
Learning quadratic neural networks in high dimensions: SGD dynamics and scaling laws

G rard Ben Arous¹, Murat A. Erdogdu^{2,3}, Nuri Mert Vural^{2,3}, Denny Wu^{1,4}

¹New York University, ²University of Toronto, ³Vector Institute, ⁴Flatiron Institute
gba@cims.nyu.edu, {erdogdu,vural}@cs.toronto.edu, dennywu@nyu.edu

Abstract

We study the optimization and sample complexity of gradient-based training of a two-layer neural network with quadratic activation function in the high-dimensional regime, where the data is generated as $f_*(\mathbf{x}) \propto \sum_{j=1}^r \lambda_j \sigma(\langle \boldsymbol{\theta}_j, \mathbf{x} \rangle)$, $\mathbf{x} \sim \mathcal{N}(0, \mathbf{I}_d)$, σ is the 2nd Hermite polynomial, and $\{\boldsymbol{\theta}_j\}_{j=1}^r \subset \mathbb{R}^d$ are orthonormal signal directions. We consider the extensive-width regime $r \asymp d^\beta$ for $\beta \in [0, 1)$, and assume a power-law decay on the (non-negative) second-layer coefficients $\lambda_j \asymp j^{-\alpha}$ for $\alpha \geq 0$. We present a sharp analysis of the SGD dynamics in the feature learning regime, for both the population limit and the finite-sample (on-line) discretization, and derive scaling laws for the prediction risk that highlight the power-law dependencies on the optimization time, sample size, and model width. Our analysis combines a precise characterization of the associated matrix Riccati differential equation with novel matrix monotonicity arguments to establish convergence guarantees for the infinite-dimensional effective dynamics.

1 Introduction

We study the problem of learning a two-layer neural network (NN) with quadratic activation on isotropic Gaussian data. The target function (or the “teacher” model) is defined as

$$y = f_*(\mathbf{x}) + \epsilon \text{ with } f_*(\mathbf{x}) = \frac{1}{\|\boldsymbol{\Lambda}\|_F} \sum_{j=1}^r \lambda_j \sigma(\langle \boldsymbol{\theta}_j, \mathbf{x} \rangle) \text{ and } \mathbf{x} \sim \mathcal{N}(0, \mathbf{I}_d), \quad (1.1)$$

where $\sigma(z) = z^2 - 1$ is the 2nd Hermite polynomial; ϵ is zero-mean, independent sub-Gaussian noise; $\{\boldsymbol{\theta}_j\}_{j=1}^r \subset \mathbb{R}^d$ are unknown signal directions (index features) which we assume to be orthonormal; $\lambda_1 > \lambda_2 > \dots > \lambda_r > 0$ are their respective contributions; and $\boldsymbol{\Lambda} = \text{diag}(\lambda_1, \dots, \lambda_r)$ collects the second-layer coefficients. Our goal is to learn this target network using a “student” two-layer neural network with quadratic activation and r_s neurons, trained via a gradient-based optimization algorithm. This setting encompasses several well-known problems:

- *Phase retrieval* ($r = 1$). The learning of one quadratic neuron has been studied extensively [Fie82, CC15, TV23]. The quadratic σ has information exponent $k = 2$ (defined as the index of the lowest non-zero Hermite coefficient [DH18, BAGJ21]). This entails that randomly initialized parameters are close to a saddle point in high dimensions; hence the SGD dynamics exhibit a plateau (“search” phase) of length $\log d$ before the loss decreases sharply (“descent” phase).
- *Multi-spike PCA* ($r = \Theta_d(1)$). The target (1.1) is a subclass of Gaussian multi-index models, for which various algorithms have been proposed for the finite-rank case $r_s = \Theta_d(1)$ [CM20, DLS22, BBPV23]. The setting also closely relates to the multi-spike PCA problem, for which online SGD [AGP24] and other streaming algorithms has been studied [OK85, JJK⁺16, AZL17].
- *Linear-width quadratic NN* ($r \asymp d$). The regime where the teacher width r grows proportionally with dimensionality d has also been studied, typically in the well-conditioned setting (e.g., identical λ_j ’s). Recent works characterized the landscape [SJL18, DL18, VBB19, GKZ19, GMMM19], optimization dynamics [MVEZ20, MBB23], and statistical efficiency [MTM⁺24, ETZK25, BCN⁺25].

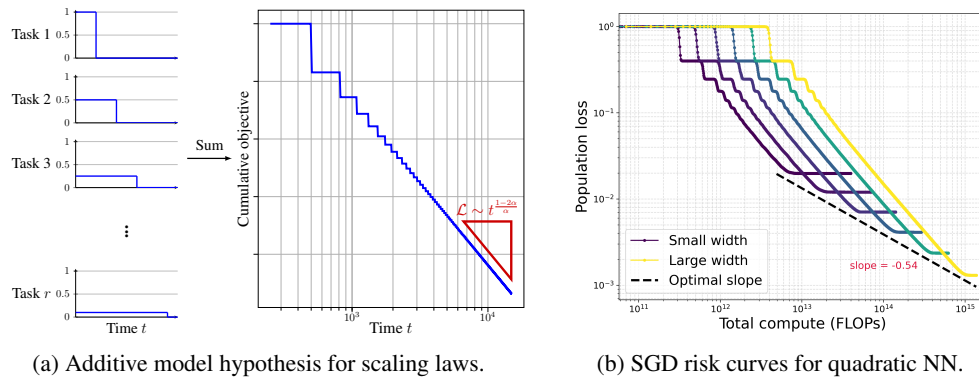


Figure 1: (a) Illustration of the additive model hypothesis, i.e., sum of emergent learning curves at different timescales yields a power law cumulative loss. (b) Population loss vs. compute for two-layer quadratic NNs trained with online SGD with batch size d on MSE loss. We set $d = 3200$, and for the teacher $r = 2400$, $\alpha = 1$.

In this work we focus on the “extensive-rank” regime where $r \asymp d^\beta$ for $\beta \in (0, 1)$ and $r_s \asymp d^\gamma$ for $\gamma \in [0, 1)$, and place a power-law assumption on the second-layer coefficients: $\lambda_j \asymp j^{-\alpha}$ for $\alpha \geq 0$. Our setting is motivated by the following lines of research.

Neural scaling laws & emergence. Recent empirical studies on large language models (LLMs) reveal that increasing the model or training data size often results in a predictable, power-law decrease in the loss known as *neural scaling laws* [HNA⁺17, KMH⁺20, HBM⁺22]. While such scaling of generalization error has been derived for sketched linear models [MRS22, BAP24, PPXP24, LWK⁺24, DLM24], these analyses assume random projection with no *feature learning*, and hence cannot capture the NN’s ability to learn useful features [GDDM14, DCLT18] that adapt to the underlying data structure. We aim to investigate a setting where the training of a nonlinear NN beyond the “lazy” regime exhibits a nontrivial scaling law.

Feature learning in neural networks is often studied theoretically through the learning of *multi-index models*, where the target function depends on a small number of latent directions (see [BH25] and references therein). For these low-dimensional targets, it is known that the training dynamics typically exhibit *emergent* (or staircase-like) behavior — long plateaus followed by sharp drops in loss [BAGJ22, AAM23]. To reconcile this emergent loss curve with smooth power-law decay, recent works hypothesized that the pretraining objective can be decomposed into a sum of losses on individual tasks [MLGT24, NFL24], the learning of each exhibits a sharp transition, and the superposition of numerous emergent risk curves at different timescales yields a power-law scaling of the cumulative loss (see Figure 1(a)). In this context, the two-layer network (1.1) can be viewed as a sum of single-index phase retrieval tasks, where the length of each $\sim \log d$ plateau in the risk trajectory can be modulated by the second-layer coefficient λ_j . This motivates the following question:

Q1: Does gradient-based training of a two-layer quadratic network yield power-law loss scaling, when the target function is an additive model with varying second-layer coefficients $\{\lambda_j\}_{j=1}^r$?

In Figure 1(b) we empirically observe the affirmative: when the target function has smoothly decaying second-layer weights, online SGD training yields a power-law risk curve that resembles the scaling laws in [KMH⁺20, HBM⁺22]. The goal of this work is to rigorously establish such scaling laws.

Learning extensive-width neural networks. Prior works on multi-index models have shown that when $r = \Theta_d(1)$, gradient-based training succeeds with polynomial sample complexity depending on properties of the link function [AAM22, DLS22, BBSS22]. The “extensive-rank” regime where $r \asymp d^\beta$ for $\beta > 0$ is relatively under-explored (except for the linear width regime $r \asymp d$ [MBB23, MTM⁺24]); this setting is arguably closer to the practical neural network training (compared to the narrow-width setting), and also bears connections to several observations in the LLM literature such as *superposition* [EHO⁺22] and *skill localization* [DDH⁺21, WWZ⁺22, PSZA23], where the model simultaneously acquires a large number of “skills” during pretraining (see e.g., [OSSW24]).

The learning dynamics of (1.1) with divergingly many neurons is challenging to analyze primarily due to the fact that the effective dynamics may not be captured by a finite set of *summary statistics* [BAGJ22] (as in the finite- r case). Recent works [OSSW24, SBH24] addressed this challenge by assuming that the activation σ has information exponent $k \geq 3$, which allows the learning dynamics

Algorithm	Decay rate (λ_j)	Risk scaling law	Result
Gradient flow	$\alpha > 0.5$	$\bar{t}^{-\frac{2\alpha-1}{\alpha}} + r_s^{-(2\alpha-1)}$	Theorem 1
	$\alpha < 0.5$	$(1 - \bar{t}^{\frac{1-2\alpha}{\alpha}})_+ + (1 - (r_s/r)^{1-2\alpha})_+$	
Online SGD (Stiefel)	$\alpha > 0.5$	$(\eta\bar{t})^{-\frac{2\alpha-1}{\alpha}} + r_s^{-(2\alpha-1)}$	Theorem 2
	$\alpha < 0.5$	$(1 - (\eta\bar{t})^{\frac{1-2\alpha}{\alpha}})_+ + (1 - (r_s/r)^{1-2\alpha})_+$	

Table 1: Scaling laws for learning quadratic neural network (1.1) using population gradient flow and its online SGD discretization. We omit constant factors in the risk scaling for ease of presentation.

- In $\alpha > 0.5$, for population gradient flow, $\bar{t} \sim t \cdot \log d$ is the rescaled time; for online SGD, $\bar{t} \sim t \cdot \log d$ where t is the number of gradient steps, which is equal to the sample size, and $\eta \sim 1/(d \text{ polylog}(d))$ is the step size.
- In $\alpha < 0.5$, for population gradient flow, $\bar{t} \sim t \cdot r \log d$ is the rescaled time; for online SGD, $\bar{t} \sim t \cdot r \log d$ where t is the number of gradient steps and $\eta \sim 1/(dr^\alpha \text{ polylog}(d))$ is the step size.

to decouple across feature directions. However, the case $k \leq 2$, which includes the quadratic activation studied in this work, remained open: existing analyses either assumed “isotropic” feature contributions ($\lambda_1 = \lambda_r$) [RL24, SBH24], or established a computational complexity for SGD that scales with $d^{\Theta(\lambda_1/\lambda_r)}$ [LMZ20], which leads to pessimistic *exponential* dimension dependency in the power-law setting we consider. We therefore ask the following question.

Q2: *Can we establish optimization and sample complexity of learning an extensive-width quadratic neural network (1.1) with anisotropic, power-decaying feature contributions?*

Finally, although our problem setup does not directly encompass commonly used activation functions such as ReLU, for SGD on multi-index models it is known that the Hermite-2 component is the first harmonic capturing multi-dimensional structure in the target function [DLS22, DKL⁺23b]. Consequently, our quadratic setting represents the first nontrivial mode of feature learning in SGD dynamics. Consistent with this view, Figure 3 shows that the risk trajectory of ReLU networks closely follows our theoretical characterization of quadratic networks over a substantial portion of training.

1.1 Our Contributions

We analyze the risk trajectory of learning (1.1) with both gradient flow on the mean squared error (MSE) loss and its online SGD discretization on Stiefel manifold, covering the extensive-width and power-law settings. We derive scaling laws for feature recovery and population risk as a function of teacher and student network widths r_s, r , the decay exponent α , the optimization time, and the sample size (for the discretized dynamics). Our contributions are summarized as follow (see Table 1).

1. In Section 3, we analyze the population gradient flow and tightly characterize the loss decay with respect to time and the student width r_s . We show that signal directions are recovered sequentially, and the population MSE follows a smooth power law specified by the decay rate $\alpha > 0$.
2. In Section 4, we consider the online stochastic gradient descent (SGD) dynamics on the Stiefel manifold and derive scaling laws with respect to sample size. When specializing to the isotropic setting $\alpha = 0$, our sample complexity improves upon [RL24] in the extensive-width setting and matches the information theoretic limit (in terms of d, r dependence) up to polylogarithmic factors.

The following technical challenges in the extensive-width regime are central to our analysis:

- *Coupled population dynamics.* As $r, r_s \rightarrow \infty$, we must track infinitely many overlapping student and teacher neurons. [OSSW24, SBH24] assumed high information exponent $k > 2$, to decouple the dynamics into r independent single-index models, but such property does not hold in our quadratic case ($k = 2$). We address this by leveraging the closed-form solution of the quadratic problem [MBB23], which satisfies a *Matrix Riccati ODE*. A key ingredient in our analysis is its *monotonicity with respect to its initialization*, illustrated in Figures 4(a), which enables sharp risk bounds via comparisons to decoupled models.
- *Operator norm discretization error.* Prior works [BAGJ21, BBPV23, AGP24] focused on finite- r settings, where Frobenius norm control of the SGD noise was sufficient and natural: it allows bounding error direction-wise without incurring additional dimension dependence. However, in the extensive-width regime, such bounds become pessimistic and lead to suboptimal r -dependent rates. Hence we need to establish *operator norm* concentration around the population dynamics.

- *Matrix-monotone comparison framework.* To control discretization error in operator norm, we extend the monotonicity-based argument to discrete time and introduce a novel comparison-based discretization technique. Our approach constructs matrix-valued reference sequences corresponding to decoupled dynamics that tightly bound the discrete evolution from above and below. This yields sharp operator norm control even when the true trajectories are non-monotone (see Figure 4), as the analysis avoids relying on the trajectory itself by comparing against simpler bounding sequences.

1.2 Additional Related Works

Learning multi-index models with SGD. When $r = 1$, the target is a *single-index model* with quadratic link function. The SGD learning of single-index models has been extensively studied in the feature learning literature [BAGJ21, BES⁺22, MHPG⁺22, BES⁺23, MHWSE23, MLHD23, MZD⁺23, BMZ24, DNGL24, GWB25]; while this model has d parameters to be estimated, the quadratic link (with information exponent $k = 2$) incurs an additional $\log d$ factor in the complexity of online SGD. More generally, the setting where $r = \Theta_d(1)$ is covered by recent analyses of *multi-index models* [AAM22, AAM23, BBPV23, DKL⁺23a, CWPPS23, AGP24, VE24, MHWE24]; however, these learning guarantees for multi-index models typically yield superpolynomial complexity when the target function is rank-extensive. The sample complexity of gradient-based learning is also connected to statistical query lower bounds [DPVLB24, DTA⁺24, LOSW24, ADK⁺24].

Quadratic NNs and additive models. Prior theoretical works on learning two-layer neural network with quadratic activation function have studied the loss landscape [SJL18, DL18, VBB19, GKZ19, GMMM19] and the optimization dynamics [MVEZ20, AKLS23, MBB23, RL24]. While existing optimization and statistical guarantees may cover the extensive-width regime (see e.g., [DL18, MBB23, RL24]), to our knowledge, precise scaling laws have not been established in our extensive-rank and power-law setting. (1.1) is also an instance of the *additive model* [Sto85, HT87, Bac17] where the individual functions are given as (orthogonal) single-index models with *unknown* index features. For this model, [OSSW24, SBH24] established learning guarantees in the well-conditioned regime, under the assumption that the link function σ has information exponent $k > 2$.

2 Background and Problem Setting

2.1 Student-teacher Setting

Teacher Network. We consider the task of learning a teacher network with a quadratic (second-order Hermite) activation function written as

$$y = f_*(\mathbf{x}) + \epsilon \quad \text{with} \quad f_*(\mathbf{x}) := \frac{1}{\|\mathbf{\Lambda}\|_F} \sum_{j=1}^r \lambda_j (\langle \boldsymbol{\theta}_j, \mathbf{x} \rangle^2 - 1) \quad \text{and} \quad \mathbf{x} \sim \mathcal{N}(0, \mathbf{I}_d), \quad (2.1)$$

where $\mathbf{x} \in \mathbb{R}^d$ is the input; ϵ is zero-mean, bounded-variance sub-Gaussian noise; r is the teacher network width; and $\{\boldsymbol{\theta}_j\}_{j=1}^r \subset \mathbb{R}^d$ is an orthonormal set of unknown signal vectors. We collect these as columns of the matrix $\boldsymbol{\Theta} \in \mathbb{R}^{d \times r}$. The contributions of these vectors are determined by the unknown second-layer coefficients $\lambda_1 > \lambda_2 > \dots > \lambda_r > 0$ with a power-law decay $\lambda_j \asymp j^{-\alpha}$ for $\alpha \geq 0$, and $\mathbf{\Lambda}$ is a diagonal matrix whose j -th diagonal entry is λ_j . The normalization in front of summation ensures $\mathbb{E}[y^2]$ is constant. We focus on the regime where $r \asymp d^\beta$ for $\beta \in (0, 1)$.

Remark 1. The orthogonality of $\{\boldsymbol{\theta}_j\}_{j=1}^r$ can be assumed without loss of generality. Specifically, consider (2.1) with arbitrary first-layer weights $\boldsymbol{\Theta}$ and normalization $\mathbb{E}[y] = 0$, the output can be written as $y \propto \text{Tr}(\mathbf{x}\mathbf{x}^\top \boldsymbol{\Theta} \mathbf{\Lambda} \boldsymbol{\Theta}^\top) + \text{constant} + \text{noise}$; hence we may redefine $(\lambda_j, \boldsymbol{\theta}_j)$ via the spectral decomposition.

Student Network. We learn the target model with a quadratic student network defined as

$$\hat{y}(\mathbf{x}, \mathbf{W}) = \frac{1}{\sqrt{r_s}} \sum_{j=1}^{r_s} \langle \mathbf{w}_j, \mathbf{x} \rangle^2 - \|\mathbf{w}_j\|_2^2, \quad (2.2)$$

where r_s is the width of the student network, and $\{\mathbf{w}_j\}_{j=1}^{r_s} \subset \mathbb{R}^d$ denotes the set of trainable weights. We collect these weights as the columns of the matrix $\mathbf{W} \in \mathbb{R}^{d \times r_s}$, and omit the dependence on $\mathbf{x}$

in $\hat{y}(\mathbf{x}, \mathbf{W})$ when clear from the context. Note that the norm subtraction ensures $\mathbb{E}_{\mathbf{x}}[\hat{y}(\mathbf{x}, \mathbf{W})] = 0$. We may equivalently write the student network as $\hat{y}(\mathbf{x}, \mathbf{W}) = \frac{1}{\sqrt{r_s}} \sum_{j=1}^{r_s} \|\mathbf{w}_j\|_2^2 \cdot (\langle \bar{\mathbf{w}}_j, \mathbf{x} \rangle^2 - 1)$ where $\bar{\mathbf{w}}_j$ is unit-norm; since our student does not have trainable second-layer, the norm component $\|\mathbf{w}_j\|_2^2$ allows the model to adapt to the target second-layer λ_j ; this homogeneous parameterization has been studied in prior works [CB20, GRWZ21].

2.2 Training Objective

Training constitutes to minimizing the squared loss; we define the instantaneous loss on $(\mathbf{x}, y)$ as

$$\mathcal{L}(\mathbf{W}; (\mathbf{x}, y)) := \frac{1}{16} (y - \hat{y}(\mathbf{x}, \mathbf{W}))^2,$$

where the prefactor is included for notational convenience in the gradient computation. We omit the dependence on $(\mathbf{x}, y)$ when clear from context. The population risk can be written as

$$R(\mathbf{W}) := \mathbb{E}_{(\mathbf{x}, y)}[\mathcal{L}(\mathbf{W})] = \frac{1}{8} \left\| \frac{1}{\sqrt{r_s}} \mathbf{W} \mathbf{W}^\top - \frac{1}{\|\mathbf{\Lambda}\|_F} \mathbf{\Theta} \mathbf{\Lambda} \mathbf{\Theta}^\top \right\|_F^2 + \frac{1}{16} \mathbb{E}[\epsilon^2]. \quad (2.3)$$

Alignment. Observe that the student network is invariant to right-multiplication of its weight matrix by an orthonormal matrix, i.e., $\hat{y}(\mathbf{x}, \mathbf{W}) = \hat{y}(\mathbf{x}, \mathbf{W} \mathbf{O})$ for any $\mathbf{O} \in \mathbb{R}^{r_s \times r_s}$ with $\mathbf{O}^\top \mathbf{O} = \mathbf{I}$. Consequently, any notion of alignment that depends on individual directions in $\mathbf{W}$ may not be informative. To capture directional learning in a way that respects this symmetry, we define alignment in terms of the subspace spanned by student weights. We formalize this using the polar decomposition:

$$\mathbf{W} := \mathbf{U} \mathbf{Q}^{1/2}, \quad \text{where } \mathbf{Q} := \mathbf{W}^\top \mathbf{W} \quad \text{and} \quad \mathbf{U}^\top \mathbf{U} = \mathbf{I}_{r_s}. \quad (2.4)$$

Here, $\mathbf{Q}$ denotes the radial component of the student weights, while $\mathbf{U}$ is an orthonormal matrix that encodes their directional component. We quantify the alignment between the student network and the j th teacher feature by the squared norm of the projection of $\boldsymbol{\theta}_j$ onto the column space of $\mathbf{W}$:

$$\text{Alignment}(\mathbf{W}, \boldsymbol{\theta}_j) := \|\mathbf{U}^\top \boldsymbol{\theta}_j\|_2^2. \quad (2.5)$$

$\text{Alignment}(\mathbf{W}, \boldsymbol{\theta}_j)$ takes values in the interval $[0, 1]$; it is 0 if $\boldsymbol{\theta}_j$ is orthogonal to $\mathbf{W}$ (no alignment), while it is 1 if $\boldsymbol{\theta}_j$ is in the column space of $\mathbf{W}$ (perfect alignment)¹.

3 Continuous Dynamics: Population Gradient Flow

We first analyze the continuous-time population gradient flow dynamics for (2.3), given as

$$\partial_t \mathbf{W}_t = -\nabla R(\mathbf{W}_t), \quad \text{where } \mathbf{W}_0 \in \mathbb{R}^{d \times r_s}, \quad \mathbf{W}_{0,ij} \sim_{iid} \mathcal{N}(0, 1/d), \quad (\text{GF})$$

and the population gradient reads

$$\nabla R(\mathbf{W}_t) = -\frac{1}{2\sqrt{r_s} \|\mathbf{\Lambda}\|_F} \left(\mathbf{\Theta} \mathbf{\Lambda} \mathbf{\Theta}^\top - \frac{\|\mathbf{\Lambda}\|_F}{\sqrt{r_s}} \mathbf{W}_t \mathbf{W}_t^\top \right) \mathbf{W}_t.$$

For notational convenience, we define

$$\mathcal{R}(t) := \left\| \frac{1}{\sqrt{r_s}} \mathbf{W}_t \mathbf{W}_t^\top - \frac{1}{\|\mathbf{\Lambda}\|_F} \mathbf{\Theta} \mathbf{\Lambda} \mathbf{\Theta}^\top \right\|_F^2 \quad \text{and} \quad \mathcal{A}(t, \boldsymbol{\theta}_j) := \text{Alignment}(\mathbf{W}_t, \boldsymbol{\theta}_j).$$

The following theorem sharply characterizes the timescale for alignment and the limiting risk curve.

Theorem 1. Let $\lambda_j = j^{-\alpha}$ and $r \asymp d^\beta$ for some $\alpha \geq 0$ and $\beta \in (0, 1)$. Consider the regime

$$\begin{cases} \frac{r_s}{r} \rightarrow \varphi \in (0, \infty) \quad \text{and} \quad d \geq \Omega_{\alpha, \beta, \varphi}(1), & \text{if } \alpha \in [0, 0.5), \\ r_s \asymp 1, & \text{and } d \geq \Omega_{\alpha, r_s}(1), \quad \text{if } \alpha > 0.5. \end{cases} \quad (3.1)$$

Define the effective student width and effective timescale as

$$r_{\text{eff}} := \begin{cases} \lfloor r_s (1 - \log^{-1/8} d) \wedge r \rfloor, & \text{if } \alpha \in [0, 0.5) \\ r_s, & \text{if } \alpha > 0.5. \end{cases} \quad \text{and} \quad T_{\text{eff}} := \sqrt{r_s} \|\mathbf{\Lambda}\|_F \log d / r_s.$$

Then, the population (GF) dynamics satisfy the following with probability $1 - o(1/d^2) - \Omega(1/r_s^2)$:

¹The definition in (2.5) may fail to converge to 1 when $\alpha = 0$ and $r_s < r$, due to rotational symmetry in the teacher network. In this case, a more suitable notion of alignment can be defined using the principal angles between the subspaces spanned by $\mathbf{W}$ and $\mathbf{\Theta}$, which provides a rotation-invariant characterization of directional overlap. Specifically, for $\alpha = 0$, we define $\text{Alignment}(\mathbf{W}, \boldsymbol{\theta}_j)$ as the j th largest eigenvalue of $\mathbf{\Theta}^\top \mathbf{U} \mathbf{U}^\top \mathbf{\Theta}$.

1. **Alignment:** For $j \leq r_{\text{eff}}$ and $t > 0$ satisfying $t \asymp r^\alpha$ when $\alpha \in [0, 0.5)$ and $t \asymp 1$ when $\alpha > 0.5$,

$$\mathcal{A}(t\mathbf{T}_{\text{eff}}, \boldsymbol{\theta}_j) = \mathbb{1}\{t \geq \frac{1}{\lambda_j}\} + o_d(1). \quad (3.2)$$

2. **Risk curve:** Under the same time scaling,

$$\mathcal{R}(t\mathbf{T}_{\text{eff}}) = 1 - \frac{1}{\|\boldsymbol{\Lambda}\|_F^2} \sum_{j=1}^{r_{\text{eff}}} \lambda_j^2 \mathbb{1}\{t \geq \frac{1}{\lambda_j}\} + o_d(1). \quad (3.3)$$

Remark 2. We make the following remarks about our result in Theorem 1:

- The spectral decay α determines both the choice of student width r_s and the learning timescale in Theorem 1. Specifically, when $\alpha > 1/2$ (i.e., light-tailed regime), the coefficients $\{\lambda_j\}_{j=1}^r$ are square-summable, making the teacher model effectively finite-dimensional. Hence only finitely many directions need to be learned to achieve small loss, and a timescale of order $\log d$ and finite-width student suffices. In contrast, for the heavy-tailed regime $\alpha < 1/2$, we need to recover linear-in- r directions, which require proportional student width $r_s/r \rightarrow \varphi$ and a longer timescale $r \log d$. This difference will be made explicit in Corollary 1.
- Theorem 1 verifies the additive model hypothesis [MLGT24] for quadratic neural networks in the feature learning regime; specifically, (3.2) identifies sharp transition time in alignment between student weights and the j -th teacher direction, and (3.3) suggests that the cumulative loss can be decomposed into individual emergent risk curves where the timescale is decided by the signal strength λ_j .

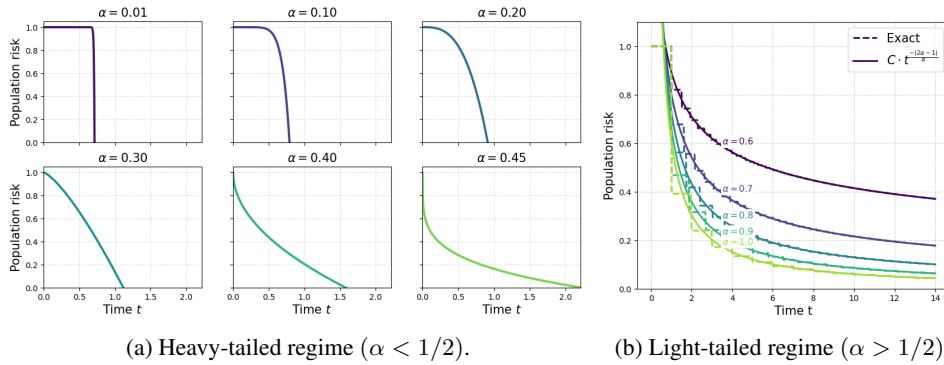


Figure 2: Illustration of the limiting risk trajectories and scaling behavior given in Corollary 1.

Neural scaling laws. As a corollary of Theorem 1, we obtain the following risk characterization.

Corollary 1. By Theorem 1, the asymptotic risk of (GF) is given as follows:

- Heavy-tailed regime ($\alpha \in [0, 0.5)$): Almost surely, for all $t > 0$

$$\mathcal{R}(tr \log d) \xrightarrow{d \rightarrow \infty} (1 - Ct^{\frac{1-2\alpha}{\alpha}})_+ \vee (1 - \varphi^{1-2\alpha})_+.$$

- Light-tailed regime ($\alpha > 0.5$): With probability $1 - \Omega(1/r_s)$, for all $t > 0$, the risk $\mathcal{R}(t \log d)$ converges as $d \rightarrow \infty$ to a deterministic limit satisfying

$$\mathcal{R}(t \log d) \xrightarrow{d \rightarrow \infty} \Theta\left(t^{-\frac{2\alpha-1}{\alpha}} + r_s^{-(2\alpha-1)}\right).$$

Corollary 1 shows that, over appropriate timescales, the cumulative effect of these emergent transitions yields a smoothly decaying risk curve. Intuitively speaking, the power-law exponent arises from the Riemann integral approximation of the infinite sum (3.3) – see Appendix D.5 for details.

The asymptotic risk behavior in Corollary 1 is visualized in Figure 2 (see also Figure 1(b) for empirical simulation). The figure illustrates how the sharp, step-like emergent curve at $\alpha = 0$ (as observed in earlier works on multi-index learning [BAGJ21, AAM23]) gradually transitions into a smooth curve as α increases. Notably, in the light-tailed regime $\alpha > 1/2$, our risk curve resembles the neural scaling laws in [KMH⁺20, HBM⁺22] which takes the form of $\mathcal{R} \sim 1/(\text{Data size})^a + 1/(\text{Model size})^b$, where the data size can be connected to optimization time under the one-pass discretization, which we analyze in the ensuing section.

4 Discrete Dynamics: Online Stochastic Gradient Descent

Now we analyze the finite-sample, discrete-time counterpart of the population dynamics (GF) and establish computational and statistical guarantees. We first discretize the directional component of the dynamics via online SGD with Stiefel constraint (see Proposition 2), and then introduce a fine-tuning step with negligible statistical and computational cost to fit the radial component; this mirrors the layer-wise training paradigm commonly used in theoretical analyses of gradient-based feature learning [AAM22, DLS22, BES⁺22, BEG⁺22]. The procedure is summarized in Algorithm 1.

Algorithm 1 Online Stochastic Gradient Descent (Stiefel)

```

1: for  $t = 1, 2, \dots$  do
2:    $\widetilde{\mathbf{W}}_t = \mathbf{W}_{t-1} - \eta \nabla_{\text{St}} \mathcal{L}(\mathbf{W}_{t-1})$ 
3:    $\mathbf{W}_t = \widetilde{\mathbf{W}}_t \left( \widetilde{\mathbf{W}}_t^\top \widetilde{\mathbf{W}}_t \right)^{-1/2}$  ▷ Feature learning
4: end for
5:  $\mathbf{W}_t^{\text{final}} = \mathbf{W}_t \Omega_*$  where  $\Omega_* = \arg \min_{\Omega \in \mathbb{R}^{r_s \times r_s}} \sum_{j=1}^{N_{\text{Ft}}} \mathcal{L}(\mathbf{W}_t \Omega; (\mathbf{x}_{t+j}, y_{t+j}))$  ▷ Fine-tuning

```

In the feature learning step, we update the first-layer weights $\mathbf{W}_t$ to recover the subspace spanned by the teacher directions. To this end, we use online SGD on Stiefel manifold [AGP24] with polar retraction. The Riemannian gradient on the Stiefel manifold is given by:

$$\nabla_{\text{St}} \mathcal{L}(\mathbf{W}_{t-1}) := \nabla \mathcal{L}(\mathbf{W}_{t-1}) - \frac{1}{2} \mathbf{W}_{t-1} \left(\mathbf{W}_{t-1}^\top \nabla \mathcal{L}(\mathbf{W}_{t-1}) + \nabla \mathcal{L}(\mathbf{W}_{t-1})^\top \mathbf{W}_{t-1} \right),$$

where the instantaneous loss is defined for sample $(\mathbf{x}_t, y_t)$. Since the goal is to ensure subspace alignment (2.5), the overlap of individual student-teacher weights is not relevant during this phase.

After the feature learning phase, we perform a fine-tuning step to rotate $\mathbf{W}_t$ so that each w_j aligns with the corresponding teacher direction θ_j . This is achieved by solving an empirical risk minimization problem over N_{Ft} fresh samples. The optimal fine-tuning matrix Ω_* admits a closed-form solution that is also numerically easy to compute. Importantly, the computational and statistical complexity of this step scales only quadratically with the student width r_s , which is negligible compared to the cost of feature learning. The complexity analysis for this phase is provided in Appendix E.

Remark 3. Recall that the stage-wise training procedure is not required in our continuous analysis in Section 3. This is because we employ a Stiefel gradient similar to [BBPV23, AGP24] – which cannot fit the radial component – to simplify the discretization analysis. We conjecture that standard Euclidean discretization can achieve the same risk scaling; see Figure 1(b) for empirical evidence.

We define the population risk of the output of Algorithm 1, the alignment with a teacher direction θ_j , and the optimal risk achievable by a student neural network with width r_s respectively as

$$\mathcal{R}(t) := R(\mathbf{W}_t^{\text{final}}), \quad \mathcal{A}(t, \theta_j) := \text{Alignment}(\mathbf{W}_t, \theta_j), \quad \mathcal{R}_{\text{opt}} := \frac{1}{\|\mathbf{A}\|_{\text{F}}^2} \sum_{j=(r_s \wedge r)+1}^r \lambda_j^2.$$

Intuitively, $\mathcal{R}_{\text{opt}}$ is the risk achieved by exactly fitting the top $r_s \leq r$ components of the teacher model. Note that the alignment $\mathcal{A}(t, \theta_j)$ depends only on the directional component of $\mathbf{W}_t$; thus, this quantity remains unchanged during fine-tuning. The following theorem characterizes the alignment and risk curve for the discrete-time Algorithm 1.

Theorem 2. Let the parameters $\{\lambda_j\}_{j=1}^r$, r, r_s , r_{eff} and T_{eff} , and the scaling regime (3.1) be as in Theorem 1. Suppose the student weights are initialized uniformly on the Stiefel manifold, and that the step size η and fine-tuning sample size N_{Ft} satisfy

$$\eta \asymp \frac{1}{d} \begin{cases} \frac{1}{r^\alpha \log^{C_\alpha} (1+d/r_s)}, & \alpha \in [0, 0.5) \\ \frac{1}{\log^{C_\alpha} d}, & \alpha > 0.5 \end{cases} \quad \text{and} \quad N_{\text{Ft}} \asymp r_s^2 \log^5 d,$$

for some constant $C_\alpha > 0$ depending only on α . Then with probability $1 - o_d(1/d^2) - \Omega(1/r_s^2)$,

1. **Runtime and sample complexity:** If

$$T \geq \begin{cases} dr^{1+\alpha} \log^{C_\alpha+1} (1+d/r_s), & \alpha \in [0, 0.5) \\ d \log^{C_\alpha+1} d, & \alpha > 0.5. \end{cases} \quad (4.1)$$

we have $\mathcal{R}(T) = \mathcal{R}_{\text{opt}} + o_d(1)$.

2. **Alignment & Risk curve:** For $t > 0$ satisfying $t \asymp r^\alpha/\eta$ for $\alpha \in [0, 0.5)$ and $t \asymp 1/\eta$ for $\alpha > 0.5$,

$$\bullet \mathcal{A}(tT_{\text{eff}}, \theta_j) = \mathbb{1}\{\eta t \geq \frac{1}{\lambda_j}\} + o_d(1). \quad \bullet \mathcal{R}(tT_{\text{eff}}) = 1 - \frac{1}{\|\mathbf{\Lambda}\|_F^2} \sum_{j=1}^{r_{\text{eff}}} \lambda_j^2 \mathbb{1}\{\eta t \geq \frac{1}{\lambda_j}\} + o_d(1).$$

Remark 4. We make the following remarks on the sample complexity.

- The bound in (4.1) implies a complexity of $n \asymp T \simeq dr^{1+\alpha} \text{polylog}(1 + d/r_s)$ in the heavy-tailed case, and $T \simeq d \text{polylog}(d)$ in the light-tailed case. Note that due to the one-pass nature of the algorithm, the runtime and sample complexity are identical (up to the negligible fine-tuning step).
- In the light-tailed regime ($\alpha > 1/2$), the required sample size $n \simeq d \text{polylog}(d)$ is information theoretically optimal up to logarithmic factors. Note that kernel methods and neural networks in the lazy regime [JGH18, COB19] requires $n \gtrsim d^2$ samples to learn a quadratic target function; thus our sample complexity bound illustrates the benefit of feature learning.
- In the heavy-tailed regime ($\alpha < 1/2$), we obtain (nearly) information theoretically optimal sample complexity when $\alpha = 0$. For the intermediate regime $\alpha \in (0, 1/2)$, we conjecture that the optimal sample complexity is $T \simeq dr$, which implies our current bound is suboptimal by r^α .

Isotropic Setting ($\alpha = 0$). In the isotropic case, where the goal is to estimate the r -dimensional subspace spanned by the teacher weights, the above theorem yields a sample and runtime complexity $n \asymp T \asymp dr \text{polylog}(1 + d/r_s)$. This interpolates between the $n \simeq d \text{polylog}(d)$ rate for phase retrieval $r = 1$ [TV23, BAGJ21], and $n \simeq d^2$ as $r \rightarrow d$, which matches the sample complexity in the linear-width regime [MTM⁺24, ETZK25]. Notably, our r -dependence improves upon the recent work of [RL24], which established a sufficient sample size of $n \gtrsim d \text{poly}(r)$ for a similar quadratic setting. We expect our result to be optimal up to polylogarithmic factors due to the intrinsic dr -dimensional nature of the subspace recovery problem.

Scaling laws in discrete time. As indicated by the alignment and risk expressions in Theorem 2, a sufficiently small learning rate η ensures that running online SGD for t steps closely tracks the population gradient flow trajectory (GF) at time ηt , exhibiting the same scaling behavior. The following corollary formalizes the discrete-time counterpart of Corollary 1.

Corollary 2. We consider $\eta t \xrightarrow{d \rightarrow \infty} t_c > 0$. By Theorem 2, we have

- Heavy-tailed case ($\alpha \in [0, 0.5)$): Almost surely,

$$\mathcal{R}(tr \log d) \xrightarrow{d \rightarrow \infty} (1 - Ct_c^{\frac{1-2\alpha}{\alpha}}) \vee (1 - \varphi^{1-2\alpha})_+.$$

- Light-tailed case ($\alpha > 0.5$): With probability $1 - \Omega(1/r_s)$, $\mathcal{R}(t \log d)$ has an asymptotic limit

$$\mathcal{R}(t \log d) \xrightarrow{d \rightarrow \infty} \Theta(t_c^{-\frac{2\alpha-1}{\alpha}} + r_s^{-(2\alpha-1)}).$$

5 Overview of Proof Techniques

To avoid notational confusion between discrete-time and continuous-time dynamics, we adopt the following convention throughout this section. Subscripts (e.g., $\mathbf{W}_t$) denote discrete-time quantities, while parentheses (e.g., $\mathbf{W}(t)$) denote continuous-time trajectories. Specifically, $\{\mathbf{W}_t\}_{t \in \mathbb{N}}$ refers to the iterates of online SGD; $\{\mathbf{W}(t)\}_{t \geq 0}$ denotes the continuous-time gradient flow governed by (GF). Since our proof strategy heavily relies on the matrix (Loewner) order for symmetric matrices, we introduce the following notations. For symmetric matrices $\mathbf{G}_1, \mathbf{G}_2 \in \mathbb{R}^{d \times d}$, we write $\mathbf{G}_1 \prec \mathbf{G}_2$ if $\mathbf{G}_2 - \mathbf{G}_1$ is positive definite. The reverse relations are denoted by $\mathbf{G}_1 \succ \mathbf{G}_2$ and $\mathbf{G}_1 \succeq \mathbf{G}_2$.

5.1 Proof Sketch of Theorem 1

We first observe that both the population risk $R(\mathbf{W}(t))$ and the alignment $\text{Alignment}(\mathbf{W}(t), \theta_j)$ depend on $\mathbf{W}(t)$ through two Gram matrices: the weight Gram matrix $\mathbf{G}_W(t) := \mathbf{W}(t)\mathbf{W}(t)^\top$, and the alignment Gram matrix $\mathbf{G}_U(t) := \mathbf{\Theta}^\top \mathbf{U}(t)\mathbf{U}(t)^\top \mathbf{\Theta}$, where $\{\mathbf{U}(t)\}_{t \geq 0}$ denotes the directional component of $\mathbf{W}(t)$, as defined in (2.4). The proof proceeds by analyzing the evolution of these matrices, each governed by an autonomous ODE; in particular, a matrix Riccati differential equation.

Proposition 1. *The Gram matrices defined above satisfy the following matrix Riccati ODEs:*

- **Weight Gram matrix:** $\partial_t \mathbf{G}_W(t) = \frac{0.5}{\|\mathbf{\Lambda}\|_F \sqrt{r_s}} \left(\mathbf{\Theta} \mathbf{\Lambda} \mathbf{\Theta}^\top \mathbf{G}_W(t) + \mathbf{G}_W(t) \mathbf{\Theta} \mathbf{\Lambda} \mathbf{\Theta}^\top - \frac{2\|\mathbf{\Lambda}\|_F}{\sqrt{r_s}} \mathbf{G}_W^2(t) \right).$
- **Alignment Gram matrix:** $\partial_t \mathbf{G}_U(t) = \frac{0.5}{\|\mathbf{\Lambda}\|_F \sqrt{r_s}} (\mathbf{\Lambda} \mathbf{G}_U(t) + \mathbf{G}_U(t) \mathbf{\Lambda} - 2\mathbf{G}_U(t) \mathbf{\Lambda} \mathbf{G}_U(t)).$

Both equations in Proposition 1 take the form of matrix Riccati ODEs [BLW91], whose structural properties play a central role in the proof. To illustrate the core idea, we focus on the alignment dynamics. For simplicity, we write $\mathbf{G}(t) := \mathbf{G}_U(t)$ and consider

$$\partial_t \mathbf{G}(t) = \frac{0.5}{\|\mathbf{\Lambda}\|_F \sqrt{r_s}} (\mathbf{\Lambda} \mathbf{G}(t) + \mathbf{G}(t) \mathbf{\Lambda} - 2\mathbf{G}(t) \mathbf{\Lambda} \mathbf{G}(t)). \quad (5.1)$$

Note that $\text{Alignment}(\mathbf{W}(t), \theta_j)$ corresponds to the j^{th} diagonal entry of $\mathbf{G}(t)$. To characterize its trajectory, we leverage the monotonicity of the matrix Riccati flow with respect to its initialization, i.e., if $\mathbf{G}_0^+ \succeq \mathbf{G}_0^-$, the corresponding solutions satisfy $\mathbf{G}(t, \mathbf{G}_0^+) \succeq \mathbf{G}(t, \mathbf{G}_0^-)$ for all $t \geq 0$, where $\mathbf{G}(t, \mathbf{G}_0)$ denotes the solution to (5.1) with initial condition $\mathbf{G}_0$. Our proof strategy builds on this monotonicity and proceeds as follows:

1. **Diagonalization & decoupling.** If $\mathbf{G}_0$ is diagonal, the solution $\{\mathbf{G}(t)\}_{t \geq 0}$ remains diagonal under (5.1), reducing the dynamics to independent scalar ODEs that govern each diagonal entry. Moreover, each scalar ODE admits a closed-form solution.
2. **Asymptotic characterization.** For general $\mathbf{G}_0$, we construct diagonal matrices $\mathbf{G}_0^+ \succeq \mathbf{G}_0 \succeq \mathbf{G}_0^-$. By monotonicity, the corresponding trajectories upper and lower bound $\{\mathbf{G}(t)\}_{t \geq 0}$. These bounding systems are diagonal and decoupled, and as $d \rightarrow \infty$, they converge to the same limit.

We apply this strategy in Appendix D.3 to derive the exact asymptotics stated in Theorem 1.

Remark 5. *We note that while the Riccati flow is monotone with respect to its initialization, this does not imply that its solution is monotone in time. That is, the trajectory $\mathbf{G}(t)$ may not evolve monotonically in matrix order, even though a larger initialization yields a trajectory that remains above that of a smaller one for all $t \geq 0$. This distinction is illustrated in Figure 4.*

5.2 Proof Sketch of Theorem 2

Extending Monotonicity Arguments to Discrete Dynamics

We begin by observing that online SGD on the Stiefel manifold approximates the directional component of the continuous-time gradient flow, with stochastic gradients arising from online sampling. The proposition below formalizes the idea that online SGD approximates the directional dynamics of the continuous gradient flow at the population level. For the statement, recall that $\mathbf{U}(t)$ denotes the directional component of the gradient flow solution $\mathbf{W}(t)$ from (GF), as defined in (2.4).

Proposition 2. *Let $\widehat{\mathbf{W}}(t)$ be the solution to the continuous-time gradient flow on the Stiefel manifold, initialized with $\mathbf{U}(0)$. Then for all $t \geq 0$, the column spaces of $\widehat{\mathbf{W}}(t)$ and $\mathbf{U}(t)$ coincide.*

This result justifies studying the online SGD on Stiefel manifold via the directional dynamics of (GF). To this end, we introduce the discrete analog of $\mathbf{G}(t)$ above as $\mathbf{G}_t = \mathbf{\Theta}^\top \mathbf{W}_t \mathbf{W}_t^\top \mathbf{\Theta}$. Extending the analysis to discrete time is non-trivial due to the *loss of monotonicity* in the Euler discretization of the Riccati dynamics (5.1). In particular, the update

$$\mathbf{G}_t = \mathbf{G}_{t-1} + \underbrace{\frac{0.5\eta}{\|\mathbf{\Lambda}\|_F \sqrt{r_s}} (\mathbf{\Lambda} \mathbf{G}_{t-1} + \mathbf{G}_{t-1} \mathbf{\Lambda} - 2\mathbf{G}_{t-1} \mathbf{\Lambda} \mathbf{G}_{t-1})}_{\text{non-monotone dynamics}} + (2\text{nd-order terms and noise}) \quad (5.2)$$

no longer preserves the matrix order structure crucial to the continuous-time argument.

To overcome this, we construct an auxiliary discrete system that approximates (5.2) up to second-order terms while preserving monotonicity. Specifically, we define the map

$$\mathbf{G}(\mathbf{G}_t, \eta) := \mathbf{G}_t - \frac{\eta}{2} (2\mathbf{G}_t - \mathbf{I}_r) \mathbf{\Lambda} (2\mathbf{G}_t - \mathbf{I}_r) (\mathbf{I}_r + \eta \mathbf{\Lambda} (2\mathbf{G}_t - \mathbf{I}_r))^{-1} + \eta \mathbf{\Lambda} \quad (5.3)$$

which matches (5.2) up to second-order terms. Indeed, expanding the inverse term gives

$$\begin{aligned} \mathbf{G}(\mathbf{G}_t, \eta) &= \mathbf{G}_t - \frac{\eta}{2} (2\mathbf{G}_t - \mathbf{I}_r) \mathbf{\Lambda} (2\mathbf{G}_t - \mathbf{I}_r) + \eta \mathbf{\Lambda} + 2\text{nd-order terms} \\ &= \underbrace{\mathbf{G}_t + \eta (\mathbf{\Lambda} \mathbf{G}_t + \mathbf{G}_t \mathbf{\Lambda} - 2\mathbf{G}_t \mathbf{\Lambda} \mathbf{G}_t)}_{\text{Riccati dynamics}} \end{aligned}$$

The key advantage of the iteration (5.3) is that it preserves matrix order:

Proposition 3. For $\eta > 0$, if $\mathbf{G}_t^+ \succeq \mathbf{G}_t^- \succeq 0$, we have $\mathbf{G}(\mathbf{G}_t^+, \eta) \succeq \mathbf{G}(\mathbf{G}_t^-, \eta)$.

