

Efficient Federated Low Rank Matrix Recovery via Alternating GD and Minimization: A Simple Proof

Namrata Vaswani[✉], *Fellow, IEEE*

Abstract—This note provides a significantly simpler and shorter proof of our sample complexity guarantee for solving the low rank column-wise sensing problem using the Alternating Gradient Descent (GD) and Minimization (AltGDmin) algorithm. AltGDmin was developed and analyzed for solving this problem in our recent work. We also provide an improved guarantee.

Index Terms—Low rank column-wise sensing (LRCS), federated learning, multi-task representation learning.

I. INTRODUCTION

WE STUDY the low rank column-wise sensing (LRCS) problem which involves recovering a low rank matrix from independent compressive measurements of each of its columns. This problem occurs in dynamic MRI [1] and in multi-task linear representation learning for few-shot learning [2]. The alternating gradient descent (GD) and minimization (AltGDmin) algorithm for solving it in a fast, communication-efficient and private fashion was developed and analyzed in our recent work [3]. This short paper provides a significantly simpler and shorter proof of our sample complexity guarantee for AltGDmin. In fact, it also improves the sample complexity needed by the AltGDmin iterations by a factor of r .

II. PROBLEM STATEMENT, NOTATION, AND ALGORITHM

A. Problem Statement and Assumption and Notation

The goal is to recover an $n \times q$ rank- r matrix $\mathbf{X}^* = [\mathbf{x}_1^*, \mathbf{x}_2^*, \dots, \mathbf{x}_q^*]$ from m linear projections (sketches) of each of its q columns, i.e. from

$$\mathbf{y}_k := \mathbf{A}_k \mathbf{x}_k^*, \quad k \in [q] \quad (1)$$

where each $\mathbf{y}_k$ is an m -length vector, $[q] := \{1, 2, \dots, q\}$, and the measurement/sketching matrices $\mathbf{A}_k$ are mutually independent and known. The setting of interest is low-rank (LR), $r \ll \min(n, q)$, and undersampled measurements, $m < n$. Each $\mathbf{A}_k$ is assumed to be random-Gaussian: each entry of it is independent and identically distributed (i.i.d.) standard Gaussian. Let $\mathbf{X}^* \stackrel{\text{SVD}}{=} \mathbf{U}^* \mathbf{\Sigma}^* \mathbf{V}^{*\top} := \mathbf{U}^* \mathbf{B}^*$ denote its

reduced (rank r) SVD, and $\kappa := \sigma_{\max}^* / \sigma_{\min}^*$ the condition number of $\mathbf{\Sigma}^*$. We let $\mathbf{B}^* := \mathbf{\Sigma}^* \mathbf{V}^{*\top}$.

Since no measurement $\mathbf{y}_{ki}$ is a global function of the entire matrix, $\mathbf{X}^*$, we need the following assumption, borrowed from LR matrix completion literature, to make our problem well-posed (allow for correct interpolation across columns).

Assumption 1 (Incoherence of Right Singular Vectors): Assume that $\|\mathbf{b}_k^*\|^2 \leq \mu^2 r \sigma_{\max}^{*2} / q$ for a numerical constant μ .

In our discussion of communication complexity and privacy, we assume a vertically federated setting: different subsets of $\mathbf{y}_k, \mathbf{A}_k$ are available at different nodes.

1) Notation: We use $\|\cdot\|_F$ to denote the Frobenius norm, $\|\cdot\|$ without a subscript for the (induced) l_2 norm, $^\top$ to denote matrix or vector transpose, $\mathbf{e}_k$ to denote the k -th canonical basis vector (k -th column of $\mathbf{I}$), and $\mathbf{M}^\dagger := (\mathbf{M}^\top \mathbf{M})^{-1} \mathbf{M}^\top$. For two $n \times r$ matrices $\mathbf{U}_1, \mathbf{U}_2$ that have orthonormal columns, we use

$$\text{SD}_2(\mathbf{U}_1, \mathbf{U}_2) := \|(\mathbf{I} - \mathbf{U}_1 \mathbf{U}_1^\top) \mathbf{U}_2\|$$

as the Subspace Distance (SD) measure. In our previous work [3], we used the Frobenius norm SD,

$$\text{SD}_F(\mathbf{U}_1, \mathbf{U}_2) := \|(\mathbf{I} - \mathbf{U}_1 \mathbf{U}_1^\top) \mathbf{U}_2\|_F.$$

Clearly, $\text{SD}_F(\mathbf{U}_1, \mathbf{U}_2) \leq \sqrt{r} \text{SD}_2(\mathbf{U}_1, \mathbf{U}_2)$. We reuse the letters c, C to denote different numerical constants in each use with the convention that $c < 1$ and $C \geq 1$. We use $\sum_k$ as a shortcut for the summation over $k = 1$ to q and $\sum_{ki}$ for the summation over $i = 1$ to m and $k = 1$ to q . We use *whp* to refer to “with high probability” and this means that the claim holds with probability (w.p.) at least $1 - n^{-10}$.

B. Review of AltGDmin Algorithm [3]

AltGDmin, summarized in Algorithm 1, imposes the LR constraint by factorizing the unknown matrix $\mathbf{X}$ as $\mathbf{X} = \mathbf{U} \mathbf{B}$ with $\mathbf{U}$ being an $n \times r$ matrix and $\mathbf{B}$ an $r \times q$ matrix. It minimizes $f(\mathbf{U}, \mathbf{B}) := \sum_{k=1}^q \|\mathbf{y}_k - \mathbf{U} \mathbf{b}_k\|^2$ as follows:

- 1) *Truncated spectral initialization:* Initialize $\mathbf{U}$ (see below).
- 2) At each iteration, update $\mathbf{B}$ and $\mathbf{U}$ as follows:
 - a) *Minimization for $\mathbf{B}$:* keeping $\mathbf{U}$ fixed, update $\mathbf{B}$ by solving $\min_{\mathbf{B}} f(\mathbf{U}, \mathbf{B})$. Due to the form of the LRCS model, this minimization decouples across columns, making it a cheap least squares problem of recovering q different r length vectors. It is solved as $\mathbf{b}_k = (\mathbf{A}_k \mathbf{U})^\dagger \mathbf{y}_k$ for each $k \in [q]$.

Manuscript received 3 July 2023; revised 26 October 2023; accepted 10 February 2024. Date of publication 1 March 2024; date of current version 18 June 2024. This work was supported in part by NSF under Grant CCF-2115200.

The author is with the Department of Electrical and Computer Engineering, Iowa State University, Ames, IA 50011 USA (e-mail: namrata@iastate.edu). Communicated by Y. Chi, Associate Editor for Machine Learning and Statistics.

Digital Object Identifier 10.1109/TIT.2024.3365795

0018-9448 © 2024 IEEE. Personal use is permitted, but republication/redistribution requires IEEE permission. See <https://www.ieee.org/publications/rights/index.html> for more information.

- b) *Projected-GD for U* : keeping B fixed, update U by a GD step, followed by orthonormalizing its columns: $U^+ = QR(U - \eta \nabla_U f(U, B))$. Here $QR(\cdot)$ orthonormalizes the columns of its input.

We initialize U by computing the top r singular vectors of

$$X_0 := \sum_k A_k^\top y_{k, \text{trunc}} e_k^\top, \quad y_{k, \text{trunc}} := \text{trunc}(y_k, \alpha)$$

Here $\alpha := \tilde{C} \sum_k \|y_k\|^2 / mq$ with $\tilde{C} := 9\kappa^2 \mu^2$. The function trunc truncates (zeroes out) all entries of y_k with magnitude greater than $\sqrt{\alpha}$, i.e., for all $j \in [n]$, $\text{trunc}(y, \alpha)_j = (y)_j \mathbb{1}_{|y_j| \leq \sqrt{\alpha}}$, with $\mathbb{1}$ being the indicator function.

