounds given by Theorem 4 under $\tau = 0.2$ for the distance $\|\hat{X}\hat{X}^T - M^*\|_F$ with $\hat{X}$ being an arbitrary local minimizer

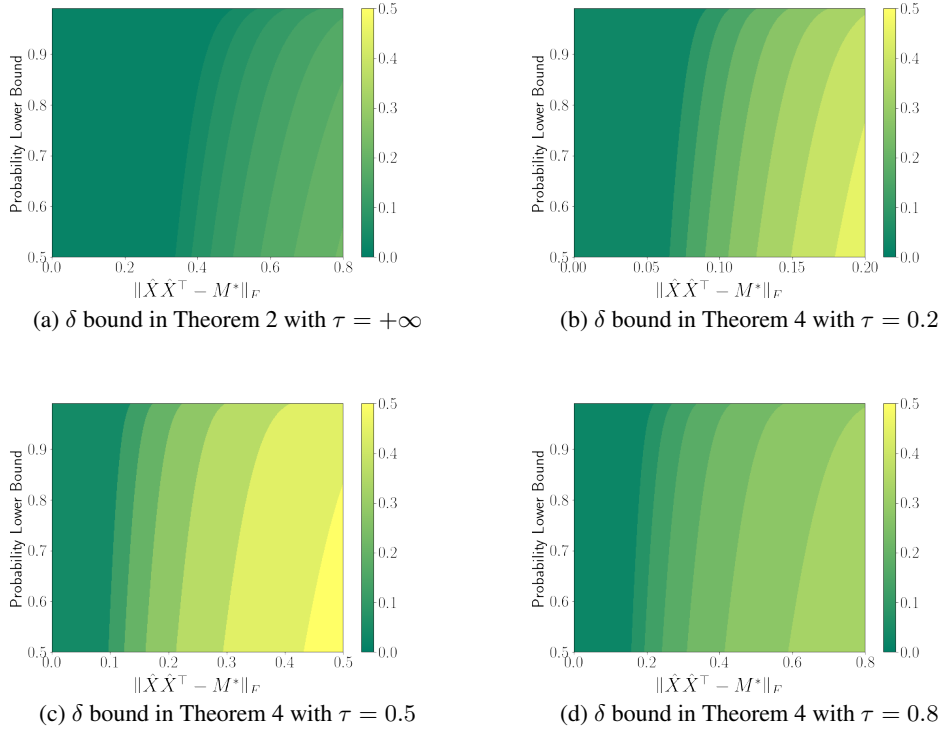


Figure 4: Comparison of the maximum RIP constants δ allowed by Theorem 2 and Theorem 4 to guarantee a given maximum distance $\|\hat{X}\hat{X}^T - M^*\|_F$ for an arbitrary local minimizer $\hat{X}$ satisfying (5) with a given probability

5 Conclusion

In this paper, we develop global and local analyses for the locations of the local minima of the low-rank matrix recovery problem with noisy linear measurements in both the exact parameterized and the overparameterized regimes. For the class of noisy problems, regardless of their RIP constants, it is now possible to characterize the worst-case quality of the local minimizers. The major innovation of this work lies in the new proof techniques developed to deal with the overparameterization and the handling of the random noise via an easy-to-compute concentration bound. Unlike the existing

results, the guarantees in our results are distribution-agnostic, meaning that the distribution can be unknown as long as the concentration bound is possible to obtain. The developed results encompass the state-of-the-art results on the non-existence of spurious solutions in the noiseless case. Last but not least, we prove a certain form of the strict saddle property, which guarantees the global convergence of the perturbed gradient descent method in polynomial time regardless of parameterization. Our analyses show how the value of the RIP constant and the intensity of noise affect the landscape of the non-convex learning problem and the locations of the local minima relative to the ground truth.