We use this to bound the non-monotone dynamics (5.2) via monotone iterates. Roughly, we show that for small enough step size η , the following holds:

$$\mathbf{G}(\mathbf{G}_{t-1}, (1 + \varepsilon)\eta) + \text{Noise} \succeq \mathbf{G}_t \succeq \mathbf{G}(\mathbf{G}_{t-1}, (1 - \varepsilon)\eta) + \text{Noise}$$

for some $\varepsilon = o_d(1)$, where we denote the effective learning rate in (5.2) with $\eta = \frac{\eta}{\sqrt{r_s} \|\mathbf{A}\|_F}$. We then follow the same bounding argument used in the continuous case by defining

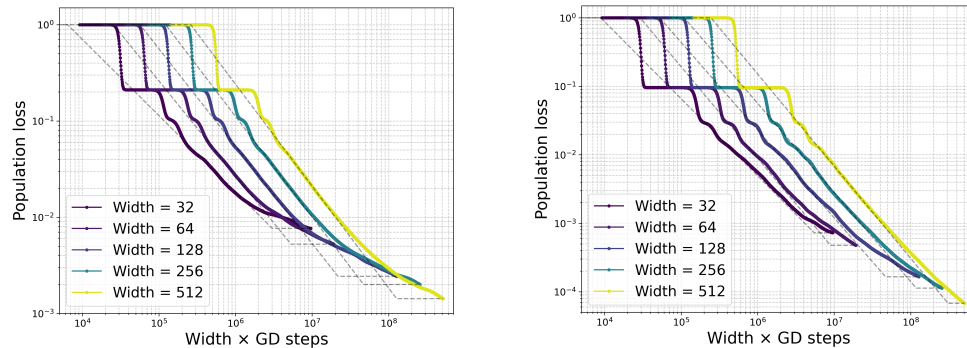
$$\mathbf{G}_t^\pm = \mathbf{G}(\mathbf{G}_{t-1}^\pm, (1 \pm \varepsilon)\eta) + \text{Noise}, \text{ where } \mathbf{G}_0^+ \succeq \mathbf{G}_0 \succeq \mathbf{G}_0^- \succeq 0,$$

and show that $\mathbf{G}_t^+ \succeq \mathbf{G}_t \succeq \mathbf{G}_t^-$ for all $t \in \mathbb{N}$. Finally, by choosing $\mathbf{G}_0^\pm$ to be diagonal, the bounding dynamics reduce to decoupled scalar recursions, which can be analyzed explicitly. This allows us to establish concentration of the original iterates $\{\mathbf{G}_t\}_{t \in \mathbb{N}}$ around the bounding sequences, leading to operator-norm convergence of the discrete-time dynamics to their continuous-time counterparts.

6 Conclusion

In this work, we presented a comprehensive theoretical analysis of gradient-based learning in high-dimensional, extensive-width two-layer neural networks with quadratic activation. We established precise scaling laws that characterize both the population gradient flow and its empirical, discrete-time approximation. These results demonstrate how anisotropic signal strengths in the target function fundamentally shapes the convergence behavior and sample efficiency of gradient-based learning.

Beyond quadratic activations. An immediate direction for future research is to extend our analysis to more general activation functions. Link functions with higher information exponent is studied in a companion work [RNWL25], where the precise risk scaling is established by exploiting a decoupling structure that is unique to the information exponent $k > 2$ setting. Importantly, many commonly-used activation functions (ReLU, GeLU, etc.) have information exponent $k = 1$ and also contain a nonzero He_2 component. For such nonlinearities, we conjecture that SGD dynamics exhibits a multi-phase risk curve (analogous to the incremental learning phenomenon in [AAM23, BBPV23]), where the higher Hermite modes affects the learning dynamics after the low-order terms are learned. In Figure 3 we report the SGD risk curves for ReLU networks, in which we observe (i) an initial loss drop driven by the He_1 component (which finds a degenerate rank-1 subspace), followed by (ii) a power-law decay phase driven by the quadratic He_2 component where the empirical scaling exponent align closely with our theoretical predictions, and finally (iii) a slope change late in training likely due to higher Hermite terms (in Figure 5 we confirm that this “late” phase is absent if we remove these higher-order components). Understanding this complex multi-phase learning dynamics remains an interesting challenge for future work.



(a) $\alpha = 1$. Ideal exponent: -1 ; empirical: -1.01 . (b) $\alpha = \frac{3}{2}$. Ideal exponent: $-\frac{4}{3}$, empirical: -1.26 .

Figure 3: Population loss vs. compute for two-layer ReLU network (power-law second-layer with exponent α) trained with population gradient descent. The student network adopts the 2-homogeneous parameterization as in (2.2). Observe that after the initial loss drop due to the He_1 component, the risk curves follow a power-law scaling where the exponent (dashed) nearly matches our theoretical prediction for the quadratic setting $\frac{1-2\alpha}{\alpha}$.

Acknowledgment

The authors thank Florent Krzakala, Jason D. Lee, and Lenka Zdeborová for discussion and feedback. The research of GBA was supported in part by NSF grant 2134216. MAE was partially supported by the NSERC Grant [2019-06167], the CIFAR AI Chairs program, and the CIFAR Catalyst grant. Part of this work was completed when NMV interned at the Flatiron Institute.

References

- [AAM22] Emmanuel Abbe, Enric Boix Adsera, and Theodor Misiakiewicz. The merged-staircase property: a necessary and nearly sufficient condition for sgd learning of sparse functions on two-layer neural networks. In *Conference on Learning Theory*, pages 4782–4887. PMLR, 2022.
- [AAM23] Emmanuel Abbe, Enric Boix Adsera, and Theodor Misiakiewicz. SGD learning on neural networks: leap complexity and saddle-to-saddle dynamics. In *The Thirty Sixth Annual Conference on Learning Theory*, pages 2552–2623. PMLR, 2023.
- [ADK⁺24] Luca Arnaboldi, Yatin Dandi, Florent Krzakala, Luca Pesce, and Ludovic Stephan. Repetita iuvant: Data repetition allows sgd to learn high-dimensional multi-index functions. *arXiv preprint arXiv:2405.15459*, 2024.
- [AGP24] Gérard Ben Arous, Cédric Gerbelot, and Vanessa Piccolo. High-dimensional optimization for multi-spiked tensor pca. *arXiv preprint arXiv:2408.06401*, 2024.
- [AKLS23] Luca Arnaboldi, Florent Krzakala, Bruno Loureiro, and Ludovic Stephan. Escaping mediocrity: how two-layer networks learn hard generalized linear models with sgd. *arXiv preprint arXiv:2305.18502*, 2023.
- [AZL17] Zeyuan Allen-Zhu and Yuanzhi Li. First efficient convergence for streaming k-pca: a global, gap-free, and near-optimal rate. In *2017 IEEE 58th Annual Symposium on Foundations of Computer Science (FOCS)*, pages 487–492. IEEE, 2017.
- [Bac17] Francis Bach. Breaking the curse of dimensionality with convex neural networks. *The Journal of Machine Learning Research*, 18(1):629–681, 2017.
- [BAGJ21] Gerard Ben Arous, Reza Gheissari, and Aukosh Jagannath. Online stochastic gradient descent on non-convex losses from high-dimensional inference. *The Journal of Machine Learning Research*, 22(1):4788–4838, 2021.
- [BAGJ22] Gerard Ben Arous, Reza Gheissari, and Aukosh Jagannath. High-dimensional limit theorems for sgd: Effective dynamics and critical scaling. *Advances in Neural Information Processing Systems*, 35:25349–25362, 2022.
- [BAP24] Blake Bordelon, Alexander Atanasov, and Cengiz Pehlevan. A dynamical model of neural scaling laws. *arXiv preprint arXiv:2402.01092*, 2024.
- [BBPV23] Alberto Bietti, Joan Bruna, and Loucas Pillaud-Vivien. On learning Gaussian multi-index models with gradient flow. *arXiv preprint arXiv:2310.19793*, 2023.
- [BBSS22] Alberto Bietti, Joan Bruna, Clayton Sanford, and Min Jae Song. Learning single-index models with shallow neural networks. *Advances in Neural Information Processing Systems*, 35:9768–9783, 2022.
- [BCN⁺25] Jean Barbier, Francesco Camilli, Minh-Toan Nguyen, Mauro Pastore, and Rudy Skerk. Statistical mechanics of extensive-width bayesian neural networks near interpolation. *arXiv preprint arXiv:2505.24849*, 2025.
- [BEG⁺22] Boaz Barak, Benjamin Edelman, Surbhi Goel, Sham Kakade, Eran Malach, and Cyril Zhang. Hidden progress in deep learning: Sgd learns parities near the computational limit. *Advances in Neural Information Processing Systems*, 35:21750–21764, 2022.

- [BES⁺22] Jimmy Ba, Murat A Erdogdu, Taiji Suzuki, Zhichao Wang, Denny Wu, and Greg Yang. High-dimensional asymptotics of feature learning: How one gradient step improves the representation. *Advances in Neural Information Processing Systems*, 35:37932–37946, 2022.
- [BES⁺23] Jimmy Ba, Murat A Erdogdu, Taiji Suzuki, Zhichao Wang, and Denny Wu. Learning in the presence of low-dimensional structure: A spiked random matrix perspective. In *Thirty-seventh Conference on Neural Information Processing Systems*, 2023.
- [BH25] Joan Bruna and Daniel Hsu. Survey on algorithms for multi-index models. *arXiv preprint arXiv:2504.05426*, 2025.
- [BLW91] Sergio Bittanti, Alan J. Laub, and Jan C. Willems. The riccati equation. 1991.
- [BMZ24] Raphaël Berthier, Andrea Montanari, and Kangjie Zhou. Learning time-scales in two-layers neural networks. *Foundations of Computational Mathematics*, pages 1–84, 2024.
- [BR14] Davide Barilari and Luca Rizzi. Comparison theorems for conjugate points in sub-riemannian geometry. *ESAIM: Control, Optimisation and Calculus of Variations*, 22:439–472, 2014.
- [Bub14] Sébastien Bubeck. Convex optimization: Algorithms and complexity. *Found. Trends Mach. Learn.*, 8:231–357, 2014.
- [CB20] Lenaïc Chizat and Francis Bach. Implicit bias of gradient descent for wide two-layer neural networks trained with the logistic loss. In *Conference on learning theory*, pages 1305–1338. PMLR, 2020.
- [CC15] Yuxin Chen and Emmanuel Candes. Solving random quadratic systems of equations is nearly as easy as solving linear systems. *Advances in Neural Information Processing Systems*, 28, 2015.
- [CM20] Sitan Chen and Raghu Meka. Learning polynomials in few relevant dimensions. In *Conference on Learning Theory*, pages 1161–1227. PMLR, 2020.
- [COB19] Lenaïc Chizat, Edouard Oyallon, and Francis Bach. On lazy training in differentiable programming. *Advances in Neural Information Processing Systems*, 32, 2019.
- [CWPPS23] Elizabeth Collins-Woodfin, Courtney Paquette, Elliot Paquette, and Inbar Seroussi. Hitting the high-dimensional notes: An ode for sgd learning dynamics on glms and multi-index models. *arXiv preprint arXiv:2308.08977*, 2023.
- [DCLT18] Jacob Devlin, Ming-Wei Chang, Kenton Lee, and Kristina Toutanova. Bert: Pre-training of deep bidirectional transformers for language understanding. *arXiv preprint arXiv:1810.04805*, 2018.
- [DDH⁺21] Damai Dai, Li Dong, Yaru Hao, Zhifang Sui, Baobao Chang, and Furu Wei. Knowledge neurons in pretrained transformers. *arXiv preprint arXiv:2104.08696*, 2021.
- [DH18] Rishabh Dudeja and Daniel Hsu. Learning single-index models in gaussian space. In *Conference On Learning Theory*, pages 1887–1930. PMLR, 2018.
- [DKL⁺23a] Yatin Dandi, Florent Krzakala, Bruno Loureiro, Luca Pesce, and Ludovic Stephan. How two-layer neural networks learn, one (giant) step at a time. *arXiv preprint arXiv:2305.18270*, 2023.
- [DKL⁺23b] Yatin Dandi, Florent Krzakala, Bruno Loureiro, Luca Pesce, and Ludovic Stephan. Learning two-layer neural networks, one (giant) step at a time. *arXiv preprint arXiv:2305.18270*, 2023.
- [DL18] Simon Du and Jason Lee. On the power of over-parametrization in neural networks with quadratic activation. In *International conference on machine learning*, pages 1329–1338. PMLR, 2018.

- [DLM24] Leonardo Defilippis, Bruno Loureiro, and Theodor Misiakiewicz. Dimension-free deterministic equivalents for random feature regression. *arXiv preprint arXiv:2405.15699*, 2024.
- [DLS22] Alexandru Damian, Jason Lee, and Mahdi Soltanolkotabi. Neural networks can learn representations with gradient descent. In *Conference on Learning Theory*, pages 5413–5452. PMLR, 2022.
- [DNGL24] Alex Damian, Eshaan Nichani, Rong Ge, and Jason D Lee. Smoothing the landscape boosts the signal for sgd: Optimal sample complexity for learning single index models. *Advances in Neural Information Processing Systems*, 36, 2024.
- [DPVLB24] Alex Damian, Loucas Pillaud-Vivien, Jason D Lee, and Joan Bruna. The computational complexity of learning gaussian single-index models. *arXiv preprint arXiv:2403.05529*, 2024.
- [DTA⁺24] Yatin Dandi, Emanuele Troiani, Luca Arnaboldi, Luca Pesce, Lenka Zdeborová, and Florent Krzakala. The benefits of reusing batches for gradient descent in two-layer networks: Breaking the curse of information and leap exponents. *arXiv preprint arXiv:2402.03220*, 2024.
- [Dur93] Richard Durrett. *Probability: Theory and examples*. 1993.
- [EHO⁺22] Nelson Elhage, Tristan Hume, Catherine Olsson, Nicholas Schiefer, Tom Henighan, Shauna Kravec, Zac Hatfield-Dodds, Robert Lasenby, Dawn Drain, Carol Chen, et al. Toy models of superposition. *arXiv preprint arXiv:2209.10652*, 2022.
- [ETZK25] Vittorio Erba, Emanuele Troiani, Lenka Zdeborová, and Florent Krzakala. The nuclear route: Sharp asymptotics of erm in overparameterized quadratic networks. *arXiv preprint arXiv:2505.17958*, 2025.
- [Fie82] James R Fienup. Phase retrieval algorithms: a comparison. *Applied optics*, 21(15):2758–2769, 1982.
- [GDDM14] Ross Girshick, Jeff Donahue, Trevor Darrell, and Jitendra Malik. Rich feature hierarchies for accurate object detection and semantic segmentation. In *Proceedings of the IEEE conference on computer vision and pattern recognition*, pages 580–587, 2014.
- [GKZ19] David Gamarnik, Eren C Kızıldağ, and Ilias Zadik. Stationary points of shallow neural networks with quadratic activation function. *arXiv preprint arXiv:1912.01599*, 2019.
- [GMMM19] Behrooz Ghorbani, Song Mei, Theodor Misiakiewicz, and Andrea Montanari. Limitations of lazy training of two-layers neural network. *Advances in Neural Information Processing Systems*, 32, 2019.
- [GRWZ21] 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.
- [GWB25] Margalit Glasgow, Denny Wu, and Joan Bruna. Propagation of chaos in one-hidden-layer neural networks beyond logarithmic time. *arXiv preprint arXiv:2504.13110*, 2025.
- [HBM⁺22] Jordan Hoffmann, Sebastian Borgeaud, Arthur Mensch, Elena Buchatskaya, Trevor Cai, Eliza Rutherford, Diego de Las Casas, Lisa Anne Hendricks, Johannes Welbl, Aidan Clark, et al. Training compute-optimal large language models. *arXiv preprint arXiv:2203.15556*, 2022.
- [HNA⁺17] Joel Hestness, Sharan Narang, Newsha Ardalani, Gregory Diamos, Heewoo Jun, Hassan Kianinejad, Md Patwary, Mostofa Ali, Yang Yang, and Yanqi Zhou. Deep learning scaling is predictable, empirically. *arXiv preprint arXiv:1712.00409*, 2017.
- [HT87] Trevor Hastie and Robert Tibshirani. Generalized additive models: some applications. *Journal of the American Statistical Association*, 82(398):371–386, 1987.

- [JGH18] Arthur Jacot, Franck Gabriel, and Clément Hongler. Neural tangent kernel: Convergence and generalization in neural networks. In *Advances in neural information processing systems*, pages 8571–8580, 2018.
- [JJK⁺16] Prateek Jain, Chi Jin, Sham M Kakade, Praneeth Netrapalli, and Aaron Sidford. Streaming pca: Matching matrix bernstein and near-optimal finite sample guarantees for oja’s algorithm. In *Conference on learning theory*, pages 1147–1164. PMLR, 2016.
- [KMH⁺20] Jared Kaplan, Sam McCandlish, Tom Henighan, Tom B Brown, Benjamin Chess, Rewon Child, Scott Gray, Alec Radford, Jeffrey Wu, and Dario Amodei. Scaling laws for neural language models. *arXiv preprint arXiv:2001.08361*, 2020.
- [LM00] B. Laurent and P. Massart. Adaptive estimation of a quadratic functional by model selection. *The Annals of Statistics*, 28(5):1302–1338, 2000.
- [LMZ20] Yuanzhi Li, Tengyu Ma, and Hongyang R Zhang. Learning over-parametrized two-layer neural networks beyond ntk. In *Conference on learning theory*, pages 2613–2682. PMLR, 2020.
- [LOSW24] Jason D Lee, Kazusato Oko, Taiji Suzuki, and Denny Wu. Neural network learns low-dimensional polynomials with sgd near the information-theoretic limit. *arXiv preprint arXiv:2406.01581*, 2024.
- [LWK⁺24] Licong Lin, Jingfeng Wu, Sham M Kakade, Peter L Bartlett, and Jason D Lee. Scaling laws in linear regression: Compute, parameters, and data. *arXiv preprint arXiv:2406.08466*, 2024.
- [MBB23] Simon Martin, Francis Bach, and Giulio Biroli. On the impact of overparameterization on the training of a shallow neural network in high dimensions. *arXiv preprint arXiv:2311.03794*, 2023.
- [MHPG⁺22] Alireza Mousavi-Hosseini, Sejun Park, Manuela Girotti, Ioannis Mitliagkas, and Murat A Erdogdu. Neural networks efficiently learn low-dimensional representations with sgd. In *The Eleventh International Conference on Learning Representations*, 2022.
- [MHWE24] Alireza Mousavi-Hosseini, Denny Wu, and Murat A Erdogdu. Learning multi-index models with neural networks via mean-field langevin dynamics. *arXiv preprint arXiv:2408.07254*, 2024.
- [MHWSE23] Alireza Mousavi-Hosseini, Denny Wu, Taiji Suzuki, and Murat A. Erdogdu. Gradient-based feature learning under structured data. In *Thirty-seventh Conference on Neural Information Processing Systems (NeurIPS 2023)*, 2023.
- [MLGT24] Eric Michaud, Ziming Liu, Uzey Girit, and Max Tegmark. The quantization model of neural scaling. *Advances in Neural Information Processing Systems*, 36, 2024.
- [MLHD23] Behrad Moniri, Donghwan Lee, Hamed Hassani, and Edgar Dobriban. A theory of non-linear feature learning with one gradient step in two-layer neural networks. *arXiv preprint arXiv:2310.07891*, 2023.
- [MRS22] Alexander Maloney, Daniel A Roberts, and James Sully. A solvable model of neural scaling laws. *arXiv preprint arXiv:2210.16859*, 2022.
- [MTM⁺24] Antoine Maillard, Emanuele Troiani, Simon Martin, Florent Krzakala, and Lenka Zdeborová. Bayes-optimal learning of an extensive-width neural network from quadratically many samples. *arXiv preprint arXiv:2408.03733*, 2024.
- [MVEZ20] Sarao Stefano Mannelli, Eric Vanden-Eijnden, and Lenka Zdeborová. Optimization and generalization of shallow neural networks with quadratic activation functions. *Advances in Neural Information Processing Systems*, 33:13445–13455, 2020.

- [MZD⁺23] Arvind Mahankali, Haochen Zhang, Kefan Dong, Margalit Glasgow, and Tengyu Ma. Beyond ntk with vanilla gradient descent: A mean-field analysis of neural networks with polynomial width, samples, and time. *Advances in Neural Information Processing Systems*, 36, 2023.
- [NFL24] Yoonsoo Nam, Nayara Fonseca, Seok Hyeong Lee, and Ard Louis. An exactly solvable model for emergence and scaling laws. *arXiv preprint arXiv:2404.17563*, 2024.
- [OK85] Erkki Oja and Juha Karhunen. On stochastic approximation of the eigenvectors and eigenvalues of the expectation of a random matrix. *Journal of mathematical analysis and applications*, 106(1):69–84, 1985.
- [OSSW24] Kazusato Oko, Yujin Song, Taiji Suzuki, and Denny Wu. Learning sum of diverse features: computational hardness and efficient gradient-based training for ridge combinations. In *Conference on Learning Theory*. PMLR, 2024.
- [PPXP24] Elliot Paquette, Courtney Paquette, Lechao Xiao, and Jeffrey Pennington. 4+ 3 phases of compute-optimal neural scaling laws. *arXiv preprint arXiv:2405.15074*, 2024.
- [PSZA23] Abhishek Panigrahi, Nikunj Saunshi, Haoyu Zhao, and Sanjeev Arora. Task-specific skill localization in fine-tuned language models. *arXiv preprint arXiv:2302.06600*, 2023.
- [RL24] Yunwei Ren and Jason D Lee. Learning orthogonal multi-index models: A fine-grained information exponent analysis. *arXiv preprint arXiv:2410.09678*, 2024.
- [RNWL25] Yunwei Ren, Eshaan Nichani, Denny Wu, and Jason D Lee. Emergence and scaling laws in sgd learning of shallow neural networks. *arXiv preprint arXiv:2504.19983*, 2025.
- [SBH24] Berfin Simsek, Amire Bendjeddou, and Daniel Hsu. Learning gaussian multi-index models with gradient flow: Time complexity and directional convergence. *arXiv preprint arXiv:2411.08798*, 2024.
- [SJL18] Mahdi Soltanolkotabi, Adel Javanmard, and Jason D Lee. Theoretical insights into the optimization landscape of over-parameterized shallow neural networks. *IEEE Transactions on Information Theory*, 65(2):742–769, 2018.
- [Sto85] Charles J Stone. Additive regression and other nonparametric models. *The annals of Statistics*, 13(2):689–705, 1985.
- [Tro10] Joel A. Tropp. User-friendly tail bounds for sums of random matrices. *Foundations of Computational Mathematics*, 12:389–434, 2010.
- [TV23] Yan Shuo Tan and Roman Vershynin. Online stochastic gradient descent with arbitrary initialization solves non-smooth, non-convex phase retrieval. *Journal of Machine Learning Research*, 24(58):1–47, 2023.
- [VBB19] Luca Venturi, Afonso S Bandeira, and Joan Bruna. Spurious valleys in one-hidden-layer neural network optimization landscapes. *Journal of Machine Learning Research*, 20(133):1–34, 2019.
- [VE24] Nuri Mert Vural and Murat A Erdogdu. Pruning is optimal for learning sparse features in high-dimensions. *arXiv preprint arXiv:2406.08658*, 2024.
- [Ver10] Roman Vershynin. Introduction to the non-asymptotic analysis of random matrices. *arXiv preprint arXiv:1011.3027*, 2010.
- [WWZ⁺22] Xiaozhi Wang, Kaiyue Wen, Zhengyan Zhang, Lei Hou, Zhiyuan Liu, and Juanzi Li. Finding skill neurons in pre-trained transformer-based language models. *arXiv preprint arXiv:2211.07349*, 2022.

Contents

1	Introduction	1
1.1	Our Contributions	3
1.2	Additional Related Works	4
2	Background and Problem Setting	4
2.1	Student-teacher Setting	4
2.2	Training Objective	5
3	Continuous Dynamics: Population Gradient Flow	5
4	Discrete Dynamics: Online Stochastic Gradient Descent	7
5	Overview of Proof Techniques	8
5.1	Proof Sketch of Theorem 1	8
5.2	Proof Sketch of Theorem 2	9
6	Conclusion	10
A	Additional Figures	18
B	Preliminaries for Proofs	19
C	Background: Matrix Riccati Dynamical Systems	21
C.1	Continuous-time Matrix Riccati ODE	21
C.2	Discrete-time Matrix Riccati Difference Equations	21
D	Proofs for Main Results	24
D.1	Proof of Propositions 2 and 3	24
D.2	Decomposition of the population risk	24
D.3	Proof of Theorem 1	25
D.3.1	High-dimensional limit for the alignment	25
D.3.2	High-dimensional limit for the risk curve	26
D.4	Proof of Theorem 2	30
D.5	Proof of Corollary 1 and Corollary 2	31
E	Details of the Fine-tuning Step	33
E.1	Characterizing the Minimum	33
E.2	Computing the Minimum	34
E.2.1	Proof of Proposition 11	34
F	Deferred Proofs for Online SGD	35
F.1	Preliminaries	35
F.2	Including second-order terms and monotone bounds	37
F.2.1	Heavy tailed case - $\alpha \in [0, 0.5)$	37
F.2.2	Light tailed case - $\alpha > 0.5$	38

F.3	Definitions and bounding systems	40
F.3.1	Proof of Proposition 14	42
F.4	Analysis of the bounding systems	46
F.4.1	Lower bounding system	46
F.4.2	Upper bounding system	51
F.5	Bounds for the second-order terms	55
F.6	Noise characterization	57
F.7	Stability near minima	67
G	Auxiliary Statements	70
G.1	Matrix bounds	70
G.1.1	Additional bounds for continuous-time analysis	71
G.1.2	Additional bounds for discrete-time analysis	73
G.2	Some moment bounds and concentration inequalities	76
G.3	Miscellaneous	78

A Additional Figures

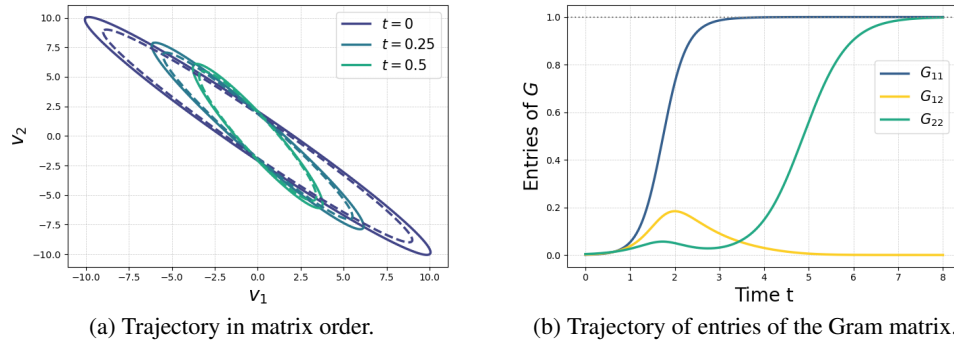


Figure 4: Solutions of the matrix Riccati ODE in (5.1) with $\lambda_1 = 2$, $\lambda_2 = 1$, $r_s = 2$. (a) To visualize the dynamics under matrix order, we plot the level sets of $G(t)$ at times $t \in \{0, 0.25, 0.5\}$ for two initializations: $G(0)$ (solid) and a scaled version $1.25 G(0)$ (dashed). The dashed ellipses remain enclosed within the solid ones at all times, illustrating monotonicity of the Riccati flow with respect to initialization. However, note that $G(t)$ is not monotone in Loewner order over time, as seen from the lack of nesting among the solid ellipses. (b) Entry-wise evolution of $G(t)$ under a random initialization with $d = 1024$. The diagonal entry $G_{22}(t)$ exhibits non-monotonic behavior, illustrating that the solution trajectory $G(t)$ need not be monotone in time.

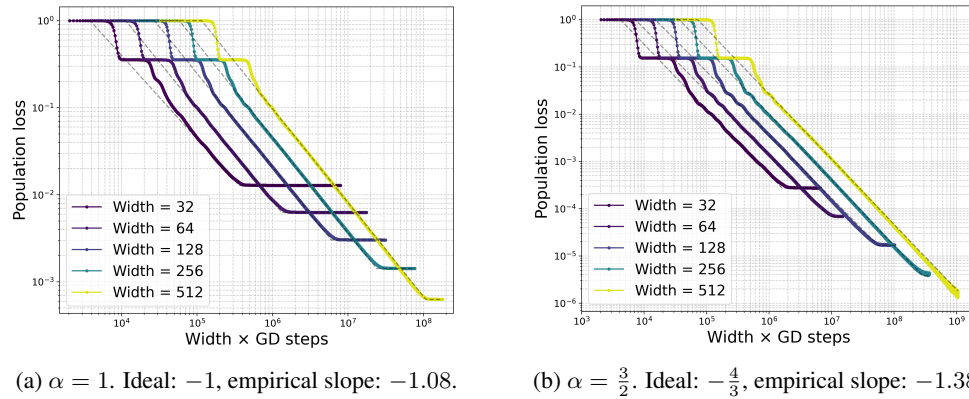


Figure 5: Population loss vs. compute for two-layer neural network with activation function $\sigma \propto \text{He}_1 + \text{He}_2$, trained with population gradient descent. The student network adopts the 2-homogeneous parameterization as in (2.2). Observe that after the initial loss drop due to the He_1 component, the risk curves exhibit a power-law scaling where the exponent (dashed lines) nearly matches our theoretical prediction for the quadratic setting $\frac{1-2\alpha}{\alpha}$; and unlike the ReLU setting (Figure 3), the loss immediately plateaus after the power-law phase.

Experiment Setting. In Figures 3 and 5, we plot the mean squared error loss for gradient descent with a constant step size on the population loss, using activations $\sigma = \text{ReLU}$ and $\sigma = \text{He}_1 + \text{He}_2$. The teacher model has orthogonal first-layer neurons and power-law decay in the second-layer coefficients with $\alpha \in \{1.0, 1.5\}$. Both teacher and student networks use the same activation function, which we normalize to have zero-mean and unit L^2 norm. The student network uses the 2-homogeneous parameterization:

$$\hat{y}(\mathbf{W}) = \frac{1}{\sqrt{r_s}} \sum_{i=1}^{r_s} \|\mathbf{w}_i\|_2^2 \cdot \sigma(\langle \mathbf{w}_i, \mathbf{x} \rangle) \quad \text{where } \sigma \in \left\{ \frac{\text{ReLU}-1/\sqrt{2\pi}}{0.5}, \frac{\text{He}_1+\text{He}_2}{3} \right\}.$$

We set dimension $d = 5000$, number of teacher neurons $r = 2400$, student widths $r_s \in \{32, 64, 128, 256, 512\}$, and learning rate $\eta = 0.5/\sqrt{r}$. To estimate the scaling exponents, we first identify the range of compute exhibiting a linear trend by visual inspection, and then fit the exponent via least squares. The dashed lines in the plots correspond to these fitted lines, and the reported empirical exponents represent the median values across different student widths.

B Preliminaries for Proofs

Proof organization. Section B introduces the notations and definitions used throughout the paper. In Section C, we provide a brief review of matrix Riccati ODEs and difference equations, along with the necessary supporting statements. The main results are proved in Section D. In Section E we discuss the fine-tuning phase for the discretized algorithm. Additional proofs related to online SGD and auxiliary lemmas are deferred to Sections F and G, respectively.

Notation and Definitions. We use $[n] := \{1, 2, \dots, n\}$ to denote the first n natural numbers. The Euclidean inner product and norm are denoted by $\langle \cdot, \cdot \rangle$ and $\|\cdot\|_2$, respectively. For matrices, $\|\cdot\|_2$ and $\|\cdot\|_F$ denote the operator norm and Frobenius norm. The positive part is denoted by $(x)_+ := \max\{x, 0\}$. We write $f_d = o_d(1)$ if $f_d \rightarrow 0$ as $d \rightarrow \infty$, and $f_d \ll g_d$ if $f_d/g_d \rightarrow 0$. We use $O(\cdot)$ or $\Omega(\cdot)$ to suppress constants in upper and lower bounds respectively, and we use subscript to indicate parameter dependence, e.g., $O_\alpha(\cdot)$.

The symmetric part of a square matrix $M \in \mathbb{R}^{d \times d}$ is given by $\text{Sym}(M) := \frac{1}{2}(M + M^\top)$. For symmetric matrices $A, B \in \mathbb{R}^{d \times d}$, we write $A \prec B$ (or $A \preceq B$) if $B - A$ is positive definite (or positive semidefinite). Moreover, if A and B are mutually diagonalizable, we write $AB^{-1} = \frac{A}{B}$.

We follow the convention that subscripts (e.g., W_t) refer to discrete-time quantities, and parentheses (e.g., $W(t)$) refer to continuous-time quantities. The overlap matrices of interest are defined as

$$\underbrace{G_W(t) := W(t)W(t)^\top}_{\text{Weight Gram matrix}}, \quad \underbrace{G_U(t) := \Theta^\top U(t)U(t)^\top \Theta}_{\text{Alignment Gram matrix}}, \quad \underbrace{G_t := \Theta^\top W_t W_t^\top \Theta}_{\text{Discrete alignment Gram matrix}}.$$

Let $Z \in \mathbb{R}^{d \times r_s}$ be a Gaussian matrix with i.i.d. entries distributed as $\mathcal{N}(0, 1/d)$. We define $Z_{1:m} \in \mathbb{R}^{m \times r_s}$ as the submatrix formed by the first m rows:

$$Z = \begin{bmatrix} Z_{1:m} \\ Z_{\text{rest}} \end{bmatrix}.$$

Without loss of generality, we assume the teacher directions coincide with the standard basis vectors, i.e., $\theta_j = e_j$. With this, the initialization satisfies:

$$G_W(0) = ZZ^\top, \quad G_U(0) = G_0 = Z_{1:r}(Z^\top Z)^{-1}Z_{1:r}^\top. \quad (\text{B.1})$$

We start with characterizing “good events” for initial matrices given by the following lemma:

Lemma 1.

Both cases ($\alpha \in [0, 0.5) \cup (0.5, \infty)$). For $d \geq \Omega(1)$, the following holds:

$$(E.1) \quad \frac{1}{1.05} \leq \lambda_{\min}(Z^\top Z) \leq \lambda_{\max}(Z^\top Z) \leq 1.05.$$

$$(E.2) \quad 1.05 Z_{1:r} Z_{1:r}^\top \succeq G_U(0) = G_0 \succeq \frac{1}{1.05} Z_{1:r} Z_{1:r}^\top.$$

Heavy-tailed case ($\alpha \in [0, 0.5)$). For $d \geq \Omega_\varphi(1)$, the following holds:

$$(H.1) \quad \text{For } m \leq r_s(1 - \log^{-1/2} d) \wedge r \text{ uniformly, we consider } \lambda_{\min}(Z_{1:m} Z_{1:m}^\top) \geq \frac{r_s}{5d} \left(1 - \frac{m}{r_s}\right)^2.$$

$$(H.2) \quad \text{For all } m \leq r_s(1 - \log^{-1/2} d) \wedge r \text{ uniformly, } \lambda_m(Z_{1:r} Z_{1:r}^\top) \geq \frac{r_s}{5d} \left(1 - \frac{m}{r_s}\right)^2.$$

$$(H.3) \quad \lambda_{\max}(Z_{1:r} Z_{1:r}^\top) \leq \frac{2r_s}{d} \left(1 + \frac{1}{\sqrt{\varphi}}\right)^2.$$

Light-tailed case ($\alpha \in [0, 0.5)$). For $d \geq \Omega(1)$, the following holds:

$$(L.1) \quad \frac{1}{r_s^2 d} \leq \lambda_{\min}(Z_{1:r_s} Z_{1:r_s}^\top)$$

$$(L.2) \quad \text{For } m \in \{1, 2, \dots, 5r_s, \lceil \log^{2.5} d \rceil, \lceil \log^6 d \rceil, r\} \text{ uniformly, } \lambda_{\max}(Z_{1:m} Z_{1:m}^\top) \leq \frac{5(r_s \vee m)}{d}$$

We define

$$\mathcal{G}_{\text{init}} \equiv \begin{cases} (E.1) \cap (E.2) \cap (H.1) \cap (H.2) \cap (H.3), & \alpha \in [0, 0.5) \\ (E.1) \cap (E.2) \cap (L.1) \cap (L.2), & \alpha \in (0.5, \infty) \end{cases}$$

We have

$$\mathbb{P}[\mathcal{G}_{init}] \geq \begin{cases} 1 - 3r_s \exp\left(\frac{-r_s}{2\log^2 d}\right), & \alpha \in [0, 0.5) \\ 1 - \Omega(1/r_s^2), & \alpha > 0.5. \end{cases}$$

Proof. We will use the following:

(S.1) By [Ver10, Corollary 5.35], for $m \leq r_s$ and $\sqrt{r_s} - \sqrt{m} \geq t > 0$

$$\mathbb{P}\left[\lambda_{\min}(\mathbf{Z}_{1:m}\mathbf{Z}_{1:m}^\top) \geq \frac{r_s}{d} \left(1 - \sqrt{\frac{m}{r_s}} - \frac{t}{\sqrt{r_s}}\right)^2\right] \geq 1 - 2e^{-t^2},$$

and for $m \geq r_s$ and $\sqrt{m} - \sqrt{r_s} \geq t > 0$

$$\mathbb{P}\left[\lambda_{\min}(\mathbf{Z}_{1:m}^\top\mathbf{Z}_{1:m}) \geq \frac{m}{d} \left(1 - \sqrt{\frac{r_s}{m}} - \frac{t}{\sqrt{m}}\right)^2\right] \geq 1 - 2e^{-t^2}.$$

(S.2) By [Ver10, Corollary 5.35], for any fixed m

$$\mathbb{P}\left[\frac{m}{d} \left(1 + \sqrt{\frac{r_s}{m}} + \frac{t}{\sqrt{m}}\right)^2 \geq \lambda_{\max}(\mathbf{Z}_{1:m}^\top\mathbf{Z}_{1:m})\right] \geq 1 - 2e^{-t^2}.$$

(S.3) By [Ver10, Theorem 5.38], there exists $C, c > 0$ such that

$$\mathbb{P}\left[\lambda_{\min}(\mathbf{Z}_{1:r_s}\mathbf{Z}_{1:r_s}^\top) \geq \frac{\varepsilon^2}{4dr_s}\right] \geq 1 - C\varepsilon - e^{-cr_s}.$$

For the heavy tailed case, we consider d is large enough to guarantee $|\frac{r_s}{r} - \varphi| \leq \frac{\varphi}{2}$. We have

- By using (S.1) and (S.2) with $m = d$, $t = \sqrt{\frac{d}{\log d}}$, we can show that $\mathbb{P}[(E.1)] \geq 1 - e^{-\frac{d}{\log d}}$ for $d \geq \Omega(1)$.
- By (B.1) and (E.1), (E.2) follows.
- For (H.1), by using (S.1) with $t = \frac{\sqrt{r_s} - \sqrt{m}}{\sqrt{\log d}} \geq \sqrt{\frac{r_s}{2\log^2 d}}$, we have with probability $1 - 2r_s \exp\left(\frac{-r_s}{2\log^2 d}\right)$, for $m \leq r_s(1 - \log^{-1/2} d) \wedge r$ uniformly:

$$\lambda_{\min}(\mathbf{Z}_{1:m}\mathbf{Z}_{1:m}^\top) > \frac{r_s}{d} \left(1 - \frac{1}{\sqrt{\log d}}\right)^2 \left(1 - \sqrt{\frac{m}{r_s}}\right)^2 \geq \frac{r_s}{5d} \left(1 - \frac{m}{r_s}\right)^2.$$

Therefore, $\mathbb{P}[(H.1)] \geq 1 - 2r_s \exp\left(\frac{-r_s}{2\log d}\right)$.

- By Cauchy's eigenvalue interlacing theorem, $\lambda_m(\mathbf{Z}_{1:r}\mathbf{Z}_{1:r}^\top) \geq \lambda_{\min}(\mathbf{Z}_{1:m}\mathbf{Z}_{1:m}^\top)$. Therefore, by (H.1), (H.2) follows.
- For (H.3), by using $m = r$ and $t = 0.4\sqrt{r_s}$ in (S.2), we have $\mathbb{P}[(H.3)] \geq 1 - 2e^{-0.16r_s}$.
- For (L.1), by using (S.3) with $\varepsilon = \frac{2}{r_s^2}$, we have $\mathbb{P}[(L.1)] \geq 1 - \Omega(1/r_s^2)$.
- For (L.2), by using (S.2) with $t = 0.4\sqrt{r_s}$, we have with probability $1 - (10r_s + 6)e^{-0.16r_s}$ for $m \in [5r_s] \cup \{\lceil \log^{2.5} d \rceil, \lceil \log^6 d \rceil, r\}$ uniformly:

$$\lambda_{\max}(\mathbf{Z}_{1:m}\mathbf{Z}_{1:m}^\top) \leq \frac{r_s}{d} \left(1.4 + \sqrt{\frac{m}{r_s}}\right)^2 \leq \frac{5(r_s \vee m)}{d}.$$

By union bound, we have the result. $\square$

C Background: Matrix Riccati Dynamical Systems

We begin by reviewing Riccati dynamical systems in both continuous and discrete time, establishing the necessary background for the arguments that follow. For the following, we define

$$\Lambda_e := \begin{bmatrix} \Lambda & 0 \\ 0 & 0 \end{bmatrix} \text{ and } \tilde{\Lambda} := \frac{\sqrt{r_s}}{\|\Lambda\|_F} \Lambda_e.$$

For notational convenience, we adapt the abuse of notation:

$$\frac{\Lambda_e}{I_d - \exp(-t\Lambda_e)} = \lim_{\varepsilon \rightarrow 0} \frac{(\Lambda_e + \varepsilon I_d)}{I_d - \exp(-t(\Lambda_e + \varepsilon I_d))} = \begin{bmatrix} \frac{\Lambda}{I_r - \exp(-t\Lambda)} & 0 \\ 0 & \frac{1}{t} I_{d-r} \end{bmatrix}.$$

C.1 Continuous-time Matrix Riccati ODE

In this paper, we study continuous-time matrix Riccati differential equations of the following form:

- Weight Gram matrix: For $T_W = r_s$

$$\partial_t G_W(t) = \frac{0.5}{T_W} \left(\tilde{\Lambda} G_W(t) + G_W(t) \tilde{\Lambda} - 2G_W^2(t) \right). \quad (\text{C.1})$$

- Alignment Gram matrix: For $T_U = \|\Lambda\|_F \sqrt{r_s}$,

$$\partial_t G_U(t) = \frac{0.5}{T_U} \left(\Lambda G_U(t) + G_U(t) \Lambda - 2G_U(t) \Lambda G_U(t) \right). \quad (\text{C.2})$$

For $\alpha = 0$, we assume that the ODEs are expressed in the eigenbasis of $G_W(0)$ or $G_U(0)$, ensuring that the trajectories remain diagonal. The solutions of these ODEs are characterized in the following statement:

Lemma 2. (C.1) and (C.2) admit the following solutions:

$$\begin{aligned} G_W(t) &= \frac{\tilde{\Lambda}}{I_d - \exp(-t\tilde{\Lambda}/T_W)} \\ &\quad - \frac{\tilde{\Lambda} \exp(-0.5t\tilde{\Lambda}/T_W)}{I_d - \exp(-t\tilde{\Lambda}/T_W)} \left(G_W(0) + \frac{\tilde{\Lambda} \exp(-t\tilde{\Lambda}/T_W)}{I_d - \exp(-t\tilde{\Lambda}/T_W)} \right)^{-1} \frac{\tilde{\Lambda} \exp(-0.5t\tilde{\Lambda}/T_W)}{I_d - \exp(-t\tilde{\Lambda}/T_W)} \\ G_U(t) &= \frac{I_r}{I_r - \exp(-t\Lambda/T_U)} \\ &\quad - \frac{\exp(-0.5t\Lambda/T_U)}{I_r - \exp(-t\Lambda/T_U)} \left(G_U(0) + \frac{\exp(-t\Lambda/T_U)}{I_r - \exp(-t\Lambda/T_U)} \right)^{-1} \frac{\exp(-0.5t\Lambda/T_U)}{I_r - \exp(-t\Lambda/T_U)} \end{aligned}$$

Moreover, $(G_W(t))_{t \geq 0}$ and $(G_U(t))_{t \geq 0}$ are monotone with respect to $G_W(0) \succeq 0$ and $G_U(0) \succeq 0$ respectively.

Proof. One can check by direct differentiation that the given closed-form expressions satisfy the ODEs above. The uniqueness of the solutions follow the local Lipschitzness of the drifts. Monotonicity is a consequence of Proposition 25. $\square$

C.2 Discrete-time Matrix Riccati Difference Equations

In this section, we will study a particular discretization of Alignment Gram matrix ODE, given as

$$G_{t+1} = G_t - \frac{\eta}{2} (2G_t - I_r) \Lambda (2G_t - I_r) (I_r + \eta \Lambda (2G_t - I_r))^{-1} + \eta \Lambda. \quad (\text{C.3})$$

For convenience, we will make a change of variable and define $V_t := 2\Lambda^{\frac{1}{2}} G_t \Lambda^{\frac{1}{2}} - \Lambda$. We write (C.3) in terms of V_t as follows:

$$V_{t+1} = V_t - \eta V_t^2 (I_r + \eta V_t)^{-1} + \eta \Lambda^2.$$

We characterize the dynamics of $(V_t)_{t \in \mathbb{N}}$ as follows:

Lemma 3. We consider

$$\begin{bmatrix} \mathbf{X}_{t+1,1} \\ \mathbf{X}_{t+1,2} \end{bmatrix} = \begin{bmatrix} \mathbf{X}_{t,1} \\ \mathbf{X}_{t,2} \end{bmatrix} + \eta \mathbf{H} \begin{bmatrix} \mathbf{X}_{t,1} \\ \mathbf{X}_{t,2} \end{bmatrix} \text{ where } \begin{bmatrix} \mathbf{X}_{0,1} \\ \mathbf{X}_{0,2} \end{bmatrix} = \begin{bmatrix} \mathbf{I}_r \\ \mathbf{V}_0 \end{bmatrix} \text{ and } \mathbf{H} := \begin{bmatrix} 0 & \mathbf{I}_r \\ \mathbf{\Lambda}^2 & \eta \mathbf{\Lambda}^2 \end{bmatrix}.$$

The following hold for all $n \in \mathbb{N}$:

(R.1) We have

$$\begin{bmatrix} \mathbf{A}_{t,11} & \mathbf{\Lambda}^{-1} \mathbf{A}_{t,12} \\ \mathbf{\Lambda} \mathbf{A}_{t,12} & \mathbf{A}_{t,22} \end{bmatrix} := (\mathbf{I}_{2r} + \eta \mathbf{H})^t. \quad (\text{C.4})$$

where $\mathbf{A}_{t,11}$, $\mathbf{A}_{t,12}$, $\mathbf{A}_{t,22}$ are positive definite diagonal matrices.

(R.2) $\mathbf{A}_{t,11} + \eta \mathbf{\Lambda} \mathbf{A}_{t,12} = \mathbf{A}_{t,22}$ and $\mathbf{A}_{t,22} \mathbf{A}_{t,11} - \mathbf{A}_{t,12}^2 = \mathbf{I}_r$.

(R.3) For $\eta \leq 1$, we have $\mathbf{A}_{t,11} \mathbf{A}_{t,12}^{-1} \succ (\mathbf{I}_r + \frac{\eta^2}{4} \mathbf{\Lambda}^2)^{1/2} - \frac{\eta}{2} \mathbf{\Lambda}$ and $\mathbf{A}_{t,22} \mathbf{A}_{t,12}^{-1} \succ (\mathbf{I}_r + \frac{\eta^2}{4} \mathbf{\Lambda}^2)^{1/2} + \frac{\eta}{2} \mathbf{\Lambda}$.