Sample-splitting is assumed, i.e., each new update of U and B uses a new independent set of measurements and measurement matrices, y_k, A_k .

The use of minimization to update B at each iteration is what helps ensure that we can show exponential error decay with a constant step size. At the same time, due to the column-wise decoupled nature of LRCS, the time complexity for this step is only as much as that of computing one gradient w.r.t. U . Both steps need time¹ of order mqr . This is only r times more than “linear time” (time needed to read the algorithm inputs, here y_k, A_k ’s). To our knowledge, r -times linear-time is the best known time complexity for any algorithm for any LR matrix recovery problem. Moreover, due to the use of the $X = UB$ factorization, AltGDmin is also communication-efficient. Each node needs to only send nr scalars (gradients w.r.t. U) at each iteration.

III. NEW GUARANTEE

Let m_0 denote the total number of samples per column needed for initialization and let m_1 denote this number for each GDmin iteration. Then, the total sample complexity per column is $m = m_0 + m_1 T$. Our guarantee given next provides the required minimum value of m .

Theorem 2: Assume that Assumption 1 holds. Set $\eta = 0.4/m\sigma_{\max}^2$ and $T = C\kappa^2 \log(1/\epsilon)$. If

$$mq \geq C\kappa^4 \mu^2 (n + q)r(\kappa^4 r + \log(1/\epsilon))$$

and $m \geq C \max(\log n, \log q, r) \log(1/\epsilon)$, then, with probability (w.p.) at least $1 - n^{-10}$,

$$\text{SD}_2(U, U^*) \leq \epsilon \text{ and } \|x_k - x_k^*\| \leq \epsilon \|x_k^*\| \text{ for all } k \in [q].$$

The time complexity is $mqr \cdot T = mqr \cdot \kappa^2 \log(1/\epsilon)$. The communication complexity is $nr \cdot T = nr \cdot \kappa^2 \log(1/\epsilon)$ per node.

Proof: We prove the three results needed for proving this in Sections IV-B, IV-D, and V below. We use these to prove the above result in Appendix A. $\square$

¹The LS step time is $\max(q \cdot mnr, q \cdot mr^2) = mqr$ (maximum of the time needed for computing $A_k U$ for all k , and that for obtaining b_k for all k) while the GD step time is $\max(q \cdot mnr, nr^2) = mqr$ (maximum of the time needed for computing the gradient w.r.t. U , and time for the QR step).

Algorithm 1 The AltGD-Min Algorithm

- 1: **Input**: $y_k, A_k, k \in [q]$
- 2: **Sample-split**: Partition the measurements and measurement matrices into $2T + 2$ equal-sized disjoint sets: two sets for initialization and $2T$ sets for the iterations. Denote these by $y_k^{(\tau)}, A_k^{(\tau)}, \tau = 0, 1, \dots, 2T$.
- 3: **Initialization**:
- 4: Using $y_k \equiv y_k^{(00)}, A_k \equiv A_k^{(00)}$,
- 5: set $\alpha \leftarrow \tilde{C} \frac{1}{mq} \sum_{ki} |y_{ki}|^2$,
- 6: Using $y_k \equiv y_k^{(0)}, A_k \equiv A_k^{(0)}$,
- 7: set $y_{k, \text{trunc}}(\alpha) \leftarrow y_{k, \text{trunc}} := \text{trunc}(y_k, \alpha)$,
- 8: set $X_0 \leftarrow (1/m) \sum_{k \in [q]} A_k^\top y_{k, \text{trunc}}(\alpha) e_k^\top$
- 9: set $U_0 \leftarrow$ top- r -singular-vectors of X_0
- 10: **GDmin iterations**:
- 11: **for** $t = 1$ **to** T **do**
- 12: Let $U \leftarrow U_{t-1}$.
- 13: Using $y_k \equiv y_k^{(t)}, A_k \equiv A_k^{(t)}$,
- 14: set $b_k \leftarrow (A_k U)^\dagger y_k, x_k \leftarrow U b_k$ for all $k \in [q]$
- 15: Using $y_k \equiv y_k^{(T+t)}, A_k \equiv A_k^{(T+t)}$, compute
- 16: set $\nabla_U f(U, B) = \sum_k A_k^\top (A_k U b_k - y_k) b_k^\top$
- 17: set $\hat{U}^+ \leftarrow U - (\eta/m) \nabla_U f(U, B_t)$.
- 18: compute $\hat{U}^+ \stackrel{\text{QR}}{=} U^+ R^+$.
- 19: Set $U_t \leftarrow U^+$.
- 20: **end for**

A. Discussion

We use $a \gtrsim b$ to mean that $a \geq C_{\kappa, \mu} b$ where $C_{\kappa, \mu}$ includes terms dependent on κ, μ . Most of our discussion treats κ, μ as numerical constants that do not grow with n, q, r . This assumption is borrowed from the rich past literature on various other LR matrix recovery problems, e.g., see [4] and [5]. Also, below, whp means w.p. at least $1 - n^{-10}$.

In this work (as well as in other works that use matrix factorization to solve LR recovery problems), the goal is to obtain a bound of the form $\text{SD}_2(U, U^*) \leq \epsilon$. If a bound on SD_F is needed, one can use $\text{SD}_F(U, U^*) \leq \sqrt{r} \text{SD}_2(U, U^*) \leq \sqrt{r} \epsilon$. Using Lemma 3 given below and triangle inequality, the bound on SD_2 helps get the bound $\|x_k - x_k^*\| \lesssim \epsilon \|x_k^*\|$.

Our older result from [3] used $\text{SD}_F(U, U^*)$ in its analysis and showed that, to guarantee $\text{SD}_F(U, U^*) \leq \epsilon$ whp, we need

$$mq \geq C\kappa^4 \mu^2 (n + q)r^2(\kappa^4 + \log(1/\epsilon))$$

Since $\text{SD}_2(U, U^*) \leq \text{SD}_F(U, U^*)$, this implies that we need this same complexity also to guarantee $\text{SD}_2(U, U^*) \leq \epsilon$. Our new result needs $mq \geq C\kappa^4 \mu^2 (n + q)r(\kappa^4 r + \log(1/\epsilon))$ to guarantee $\text{SD}_2(U, U^*) \leq \epsilon$. Thus, this new result improves the dependence on r, ϵ from order $r^2 \log(1/\epsilon)$ to order $r \max(r, \log(1/\epsilon))$. This is an improvement over the old result by a factor of $\min(r, \log(1/\epsilon))$. This improvement is obtained because our new guarantee for the GD step (Theorem 6) only needs $mq \gtrsim nr$ at each iteration. On the other hand, the older result needed $mq \gtrsim nr^2$. Both guarantees need $mq \gtrsim nr^2$ for initialization, see Theorem 13.

