

On the Benefits of Weight Normalization for Overparameterized Matrix Sensing

Yudong Wei

Liang Zhang

Bingcong Li*

Niao He*

ETH Zurich, Switzerland

YUDWEI@ETHZ.CH

LIANG.ZHANG@INF.ETHZ.CH

BINGCONG.LI@INF.ETHZ.CH

NIAO.HE@INF.ETHZ.CH

Abstract

While normalization techniques are widely used in deep learning, their theoretical understanding remains relatively limited. In this work, we establish the benefits of (generalized) weight normalization (WN) applied to the overparameterized matrix sensing problem. We prove that WN with Riemannian optimization achieves linear convergence, yielding an *exponential* speedup over standard methods that do not use WN. Our analysis further demonstrates that both iteration and sample complexity improve polynomially as the level of overparameterization increases. To the best of our knowledge, this work provides the first characterization of how WN leverages overparameterization for faster convergence in matrix sensing.

1. Introduction

Normalization schemes, such as layer, batch, and weight normalization, are essential in modern deep networks and have proven highly effective for stabilizing training in both vision and language models [4, 19, 34]. Despite their practical success, theoretical explanations of why they work remain elusive, even for relatively simple problems.

This work focuses on weight normalization (WN), which decouples parameters (i.e., variables) into directions and magnitudes, and then optimizes them separately. It has recently regained considerable attention because of the seamless integration with LoRA [18], leading to several powerful strategies for parameter-efficient fine-tuning of large language models; see e.g., [27, 28]. Yet, theoretical support for WN remains relatively limited. Prior results in [44] show that WN applied to overparameterized least squares induces implicit regularization towards the minimum ℓ_2 -norm solution. The implicit regularization of WN on diagonal linear neural networks is studied in [13]. WN is also observed to reduce Hessian spectral norm and improve generalization in deep networks [14].

Our work broadens the understanding of WN by establishing its merits in overparameterized matrix sensing. The goal here is to recover a low-rank positive semi-definite (PSD) matrix $\mathbf{A} \in \mathbb{R}^{m \times m}$ from linear measurements. In the vanilla formulation without WN, one can exploit the low-rankness of ground-truth matrix, i.e., $r_A := \text{rank}(\mathbf{A}) \ll m$ for efficient parameterization. Specifically, we can optimize on $\mathbf{Y} \in \mathbb{R}^{m \times r}$ such that $\mathbf{Y}\mathbf{Y}^\top \approx \mathbf{A}$ [7]. The overparameterized regime $r > r_A$ is of interest due to the need of exact recovery without knowing r_A a priori. This problem has

* Equal supervision.

Table 1: Comparison with existing algorithms for overparameterized matrix sensing. “E.C.” denotes “exact convergence”, i.e., whether the reconstruction error bound will go to zero when $t \rightarrow \infty$. UB, LB and OP are short for upper bound, lower bound, and overparameterization, respectively.

Algorithm	WN	E.C.	Initialization	Convergence Rate	Faster with OP
GD (UB) [38]	✗	✗	Small & random	N/A	-
GD (LB) [45]	✗	✓	Small & random	$\Omega(\frac{\kappa^2}{\log(mr_A^2)t})$	✗
RGD (Theorem 2)	✓	✓	Random	$\exp\left(-\mathcal{O}\left(\frac{(r-r_A)^4}{\kappa^4 m^2 r^2 r_A} t\right)\right)$	✓

wide applications in machine learning and signal processing [9], and serves as a popular testbed for theoretical deep learning given its non-convexity and rich loss landscape; see e.g., [3, 20, 25].

Without WN, prior work [45] establishes a sublinear lower bound on the convergence rate when the above sensing problem is optimized via gradient descent (GD), even with infinite data samples. We circumvent this lower bound by i) extending WN for coping with matrix variables; and, ii) proving that applying this generalized WN with Riemannian gradient descent (RGD) enables a *linear* convergence rate in the finite sample regime, leading to an *exponential* improvement. Remarkably, *WN leverages higher level of overparameterization to achieve both faster convergence and lower sample complexity*. To the best of our knowledge, this is the first theoretical result demonstrating that normalization benefits from overparameterization.

More concretely, our contributions are summarized as follows:

❖ **Exponentially faster rate.** For overparameterized matrix sensing problems, we prove that randomly initialized WN achieves a linear convergence rate of $\exp(-\mathcal{O}(\frac{(r-r_A)^4}{\kappa^4 m^2 r^2 r_A} t))$, where κ is the condition number of the ground-truth matrix $\mathbf{A}$. This linear rate is exponentially faster than the sublinear lower bound $\Omega(\frac{\kappa^2}{\log(mr_A^2)t})$ obtained without WN. Moreover, additional overparameterization in WN provides quantifiable benefits: the iteration complexity scales down polynomially as the overparameterization level r increases; see Table 1 for a summary.

❖ **Empirical validation.** We conduct experiments on overparameterized matrix sensing and the numerical results corroborate our theoretical findings.

Notation. Bold lowercase (capital) letters denote column vectors (matrices); $(\cdot)^\top$ and $\|\cdot\|_F$ refer to transpose and Frobenius norm of a matrix; $\|\cdot\|$ denotes the ℓ_2 norm for vectors and the spectral norm for matrices; $\langle \mathbf{A}, \mathbf{B} \rangle = \text{Tr}(\mathbf{A}^\top \mathbf{B})$ represents the standard matrix inner product; $\sigma_i(\cdot)$ and $\lambda_i(\cdot)$ denote the i -th largest singular value and eigenvalue, respectively. $\mathbb{S}^m$ and $\mathbb{S}_+^m$ denote symmetric and positive semi-definite (PSD) matrices of size $m \times m$, respectively.

2. WN for overparameterized matrix sensing

We focus on applying WN to the symmetric low-rank matrix sensing problem. The objective is to recover a low-rank and positive semi-definite (PSD) matrix $\mathbf{A} \in \mathbb{S}_+^m$ from a collection of n data $\{(\mathbf{M}_i, y_i)\}_{i=1}^n$, where each feature matrix $\mathbf{M}_i \in \mathbb{S}^m$ is symmetric and the corresponding label is $y_i = \text{Tr}(\mathbf{M}_i^\top \mathbf{A})$. For notational conciseness, we let $\mathbf{y} = [y_1, \dots, y_n]^\top \in \mathbb{R}^n$ and define a linear

mapping $\mathcal{M} : \mathbb{S}^m \mapsto \mathbb{R}^n$ with $[\mathcal{M}(\mathbf{A})]_i = \text{Tr}(\mathbf{M}_i^\top \mathbf{A})$. As mentioned in the introduction, for the vanilla formulation, we optimize on $\mathbf{Y} \in \mathbb{R}^{m \times r}$ such that $\mathbf{Y}\mathbf{Y}^\top \approx \mathbf{A}$ [7]. This leads to

$$\min_{\mathbf{Y} \in \mathbb{R}^{m \times r}} \frac{1}{4} \|\mathcal{M}(\mathbf{Y}\mathbf{Y}^\top) - \mathbf{y}\|^2. \quad (1)$$

Despite its seemingly simple formulation, the loss landscape contains saddle points, hence achieving a *global* optimum from a random initialization is nontrivial. Moreover, overparameterization, i.e., $r > r_A$, is often considered in practice to ensure exact recovery of $\mathbf{A}$ without prior knowledge of its rank. It is established in [45] that such overparameterization induces optimization challenges even in the population setting ($n \rightarrow \infty$). In particular, a lower bound of GD shows that $\|\mathbf{Y}_t \mathbf{Y}_t^\top - \mathbf{A}\|_F$ converges no faster than $\Omega(1/t)$, where t is the iteration number. This rate is exponentially slower than the linear one when r_A is known to employ $r = r_A$ [48].

Applying WN to problem (1). For a vector variable, WN decouples it into direction and magnitude, and optimizes them separately. Extending this idea to matrix variables in (1), we leverage polar decomposition to write $\mathbf{Y} = \mathbf{X}\tilde{\Theta}$, where $\mathbf{X} \in \text{St}(m, r)$ lies in a Stiefel manifold and $\tilde{\Theta} \in \mathbb{S}_+^r$. Here, the Stiefel manifold $\text{St}(m, r)$ is defined as $\{\mathbf{X} \in \mathbb{R}^{m \times r} | \mathbf{X}^\top \mathbf{X} = \mathbf{I}_r\}$. One can geometrically interpret $\mathbf{X}$ as orthonormal bases for an r -dimensional subspace, thus representing “directions”, and $\tilde{\Theta}$ captures the “magnitude” of a matrix. Substituting $\mathbf{Y}$ in (1), we arrive at

$$\min_{\mathbf{X}, \tilde{\Theta}} \frac{1}{4} \|\mathcal{M}(\mathbf{X}\tilde{\Theta}\tilde{\Theta}^\top \mathbf{X}^\top) - \mathbf{y}\|^2 \quad \text{s.t.} \quad \mathbf{X} \in \text{St}(m, r), \tilde{\Theta} \in \mathbb{S}_+^r.$$

The above problem can be further simplified by i) merging $\tilde{\Theta}\tilde{\Theta}^\top$ into a single matrix $\Theta \in \mathbb{S}_+^r$; and ii) relaxing the PSD constraint on Θ to only symmetry, i.e., $\Theta \in \mathbb{S}^r$. This relaxation achieves the same global objective in the overparameterized regime, yet significantly improves computational efficiency by avoiding SVDs or matrix exponentials needed for optimizing over PSD cones [39, 41]. In sum, applying WN gives the objective

$$\min_{\mathbf{X}, \Theta} f(\mathbf{X}, \Theta) := \frac{1}{4} \|\mathcal{M}(\mathbf{X}\Theta \mathbf{X}^\top) - \mathbf{y}\|^2 \quad \text{s.t.} \quad \mathbf{X} \in \text{St}(m, r), \Theta \in \mathbb{S}^r. \quad (2)$$

For convenience, we continue to refer to this generalized variant as WN, since it aligns with the direction-magnitude decomposition paradigm. Similar reformulations of (1) have appeared in [23, 31]. The former empirically studies the faster convergence on matrix completion problems, while the latter tackles local geometry around stationary points. Our work, on the other hand, characterizes the optimization benefits of WN and clarifies its interaction with overparameterization.

2.1. Solving WN via Riemannian optimization

Generalizing the vector WN¹ on matrix problems, Riemannian optimization is adopted for coping with the manifold constraint $\mathbf{X} \in \text{St}(m, r)$. We simply treat the manifold as an embedded one in Euclidean space. Extensions to other geometry are straightforward. To optimize the direction variable $\mathbf{X}_t$, let $\tilde{\mathbf{G}}_t := \nabla_{\mathbf{X}} f(\mathbf{X}_t, \Theta_t)$ denote the Euclidean gradient on $\mathbf{X}_t$ (a detailed expression is given in (5) of Appendix A). The Riemannian gradient for $\mathbf{X}_t$ can be written as

$$\mathbf{G}_t := (\mathbf{I}_m - \mathbf{X}_t \mathbf{X}_t^\top) \tilde{\mathbf{G}}_t + \frac{\mathbf{X}_t}{2} (\mathbf{X}_t^\top \tilde{\mathbf{G}}_t - \tilde{\mathbf{G}}_t^\top \mathbf{X}_t). \quad (3)$$

1. While the practical update rule of WN [34, eq. (4)] lies between Riemannian and Euclidean optimization, [44, Lemma 2.2] shows that the limiting flow is Riemannian flow.

Algorithm 1 Riemannian gradient descent (RGD) for solving WN (2)

Input: Initial point $\mathbf{X}_0 \in \text{St}(m, r)$, $\Theta_0 \in \mathbb{S}^r$, step sizes η, μ

```

for  $t = 0, 1, \dots, T$  do
    Compute  $\mathbf{G}_t$  (Riemannian gradient of  $\mathbf{X}_t$ ) via (3)
    Update  $\mathbf{X}_{t+1}$  via (4)                                // direction variable
    Compute  $\mathbf{K}_t$  (gradient of  $\Theta_t$ ) via (6)
    Update  $\Theta_{t+1}$  via (7)                                // magnitude variable
end
Output:  $\mathbf{X}_{T+1}, \Theta_{T+1}$ 
    
```

Further applying the polar retraction² to ensure feasibility, the update for $\mathbf{X}_t$ is given by

$$\mathbf{X}_{t+1} = (\mathbf{X}_t - \eta \mathbf{G}_t)(\mathbf{I}_r + \eta^2 \mathbf{G}_t^\top \mathbf{G}_t)^{-1/2} \quad (4)$$

where $\eta > 0$ is the stepsize. Detailed derivations of (3) and (4) are deferred to Appendix A. Note that polar retraction is used here for theoretical simplicity. Shown in Appendix C, other retractions for Stiefel manifolds such as QR and Cayley³ share almost identical performance numerically.

An alternative update method is adopted for the magnitude variable Θ_t . Denote the gradient as $\mathbf{K}_t := \nabla_{\Theta} f(\mathbf{X}_{t+1}, \Theta_t)$, whose expression can be found in (6) in Appendix A. We use GD with a step size $\mu > 0$ to optimize Θ_t

$$\Theta_{t+1} = \Theta_t - \mu \mathbf{K}_t.$$

This update ensures feasibility of the symmetric constraint $\Theta_t \in \mathbb{S}^r, \forall t$ whenever initialized with $\Theta_0 \in \mathbb{S}^r$; see a proof in Lemma 22. The step-by-step procedure for solving (2) is summarized in Algorithm 1, and it is termed as RGD for future reference.

3. On the benefits of WN

This section demonstrates that WN delivers exact convergence at a linear rate for overparameterized matrix sensing (2) and leverages additional overparameterization to yield faster optimization and lower sample complexity. Recall that the rank of $\mathbf{A}$ is denoted by r_A . Let the compact SVD of $\mathbf{A}$ be $\mathbf{A} = \mathbf{U}\Sigma\mathbf{U}^\top$, where $\mathbf{U} \in \mathbb{R}^{m \times r_A}$ and $\Sigma \in \mathbb{S}_+^{r_A}$. Without loss of generality, we assume $\sigma_1(\Sigma) = 1$ and $\sigma_{r_A}(\Sigma) = 1/\kappa$ with $\kappa \geq 1$ denoting the condition number. We will use the restricted isometry property (RIP) [32], a standard assumption in matrix sensing, in our proofs; see more in, e.g., [38, 45, 47, 49].

Definition 1 (Restricted Isometry Property (RIP)) *The linear mapping $\mathcal{M}(\cdot)$ is (r, δ) -RIP, with $\delta \in [0, 1)$, if for all matrices $\mathbf{A} \in \mathbb{S}^m$ of rank at most r , it satisfies*

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

RIP ensures that the linear measurement approximately preserves the Frobenius norm of low-rank matrices. A detailed discussion and illustrative examples of RIP are provided in Appendix D.2. With these preparations, we are ready to uncover the merits of WN.

2. Let $\mathbf{X} \in \text{St}(m, r)$ and a point in its tangent space $\mathbf{G} \in \mathcal{T}_{\mathbf{X}}\text{St}(m, r)$. The polar retraction for $\mathbf{X} + \mathbf{G}$ is given by $\mathcal{R}_{\mathbf{X}}(\mathbf{G}) = (\mathbf{X} + \mathbf{G})(\mathbf{I}_r + \mathbf{G}^\top \mathbf{G})^{-1/2}$.
 3. See e.g., [1], for more detailed discussions on retractions.

3.1. Main results

We consider WN under random initialization, meaning that $\mathbf{X}_0$ is chosen uniformly at random from the manifold $\text{St}(m, r)$. One possible approach is to set $\mathbf{X}_0 = \mathbf{Z}_0(\mathbf{Z}_0^\top \mathbf{Z}_0)^{-1/2}$, where the entries of $\mathbf{Z}_0 \in \mathbb{R}^{m \times r}$ are i.i.d. Gaussian random variables $\mathcal{N}(0, 1)$ [12].

Theorem 2 *Consider solving the WN-aided sensing problem (2) initialized with random $\mathbf{X}_0 \in \text{St}(m, r)$ and $\Theta_0 \in \mathbb{S}^r$ satisfying $\|\Theta_0\| \leq 2$. Assume that $r_A \leq \frac{m}{2}$ and $\mathcal{M}(\cdot)$ is $(r + r_A + 1, \delta)$ -RIP with $\delta = \mathcal{O}(\frac{(r-r_A)^6}{\kappa^2 m^3 r^4 r_A})$. Algorithm 1 using stepsizes $\eta = \mathcal{O}(\frac{(r-r_A)^4}{\kappa^2 m^2 r^2 r_A})$ and $\mu = 2$ generates a sequence $\{\mathbf{X}_t, \Theta_t\}_{t=0}^\infty$. With high probability over the initialization, this sequence proceeds in two distinct phases, separated by a burn-in time t_0 with an upper bound $\mathcal{O}(\frac{\kappa^4 m^4 r^4 r_A^2}{(r-r_A)^8})$:*

i) *Initial phase. For some universal constant $c_2 \in (0, 1)$, it follows that*

$$\|\mathbf{X}_t \Theta_t \mathbf{X}_t^\top - \mathbf{A}\|_F \leq 2 \sqrt{r_A - \frac{c_2 (r - r_A)^8 t}{\kappa^4 m^4 r^4 r_A}} + 1, \quad 1 \leq t \leq t_0.$$

ii) *Linearly convergent phase. For some universal constant $c_3 \in (0, 1)$, it is guaranteed that*

$$\|\mathbf{X}_t \Theta_t \mathbf{X}_t^\top - \mathbf{A}\|_F \leq 3 \left(1 - \frac{c_3 (r - r_A)^4}{\kappa^4 m^2 r^2 r_A} \right)^{t-t_0}, \quad \forall t \geq t_0 + 1.$$

We refer to $\|\mathbf{X}_t \Theta_t \mathbf{X}_t^\top - \mathbf{A}\|_F$ as the reconstruction error since it measures the distance to our target matrix. Next, we break down Theorem 2 to demonstrate the benefits of generalized WN for the overparameterized matrix sensing problem from two different perspectives.

Optimization benefits of WN include i) faster convergence rate, and ii) less stringent initialization requirements. Theorem 2 shows that WN achieves exact convergence with a linear rate. In contrast, without WN, the convergence behavior of randomly initialized GD on (1) is weaker. Specifically, [38] shows that GD can only attain a constant reconstruction error with early stopping, but not guarantee last-iteration convergence. On the other hand, [45] establishes a lower bound for exact recovery of GD, giving a sublinear dependence on t ; see a detailed comparison in Table 1. In addition, our guarantee of this linear rate is obtained without strict requirements on initialization, which stands in stark contrast to the non-WN setting, where the magnitude of random initialization must be carefully controlled, often inversely proportional to κ [20, 38, 47].

WN makes overparameterization a friend. Because the additional parameters induce computation and memory overheads, it is natural to expect more gains from overparameterization. It can be seen from Table 1 that GD does not benefit from overparameterization, while the benefits of overparameterization for WN are twofold. Setting $r = pr_A$ for some $p > 1$, one can rewrite the upper bound of the burn-in time t_0 as $\mathcal{O}(\frac{\kappa^4 m^4 p^4}{(p-1)^8 r_A^2})$, which decreases polynomially with p . In the linearly convergent phase, WN achieves a convergence rate of $\exp(-\mathcal{O}(\frac{(p-1)^4 r_A}{\kappa^4 m^2 p^2 t}))$, which is also faster with a larger p . In terms of iteration complexity, this translates into a polynomial improvement with the level of overparameterization. To quantitatively understand the merits of overparameterization, we consider two cases. In the mildly overparameterized regime, where $r = r_A + c$ for some constant $c = \mathcal{O}(1)$, the convergence rate reads $\exp(-\mathcal{O}(\frac{t}{\kappa^4 m^2 r_A^3}))$. When the level of overparameterization increases to $r = cr_A$, the rate improves to $\exp(-\mathcal{O}(\frac{r_A t}{\kappa^4 m^2}))$. Through comparison, we see that

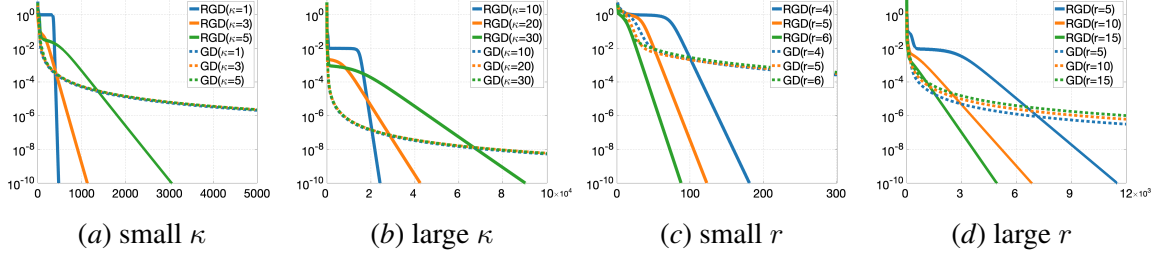


Figure 1: Comparison of RGD on WN and GD on (1) (squared reconstruction error vs. iteration).

additional overparameterization yields up to a factor of $\mathcal{O}(r_A^4)$ improvement in the exponent. On the statistical side, the sample complexity of WN is determined by the RIP assumption on $\mathcal{M}(\cdot)$. Under the Gaussian design, as detailed in Appendix D.2, the RIP holds w.h.p. when $n = \mathcal{O}\left(\frac{\kappa^4 m^7 r_A^2}{(r - r_A)^{12}}\right)$. Notably, the sample complexity n reduces polynomially as r increases. In particular, following the same analysis as for the convergence rate, this reduction can reach up to a factor of $\mathcal{O}(r_A^{12})$.

4. Numerical experiments

Numerical experiments are conducted to validate our theoretical findings for WN on overparameterized matrix sensing problems. More details on experiments setup are deferred to Appendix I.1.

Faster convergence of WN. We compare RGD on WN with vanilla GD on (1) under different condition numbers κ . For small κ instances with $m = 10$, $r = 5$, $r_A = 3$, and $n = 1000$, we consider $\kappa \in \{1, 3, 5\}$. For large κ instances where m, r, r_A remain fixed, and $n = 3000$, we test $\kappa \in \{10, 20, 30\}$. Figures 1(a), (b) demonstrate that WN converges to 0 in a linear convergence rate after a small initial phase, while vanilla GD slows down to a sub-linear rate.

On the benefit of overparameterization. We further examine the impact of the level of overparameterization r . For small r instances with $m = 10$, $r_A = 3$, $\kappa = 1$, and $n = 1000$, we test $r \in \{4, 5, 6\}$. For large r instances with $m = 20$, $r_A = 3$, $\kappa = 10$, and $n = 3000$, we consider $r \in \{5, 10, 15\}$. Figures 1(c), (d) show that larger r leads to a quicker escape from the initial phase and a higher convergence rate for WN, consistent with our theoretical findings.

5. Conclusion

This work provides new theoretical insights into the role of weight normalization (WN) in overparameterized matrix sensing. We prove that randomly initialized WN with proper Riemannian optimization guarantees a linear rate, yielding an exponential improvement on overparameterized sensing problems without WN. Moreover, we show that overparameterization can be exploited under WN to achieve faster optimization and lower sample complexity. Numerical experiments further validate our findings. Future work includes extending these results to broader non-convex learning settings, such as tensor problems [40], and developing new algorithms that build on WN.

References

- [1] P-A Absil, Robert Mahony, and Rodolphe Sepulchre. *Optimization Algorithms on Matrix Manifolds*. Princeton University Press, 2008.
- [2] Animashree Anandkumar, Rong Ge, Daniel J Hsu, Sham M Kakade, and Matus Telgarsky. Tensor decompositions for learning latent variable models. *Journal of Machine Learning Research*, 15(1):2773–2832, 2014.
- [3] Sanjeev Arora, Nadav Cohen, Wei Hu, and Yuping Luo. Implicit regularization in deep matrix factorization. *Advances in Neural Information Processing Systems*, 32, 2019.
- [4] Jimmy Lei Ba, Jamie Ryan Kiros, and Geoffrey E Hinton. Layer normalization. *arXiv preprint arXiv:1607.06450*, 2016.
- [5] Åke Björck and Gene H. Golub. Numerical methods for computing angles between linear subspaces. *Mathematics of Computation*, 27(123):579–594, 1973.
- [6] Nicolas Boumal. *An Introduction to Optimization on Smooth Manifolds*. Cambridge University Press, 2023.
- [7] Samuel Burer and Renato D. C. Monteiro. Local minima and convergence in low-rank semidefinite programming. *Mathematical Programming*, 103(3):427–444, 2005.
- [8] Emmanuel J Candès and Xiaodong Li. Solving quadratic equations via PhaseLift when there are about as many equations as unknowns. *Foundations of Computational Mathematics*, 14(5):1017–1026, 2014.
- [9] Emmanuel J Candès, Thomas Strohmer, and Vladislav Voroninski. PhaseLift: Exact and stable signal recovery from magnitude measurements via convex programming. *Communications on Pure and Applied Mathematics*, 66(8):1241–1274, 2013.
- [10] Emmanuel J. Candès and Yaniv Plan. Tight oracle inequalities for low-rank matrix recovery from a minimal number of noisy random measurements. *IEEE Transactions on Information Theory*, 57(4):2342–2359, 2011.
- [11] Cheng Cheng and Ziping Zhao. Accelerating gradient descent for over-parameterized asymmetric low-rank matrix sensing via preconditioning. In *IEEE International Conference on Acoustics, Speech and Signal Processing*, pages 7705–7709, 2024.
- [12] Yasuko Chikuse. *Statistics on Special Manifolds*, volume 174. Springer Science & Business Media, 2012.
- [13] Hung-Hsu Chou, Holger Rauhut, and Rachel Ward. Robust implicit regularization via weight normalization. *Information and Inference: A Journal of the IMA*, 13(3):iaae022, 2024.
- [14] Pedro Cisneros-Velarde, Zhijie Chen, Sanmi Koyejo, and Arindam Banerjee. Optimization and generalization guarantees for weight normalization. *Transactions on Machine Learning Research*, 2025. ISSN 2835-8856.

- [15] John C Duchi, Oliver Hinder, Andrew Naber, and Yinyu Ye. Conic descent and its application to memory-efficient optimization over positive semidefinite matrices. *Advances in Neural Information Processing Systems*, 33:8308–8317, 2020.
- [16] Rong Ge, Yunwei Ren, Xiang Wang, and Mo Zhou. Understanding deflation process in over-parametrized tensor decomposition. *Advances in Neural Information Processing Systems*, 34: 1299–1311, 2021.
- [17] Gene H. Golub and Charles F. Van Loan. *Matrix Computations - 4th Edition*. Johns Hopkins University Press, 2013.
- [18] Edward J Hu, yelong shen, Phillip Wallis, Zeyuan Allen-Zhu, Yuanzhi Li, Shean Wang, Lu Wang, and Weizhu Chen. LoRA: Low-rank adaptation of large language models. In *International Conference on Learning Representations*, 2022.
- [19] Sergey Ioffe and Christian Szegedy. Batch normalization: Accelerating deep network training by reducing internal covariate shift. In *International Conference on Machine Learning*, pages 448–456. PMLR, 2015.
- [20] Jikai Jin, Zhiyuan Li, Kaifeng Lyu, Simon Shaolei Du, and Jason D Lee. Understanding incremental learning of gradient descent: A fine-grained analysis of matrix sensing. In *International Conference on Machine Learning*, pages 15200–15238. PMLR, 2023.
- [21] Yuhi Kawakami and Mahito Sugiyama. Investigating overparameterization for non-negative matrix factorization in collaborative filtering. In *ACM Conference on Recommender Systems*, pages 685–690, 2021.
- [22] Alex Krizhevsky and Geoffrey Hinton. Learning multiple layers of features from tiny images. Technical report, University of Toronto, 2009.
- [23] Eitan Levin, Joe Kileel, and Nicolas Boumal. The effect of smooth parametrizations on non-convex optimization landscapes. *Mathematical Programming*, 209(1):63–111, 2025.
- [24] Bingcong Li, Liang Zhang, Aryan Mokhtari, and Niao He. On the crucial role of initialization for matrix factorization. In *International Conference on Learning Representations*, 2025.
- [25] Yuanzhi Li, Tengyu Ma, and Hongyang Zhang. Algorithmic regularization in over-parameterized matrix sensing and neural networks with quadratic activations. In *Conference on Learning Theory*, pages 2–47. PMLR, 2018.
- [26] Zhiyuan Li, Yuping Luo, and Kaifeng Lyu. Towards resolving the implicit bias of gradient descent for matrix factorization: Greedy low-rank learning. In *International Conference on Learning Representations*, 2021.
- [27] Kai Lion, Liang Zhang, Bingcong Li, and Niao He. PoLAR: Polar-decomposed low-rank adapter representation. *arXiv preprint arXiv:2506.03133*, 2025.
- [28] Shih-Yang Liu, Chien-Yi Wang, Hongxu Yin, Pavlo Molchanov, Yu-Chiang Frank Wang, Kwang-Ting Cheng, and Min-Hung Chen. DoRA: Weight-decomposed low-rank adaptation. In *International Conference on Machine Learning*, 2024.

- [29] Cong Ma, Xingyu Xu, Tian Tong, and Yuejie Chi. Provably accelerating ill-conditioned low-rank estimation via scaled gradient descent, even with overparameterization. *arXiv preprint arXiv:2310.06159*, 2023.
- [30] Jianhao Ma and Salar Fattahi. Convergence of gradient descent with small initialization for unregularized matrix completion. In *Conference on Learning Theory*, pages 3683–3742. PMLR, 2024.
- [31] Bamdev Mishra, Gilles Meyer, Silvere Bonnabel, and Rodolphe Sepulchre. Fixed-rank matrix factorizations and Riemannian low-rank optimization. *Computational Statistics*, 29(3):591–621, 2014.
- [32] Benjamin Recht, Maryam Fazel, and Pablo A. Parrilo. Guaranteed minimum-rank solutions of linear matrix equations via nuclear norm minimization. *SIAM Review*, 52(3):471–501, 2010.
- [33] Mark Rudelson and Roman Vershynin. Smallest singular value of a random rectangular matrix. *Communications on Pure and Applied Mathematics: A Journal Issued by the Courant Institute of Mathematical Sciences*, 62(12):1707–1739, 2009.
- [34] Tim Salimans and Durk P Kingma. Weight normalization: A simple reparameterization to accelerate training of deep neural networks. *Advances in Neural Information Processing Systems*, 29, 2016.
- [35] J Ben Schafer, Dan Frankowski, Jon Herlocker, and Shilad Sen. Collaborative filtering recommender systems. In *The Adaptive Web: Methods and Strategies of Web Personalization*, pages 291–324. Springer, 2007.
- [36] Mohamed El Amine Seddik, Mohammed Mahfoud, and Merouane Debbah. Optimizing orthogonalized tensor deflation via random tensor theory. *arXiv preprint arXiv:2302.05798*, 2023.
- [37] Nathan Srebro and Russ R Salakhutdinov. Collaborative filtering in a non-uniform world: Learning with the weighted trace norm. *Advances in Neural Information Processing Systems*, 23, 2010.
- [38] Dominik Stöger and Mahdi Soltanolkotabi. Small random initialization is akin to spectral learning: Optimization and generalization guarantees for overparameterized low-rank matrix reconstruction. *Advances in Neural Information Processing Systems*, 34:23831–23843, 2021.
- [39] M. J. Todd. Semidefinite optimization. *Acta Numerica*, page 515–560, 2001.
- [40] Tian Tong, Cong Ma, Ashley Prater-Bennette, Erin Tripp, and Yuejie Chi. Scaling and scalability: Provable nonconvex low-rank tensor estimation from incomplete measurements. *Journal of Machine Learning Research*, 23(163):1–77, 2022.
- [41] Lieven Vandenberghe and Stephen Boyd. Semidefinite programming. *SIAM review*, 38(1):49–95, 1996.
- [42] Roman Vershynin. Introduction to the non-asymptotic analysis of random matrices. *arXiv preprint arXiv:1011.3027*, 2010.

- [43] Rachel Ward and Tamara Kolda. Convergence of Alternating Gradient Descent for Matrix Factorization. *Advances in Neural Information Processing Systems*, 36:22369–22382, 2023.
- [44] Xiaoxia Wu, Edgar Dobriban, Tongzheng Ren, Shanshan Wu, Zhiyuan Li, Suriya Gunasekar, Rachel Ward, and Qiang Liu. Implicit regularization and convergence for weight normalization. *Advances in Neural Information Processing Systems*, 33:2835–2847, 2020.
- [45] Nuoya Xiong, Lijun Ding, and Simon Shaolei Du. How over-parameterization slows down gradient descent in matrix sensing: The curses of symmetry and initialization. In *International Conference on Learning Representations*, 2024.
- [46] Weihang Xu and Simon Du. Over-parameterization exponentially slows down gradient descent for learning a single neuron. In *Conference on Learning Theory*, pages 1155–1198. PMLR, 2023.
- [47] Xingyu Xu, Yandi Shen, Yuejie Chi, and Cong Ma. The power of preconditioning in over-parameterized low-rank matrix sensing. In *International Conference on Machine Learning*, pages 38611–38654. PMLR, 2023.
- [48] Tian Ye and Simon S Du. Global convergence of gradient descent for asymmetric low-rank matrix factorization. *Advances in Neural Information Processing Systems*, 34:1429–1439, 2021.
- [49] Jialun Zhang, Salar Fattahi, and Richard Y Zhang. Preconditioned gradient descent for over-parameterized nonconvex matrix factorization. In *Advances in Neural Information Processing Systems*, volume 34, pages 5985–5996, 2021.
- [50] Yilang Zhang, Bingcong Li, and Georgios B Giannakis. Reflora: Refactored low-rank adaptation for efficient fine-tuning of large models. *arXiv preprint arXiv:2505.18877*, 2025.
- [51] Mo Zhou, Weihang Xu, Maryam Fazel, and Simon S Du. Global convergence of gradient EM for over-parameterized gaussian mixtures. *arXiv preprint arXiv:2506.06584*, 2025.
- [52] Jiacheng Zhuo, Jeongyeol Kwon, Nhat Ho, and Constantine Caramanis. On the computational and statistical complexity of over-parameterized matrix sensing. *Journal of Machine Learning Research*, 25(169):1–47, 2024.