(R.4) If $\|\eta \mathbf{\Lambda}\|_2 < 1$,

$$\mathbf{A}_{t,22} \mathbf{A}_{t,12}^{-1} \succ \frac{(\mathbf{I}_r + \eta \mathbf{\Lambda})^t + (\mathbf{I}_r - \eta \mathbf{\Lambda})^t}{(\mathbf{I}_r + \eta \mathbf{\Lambda})^t - (\mathbf{I}_r - \eta \mathbf{\Lambda})^t} \succeq \mathbf{A}_{t,11} \mathbf{A}_{t,12}^{-1}.$$

Moreover, if $\mathbf{X}_{t,1}$ and $\mathbf{X}_{t+1,1}$ are invertible:

(R.7) For $\mathbf{V}_{t+1} := \mathbf{X}_{t+1,2} \mathbf{X}_{t+1,1}^{-1}$, and $\mathbf{V}_t := \mathbf{X}_{t,2} \mathbf{X}_{t,1}^{-1}$, we have

$$\mathbf{V}_{t+1} = \mathbf{V}_t - \eta \mathbf{V}_t^2 (\mathbf{I}_r + \eta \mathbf{V}_t)^{-1} + \eta \mathbf{\Lambda}^2.$$

Proof of Lemma 3. We have

$$\mathbf{I}_{2r} + \eta \mathbf{H} = \begin{bmatrix} \mathbf{I}_r & \eta \mathbf{I}_r \\ \eta \mathbf{\Lambda}^2 & \mathbf{I}_r + \eta^2 \mathbf{\Lambda}^2 \end{bmatrix}. \quad (\text{C.5})$$

Let

$$(\mathbf{I}_{2r} + \eta \mathbf{H})^t =: \begin{bmatrix} \tilde{\mathbf{A}}_{t,11} & \tilde{\mathbf{A}}_{t,12} \\ \tilde{\mathbf{A}}_{21,t} & \tilde{\mathbf{A}}_{t,22} \end{bmatrix}.$$

Since each submatrix in (C.5) is diagonal positive definite, the matrices in (C.4) are also diagonal positive definite. To prove (R.1) and the first part of (R.2) we use proof by induction. We assume $\tilde{\mathbf{A}}_{t,11} + \eta \tilde{\mathbf{A}}_{21,t} = \tilde{\mathbf{A}}_{t,22}$ and $\tilde{\mathbf{A}}_{21,t} \tilde{\mathbf{A}}_{t,22}^{-1} = \mathbf{\Lambda}^2$. We have

$$\begin{aligned} \tilde{\mathbf{A}}_{12,t+1} &= \tilde{\mathbf{A}}_{t,12} + \eta \tilde{\mathbf{A}}_{t,12} \stackrel{(a)}{=} \tilde{\mathbf{A}}_{t,12} + \eta (\tilde{\mathbf{A}}_{t,11} + \eta \tilde{\mathbf{A}}_{21,t}) \stackrel{(b)}{=} \eta \tilde{\mathbf{A}}_{t,11} + (\mathbf{I}_r + \eta^2 \mathbf{\Lambda}^2) \tilde{\mathbf{A}}_{t,12} \\ &\stackrel{(c)}{=} \eta \tilde{\mathbf{A}}_{t,11} + \mathbf{\Lambda}^{-2} (\mathbf{I}_r + \eta^2 \mathbf{\Lambda}^2) \tilde{\mathbf{A}}_{21,t} \\ &= \mathbf{\Lambda}^{-2} \tilde{\mathbf{A}}_{21,t+1}. \end{aligned}$$

where (a) follows the first assumption, (b) and (c) follow the second assumption. Moreover,

$$\begin{aligned} \tilde{\mathbf{A}}_{11,t+1} + \eta \tilde{\mathbf{A}}_{21,t+1} &= \tilde{\mathbf{A}}_{t,11} + \eta \tilde{\mathbf{A}}_{21,t} + \eta^2 \mathbf{\Lambda}^2 \tilde{\mathbf{A}}_{t,11} + \eta (\mathbf{I}_r + \eta^2 \mathbf{\Lambda}^2) \tilde{\mathbf{A}}_{21,t} \\ &= (\mathbf{I}_r + \eta^2 \mathbf{\Lambda}^2) (\tilde{\mathbf{A}}_{t,11} + \eta \tilde{\mathbf{A}}_{21,t}) + \eta \tilde{\mathbf{A}}_{21,t} \\ &\stackrel{(d)}{=} (\mathbf{I}_r + \eta^2 \mathbf{\Lambda}^2) \tilde{\mathbf{A}}_{t,22} + \eta \mathbf{\Lambda}^2 \tilde{\mathbf{A}}_{t,12} = \tilde{\mathbf{A}}_{22,t+1}. \end{aligned}$$

where (d) follows the first and second assumptions. For the second part of (R.2) we again use proof by induction. We assume $\mathbf{A}_{t,22} \mathbf{A}_{t,11} - \mathbf{A}_{t,12}^2 = \mathbf{I}_r$. We have

$$\begin{aligned} \tilde{\mathbf{A}}_{11,t+1} \tilde{\mathbf{A}}_{22,t+1} - \tilde{\mathbf{A}}_{12,t+1} \tilde{\mathbf{A}}_{21,t+1} &= \left(\tilde{\mathbf{A}}_{t,11} + \eta \tilde{\mathbf{A}}_{21,t} \right) \left(\eta \mathbf{\Lambda}^2 \tilde{\mathbf{A}}_{t,12} + (\mathbf{I}_r + \eta^2 \mathbf{\Lambda}^2) \tilde{\mathbf{A}}_{t,12} \right) \\ &\quad - \left(\eta \mathbf{\Lambda}^2 \tilde{\mathbf{A}}_{t,11} + (\mathbf{I}_r + \eta^2 \mathbf{\Lambda}^2) \tilde{\mathbf{A}}_{21,t} \right) \left(\tilde{\mathbf{A}}_{t,12} + \eta \tilde{\mathbf{A}}_{t,12} \right) \end{aligned}$$

$$= \tilde{\mathbf{A}}_{t,11} \tilde{\mathbf{A}}_{t,12} - \tilde{\mathbf{A}}_{t,12} \tilde{\mathbf{A}}_{21,t} = \mathbf{I}_r.$$

For (R.3), by using (R.2), we have

$$\begin{aligned} \mathbf{A}_{t,11} (\mathbf{A}_{t,11} + \eta \mathbf{\Lambda} \mathbf{A}_{t,12}) - \mathbf{A}_{t,12}^2 &= \mathbf{I}_r \Rightarrow (\mathbf{A}_{t,11} \mathbf{A}_{t,12}^{-1})^2 + \eta \mathbf{\Lambda} (\mathbf{A}_{t,11} \mathbf{A}_{t,12}^{-1}) - \mathbf{I}_r \succ 0 \\ &\Rightarrow \mathbf{A}_{t,11} \mathbf{A}_{t,12}^{-1} \succ \left(\mathbf{I}_r + \frac{\eta^2}{4} \mathbf{\Lambda} \right)^{1/2} - \frac{\eta}{2} \mathbf{\Lambda}. \end{aligned}$$

The second part follows (R.2). For (R.4), we recall that

$$\mathbf{A}_{t+1,12} = \mathbf{A}_{t,12} + \eta \mathbf{\Lambda} \mathbf{A}_{t,22} = (\mathbf{I}_r + \eta^2 \mathbf{\Lambda}^2) \mathbf{A}_{t,12} + \eta \mathbf{\Lambda} \mathbf{A}_{t,11} \quad (\text{C.6})$$

$$\mathbf{A}_{t+1,22} = \eta \mathbf{\Lambda} \mathbf{A}_{t,12} + (\mathbf{I}_r + \eta^2 \mathbf{\Lambda}^2) \mathbf{A}_{t,22} \succ \eta \mathbf{\Lambda} \mathbf{A}_{t,12} + \mathbf{A}_{t,22}. \quad (\text{C.7})$$

We use proof by induction. Suppose the lower bound for $\frac{\mathbf{A}_{t,22}}{\mathbf{A}_{t,12}}$ holds. We have

$$\begin{aligned} \frac{\mathbf{A}_{t+1,22}}{\mathbf{A}_{t+1,12}} &\stackrel{(e)}{\succ} \frac{\eta \mathbf{\Lambda} \mathbf{A}_{t,12} + \mathbf{A}_{t,22}}{\mathbf{A}_{t,12} + \eta \mathbf{\Lambda} \mathbf{A}_{t,22}} \\ &\stackrel{(f)}{\succ} \left(\eta \mathbf{\Lambda} + \frac{(\mathbf{I}_r + \eta \mathbf{\Lambda})^t + (\mathbf{I}_r - \eta \mathbf{\Lambda})^t}{(\mathbf{I}_r + \eta \mathbf{\Lambda})^t - (\mathbf{I}_r - \eta \mathbf{\Lambda})^t} \right) \left(\mathbf{I}_r + \eta \mathbf{\Lambda} \frac{(\mathbf{I}_r + \eta \mathbf{\Lambda})^t + (\mathbf{I}_r - \eta \mathbf{\Lambda})^t}{(\mathbf{I}_r + \eta \mathbf{\Lambda})^t - (\mathbf{I}_r - \eta \mathbf{\Lambda})^t} \right)^{-1} \\ &= \frac{(\mathbf{I}_r + \eta \mathbf{\Lambda})^{t+1} + (\mathbf{I}_r - \eta \mathbf{\Lambda})^{t+1}}{(\mathbf{I}_r + \eta \mathbf{\Lambda})^{t+1} - (\mathbf{I}_r - \eta \mathbf{\Lambda})^{t+1}}. \end{aligned}$$

where (e) follows (C.7) and (f) follows the induction hypothesis with that $x \rightarrow \frac{x+\eta\lambda}{1+\eta\lambda x}$ is monotonic increasing for $\eta\lambda < 1$. For the upper bound, suppose the lower bound for $\frac{\mathbf{A}_{t,11}}{\mathbf{A}_{t,12}}$ holds. We have

$$\begin{aligned} \frac{\mathbf{A}_{t+1,11}}{\mathbf{A}_{t+1,12}} &\stackrel{(g)}{\preceq} \frac{\mathbf{A}_{t,11} + \eta \mathbf{\Lambda} \mathbf{A}_{t,12}}{\mathbf{A}_{t,12} + \eta \mathbf{\Lambda} \mathbf{A}_{t,11}} \\ &\stackrel{(h)}{\preceq} \left(\eta \mathbf{\Lambda} + \frac{(\mathbf{I}_r + \eta \mathbf{\Lambda})^t + (\mathbf{I}_r - \eta \mathbf{\Lambda})^t}{(\mathbf{I}_r + \eta \mathbf{\Lambda})^t - (\mathbf{I}_r - \eta \mathbf{\Lambda})^t} \right) \left(\mathbf{I}_r + \eta \mathbf{\Lambda} \frac{(\mathbf{I}_r + \eta \mathbf{\Lambda})^t + (\mathbf{I}_r - \eta \mathbf{\Lambda})^t}{(\mathbf{I}_r + \eta \mathbf{\Lambda})^t - (\mathbf{I}_r - \eta \mathbf{\Lambda})^t} \right)^{-1} \\ &= \frac{(\mathbf{I}_r + \eta \mathbf{\Lambda})^{t+1} + (\mathbf{I}_r - \eta \mathbf{\Lambda})^{t+1}}{(\mathbf{I}_r + \eta \mathbf{\Lambda})^{t+1} - (\mathbf{I}_r - \eta \mathbf{\Lambda})^{t+1}}. \end{aligned}$$

where (g) follows (C.6), and (h) follows the induction hypothesis.

Lastly if $\mathbf{X}_{t,1}$ and $\mathbf{X}_{t+1,1}$ are invertible,

$$\begin{aligned} \mathbf{X}_{t+1,2} \mathbf{X}_{t+1,1}^{-1} &= (\mathbf{X}_{t,2} \mathbf{X}_{t,1}^{-1} + \eta^2 \mathbf{\Lambda}^2 \mathbf{X}_{t,2} \mathbf{X}_{t,1}^{-1} + \eta \mathbf{\Lambda}^2) (\mathbf{I}_r + \eta \mathbf{X}_{t,2} \mathbf{X}_{t,2}^{-1})^{-1} \\ &= \mathbf{X}_{t,2} \mathbf{X}_{t,2}^{-1} (\mathbf{I}_r + \eta \mathbf{X}_{t,2} \mathbf{X}_{t,2}^{-1})^{-1} + \eta \mathbf{\Lambda}^2 \\ &= \mathbf{X}_{t,2} \mathbf{X}_{t,2}^{-1} - \eta \mathbf{X}_{t,2} \mathbf{X}_{t,1}^{-1} \mathbf{X}_{t,2} \mathbf{X}_{t,1}^{-1} (\mathbf{I}_r + \eta \mathbf{X}_{t,2} \mathbf{X}_{t,1}^{-1})^{-1} + \eta \mathbf{\Lambda}^2. \end{aligned}$$

□

Corollary 3. For $\mathbf{V}_0 = 2\mathbf{\Lambda}_2^{\frac{1}{2}} \mathbf{G}_0 \mathbf{\Lambda}_2^{\frac{1}{2}} - \mathbf{\Lambda}_1$, we define

$$\mathbf{V}_{t+1} = \mathbf{V}_t - \eta \mathbf{V}_t^2 (\mathbf{I}_r + \eta \mathbf{V}_t)^{-1} + \eta \hat{\mathbf{\Lambda}}^2.$$

If $\mathbf{\Lambda}_1$, $\mathbf{\Lambda}_2$ and $\hat{\mathbf{\Lambda}}$ are mutually diagonalizable, and $\mathbf{X}_{1,t}$ is invertible for $t \leq t^* \in \mathbb{N}$, we have for $t \leq t^*$

$$\mathbf{G}_t = \frac{\frac{\mathbf{\Lambda}_1}{\mathbf{\Lambda}_2} + \frac{\mathbf{A}_{t,22}}{\mathbf{A}_{t,12}} \frac{\hat{\mathbf{\Lambda}}}{\mathbf{\Lambda}_2}}{2} - \frac{1}{4} \frac{\mathbf{A}_{t,12}^{-1} \hat{\mathbf{\Lambda}}}{\mathbf{\Lambda}_2} \left(\frac{\frac{\hat{\mathbf{\Lambda}}}{\mathbf{\Lambda}_2} \frac{\mathbf{A}_{t,11}}{\mathbf{A}_{t,12}} - \frac{\mathbf{\Lambda}_1}{\mathbf{\Lambda}_2}}{2} + \mathbf{G}_0 \right)^{-1} \frac{\hat{\mathbf{\Lambda}} \mathbf{A}_{t,12}^{-1}}{\mathbf{\Lambda}_2},$$

where $\mathbf{A}_{11t}$, $\mathbf{A}_{t,12}$, $\mathbf{A}_{t,22}$ are defined with $\hat{\mathbf{\Lambda}}$.

Proof. By using (C.4), we can write that

$$\begin{aligned} V_t &= \left(\hat{\Lambda} A_{t,12} + A_{t,22} V_0 \right) \left(\hat{\Lambda} A_{t,12}^{-1} A_{t,11} + V_0 \right)^{-1} \hat{\Lambda} A_{t,12}^{-1} \\ &= \hat{\Lambda} A_{t,12} \left(\hat{\Lambda} A_{t,12}^{-1} A_{t,11} + V_0 \right)^{-1} \hat{\Lambda} A_{t,12}^{-1} \\ &\quad + A_{t,22} \left(I_r - A_{t,11} A_{t,12}^{-1} \hat{\Lambda} \left(\hat{\Lambda} A_{t,12}^{-1} A_{t,11} + V_0 \right)^{-1} \right) \hat{\Lambda} A_{t,12}^{-1} \\ &= A_{t,22} A_{t,12}^{-1} \hat{\Lambda} - A_{t,12}^{-1} \hat{\Lambda} \left(\hat{\Lambda} A_{t,12}^{-1} A_{t,11} + V_0 \right)^{-1} \hat{\Lambda} A_{t,12}^{-1}. \end{aligned}$$

Therefore,

$$G_t = \frac{\frac{\Lambda_1}{\Lambda_2} + \frac{A_{t,22}}{A_{t,12}} \frac{\hat{\Lambda}}{\Lambda_2}}{2} - \frac{1}{4} \frac{A_{t,12}^{-1} \hat{\Lambda}}{\Lambda_2} \left(\frac{\frac{\hat{\Lambda}}{\Lambda_2} \frac{A_{t,11}}{A_{t,12}} - \frac{\Lambda_1}{\Lambda_2}}{2} + G_0 \right)^{-1} \frac{\hat{\Lambda} A_{t,12}^{-1}}{\Lambda_2}.$$

□

Proposition 4. For some symmetric matrix S , we consider

$$V_1 = V_0 + \eta S - \eta V_0^2 (I_r + \eta V_0)^{-1}. \quad (\text{C.8})$$

If $V_0^+ \succeq V_0 \succ \frac{-1}{\eta} I_r$, we have $V_1^+ \succeq V_1$, where V_1^+ is the next iterate if we use V_0^+ in (C.8).

Proof. We have

$$V_1 = \frac{1}{\eta} \left(I_r - (I_r + \eta V_0)^{-1} \right) + \eta S.$$

The statement follows by Proposition 25. □

D Proofs for Main Results

D.1 Proof of Propositions 2 and 3

For Proposition 2, we observe that

$$\hat{G}(t) := \widehat{W}(t) \widehat{W}(t)^\top \text{ and } \tilde{G}(t) := U(t) U(t)^\top = W(t) (W(t)^\top W(t))^{-1} W(t)^\top$$

have the exact same dynamics. Therefore the statement follows. Proposition 3 follows Proposition 25.

D.2 Decomposition of the population risk

The fine-tuning step relies on the following decomposition of the population risk:

Proposition 5. For any $\Omega \in \mathbb{R}^{r_s \times r_s}$, the population risk defined in (2.3) can be written as:

$$R(W_t \Omega) = \frac{1}{r_s} \left\| \Omega \Omega^\top - \frac{\sqrt{r_s}}{\|\Lambda\|_F} W_t^\top \Theta \Lambda \Theta^\top W_t \right\|_F^2 + \frac{1}{\|\Lambda\|_F^2} \left(\|\Lambda\|_F^2 - \|\Lambda^{\frac{1}{2}} G_t \Lambda^{\frac{1}{2}}\|_F^2 \right),$$

where $G_t = \Theta^\top W_t W_t^\top \Theta$ is the discrete alignment Gram matrix defined in the previous part.

We observe that both the second term and the matrix $W_t^\top \Theta \Lambda \Theta^\top W_t$ are independent of Ω . Hence, the fine-tuning step reduces to a least squares problem in the matrix $\Omega \Omega^\top$ in population, which is approximated via empirical risk minimization over a fresh batch of samples. By standard concentration arguments, a sample size of $N_{\text{Ft}} \geq r_s^2 \text{polylog} d$ suffices to ensure that the empirical minimizer approximates the population solution with high probability.

Proof. We begin by noting that W_t is an orthonormal matrix. Using this, we can express the population risk as:

$$R(W_t \Omega) = \left\| \frac{1}{\sqrt{r_s}} W_t \Omega \Omega^\top W_t^\top - \frac{1}{\|\Lambda\|_F} \Theta \Lambda \Theta^\top \right\|_F^2$$

$$\begin{aligned}
&= \frac{1}{\|\Lambda\|_F^2} \left(\|\Lambda\|_F^2 + \frac{\|\Lambda\|_F^2}{r_s} \|\Omega\Omega^\top\|_F^2 - \frac{2\|\Lambda\|_F}{\sqrt{r_s}} \text{Tr}(\Omega\Omega^\top \mathbf{W}_t^\top \Theta \Lambda \Theta^\top \mathbf{W}_t) \pm \|\mathbf{W}_t^\top \Theta \Lambda \Theta^\top \mathbf{W}_t\|_F^2 \right) \\
&= \frac{1}{\|\Lambda\|_F^2} \left(\|\Lambda\|_F^2 - \|\mathbf{W}_t^\top \Theta \Lambda \Theta^\top \mathbf{W}_t\|_F^2 + \left\| \frac{\|\Lambda\|_F}{\sqrt{r_s}} \Omega\Omega^\top - \mathbf{W}_t^\top \Theta \Lambda \Theta^\top \mathbf{W}_t \right\|_F^2 \right)
\end{aligned}$$

By observing that $\|\mathbf{W}^\top \Theta \Lambda \Theta^\top \mathbf{W}\|_F^2 = \|\Lambda^{\frac{1}{2}} \mathbf{G}_t \Lambda^{\frac{1}{2}}\|_F^2$, we have the statement. $\square$

D.3 Proof of Theorem 1

We let

$$t_{\text{sc}} := t\sqrt{r_s}\|\Lambda\|_F, \quad \kappa_{\text{eff}} := \begin{cases} r^\alpha, & \alpha \in [0, 0.5) \\ 1, & \alpha > 0.5 \end{cases}, \quad \mathsf{T}_{\text{eff}} := \kappa_{\text{eff}}\sqrt{r_s}\|\Lambda\|_F \log d/r_s.$$

and

$$r_u := \begin{cases} r, & \alpha \in [0, 0.5) \\ \lceil \log^{2.5} d \rceil, & \alpha > 0.5 \end{cases} \quad r_{u_*} := \begin{cases} \lfloor r_s(1 - \log^{-1/8} d) \wedge r \rfloor, & \alpha \in [0, 0.5) \\ r_s, & \alpha > 0.5. \end{cases}$$

In the following part, we will establish the high-dimensional limit of the risk curve and the alignment.

D.3.1 High-dimensional limit for the alignment

By Lemma 2, we have

$$\mathbf{G}_U(t_{\text{sc}}) = \frac{\mathbf{I}_r}{\mathbf{I}_r - \exp(-t\Lambda)} - \frac{\exp(-0.5t\Lambda)}{\mathbf{I}_r - \exp(-t\Lambda)} \left(\mathbf{G}_U(0) + \frac{\exp(-t\Lambda)}{\mathbf{I}_r - \exp(-t\Lambda)} \right)^{-1} \frac{\exp(-0.5t\Lambda)}{\mathbf{I}_r - \exp(-t\Lambda)}.$$

We define the block matrix forms

$$\mathbf{G}_U(t) =: \begin{bmatrix} \mathbf{G}_{U,11}(t) & \mathbf{G}_{U,12}(t) \\ \mathbf{G}_{U,12}^\top(t) & \mathbf{G}_{U,22}(t) \end{bmatrix}, \quad \Lambda = \begin{bmatrix} \Lambda_{\text{eff}} & 0 \\ 0 & \Lambda_{22} \end{bmatrix}, \quad \Lambda_{\text{e},11} := \Lambda_{\text{eff}}, \quad \Lambda_{\text{e},22} := \begin{bmatrix} \Lambda_{22} & 0 \\ 0 & 0 \end{bmatrix},$$

where $\mathbf{G}_{U,11}(t), \Lambda_{\text{eff}} \in \mathbb{R}^{r_{u_*} \times r_{u_*}}$. The following statement characterizes the time-scales for the alignment terms.

Proposition 6. $\mathcal{G}_{\text{init}}$ implies that $\mathcal{A}(t\mathsf{T}_{\text{eff}}, \theta_j) = \mathbb{1}\{t\kappa_{\text{eff}} \geq \frac{1}{\lambda_j}\} + o_d(1)$ for $t \neq \lim_{d \rightarrow \infty} \frac{1}{\lambda_j \kappa_{\text{eff}}}$ and $j \leq r_{u_*}$.

Proof. For $\alpha = 0$, since the trajectory stays diagonal and the diagonal entries are monotonically increasing, by using the events (E.2) and (H.2) with Lemma 10 we have the result.

In the following, we will prove the result for $\alpha > 0$. By using Proposition 22 with (L.2)

$$\mathbf{G}_U(0) \preceq \begin{bmatrix} 2.1\mathbf{Z}_{1:r_{u_*}} \mathbf{Z}_{1:r_{u_*}}^\top & 0 \\ 0 & 2.1\mathbf{Z}_2 \mathbf{Z}_2^\top \end{bmatrix},$$

where

$$2.1\lambda_{\max}(\mathbf{Z}_{1:r_{u_*}} \mathbf{Z}_{1:r_{u_*}}^\top) \leq \begin{cases} 5(1 + \frac{1}{\sqrt{\varphi}})^2, & \alpha \in [0, 0.5) \\ 15, & \alpha > 0.5. \end{cases}$$

Therefore,

$$\begin{aligned}
\mathbf{G}_{U,11}(t_{\text{sc}}) &\preceq \frac{\mathbf{I}_{r_{u_*}}}{\mathbf{I}_{r_{u_*}} - \exp(-t\Lambda_{\text{eff}})} \\
&\quad - \frac{\exp(-0.5t\Lambda_{\text{eff}})}{\mathbf{I}_{r_{u_*}} - \exp(-t\Lambda_{\text{eff}})} \left(\frac{O(r_s)}{d} \mathbf{I}_{r_{u_*}} + \frac{\exp(-t\Lambda_{\text{eff}})}{\mathbf{I}_{r_{u_*}} - \exp(-t\Lambda_{\text{eff}})} \right)^{-1} \frac{\exp(-0.5t\Lambda_{\text{eff}})}{\mathbf{I}_{r_{u_*}} - \exp(-t\Lambda_{\text{eff}})}.
\end{aligned}$$

Therefore, by Proposition 36, for $j \leq r_{u_*}$,

$$\mathcal{A}(t\mathsf{T}_{\text{eff}}, \theta_j) \leq \frac{1}{1 + \left(\frac{d}{r_s} \frac{1}{\log^3 d} - 1 \right) \frac{d^{-t\kappa_{\text{eff}}j^{-\alpha}}}{r_s}} = \mathbb{1}\{t\kappa_{\text{eff}} \geq \frac{1}{\lambda_j}\} + o_d(1).$$

Moreover, for $t \leq (r_{u_*} + 1)^\alpha \log \frac{d}{r_s}$, by using the events (H.3) and (L.2), we have

$$\mathbf{Z}_2^\top \exp(t\mathbf{\Lambda}_{22})\mathbf{Z}_2 \preceq \begin{cases} O_\varphi(1)\mathbf{I}_{r_s}, & \alpha \in (0, 0.5) \\ O(\log^{2.5} d)\mathbf{I}_{r_s}, & \alpha > 0.5. \end{cases}$$

Therefore, for $t \leq (r_{u_*} + 1)^\alpha \log \frac{d}{r_s}$, we have $\mathbf{G}_{U,11}(t_{sc}) \succeq \underline{\mathbf{G}}(t)$ where

$$\begin{aligned} \underline{\mathbf{G}}(t) &:= \frac{\mathbf{I}_{r_{u_*}}}{\mathbf{I}_{r_{u_*}} - \exp(-t\mathbf{\Lambda}_{\text{eff}})} \\ &\quad - \frac{\exp(-0.5t\mathbf{\Lambda}_{\text{eff}})}{\mathbf{I}_{r_{u_*}} - \exp(-t\mathbf{\Lambda}_{\text{eff}})} \left(\frac{r_s}{d} \frac{O(1)}{\log^4 d} \mathbf{I}_{r_{u_*}} + \frac{\exp(-t\mathbf{\Lambda}_{\text{eff}})}{\mathbf{I}_{r_{u_*}} - \exp(-t\mathbf{\Lambda}_{\text{eff}})} \right)^{-1} \frac{\exp(-0.5t\mathbf{\Lambda}_{\text{eff}})}{\mathbf{I}_{r_{u_*}} - \exp(-t\mathbf{\Lambda}_{\text{eff}})}, \end{aligned}$$

which implies that for $t < (r_{u_*} + 1)^\alpha \log \frac{d}{r_s}$

$$\mathcal{A}(t\mathbf{T}_{\text{eff}}, \boldsymbol{\theta}_j) \geq \frac{1}{1 + O(\log^4 d) \frac{d}{r_s}^{1-t\kappa_{\text{eff}}j^{-\alpha}}} = \mathbb{I}\{t\kappa_{\text{eff}} \geq \frac{1}{\lambda_j}\} + o_d(1).$$

To extend the lower bound for $t > (r_{u_*} + 1)^\alpha \log \frac{d}{r_s}$, let us define

$$t_0 := (r_{u_*} + 1)^\alpha \log \frac{d}{r_s} \quad \text{and} \quad \mathbf{\Lambda}_{\text{eff}}^- := \mathbf{\Lambda}_{\text{eff}} - (r_{u_*} + 1)^{-\alpha} \mathbf{I}_{r_{u_*}}.$$

We have for $t > t_0$,

$$\begin{aligned} \partial_t \mathbf{G}_{U,11}(t) &= \frac{0.5}{T_U} \left(\mathbf{\Lambda}_{\text{eff}}^- \mathbf{G}_{U,11}(t) + \mathbf{G}_{U,11}(t) \mathbf{\Lambda}_{\text{eff}}^- - 2\mathbf{G}_{U,11}(t) \mathbf{\Lambda}_{\text{eff}}^- \mathbf{G}_{U,11}(t) \right) \\ &\quad + \underbrace{\frac{1}{T_U} \left((r_{u_*} + 1)^{-\alpha} \mathbf{G}_{U,11}(t) (\mathbf{I}_{r_{u_*}} - \mathbf{G}_{U,11}(t)) - \mathbf{G}_{U,12}(t) \mathbf{\Lambda}_{22} \mathbf{G}_{U,12}^\top(t) \right)}_{\succeq \frac{(r_{u_*} + 1)^{-\alpha}}{T_U} \left(\mathbf{G}_{U,11}(t) (\mathbf{I}_{r_{u_*}} - \mathbf{G}_{U,11}(t)) - \mathbf{G}_{U,12}(t) \mathbf{G}_{U,12}^\top(t) \right)} \succeq 0 \end{aligned}$$

Therefore, for $t > t_0$, by monotonicity and [BR14, Theorem 38], $\mathbf{G}_{U,11}(t_{sc}) \succeq \underline{\mathbf{G}}(t) \succeq \underline{\mathbf{G}}(t_0)$, where

$$\begin{aligned} \underline{\mathbf{G}}(t) &= \frac{\mathbf{I}_{r_{u_*}}}{\mathbf{I}_{r_{u_*}} - \exp(-(t - t_0)\mathbf{\Lambda}_{\text{eff}}^-)} \\ &\quad - \frac{\exp(-0.5(t - t_0)\mathbf{\Lambda}_{\text{eff}}^-)}{\mathbf{I}_{r_{u_*}} - \exp(-(t - t_0)\mathbf{\Lambda}_{\text{eff}}^-)} \left(\underline{\mathbf{G}}(t_0) + \frac{\exp(-(t - t_0)\mathbf{\Lambda}_{\text{eff}}^-)}{\mathbf{I}_{r_{u_*}} - \exp(-(t - t_0)\mathbf{\Lambda}_{\text{eff}}^-)} \right)^{-1} \frac{\exp(-0.5(t - t_0)\mathbf{\Lambda}_{\text{eff}}^-)}{\mathbf{I}_{r_{u_*}} - \exp(-(t - t_0)\mathbf{\Lambda}_{\text{eff}}^-)}. \end{aligned}$$

Therefore, the result extends to $t > t_0$ as well. $\square$

D.3.2 High-dimensional limit for the risk curve

For $\text{Err}(t) := \|\mathbf{\Lambda}\|_F \left(\frac{\mathbf{\Lambda}_e}{\|\mathbf{\Lambda}\|_F} - \frac{\mathbf{G}_W(t)}{\sqrt{r_s}} \right)$, by Lemma 2, we have

$$\begin{aligned} \text{Err}(t_{sc}) &= \frac{-\mathbf{\Lambda}_e \exp(-t\mathbf{\Lambda}_e)}{\mathbf{I}_d - \exp(-t\mathbf{\Lambda}_e)} \\ &\quad + \frac{\mathbf{\Lambda}_e \exp(-0.5t\mathbf{\Lambda}_e)}{\mathbf{I}_d - \exp(-t\mathbf{\Lambda}_e)} \left(\frac{\|\mathbf{\Lambda}\|_F}{\sqrt{r_s}} \mathbf{G}_W(0) + \frac{\mathbf{\Lambda}_e \exp(-t\mathbf{\Lambda}_e)}{\mathbf{I}_d - \exp(-t\mathbf{\Lambda}_e)} \right)^{-1} \frac{\mathbf{\Lambda}_e \exp(-0.5t\mathbf{\Lambda}_e)}{\mathbf{I}_d - \exp(-t\mathbf{\Lambda}_e)}, \end{aligned} \quad (\text{D.1})$$

We define the block matrix forms

$$\mathbf{G}_W(t) = \begin{bmatrix} \mathbf{G}_{W,11}(t) & \mathbf{G}_{W,12}(t) \\ \mathbf{G}_{W,12}^\top(t) & \mathbf{G}_{W,22}(t) \end{bmatrix}, \quad \mathbf{\Lambda} = \begin{bmatrix} \mathbf{\Lambda}_{\text{eff}} & 0 \\ 0 & \mathbf{\Lambda}_{22} \end{bmatrix}, \quad \mathbf{\Lambda}_{e,11} := \mathbf{\Lambda}_{\text{eff}}, \quad \mathbf{\Lambda}_{e,22} := \begin{bmatrix} \mathbf{\Lambda}_{22} & 0 \\ 0 & 0 \end{bmatrix},$$

where $\mathbf{G}_{W,11}(t), \mathbf{\Lambda}_{\text{eff}} \in \mathbb{R}^{r_u \times r_u}$. Our proof strategy is as follows: In Proposition 7, we show that the off-diagonal and lower-right terms in (D.1) does not contribute to the high-dimensional limit. Then, in Proposition 8, we characterize the limit of the left-top terms. Finally, in Proposition 9, we prove the asymptotic behaviour of the risk curve.

Proposition 7. $\mathcal{G}_{\text{init}}$ implies that $\|G_{W,12}(t\mathbb{T}_{\text{eff}})\|_F^2 = o_d(r_s)$ and $\|G_{W,22}(t\mathbb{T}_{\text{eff}})\|_F^2 = o_d(r_s)$.

Proof. We let

$$D_1 := \frac{\Lambda_{e,11} \exp(-t\Lambda_{e,11})}{I_{r_u} - \exp(-t\Lambda_{e,11})}, \quad D_2 := \frac{\Lambda_{e,22} \exp(-t\Lambda_{e,22})}{I_{d-r_u} - \exp(-t\Lambda_{e,22})}, \quad Z := \begin{bmatrix} Z_{1:r_u} \\ Z_2 \end{bmatrix}.$$

and

$$\begin{bmatrix} S_{11} & S_{12} \\ S_{12}^\top & S_{22} \end{bmatrix} := \left(\begin{bmatrix} \frac{\|\Lambda\|_F}{\sqrt{r_s}} Z_{1:r_u} Z_{1:r_u}^\top + D_1 & \frac{\|\Lambda\|_F}{\sqrt{r_s}} Z_{1:r_u} Z_2^\top \\ \frac{\|\Lambda\|_F}{\sqrt{r_s}} Z_2 Z_{1:r_u}^\top & \frac{\|\Lambda\|_F}{\sqrt{r_s}} Z_2 Z_2^\top + D_2 \end{bmatrix} \right)^{-1}.$$

where

$$\begin{aligned} S_{11} &= \left(D_1 + Z_{1:r_u} \left(\frac{\sqrt{r_s}}{\|\Lambda\|_F} I_{r_s} + Z_2^\top D_2^{-1} Z_2 \right)^{-1} Z_{1:r_u}^\top \right)^{-1}, \\ S_{12} &= - \left(D_1 + Z_{1:r_u} \left(\frac{\sqrt{r_s}}{\|\Lambda\|_F} I_{r_s} + Z_2^\top D_2^{-1} Z_2 \right)^{-1} Z_{1:r_u}^\top \right)^{-1} Z_{1:r_u} Z_2^\top (Z_2 Z_2^\top + D_2)^{-1}, \\ S_{22} &= \left(D_2 + Z_2 \left(\frac{\sqrt{r_s}}{\|\Lambda\|_F} I_{r_s} + Z_{1:r_u}^\top D_1^{-1} Z_{1:r_u} \right)^{-1} Z_2^\top \right)^{-1}. \end{aligned}$$

Off-diagonal terms: By Proposition 26

$$\begin{aligned} \tilde{G}_{W,12}(t) &:= \frac{\Lambda_{e,11} \exp(-0.5t\Lambda_{e,11})}{I_{r_u} - \exp(-t\Lambda_{e,11})} S_{12} \frac{\Lambda_{e,22} \exp(-0.5t\Lambda_{e,11})}{I_{d-r_u} - \exp(-t\Lambda_{e,22})} \\ &= \exp(0.5t\Lambda_{e,11}) Z_{1:r_u} \left(\frac{\sqrt{r_s}}{\|\Lambda\|_F} I_{r_s} + Z_2^\top D_2^{-1} Z_2 + Z_{1:r_u}^\top D_1^{-1} Z_{1:r_u} \right)^{-1} Z_2^\top \exp(0.5t\Lambda_{e,22}). \end{aligned}$$

We observe that $\lambda_{\max}(D_2) \leq \frac{1}{t}$ and $\frac{\sqrt{r_s}}{\|\Lambda\|_F} \asymp \kappa_{\text{eff}}$. By using Proposition 27 with $\mathcal{G}_{\text{init}}$ and $\tilde{t} := t\kappa_{\text{eff}} \log \frac{d}{r_s}$, we write

$$\begin{aligned} \frac{1}{r_s} \|G_{W,12}(t\mathbb{T}_{\text{eff}})\|_F^2 &= \frac{1}{\|\Lambda\|_F^2} \|\tilde{G}_{W,12}(\tilde{t})\|_F^2 \\ &\leq \frac{1}{\|\Lambda\|_F^2} \frac{O(1)}{(\kappa_{\text{eff}} + \tilde{t})^2} \sum_{i=1}^{r_u \wedge r_s} \left(\lambda_{\max}(Z_{1:r_u} Z_{1:r_u}^\top) \exp(\tilde{t}\lambda_i) \wedge (\kappa_{\text{eff}} + \tilde{t}) \frac{\lambda_i \exp(\tilde{t}\lambda_i)}{\exp(\tilde{t}\lambda_i) - 1} \right). \quad (\text{D.2}) \end{aligned}$$

For the heavy-tailed case ($\alpha \in [0, 0.5)$),

$$(\text{D.2}) \leq \frac{O_{\alpha, \varphi, \beta}(1)}{r \log^2 d} \sum_{i \leq r} \mathbb{1}\{\tilde{t}\lambda_i \leq \log \frac{d}{r_s}\} + \frac{O_{\alpha, \varphi, \beta}(1)}{r^{1-\alpha} \log d} \sum_{i \leq r} \lambda_i \mathbb{1}\{\tilde{t}\lambda_i > \log \frac{d}{r_s}\} = o_d(1).$$

For the light-tailed case ($\alpha > 0.5$),

$$(\text{D.2}) \leq \frac{O_{\alpha, r_s, \beta}(1)}{\log^2 d} \sum_{i \leq r_s} \mathbb{1}\{\tilde{t}\lambda_i \leq \log \frac{d}{r_u r_s}\} + \frac{O_{\alpha, \varphi, \beta}(1)}{\log d} \sum_{i \leq r_s} \lambda_i \mathbb{1}\{\tilde{t}\lambda_i > \log \frac{d}{r_u r_s}\} = o_d(1).$$

Lower-right terms: By using Matrix-Inversion lemma, we have

$$-D_2 + D_2 S_{22} D_2 = -Z_2 \left(\frac{\sqrt{r_s}}{\|\Lambda\|_F} I_{r_s} + Z_{1:r_u}^\top D_1^{-1} Z_{1:r_u} + Z_2^\top D_2^{-1} Z_2 \right)^{-1} Z_2^\top.$$

We observe that $\lambda_{\max}(D_2) \leq \frac{1}{t}$ and $\frac{\sqrt{r_s}}{\|\Lambda\|_F} \asymp \kappa_{\text{eff}}$. By using $\tilde{t} := t\kappa_{\text{eff}} \log \frac{d}{r_s}$, we have

$$\frac{1}{\sqrt{r_s}} G_{W,22}(t\mathbb{T}_{\text{eff}}) = \frac{1}{\|\Lambda\|_F} Z_2 \left(\frac{\sqrt{r_s}}{\|\Lambda\|_F} I_{r_s} + Z_{1:r_u}^\top D_1^{-1} Z_{1:r_u} + Z_2^\top D_2^{-1} Z_2 \right)^{-1} Z_2^\top$$

$$\preceq \frac{O(1)}{\|\mathbf{\Lambda}\|_F} \mathbf{Z}_2 (\kappa_{\text{eff}} \mathbf{I}_{r_s} + \tilde{t} \mathbf{Z}_2^\top \mathbf{Z}_2)^{-1} \mathbf{Z}_2^\top. \quad (\text{D.3})$$

For the heavy-tailed case ($\alpha \in [0, 0.5)$),

$$\|(\text{D.3})\|_F^2 \leq \frac{O_{\alpha, \varphi, \beta}(1)}{r^{-2\alpha} \log^2 d} \frac{1}{t^2 r^{2\alpha} \log^2 d} = o_d(1).$$

For the light-tailed case ($\alpha > 0.5$),

$$\|(\text{D.3})\|_F^2 \leq O_{\alpha, r_s, \beta}(1) \frac{1}{t^2 \log^2 d} = o_d(1).$$

□

Proposition 8. For some $c > 0$, let $\mathbf{G}(0) := \frac{c}{t} \mathbf{Z}_{1:r_u} \mathbf{Z}_{1:r_u}^\top$ and

$$t \in \begin{cases} (0, \frac{(r_{u_*}+1)^\alpha}{\kappa_{\text{eff}}}), & \alpha \in [0, 0.5) \\ (0, \frac{(r_{u_*}+1)^\alpha}{\kappa_{\text{eff}}}) \setminus \{j^\alpha : j \in \mathbb{N}\}, & \alpha > 0.5, \end{cases} \quad d \geq \begin{cases} \Omega_{\varphi, \alpha}(1) & \alpha \in [0, 0.5) \\ \Omega_{r_s, \alpha}(1), & \alpha > 0.5. \end{cases}$$

We define

$$\begin{aligned} \text{Err}_{r_u}(t_{\text{sc}}) &:= \frac{-\mathbf{\Lambda}_{\text{eff}} \exp(-t\mathbf{\Lambda}_{\text{eff}})}{\mathbf{I}_{r_u} - \exp(-t\mathbf{\Lambda}_{\text{eff}})} \\ &+ \frac{\mathbf{\Lambda}_{\text{eff}} \exp(-0.5t\mathbf{\Lambda}_{\text{eff}})}{\mathbf{I}_{r_u} - \exp(-t\mathbf{\Lambda}_{\text{eff}})} \left(\frac{\mathbf{\Lambda}_{\text{eff}} \exp(-t\mathbf{\Lambda}_{\text{eff}})}{\mathbf{I}_{r_u} - \exp(-t\mathbf{\Lambda}_{\text{eff}})} + \mathbf{G}(0) \right)^{-1} \frac{\mathbf{\Lambda}_{\text{eff}} \exp(-0.5t\mathbf{\Lambda}_{\text{eff}})}{\mathbf{I}_{r_u} - \exp(-t\mathbf{\Lambda}_{\text{eff}})}. \end{aligned}$$

$\mathcal{G}_{\text{init}}$ implies that

$$\frac{\|\text{Err}_{r_u}(t_{\text{sc}})\|_F^2}{\|\mathbf{\Lambda}\|_F^2} = 1 - \frac{1}{\|\mathbf{\Lambda}\|_F^2} \sum_{j=1}^{r_{u_*}} \lambda_j^2 \mathbb{1}\{t\kappa_{\text{eff}} \geq \frac{1}{\lambda_j}\} + o_d(1).$$

Proof. We let

$$\mathbf{Z}_{1:r_u} \mathbf{Z}_{1:r_u}^\top := \begin{bmatrix} \mathbf{Z}_{1:r_{u_*}} \mathbf{Z}_{1:r_{u_*}}^\top & \mathbf{Z}_1 \mathbf{Z}_2^\top \\ \mathbf{Z}_2 \mathbf{Z}_1^\top & \mathbf{Z}_2 \mathbf{Z}_2^\top \end{bmatrix}, \quad \mathbf{\Lambda}_{\text{eff}} := \begin{bmatrix} \mathbf{\Lambda}_{\text{eff},11} & 0 \\ 0 & \mathbf{\Lambda}_{\text{eff},22} \end{bmatrix},$$

where $\mathbf{\Lambda}_{\text{eff},11} \in \mathbb{R}^{r_{u_*} \times r_{u_*}}$, $\mathbf{Z}_2 \in \mathbb{R}^{(r_u - r_{u_*}) \times r_s}$. Let

$$\mathbf{\Gamma}(t_{\text{sc}}) := \left(\frac{\mathbf{\Lambda}_{\text{eff}} \exp(-t\mathbf{\Lambda}_{\text{eff}})}{\mathbf{I}_{r_u} - \exp(-t\mathbf{\Lambda}_{\text{eff}})} + \frac{c}{t} \mathbf{Z}_{1:r_u} \mathbf{Z}_{1:r_u}^\top \right)^{-1} \quad \text{and} \quad \text{Err}(t_{\text{sc}}, \mathbf{\Gamma}) := \text{Err}_{r_u}(t_{\text{sc}}).$$

By using Proposition 22 and the events (H.3) and (L.2),

$$\mathbf{\Gamma}(t_{\text{sc}}) \succeq \underline{\mathbf{\Gamma}}(t_{\text{sc}}) := \begin{bmatrix} \frac{10c}{t} \frac{r_s}{d} \mathbf{I}_{r_{u_*}} + \frac{\mathbf{\Lambda}_{\text{eff},11} \exp(-t\mathbf{\Lambda}_{\text{eff},11})}{\mathbf{I}_{r_{u_*}} - \exp(-t\mathbf{\Lambda}_{\text{eff},11})} & 0 \\ 0 & \frac{2c}{t} \frac{r_s \log^{2.5} d}{d} \mathbf{I}_{r_u - r_{u_*}} + \frac{\mathbf{\Lambda}_{\text{eff},22} \exp(-t\mathbf{\Lambda}_{\text{eff},22})}{\mathbf{I}_{r_u - r_{u_*}} - \exp(-t\mathbf{\Lambda}_{\text{eff},22})} \end{bmatrix}^{-1}.$$

For the upper bound, by using $\varepsilon = \frac{1}{c \log^3 d}$ in Proposition 24, and the events (H.1), (L.1) and (H.3), (L.2), we have for $t \leq (r_{u_*} + 1)^\alpha \log \frac{d}{r_s}$,

$$\mathbf{\Gamma}(t_{\text{sc}}) \preceq \bar{\mathbf{\Gamma}}(t_{\text{sc}}) := \begin{bmatrix} \frac{0.2/t}{\log^4 d} \frac{r_s}{d} \mathbf{I}_{r_{u_*}} + \frac{\mathbf{\Lambda}_{\text{eff},11} \exp(-t\mathbf{\Lambda}_{\text{eff},11})}{\mathbf{I}_{r_{u_*}} - \exp(-t\mathbf{\Lambda}_{\text{eff},11})} & 0 \\ 0 & \frac{-1.1/t}{\log^{1/2} d} \frac{r_s}{d} \mathbf{I}_{r_u - r_{u_*}} + \frac{\mathbf{\Lambda}_{\text{eff},22} \exp(-t\mathbf{\Lambda}_{\text{eff},22})}{\mathbf{I}_{r_u - r_{u_*}} - \exp(-t\mathbf{\Lambda}_{\text{eff},22})} \end{bmatrix}^{-1}.$$

Therefore, for $t < \frac{(r_{u_*}+1)^\alpha}{\kappa_{\text{eff}}}$, by Corollary 8, we have

$$\begin{aligned} \frac{\|\text{Err}(t\mathbf{T}_{\text{eff}}, \mathbf{\Gamma})\|_F^2}{\|\mathbf{\Lambda}\|_F^2} &\geq \frac{\|\text{Err}(t\mathbf{T}_{\text{eff}}, \mathbf{\underline{\Gamma}})\|_F^2}{\|\mathbf{\Lambda}\|_F^2} = \frac{1}{\|\mathbf{\Lambda}\|_F^2} \sum_{j=1}^{r_u} \lambda_j^2 \mathbb{1}\{\frac{1}{\lambda_j} > t\kappa_{\text{eff}}\} + o_d(1) \\ &= 1 - \frac{1}{\|\mathbf{\Lambda}\|_F^2} \sum_{j=1}^{r_{u_*}} \lambda_j^2 \mathbb{1}\{t\kappa_{\text{eff}} \geq \frac{1}{\lambda_j}\} + o_d(1). \end{aligned}$$

On the other hand, by Corollary 8, we have

$$\begin{aligned} \frac{\|\text{Err}(t\mathbf{T}_{\text{eff}}, \mathbf{\Gamma})\|_F^2}{\|\mathbf{\Lambda}\|_F^2} &\leq \frac{\|\text{Err}(t\mathbf{T}_{\text{eff}}, \mathbf{\bar{\Gamma}})\|_F^2}{\|\mathbf{\Lambda}\|_F^2} \\ &= \frac{1}{\|\mathbf{\Lambda}\|_F^2} \sum_{j=1}^{r_{u_*}} \lambda_j^2 \mathbb{1}\{\frac{1}{\lambda_j} > t\kappa_{\text{eff}}\} + \frac{1}{\|\mathbf{\Lambda}\|_F^2} \sum_{j=r_{u_*}+1}^{r_u} \lambda_j^2 + o_d(1) \\ &= 1 - \frac{1}{\|\mathbf{\Lambda}\|_F^2} \sum_{j=1}^{r_{u_*}} \lambda_j^2 \mathbb{1}\{t\kappa_{\text{eff}} \geq \frac{1}{\lambda_j}\} + o_d(1). \end{aligned}$$

□

Proposition 9. $\mathcal{G}_{\text{init}}$ implies that

$$\frac{\|\text{Err}(t\mathbf{T}_{\text{eff}})\|_F^2}{\|\mathbf{\Lambda}\|_F^2} = 1 - \frac{1}{\|\mathbf{\Lambda}\|_F^2} \sum_{j=1}^{r_{u_*}} \lambda_j^2 \mathbb{1}\{t\kappa_{\text{eff}} \geq \frac{1}{\lambda_j}\} + o_d(1),$$

for

$$t \in \begin{cases} (0, \infty), & \alpha \in [0, 0.5) \\ (0, \infty) \setminus \{j^\alpha : j \in \mathbb{N}\}, & \alpha > 0.5, \end{cases} \quad d \geq \begin{cases} \Omega_{\varphi, \alpha}(1) & \alpha \in [0, 0.5) \\ \Omega_{r_s, \alpha}(1), & \alpha > 0.5. \end{cases}$$