We are able to improve our result because we now use a simpler proof technique that works for the LRCS problem, but

not for its LR phase retrieval (LRPR) generalization. LRPR involves recovering $\mathbf{X}^*$ from $\mathbf{z}_k := [\mathbf{y}_k], k \in [q]$. In [3], we were attempting to solve both problems. There are two differences in our new proof compared with that of [3]: (i) we use the 2-norm subspace distance $\text{SD}_2(\mathbf{U}, \mathbf{U}^*)$ instead of $\text{SD}_F(\mathbf{U}, \mathbf{U}^*)$, and (ii) we do not use the fundamental theorem of calculus [6], [7] for analyzing the GD step, but instead use a much simpler direct approach. If we only use (ii) but not (i), we will still get a simpler proof, but we will not get the sample complexity gain. In [3], we used the Frobenius norm SD because it helped obtain the desired nr^3 guarantee for LRPR.² Also, in hindsight, the use of the fundamental theorem of calculus was unnecessary. It has been used in earlier work [7] for analyzing a GD based algorithm for standard PR and for LR matrix completion and that originally motivated us to adapt the same approach for LRCS and LRPR.

Both the current result and our old one have the same dependence on κ . After initialization, the sample complexity grows roughly as κ^4 . A similar dependence on κ also exists in all past sample complexity guarantees for iterative – AMin or GD – algorithms for LR matrix sensing, LR matrix completion or robust PCA, e.g., see [4] and [5]; and also for our older work on AltMin for a generalization of LRCS, LR phase retrieval [8]. One way to improve this dependence is to develop a stage-wise algorithm similar to that introduced in [4] and [9] for LR matrix sensing and completion respectively. Developing this for the current LRCS problem is an open future work question. A second option is to use the convex relaxation approaches which usually depend on lower powers of κ . However, these need a much higher iteration complexity, making them much slower.

IV. NEW PROOF: GDMIN ITERATIONS

A. Definitions and Preliminaries

Let $\mathbf{U}$ be the estimate at the t -th iteration. Define

$$\mathbf{g}_k := \mathbf{U}^\top \mathbf{x}_k^*, k \in [q], \text{ and } \mathbf{G} := \mathbf{U}^\top \mathbf{X}^*,$$

$$\mathbf{P} := \mathbf{I} - \mathbf{U}^* \mathbf{U}^{*\top},$$

$$\text{GradU} := \nabla_{\mathbf{U}} f(\mathbf{U}, \mathbf{B}) = \sum_k \mathbf{A}_k^\top (\mathbf{A}_k \mathbf{U} \mathbf{b}_k - \mathbf{y}_k) \mathbf{b}_k^\top$$

$$= \sum_{ki} (\mathbf{y}_{ki} - \mathbf{a}_{ki}^\top \mathbf{U} \mathbf{b}_k) \mathbf{a}_{ki} \mathbf{b}_k^\top$$

For an $n_1 \times n_2$ matrix, $\mathbf{Z}$, $\sigma_{\min}(\mathbf{Z}) = \sigma_{\min}(\mathbf{Z}^\top) = \sigma_{\min(n_1, n_2)}(\mathbf{Z})$. Thus, if $\mathbf{A}$ is tall, then $\sigma_{\min}(\mathbf{A}) = \sqrt{\lambda_{\min}(\mathbf{A}^\top \mathbf{A})}$. Using this, it follows that, if $\mathbf{A} = \mathbf{BC}$ and $\mathbf{A}$ and $\mathbf{B}$ are tall (or square), then $\sigma_{\min}(\mathbf{A}) \geq \sigma_{\min}(\mathbf{B})\sigma_{\min}(\mathbf{C})$.

B. Minimization Step

Assume $\text{SD}_2(\mathbf{U}, \mathbf{U}^*) \leq \delta_t$ with $\delta_t < 0.02$.

We use the following lemma from [3].

Lemma 3 [3]: Let $\mathbf{g}_k := \mathbf{U}^\top \mathbf{x}_k^*$. Then, w.p. at least $1 - \exp(\log q + r - c\epsilon m)$, for all $k \in [q]$,

$$\|\mathbf{g}_k - \mathbf{b}_k\| \leq 1.2\epsilon \|(\mathbf{I}_n - \mathbf{U}\mathbf{U}^\top) \mathbf{U}^* \mathbf{b}_k^*\|$$

²Our older guarantee for LRPR [3] needed $mq \gtrsim nr^2(r + \log(1/\epsilon))$. If we use the new approach developed here, it will need $mq \gtrsim nr(r^3 + \log(1/\epsilon))$.

Proof: We provide a proof of this lemma in Appendix B to emphasise a slightly more general version of this lemma. In particular, our proof shows that we can replace 0.4 in the previous version of this lemma by any $\epsilon > 0$, and the bound holds w.p. at least $1 - \exp(\log q + r - c\epsilon m)$. $\square$

By Lemma 3 with $\epsilon = 0.3$, if $m \gtrsim \max(\log q, \log n, r)$, then, whp, for all $k \in [q]$,

$$\|\mathbf{b}_k - \mathbf{g}_k\| \leq 0.4 \|(\mathbf{I} - \mathbf{U}\mathbf{U}^\top) \mathbf{U}^* \mathbf{b}_k^*\|$$

Using $\text{SD}_2(\mathbf{U}, \mathbf{U}^*) \leq \delta_t$, this directly implies that

- 1) $\|\mathbf{b}_k - \mathbf{g}_k\| \leq 0.4\delta_t \|\mathbf{b}_k^*\|$
- 2) $\|\mathbf{b}_k\| \leq \|\mathbf{g}_k\| + 0.4 \cdot 0.02 \|\mathbf{b}_k^*\| \leq 1.1 \|\mathbf{b}_k^*\|$
- 3) $\|\mathbf{x}_k - \mathbf{x}_k^*\| \leq 1.4\delta_t \|\mathbf{b}_k^*\|$

(to obtain the last bound, we need to add and subtract $\mathbf{U}\mathbf{g}_k$ and then use triangle inequality). Using these bounds,

$$\|\mathbf{B} - \mathbf{G}\|_F \leq 0.4\delta_t \sqrt{\sum_k \|\mathbf{b}_k^*\|^2} = 0.4\delta_t \sqrt{r} \sigma_{\max}^*$$

and we can get a similar bound on $\|\mathbf{X} - \mathbf{X}^*\|_F$. Thus,

- 1) $\|\mathbf{B} - \mathbf{G}\|_F \leq 0.4\delta_t \|\mathbf{B}^*\|_F \leq 0.4\sqrt{r}\delta_t \sigma_{\max}^*$
- 2) $\|\mathbf{X} - \mathbf{X}^*\|_F \leq 1.4\sqrt{r}\delta_t \sigma_{\max}^*$

Furthermore,

$$\sigma_{\min}(\mathbf{B}) \geq \sigma_{\min}(\mathbf{G}) - \|\mathbf{B} - \mathbf{G}\| \geq \sigma_{\min}(\mathbf{G}) - \|\mathbf{B} - \mathbf{G}\|_F$$

To lower bound $\sigma_{\min}(\mathbf{G})$, observe that

$$\sigma_{\min}(\mathbf{G}) = \sigma_{\min}(\mathbf{G}^\top) \geq \sigma_{\min}^* \sigma_{\min}(\mathbf{U}^{*\top} \mathbf{U})$$

and

$$\sigma_{\min}(\mathbf{U}^{*\top} \mathbf{U}) = \sqrt{1 - \|\mathbf{P}\mathbf{U}\|^2} \geq \sqrt{1 - \delta_t^2}.$$

This follows using $\sigma_{\min}^2(\mathbf{U}^{*\top} \mathbf{U}) = \lambda_{\min}(\mathbf{U}^\top \mathbf{U}^* \mathbf{U}^{*\top} \mathbf{U}) = \lambda_{\min}(\mathbf{U}^\top (\mathbf{I} - \mathbf{P}) \mathbf{U}) = \lambda_{\min}(\mathbf{I} - \mathbf{U}^\top \mathbf{P} \mathbf{U}) = \lambda_{\min}(\mathbf{I} - \mathbf{U}^\top \mathbf{P}^2 \mathbf{U}) = 1 - \lambda_{\max}(\mathbf{U}^\top \mathbf{P}^2 \mathbf{U}) = 1 - \|\mathbf{P}\mathbf{U}\|^2$.