Appendix A. Algorithm 1 derivation

We consider the overparameterized setting $r > r_A$ and apply a joint update on both $\mathbf{X}_t$ and Θ_t in an alternating manner. Let $\mathcal{M}^* : \mathbb{R}^n \mapsto \mathbb{S}^m$ denote the adjoint of $\mathcal{M}$ with explicit form $\mathcal{M}^*(\mathbf{y}) = \sum_{i=1}^n y_i \mathbf{M}_i$. The Stiefel manifold $\text{St}(m, r)$ is embedded in the Euclidean space, then we first compute the Euclidean gradient of $\mathbf{X}_t$ as

$$\begin{aligned} \tilde{\mathbf{G}}_t &= [\mathcal{M}^* \mathcal{M}(\mathbf{X}_t \Theta_t \mathbf{X}_t^\top - \mathbf{A})] \mathbf{X}_t \Theta_t \\ &= (\mathbf{X}_t \Theta_t \mathbf{X}_t^\top - \mathbf{A}) \mathbf{X}_t \Theta_t + [(\mathcal{M}^* \mathcal{M} - \mathcal{I})(\mathbf{X}_t \Theta_t \mathbf{X}_t^\top - \mathbf{A})] \mathbf{X}_t \Theta_t. \end{aligned} \quad (5)$$

Projecting it onto the tangent space of $\text{St}(m, r)$ yields the Riemannian gradient

$$\mathbf{G}_t := (\mathbf{I}_m - \mathbf{X}_t \mathbf{X}_t^\top) \tilde{\mathbf{G}}_t + \frac{\mathbf{X}_t}{2} (\mathbf{X}_t^\top \tilde{\mathbf{G}}_t - \tilde{\mathbf{G}}_t^\top \mathbf{X}_t).$$

Using polar retraction, the update of $\mathbf{X}_t$ along the direction $\mathbf{G}_t$ with stepsize η is given by

$$\mathbf{X}_{t+1} = (\mathbf{X}_t - \eta \mathbf{G}_t) (\mathbf{I}_r + \eta^2 \mathbf{G}_t^\top \mathbf{G}_t)^{-1/2}.$$

For the magnitude variable Θ_t , the Euclidean gradient is

$$\mathbf{K}_t = \frac{1}{2} \mathbf{X}_{t+1}^\top [\mathcal{M}^* \mathcal{M}(\mathbf{X}_{t+1} \Theta_t \mathbf{X}_{t+1}^\top - \mathbf{A})] \mathbf{X}_{t+1}. \quad (6)$$

Denoting the identity mapping by $\mathcal{I}$, the update of Θ_t with stepsize μ becomes

$$\begin{aligned} \Theta_{t+1} &= \Theta_t - \frac{\mu}{2} \mathbf{X}_{t+1}^\top [\mathcal{M}^* \mathcal{M}(\mathbf{X}_{t+1} \Theta_t \mathbf{X}_{t+1}^\top - \mathbf{A})] \mathbf{X}_{t+1} \\ &= \mathbf{X}_{t+1}^\top \mathbf{A} \mathbf{X}_{t+1} - \mathbf{X}_{t+1}^\top [(\mathcal{M}^* \mathcal{M} - \frac{\mu}{2} \mathcal{I})(\mathbf{X}_{t+1} \Theta_t \mathbf{X}_{t+1}^\top - \mathbf{A})] \mathbf{X}_{t+1}. \end{aligned} \quad (7)$$

Appendix B. Diving deeper into the initial phase

In this section, we take a closer look at the convergence of RGD on WN in the initial phase, that is $t \leq t_0$, or equivalently $\text{Tr}(\Phi_t \Phi_t^\top) \leq r_A - 0.5$. Here, $\Phi_t := \mathbf{U}^\top \mathbf{X}_t$ depicts the alignment between $\text{span}(\mathbf{U})$ and $\text{span}(\mathbf{X}_t)$ and its singular values coincide with the cosine of the principle angles between these two subspaces [5]. Our numerical experiments in Figure 2 indicate that RGD traverse a sequence of saddles. The saddle-to-saddle behavior is known for GD on (1) [20, 26]. This section shows that this behavior persists for (2), yet can be faster with a higher level of overparameterization. To bypass the randomness associated with $\mathbf{M}_i$, we begin by pinpointing the saddles for the population loss, i.e., problem (2) in the infinite data limit $n \rightarrow \infty$. More precisely, the objective is given by $f_\infty(\mathbf{X}, \Theta) = \frac{1}{4} \|\mathbf{X} \Theta \mathbf{X}^\top - \mathbf{A}\|_F^2$.

Lemma 3 *For a given $\rho \in \{0, 1, \dots, r_A - 1\}$, let $\mathbf{A}_\rho$ be the best rank- ρ approximation of $\mathbf{A}$, i.e., $\mathbf{A}_\rho = \arg \min_{\text{rank}(\hat{\mathbf{A}}) \leq \rho} \|\hat{\mathbf{A}} - \mathbf{A}\|_F^2$. In particular, we let $\mathbf{A}_0 = \mathbf{0}$. A point $(\mathbf{X}, \Theta)$ is a saddle of the population loss f_∞ if $\mathbf{X} \Theta \mathbf{X}^\top = \mathbf{A}_\rho$ and $\text{Tr}(\mathbf{X}^\top \mathbf{U} \mathbf{U}^\top \mathbf{X}) = \rho$.*

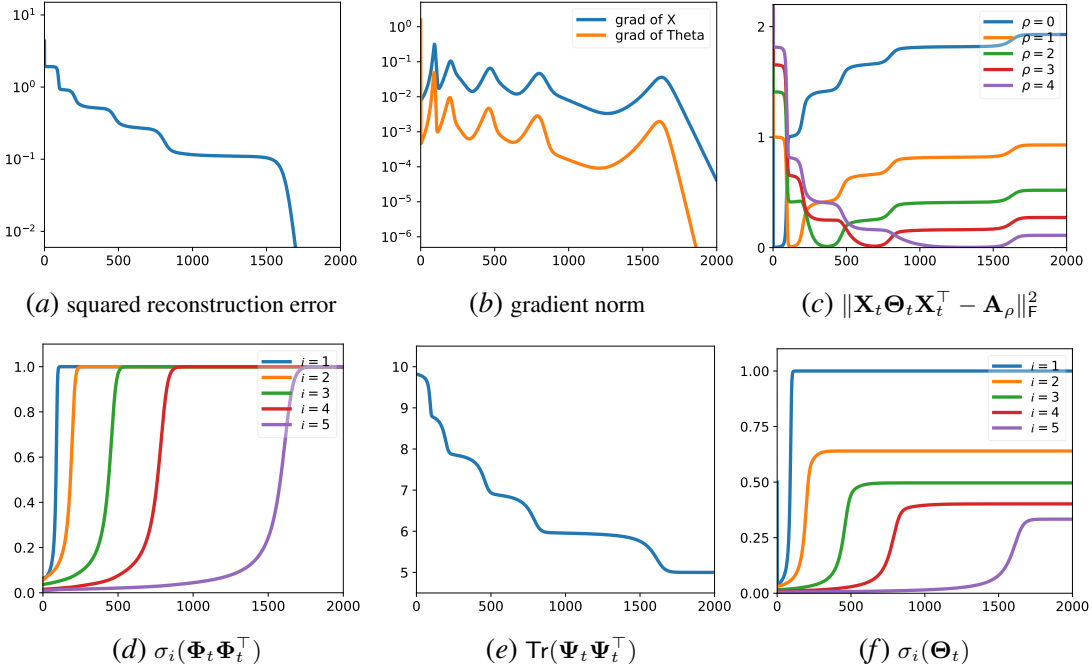


Figure 2: The saddle-to-saddle (i.e., sequential learning) behaviors in WN. The x-axis corresponds to the iteration number. (a) Each plateau signifies a saddle; (b) gradient norm at saddles drops by orders; (c) saddles strongly relate to the best rank- ρ approximation of $\mathbf{A}$; (d) sequential learning in the alignment between $\mathbf{X}_t$ and $\mathbf{U}$; (e) sequential learning in the alignment between $\mathbf{X}_t$ and $\mathbf{U}_\perp$; and, (f) sequential pattern in the magnitude variable Θ_t .

Lemma 3 indicates that the saddles of f_∞ are closely related to the best rank- ρ approximation of $\mathbf{A}$. It further suggests that a saddle-to-saddle dynamic is aligned with incremental learning⁴: the algorithm successively learns $\mathbf{A}_\rho$ for increasing ρ until the ground-truth matrix is recovered. Lemma 4 below shows that in the finite-sample regime, the saddles of f_∞ also have small gradient norm on f , i.e., no larger than $\mathcal{O}(\frac{(r-r_A)^6}{\kappa^2 m^2 r^4 r_A})$ under the parameter choices of Theorem 2.

Lemma 4 Assume that $\mathcal{M}(\cdot)$ is $(r + r_A + 1, \delta)$ -RIP, and $\|\Theta\| \leq 2$, the finite sample loss in (2) satisfies $\|\nabla_{\mathbf{X}}^R f_\infty(\mathbf{X}, \Theta) - \nabla_{\mathbf{X}}^R f(\mathbf{X}, \Theta)\|_F \leq 12m\delta$ and $\|\nabla_{\Theta} f_\infty(\mathbf{X}, \Theta) - \nabla_{\Theta} f(\mathbf{X}, \Theta)\|_F \leq \frac{3}{2}m\delta$. Here, $\nabla_{\mathbf{X}}^R$ denotes the Riemannian gradient with respect to $\mathbf{X}$.

Having characterized the saddles, we now turn to the saddle-to-saddle trajectory in Figure 2. This figure traces the optimization trajectory of Algorithm 1 on WN with $m = 300$, $r_A = 5$, $r = 10$, and $\kappa = 3$, with more details shown in Appendix I.2. Figure 2(a) plots the squared reconstruction error across iterations. Each plateau marks escape from a saddle, as confirmed by the small gradient norm shown in Figure 2(b). Figure 2(c) further shows that these saddles are exactly those characterized in Lemma 3, where $\|\mathbf{X}_t \Theta_t \mathbf{X}_t^\top - \mathbf{A}_\rho\|_F^2$ for $\rho \in \{0, 1, \dots, r_A - 1\}$ stays close to 0 sequentially. In other words, each saddle escape corresponds to leaving the neighborhood of $\mathbf{A}_\rho$.

4. Also known as deflation; see e.g., [2, 16, 36]

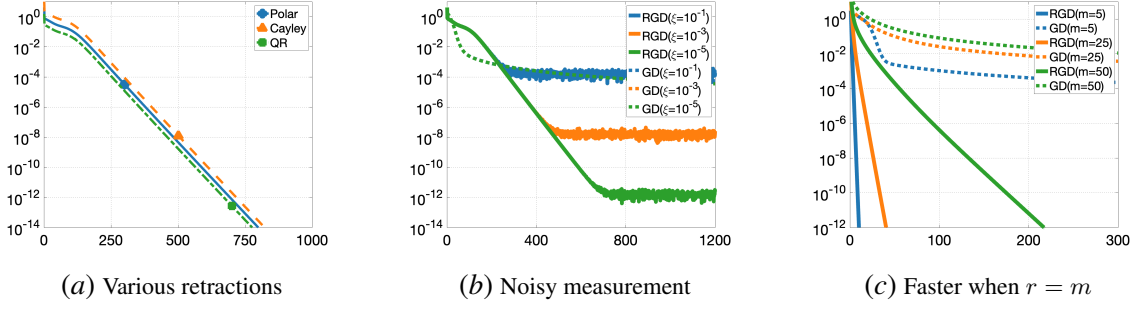


Figure 3: Additional numerical results. X-axis: squared reconstruction error; Y-axis: iteration.

In addition, the optimization variables, geometrically interpretable as direction and magnitude, also exhibit a sequential learning behavior. For the direction variable $\mathbf{X}_t$, the singular values of $\Phi_t \Phi_t^\top$ (which characterize the squared cosine of the principle angles between $\mathbf{X}_t$ and $\mathbf{U}$) are visualized in Figure 2(d). Further, let $\mathbf{U}_\perp \in \mathbb{R}^{m \times (m-r_A)}$ be an orthonormal basis for the orthogonal complement of $\text{span}(\mathbf{U})$. The alignment of $\mathbf{X}_t$ and $\mathbf{U}_\perp$ is plotted in Figure 2(e), with the alignment matrix defined as $\Psi_t := \mathbf{U}_\perp^\top \mathbf{X}_t$. The singular values of the magnitude variable Θ_t are plotted in Figure 2(f). A clear sequential learning pattern is observed among all these figures.

Lastly, we highlight that *polynomial* time is needed to escape all saddles: Theorem 2 bounds the duration of this phase to be at most $\mathcal{O}\left(\frac{\kappa^4 m^4 r^4 r_A^2}{(r-r_A)^8}\right)$ iterations. This bound decreases with larger r , indicating that overparameterization facilitates saddle escape under WN.

Appendix C. Additional experiments

In this section, we take additional experiments to reveal other interesting behaviors of WN with both synthetic and real-world data. More details on experiments setup are deferred to Appendix I.3.

Alternative manners of retraction. Although our algorithm for WN tackles only the polar retraction, other popular retractions share similar performance. In Figure 3(a), we plot the performance of RGD with different manners for retraction, such as Cayley and QR, on an instance of (2) with $m = 10, r = 5, r_A = 3, \kappa = 2, n = 1000$. The three curves of squared reconstruction errors nearly coincide. For better visualization, we scale the errors of Cayley and QR by 3 and $1/3$, respectively.

Noisy measurements. To examine the robustness of WN, we consider a setting with corrupted labels, i.e., $y_i = \text{Tr}(\mathbf{M}_i^\top \mathbf{A}) + b_i$ for i.i.d. Gaussian noise $b_i \sim \mathcal{N}(0, \xi^2)$. Figure 3(b) compares WN with the vanilla problem (1) under the choices of $\xi = 10^{-1}, \xi = 10^{-3}$, and $\xi = 10^{-5}$. It can be seen that RGD holds a linear rate under all choices of ξ , and the final squared reconstruction error stabilizes around $\mathcal{O}(\xi^2)$. On the other hand, the error of GD is mainly confined by its slow convergence rate. This demonstrates that the power of WN carries to noisy settings as well.

Full rank case with $r = m$. WN shows remarkable effectiveness in the special setting with $r = m$. Instances with three different choices of $m = r \in \{5, 25, 50\}, r_A \in \{2, 10, 20\}, \kappa \in \{1, 15, 50\}$, and $n = 5000$ are plotted in Figure 3(c). The faster convergence arises from the fact that at initialization, $\mathbf{X}_0 \in \text{St}(m, m)$ already aligns with the target subspace spanned by $\mathbf{U}$, i.e., $\text{Tr}(\mathbf{I}_{r_A} - \Phi_0 \Phi_0^\top) = 0$.

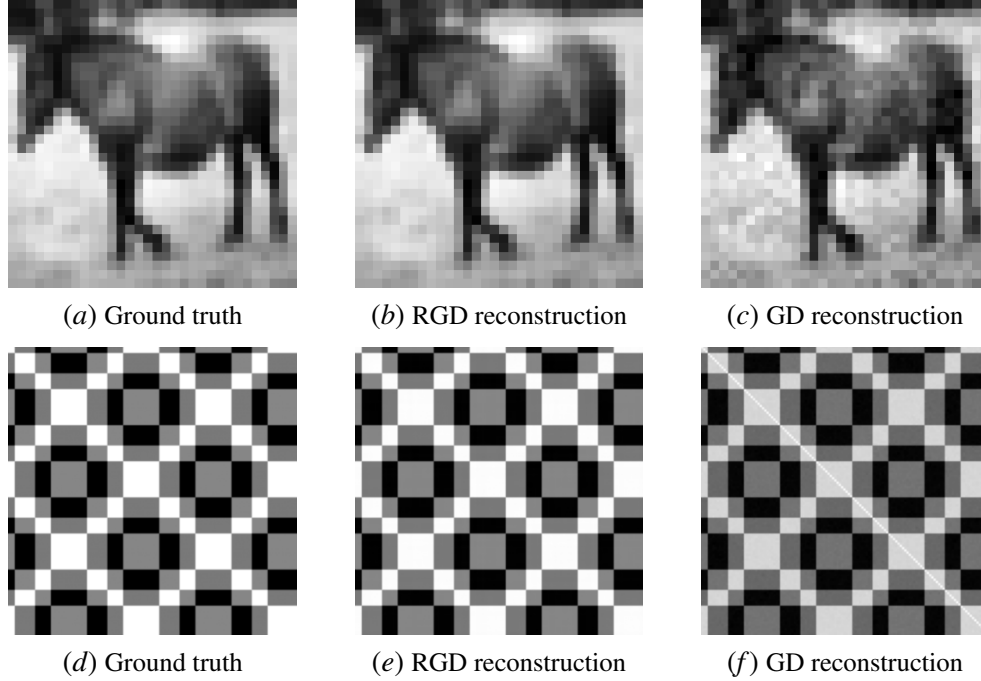


Figure 4: The advantages of WN on image reconstruction tasks.

Equivalently, this is the case where only the magnitude Θ is optimized. The faster convergence in this case implies that learning the correct direction (i.e., $\mathbf{U}$) is more challenging than magnitude.

Image reconstruction experiments. Beyond the above synthetic experiments, we further evaluate the advantages of WN with real-world data on two image reconstruction tasks.

The first experiment follows [15] to consider a generalized phase retrieval problem on a 32×32 horse image from the CIFAR-10 dataset [22]. The image is converted to grayscale and vectorized as $\mathbf{a} \in \mathbb{R}^{1024}$. Standard lifting reformulation converts this problem to a sensing problem on a rank-one ground-truth matrix $\mathbf{A} = \mathbf{a}\mathbf{a}^\top \in \mathbb{S}_+^{1024}$; see [8]. The second considers direct matrix sensing of a structured image given by $\mathbf{A} \in \mathbb{S}_+^{128}$ with $r_A = 2$. In both cases, we set the overparameterization level to $r = 100$ and use $n = 50000$ feature matrices. RGD and GD are randomly initialized and run for $t_{\text{RGD}} = 100, t_{\text{GD}} = 200$ iterations in both experiments to make the overall runtime comparable; see Appendix I.3.2 for details.

The reconstructions from the two experiments are presented in Figure 4. As shown, WN enables RGD to achieve more accurate recovery of the ground truth compared to GD. These results demonstrate that WN provides a significant improvement for image reconstruction problems.

Appendix D. More on backgrounds

D.1. Polar decomposition

The definition of the polar decomposition is provided below; see [17, Section 9.4.3] for a detailed discussion and theoretical background.

Definition 5 *The polar decomposition of a matrix $\mathbf{X} \in \mathbb{R}^{m \times r}$ with $m \geq r$ is defined as*

$$\mathbf{X} = \mathbf{U}\mathbf{P},$$

where $\mathbf{U} \in \mathbb{R}^{m \times r}$ has orthonormal columns and $\mathbf{P} \in \mathbb{S}_+^r$ is a positive semi-definite matrix.

This decomposition can be interpreted as expressing $\mathbf{X}$ as the product of directions ($\mathbf{U}$) and a magnitude part ($\mathbf{P}$). It is unique when $\mathbf{X}$ has full column rank.

D.2. Restricted Isometry Property (RIP)

The RIP condition [32] in Definition 1 is a standard assumption in matrix sensing, ensuring that the linear measurement operator approximately preserves the Frobenius norm of low-rank matrices. This property has been verified to hold with high probability for a wide variety of measurement operators. The following lemma establishes RIP for Gaussian design measurements.

Lemma 6 [10] *If $\mathcal{M}(\cdot)$ is a Gaussian random measurement ensemble, i.e., the entries of $\{\mathbf{M}_i\}_{i=1}^n \subset \mathbb{S}^m$ are independent up to symmetry with diagonal elements sampled from $\mathcal{N}(0, 1/n)$ and off-diagonal elements from $\mathcal{N}(0, 1/2n)$, then with high probability, $\mathcal{M}(\cdot)$ is (r, δ_r) -RIP, as long as $n \geq Cmr/\delta_r^2$ for some sufficiently large universal constant $C > 0$.*

Appendix E. Related work

Overparameterized matrix sensing. Overparameterized matrix sensing arises from many machine learning and signal processing applications such as collaborative filtering and phase retrieval [9, 15, 35, 37]. The problem is now a canonical benchmark in theoretical deep learning, mainly because the loss landscape is riddled with saddle points and lacks global smoothness or a global PL condition. Convergence analyses for various algorithms on its population loss, i.e., matrix factorization, can be found in [21, 24, 43, 50]. Small random initialization in overparameterized matrix sensing has been studied in [20, 38, 45, 47], while [11, 29, 52] are based on spectral initialization. Besides saddle escaping under small initialization, another intriguing phenomenon is that overparameterization can exponentially slow the convergence of GD compared to the exactly parameterized case [45, 52]. Our work proves that WN avoids this slowdown and achieves an improved rate. Moreover, additional overparameterization leads to faster convergence and lower sample complexity.

Overparameterization in other nonconvex estimation problems. Beyond matrix sensing, the role of overparameterization has also been examined in several nonconvex estimation problems. For matrix completion, [30] proves that the vanilla gradient descent with small initialization converges to the ground truth without requiring any explicit regularization, even in the overparameterized scenario. In Gaussian mixture learning, [51] establishes that Gradient EM achieves global convergence at a polynomial rate with polynomial samples, when the model is mildly overparameterized. For neural network training, [46] shows that in the problem of learning a single neuron with ReLU activation, randomly initialized gradient descent can suffer from an exponential slowdown when the model is overparameterized. These studies illustrate that overparameterization appears in diverse problem settings, while its precise influence on the convergence behavior is problem-dependent.

Riemannian optimization. Riemannian optimization is naturally connected to WN because the set of directions of a vector forms a unit sphere, which is a smooth manifold. It extends gradient-based

methods to problems with smooth manifold constraints [1, 6]. We rely on standard notions. In its simplest form, Riemannian gradient descent (RGD) iteratively moves along the negative direction of Riemannian gradient, which can be thought as gradient projected to the tangent space, and then maps the iterate back to the manifold via a retraction.

Appendix F. Proof strategies and supporting lemmas

F.1. Proof strategies

To establish convergence of Theorem 2, we analyze the evolution of the principle angles between $\text{span}(\mathbf{U})$ and $\text{span}(\mathbf{X}_t)$. Specifically, we track the quantity $\text{Tr}(\mathbf{I}_{r_A} - \Phi_t \Phi_t^\top)$. This term reflects the subspace alignment error between $\text{span}(\mathbf{U})$ and $\text{span}(\mathbf{X}_t)$. For notational convenience, we set $\mu = 2$, which is consistent with our choice in Theorem 2.

Our proof is structured into two phases:

- Phase I (Initial phase): When the alignment error is large, i.e., $\text{Tr}(\mathbf{I}_{r_A} - \Phi_t \Phi_t^\top) \geq 0.5$, we rely on the fact that $\sigma_{r_A}^2(\Phi_t)$ remains bounded away from zero. This property guarantees that the alignment error decreases by at least a constant amount at each iteration.
- Phase II (Linearly convergent phase): Once $\text{Tr}(\mathbf{I}_{r_A} - \Phi_t \Phi_t^\top) < 0.5$, we enter a contraction regime. In this regime, we establish that the reconstruction error and the alignment error decrease jointly, governed by a coupled inequality system.

Throughout both phases, two error terms caused by the limited number of measurements must be carefully controlled. Formally, we introduce the following definitions:

$$\begin{aligned}\Delta_t &:= (\mathcal{M}^* \mathcal{M} - \mathcal{I})(\mathbf{X}_{t+1} \Theta_t \mathbf{X}_{t+1}^\top - \mathbf{A}), \\ \Xi_t &:= (\mathcal{M}^* \mathcal{M} - \mathcal{I})(\mathbf{X}_t \Theta_t \mathbf{X}_t^\top - \mathbf{A}).\end{aligned}$$

Incorporating these two error terms, we can rewrite $\tilde{\mathbf{G}}_t$ and Θ_{t+1} as follows:

$$\begin{aligned}\tilde{\mathbf{G}}_t &= (\mathbf{X}_t \Theta_t \mathbf{X}_t^\top - \mathbf{A}) \mathbf{X}_t \Theta_t + \Xi_t \mathbf{X}_t \Theta_t, \\ \Theta_{t+1} &= \mathbf{X}_{t+1}^\top \mathbf{A} \mathbf{X}_{t+1} - \mathbf{X}_{t+1}^\top \Delta_t \mathbf{X}_{t+1}.\end{aligned}$$

These two terms will be used repeatedly throughout the proofs in the following sections.

F.2. Supporting lemmas

Since Theorem 2 considers random initialization, it is conditioned on the following high-probability event F , which gives a lower bound on the smallest singular value of $\Phi_0 = \mathbf{U}^\top \mathbf{X}_0$:

$$F = \left\{ \sigma_{r_A}^2(\mathbf{U}^\top \mathbf{X}_0) \geq \frac{(r - r_A)^2}{c_1 m r} \right\},$$

where $c_1 > \max\{1, 36C_1^2\}$ is a universal constant, with C_1 given in Lemma 20.

Lemma 7 *With respect to the randomness in $\mathbf{X}_0$, event F occurs with probability at least*

$$1 - \exp(-m/2) - C_3^{r-r_A+1} - \exp(-C_2 r),$$

where $C_2 > 0$ and $C_3 = \frac{6C_1}{\sqrt{c_1}} \in (0, 1)$ are universal constants.

This lemma ensures that the smallest singular value of the initial alignment between $\mathbf{U}$ and $\mathbf{X}_0$ is bounded away from zero with high probability, which is critical to initialize Phase I.

Lemma 8 *Suppose that at iteration t , the alignment error satisfies that*

$$\text{Tr}(\mathbf{I}_{r_A} - \Phi_t \Phi_t^\top) \leq \rho,$$

then the reconstruction error at iteration t satisfies that

$$\|\mathbf{X}_t \Theta_t \mathbf{X}_t^\top - \mathbf{A}\|_F \leq 2\sqrt{\rho} + \|\Delta_{t-1}\|_F.$$

The lemma above connects the reconstruction error $\|\mathbf{X}_t \Theta_t \mathbf{X}_t^\top - \mathbf{A}\|_F$ with the alignment error $\text{Tr}(\mathbf{I}_{r_A} - \Phi_t \Phi_t^\top)$ and the measurement error $\|\Delta_{t-1}\|_F$. It means that the reconstruction error is small once $\mathbf{X}_t$ and $\mathbf{U}$ are sufficiently aligned and the measurement error is small.

Lemma 9 *Assuming $\eta < \frac{1}{300\kappa^2 r_A}$, $\mathcal{M}(\cdot)$ is $(r + r_A + 1, \delta)$ -RIP with $\delta = \frac{\xi}{\sqrt{mr}}$, $\xi \in [0, 1]$, and $\|\Theta_t\| \leq 2$. Then, the measurement errors satisfy that*

$$\begin{aligned} \|\Delta_t\|_F &\leq \xi \|\mathbf{X}_t \Theta_t \mathbf{X}_t^\top - \mathbf{A}\|_F, \\ \|\Xi_t\|_F &\leq \xi \|\mathbf{X}_t \Theta_t \mathbf{X}_t^\top - \mathbf{A}\|_F. \end{aligned}$$

This provides upper bounds on the norm of the measurement error terms Δ_t, Ξ_t by the reconstruction error $\|\mathbf{X}_t \Theta_t \mathbf{X}_t^\top - \mathbf{A}\|_F$, which is guaranteed by the RIP property of $\mathcal{M}(\cdot)$.

Lemma 10 *Let $\chi_t := (\|\Delta_{t-1}\| + \|\Xi_t\|)^2 + \sqrt{\text{Tr}(\mathbf{I}_{r_A} - \Phi_t \Phi_t^\top)}(\|\Delta_{t-1}\| + \|\Xi_t\|)$,*

$$\begin{aligned} \beta_t &:= \sigma_1(\mathbf{I}_{r_A} - \Phi_t \Phi_t^\top), \\ \mathbf{H}_t &:= (\mathbf{I}_m - \mathbf{X}_t \mathbf{X}_t^\top)(\mathbf{A} \mathbf{X}_t \mathbf{X}_t^\top \Delta_{t-1} \mathbf{X}_t + \Xi_t \mathbf{X}_t \Theta_t) \\ &\quad + \frac{1}{2}(\mathbf{X}_t \mathbf{X}_t^\top \Xi_t \mathbf{X}_t \Theta_t - \mathbf{X}_t \Theta_t \mathbf{X}_t^\top \Xi_t \mathbf{X}_t) \\ &\quad + \frac{1}{2}(\mathbf{X}_t \mathbf{X}_t^\top \mathbf{A} \mathbf{X}_t \mathbf{X}_t^\top \Delta_{t-1} \mathbf{X}_t - \mathbf{X}_t \mathbf{X}_t^\top \Delta_{t-1} \mathbf{X}_t \mathbf{X}_t^\top \mathbf{A} \mathbf{X}_t). \end{aligned}$$

Assuming $\|\Delta_{t-1}\|_F, \|\Xi_t\|_F \leq 1$, $\eta \leq \frac{1}{10r_A}$, and $\|\Theta_t\| \leq 2$, then the following inequality holds:

$$\begin{aligned} &\text{Tr}(\mathbf{I}_{r_A} - \Phi_{t+1} \Phi_{t+1}^\top) - \text{Tr}(\mathbf{I}_{r_A} - \Phi_t \Phi_t^\top) \\ &\leq \eta^2(\beta_t + 16\chi_t) \text{Tr}(\Phi_t \Phi_t^\top) - \frac{2\eta(1 - \eta^2\beta_t - 16\eta^2\chi_t)\sigma_{r_A}^2(\Phi_t)}{\kappa^2} \text{Tr}((\mathbf{I}_{r_A} - \Phi_t \Phi_t^\top) \Phi_t \Phi_t^\top) \\ &\quad + 2\eta \sqrt{\text{Tr}(\mathbf{I}_{r_A} - \Phi_t \Phi_t^\top)}(\|\Delta_{t-1}\|_F + 2\|\Xi_t\|_F) \\ &\quad + 2\eta^2 \sqrt{\text{Tr}(\mathbf{I}_{r_A} - \Phi_t \Phi_t^\top)} \|\mathbf{H}_t\|_F. \end{aligned} \tag{8}$$

This lemma quantifies how the alignment error $\text{Tr}(\mathbf{I}_{r_A} - \Phi_t \Phi_t^\top)$ evolves between iterations. This is the key lemma that drives the reduction of the alignment error.

Lemma 11 *Assuming $\mathcal{M}(\cdot)$ is $(r + r_A + 1, \delta)$ -RIP with $\delta \leq \frac{1}{3\sqrt{m}}$. If $\|\Theta_t\| \leq 2$, then it is guaranteed that*

$$\|\Theta_{t+1}\| \leq 2.$$

As shown in Lemmas 9 and Lemma 10, the analyses require that $\|\Theta_t\|$ is upper bounded by 2. This condition has already been guaranteed at initialization. Moreover, based on the update rule of Θ_t given in (7), we observe that $\|\Theta_t\|$ remains close to $\|\mathbf{X}_t^\top \mathbf{A} \mathbf{X}_t\|$ in each iteration.

Lemma 12 *Assuming $\eta \leq 1$, $\mathcal{M}(\cdot)$ is $(r + r_A + 1, \delta)$ -RIP, and $\|\Theta_{t-1}\|, \|\Theta_t\| \leq 2$, we have that*

$$\Psi_{t+1} \Psi_{t+1}^\top \preceq (1 + 6\eta(\sqrt{r_A} + 2\sqrt{r})\delta)^2 \Psi_t \Psi_t^\top + (4\eta(\sqrt{r_A} + 2\sqrt{r})\delta + 28\eta^2(\sqrt{r_A} + 2\sqrt{r})^2 \delta^2) \mathbf{I}_{m-r_A}.$$

Moreover, it is also guaranteed that

$$\Psi_1 \Psi_1^\top \preceq (1 + 2\eta + 2\eta(\sqrt{r_A} + 2\sqrt{r})\delta)^2 \Psi_0 \Psi_0^\top + (12\eta(\sqrt{r_A} + 2\sqrt{r})\delta + 8\eta^2(\sqrt{r_A} + 2\sqrt{r})^2 \delta^2) \mathbf{I}_{m-r_A}.$$

This lemma establishes an upper bound on the growth of $\Psi_t \Psi_t^\top$. Together with Lemma 15, we can ensure that $\sigma_{r_A}^2(\Phi_t)$ remains adequately large throughout Phase I.

Lemma 13 *Assuming $\eta \leq \frac{1}{500r_A}$, $\mathcal{M}(\cdot)$ is $(r + r_A + 1, \delta)$ -RIP with $\delta \leq \frac{1}{\sqrt{m}}$, and $\|\Theta_t\| \leq 2$. Then for any $t \geq 0$, the alignment error satisfies that*

$$\text{Tr}(\mathbf{I}_{r_A} - \Phi_{t+1} \Phi_{t+1}^\top) \leq \text{Tr}(\mathbf{I}_{r_A} - \Phi_t \Phi_t^\top) + 0.1.$$

This guarantees that the alignment error $\text{Tr}(\mathbf{I}_{r_A} - \Phi_t \Phi_t^\top)$ does not increase too much in one step when we choose suitable stepsize η , which is crucial for bridging Phase I and Phase II.