Proof. We recall that

$$\mathbf{D}_1 := \frac{\mathbf{\Lambda}_{e,11} \exp(-t\mathbf{\Lambda}_{e,11})}{\mathbf{I}_{r_u} - \exp(-t\mathbf{\Lambda}_{e,11})}, \quad \mathbf{D}_2 := \frac{\mathbf{\Lambda}_{e,22} \exp(-t\mathbf{\Lambda}_{e,22})}{\mathbf{I}_{d-r_u} - \exp(-t\mathbf{\Lambda}_{e,22})}, \quad \mathbf{Z} := \begin{bmatrix} \mathbf{Z}_{1:r_u} \\ \mathbf{Z}_2 \end{bmatrix},$$

and

$$\begin{aligned} \text{Err}(t_{\text{sc}}) &= \frac{-\mathbf{\Lambda}_e \exp(-t\mathbf{\Lambda}_e)}{\mathbf{I}_d - \exp(-t\mathbf{\Lambda}_e)} \\ &\quad + \frac{\mathbf{\Lambda}_e \exp(-0.5t\mathbf{\Lambda}_e)}{\mathbf{I}_d - \exp(-t\mathbf{\Lambda}_e)} \left(\frac{\|\mathbf{\Lambda}\|_F}{\sqrt{r_s}} \mathbf{Z} \mathbf{Z}^\top + \frac{\mathbf{\Lambda}_e \exp(-t\mathbf{\Lambda}_e)}{\mathbf{I}_d - \exp(-t\mathbf{\Lambda}_e)} \right)^{-1} \frac{\mathbf{\Lambda}_e \exp(-0.5t\mathbf{\Lambda}_e)}{\mathbf{I}_d - \exp(-t\mathbf{\Lambda}_e)} \\ &= \begin{bmatrix} \text{Err}_{r_u}(t_{\text{sc}}) & \frac{-1}{\sqrt{r_s}} \mathbf{G}_{W,12}(t_{\text{sc}}) \\ \frac{-1}{\sqrt{r_s}} \mathbf{G}_{W,12}(t_{\text{sc}}) & \frac{-1}{\sqrt{r_s}} \mathbf{G}_{W,22}(t_{\text{sc}}) \end{bmatrix}. \end{aligned}$$

Note that by Proposition 35, in the time scale we consider we have $\frac{1-o_d(1)}{t\mathbf{T}_{\text{eff}}} \leq \lambda_{\min}(\mathbf{D}_2) \leq \lambda_{\max}(\mathbf{D}_2) \leq \frac{1}{t\mathbf{T}_{\text{eff}}}$. The by using (E.1),

$$\begin{aligned} \text{Err}_{r_u}(t_{\text{sc}}) &= \frac{-\mathbf{\Lambda}_{\text{eff}} \exp(-t\mathbf{\Lambda}_{\text{eff}})}{\mathbf{I}_{r_u} - \exp(-t\mathbf{\Lambda}_{\text{eff}})} \\ &\quad + \frac{\mathbf{\Lambda}_{\text{eff}} \exp(-0.5t\mathbf{\Lambda}_{\text{eff}})}{\mathbf{I}_{r_u} - \exp(-t\mathbf{\Lambda}_{\text{eff}})} \left(\frac{\Theta(1)}{\frac{\|\mathbf{\Lambda}\|_F}{\sqrt{r_s}} + t} \mathbf{Z}_{1:r_u} \mathbf{Z}_{1:r_u}^\top + \frac{\mathbf{\Lambda}_e \exp(-t\mathbf{\Lambda}_{\text{eff}})}{\mathbf{I}_{r_u} - \exp(-t\mathbf{\Lambda}_{\text{eff}})} \right)^{-1} \frac{\mathbf{\Lambda}_e \exp(-0.5t\mathbf{\Lambda}_{\text{eff}})}{\mathbf{I}_{r_u} - \exp(-t\mathbf{\Lambda}_{\text{eff}})}. \end{aligned}$$

By Propositions 7 and 8, we have

$$\frac{\|\text{Err}(t\mathbf{T}_{\text{eff}})\|_F^2}{\|\mathbf{\Lambda}\|_F^2} = 1 - \frac{1}{\|\mathbf{\Lambda}\|_F^2} \sum_{j=1}^{(r_s \wedge r)} \lambda_j^2 \mathbb{1}\{t\kappa_{\text{eff}} \geq \frac{1}{\lambda_j}\} + o_d(1),$$

for

$$t \in \begin{cases} (0, \frac{(r_{u_*}+1)^\alpha}{\kappa_{\text{eff}}}), & \alpha \in [0, 0.5) \\ (0, \frac{(r_{u_*}+1)^\alpha}{\kappa_{\text{eff}}}) \setminus \{j^\alpha : j \in \mathbb{N}\}, & \alpha > 0.5. \end{cases} \quad (\text{D.4})$$

To extend the limit for $t > \frac{(r_{u_*}+1)^\alpha}{\kappa_{\text{eff}}}$, we observe that

- $\|\text{Err}(t)\|_F^2$ non increasing since it corresponds to the objective under (GF).
- The global optimum of (GF) and the previous item with (D.4) guarantees that for $t > \frac{(r_{u_*}+1)^\alpha}{\kappa_{\text{eff}}}$,

$$1 - \frac{1}{\|\mathbf{\Lambda}\|_F^2} \sum_{j=1}^{(r_s \wedge r)} \lambda_j^2 \mathbb{1}\{t\kappa_{\text{eff}} \geq \frac{1}{\lambda_j}\} \leq \frac{\|\text{Err}(t\mathbf{T}_{\text{eff}})\|_F^2}{\|\mathbf{\Lambda}\|_F^2} \leq 1 - \frac{1}{\|\mathbf{\Lambda}\|_F^2} \sum_{j=1}^{(r_s \wedge r)} \lambda_j^2 \mathbb{1}\{t\kappa_{\text{eff}} \geq \frac{1}{\lambda_j}\} + o_d(1).$$

Therefore, the statement extends to $t > \frac{(r_{u_*}+1)^\alpha}{\kappa_{\text{eff}}}$. $\square$

D.4 Proof of Theorem 2

We redefine the time-scale and effective-width as:

$$t_{\text{sc}} = t\sqrt{r_s}\|\mathbf{\Lambda}\|_F, \quad \kappa_{\text{eff}} = \begin{cases} r^\alpha/\eta, & \alpha \in [0, 0.5) \\ 1/\eta, & \alpha > 0.5 \end{cases}, \quad \mathbf{T}_{\text{eff}} = \kappa_{\text{eff}}\sqrt{r_s}\|\mathbf{\Lambda}\|_F \log d/r_s.$$

and

$$r_u = \begin{cases} r, & \alpha \in [0, 0.5) \\ \lceil \log^{2.5} d \rceil, & \alpha > 0.5 \end{cases}, \quad r_{u_*} := \begin{cases} \lfloor r_s(1 - \log^{-1/8} d) \wedge r \rfloor, & \alpha \in [0, 0.5) \\ r_s, & \alpha > 0.5. \end{cases}$$

We consider the learning rate and fine-tuning sample size given as

$$\eta \asymp \frac{1}{d} \begin{cases} \frac{1}{r^\alpha \log^{20}(1+d/r_s)}, & \alpha \in [0, 0.5) \\ \frac{1}{r_u^{4\alpha+3} \log^{18} d}, & \alpha > 0.5 \end{cases} \quad \text{and} \quad N_{\text{Ft}} \asymp r_s^2 \log^5 d.$$

We define the effective learning rate η and the hitting time $\mathcal{T}_{\text{hit}}$ as follows:

$$\eta := \frac{\eta/2}{\|\mathbf{\Lambda}\|_F \sqrt{r_s}}, \quad \mathcal{T}_{\text{hit}} := \left\{ t \geq 0 \mid 1 - \frac{\|\mathbf{\Lambda}^{\frac{1}{2}} \mathbf{G}_t \mathbf{\Lambda}^{\frac{1}{2}}\|_F^2}{\|\mathbf{\Lambda}\|_F^2} \leq \frac{1}{\|\mathbf{\Lambda}\|_F^2} \sum_{j=(r_s \wedge 1)+1}^r \lambda_j^2 + \frac{10}{\log^{\frac{1}{8}} d} \right\}.$$

We note that bounding $\mathcal{T}_{\text{hit}}$ suffices to derive sample complexity since by Proposition 11, we have

$$R(\mathbf{W}_t^{\text{final}}) \leq 1 - \frac{\|\mathbf{\Lambda}^{\frac{1}{2}} \mathbf{G}_t \mathbf{\Lambda}^{\frac{1}{2}}\|_F^2}{\|\mathbf{\Lambda}\|_F^2} + \frac{O(1)}{\log d}.$$

The main statement of this part is as follows:

Proposition 10. *The intersection of the following events hold with probability $1 - o_d(1/d^2) - \Omega(1/r_s^2)$:*

1. We have

$$\mathcal{T}_{\text{hit}} \leq \begin{cases} \frac{1}{2\eta} \left(r_s(1 - \log^{-1/8} d) \wedge r \right)^\alpha \log \left(\frac{20d \log^{\frac{3}{4}}(1+d/r_s)}{r_s} \right), & \alpha \in [0, 0.5) \\ \frac{1}{2\eta} r_s^\alpha \log \left(20 \frac{d \log^{3/4} d}{r_s} \right), & \alpha > 0.5. \end{cases}$$

2. For $t > 0$,

- $\mathcal{A}(t\mathbf{T}_{\text{eff}}, \boldsymbol{\theta}_j) = \mathbb{1}\{\eta t\kappa_{\text{eff}} \geq \frac{1}{\lambda_j}\} + o_d(1)$ for $t \neq \lim_{d \rightarrow \infty} \frac{1}{\eta\kappa_{\text{eff}}\lambda_j}$ and $j \leq r_{u_*}$.
- $\|\mathbf{\Lambda}\|_F^2 - \|\mathbf{\Lambda}^{\frac{1}{2}} \mathbf{G}_{t\mathbf{T}_{\text{eff}}} \mathbf{\Lambda}^{\frac{1}{2}}\|_F^2 = 1 - \sum_{j=1}^{r_{u_*}} \lambda_j^2 \mathbb{1}\{\eta t\kappa_{\text{eff}} \geq \frac{1}{\lambda_j}\} + o_d(\|\mathbf{\Lambda}\|_F^2)$.

Proof. By using Lemma 1 and Corollary 5, we have with probability $1 - o_d(1/d^2) - \Omega(1/r_s^2)$:

$$\mathcal{T}_{\text{bad}} \geq \begin{cases} \frac{1}{2\eta} \left(r_s (1 - \log^{-\frac{1}{2}} d) \wedge r \right)^\alpha \log \left(\frac{d \log^{1.5} d}{r_s} \right), & \alpha \in [0, 0.5) \\ \frac{1}{2\eta} r_s^\alpha \log \left(\frac{d \log^{1.5} d}{r_s} \right), & \alpha > 0.5, \end{cases}$$

where $\mathcal{T}_{\text{bad}}$ is defined in (F.14). Given the lower bound, by Proposition 14, and the third item of Proposition 15, we have the first item.

For the second item, by Proposition 14, and Proposition 15 (for the lower bound) and Proposition 16 (for the upper bound), we have for $r_{u_*} \times r_{u_*}$ dimensional top left submatrices $\mathbf{G}_{t,11}$ and $\mathbf{\Lambda}_{11}$,

$$\frac{1}{\frac{1.2}{C_{\text{lb}}} \frac{d}{r_s} \exp(-2t\eta\mathbf{\Lambda}_{11}) + 1} - o_d(1) \preceq \mathbf{G}_{t,11} \preceq \left(\frac{C_{\text{ub}} r_s}{d} \exp(2\eta t\mathbf{\Lambda}_{11}) \wedge 1 \right) + o_d(1), \quad (\text{D.5})$$

and

$$\begin{aligned} \|\mathbf{\Lambda}\|_F^2 - \|\mathbf{\Lambda}^{\frac{1}{2}} \mathbf{G}_t \mathbf{\Lambda}^{\frac{1}{2}}\|_F^2 &\geq \sum_{i=1}^{r_u} \lambda_i^2 \left(1 - \frac{C_{\text{ub}} r_s}{d} \exp(2\eta t \lambda_i) \right)_+ - o_d(\|\mathbf{\Lambda}\|_F^2) \\ \|\mathbf{\Lambda}\|_F^2 - \|\mathbf{\Lambda}^{\frac{1}{2}} \mathbf{G}_t \mathbf{\Lambda}^{\frac{1}{2}}\|_F^2 &\leq \sum_{i=(r_{u_*} \wedge r)+1}^r \lambda_i^2 + \sum_{i=1}^{r_{u_*}} \lambda_i^2 \left(1 - \frac{1}{\frac{1.2}{C_{\text{lb}}} \frac{d}{r_s} \exp(-2t\eta \lambda_i) + 1} \right)^2 + o_d(\|\mathbf{\Lambda}\|_F^2) \end{aligned} \quad (\text{D.6})$$

for

$$t \leq \begin{cases} \frac{1}{2\eta} \left(r_s (1 - \log^{-1/8} d) \wedge r \right)^\alpha \log \left(\frac{d \log^{1.5} d}{r_s} \right), & \alpha \in [0, 0.5) \\ \frac{1}{2\eta} r_s^\alpha \log \left(\frac{d \log^{1.5} d}{r_s} \right), & \alpha > 0.5. \end{cases} \quad (\text{D.7})$$

where

$$C_{\text{ub}} = \begin{cases} 2.5 \left(1 + \frac{1}{\sqrt{\varphi}} \right)^2, & \alpha \in [0, 0.5) \\ 15, & \alpha > 0.5, \end{cases} \quad C_{\text{lb}} = \frac{1}{15} \begin{cases} \log^{-1/2} d, & \alpha \in [0, 0.5) \\ r_s^{-6}, & \alpha > 0.5. \end{cases}$$

The high-dimensional limits of the alignment and risk up to the time horizon in (D.7) follow from (D.5) (for the alignment), and from (D.6) (for the risk) by Proposition 36. Proposition 21 then allows us to extend these results beyond the time limit in (D.7), yielding the full statement. $\square$

D.5 Proof of Corollary 1 and Corollary 2

Finally, we derive the scaling of prediction risk under power-law second-layer coefficients. Since Corollary 2 is a rescaled version of Corollary 1, we will only consider the latter.

Proof of Corollary 1. We will prove heavy and light-tailed cases separately.

Heavy-tailed case ($\alpha \in [0, 0.5)$): We define $C := \left(\frac{(1-\beta)\sqrt{\varphi}}{\sqrt{1-2\alpha}} \right)^{\frac{1}{\alpha}}$. We first fix a $(C\varphi)^\alpha > t > 0$. By Proposition 9, for any $d \geq \Omega_{\varphi,\alpha}(1)$, we have with probability at least $1 - o(1/d^2)$

$$\mathcal{R}(tr \log d) = 1 - \underbrace{\frac{1}{\|\mathbf{\Lambda}\|_F^2} \sum_{i=1}^{r_s} \lambda_j^2 \mathbb{1}\left\{ \frac{tr^\alpha}{C^\alpha} \geq \frac{1 \pm o_d(1)}{\lambda_j} \right\}}_{:= \mathcal{R}_d((Ct)^{\frac{1}{\alpha}})} + o_d(1)$$

where we define $\mathcal{R}_d((Ct)^{\frac{1}{\alpha}})$ to isolate the main term and make the dependence on the ambient dimension explicit. By using $\lambda_j = j^{-\alpha}$ in the indicator function, we can rewrite

$$\mathcal{R}_d(t) = 1 - \frac{1}{\|\mathbf{\Lambda}\|_F^2} \sum_{i=1}^{r_s} \lambda_j^2 \mathbb{1}\left\{ (1 \pm o_d(1))t \geq \frac{j}{r} \right\}$$

We define a sequence of measures supported on $\{j/r : j \in [r]\}$, where $\mu_d\{\frac{j}{r}\} \propto j^{-2\alpha}$ for $j = 1, \dots, r$. We observe the following:

- μ_d converges weakly to a limiting probability measure μ supported on $[0, 1]$, with cumulative distribution function

$$\mu\{[0, c)\} = \begin{cases} c^{1-2\alpha}, & c < 1 \\ 1 & c \geq 1 \end{cases}$$

- Moreover, the risk can be expressed as

$$\mathcal{R}_d(t) = 1 - (1 \pm o_d(1)) \mathbb{E}_{X \sim \mu_d}[\mathbb{1}\{(1 \pm o_d(1))t \wedge \varphi \geq X\}]$$

By the Portmanteau theorem [Dur93], it follows that for any fixed $t \in (0, \varphi)$,

$$\mathcal{R}_d(t) \rightarrow 1 - t^{1-2\alpha}.$$

almost surely as $d \rightarrow \infty$. The almost sure convergence follows from the Borel-Cantelli lemma [Dur93] applied to the failure probabilities.

To extend this result to $t \geq \varphi$, we observe that by (GF), $\mathcal{R}_d(t)$ is non-increasing and $\inf_{t \geq 0} \mathcal{R}_d(t) \geq (1 - \varphi^{1-2\alpha})_+ - o_d(1)$. Hence, for all $t > 0$, we obtain

$$\mathcal{R}_d(t) = (1 - t^{1-2\alpha})_+ \vee (1 - \varphi^{1-2\alpha})_+$$

The desired result for a fixed $t > 0$ follows by a change of variable. Finally, since the risk curves are continuous in t , the almost sure convergence extends to all $t > 0$ pointwise.

Light-tailed case ($\alpha > 0.5$): For this part, we consider the probability space conditioned on $\mathcal{G}_{init}$ which holds with probability at least $1 - o(1/r_s^2)$. We define

$$\mathcal{Z} := \sum_{j=1}^{\infty} j^{-2\alpha}, \quad C := (r_s \mathcal{Z})^{\frac{1}{2\alpha}}.$$

We first fix a $t \in (0, (Cr_s)^\alpha) \setminus \{j^\alpha : j \in \mathbb{N}\}$. By Proposition 9, for any $d \geq \Omega_{r_s, \alpha}(1)$, we have,

$$\mathcal{R}(t \log d) = 1 - \underbrace{\frac{1}{\|\mathbf{A}\|_F^2} \sum_{i=1}^{r_s} \lambda_j^2 \mathbb{1}\{\frac{t}{C^\alpha} \geq \frac{1 \pm o_d(1)}{\lambda_j}\}}_{:= \mathcal{R}_d((Ct)^{\frac{1}{\alpha}})} + o_d(1).$$

By using $\lambda_j = j^{-\alpha}$ in the indicator function, we rewrite

$$\mathcal{R}_d(t) = 1 - \frac{1}{\|\mathbf{A}\|_F^2} \sum_{i=1}^{r_s} \lambda_j^2 \mathbb{1}\{(1 \pm o_d(1))t \geq j\}$$

We define a sequence of measures supported on $\mathbb{N}$, where $\mu_d\{j\} \propto j^{-2\alpha}$ for $j = 1, \dots, r_s$. We observe the following:

- μ_d converges weakly to a limiting probability measure μ supported on $\mathbb{N}$, such that $\mu\{j\} = \frac{j^{-2\alpha}}{\mathcal{Z}}$.
- Moreover, the risk can be expressed as

$$\mathcal{R}_d(t) = \mathbb{E}_{X \sim \mu_d}[\mathbb{1}\{(1 \pm o_d(1))t \vee r_s < X\}]$$

Since $t \notin \mathbb{N}$, we have

$$\mathbb{R}_d(t) \rightarrow \mu([t \vee r_s, \infty)).$$

By observing that $\mu([t, \infty)) \in \Theta(t^{1-2\alpha})$, the result follows for a fixed $t \in (0, (Cr_s)^\alpha) \setminus \{j^\alpha : j \in \mathbb{N}\}$. Since the limit is piecewise continuous and non increasing, it is sufficient to take a union over $t \in \{0.5, 1.5, \dots, r_s + 0.5\}$ to extend the result for all $t > 0$. $\square$

E Details of the Fine-tuning Step

In this part, we describe how to efficiently solve the empirical risk minimization problem used in the fine-tuning step of Algorithm 1. Recall that this step aims to find a rotation matrix $\Omega \in \mathbb{R}^{r_s \times r_s}$ that aligns the learned features with the teacher directions by minimizing the empirical loss over N_{Ft} fresh samples:

$$\Omega_* = \arg \min_{\Omega \in \mathbb{R}^{r_s \times r_s}} \sum_{j=1}^{N_{\text{Ft}}} \mathcal{L}(\mathbf{W}_t \Omega; (\mathbf{x}_{t+j}, y_{t+j})), \quad (\text{E.1})$$

where each sample loss is given by

$$\mathcal{L}(\mathbf{W}_t \Omega; (\mathbf{x}_{t+j}, y_{t+j})) = \frac{1}{16} \left(y_{t+j} - \frac{1}{\sqrt{r_s}} \text{Tr}(\Omega \Omega^\top \mathbf{W}_t^\top (\mathbf{x}_{t+j} \mathbf{x}_{t+j}^\top - \mathbf{I}_d) \mathbf{W}_t) \right)^2.$$

Let us define $\mathbf{A}_j := \mathbf{W}_t^\top (\mathbf{x}_{t+j} \mathbf{x}_{t+j}^\top - \mathbf{I}_d) \mathbf{W}_t$. We observe that the loss becomes quadratic in the symmetric matrix positive semidefinite matrix $\mathbf{S} := \Omega \Omega^\top$. Then, the fine-tuning objective reduces to a standard least squares problem over the cone of symmetric matrix positive semidefinite matrices:

$$\mathbf{S}_* := \arg \min_{\substack{\mathbf{S} \in \mathbb{R}^{r_s \times r_s} \\ \mathbf{S} = \mathbf{S}^\top, \mathbf{S} \succeq 0}} \frac{1}{2N_{\text{Ft}}} \sum_{j=1}^{N_{\text{Ft}}} \underbrace{\left(\sqrt{r_s} y_{t+j} - \text{Tr}(\mathbf{S} \mathbf{A}_j) \right)^2}_{:= \text{Ft}(\mathbf{S})}. \quad (\text{E.2})$$

For the following, we also define the global minimum of the least square objective in (E.2) as:

$$\mathbf{S}_{\text{glob}} := \arg \min_{\substack{\mathbf{S} \in \mathbb{R}^{r_s \times r_s} \\ \mathbf{S} = \mathbf{S}^\top}} \text{Ft}(\mathbf{S}). \quad (\text{E.3})$$

E.1 Characterizing the Minimum

Since the fine-tuning objective reduces to a least squares regression problem over symmetric matrices, we can write

$$\text{Ft}(\mathbf{S}) = \text{Ft}(\mathbf{S}_{\text{glob}}) + \text{Tr}((\mathbf{S} - \mathbf{S}_{\text{glob}}) \mathbf{L}(\mathbf{S} - \mathbf{S}_{\text{glob}}))$$

where $\mathbf{L}$ is defined as the linear operator acting on symmetric matrices via

$$\mathbf{L}(\mathbf{S}) := \frac{1}{2N_{\text{Ft}}} \sum_{j=1}^{N_{\text{Ft}}} \text{Tr}(\mathbf{S} \mathbf{A}_j) \mathbf{A}_j,$$

which corresponds to the empirical second moment operator associated with the covariates $\mathbf{A}_j$. We note that the operator $\mathbf{L}$ is self-adjoint and positive semi-definite on the space of symmetric matrices, and we can write the characterization in (E.2) equivalently

$$\mathbf{S}_* := \arg \min_{\substack{\mathbf{S} \in \mathbb{R}^{r_s \times r_s} \\ \mathbf{S} = \mathbf{S}^\top, \mathbf{S} \succeq 0}} \text{Tr}((\mathbf{S} - \mathbf{S}_{\text{glob}}) \mathbf{L}(\mathbf{S} - \mathbf{S}_{\text{glob}})).$$

We define the projection on the cone of symmetric positive semi-definite matrices as:

$$\Pi(\tilde{\mathbf{S}}) := \arg \min_{\substack{\mathbf{S} \in \mathbb{R}^{r_s \times r_s} \\ \mathbf{S} = \mathbf{S}^\top, \mathbf{S} \succeq 0}} \|\mathbf{S} - \tilde{\mathbf{S}}\|_F^2.$$

In the following, we will show that the operator $\mathbf{L}$ is close to the identity, and thus, $\mathbf{S}_*$ is close to $\Pi \circ \mathbf{L}(\mathbf{S}_{\text{glob}})$. Before proceeding, we make the following observations:

- We observe that by the first-order optimality condition applied in (E.3), we have

$$\mathbf{L}(\mathbf{S}_{\text{glob}}) = \frac{\sqrt{r_s}}{2N_{\text{Ft}}} \sum_{j=1}^{N_{\text{Ft}}} y_{t+j} \mathbf{A}_j. \quad (\text{E.4})$$

- By the generalized Pythagorean theorem [Bub14, Lemma 3.1], we have

$$\begin{aligned} \|\mathbf{S}_* - \Pi \circ \mathbf{L}(\mathbf{S}_{\text{glob}})\|_F^2 &\leq \|\mathbf{S}_* - \mathbf{L}(\mathbf{S}_{\text{glob}})\|_F^2 - \|\Pi \circ \mathbf{L}(\mathbf{S}_{\text{glob}}) - \mathbf{L}(\mathbf{S}_{\text{glob}})\|_F^2 \\ &= \text{Ft}(\mathbf{S}_*) - \text{Ft}(\Pi \circ \mathbf{L}(\mathbf{S}_{\text{glob}})) \\ &\quad - \text{Tr}((\mathbf{S}_* - \Pi \circ \mathbf{L}(\mathbf{S}_{\text{glob}}))(\mathbf{L} - \text{Id})(\mathbf{S}_* + \Pi \circ \mathbf{L}(\mathbf{S}_{\text{glob}}))), \end{aligned} \quad (\text{E.5})$$

where we use Id to denote the identity map on symmetric matrices.

E.2 Computing the Minimum

We define the approximate solution for (E.1) as:

$$\hat{\Omega} := \left(\Pi \circ \mathbf{L}(\mathbf{S}_{\text{glob}}) \right)^{\frac{1}{2}}, \quad (\text{E.6})$$

where $\mathbf{S} \rightarrow \mathbf{S}^{1/2}$ denotes the square root operator on symmetric positive semidefinite matrices. Note that the approximation in (E.6) can be computed by taking the spectral decomposition of $\mathbf{L}(\mathbf{S}_{\text{glob}})$ given in (E.4), which requires $\tilde{O}(dr_s^3)$ including the computation of $\mathbf{L}(\mathbf{S}_{\text{glob}})$. This is negligible compared to the feature learning phase, whose complexity scales as $O(Tdr_s)$. The following statement shows that $\hat{\Omega}$ is sufficiently close to the fine-tuning solution Ω^* :

Proposition 11. *Suppose $N_{\text{Ft}} \geq r_s^2 \log^5 d$. Then, with probability at least $1 - 2d^{-3}$, the final risk incurred by $\mathbf{W}_t \hat{\Omega}$ is close to that of the optimal fine-tuning solution:*

$$R(\mathbf{W}_t \hat{\Omega}) \leq R(\mathbf{W}_t \Omega_*) + \frac{1}{\log d} \leq 1 - \frac{\|\Lambda^{\frac{1}{2}} \mathbf{G}_t \Lambda^{\frac{1}{2}}\|_F^2}{\|\Lambda\|_F^2} + \frac{O(1)}{\log d}.$$

E.2.1 Proof of Proposition 11

We define the operator norm of $\mathbf{L}$ as

$$\|\mathbf{L}\|_2 = \sup_{\substack{\mathbf{S} \in \mathbb{R}^{r_s \times r_s} \\ \mathbf{S} = \mathbf{S}^\top}} \|\mathbf{L}(\mathbf{S})\|_F.$$

We consider the intersection of the following events:

- $\|\mathbf{L} - \text{Id}\|_2 \leq \frac{6}{\sqrt{\log d}}$
- $\left\| \frac{1}{2N_{\text{Ft}}} \sum_{j=1}^{N_{\text{Ft}}} y_{t+j} \mathbf{A}_j - \frac{1}{\|\Lambda\|_F} \mathbf{W}_t^\top \Theta \Lambda \Theta^\top \mathbf{W}_t \right\|_F^2 \leq \frac{1}{\log d}.$

We note that for $d \geq \Omega(1)$ the first item holds with probability $1 - d^{-3}$ by Proposition 31, where we choose $C = 5$ and $u = \log d$, and the second item holds follows with probability $1 - d^{-3}$ by Proposition 32 where we choose $C = 16$. Given the events, we have

$$\begin{aligned} R(\mathbf{W}_t \hat{\Omega}) &= \frac{1}{r_s} \left\| \Pi \circ \mathbf{L}(\mathbf{S}_{\text{glob}}) - \frac{\sqrt{r_s}}{\|\Lambda\|_F} \mathbf{W}_t^\top \Theta \Lambda \Theta^\top \mathbf{W}_t \right\|_F^2 + \left(1 - \frac{\|\Lambda^{\frac{1}{2}} \mathbf{G}_t \Lambda^{\frac{1}{2}}\|_F^2}{\|\Lambda\|_F^2} \right) \\ &\stackrel{(a)}{\leq} \frac{1}{r_s} \left\| \mathbf{L}(\mathbf{S}_{\text{glob}}) - \frac{\sqrt{r_s}}{\|\Lambda\|_F} \mathbf{W}_t^\top \Theta \Lambda \Theta^\top \mathbf{W}_t \right\|_F^2 + \left(1 - \frac{\|\Lambda^{\frac{1}{2}} \mathbf{G}_t \Lambda^{\frac{1}{2}}\|_F^2}{\|\Lambda\|_F^2} \right) \\ &\stackrel{(b)}{\leq} \frac{1}{\log d} + R(\mathbf{W}_t \Omega_*). \end{aligned}$$

where we use the convexity of the cone of symmetric positive semi-definite matrices in (a) and the second event above in (b). By using (E.5), we have

$$\|\mathbf{S}_* - \Pi \circ \mathbf{L}(\mathbf{S}_{\text{glob}})\|_F \leq \|\mathbf{L} - \text{Id}\|_2 \|\mathbf{S}_* + \Pi \circ \mathbf{L}(\mathbf{S}_{\text{glob}})\|_F.$$

Therefore,

$$\|\mathbf{S}_*\|_F \leq \frac{1 + \|\mathbf{L} - \text{Id}\|_2}{1 - \|\mathbf{L} - \text{Id}\|_2} \|\Pi \circ \mathbf{L}(\mathbf{S}_{\text{glob}})\|_F,$$

and thus,

$$\|\mathbf{S}_* - \Pi \circ \mathbf{L}(\mathbf{S}_{\text{glob}})\|_F \leq \frac{2\|\mathbf{L} - \text{Id}\|_2 \|\Pi \circ \mathbf{L}(\mathbf{S}_{\text{glob}})\|_F}{1 - \|\mathbf{L} - \text{Id}\|_2} \stackrel{(c)}{\leq} \frac{15r_s}{\sqrt{\log d}}$$

where we followed the reasoning in (a)-(b) to bound $\|\Pi \circ \mathbf{L}(\mathbf{S}_{\text{glob}})\|_F$ in (c). Therefore,

$$R(\mathbf{W}_t \Omega_*) = \frac{1}{r_s} \left\| \mathbf{S}_* - \frac{\sqrt{r_s}}{\|\Lambda\|_F} \mathbf{W}_t^\top \Theta \Lambda \Theta^\top \mathbf{W}_t \right\|_F^2 + \left(1 - \frac{\|\Lambda^{\frac{1}{2}} \mathbf{G}_t \Lambda^{\frac{1}{2}}\|_F^2}{\|\Lambda\|_F^2} \right)$$

$$\begin{aligned}
&\leq \frac{2}{r_s} \|\mathbf{S}_* - \Pi \circ \mathbf{L}(\mathbf{S}_{\text{glob}})\|_F^2 + \frac{2}{r_s} \left\| \Pi \circ \mathbf{L}(\mathbf{S}_{\text{glob}}) - \frac{\sqrt{r_s}}{\|\mathbf{\Lambda}\|_F} \mathbf{W}_t^\top \mathbf{\Theta} \mathbf{\Lambda} \mathbf{\Theta}^\top \mathbf{W}_t \right\|_F^2 \\
&\quad + \left(1 - \frac{\|\mathbf{\Lambda}^{\frac{1}{2}} \mathbf{G}_t \mathbf{\Lambda}^{\frac{1}{2}}\|_F^2}{\|\mathbf{\Lambda}\|_F^2} \right) \\
&\leq \frac{O(1)}{\log d} + \left(1 - \frac{\|\mathbf{\Lambda}^{\frac{1}{2}} \mathbf{G}_t \mathbf{\Lambda}^{\frac{1}{2}}\|_F^2}{\|\mathbf{\Lambda}\|_F^2} \right).
\end{aligned}$$

F Deferred Proofs for Online SGD

F.1 Preliminaries

We consider

$$y_{t+1} = \frac{1}{\|\mathbf{\Lambda}\|_F} \sum_{j=1}^r \lambda_j (\langle \boldsymbol{\theta}_j, \mathbf{x}_{t+1} \rangle^2 - 1) + \epsilon_{t+1} \text{ and } \hat{y}(\mathbf{W}_t; \mathbf{x}_{t+1}) = \frac{1}{\sqrt{r_s}} \sum_{j=1}^{r_s} \langle \mathbf{w}_{t,j}, \mathbf{x}_{t+1} \rangle^2 - 1,$$

where $\|\epsilon_{t+1}\|_{\psi_2} \leq \sigma$. We use $\hat{y}_{t+1} := \hat{y}(\mathbf{W}_t; \mathbf{x}_{t+1})$ and consider

- The loss function is $\mathcal{L}(\mathbf{W}_t; (\mathbf{x}_{t+1}, y_{t+1})) = \frac{1}{16} (y_{t+1} - \hat{y}_{t+1})^2$
- The Euclidean gradient is $\nabla \mathcal{L}(\mathbf{W}_t) = \frac{-1}{4\sqrt{r_s}} (y_{t+1} - \hat{y}_{t+1}) \mathbf{x}_{t+1} \mathbf{x}_{t+1}^\top \mathbf{W}_t$. Therefore, we have

$$\nabla_{\text{St}} \mathcal{L}(\mathbf{W}_t) = \frac{-1/4}{\sqrt{r_s}} (\mathbf{I}_d - \mathbf{W}_t \mathbf{W}_t^\top) (y_{t+1} - \hat{y}_{t+1}) \mathbf{x}_{t+1} \mathbf{x}_{t+1}^\top \mathbf{W}_t.$$

- We recall that $\mathbf{G}_t = \mathbf{\Theta}^\top \mathbf{W}_t \mathbf{W}_t^\top \mathbf{\Theta}$.

Then, (SGD) reads

$$\begin{aligned}
\widetilde{\mathbf{W}}_{t+1} &= \mathbf{W}_t + \frac{\eta/4}{\sqrt{r_s}} \underbrace{(\mathbf{I}_d - \mathbf{W}_t \mathbf{W}_t^\top) (y_{t+1} - \hat{y}_{t+1}) \mathbf{x}_{t+1} \mathbf{x}_{t+1}^\top \mathbf{W}_t}_{:= \nabla_{\text{St}} \mathbf{L}_{t+1}} \\
\mathbf{W}_{t+1} &= \widetilde{\mathbf{W}}_{t+1} \left(\mathbf{I}_{r_s} + \frac{\eta^2/16}{r_s} \underbrace{\nabla_{\text{St}} \mathbf{L}_{t+1}^\top \nabla_{\text{St}} \mathbf{L}_{t+1}}_{:= \mathcal{P}_{t+1}} \right)^{-1/2}. \tag{SGD}
\end{aligned}$$

We observe that

$$\begin{aligned}
\frac{\eta^2/16}{r_s} \mathcal{P}_{t+1} &= \frac{\eta^2/16}{r_s} (y_{t+1} - \hat{y}_{t+1})^2 \mathbf{W}_t^\top \mathbf{x}_{t+1} \mathbf{x}_{t+1}^\top (\mathbf{I}_d - \mathbf{W}_t \mathbf{W}_t^\top) \mathbf{x}_{t+1} \mathbf{x}_{t+1}^\top \mathbf{W}_t \\
&= \frac{\eta^2/16}{r_s} (y_{t+1} - \hat{y}_{t+1})^2 \|(\mathbf{I}_d - \mathbf{W}_t \mathbf{W}_t^\top) \mathbf{x}_{t+1}\|_2^2 \mathbf{W}_t^\top \mathbf{x}_{t+1} \mathbf{x}_{t+1}^\top \mathbf{W}_t.
\end{aligned}$$

Let

$$c_{t+1}^2 := \frac{\eta^2/16}{r_s} \|\mathcal{P}_{t+1}\|_2 = \frac{\eta^2/16}{r_s} (y_{t+1} - \hat{y}_{t+1})^2 \|(\mathbf{I}_d - \mathbf{W}_t \mathbf{W}_t^\top) \mathbf{x}_{t+1}\|_2^2 \|\mathbf{W}_t^\top \mathbf{x}_{t+1}\|_2^2.$$

We define $\mathbf{P}_{t+1} := \left(\mathbf{I}_{r_s} + \frac{\eta^2/16}{r_s} \mathcal{P}_{t+1} \right)^{-1/2}$ and since $\mathcal{P}_{t+1}$ is 1-rank, we have

$$\mathbf{P}_{t+1}^2 = \mathbf{I}_{r_s} - \frac{\eta^2/16}{r_s} \frac{\mathcal{P}_{t+1}}{1 + c_{t+1}^2}.$$

We let

$$\mathbf{M}_t := \mathbf{\Theta}^\top \mathbf{W}_t \text{ and } \hat{\mathbf{M}}_{t+1} := \mathbf{\Theta}^\top \hat{\mathbf{W}}_{t+1}.$$

We have

$$\hat{\mathbf{M}}_{t+1} = \mathbf{M}_t + \frac{\eta/4}{\sqrt{r_s}} \boldsymbol{\Theta}^\top \nabla_{\text{St}} \mathbf{L}_{t+1}.$$

By recalling that $\mathbf{G}_t = \mathbf{M}_t \mathbf{M}_t^\top$, we have

$$\begin{aligned} \mathbf{G}_{t+1} &= \hat{\mathbf{M}}_{t+1} \hat{\mathbf{M}}_{t+1}^\top + \hat{\mathbf{M}}_{t+1} (\mathbf{P}_{t+1}^2 - \mathbf{I}_{r_s}) \hat{\mathbf{M}}_{t+1}^\top \\ &= \mathbf{G}_t + \frac{\eta/4}{\sqrt{r_s}} \mathbf{M}_t \nabla_{\text{St}} \mathbf{L}_{t+1}^\top \boldsymbol{\Theta} + \frac{\eta/4}{\sqrt{r_s}} \boldsymbol{\Theta}^\top \nabla_{\text{St}} \mathbf{L}_{t+1} \mathbf{M}_t^\top + \frac{\eta^2}{16r_s} \boldsymbol{\Theta}^\top \nabla_{\text{St}} \mathbf{L}_{t+1} \nabla_{\text{St}} \mathbf{L}_{t+1}^\top \boldsymbol{\Theta} \\ &\quad - \frac{\eta^2}{16r_s} \frac{\hat{\mathbf{M}}_{t+1} \mathbf{P}_{t+1} \hat{\mathbf{M}}_{t+1}^\top}{1 + c_{t+1}^2}. \end{aligned}$$

We have

$$\nabla_{\text{St}} \mathbf{L}_{t+1} = \frac{2}{\|\boldsymbol{\Lambda}\|_F} (\mathbf{I}_d - \mathbf{W}_t \mathbf{W}_t^\top) \boldsymbol{\Theta} \boldsymbol{\Lambda} \boldsymbol{\Theta}^\top \mathbf{W}_t + (\nabla_{\text{St}} \mathbf{L}_{t+1} - \mathbb{E}_t [\nabla_{\text{St}} \mathbf{L}_{t+1}]),$$

Therefore,

$$\begin{aligned} \boldsymbol{\Theta}^\top \nabla_{\text{St}} \mathbf{L}_{t+1} \mathbf{M}_t^\top &= \frac{2}{\|\boldsymbol{\Lambda}\|_F} \boldsymbol{\Theta}^\top (\mathbf{I}_d - \mathbf{W}_t \mathbf{W}_t^\top) \boldsymbol{\Theta} \boldsymbol{\Lambda} \boldsymbol{\Theta}^\top \mathbf{W}_t \mathbf{W}_t^\top \boldsymbol{\Theta} \\ &\quad + \boldsymbol{\Theta}^\top (\nabla_{\text{St}} \mathbf{L}_{t+1} - \mathbb{E}_t [\nabla_{\text{St}} \mathbf{L}_{t+1}]) \mathbf{M}_t^\top \\ &= \frac{2}{\|\boldsymbol{\Lambda}\|_F} (\mathbf{I}_r - \mathbf{G}_t) \boldsymbol{\Lambda} \mathbf{G}_t + \boldsymbol{\Theta}^\top (\nabla_{\text{St}} \mathbf{L}_{t+1} - \mathbb{E}_t [\nabla_{\text{St}} \mathbf{L}_{t+1}]) \mathbf{M}_t^\top. \end{aligned}$$

Hence, we have

$$\begin{aligned} \mathbf{G}_{t+1} &= \mathbf{G}_t + \frac{\eta/2}{\|\boldsymbol{\Lambda}\|_F \sqrt{r_s}} (\boldsymbol{\Lambda} \mathbf{G}_t + \mathbf{G}_t \boldsymbol{\Lambda} - 2\mathbf{G}_t \boldsymbol{\Lambda} \mathbf{G}_t) \\ &\quad + \frac{\eta/2}{\sqrt{r_s}} \text{Sym} (\boldsymbol{\Theta}^\top (\nabla_{\text{St}} \mathbf{L}_{t+1} - \mathbb{E}_t [\nabla_{\text{St}} \mathbf{L}_{t+1}]) \mathbf{M}_t^\top) \\ &\quad + \frac{\eta^2}{16r_s} \boldsymbol{\Theta}^\top \nabla_{\text{St}} \mathbf{L}_{t+1} \nabla_{\text{St}} \mathbf{L}_{t+1}^\top \boldsymbol{\Theta} - \frac{\eta^2}{16r_s} \frac{\hat{\mathbf{M}}_{t+1} \mathbf{P}_{t+1} \hat{\mathbf{M}}_{t+1}^\top}{1 + c_{t+1}^2} \end{aligned}$$

On the other hand,

$$\begin{aligned} \hat{\mathbf{M}}_{t+1} \mathbf{P}_{t+1} \hat{\mathbf{M}}_{t+1}^\top &= \left(\mathbf{M}_t + \frac{\eta/4}{\sqrt{r_s}} \boldsymbol{\Theta}^\top \nabla_{\text{St}} \mathbf{L}_{t+1} \right) \mathbf{P}_{t+1} \left(\mathbf{M}_t + \frac{\eta/4}{\sqrt{r_s}} \boldsymbol{\Theta}^\top \nabla_{\text{St}} \mathbf{L}_{t+1} \right)^\top \\ &= \mathbf{M}_t \mathbf{P}_{t+1} \mathbf{M}_t^\top + \frac{\eta/2}{\sqrt{r_s}} \text{Sym} (\boldsymbol{\Theta}^\top \nabla_{\text{St}} \mathbf{L}_{t+1} \mathbf{P}_{t+1} \mathbf{M}_t^\top) + \frac{\eta^2}{16r_s} \boldsymbol{\Theta}^\top \nabla_{\text{St}} \mathbf{L}_{t+1} \mathbf{P}_{t+1} \nabla_{\text{St}} \mathbf{L}_{t+1}^\top \boldsymbol{\Theta}. \end{aligned}$$

We collect the higher order terms in a single term defined as follows:

$$\begin{aligned} R_{\text{so}}[\mathbf{G}_t] &:= \frac{\eta^2}{16r_s} \boldsymbol{\Theta}^\top \mathbb{E}_t [\nabla_{\text{St}} \mathbf{L}_{t+1} \nabla_{\text{St}} \mathbf{L}_{t+1}^\top] \boldsymbol{\Theta} - \frac{\eta^2}{16r_s} \mathbf{M}_t \mathbb{E}_t \left[\frac{\mathbf{P}_{t+1}}{1 + c_{t+1}^2} \right] \mathbf{M}_t^\top \\ &\quad - \frac{\eta^3}{32r_s^{3/2}} \text{Sym} \left(\boldsymbol{\Theta}^\top \mathbb{E}_t \left[\frac{\nabla_{\text{St}} \mathbf{L}_{t+1} \mathbf{P}_{t+1}}{1 + c_{t+1}^2} \right] \mathbf{M}_t^\top \right) \\ &\quad - \frac{\eta^4}{256r_s^2} \boldsymbol{\Theta}^\top \mathbb{E}_t \left[\frac{\nabla_{\text{St}} \mathbf{L}_{t+1} \mathbf{P}_{t+1} \nabla_{\text{St}} \mathbf{L}_{t+1}^\top}{1 + c_{t+1}^2} \right] \boldsymbol{\Theta}. \end{aligned}$$

We collect the noise terms in a single term defined as follows:

$$\begin{aligned} \frac{\eta/2}{\sqrt{r_s}} \boldsymbol{\nu}_{t+1} &:= \frac{\eta/2}{\sqrt{r_s}} \text{Sym} \left(\boldsymbol{\Theta}^\top (\nabla_{\text{St}} \mathbf{L}_{t+1} - \mathbb{E}_t [\nabla_{\text{St}} \mathbf{L}_{t+1}]) \mathbf{M}_t^\top \right) \\ &\quad - \frac{\eta^2}{16r_s} \mathbf{M}_t \left(\frac{\mathbf{P}_{t+1}}{1 + c_{t+1}^2} - \mathbb{E}_t \left[\frac{\mathbf{P}_{t+1}}{1 + c_{t+1}^2} \right] \right) \mathbf{M}_t^\top \end{aligned}$$

$$\begin{aligned}
& + \frac{\eta^2}{16r_s} \Theta^\top (\nabla_{\text{St}} \mathbf{L}_{t+1} \nabla_{\text{St}} \mathbf{L}_{t+1}^\top - \mathbb{E}_t [\nabla_{\text{St}} \mathbf{L}_{t+1} \nabla_{\text{St}} \mathbf{L}_{t+1}^\top]) \Theta \\
& - \frac{\eta^3}{32r_s^{3/2}} \text{Sym} \left(\Theta^\top \left(\frac{\nabla_{\text{St}} \mathbf{L}_{t+1} \mathbf{P}_{t+1}}{1 + c_{t+1}^2} - \mathbb{E}_t \left[\frac{\nabla_{\text{St}} \mathbf{L}_{t+1} \mathbf{P}_{t+1}}{1 + c_{t+1}^2} \right] \right) \mathbf{M}_t^\top \right) \\
& - \frac{\eta^4}{256r_s^2} \Theta^\top \left(\frac{\nabla_{\text{St}} \mathbf{L}_{t+1} \mathbf{P}_{t+1} \nabla_{\text{St}} \mathbf{L}_{t+1}^\top}{1 + c_{t+1}^2} - \mathbb{E}_t \left[\frac{\nabla_{\text{St}} \mathbf{L}_{t+1} \mathbf{P}_{t+1} \nabla_{\text{St}} \mathbf{L}_{t+1}^\top}{1 + c_{t+1}^2} \right] \right) \Theta.
\end{aligned}$$

With these definitions in hand, we have

$$\mathbf{G}_{t+1} = \mathbf{G}_t + \frac{\eta/2}{\|\mathbf{\Lambda}\|_{\text{F}} \sqrt{r_s}} (\mathbf{\Lambda} \mathbf{G}_t + \mathbf{G}_t \mathbf{\Lambda} - 2\mathbf{G}_t \mathbf{\Lambda} \mathbf{G}_t) + R_{\text{so}}[\mathbf{G}_t] + \frac{\eta/2}{\sqrt{r_s}} \boldsymbol{\nu}_{t+1}.$$

F.2 Including second-order terms and monotone bounds

For $C > 1$, we define

$$\mathbf{\Lambda}_{\ell_1} := \mathbf{\Lambda} - C \|\mathbf{\Lambda}\|_{\text{F}} \frac{\eta d}{\sqrt{r_s}} \mathbf{I}_r \text{ and } \mathbf{\Lambda}_{u_1} := \mathbf{\Lambda} + C \|\mathbf{\Lambda}\|_{\text{F}} \frac{\eta d}{\sqrt{r_s}} \mathbf{I}_r. \quad (\text{F.2})$$

We recall the definition of effective learning rate $\eta = \frac{\eta/2}{\|\mathbf{\Lambda}\|_{\text{F}} \sqrt{r_s}}$. By Proposition 17, we have

$$\mathbf{G}_{t+1} \succeq \mathbf{G}_t + \eta (\mathbf{\Lambda}_{\ell_1} \mathbf{G}_t + \mathbf{G}_t \mathbf{\Lambda}_{\ell_1} - 2\mathbf{G}_t \mathbf{\Lambda} \mathbf{G}_t) - \frac{C}{2} \eta^2 \|\mathbf{\Lambda}\|_{\text{F}}^2 r_s \mathbf{I}_r + \frac{\eta/2}{\sqrt{r_s}} \boldsymbol{\nu}_{t+1}, \quad (\text{F.3})$$

$$\mathbf{G}_{t+1} \preceq \mathbf{G}_t + \eta (\mathbf{\Lambda}_{u_1} \mathbf{G}_t + \mathbf{G}_t \mathbf{\Lambda}_{u_1} - 2\mathbf{G}_t \mathbf{\Lambda} \mathbf{G}_t) + \frac{C}{2} \eta^2 \|\mathbf{\Lambda}\|_{\text{F}}^2 r_s \mathbf{I}_r + \frac{\eta/2}{\sqrt{r_s}} \boldsymbol{\nu}_{t+1}. \quad (\text{F.4})$$

F.2.1 Heavy tailed case - $\alpha \in [0, 0.5)$

Proposition 12. We consider $\alpha \in [0, 0.5)$, $\frac{r_s}{r} \rightarrow (0, \infty]$ and $\eta \ll \frac{1}{d \log^4 d} \sqrt{\frac{r_s}{r}}$. We define

$$\mathbf{V}_t^- := 2\mathbf{\Lambda}^{\frac{1}{2}} \mathbf{G}_t \mathbf{\Lambda}^{\frac{1}{2}} - \mathbf{\Lambda}_{\ell_1} \text{ and } \mathbf{V}_t^+ := 2\mathbf{\Lambda}^{\frac{1}{2}} \mathbf{G}_t \mathbf{\Lambda}^{\frac{1}{2}} - \mathbf{\Lambda}_{u_1}.$$

For $d \geq \Omega(1)$, we have

$$\mathbf{\Lambda} + \frac{0.1r^{-\alpha}}{\log^4 d} \mathbf{I}_r \succ \mathbf{\Lambda}_{u_1} \succ \mathbf{\Lambda} \succ \mathbf{\Lambda}_{\ell_1} \succ \mathbf{\Lambda} - \frac{0.1r^{-\alpha}}{\log^4 d} \mathbf{I}_r \quad (\text{F.5})$$

and

$$\begin{aligned}
\mathbf{V}_{t+1}^- & \succeq \mathbf{V}_t^- \left(\mathbf{I}_r + \frac{\eta}{1 - 1.1\eta} \mathbf{V}_t^- \right)^{-1} + \eta \mathbf{\Lambda}_{\ell_1}^2 - C\eta^2 \|\mathbf{\Lambda}\|_{\text{F}}^2 r_s \mathbf{\Lambda} + \frac{\eta}{\sqrt{r_s}} \mathbf{\Lambda}^{\frac{1}{2}} \boldsymbol{\nu}_{t+1} \mathbf{\Lambda}^{\frac{1}{2}} \\
\mathbf{V}_{t+1}^+ & \preceq \mathbf{V}_t^+ \left(\mathbf{I}_r + \frac{\eta}{1 + 1.1\eta} \mathbf{V}_t^+ \right)^{-1} + \eta \mathbf{\Lambda}_{u_1}^2 + C\eta^2 \|\mathbf{\Lambda}\|_{\text{F}}^2 r_s \mathbf{\Lambda} + \frac{\eta}{\sqrt{r_s}} \mathbf{\Lambda}^{\frac{1}{2}} \boldsymbol{\nu}_{t+1} \mathbf{\Lambda}^{\frac{1}{2}}
\end{aligned} \quad (\text{F.6})$$

where the bounding iterations are monotone in the sense defined in Proposition 4.