Combining the above three bounds and using the bound on $\|\mathbf{B} - \mathbf{G}\|_F$, if $\delta_t < 0.02/\sqrt{r}\kappa$, then

$$\sigma_{\min}(\mathbf{B}) \geq \sqrt{1 - \delta_t^2} \sigma_{\min}^* - 0.4\sqrt{r}\delta_t \sigma_{\max}^* \geq 0.9\sigma_{\min}^*$$

and

$$\sigma_{\max}(\mathbf{B}) \leq \|\mathbf{G}\| + 0.4\sqrt{r}\delta_t \sigma_{\max}^* \leq 1.1\sigma_{\max}^*$$

since $\|\mathbf{G}\| \leq \|\mathbf{B}^*\| = \sigma_{\max}^*$. Thus, we have proved the following claim:

Theorem 4: Assume that $\text{SD}_2(\mathbf{U}, \mathbf{U}^*) \leq \delta_t$. If $\delta_t \leq 0.02/\sqrt{r}\kappa$, and if $m \gtrsim \max(\log q, \log n, r)$, then whp,

- 1) $\|\mathbf{b}_k - \mathbf{g}_k\| \leq 0.4\delta_t \|\mathbf{b}_k^*\|$
- 2) $\|\mathbf{b}_k\| \leq \|\mathbf{g}_k\| + 0.4 \cdot 0.02 \|\mathbf{b}_k^*\| \leq 1.1 \|\mathbf{b}_k^*\|$
- 3) $\|\mathbf{B} - \mathbf{G}\|_F \leq 0.4\delta_t \|\mathbf{B}^*\|_F \leq 0.4\sqrt{r}\delta_t \sigma_{\max}^*$
- 4) $\|\mathbf{x}_k - \mathbf{x}_k^*\| \leq 1.4\delta_t \|\mathbf{b}_k^*\|$
- 5) $\|\mathbf{X} - \mathbf{X}^*\|_F \leq 1.4\sqrt{r}\delta_t \sigma_{\max}^*$
- 6) $\sigma_{\min}(\mathbf{B}) \geq 0.9\sigma_{\min}^*$
- 7) $\sigma_{\max}(\mathbf{B}) \leq 1.1\sigma_{\max}^*$

(only the last two bounds require the upper bound on δ_t).

C. New Bounds on the Expected Gradient and Deviation From it

Using independence of $\mathbf{A}_k$ and $\{\mathbf{U}, \mathbf{b}_k\}$ (due to sample splitting),

$$\mathbb{E}[\text{GradU}] = \sum_k m(\mathbf{x}_k - \mathbf{x}_k^*) \mathbf{b}_k^\top$$

Using bounds on $\|\mathbf{B}\|$ and $\|\mathbf{X}^* - \mathbf{X}\|_F$ from Theorem 4, if $\delta_t < \frac{c}{\sqrt{r\kappa}}$,

$$\begin{aligned} \|\mathbb{E}[\text{GradU}]\| &= \left\| \sum_k m(\mathbf{x}_k - \mathbf{x}_k^*) \mathbf{b}_k^\top \right\| = m \|(\mathbf{X} - \mathbf{X}^*) \mathbf{B}^\top \| \\ &\leq m \|\mathbf{X} - \mathbf{X}^*\| \cdot \|\mathbf{B}\| \\ &\leq m \|\mathbf{X} - \mathbf{X}^*\|_F \cdot \|\mathbf{B}\| \leq 1.1 m \delta_t \sqrt{r} \sigma_{\max}^*{}^2. \end{aligned}$$

w.p. $1 - \exp(\log q + r - cm)$.

Next, we bound $\|\text{GradU} - \mathbb{E}[\text{GradU}]\| = \max_{\|\mathbf{w}\|=1, \|\mathbf{z}\|=1} \mathbf{w}^\top (\sum_k \sum_i \mathbf{a}_{ki} \mathbf{a}_{ki}^\top (\mathbf{x}_k - \mathbf{x}_k^*) \mathbf{b}_k^\top - \mathbb{E}[\cdot]) \mathbf{z}$. This also uses independence of $\mathbf{A}_k$ and $\{\mathbf{U}, \mathbf{b}_k\}$.

We bound the above for fixed unit norm $\mathbf{w}, \mathbf{z}$ using sub-exponential Bernstein inequality (Theorem 2.8.1 of [10]). We extend the bound to all unit norm $\mathbf{w}, \mathbf{z}$ by using a standard epsilon-net argument. For fixed unit norm $\mathbf{w}, \mathbf{z}$, consider

$$\sum_k \sum_i \left((\mathbf{w}^\top \mathbf{a}_{ki}) (\mathbf{b}_k^\top \mathbf{z}) \mathbf{a}_{ki}^\top (\mathbf{x}_k - \mathbf{x}_k^*) - \mathbb{E}[\cdot] \right)$$

Observe that the summands are independent, zero mean, sub-exponential r.v.s with sub-exponential norm $K_{ki} \leq C \|\mathbf{w}\| \cdot |\mathbf{z}^\top \mathbf{b}_k| \cdot \|\mathbf{x}_k - \mathbf{x}_k^*\| = C |\mathbf{z}^\top \mathbf{b}_k| \cdot \|\mathbf{x}_k - \mathbf{x}_k^*\|$. We apply the sub-exponential Bernstein inequality, Theorem 2.8.1 of [10], with $t = \epsilon_1 \delta_t m \sigma_{\min}^*{}^2$. We have

$$\begin{aligned} \frac{t^2}{\sum_{ki} K_{ki}^2} &\geq \frac{\epsilon_1^2 \delta_t^2 m^2 \sigma_{\min}^*{}^4}{m \sum_k \|\mathbf{x}_k - \mathbf{x}_k^*\|^2 (\mathbf{z}^\top \mathbf{b}_k)^2} \\ &\geq \frac{\epsilon_1^2 \delta_t^2 m^2 \sigma_{\min}^*{}^4}{m \max_k \|\mathbf{x}_k - \mathbf{x}_k^*\|^2 \sum_k (\mathbf{z}^\top \mathbf{b}_k)^2} \\ &= \frac{\epsilon_1^2 \delta_t^2 m^2 \sigma_{\min}^*{}^4}{m \max_k \|\mathbf{x}_k - \mathbf{x}_k^*\|^2 \|\mathbf{z}^\top \mathbf{B}\|^2} \\ &\geq \frac{\epsilon_1^2 \delta_t^2 m^2 \sigma_{\min}^*{}^4 q}{1.4^2 m \mu^2 \delta_t^2 r \sigma_{\max}^*{}^2 \|\mathbf{B}\|^2} \\ &\geq \frac{\epsilon_1^2 \delta_t^2 m^2 \sigma_{\min}^*{}^4 q}{1.4^2 m \mu^2 \delta_t^2 r \sigma_{\max}^*{}^2 1.1 \sigma_{\max}^*{}^2} = c \frac{\epsilon_1^2 m q}{\kappa^4 \mu^2 r} \\ \frac{t}{\max_{ki} K_{ki}} &\geq \frac{\epsilon_1 \delta_t m \sigma_{\min}^*{}^2}{\max_k \|\mathbf{b}_k\| \max_k \|\mathbf{x}_k - \mathbf{x}_k^*\|} \geq c \frac{\epsilon_1 m q}{\kappa^2 \mu^2 r} \end{aligned}$$

In the above, we used (i) $\sum_k (\mathbf{z}^\top \mathbf{b}_k)^2 = \|\mathbf{z}^\top \mathbf{B}\|^2 \leq \|\mathbf{B}\|^2$ since $\mathbf{z}$ is unit norm, (ii) Theorem 4 to bound $\|\mathbf{B}\| \leq 1.1 \sigma_{\max}^*$, and (iii) Theorem 4 followed by Assumption 1 (right incoherence) to bound $\|\mathbf{x}_k - \mathbf{x}_k^*\| \leq \delta_t \cdot \mu \sigma_{\max}^* \sqrt{r/q}$ and $|\mathbf{z}^\top \mathbf{b}_k| \leq \|\mathbf{b}_k\| \leq 1.1 \|\mathbf{b}_k^*\| \leq 1.1 \mu \sigma_{\max}^* \sqrt{r/q}$.