Appendix G. Proofs

G.1. Proof of Lemma 7

Proof Since the initialization $\mathbf{X}_0$ satisfies the conditions stated in Lemma 21, we can apply the lemma directly. In particular, substituting $\tau = \frac{6}{\sqrt{c_1}}$ yields the desired result. $\blacksquare$

G.2. Proof of Lemma 8

Proof Directly substituting the expression of Θ_t into the Frobenius norm term, we have that

$$\begin{aligned} \|\mathbf{X}_t \Theta_t \mathbf{X}_t^\top - \mathbf{A}\|_F &= \|\mathbf{X}_t \mathbf{X}_t^\top \mathbf{A} \mathbf{X}_t \mathbf{X}_t^\top - \mathbf{A} - \mathbf{X}_t \mathbf{X}_t^\top \Delta_{t-1} \mathbf{X}_t \mathbf{X}_t^\top\|_F \\ &\leq \|\mathbf{X}_t \mathbf{X}_t^\top \mathbf{A} \mathbf{X}_t \mathbf{X}_t^\top - \mathbf{A}\|_F + \|\mathbf{X}_t \mathbf{X}_t^\top \Delta_{t-1} \mathbf{X}_t \mathbf{X}_t^\top\|_F \\ &\leq \|\mathbf{X}_t \mathbf{X}_t^\top \mathbf{A} \mathbf{X}_t \mathbf{X}_t^\top - \mathbf{A} \mathbf{X}_t \mathbf{X}_t^\top\|_F + \|\mathbf{A} \mathbf{X}_t \mathbf{X}_t^\top - \mathbf{A}\|_F + \|\mathbf{X}_t \mathbf{X}_t^\top \Delta_{t-1} \mathbf{X}_t \mathbf{X}_t^\top\|_F \\ &\stackrel{(a)}{\leq} 2\|\Sigma\| \|(\mathbf{I}_m - \mathbf{X}_t \mathbf{X}_t^\top) \mathbf{U}\|_F + \|\Delta_{t-1}\|_F \end{aligned}$$

$$\begin{aligned}
 &= 2\|\Sigma\|\sqrt{\text{Tr}(\mathbf{I}_{r_A} - \Phi_t \Phi_t^\top)} + \|\Delta_{t-1}\|_F \\
 &\leq 2\|\Sigma\|\sqrt{\rho} + \|\Delta_{t-1}\|_F \\
 &= 2\sqrt{\rho} + \|\Delta_{t-1}\|_F,
 \end{aligned}$$

where (a) is by the inequality $\|\mathbf{AB}\|_F \leq \|\mathbf{A}\| \|\mathbf{B}\|_F$ that is valid for any conformable matrices. $\blacksquare$

G.3. Proof of Lemma 9

Proof We first prove that $\|\mathbf{G}_t\|_F \leq 2\|\tilde{\mathbf{G}}_t\|_F$. Indeed,

$$\begin{aligned}
 \|\mathbf{G}_t\|_F &= \|(\mathbf{I}_m - \mathbf{X}_t \mathbf{X}_t^\top) \tilde{\mathbf{G}}_t + \frac{\mathbf{X}_t}{2} (\mathbf{X}_t^\top \tilde{\mathbf{G}}_t - \tilde{\mathbf{G}}_t^\top \mathbf{X}_t)\|_F \\
 &\leq \|\mathbf{I}_m - \mathbf{X}_t \mathbf{X}_t^\top\| \|\tilde{\mathbf{G}}_t\|_F + \|\mathbf{X}_t \mathbf{X}_t^\top\| \|\tilde{\mathbf{G}}_t\|_F \\
 &\leq 2\|\tilde{\mathbf{G}}_t\|_F.
 \end{aligned}$$

We now proceed to estimate the update distance $\|\mathbf{X}_{t+1} - \mathbf{X}_t\|_F$.

$$\begin{aligned}
 \|\mathbf{X}_{t+1} - \mathbf{X}_t\|_F &= \|(\mathbf{X}_t - \eta \mathbf{G}_t)(\mathbf{I}_r + \eta^2 \mathbf{G}_t^\top \mathbf{G}_t)^{-1/2} - \mathbf{X}_t\|_F \\
 &\leq \|\mathbf{X}_t\|_F \|(\mathbf{I}_r + \eta^2 \mathbf{G}_t^\top \mathbf{G}_t)^{-1/2} - \mathbf{I}_r\|_F + \|\eta \mathbf{G}_t(\mathbf{I}_r + \eta^2 \mathbf{G}_t^\top \mathbf{G}_t)^{-1/2}\|_F \\
 &\leq \|\mathbf{X}_t\|_F \|(\mathbf{I}_r + \eta^2 \mathbf{G}_t^\top \mathbf{G}_t)^{-1/2} - \mathbf{I}_r\|_F + \eta \|(\mathbf{I}_r + \eta^2 \mathbf{G}_t^\top \mathbf{G}_t)^{-1/2}\| \|\mathbf{G}_t\|_F \\
 &\leq \sqrt{r} \|(\mathbf{I}_r + \eta^2 \mathbf{G}_t^\top \mathbf{G}_t)^{-1/2} - \mathbf{I}_r\|_F + \eta \|\mathbf{G}_t\|_F \\
 &\leq \sqrt{r} \|(\mathbf{I}_r + \eta^2 \mathbf{G}_t^\top \mathbf{G}_t)^{-1/2} - \mathbf{I}_r\|_F + 2\eta \|\tilde{\mathbf{G}}_t\|_F \\
 &\leq \sqrt{r} (1 - (1 + \eta^2 \sigma_1(\mathbf{G}_t^\top \mathbf{G}_t))^{-1/2}) + 2\eta \|\tilde{\mathbf{G}}_t\|_F \\
 &\stackrel{(a)}{\leq} \sqrt{r} (1 - \frac{1}{1 + (\eta^2 \|\mathbf{G}_t\|_F^2)^{1/2}}) + 2\eta \|\tilde{\mathbf{G}}_t\|_F \\
 &\leq \sqrt{r} \eta \|\mathbf{G}_t\|_F + 2\eta \|\tilde{\mathbf{G}}_t\|_F,
 \end{aligned}$$

where (a) is by $\sqrt{1+x} \leq 1 + \sqrt{x}$ for any $x \geq 0$. Since $\|\mathbf{G}_t\|_F \leq 2\|\tilde{\mathbf{G}}_t\|_F$, we arrive at

$$\begin{aligned}
 \|\mathbf{X}_{t+1} - \mathbf{X}_t\|_F &\leq 2\eta(\sqrt{r} + 1) \|\tilde{\mathbf{G}}_t\|_F \\
 &= 2\eta(\sqrt{r} + 1) \|(\mathbf{X}_t \Theta_t \mathbf{X}_t^\top - \mathbf{A}) \mathbf{X}_t \Theta_t + \Xi_t \mathbf{X}_t \Theta_t\|_F \\
 &\stackrel{(b)}{\leq} 2\eta(\sqrt{r} + 1) \|\Theta_t\| \|\mathbf{X}_t\| \|(\mathbf{X}_t \Theta_t \mathbf{X}_t^\top - \mathbf{A}) + \Xi_t\|_F \\
 &\stackrel{(c)}{\leq} 4\eta(\sqrt{r} + 1) (\|\Xi_t\|_F + \|\mathbf{X}_t \Theta_t \mathbf{X}_t^\top - \mathbf{A}\|_F) \\
 &\leq 4\eta(\sqrt{r} + 1) (\sqrt{m} \|(\mathcal{M}^* \mathcal{M} - \mathcal{I})(\mathbf{X}_t \Theta_t \mathbf{X}_t^\top - \mathbf{A})\| + \|\mathbf{X}_t \Theta_t \mathbf{X}_t^\top - \mathbf{A}\|_F) \\
 &\stackrel{(d)}{\leq} 4\eta(\sqrt{r} + 1) (\sqrt{m} \delta + 1) \|\mathbf{X}_t \Theta_t \mathbf{X}_t^\top - \mathbf{A}\|_F,
 \end{aligned}$$

where (b) is from $\|\mathbf{AB}\|_F \leq \|\mathbf{A}\| \|\mathbf{B}\|_F$; (c) is due to $\|\Theta_t\| \leq 2$, $\|\mathbf{X}_t\| \leq 1$; and (d) follows from Lemma 24 and $\text{rank}(\mathbf{X}_t \Theta_t \mathbf{X}_t^\top - \mathbf{A}) \leq \text{rank}(\mathbf{X}_t \Theta_t \mathbf{X}_t^\top) + \text{rank}(\mathbf{A}) \leq r + r_A$.

Finally, we turn to estimating $\|\Delta_t\|_F$ and $\|\Xi_t\|_F$.

$$\begin{aligned}
 \|\Delta_t\|_F &= \|(\mathcal{M}^* \mathcal{M} - \mathcal{I})(\mathbf{X}_{t+1} \Theta_t \mathbf{X}_{t+1}^\top - \mathbf{A})\|_F \\
 &\leq \sqrt{m} \|(\mathcal{M}^* \mathcal{M} - \mathcal{I})(\mathbf{X}_{t+1} \Theta_t \mathbf{X}_{t+1}^\top - \mathbf{A})\| \\
 &\stackrel{(e)}{\leq} \sqrt{m} \delta \|\mathbf{X}_{t+1} \Theta_t \mathbf{X}_{t+1}^\top - \mathbf{A}\|_F \\
 &\leq \sqrt{m} \delta (\|\mathbf{X}_t \Theta_t \mathbf{X}_t^\top - \mathbf{A}\|_F + \|\mathbf{X}_{t+1} \Theta_t (\mathbf{X}_{t+1}^\top - \mathbf{X}_t^\top)\|_F + \|(\mathbf{X}_{t+1} - \mathbf{X}_t) \Theta_t \mathbf{X}_t^\top\|_F) \\
 &\stackrel{(f)}{\leq} \sqrt{m} \delta (\|\mathbf{X}_t \Theta_t \mathbf{X}_t^\top - \mathbf{A}\|_F + 4\|\mathbf{X}_{t+1} - \mathbf{X}_t\|_F) \\
 &\leq \sqrt{m} \delta (1 + 16\eta(\sqrt{r} + 1)(\sqrt{m} \delta + 1)) \|\mathbf{X}_t \Theta_t \mathbf{X}_t^\top - \mathbf{A}\|_F \\
 &\stackrel{(g)}{\leq} \frac{\xi \sqrt{m}}{\sqrt{mr}} (1 + \frac{64}{300} \sqrt{r}) \|\mathbf{X}_t \Theta_t \mathbf{X}_t^\top - \mathbf{A}\|_F \\
 &\leq \xi \|\mathbf{X}_t \Theta_t \mathbf{X}_t^\top - \mathbf{A}\|_F, \\
 \|\Xi_t\|_F &= \|(\mathcal{M}^* \mathcal{M} - \mathcal{I})(\mathbf{X}_t \Theta_t \mathbf{X}_t^\top - \mathbf{A})\|_F \\
 &\leq \sqrt{m} \|(\mathcal{M}^* \mathcal{M} - \mathcal{I})(\mathbf{X}_t \Theta_t \mathbf{X}_t^\top - \mathbf{A})\| \\
 &\stackrel{(h)}{\leq} \sqrt{m} \delta \|\mathbf{X}_t \Theta_t \mathbf{X}_t^\top - \mathbf{A}\|_F \\
 &\stackrel{(g)}{\leq} \xi \|\mathbf{X}_t \Theta_t \mathbf{X}_t^\top - \mathbf{A}\|_F,
 \end{aligned}$$

where (e) is by Lemma 24 and $\text{rank}(\mathbf{X}_{t+1} \Theta_t \mathbf{X}_{t+1}^\top - \mathbf{A}) \leq \text{rank}(\mathbf{X}_{t+1} \Theta_t \mathbf{X}_{t+1}^\top) + \text{rank}(\mathbf{A}) \leq r + r_A$; (f) is from $\|\mathbf{X}_{t+1}\| \leq 1$ and $\|\Theta_t\| \leq 2$; (g) is due to $\eta \leq \frac{1}{300\kappa^2 r_A}$ and $\delta \leq \frac{\xi}{\sqrt{mr}}$; and (h) follows from Lemma 24 and $\text{rank}(\mathbf{X}_t \Theta_t \mathbf{X}_t^\top - \mathbf{A}) \leq \text{rank}(\mathbf{X}_t \Theta_t \mathbf{X}_t^\top) + \text{rank}(\mathbf{A}) \leq r + r_A$. ■

G.4. Proof of Lemma 10

Proof Noting that $\|\mathbf{X}_t\| \leq 1$, $\|\mathbf{A}\| \leq 1$, $\|\Theta_t\| \leq 2$, $\|\mathbf{I}_m - \mathbf{X}_t \mathbf{X}_t^\top\| \leq 1$, we obtain

$$\begin{aligned}
 \|\mathbf{H}_t\|_F &\leq \|(\mathbf{I}_m - \mathbf{X}_t \mathbf{X}_t^\top) \mathbf{A} \mathbf{X}_t \mathbf{X}_t^\top \Delta_{t-1} \mathbf{X}_t\|_F + \|(\mathbf{I}_m - \mathbf{X}_t \mathbf{X}_t^\top) \Xi_t \mathbf{X}_t \Theta_t\|_F \\
 &\quad + \frac{1}{2} (\|\mathbf{X}_t \mathbf{X}_t^\top \Xi_t \mathbf{X}_t \Theta_t\|_F + \|\mathbf{X}_t \Theta_t \mathbf{X}_t^\top \Xi_t \mathbf{X}_t\|_F) \\
 &\quad + \frac{1}{2} (\|\mathbf{X}_t \mathbf{X}_t^\top \mathbf{A} \mathbf{X}_t \mathbf{X}_t^\top \Delta_{t-1} \mathbf{X}_t\|_F + \|\mathbf{X}_t \mathbf{X}_t^\top \Delta_{t-1} \mathbf{X}_t \mathbf{X}_t^\top \mathbf{A} \mathbf{X}_t\|_F) \\
 &\leq 2\|\Delta_{t-1}\|_F + 4\|\Xi_t\|_F.
 \end{aligned} \tag{9}$$

In the same way, it follows that $\|\mathbf{H}_t\| \leq 2\|\Delta_{t-1}\| + 4\|\Xi_t\|$.

From the update of $\mathbf{X}_t$, we have $\mathbf{X}_{t+1} \mathbf{X}_{t+1}^\top = (\mathbf{X}_t - \eta \mathbf{G}_t)(\mathbf{I}_r + \eta^2 \mathbf{G}_t^\top \mathbf{G}_t)^{-1}(\mathbf{X}_t - \eta \mathbf{G}_t)^\top$. Pre-multiplying by $\mathbf{U}^\top$ and postmultiplying by $\mathbf{U}$, it follows that

$$\begin{aligned}
 &\Phi_{t+1} \Phi_{t+1}^\top \\
 &= (\Phi_t - \eta \mathbf{U}^\top \mathbf{G}_t)(\mathbf{I}_r + \eta^2 \mathbf{G}_t^\top \mathbf{G}_t)^{-1}(\Phi_t^\top - \eta \mathbf{G}_t^\top \mathbf{U}) \\
 &\stackrel{(a)}{=} \left([\mathbf{I}_{r_A} + \eta(\mathbf{I}_{r_A} - \Phi_t \Phi_t^\top) \Sigma \Phi_t \Phi_t^\top \Sigma] \Phi_t - \eta \mathbf{U}^\top \mathbf{H}_t \right) (\mathbf{I}_r + \eta^2 \mathbf{G}_t^\top \mathbf{G}_t)^{-1}
 \end{aligned}$$

$$\begin{aligned}
 & \left([\mathbf{I}_{r_A} + \eta(\mathbf{I}_{r_A} - \Phi_t \Phi_t^\top) \Sigma \Phi_t \Phi_t^\top \Sigma] \Phi_t - \eta \mathbf{U}^\top \mathbf{H}_t \right)^\top \\
 \stackrel{(b)}{\succeq} & \left([\mathbf{I}_{r_A} + \eta(\mathbf{I}_{r_A} - \Phi_t \Phi_t^\top) \Sigma \Phi_t \Phi_t^\top \Sigma] \Phi_t - \eta \mathbf{U}^\top \mathbf{H}_t \right) (\mathbf{I}_r - \eta^2 \mathbf{G}_t^\top \mathbf{G}_t) \\
 & \left([\mathbf{I}_{r_A} + \eta(\mathbf{I}_{r_A} - \Phi_t \Phi_t^\top) \Sigma \Phi_t \Phi_t^\top \Sigma] \Phi_t - \eta \mathbf{U}^\top \mathbf{H}_t \right)^\top,
 \end{aligned}$$

where (a) is from directly expanding $\mathbf{G}_t$; and (b) is by Lemma 14.

We next derive an upper bound for $\mathbf{G}_t^\top \mathbf{G}_t$. Substituting the expression of $\mathbf{G}_t$, we obtain

$$\begin{aligned}
 \mathbf{G}_t^\top \mathbf{G}_t &= \mathbf{X}_t^\top \mathbf{A} \mathbf{X}_t \mathbf{X}_t^\top \mathbf{A} (\mathbf{I}_m - \mathbf{X}_t \mathbf{X}_t^\top) \mathbf{A} \mathbf{X}_t \mathbf{X}_t^\top \mathbf{A} \mathbf{X}_t + \mathbf{H}_t^\top \mathbf{H}_t \\
 &\quad - \mathbf{X}_t^\top \mathbf{A} \mathbf{X}_t \mathbf{X}_t^\top \mathbf{A} (\mathbf{I}_m - \mathbf{X}_t \mathbf{X}_t^\top) \mathbf{H}_t \\
 &\quad - \mathbf{H}_t^\top (\mathbf{I}_m - \mathbf{X}_t \mathbf{X}_t^\top) \mathbf{A} \mathbf{X}_t \mathbf{X}_t^\top \mathbf{A} \mathbf{X}_t \\
 &\stackrel{(c)}{\leq} \sigma_1 (\mathbf{I}_{r_A} - \Phi_t \Phi_t^\top) \mathbf{I}_r + (\|\mathbf{H}_t\|^2 + 2\|(\mathbf{I}_m - \mathbf{X}_t \mathbf{X}_t^\top) \mathbf{U}\| \|\mathbf{H}_t\|) \mathbf{I}_r \\
 &\stackrel{(d)}{\leq} \sigma_1 (\mathbf{I}_{r_A} - \Phi_t \Phi_t^\top) \mathbf{I}_r + 16\chi_t \mathbf{I}_r,
 \end{aligned}$$

where (c) follows from Lemma 25, $\|\mathbf{X}_t\| \leq 1$ and $\|\mathbf{A}\| \leq 1$; and (d) is due to $\|(\mathbf{I}_m - \mathbf{X}_t \mathbf{X}_t^\top) \mathbf{U}\| \leq \|(\mathbf{I}_m - \mathbf{X}_t \mathbf{X}_t^\top) \mathbf{U}\|_F = \sqrt{\text{Tr}(\mathbf{I}_{r_A} - \Phi_t \Phi_t^\top)}$, and $\|\mathbf{H}_t\| \leq 2\|\Delta_{t-1}\| + 4\|\Xi_t\| \leq 6$.

Combining the lower bound on $\Phi_{t+1} \Phi_{t+1}^\top$, the upper bound on $\mathbf{G}_t^\top \mathbf{G}_t$ derived above, and the inequality $1 - \eta^2 \beta_t - 16\eta^2 \chi_t \geq 1 - \frac{1}{100}(1 + 96) > 0$, we derive

$$\begin{aligned}
 & \frac{1}{1 - \eta^2 \beta_t - 16\eta^2 \chi_t} \Phi_{t+1} \Phi_{t+1}^\top \\
 \succeq & \left([\mathbf{I}_{r_A} + \eta(\mathbf{I}_{r_A} - \Phi_t \Phi_t^\top) \Sigma \Phi_t \Phi_t^\top \Sigma] \Phi_t - \eta \mathbf{U}^\top \mathbf{H}_t \right) \\
 & \left([\mathbf{I}_{r_A} + \eta(\mathbf{I}_{r_A} - \Phi_t \Phi_t^\top) \Sigma \Phi_t \Phi_t^\top \Sigma] \Phi_t - \eta \mathbf{U}^\top \mathbf{H}_t \right)^\top.
 \end{aligned} \tag{10}$$

Let the compact SVD of Φ_t be $\mathbf{Q}_t \mathbf{\Lambda}_t \mathbf{P}_t^\top$, where $\mathbf{Q}_t \in \mathbb{R}^{r_A \times r_A}$, $\mathbf{\Lambda}_t \in \mathbb{R}^{r_A \times r_A}$, and $\mathbf{P}_t \in \mathbb{R}^{r \times r_A}$. Denote $\mathbf{S}_t := \mathbf{Q}_t^\top \Sigma \mathbf{Q}_t$. It is a positive definite matrix. This gives that

$$\begin{aligned}
 & \text{Tr} \left([\mathbf{I}_{r_A} + \eta(\mathbf{I}_{r_A} - \Phi_t \Phi_t^\top) \Sigma \Phi_t \Phi_t^\top \Sigma] \Phi_t \Phi_t^\top [\mathbf{I}_{r_A} + \eta(\mathbf{I}_{r_A} - \Phi_t \Phi_t^\top) \Sigma \Phi_t \Phi_t^\top \Sigma]^\top \right) \\
 &= \text{Tr} \left(\mathbf{Q}_t [\mathbf{I}_{r_A} + \eta(\mathbf{I}_{r_A} - \mathbf{\Lambda}_t^2) \mathbf{S}_t \mathbf{\Lambda}_t^2 \mathbf{S}_t] \mathbf{\Lambda}_t^2 [\mathbf{I}_{r_A} + \eta \mathbf{S}_t \mathbf{\Lambda}_t^2 \mathbf{S}_t (\mathbf{I}_{r_A} - \mathbf{\Lambda}_t^2)] \mathbf{Q}_t^\top \right) \\
 &\stackrel{(e)}{\geq} \text{Tr} \left(\mathbf{Q}_t [\mathbf{\Lambda}_t^2 + \eta(\mathbf{I}_{r_A} - \mathbf{\Lambda}_t^2) \mathbf{S}_t \mathbf{\Lambda}_t^2 \mathbf{S}_t \mathbf{\Lambda}_t^2 + \eta \mathbf{\Lambda}_t^2 \mathbf{S}_t \mathbf{\Lambda}_t^2 \mathbf{S}_t (\mathbf{I}_{r_A} - \mathbf{\Lambda}_t^2)] \mathbf{Q}_t^\top \right) \\
 &= \text{Tr}(\mathbf{Q}_t \mathbf{\Lambda}_t^2 \mathbf{Q}_t^\top) + \eta \text{Tr}((\mathbf{I}_{r_A} - \mathbf{\Lambda}_t^2) \mathbf{S}_t \mathbf{\Lambda}_t^2 \mathbf{S}_t \mathbf{\Lambda}_t^2 + \mathbf{\Lambda}_t^2 \mathbf{S}_t \mathbf{\Lambda}_t^2 \mathbf{S}_t (\mathbf{I}_{r_A} - \mathbf{\Lambda}_t^2)) \\
 &\stackrel{(f)}{\geq} \text{Tr}(\mathbf{Q}_t \mathbf{\Lambda}_t^2 \mathbf{Q}_t^\top) + \frac{2\eta \sigma_{r_A}(\mathbf{\Lambda}_t^2)}{\kappa^2} \text{Tr}((\mathbf{I}_{r_A} - \mathbf{\Lambda}_t^2) \mathbf{\Lambda}_t^2) \\
 &= \text{Tr}(\Phi_t \Phi_t^\top) + \frac{2\eta \sigma_{r_A}(\mathbf{\Lambda}_t^2)}{\kappa^2} \text{Tr}((\mathbf{I}_{r_A} - \Phi_t \Phi_t^\top) \Phi_t \Phi_t^\top),
 \end{aligned}$$

where (e) follows from the fact that $\eta^2 \mathbf{Q}_t(\mathbf{I}_{r_A} - \Lambda_t^2) \mathbf{S}_t \Lambda_t^2 \mathbf{S}_t \Lambda_t^2 \mathbf{S}_t \Lambda_t^2 \mathbf{S}_t (\mathbf{I}_{r_A} - \Lambda_t^2) \mathbf{Q}_t^\top$ is PSD; and (f) is by Lemma 16 and Lemma 17. More precisely, we use $\sigma_{r_A}(\mathbf{S}_t \Lambda_t^2 \mathbf{S}_t) \geq \sigma_{r_A}^2(\mathbf{S}_t) \sigma_{r_A}(\Lambda_t^2) = \sigma_{r_A}(\Lambda_t^2)/\kappa^2$.

Taking trace on both sides of (10), we arrive at

$$\begin{aligned}
 & \frac{1}{1 - \eta^2 \beta_t - 16\eta^2 \chi_t} \text{Tr}(\Phi_{t+1} \Phi_{t+1}^\top) \\
 & \geq \text{Tr}(\Phi_t \Phi_t^\top) + \frac{2\eta \sigma_{r_A}(\Lambda_t^2)}{\kappa^2} \text{Tr}((\mathbf{I}_{r_A} - \Phi_t \Phi_t^\top) \Phi_t \Phi_t^\top) \\
 & \quad - 2\eta \text{Tr}\left(\left[\mathbf{I}_{r_A} + \eta(\mathbf{I}_{r_A} - \Phi_t \Phi_t^\top) \Sigma \Phi_t \Phi_t^\top \Sigma\right] \Phi_t \mathbf{H}_t^\top \mathbf{U}\right) \\
 & \quad + \eta^2 \text{Tr}(\mathbf{U}^\top \mathbf{H}_t \mathbf{H}_t^\top \mathbf{U}) \\
 & \geq \text{Tr}(\Phi_t \Phi_t^\top) + \frac{2\eta \sigma_{r_A}(\Lambda_t^2)}{\kappa^2} \text{Tr}((\mathbf{I}_{r_A} - \Phi_t \Phi_t^\top) \Phi_t \Phi_t^\top) \\
 & \quad - 2\eta \text{Tr}\left(\Phi_t \mathbf{H}_t^\top \mathbf{U} + \eta(\mathbf{I}_{r_A} - \Phi_t \Phi_t^\top) \Sigma \Phi_t \Phi_t^\top \Sigma \Phi_t \mathbf{H}_t^\top \mathbf{U}\right) \\
 & \stackrel{(g)}{=} \text{Tr}(\Phi_t \Phi_t^\top) + \frac{2\eta \sigma_{r_A}(\Lambda_t^2)}{\kappa^2} \text{Tr}((\mathbf{I}_{r_A} - \Phi_t \Phi_t^\top) \Phi_t \Phi_t^\top) \\
 & \quad - 2\eta \text{Tr}(\mathbf{U}^\top (\mathbf{I}_m - \mathbf{X}_t \mathbf{X}_t^\top) \mathbf{A} \mathbf{X}_t \mathbf{X}_t^\top \Delta_{t-1} \mathbf{X}_t \Phi_t^\top) \\
 & \quad - 2\eta \text{Tr}(\mathbf{U}^\top (\mathbf{I}_m - \mathbf{X}_t \mathbf{X}_t^\top) \Xi_t \mathbf{X}_t \Theta_t \Phi_t^\top) \\
 & \quad - \eta \text{Tr}(\Phi_t \mathbf{X}_t^\top \Xi_t \mathbf{X}_t \Theta_t \Phi_t^\top - \Phi_t \Theta_t \mathbf{X}_t^\top \Xi_t \mathbf{X}_t \Phi_t^\top) \\
 & \quad - \eta \text{Tr}(\Phi_t \mathbf{X}_t^\top \mathbf{A} \mathbf{X}_t \mathbf{X}_t^\top \Delta_{t-1} \mathbf{X}_t \Phi_t^\top - \Phi_t \mathbf{X}_t^\top \Delta_{t-1} \mathbf{X}_t \mathbf{X}_t^\top \mathbf{A} \mathbf{X}_t \Phi_t^\top) \\
 & \quad - 2\eta^2 \text{Tr}((\mathbf{I}_{r_A} - \Phi_t \Phi_t^\top) \Sigma \Phi_t \Phi_t^\top \Sigma \Phi_t \mathbf{H}_t^\top \mathbf{U}) \\
 & \stackrel{(h)}{=} \text{Tr}(\Phi_t \Phi_t^\top) + \frac{2\eta \sigma_{r_A}(\Lambda_t^2)}{\kappa^2} \text{Tr}((\mathbf{I}_{r_A} - \Phi_t \Phi_t^\top) \Phi_t \Phi_t^\top) \\
 & \quad - 2\eta \text{Tr}(\mathbf{U}^\top (\mathbf{I}_m - \mathbf{X}_t \mathbf{X}_t^\top) \mathbf{U} \Sigma \mathbf{U}^\top \mathbf{X}_t \mathbf{X}_t^\top \Delta_{t-1} \mathbf{X}_t \Phi_t^\top) \\
 & \quad - 2\eta \text{Tr}(\mathbf{U}^\top (\mathbf{I}_m - \mathbf{X}_t \mathbf{X}_t^\top) \Xi_t \mathbf{X}_t \Theta_t \Phi_t^\top) \\
 & \quad - 2\eta^2 \text{Tr}((\mathbf{I}_{r_A} - \Phi_t \Phi_t^\top) \Sigma \Phi_t \Phi_t^\top \Sigma \Phi_t \mathbf{H}_t^\top \mathbf{U})
 \end{aligned} \tag{11}$$

where (g) is by substituting $\mathbf{H}_t$ in; and (h) arises from $\text{Tr}(\mathbf{M}) = \text{Tr}(\mathbf{M}^\top)$ for any $\mathbf{M} \in \mathbb{R}^{r_A \times r_A}$. By the Cauchy–Schwarz inequality, we can upper bound the three trace terms as follows.