Proof. We first note that since $\|\mathbf{\Lambda}\|_{\text{F}} \asymp r^{\frac{1}{2}-\alpha}$ for $\alpha \in [0, 0.5)$, we have

$$\|\mathbf{\Lambda}\|_{\text{F}} \frac{\eta d}{\sqrt{r_s}} \ll \frac{r^{-\alpha}}{\log^4 d}.$$

Therefore, (F.5) holds for $d \geq \Omega(1)$, which implies

$$\|\mathbf{V}_t^-\|_2 \vee \|\mathbf{V}_t^+\|_2 \leq 1 + \frac{0.1r^{-\alpha}}{\log^4 d}, \text{ for all } t \in \mathbb{N}. \quad (\text{F.7})$$

Therefore, the monotonicity follows from Proposition 4.

For the remaining part, we introduce the following notation, $\mathbf{K}_t := \mathbf{\Lambda}^{\frac{1}{2}} \mathbf{G}_t \mathbf{\Lambda}^{\frac{1}{2}}$. For the lower bound, by (F.3), we have

$$\mathbf{K}_{t+1} \succeq \mathbf{K}_t + \frac{\eta}{2} (\mathbf{\Lambda}_{\ell_1}^2 - (2\mathbf{K}_t - \mathbf{\Lambda}_{\ell_1})^2) - \frac{C}{2} \eta^2 \|\mathbf{\Lambda}\|_{\text{F}}^2 r_s \mathbf{\Lambda} + \frac{\eta/2}{\sqrt{r_s}} \mathbf{\Lambda}^{\frac{1}{2}} \boldsymbol{\nu}_{t+1} \mathbf{\Lambda}^{\frac{1}{2}}.$$

By multiplying both sides with 2 and subtracting Λ_{ℓ_1} from both sides, we have

$$\begin{aligned} V_{t+1}^- &\succeq V_t^- - \eta(V_t^-)^2 + \eta\Lambda_{\ell_1}^2 - C\eta^2\|\Lambda\|_{\mathbb{F}}^2 r_s \Lambda + \frac{\eta}{\sqrt{r_s}} \Lambda^{\frac{1}{2}} \nu_{t+1} \Lambda^{\frac{1}{2}} \\ &\stackrel{(a)}{\succeq} V_t^- - \frac{\eta}{1-1.1\eta} (V_t^-)^2 \left(I_r + \frac{\eta}{1-1.1\eta} V_t^- \right)^{-1} + \eta\Lambda_{\ell_1}^2 \\ &\quad - C\eta^2\|\Lambda\|_{\mathbb{F}}^2 r_s \Lambda + \frac{\eta}{\sqrt{r_s}} \Lambda^{\frac{1}{2}} \nu_{t+1} \Lambda^{\frac{1}{2}} \\ &= V_t^- \left(I_r + \frac{\eta}{1-1.1\eta} V_t^- \right)^{-1} + \eta\Lambda_{\ell_1}^2 - C\eta^2\|\Lambda\|_{\mathbb{F}}^2 r_s \Lambda + \frac{\eta}{\sqrt{r_s}} \Lambda^{\frac{1}{2}} \nu_{t+1} \Lambda^{\frac{1}{2}}, \end{aligned}$$

where we used (F.7) for (a).

For the upper bound, by (F.4), we have

$$K_{t+1} \preceq K_t + \frac{\eta}{2} (\Lambda_{u_1}^2 - (2K_t - \Lambda_{u_1})^2) + \frac{C}{2} \eta^2 \|\Lambda\|_{\mathbb{F}}^2 r_s \Lambda + \frac{\eta/2}{\sqrt{r_s}} \Lambda^{\frac{1}{2}} \nu_{t+1} \Lambda^{\frac{1}{2}}.$$

By multiplying both sides with 2 and subtracting Λ_{u_1} from both sides, we get

$$\begin{aligned} V_{t+1}^+ &\preceq V_t^+ - \eta(V_t^+)^2 + \eta\Lambda_{u_1}^2 + C\eta^2\|\Lambda\|_{\mathbb{F}}^2 r_s \Lambda + \frac{\eta}{\sqrt{r_s}} \Lambda^{\frac{1}{2}} \nu_{t+1} \Lambda^{\frac{1}{2}} \\ &\stackrel{(b)}{\preceq} V_t^+ - \frac{\eta}{1+1.1\eta} (V_t^+)^2 \left(I_r + \frac{\eta}{1+1.1\eta} V_t^+ \right)^{-1} + \eta\Lambda_{u_1}^2 \\ &\quad + C\eta^2\|\Lambda\|_{\mathbb{F}}^2 r_s \Lambda + \frac{\eta}{\sqrt{r_s}} \Lambda^{\frac{1}{2}} \nu_{t+1} \Lambda^{\frac{1}{2}} \\ &= V_t^+ \left(I_r + \frac{\eta}{1+1.1\eta} V_t^+ \right)^{-1} + \eta\Lambda_{u_1}^2 + C\eta^2\|\Lambda\|_{\mathbb{F}}^2 r_s \Lambda + \frac{\eta}{\sqrt{r_s}} \Lambda^{\frac{1}{2}} \nu_{t+1} \Lambda^{\frac{1}{2}}, \end{aligned}$$

where we used (F.7) for (b). $\square$

F.2.2 Light tailed case - $\alpha > 0.5$

We introduce the submatrix notation

$$G_t =: \begin{bmatrix} G_{t,11} & G_{t,12} \\ G_{t,12}^\top & G_{t,22} \end{bmatrix} \quad \nu_t =: \begin{bmatrix} \nu_{t,11} & \nu_{t,12} \\ \nu_{t,12}^\top & \nu_{t,22} \end{bmatrix} \quad \Lambda =: \begin{bmatrix} \Lambda_{11} & 0 \\ 0 & \Lambda_{22} \end{bmatrix} \quad \Lambda_{\ell_1} =: \begin{bmatrix} \Lambda_{\ell_1,11} & 0 \\ 0 & \Lambda_{\ell_1,22} \end{bmatrix},$$

where $G_{t,11}, \nu_{t,11}, \Lambda_{11}, \Lambda_{\ell_1,11} \in \mathbb{R}^{r_u \times r_u}$ for $r_u < r$. Similarly, we define the block matrices of Λ_{u_1} as $\Lambda_{u_1,11} \in \mathbb{R}^{r_u \times r_u}$ and $\Lambda_{u_1,22}$. We can write iterations (F.3) and (F.4) for the left-top submatrix as:

$$\begin{aligned} G_{t+1,11} &\succeq G_{t,11} + \eta \left(\Lambda_{\ell_1,11} G_{t,11} + G_{t,11} \Lambda_{\ell_1,11} - 2G_{t,11} \Lambda_{11} G_{t,11} - 2G_{t,12} \Lambda_{22} G_{t,12}^\top \right) \\ &\quad - \frac{C}{2} \eta^2 \|\Lambda\|_{\mathbb{F}}^2 r_s I_{r_u} + \frac{\eta/2}{\sqrt{r_s}} \nu_{t+1,11} \end{aligned} \quad (\text{F.8})$$

$$\begin{aligned} G_{t+1,11} &\preceq G_{t,11} + \eta \left(\Lambda_{u_1,11} G_{t,11} + G_{t,11} \Lambda_{u_1,11} - 2G_{t,11} \Lambda_{11} G_{t,11} - 2G_{t,12} \Lambda_{22} G_{t,12}^\top \right) \\ &\quad + \frac{C}{2} \eta^2 \|\Lambda\|_{\mathbb{F}}^2 r_s I_{r_u} + \frac{\eta/2}{\sqrt{r_s}} \nu_{t+1,11}. \end{aligned} \quad (\text{F.9})$$

The following statement is analogous to Proposition 12 in the case $\alpha > 0.5$.

Proposition 13. We consider $\alpha > 0.5$, $r_s \asymp 1$, and

$$\eta \ll \frac{1}{d \log^3 d} \frac{1}{r_u^{2+\alpha}} \quad \text{and} \quad r_u = \lceil \log^{2.5} d \rceil.$$

We define $V_t^+ := 2\Lambda_{11}^{\frac{1}{2}} G_{t,11} \Lambda_{11}^{\frac{1}{2}} - \Lambda_{u_1,11}$ and

$$V_t^- := 2 \left(\Lambda_{11} - \frac{1}{(r_u+1)^\alpha} I_{r_u} \right)^{\frac{1}{2}} G_{t,11} \left(\Lambda_{11} - \frac{1}{(r_u+1)^\alpha} I_{r_u} \right)^{\frac{1}{2}} - \left(\Lambda_{\ell_1,11} - \frac{1}{(r_u+1)^\alpha} I_{r_u} \right).$$

For $d \geq \Omega(1)$, we have

$$\mathbf{\Lambda}_{11} + \frac{0.1}{r_u^{2+\alpha} \log^3 d} \mathbf{I}_{r_u} \succ \mathbf{\Lambda}_{u_1,11} \succ \mathbf{\Lambda}_{11} \succ \mathbf{\Lambda}_{\ell_1,11} \succ \mathbf{\Lambda}_{11} - \frac{0.1}{r_u^{2+\alpha} \log^3 d} \mathbf{I}_{r_u} \quad (\text{F.10})$$

and

$$\begin{aligned} \mathbf{V}_{t+1}^- &\succeq \mathbf{V}_t^- \left(\mathbf{I}_{r_u} + \frac{\eta}{1-1.1\eta} \mathbf{V}_t^- \right)^{-1} + \eta \left(\mathbf{\Lambda}_{\ell_1,11} - \frac{1}{(r_u+1)^\alpha} \mathbf{I}_{r_u} \right)^2 - C\eta^2 \|\mathbf{\Lambda}\|_{\text{F}}^2 r_s \left(\mathbf{\Lambda}_{11} - \frac{1}{(r_u+1)^\alpha} \mathbf{I}_{r_u} \right) \\ &\quad + \frac{\eta}{\sqrt{r_s}} \left(\mathbf{\Lambda}_{11} - \frac{1}{(r_u+1)^\alpha} \mathbf{I}_{r_u} \right)^{\frac{1}{2}} \boldsymbol{\nu}_{t+1,11} \left(\mathbf{\Lambda}_{11} - \frac{1}{(r_u+1)^\alpha} \mathbf{I}_{r_u} \right)^{\frac{1}{2}} \\ \mathbf{V}_{t+1}^+ &\preceq \mathbf{V}_t^+ \left(\mathbf{I}_{r_u} + \frac{\eta}{1+1.1\eta} \mathbf{V}_t^+ \right)^{-1} + \eta \mathbf{\Lambda}_{u_1,11}^2 + C\eta^2 \|\mathbf{\Lambda}\|_{\text{F}}^2 r_s \mathbf{\Lambda}_{11} + \frac{\eta}{\sqrt{r_s}} \mathbf{\Lambda}_{11}^{\frac{1}{2}} \boldsymbol{\nu}_{t+1,11} \mathbf{\Lambda}_{11}^{\frac{1}{2}}. \end{aligned}$$

where the bounding iterations are monotone in the sense defined in Proposition 4.

Proof. We first note that since $r_s \asymp 1$ and $\|\mathbf{\Lambda}\|_{\text{F}} \asymp 1$ for $\alpha > 0.5$, we have

$$\|\mathbf{\Lambda}\|_{\text{F}} \frac{\eta d}{\sqrt{r_s}} \ll \frac{1}{r_u^{2+\alpha} \log^3 d}.$$

Therefore, (F.10) holds for $d \geq \Omega(1)$, which implies

$$\|\mathbf{V}_t^-\|_2 \vee \|\mathbf{V}_t^+\|_2 \leq 1 + \frac{0.1}{r_u^{2+\alpha} \log^3 d} + \frac{1}{(r_u+1)^\alpha}, \quad \text{for all } t \in \mathbb{N}. \quad (\text{F.11})$$

For the remaining part, we introduce the following notation,

$$\mathbf{K}_t^- := \left(\mathbf{\Lambda}_{11} - \frac{1}{(r_u+1)^\alpha} \mathbf{I}_{r_u} \right)^{\frac{1}{2}} \mathbf{G}_{t,11} \left(\mathbf{\Lambda}_{11} - \frac{1}{(r_u+1)^\alpha} \mathbf{I}_{r_u} \right)^{\frac{1}{2}} \quad \text{and} \quad \mathbf{K}_t^+ := \mathbf{\Lambda}_{11}^{\frac{1}{2}} \mathbf{G}_{t,11} \mathbf{\Lambda}_{11}^{\frac{1}{2}}.$$

For the upper bound, since $\mathbf{\Lambda}_{22} \succ 0$, by (F.9) we have

$$\mathbf{K}_{t+1}^+ \preceq \mathbf{K}_t^+ + \frac{\eta}{2} (\mathbf{\Lambda}_{u_1,11}^2 - (2\mathbf{K}_t^+ - \mathbf{\Lambda}_{u_1,11})^2) + \frac{C}{2} \eta^2 \|\mathbf{\Lambda}\|_{\text{F}}^2 r_s \mathbf{\Lambda}_{11} + \frac{\eta/2}{\sqrt{r_s}} \mathbf{\Lambda}_{11}^{\frac{1}{2}} \boldsymbol{\nu}_{t+1,11} \mathbf{\Lambda}_{11}^{\frac{1}{2}}.$$

By multiplying both sides with 2 and subtracting $\mathbf{\Lambda}_{u_1,11}$ from both sides, we get

$$\begin{aligned} \mathbf{V}_{t+1}^+ &\preceq \mathbf{V}_t^+ - \eta(\mathbf{V}_t^+)^2 + \eta \mathbf{\Lambda}_{u_1,11}^2 + C\eta^2 \|\mathbf{\Lambda}\|_{\text{F}}^2 r_s \mathbf{\Lambda}_{11} + \frac{\eta}{\sqrt{r_s}} \mathbf{\Lambda}_{11}^{\frac{1}{2}} \boldsymbol{\nu}_{t+1,11} \mathbf{\Lambda}_{11}^{\frac{1}{2}} \\ &\stackrel{(a)}{\preceq} \mathbf{V}_t^+ - \frac{\eta}{1+1.1\eta} (\mathbf{V}_t^+)^2 \left(\mathbf{I}_{r_u} + \frac{\eta}{1+1.1\eta} \mathbf{V}_t^+ \right)^{-1} + \eta \mathbf{\Lambda}_{u_1,11}^2 \\ &\quad + C\eta^2 \|\mathbf{\Lambda}\|_{\text{F}}^2 r_s \mathbf{\Lambda}_{11} + \frac{\eta}{\sqrt{r_s}} \mathbf{\Lambda}_{11}^{\frac{1}{2}} \boldsymbol{\nu}_{t+1,11} \mathbf{\Lambda}_{11}^{\frac{1}{2}} \\ &= \mathbf{V}_t^+ \left(\mathbf{I}_{r_u} + \frac{\eta}{1+1.1\eta} \mathbf{V}_t^+ \right)^{-1} + \eta \mathbf{\Lambda}_{u_1,11}^2 + C\eta^2 \|\mathbf{\Lambda}\|_{\text{F}}^2 r_s \mathbf{\Lambda}_{11} + \frac{\eta}{\sqrt{r_s}} \mathbf{\Lambda}_{11}^{\frac{1}{2}} \boldsymbol{\nu}_{t+1,11} \mathbf{\Lambda}_{11}^{\frac{1}{2}}, \end{aligned}$$

where we used (F.11) for (a).

For the lower bound, we first observe that $\mathbf{G}_{t,11}(\mathbf{I}_{r_u} - \mathbf{G}_{t,11}) - \mathbf{G}_{t,12} \mathbf{G}_{t,12}^\top \succeq 0$ since it corresponds to the left-top submatrix of $\mathbf{G}_t(\mathbf{I}_r - \mathbf{G}_t)$. Therefore, by (F.8)

$$\begin{aligned} \mathbf{G}_{t+1,11} &\succeq \mathbf{G}_{t,11} + \eta \left(\mathbf{\Lambda}_{\ell_1,11} \mathbf{G}_{t,11} + \mathbf{G}_{t,11} \mathbf{\Lambda}_{\ell_1,11} - 2\mathbf{G}_{t,11} \mathbf{\Lambda}_{11} \mathbf{G}_{t,11} - 2\mathbf{G}_{t,12} \mathbf{\Lambda}_{22} \mathbf{G}_{t,12}^\top \right) \\ &\quad - \frac{C}{2} \eta^2 \|\mathbf{\Lambda}\|_{\text{F}}^2 r_s \mathbf{I}_{r_u} + \frac{\eta/2}{\sqrt{r_s}} \boldsymbol{\nu}_{t+1,11} - \frac{2\eta}{(r_u+1)^\alpha} (\mathbf{G}_{t,11}(\mathbf{I}_{r_u} - \mathbf{G}_{t,11}) - \mathbf{G}_{t,12} \mathbf{G}_{t,12}^\top) \\ &\stackrel{(b)}{\succeq} \mathbf{G}_{t,11} + \eta \left(\left(\mathbf{\Lambda}_{\ell_1,11} - \frac{1}{(r_u+1)^\alpha} \mathbf{I}_{r_u} \right) \mathbf{G}_{t,11} + \mathbf{G}_{t,11} \left(\mathbf{\Lambda}_{\ell_1,11} - \frac{1}{(r_u+1)^\alpha} \mathbf{I}_{r_u} \right) \right) \\ &\quad - 2\eta \mathbf{G}_{t,11} \left(\mathbf{\Lambda}_{11} - \frac{1}{(r_u+1)^\alpha} \mathbf{I}_{r_u} \right) \mathbf{G}_{t,11} - \frac{C}{2} \eta^2 \|\mathbf{\Lambda}\|_{\text{F}}^2 r_s \mathbf{I}_{r_u} + \frac{\eta/2}{\sqrt{r_s}} \boldsymbol{\nu}_{t+1,11}, \end{aligned}$$

where (b) follows by $\Lambda_{22} \preceq \frac{1}{(r_u+1)^\alpha} \mathbf{I}_{r-r_u}$. Therefore, we have

$$\begin{aligned} \mathbf{K}_{t+1}^- &\succeq \mathbf{K}_t^- + \frac{\eta}{2} \left((\Lambda_{\ell_1,11} - \frac{1}{(r_u+1)^\alpha} \mathbf{I}_{r_u})^2 - (2\mathbf{K}_t^- - (\Lambda_{\ell_1,11} - \frac{1}{(r_u+1)^\alpha} \mathbf{I}_{r_u}))^2 \right) \\ &\quad - \frac{C}{2} \eta^2 \|\Lambda\|_{\mathbb{F}}^2 r_s (\Lambda_{11} - \frac{1}{(r_u+1)^\alpha} \mathbf{I}_{r_u}) \\ &\quad + \frac{\eta/2}{\sqrt{r_s}} (\Lambda_{11} - \frac{1}{(r_u+1)^\alpha} \mathbf{I}_{r_u})^{\frac{1}{2}} \nu_{t+1,11} (\Lambda_{11} - \frac{1}{(r_u+1)^\alpha} \mathbf{I}_{r_u})^{\frac{1}{2}}. \end{aligned}$$

By multiplying both sides with 2 and subtracting $\Lambda_{\ell_1,11}$ from both sides, we get

$$\begin{aligned} \mathbf{V}_{t+1}^- &\succeq \mathbf{V}_t^- - \eta(\mathbf{V}_t^-)^2 + \eta(\Lambda_{\ell_1,11} - \frac{1}{(r_u+1)^\alpha} \mathbf{I}_{r_u})^2 - C\eta^2 \|\Lambda\|_{\mathbb{F}}^2 r_s (\Lambda_{11} - \frac{1}{(r_u+1)^\alpha} \mathbf{I}_{r_u}) \\ &\quad + \frac{\eta}{\sqrt{r_s}} (\Lambda_{11} - \frac{1}{(r_u+1)^\alpha} \mathbf{I}_{r_u})^{\frac{1}{2}} \nu_{t+1,11} (\Lambda_{11} - \frac{1}{(r_u+1)^\alpha} \mathbf{I}_{r_u})^{\frac{1}{2}} \\ &\stackrel{(c)}{\succeq} \mathbf{V}_t^- - \frac{\eta}{1-1.1\eta} (\mathbf{V}_t^-)^2 \left(\mathbf{I}_{r_u} + \frac{\eta}{1-1.1\eta} \mathbf{V}_t^- \right)^{-1} + \eta(\Lambda_{\ell_1,11} - \frac{1}{(r_u+1)^\alpha} \mathbf{I}_{r_u})^2 \\ &\quad - C\eta^2 \|\Lambda\|_{\mathbb{F}}^2 r_s (\Lambda_{11} - \frac{1}{(r_u+1)^\alpha} \mathbf{I}_{r_u}) \\ &\quad + \frac{\eta}{\sqrt{r_s}} (\Lambda_{11} - \frac{1}{(r_u+1)^\alpha} \mathbf{I}_{r_u})^{\frac{1}{2}} \nu_{t+1,11} (\Lambda_{11} - \frac{1}{(r_u+1)^\alpha} \mathbf{I}_{r_u})^{\frac{1}{2}} \\ &= \mathbf{V}_t^- \left(\mathbf{I}_{r_u} + \frac{\eta}{1-1.1\eta} \mathbf{V}_t^- \right)^{-1} + \eta(\Lambda_{\ell_1,11} - \frac{1}{(r_u+1)^\alpha} \mathbf{I}_{r_u})^2 \\ &\quad - C\eta^2 \|\Lambda\|_{\mathbb{F}}^2 r_s (\Lambda_{11} - \frac{1}{(r_u+1)^\alpha} \mathbf{I}_{r_u}) \\ &\quad + \frac{\eta}{\sqrt{r_s}} (\Lambda_{11} - \frac{1}{(r_u+1)^\alpha} \mathbf{I}_{r_u})^{\frac{1}{2}} \nu_{t+1,11} (\Lambda_{11} - \frac{1}{(r_u+1)^\alpha} \mathbf{I}_{r_u})^{\frac{1}{2}}, \end{aligned}$$

where we used (F.11) for (c). The monotonicity of the update follows the same argument in the heavy-tailed case. $\square$

F.3 Definitions and bounding systems

To avoid repetition in the derivations, we introduce the following unified notation:

$$\mathbf{rk} \in \{r, r_u\}, \quad \mathbf{G}_t \in \{\mathbf{G}_t, \mathbf{G}_{t,11}\}, \quad \nu_t \in \{\nu_t, \nu_{t,11}\}.$$

where each variable will take its first value in the heavy-tailed case and its second value in the light-tailed case. To avoid repetition in the following sections, we make the following simplifications by slight abuse of notation:

$$\Lambda_{\ell_1} \leftarrow \begin{cases} \Lambda_{\ell_1}, & \alpha \in [0, 0.5) \\ \Lambda_{\ell_1,11} - \frac{1}{(r_u+1)^\alpha} \mathbf{I}_{r_u}, & \alpha > 0.5 \end{cases} \quad \text{and} \quad \Lambda_{\ell_2} \leftarrow \begin{cases} \Lambda, & \alpha \in [0, 0.5) \\ \Lambda_{11} - \frac{1}{(r_u+1)^\alpha} \mathbf{I}_{r_u}, & \alpha > 0.5 \end{cases}$$

and

$$\Lambda_{u_1} \leftarrow \begin{cases} \Lambda_{u_1}, & \alpha \in [0, 0.5) \\ \Lambda_{u_1,11}, & \alpha > 0.5 \end{cases} \quad \text{and} \quad \Lambda_{u_2} \leftarrow \begin{cases} \Lambda, & \alpha \in [0, 0.5) \\ \Lambda_{11}, & \alpha > 0.5. \end{cases}$$

The dimension of each block is $r_u < r$ for $\alpha > 0.5$ and r for $\alpha \in [0, 0.5)$, from which readers can distinguish the light tailed case from the heavy tailed case. Throughout the proof, we will also use constants $\kappa_d \in o_d(1)$ and $\tilde{C} \in O(1)$ that will be specified later. Moreover, we make the following definitions:

- **Noise sequence.** For $\nu_0 = 0$, we define the noise sequence $\underline{\nu}_{t+1} := \underline{\nu}_t + \frac{\eta/2}{\sqrt{r_s}} \nu_{t+1}$.
- **Reference sequence.** For $\mathbf{T}_0 = \frac{\kappa_d r_s}{d} \mathbf{I}_{\mathbf{rk}}$, we define the reference sequence

$$\mathbf{T}_{t+1} = \mathbf{T}_t + 2(1 - 2\kappa_d)\eta \left(\Lambda_{\ell_1} \mathbf{T}_t - \frac{3\kappa_d + 1}{\kappa_d(1 - 2\kappa_d)} \Lambda_{u_2} \mathbf{T}_t^2 \right).$$

• **Bounding systems.** We define the lower and upper bounding recursions as

$$\underline{V}_{t+1} = \underline{V}_t \left(I_{rk} + \frac{\eta(1+2\kappa_d)}{1-1.2\eta} \underline{V}_t \right)^{-1} + \frac{\eta(1+2\kappa_d)}{1-1.2\eta} \left(\frac{\Lambda_{\ell_1}^2}{(1+2\kappa_d)^2} - \tilde{C}\eta \|\Lambda\|_{\mathbb{F}^{r_s} \Lambda_{\ell_1}}^2 \right) \quad (\text{F.12})$$

$$\bar{V}_{t+1} = \bar{V}_t \left(I_{rk} + \frac{\eta(1-2\kappa_d)}{1+1.2\eta} \bar{V}_t \right)^{-1} + \frac{\eta(1-2\kappa_d)}{1+1.2\eta} \left(\frac{\Lambda_{u_1}^2}{(1-2\kappa_d)^2} + \tilde{C}\eta \|\Lambda\|_{\mathbb{F}^{r_s} \Lambda_{u_1}}^2 \right). \quad (\text{F.13})$$

where the iterates $\{\underline{V}_t\}_{t \in \mathbb{N}}$ and $\{\bar{V}_t\}_{t \in \mathbb{N}}$ are functions of the bounding sequences $\{\underline{G}_t\}_{t \in \mathbb{N}}$ and $\{\bar{G}_t\}_{t \in \mathbb{N}}$ as following:

$$\underline{K}_t := \Lambda_{\ell_2}^{\frac{1}{2}} \underline{G}_t \Lambda_{\ell_2}^{\frac{1}{2}} \quad \text{and} \quad \underline{V}_t = 2\underline{K}_t - \frac{\Lambda_{\ell_1}}{1+2\kappa_d} \quad \text{and} \quad \underline{G}_0 \preceq G_0 - T_0,$$

$$\bar{K}_t := \Lambda_{u_2}^{\frac{1}{2}} \bar{G}_t \Lambda_{u_2}^{\frac{1}{2}} \quad \text{and} \quad \bar{V}_t = 2\bar{K}_t - \frac{\Lambda_{u_1}}{1-2\kappa_d} \quad \text{and} \quad \bar{G}_0 \succeq G_0 + T_0.$$

• **Stopping times.** We define a sequence of events $\{\mathcal{E}_t\}_{t \geq 0}$

$$\mathcal{E}_t := \begin{cases} \left\{ -\kappa_d r^{-\frac{\alpha}{2}} T_t \preceq \underline{V}_t \preceq \kappa_d r^{-\frac{\alpha}{2}} T_t \right\} \cap \left\{ -\kappa_d^2 r^{-\alpha} T_t \preceq \Lambda^{\frac{1}{2}} \underline{V}_t \Lambda^{\frac{1}{2}} \preceq \kappa_d^2 r^{-\alpha} T_t \right\}, & \alpha \in [0, 0.5) \\ \left\{ -\kappa_d r_u^{-\frac{\alpha}{2}} T_t \preceq \underline{V}_t \preceq \kappa_d r_u^{-\frac{\alpha}{2}} T_t \right\} \cap \left\{ -\frac{\kappa_d^2}{4} r_u^{-\alpha} T_t \preceq \Lambda_{11}^{\frac{1}{2}} \underline{V}_t \Lambda_{11}^{\frac{1}{2}} \preceq \frac{\kappa_d^2}{4} r_u^{-\alpha} T_t \right\}, & \alpha > 0.5. \end{cases}$$

We define the stopping times

$$\mathcal{T}_{\text{noise}}(\omega) := \inf \{t \geq 0 \mid \omega \notin \mathcal{E}_t\} \wedge d^3 \quad \text{and} \quad \mathcal{T}_{\text{bounded}} := \inf \{t \geq 0 \mid \|\underline{G}_t\|_2 \vee \|\bar{G}_t\|_2 \geq 1.5\},$$

and

$$\mathcal{T}_{\text{bad}} := \mathcal{T}_{\text{noise}} \wedge \mathcal{T}_{\text{bounded}} \wedge \{t \geq 0 : \|T_t\|_2 > 1.2\kappa_d\}. \quad (\text{F.14})$$

The main result of this section is the following:

Proposition 14. Let κ_d satisfy (F.15) and consider d large enough so that $\kappa_d \leq \frac{1}{50}$. Under the learning rate conditions considered in Propositions 12 and 13, we have for $\tilde{C} > 1 \vee \Omega(1)$

$$\underline{G}_{t \wedge \mathcal{T}_{\text{bad}}} + T_{t \wedge \mathcal{T}_{\text{bad}}} + \underline{V}_{t \wedge \mathcal{T}_{\text{bad}}} \preceq G_{t \wedge \mathcal{T}_{\text{bad}}} \preceq \bar{G}_{t \wedge \mathcal{T}_{\text{bad}}} - T_{t \wedge \mathcal{T}_{\text{bad}}} + \underline{V}_{t \wedge \mathcal{T}_{\text{bad}}}.$$

Before starting the proof, we provide an auxiliary statement.

Lemma 4. We consider the learning rate conditions considered in Propositions 12 and 13 with

$$\kappa_d \ll \begin{cases} \frac{1}{\log d}, & \alpha \in [0, 0.5) \\ \frac{r_u}{\log d}, & \alpha > 0.5. \end{cases} \quad (\text{F.15})$$

The event $\mathcal{E}_t$ implies for $d \geq \Omega(1)$ and $t \leq \mathcal{T}_{\text{bounded}} \wedge \{t : \|T_t\|_2 > 1.2\kappa_d\}$ that

1. $-3\kappa_d \Lambda_{\ell_1}^{\frac{1}{2}} T_t \Lambda_{\ell_1}^{\frac{1}{2}} \preceq \Lambda_{\ell_1} \underline{V}_t + \underline{V}_t \Lambda_{\ell_1}$
2. $\Lambda_{u_1} \underline{V}_t + \underline{V}_t \Lambda_{u_1} \preceq 3\kappa_d \Lambda_{u_1}^{\frac{1}{2}} T_t \Lambda_{u_1}^{\frac{1}{2}}$
3. $(\Lambda_{\ell_2}^{\frac{1}{2}} \underline{V}_t \Lambda_{\ell_2}^{\frac{1}{2}})^2 \preceq \frac{\kappa_d^2}{4} \Lambda_{\ell_1} T_t \Lambda_{\ell_1}$
4. $(\Lambda_{u_2}^{\frac{1}{2}} \underline{V}_t \Lambda_{u_2}^{\frac{1}{2}})^2 \preceq \frac{\kappa_d^2}{4} \Lambda_{u_1} T_t \Lambda_{u_1}$

Proof. For notational convenience, we define $\tilde{V}_t := T_t^{-\frac{1}{2}} \underline{V}_t T_t^{-\frac{1}{2}}$. We initially observe that $\mathcal{E}_t$ implies for $\alpha \in [0, 0.5)$ that

$$\begin{aligned} \Lambda \tilde{V}_t + \tilde{V}_t \Lambda &\preceq \kappa_d^2 \Lambda + \frac{1}{\kappa_d^2} \Lambda^{-\frac{1}{2}} \left(\Lambda^{\frac{1}{2}} \tilde{V}_t \Lambda^{\frac{1}{2}} \right)^2 \Lambda^{-\frac{1}{2}} \preceq \kappa_d^2 \Lambda + \kappa_d^2 r^{-\alpha} I_r \\ \Lambda \tilde{V}_t + \tilde{V}_t \Lambda &\succeq -\kappa_d^2 \Lambda - \frac{1}{\kappa_d^2} \Lambda^{-\frac{1}{2}} \left(\Lambda^{\frac{1}{2}} \tilde{V}_t \Lambda^{\frac{1}{2}} \right)^2 \Lambda^{-\frac{1}{2}} \succeq -\kappa_d^2 \Lambda - \kappa_d^2 r^{-\alpha} I_r. \end{aligned}$$

For $\alpha > 0.5$, these bounds become

$$\begin{aligned}\Lambda_{11}\tilde{\nu}_t + \tilde{\nu}_t\Lambda_{11} &\preceq \frac{\kappa_d^2}{4}\Lambda_{11} + \frac{4}{\kappa_d^2}\Lambda_{11}^{-\frac{1}{2}}\left(\Lambda_{11}^{\frac{1}{2}}\tilde{\nu}_t\Lambda_{11}^{\frac{1}{2}}\right)^2\Lambda_{11}^{-\frac{1}{2}} \preceq \frac{\kappa_d^2}{4}\Lambda_{11} + \frac{\kappa_d^2}{4}r_u^{-\alpha}\mathbf{I}_{r_u} \\ \Lambda_{11}\tilde{\nu}_t + \tilde{\nu}_t\Lambda_{11} &\succeq -\frac{\kappa_d^2}{4}r_u^{-1}\Lambda_{11} - \frac{4}{\kappa_d^2}\Lambda_{11}^{-\frac{1}{2}}\left(\Lambda_{11}^{\frac{1}{2}}\tilde{\nu}_t\Lambda_{11}^{\frac{1}{2}}\right)^2\Lambda_{11}^{-\frac{1}{2}} \succeq -\frac{\kappa_d^2}{4}r_u^{-1}\Lambda_{11} - \frac{\kappa_d^2}{4}r_u^{-\alpha}\mathbf{I}_{r_u}.\end{aligned}$$

In the following, we will use these bounds. For the first item, we have

$$\begin{aligned}\Lambda_{\ell_1}\tilde{\nu}_t + \tilde{\nu}_t\Lambda_{\ell_1} &\succeq \begin{cases} -\kappa_d\left(\kappa_d\Lambda + \kappa_dr^{-\alpha}\mathbf{I}_r + r^{-\frac{\alpha}{2}}C\|\Lambda\|_F\frac{\eta d}{\sqrt{r_s}}\mathbf{I}_r\right), & \alpha \in [0, 0.5) \\ -\kappa_d\left(\frac{\kappa_d}{4}\Lambda_{11} + \frac{\kappa_d}{4}r_u^{-\alpha}\mathbf{I}_{r_u} + r_u^{-\frac{\alpha}{2}}C\|\Lambda\|_F\frac{\eta d}{\sqrt{r_s}}\mathbf{I}_{r_u}\right), & \alpha > 0.5 \end{cases} \\ &\succeq -3\kappa_d\Lambda_{\ell_1}.\end{aligned}$$

For the second item, we have

$$\begin{aligned}\Lambda_{u_1}\tilde{\nu}_t + \tilde{\nu}_t\Lambda_{u_1} &\preceq \begin{cases} \kappa_d\left(\kappa_d\Lambda + \kappa_dr^{-\alpha}\mathbf{I}_r + r^{-\frac{\alpha}{2}}C\|\Lambda\|_F\frac{\eta d}{\sqrt{r_s}}\mathbf{I}_r\right), & \alpha \in [0, 0.5) \\ \kappa_d\left(\frac{\kappa_d}{4}\Lambda_{11} + \frac{\kappa_d}{4}r_u^{-\alpha}\mathbf{I}_{r_u} + r_u^{-\frac{\alpha}{2}}C\|\Lambda\|_F\frac{\eta d}{\sqrt{r_s}}\mathbf{I}_{r_u}\right), & \alpha > 0.5 \end{cases} \\ &\preceq 3\kappa_d\Lambda_{u_1}.\end{aligned}$$

For the third item, we immediately observe that $(\Lambda_{\ell_2}^{\frac{1}{2}}\underline{\nu}_t\Lambda_{\ell_2}^{\frac{1}{2}})^2 \preceq \Lambda_{\ell_2}^{\frac{1}{2}}\underline{\nu}_t^2\Lambda_{\ell_2}^{\frac{1}{2}}$. Therefore,

$$T_t^{-\frac{1}{2}}\Lambda_{\ell_2}^{\frac{1}{2}}\underline{\nu}_t^2\Lambda_{\ell_2}^{\frac{1}{2}}T_t^{-\frac{1}{2}} \preceq 1.2\kappa_d\Lambda_{\ell_2}^{\frac{1}{2}}\tilde{\nu}_t^2\Lambda_{\ell_2}^{\frac{1}{2}} \preceq 1.2\kappa_d^3\text{rk}^{-\alpha}\Lambda_{\ell_2} \preceq \frac{\kappa_d^2}{4}\Lambda_{\ell_1}^2.$$

For the fourth item, we observe that

$$(\Lambda_{u_2}^{\frac{1}{2}}\underline{\nu}_t\Lambda_{u_2}^{\frac{1}{2}})^2 = \Lambda_{u_2}^{\frac{1}{2}}\underline{\nu}_t\Lambda_{u_2}\underline{\nu}_t\Lambda_{u_2}^{\frac{1}{2}} \preceq \left(1 + \frac{0.1\text{rk}^{-\alpha}}{\log^3 d}\right)\Lambda_{u_2}^{\frac{1}{2}}\underline{\nu}_t^2\Lambda_{u_2}^{\frac{1}{2}}.$$

Therefore

$$\left(1 + \frac{0.1\text{rk}^{-\alpha}}{\log^3 d}\right)T_t^{-\frac{1}{2}}\Lambda_{u_2}^{\frac{1}{2}}\underline{\nu}_t^2\Lambda_{u_2}^{\frac{1}{2}}T_t^{-\frac{1}{2}} \preceq 1.25\kappa_d\Lambda_{u_2}^{\frac{1}{2}}\tilde{\nu}_t^2\Lambda_{u_2}^{\frac{1}{2}} \preceq 1.25\kappa_d^3\text{rk}^{-\alpha}\Lambda_{u_2} \preceq \frac{\kappa_d^2}{4}\Lambda_{u_1}^2.$$

□

F.3.1 Proof of Proposition 14

Proof. For the proof, we introduce the following notations

$$\underline{\zeta}_t := 2\Lambda_{\ell_2}^{\frac{1}{2}}\underline{\nu}_t\Lambda_{\ell_2}^{\frac{1}{2}} \text{ and } \bar{\zeta}_t := 2\Lambda_{u_2}^{\frac{1}{2}}\underline{\nu}_t\Lambda_{u_2}^{\frac{1}{2}} \text{ and } \underline{B}_t := 2\Lambda_{\ell_2}^{\frac{1}{2}}T_t\Lambda_{\ell_2}^{\frac{1}{2}} \text{ and } \bar{B}_t := 2\Lambda_{u_2}^{\frac{1}{2}}T_t\Lambda_{u_2}^{\frac{1}{2}}.$$

Using this notation, we obtain:

$$\underline{B}_{t+1} \preceq \underline{B}_t + \frac{2(1-2\kappa_d)\eta}{1-1.1\eta}\left(\Lambda_{\ell_1}\underline{B}_t - \frac{1.5\kappa_d+0.5}{\kappa_d(1-2\kappa_d)}\underline{B}_t^2\right) + 10\eta^2\kappa_d\Lambda_{\ell_1} \quad (\text{F.16})$$

$$\bar{B}_{t+1} \preceq \bar{B}_t + \frac{2(1-2\kappa_d)\eta}{1+1.1\eta}\left(\Lambda_{u_1}\bar{B}_t - \frac{1.5\kappa_d+0.5}{\kappa_d(1-2\kappa_d)}\bar{B}_t^2\right) + 10\eta^2\kappa_d\Lambda_{u_1} \quad (\text{F.17})$$

Before proceeding with the proof, we observe that the following inequalities hold:

$$\|\Lambda_{\ell_2}^{-1}\Lambda_{\ell_1}\|_2 \leq 1, \|\Lambda_{\ell_1}^{-1}\Lambda_{\ell_2}\|_2 \leq \frac{1}{1-\frac{0.1}{\log^4 d}} \text{ and } \|\Lambda_{u_2}^{-1}\Lambda_{u_1}\|_2 \leq 1, \|\Lambda_{u_1}^{-1}\Lambda_{u_2}\|_2 \leq \frac{1}{1-\frac{0.1}{\log^4 d}}.$$

These bounds will be used in the following whenever we apply Propositions 29 and 30, without explicitly restating them each time. We will establish the upper and lower bounds simultaneously for $\text{rk} \in \{r, r_u\}$.

Upper bound proof: We will use proof by induction. Specifically, we will show that for $t < \mathcal{T}_{\text{bad}}$,

$$\mathbf{V}_t^+ \preceq \bar{\mathbf{V}}_t + \bar{\boldsymbol{\zeta}}_t - \bar{\mathbf{B}}_t + \frac{2\kappa_d \boldsymbol{\Lambda}_{u_1}}{1 - 2\kappa_d} \Rightarrow \mathbf{V}_{t+1}^+ \preceq \bar{\mathbf{V}}_{t+1} + \bar{\boldsymbol{\zeta}}_{t+1} - \bar{\mathbf{B}}_{t+1} + \frac{2\kappa_d \boldsymbol{\Lambda}_{u_1}}{1 - 2\kappa_d}. \quad (\text{F.18})$$

Since the base case holds at $t = 0$ and $\mathcal{T}_{\text{bad}} > 0$, it remains to prove (F.18). By (F.6), we have

$$\begin{aligned} \mathbf{V}_{t+1}^+ &\preceq \left(\bar{\mathbf{V}}_t + \bar{\boldsymbol{\zeta}}_t - \bar{\mathbf{B}}_t + \frac{2\kappa_d \boldsymbol{\Lambda}_{u_1}}{1 - 2\kappa_d} \right) \left(\mathbf{I}_{\text{rk}} + \frac{\eta}{1 + 1.1\eta} (\bar{\mathbf{V}}_t + \bar{\boldsymbol{\zeta}}_t - \bar{\mathbf{B}}_t + \frac{2\kappa_d \boldsymbol{\Lambda}_{u_1}}{1 - 2\kappa_d}) \right)^{-1} \\ &\quad + \eta \boldsymbol{\Lambda}_{u_1}^2 + C\eta^2 \|\boldsymbol{\Lambda}\|_{\text{F}}^2 r_s \boldsymbol{\Lambda}_{u_2} + \frac{\eta}{\sqrt{r_s}} \boldsymbol{\Lambda}_{u_2}^{\frac{1}{2}} \mathbf{v}_{t+1} \boldsymbol{\Lambda}_{u_2}^{\frac{1}{2}} \\ &= \bar{\mathbf{V}}_t \left(\mathbf{I}_{\text{rk}} + \frac{\eta}{1 + 1.1\eta} \bar{\mathbf{V}}_t \right)^{-1} + \eta \boldsymbol{\Lambda}_{u_1}^2 + C\eta^2 \|\boldsymbol{\Lambda}\|_{\text{F}}^2 r_s \boldsymbol{\Lambda}_{u_2} + \frac{\eta}{\sqrt{r_s}} \boldsymbol{\Lambda}_{u_2}^{\frac{1}{2}} \mathbf{v}_{t+1} \boldsymbol{\Lambda}_{u_2}^{\frac{1}{2}} \\ &\quad + \left(\mathbf{I}_{\text{rk}} + \frac{\eta}{1 + 1.1\eta} \bar{\mathbf{V}}_t \right)^{-1} (\bar{\boldsymbol{\zeta}}_t - \bar{\mathbf{B}}_t + \frac{2\kappa_d \boldsymbol{\Lambda}_{u_1}}{1 - 2\kappa_d}) \left(\mathbf{I}_{\text{rk}} + \frac{\eta}{1 + 1.1\eta} (\bar{\mathbf{V}}_t + \bar{\boldsymbol{\zeta}}_t - \bar{\mathbf{B}}_t + \frac{2\kappa_d \boldsymbol{\Lambda}_{u_1}}{1 - 2\kappa_d}) \right)^{-1}. \end{aligned} \quad (\text{F.19})$$

By using Proposition 29, we have for $t < \mathcal{T}_{\text{bad}}$

$$\begin{aligned} &\left(\mathbf{I}_{\text{rk}} + \frac{\eta}{1 + 1.1\eta} \bar{\mathbf{V}}_t \right)^{-1} (\bar{\boldsymbol{\zeta}}_t - \bar{\mathbf{B}}_t + \frac{2\kappa_d \boldsymbol{\Lambda}_{u_1}}{1 - 2\kappa_d}) \left(\mathbf{I}_{\text{rk}} + \frac{\eta}{1 + 1.1\eta} (\bar{\mathbf{V}}_t + \bar{\boldsymbol{\zeta}}_t - \bar{\mathbf{B}}_t + \frac{2\kappa_d \boldsymbol{\Lambda}_{u_1}}{1 - 2\kappa_d}) \right)^{-1} \\ &\preceq \left(\bar{\boldsymbol{\zeta}}_t - \bar{\mathbf{B}}_t + \frac{2\kappa_d \boldsymbol{\Lambda}_{u_1}}{1 - 2\kappa_d} \right) - \frac{\eta}{1 + 1.1\eta} \bar{\mathbf{V}}_t \left(\bar{\boldsymbol{\zeta}}_t - \bar{\mathbf{B}}_t + \frac{2\kappa_d \boldsymbol{\Lambda}_{u_1}}{1 - 2\kappa_d} \right) \\ &\quad - \frac{\eta}{1 + 1.1\eta} \left(\bar{\boldsymbol{\zeta}}_t - \bar{\mathbf{B}}_t + \frac{2\kappa_d \boldsymbol{\Lambda}_{u_1}}{1 - 2\kappa_d} \right) \bar{\mathbf{V}}_t - \frac{\eta}{1 + 1.1\eta} \left(\bar{\boldsymbol{\zeta}}_t - \bar{\mathbf{B}}_t + \frac{2\kappa_d \boldsymbol{\Lambda}_{u_1}}{1 - 2\kappa_d} \right)^2 + \frac{\eta^2 \kappa_d^2}{(1 + 1.1\eta)^2} \bar{\mathbf{V}}_t^4 \\ &\quad + \frac{2\eta^2 / \kappa_d^2}{(1 + 1.1\eta)^2} \left(\bar{\boldsymbol{\zeta}}_t - \bar{\mathbf{B}}_t + \frac{2\kappa_d \boldsymbol{\Lambda}_{u_1}}{1 - 2\kappa_d} \right)^2 + \frac{\eta^2 \kappa_d^2}{(1 + 1.1\eta)^2} \bar{\mathbf{V}}_t \left(\bar{\boldsymbol{\zeta}}_t - \bar{\mathbf{B}}_t + \frac{2\kappa_d \boldsymbol{\Lambda}_{u_1}}{1 - 2\kappa_d} \right)^2 \bar{\mathbf{V}}_t \\ &\quad + \frac{\eta^2}{(1 + 1.1\eta)^2} \left(\bar{\boldsymbol{\zeta}}_t - \bar{\mathbf{B}}_t + \frac{2\kappa_d \boldsymbol{\Lambda}_{u_1}}{1 - 2\kappa_d} \right) \bar{\mathbf{V}}_t \left(\bar{\boldsymbol{\zeta}}_t - \bar{\mathbf{B}}_t + \frac{2\kappa_d \boldsymbol{\Lambda}_{u_1}}{1 - 2\kappa_d} \right) + \eta^3 \tilde{C}_1 \boldsymbol{\Lambda}_{u_1} \\ &\quad + \frac{\eta^2}{(1 + 1.1\eta)^2} \left(\bar{\boldsymbol{\zeta}}_t - \bar{\mathbf{B}}_t + \frac{2\kappa_d \boldsymbol{\Lambda}_{u_1}}{1 - 2\kappa_d} \right)^3 + \frac{\eta^2}{(1 + 1.1\eta)^2} \bar{\mathbf{V}}_t \left(\bar{\boldsymbol{\zeta}}_t - \bar{\mathbf{B}}_t + \frac{2\kappa_d \boldsymbol{\Lambda}_{u_1}}{1 - 2\kappa_d} \right) \bar{\mathbf{V}}_t \end{aligned}$$

for some $\tilde{C}_1 = O(1)$. We have the following: First:

$$\begin{aligned} \kappa_d^2 \bar{\mathbf{V}}_t^4 + \kappa_d^2 \bar{\mathbf{V}}_t \left(\bar{\boldsymbol{\zeta}}_t - \bar{\mathbf{B}}_t + \frac{2\kappa_d \boldsymbol{\Lambda}_{u_1}}{1 - 2\kappa_d} \right)^2 \bar{\mathbf{V}}_t + \bar{\mathbf{V}}_t \left(\bar{\boldsymbol{\zeta}}_t - \bar{\mathbf{B}}_t + \frac{2\kappa_d \boldsymbol{\Lambda}_{u_1}}{1 - 2\kappa_d} \right) \bar{\mathbf{V}}_t &\stackrel{(a)}{\preceq} 4\kappa_d \bar{\mathbf{V}}_t^2 \\ &\stackrel{(b)}{\preceq} \frac{1}{15} \frac{1 + 1.1\eta}{1 + 1.2\eta} \bar{\mathbf{V}}_t^2, \end{aligned}$$

where (a) follows by $\|\bar{\mathbf{V}}_t\|_2 \leq 5$ and $\left\| \bar{\boldsymbol{\zeta}}_t - \bar{\mathbf{B}}_t + \frac{2\kappa_d \boldsymbol{\Lambda}_{u_1}}{1 - 2\kappa_d} \right\|_2 \leq 2.5\kappa_d$, and (b) follows by $\kappa_d \leq \frac{1}{50}$.