For $\epsilon_1 < 1$, the first term above is smaller (since $1/\kappa^4 \leq 1/\kappa^2$), i.e., $\min(\frac{t^2}{\sum_{ki} K_{ki}^2}, \frac{t}{\max_{ki} K_{ki}}) = c \frac{\epsilon_1^2 m q}{\kappa^4 \mu^2 r}$. Thus, by sub-exponential Bernstein, w.p. at least $1 - \exp(-c \frac{\epsilon_1^2 m q}{\kappa^4 \mu^2 r}) - \exp(\log q + r - cm)$, for a given $\mathbf{w}, \mathbf{z}$,

$$\mathbf{w}^\top (\text{GradU} - \mathbb{E}[\text{GradU}]) \mathbf{z} \leq \epsilon_1 \delta_t m \sigma_{\min}^*{}^2$$

Using a standard epsilon-net argument to bound the maximum of the above over all unit norm $\mathbf{w}, \mathbf{z}$, e.g., using [3, Proposition 4.7], we can conclude that

$$\|\text{GradU} - \mathbb{E}[\text{GradU}]\| \leq 1.1 \epsilon_1 \delta_t m \sigma_{\min}^*{}^2$$

w.p. at least $1 - \exp(C(n+r) - c \frac{\epsilon_1^2 m q}{\kappa^4 \mu^2 r}) - \exp(\log q + r - cm)$. The factor of $\exp(C(n+r))$ is due to the epsilon-net over $\mathbf{w}$ and that over $\mathbf{z}$: $\mathbf{w}$ is an n -length unit norm vector while $\mathbf{z}$ is an r -length unit norm vector. The smallest epsilon net covering the hyper-sphere of all $\mathbf{w}$ s is of size $(1 + 2/\epsilon_{\text{net}})^n = C^n$ with $\epsilon_{\text{net}} = c$ while that for $\mathbf{z}$ is of size C^r . Union bounding over both thus gives a factor of C^{n+r} . By replacing ϵ_1 by $\epsilon_1/1.1$, our bound becomes simpler (and $1/1.1^2$ gets incorporated into the factor c). We have thus proved the following.

Lemma 5: Assume that $\text{SD}_2(\mathbf{U}, \mathbf{U}^*) \leq \delta_t$. The following hold:

- 1) $\mathbb{E}[\text{GradU}] = m(\mathbf{X} - \mathbf{X}^*) \mathbf{B}^\top = m(\mathbf{U} \mathbf{B} \mathbf{B}^\top - \mathbf{X}^* \mathbf{B}^\top)$
- 2) $\|\mathbb{E}[\text{GradU}]\| \leq 1.1 m \delta_t \sqrt{r} \sigma_{\max}^*{}^2$
- 3) If $\delta_t < \frac{c}{\sqrt{r\kappa}}$, then, w.p. at least $1 - \exp(C(n+r) - c \frac{\epsilon_1^2 m q}{\kappa^4 \mu^2 r}) - \exp(\log q + r - cm)$,

$$\|\text{GradU} - \mathbb{E}[\text{GradU}]\| \leq \epsilon_1 \delta_t m \sigma_{\min}^*{}^2$$

The above lemma is an improvement over the bounds given in [3] because δ_t is now the bound on the 2-norm SD, and still it only needs $m q \gtrsim nr/\epsilon_1^2$.

D. GD Step

Assume that $\text{SD}_2(\mathbf{U}, \mathbf{U}^*) \leq \delta_t$ with $\delta_t < 0.02$.

Recall the Projected GD step for $\mathbf{U}$:

$$\tilde{\mathbf{U}}^+ = \mathbf{U} - \eta \text{GradU} \text{ and } \tilde{\mathbf{U}}^+ \stackrel{\text{QR}}{=} \mathbf{U}^+ \mathbf{R}^+$$

Since $\mathbf{U}^+ = \tilde{\mathbf{U}}^+ (\mathbf{R}^+)^{-1}$ and since $\|(\mathbf{R}^+)^{-1}\| = 1/\sigma_{\min}(\mathbf{R}^+) = 1/\sigma_{\min}(\tilde{\mathbf{U}}^+)$, thus, $\text{SD}_2(\mathbf{U}^+, \mathbf{U}^*) = \|\mathbf{P} \mathbf{U}^+\|$ can be bounded as

$$\text{SD}_2(\mathbf{U}^+, \mathbf{U}^*) \leq \frac{\|\mathbf{P} \tilde{\mathbf{U}}^+\|}{\sigma_{\min}(\tilde{\mathbf{U}}^+)} \leq \frac{\|\mathbf{P} \tilde{\mathbf{U}}^+\|}{\sigma_{\min}(\mathbf{U}) - \eta \|\text{GradU}\|} \quad (2)$$

Consider the numerator. Adding/subtracting $\mathbb{E}[\text{GradU}]$, left multiplying both sides by $\mathbf{P}$, and using Lemma 5 (first part),

$$\begin{aligned} \tilde{\mathbf{U}}^+ &= \mathbf{U} - \eta \mathbb{E}[\text{GradU}] + \eta (\mathbb{E}[\text{GradU}] - \text{GradU}), \text{ thus,} \\ \mathbf{P} \tilde{\mathbf{U}}^+ &= \mathbf{P} \mathbf{U} - \eta m \mathbf{P} \mathbf{U} \mathbf{B} \mathbf{B}^\top + \eta \mathbf{P} (\mathbb{E}[\text{GradU}] - \text{GradU}) \end{aligned}$$

The last row used $\mathbf{P} \mathbf{X}^* = 0$. Thus,

$$\|\mathbf{P} \tilde{\mathbf{U}}^+\| \leq \|\mathbf{P} \mathbf{U}\| \|\mathbf{I} - \eta m \mathbf{B} \mathbf{B}^\top\| + \eta \|\mathbb{E}[\text{GradU}] - \text{GradU}\| \quad (3)$$

Using Theorem 4, we get

$$\lambda_{\min}(\mathbf{I} - \eta m \mathbf{B} \mathbf{B}^\top) = 1 - \eta m \|\mathbf{B}\|^2 \geq 1 - 1.2 \eta m \sigma_{\max}^*{}^2$$

Thus, if $\eta < 0.5/m \sigma_{\max}^*{}^2$, then the above matrix is p.s.d. This along with Theorem 4 then implies that

$$\|\mathbf{I} - \eta m \mathbf{B} \mathbf{B}^\top\| = \lambda_{\max}(\mathbf{I} - \eta m \mathbf{B} \mathbf{B}^\top) \leq 1 - 0.9 \eta m \sigma_{\min}^*{}^2$$