For the first term, we have that

$$\begin{aligned}
 & \text{Tr}(\mathbf{U}^\top (\mathbf{I}_m - \mathbf{X}_t \mathbf{X}_t^\top) \mathbf{U} \Sigma \mathbf{U}^\top \mathbf{X}_t \mathbf{X}_t^\top \Delta_{t-1} \mathbf{X}_t \Phi_t^\top) \\
 & \leq \|\mathbf{U}^\top (\mathbf{I}_m - \mathbf{X}_t \mathbf{X}_t^\top) \mathbf{U} \Sigma \mathbf{U}^\top\|_F \|\mathbf{X}_t \mathbf{X}_t^\top \Delta_{t-1} \mathbf{X}_t \Phi_t^\top\|_F \\
 & \stackrel{(i)}{\leq} \|(\mathbf{I}_m - \mathbf{X}_t \mathbf{X}_t^\top) \mathbf{U}\|_F \|\Delta_{t-1}\|_F \\
 & = \sqrt{\text{Tr}(\mathbf{I}_{r_A} - \Phi_t \Phi_t^\top)} \|\Delta_{t-1}\|_F.
 \end{aligned} \tag{12}$$

For the second term, we can obtain that

$$\text{Tr}(\mathbf{U}^\top (\mathbf{I}_m - \mathbf{X}_t \mathbf{X}_t^\top) \Xi_t \mathbf{X}_t \Theta_t \Phi_t^\top) \tag{13}$$

$$\begin{aligned}
 &\leq \|\mathbf{U}^\top (\mathbf{I}_m - \mathbf{X}_t \mathbf{X}_t^\top)\|_F \|\boldsymbol{\Xi}_t \mathbf{X}_t \boldsymbol{\Theta}_t \boldsymbol{\Phi}_t^\top\|_F \\
 &\stackrel{(i)}{\leq} 2 \|\mathbf{U}^\top (\mathbf{I}_m - \mathbf{X}_t \mathbf{X}_t^\top)\|_F \|\boldsymbol{\Xi}_t\|_F \\
 &= 2 \sqrt{\text{Tr}(\mathbf{I}_{r_A} - \boldsymbol{\Phi}_t \boldsymbol{\Phi}_t^\top)} \|\boldsymbol{\Xi}_t\|_F.
 \end{aligned}$$

For the third term, it holds that

$$\begin{aligned}
 &\text{Tr}((\mathbf{I}_{r_A} - \boldsymbol{\Phi}_t \boldsymbol{\Phi}_t^\top) \boldsymbol{\Sigma} \boldsymbol{\Phi}_t \boldsymbol{\Phi}_t^\top \boldsymbol{\Sigma} \boldsymbol{\Phi}_t \mathbf{H}_t^\top \mathbf{U}) \\
 &\leq \|\mathbf{I}_{r_A} - \boldsymbol{\Phi}_t \boldsymbol{\Phi}_t^\top\|_F \|\boldsymbol{\Sigma} \boldsymbol{\Phi}_t \boldsymbol{\Phi}_t^\top \boldsymbol{\Sigma} \boldsymbol{\Phi}_t \mathbf{H}_t^\top \mathbf{U}\|_F \\
 &= \|\mathbf{U}^\top (\mathbf{I}_m - \mathbf{X}_t \mathbf{X}_t^\top) \mathbf{U}\|_F \|\boldsymbol{\Sigma} \boldsymbol{\Phi}_t \boldsymbol{\Phi}_t^\top \boldsymbol{\Sigma} \boldsymbol{\Phi}_t \mathbf{H}_t^\top \mathbf{U}\|_F \\
 &\stackrel{(i)}{\leq} \|(\mathbf{I}_m - \mathbf{X}_t \mathbf{X}_t^\top) \mathbf{U}\|_F \|\mathbf{H}_t\|_F \\
 &= \sqrt{\text{Tr}(\mathbf{I}_{r_A} - \boldsymbol{\Phi}_t \boldsymbol{\Phi}_t^\top)} \|\mathbf{H}_t\|_F.
 \end{aligned} \tag{14}$$

Here (i) is from $\|\mathbf{U}\| \leq 1$, $\|\boldsymbol{\Sigma}\| \leq 1$, $\|\mathbf{X}_t\| \leq 1$, $\|\boldsymbol{\Phi}_t\| \leq 1$, and $\|\boldsymbol{\Theta}_t\| \leq 2$. Combining inequalities (11), (12), (13), and (14), it follows that

$$\begin{aligned}
 \frac{1}{1 - \eta^2 \beta_t - 16\eta^2 \chi_t} \text{Tr}(\boldsymbol{\Phi}_{t+1} \boldsymbol{\Phi}_{t+1}^\top) &\geq \text{Tr}(\boldsymbol{\Phi}_t \boldsymbol{\Phi}_t^\top) + \frac{2\eta \sigma_{r_A}(\Lambda_t^2)}{\kappa^2} \text{Tr}((\mathbf{I}_{r_A} - \boldsymbol{\Phi}_t \boldsymbol{\Phi}_t^\top) \boldsymbol{\Phi}_t \boldsymbol{\Phi}_t^\top) \\
 &\quad - 2\eta \sqrt{\text{Tr}(\mathbf{I}_{r_A} - \boldsymbol{\Phi}_t \boldsymbol{\Phi}_t^\top)} (\|\boldsymbol{\Delta}_{t-1}\|_F + 2\|\boldsymbol{\Xi}_t\|_F) \\
 &\quad - 2\eta^2 \sqrt{\text{Tr}(\mathbf{I}_{r_A} - \boldsymbol{\Phi}_t \boldsymbol{\Phi}_t^\top)} \|\mathbf{H}_t\|_F.
 \end{aligned}$$

Reorganizing the terms, we arrive at

$$\begin{aligned}
 &\text{Tr}(\mathbf{I}_{r_A} - \boldsymbol{\Phi}_{t+1} \boldsymbol{\Phi}_{t+1}^\top) - \text{Tr}(\mathbf{I}_{r_A} - \boldsymbol{\Phi}_t \boldsymbol{\Phi}_t^\top) \\
 &\leq \eta^2 (\beta_t + 16\chi_t) \text{Tr}(\boldsymbol{\Phi}_t \boldsymbol{\Phi}_t^\top) - \frac{2\eta(1 - \eta^2 \beta_t - 16\eta^2 \chi_t) \sigma_{r_A}(\Lambda_t^2)}{\kappa^2} \text{Tr}((\mathbf{I}_{r_A} - \boldsymbol{\Phi}_t \boldsymbol{\Phi}_t^\top) \boldsymbol{\Phi}_t \boldsymbol{\Phi}_t^\top) \\
 &\quad + (2\eta - 2\eta^3 \beta_t - 32\eta^3 \chi_t) \sqrt{\text{Tr}(\mathbf{I}_{r_A} - \boldsymbol{\Phi}_t \boldsymbol{\Phi}_t^\top)} (\|\boldsymbol{\Delta}_{t-1}\|_F + 2\|\boldsymbol{\Xi}_t\|_F) \\
 &\quad + (2\eta^2 - 2\eta^4 \beta_t - 32\eta^4 \chi_t) \sqrt{\text{Tr}(\mathbf{I}_{r_A} - \boldsymbol{\Phi}_t \boldsymbol{\Phi}_t^\top)} \|\mathbf{H}_t\|_F \\
 &\leq \eta^2 (\beta_t + 16\chi_t) \text{Tr}(\boldsymbol{\Phi}_t \boldsymbol{\Phi}_t^\top) - \frac{2\eta(1 - \eta^2 \beta_t - 16\eta^2 \chi_t) \sigma_{r_A}(\Lambda_t^2)}{\kappa^2} \text{Tr}((\mathbf{I}_{r_A} - \boldsymbol{\Phi}_t \boldsymbol{\Phi}_t^\top) \boldsymbol{\Phi}_t \boldsymbol{\Phi}_t^\top) \\
 &\quad + 2\eta \sqrt{\text{Tr}(\mathbf{I}_{r_A} - \boldsymbol{\Phi}_t \boldsymbol{\Phi}_t^\top)} (\|\boldsymbol{\Delta}_{t-1}\|_F + 2\|\boldsymbol{\Xi}_t\|_F) \\
 &\quad + 2\eta^2 \sqrt{\text{Tr}(\mathbf{I}_{r_A} - \boldsymbol{\Phi}_t \boldsymbol{\Phi}_t^\top)} \|\mathbf{H}_t\|_F.
 \end{aligned}$$

Together with $\sigma_{r_A}(\Lambda_t^2) = \sigma_{r_A}(\mathbf{Q}_t^\top \boldsymbol{\Phi}_t \boldsymbol{\Phi}_t^\top \mathbf{Q}_t) = \sigma_{r_A}^2(\boldsymbol{\Phi}_t)$, we conclude the proof. $\blacksquare$

G.5. Proof of Lemma 11

Proof From the update formula of $\boldsymbol{\Theta}_t$, we obtain

$$\|\boldsymbol{\Theta}_{t+1}\| = \|\mathbf{X}_{t+1}^\top \mathbf{A} \mathbf{X}_{t+1} - \mathbf{X}_{t+1}^\top \boldsymbol{\Delta}_t \mathbf{X}_{t+1}\|$$

$$\begin{aligned}
 &\leq \|\mathbf{X}_{t+1}^\top \mathbf{A} \mathbf{X}_{t+1}\| + \|\mathbf{X}_{t+1}^\top \boldsymbol{\Delta}_t \mathbf{X}_{t+1}\| \\
 &\stackrel{(a)}{\leq} 1 + \|\boldsymbol{\Delta}_t\| \\
 &= 1 + \|(\mathcal{M}^* \mathcal{M} - \mathcal{I})(\mathbf{X}_{t+1} \boldsymbol{\Theta}_t \mathbf{X}_{t+1}^\top - \mathbf{A})\| \\
 &\stackrel{(b)}{\leq} 1 + \frac{1}{3\sqrt{m}} \|\mathbf{X}_{t+1} \boldsymbol{\Theta}_t \mathbf{X}_{t+1}^\top - \mathbf{A}\|_F \\
 &\leq 1 + \frac{\sqrt{m}}{3\sqrt{m}} \|\mathbf{X}_{t+1} \boldsymbol{\Theta}_t \mathbf{X}_{t+1}^\top - \mathbf{A}\| \\
 &\leq 1 + \frac{1}{3} (\|\boldsymbol{\Theta}_t\| + \|\mathbf{A}\|) \\
 &\leq 2,
 \end{aligned}$$

where (a) is by $\|\mathbf{X}_t\|, \|\mathbf{A}\| \leq 1$; and (b) follows from Lemma 24 and $\text{rank}(\mathbf{X}_{t+1} \boldsymbol{\Theta}_t \mathbf{X}_{t+1}^\top - \mathbf{A}) \leq \text{rank}(\mathbf{X}_{t+1} \boldsymbol{\Theta}_t \mathbf{X}_{t+1}^\top) + \text{rank}(\mathbf{A}) \leq r + r_A$. $\blacksquare$

G.6. Proof of Lemma 12

Proof Let $\mathbf{L}_t := \mathbf{X}_t^\top \mathbf{A} \mathbf{X}_t \mathbf{X}_t^\top \boldsymbol{\Delta}_{t-1} \mathbf{X}_t + \mathbf{X}_t^\top \boldsymbol{\Xi}_t \mathbf{X}_t \boldsymbol{\Theta}_t + \frac{1}{2}(\boldsymbol{\Theta}_t \mathbf{X}_t^\top \boldsymbol{\Xi}_t \mathbf{X}_t - \mathbf{X}_t^\top \boldsymbol{\Xi}_t \mathbf{X}_t \boldsymbol{\Theta}_t) + \frac{1}{2}(\mathbf{X}_t^\top \boldsymbol{\Delta}_{t-1} \mathbf{X}_t \mathbf{X}_t^\top \mathbf{A} \mathbf{X}_t - \mathbf{X}_t^\top \mathbf{A} \mathbf{X}_t \mathbf{X}_t^\top \boldsymbol{\Delta}_{t-1} \mathbf{X}_t)$.

Applying the triangular inequality, we obtain

$$\begin{aligned}
 \|\mathbf{L}_t\| &\leq \|\mathbf{X}_t^\top \mathbf{A} \mathbf{X}_t \mathbf{X}_t^\top \boldsymbol{\Delta}_{t-1} \mathbf{X}_t\| + \|\mathbf{X}_t^\top \boldsymbol{\Xi}_t \mathbf{X}_t \boldsymbol{\Theta}_t\| + \frac{1}{2}(\|\boldsymbol{\Theta}_t \mathbf{X}_t^\top \boldsymbol{\Xi}_t \mathbf{X}_t\| + \|\mathbf{X}_t^\top \boldsymbol{\Xi}_t \mathbf{X}_t \boldsymbol{\Theta}_t\|) \\
 &\quad + \frac{1}{2}(\|\mathbf{X}_t^\top \boldsymbol{\Delta}_{t-1} \mathbf{X}_t \mathbf{X}_t^\top \mathbf{A} \mathbf{X}_t\| + \|\mathbf{X}_t^\top \mathbf{A} \mathbf{X}_t \mathbf{X}_t^\top \boldsymbol{\Delta}_{t-1} \mathbf{X}_t\|) \\
 &\stackrel{(a)}{\leq} 2\|\boldsymbol{\Delta}_{t-1}\| + 4\|\boldsymbol{\Xi}_t\|,
 \end{aligned} \tag{15}$$

where (a) is from $\|\mathbf{X}_t\| \leq 1, \|\mathbf{A}\| \leq 1$ and $\|\boldsymbol{\Theta}_t\| \leq 2$. Multiplying the update formula (4) on the left by $\mathbf{U}_\perp^\top$, we have that

$$\begin{aligned}
 \boldsymbol{\Psi}_{t+1} &= \mathbf{U}_\perp^\top (\mathbf{X}_t - \eta \mathbf{G}_t) (\mathbf{I}_r + \eta^2 \mathbf{G}_t^\top \mathbf{G}_t)^{-1/2} \\
 &\stackrel{(b)}{=} \left(\boldsymbol{\Psi}_t - \eta \boldsymbol{\Psi}_t \mathbf{X}_t^\top \mathbf{A} \mathbf{X}_t \mathbf{X}_t^\top \mathbf{A} \mathbf{X}_t + \eta \boldsymbol{\Psi}_t \mathbf{L}_t - \eta \mathbf{U}_\perp^\top \boldsymbol{\Xi}_t \mathbf{X}_t \boldsymbol{\Theta}_t \right) (\mathbf{I}_r + \eta^2 \mathbf{G}_t^\top \mathbf{G}_t)^{-1/2} \\
 &= \left(\boldsymbol{\Psi}_t \left(\mathbf{I}_r - \eta \mathbf{X}_t^\top \mathbf{A} \mathbf{X}_t \mathbf{X}_t^\top \mathbf{A} \mathbf{X}_t + \eta \mathbf{L}_t \right) - \eta \mathbf{U}_\perp^\top \boldsymbol{\Xi}_t \mathbf{X}_t \boldsymbol{\Theta}_t \right) (\mathbf{I}_r + \eta^2 \mathbf{G}_t^\top \mathbf{G}_t)^{-1/2},
 \end{aligned}$$

where (b) is by expanding $\mathbf{G}_t$ directly. Consequently, we have the following upper bound for $\boldsymbol{\Psi}_{t+1} \boldsymbol{\Psi}_{t+1}^\top$:

$$\begin{aligned}
 \boldsymbol{\Psi}_{t+1} \boldsymbol{\Psi}_{t+1}^\top &= \left(\boldsymbol{\Psi}_t \left(\mathbf{I}_r - \eta \mathbf{X}_t^\top \mathbf{A} \mathbf{X}_t \mathbf{X}_t^\top \mathbf{A} \mathbf{X}_t + \eta \mathbf{L}_t \right) - \eta \mathbf{U}_\perp^\top \boldsymbol{\Xi}_t \mathbf{X}_t \boldsymbol{\Theta}_t \right) (\mathbf{I}_r + \eta^2 \mathbf{G}_t^\top \mathbf{G}_t)^{-1} \\
 &\quad \left(\boldsymbol{\Psi}_t \left(\mathbf{I}_r - \eta \mathbf{X}_t^\top \mathbf{A} \mathbf{X}_t \mathbf{X}_t^\top \mathbf{A} \mathbf{X}_t + \eta \mathbf{L}_t \right) - \eta \mathbf{U}_\perp^\top \boldsymbol{\Xi}_t \mathbf{X}_t \boldsymbol{\Theta}_t \right)^\top \\
 &\stackrel{(c)}{\preceq} \left(\boldsymbol{\Psi}_t \left(\mathbf{I}_r - \eta \mathbf{X}_t^\top \mathbf{A} \mathbf{X}_t \mathbf{X}_t^\top \mathbf{A} \mathbf{X}_t + \eta \mathbf{L}_t \right) - \eta \mathbf{U}_\perp^\top \boldsymbol{\Xi}_t \mathbf{X}_t \boldsymbol{\Theta}_t \right)
 \end{aligned}$$

$$\begin{aligned}
 & \left(\Psi_t \left(\mathbf{I}_r - \eta \mathbf{X}_t^\top \mathbf{A} \mathbf{X}_t \mathbf{X}_t^\top \mathbf{A} \mathbf{X}_t + \eta \mathbf{L}_t \right) - \eta \mathbf{U}_\perp^\top \Xi_t \mathbf{X}_t \Theta_t \right)^\top \\
 = & \Psi_t \left(\mathbf{I}_r - \eta \mathbf{X}_t^\top \mathbf{A} \mathbf{X}_t \mathbf{X}_t^\top \mathbf{A} \mathbf{X}_t + \eta \mathbf{L}_t \right) \\
 & \left(\mathbf{I}_r - \eta \mathbf{X}_t^\top \mathbf{A} \mathbf{X}_t \mathbf{X}_t^\top \mathbf{A} \mathbf{X}_t + \eta \mathbf{L}_t \right)^\top \Psi_t^\top \\
 & - \eta \mathbf{U}_\perp^\top \Xi_t \mathbf{X}_t \Theta_t \left(\mathbf{I}_r - \eta \mathbf{X}_t^\top \mathbf{A} \mathbf{X}_t \mathbf{X}_t^\top \mathbf{A} \mathbf{X}_t + \eta \mathbf{L}_t \right)^\top \Psi_t^\top \\
 & - \eta \Psi_t \left(\mathbf{I}_r - \eta \mathbf{X}_t^\top \mathbf{A} \mathbf{X}_t \mathbf{X}_t^\top \mathbf{A} \mathbf{X}_t + \eta \mathbf{L}_t \right) \Theta_t \mathbf{X}_t^\top \Xi_t \mathbf{U}_\perp \\
 & + \eta^2 \mathbf{U}_\perp^\top \Xi_t \mathbf{X}_t \Theta_t^2 \mathbf{X}_t^\top \Xi_t \mathbf{U}_\perp,
 \end{aligned} \tag{16}$$

where (c) is from that $(\mathbf{I}_r + \eta^2 \mathbf{G}_t^\top \mathbf{G}_t)^{-1}$ is PSD and all of its eigenvalues are smaller than 1. Since $\mathbf{Y} \mathbf{Y}^\top \preceq \|\mathbf{Y}\|^2 \mathbf{I}_r$ holds for any symmetric matrix $\mathbf{Y} \in \mathbb{R}^{r \times r}$ and by Lemma 25, we can upper bound the three terms as follows.

For the first term, we can obtain that

$$\begin{aligned}
 & \Psi_t \left(\mathbf{I}_r - \eta \mathbf{X}_t^\top \mathbf{A} \mathbf{X}_t \mathbf{X}_t^\top \mathbf{A} \mathbf{X}_t + \eta \mathbf{L}_t \right) \\
 & \left(\mathbf{I}_r - \eta \mathbf{X}_t^\top \mathbf{A} \mathbf{X}_t \mathbf{X}_t^\top \mathbf{A} \mathbf{X}_t + \eta \mathbf{L}_t \right)^\top \Psi_t^\top \\
 & \preceq \|\mathbf{I}_r - \eta \mathbf{X}_t^\top \mathbf{A} \mathbf{X}_t \mathbf{X}_t^\top \mathbf{A} \mathbf{X}_t + \eta \mathbf{L}_t\|^2 \Psi_t \Psi_t^\top.
 \end{aligned} \tag{17}$$

For the second term, it holds that

$$\begin{aligned}
 & \mathbf{U}_\perp^\top \Xi_t \mathbf{X}_t \Theta_t \left(\mathbf{I}_r - \eta \mathbf{X}_t^\top \mathbf{A} \mathbf{X}_t \mathbf{X}_t^\top \mathbf{A} \mathbf{X}_t + \eta \mathbf{L}_t \right)^\top \Psi_t^\top \\
 & + \Psi_t \left(\mathbf{I}_r - \eta \mathbf{X}_t^\top \mathbf{A} \mathbf{X}_t \mathbf{X}_t^\top \mathbf{A} \mathbf{X}_t + \eta \mathbf{L}_t \right) \Theta_t \mathbf{X}_t^\top \Xi_t \mathbf{U}_\perp \\
 & \preceq 2 \|\Psi_t \left(\mathbf{I}_r - \eta \mathbf{X}_t^\top \mathbf{A} \mathbf{X}_t \mathbf{X}_t^\top \mathbf{A} \mathbf{X}_t + \eta \mathbf{L}_t \right)\| \|\Theta_t \mathbf{X}_t^\top \Xi_t \mathbf{U}_\perp\| \mathbf{I}_{m-r_A}.
 \end{aligned} \tag{18}$$

For the third term, we have that

$$\mathbf{U}_\perp^\top \Xi_t \mathbf{X}_t \Theta_t^2 \mathbf{X}_t^\top \Xi_t \mathbf{U}_\perp \preceq \|\mathbf{U}_\perp^\top \Xi_t \mathbf{X}_t \Theta_t^2 \mathbf{X}_t^\top \Xi_t \mathbf{U}_\perp\| \mathbf{I}_{m-r_A}. \tag{19}$$

Combining inequalities (16), (17), (18) and (19), it follows that

$$\begin{aligned}
 \Psi_{t+1} \Psi_{t+1}^\top & \preceq \|\mathbf{I}_r - \eta \mathbf{X}_t^\top \mathbf{A} \mathbf{X}_t \mathbf{X}_t^\top \mathbf{A} \mathbf{X}_t + \eta \mathbf{L}_t\|^2 \Psi_t \Psi_t^\top \\
 & + 2\eta \|\Psi_t \left(\mathbf{I}_r - \eta \mathbf{X}_t^\top \mathbf{A} \mathbf{X}_t \mathbf{X}_t^\top \mathbf{A} \mathbf{X}_t + \eta \mathbf{L}_t \right)\| \|\Theta_t \mathbf{X}_t^\top \Xi_t \mathbf{U}_\perp\| \mathbf{I}_{m-r_A} \\
 & + \eta^2 \|\mathbf{U}_\perp^\top \Xi_t \mathbf{X}_t \Theta_t^2 \mathbf{X}_t^\top \Xi_t \mathbf{U}_\perp\| \mathbf{I}_{m-r_A} \\
 & \stackrel{(d)}{\preceq} (\|\mathbf{I}_r - \eta \mathbf{X}_t^\top \mathbf{A} \mathbf{X}_t \mathbf{X}_t^\top \mathbf{A} \mathbf{X}_t\| + \eta \|\mathbf{L}_t\|)^2 \Psi_t \Psi_t^\top \\
 & + 2\eta \|\Psi_t\| (\|\mathbf{I}_r - \eta \mathbf{X}_t^\top \mathbf{A} \mathbf{X}_t \mathbf{X}_t^\top \mathbf{A} \mathbf{X}_t\| + \eta \|\mathbf{L}_t\|) \|\Theta_t \mathbf{X}_t^\top \Xi_t \mathbf{U}_\perp\| \mathbf{I}_{m-r_A} \\
 & + \eta^2 \|\mathbf{U}_\perp^\top \Xi_t \mathbf{X}_t \Theta_t^2 \mathbf{X}_t^\top \Xi_t \mathbf{U}_\perp\| \mathbf{I}_{m-r_A} \\
 & \stackrel{(e)}{\preceq} (1 + \eta \|\mathbf{L}_t\|)^2 \Psi_t \Psi_t^\top + 4\eta (1 + \eta \|\mathbf{L}_t\|) \|\Xi_t\| \mathbf{I}_{m-r_A} + 4\eta^2 \|\Xi_t\|^2 \mathbf{I}_{m-r_A},
 \end{aligned}$$

where (d) is by triangular inequality; and (e) is from that all the eigenvalues of the PSD matrix $\mathbf{X}_t^\top \mathbf{A} \mathbf{X}_t \mathbf{X}_t^\top \mathbf{A} \mathbf{X}_t$ are smaller than 1, along with $\|\Psi_t\| \leq 1$, $\|\mathbf{X}_t\| \leq 1$, $\|\mathbf{U}_\perp\| \leq 1$ and $\|\Theta_t\| \leq 2$.

From (15), we obtain $\|\mathbf{L}_t\| \leq 2\|\Delta_{t-1}\| + 4\|\Xi_t\|$. Then, we can further simplify the inequality as

$$\begin{aligned} \Psi_{t+1} \Psi_{t+1}^\top &\preceq (1 + 2\eta(\|\Delta_{t-1}\| + 2\|\Xi_t\|))^2 \Psi_t \Psi_t^\top \\ &\quad + (4\eta\|\Xi_t\| + 4\eta^2(5\|\Xi_t\|^2 + 2\|\Delta_{t-1}\|\|\Xi_t\|)) \mathbf{I}_{m-r_A}. \end{aligned} \quad (20)$$

From Lemma 24 and our assumption of the RIP property of $\mathcal{M}(\cdot)$, we obtain upper bounds for the two error terms.

$$\begin{aligned} \|\Delta_{t-1}\| &= \|(\mathcal{M}^* \mathcal{M} - \mathcal{I})(\mathbf{X}_t \Theta_{t-1} \mathbf{X}_t^\top - \mathbf{A})\| \\ &\leq \delta \|\mathbf{X}_t \Theta_{t-1} \mathbf{X}_t^\top - \mathbf{A}\|_F \\ &\leq \delta(\|\mathbf{X}_t \Theta_{t-1} \mathbf{X}_t^\top\|_F + \|\mathbf{A}\|_F) \\ &\stackrel{(f)}{\leq} (2\sqrt{r} + \sqrt{r_A})\delta, \end{aligned}$$

$$\begin{aligned} \|\Xi_t\| &= \|(\mathcal{M}^* \mathcal{M} - \mathcal{I})(\mathbf{X}_t \Theta_t \mathbf{X}_t^\top - \mathbf{A})\| \\ &\leq \delta \|\mathbf{X}_t \Theta_t \mathbf{X}_t^\top - \mathbf{A}\|_F \\ &\leq \delta(\|\mathbf{X}_t \Theta_t \mathbf{X}_t^\top\|_F + \|\mathbf{A}\|_F) \\ &\stackrel{(f)}{\leq} (2\sqrt{r} + \sqrt{r_A})\delta, \end{aligned}$$

where (f) is from $\|\mathbf{X}_t\| \leq 1$, $\|\Theta_{t-1}\|_F \leq \sqrt{r}\|\Theta_{t-1}\| \leq 2\sqrt{r}$, $\|\Theta_t\|_F \leq \sqrt{r}\|\Theta_t\| \leq 2\sqrt{r}$, and $\|\mathbf{A}\|_F \leq \sqrt{r_A}\|\mathbf{A}\| \leq \sqrt{r_A}$. Plugging these two upper bounds into (20), we arrive at

$$\Psi_{t+1} \Psi_{t+1}^\top \preceq (1 + 6\eta(\sqrt{r_A} + 2\sqrt{r})\delta)^2 \Psi_t \Psi_t^\top + (4\eta(\sqrt{r_A} + 2\sqrt{r})\delta + 28\eta^2(\sqrt{r_A} + 2\sqrt{r})^2 \delta^2) \mathbf{I}_{m-r_A}.$$

We now consider the relationship between $\Psi_1 \Psi_1^\top$ and $\Psi_0 \Psi_0^\top$.

Let $\tilde{\mathbf{L}}_0 := \frac{1}{2}(\mathbf{X}_0^\top \mathbf{A} \mathbf{X}_0 \Theta_0 + \Theta_0 \mathbf{X}_0^\top \mathbf{A} \mathbf{X}_0) - \frac{1}{2}(\mathbf{X}_0^\top \Xi_0 \mathbf{X}_0 \Theta_0 + \Theta_0 \mathbf{X}_0^\top \Xi_0 \mathbf{X}_0)$.

Multiplying the update formula (4) at $t = 0$ on the left by $\mathbf{U}_\perp^\top$, we have that

$$\Psi_1 = \mathbf{U}_\perp^\top (\mathbf{X}_0 - \eta \mathbf{G}_0) (\mathbf{I}_r + \eta^2 \mathbf{G}_0^\top \mathbf{G}_0)^{-1/2}.$$

Consequently, we derive the following upper bound on $\Psi_1 \Psi_1^\top$:

$$\begin{aligned} \Psi_1 \Psi_1^\top &= \mathbf{U}_\perp^\top (\mathbf{X}_0 - \eta \mathbf{G}_0) (\mathbf{I}_r + \eta^2 \mathbf{G}_0^\top \mathbf{G}_0)^{-1} (\mathbf{X}_0 - \eta \mathbf{G}_0)^\top \mathbf{U}_\perp \\ &\stackrel{(g)}{\preceq} \mathbf{U}_\perp^\top (\mathbf{X}_0 - \eta \mathbf{G}_0) (\mathbf{X}_0 - \eta \mathbf{G}_0)^\top \mathbf{U}_\perp \\ &\stackrel{(h)}{=} (\Psi_0 (\mathbf{I}_r - \eta \tilde{\mathbf{L}}_0) - \eta \mathbf{U}_\perp^\top \Xi_0 \mathbf{X}_0 \Theta_0) (\Psi_0 (\mathbf{I}_r - \eta \tilde{\mathbf{L}}_0) - \eta \mathbf{U}_\perp^\top \Xi_0 \mathbf{X}_0 \Theta_0)^\top \\ &= \Psi_0 (\mathbf{I}_r - \eta \tilde{\mathbf{L}}_0) (\mathbf{I}_r - \eta \tilde{\mathbf{L}}_0)^\top \Psi_0^\top - \eta \Psi_0 (\mathbf{I}_r - \eta \tilde{\mathbf{L}}_0) \Theta_0 \mathbf{X}_0^\top \Xi_0 \mathbf{U}_\perp \\ &\quad - \eta \mathbf{U}_\perp^\top \Xi_0 \mathbf{X}_0 \Theta_0 (\mathbf{I}_r - \eta \tilde{\mathbf{L}}_0)^\top \Psi_0^\top + \eta^2 \mathbf{U}_\perp^\top \Xi_0 \mathbf{X}_0 \Theta_0^2 \mathbf{X}_0^\top \Xi_0 \mathbf{U}_\perp, \end{aligned} \quad (21)$$

where (g) is from that $(\mathbf{I}_r + \eta^2 \mathbf{G}_0^\top \mathbf{G}_0)^{-1}$ is PSD and all of its eigenvalues are smaller than 1; and (h) is by expanding the expression of $\mathbf{G}_0$ directly. Since $\mathbf{Y}\mathbf{Y}^\top \preceq \|\mathbf{Y}\|^2 \mathbf{I}_r$ holds for any symmetric matrix $\mathbf{Y} \in \mathbb{R}^{r \times r}$ and by Lemma 25, we can upper bound the three terms as follows.

For the first term, it holds that

$$\Psi_0(\mathbf{I}_r - \eta \tilde{\mathbf{L}}_0)(\mathbf{I}_r - \eta \tilde{\mathbf{L}}_0)^\top \Psi_0^\top \preceq \|\mathbf{I}_r - \eta \tilde{\mathbf{L}}_0\|^2 \Psi_0 \Psi_0^\top. \quad (22)$$

For the second term, we have that

$$\begin{aligned} & \Psi_0(\mathbf{I}_r - \eta \tilde{\mathbf{L}}_0) \Theta_0 \mathbf{X}_0^\top \Xi_0 \mathbf{U}_\perp + \mathbf{U}_\perp^\top \Xi_0 \mathbf{X}_0 \Theta_0 (\mathbf{I}_r - \eta \tilde{\mathbf{L}}_0)^\top \Psi_0^\top \\ & \preceq 2 \|\Psi_0(\mathbf{I}_r - \eta \tilde{\mathbf{L}}_0)\| \|\Theta_0 \mathbf{X}_0^\top \Xi_0 \mathbf{U}_\perp\| \mathbf{I}_{m-r_A}. \end{aligned} \quad (23)$$

For the third term, we can obtain that

$$\mathbf{U}_\perp^\top \Xi_0 \mathbf{X}_0 \Theta_0^2 \mathbf{X}_0^\top \Xi_0 \mathbf{U}_\perp \preceq \|\mathbf{U}_\perp^\top \Xi_0 \mathbf{X}_0 \Theta_0^2 \mathbf{X}_0^\top \Xi_0 \mathbf{U}_\perp\| \mathbf{I}_{m-r_A}. \quad (24)$$

Combining inequalities (21), (22), (23) and (24), it follows that

$$\begin{aligned} \Psi_1 \Psi_1^\top & \preceq \|\mathbf{I}_r - \eta \tilde{\mathbf{L}}_0\|^2 \Psi_0 \Psi_0^\top + 2\eta \|\Psi_0(\mathbf{I}_r - \eta \tilde{\mathbf{L}}_0)\| \|\Theta_0 \mathbf{X}_0^\top \Xi_0 \mathbf{U}_\perp\| \mathbf{I}_{m-r_A} \\ & \quad + \eta^2 \|\mathbf{U}_\perp^\top \Xi_0 \mathbf{X}_0 \Theta_0^2 \mathbf{X}_0^\top \Xi_0 \mathbf{U}_\perp\| \mathbf{I}_{m-r_A} \\ & \stackrel{(i)}{\preceq} (1 + \eta \|\tilde{\mathbf{L}}_0\|)^2 \Psi_0 \Psi_0^\top + 4\eta(1 + \eta \|\tilde{\mathbf{L}}_0\|) \|\Xi_0\| \mathbf{I}_{m-r_A} + 4\eta^2 \|\Xi_0\|^2 \mathbf{I}_{m-r_A}, \end{aligned} \quad (25)$$

where (i) is by $\|\mathbf{X}_0\| \leq 1$, $\|\mathbf{U}_\perp\| \leq 1$, and $\|\Theta_0\| \leq 2$.