Second,

$$\begin{aligned} &\frac{2}{\kappa_d^2} \left(\bar{\boldsymbol{\zeta}}_t - \bar{\mathbf{B}}_t + \frac{2\kappa_d \boldsymbol{\Lambda}_{u_1}}{1 - 2\kappa_d} \right)^2 + \left(\bar{\boldsymbol{\zeta}}_t - \bar{\mathbf{B}}_t + \frac{2\kappa_d \boldsymbol{\Lambda}_{u_1}}{1 - 2\kappa_d} \right) \bar{\mathbf{V}}_t \left(\bar{\boldsymbol{\zeta}}_t - \bar{\mathbf{B}}_t + \frac{2\kappa_d \boldsymbol{\Lambda}_{u_1}}{1 - 2\kappa_d} \right) \\ &\quad + \left(\bar{\boldsymbol{\zeta}}_t + \bar{\mathbf{B}}_t + \frac{2\kappa_d \boldsymbol{\Lambda}_{u_1}}{1 - 2\kappa_d} \right)^3 \stackrel{(c)}{\preceq} \frac{3}{\kappa_d^2} \left(\bar{\boldsymbol{\zeta}}_t - \bar{\mathbf{B}}_t + \frac{2\kappa_d \boldsymbol{\Lambda}_{u_1}}{1 - 2\kappa_d} \right)^2 \preceq \frac{9}{\kappa_d^2} (\bar{\boldsymbol{\zeta}}_t^2 + \bar{\mathbf{B}}_t^2) + \frac{36}{(1 - 2\kappa_d)^2} \boldsymbol{\Lambda}_{u_1}^2, \end{aligned}$$

where (c) follows by $\|\bar{\mathbf{V}}_t\|_2 \leq 5$ and $\left\| \bar{\boldsymbol{\zeta}}_t - \bar{\mathbf{B}}_t + \frac{2\kappa_d \boldsymbol{\Lambda}_{u_1}}{1 - 2\kappa_d} \right\|_2 \leq 2.5\kappa_d$. Third:

$$\begin{aligned} &-\bar{\mathbf{V}}_t \left(\bar{\boldsymbol{\zeta}}_t - \bar{\mathbf{B}}_t + \frac{2\kappa_d \boldsymbol{\Lambda}_{u_1}}{1 - 2\kappa_d} \right) - \left(\bar{\boldsymbol{\zeta}}_t - \bar{\mathbf{B}}_t + \frac{2\kappa_d \boldsymbol{\Lambda}_{u_1}}{1 - 2\kappa_d} \right) \bar{\mathbf{V}}_t - \left(\bar{\boldsymbol{\zeta}}_t - \bar{\mathbf{B}}_t + \frac{2\kappa_d \boldsymbol{\Lambda}_{u_1}}{1 - 2\kappa_d} \right)^2 \\ &\quad + \frac{9\eta / \kappa_d^2}{1 + 1.1\eta} (\bar{\boldsymbol{\zeta}}_t^2 + \bar{\mathbf{B}}_t^2) \end{aligned}$$

$$\begin{aligned}
&\preceq -2(\bar{K}_t \bar{\zeta}_t + \bar{\zeta}_t \bar{K}_t) + 2(\bar{K}_t \bar{B}_t + \bar{B}_t \bar{K}_t) - \frac{4\kappa_d}{1-2\kappa_d}(\bar{K}_t \Lambda_{u_1} + \Lambda_{u_1} \bar{K}_t) + 3(\bar{\zeta}_t^2 + \bar{B}_t^2) \\
&+ (\Lambda_{u_1} \bar{\zeta}_t + \bar{\zeta}_t \Lambda_{u_1}) - (\Lambda_{u_1} \bar{B}_t + \bar{B}_t \Lambda_{u_1}) + \frac{4\kappa_d(1-\kappa_d)}{(1-2\kappa_d)^2} \Lambda_{u_1}^2 \\
&\stackrel{(d)}{\preceq} 8\kappa_d \bar{K}_t^2 - \frac{4\kappa_d}{1-2\kappa_d}(\bar{K}_t \Lambda_{u_1} + \Lambda_{u_1} \bar{K}_t) + \frac{4\kappa_d(1-\kappa_d)}{(1-2\kappa_d)^2} \Lambda_{u_1}^2 \\
&- (2-4\kappa_d) \Lambda_{u_1} \bar{B}_t + \left(3 + \frac{1}{\kappa_d}\right) \bar{B}_t^2 \\
&= 2\kappa_d \bar{V}_t^2 + \frac{2\kappa_d}{1-2\kappa_d} \Lambda_{u_1}^2 - (2-4\kappa_d) \Lambda_{u_1} \bar{B}_t + \left(3 + \frac{1}{\kappa_d}\right) \bar{B}_t^2,
\end{aligned}$$

where we used Proposition 22, and the second and fourth items in Lemma 4 in (d). Therefore, we have

$$\begin{aligned}
(\text{F.19}) &\preceq \bar{V}_t \left(I_{rk} + \frac{\eta}{1+1.1\eta} \bar{V}_t \right)^{-1} + \frac{2\kappa_d \eta}{1+1.1\eta} \bar{V}_t^2 + \frac{1}{10} \frac{(1-2\kappa_d)\eta^2}{(1+1.1\eta)(1+1.2\eta)} \bar{V}_t^2 + \frac{2\kappa_d \Lambda_{u_1}}{1-2\kappa_d} \\
&+ \frac{\eta}{1+1.1\eta} \left(\frac{\Lambda_{u_1}^2}{1-2\kappa_d} + \tilde{C}\eta \|\Lambda\|_{\mathbb{F}}^2 r_s \Lambda_{u_2} \right) + \bar{\zeta}_{t+1} - \bar{B}_t \\
&- \frac{2(1-2\kappa_d)\eta}{1+1.1\eta} \left(\Lambda_{u_1} \bar{B}_t - \frac{1.5\kappa_d + 0.5}{\kappa_d(1-2\kappa_d)} \bar{B}_t^2 \right) \\
&\stackrel{(e)}{\preceq} \bar{V}_t \left(I_{rk} + \frac{\eta(1-2\kappa_d)}{1+1.2\eta} \bar{V}_t \right)^{-1} + \frac{\eta}{1+1.2\eta} \left(\frac{\Lambda_{u_1}^2}{1-2\kappa_d} + \tilde{C}\eta \|\Lambda\|_{\mathbb{F}}^2 r_s \Lambda_{u_1} \right) \\
&+ \bar{\zeta}_{t+1} + \frac{2\kappa_d \Lambda_{u_1}}{1-2\kappa_d} - \bar{B}_{t+1} \\
&\preceq \bar{V}_{t+1} + \bar{\zeta}_{t+1} + \frac{2\kappa_d \Lambda_{u_1}}{1-2\kappa_d} - \bar{B}_{t+1},
\end{aligned}$$

where we used Proposition 30 and (F.17) in (e).

Lower bound proof: Similar to the upper bound proof, here we will show that for $t < \mathcal{T}_{\text{bad}}$,

$$\underline{V}_t + \underline{\zeta}_t + \underline{B}_t - \frac{2\kappa_d \Lambda_{\ell_1}}{1+2\kappa_d} \preceq V_t^- \Rightarrow \underline{V}_{t+1} + \underline{\zeta}_{t+1} + \underline{B}_{t+1} - \frac{2\kappa_d}{1+2\kappa_d} \Lambda_{\ell_1} \preceq V_{t+1}^-. \quad (\text{F.20})$$

Since the base case holds at $t = 0$ and $\mathcal{T}_{\text{bad}} > 0$, it remains to prove (F.20). By (F.6), we have

$$\begin{aligned}
V_{t+1}^- &\succeq \left(\underline{V}_t + \underline{\zeta}_t + \underline{B}_t - \frac{2\kappa_d \Lambda_{\ell_1}}{1+2\kappa_d} \right) \left(I_{rk} + \frac{\eta}{1-1.1\eta} \left(\underline{V}_t + \underline{\zeta}_t + \underline{B}_t - \frac{2\kappa_d \Lambda_{\ell_1}}{1+2\kappa_d} \right) \right)^{-1} + \eta \Lambda_{\ell_1}^2 \\
&- C\eta^2 \|\Lambda\|_{\mathbb{F}}^2 r_s \Lambda_{\ell_2} + \frac{\eta}{\sqrt{r_s}} \Lambda_{\ell_2}^{\frac{1}{2}} v_{t+1} \Lambda_{\ell_2}^{\frac{1}{2}} \\
&= \underline{V}_t \left(I_{rk} + \frac{\eta}{1-1.1\eta} \underline{V}_t \right)^{-1} + \eta \Lambda_{\ell_1}^2 - C\eta^2 \|\Lambda\|_{\mathbb{F}}^2 r_s \Lambda_{\ell_2} + \frac{\eta}{\sqrt{r_s}} \Lambda_{\ell_2}^{\frac{1}{2}} v_{t+1} \Lambda_{\ell_2}^{\frac{1}{2}} \\
&+ \left(I_{rk} + \frac{\eta}{1-1.1\eta} \underline{V}_t \right)^{-1} \left(\underline{\zeta}_t + \underline{B}_t - \frac{2\kappa_d \Lambda_{\ell_1}}{1+2\kappa_d} \right) \left(I_{rk} + \frac{\eta}{1-1.1\eta} \left(\underline{V}_t + \underline{\zeta}_t + \underline{B}_t - \frac{2\kappa_d \Lambda_{\ell_1}}{1+2\kappa_d} \right) \right)^{-1}. \quad (\text{F.21})
\end{aligned}$$

By using Proposition 29, we have for $t < \mathcal{T}_{\text{bad}}$

$$\begin{aligned}
&\left(I_{rk} + \frac{\eta}{1-1.1\eta} \underline{V}_t \right)^{-1} \left(\underline{\zeta}_t + \underline{B}_t - \frac{2\kappa_d \Lambda_{\ell_1}}{1+2\kappa_d} \right) \left(I_{rk} + \frac{\eta}{1-1.1\eta} \left(\underline{V}_t + \underline{\zeta}_t + \underline{B}_t - \frac{2\kappa_d \Lambda_{\ell_1}}{1+2\kappa_d} \right) \right)^{-1} \\
&\succeq \left(\underline{\zeta}_t + \underline{B}_t - \frac{2\kappa_d \Lambda_{\ell_1}}{1+2\kappa_d} \right) - \frac{\eta}{1-1.1\eta} \underline{V}_t \left(\underline{\zeta}_t + \underline{B}_t - \frac{2\kappa_d \Lambda_{\ell_1}}{1+2\kappa_d} \right) - \frac{\eta^2 \kappa_d^2}{(1-1.1\eta)^2} \underline{V}_t^4 \\
&- \frac{\eta}{1-1.1\eta} \left(\underline{\zeta}_t + \underline{B}_t - \frac{2\kappa_d \Lambda_{\ell_1}}{1+2\kappa_d} \right) \underline{V}_t - \frac{\eta}{1-1.1\eta} \left(\underline{\zeta}_t + \underline{B}_t - \frac{2\kappa_d \Lambda_{\ell_1}}{1+2\kappa_d} \right)^2
\end{aligned}$$

$$\begin{aligned}
& -\frac{2\eta^2/\kappa_d^2}{(1-1.1\eta)^2} \left(\underline{\zeta}_t + \underline{B}_t - \frac{2\kappa_d \Lambda_{\ell_1}}{1+2\kappa_d} \right)^2 - \frac{\eta^2 \kappa_d^2}{(1-1.1\eta)^2} \underline{V}_t \left(\underline{\zeta}_t + \underline{B}_t - \frac{2\kappa_d \Lambda_{\ell_1}}{1+2\kappa_d} \right)^2 \underline{V}_t \\
& + \frac{\eta^2}{(1-1.1\eta)^2} \left(\underline{\zeta}_t + \underline{B}_t - \frac{2\kappa_d \Lambda_{\ell_1}}{1+2\kappa_d} \right) \underline{V}_t \left(\underline{\zeta}_t + \underline{B}_t - \frac{2\kappa_d \Lambda_{\ell_1}}{1+2\kappa_d} \right) - \eta^3 \tilde{C}_2 \Lambda_{\ell_1} \\
& + \frac{\eta^2}{(1-1.1\eta)^2} \left(\underline{\zeta}_t + \underline{B}_t - \frac{2\kappa_d \Lambda_{\ell_1}}{1+2\kappa_d} \right)^3 + \frac{\eta^2}{(1-1.1\eta)^2} \underline{V}_t \left(\underline{\zeta}_t + \underline{B}_t - \frac{2\kappa_d \Lambda_{\ell_1}}{1+2\kappa_d} \right) \underline{V}_t
\end{aligned}$$

for some $\tilde{C}_2 = O(1)$. We have the following: First:

$$\begin{aligned}
& -\kappa_d^2 \underline{V}_t^4 - \kappa_d^2 \underline{V}_t \left(\underline{\zeta}_t + \underline{B}_t - \frac{2\kappa_d \Lambda_{\ell_1}}{1+2\kappa_d} \right)^2 \underline{V}_t + \underline{V}_t \left(\underline{\zeta}_t + \underline{B}_t - \frac{2\kappa_d \Lambda_{\ell_1}}{1+2\kappa_d} \right) \underline{V}_t \stackrel{(f)}{\geq} -3.2\kappa_d \underline{V}_t^2 \\
& \stackrel{(g)}{\geq} -\frac{1}{15} \frac{1-1.1\eta}{1-1.2\eta} \underline{V}_t^2.
\end{aligned}$$

where (f) follows by $\|\underline{V}_t\|_2 \leq 5$ and $\left\| \underline{\zeta}_t + \underline{B}_t - \frac{2\kappa_d \Lambda_{\ell_1}}{1+2\kappa_d} \right\|_2 \leq 2.5\kappa_d$, and (g) follows by $\kappa_d \leq \frac{1}{50}$.

Second:

$$\begin{aligned}
& -\frac{2}{\kappa_d^2} \left(\underline{\zeta}_t + \underline{B}_t - \frac{2\kappa_d \Lambda_{\ell_1}}{1+2\kappa_d} \right)^2 + \left(\underline{\zeta}_t + \underline{B}_t - \frac{2\kappa_d \Lambda_{\ell_1}}{1+2\kappa_d} \right) \underline{V}_t \left(\underline{\zeta}_t + \underline{B}_t - \frac{2\kappa_d \Lambda_{\ell_1}}{1+2\kappa_d} \right) \\
& + \left(\underline{\zeta}_t + \underline{B}_t - \frac{2\kappa_d \Lambda_{\ell_1}}{1+2\kappa_d} \right)^3 \stackrel{(h)}{\geq} \frac{-3}{\kappa_d^2} \left(\underline{\zeta}_t + \underline{B}_t - \frac{2\kappa_d \Lambda_{\ell_1}}{1+2\kappa_d} \right)^2 \geq \frac{-9}{\kappa_d^2} (\underline{\zeta}_t^2 + \underline{B}_t^2) - 36\Lambda_{\ell_1}^2,
\end{aligned}$$

where (h) follows by $\|\underline{V}_t\|_2 \leq 5$ and $\left\| \underline{\zeta}_t + \underline{B}_t - \frac{2\kappa_d \Lambda_{\ell_1}}{1+2\kappa_d} \right\|_2 \leq 2.5\kappa_d$. Third:

$$\begin{aligned}
& -\underline{V}_t \left(\underline{\zeta}_t + \underline{B}_t - \frac{2\kappa_d \Lambda_{\ell_1}}{1+2\kappa_d} \right) - \left(\underline{\zeta}_t + \underline{B}_t - \frac{2\kappa_d \Lambda_{\ell_1}}{1+2\kappa_d} \right) \underline{V}_t - \left(\underline{\zeta}_t + \underline{B}_t - \frac{2\kappa_d \Lambda_{\ell_1}}{1+2\kappa_d} \right)^2 \\
& - \frac{9\eta/\kappa_d^2}{1-1.1\eta} (\underline{\zeta}_t^2 + \underline{B}_t^2) \\
& \geq -2 \left(\underline{K}_t \underline{\zeta}_t + \underline{\zeta}_t \underline{K}_t \right) - 2 \left(\underline{K}_t \underline{B}_t + \underline{B}_t \underline{K}_t \right) + \frac{4\kappa_d}{1+2\kappa_d} (\underline{K}_t \Lambda_{\ell_1} + \Lambda_{\ell_1} \underline{K}_t) - 3(\underline{\zeta}_t^2 + \underline{B}_t^2) \\
& + \left(\Lambda_{\ell_1} \underline{\zeta}_t + \underline{\zeta}_t \Lambda_{\ell_1} \right) + \left(\Lambda_{\ell_1} \underline{B}_t + \underline{B}_t \Lambda_{\ell_1} \right) - \frac{4\kappa_d(1+\kappa_d)\Lambda_{\ell_1}^2}{(1+2\kappa_d)^2} \\
& \stackrel{(i)}{\geq} -8c_d \underline{K}_t^2 + \frac{4\kappa_d}{1+2\kappa_d} (\underline{K}_t \Lambda_{\ell_1} + \Lambda_{\ell_1} \underline{K}_t) - \frac{4\kappa_d(1+\kappa_d)\Lambda_{\ell_1}^2}{(1+2\kappa_d)^2} \\
& + (2-4\kappa_d)\Lambda_{\ell_1} \underline{B}_t - \left(3 + \frac{1}{\kappa_d} \right) \underline{B}_t^2 \\
& = -2\kappa_d \underline{V}_t^2 - \frac{2\kappa_d \Lambda_{\ell_1}^2}{(1+2\kappa_d)} + (2-4\kappa_d)\underline{B}_t \Lambda_{\ell_1} - \left(3 + \frac{1}{\kappa_d} \right) \underline{B}_t^2,
\end{aligned}$$

where we used Proposition 22, and the first and third items in Lemma 4 in (i). Therefore, we have

$$\begin{aligned}
(F.21) & \geq \underline{V}_t \left(\underline{I}_r + \frac{\eta}{1-1.1\eta} \underline{V}_t \right)^{-1} - \frac{2\kappa_d \eta}{1-1.1\eta} \underline{V}_t^2 - \frac{(1+2\kappa_d)\eta^2/15}{(1-1.1\eta)(1-1.2\eta)} \underline{V}_t^2 - \frac{2\kappa_d}{1+2\kappa_d} \Lambda_{\ell_1} \\
& + \frac{\eta}{1-1.1\eta} \left(\frac{\Lambda_{\ell_1}^2}{1+2\kappa_d} - \tilde{C}\eta \|\Lambda\|_{\mathbb{F}}^2 r_s \Lambda_{\ell_2} \right) + \underline{\zeta}_{t+1} + \underline{B}_t \\
& + \frac{2(1-2\kappa_d)\eta}{1-1.1\eta} \left(\Lambda_{\ell_1} \underline{B}_t - \frac{1.5\kappa_d + 0.5}{\kappa_d(1-2\kappa_d)} \underline{B}_t^2 \right) \\
& \stackrel{(j)}{\geq} \underline{V}_t \left(\underline{I}_r + \frac{\eta(1+2\kappa_d)}{1-1.2\eta} \underline{V}_t \right)^{-1} + \frac{\eta}{1-1.1\eta} \left(\frac{\Lambda_{\ell_1}^2}{1+2\kappa_d} - \tilde{C}\eta \|\Lambda\|_{\mathbb{F}}^2 r_s \Lambda_{\ell_1} \right) + \underline{\zeta}_{t+1} \\
& + \underline{B}_{t+1} - \frac{2\kappa_d}{1+2\kappa_d} \Lambda_{\ell_1}
\end{aligned}$$

$$= \underline{\mathbf{V}}_{t+1} + \underline{\zeta}_{t+1} + \underline{\mathbf{B}}_{t+1} - \frac{2\kappa_d}{1+2\kappa_d} \mathbf{\Lambda}_{\ell_1},$$

where we used Proposition 30 and (F.16) in (j). $\square$

F.4 Analysis of the bounding systems

F.4.1 Lower bounding system

In this section, we consider (F.12). For notational convenience, we multiply both sides by the factor $(1+2\kappa_d)$ and use a generic learning rate η , i.e.,

$$\underline{\mathbf{V}}_{t+1} = \underline{\mathbf{V}}_t (\mathbf{I}_{\text{rk}} + \eta \underline{\mathbf{V}}_t)^{-1} + \eta \left(\mathbf{\Lambda}_{\ell_1}^2 - \tilde{C} \eta \|\mathbf{\Lambda}\|_{\mathbb{F}}^2 r_s \mathbf{\Lambda}_{\ell_1} \right), \text{ where } \underline{\mathbf{V}}_t = 2\mathbf{\Lambda}_{\ell_2}^{\frac{1}{2}} \underline{\mathbf{G}}_t \mathbf{\Lambda}_{\ell_2}^{\frac{1}{2}} - \mathbf{\Lambda}_{\ell_1}.$$

The main result of this section is stated in Proposition 15. To establish it, we first prove an auxiliary result, Lemma 5. For the following, we define

$$\hat{\mathbf{\Lambda}} := \sqrt{\mathbf{\Lambda}_{\ell_1}^2 - \tilde{C} \eta \|\mathbf{\Lambda}\|_{\mathbb{F}}^2 r_s \mathbf{\Lambda}_{\ell_1}} = \text{diag}(\{\hat{\lambda}_i\}_{i=1}^r), \quad \mathbf{D}_t := \frac{\mathbf{\Lambda}_{\ell_2}^{-1} \hat{\mathbf{\Lambda}} \left(\frac{\mathbf{A}_{t,11}}{\mathbf{A}_{t,12}} - \mathbf{I}_{\text{rk}} \right)}{2} - \frac{1.1\kappa_d r_s}{d} \mathbf{I}_{\text{rk}}.$$

By Corollary 3, we have

$$\underline{\mathbf{G}}_t = \frac{1}{2} \left(\frac{\mathbf{\Lambda}_{\ell_1}}{\mathbf{\Lambda}_{\ell_2}} + \frac{\mathbf{A}_{t,22}}{\mathbf{A}_{t,12}} \frac{\hat{\mathbf{\Lambda}}}{\mathbf{\Lambda}_{\ell_2}} \right) - \frac{1}{4} \frac{\mathbf{\Lambda}_{\ell_2}^{-1} \hat{\mathbf{\Lambda}}}{\mathbf{\Lambda}_{\ell_2}} \left(\frac{\frac{\hat{\mathbf{\Lambda}}}{\mathbf{\Lambda}_{\ell_2}} \left(\frac{\mathbf{A}_{t,11}}{\mathbf{A}_{t,12}} - \mathbf{I}_{\text{rk}} \right)}{2} + \frac{(\hat{\mathbf{\Lambda}} - \mathbf{\Lambda}_{\ell_1})}{2\mathbf{\Lambda}_{\ell_2}} + \underline{\mathbf{G}}_0 \right)^{-1} \frac{\hat{\mathbf{\Lambda}} \mathbf{\Lambda}_{\ell_2}^{-1}}{\mathbf{\Lambda}_{\ell_2}}, \quad (\text{F.22})$$

where $\mathbf{A}_{t,11}$, $\mathbf{A}_{t,12}$, and $\mathbf{A}_{t,22}$ are defined as in (R.1) with $\hat{\mathbf{\Lambda}}$. For $\alpha = 0$, we will consider $\{\underline{\mathbf{G}}_t\}_{t \in \mathbb{N}}$ in the basis of $\underline{\mathbf{G}}_0$ without writing explicitly, which will imply that $\{\underline{\mathbf{G}}_t\}_{t \in \mathbb{N}}$ is diagonal due to the rotational symmetry for $\alpha = 0$.

We further decompose $\{\underline{\mathbf{G}}_t\}_{t \in \mathbb{N}}$ and related matrices to isolate their top-left submatrices of dimension $\text{rk}_* \in \{r_*, r_{u_*}\}$, where $r_* < r$ and $r_{u_*} < r_u$ which we will denote as $\text{rk}_* < \text{rk}$. The decompositions are as follows:

$$\underline{\mathbf{G}}_t := \begin{bmatrix} \underline{\mathbf{G}}_{t,11} & \underline{\mathbf{G}}_{t,12} \\ \underline{\mathbf{G}}_{t,12}^\top & \underline{\mathbf{G}}_{t,22} \end{bmatrix}, \quad \hat{\mathbf{\Lambda}} := \begin{bmatrix} \hat{\mathbf{\Lambda}}_{11} & 0 \\ 0 & \hat{\mathbf{\Lambda}}_{22} \end{bmatrix}, \quad \mathbf{D}_t := \begin{bmatrix} \mathbf{D}_{t,1} & 0 \\ 0 & \mathbf{D}_{t,2} \end{bmatrix}, \quad \mathbf{Z}_{1:\text{rk}_*} := \begin{bmatrix} \mathbf{Z}_{1:\text{rk}_*} \\ \mathbf{Z}_2 \end{bmatrix},$$

where $\underline{\mathbf{G}}_{t,11}$, $\mathbf{D}_{t,1}$, $\hat{\mathbf{\Lambda}}_{11} \in \mathbb{R}^{\text{rk}_* \times \text{rk}_*}$. We define

$$\mathbf{\Gamma}_t := \mathbf{D}_t + \frac{1}{1.05} \mathbf{Z}_{1:\text{rk}_*} \mathbf{Z}_{1:\text{rk}_*}^\top = \begin{bmatrix} \frac{1}{1.05} \mathbf{Z}_{1:\text{rk}_*} \mathbf{Z}_{1:\text{rk}_*}^\top + \mathbf{D}_{t,1} & \frac{1}{1.05d} \mathbf{Z}_{1:\text{rk}_*} \mathbf{Z}_2^\top \\ \frac{1}{1.05} \mathbf{Z}_2 \mathbf{Z}_{1:\text{rk}_*}^\top & \frac{1}{1.05} \mathbf{Z}_2 \mathbf{Z}_2^\top + \mathbf{D}_{t,2} \end{bmatrix}$$

and

$$\mathbf{\Gamma}_t^{-1} := \begin{bmatrix} (\mathbf{\Gamma}_t^{-1})_{11} & (\mathbf{\Gamma}_t^{-1})_{12} \\ (\mathbf{\Gamma}_t^{-1})_{12}^\top & (\mathbf{\Gamma}_t^{-1})_{22} \end{bmatrix}$$

whenever $\mathbf{\Gamma}_t$ is invertible. Lemma 5 is stated as follows:

Lemma 5. *We consider the following setting:*

$$\begin{aligned} \alpha \in [0, 0.5]: \quad & \frac{r_s}{r} \rightarrow \varphi \in (0, \infty), \quad \eta \ll \frac{1}{d r^{1-\alpha} \log^4 d}, \quad \kappa_d = \frac{1}{\log^{3.5} d}, \\ \alpha > 0.5: \quad & r_s \asymp 1, \quad \eta \ll \frac{1}{d r_u^{2+\alpha} \log^3 d}, \quad \kappa_d = \frac{1}{r_u \log^{2.5} d}. \end{aligned}$$

$\mathcal{G}_{\text{init}}$ implies the following:

- For $\alpha \geq 0$ and $K \leq \text{rk}_* \leq \text{rk}$, we have for $\eta t \leq \frac{1}{2}(K+1)^\alpha \log \left(\frac{d \log^{1.5} d}{r_s} \right)$,

$$\mathbf{D}_t \succeq \frac{\mathbf{\Lambda}_{\ell_2}^{-1} \hat{\mathbf{\Lambda}} \left(\frac{\mathbf{A}_{t,22}}{\mathbf{A}_{t,12}} - \mathbf{I}_{\text{rk}} \right)}{2} - \frac{1.2\kappa_d r_s}{d} \mathbf{I}_{\text{rk}}, \quad \mathbf{D}_{t,2} \succeq \frac{\log^3 d - 1}{\log^3 d} \left(\frac{0.5r_s}{d \log^{1.5} d} \right)^{\left(\frac{K+1}{\text{rk}_*+1} \right)^\alpha} \mathbf{I}_{\text{rk}-\text{rk}_*}.$$

- For $r_\star = \lfloor r_s(1 - \log^{-\frac{1}{2}} d) \wedge r \rfloor$ and $r_{u_\star} = r_s$, and $\eta t \leq \frac{1}{2}(\text{rk}_\star + 1)^\alpha \log\left(\frac{d \log^{1.5} d}{r_s}\right)$, we have

$$\Gamma_t \succeq \frac{\Lambda_{\ell_2}^{-1} \hat{\Lambda} \left(\frac{\mathbf{A}_{t,22}}{\mathbf{A}_{t,12}} - \mathbf{I}_{\text{rk}} \right)}{2} + \begin{bmatrix} \frac{C_1 r_s}{d \log^{4.5} d} \mathbf{I}_{\text{rk}_\star} & 0 \\ 0 & -\frac{C_2 r_s}{d \log^2 d} \mathbf{I}_{\text{rk} - \text{rk}_\star} \end{bmatrix} \succ 0, \quad (\text{F.23})$$

where

$$C_1 = \begin{cases} \frac{1}{10}, & \alpha \in [0, 0.5) \\ \left(\frac{1}{1.1 r_s^\alpha} - \frac{1.3}{\sqrt{\log d}} \right), & \alpha > 0.5 \end{cases} \quad \text{and} \quad C_2 = \begin{cases} 2.1 \left(1 + \frac{1}{\sqrt{\varphi}} \right)^2, & \alpha \in [0, 0.5) \\ 2, & \alpha > 0.5. \end{cases}$$

Within the same time interval, we have

$$\Gamma_t^{-1} \preceq \left(\frac{\Lambda_{\ell_2}^{-1} \hat{\Lambda} \left(\frac{\mathbf{A}_{t,22}}{\mathbf{A}_{t,12}} - \mathbf{I}_{\text{rk}} \right)}{2} + \begin{bmatrix} \frac{C_1 r_s}{d \log^{4.5} d} \mathbf{I}_{\text{rk}_\star} & 0 \\ 0 & -\frac{C_2 r_s}{d \log^2 d} \mathbf{I}_{\text{rk} - \text{rk}_\star} \end{bmatrix} \right)^{-1}. \quad (\text{F.24})$$

- For $\alpha > 0$, we have

$$(\Gamma_t^{-1})_{11} \preceq \left(\mathbf{D}_{t,1} + \frac{1}{2} \mathbf{Z}_{1:\text{rk}_\star} \mathbf{Z}_{1:\text{rk}_\star}^\top \right)^{-1},$$

for $0.001 > \delta \geq \log^{-\frac{1}{4}} d$

$$\begin{cases} r_\star = \lfloor r_s(1 - \delta) \wedge r \rfloor \text{ and } \eta t \leq \frac{1}{2} \left(r_s(1 - \sqrt{\delta}) \wedge r \right)^\alpha \log\left(\frac{d \log^{1.5} d}{r_s}\right), & \alpha \in (0, 0.5) \\ r_{u_\star} = r_s \text{ and } \eta t \leq \frac{1}{2} r_s^\alpha \log\left(\frac{d \log^{1.5} d}{r_s}\right), & \alpha > 0.5 \end{cases}$$

provided that

$$\begin{cases} d \geq \Omega_\beta(1) \vee \exp(2.5\alpha^{-8}), & \alpha \in (0, 0.5) \\ d \geq \Omega_{r_s}(1), & \alpha > 0.5. \end{cases}$$

Proof. For the first part of the first item, by (R.2), we have

$$\mathbf{D}_t = \frac{\Lambda_{\ell_2}^{-1} \hat{\Lambda} \left(\frac{\mathbf{A}_{t,22}}{\mathbf{A}_{t,12}} - \mathbf{I}_{\text{rk}} \right)}{2} - \frac{\eta}{2} \Lambda_{\ell_2}^{-1} \hat{\Lambda}^2 - \frac{1.1 \kappa_d r_s}{d} \mathbf{I}_{\text{rk}} \stackrel{(a)}{\succeq} \frac{\Lambda_{\ell_2}^{-1} \hat{\Lambda} \left(\frac{\mathbf{A}_{t,22}}{\mathbf{A}_{t,12}} - \mathbf{I}_{\text{rk}} \right)}{2} - \frac{1.2 \kappa_d r_s}{d} \mathbf{I}_{\text{rk}},$$

where (a) follows $\eta \ll \frac{\kappa_d r_s}{d}$. Moreover, since $\Lambda_{\ell_2}^{-1} \hat{\Lambda} \succeq (1 - \frac{1}{\log^3 d}) \mathbf{I}_{\text{rk}}$, by (R.4), we have

$$\mathbf{D}_{t,2} \succeq \left(1 - \frac{1}{\log^3 d} \right) \left[\frac{\left(\mathbf{I}_{\text{rk} - \text{rk}_\star} - \eta \hat{\Lambda}_{22} \right)^t}{\left(\mathbf{I}_{\text{rk} - \text{rk}_\star} + \eta \hat{\Lambda}_{22} \right)^t - \left(\mathbf{I}_{\text{rk} - \text{rk}_\star} - \eta \hat{\Lambda}_{22} \right)^t} - \frac{1.3 \kappa_d r_s}{d} \mathbf{I}_{\text{rk} - \text{rk}_\star} \right] \quad (\text{F.25})$$

We observe that

$$\begin{aligned} \frac{\left(\mathbf{I}_{\text{rk} - \text{rk}_\star} - \eta \hat{\Lambda}_{22} \right)^t}{\left(\mathbf{I}_{\text{rk} - \text{rk}_\star} + \eta \hat{\Lambda}_{22} \right)^t - \left(\mathbf{I}_{\text{rk} - \text{rk}_\star} - \eta \hat{\Lambda}_{22} \right)^t} &\succeq \frac{\mathbf{I}_{\text{rk} - \text{rk}_\star}}{\exp\left(\frac{2t\eta \hat{\Lambda}_{22}}{1 - \eta \hat{\Lambda}_{22}}\right) - \mathbf{I}_{\text{rk} - \text{rk}_\star}} \\ &\stackrel{(b)}{\succeq} \left(\frac{r_s}{d \log^{1.5} d} \right)^{(1 + 2\eta \text{rk}_\star^{-\alpha}) \left(\frac{K+1}{\text{rk}_\star + 1} \right)^\alpha} \mathbf{I}_{\text{rk} - \text{rk}_\star} \\ &\stackrel{(c)}{\succeq} \left(\frac{0.9 r_s}{d \log^{1.5} d} \right)^{\left(\frac{K+1}{\text{rk}_\star + 1} \right)^\alpha} \mathbf{I}_{\text{rk} - \text{rk}_\star}, \end{aligned}$$

where we use $\eta t \leq \frac{1}{2}(K+1)^\alpha \log\left(\frac{d \log^{1.5} d}{r_s}\right)$ in (b) and $d \geq \Omega(1)$ in (c). By (F.25), the first item follows.

For the second item, by Proposition 24 with $\varepsilon = \frac{1}{\log^2 d}$ for $\alpha \in [0, 0.5)$ and $\varepsilon = \frac{1}{r_u \log^2 d}$ for $\alpha > 0.5$, we have

$$\mathbf{\Gamma}_t \succeq \frac{\Lambda_{\ell_2}^{-1} \hat{\mathbf{\Lambda}} \left(\frac{\mathbf{A}_{t,22}}{\mathbf{A}_{t,12}} - \mathbf{I}_{rk} \right)}{2} - \frac{1.2\kappa_d r_s}{d} \mathbf{I}_{rk} + \frac{1}{1.05} \begin{bmatrix} \varepsilon \mathbf{Z}_{1:rk_*} \mathbf{Z}_{1:rk_*}^\top & 0 \\ 0 & \frac{-\varepsilon}{1-\varepsilon} \mathbf{Z}_2 \mathbf{Z}_2^\top \end{bmatrix}.$$

For $\alpha \in [0, 0.5)$, since $\kappa_d = \frac{1}{\log^{3.5} d}$, by (H.1), we have

$$\frac{\varepsilon}{1.05} \mathbf{Z}_{1:r_*} \mathbf{Z}_{1:r_*}^\top - \frac{1.2\kappa_d r_s}{d} \mathbf{I}_{r_*} \succeq \frac{r_s}{d \log^3 d} \frac{1}{6.25} \mathbf{I}_{r_*} - \frac{1.2\kappa_d r_s}{d} \mathbf{I}_{r_*} \succ \frac{1}{10} \frac{r_s}{d \log^3 d} \mathbf{I}_{r_*}.$$

Similarly by (H.3), we have

$$\frac{-\varepsilon}{1-\varepsilon} \frac{1}{1.05} \mathbf{Z}_2 \mathbf{Z}_2^\top - \frac{1.2\kappa_d r_s}{d} \mathbf{I}_{r-r_*} \succeq - \left(1 + \frac{1}{\sqrt{\varphi}} \right)^2 \frac{2.1r_s}{d \log^2 d} \mathbf{I}_{r-r_*}.$$

For $\alpha > 0.5$, since $\kappa_d = \frac{1}{r_u \log^{2.5} d}$ and $r_u = \lceil \log^{2.5} d \rceil$, by (L.1), we have

$$\frac{\varepsilon}{1.05} \mathbf{Z}_{1:r_{u_*}} \mathbf{Z}_{1:r_{u_*}}^\top - \frac{1.2\kappa_d r_s}{d} \mathbf{I}_{r_{u_*}} \succ \frac{r_s}{d \log^{4.5} d} \left(\frac{1}{1.1r_s^6} - \frac{1.3}{\sqrt{\log d}} \right) \mathbf{I}_{r_{u_*}}.$$

Similarly by (L.2),

$$\frac{-\varepsilon}{1-\varepsilon} \frac{1}{1.05d} \mathbf{Z}_2 \mathbf{Z}_2^\top - \frac{1.2\kappa_d r_s}{d} \mathbf{I}_{r_u-r_{u_*}} \succeq \frac{-2r_s}{d \log^2 d} \mathbf{I}_{r_u-r_{u_*}}.$$

Therefore, we have (F.23). By Proposition 25, we have (F.24).

For the last item, we have

$$(\mathbf{\Gamma}_t^{-1})_{11} = \left(\mathbf{D}_{t,1} + \frac{1}{1.05} \mathbf{Z}_{1:rk_*} \left(\mathbf{I}_{r_s} + \frac{1}{1.05} \mathbf{Z}_2^\top \mathbf{D}_{t,2}^{-1} \mathbf{Z}_2 \right)^{-1} \mathbf{Z}_{1:rk_*}^\top \right)^{-1}. \quad (\text{F.26})$$

For $\alpha \in (0, 0.5)$, if $r_* = r$, the statement follows. If not by the first item, for $K = \lfloor (1 - \sqrt{\delta}) r_s \rfloor$ and $r_* = \lfloor r_s (1 - \delta) \rfloor$, we have

$$\frac{K+1}{r_*+1} \leq \frac{1-\sqrt{\delta}}{1-\delta} + \frac{2}{r_s} \leq 1 - 0.9\sqrt{\delta} \Rightarrow \left(\frac{K+1}{r_*+1} \right)^\alpha \leq 1 - \alpha 0.9\sqrt{\delta}.$$

Therefore,

$$\mathbf{D}_{t,2} \succeq \left(\frac{0.5r_s}{d \log^{1.5} d} \right)^{1-\alpha 0.9\sqrt{\delta}} \mathbf{I}_{r-r_*} \stackrel{(d)}{\succeq} \frac{0.5r_s}{d \log^{1.5} d} \left(\frac{d}{r_s} \right)^{\log^{-1/4} d} \mathbf{I}_{r-r_*} \stackrel{(e)}{\succeq} \frac{r_s \log d}{d} \mathbf{I}_{r-r_*},$$

where we used $d \geq \Omega(1) \vee \exp(2.5\alpha^{-8})$ in (d) and $d \geq \Omega_\beta(1)$ in (e). By (F.26) and (H.3), we have the statement for $\alpha \in (0, 0.5)$. For $\alpha > 0.5$, $K = r_* = r_s$, we have

$$\left(\frac{K+1}{r_*+1} \right)^\alpha \leq \left(1 + \frac{1}{r_s+1} \right)^{0.5} \leq 1 - \frac{1}{2(r_s+1)}.$$

Therefore,

$$\mathbf{D}_{t,2} \succeq \left(\frac{0.5r_s}{d \log^{1.5} d} \right)^{1-\frac{1}{2(r_s+1)}} \mathbf{I}_{r_u-r_{u_*}} \succeq \frac{r_s \log^8 d}{d} \mathbf{I}_{r-r_*}$$

for $d \geq \Omega_{r_s}(1)$. By (F.26) and (L.2), we have the statement for $\alpha > 0.5$. $\square$

Proposition 15. *Let*

$$\underline{\mathbf{G}}_0 = (1 + 2\kappa_d) \left(\mathbf{G}_0 - \frac{\kappa_d r_s}{d} \mathbf{I}_{rk} \right),$$

Under the parameter choice in Lemma 5, $\mathcal{G}_{\text{init}}$ guarantees that:

- We have $\Omega(-\log^{\frac{-1}{2}} d) \mathbf{I}_{\text{rk}} \preceq \underline{\mathbf{G}}_t$ whenever

$$\eta t \leq \begin{cases} \frac{1}{2} \left(r_s (1 - \log^{\frac{-1}{2}} d) \wedge r \right)^\alpha \log \left(\frac{d \log^{1.5} d}{r_s} \right), & \alpha \in [0, 0.5) \\ \frac{1}{2} r_s^\alpha \log \left(\frac{d \log^{1.5} d}{r_s} \right), & \alpha > 0.5 \end{cases}.$$

- Let Λ_{11} be the $\text{rk}_* \times \text{rk}_*$ dimensional top-left sub-matrix of Λ . Given $0.001 \geq \delta \geq \log^{\frac{-1}{4}} d$ and $\text{rk}_* = \{r_* = \lfloor r_s (1 - \delta) \wedge r \rfloor, r_{u_*} = r_s\}$, we have

$$\underline{\mathbf{G}}_{t,11} \succeq \frac{1 - \frac{10}{\log^3 d}}{\frac{1.2}{C_{lb}} \frac{d}{r_s} \exp(-2\eta t \Lambda_{11}) + 1}$$

and

$$\|\hat{\Lambda}\|_F^2 - \|\hat{\Lambda}_1^{\frac{1}{2}} \underline{\mathbf{G}}_{t,11} \hat{\Lambda}_1^{\frac{1}{2}}\|_F^2 \leq \sum_{i=(\text{rk}_* \wedge r)+1}^r \hat{\lambda}_i^2 + \sum_{i=1}^{\text{rk}_*} \hat{\lambda}_i^2 \left(1 - \frac{1 - \frac{10}{\log^3 d}}{\frac{1.2}{C_{lb}} \frac{d}{r_s} \exp(-2\eta t \lambda_i) + 1} \right)^2,$$

for

$$C_{lb} = \frac{1}{15} \begin{cases} \delta^2, & \alpha \in [0, 0.5) \\ \frac{1}{r_s^6}, & \alpha > 0.5 \end{cases} \quad \text{and} \quad \eta t \leq \begin{cases} \frac{1}{2} \left(r_s (1 - \sqrt{\delta}) \wedge r \right)^\alpha \log \left(\frac{d \log^{1.5} d}{r_s} \right), & \alpha \in [0, 0.5) \\ \frac{1}{2} r_s^\alpha \log \left(\frac{d \log^{1.5} d}{r_s} \right), & \alpha > 0.5. \end{cases}$$

- For $\delta = \log^{\frac{-1}{4}} d$, we define

$$\mathcal{T}_{lb} := \inf \left\{ n \geq 0 \mid \|\hat{\Lambda}\|_F^2 - \|\hat{\Lambda}_1^{\frac{1}{2}} \underline{\mathbf{G}}_{t,11} \hat{\Lambda}_1^{\frac{1}{2}}\|_F^2 \leq \sum_{j=(r_s \wedge r)+1}^r \lambda_j^2 + \frac{3\|\hat{\Lambda}\|_F^2}{\log^{\frac{1}{8}} d} \right\}.$$

Then,

$$\mathcal{T}_{lb} \leq \begin{cases} \frac{1}{2\eta} \left(r_s (1 - \log^{\frac{-1}{8}} d) \wedge r \right)^\alpha \log \left(\frac{20d \log^{\frac{3}{4}}(1+d/r_s)}{r_s} \right), & \alpha \in [0, 0.5) \\ \frac{1}{2\eta} r_s^\alpha \log \left(\frac{20d \log^{\frac{3}{4}} d}{r_s} \right), & \alpha > 0.5. \end{cases}$$

Proof. For $\alpha > 0.5$, we assume that d is large enough to guarantee that $\left(\frac{1}{1.1r_s^6} - \frac{1.3}{\sqrt{\log d}} \right) > 0$. We observe that

$$\frac{\Lambda_{\ell_2}^{-1} \hat{\Lambda} \left(\frac{\mathbf{A}_{t,11}}{\mathbf{A}_{t,12}} - \mathbf{I}_{\text{rk}} \right)}{2} + \frac{\Lambda_{\ell_2}^{-1} (\hat{\Lambda} - \Lambda_{\ell_1})}{2} + \underline{\mathbf{G}}_0 \succeq \mathbf{\Gamma}_t,$$

where we used (E.1), $\Lambda_{\ell_2} \succeq \Lambda_{\ell_1}$ and $\eta \|\Lambda\|_F^2 r_s \mathbf{I}_{\text{rk}} \ll \frac{\kappa_d r_s}{d} \Lambda_{\ell_1}$.