Using the above, (3), and the bound on $\|\mathbb{E}[\text{GradU}] - \text{GradU}\|$ from Lemma 5, we conclude the following: If $\eta \leq 0.5/m\sigma_{\max}^*$, and $\delta_t \leq c/\sqrt{r}\kappa$, then

$$\|\mathbf{P}\tilde{\mathbf{U}}^+\| \leq \|\mathbf{P}\mathbf{U}\| \|\mathbf{I} - \eta\mathbf{m}\mathbf{B}\mathbf{B}^\top\| + \eta\|\mathbb{E}[\text{GradU}] - \text{GradU}\| \leq \delta_t(1 - 0.9\eta m\sigma_{\min}^*) + \eta m\epsilon_1\delta_t\sigma_{\min}^* \quad (4)$$

w.p. at least $1 - \exp(-C(n+r) - c\frac{\epsilon_1^2 m q}{\kappa^4 \mu^2 r}) - \exp(\log q + r - cm)$. This probability is at least $1 - n^{-10}$ if $m q \gtrsim \kappa^4 \mu^2 n r / \epsilon_1^2$ and $m \gtrsim \max(\log n, \log q, r)$.

Next we use (4) with $\epsilon_1 = 0.1$ and Lemma 5 to bound the right hand side of (2). Set $\eta = c_\eta/m\sigma_{\max}^*$. If $c_\eta \leq 0.5$, if $\delta_t \leq c/\sqrt{r}\kappa^2$, and lower bounds on m from above hold, (2) implies that, whp,

$$\begin{aligned} \text{SD}_2(\mathbf{U}^+, \mathbf{U}^*) &\leq \frac{\|\mathbf{P}\tilde{\mathbf{U}}^+\|}{\sigma_{\min}(\tilde{\mathbf{U}}^+)} \\ &\leq \frac{\|\mathbf{P}\tilde{\mathbf{U}}^+\|}{\sigma_{\min}(\mathbf{U}) - \eta\|\text{GradU}\|} \\ &\leq \frac{\|\mathbf{P}\mathbf{U}\| \|\mathbf{I} - \eta\mathbf{m}\mathbf{B}\mathbf{B}^\top\| + \eta\|\mathbb{E}[\text{GradU}] - \text{GradU}\|}{1 - \eta\|\mathbb{E}[\text{GradU}]\| - \eta\|\text{GradU} - \mathbb{E}[\text{GradU}]\|} \\ &\leq \frac{\delta_t(1 - \eta m\sigma_{\min}^*(0.9 - 0.1))}{1 - \eta\|\mathbb{E}[\text{GradU}]\| - \eta\|\text{GradU} - \mathbb{E}[\text{GradU}]\|} \\ &\leq \frac{\delta_t(1 - 0.8\eta m\sigma_{\min}^*)}{1 - \eta m\delta_t\sqrt{r}\sigma_{\max}^*(1.4 + \frac{0.1}{\kappa^2\sqrt{r}})} \\ &\leq \frac{\delta_t(1 - 0.8\eta m\sigma_{\min}^*)}{1 - 1.5\eta m\delta_t\sqrt{r}\sigma_{\max}^*} \\ &\leq \delta_t(1 - 0.8\eta m\sigma_{\min}^*)(1 + 1.5\eta m\delta_t\sqrt{r}\sigma_{\max}^*) \\ &\leq \delta_t(1 - 0.8\eta m\sigma_{\min}^* + 1.5\eta m\delta_t\sqrt{r}\sigma_{\max}^*) \\ &= \delta_t(1 - \eta m\sigma_{\min}^*(0.8 - 1.5\delta_t\sqrt{r}\kappa^2)) \\ &\leq \delta_t(1 - \eta m\sigma_{\min}^*(0.8 - 0.15)) \\ &\leq \delta_t(1 - 0.6\eta m\sigma_{\min}^*/\kappa^2) = \delta_t(1 - 0.6c_\eta/\kappa^2) \end{aligned}$$

In the above we used $\kappa^2\sqrt{r} > 1$, $(1-x)^{-1} < (1+x)$ if $|x| < 1$, and $\delta_t < 0.1/\sqrt{r}\kappa^2$ (used in second-last inequality).

Thus, we have proved the following result.

Theorem 6: Assume that Assumption 1 holds and $\text{SD}_2(\mathbf{U}, \mathbf{U}^*) \leq \delta_t$. If $\delta_t \leq 0.02/\sqrt{r}\kappa^2$, if $\eta = c_\eta/m\sigma_{\max}^*$ with $c_\eta \leq 0.5$, and if $m q \gtrsim \kappa^4 \mu^2 n r$ and $m \gtrsim \max(\log n, \log q, r)$, then, whp

$$\text{SD}_2(\mathbf{U}^+, \mathbf{U}^*) \leq \delta_{t+1} := \delta_t(1 - 0.6c_\eta/\kappa^2)$$

V. NEW PROOF: INITIALIZATION

We need a new result for the initialization step because we need a tight bound on $\text{SD}_2(\mathbf{U}_0, \mathbf{U}^*)$.

A. Results Taken From [3]

Recall from Algorithm 1 that α uses a different set of measurements that is independent of those used for $\mathbf{X}_0$. We use the following four results from [3].

Lemma 7 [3]: Conditioned on α , we have the following conclusions. Let ζ be a scalar standard Gaussian r.v.. Define

$$\beta_k(\alpha) := \mathbb{E}[\zeta^2 \mathbb{1}_{\{\|\mathbf{x}_k^*\|^2 \zeta^2 \leq \alpha\}}].$$

Then,

$$\mathbb{E}[\mathbf{X}_0|\alpha] = \mathbf{X}^* \mathbf{D}(\alpha),$$

$$\text{where } \mathbf{D}(\alpha) := \text{diagonal}(\beta_k(\alpha), k \in [q])$$

i.e. $\mathbf{D}(\alpha)$ is a diagonal matrix of size $q \times q$ with diagonal entries β_k defined above.

Fact 8 [3]: Let

$$\mathcal{E} := \left\{ \tilde{C}(1 - \epsilon_1) \frac{\|\mathbf{X}^*\|_F^2}{q} \leq \alpha \leq \tilde{C}(1 + \epsilon_1) \frac{\|\mathbf{X}^*\|_F^2}{q} \right\}.$$

$\Pr(\alpha \in \mathcal{E}) \geq 1 - \exp(-\tilde{c}mq\epsilon_1^2)$. Here $\tilde{c} = c/\tilde{C} = c/\kappa^2\mu^2$.

Fact 9 [3]: For any $\epsilon_1 \leq 0.1$,

$$\min_k \mathbb{E} \left[\zeta^2 \mathbb{1}_{\left\{ |\zeta| \leq \tilde{C} \frac{\sqrt{1-\epsilon_1} \|\mathbf{x}^*\|_F}{\sqrt{q} \|\mathbf{x}_k^*\|} \right\}} \right] \geq 0.92.$$

Lemma 10 [3]: Fix $0 < \epsilon_1 < 1$. Then, w.p. at least $1 - \exp[-(n+q) - c\epsilon_1^2 m q / \mu^2 \kappa^2]$, conditioned on α , for an $\alpha \in \mathcal{E}$,

$$\|\mathbf{X}_0 - \mathbb{E}[\mathbf{X}_0|\alpha]\| \leq 1.1\epsilon_1 \|\mathbf{X}^*\|_F$$

Facts 8 and 9 together imply that, w.p. at least $1 - \exp(-\tilde{c}mq\epsilon_1^2)$,

$$\min_k \beta_k(\alpha) \geq \min_k \mathbb{E} \left[\zeta^2 \mathbb{1}_{\left\{ |\zeta| \leq \tilde{C} \frac{\sqrt{1-\epsilon_1} \|\mathbf{x}^*\|_F}{\sqrt{q} \|\mathbf{x}_k^*\|} \right\}} \right] \geq 0.92. \quad (5)$$

The first inequality is an immediate consequence of Fact 8 ($\alpha \in \mathcal{E}$) and the second follows by Fact 9.