From Lemma 24 and our assumption of the RIP property of $\mathcal{M}(\cdot)$, we have that

$$\begin{aligned} \|\Xi_0\| & = \|(\mathcal{M}^* \mathcal{M} - \mathcal{I})(\mathbf{X}_0 \Theta_0 \mathbf{X}_0^\top - \mathbf{A})\| \\ & \leq \delta \|\mathbf{X}_0 \Theta_0 \mathbf{X}_0^\top - \mathbf{A}\|_F \\ & \leq \delta (\|\mathbf{X}_0 \Theta_0 \mathbf{X}_0^\top\|_F + \|\mathbf{A}\|_F) \\ & \leq (2\sqrt{r} + \sqrt{r_A})\delta. \end{aligned}$$

Then, we can bound $\|\tilde{\mathbf{L}}_0\|$ as follows:

$$\begin{aligned} \|\tilde{\mathbf{L}}_0\| & = \frac{1}{2} \|\mathbf{X}_0^\top \mathbf{A} \mathbf{X}_0 \Theta_0 + \Theta_0 \mathbf{X}_0^\top \mathbf{A} \mathbf{X}_0 - (\mathbf{X}_0^\top \Xi_0 \mathbf{X}_0 \Theta_0 + \Theta_0 \mathbf{X}_0 \Xi_0 \mathbf{X}_0)\| \\ & \leq 2 + 2\|\Xi_0\| \\ & \leq 2 + 2(2\sqrt{r} + \sqrt{r_A})\delta. \end{aligned}$$

Plugging theses two upper bounds into inequality (25), we finally arrive at

$$\Psi_1 \Psi_1^\top \preceq (1 + 2\eta + 2\eta(\sqrt{r_A} + 2\sqrt{r})\delta)^2 \Psi_0 \Psi_0^\top + (12\eta(\sqrt{r_A} + 2\sqrt{r})\delta + 8\eta^2(\sqrt{r_A} + 2\sqrt{r})^2 \delta^2) \mathbf{I}_{m-r_A}.$$

■

G.7. Proof of Lemma 13

Proof We first estimate $\|\tilde{\mathbf{G}}_t\|$ and $\|\mathbf{G}_t\|$. From the expression of $\tilde{\mathbf{G}}_t$, we have that

$$\begin{aligned}
 \|\tilde{\mathbf{G}}_t\| &= \|\mathcal{M}^* \mathcal{M}(\mathbf{X}_t \boldsymbol{\Theta}_t \mathbf{X}_t^\top - \mathbf{A})\| \mathbf{X}_t \boldsymbol{\Theta}_t \\
 &\stackrel{(a)}{\leq} 2\|\mathcal{M}^* \mathcal{M}(\mathbf{X}_t \boldsymbol{\Theta}_t \mathbf{X}_t^\top - \mathbf{A})\| \\
 &\leq 2(\|(\mathcal{M}^* \mathcal{M} - \mathcal{I})(\mathbf{X}_t \boldsymbol{\Theta}_t \mathbf{X}_t^\top - \mathbf{A})\| + \|\mathbf{X}_t \boldsymbol{\Theta}_t \mathbf{X}_t^\top - \mathbf{A}\|) \\
 &\stackrel{(b)}{\leq} \frac{2}{\sqrt{m}} \|\mathbf{X}_t \boldsymbol{\Theta}_t \mathbf{X}_t^\top - \mathbf{A}\|_F + 2\|\mathbf{X}_t \boldsymbol{\Theta}_t \mathbf{X}_t^\top - \mathbf{A}\| \\
 &\leq 4\|\mathbf{X}_t \boldsymbol{\Theta}_t \mathbf{X}_t^\top - \mathbf{A}\| \\
 &\leq 4(\|\mathbf{X}_t \boldsymbol{\Theta}_t \mathbf{X}_t^\top\| + \|\mathbf{A}\|) \\
 &\stackrel{(a)}{\leq} 12,
 \end{aligned}$$

where (a) is due to $\|\mathbf{X}_t\| \leq 1$, $\|\boldsymbol{\Theta}_t\| \leq 2$, and $\|\mathbf{A}\| \leq 1$; and (b) is from Lemma 24. Analogously, we can upper bound $\|\mathbf{G}_t\|$ as follows:

$$\begin{aligned}
 \|\mathbf{G}_t\| &= \|(\mathbf{I}_m - \mathbf{X}_t \mathbf{X}_t^\top) \tilde{\mathbf{G}}_t + \frac{\mathbf{X}_t}{2} (\mathbf{X}_t^\top \tilde{\mathbf{G}}_t - \tilde{\mathbf{G}}_t^\top \mathbf{X}_t)\| \\
 &\leq \|\mathbf{I}_m - \mathbf{X}_t \mathbf{X}_t^\top\| \|\tilde{\mathbf{G}}_t\| + \|\mathbf{X}_t\| \|\tilde{\mathbf{G}}_t^\top \mathbf{X}_t\| \\
 &\stackrel{(c)}{\leq} 2\|\tilde{\mathbf{G}}_t\| \\
 &\leq 24,
 \end{aligned}$$

where (c) follows from $\|\mathbf{X}_t\| \leq 1$ and the fact that all the eigenvalues of the PSD matrix $\mathbf{I}_m - \mathbf{X}_t \mathbf{X}_t^\top$ are less than 1. Multiplying the update formula (4) on the left by $\mathbf{U}^\top$, we obtain

$$\begin{aligned}
 \Phi_{t+1} \Phi_{t+1}^\top &= \mathbf{U}^\top \mathbf{X}_{t+1} \mathbf{X}_{t+1}^\top \mathbf{U} \\
 &= (\Phi_t - \eta \mathbf{U}^\top \mathbf{G}_t) (\mathbf{I}_r + \eta^2 \mathbf{G}_t^\top \mathbf{G}_t)^{-1} (\Phi_t^\top - \eta \mathbf{G}_t^\top \mathbf{U}) \\
 &\stackrel{(d)}{\succeq} (\Phi_t - \eta \mathbf{U}^\top \mathbf{G}_t) (\mathbf{I}_r - \eta^2 \mathbf{G}_t^\top \mathbf{G}_t) (\Phi_t^\top - \eta \mathbf{G}_t^\top \mathbf{U}) \\
 &= \Phi_t \Phi_t^\top - \eta^2 \Phi_t \mathbf{G}_t^\top \mathbf{G}_t \Phi_t^\top - \eta (\Phi_t \mathbf{G}_t^\top \mathbf{U} + \mathbf{U}^\top \mathbf{G}_t \Phi_t^\top) \\
 &\quad + \eta^3 (\Phi_t \mathbf{G}_t^\top \mathbf{G}_t \mathbf{G}_t^\top \mathbf{U} + \mathbf{U}^\top \mathbf{G}_t \mathbf{G}_t^\top \mathbf{G}_t \Phi_t^\top) \\
 &\quad - \eta^4 \mathbf{U}^\top \mathbf{G}_t \mathbf{G}_t^\top \mathbf{G}_t \mathbf{G}_t^\top \mathbf{U} + \eta^2 \mathbf{U}^\top \mathbf{G}_t \mathbf{G}_t^\top \mathbf{U} \\
 &\stackrel{(e)}{\succeq} \Phi_t \Phi_t^\top - (\eta^2 \|\Phi_t \mathbf{G}_t^\top \mathbf{G}_t \Phi_t^\top\| + 2\eta \|\Phi_t \mathbf{G}_t^\top \mathbf{U}\| \\
 &\quad + 2\eta^3 \|\Phi_t \mathbf{G}_t^\top \mathbf{G}_t \mathbf{G}_t^\top \mathbf{U}\| + \eta^4 \|\mathbf{U}^\top \mathbf{G}_t \mathbf{G}_t^\top \mathbf{G}_t \mathbf{G}_t^\top \mathbf{U}\|) \mathbf{I}_{r_A} \\
 &\stackrel{(f)}{\succeq} \Phi_t \Phi_t^\top - \frac{1}{10r_A} \mathbf{I}_{r_A},
 \end{aligned}$$

where (d) is from Lemma 14; (e) is by Lemma 25; and (f) is due to $\|\Phi_t\| \leq 1$, $\|\mathbf{U}\| \leq 1$, $\|\mathbf{G}_t\| \leq 24$ and $\eta \leq \frac{1}{500r_A}$. By subtracting the inequality from $\mathbf{I}_{r_A}$, it follows that

$$\mathbf{I}_{r_A} - \Phi_{t+1} \Phi_{t+1}^\top \preceq \mathbf{I}_{r_A} - \Phi_t \Phi_t^\top + \frac{1}{10r_A} \mathbf{I}_{r_A}.$$

Taking trace on both sides yields

$$\text{Tr}(\mathbf{I}_{r_A} - \Phi_{t+1}\Phi_{t+1}^\top) \leq \text{Tr}(\mathbf{I}_{r_A} - \Phi_t\Phi_t^\top) + 0.1.$$

■

G.8. Proof of Theorem 2

Proof For the proof, we take $\eta = \frac{(r-r_A)^4}{975c_1^2\kappa^2m^2r^2r_A}$ and $\delta = \frac{c_4(r-r_A)^6}{\kappa^2m^3r^4r_A}$, where $c_4 = \mathcal{O}(\frac{1}{c_1^3})$. From Lemma 11, we have that $\|\Theta_t\| \leq 2$ holds for all $t \geq 0$ by mathematical induction. For later use, we define the following three terms in the same way as in Lemma 10:

$$\begin{aligned} \beta_t &:= \sigma_1(\mathbf{I}_{r_A} - \Phi_t\Phi_t^\top) \leq 1, \\ \chi_t &:= (\|\Delta_{t-1}\| + \|\Xi_t\|)^2 + \sqrt{\text{Tr}(\mathbf{I}_{r_A} - \Phi_t\Phi_t^\top)(\|\Delta_{t-1}\| + \|\Xi_t\|)}, \\ \mathbf{H}_t &:= (\mathbf{I}_m - \mathbf{X}_t\mathbf{X}_t^\top)(\mathbf{A}\mathbf{X}_t\mathbf{X}_t^\top\Delta_{t-1}\mathbf{X}_t + \Xi_t\mathbf{X}_t\Theta_t) \\ &\quad + \frac{1}{2}(\mathbf{X}_t\mathbf{X}_t^\top\Xi_t\mathbf{X}_t\Theta_t - \mathbf{X}_t\Theta_t\mathbf{X}_t^\top\Xi_t\mathbf{X}_t) \\ &\quad + \frac{1}{2}(\mathbf{X}_t\mathbf{X}_t^\top\mathbf{A}\mathbf{X}_t\mathbf{X}_t^\top\Delta_{t-1}\mathbf{X}_t - \mathbf{X}_t\mathbf{X}_t^\top\Delta_{t-1}\mathbf{X}_t\mathbf{X}_t^\top\mathbf{A}\mathbf{X}_t). \end{aligned}$$

Lemma 9 with the RIP property of $\mathcal{M}(\cdot)$ implies that $\|\Delta_{t-1}\|_F, \|\Xi_t\|_F \leq 1$ for all $t \geq 1$. Thus, the assumptions of Lemma 10 are met, guaranteeing that inequality (8) holds for all iterations. Building on inequality (8), we divide the convergence analysis into two phases.

Phase I (Initial phase). $\text{Tr}(\mathbf{I}_{r_A} - \Phi_t\Phi_t^\top) \geq 0.5$.

We assume for now that $\sigma_{r_A}^2(\Phi_t) \geq (r-r_A)^2/(2c_1mr)$ holds in Phase I, which will be proved later. Let the compact SVD of Φ_t be $\mathbf{Q}_t\mathbf{\Lambda}_t\mathbf{P}_t^\top$, where $\mathbf{Q}_t \in \mathbb{R}^{r_A \times r_A}$, $\mathbf{\Lambda}_t \in \mathbb{R}^{r_A \times r_A}$, and $\mathbf{P}_t \in \mathbb{R}^{r \times r_A}$. We can simplify (8) as follows:

$$\begin{aligned} &\text{Tr}(\mathbf{I}_{r_A} - \Phi_{t+1}\Phi_{t+1}^\top) - \text{Tr}(\mathbf{I}_{r_A} - \Phi_t\Phi_t^\top) \\ &\leq \eta^2(1 + 16\chi_t)\text{Tr}(\Phi_t\Phi_t^\top) - \frac{2\eta(1 - \eta^2 - 16\eta^2\chi_t)\sigma_{r_A}^2(\Phi_t)}{\kappa^2}\text{Tr}((\mathbf{I}_{r_A} - \Phi_t\Phi_t^\top)\Phi_t\Phi_t^\top) \\ &\quad + 2\eta\sqrt{r_A}(\|\Delta_{t-1}\|_F + 2\|\Xi_t\|_F) + 2\eta^2\sqrt{r_A}\|\mathbf{H}_t\|_F \\ &\stackrel{(a)}{\leq} \eta^2(1 + 16\chi_t)r_A - \frac{2\eta(1 - \eta^2 - 16\eta^2\chi_t)\sigma_{r_A}^4(\Phi_t)}{\kappa^2}\text{Tr}(\mathbf{I}_{r_A} - \Phi_t\Phi_t^\top) \\ &\quad + 2\eta\sqrt{r_A}(\|\Delta_{t-1}\|_F + 2\|\Xi_t\|_F) + 2\eta^2\sqrt{r_A}\|\mathbf{H}_t\|_F \\ &\stackrel{(b)}{\leq} \eta^2(1 + 16\chi_t)r_A - \frac{\eta(1 - \eta^2 - 16\eta^2\chi_t)(r-r_A)^4}{2c_1^2\kappa^2m^2r^2}\text{Tr}(\mathbf{I}_{r_A} - \Phi_t\Phi_t^\top) \\ &\quad + 2\eta\sqrt{r_A}(\|\Delta_{t-1}\|_F + 2\|\Xi_t\|_F) + 2\eta^2\sqrt{r_A}\|\mathbf{H}_t\|_F, \end{aligned} \tag{26}$$

where (a) is by Lemma 16 and $\text{Tr}((\mathbf{I}_{r_A} - \Phi_t\Phi_t^\top)\Phi_t\Phi_t^\top) = \text{Tr}((\mathbf{I}_{r_A} - \mathbf{\Lambda}_t^2)\mathbf{\Lambda}_t^2)$; and (b) is from our assumption that $\sigma_{r_A}^2(\Phi_t) \geq (r-r_A)^2/(2c_1mr)$.

Using Lemma 24 and the RIP property of $\mathcal{M}(\cdot)$, we can control the quantities of the two error terms. In particular, following inequalities imply that both $\|\Delta_{t-1}\|_F$ and $\|\Xi_t\|_F$ are uniformly bounded by a constant that depends only on m, r and r_A but is independent of t .

Expanding the expression of Δ_{t-1} and applying Lemma 24, we have that

$$\begin{aligned}
 \|\Delta_{t-1}\|_F &\leq \sqrt{m}\|(\mathcal{M}^*\mathcal{M} - \mathcal{I})(\mathbf{X}_t\Theta_{t-1}\mathbf{X}_t^\top - \mathbf{A})\| \\
 &\leq \frac{c_4(r-r_A)^6}{\kappa^2 m^{5/2} r^4 r_A} \|\mathbf{X}_t\Theta_{t-1}\mathbf{X}_t^\top - \mathbf{A}\|_F \\
 &\leq \frac{c_4(r-r_A)^6}{\kappa^2 m^{5/2} r^4 r_A} (2\sqrt{r} + \sqrt{r_A}) \\
 &\leq \frac{3c_4(r-r_A)^6}{\kappa^2 m^{5/2} r^{7/2} r_A} \\
 &\stackrel{(c)}{\leq} \min\left\{\frac{(r-r_A)^4}{48c_1^2 \kappa^2 m^2 r^2 r_A}, \frac{1}{48\sqrt{r_A}}\right\}.
 \end{aligned} \tag{27}$$

Applying the same reasoning to Ξ_t , it follows that

$$\begin{aligned}
 \|\Xi_t\|_F &\leq \sqrt{m}\|(\mathcal{M}^*\mathcal{M} - \mathcal{I})(\mathbf{X}_t\Theta_t\mathbf{X}_t^\top - \mathbf{A})\| \\
 &\leq \frac{c_4(r-r_A)^6}{\kappa^2 m^{5/2} r^4 r_A} \|\mathbf{X}_t\Theta_t\mathbf{X}_t^\top - \mathbf{A}\|_F \\
 &\leq \frac{c_4(r-r_A)^6}{\kappa^2 m^{5/2} r^4 r_A} (2\sqrt{r} + \sqrt{r_A}) \\
 &\leq \frac{3c_4(r-r_A)^6}{\kappa^2 m^{5/2} r^{7/2} r_A} \\
 &\stackrel{(c)}{\leq} \min\left\{\frac{(r-r_A)^4}{48c_1^2 \kappa^2 m^2 r^2 r_A}, \frac{1}{48\sqrt{r_A}}\right\}.
 \end{aligned} \tag{28}$$

Here, (c) is from $c_4 = \mathcal{O}(\frac{1}{c_1^3})$, $c_1 > 1$ and $r - r_A \leq r \leq m$. Since $\text{Tr}(\mathbf{I}_{r_A} - \Phi_t\Phi_t^\top) \leq r_A$, together with (27) and (28), we can upper bound χ_t as follows:

$$\begin{aligned}
 \chi_t &= (\|\Delta_{t-1}\| + \|\Xi_t\|)^2 + \sqrt{\text{Tr}(\mathbf{I}_{r_A} - \Phi_t\Phi_t^\top)}(\|\Delta_{t-1}\| + \|\Xi_t\|) \\
 &\leq (\|\Delta_{t-1}\|_F + \|\Xi_t\|_F)^2 + \sqrt{r_A}(\|\Delta_{t-1}\|_F + \|\Xi_t\|_F) \\
 &\leq \left(\frac{1}{48} + \frac{1}{48}\right)^2 + \sqrt{r_A}\left(\frac{1}{48\sqrt{r_A}} + \frac{1}{48\sqrt{r_A}}\right) \\
 &\leq \frac{1}{16}.
 \end{aligned}$$

From inequalities (9), (27), and (28), we obtain the following upper bound on $\|\mathbf{H}_t\|_F$.

$$\begin{aligned}
 \|\mathbf{H}_t\|_F &\leq 2\|\Delta_{t-1}\|_F + 4\|\Xi_t\|_F \\
 &\leq 2 \times \frac{1}{48\sqrt{r_A}} + 4 \times \frac{1}{48\sqrt{r_A}} \\
 &\leq \frac{1}{2\sqrt{r_A}}.
 \end{aligned}$$

With these upper bounds, inequality (26) can be simplified as follows:

$$\begin{aligned}
 & \text{Tr}(\mathbf{I}_{r_A} - \Phi_{t+1}\Phi_{t+1}^\top) - \text{Tr}(\mathbf{I}_{r_A} - \Phi_t\Phi_t^\top) \\
 & \leq 2\eta^2 r_A - \frac{\eta(1-2\eta^2)(r-r_A)^4}{2c_1^2\kappa^2m^2r^2} \text{Tr}(\mathbf{I}_{r_A} - \Phi_t\Phi_t^\top) + \frac{\eta(r-r_A)^4}{8c_1^2\kappa^2m^2r^2} + \eta^2 \\
 & \stackrel{(d)}{\leq} \left(-\frac{\eta(r-r_A)^4}{2c_1^2\kappa^2m^2r^2} + \frac{\eta(r-r_A)^4}{4c_1^2\kappa^2m^2r^2} + 6\eta^2 r_A + \frac{\eta^3(r-r_A)^4}{c_1^2\kappa^2m^2r^2} \right) \text{Tr}(\mathbf{I}_{r_A} - \Phi_t\Phi_t^\top) \\
 & = \left(-\frac{\eta(r-r_A)^4}{4c_1^2\kappa^2m^2r^2} + 6\eta^2 r_A + \frac{\eta^3(r-r_A)^4}{c_1^2\kappa^2m^2r^2} \right) \text{Tr}(\mathbf{I}_{r_A} - \Phi_t\Phi_t^\top) \\
 & \stackrel{(e)}{\leq} \frac{1}{2} \left(-\frac{\eta(r-r_A)^4}{4c_1^2\kappa^2m^2r^2} + 6\eta^2 r_A + \frac{\eta^3(r-r_A)^4}{c_1^2\kappa^2m^2r^2} \right),
 \end{aligned}$$

where (d) is by $\text{Tr}(\mathbf{I}_{r_A} - \Phi_t\Phi_t^\top) \geq 0.5$; and (e) holds if the expression in bracket is less than zero.

Recall that $\eta = \frac{(r-r_A)^4}{975c_1^2\kappa^2m^2r^2r_A}$. The summation of the terms in bracket is negative, which implies that at each step, $\text{Tr}(\mathbf{I}_{r_A} - \Phi_t\Phi_t^\top)$ decreases at least by $\Delta := \frac{(r-r_A)^8}{7000c_1^4\kappa^4m^4r^4r_A}$. Consequently, after at most $(r_A - 0.5)/\Delta \leq \frac{7000c_1^4\kappa^4m^4r^4r_A}{(r-r_A)^8}$ iterations, RGD leaves Phase I.

Let $c_2 := \frac{1}{7000c_1^4} \in (0, 1)$. Denote $t_0 \geq 1$ as the last iteration in this phase. The analysis above implies that $\text{Tr}(\mathbf{I}_{r_A} - \Phi_t\Phi_t^\top) \leq r_A - \frac{c_2(r-r_A)^8t}{\kappa^4m^4r^4r_A}$ for all $1 \leq t \leq t_0$ and $t_0 \leq \frac{7000c_1^4\kappa^4m^4r^4r_A}{(r-r_A)^8}$.

From Lemma 8 and inequality (27), we obtain the following bound for $1 \leq t \leq t_0$:

$$\begin{aligned}
 \|\mathbf{X}_t\Theta_t\mathbf{X}_t^\top - \mathbf{A}\|_F & \leq 2\sqrt{\text{Tr}(\mathbf{I}_{r_A} - \Phi_t\Phi_t^\top)} + \|\Delta_{t-1}\|_F \\
 & \leq 2\sqrt{r_A - \frac{c_2(r-r_A)^8t}{\kappa^4m^4r^4r_A}} + \|\Delta_{t-1}\|_F \\
 & \leq 2\sqrt{r_A - \frac{c_2(r-r_A)^8t}{\kappa^4m^4r^4r_A}} + 1.
 \end{aligned}$$

We now prove that $\sigma_{r_A}^2(\Phi_t) \geq (r-r_A)^2/(2c_1mr)$ holds in Phase I. By Lemma 7, it holds w.h.p.,

$$\sigma_{r_A}^2(\Phi_0) = \sigma_{r_A}^2(\mathbf{U}^\top \mathbf{X}_0) \geq \frac{(r-r_A)^2}{c_1mr}.$$

Moreover, by Lemma 15 and the assumption $r_A \leq \frac{m}{2}$, it follows that

$$\sigma_{r_A}^2(\Psi_0) = 1 - \sigma_{r_A}^2(\Phi_0) \leq 1 - \frac{(r-r_A)^2}{c_1mr}.$$

Since $\eta = \frac{(r-r_A)^4}{975c_1^2\kappa^2m^2r^2r_A}$ and $\delta = \frac{c_4(r-r_A)^6}{\kappa^2m^3r^4r_A}$, we can deduce that

$$\eta(\sqrt{r_A} + 2\sqrt{r})\delta \leq \frac{c_4(r-r_A)^{10}}{325c_1^2\kappa^4m^5r^5r_A^2}.$$

From Lemma 12 and the upper bound on $\eta(\sqrt{r_A} + 2\sqrt{r})\delta$, we obtain the following inequality

$$\Psi_1 \Psi_1^\top \preceq \left(1 + \frac{4(r - r_A)^4}{975c_1^2 \kappa^2 m^2 r^2 r_A}\right)^2 \Psi_0 \Psi_0^\top + \frac{4c_4(r - r_A)^{10}}{65c_1^2 \kappa^4 m^5 r^5 r_A^2} \mathbf{I}_{m-r_A}.$$

Using Weyl's inequality and $c_4 = \mathcal{O}(\frac{1}{c_1^3})$, we have the following upper bound on $\sigma_{r_A}^2(\Psi_1)$

$$\begin{aligned} \sigma_{r_A}^2(\Psi_1) &\leq \left(1 + \frac{4(r - r_A)^4}{975c_1^2 \kappa^2 m^2 r^2 r_A}\right)^2 \sigma_{r_A}^2(\Psi_0) + \frac{4c_4(r - r_A)^{10}}{65c_1^2 \kappa^4 m^5 r^5 r_A^2} \\ &\leq \left(1 + \frac{4(r - r_A)^4}{975c_1^2 \kappa^2 m^2 r^2 r_A}\right)^2 \left(1 - \frac{(r - r_A)^2}{c_1 m r}\right) + \frac{4c_4(r - r_A)^{10}}{65c_1^2 \kappa^4 m^5 r^5 r_A^2} \\ &\stackrel{(f)}{\leq} 1 + \frac{16(r - r_A)^4}{195c_1^2 \kappa^2 m^2 r^2 r_A} - \frac{(r - r_A)^2}{c_1 m r} \\ &\stackrel{(f)}{\leq} 1 - \frac{2(r - r_A)^2}{3c_1 m r}, \end{aligned}$$

where (f) is by $r - r_A \leq r \leq m$ and $c_1, \kappa, r_A \geq 1$. Applying Lemma 12 with the upper bound on $\eta(\sqrt{r_A} + 2\sqrt{r})\delta$, we obtain

$$\Psi_{t+1} \Psi_{t+1}^\top \preceq \left(1 + \frac{c_4(r - r_A)^{10}}{40c_1^2 \kappa^4 m^5 r^5 r_A^2}\right)^2 \Psi_t \Psi_t^\top + \frac{c_4(r - r_A)^{10}}{40c_1^2 \kappa^4 m^5 r^5 r_A^2} \mathbf{I}_{m-r_A}, \quad t \geq 1.$$

Using Weyl's inequality, we have the following relationship between $\sigma_{r_A}^2(\Psi_{t+1})$ and $\sigma_{r_A}^2(\Psi_t)$

$$\begin{aligned} \sigma_{r_A}^2(\Psi_{t+1}) &= \sigma_{r_A}(\Psi_{t+1} \Psi_{t+1}^\top) \leq \left(1 + \frac{c_4(r - r_A)^{10}}{40c_1^2 \kappa^4 m^5 r^5 r_A^2}\right)^2 \sigma_{r_A}(\Psi_t \Psi_t^\top) + \frac{c_4(r - r_A)^{10}}{40c_1^2 \kappa^4 m^5 r^5 r_A^2} \\ &= \left(1 + \frac{c_4(r - r_A)^{10}}{40c_1^2 \kappa^4 m^5 r^5 r_A^2}\right)^2 \sigma_{r_A}^2(\Psi_t) + \frac{c_4(r - r_A)^{10}}{40c_1^2 \kappa^4 m^5 r^5 r_A^2}. \end{aligned}$$

Denote $\zeta := \frac{c_4(r - r_A)^{10}}{40c_1^2 \kappa^4 m^5 r^5 r_A^2}$. By iterating the recursive inequality, the following upper bound holds

$$\begin{aligned} \sigma_{r_A}^2(\Psi_t) &\leq (1 + \zeta)^{2(t-1)} \sigma_{r_A}^2(\Psi_1) + \zeta \sum_{i=0}^{t-2} (1 + \zeta)^{2i} \\ &\leq (1 + \zeta)^{2t} \sigma_{r_A}^2(\Psi_1) + \zeta \sum_{i=0}^{t-1} (1 + \zeta)^{2i} \\ &= (1 + \zeta)^{2t} \sigma_{r_A}^2(\Psi_1) + \zeta \left[(1 + \zeta)^{2t} - 1 \right] / \left[(1 + \zeta)^2 - 1 \right] \\ &\leq (1 + \zeta)^{2t} \sigma_{r_A}^2(\Psi_1) + (1 + \zeta)^{2t} - 1, \end{aligned}$$

for all $1 \leq t \leq \frac{7000c_1^4 \kappa^4 m^4 r^4 r_A^2}{(r - r_A)^8}$. Invoking Lemma 26 and noting that $\zeta \leq \frac{1}{2t}$, which is ensured by $c_4 = \mathcal{O}(\frac{1}{c_1^3})$, we obtain

$$\sigma_{r_A}^2(\Psi_t) \leq (1 + 6t\zeta) \sigma_{r_A}^2(\Psi_1) + 6t\zeta$$

$$\begin{aligned}
 &\stackrel{(g)}{\leq} \sigma_{r_A}^2(\Psi_1) + \frac{2100c_1^2c_4(r-r_A)^2}{mr} \\
 &\leq 1 - \frac{2(r-r_A)^2}{3c_1mr} + \frac{2100c_1^2c_4(r-r_A)^2}{mr} \\
 &\stackrel{(h)}{\leq} 1 - \frac{(r-r_A)^2}{2c_1mr},
 \end{aligned}$$

where (g) is by $t\zeta \leq \frac{175c_1^2c_4(r-r_A)^2}{mr}$ and $\sigma_{r_A}^2(\Psi_1) \leq 1$; and (h) holds by $c_4 = \mathcal{O}(\frac{1}{c_1^3})$. By Lemma 15 and the assumption $r_A \leq \frac{m}{2}$, it can be seen that $\sigma_{r_A}^2(\Phi_t) = 1 - \sigma_{r_A}^2(\Psi_t) \geq \frac{(r-r_A)^2}{2c_1mr}$ holds for all $t \leq t_0 \leq \frac{7000c_1^2\kappa^4m^4r_A^2}{(r-r_A)^8}$, i.e., throughout Phase I.

Phase II (Linearly convergent phase). $\text{Tr}(\mathbf{I}_{r_A} - \Phi_t \Phi_t^\top) < 0.5$.

This corresponds to a near-optimal regime. An immediate implication of this phase is that $\text{Tr}(\Phi_t \Phi_t^\top) \geq r_A - 0.5 \geq r_A - 0.6$. Recall that $t_0 \geq 1$ is the last iteration in the first phase. We assume that $\text{Tr}(\Phi_t \Phi_t^\top) \geq r_A - 0.6$ for all $t \geq t_0 + 1$, and we will prove this later.

Given that the singular values of $\Phi_t \Phi_t^\top$ lie in $[0, 1]$, we have

$$0.4 \leq \sigma_{r_A}^2(\Phi_t) = \sigma_{r_A}(\Phi_t \Phi_t^\top) \leq \sigma_1^2(\Phi_t) \leq 1.$$

Moreover, since $\beta_t = \sigma_1(\mathbf{I}_{r_A} - \Phi_t \Phi_t^\top) \leq \text{Tr}(\mathbf{I}_{r_A} - \Phi_t \Phi_t^\top)$ and $\beta_t \leq 1$, it follows that

$$\beta_t \text{Tr}(\Phi_t \Phi_t^\top) \leq r_A \text{Tr}(\mathbf{I}_{r_A} - \Phi_t \Phi_t^\top), \quad \chi_t \text{Tr}(\Phi_t \Phi_t^\top) \leq r_A \chi_t.$$

In addition, it can be derived that

$$\begin{aligned}
 \frac{4}{25} \text{Tr}(\mathbf{I}_{r_A} - \Phi_t \Phi_t^\top) &\leq \sigma_{r_A}^2(\Phi_t) \text{Tr}((\mathbf{I}_{r_A} - \Phi_t \Phi_t^\top) \Phi_t \Phi_t^\top) \\
 &\leq \sigma_1^2(\Phi_t) \text{Tr}(\mathbf{I}_{r_A} - \Phi_t \Phi_t^\top) \\
 &\leq \text{Tr}(\mathbf{I}_{r_A} - \Phi_t \Phi_t^\top).
 \end{aligned}$$

With the inequalities above, we can simplify (8) as follows:

$$\begin{aligned}
 \text{Tr}(\mathbf{I}_{r_A} - \Phi_{t+1} \Phi_{t+1}^\top) &\leq (1 - \frac{8\eta}{25\kappa^2} + \eta^2 r_A + \frac{2\eta^3}{\kappa^2} + \frac{32\eta^3}{\kappa^2} \chi_t) \text{Tr}(\mathbf{I}_{r_A} - \Phi_t \Phi_t^\top) \\
 &\quad + 16\eta^2 r_A \chi_t + 2\eta \sqrt{\text{Tr}(\mathbf{I}_{r_A} - \Phi_t \Phi_t^\top)} (\|\Delta_{t-1}\|_F + 2\|\Xi_t\|_F + \eta\|\mathbf{H}_t\|_F).
 \end{aligned} \tag{29}$$

Recall that $\chi_t = (\|\Delta_{t-1}\| + \|\Xi_t\|)^2 + \sqrt{\text{Tr}(\mathbf{I}_{r_A} - \Phi_t \Phi_t^\top)} (\|\Delta_{t-1}\| + \|\Xi_t\|)$ and $\|\mathbf{H}_t\|_F \leq 2\|\Delta_{t-1}\|_F + 4\|\Xi_t\|_F$. Since $\|\Delta_{t-1}\| \leq 1$ and $\|\Xi_t\| \leq 1$, (29) can be written as

$$\begin{aligned}
 \text{Tr}(\mathbf{I}_{r_A} - \Phi_{t+1} \Phi_{t+1}^\top) &\leq (1 - \frac{8\eta}{25\kappa^2} + \eta^2 r_A + \frac{194\eta^3}{\kappa^2}) \text{Tr}(\mathbf{I}_{r_A} - \Phi_t \Phi_t^\top) \\
 &\quad + 16\eta^2 r_A (\|\Delta_{t-1}\| + \|\Xi_t\|)^2 \\
 &\quad + \sqrt{\text{Tr}(\mathbf{I}_{r_A} - \Phi_t \Phi_t^\top)} \cdot (16\eta^2 r_A (\|\Delta_{t-1}\| + \|\Xi_t\|)) \\
 &\quad + \sqrt{\text{Tr}(\mathbf{I}_{r_A} - \Phi_t \Phi_t^\top)} \cdot (4\eta^2 + 2\eta) (\|\Delta_{t-1}\|_F + 2\|\Xi_t\|_F).
 \end{aligned}$$

Substituting $\eta = \frac{(r-r_A)^4}{975c_1^2\kappa^2m^2r^2r_A}$ into the inequality above, it follows that

$$\begin{aligned} \text{Tr}(\mathbf{I}_{r_A} - \Phi_{t+1}\Phi_{t+1}^\top) &\leq q\text{Tr}(\mathbf{I}_{r_A} - \Phi_t\Phi_t^\top) + \frac{1}{\kappa^4r_A}(\|\Delta_{t-1}\| + \|\Xi_t\|)^2 \\ &\quad + \sqrt{\text{Tr}(\mathbf{I}_{r_A} - \Phi_t\Phi_t^\top)} \cdot \frac{1}{\kappa^4r_A}(\|\Delta_{t-1}\| + \|\Xi_t\|) \\ &\quad + \sqrt{\text{Tr}(\mathbf{I}_{r_A} - \Phi_t\Phi_t^\top)} \cdot \left(\frac{1}{\kappa^4r_A^2} + \frac{1}{\kappa^2r_A}\right)(\|\Delta_{t-1}\|_F + 2\|\Xi_t\|_F), \end{aligned} \quad (30)$$

where $q := 1 - \frac{(r-r_A)^4}{8125c_1^2\kappa^4m^2r^2r_A}$ is a constant in $(\frac{1}{2}, 1)$.

From Lemma 9 and the RIP property of $\mathcal{M}(\cdot)$ with $\delta = \frac{c_4(r-r_A)^6}{\kappa^2m^3r^4r_A}$, it guarantees that

$$\begin{aligned} \|\Delta_{t-1}\| &\leq \|\Delta_{t-1}\|_F \leq \frac{c_4(r-r_A)^6}{\kappa^2m^{5/2}r^{7/2}r_A} \|\mathbf{X}_{t-1}\Theta_{t-1}\mathbf{X}_{t-1}^\top - \mathbf{A}\|_F \\ &\leq \frac{c_4(r-r_A)^4}{\kappa^2m^2r^2} \|\mathbf{X}_{t-1}\Theta_{t-1}\mathbf{X}_{t-1}^\top - \mathbf{A}\|_F, \\ \|\Xi_t\| &\leq \|\Xi_t\|_F \leq \frac{c_4(r-r_A)^6}{\kappa^2m^{5/2}r^{7/2}r_A} \|\mathbf{X}_t\Theta_t\mathbf{X}_t^\top - \mathbf{A}\|_F \\ &\leq \frac{c_4(r-r_A)^4}{\kappa^2m^2r^2} \|\mathbf{X}_t\Theta_t\mathbf{X}_t^\top - \mathbf{A}\|_F. \end{aligned}$$

Together with $c_4 = \mathcal{O}(\frac{1}{c_1^3})$, we can rewrite inequality (30) as

$$\begin{aligned} \text{Tr}(\mathbf{I}_{r_A} - \Phi_{t+1}\Phi_{t+1}^\top) &\leq q\text{Tr}(\mathbf{I}_{r_A} - \Phi_t\Phi_t^\top) + \frac{1-q}{180} \left(\|\mathbf{X}_{t-1}\Theta_{t-1}\mathbf{X}_{t-1}^\top - \mathbf{A}\|_F^2 + \|\mathbf{X}_t\Theta_t\mathbf{X}_t^\top - \mathbf{A}\|_F^2 \right. \\ &\quad \left. + 2\|\mathbf{X}_{t-1}\Theta_{t-1}\mathbf{X}_{t-1}^\top - \mathbf{A}\|_F \|\mathbf{X}_t\Theta_t\mathbf{X}_t^\top - \mathbf{A}\|_F \right. \\ &\quad \left. + \sqrt{\text{Tr}(\mathbf{I}_{r_A} - \Phi_t\Phi_t^\top)} (\|\mathbf{X}_{t-1}\Theta_{t-1}\mathbf{X}_{t-1}^\top - \mathbf{A}\|_F + \|\mathbf{X}_t\Theta_t\mathbf{X}_t^\top - \mathbf{A}\|_F) \right). \end{aligned} \quad (31)$$

Denote $b_t := \text{Tr}(\mathbf{I}_{r_A} - \Phi_t\Phi_t^\top)$, $a_t := \|\mathbf{X}_t\Theta_t\mathbf{X}_t^\top - \mathbf{A}\|_F$. Inequality (31) can be expressed as

$$b_{t+1} \leq qb_t + \frac{1-q}{180} (a_{t-1}^2 + a_t^2 + 2a_{t-1}a_t + \sqrt{b_t}(a_{t-1} + a_t)). \quad (32)$$

Combining Lemma 8, Lemma 9 and the RIP property of $\mathcal{M}(\cdot)$ with $\delta = \frac{c_4(r-r_A)^6}{\kappa^2m^3r^4r_A}$, we obtain

$$a_t \leq 2\sqrt{b_t} + \frac{1}{6}a_{t-1}. \quad (33)$$

Since $t_0 + 1$ is the first iteration in Phase II, we have $\text{Tr}(\mathbf{I}_{r_A} - \Phi_{t_0+1}\Phi_{t_0+1}^\top) \leq 0.5$. From Lemma 13, it follows that $\text{Tr}(\mathbf{I}_{r_A} - \Phi_{t_0+2}\Phi_{t_0+2}^\top) \leq 0.6$. Hence, $b_{t_0+1}, b_{t_0+2} \in [0, 0.6]$.