For the first item, by using $\text{rk}_* = \{r_* = \lfloor r_s (1 - \log^{\frac{-1}{2}} d) \wedge r \rfloor, r_{u_*} = r_s\}$, we define

$$\mathbf{D}_{lb} := \begin{bmatrix} \frac{C_1 r_s}{d \log^{4.5} d} \mathbf{I}_{\text{rk}_*} & 0 \\ 0 & \frac{-C_2 r_s}{d \log^2 d} \mathbf{I}_{\text{rk} - \text{rk}_*} \end{bmatrix} \quad \text{and} \quad \tilde{\mathbf{D}}_{lb} := \frac{\Lambda_{\ell_2}}{\hat{\Lambda}} \mathbf{D}_{lb}.$$

We introduce submatrix notation for block-diagonal matrices. Specifically, we write

$$\tilde{\mathbf{D}}_{lb} = \begin{bmatrix} \tilde{\mathbf{D}}_{lb,1} & 0 \\ 0 & \tilde{\mathbf{D}}_{lb,2} \end{bmatrix} \quad \text{and} \quad \frac{\mathbf{A}_{t,22}}{\mathbf{A}_{t,12}} \pm \mathbf{I}_{\text{rk}} = \begin{bmatrix} \left(\frac{\mathbf{A}_{t,22}}{\mathbf{A}_{t,12}} \pm \mathbf{I}_{\text{rk}} \right)_{11} & 0 \\ 0 & \left(\frac{\mathbf{A}_{t,22}}{\mathbf{A}_{t,12}} \pm \mathbf{I}_{\text{rk}} \right)_{22} \end{bmatrix},$$

where the block dimensions of each submatrix match those of $\mathbf{D}_{lb}$. We start with proving the lower bound part. By the second item in Lemma 5, we have

$$\underline{\mathbf{G}}_t \succeq \frac{1}{2} \sqrt{\frac{\hat{\Lambda}}{\Lambda_{\ell_2}}} \left(\left(\frac{\mathbf{A}_{t,22}}{\mathbf{A}_{t,12}} + \mathbf{I}_{\text{rk}} \right) - \mathbf{A}_{t,12}^{-1} \left(\left(\frac{\mathbf{A}_{t,22}}{\mathbf{A}_{t,12}} - \mathbf{I}_{\text{rk}} \right) + 2\tilde{\mathbf{D}}_{lb} \right)^{-1} \mathbf{A}_{t,12}^{-1} \right) \sqrt{\frac{\hat{\Lambda}}{\Lambda_{\ell_2}}}$$

$$\begin{aligned}
& + \frac{\frac{\Lambda_{\ell_1}}{\Lambda_{\ell_2}} - \frac{\hat{\Lambda}}{\Lambda_{\ell_2}}}{2} \\
& \succeq \frac{1}{2} \sqrt{\frac{\hat{\Lambda}}{\Lambda_{\ell_2}}} \left(\left(\frac{\mathbf{A}_{t,22}}{\mathbf{A}_{t,12}} + \mathbf{I}_{rk} \right) - \mathbf{A}_{t,12}^{-1} \left(\left(\frac{\mathbf{A}_{t,22}}{\mathbf{A}_{t,12}} - \mathbf{I}_{rk} \right) + 2\tilde{\mathbf{D}}_{lb} \right)^{-1} \mathbf{A}_{t,12}^{-1} \right) \sqrt{\frac{\hat{\Lambda}}{\Lambda_{\ell_2}}} \quad (\text{F.27})
\end{aligned}$$

where we used $\Lambda_{\ell_1} \succ \hat{\Lambda}$ in the second line. We have

$$\begin{aligned}
& \left(\frac{\mathbf{A}_{t,22}}{\mathbf{A}_{t,12}} + \mathbf{I}_{rk} \right) - \mathbf{A}_{t,12}^{-1} \left(\left(\frac{\mathbf{A}_{t,22}}{\mathbf{A}_{t,12}} - \mathbf{I}_{rk} \right) + 2\tilde{\mathbf{D}}_{lb} \right)^{-1} \mathbf{A}_{t,12}^{-1} \\
& = \frac{(\mathbf{A}_{t,22} + \mathbf{A}_{t,12})(\mathbf{A}_{t,22} - \mathbf{A}_{t,12} + 2\tilde{\mathbf{D}}_{lb}\mathbf{A}_{t,12}) - \mathbf{I}_{rk}}{\mathbf{A}_{t,12}(\mathbf{A}_{t,22} - \mathbf{A}_{t,12} + 2\tilde{\mathbf{D}}_{lb}\mathbf{A}_{t,12})} \\
& \stackrel{(a)}{\succeq} \frac{2\tilde{\mathbf{D}}_{lb} \left(\frac{\mathbf{A}_{t,22}}{\mathbf{A}_{t,12}} + \mathbf{I}_{rk} \right)}{\frac{\mathbf{A}_{t,22}}{\mathbf{A}_{t,12}} - \mathbf{I}_{rk} + 2\tilde{\mathbf{D}}_{lb}} \\
& = \begin{bmatrix} \frac{2\tilde{\mathbf{D}}_{lb,1} \left(\frac{\mathbf{A}_{t,22}}{\mathbf{A}_{t,12}} + \mathbf{I}_{rk} \right)_{11}}{\left(\frac{\mathbf{A}_{t,22}}{\mathbf{A}_{t,12}} - \mathbf{I}_{rk} \right)_{11} + 2\tilde{\mathbf{D}}_{lb,1}} & 0 \\ 0 & \frac{2\tilde{\mathbf{D}}_{lb,2} \left(\frac{\mathbf{A}_{t,22}}{\mathbf{A}_{t,12}} + \mathbf{I}_{rk} \right)_{22}}{\left(\frac{\mathbf{A}_{t,22}}{\mathbf{A}_{t,12}} - \mathbf{I}_{rk} \right)_{22} + 2\tilde{\mathbf{D}}_{lb,2}} \end{bmatrix}, \quad (\text{F.28})
\end{aligned}$$

where we used $\mathbf{A}_{t,22}^2 - \mathbf{A}_{t,12}^2 \succ \mathbf{I}_{rk}$ (by (R.2)) and $\mathbf{A}_{t,22} - \mathbf{A}_{t,12} + 2\tilde{\mathbf{D}}_{lb}\mathbf{A}_{t,12} \succ 0$ (by (F.23)) in (a). Since $\frac{\mathbf{A}_{t,22}}{\mathbf{A}_{t,12}} - \mathbf{I}_{rk} \succ 0$ and $\tilde{\mathbf{D}}_{lb,1} \succ 0$, it is enough to look at the bottom-right submatrix in (F.28) for the lower bound part. We have

$$\frac{2\tilde{\mathbf{D}}_{lb,2} \left(\frac{\mathbf{A}_{t,22}}{\mathbf{A}_{t,12}} + \mathbf{I}_{rk} \right)_{22}}{\left(\frac{\mathbf{A}_{t,22}}{\mathbf{A}_{t,12}} - \mathbf{I}_{rk} \right)_{22} + 2\tilde{\mathbf{D}}_{lb,2}} = \frac{2\tilde{\mathbf{D}}_{lb,2} \left(\frac{\mathbf{A}_{t,22}}{\mathbf{A}_{t,12}} + \mathbf{I}_{rk} \right)_{22}}{\left(\frac{\mathbf{A}_{t,22}}{\mathbf{A}_{t,12}} + \mathbf{I}_{rk} \right)_{22} - 2\mathbf{I}_{rk-rk_*} + 2\tilde{\mathbf{D}}_{lb,2}}. \quad (\text{F.29})$$

Note that by (R.4),

$$\begin{aligned}
& \left(\frac{\mathbf{A}_{t,22}}{\mathbf{A}_{t,12}} + \mathbf{I}_{rk} \right)_2 \succeq \frac{2(\mathbf{I}_{rk-rk_*} + \eta\hat{\Lambda}_2)^t}{(\mathbf{I}_{rk-rk_*} + \eta\hat{\Lambda}_2)^t - (\mathbf{I}_{rk-rk_*} - \eta\hat{\Lambda}_2)^t} \\
& \succeq 2\mathbf{I}_{rk-rk_*} + \frac{2(\mathbf{I}_{rk-rk_*} - \eta^2\hat{\Lambda}_2^2)^t \exp(-2t\eta\hat{\Lambda}_2)}{\mathbf{I}_{rk-rk_*} - (\mathbf{I}_{rk-rk_*} - \eta^2\hat{\Lambda}_2^2)^t \exp(-2t\eta\hat{\Lambda}_2)} \\
& \stackrel{(b)}{\succeq} \left(2 + \frac{0.9r_s}{d \log^{1.5} d} \right) \mathbf{I}_{rk-rk_*},
\end{aligned}$$

where we use $\eta t \leq \frac{1}{2}(rk_* + 1)^\alpha \log \left(\frac{d \log^{1.5} d}{r_s} \right)$ in (b). Hence, for $d \geq \Omega(1)$

$$(\text{F.29}) \succeq \frac{2 \left(2 + \frac{0.9r_s}{d \log^{1.5} d} \right) \tilde{\mathbf{D}}_{lb,2}}{\frac{0.9r_s}{d \log^{1.5} d} \mathbf{I}_{rk-rk_*} + \tilde{\mathbf{D}}_{lb,2}} \stackrel{(c)}{\succeq} \frac{12\tilde{\mathbf{D}}_{lb,2}}{\frac{r_s}{d \log^{1.5} d} \mathbf{I}_{rk-rk_*}} \stackrel{(d)}{\succeq} \frac{-15C_2}{\log^{0.5} d} \mathbf{I}_{rk-rk_*}, \quad (\text{F.30})$$

where we used $\tilde{\mathbf{D}}_{lb,2} \succeq \frac{-1.1C_2r_s}{d \log^2 d} \mathbf{I}_{rk}$ in (c) and (d). The first item follows from (F.30).

For the second and third items, let $\left(\frac{\mathbf{A}_{t,22}}{\mathbf{A}_{t,12}} \pm \mathbf{I}_{rk} \right)_{11}$ denote the $rk_* \times rk_*$ dimensional top-left submatrices with $rk_* = \{r_* = \lfloor r_s(1 - \delta) \wedge r \rfloor, r_{u_*} = r_s\}$. By using the third item in Lemma 5, we immediately observe that for $\alpha > 0$, $\underline{\mathbf{G}}_{t,11} \succeq 0$ and

$$\underline{\mathbf{G}}_{t,11} \stackrel{(e)}{\succeq} \left(1 - \frac{10}{\log^3 d} \right) \frac{1}{2} \frac{\frac{2C_{lb}r_s}{d} \left(\frac{\mathbf{A}_{t,22}}{\mathbf{A}_{t,12}} + \mathbf{I}_{rk} \right)_{11}}{\left(\frac{\mathbf{A}_{t,11}}{\mathbf{A}_{t,12}} - \mathbf{I}_{rk} \right)_{11} + \frac{2C_{lb}r_s}{d} \mathbf{I}_{rk_*}}, \quad (\text{F.31})$$

for

$$C_{\text{lb}} = \frac{1}{15} \begin{cases} \delta^2, & \alpha \in (0, 0.5) \\ \frac{1}{r_s^6}, & \alpha > 0.5 \end{cases} \text{ and } \eta t \leq \begin{cases} \frac{1}{2} \left(r_s (1 - \sqrt{\delta}) \wedge r \right)^\alpha \log \left(\frac{d \log^{1.5} d}{r_s} \right), & \alpha \in (0, 0.5) \\ \frac{1}{2} r_s^\alpha \log \left(\frac{d \log^{1.5} d}{r_s} \right), & \alpha > 0.5, \end{cases}$$

where we used $\Lambda_{\ell_2} \succeq \hat{\Lambda} \succeq \left(1 - \frac{0.5}{\log^4 d}\right) \Lambda_{\ell_2}$, and followed the steps in (F.27)-(F.28) with (H.1) and (L.1) to obtain (e). Then, by (R.4), we have

$$\begin{aligned} \frac{1}{2} \frac{\frac{2C_{\text{lb}} r_s}{d} \left(\frac{\mathbf{A}_{t,22}}{\mathbf{A}_{t,12}} + \mathbf{I}_{\text{rk}} \right)_{11}}{\left(\frac{\mathbf{A}_{t,11}}{\mathbf{A}_{t,12}} - \mathbf{I}_{\text{rk}} \right)_{11} + \frac{2C_{\text{lb}} r_s}{d} \mathbf{I}_{\text{rk}_*}} &\succeq \frac{\mathbf{I}_{\text{rk}_*}}{\left(\frac{1}{C_{\text{lb}}} \frac{d}{r_s} - 1 \right) \frac{(\mathbf{I}_{\text{rk}_*} - \eta \hat{\Lambda}_1)^t}{(\mathbf{I}_{\text{rk}_*} + \eta \hat{\Lambda}_1)^t} + \mathbf{I}_{\text{rk}_*}} \\ &\succeq \frac{\mathbf{I}_{\text{rk}_*}}{\frac{1}{C_{\text{lb}}} \frac{d}{r_s} \exp(-2\eta t \hat{\Lambda}_1) + \mathbf{I}_{\text{rk}_*}}. \end{aligned}$$

Consequently, by observing $\hat{\Lambda} \succeq \Lambda_{\ell_1} - \tilde{C}\eta \mathbf{I}_{\text{rk}}$ and using the lower bounds for Λ_{ℓ_1} in Propositions 12 and 13, we have

$$\underline{\mathbf{G}}_{t,11} \succeq \frac{1 - \frac{10}{\log^3 d}}{\frac{1.2}{C_{\text{lb}}} \frac{d}{r_s} \exp(-2\eta t \Lambda_{11}) + 1},$$

where Λ_{11} denotes the $\text{rk}_* \times \text{rk}_*$ dimensional top-left sub-matrix of Λ . Therefore,

$$\|\hat{\Lambda}\|_F^2 - \|\hat{\Lambda}_1^{\frac{1}{2}} \underline{\mathbf{G}}_{t,11} \hat{\Lambda}_1^{\frac{1}{2}}\|_F^2 \leq \sum_{i=(\text{rk}_* \wedge r)+1}^r \hat{\lambda}_i^2 + \sum_{i=1}^{\text{rk}_*} \hat{\lambda}_i^2 \left(1 - \frac{1 - \frac{10}{\log^3 d}}{\frac{1.2}{C_{\text{lb}}} \frac{d}{r_s} \exp(-2\eta t \lambda_i) + 1} \right)^2, \quad (\text{F.32})$$

which proves the second item for $\alpha > 0$. Moreover, since (F.22) is in the eigenbasis of $\underline{\mathbf{G}}_0$, the arguments in (F.27)-(F.28) and the condition in (H.2) extend (F.31) to $\alpha = 0$ in the eigenbasis of $\underline{\mathbf{G}}_0$ for $d \geq \Omega(1)$. Given (F.31), we can extend (F.32) to $\alpha = 0$ as the Frobenious norm is basis independent.

For the third item, for $\alpha > 0.5$ and $t \geq \frac{1}{2\eta} r_s^\alpha \log \left(\frac{20d \log^{\frac{3}{2}} d}{r_s} \right)$, we have

$$(\text{F.32}) \leq \sum_{i=(r_s \wedge r)+1}^r \hat{\lambda}_i^2 + \frac{\|\hat{\Lambda}\|_F^2}{\log^{\frac{1}{2}} d} \leq \sum_{i=(r_s \wedge r)+1}^r \lambda_i^2 + \frac{\|\hat{\Lambda}\|_F^2}{\log^{\frac{1}{2}} d},$$

which gives us the corresponding bound for $\mathcal{T}_{\text{lb}}$.

For $\alpha \in [0, 0.5)$ and $t \geq \frac{1}{2\eta} (r_s (1 - \log^{-\frac{1}{8}} d) \wedge r)^\alpha \log \left(\frac{20d \log^{\frac{3}{4}} (1+d/r_s)}{r_s} \right)$, we have

$$\begin{aligned} (\text{F.32}) &\leq \sum_{i=(r_s \wedge r)+1}^r \hat{\lambda}_i^2 + \sum_{i=\lfloor r_s (1 - \log^{-\frac{1}{8}} d) \wedge r \rfloor + 1}^{r_s \wedge r} \hat{\lambda}_i^2 \\ &\quad + \sum_{i=1}^{\lfloor r_s (1 - \log^{-\frac{1}{8}} d) \wedge r \rfloor} \hat{\lambda}_i^2 \left(1 - \frac{1 - \frac{10}{\log^3 d}}{\frac{1.2}{C_{\text{lb}}} \frac{d}{r_s} \exp(-2\eta t \lambda_i) + 1} \right)^2 \\ &\leq \sum_{i=(r_s \wedge r)+1}^r \lambda_i^2 + \frac{3\|\hat{\Lambda}\|_F^2}{\log^{\frac{1}{8}} d}, \end{aligned}$$

which gives us its bound for $\mathcal{T}_{\text{lb}}$. □

F.4.2 Upper bounding system

In this section, we consider (F.13). For notational convenience, we multiply both sides by the factor $(1 - 2\kappa_d)$ and use a generic learning rate η , i.e.,

$$\bar{\mathbf{V}}_{t+1} = \bar{\mathbf{V}}_t (\mathbf{I}_{\text{rk}} + \eta \bar{\mathbf{V}}_t)^{-1} + \eta \left(\Lambda_{u_1}^2 + \tilde{C}\eta \|\Lambda\|_F^2 r_s \Lambda_{u_1} \right), \text{ where } \bar{\mathbf{V}}_t = 2\Lambda_{u_2}^{\frac{1}{2}} \bar{\mathbf{G}}_t \Lambda_{u_2}^{\frac{1}{2}} - \Lambda_{u_1}.$$

The main result of this section is stated in Proposition 16. To establish it, we first prove an auxiliary result:

Lemma 6. *The following statement holds:*

- The reference sequence satisfies $\mathbf{T}_t \succeq \frac{\kappa_d r_s}{d} \mathbf{I}_{rk}$ and $\{t \geq 0 : \|\mathbf{T}_t\|_2 > 1.2\kappa_d\} = \infty$.
- For $r_{u_*} = 2r_s$, we have

$$\begin{cases} \bar{\mathbf{G}}_0 = \frac{2.2(1+\frac{1}{\sqrt{\varphi}})^2 r_s}{d} \mathbf{I}_r \succeq (1 - 2\kappa_d) (\mathbf{G}_0 + \frac{\kappa_d r_s}{d} \mathbf{I}_r), & \alpha \in [0, 0.5) \\ \bar{\mathbf{G}}_0 = \frac{5.5}{d} \begin{bmatrix} 2r_s \mathbf{I}_{r_{u_*}} & 0 \\ 0 & r_u \mathbf{I}_{r-r_{u_*}} \end{bmatrix} \succeq (1 - 2\kappa_d) (\mathbf{G}_{0,11} + \frac{\kappa_d r_s}{d} \mathbf{I}_{r_u}), & \alpha > 0.5 \end{cases} \quad (\text{F.33})$$

provided that $\mathcal{G}_{\text{init}}$ holds.

For the following, we introduce $\hat{\mathbf{T}}_t := \frac{\Lambda_{u_2}}{\Lambda_{\ell_1}} \frac{(3\kappa_d+1)}{\kappa_d(1-\kappa_d)} \mathbf{T}_t$. Note that for $d \geq \Omega(1)$, we have

$$\hat{\mathbf{T}}_{t+1} = \hat{\mathbf{T}}_t + 2(1 - 2\kappa_d)\eta \Lambda_{\ell_1} \hat{\mathbf{T}}_t (\mathbf{I}_{rk} - \hat{\mathbf{T}}_t) \quad \text{and} \quad \frac{\kappa_d}{1.1} \frac{\Lambda_{\ell_1}}{\Lambda_{u_2}} \preceq \frac{\mathbf{T}_t}{\hat{\mathbf{T}}_t} \preceq \kappa_d \mathbf{I}_{rk}. \quad (\text{F.34})$$

By Proposition 34, we have

$$\begin{aligned} 1.1 \wedge \hat{\mathbf{T}}_{0,ii} \exp(2\eta t \lambda_i) &\geq \hat{\mathbf{T}}_{t,ii} \\ &\geq \frac{1}{2} \begin{cases} 1 \wedge \hat{\mathbf{T}}_{0,ii} \exp\left(\frac{(1-2\kappa_d)2\eta t \left(\lambda_i - \frac{0.1r^{-\alpha}}{\log^4 d}\right)}{1+2(1-2\kappa_d)\eta \lambda_i}\right), & \alpha \in [0, 0.5) \\ 1 \wedge \hat{\mathbf{T}}_{0,ii} \exp\left(\frac{(1-2\kappa_d)2\eta t \left(\lambda_i - \frac{1}{(r_u+1)^\alpha} - \frac{0.1}{r_u^{2+\alpha} \log^4 d}\right)}{1+2(1-2\kappa_d)\eta \lambda_i}\right), & \alpha > 0.5. \end{cases} \end{aligned} \quad (\text{F.35})$$

Proof of Lemma 6. For the first item, by Proposition 34, we have

$$\hat{\mathbf{T}}_t \succeq \hat{\mathbf{T}}_0 \stackrel{(a)}{\Rightarrow} \mathbf{T}_t \succeq \mathbf{T}_0 = \frac{\kappa_d r_s}{d} \mathbf{I}_{rk},$$

where we multiplied each side with $\frac{\kappa_d(1-\kappa_d)}{3\kappa_d+1} \frac{\Lambda_{\ell_1}}{\Lambda_{u_2}}$ for (a). Moreover, by (F.34)-(F.35), we have

$$\mathbf{T}_t \preceq \kappa_d \hat{\mathbf{T}}_t \preceq 1.1 \mathbf{I}_{rk} \Rightarrow \{t \geq 0 : \|\mathbf{T}_t\|_2 > 1.2\kappa_d\} = \infty.$$

The second item follows (E.2) and (H.3) (for $\alpha \in [0, 0.5)$) and (L.2) (for $\alpha > 0.5$). $\square$

Proposition 16. *We consider $rk \in \{r, r_u\}$, where $r_u = \lceil \log^{2.5} d \rceil$, and*

$$\begin{aligned} \alpha \in [0, 0.5) : \quad & \frac{r_s}{r} \rightarrow \varphi \in (0, \infty), \quad \eta \ll \frac{1}{d r^{1-\alpha} \log^4 d}, \quad \kappa_d = \frac{1}{\log^{3.5} d}, \\ \alpha > 0.5 : \quad & r_s \asymp 1, \quad \eta \ll \frac{1}{d r_u^{2+\alpha} \log^3 d}, \quad \kappa_d = \frac{1}{r_u \log^{2.5} d}. \end{aligned}$$

If $\bar{\mathbf{G}}_0$ are taken as in (F.33), we have the following:

- $\{\bar{\mathbf{G}}_t\}_{n \in \mathbb{N}}$ is diagonal and satisfies

$$\frac{r_s}{d} \mathbf{I}_{rk} \preceq \bar{\mathbf{G}}_{t+1} \preceq \bar{\mathbf{G}}_t + \eta((1 + \kappa_d) \Lambda_{u_1} \bar{\mathbf{G}}_t + (1 + \kappa_d) \bar{\mathbf{G}}_t \Lambda_{u_1} - 2\bar{\mathbf{G}}_t \Lambda_{u_2} \bar{\mathbf{G}}_t) \preceq 1.1 \mathbf{I}_{rk}.$$

- For $\alpha \in [0, 0.5)$ and $d \geq \Omega(1)$, we have for $t \leq \frac{1}{2\eta} r^\alpha \log\left(\frac{d \log^{1.5} d}{r_s}\right)$:

$$\begin{aligned} - \mathbf{T}_t^{-\frac{1}{2}} \bar{\mathbf{G}}_j \mathbf{T}_t^{-\frac{1}{2}} &\preceq \frac{11(1+\frac{1}{\sqrt{\varphi}})^2}{\kappa_d} \mathbf{I}_r \text{ for } 0 \leq j \leq t. \\ - \mathbf{T}_t^{-\frac{1}{2}} \left(\eta \sum_{j=1}^t \bar{\mathbf{G}}_{j-1}\right) \mathbf{T}_t^{-\frac{1}{2}} &\preceq \frac{5.5(1+\frac{1}{\sqrt{\varphi}})^2}{\kappa_d} (2\eta t \vee r^\alpha) \mathbf{I}_r. \\ - \bar{\mathbf{G}}_t &\preceq (1.1 \bar{\mathbf{G}}_0 \exp(2\eta t \Lambda) \wedge \mathbf{I}_r) + o_d(1) \end{aligned}$$

$$- \|\Lambda\|_F^2 - \text{Tr}(\Lambda \bar{G}_t \Lambda) \geq \sum_{i=1}^r \lambda_i^2 \left(1 - \frac{2.5 \left(1 + \frac{1}{\sqrt{\varphi}} \right)^2 r_s}{d} \exp(2\eta t \lambda_i) \right)_+ - o_d(1).$$

• For $\alpha > 0.5$ and $d \geq \Omega_{r_s}(1)$, we have for $t \leq \frac{1}{2\eta} r_s^\alpha \log \left(\frac{d \log^{1.5} d}{r_s} \right)$:

$$\begin{aligned} - \mathbf{T}_t^{-\frac{1}{2}} \bar{\mathbf{G}}_j \mathbf{T}_t^{-\frac{1}{2}} &\preceq \frac{26.4 r_u}{\kappa_d} \mathbf{I}_{r_u} \text{ for } 0 \leq j \leq t. \\ - \mathbf{T}_t^{-\frac{1}{2}} \left(\eta \sum_{j=1}^t \bar{\mathbf{G}}_{j-1} \right) \mathbf{T}_t^{-\frac{1}{2}} &\preceq \frac{15 r_u}{\kappa_d} (2r_s)^\alpha \log d \mathbf{I}_{r_u}. \\ - \bar{\mathbf{G}}_t &\preceq (1.1 \bar{\mathbf{G}}_0 \exp(2\eta t \Lambda_{11}) \wedge \mathbf{I}_{r_u}) + o_d(1). \\ - \|\Lambda_{11}\|_F^2 - \text{Tr}(\Lambda_{11} \bar{\mathbf{G}}_{t,11} \Lambda_{11}) &\geq \sum_{i=1}^{r_s} \lambda_i^2 \left(1 - \frac{12.1 r_s}{d} \exp(2\eta t \lambda_i) \right)_+ + \sum_{i=r_s+1}^{r_u} \lambda_i^2 - o_d(1). \end{aligned}$$

Proof. Given that $\frac{r_s}{d} \mathbf{I}_{rk} \preceq \bar{\mathbf{G}}_t \preceq 1.1 \mathbf{I}_{rk}$, we have

$$\begin{aligned} \bar{\mathbf{G}}_{t+1} &\stackrel{(a)}{\preceq} \bar{\mathbf{G}}_t + \eta (\Lambda_{u_1} \bar{\mathbf{G}}_t + \bar{\mathbf{G}}_t \Lambda_{u_1} - 2 \bar{\mathbf{G}}_t \Lambda_{u_2} \bar{\mathbf{G}}_t) + 1.1 \tilde{C} \eta^2 \|\Lambda\|_F^2 r_s \mathbf{I}_{rk} \\ &\stackrel{(b)}{\preceq} \bar{\mathbf{G}}_t + \eta ((1 + \kappa_d) \Lambda_{u_1} \bar{\mathbf{G}}_t + (1 + \kappa_d) \bar{\mathbf{G}}_t \Lambda_{u_1} - 2 \bar{\mathbf{G}}_t \Lambda_{u_2} \bar{\mathbf{G}}_t), \\ \bar{\mathbf{G}}_{t+1} &\stackrel{(c)}{\succeq} \bar{\mathbf{G}}_t + \eta (\Lambda_{u_1} \bar{\mathbf{G}}_t + \bar{\mathbf{G}}_t \Lambda_{u_1} - 2 \bar{\mathbf{G}}_t \Lambda_{u_2} \bar{\mathbf{G}}_t) - 1.1 \tilde{C} \|\Lambda\|_F^2 r_s \eta^2 \mathbf{I}_{rk} \\ &\stackrel{(d)}{\succeq} \bar{\mathbf{G}}_t + \eta ((1 - \kappa_d) \Lambda_{u_1} \bar{\mathbf{G}}_t + (1 - \kappa_d) \bar{\mathbf{G}}_t \Lambda_{u_1} - 2 \bar{\mathbf{G}}_t \Lambda_{u_2} \bar{\mathbf{G}}_t) \end{aligned}$$

where we use $-2\Lambda_{u_2} \preceq \bar{\mathbf{V}}_t^3 (\mathbf{I}_r + \eta \bar{\mathbf{V}}_t)^{-1} \preceq 2\Lambda_{u_2}$ in (a) and (c), and we use $\eta \|\Lambda\|_F^2 r_s \ll \kappa_d r k^{-\alpha} \frac{r_s}{d}$ in (b) and (d). By Proposition 34, we have $\frac{r_s}{d} \mathbf{I}_{rk} \preceq \bar{\mathbf{G}}_{t+1} \preceq 1.1 \mathbf{I}_{rk}$, hence, the induction hypothesis holds. Therefore, we have the first item.

By using the first item and Proposition 34, we can write for the given time horizons in second and third items that

$$\begin{aligned} \bar{\mathbf{G}}_t &\preceq \left(\bar{\mathbf{G}}_0 \exp(2(1 + \kappa_d) \eta t \Lambda_{u_1}) \wedge (1 + (1 + \kappa_d)^2 \eta^2 \Lambda_{u_1}^2) \mathbf{I}_{rk} \right) \\ &\preceq \begin{cases} (1.1 \bar{\mathbf{G}}_0 \exp(2\eta t \Lambda) \wedge \mathbf{I}_{rk}) + 2\eta^2 \mathbf{I}_{rk} \\ 1.2 (\bar{\mathbf{G}}_0 \exp(2\eta t \Lambda) \wedge \mathbf{I}_{rk}), \end{cases} \end{aligned} \quad (\text{F.36})$$

where both upper bounds in (F.36) are valid and will be used in different parts of the proof. The third sub-items immediately follow from the first bound.

For $\alpha \in [0, 0.5)$, we have

$$\begin{aligned} \hat{\mathbf{T}}_{t,ii} &\geq \frac{1}{3} (1 \wedge \hat{\mathbf{T}}_{0,ii} \exp(2\eta t \lambda_i)) \Rightarrow \mathbf{T}_{t,ii} \geq \frac{1}{4} (\kappa_d \wedge \mathbf{T}_{0,ii} \exp(2\eta t \lambda_i)) \\ &\stackrel{(e)}{\Rightarrow} \mathbf{T}_{t,ii} \geq \frac{\kappa_d}{4} (1 \wedge \frac{r_s}{d} \exp(2\eta t \lambda_i)), \end{aligned}$$

where we used $\mathbf{T}_0 = \frac{\kappa_d r_s}{d} \mathbf{I}_{rk}$ in (e). Therefore by (F.33) and the second bound in (F.36), we have for $j \leq t$

$$\frac{\bar{\mathbf{G}}_{j,ii}}{\mathbf{T}_{t,ii}} \leq \frac{1.2 \left(1 \wedge \left(1 + \frac{1}{\sqrt{\varphi}} \right)^2 \frac{2.2 r_s}{d} \exp(2\eta t \lambda_i) \right)}{0.25 \kappa_d (1 \wedge \frac{r_s}{d} \exp(2\eta t \lambda_i))} \leq \frac{11}{\kappa_d} \left(1 + \frac{1}{\sqrt{\varphi}} \right)^2.$$

On the other hand, by using the second bound in (F.36),

$$\begin{aligned} \frac{\eta \sum_{j=0}^{t-1} \bar{\mathbf{G}}_{j,ii}}{\mathbf{T}_{t,ii}} &\leq \frac{11 \left(1 + \frac{1}{\sqrt{\varphi}} \right)^2}{\kappa_d} \frac{\eta \left(t \wedge \frac{r_s}{d} \sum_{j=0}^{t-1} \exp(2\eta j \lambda_i) \right)}{(1 \wedge \frac{r_s}{d} \exp(2\eta t \lambda_i))} \\ &\leq \frac{5.5 \left(1 + \frac{1}{\sqrt{\varphi}} \right)^2}{\kappa_d} \begin{cases} \frac{1}{\lambda_i}, & 2\eta t \leq \frac{\log \frac{d}{r_s}}{\lambda_i} \\ 2\eta t, & 2\eta t > \frac{\log \frac{d}{r_s}}{\lambda_i} \end{cases} \end{aligned}$$

$$\leq \frac{5.5 \left(1 + \frac{1}{\sqrt{\varphi}}\right)^2}{\kappa_d} (2\eta t \vee r^\alpha).$$

Lastly, by using the first bound in (F.36), we get

$$\|\Lambda\|_{\mathbb{F}}^2 - \text{Tr}(\Lambda \bar{\mathbf{G}}_t \Lambda) \geq \sum_{i=1}^r \lambda_i^2 \left(1 - \frac{2.5 \left(1 + \frac{1}{\sqrt{\varphi}}\right)^2 r_s}{d} \exp(2\eta t \lambda_i)\right)_+ - 2\eta^2 \|\Lambda\|_{\mathbb{F}}^2.$$

For $\alpha > 0.5$, we have for $i \leq 2r_s \log^{\frac{1}{\alpha}} d$ and $d \geq \Omega_{r_s}(1)$,

$$\hat{\mathbf{T}}_{t,ii} \geq \frac{1}{3} (1 \wedge \hat{\mathbf{T}}_{0,ii} \exp(2\eta t \lambda_i)) \Rightarrow \mathbf{T}_{t,ii} \geq \frac{1}{4} (\kappa_d \wedge \mathbf{T}_{0,ii} \exp(2\eta t \lambda_i)).$$

Therefore, we have

$$\mathbf{T}_{t,ii} \geq \kappa_d \begin{cases} 0.25 (1 \wedge \frac{r_s}{d} \exp(2\eta t \lambda_i)), & i \leq 2r_s \log^{\frac{1}{\alpha}} d \\ \frac{r_s}{d}, & i > 2r_s \log^{\frac{1}{\alpha}} d. \end{cases}$$

On the other hand, for $\eta t \leq \frac{1}{2} r_s^\alpha \log(\frac{d \log^{1.5} d}{r_s})$ and $i > 2r_s \log^{\frac{1}{\alpha}} d$, we have for $d \geq \Omega(1)$.

$$\bar{\mathbf{G}}_{t,ii} \leq 1.2 (\bar{\mathbf{G}}_{0,ii} \exp(2\eta t \lambda_i) \wedge 1) \leq 1.2 \left(\bar{\mathbf{G}}_{0,ii} \exp\left(\frac{r_s^\alpha \log(\frac{d \log d}{r_s})}{2^\alpha r_s^\alpha \log d}\right) \wedge 1 \right) \leq 1.5 \bar{\mathbf{G}}_{0,ii}.$$

Therefore, for $\eta t \leq \frac{1}{2} r_s^\alpha \log(\frac{d \log^{1.5} d}{r_s})$,

$$\begin{aligned} \frac{\bar{\mathbf{G}}_{j,ii}}{\mathbf{T}_{t,ii}} &\leq \begin{cases} \frac{1.2 (1 \wedge \frac{5.5 r_u r_s}{d} \exp(2\eta t \lambda_i))}{0.25 \kappa_d (1 \wedge \frac{r_s}{d} \exp(2\eta t \lambda_i))}, & i \leq 2r_s \log^{\frac{1}{\alpha}} d \\ \frac{d}{\kappa_d r_s} \frac{8.25 r_u r_s}{d}, & i > 2r_s \log^{\frac{1}{\alpha}} d \end{cases} \\ &\leq \frac{26.4 r_u}{\kappa_d}. \end{aligned}$$

Moreover, for $\eta t \leq \frac{1}{2} r_s^\alpha \log(\frac{d \log^{1.5} d}{r_s})$,

$$\begin{aligned} \frac{\eta \sum_{j=0}^{t-1} \bar{\mathbf{G}}_{j,ii}}{\mathbf{T}_{t,ii}} &\leq \begin{cases} \frac{26.4 r_u}{\kappa_d} \frac{\eta \left(t \wedge \frac{r_s}{d} \sum_{j=0}^{t-1} \exp(2\eta j \lambda_i)\right)}{(1 \wedge \frac{r_s}{d} \exp(2\eta t \lambda_i))}, & i \leq 2r_s \log^{\frac{1}{\alpha}} d \\ \frac{d}{\kappa_d r_s} \frac{8.25 r_u r_s}{d} \eta t, & i > 2r_s \log^{\frac{1}{\alpha}} d \end{cases} \\ &\leq \frac{13.2 r_u}{\kappa_d} \begin{cases} \frac{1}{\lambda_i}, & 2\eta t \leq \frac{\log \frac{d}{r_s}}{\lambda_i} \text{ and } i \leq 2r_s \log^{\frac{1}{\alpha}} d \\ 2\eta t, & \text{otherwise} \end{cases} \\ &\leq \frac{13.2 r_u}{\kappa_d} (2\eta t \vee (2r_s)^\alpha \log d) \leq \frac{15 r_u}{\kappa_d} (2r_s)^\alpha \log d. \end{aligned}$$

Finally for $\eta t \leq \frac{1}{2} r_s^\alpha \log(\frac{d \log^{1.5} d}{r_s})$ and $d \geq \Omega_{r_s}(1)$, by using the first bound in (F.36),

$$\begin{aligned} \|\Lambda_{11}\|_{\mathbb{F}}^2 - \text{Tr}(\Lambda_{11} \bar{\mathbf{G}}_{t,11} \Lambda_{11}) &\geq \sum_{i=1}^{r_s} \lambda_i^2 \left(1 - \frac{12.1 r_s}{d} \exp(2\eta t \lambda_i)\right)_+ \\ &+ \sum_{i=r_s+1}^{2r_s} \lambda_i^2 \left(1 - \frac{12.1 r_s}{d} \exp(2\eta t \lambda_i)\right)_+ + \sum_{i=2r_s+1}^{r_u} \lambda_i^2 \left(1 - \frac{6.05 r_u r_s}{d} \exp(2\eta t \lambda_i)\right)_+ - 2\eta^2 \|\Lambda\|_{\mathbb{F}}^2 \\ &\stackrel{(f)}{\geq} \sum_{i=1}^{r_s} \lambda_i^2 \left(1 - \frac{12.1 r_s}{d} \exp(2\eta t \lambda_i)\right)_+ + \left(1 - \frac{1}{\log d}\right) \sum_{i=r_s+1}^{2r_s} \lambda_i^2 \end{aligned}$$

$$\begin{aligned}
& + \left(1 - 6.05r_u \log^{\frac{1.5}{\sqrt{2}}} d \left(\frac{r_s}{d}\right)^{1-\frac{1}{\sqrt{2}}}\right) \sum_{i=2r_s+1}^{r_u} \lambda_i^2 - 2\eta^2 \|\mathbf{A}\|_{\mathbb{F}}^2 \\
& \geq \sum_{i=1}^{r_s} \lambda_i^2 \left(1 - \frac{12.1r_s}{d} \exp(2\eta t \lambda_i)\right)_+ + \left(1 - 6.05r_u \log^{\frac{1.5}{\sqrt{2}}} d \left(\frac{r_s}{d}\right)^{1-\frac{1}{\sqrt{2}}}\right) \sum_{i=r_s+1}^{r_u} \lambda_i^2 \\
& \quad - \frac{(r_s+1)^{1+2\alpha}}{\log d} - 2\eta^2 \|\mathbf{A}\|_{\mathbb{F}}^2,
\end{aligned}$$

where we used the bounds for t, d in (f). $\square$

F.5 Bounds for the second-order terms

We recall

$$\begin{aligned}
R_{\text{so}}[\mathbf{G}_t] &= \frac{\eta^2}{16r_s} \mathbf{\Theta}^\top \mathbb{E}_t [\nabla_{\text{St}} \mathbf{L}_{t+1} \nabla_{\text{St}} \mathbf{L}_{t+1}^\top] \mathbf{\Theta} \\
&\quad - \frac{\eta^2}{16r_s} \mathbf{M}_t \mathbb{E}_t \left[\frac{\mathcal{P}_{t+1}}{1 + c_{t+1}^2} \right] \mathbf{M}_t^\top - \frac{\eta^3}{32r_s^{3/2}} \text{Sym} \left(\mathbf{\Theta}^\top \mathbb{E}_t \left[\frac{\nabla_{\text{St}} \mathbf{L}_{t+1} \mathcal{P}_{t+1}}{1 + c_{t+1}^2} \right] \mathbf{M}_t^\top \right) \\
&\quad - \frac{\eta^4}{256r_s^2} \mathbf{\Theta}^\top \mathbb{E}_t \left[\frac{\nabla_{\text{St}} \mathbf{L}_{t+1} \mathcal{P}_{t+1} \nabla_{\text{St}} \mathbf{L}_{t+1}^\top}{1 + c_{t+1}^2} \right] \mathbf{\Theta}. \tag{F.37}
\end{aligned}$$

Proposition 17. For $\eta \ll d^{-1/2}$, there exists a universal constant $C > 0$ such that

$$-C \left(\frac{\eta^2 d}{r_s} \mathbf{G}_t + \eta^2 \mathbf{I}_r \right) \preceq R_{\text{so}}[\mathbf{G}_t] \preceq C \left(\frac{\eta^2 d}{r_s} \mathbf{G}_t + \eta^2 \mathbf{I}_r \right).$$

Proof. We bound each term in (F.37). In the following, $\mathbf{v}$ denotes a generic unit norm vector with proper dimensionality. For the first term,

$$\begin{aligned}
& \mathbf{\Theta}^\top \mathbb{E}_t [\nabla_{\text{St}} \mathbf{L}_{t+1} \nabla_{\text{St}} \mathbf{L}_{t+1}^\top] \mathbf{\Theta} \\
&= \mathbf{\Theta}^\top (\mathbf{I}_d - \mathbf{W}_t \mathbf{W}_t^\top) \mathbb{E}_t [(y_{t+1} - \hat{y}_{t+1})^2 \|\mathbf{W}_t^\top \mathbf{x}_{t+1}\|_2^2 \mathbf{x}_{t+1} \mathbf{x}_{t+1}^\top] (\mathbf{I}_d - \mathbf{W}_t \mathbf{W}_t^\top) \mathbf{\Theta}.
\end{aligned}$$

We have

$$\mathbb{E}_t [(y_{t+1} - \hat{y}_{t+1})^2 \|\mathbf{W}_t^\top \mathbf{x}_{t+1}\|_2^2 \langle \mathbf{v}, \mathbf{x}_{t+1} \rangle^2] \leq Cr_s.$$

Therefore,

$$0 \preceq \frac{\eta^2}{16r_s} \mathbf{\Theta}^\top \mathbb{E}_t [\nabla_{\text{St}} \mathbf{L}_{t+1} \nabla_{\text{St}} \mathbf{L}_{t+1}^\top] \mathbf{\Theta} \preceq C\eta^2 (\mathbf{I}_r - \mathbf{G}_t).$$

For the second term,

$$\begin{aligned}
& \mathbf{M}_t \mathbb{E}_t \left[\frac{\mathcal{P}_{t+1}}{1 + c_{t+1}^2} \right] \mathbf{M}_t^\top \\
&= \mathbf{M}_t \mathbb{E}_t \left[\frac{(y_{t+1} - \hat{y}_{t+1})^2 \|(\mathbf{I}_d - \mathbf{W}_t \mathbf{W}_t^\top) \mathbf{x}_{t+1}\|_2^2 \mathbf{W}_t^\top \mathbf{x}_{t+1} \mathbf{x}_{t+1}^\top \mathbf{W}_t}{1 + c_{t+1}^2} \right] \mathbf{M}_t^\top.
\end{aligned}$$

We have

$$\mathbb{E}_t \left[\frac{(y_{t+1} - \hat{y}_{t+1})^2 \|(\mathbf{I}_d - \mathbf{W}_t \mathbf{W}_t^\top) \mathbf{x}_{t+1}\|_2^2 \langle \mathbf{v}, \mathbf{W}_t^\top \mathbf{x}_{t+1} \rangle^2}{1 + c_{t+1}^2} \right] \leq Cd.$$

Therefore,

$$0 \preceq \frac{\eta^2}{16r_s} \mathbf{M}_t \mathbb{E}_t \left[\frac{\mathcal{P}_{t+1}}{1 + c_{t+1}^2} \right] \mathbf{M}_t^\top \preceq C \frac{\eta^2 d}{r_s} \mathbf{G}_t.$$

For the third term by using Proposition 22,

$$\begin{aligned} & \frac{\eta^3}{32r_s^{3/2}} \text{Sym} \left(\Theta^\top \mathbb{E}_t \left[\frac{\nabla_{\text{St}} \mathbf{L}_{t+1} \mathcal{P}_{t+1}}{1 + c_{t+1}^2} \right] \mathbf{M}_t^\top \right) \\ & \preceq C \left(\frac{\eta^4}{r_s^2 d} \Theta^\top \mathbb{E}_t \left[\frac{\nabla_{\text{St}} \mathbf{L}_{t+1} \mathcal{P}_{t+1}}{1 + c_{t+1}^2} \right] \mathbb{E}_t \left[\frac{\mathcal{P}_{t+1} \nabla_{\text{St}} \mathbf{L}_{t+1}^\top}{1 + c_{t+1}^2} \right] \Theta + \frac{\eta^2 d}{r_s} \mathbf{G}_t \right) \end{aligned}$$

We have

$$\begin{aligned} & \Theta^\top \mathbb{E}_t \left[\frac{\nabla_{\text{St}} \mathbf{L}_{t+1} \mathcal{P}_{t+1}}{1 + c_{t+1}^2} \right] \\ & = \Theta^\top (\mathbf{I}_d - \mathbf{W}_t \mathbf{W}_t^\top) \\ & \quad \times \mathbb{E}_t \left[\frac{(y_{t+1} - \hat{y}_{t+1})^3}{1 + c_{t+1}^2} \|(\mathbf{I}_d - \mathbf{W}_t \mathbf{W}_t^\top) \mathbf{x}_{t+1}\|_2^2 \|\mathbf{W}_t^\top \mathbf{x}_{t+1}\|_2^2 \mathbf{x}_{t+1} \mathbf{x}_{t+1}^\top \mathbf{W}_t \right]. \end{aligned}$$

Then, by using Cauchy-Schwartz inequality, we can show that

$$\left\| \mathbb{E}_t \left[\frac{(y_{t+1} - \hat{y}_{t+1})^3}{1 + c_{t+1}^2} \|(\mathbf{I}_d - \mathbf{W}_t \mathbf{W}_t^\top) \mathbf{x}_{t+1}\|_2^2 \|\mathbf{W}_t^\top \mathbf{x}_{t+1}\|_2^2 \mathbf{x}_{t+1} \mathbf{x}_{t+1}^\top \mathbf{W}_t \right] \right\|_2 \leq C d r_s.$$

Therefore,

$$\frac{\eta^4}{r_s^2 d} \Theta^\top \mathbb{E}_t \left[\frac{\nabla_{\text{St}} \mathbf{L}_{t+1} \mathcal{P}_{t+1}}{1 + c_{t+1}^2} \right] \mathbb{E}_t \left[\frac{\mathcal{P}_{t+1} \nabla_{\text{St}} \mathbf{L}_{t+1}^\top}{1 + c_{t+1}^2} \right] \Theta \leq C \eta^4 d (\mathbf{I}_r - \mathbf{G}_t).$$

We get

$$\frac{\eta^3}{32r_s^{3/2}} \text{Sym} \left(\Theta^\top \mathbb{E}_t \left[\frac{\nabla_{\text{St}} \mathbf{L}_{t+1} \mathcal{P}_{t+1}}{1 + c_{t+1}^2} \right] \mathbf{M}_t^\top \right) \preceq C \left(\eta^4 d (\mathbf{I}_r - \mathbf{G}_t) + \frac{\eta^2 d}{r_s} \mathbf{G}_t \right).$$

By repeating the argument with the lower bound in Proposition 22, we can also show

$$\frac{\eta^3}{32r_s^{3/2}} \text{Sym} \left(\Theta^\top \mathbb{E}_t \left[\frac{\nabla_{\text{St}} \mathbf{L}_{t+1} \mathcal{P}_{t+1}}{1 + c_{t+1}^2} \right] \mathbf{M}_t^\top \right) \succeq -C \left(\eta^4 d (\mathbf{I}_r - \mathbf{G}_t) + \frac{\eta^2 d}{r_s} \mathbf{G}_t \right).$$

For the last term, we write

$$\begin{aligned} & \Theta^\top \mathbb{E}_t \left[\frac{\nabla_{\text{St}} \mathbf{L}_{t+1} \mathcal{P}_{t+1} \nabla_{\text{St}} \mathbf{L}_{t+1}^\top}{1 + c_{t+1}^2} \right] \Theta \\ & = \Theta^\top (\mathbf{I}_d - \mathbf{W}_t \mathbf{W}_t^\top) \\ & \quad \times \mathbb{E}_t \left[\frac{(y_{t+1} - \hat{y}_{t+1})^4}{1 + c_{t+1}^2} \|\mathbf{W}_t^\top \mathbf{x}_{t+1}\|_2^4 \|(\mathbf{I}_d - \mathbf{W}_t \mathbf{W}_t^\top) \mathbf{x}_{t+1}\|_2^2 \mathbf{x}_{t+1} \mathbf{x}_{t+1}^\top \right] (\mathbf{I}_d - \mathbf{W}_t \mathbf{W}_t^\top) \Theta. \end{aligned}$$

We have

$$\mathbb{E}_t \left[\frac{(y_{t+1} - \hat{y}_{t+1})^4}{1 + c_{t+1}^2} \|\mathbf{W}_t^\top \mathbf{x}_{t+1}\|_2^4 \|(\mathbf{I}_d - \mathbf{W}_t \mathbf{W}_t^\top) \mathbf{x}_{t+1}\|_2^2 \langle \mathbf{v}, \mathbf{x}_{t+1} \rangle^2 \right] \leq C d r_s^2.$$

Therefore,

$$0 \preceq \frac{\eta^4}{r_s^2} \Theta^\top \mathbb{E}_t \left[\frac{\nabla_{\text{St}} \mathbf{L}_{t+1} \mathcal{P}_{t+1} \nabla_{\text{St}} \mathbf{L}_{t+1}^\top}{1 + c_{t+1}^2} \right] \Theta \preceq C \eta^4 d (\mathbf{I}_r - \mathbf{G}_t).$$

By using $\mathbf{G}_t \succeq 0$ and $\eta \ll d^{-1/2}$, the result follows. $\square$