By setting $\epsilon_1 = \epsilon_0/1.1\sqrt{r}\kappa$ in Lemma 10, and using Fact 8, we get the following corollary.

Corollary 11: Fix $0 < \epsilon_1 < 1$. Then, w.p. at least $1 - \exp[(n+q) - c\frac{\epsilon_0^2 m q}{\mu^2 \kappa^4 r}] - \exp(-cmq\epsilon_1^2/\kappa^2\mu^2)$,

$$\|\mathbf{X}_0 - \mathbb{E}[\mathbf{X}_0|\alpha]\| \leq \epsilon_0\sigma_{\min}^*$$

with $\mathbb{E}[\mathbf{X}_0|\alpha]$ being as given on Lemma 7.

B. Obtaining the SD Error Bound for Initialization

By Lemma 7, $\mathbb{E}[\mathbf{X}_0|\alpha] = \mathbf{X}^* \mathbf{D}(\alpha)$ with $\mathbf{D}(\alpha)$ as defined there. Clearly, its rank is r or less. To obtain the bound, we apply Wedin's $\sin \theta$ theorem [11] for SD_2 with $\mathbf{M} = \mathbf{X}_0$, $\mathbf{M}^* = \mathbb{E}[\mathbf{X}_0|\alpha] = \mathbf{X}^* \mathbf{D}$. Recall that $\mathbf{U}_0$ is the matrix of top r singular vectors of $\mathbf{X}_0$. Also, the span of top r singular vectors of $\mathbb{E}[\mathbf{X}_0|\alpha] = \mathbf{X}^* \mathbf{D}$ equals the column span of $\mathbf{U}^*$. Thus applying Wedin will help us bound $\text{SD}_2(\mathbf{U}_0, \mathbf{U}^*)$. To do this, we need to define the SVD of $\mathbb{E}[\mathbf{X}_0|\alpha]$. Let

$$\mathbb{E}[\mathbf{X}_0|\alpha] = \mathbf{X}^* \mathbf{D}(\alpha) \stackrel{\text{SVD}}{=} (\mathbf{U}^* \mathbf{Q}) \tilde{\Sigma}^* \tilde{\mathbf{V}}$$

where $\mathbf{Q}$ is a $r \times r$ unitary matrix, $\tilde{\Sigma}^*$ is an $r \times r$ diagonal matrix with non-negative entries (singular values) and $\tilde{\mathbf{V}}$ is an $r \times q$ matrix with orthonormal rows, i.e. $\tilde{\mathbf{V}} \tilde{\mathbf{V}}^\top = \mathbf{I}$. Observe also that $\sigma_{r+1}(\mathbb{E}[\mathbf{X}_0|\alpha]) = 0$ since it is a rank r matrix. Also, from above,

$$\sigma_r(\mathbb{E}[\mathbf{X}_0|\alpha]) = \sigma_{\min}(\tilde{\Sigma}^*) \geq \sigma_{\min}^* \sigma_{\min}(\mathbf{D}) = \sigma_{\min}^* \min_k \beta_k(\alpha)$$

This follows since $\check{\Sigma}^* = \mathbf{Q}^\top \mathbf{B}^* \mathbf{D} \check{\mathbf{V}}^\top$ and so $\sigma_{\min}(\check{\Sigma}^*) \geq \sigma_{\min}(\mathbf{B}^* \mathbf{D} \check{\mathbf{V}}^\top) \geq \sigma_{\min}(\mathbf{B}^* \mathbf{D}) \cdot 1 \geq \sigma_{\min}(\mathbf{D} \mathbf{B}^*) \geq \sigma_{\min}(\mathbf{D}) \cdot \sigma_{\min}^*$ and $\sigma_{\min}(\mathbf{D}) = \min_k \beta_k$. The last equality follows since $\mathbf{D}$ is diagonal with entries β_k . Thus,

Fact 12: $\sigma_r(\mathbb{E}[\mathbf{X}_0|\alpha]) \geq \sigma_{\min}^* \min_k \beta_k(\alpha)$, $\sigma_{r+1}(\mathbb{E}[\mathbf{X}_0|\alpha]) = 0$.

Applying Wedin's $\sin \theta$ theorem, and then using Corollary 11, Fact 12 and (5), we get

$$\begin{aligned} \text{SD}_2(\mathbf{U}_0, \mathbf{U}^*) &\leq \sqrt{2} \frac{\max(\|(\mathbf{X}_0 - \mathbb{E}[\mathbf{X}_0|\alpha])^\top \mathbf{U}^*\|, \|(\mathbf{X}_0 - \mathbb{E}[\mathbf{X}_0|\alpha]) \check{\mathbf{V}}\|)}{\sigma_{\min}^* \min_k \beta_k(\alpha) - 0 - \|\mathbf{X}_0 - \mathbb{E}[\mathbf{X}_0|\alpha]\|} \\ &\lesssim \sqrt{2} \frac{\epsilon_0 \sigma_{\min}^*}{0.92 \sigma_{\min}^* - \epsilon_0 \sigma_{\min}^*} \lesssim \sqrt{2} \frac{\epsilon_0}{0.9} < 1.6 \epsilon_0 \end{aligned}$$

if $\epsilon_0 < 0.02$ and $m q \gtrsim (n+q)r/\epsilon_0^2$. In the above we used $\|(\mathbf{X}_0 - \mathbb{E}[\mathbf{X}_0|\alpha])^\top \mathbf{U}^*\| \leq \|\mathbf{X}_0 - \mathbb{E}[\mathbf{X}_0|\alpha]\| \cdot 1$ and $\|(\mathbf{X}_0 - \mathbb{E}[\mathbf{X}_0|\alpha]) \check{\mathbf{V}}\| \leq \|\mathbf{X}_0 - \mathbb{E}[\mathbf{X}_0|\alpha]\| \cdot 1$. Setting $\epsilon_0 = 0.5 \delta_0$ with $\delta_0 < 0.02$, we obtain our desired result.

Theorem 13: Assume that Assumption 1 holds. Pick a $\delta_0 \leq 0.02$. If $m q \gtrsim \kappa^4 \mu^2 (n+q)r/\delta_0^2$ then whp

$$\text{SD}_2(\mathbf{U}_0, \mathbf{U}^*) \leq \delta_0$$

APPENDIX A PROOF OF THEOREM 2

Theorem 13 tells us that $\text{SD}_2(\mathbf{U}_0, \mathbf{U}^*) \leq \delta_0$ whp if $m_0 q \gtrsim (n+q)r/\delta_0^2$. Theorem 6 tells us that if $\delta_t \leq 0.02/\sqrt{r}\kappa^2$, and if $m_1 q \gtrsim nr$, then, whp, δ_t reduces by a factor of $(1 - 0.6c_\eta/\kappa^2)$ at each iteration. In particular, this implies that $\delta_t \leq \delta_0$. Thus, in order to apply Theorem 6, it suffices to require $\delta_0 = 0.02/\sqrt{r}\kappa^2$.