From Lemma 27, to establish the linear convergence rate of a_t , it suffices to analyze the following equality system of $\{\tilde{b}_t\}_{t=t_0+1}^\infty$ and $\{\tilde{a}_t\}_{t=t_0+1}^\infty$:

$$\tilde{b}_{t+1} = q\tilde{b}_t + \frac{1-q}{180} (\tilde{a}_{t-1}^2 + \tilde{a}_t^2 + 2\tilde{a}_{t-1}\tilde{a}_t + \sqrt{\tilde{b}_t}(\tilde{a}_{t-1} + \tilde{a}_t)),$$

$$\begin{aligned}\tilde{a}_t &= 2\sqrt{\tilde{b}_t} + \frac{1}{6}\tilde{a}_{t-1}, \quad t = t_0 + 2, t_0 + 3, \dots, \\ \tilde{a}_{t_0+1} &= a_{t_0+1}, \tilde{b}_{t_0+1} = 0.6, \tilde{b}_{t_0+2} = 0.6.\end{aligned}$$

By Lemma 8 and the RIP property of $\mathcal{M}(\cdot)$ with $\delta = \frac{c_4(r-r_A)^6}{\kappa^2 m^3 r^4 r_A}$, we derive

$$\tilde{a}_{t_0+1} = a_{t_0+1} \leq 2\sqrt{b_{t_0+1}} + \|\Delta_{t_0}\|_F \stackrel{(i)}{\leq} 2\sqrt{b_{t_0+1}} + \frac{1}{48\sqrt{r_A}} \leq 3\sqrt{\tilde{b}_{t_0+1}} \leq \frac{3\sqrt{2}}{\sqrt{1+q}}\sqrt{\tilde{b}_{t_0+2}},$$

where (i) is from inequality (27). From the update of $\tilde{a}_t$ at $t = t_0 + 2$, it follows that

$$\tilde{a}_{t_0+2} = 2\sqrt{\tilde{b}_{t_0+2}} + \frac{1}{6}\tilde{a}_{t_0+1} \leq 2\sqrt{\tilde{b}_{t_0+2}} + \frac{1}{3}\sqrt{b_{t_0+1}} + \frac{1}{288\sqrt{r_A}} \leq 3\sqrt{\tilde{b}_{t_0+2}}.$$

Therefore, applying Lemma 27 and Lemma 28, we arrive at

$$a_{t_0+1+t} \leq \tilde{a}_{t_0+1+t} \leq 3\sqrt{\tilde{b}_{t_0+1}} \left(\frac{1+q}{2}\right)^{t/2} \leq 3 \left(1 - \frac{1-q}{4}\right)^{t+1} = 3 \left(1 - \frac{c_3(r-r_A)^4}{\kappa^4 m^2 r^2 r_A}\right)^{t+1},$$

for all $t \geq 0$, with $c_3 := \frac{1}{32500c_1^2} \in (0, 1)$. This establishes the linear convergence rate of a_t .

We now prove that $\text{Tr}(\Phi_t \Phi_t^\top) \geq r_A - 0.6$ for all $t \geq t_0 + 1$. This amounts to proving that $\text{Tr}(\mathbf{I}_{r_A} - \Phi_t \Phi_t^\top) = b_t \leq 0.6$ for all $t \geq t_0 + 1$.

Since $b_{t_0+2} \leq 0.6$, inequality (32) holds for $t = t_0 + 2$. Hence,

$$\begin{aligned}b_{t_0+3} &\leq qb_{t_0+2} + \frac{1-q}{180}(a_{t_0+1}^2 + a_{t_0+2}^2 + 2a_{t_0+1}a_{t_0+2} + \sqrt{b_{t_0+2}}(a_{t_0+1} + a_{t_0+2})) \\ &\stackrel{(j)}{\leq} 0.6q + \frac{1-q}{180}(9 \times 0.6 + 9 \times 0.6 + 18 \times 0.6 + 6 \times 0.6) \\ &\leq \frac{1+q}{2} \times 0.6 \\ &\leq 0.6,\end{aligned}$$

where (j) is from the fact that $a_{t_0+1} \leq 3\sqrt{\tilde{b}_{t_0+1}} = 3\sqrt{0.6}$, $a_{t_0+2} \leq \tilde{a}_{t_0+2} \leq 3\sqrt{\tilde{b}_{t_0+2}} = 3\sqrt{0.6}$. Therefore, inequality (32) holds for $t = t_0 + 3$. From inequality (33), we have $a_{t_0+3} \leq 2\sqrt{b_{t_0+3}} + \frac{1}{6}a_{t_0+2} \leq 3\sqrt{0.6}$. By recursion, it follows that $\text{Tr}(\mathbf{I}_{r_A} - \Phi_t \Phi_t^\top) = b_t \leq 0.6$ for all $t \geq t_0 + 1$.

To conclude, by choosing stepsizes $\eta = \frac{(r-r_A)^4}{975c_1^2\kappa^2 m^2 r^2 r_A}$ and $\mu = 2$, we have that $\|\mathbf{X}_t \Theta_t \mathbf{X}_t^\top - \mathbf{A}\|_F \leq 3 \left(1 - \frac{c_3(r-r_A)^4}{\kappa^4 m^2 r^2 r_A}\right)^{t-t_0}$ for all $t \geq t_0 + 1$, with high probability over the initialization. $\blacksquare$

G.9. Proof of Lemma 3

Proof Let $\mathbf{U}\Sigma\mathbf{U}^\top$ be the compact SVD of $\mathbf{A}$, where $\mathbf{U} = [\mathbf{u}_1, \mathbf{u}_2, \dots, \mathbf{u}_{r_A}]$ and $\Sigma = \text{diag}(\lambda_1, \lambda_2, \dots, \lambda_{r_A})$, with $\lambda_1 \geq \lambda_2 \geq \dots \geq \lambda_{r_A} > 0$. Here, $\text{diag}(\lambda_1, \lambda_2, \dots, \lambda_{r_A})$ denotes the diagonal matrix whose diagonal entries are $\lambda_1, \lambda_2, \dots, \lambda_{r_A}$.

We first consider $\rho \geq 1$. From the Eckart–Young–Mirsky theorem, we have that the best rank- ρ approximation of $\mathbf{A}$ under the Frobenius norm is $\mathbf{A}_\rho = \mathbf{U}_1 \mathbf{\Sigma}_1 \mathbf{U}_1^\top$, where $\mathbf{U}_1 = [\mathbf{u}_1, \mathbf{u}_2, \dots, \mathbf{u}_\rho]$ and $\mathbf{\Sigma} = \text{diag}(\lambda_1, \lambda_2, \dots, \lambda_\rho)$, without considering the ordering of the eigenvalues.

We begin by analyzing the form of $\mathbf{X}$ and $\mathbf{\Theta}$. Since $\text{rank}(\mathbf{A}_\rho) = \text{rank}(\mathbf{U}_1) = \rho$ and $\text{range}(\mathbf{A}_\rho) \subseteq \text{range}(\mathbf{U}_1)$, it follows that $\text{range}(\mathbf{A}_\rho) = \text{range}(\mathbf{U}_1)$. Together with $\text{range}(\mathbf{A}_\rho) = \text{range}(\mathbf{X}\mathbf{\Theta}\mathbf{X}^\top) \subseteq \text{range}(\mathbf{X})$, we can obtain $\text{range}(\mathbf{U}_1) \subseteq \text{range}(\mathbf{X})$. Therefore, there exists a matrix $\mathbf{Q} \in \mathbb{R}^{r \times \rho}$, such that $\mathbf{U}_1 = \mathbf{X}\mathbf{Q}$. By the definition of $\mathbf{U}_1$, we derive that

$$\mathbf{U}_1^\top \mathbf{U}_1 = \mathbf{Q}^\top \mathbf{X}^\top \mathbf{X} \mathbf{Q} = \mathbf{Q}^\top \mathbf{Q} = \mathbf{I}_\rho,$$

which implies that $\mathbf{Q}$ is a column-orthonormal matrix.

We extend $\mathbf{Q}$ to an $r \times r$ orthogonal matrix $\tilde{\mathbf{Q}} = [\mathbf{Q}, \mathbf{P}]$. Let $\mathbf{V}_1 = \mathbf{X}\mathbf{P}$, then $[\mathbf{U}_1, \mathbf{V}_1] = [\mathbf{X}\mathbf{Q}, \mathbf{X}\mathbf{P}] = \mathbf{X}\tilde{\mathbf{Q}}$. Since $\tilde{\mathbf{Q}}^\top \mathbf{X}^\top \mathbf{X} \tilde{\mathbf{Q}} = \tilde{\mathbf{Q}}^\top \tilde{\mathbf{Q}} = \mathbf{I}_r$, then $[\mathbf{U}_1, \mathbf{V}_1]$ is also a column-orthonormal matrix, which means that

$$\mathbf{V}_1 = [\mathbf{v}_1, \mathbf{v}_2, \dots, \mathbf{v}_{r-\rho}], \text{ with } \mathbf{v}_1, \mathbf{v}_2, \dots, \mathbf{v}_{r-\rho} \in \mathbf{U}_1^\perp; \mathbf{V}_1^\top \mathbf{V}_1 = \mathbf{I}_{r-\rho}.$$

Let $\mathbf{U}_2 = [\mathbf{u}_{\rho+1}, \mathbf{u}_{\rho+2}, \dots, \mathbf{u}_{r_A}]$, and then $\mathbf{U} = [\mathbf{U}_1, \mathbf{U}_2]$. By substituting $\mathbf{U}$ and $\mathbf{X}$, we obtain

$$\begin{aligned} \mathbf{X}^\top \mathbf{U} &= \tilde{\mathbf{Q}} \begin{bmatrix} \mathbf{U}_1^\top \\ \mathbf{V}_1^\top \end{bmatrix} [\mathbf{U}_1, \mathbf{U}_2] \\ &\stackrel{(a)}{=} \tilde{\mathbf{Q}} \begin{bmatrix} \mathbf{I}_\rho & \mathbf{0} \\ \mathbf{0} & \mathbf{V}_1^\top \mathbf{U}_2 \end{bmatrix}, \end{aligned}$$

where (a) is from $\mathbf{v}_1, \mathbf{v}_2, \dots, \mathbf{v}_{r-\rho} \in \mathbf{U}_1^\perp$.

Since $\text{Tr}(\mathbf{X}^\top \mathbf{U} \mathbf{U}^\top \mathbf{X}) = \rho$, it follows that

$$\begin{aligned} \rho &= \text{Tr}(\mathbf{X}^\top \mathbf{U} \mathbf{U}^\top \mathbf{X}) \\ &= \text{Tr} \left(\tilde{\mathbf{Q}} \begin{bmatrix} \mathbf{I}_\rho & \mathbf{0} \\ \mathbf{0} & \mathbf{V}_1^\top \mathbf{U}_2 \end{bmatrix} \begin{bmatrix} \mathbf{I}_\rho & \mathbf{0} \\ \mathbf{0} & \mathbf{U}_2^\top \mathbf{V}_1 \end{bmatrix} \tilde{\mathbf{Q}}^\top \right) \\ &= \text{Tr} \left(\begin{bmatrix} \mathbf{I}_\rho & \mathbf{0} \\ \mathbf{0} & \mathbf{V}_1^\top \mathbf{U}_2 \mathbf{U}_2^\top \mathbf{V}_1 \end{bmatrix} \right) \\ &= \rho + \text{Tr}(\mathbf{V}_1^\top \mathbf{U}_2 \mathbf{U}_2^\top \mathbf{V}_1). \end{aligned}$$

After cancelling the term ρ on both sides, we obtain $\text{Tr}(\mathbf{V}_1^\top \mathbf{U}_2 \mathbf{U}_2^\top \mathbf{V}_1) = \|\mathbf{U}_2^\top \mathbf{V}_1\|_F^2 = 0$. Hence, we have that $\mathbf{U}_2^\top \mathbf{V}_1 = \mathbf{0}$, which implies that $\mathbf{v}_1, \mathbf{v}_2, \dots, \mathbf{v}_{r-\rho} \in \mathbf{U}_2^\perp$. Moreover, since $\mathbf{v}_1, \mathbf{v}_2, \dots, \mathbf{v}_{r-\rho} \in \mathbf{U}_1^\perp$ as well, we conclude that $\mathbf{v}_1, \mathbf{v}_2, \dots, \mathbf{v}_{r-\rho} \in \mathbf{U}^\perp$.

Substituting $\mathbf{X} = [\mathbf{U}_1, \mathbf{V}_1] \tilde{\mathbf{Q}}^\top$ into $\mathbf{X} \mathbf{\Theta} \mathbf{X}^\top = \mathbf{A}_\rho$, we can obtain

$$\begin{aligned} [\mathbf{U}_1, \mathbf{V}_1] \tilde{\mathbf{Q}}^\top \mathbf{\Theta} \tilde{\mathbf{Q}} [\mathbf{U}_1, \mathbf{V}_1]^\top &= \mathbf{A}_\rho \\ &= \mathbf{U}_1 \text{diag}(\lambda_1, \lambda_2, \dots, \lambda_\rho) \mathbf{U}_1^\top \\ &= [\mathbf{U}_1, \mathbf{V}_1] \text{diag}(\lambda_1, \lambda_2, \dots, \lambda_\rho, 0, \dots, 0) [\mathbf{U}_1, \mathbf{V}_1]^\top. \end{aligned}$$

Expanding both sides of the equation, together with $\mathbf{v}_1, \mathbf{v}_2, \dots, \mathbf{v}_{r-\rho} \in \mathbf{U}^\perp$, we can obtain

$$\tilde{\mathbf{Q}}^\top \boldsymbol{\Theta} \tilde{\mathbf{Q}} = \text{diag}(\lambda_1, \lambda_2, \dots, \lambda_\rho, 0, \dots, 0).$$

This implies that

$$\boldsymbol{\Theta} = \tilde{\mathbf{Q}} \text{diag}(\lambda_1, \lambda_2, \dots, \lambda_\rho, 0, \dots, 0) \tilde{\mathbf{Q}}^\top.$$

To proceed, we first verify that $(\tilde{\mathbf{X}}, \tilde{\boldsymbol{\Theta}}) := ([\mathbf{U}_1, \mathbf{V}_1], \text{diag}(\lambda_1, \lambda_2, \dots, \lambda_\rho, 0, \dots, 0))$ is indeed a saddle point and then prove $(\mathbf{X}, \boldsymbol{\Theta})$ is also a saddle point.

We compute the Euclidean gradients of f_∞ with respect to $\mathbf{X}$ and $\boldsymbol{\Theta}$ as follows:

$$\begin{aligned} \nabla_{\mathbf{X}} f_\infty(\tilde{\mathbf{X}}, \tilde{\boldsymbol{\Theta}}) &= (\tilde{\mathbf{X}} \tilde{\boldsymbol{\Theta}} \tilde{\mathbf{X}}^\top - \mathbf{A}) \tilde{\mathbf{X}} \tilde{\boldsymbol{\Theta}} = \tilde{\mathbf{X}} \tilde{\boldsymbol{\Theta}}^2 - \mathbf{A} \tilde{\mathbf{X}} \tilde{\boldsymbol{\Theta}}, \\ \nabla_{\boldsymbol{\Theta}} f_\infty(\tilde{\mathbf{X}}, \tilde{\boldsymbol{\Theta}}) &= \frac{1}{2} \tilde{\mathbf{X}}^\top (\tilde{\mathbf{X}} \tilde{\boldsymbol{\Theta}} \tilde{\mathbf{X}}^\top - \mathbf{A}) \tilde{\mathbf{X}} = \frac{1}{2} (\tilde{\boldsymbol{\Theta}} - \tilde{\mathbf{X}}^\top \mathbf{A} \tilde{\mathbf{X}}). \end{aligned}$$

By plugging the expression of $(\tilde{\mathbf{X}}, \tilde{\boldsymbol{\Theta}})$ in, we obtain

$$\begin{aligned} \nabla_{\mathbf{X}} f_\infty(\tilde{\mathbf{X}}, \tilde{\boldsymbol{\Theta}}) &= \tilde{\mathbf{X}} \tilde{\boldsymbol{\Theta}}^2 - \mathbf{A} \tilde{\mathbf{X}} \tilde{\boldsymbol{\Theta}} \\ &= [\lambda_1^2 \mathbf{u}_1, \lambda_2^2 \mathbf{u}_2, \dots, \lambda_\rho^2 \mathbf{u}_\rho, \mathbf{0}, \dots, \mathbf{0}] - [\lambda_1^2 \mathbf{u}_1, \lambda_2^2 \mathbf{u}_2, \dots, \lambda_\rho^2 \mathbf{u}_\rho, \mathbf{0}, \dots, \mathbf{0}] \\ &= \mathbf{0}, \\ \nabla_{\boldsymbol{\Theta}} f_\infty(\tilde{\mathbf{X}}, \tilde{\boldsymbol{\Theta}}) &= \frac{1}{2} (\tilde{\boldsymbol{\Theta}} - \tilde{\mathbf{X}}^\top \mathbf{A} \tilde{\mathbf{X}}) \\ &= \frac{1}{2} (\text{diag}(\lambda_1, \lambda_2, \dots, \lambda_\rho, 0, \dots, 0) - \text{diag}(\lambda_1, \lambda_2, \dots, \lambda_\rho, 0, \dots, 0)) \\ &= \mathbf{0}. \end{aligned}$$

Then, the Riemannian gradient is

$$(\mathbf{I}_m - \tilde{\mathbf{X}} \tilde{\mathbf{X}}^\top) \nabla_{\mathbf{X}} f_\infty(\tilde{\mathbf{X}}, \tilde{\boldsymbol{\Theta}}) + \frac{\tilde{\mathbf{X}}}{2} (\tilde{\mathbf{X}}^\top \nabla_{\mathbf{X}} f_\infty(\tilde{\mathbf{X}}, \tilde{\boldsymbol{\Theta}}) - \nabla_{\mathbf{X}} f_\infty(\tilde{\mathbf{X}}, \tilde{\boldsymbol{\Theta}})^\top \tilde{\mathbf{X}}) = \mathbf{0}.$$

Therefore, $(\tilde{\mathbf{X}}, \tilde{\boldsymbol{\Theta}})$ is a stationary point in the Riemannian sense.

We now show that $(\tilde{\mathbf{X}}, \tilde{\boldsymbol{\Theta}})$ is neither a local minimum nor a local maximum of the objective function.

For any $0 < \nu < \lambda_{r_A}$, we will construct a pair $(\tilde{\mathbf{X}}_+, \tilde{\boldsymbol{\Theta}}_+)$, such that $f_\infty(\tilde{\mathbf{X}}_+, \tilde{\boldsymbol{\Theta}}_+) > f_\infty(\tilde{\mathbf{X}}, \tilde{\boldsymbol{\Theta}})$, $d((\tilde{\mathbf{X}}_+, \tilde{\boldsymbol{\Theta}}_+), (\tilde{\mathbf{X}}, \tilde{\boldsymbol{\Theta}})) := \sqrt{\|\tilde{\mathbf{X}}_+ - \tilde{\mathbf{X}}\|_F^2 + \|\tilde{\boldsymbol{\Theta}}_+ - \tilde{\boldsymbol{\Theta}}\|_F^2} \leq \nu$ and $\tilde{\mathbf{X}}_+^\top \tilde{\mathbf{X}}_+ = \mathbf{I}_r$.

Let $\tilde{\mathbf{X}}_+ = \tilde{\mathbf{X}} = [\mathbf{U}_1, \mathbf{V}_1]$ and $\tilde{\boldsymbol{\Theta}}_+ = \text{diag}(\lambda_1 - \nu, \lambda_2, \dots, \lambda_\rho, 0, \dots, 0)$. By construction, $\tilde{\mathbf{X}}_+^\top \tilde{\mathbf{X}}_+ = \mathbf{I}_r$ and $d((\tilde{\mathbf{X}}_+, \tilde{\boldsymbol{\Theta}}_+), (\tilde{\mathbf{X}}, \tilde{\boldsymbol{\Theta}})) = \sqrt{\nu^2} \leq \nu$ hold. The value of the objective function is

$$\begin{aligned} f_\infty(\tilde{\mathbf{X}}_+, \tilde{\boldsymbol{\Theta}}_+) &= \frac{1}{4} \|\tilde{\mathbf{X}}_+ \tilde{\boldsymbol{\Theta}}_+ \tilde{\mathbf{X}}_+^\top - \mathbf{A}\|_F^2 \\ &= \frac{1}{4} \|(\lambda_1 - \nu) \mathbf{u}_1 \mathbf{u}_1^\top + \sum_{i=2}^{\rho} \lambda_i \mathbf{u}_i \mathbf{u}_i^\top - \sum_{i=1}^{r_A} \lambda_i \mathbf{u}_i \mathbf{u}_i^\top\|_F^2 \end{aligned}$$

$$\begin{aligned}
 &= \frac{1}{4} \|\nu \mathbf{u}_1 \mathbf{u}_1^\top + \sum_{i=\rho+1}^{r_A} \lambda_i \mathbf{u}_i \mathbf{u}_i^\top\|_{\text{F}}^2 \\
 &\stackrel{(b)}{=} \frac{1}{4} (\nu^2 \|\mathbf{u}_1 \mathbf{u}_1^\top\|_{\text{F}}^2 + \|\tilde{\mathbf{X}} \tilde{\mathbf{\Theta}} \tilde{\mathbf{X}}^\top - \mathbf{A}\|_{\text{F}}^2) \\
 &> f_\infty(\tilde{\mathbf{X}}, \tilde{\mathbf{\Theta}}),
 \end{aligned}$$

where (b) is by the orthogonality of $\{\mathbf{u}_1, \mathbf{u}_2, \dots, \mathbf{u}_{r_A}\}$.

We now try to construct a pair $(\tilde{\mathbf{X}}_-, \tilde{\mathbf{\Theta}}_-)$, such that $f_\infty(\tilde{\mathbf{X}}_-, \tilde{\mathbf{\Theta}}_-) < f_\infty(\tilde{\mathbf{X}}, \tilde{\mathbf{\Theta}})$, $d((\tilde{\mathbf{X}}_-, \tilde{\mathbf{\Theta}}_-), (\tilde{\mathbf{X}}, \tilde{\mathbf{\Theta}})) := \sqrt{\|\tilde{\mathbf{X}}_- - \tilde{\mathbf{X}}\|_{\text{F}}^2 + \|\tilde{\mathbf{\Theta}}_- - \tilde{\mathbf{\Theta}}\|_{\text{F}}^2} \leq \nu$, and $\tilde{\mathbf{X}}_-^\top \tilde{\mathbf{X}}_- = \mathbf{I}_r$.

Since $\mathbf{v}_i \in \mathbf{U}^\perp$ for any $i \in \{1, 2, \dots, r - \rho\}$, it follows that $\mathbf{v}_i \in \text{span}\{\mathbf{u}_{r_A+1}, \mathbf{u}_{r_A+2}, \dots, \mathbf{u}_m\}$. Accordingly, we consider

$$\begin{aligned}
 \tilde{\mathbf{X}}_- &= [\mathbf{U}_1, k\mathbf{v}_1 + s\mathbf{u}_{\rho+1}, \mathbf{v}_2, \dots, \mathbf{v}_{r-\rho}], \\
 \tilde{\mathbf{\Theta}}_- &= \text{diag}(\lambda_1, \lambda_2, \dots, \lambda_\rho, \nu_0, 0, \dots, 0),
 \end{aligned}$$

where $k, s, \nu_0 > 0$, $k^2 + s^2 = 1$ and k, s, ν_0 will be given later. We can easily verify that $\tilde{\mathbf{X}}_-^\top \tilde{\mathbf{X}}_- = \mathbf{I}_r$ holds. The distance is

$$\begin{aligned}
 d((\tilde{\mathbf{X}}_-, \tilde{\mathbf{\Theta}}_-), (\tilde{\mathbf{X}}, \tilde{\mathbf{\Theta}})) &= \sqrt{\|\tilde{\mathbf{X}}_- - \tilde{\mathbf{X}}\|_{\text{F}}^2 + \|\tilde{\mathbf{\Theta}}_- - \tilde{\mathbf{\Theta}}\|_{\text{F}}^2} \\
 &= \sqrt{\|(k-1)\mathbf{v}_1 + s\mathbf{u}_{\rho+1}\|^2 + \nu_0^2} \\
 &= \sqrt{(k-1)^2 + s^2 + \nu_0^2} \\
 &= \sqrt{2 - 2k + \nu_0^2}.
 \end{aligned}$$

Let $k = 1 - \frac{\nu^2}{4}$, $s = \sqrt{1 - k^2}$ and $\nu_0 \leq \frac{\nu}{2}$, then $d((\tilde{\mathbf{X}}_-, \tilde{\mathbf{\Theta}}_-), (\tilde{\mathbf{X}}, \tilde{\mathbf{\Theta}})) \leq \sqrt{\frac{\nu^2}{2} + \frac{\nu^2}{4}} \leq \nu$. The value of the objective function is

$$\begin{aligned}
 &f_\infty(\tilde{\mathbf{X}}_-, \tilde{\mathbf{\Theta}}_-) \\
 &= \frac{1}{4} \|\tilde{\mathbf{X}}_- \tilde{\mathbf{\Theta}}_- \tilde{\mathbf{X}}_-^\top - \mathbf{A}\|_{\text{F}}^2 \\
 &= \frac{1}{4} \|\nu_0(k^2 \mathbf{v}_1 \mathbf{v}_1^\top + ks \mathbf{u}_{\rho+1} \mathbf{v}_1^\top + ks \mathbf{v}_1 \mathbf{u}_{\rho+1}^\top) + (\nu_0 s^2 - \lambda_\rho) \mathbf{u}_{\rho+1} \mathbf{u}_{\rho+1}^\top - \sum_{i=\rho+2}^{r_A} \lambda_i \mathbf{u}_i \mathbf{u}_i^\top\|_{\text{F}}^2 \\
 &\stackrel{(c)}{=} \frac{1}{4} (\nu_0^2 (k^4 + k^2 s^2 \|\mathbf{u}_{\rho+1} \mathbf{v}_1^\top\|_{\text{F}}^2 + k^2 s^2 \|\mathbf{v}_1 \mathbf{u}_{\rho+1}^\top\|_{\text{F}}^2) + (\nu_0 s^2 - \lambda_\rho)^2) + f_\infty(\tilde{\mathbf{X}}, \tilde{\mathbf{\Theta}}) - \frac{1}{4} \lambda_\rho^2 \\
 &= \frac{1}{4} \nu_0^2 (k^4 + 2k^2 s^2 + s^4) - \frac{1}{2} \nu_0 \lambda_\rho s^2 + f_\infty(\tilde{\mathbf{X}}, \tilde{\mathbf{\Theta}}),
 \end{aligned}$$

where (c) is from the orthogonality of $\{\mathbf{u}_1, \mathbf{u}_2, \dots, \mathbf{u}_{r_A}, \mathbf{v}_1\}$. Let $\nu_0 > 0$ be sufficiently small. Then $\frac{1}{4} \nu_0^2 (k^4 + 2k^2 s^2 + s^4) - \frac{1}{2} \nu_0 \lambda_\rho s^2 < 0$. This ensures that the perturbed pair leads to a strictly smaller objective value, i.e., $f_\infty(\tilde{\mathbf{X}}_-, \tilde{\mathbf{\Theta}}_-) < f_\infty(\tilde{\mathbf{X}}, \tilde{\mathbf{\Theta}})$.

Therefore, we have verified that $(\tilde{\mathbf{X}}, \tilde{\boldsymbol{\Theta}})$ is a saddle point. Building upon this result, we now proceed to show that $(\mathbf{X}, \boldsymbol{\Theta}) = (\tilde{\mathbf{X}}\tilde{\mathbf{Q}}^\top, \tilde{\mathbf{Q}}\tilde{\boldsymbol{\Theta}}\tilde{\mathbf{Q}}^\top)$ is also a saddle point.

Plugging in the expression of $(\mathbf{X}, \boldsymbol{\Theta})$, we obtain the Euclidean gradients as follows:

$$\begin{aligned}\nabla_{\mathbf{X}} f_{\infty}(\mathbf{X}, \boldsymbol{\Theta}) &= (\mathbf{X}\boldsymbol{\Theta}\mathbf{X}^\top - \mathbf{A})\mathbf{X}\boldsymbol{\Theta} \\ &= (\tilde{\mathbf{X}}\tilde{\mathbf{Q}}^\top\tilde{\mathbf{Q}}\tilde{\boldsymbol{\Theta}}\tilde{\mathbf{Q}}^\top\tilde{\mathbf{Q}}\tilde{\mathbf{X}}^\top - \mathbf{A})\tilde{\mathbf{X}}\tilde{\mathbf{Q}}^\top\tilde{\mathbf{Q}}\tilde{\boldsymbol{\Theta}}\tilde{\mathbf{Q}}^\top \\ &= (\tilde{\mathbf{X}}\tilde{\boldsymbol{\Theta}}^2 - \mathbf{A}\tilde{\mathbf{X}}\tilde{\boldsymbol{\Theta}})\tilde{\mathbf{Q}}^\top \\ &= \mathbf{0}, \\ \nabla_{\boldsymbol{\Theta}} f_{\infty}(\mathbf{X}, \boldsymbol{\Theta}) &= \frac{1}{2}\mathbf{X}^\top(\mathbf{X}\boldsymbol{\Theta}\mathbf{X}^\top - \mathbf{A})\mathbf{X} \\ &= \frac{1}{2}\tilde{\mathbf{Q}}\tilde{\mathbf{X}}^\top(\tilde{\mathbf{X}}\tilde{\mathbf{Q}}^\top\tilde{\mathbf{Q}}\tilde{\boldsymbol{\Theta}}\tilde{\mathbf{Q}}^\top\tilde{\mathbf{Q}}\tilde{\mathbf{X}}^\top - \mathbf{A})\tilde{\mathbf{X}}\tilde{\mathbf{Q}}^\top \\ &= \frac{1}{2}\tilde{\mathbf{Q}}(\tilde{\boldsymbol{\Theta}} - \tilde{\mathbf{X}}^\top\mathbf{A}\tilde{\mathbf{X}})\tilde{\mathbf{Q}}^\top \\ &= \mathbf{0}.\end{aligned}$$

Then, the Riemannian gradient is

$$(\mathbf{I}_m - \mathbf{X}\mathbf{X}^\top)\nabla_{\mathbf{X}} f_{\infty}(\mathbf{X}, \boldsymbol{\Theta}) + \frac{\mathbf{X}}{2}(\mathbf{X}^\top\nabla_{\mathbf{X}} f_{\infty}(\mathbf{X}, \boldsymbol{\Theta}) - \nabla_{\mathbf{X}} f_{\infty}(\mathbf{X}, \boldsymbol{\Theta})^\top\mathbf{X}) = \mathbf{0}.$$

Therefore, $(\mathbf{X}, \boldsymbol{\Theta})$ is a stationary point in the Riemannian sense.