References

- [1] E. J. Candès and B. Recht, “Exact matrix completion via convex optimization,” *Foundations of Computational Mathematics*, vol. 9, no. 6, pp. 717–772, 2009.
- [2] B. Recht, M. Fazel, and P. A. Parrilo, “Guaranteed minimum-rank solutions of linear matrix equations via nuclear norm minimization,” *SIAM Review*, vol. 52, no. 3, pp. 471–501, 2010.
- [3] A. Singer, “Angular synchronization by eigenvectors and semidefinite programming,” *Applied and Computational Harmonic Analysis*, vol. 30, no. 1, pp. 20–36, 2011.
- [4] N. Boumal, “Nonconvex phase synchronization,” *SIAM Journal on Optimization*, vol. 26, no. 4, pp. 2355–2377, 2016.
- [5] Y. Shechtman, Y. C. Eldar, O. Cohen, H. N. Chapman, J. Miao, and M. Segev, “Phase retrieval with application to optical imaging: A contemporary overview,” *IEEE Signal Processing Magazine*, vol. 32, no. 3, pp. 87–109, 2015.
- [6] R. Ge, C. Jin, and Y. Zheng, “No spurious local minima in nonconvex low rank problems: A unified geometric analysis,” in *Proceedings of the 34th International Conference on Machine Learning*, vol. 70 of *Proceedings of Machine Learning Research*, pp. 1233–1242, 2017.
- [7] 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 Processing Magazine*, vol. 35, no. 4, pp. 14–31, 2018.
- [8] Y. Chi, Y. M. Lu, and Y. Chen, “Nonconvex optimization meets low-rank matrix factorization: An overview,” *IEEE Transactions on Signal Processing*, vol. 67, no. 20, pp. 5239–5269, 2019.
- [9] E. J. Candès and T. Tao, “The power of convex relaxation: Near-optimal matrix completion,” *IEEE Transactions on Information Theory*, vol. 56, no. 5, pp. 2053–2080, 2010.
- [10] S. Burer and R. D. Monteiro, “A nonlinear programming algorithm for solving semidefinite programs via low-rank factorization,” *Mathematical Programming*, vol. 95, no. 2, pp. 329–357, 2003.
- [11] S. Bhojanapalli, B. Neyshabur, and N. Srebro, “Global optimality of local search for low rank matrix recovery,” in *Advances in Neural Information Processing Systems*, vol. 29, 2016.
- [12] D. Park, A. Kyrillidis, C. Carmanis, and S. Sanghavi, “Non-square matrix sensing without spurious local minima via the Burer–Monteiro approach,” in *Proceedings of the 20th International Conference on Artificial Intelligence and Statistics*, vol. 54 of *Proceedings of Machine Learning Research*, pp. 65–74, 2017.
- [13] X. Zhang, L. Wang, Y. Yu, and Q. Gu, “A primal-dual analysis of global optimality in nonconvex low-rank matrix recovery,” in *Proceedings of the 35th International Conference on Machine Learning*, vol. 80 of *Proceedings of Machine Learning Research*, pp. 5862–5871, 2018.
- [14] Z. Zhu, Q. Li, G. Tang, and M. B. Wakin, “Global optimality in low-rank matrix optimization,” *IEEE Transactions on Signal Processing*, vol. 66, no. 13, pp. 3614–3628, 2018.
- [15] R. Y. Zhang, S. Sojoudi, and J. Lavaei, “Sharp restricted isometry bounds for the inexistence of spurious local minima in nonconvex matrix recovery,” *Journal of Machine Learning Research*, vol. 20, no. 114, pp. 1–34, 2019.

- [16] G. Zhang and R. Y. Zhang, “How many samples is a good initial point worth in low-rank matrix recovery?,” in *Advances in Neural Information Processing Systems*, vol. 33, pp. 12583–12592, 2020.
- [17] Y. Bi and J. Lavaei, “Global and local analyses of nonlinear low-rank matrix recovery problems,” 2020. arXiv:2010.04349.
- [18] W. Ha, H. Liu, and R. F. Barber, “An equivalence between critical points for rank constraints versus low-rank factorizations,” *SIAM Journal on Optimization*, vol. 30, no. 4, pp. 2927–2955, 2020.
- [19] Z. Zhu, Q. Li, G. Tang, and M. B. Wakin, “The global optimization geometry of low-rank matrix optimization,” *IEEE Transactions on Information Theory*, vol. 67, no. 2, pp. 1308–1331, 2021.
- [20] R. Y. Zhang, “Sharp global guarantees for nonconvex low-rank matrix recovery in the overparameterized regime,” 2021. arXiv:2104.10790.
- [21] H. Zhang, Y. Bi, and J. Lavaei, “General low-rank matrix optimization: Geometric analysis and sharper bounds,” in *Advances in Neural Information Processing Systems*, 2021.
- [22] L. Wang, X. Zhang, and Q. Gu, “A unified computational and statistical framework for non-convex low-rank matrix estimation,” in *Proceedings of the 20th International Conference on Artificial Intelligence and Statistics*, vol. 54 of *Proceedings of Machine Learning Research*, pp. 981–990, 2017.
- [23] D. Park, A. Kyrillidis, C. Caramanis, and S. Sanghavi, “Finding low-rank solutions via nonconvex matrix factorization, efficiently and provably,” *SIAM Journal on Imaging Sciences*, vol. 11, no. 4, pp. 2165–2204, 2018.
- [24] R. Y. Zhang, C. Josz, S. Sojoudi, and J. Lavaei, “How much restricted isometry is needed in nonconvex matrix recovery?,” in *Advances in Neural Information Processing Systems*, vol. 31, 2018.
- [25] X. Li, Z. Zhu, A. Man-Cho So, and R. Vidal, “Nonconvex robust low-rank matrix recovery,” *SIAM Journal on Optimization*, vol. 30, no. 1, pp. 660–686, 2020.
- [26] C. Jin, R. Ge, P. Netrapalli, S. M. Kakade, and M. I. Jordan, “How to escape saddle points efficiently,” in *Proceedings of the 34th International Conference on Machine Learning*, vol. 70 of *Proceedings of Machine Learning Research*, pp. 1724–1732, Aug. 2017.
- [27] Y. Bi, H. Zhang, and J. Lavaei, “Local and global linear convergence of general low-rank matrix recovery problems,” *arXiv preprint arXiv:2104.13348*, 2021.
- [28] E. J. Candès and Y. Plan, “Tight oracle inequalities for low-rank matrix recovery from a minimal number of noisy random measurements,” *IEEE Transactions on Information Theory*, vol. 57, pp. 2342–2359, Apr. 2011.
- [29] M. Jin, J. Lavaei, S. Sojoudi, and R. Baldick, “Boundary defense against cyber threat for power system state estimation,” *IEEE Transactions on Information Forensics and Security*, vol. 16, pp. 1752–1767, 2021.
- [30] C. Jin, P. Netrapalli, R. Ge, S. M. Kakade, and M. I. Jordan, “A short note on concentration inequalities for random vectors with subGaussian norm,” 2019. arXiv:1902.03736.