Combining both results, we have shown that if $m_0 q \geq C\kappa^8 \mu^2 (n+q)r^2$ and if $m_1 q \geq C\kappa^4 \mu^2 nr$ and $m_1 \geq C \max(r, \log n, \log r)$, and if $\eta = c_\eta/(m\sigma_{\max}^*)^2$ with $c_\eta < 0.5$, then, whp, at each $t \geq 0$,

$$\text{SD}_2(\mathbf{U}_t, \mathbf{U}^*) \leq \delta_t := (1 - 0.6 \frac{c_\eta}{\kappa^2})^t \delta_0 = (1 - 0.6 \frac{c_\eta}{\kappa^2})^t \frac{0.02}{\sqrt{r}\kappa^2}$$

and all bounds of Theorem 4 hold with δ_t as above.

Thus, to guarantee $\text{SD}_2(\mathbf{U}_T, \mathbf{U}^*) \leq \epsilon$, we need

$$T \geq C \frac{\kappa^2}{c_\eta} \log(1/\epsilon)$$

This follows by using $\log(1-x) < -x$ for $|x| < 1$ and using $\kappa^2 \sqrt{r} \geq 1$. Thus, setting $c_\eta = 0.4$, our sample complexity $m = m_0 + m_1 T$ becomes $m q \geq C\kappa^8 \mu^2 (n+q)r(1 + \log(1/\epsilon))$, and $m \geq C \max(r, \log q, \log n) \log(1/\epsilon)$.

APPENDIX B PROOF OF LEMMA 3

Using the expression for $\mathbf{b}_k$ and simplifying it,

$$\mathbf{b}_k - \mathbf{g}_k = \mathbf{M}^{-1} \mathbf{U}^\top \mathbf{A}_k^\top \mathbf{A}_k (\mathbf{I} - \mathbf{U} \mathbf{U}^\top) \mathbf{U}^* \mathbf{b}_k^*$$

with $\mathbf{M} := \mathbf{U}^\top \mathbf{A}_k^\top \mathbf{A}_k \mathbf{U}$.

Clearly,

$$\mathbb{E}[\mathbf{M}] = m \mathbf{U}^\top \mathbf{U} = m \mathbf{I}_r, \quad \mathbb{E}[\mathbf{U}^\top \mathbf{A}_k^\top \mathbf{A}_k (\mathbf{I} - \mathbf{U} \mathbf{U}^\top) \mathbf{U}^* \mathbf{b}_k^*] = 0$$

Thus, using the standard technique (sub-expo Bern ineq followed by an epsilon-net argument), one can show that

- 1) w.p. at least $1 - \exp(r - \epsilon_2 m)$,

$$\|\mathbf{M} - \mathbf{I}_r\| \leq 1.1 \epsilon_2 m$$

This implies of course that $\sigma_{\min}^*(\mathbf{M}) \geq (1 - 1.1 \epsilon_2)m$ and thus $\|\mathbf{M}^{-1}\| \leq 1/((1 - 1.1 \epsilon_2)m)$.

- 2) w.p. at least $1 - \exp(r - \epsilon_3 m)$,

$$\begin{aligned} \|\mathbf{U}^\top \mathbf{A}_k^\top \mathbf{A}_k (\mathbf{I} - \mathbf{U} \mathbf{U}^\top) \mathbf{U}^* \mathbf{b}_k^*\| &\leq 1.1 \epsilon_3 m \|(\mathbf{I} - \mathbf{U} \mathbf{U}^\top) \mathbf{U}^* \mathbf{b}_k^*\| \end{aligned}$$

Thus, setting $\epsilon_2 = 0.1$, and taking union bound over all q vectors, w.p. at least $1 - q \exp(r - \epsilon_3 m)$,

$$\|\mathbf{b}_k - \mathbf{g}_k\| \leq 1.2 \epsilon_3 \|(\mathbf{I} - \mathbf{U} \mathbf{U}^\top) \mathbf{U}^* \mathbf{b}_k^*\| \text{ for all } k \in [q]$$

REFERENCES

- [1] S. Babu, S. G. Lingala, and N. Vaswani, "Fast low rank compressive sensing for accelerated dynamic MRI," *IEEE Trans. Comput. Imag.*, vol. 9, pp. 409–424, 2023, doi: [10.1109/TCI.2023.3263810](https://doi.org/10.1109/TCI.2023.3263810).
- [2] S. S. Du, W. Hu, S. M. Kakade, J. D. Lee, and Q. Lei, "Few-shot learning via learning the representation, provably," in *Proc. Int. Conf. Learn. Represent. (ICLR)*, 2021.
- [3] S. Nayer and N. Vaswani, "Fast and sample-efficient federated low rank matrix recovery from column-wise linear and quadratic projections," *IEEE Trans. Inf. Theory*, vol. 69, no. 2, pp. 1177–1202, Feb. 2023.
- [4] P. Jain, P. Netrapalli, and S. Sanghavi, "Low-rank matrix completion using alternating minimization," in *Proc. 45th Annu. ACM Symp. Theory Comput.*, Jun. 2013.
- [5] X. Yi, D. Park, Y. Chen, and C. Caramanis, "Fast algorithms for robust PCA via gradient descent," in *Proc. NeurIPS*, 2016.
- [6] S. Lang, *Real and Functional Analysis*, vol. 10. New York, NY, USA: Springer, 1993, pp. 11–13.
- [7] C. Ma, K. Wang, Y. Chi, and Y. Chen, "Implicit regularization in nonconvex statistical estimation: Gradient descent converges linearly for phase retrieval, matrix completion and blind deconvolution," in *Proc. Int. Conf. Mach. Learn. (ICML)*, 2018.
- [8] S. Nayer and N. Vaswani, "Sample-efficient low rank phase retrieval," in *Proc. IEEE Int. Symp. Inf. Theory (ISIT)*, Jul. 2021, pp. 2244–2249.
- [9] Y. Cherapanamjeri, K. Gupta, and P. Jain, "Nearly-optimal robust matrix completion," in *Proc. Int. Conf. Mach. Learn. (ICML)*, 2016.
- [10] R. Vershynin, *High-Dimensional Probability: An Introduction With Applications in Data Science*, vol. 47. Cambridge, U.K.: Cambridge Univ. Press, 2018.
- [11] P. Wedin, "Perturbation bounds in connection with singular value decomposition," *BIT Numer. Math.*, vol. 12, no. 1, pp. 99–111, 1972.

Namrata Vaswani (Fellow, IEEE) received the B.Tech. degree from IIT Delhi, India, in 1999, and the Ph.D. degree from the University of Maryland, College Park, in 2004.

She is currently the Anderlik Professor in electrical and computer engineering with Iowa State University, where she is also the Director of the CyMath Program, in which graduate students provided school-year-long math tutoring for under-served grade school students. Her research interests are in data science, with a particular focuses on statistical machine learning for signal processing and computational imaging. She was a recipient of the 2014 IEEE Signal Processing Society Best Paper Award, the Iowa State Early Career Engineering Faculty Research Award in 2014, and the Iowa State Mid-Career Achievement in Research Award in 2019. She has served as an Associate Editor or an Area Editor for IEEE TRANSACTIONS ON INFORMATION THEORY, IEEE TRANSACTIONS ON SIGNAL PROCESSING, and the *Signal Processing Magazine*. She has also guest-edited special issues for PROCEEDINGS OF THE IEEE and the IEEE JOURNAL OF SELECTED TOPICS IN SIGNAL PROCESSING.