Let $(\mathbf{X}_+, \boldsymbol{\Theta}_+) = (\tilde{\mathbf{X}}_+\tilde{\mathbf{Q}}^\top, \tilde{\mathbf{Q}}\tilde{\boldsymbol{\Theta}}_+\tilde{\mathbf{Q}}^\top)$, $(\mathbf{X}_-, \boldsymbol{\Theta}_-) = (\tilde{\mathbf{X}}_-\tilde{\mathbf{Q}}^\top, \tilde{\mathbf{Q}}\tilde{\boldsymbol{\Theta}}_-\tilde{\mathbf{Q}}^\top)$. The distance is

$$\begin{aligned}d((\mathbf{X}_+, \boldsymbol{\Theta}_+), (\mathbf{X}, \boldsymbol{\Theta})) &= \sqrt{\|\mathbf{X}_+ - \mathbf{X}\|_{\mathbb{F}}^2 + \|\boldsymbol{\Theta}_+ - \boldsymbol{\Theta}\|_{\mathbb{F}}^2} \\ &= \sqrt{\|(\tilde{\mathbf{X}}_+ - \tilde{\mathbf{X}})\tilde{\mathbf{Q}}^\top\|_{\mathbb{F}}^2 + \|\tilde{\mathbf{Q}}(\tilde{\boldsymbol{\Theta}}_+ - \tilde{\boldsymbol{\Theta}})\tilde{\mathbf{Q}}^\top\|_{\mathbb{F}}^2} \\ &= \sqrt{\|\tilde{\mathbf{X}}_+ - \tilde{\mathbf{X}}\|_{\mathbb{F}}^2 + \|\tilde{\boldsymbol{\Theta}}_+ - \tilde{\boldsymbol{\Theta}}\|_{\mathbb{F}}^2} \\ &= d((\tilde{\mathbf{X}}_+, \tilde{\boldsymbol{\Theta}}_+), (\tilde{\mathbf{X}}, \tilde{\boldsymbol{\Theta}})).\end{aligned}$$

In the same manner, we can obtain that $d((\mathbf{X}_-, \boldsymbol{\Theta}_-), (\mathbf{X}, \boldsymbol{\Theta})) = d((\tilde{\mathbf{X}}_-, \tilde{\boldsymbol{\Theta}}_-), (\tilde{\mathbf{X}}, \tilde{\boldsymbol{\Theta}}))$. By the orthogonality of $\tilde{\mathbf{Q}}$, the following three identities hold:

$$\begin{aligned}\mathbf{X}\boldsymbol{\Theta}\mathbf{X}^\top &= \tilde{\mathbf{X}}\tilde{\mathbf{Q}}^\top\tilde{\mathbf{Q}}\tilde{\boldsymbol{\Theta}}\tilde{\mathbf{Q}}^\top\tilde{\mathbf{Q}}\tilde{\mathbf{X}}^\top = \tilde{\mathbf{X}}\tilde{\boldsymbol{\Theta}}\tilde{\mathbf{X}}^\top, \\ \mathbf{X}_+\boldsymbol{\Theta}_+\mathbf{X}_+^\top &= \tilde{\mathbf{X}}_+\tilde{\mathbf{Q}}^\top\tilde{\mathbf{Q}}\tilde{\boldsymbol{\Theta}}_+\tilde{\mathbf{Q}}^\top\tilde{\mathbf{Q}}\tilde{\mathbf{X}}_+^\top = \tilde{\mathbf{X}}_+\tilde{\boldsymbol{\Theta}}_+\tilde{\mathbf{X}}_+^\top, \\ \mathbf{X}_-\boldsymbol{\Theta}_-\mathbf{X}_-^\top &= \tilde{\mathbf{X}}_-\tilde{\mathbf{Q}}^\top\tilde{\mathbf{Q}}\tilde{\boldsymbol{\Theta}}_-\tilde{\mathbf{Q}}^\top\tilde{\mathbf{Q}}\tilde{\mathbf{X}}_-^\top = \tilde{\mathbf{X}}_-\tilde{\boldsymbol{\Theta}}_-\tilde{\mathbf{X}}_-^\top.\end{aligned}$$

Then, we have $f(\mathbf{X}, \boldsymbol{\Theta}) = f(\tilde{\mathbf{X}}, \tilde{\boldsymbol{\Theta}})$, $f(\mathbf{X}_+, \boldsymbol{\Theta}_+) = f(\tilde{\mathbf{X}}_+, \tilde{\boldsymbol{\Theta}}_+)$ and $f(\mathbf{X}_-, \boldsymbol{\Theta}_-) = f(\tilde{\mathbf{X}}_-, \tilde{\boldsymbol{\Theta}}_-)$. Thus, we obtain the strict inequality $f(\mathbf{X}_-, \boldsymbol{\Theta}_-) < f(\mathbf{X}, \boldsymbol{\Theta}) < f(\mathbf{X}_+, \boldsymbol{\Theta}_+)$. Therefore, $(\mathbf{X}, \boldsymbol{\Theta})$ is also a saddle point.

We now turn to the case $\rho = 0$, i.e., $\mathbf{X}\boldsymbol{\Theta}\mathbf{X}^\top = \mathbf{A}_0 = \mathbf{0}$. Consequently, $\boldsymbol{\Theta} = \mathbf{X}^\top\mathbf{A}_0\mathbf{X} = \mathbf{0}$. Let $\mathbf{X}$ be expressed as $\mathbf{X} = [\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_r]$, where each $\mathbf{x}_i$ is a column vector. Since $\text{Tr}(\mathbf{X}^\top\mathbf{U}\mathbf{U}^\top\mathbf{X}) = \|\mathbf{U}^\top\mathbf{X}\|_{\mathbb{F}}^2 = 0$, it follows that $\mathbf{U}^\top\mathbf{X} = \mathbf{0}$. Hence, each $\mathbf{x}_i$ lies in $\mathbf{U}^\perp$ for $i \in \{1, 2, \dots, r\}$.

We compute the Euclidean gradient of the objective function f_∞ with respect to $\mathbf{X}$ and Θ

$$\begin{aligned}\nabla_{\mathbf{X}} f_\infty(\mathbf{X}, \Theta) &= (\mathbf{X}\Theta\mathbf{X}^\top - \mathbf{A})\mathbf{X}\Theta = \mathbf{X}\Theta^2 - \mathbf{A}\mathbf{X}\Theta = \mathbf{0}, \\ \nabla_{\Theta} f_\infty(\mathbf{X}, \Theta) &= \frac{1}{2}\mathbf{X}^\top(\mathbf{X}\Theta\mathbf{X}^\top - \mathbf{A})\mathbf{X} = \frac{1}{2}(\Theta - \mathbf{X}^\top\mathbf{A}\mathbf{X}) = \mathbf{0}.\end{aligned}$$

Then, the Riemannian gradient is

$$(\mathbf{I}_m - \mathbf{X}\mathbf{X}^\top)\nabla_{\mathbf{X}} f_\infty(\mathbf{X}, \Theta) + \frac{\mathbf{X}}{2}(\mathbf{X}^\top\nabla_{\mathbf{X}} f_\infty(\mathbf{X}, \Theta) - \nabla_{\mathbf{X}} f_\infty(\mathbf{X}, \Theta)^\top\mathbf{X}) = \mathbf{0}.$$

Therefore, $(\mathbf{X}, \Theta)$ is a stationary point in the Riemannian sense.

For any $0 < \nu < \lambda_{r_A}$, we construct the pair $(\mathbf{X}_+, \Theta_+)$ as follows:

$$\begin{aligned}\mathbf{X}_+ &= [k\mathbf{x}_1 + s\mathbf{u}_1, \mathbf{x}_2, \dots, \mathbf{x}_r], \\ \Theta_+ &= \text{diag}(-\nu_1, 0, \dots, 0),\end{aligned}$$

where $k = 1 - \frac{\nu^2}{4}$, $s = \sqrt{1 - k^2}$, and $0 < \nu_1 \leq \frac{\nu}{2}$. We can easily verify that $\mathbf{X}_+^\top\mathbf{X}_+ = \mathbf{I}_r$ and the distance is

$$\begin{aligned}d((\mathbf{X}_+, \Theta_+), (\mathbf{X}, \Theta)) &= \sqrt{\|\mathbf{X}_+ - \mathbf{X}\|_F^2 + \|\Theta_+ - \Theta\|_F^2} \\ &= \sqrt{\|(k-1)\mathbf{x}_1 + s\mathbf{u}_1\|^2 + \nu_1^2} \\ &= \sqrt{(k-1)^2 + s^2 + \nu_1^2} \\ &\leq \sqrt{\frac{\nu^2}{2} + \frac{\nu^2}{4}} \\ &\leq \nu.\end{aligned}$$

The value of the objective function is

$$\begin{aligned}f_\infty(\mathbf{X}_+, \Theta_+) &= \frac{1}{4}\|\mathbf{X}_+\Theta_+\mathbf{X}_+^\top - \mathbf{A}\|_F^2 \\ &= \frac{1}{4}\|-\nu_1(k^2\mathbf{x}_1\mathbf{x}_1^\top + s^2\mathbf{u}_1\mathbf{u}_1^\top) - \sum_{i=1}^{r_A}\lambda_i\mathbf{u}_i\mathbf{u}_i^\top\|_F^2 \\ &= \frac{1}{4}\|\nu_1k^2\mathbf{x}_1\mathbf{x}_1^\top + \nu_1s^2\mathbf{u}_1\mathbf{u}_1^\top + \sum_{i=1}^{r_A}\lambda_i\mathbf{u}_i\mathbf{u}_i^\top\|_F^2 \\ &\stackrel{(d)}{=} \frac{1}{4}\left(\nu_1^2k^4 + \nu_1^2s^4 + 2\nu_1\lambda_1s^2 + \|\mathbf{X}\Theta\mathbf{X}^\top - \mathbf{A}\|_F^2\right) \\ &> f_\infty(\mathbf{X}, \Theta),\end{aligned}$$

where (d) is due to the orthogonality of $\{\mathbf{u}_1, \mathbf{u}_2, \dots, \mathbf{u}_{r_A}, \mathbf{x}_1\}$. Now consider the pair $(\mathbf{X}_-, \Theta_-)$ defined as:

$$\mathbf{X}_- = [k\mathbf{x}_1 + s\mathbf{u}_1, \mathbf{x}_2, \dots, \mathbf{x}_r],$$

$$\Theta_- = \text{diag}(\nu_2, 0, \dots, 0),$$

where $k = 1 - \frac{\nu^2}{4}$, $s = \sqrt{1 - k^2}$, and $0 < \nu_2 \leq \frac{\nu}{2}$. It can be verified that $\mathbf{X}_-^\top \mathbf{X}_- = \mathbf{I}_r$, and the distance is

$$\begin{aligned} d((\mathbf{X}_-, \Theta_-), (\mathbf{X}, \Theta)) &= \sqrt{\|\mathbf{X}_- - \mathbf{X}\|_F^2 + \|\Theta_- - \Theta\|_F^2} \\ &= \sqrt{\|(k-1)\mathbf{x}_1 + s\mathbf{u}_1\|^2 + \nu_2^2} \\ &= \sqrt{(k-1)^2 + s^2 + \nu_2^2} \\ &\leq \sqrt{\frac{\nu^2}{2} + \frac{\nu^2}{4}} \\ &\leq \nu. \end{aligned}$$

The value of the objective function is

$$\begin{aligned} f_\infty(\mathbf{X}_-, \Theta_-) &= \frac{1}{4} \|\mathbf{X}_- \Theta_- \mathbf{X}_-^\top - \mathbf{A}\|_F^2 \\ &= \frac{1}{4} \|\nu_2 (k^2 \mathbf{x}_1 \mathbf{x}_1^\top + s^2 \mathbf{u}_1 \mathbf{u}_1^\top) - \sum_{i=1}^{r_A} \lambda_i \mathbf{u}_i \mathbf{u}_i^\top\|_F^2 \\ &= \frac{1}{4} \|\nu_2 k^2 \mathbf{x}_1 \mathbf{x}_1^\top + \nu_2 s^2 \mathbf{u}_1 \mathbf{u}_1^\top - \sum_{i=1}^{r_A} \lambda_i \mathbf{u}_i \mathbf{u}_i^\top\|_F^2 \\ &\stackrel{(e)}{=} \frac{1}{4} (\nu_2^2 k^4 + \nu_2^2 s^4 - 2\nu_2 \lambda_1 s^2 + \|\mathbf{X} \Theta \mathbf{X}^\top - \mathbf{A}\|_F^2) \\ &= \frac{1}{4} (\nu_2^2 (k^4 + s^4) - 2\nu_2 \lambda_1 s^2) + f_\infty(\mathbf{X}, \Theta), \end{aligned}$$

where (e) is by the orthogonality of $\{\mathbf{u}_1, \mathbf{u}_2, \dots, \mathbf{u}_{r_A}, \mathbf{x}_1\}$. Let $\nu_2 > 0$ be sufficiently small. Then $\frac{1}{4} (\nu_2^2 (k^4 + s^4) - 2\nu_2 \lambda_1 s^2) < 0$. This guarantees that $f_\infty(\mathbf{X}_-, \Theta_-) < f_\infty(\mathbf{X}, \Theta)$. Therefore, $(\mathbf{X}, \Theta)$ is also a saddle point when $\rho = 0$. $\blacksquare$

G.10. Proof of Lemma 4

Proof We begin by computing the Euclidean gradients of f_∞ and f with respect to $\mathbf{X}$ and Θ :

$$\begin{aligned} \nabla_{\mathbf{X}} f_\infty(\mathbf{X}, \Theta) &= (\mathbf{X} \Theta \mathbf{X}^\top - \mathbf{A}) \mathbf{X} \Theta, \\ \nabla_{\Theta} f_\infty(\mathbf{X}, \Theta) &= \frac{1}{2} \mathbf{X}^\top (\mathbf{X} \Theta \mathbf{X}^\top - \mathbf{A}) \mathbf{X}, \\ \nabla_{\mathbf{X}} f(\mathbf{X}, \Theta) &= \mathcal{M}^* \mathcal{M} (\mathbf{X} \Theta \mathbf{X}^\top - \mathbf{A}) \mathbf{X} \Theta, \\ \nabla_{\Theta} f(\mathbf{X}, \Theta) &= \frac{1}{2} \mathbf{X}^\top \mathcal{M}^* \mathcal{M} (\mathbf{X} \Theta \mathbf{X}^\top - \mathbf{A}) \mathbf{X}. \end{aligned}$$

Then, we can obtain that the gap between population gradient and sensing gradient is

$$\|\nabla_{\mathbf{X}} f_\infty(\mathbf{X}, \Theta) - \nabla_{\mathbf{X}} f(\mathbf{X}, \Theta)\|_F = \|(\mathcal{M}^* \mathcal{M} - \mathcal{I})(\mathbf{X} \Theta \mathbf{X}^\top - \mathbf{A}) \mathbf{X} \Theta\|_F$$

$$\begin{aligned}
 &\leq \|(\mathcal{M}^* \mathcal{M} - \mathcal{I})(\mathbf{X} \boldsymbol{\Theta} \mathbf{X}^\top - \mathbf{A})\|_F \|\mathbf{X}\| \|\boldsymbol{\Theta}\| \\
 &\stackrel{(f)}{\leq} 2 \|(\mathcal{M}^* \mathcal{M} - \mathcal{I})(\mathbf{X} \boldsymbol{\Theta} \mathbf{X}^\top - \mathbf{A})\|_F \\
 &\leq 2\sqrt{m} \|(\mathcal{M}^* \mathcal{M} - \mathcal{I})(\mathbf{X} \boldsymbol{\Theta} \mathbf{X}^\top - \mathbf{A})\| \\
 &\stackrel{(g)}{\leq} 2\sqrt{m}\delta \|\mathbf{X} \boldsymbol{\Theta} \mathbf{X}^\top - \mathbf{A}\|_F \\
 &\leq 2m\delta \|\mathbf{X} \boldsymbol{\Theta} \mathbf{X}^\top - \mathbf{A}\| \\
 &\leq 2m\delta (\|\mathbf{X}\| \|\boldsymbol{\Theta}\| \|\mathbf{X}\| + \|\mathbf{A}\|) \\
 &\stackrel{(f)}{\leq} 6m\delta,
 \end{aligned}$$

$$\begin{aligned}
 \|\nabla_{\boldsymbol{\Theta}} f_\infty(\mathbf{X}, \boldsymbol{\Theta}) - \nabla_{\boldsymbol{\Theta}} f(\mathbf{X}, \boldsymbol{\Theta})\|_F &= \frac{1}{2} \|\mathbf{X}^\top (\mathcal{M}^* \mathcal{M} - \mathcal{I})(\mathbf{X} \boldsymbol{\Theta} \mathbf{X}^\top - \mathbf{A}) \mathbf{X}\|_F \\
 &\leq \frac{1}{2} \|\mathbf{X}\| \|(\mathcal{M}^* \mathcal{M} - \mathcal{I})(\mathbf{X} \boldsymbol{\Theta} \mathbf{X}^\top - \mathbf{A})\|_F \|\mathbf{X}\| \\
 &\leq \frac{1}{2} \|(\mathcal{M}^* \mathcal{M} - \mathcal{I})(\mathbf{X} \boldsymbol{\Theta} \mathbf{X}^\top - \mathbf{A})\|_F \\
 &\leq \frac{1}{2} \sqrt{m} \|(\mathcal{M}^* \mathcal{M} - \mathcal{I})(\mathbf{X} \boldsymbol{\Theta} \mathbf{X}^\top - \mathbf{A})\| \\
 &\stackrel{(g)}{\leq} \frac{1}{2} \sqrt{m}\delta \|\mathbf{X} \boldsymbol{\Theta} \mathbf{X}^\top - \mathbf{A}\|_F \\
 &\leq \frac{1}{2} m\delta \|\mathbf{X} \boldsymbol{\Theta} \mathbf{X}^\top - \mathbf{A}\| \\
 &\leq \frac{1}{2} m\delta (\|\mathbf{X}\| \|\boldsymbol{\Theta}\| \|\mathbf{X}\| + \|\mathbf{A}\|) \\
 &\stackrel{(f)}{\leq} \frac{3}{2} m\delta,
 \end{aligned}$$

where (f) is by $\|\mathbf{X}\| \leq 1, \|\boldsymbol{\Theta}\| \leq 2, \|\mathbf{A}\| \leq 1$; and (g) is from Lemma 24. Then, the difference between the two Riemannian gradients can be bounded as

$$\begin{aligned}
 \|\nabla_{\mathbf{X}}^R f_\infty(\mathbf{X}, \boldsymbol{\Theta}) - \nabla_{\mathbf{X}}^R f(\mathbf{X}, \boldsymbol{\Theta})\|_F &= \|(\mathbf{I}_m - \mathbf{X} \mathbf{X}^\top)(\nabla_{\mathbf{X}} f_\infty(\mathbf{X}, \boldsymbol{\Theta}) - \nabla_{\mathbf{X}} f(\mathbf{X}, \boldsymbol{\Theta})) \\
 &\quad + \frac{1}{2} \mathbf{X} \mathbf{X}^\top (\nabla_{\mathbf{X}} f_\infty(\mathbf{X}, \boldsymbol{\Theta}) - \nabla_{\mathbf{X}} f(\mathbf{X}, \boldsymbol{\Theta})) \\
 &\quad - \frac{1}{2} \mathbf{X} (\nabla_{\mathbf{X}} f_\infty(\mathbf{X}, \boldsymbol{\Theta})^\top - \nabla_{\mathbf{X}} f(\mathbf{X}, \boldsymbol{\Theta})^\top) \mathbf{X}\|_F \\
 &\leq \|(\mathbf{I}_m - \mathbf{X} \mathbf{X}^\top)(\nabla_{\mathbf{X}} f_\infty(\mathbf{X}, \boldsymbol{\Theta}) - \nabla_{\mathbf{X}} f(\mathbf{X}, \boldsymbol{\Theta}))\|_F \\
 &\quad + \frac{1}{2} \|\mathbf{X} \mathbf{X}^\top (\nabla_{\mathbf{X}} f_\infty(\mathbf{X}, \boldsymbol{\Theta}) - \nabla_{\mathbf{X}} f(\mathbf{X}, \boldsymbol{\Theta}))\|_F \\
 &\quad + \frac{1}{2} \|\mathbf{X} (\nabla_{\mathbf{X}} f_\infty(\mathbf{X}, \boldsymbol{\Theta})^\top - \nabla_{\mathbf{X}} f(\mathbf{X}, \boldsymbol{\Theta})^\top) \mathbf{X}\|_F \\
 &\leq 6m\delta (\|\mathbf{I}_m - \mathbf{X} \mathbf{X}^\top\| + \frac{1}{2} \|\mathbf{X}\| \|\mathbf{X}\| + \frac{1}{2} \|\mathbf{X}\| \|\mathbf{X}\|) \\
 &\stackrel{(h)}{\leq} 12m\delta,
 \end{aligned}$$

where (h) is due to $\|\mathbf{I}_m - \mathbf{X}\mathbf{X}^\top\|, \|\mathbf{X}\| \leq 1$. ■

Appendix H. Other useful lemmas

Lemma 14 *Given a PSD matrix $\mathbf{A}$, we have that $(\mathbf{I} + \mathbf{A})^{-1} \succeq \mathbf{I} - \mathbf{A}$.*

Proof Diagonalizing both sides and using $1/(1 + \lambda) \geq 1 - \lambda, \forall \lambda \geq 0$ yields the result. ■

Lemma 15 *Let $\mathbf{X} \in \text{St}(m, r)$ and $\mathbf{U} \in \text{St}(m, r_A)$. Let $\mathbf{U}_\perp \in \mathbb{R}^{m \times (m-r_A)}$ be an orthonormal basis for the orthogonal complement of $\text{span}(\mathbf{U})$. Denote $\Phi = \mathbf{U}^\top \mathbf{X} \in \mathbb{R}^{r_A \times r}$ and $\Psi = \mathbf{U}_\perp^\top \mathbf{X} \in \mathbb{R}^{(m-r_A) \times r}$. It is guaranteed that $\sigma_i^2(\Phi) + \sigma_i^2(\Psi) = 1$ holds for $i \in \{1, 2, \dots, r\}$.*

Proof Since $\mathbf{X}$ lies in the Stiefel manifold, we have that

$$\begin{aligned} \mathbf{I}_r &= \mathbf{X}^\top \mathbf{X} = \mathbf{X}^\top \mathbf{I}_m \mathbf{X} = \mathbf{X}^\top [\mathbf{U}, \mathbf{U}_\perp] \begin{bmatrix} \mathbf{U}^\top \\ \mathbf{U}_\perp^\top \end{bmatrix} \mathbf{X} \\ &= \Phi^\top \Phi + \Psi^\top \Psi. \end{aligned} \tag{34}$$

Equation (34) shows that $\Psi^\top \Psi$ and $\Phi^\top \Phi$ commute, i.e.,

$$\begin{aligned} (\Phi^\top \Phi)(\Psi^\top \Psi) &= (\Phi^\top \Phi)(\mathbf{I}_r - \Phi^\top \Phi) = \Phi^\top \Phi - \Phi^\top \Phi \Phi^\top \Phi \\ &= (\mathbf{I}_r - \Phi^\top \Phi)(\Phi^\top \Phi) = (\Psi^\top \Psi)(\Phi^\top \Phi). \end{aligned}$$

The commutativity shows that the eigenspaces of $\Phi^\top \Phi$ and $\Psi^\top \Psi$ coincide. As a result, we have again from (34) that $\sigma_i^2(\Phi) + \sigma_i^2(\Psi) = 1$ for $i \in \{1, 2, \dots, r\}$. ■

Lemma 16 *Suppose that $\mathbf{P}$ and $\mathbf{Q}$ are $m \times m$ diagonal matrices, with non-negative diagonal entries. Let $\mathbf{S} \in \mathbb{S}^m$ be a positive definite matrix with smallest eigenvalue $\lambda_{\min}$, then we have that*

$$\text{Tr}(\mathbf{P}\mathbf{S}\mathbf{Q}) \geq \lambda_{\min} \text{Tr}(\mathbf{P}\mathbf{Q}).$$

Proof Let p_i and q_i be the (i, i) -th entry of $\mathbf{P}$ and $\mathbf{Q}$, respectively. Then we have that

$$\text{Tr}(\mathbf{P}\mathbf{S}\mathbf{Q}) = \sum_i p_i \mathbf{S}_{i,i} q_i \geq \lambda_{\min} \sum_i p_i q_i = \lambda_{\min} \text{Tr}(\mathbf{P}\mathbf{Q}),$$

where the last inequality comes from $\mathbf{S}$ being positive definite, i.e., $\mathbf{S}_{i,i} = \mathbf{e}_i^\top \mathbf{S} \mathbf{e}_i \geq \lambda_{\min}$. ■

Lemma 17 *Let $\mathbf{A} \in \mathbb{R}^{m \times n}$ be a matrix with full column rank and $\mathbf{B} \in \mathbb{R}^{n \times p}$ be a non-zero matrix. Let $\sigma_{\min}(\cdot)$ denote the smallest non-zero singular value. Then it holds that $\sigma_{\min}(\mathbf{A}\mathbf{B}) \geq \sigma_{\min}(\mathbf{A})\sigma_{\min}(\mathbf{B})$.*

Proof Using the min-max principle for singular values,

$$\begin{aligned}
 \sigma_{\min}(\mathbf{AB}) &= \min_{\|\mathbf{x}\|=1, \mathbf{x} \in \text{ColSpan}(\mathbf{B})} \|\mathbf{ABx}\| \\
 &= \min_{\|\mathbf{x}\|=1, \mathbf{x} \in \text{ColSpan}(\mathbf{B})} \left\| \mathbf{A} \frac{\mathbf{Bx}}{\|\mathbf{Bx}\|} \right\| \cdot \|\mathbf{Bx}\| \\
 &\stackrel{(a)}{=} \min_{\|\mathbf{x}\|=1, \|\mathbf{y}\|=1, \mathbf{x} \in \text{ColSpan}(\mathbf{B}), \mathbf{y} \in \text{ColSpan}(\mathbf{B})} \|\mathbf{Ay}\| \cdot \|\mathbf{Bx}\| \\
 &\geq \min_{\|\mathbf{y}\|=1, \mathbf{y} \in \text{ColSpan}(\mathbf{B})} \|\mathbf{Ay}\| \cdot \min_{\|\mathbf{x}\|=1, \mathbf{x} \in \text{ColSpan}(\mathbf{B})} \|\mathbf{Bx}\| \\
 &\geq \min_{\|\mathbf{y}\|=1} \|\mathbf{Ay}\| \cdot \min_{\|\mathbf{x}\|=1, \mathbf{x} \in \text{ColSpan}(\mathbf{B})} \|\mathbf{Bx}\| \\
 &= \sigma_{\min}(\mathbf{A}) \sigma_{\min}(\mathbf{B}),
 \end{aligned}$$

where (a) is by changing of variables, i.e., $\mathbf{y} = \mathbf{Bx}/\|\mathbf{Bx}\|$. ■

Lemma 18 (Theorem 2.2.1 of [12]) *If $\mathbf{Z} \in \mathbb{R}^{m \times r}$ has entries drawn i.i.d. from Gaussian distribution $\mathcal{N}(0, 1)$, then $\mathbf{X} = \mathbf{Z}(\mathbf{Z}^\top \mathbf{Z})^{-1/2}$ is a random matrix uniformly distributed on $\text{St}(m, r)$.*

Lemma 19 [42] *If $\mathbf{Z} \in \mathbb{R}^{m \times r}$ is a matrix whose entries are independently drawn from $\mathcal{N}(0, 1)$. Then for every $\tau \geq 0$, with probability at least $1 - \exp(-\tau^2/2)$, we have*

$$\sigma_1(\mathbf{Z}) \leq \sqrt{m} + \sqrt{r} + \tau.$$

Lemma 20 [33] *If $\mathbf{Z} \in \mathbb{R}^{m \times r}$ is a matrix whose entries are independently drawn from $\mathcal{N}(0, 1)$. Suppose that $m \geq r$. Then for every $\tau \geq 0$, we have for two universal constants $C_1 > 0$ and $C_2 > 0$ that*

$$\mathbb{P}\left(\sigma_r(\mathbf{Z}) \leq \tau(\sqrt{m} - \sqrt{r-1})\right) \leq (C_1 \tau)^{m-r+1} + \exp(-C_2 m).$$

Lemma 21 *If $\mathbf{U} \in \text{St}(m, r_A)$ is a fixed matrix, $\mathbf{X} \in \text{St}(m, r)$ is uniformly sampled from $\text{St}(m, r)$ using methods described in Lemma 18, and $r > r_A$, then we have that with probability at least $1 - \exp(-m/2) - (C_1 \tau)^{r-r_A+1} - \exp(-C_2 r)$,*

$$\sigma_{r_A}(\mathbf{U}^\top \mathbf{X}) \geq \frac{\tau(r - r_A + 1)}{6\sqrt{mr}}.$$

Proof Since $\mathbf{X} \in \text{St}(m, r)$ is uniformly sampled from $\text{St}(m, r)$ using methods described in Lemma 18, we can write $\mathbf{X} = \mathbf{Z}(\mathbf{Z}^\top \mathbf{Z})^{-1/2}$, where $\mathbf{Z} \in \mathbb{R}^{m \times r}$ has entries i.i.d. sampled from $\mathcal{N}(0, 1)$. We thus have

$$\sigma_{r_A}(\mathbf{U}^\top \mathbf{X}) = \sigma_{r_A}(\mathbf{U}^\top \mathbf{Z}(\mathbf{Z}^\top \mathbf{Z})^{-1/2}).$$

We now consider $\mathbf{U}^\top \mathbf{Z} \in \mathbb{R}^{r_A \times r}$. It is clear that the entries of $\mathbf{U}^\top \mathbf{Z}$ are also i.i.d $\mathcal{N}(0, 1)$ random variables. As a consequence of Lemma 20, we have that with probability at least $1 - (C_1 \tau)^{r-r_A+1} - \exp(-C_2 r)$,

$$\sigma_{r_A}(\mathbf{U}^\top \mathbf{Z}) \geq \tau(\sqrt{r} - \sqrt{r_A - 1}).$$

We also have from Lemma 19 that with probability at least $1 - \exp(-m/2)$,

$$\sigma_1(\mathbf{Z}^\top \mathbf{Z}) = \sigma_1^2(\mathbf{Z}) \leq (2\sqrt{m} + \sqrt{r})^2.$$

Taking union bound, we have with probability at least $1 - \exp(-m/2) - (C_1\tau)^{r-r_A+1} - \exp(-C_2r)$,

$$\sigma_{r_A}(\mathbf{U}^\top \mathbf{X}) \stackrel{(a)}{\geq} \frac{\sigma_{r_A}(\mathbf{U}^\top \mathbf{Z})}{\sigma_1(\mathbf{Z})} = \frac{\tau(\sqrt{r} - \sqrt{r_A - 1})}{2\sqrt{m} + \sqrt{r}} \geq \frac{\tau(r - r_A + 1)}{3\sqrt{m} \cdot 2\sqrt{r}} = \frac{\tau(r - r_A + 1)}{6\sqrt{mr}},$$

where (a) comes from Lemma 17. ■

Lemma 22 Suppose $\Theta_t \in \mathbb{S}^r$. Then the update rule (7) guarantees that Θ_{t+1} also belongs to $\mathbb{S}^r$.

Proof From the update rule, we have that

$$\Theta_{t+1} = \mathbf{X}_{t+1}^\top \mathbf{A} \mathbf{X}_{t+1} - \mathbf{X}_{t+1}^\top \left[(\mathcal{M}^* \mathcal{M} - \frac{\mu}{2} \mathcal{I})(\mathbf{X}_{t+1} \Theta_t \mathbf{X}_{t+1}^\top - \mathbf{A}) \right] \mathbf{X}_{t+1}.$$

Since $\Theta_t \in \mathbb{S}^r$ and $\mathbf{A} \in \mathbb{S}^m$, it follows that $\mathbf{X}_{t+1} \Theta_t \mathbf{X}_{t+1}^\top - \mathbf{A} \in \mathbb{S}^m$ and $\mathbf{X}_{t+1}^\top \mathbf{A} \mathbf{X}_{t+1} \in \mathbb{S}^r$.

By definition of $\mathcal{M}$ and $\mathcal{M}^*$, the composition $\mathcal{M}^* \mathcal{M}$ defines a self-adjoint operator in $\mathbb{S}^m$. Hence,

$$\mathbf{X}_{t+1}^\top \left[(\mathcal{M}^* \mathcal{M} - \frac{\mu}{2} \mathcal{I})(\mathbf{X}_{t+1} \Theta_t \mathbf{X}_{t+1}^\top - \mathbf{A}) \right] \mathbf{X}_{t+1} \in \mathbb{S}^r.$$

Thus, $\Theta_{t+1} \in \mathbb{S}^r$, which completes the proof. ■

Lemma 23 Let $\mathcal{M}(\cdot) : \mathbb{S}^m \rightarrow \mathbb{R}^n$ be a linear mapping that is $(r + r', \delta)$ -RIP with $\delta \in [0, 1]$. Then for any symmetric matrix $\mathbf{Z}$ of rank at most r and any symmetric matrix $\mathbf{Y}$ of rank at most r' , we have that

$$|\langle (\mathcal{M}^* \mathcal{M} - \mathcal{I})(\mathbf{Z}), \mathbf{Y} \rangle| \leq \delta \|\mathbf{Z}\|_F \|\mathbf{Y}\|_F.$$

Proof Denote $\Delta(\mathbf{Z}, \mathbf{Y}) := \langle (\mathcal{M}^* \mathcal{M} - \mathcal{I})(\mathbf{Z}), \mathbf{Y} \rangle = \langle \mathcal{M}(\mathbf{Z}), \mathcal{M}(\mathbf{Y}) \rangle - \langle \mathbf{Z}, \mathbf{Y} \rangle$. The above inequality trivially holds when $\|\mathbf{Z}\|_F = 0$ or $\|\mathbf{Y}\|_F = 0$. Without loss of generality, we assume that $\|\mathbf{Z}\|_F \neq 0$ and $\|\mathbf{Y}\|_F \neq 0$. Define $\tilde{\mathbf{Z}} := \frac{\mathbf{Z}}{\|\mathbf{Z}\|_F}$ and $\tilde{\mathbf{Y}} := \frac{\mathbf{Y}}{\|\mathbf{Y}\|_F}$. It then follows that

$$\Delta(\mathbf{Z}, \mathbf{Y}) = \Delta(\tilde{\mathbf{Z}}, \tilde{\mathbf{Y}}) \cdot \|\mathbf{Z}\|_F \|\mathbf{Y}\|_F.$$

Using the polarization identity, we obtain

$$\begin{aligned} \langle \mathcal{M}(\tilde{\mathbf{Z}}), \mathcal{M}(\tilde{\mathbf{Y}}) \rangle &= \frac{1}{4} (\|\mathcal{M}(\tilde{\mathbf{Z}} + \tilde{\mathbf{Y}})\|^2 - \|\mathcal{M}(\tilde{\mathbf{Z}} - \tilde{\mathbf{Y}})\|^2), \\ \langle \tilde{\mathbf{Z}}, \tilde{\mathbf{Y}} \rangle &= \frac{1}{4} (\|\tilde{\mathbf{Z}} + \tilde{\mathbf{Y}}\|_F^2 - \|\tilde{\mathbf{Z}} - \tilde{\mathbf{Y}}\|_F^2). \end{aligned}$$

Substituting the two equalities into the expression of $\Delta(\tilde{\mathbf{Z}}, \tilde{\mathbf{Y}})$, we have that

$$\begin{aligned}
 |\Delta(\tilde{\mathbf{Z}}, \tilde{\mathbf{Y}})| &= |\langle \mathcal{M}(\tilde{\mathbf{Z}}), \mathcal{M}(\tilde{\mathbf{Y}}) \rangle - \langle \tilde{\mathbf{Z}}, \tilde{\mathbf{Y}} \rangle| \\
 &= \frac{1}{4} |(\|\mathcal{M}(\tilde{\mathbf{Z}} + \tilde{\mathbf{Y}})\|^2 - \|\mathcal{M}(\tilde{\mathbf{Z}} - \tilde{\mathbf{Y}})\|^2) - (\|\tilde{\mathbf{Z}} + \tilde{\mathbf{Y}}\|_F^2 - \|\tilde{\mathbf{Z}} - \tilde{\mathbf{Y}}\|_F^2)| \\
 &\leq \frac{1}{4} (|\|\mathcal{M}(\tilde{\mathbf{Z}} + \tilde{\mathbf{Y}})\|^2 - \|\tilde{\mathbf{Z}} + \tilde{\mathbf{Y}}\|_F^2| + |\|\mathcal{M}(\tilde{\mathbf{Z}} - \tilde{\mathbf{Y}})\|^2 - \|\tilde{\mathbf{Z}} - \tilde{\mathbf{Y}}\|_F^2|) \\
 &\stackrel{(a)}{\leq} \frac{\delta}{4} (\|\tilde{\mathbf{Z}} + \tilde{\mathbf{Y}}\|_F^2 + \|\tilde{\mathbf{Z}} - \tilde{\mathbf{Y}}\|_F^2) \\
 &= \frac{\delta}{2} (\|\tilde{\mathbf{Z}}\|_F^2 + \|\tilde{\mathbf{Y}}\|_F^2) \\
 &= \delta,
 \end{aligned}$$

where (a) is from the facts that $\mathcal{M}(\cdot)$ is $(r + r', \delta)$ -RIP with constant δ , $\text{rank}(\tilde{\mathbf{Z}} + \tilde{\mathbf{Y}}) \leq \text{rank}(\tilde{\mathbf{Z}}) + \text{rank}(\tilde{\mathbf{Y}}) \leq r + r'$, and $\text{rank}(\tilde{\mathbf{Z}} - \tilde{\mathbf{Y}}) \leq \text{rank}(\tilde{\mathbf{Z}}) + \text{rank}(\tilde{\mathbf{Y}}) \leq r + r'$. Therefore, we have that

$$|\langle (\mathcal{M}^* \mathcal{M} - \mathcal{I})(\mathbf{Z}), \mathbf{Y} \rangle| = |\Delta(\mathbf{Z}, \mathbf{Y})| = |\Delta(\tilde{\mathbf{Z}}, \tilde{\mathbf{Y}})| \cdot \|\mathbf{Z}\|_F \|\mathbf{Y}\|_F \leq \delta \|\mathbf{Z}\|_F \|\mathbf{Y}\|_F,$$

which completes the proof. ■

Lemma 24 (Lemma 7.3 of [38]) *Let $\mathcal{M}(\cdot) : \mathbb{S}^m \rightarrow \mathbb{R}^n$ be a linear mapping that is $(r + r_A + 1, \delta)$ -RIP with $\delta \in [0, 1)$, then $\|(\mathcal{M}^* \mathcal{M} - \mathcal{I})(\mathbf{A})\| \leq \delta \|\mathbf{A}\|_F$ for all matrices $\mathbf{A} \in \mathbb{S}^m$ of rank at most $r + r_A$.*

Proof By Lemma 23, if $\mathbf{A} \in \mathbb{S}^m$ has rank at most $r + r_A$ and $\mathbf{Y} \in \mathbb{S}^m$ has rank at most 1, then it holds that

$$|\langle (\mathcal{M}^* \mathcal{M} - \mathcal{I})(\mathbf{A}), \mathbf{Y} \rangle| \leq \delta \|\mathbf{A}\|_F \|\mathbf{Y}\|_F.$$

Hence, it suffices to prove that there exists a matrix $\mathbf{Y}$ of rank 1, such that $|\langle (\mathcal{M}^* \mathcal{M} - \mathcal{I})(\mathbf{A}), \mathbf{Y} \rangle| = \|(\mathcal{M}^* \mathcal{M} - \mathcal{I})(\mathbf{A})\|$ and $\|\mathbf{Y}\|_F \leq 1$. Since $(\mathcal{M}^* \mathcal{M} - \mathcal{I})(\mathbf{A})$ is a symmetric matrix, it follows that

$$\begin{aligned}
 \|(\mathcal{M}^* \mathcal{M} - \mathcal{I})(\mathbf{A})\| &= \max_{\|\mathbf{u}\|=1} \mathbf{u}^\top (\mathcal{M}^* \mathcal{M} - \mathcal{I})(\mathbf{A}) \mathbf{u} \\
 &= \max_{\|\mathbf{u}\|=1} \text{Tr}((\mathcal{M}^* \mathcal{M} - \mathcal{I})(\mathbf{A}) \mathbf{u} \mathbf{u}^\top) \\
 &= \max_{\|\mathbf{u}\|=1} \langle (\mathcal{M}^* \mathcal{M} - \mathcal{I})(\mathbf{A}), \mathbf{u} \mathbf{u}^\top \rangle.
 \end{aligned}$$

Let $\mathbf{Y} = \tilde{\mathbf{u}} \tilde{\mathbf{u}}^\top$, where $\tilde{\mathbf{u}} \in \arg \max_{\|\mathbf{u}\|=1} \langle (\mathcal{M}^* \mathcal{M} - \mathcal{I})(\mathbf{A}), \mathbf{u} \mathbf{u}^\top \rangle$. We then have that $\text{rank}(\mathbf{Y}) = 1$, $\mathbf{Y} \in \mathbb{S}^m$, $|\langle (\mathcal{M}^* \mathcal{M} - \mathcal{I})(\mathbf{A}), \mathbf{Y} \rangle| = \|(\mathcal{M}^* \mathcal{M} - \mathcal{I})(\mathbf{A})\|$, and $\|\mathbf{Y}\|_F \leq 1$. ■

Lemma 25 *Let $\mathbf{A} \in \mathbb{R}^{n \times m}$, $\mathbf{B} \in \mathbb{R}^{m \times n}$ be two real matrices, then the following inequality holds*

$$\mathbf{A} \mathbf{B} + \mathbf{B}^\top \mathbf{A}^\top \preceq 2 \|\mathbf{A}\| \|\mathbf{B}\| \mathbf{I}_n.$$

Proof For any unit vector $\mathbf{x} \in \mathbb{R}^n$ with $\|\mathbf{x}\| = 1$, we can obtain that

$$\mathbf{x}^\top (\mathbf{A}\mathbf{B} + \mathbf{B}^\top \mathbf{A}^\top) \mathbf{x} = \mathbf{x}^\top \mathbf{A}\mathbf{B}\mathbf{x} + \mathbf{x}^\top \mathbf{B}^\top \mathbf{A}^\top \mathbf{x} \stackrel{(a)}{=} 2\mathbf{x}^\top \mathbf{A}\mathbf{B}\mathbf{x},$$

where (a) is from the fact that $\mathbf{x}^\top \mathbf{B}^\top \mathbf{A}^\top \mathbf{x}$ is a scalar. By the Cauchy–Schwarz inequality and the definition of the spectral norm, we have that

$$|\mathbf{x}^\top \mathbf{A}\mathbf{B}\mathbf{x}| \leq \|\mathbf{A}\mathbf{B}\mathbf{x}\| \cdot \|\mathbf{x}\| \leq \|\mathbf{A}\| \cdot \|\mathbf{B}\| \cdot \|\mathbf{x}\|^2 = \|\mathbf{A}\| \|\mathbf{B}\|.$$

Hence, we obtain the following inequality:

$$\mathbf{x}^\top (\mathbf{A}\mathbf{B} + \mathbf{B}^\top \mathbf{A}^\top) \mathbf{x} \leq 2\|\mathbf{A}\| \|\mathbf{B}\|.$$

Since this holds for any unit vector $\mathbf{x}$, it follows that

$$\mathbf{A}\mathbf{B} + \mathbf{B}^\top \mathbf{A}^\top \preceq 2\|\mathbf{A}\| \|\mathbf{B}\| \mathbf{I}_n.$$

■

Lemma 26 *Let $t \geq 1$ be a positive integer. For all real numbers x , satisfying $0 \leq x \leq \frac{1}{t}$, the following inequality holds:*

$$(1+x)^t \leq 1 + 3tx.$$

Proof Let $f(x) := 1 + 3tx - (1+x)^t$, $x \in [0, \frac{1}{t}]$. Then, for all $x \in [0, \frac{1}{t}]$, we obtain

$$f'(x) = 3t - t(1+x)^{t-1} \geq 3t - t(1 + \frac{1}{t})^{t-1} \geq (3-e)t > 0.$$

Therefore, for all $x \in [0, \frac{1}{t}]$, $f(x) \geq f(0) = 0$, which means $(1+x)^t \leq 1 + 3tx$ for all $x \in [0, \frac{1}{t}]$. ■

Lemma 27 *Let $k \in \mathbb{R}_{\geq 1}$, $q \in (\frac{1}{2}, 1)$. Suppose that sequences $\{a_t\}_{t=0}^\infty, \{b_t\}_{t=0}^\infty \subset \mathbb{R}_{\geq 0}$ satisfy*

$$b_{t+1} \leq qb_t + \frac{1-q}{180k^2} (a_{t-1}^2 + a_t^2 + 2a_{t-1}a_t + \sqrt{b_t}(a_{t-1} + a_t)), \quad (35)$$

$$a_t \leq 2k\sqrt{b_t} + \frac{1}{6}a_{t-1}, \quad t = 1, 2, \dots, \quad (36)$$

and another pair of sequences $\{\tilde{a}_t\}_{t=0}^\infty, \{\tilde{b}_t\}_{t=0}^\infty \subset \mathbb{R}_{\geq 0}$ satisfy

$$\tilde{b}_{t+1} = q\tilde{b}_t + \frac{1-q}{180k^2} (\tilde{a}_{t-1}^2 + \tilde{a}_t^2 + 2\tilde{a}_{t-1}\tilde{a}_t + \sqrt{\tilde{b}_t}(\tilde{a}_{t-1} + \tilde{a}_t)), \quad (37)$$

$$\tilde{a}_t = 2k\sqrt{\tilde{b}_t} + \frac{1}{6}\tilde{a}_{t-1}, \quad t = 1, 2, \dots \quad (38)$$

If the initial conditions satisfy

$$a_0 \leq \tilde{a}_0, b_0 \leq \tilde{b}_0, b_1 \leq \tilde{b}_1,$$

then $a_t \leq \tilde{a}_t$ and $b_t \leq \tilde{b}_t$ hold for all $t \geq 0$.

Proof We proceed by mathematical induction. From inequality (35), we obtain

$$\begin{aligned} a_1 &\leq 2k\sqrt{b_1} + \frac{1}{6}a_0 \\ &\stackrel{(a)}{\leq} 2k\sqrt{\tilde{b}_1} + \frac{1}{6}\tilde{a}_0 \\ &\stackrel{(b)}{=} \tilde{a}_1, \end{aligned}$$

where (a) is by initial conditions; and (b) is from equality (37). Analogously, inequality (36) implies

$$\begin{aligned} b_2 &\leq qb_1 + \frac{1-q}{180k^2}(a_0^2 + a_1^2 + 2a_0a_1 + \sqrt{b_1}(a_0 + a_1)) \\ &\stackrel{(c)}{\leq} q\tilde{b}_1 + \frac{1-q}{180k^2}(\tilde{a}_0^2 + \tilde{a}_1^2 + 2\tilde{a}_0\tilde{a}_1 + \sqrt{\tilde{b}_1}(\tilde{a}_0 + \tilde{a}_1)) \\ &\stackrel{(d)}{=} \tilde{b}_2, \end{aligned}$$

where (c) is due to initial conditions and $a_1 \leq \tilde{a}_1$; and (d) is by equality (38). By induction, we conclude that $a_t \leq \tilde{a}_t$ and $b_t \leq \tilde{b}_t$ for all $t \geq 0$, which completes the proof. $\blacksquare$

Lemma 28 Let $k \in \mathbb{R}_{\geq 1}$, $q \in (\frac{1}{2}, 1)$. Suppose that sequences $\{\tilde{a}_t\}_{t=0}^\infty, \{\tilde{b}_t\}_{t=0}^\infty \subset \mathbb{R}_{\geq 0}$ satisfy:

$$\tilde{b}_{t+1} = q\tilde{b}_t + \frac{1-q}{180k^2}(\tilde{a}_{t-1}^2 + \tilde{a}_t^2 + 2\tilde{a}_{t-1}\tilde{a}_t + \sqrt{\tilde{b}_t}(\tilde{a}_{t-1} + \tilde{a}_t)), \quad (39)$$

$$\tilde{a}_t = 2k\sqrt{\tilde{b}_t} + \frac{1}{6}\tilde{a}_{t-1}, \quad t = 1, 2, \dots \quad (40)$$

If the initial conditions satisfy

$$\tilde{a}_0, \tilde{a}_1, \tilde{b}_0, \tilde{b}_1 \in \mathbb{R}_{\geq 0}, \tilde{a}_0 \leq 3k\sqrt{\tilde{b}_0} \leq \frac{3\sqrt{2}k}{\sqrt{1+q}}\sqrt{\tilde{b}_1}, \tilde{a}_1 \leq 3k\sqrt{\tilde{b}_1},$$

then we have that $\tilde{a}_t \leq 3k\sqrt{\tilde{b}_0} \left(\frac{1+q}{2}\right)^{t/2}$ for all $t \geq 0$.

Proof We proceed by mathematical induction. We first consider the following auxiliary system:

$$\hat{b}_{t+1} = \max\{q\hat{b}_t + \frac{1-q}{180k^2}(\hat{a}_{t-1}^2 + \hat{a}_t^2 + 2\hat{a}_{t-1}\hat{a}_t + \sqrt{\hat{b}_t}(\hat{a}_{t-1} + \hat{a}_t)), \frac{1+q}{2}\hat{b}_t\}, \quad (41)$$

$$\hat{a}_t = \max\{2k\sqrt{\hat{b}_t} + \frac{1}{6}\hat{a}_{t-1}, 3k\sqrt{\hat{b}_t}\}, \quad t = 1, 2, \dots \quad (42)$$

Let $\hat{a}_0 = \tilde{a}_0, \hat{b}_0 = \tilde{b}_0$, and $\hat{b}_1 = \tilde{b}_1$. It holds that $\hat{a}_0 \leq 3k\sqrt{\hat{b}_0} \leq \frac{3\sqrt{2}k}{\sqrt{1+q}}\sqrt{\hat{b}_1}$, and thus we have

$$2k\sqrt{\hat{b}_1} + \frac{1}{6}\hat{a}_0 \leq 2k\sqrt{\hat{b}_1} + \frac{\sqrt{2}k}{2\sqrt{1+q}}\sqrt{\hat{b}_1} \leq 3k\sqrt{\hat{b}_1}.$$

From equality (42) at $t = 1$, we obtain $\hat{a}_1 = 3k\sqrt{\hat{b}_1}$. Since $\hat{a}_0 \leq \frac{3\sqrt{2k}}{\sqrt{1+q}}\sqrt{\hat{b}_1}$ and $\hat{a}_1 = 3k\sqrt{\hat{b}_1}$, it follows that

$$\begin{aligned}
 & q\hat{b}_1 + \frac{1-q}{180k^2}(\hat{a}_0^2 + \hat{a}_1^2 + 2\hat{a}_0\hat{a}_1 + \sqrt{\hat{b}_1}(\hat{a}_0 + \hat{a}_1)) \\
 & \leq q\hat{b}_1 + \frac{1-q}{180k^2}\left(\frac{18k^2}{1+q}\hat{b}_1 + 9k^2\hat{b}_1 + \frac{18\sqrt{2k^2}}{\sqrt{1+q}}\hat{b}_1 + \frac{3\sqrt{2k}}{\sqrt{1+q}}\hat{b}_1 + 3k\hat{b}_1\right) \\
 & \leq q\hat{b}_1 + \frac{1-q}{180k^2}(18k^2 + 9k^2 + 18\sqrt{2k^2} + 3\sqrt{2k^2} + 3k^2)\hat{b}_1 \\
 & \leq q\hat{b}_1 + \frac{1-q}{2}\hat{b}_1 \\
 & \leq \frac{1+q}{2}\hat{b}_1.
 \end{aligned}$$

From equality (41) at $t = 1$, we have $\hat{b}_2 = \frac{1+q}{2}\hat{b}_1$. Using the same reasoning for $t = 2$ yields

$$2k\sqrt{\hat{b}_2} + \frac{1}{6}\hat{a}_1 = 2k\sqrt{\hat{b}_2} + \frac{k}{2}\sqrt{\hat{b}_1} = 2k\sqrt{\hat{b}_2} + \frac{k}{2}\sqrt{\frac{2}{1+q}}\sqrt{\hat{b}_2} \leq 3k\sqrt{\hat{b}_2}.$$

Equality (42) at $t = 1$ implies that $\hat{a}_2 = 3k\sqrt{\hat{b}_2}$. Since $\hat{a}_1 = 3k\sqrt{\hat{b}_1}$ and $\hat{a}_2 = 3k\sqrt{\hat{b}_2}$, we obtain

$$\begin{aligned}
 & q\hat{b}_2 + \frac{1-q}{180k^2}(\hat{a}_1^2 + \hat{a}_2^2 + 2\hat{a}_1\hat{a}_2 + \sqrt{\hat{b}_2}(\hat{a}_1 + \hat{a}_2)) \\
 & = q\hat{b}_2 + \frac{1-q}{180k^2}(9k^2\hat{b}_1 + 9k^2\hat{b}_2 + 18k^2\sqrt{\hat{b}_1\hat{b}_2} + 3k(\sqrt{\hat{b}_1\hat{b}_2} + \hat{b}_2)) \\
 & = q\hat{b}_2 + \frac{1-q}{180k^2}\left(\frac{18k^2}{1+q} + 9k^2 + \frac{18\sqrt{2k^2}}{\sqrt{1+q}} + \frac{3\sqrt{2k}}{\sqrt{1+q}} + 3k\right)\hat{b}_2 \\
 & \leq q\hat{b}_2 + \frac{1-q}{180k^2}(18k^2 + 9k^2 + 18\sqrt{2k^2} + 3\sqrt{2k^2} + 3k^2)\hat{b}_2 \\
 & \leq q\hat{b}_2 + \frac{1-q}{2}\hat{b}_2 \\
 & \leq \frac{1+q}{2}\hat{b}_2.
 \end{aligned}$$

Applying equality (41) at $t = 2$, $\hat{b}_3 = \frac{1+q}{2}\hat{b}_2$ is derived. Therefore, we have that $\hat{a}_1 = 3k\sqrt{\hat{b}_1}$, $\hat{a}_2 = 3k\sqrt{\hat{b}_2}$, and $\hat{b}_3 = \frac{1+q}{2}\hat{b}_2$. Assume that $\hat{a}_{t-1} = 3k\sqrt{\hat{b}_{t-1}}$, $\hat{a}_t = 3k\sqrt{\hat{b}_t}$, and $\hat{b}_{t+1} = \frac{1+q}{2}\hat{b}_t$, we claim that $\hat{a}_{t+1} = 3k\sqrt{\hat{b}_{t+1}}$ and $\hat{b}_{t+2} = \frac{1+q}{2}\hat{b}_{t+1}$. From equality (42), we obtain

$$\begin{aligned}
 \hat{a}_{t+1} &= \max\{2k\sqrt{\hat{b}_{t+1}} + \frac{1}{6}\hat{a}_t, 3k\sqrt{\hat{b}_{t+1}}\} \\
 &= \max\{2k\sqrt{\hat{b}_{t+1}} + \frac{1}{2}k\sqrt{\hat{b}_t}, 3k\sqrt{\hat{b}_{t+1}}\} \\
 &= \max\{2k\sqrt{\hat{b}_{t+1}} + \frac{k\sqrt{\frac{2}{1+q}}}{2}\sqrt{\hat{b}_{t+1}}, 3k\sqrt{\hat{b}_{t+1}}\} \\
 &= 3k\sqrt{\hat{b}_{t+1}}.
 \end{aligned}$$

Analogously, equality (41) implies that

$$\begin{aligned}
 \hat{b}_{t+2} &= \max\{q\hat{b}_{t+1} + \frac{1-q}{180k^2}(\hat{a}_t^2 + \hat{a}_{t+1}^2 + 2\hat{a}_t\hat{a}_{t+1} + \sqrt{\hat{b}_{t+1}}(\hat{a}_t + \hat{a}_{t+1})), \frac{1+q}{2}\hat{b}_{t+1}\} \\
 &= \max\{q\hat{b}_{t+1} + \frac{1-q}{180k^2}(9k^2\hat{b}_t + 9k^2\hat{b}_{t+1} + 18k^2\sqrt{\hat{b}_t\hat{b}_{t+1}} + \sqrt{\hat{b}_{t+1}}(3k\sqrt{\hat{b}_t} + 3k\sqrt{\hat{b}_{t+1}})), \\
 &\quad \frac{1+q}{2}\hat{b}_{t+1}\} \\
 &= \max\{q\hat{b}_{t+1} + \frac{1-q}{20}(\frac{2}{1+q} + 1 + 2(\frac{2}{1+q})^{1/2} + \frac{1}{3k}(\frac{2}{1+q})^{1/2} + \frac{1}{3k})\hat{b}_{t+1}, \frac{1+q}{2}\hat{b}_{t+1}\} \\
 &= \frac{1+q}{2}\hat{b}_{t+1}.
 \end{aligned}$$

Therefore, we have that $\{\hat{b}_t\}_{t=0}^{t=\infty}$ decreases in a linear rate and that $\hat{a}_t = 3k\sqrt{\hat{b}_t}$ in the system (41) and (42), which means that $\hat{a}_t \leq 3k\sqrt{\hat{b}_0} \left(\frac{1+q}{2}\right)^{t/2} = 3k\sqrt{\hat{b}_0} \left(\frac{1+q}{2}\right)^{t/2}$ for all $t \geq 0$.

We now prove that $\tilde{a}_t \leq \hat{a}_t, \tilde{b}_t \leq \hat{b}_t$ for all $t \geq 0$. Obviously, $\tilde{a}_0 \leq \hat{a}_0, \tilde{a}_1 \leq \hat{a}_1, \tilde{b}_0 \leq \hat{b}_0$, and $\tilde{b}_1 \leq \hat{b}_1$ hold. Applying equality (39) at $t = 1$ and equality (40) at $t = 2$, we obtain

$$\begin{aligned}
 \tilde{b}_2 &= q\tilde{b}_1 + \frac{1-q}{180k^2}(\tilde{a}_0^2 + \tilde{a}_1^2 + 2\tilde{a}_0\tilde{a}_1 + \sqrt{\tilde{b}_1}(\tilde{a}_0 + \tilde{a}_1)) \\
 &\leq q\hat{b}_1 + \frac{1-q}{180k^2}(\hat{a}_0^2 + \hat{a}_1^2 + 2\hat{a}_0\hat{a}_1 + \sqrt{\hat{b}_1}(\hat{a}_0 + \hat{a}_1)) \\
 &\leq \max\{q\hat{b}_1 + \frac{1-q}{180k^2}(\hat{a}_0^2 + \hat{a}_1^2 + 2\hat{a}_0\hat{a}_1 + \sqrt{\hat{b}_1}(\hat{a}_0 + \hat{a}_1)), \frac{1+q}{2}\hat{b}_1\} \\
 &= \hat{b}_2, \\
 \tilde{a}_2 &= 2k\sqrt{\tilde{b}_2} + \frac{1}{6}\tilde{a}_1 \\
 &\leq 2k\sqrt{\hat{b}_2} + \frac{1}{6}\hat{a}_1 \\
 &\leq \max\{2k\sqrt{\hat{b}_2} + \frac{1}{6}\hat{a}_1, 3k\sqrt{\hat{b}_2}\} \\
 &= \hat{a}_2.
 \end{aligned}$$

Hence, $\tilde{a}_2 \leq \hat{a}_2, \tilde{b}_2 \leq \hat{b}_2$ and recursively, we can obtain $\tilde{a}_t \leq \hat{a}_t, \tilde{b}_t \leq \hat{b}_t$ for all $t \geq 0$. Consequently, $\{\tilde{a}_t\}_{t=0}^{t=\infty}$ achieves at least a linear convergence rate in the system (39) and (40), which means that $\tilde{a}_t \leq 3k\sqrt{\hat{b}_0} \left(\frac{1+q}{2}\right)^{t/2}$ for all $t \geq 0$. ■

Appendix I. Experimental setup

In this section, we provide experimental setup details for Section 4, Figure 2, and Appendix C.

I.1. Setup for Section 4

We apply RGD on WN and compare the convergence with vanilla GD on problem (1).

In the “Faster convergence of WN” part, we divide our experiments into two sets. In the first set of experiments, we consider target matrices with small condition numbers, i.e., $\kappa \in \{1, 3, 5\}$. Other parameters are chosen as $m = 10, r = 5, r_A = 3$, and $n = 1000$. In the second set of experiments, we consider target matrices with large condition numbers of $\kappa \in \{10, 20, 30\}$, on a problem instance with $m = 10, r = 5, r_A = 3$, and $n = 3000$.

In the “On the benefit of overparameterization” part, we also divide our experiments into two sets. The first set tests small $r \in \{4, 5, 6\}$ with $m = 10, r_A = 3$, and $\kappa = 1$. The number of measurements is fixed at $n = 1000$. The second set comes with $m = 20, r_A = 3, \kappa = 10$, and the level of overparameterization is chosen as $r \in \{5, 10, 15\}$, and $n = 3000$ is leveraged.

In these experiments, the ground truth matrix $\mathbf{A} \in \mathbb{R}^{m \times m}$ is formed as $\mathbf{A} = \mathbf{U}\mathbf{\Sigma}\mathbf{U}^\top$, where $\mathbf{U} \in \mathbb{R}^{m \times r_A}$ is a random orthonormal matrix and $\mathbf{\Sigma} \in \mathbb{S}_+^{r_A}$ is a diagonal matrix with entries evenly distributed on a logarithmic scale in the interval $[1/\kappa, 1]$. The independent feature matrices $\{\mathbf{M}_i\}_{i=1}^n \subset \mathbb{S}^m$ are generated in the following manner. For each $i \in \{1, \dots, n\}$, we sample $\mathbf{R}_i \in \mathbb{R}^{m \times m}$ with i.i.d. standard Gaussian entries and define $\mathbf{M}_i = \frac{1}{2\sqrt{n}}(\mathbf{R}_i + \mathbf{R}_i^\top)$, which ensures the symmetry of $\mathbf{M}_i$.

We initialize RGD with $\mathbf{X}_0 = \mathbf{Z}_0(\mathbf{Z}_0^\top \mathbf{Z}_0)^{-1/2}$ and $\mathbf{\Theta}_0 = \mathbf{I}_r$, where $\mathbf{Z}_0 \in \mathbb{R}^{m \times r}$ has i.i.d. standard Gaussian entries. This initialization ensures that $\mathbf{X}_0$ lies on the manifold $\text{St}(m, r)$ and $\mathbf{\Theta}_0 \in \mathbb{S}^r$. For GD, we use $\mathbf{X}_0^{\text{GD}} = 0.1\mathbf{Z}_0$ as small random initialization.

In the “Faster convergence of WN” part, we set stepsizes $\eta = 0.1$ and $\mu = 2$ for RGD and $\eta = 0.1$ for GD. In the “On the benefit of overparameterization” part, we set $\eta = 0.1, 0.12$, and 0.14 with $\mu = 2$ for RGD, and $\eta = 0.1, 0.12$, and 0.14 for GD, where η increases with r .

I.2. Setup for Figure 2

We apply RGD on WN and study the trajectory generated by the algorithm.

In this experiment, we set $m = 300, r = 10, r_A = 5, \kappa = 3$, and use $n = 50000$ feature matrices generated as in I.1. The ground-truth matrix $\mathbf{A} \in \mathbb{R}^{m \times m}$ is constructed as $\mathbf{A} = \mathbf{U}\mathbf{\Sigma}\mathbf{U}^\top$, where $\mathbf{U} \in \mathbb{R}^{m \times r_A}$ is a random orthonormal matrix and $\mathbf{\Sigma} \in \mathbb{S}_+^{r_A}$ is diagonal with entries generated by a power spacing scheme. Specifically, the j -th entry of $\mathbf{\Sigma}$ is given by $\sigma_j = \kappa^{-(\frac{j-1}{r_A-1})^p}$ for $j = 1, \dots, r_A$, where we set $p = 0.6$.

We initialize $\mathbf{X}_0 = \mathbf{Z}_0(\mathbf{Z}_0^\top \mathbf{Z}_0)^{-1/2}$ and $\mathbf{\Theta}_0 = 0.5\mathbf{I}_r$, where $\mathbf{Z}_0 \in \mathbb{R}^{m \times r}$ has i.i.d. standard Gaussian entries. For this experiment, we set the stepsizes to $\eta = 0.2$ and $\mu = 2$.

I.3. Setup for Appendix C

I.3.1. Setup for the experiments with synthetic data

We apply RGD on WN and vanilla GD on problem (1) to reveal other interesting behaviors of WN.

We set m, r, r_A, κ , and n as specified in Appendix C. The ground truth matrix and feature matrices are constructed following the procedure described in I.1.

We initialize $\mathbf{X}_0 = \mathbf{Z}_0(\mathbf{Z}_0^\top \mathbf{Z}_0)^{-1/2}$ and $\mathbf{\Theta}_0 = \mathbf{I}_r$ for RGD, where $\mathbf{Z}_0 \in \mathbb{R}^{m \times r}$ is generated with i.i.d. standard Gaussian entries. For GD, we use $\mathbf{X}_0^{\text{GD}} = 0.1\mathbf{Z}_0$ as small random initialization. In all experiments, we use step sizes $\eta = 0.1, \mu = 2$ for RGD and $\eta = 0.1$ for GD.

I.3.2. Setup for image reconstruction experiments

For the image reconstruction experiments, we conduct two setups: one bases on recovering a CIFAR-10 image from linear measurements and the other on direct matrix sensing of a structured image.

For the CIFAR-10 experiment, we take the first horse image from CIFAR-10 dataset, convert it to grayscale, and vectorize it as $\mathbf{a} \in \mathbb{R}^{1024}$. The ground-truth matrix is set as $\mathbf{A} = \mathbf{a}\mathbf{a}^\top \in \mathbb{S}_+^{1024}$. The overparameterization level is set to $r = 100$, with $n = 50000$ feature matrices generated as in I.1. RGD is initialized with $\mathbf{X}_0 = \mathbf{Z}_0(\mathbf{Z}_0^\top \mathbf{Z}_0)^{-1/2}$ and $\mathbf{\Theta}_0 = \mathbf{I}_r$, where $\mathbf{Z}_0 \in \mathbb{R}^{m \times r}$ has i.i.d. standard Gaussian entries. GD uses small random initialization: $\mathbf{X}_0^{\text{GD}} = 0.1 \mathbf{Z}_0$. We run RGD for $t_{\text{RGD}} = 100$ and GD for $t_{\text{GD}} = 200$ iterations. For RGD, we adopt stepsizes of $\eta = 0.01$ for updating $\mathbf{X}$ and $\mu = 2$ for updating $\mathbf{\Theta}$. For GD, we apply stepsize $\eta = 0.01$ to update $\mathbf{X}$.

After optimization, following the approach of [15, Section 4.1], we perform a rank-one truncated SVD on the recovered matrix $\hat{\mathbf{A}}$, and the estimate of the original signal is constructed as the leading singular vector multiplied by the square root of its corresponding singular value. The resulting vector is then reshaped into a 32×32 reconstruction image.

For the structured image experiment, we generate a grayscale matrix $\mathbf{A} \in \mathbb{S}_+^{128}$ of rank $r_A = 2$ using block-wave basis functions. Specifically, we construct r_A one-dimensional signals of length 128, where each signal is a normalized block wave taking values in ± 1 with a random period. Stacking these signals forms a matrix $\mathbf{U} \in \mathbb{R}^{128 \times r_A}$. The ground-truth image is defined as $\mathbf{A} = \mathbf{U}\mathbf{\Lambda}\mathbf{U}^\top$, where $\mathbf{\Lambda}$ is a 2×2 diagonal matrix with diagonal entries 1 and 0.9. This diagonal matrix assigns geometrically decaying weights to different block-wave modes.

We again fix $r = 100$ and use $n = 50000$ feature matrices generated as in I.1. Both RGD and GD are randomly initialized as above. We run RGD for $t_{\text{RGD}} = 100$ and GD for $t_{\text{GD}} = 200$ iterations. We adopt stepsizes of $\eta = 0.03$ and $\mu = 2$ in RGD and a stepsize of $\eta = 0.03$ in GD.

The per-iteration computational complexity of both RGD and GD is $\mathcal{O}(nm^2r)$, which is dominated by the operation of sensing. Since each RGD iteration requires performing two sensing operations while GD requires only one, we set the number of iterations as $t_{\text{GD}} = 2t_{\text{RGD}}$ to make the overall runtime roughly comparable between the two methods.