F.6 Noise characterization

To prove the noise concentration bound for both the heavy-tailed and light-tailed cases simultaneously, we introduce some new notation. Specifically, we define the submatrix notation:

$$\Theta =: [\Theta_1 \quad \Theta_2] \quad \text{and} \quad M_t =: \begin{bmatrix} M_{t,1} \\ M_{t,2} \end{bmatrix} = \begin{bmatrix} \Theta_1^\top \mathbf{W}_t \\ \Theta_2^\top \mathbf{W}_t \end{bmatrix},$$

where $\Theta_1 \in \mathbb{R}^{d \times r_u}$ and $M_{t,1} \in \mathbb{R}^{r_u \times r_s}$. We note that $\mathbf{G}_{t,11} = M_{t,1} M_{t,1}^\top$. To unify the treatment of the heavy-tailed and light-tailed cases, we use the following notation to represent both cases:

$$\Theta := \{\Theta, \Theta_1\} \quad M_t := \{M_t, M_{t,1}\}.$$

With the new notation, we have

$$\begin{aligned} \frac{\eta/2}{\sqrt{r_s}} \mathbf{v}_{t+1} &= \frac{\eta/2}{\sqrt{r_s}} \text{Sym} \left(\Theta^\top \left(\nabla_{\text{St}} \mathbf{L}_{t+1} - \mathbb{E}_t [\nabla_{\text{St}} \mathbf{L}_{t+1}] \right) M_t^\top \right) \\ &\quad - \frac{\eta^2}{16r_s} M_t \left(\frac{\mathcal{P}_{t+1}}{1 + c_{t+1}^2} - \mathbb{E}_t \left[\frac{\mathcal{P}_{t+1}}{1 + c_{t+1}^2} \right] \right) M_t^\top \\ &\quad + \frac{\eta^2}{16r_s} \Theta^\top \left(\nabla_{\text{St}} \mathbf{L}_{t+1} \nabla_{\text{St}} \mathbf{L}_{t+1}^\top - \mathbb{E}_t [\nabla_{\text{St}} \mathbf{L}_{t+1} \nabla_{\text{St}} \mathbf{L}_{t+1}^\top] \right) \Theta \\ &\quad - \frac{\eta^3}{32r_s^{3/2}} \text{Sym} \left(\Theta^\top \left(\frac{\nabla_{\text{St}} \mathbf{L}_{t+1} \mathcal{P}_{t+1}}{1 + c_{t+1}^2} - \mathbb{E}_t \left[\frac{\nabla_{\text{St}} \mathbf{L}_{t+1} \mathcal{P}_{t+1}}{1 + c_{t+1}^2} \right] \right) M_t^\top \right) \\ &\quad - \frac{\eta^4}{256r_s^2} \Theta^\top \left(\frac{\nabla_{\text{St}} \mathbf{L}_{t+1} \mathcal{P}_{t+1} \nabla_{\text{St}} \mathbf{L}_{t+1}^\top}{1 + c_{t+1}^2} - \mathbb{E}_t \left[\frac{\nabla_{\text{St}} \mathbf{L}_{t+1} \mathcal{P}_{t+1} \nabla_{\text{St}} \mathbf{L}_{t+1}^\top}{1 + c_{t+1}^2} \right] \right) \Theta. \end{aligned}$$

For $\text{rk} \in \{r, r_u\}$ and $\mathbf{T}_1, \mathbf{T}_2 \in \mathbb{R}^{\text{rk} \times \text{rk}}$ be a deterministic symmetric positive definite matrices, we define

$$\begin{aligned} \mathcal{A}_{t+1}(\mathbf{T}_1, \mathbf{T}_2) &\equiv \left\{ \left\| \mathbf{T}_1 \Theta^\top \nabla_{\text{St}} \mathbf{L}_{t+1} M_t^\top \mathbf{T}_2 \right\|_2 \leq \frac{L^2}{2} \sqrt{\text{Tr}(\mathbf{T}_1^2 (\mathbf{I}_{\text{rk}} - \mathbf{G}_t)) \text{Tr}(\mathbf{T}_2^2 \mathbf{G}_t)} \right\} \\ \mathcal{B}_{t+1}(\mathbf{T}_1, \mathbf{T}_2) &\equiv \left\{ \left\| \mathbf{T}_1 M_t \mathcal{P}_{t+1} M_t^\top \mathbf{T}_2 \right\|_2 \leq \frac{L^4 d}{2} \sqrt{\text{Tr}(\mathbf{T}_1^2 \mathbf{G}_t) \text{Tr}(\mathbf{T}_2^2 \mathbf{G}_t)} \right\} \\ \mathcal{C}_{t+1}(\mathbf{T}_1, \mathbf{T}_2) &\equiv \left\{ \left\| \mathbf{T}_1 \Theta^\top \nabla_{\text{St}} \mathbf{L}_{t+1} \nabla_{\text{St}} \mathbf{L}_{t+1}^\top \Theta \mathbf{T}_2 \right\|_2 \leq \frac{L^4 r_s}{2} \sqrt{\text{Tr}(\mathbf{T}_1^2 (\mathbf{I}_{\text{rk}} - \mathbf{G}_t)) \text{Tr}(\mathbf{T}_2^2 (\mathbf{I}_{\text{rk}} - \mathbf{G}_t))} \right\} \\ \mathcal{D}_{t+1}(\mathbf{T}_1, \mathbf{T}_2) &\equiv \left\{ \left\| \mathbf{T}_1 \Theta^\top \nabla_{\text{St}} \mathbf{L}_{t+1} \mathcal{P}_{t+1} M_t^\top \mathbf{T}_2 \right\|_2 \leq \frac{L^6 d r_s}{2} \sqrt{\text{Tr}(\mathbf{T}_1^2 (\mathbf{I}_{\text{rk}} - \mathbf{G}_t)) \text{Tr}(\mathbf{T}_2^2 \mathbf{G}_t)} \right\} \\ \mathcal{F}_{t+1}(\mathbf{T}_1, \mathbf{T}_2) &\equiv \left\{ \left\| \mathbf{T}_1 \Theta^\top \nabla_{\text{St}} \mathbf{L}_{t+1} \mathcal{P}_{t+1} \nabla_{\text{St}} \mathbf{L}_{t+1}^\top \Theta \mathbf{T}_2 \right\|_2 \leq \frac{L^8 d r_s^2}{2} \sqrt{\text{Tr}(\mathbf{T}_1^2 (\mathbf{I}_{\text{rk}} - \mathbf{G}_t)) \text{Tr}(\mathbf{T}_2^2 (\mathbf{I}_{\text{rk}} - \mathbf{G}_t))} \right\}. \end{aligned}$$

We start with the following statement:

Proposition 18. *Let $e_{t+1} := (y_{t+1} - \hat{y}_{t+1})$. There exists a universal constant $C > 0$ such that for $L \geq 2e(\sqrt{8} + \sigma)$, the following statements hold:*

1. We have $\mathbb{P}_t [\mathcal{A}_{t+1}(\mathbf{T}_1, \mathbf{T}_2) \cap \{|e_{t+1}| \leq L\}] \geq 1 - 2e^{\frac{-L/e}{(\sqrt{8} + \sigma)}}$. Moreover,

$$\begin{aligned} &\mathbb{E}_t \left[\left(\text{Sym}(\mathbf{T}_1 \Theta^\top \nabla_{\text{St}} \mathbf{L}_{t+1} M_t^\top \mathbf{T}_2) \right)^2 \right] \\ &\quad \preceq C \left(\text{Tr}(\mathbf{T}_2^2 \mathbf{G}_t) \mathbf{T}_1 (\mathbf{I}_{\text{rk}} - \mathbf{G}_t) \mathbf{T}_1 + \text{Tr}(\mathbf{T}_1^2 (\mathbf{I}_{\text{rk}} - \mathbf{G}_t)) \mathbf{T}_2 \mathbf{G}_t \mathbf{T}_2 \right). \end{aligned}$$

2. We have $\mathbb{P}_t [\mathcal{B}_{t+1}(\mathbf{T}_1, \mathbf{T}_2) \cap \{|e_{t+1}| \leq L\}] \geq 1 - 2e^{\frac{-L/e}{(\sqrt{8} + \sigma)}}$. Moreover,

$$\begin{aligned} &\mathbb{E}_t \left[\left(\text{Sym}(\mathbf{T}_1 M_t \mathcal{P}_{t+1} M_t^\top \mathbf{T}_2) \right)^2 \right] \\ &\quad \preceq C d^2 \left(\text{Tr}(\mathbf{T}_2^2 \mathbf{G}_t) \mathbf{T}_1 \mathbf{G}_t \mathbf{T}_1 + \text{Tr}(\mathbf{T}_1^2 \mathbf{G}_t) \mathbf{T}_2 \mathbf{G}_t \mathbf{T}_2 \right). \end{aligned}$$

3. We have $\mathbb{P}_t[\mathcal{C}_{t+1}(\mathbf{T}_1, \mathbf{T}_2) \cap \{|e_{t+1}| \leq L\}] \geq 1 - 2e^{\frac{-L/e}{(\sqrt{8}+\sigma)}}$. Moreover,

$$\begin{aligned} & \mathbb{E}_t \left[\left(\text{Sym} \left(\mathbf{T}_1 \Theta^\top \nabla_{S_t} \mathbf{L}_{t+1} \nabla_{S_t} \mathbf{L}_{t+1}^\top \mathbf{T}_2 \right) \right)^2 \right] \\ & \leq C r_s^2 \left(\text{Tr}(\mathbf{T}_2^2 (\mathbf{I}_{\text{rk}} - \mathbf{G}_t)) \mathbf{T}_1 (\mathbf{I}_{\text{rk}} - \mathbf{G}_t) \mathbf{T}_1 + \text{Tr}(\mathbf{T}_1^2 (\mathbf{I}_{\text{rk}} - \mathbf{G}_t)) \mathbf{T}_2 (\mathbf{I}_{\text{rk}} - \mathbf{G}_t) \mathbf{T}_2 \right). \end{aligned}$$

4. We have $\mathbb{P}_t[\mathcal{D}_{t+1}(\mathbf{T}_1, \mathbf{T}_2) \cap \{|e_{t+1}| \leq L\}] \geq 1 - 2e^{\frac{-L/e}{(\sqrt{7/2}+\sigma)}}$. Moreover,

$$\begin{aligned} & \mathbb{E}_t \left[\left(\text{Sym} \left(\mathbf{T}_1 \Theta^\top \nabla_{S_t} \mathbf{L}_{t+1} \mathcal{P}_{t+1} \mathbf{M}_t^\top \mathbf{T}_2 \right) \right)^2 \right] \\ & \leq C d^2 r_s^2 \left(\text{Tr}(\mathbf{T}_2^2 \mathbf{G}_t) \mathbf{T}_1 (\mathbf{I}_{\text{rk}} - \mathbf{G}_t) \mathbf{T}_1 + \text{Tr}(\mathbf{T}_1^2 (\mathbf{I}_{\text{rk}} - \mathbf{G}_t)) \mathbf{T}_2 \mathbf{G}_t \mathbf{T}_2 \right). \end{aligned}$$

5. We have $\mathbb{P}_t[\mathcal{F}_{t+1}(\mathbf{T}_1, \mathbf{T}_2) \cap \{|e_{t+1}| \leq L\}] \geq 1 - 2e^{\frac{-L/e}{(4\sqrt{2}+\sigma)}}$. Moreover,

$$\begin{aligned} & \mathbb{E}_t \left[\left(\text{Sym} \left(\mathbf{T}_1 \Theta^\top \nabla_{S_t} \mathbf{L}_{t+1} \mathcal{P}_{t+1} \nabla_{S_t} \mathbf{L}_{t+1}^\top \Theta \mathbf{T}_2 \right) \right)^2 \right] \\ & \leq C d^2 r_s^4 \left(\text{Tr}(\mathbf{T}_2^2 (\mathbf{I}_{\text{rk}} - \mathbf{G}_t)) \mathbf{T}_1 (\mathbf{I}_{\text{rk}} - \mathbf{G}_t) \mathbf{T}_1 + \text{Tr}(\mathbf{T}_1^2 (\mathbf{I}_{\text{rk}} - \mathbf{G}_t)) \mathbf{T}_2 (\mathbf{I}_{\text{rk}} - \mathbf{G}_t) \mathbf{T}_2 \right). \end{aligned}$$

Proof. First, we derive a concentration bound for $|e_{t+1}|$. By Corollary 6 and Proposition 33 we have

$$\mathbb{E}_t[|e_{t+1}|^p] \leq (\mathbb{E}_t[|y_{t+1}|^p]^{\frac{1}{p}} + \mathbb{E}_t[|\hat{y}_{t+1}|^p]^{\frac{1}{p}})^p \leq (\sqrt{8} + \sigma)^p p^p \text{ for } p \geq 2,$$

which implies $\mathbb{P}_t[|e_{t+1}| \geq u] \leq e^{\frac{-u/e}{(\sqrt{8}+\sigma)}}$ for $u \geq 2e(\sqrt{8} + \sigma)$. In the following, we prove each item separately.

First item. We define

$$\mathbf{T}_1 \Theta^\top \nabla_{S_t} \mathbf{L}_{t+1} \mathbf{M}_t^\top \mathbf{T}_2 = \underbrace{\mathbf{T}_1 \Theta^\top (\mathbf{I}_d - \mathbf{W}_t \mathbf{W}_t^\top) e_{t+1} \mathbf{x}_{t+1}}_{:= \mathbf{u}_{t+1}} \underbrace{\mathbf{x}_{t+1}^\top \mathbf{W}_t \mathbf{W}_t^\top \Theta \mathbf{T}_2}_{:= \mathbf{v}_{t+1}^\top}.$$

For $u, L > 0$

$$\begin{aligned} & \mathbb{P}_t \left[\|\mathbf{u}_{t+1} \mathbf{v}_{t+1}^\top\|_2 \geq uL \sqrt{\text{Tr}(\mathbf{T}_1^2 (\mathbf{I}_{\text{rk}} - \mathbf{G}_t)) \text{Tr}(\mathbf{T}_2^2 \mathbf{G}_t)} \text{ or } |e_{t+1}| \geq L \right] \\ & \leq \mathbb{P}_t \left[\|\mathbf{u}_{t+1} \mathbf{v}_{t+1}^\top\|_2 \geq uL \sqrt{\text{Tr}(\mathbf{T}_1^2 (\mathbf{I}_{\text{rk}} - \mathbf{G}_t)) \text{Tr}(\mathbf{T}_2^2 \mathbf{G}_t)} \text{ and } |e_{t+1}| \leq L \right] + \mathbb{P}_t[|e_{t+1}| \geq L] \\ & \leq \mathbb{P}_t \left[\|\mathbb{1}_{|e_{t+1}| \leq L} \mathbf{u}_{t+1} \mathbf{v}_{t+1}^\top\|_2 \geq uL \sqrt{\text{Tr}(\mathbf{T}_1^2 (\mathbf{I}_{\text{rk}} - \mathbf{G}_t)) \text{Tr}(\mathbf{T}_2^2 \mathbf{G}_t)} \right] + \mathbb{P}_t[|e_{t+1}| \geq L]. \end{aligned}$$

We have for $p \geq 2$

$$\begin{aligned} & \mathbb{E}_t \left[\|\mathbb{1}_{|e_{t+1}| \leq L} \mathbf{u}_{t+1} \mathbf{v}_{t+1}^\top\|_2^p \right] \\ & \leq L^p \mathbb{E}_t \left[\|\mathbf{T}_1 \Theta^\top (\mathbf{I}_d - \mathbf{W}_t \mathbf{W}_t^\top) \mathbf{x}_{t+1}\|_2^p \right] \mathbb{E}_t \left[\|\mathbf{T}_2 \Theta^\top \mathbf{W}_t \mathbf{W}_t^\top \mathbf{x}_{t+1}\|_2^p \right] \\ & \stackrel{(a)}{\leq} L^p \left(\frac{p}{2} \right)^p \left(3 \text{Tr}(\mathbf{T}_1^2 (\mathbf{I}_{\text{rk}} - \mathbf{G}_t)) \text{Tr}(\mathbf{T}_2^2 \mathbf{G}_t) \right)^{\frac{p}{2}}, \end{aligned}$$

where we used Corollary 7 in (a). By Proposition 33, we have for $u \geq 2e$

$$\mathbb{P}_t \left[\|\mathbb{1}_{|e_{t+1}| \leq L} \mathbf{u}_{t+1} \mathbf{v}_{t+1}^\top\|_2 \geq uL \sqrt{\text{Tr}(\mathbf{T}_1^2 (\mathbf{I}_{\text{rk}} - \mathbf{G}_t)) \text{Tr}(\mathbf{T}_2^2 \mathbf{G}_t)} \right] \leq e^{-\frac{u}{e}}.$$

By choosing $u = \frac{L}{2}$, we have the probability bound.

For the variance bound, we have

$$\mathbb{E}_t \left[\text{Sym} \left(\mathbf{T}_1 \Theta^\top \nabla_{S_t} \mathbf{L}_{t+1} \mathbf{M}_t^\top \mathbf{T}_2 \right)^2 \right] = \mathbb{E}_t \left[\text{Sym} \left(\mathbf{u}_{t+1} \mathbf{v}_{t+1}^\top \right)^2 \right].$$

By using Proposition 22, we have

$$\mathbb{E}_t \left[\text{Sym} \left(\mathbf{u}_{t+1} \mathbf{v}_{t+1}^\top \right)^2 \right]$$

$$\begin{aligned}
&\preceq T_1 \Theta^\top (I_d - \mathbf{W}_t \mathbf{W}_t^\top) \mathbb{E}_t [e_{t+1}^2 \|T_2 \Theta^\top \mathbf{W}_t \mathbf{W}_t^\top \mathbf{x}_{t+1}\|_2^2 \mathbf{x}_{t+1} \mathbf{x}_{t+1}^\top] (I_d - \mathbf{W}_t \mathbf{W}_t^\top) \Theta T_1 \\
&+ T_2 \Theta^\top \mathbf{W}_t \mathbf{W}_t^\top \mathbb{E}_t [e_{t+1}^2 \|T_1 \Theta^\top (I_d - \mathbf{W}_t \mathbf{W}_t^\top) \mathbf{x}_{t+1}\|_2^2 \mathbf{x}_{t+1} \mathbf{x}_{t+1}^\top] \mathbf{W}_t \mathbf{W}_t^\top \Theta T_2 \\
&\stackrel{(b)}{\preceq} C \left(\text{Tr}(T_2^2 \mathbf{G}_t) T_1 (I_{rk} - \mathbf{G}_t) T_1 + \text{Tr}(T_1^2 (I_{rk} - \mathbf{G}_t)) T_2 \mathbf{G}_t T_2 \right),
\end{aligned}$$

where we used the Cauchy-Schwartz inequality in (b).

Second item. We define

$$T_1 \mathbf{M}_t \mathcal{P}_{t+1} \mathbf{M}_t^\top T_2 = \underbrace{e_{t+1}^2 \|(I_d - \mathbf{W}_t \mathbf{W}_t^\top) \mathbf{x}_{t+1}\|_2^2 T_1 \Theta^\top \mathbf{W}_t \mathbf{W}_t^\top \mathbf{x}_{t+1}}_{:= \mathbf{u}_{t+1}} \underbrace{\mathbf{x}_{t+1}^\top \mathbf{W}_t \mathbf{W}_t^\top \Theta T_2}_{:= \mathbf{v}_{t+1}^\top}.$$

We have for $p \geq 2$

$$\begin{aligned}
&\mathbb{E}_t \left[\|\mathbb{1}_{|e_{t+1}| \leq L} \mathbf{u}_{t+1} \mathbf{v}_{t+1}^\top\|_2^p \right] \\
&\leq L^{2p} \mathbb{E}_t \left[\|(I_d - \mathbf{W}_t \mathbf{W}_t^\top) \mathbf{x}_{t+1}\|_2^{2p} \right] \mathbb{E}_t \left[\|T_1 \Theta^\top \mathbf{W}_t \mathbf{W}_t^\top \mathbf{x}_{t+1}\|_2^{2p} \right]^{\frac{1}{2}} \\
&\quad \times \mathbb{E}_t \left[\|T_2 \Theta^\top \mathbf{W}_t \mathbf{W}_t^\top \mathbf{x}_{t+1}\|_2^{2p} \right]^{\frac{1}{2}} \\
&\stackrel{(c)}{\leq} L^{2p} p^{2p} \left(3d \sqrt{\text{Tr}(T_1^2 \mathbf{G}_t) \text{Tr}(T_2^2 \mathbf{G}_t)} \right)^p,
\end{aligned}$$

where we use Corollary 7 in (c). By Proposition 33, we have for $u \geq (2e)^2$

$$\mathbb{P}_t \left[\|\mathbb{1}_{|e_{t+1}| \leq L} \mathbf{u}_{t+1} \mathbf{v}_{t+1}^\top\|_2 \geq u L^2 3d \sqrt{\text{Tr}(T_1^2 \mathbf{G}_t) \text{Tr}(T_2^2 \mathbf{G}_t)} \right] \leq e^{-\frac{u^{1/2}}{e}}.$$

By choosing $u = \frac{L^2}{6}$, we have the probability bound. For the variance bound, we have

$$\mathbb{E}_t \left[\left(\text{Sym} (T_1 \mathbf{M}_t \mathcal{P}_{t+1} \mathbf{M}_t^\top T_2) \right)^2 \right] = \mathbb{E}_t \left[\text{Sym} (\mathbf{u}_{t+1} \mathbf{v}_{t+1}^\top)^2 \right].$$

By using Proposition 22, we have

$$\begin{aligned}
&\mathbb{E}_t \left[\text{Sym} (\mathbf{u}_{t+1} \mathbf{v}_{t+1}^\top)^2 \right] \\
&\preceq T_1 \Theta^\top \mathbf{W}_t \mathbf{W}_t^\top \\
&\quad \times \mathbb{E}_t [e_{t+1}^4 \|(I_d - \mathbf{W}_t \mathbf{W}_t^\top) \mathbf{x}_{t+1}\|_2^4 \|T_2 \Theta^\top \mathbf{W}_t \mathbf{W}_t^\top \mathbf{x}_{t+1}\|_2^2 \mathbf{x}_{t+1} \mathbf{x}_{t+1}^\top] \mathbf{W}_t \mathbf{W}_t^\top \Theta T_1 \\
&+ T_2 \Theta^\top \mathbf{W}_t \mathbf{W}_t^\top \\
&\quad \times \mathbb{E}_t [e_{t+1}^4 \|(I_d - \mathbf{W}_t \mathbf{W}_t^\top) \mathbf{x}_{t+1}\|_2^4 \|T_1 \Theta^\top \mathbf{W}_t \mathbf{W}_t^\top \mathbf{x}_{t+1}\|_2^2 \mathbf{x}_{t+1} \mathbf{x}_{t+1}^\top] \mathbf{W}_t \mathbf{W}_t^\top \Theta T_2 \\
&\stackrel{(d)}{\preceq} C d^2 \left(\text{Tr}(T_2^2 \mathbf{G}_t) T_1 \mathbf{G}_t T_1 + \text{Tr}(T_1^2 \mathbf{G}_t) T_2 \mathbf{G}_t T_2 \right),
\end{aligned}$$

where we use the Cauchy-Schwartz inequality in (d).

Third item. We define

$$\begin{aligned}
&T_1 \Theta^\top \nabla_{\text{St}} L_{t+1} \nabla_{\text{St}} L_{t+1}^\top \Theta T_2 \\
&= \underbrace{e_{t+1}^2 \|\mathbf{W}_t^\top \mathbf{x}_{t+1}\|_2^2 T_1 \Theta^\top (I_d - \mathbf{W}_t \mathbf{W}_t^\top) \mathbf{x}_{t+1}}_{:= \mathbf{u}_{t+1}} \underbrace{\mathbf{x}_{t+1}^\top (I_d - \mathbf{W}_t \mathbf{W}_t^\top) \Theta T_2}_{:= \mathbf{v}_{t+1}^\top}.
\end{aligned}$$

We have for $p \geq 2$

$$\begin{aligned}
&\mathbb{E}_t \left[\|\mathbb{1}_{|e_{t+1}| \leq L} \mathbf{u}_{t+1} \mathbf{v}_{t+1}^\top\|_2^p \right] \\
&\leq L^{2p} \mathbb{E}_t \left[\|\mathbf{W}_t^\top \mathbf{x}_{t+1}\|_2^{2p} \right] \mathbb{E}_t \left[\|T_1 \Theta^\top (I_d - \mathbf{W}_t \mathbf{W}_t^\top) \mathbf{x}_{t+1}\|_2^{2p} \right]^{\frac{1}{2}} \\
&\quad \times \mathbb{E}_t \left[\|T_2 \Theta^\top (I_d - \mathbf{W}_t \mathbf{W}_t^\top) \mathbf{x}_{t+1}\|_2^{2p} \right]^{\frac{1}{2}}
\end{aligned}$$

$$\stackrel{(e)}{\leq} L^{2p} p^{2p} \left(3r_s \sqrt{\text{Tr}(\mathbf{T}_1^2(\mathbf{I}_{\text{rk}} - \mathbf{G}_t)) \text{Tr}(\mathbf{T}_2^2(\mathbf{I}_{\text{rk}} - \mathbf{G}_t))} \right)^p,$$

where we use Corollary 7 in (e). By Proposition 33, we have for $u \geq (2e)^2$

$$\mathbb{P}_t \left[\left\| \mathbb{1}_{|e_{t+1}| \leq L} \mathbf{u}_{t+1} \mathbf{v}_{t+1}^\top \right\|_2 \geq u L^2 3r_s \sqrt{\text{Tr}(\mathbf{T}_1^2(\mathbf{I}_{\text{rk}} - \mathbf{G}_t)) \text{Tr}(\mathbf{T}_2^2(\mathbf{I}_{\text{rk}} - \mathbf{G}_t))} \right] \leq e^{-\frac{u^{1/2}}{e}}.$$

By choosing $u = \frac{L^2}{6}$, we have the probability bound. For the variance bound, we have

$$\mathbb{E}_t \left[\left(\text{Sym}(\mathbf{T}_1 \Theta^\top \nabla_{\text{St}} \mathbf{L}_{t+1} \nabla_{\text{St}} \mathbf{L}_{t+1}^\top \Theta \mathbf{T}_2) \right)^2 \right] = \mathbb{E}_t \left[\text{Sym}(\mathbf{u}_{t+1} \mathbf{v}_{t+1}^\top)^2 \right].$$

By using Proposition 22, we have

$$\begin{aligned} & \mathbb{E}_t \left[\text{Sym}(\mathbf{u}_{t+1} \mathbf{v}_{t+1}^\top)^2 \right] \\ & \preceq \mathbf{T}_1 \Theta^\top (\mathbf{I}_d - \mathbf{W}_t \mathbf{W}_t^\top) \\ & \quad \times \mathbb{E}_t [e_{t+1}^4 \|\mathbf{W}_t^\top \mathbf{x}_{t+1}\|_2^4 \|\mathbf{T}_2 \Theta^\top (\mathbf{I}_d - \mathbf{W}_t \mathbf{W}_t^\top) \mathbf{x}_{t+1}\|_2^2 \mathbf{x}_{t+1} \mathbf{x}_{t+1}^\top] (\mathbf{I}_d - \mathbf{W}_t \mathbf{W}_t^\top) \Theta \mathbf{T}_1 \\ & + \mathbf{T}_2 \Theta^\top (\mathbf{I}_d - \mathbf{W}_t \mathbf{W}_t^\top) \\ & \quad \times \mathbb{E}_t [e_{t+1}^4 \|\mathbf{W}_t^\top \mathbf{x}_{t+1}\|_2^4 \|\mathbf{T}_1 \Theta^\top (\mathbf{I}_d - \mathbf{W}_t \mathbf{W}_t^\top) \mathbf{x}_{t+1}\|_2^2 \mathbf{x}_{t+1} \mathbf{x}_{t+1}^\top] (\mathbf{I}_d - \mathbf{W}_t \mathbf{W}_t^\top) \Theta \mathbf{T}_2 \\ & \stackrel{(f)}{\preceq} C r_s^2 (\text{Tr}(\mathbf{T}_2^2(\mathbf{I}_{\text{rk}} - \mathbf{G}_t)) \mathbf{T}_1 (\mathbf{I}_{\text{rk}} - \mathbf{G}_t) \mathbf{T}_1 + \text{Tr}(\mathbf{T}_1^2(\mathbf{I}_{\text{rk}} - \mathbf{G}_t)) \mathbf{T}_2 (\mathbf{I}_{\text{rk}} - \mathbf{G}_t) \mathbf{T}_2), \end{aligned}$$

where we used the Cauchy-Schwartz inequality in (f).

Fourth item. We define

$$\begin{aligned} & \mathbf{T}_1 \Theta^\top \nabla_{\text{St}} \mathbf{L}_{t+1} \mathcal{P}_{t+1} \mathbf{M}_t^\top \mathbf{T}_2 \\ & = \underbrace{e_{t+1}^3 \|(\mathbf{I}_d - \mathbf{W}_t \mathbf{W}_t^\top) \mathbf{x}_{t+1} \|_2^2 \| \mathbf{W}_t^\top \mathbf{x}_{t+1} \|_2^2 \mathbf{T}_1 \Theta^\top (\mathbf{I}_d - \mathbf{W}_t \mathbf{W}_t^\top) \mathbf{x}_{t+1}}_{:= \mathbf{u}_{t+1}} \underbrace{\mathbf{x}_{t+1}^\top \mathbf{W}_t \mathbf{W}_t^\top \Theta \mathbf{T}_2}_{:= \mathbf{v}_{t+1}^\top}. \end{aligned}$$

We have for $p \geq 2$

$$\begin{aligned} & \mathbb{E}_t \left[\left\| \mathbb{1}_{|e_{t+1}| \leq L} \mathbf{u}_{t+1} \mathbf{v}_{t+1}^\top \right\|_2^p \right] \\ & \leq L^{3p} \mathbb{E}_t \left[\left\| (\mathbf{I}_d - \mathbf{W}_t \mathbf{W}_t^\top) \mathbf{x}_{t+1} \right\|_2^{2p} \left\| \mathbf{T}_1 \Theta^\top (\mathbf{I}_d - \mathbf{W}_t \mathbf{W}_t^\top) \mathbf{x}_{t+1} \right\|_2^p \right] \\ & \quad \times \mathbb{E}_t \left[\left\| \mathbf{W}_t^\top \mathbf{x}_{t+1} \right\|_2^{2p} \left\| \mathbf{T}_2 \Theta^\top \mathbf{W}_t \mathbf{W}_t^\top \mathbf{x}_{t+1} \right\|_2^p \right] \\ & \leq L^{3p} (2p)^p (\sqrt{3}d)^p p^{\frac{p}{2}} \left(\sqrt{3} \text{Tr}(\mathbf{T}_1^2(\mathbf{I}_{\text{rk}} - \mathbf{G}_t)) \right)^{\frac{p}{2}} (2p)^p (\sqrt{3}r_s)^p p^{\frac{p}{2}} \left(\sqrt{3} \text{Tr}(\mathbf{T}_2^2 \mathbf{G}_t) \right)^{\frac{p}{2}} \\ & = L^{3p} (12\sqrt{3})^p p^{3p} \left(dr_s \sqrt{\text{Tr}(\mathbf{T}_1^2(\mathbf{I}_{\text{rk}} - \mathbf{G}_t)) \text{Tr}(\mathbf{T}_2^2 \mathbf{G}_t)} \right)^p. \end{aligned}$$

By Proposition 33, we have for $u \geq (2e)^3$

$$\mathbb{P}_t \left[\left\| \mathbb{1}_{|e_{t+1}| \leq L} \mathbf{u}_{t+1} \mathbf{v}_{t+1}^\top \right\|_2 \geq u L^3 12\sqrt{3} dr_s \sqrt{\text{Tr}(\mathbf{T}_1^2(\mathbf{I}_{\text{rk}} - \mathbf{G}_t)) \text{Tr}(\mathbf{T}_2^2 \mathbf{G}_t)} \right] \leq e^{-\frac{u^{1/3}}{e}}.$$

By choosing $u = \frac{L^3}{24\sqrt{3}}$, we have the probability bound. For the variance bound, we have

$$\mathbb{E}_t \left[\left(\text{Sym}(\mathbf{T}_1 \Theta^\top \nabla_{\text{St}} \mathbf{L}_{t+1} \mathcal{P}_{t+1} \mathbf{M}_t^\top \mathbf{T}_2) \right)^2 \right] = \mathbb{E}_t \left[\text{Sym}(\mathbf{u}_{t+1} \mathbf{v}_{t+1}^\top)^2 \right].$$

By using Proposition 22, we have

$$\begin{aligned} & \mathbb{E}_t \left[\text{Sym}(\mathbf{u}_{t+1} \mathbf{v}_{t+1}^\top)^2 \right] \\ & \preceq \mathbf{T}_1 \Theta^\top (\mathbf{I}_d - \mathbf{W}_t \mathbf{W}_t^\top) \\ & \quad \times \mathbb{E}_t [e_{t+1}^6 \|(\mathbf{I}_d - \mathbf{W}_t \mathbf{W}_t^\top) \mathbf{x}_{t+1} \|_2^4 \| \mathbf{W}_t^\top \mathbf{x}_{t+1} \|_2^4 \|\mathbf{T}_2 \Theta^\top \mathbf{W}_t \mathbf{W}_t^\top \mathbf{x}_{t+1}\|_2^2 \mathbf{x}_{t+1} \mathbf{x}_{t+1}^\top] \\ & \quad \times (\mathbf{I}_d - \mathbf{W}_t \mathbf{W}_t^\top) \Theta \mathbf{T}_1 \end{aligned}$$

$$\begin{aligned}
& + \mathbf{T}_2 \Theta^\top \mathbf{W}_t \mathbf{W}_t^\top \\
& \times \mathbb{E}_t \left[e_{t+1}^6 \|(\mathbf{I}_d - \mathbf{W}_t \mathbf{W}_t^\top) \mathbf{x}_{t+1}\|_2^4 \|\mathbf{W}_t^\top \mathbf{x}_{t+1}\|_2^4 \|\mathbf{T}_1 \Theta^\top (\mathbf{I}_d - \mathbf{W}_t \mathbf{W}_t^\top) \mathbf{x}_{t+1}\|_2^2 \mathbf{x}_{t+1}^\top \mathbf{x}_{t+1} \right] \\
& \times \mathbf{W}_t \mathbf{W}_t^\top \Theta \mathbf{T}_2 \\
& \preceq C d^2 r_s^2 \left(\text{Tr}(\mathbf{T}_2^2 \mathbf{G}_t) \mathbf{T}_1 (\mathbf{I}_{\text{rk}} - \mathbf{G}_t) \mathbf{T}_1 + \text{Tr}(\mathbf{T}_1^2 (\mathbf{I}_{\text{rk}} - \mathbf{G}_t)) \mathbf{T}_2 \mathbf{G}_t \mathbf{T}_2 \right).
\end{aligned}$$

Fifth item. We define

$$\begin{aligned}
& \mathbf{T}_1 \Theta^\top \nabla_{\text{St}} \mathbf{L}_{t+1} \mathcal{P}_{t+1} \nabla_{\text{St}} \mathbf{L}_{t+1}^\top \Theta \mathbf{T}_2 \\
& = \underbrace{e_{t+1}^4 \|(\mathbf{I}_d - \mathbf{W}_t \mathbf{W}_t^\top) \mathbf{x}_{t+1}\|_2^2 \|\mathbf{W}_t^\top \mathbf{x}_{t+1}\|_2^4 \mathbf{T}_1 \Theta^\top (\mathbf{I}_d - \mathbf{W}_t \mathbf{W}_t^\top) \mathbf{x}_{t+1}}_{:= \mathbf{u}_{t+1}} \\
& \times \underbrace{\mathbf{x}_{t+1}^\top (\mathbf{I}_d - \mathbf{W}_t \mathbf{W}_t^\top) \Theta \mathbf{T}_2}_{:= \mathbf{v}_{t+1}^\top}.
\end{aligned}$$

We have for $p \geq 2$

$$\begin{aligned}
& \mathbb{E}_t \left[\left\| \mathbb{1}_{|e_{t+1}| \leq L} \mathbf{u}_{t+1} \mathbf{v}_{t+1}^\top \right\|_2^p \right] \\
& \leq L^{4p} \mathbb{E}_t \left[\left\| (\mathbf{I}_d - \mathbf{W}_t \mathbf{W}_t^\top) \mathbf{x}_{t+1} \right\|_2^{2p} \|\mathbf{W}_t^\top \mathbf{x}_{t+1}\|_2^{4p} \right. \\
& \quad \left. \times \|\mathbf{T}_1 \Theta^\top (\mathbf{I}_d - \mathbf{W}_t \mathbf{W}_t^\top) \mathbf{x}_{t+1}\|_2^p \|\mathbf{T}_1 \Theta^\top (\mathbf{I}_d - \mathbf{W}_t \mathbf{W}_t^\top) \mathbf{x}_{t+1}\|_2^p \right] \\
& \leq L^{4p} \mathbb{E}_t \left[\left\| (\mathbf{I}_d - \mathbf{W}_t \mathbf{W}_t^\top) \mathbf{x}_{t+1} \right\|_2^{4p} \right]^{\frac{1}{2}} \mathbb{E}_t \left[\left\| \mathbf{T}_1 \Theta^\top (\mathbf{I}_d - \mathbf{W}_t \mathbf{W}_t^\top) \mathbf{x}_{t+1} \right\|_2^{4p} \right]^{\frac{1}{4}} \\
& \quad \times \mathbb{E}_t \left[\left\| \mathbf{T}_2 \Theta^\top (\mathbf{I}_d - \mathbf{W}_t \mathbf{W}_t^\top) \mathbf{x}_{t+1} \right\|_2^{4p} \right]^{\frac{1}{4}} \mathbb{E}_t \left[\left\| \mathbf{W}_t^\top \mathbf{x}_{t+1} \right\|_2^{4p} \right] \\
& \leq L^{4p} (2p)^p (\sqrt{3}d)^p (2p)^p \left(3 \text{Tr}(\mathbf{T}_1^2 (\mathbf{I}_{\text{rk}} - \mathbf{G}_t)) \text{Tr}(\mathbf{T}_2^2 (\mathbf{I}_{\text{rk}} - \mathbf{G}_t)) \right)^{\frac{p}{2}} (2p)^{2p} (\sqrt{3}r_s)^{2p} \\
& = L^{4p} (2\sqrt{3})^{4p} p^{4p} \left(d r_s^2 \sqrt{\text{Tr}(\mathbf{T}_1^2 (\mathbf{I}_{\text{rk}} - \mathbf{G}_t)) \text{Tr}(\mathbf{T}_2^2 (\mathbf{I}_{\text{rk}} - \mathbf{G}_t))} \right)^p.
\end{aligned}$$

By Proposition 33, we have for $u \geq (2e)^4$

$$\mathbb{P}_t \left[\left\| \mathbb{1}_{|e_{t+1}| \leq L} \mathbf{u}_{t+1} \mathbf{v}_{t+1}^\top \right\|_2 \geq u L^4 (2\sqrt{3})^4 d r_s^2 \sqrt{\text{Tr}(\mathbf{T}_1^2 (\mathbf{I}_{\text{rk}} - \mathbf{G}_t)) \text{Tr}(\mathbf{T}_2^2 (\mathbf{I}_{\text{rk}} - \mathbf{G}_t))} \right] \leq 2e^{-\frac{u^{1/4}}{\epsilon}}.$$

By choosing $u = \frac{L^4}{(2\sqrt{3})^4}$, we have the probability bound. For the variance bound, we have

$$\mathbb{E}_t \left[\left(\text{Sym} \left(\mathbf{T}_1 \Theta^\top \nabla_{\text{St}} \mathbf{L}_{t+1} \mathcal{P}_{t+1} \nabla_{\text{St}} \mathbf{L}_{t+1}^\top \Theta \mathbf{T}_2 \right) \right)^2 \right] = \mathbb{E}_t \left[\text{Sym} \left(\mathbf{u}_{t+1} \mathbf{v}_{t+1}^\top \right)^2 \right].$$

By using Proposition 22, we have

$$\begin{aligned}
& \mathbb{E}_t \left[\text{Sym} \left(\mathbf{u}_{t+1} \mathbf{v}_{t+1}^\top \right)^2 \right] \\
& \preceq \mathbf{T}_1 \Theta^\top (\mathbf{I}_d - \mathbf{W}_t \mathbf{W}_t^\top) \\
& \quad \times \mathbb{E}_t \left[e_{t+1}^8 \|(\mathbf{I}_d - \mathbf{W}_t \mathbf{W}_t^\top) \mathbf{x}_{t+1}\|_2^4 \|\mathbf{W}_t^\top \mathbf{x}_{t+1}\|_2^8 \|\mathbf{T}_2 \Theta^\top (\mathbf{I}_d - \mathbf{W}_t \mathbf{W}_t^\top) \mathbf{x}_{t+1}\|_2^2 \mathbf{x}_{t+1}^\top \mathbf{x}_{t+1} \right] \\
& \quad \times (\mathbf{I}_d - \mathbf{W}_t \mathbf{W}_t^\top) \Theta \mathbf{T}_1 \\
& + \mathbf{T}_2 \Theta^\top (\mathbf{I}_d - \mathbf{W}_t \mathbf{W}_t^\top) \\
& \quad \times \mathbb{E}_t \left[e_{t+1}^8 \|(\mathbf{I}_d - \mathbf{W}_t \mathbf{W}_t^\top) \mathbf{x}_{t+1}\|_2^4 \|\mathbf{W}_t^\top \mathbf{x}_{t+1}\|_2^8 \|\mathbf{T}_1 \Theta^\top (\mathbf{I}_d - \mathbf{W}_t \mathbf{W}_t^\top) \mathbf{x}_{t+1}\|_2^2 \mathbf{x}_{t+1}^\top \mathbf{x}_{t+1} \right] \\
& \quad \times (\mathbf{I}_d - \mathbf{W}_t \mathbf{W}_t^\top) \Theta \mathbf{T}_2 \\
& \preceq C d^2 r_s^4 \left(\text{Tr}(\mathbf{T}_2^2 (\mathbf{I}_{\text{rk}} - \mathbf{G}_t)) \mathbf{T}_1 (\mathbf{I}_{\text{rk}} - \mathbf{G}_t) \mathbf{T}_1 + \text{Tr}(\mathbf{T}_1^2 (\mathbf{I}_{\text{rk}} - \mathbf{G}_t)) \mathbf{T}_2 (\mathbf{I}_{\text{rk}} - \mathbf{G}_t) \mathbf{T}_2 \right).
\end{aligned}$$

□

By recalling the definitions $\{\mathbf{T}_t\}_{t \in \mathbb{N}}$, $\mathbf{A}$, $\mathbf{A}_{11}$ in Sections F.2 and F.3, we define the event:

$$\mathcal{A}_{t+1} := \begin{cases} \mathcal{A}_{t+1} \left(\mathbf{T}_t^{-\frac{1}{2}}, \mathbf{T}_t^{-\frac{1}{2}} \right) \cap \mathcal{A}_{t+1} \left(\mathbf{A}^{\frac{1}{2}} \mathbf{T}_t^{-\frac{1}{2}}, \mathbf{T}_t^{-\frac{1}{2}} \mathbf{A}^{\frac{1}{2}} \right) \cap \{e_{t+1} \leq L\}, & \alpha \in [0, 0.5) \\ \mathcal{A}_{t+1} \left(\mathbf{T}_t^{-\frac{1}{2}}, \mathbf{T}_t^{-\frac{1}{2}} \right) \cap \mathcal{A}_{t+1} \left(\mathbf{A}_{11}^{\frac{1}{2}} \mathbf{T}_t^{-\frac{1}{2}}, \mathbf{T}_t^{-\frac{1}{2}} \mathbf{A}_{11}^{\frac{1}{2}} \right) \cap \{e_{t+1} \leq L\}, & \alpha > 0.5. \end{cases}$$

We define the events $\mathcal{B}_{t+1}$, $\mathcal{C}_{t+1}$, $\mathcal{D}_{t+1}$, and $\mathcal{F}_{t+1}$ in the same way. Based on these events, we define the clipped versions of the noise matrices:

$$\begin{aligned} \mathbf{A}_{t+1} &:= \text{Sym} \left(\Theta^\top \nabla_{\text{St}} \mathbf{L}_{t+1} \mathbf{M}_t^\top \mathbb{1}_{\mathcal{A}_{t+1}} - \mathbb{E}_t \left[\Theta^\top \nabla_{\text{St}} \mathbf{L}_{t+1} \mathbf{M}_t^\top \mathbb{1}_{\mathcal{A}_{t+1}} \right] \right) \\ \mathbf{B}_{t+1} &:= \frac{\mathbf{M}_t \mathbf{P}_{t+1} \mathbf{M}_t^\top \mathbb{1}_{\mathcal{B}_{t+1}}}{1 + c_{t+1}^2} - \mathbb{E}_t \left[\frac{\mathbf{M}_t \mathbf{P}_{t+1} \mathbf{M}_t^\top \mathbb{1}_{\mathcal{B}_{t+1}}}{1 + c_{t+1}^2} \right] \\ \mathbf{C}_{t+1} &:= \Theta^\top \nabla_{\text{St}} \mathbf{L}_{t+1} \nabla_{\text{St}} \mathbf{L}_{t+1}^\top \Theta \mathbb{1}_{\mathcal{C}_{t+1}} - \mathbb{E}_t \left[\Theta^\top \nabla_{\text{St}} \mathbf{L}_{t+1} \nabla_{\text{St}} \mathbf{L}_{t+1}^\top \Theta \mathbb{1}_{\mathcal{C}_{t+1}} \right] \\ \mathbf{D}_{t+1} &:= \text{Sym} \left(\frac{\Theta^\top \nabla_{\text{St}} \mathbf{L}_{t+1} \mathbf{P}_{t+1} \mathbf{M}_t^\top \mathbb{1}_{\mathcal{D}_{t+1}}}{1 + c_{t+1}^2} - \mathbb{E}_t \left[\frac{\Theta^\top \nabla_{\text{St}} \mathbf{L}_{t+1} \mathbf{P}_{t+1} \mathbf{M}_t^\top \mathbb{1}_{\mathcal{D}_{t+1}}}{1 + c_{t+1}^2} \right] \right) \\ \mathbf{F}_{t+1} &:= \frac{\Theta^\top \nabla_{\text{St}} \mathbf{L}_{t+1} \mathbf{P}_{t+1} \nabla_{\text{St}} \mathbf{L}_{t+1}^\top \Theta \mathbb{1}_{\mathcal{F}_{t+1}}}{1 + c_{t+1}^2} - \mathbb{E}_t \left[\frac{\Theta^\top \nabla_{\text{St}} \mathbf{L}_{t+1} \mathbf{P}_{t+1} \nabla_{\text{St}} \mathbf{L}_{t+1}^\top \Theta \mathbb{1}_{\mathcal{F}_{t+1}}}{1 + c_{t+1}^2} \right] \quad (\text{F.38}) \end{aligned}$$

Let $\mathbf{X} \in \left\{ \frac{\eta/2}{\sqrt{r_s}} \mathbf{A}, \frac{\eta^2/16}{r_s} \mathbf{B}, \frac{\eta^2/16}{r_s} \mathbf{C}, \frac{\eta^3/32}{r_s^{3/2}} \mathbf{D}, \frac{\eta^4/256}{r_s^2} \mathbf{F} \right\}$ and

$$\Gamma_1 := \begin{cases} \mathbf{I}_r, & \alpha \in [0, 0.5) \\ \mathbf{I}_{r_u}, & \alpha > 0.5 \end{cases} \quad \Gamma_2 := \begin{cases} \mathbf{A}^{\frac{1}{2}}, & \alpha \in [0, 0.5) \\ \mathbf{A}_{11}^{\frac{1}{2}}, & \alpha > 0.5. \end{cases}$$

For $\ell \in \{1, 2\}$, we define:

$$\text{Quad}_{k,t}^{(\ell)}(\mathbf{X}) := \sum_{j=1}^k \mathbb{E}_{j-1} \left[\left(\Gamma_\ell \mathbf{T}_t^{-\frac{1}{2}} \mathbf{X}_j \mathbf{T}_t^{-\frac{1}{2}} \Gamma_\ell \right)^2 \right].$$

We have the following corollary.

Corollary 4. Let $\text{rk} \in \{r, r_u\}$, $\text{rk}_* \in \{r, r_s\}$ and

$$\mathbf{S}_t := \eta \sum_{j=1}^t \mathbf{G}_{j-1} \quad \text{and} \quad \eta = \frac{\eta}{\sqrt{r_s} \|\mathbf{A}\|_F} \quad \text{and} \quad r_u = \lceil \log^{2.5} d \rceil.$$

Assume the following conditions hold:

- $\mathbf{T}_t \succeq \frac{\kappa_d r_s}{d} \mathbf{I}_{\text{rk}}$,
- $\mathbf{T}_t^{-\frac{1}{2}} \mathbf{S}_t \mathbf{T}_t^{-\frac{1}{2}} \preceq \frac{C_{ub}}{\kappa_d} \begin{cases} (2\eta t \vee r^\alpha) \mathbf{I}_r, & \alpha \in [0, 0.5) \\ r_s^\alpha \log d \mathbf{I}_{r_u}, & \alpha > 0.5, \end{cases}$
- $\mathbf{T}_t^{-\frac{1}{2}} \mathbf{G}_j \mathbf{T}_t^{-\frac{1}{2}} \preceq \frac{C_{ub}}{\kappa_d} \mathbf{I}_{\text{rk}}$ for $j \leq t-1$.

Let

$$\begin{cases} \mathbf{p}_1 = 1 \text{ and } \mathbf{p}_2 = 1 - \alpha, & \alpha \in [0, 1) \\ \mathbf{p}_1 = 1 \text{ and } \mathbf{p}_2 = \frac{2 \log \log r_u}{\log r_u} & \alpha = 1 \\ \mathbf{p}_1 = 1 \text{ and } \mathbf{p}_2 = 0 & \alpha > 1. \end{cases}$$

For $\eta t \leq \frac{1}{2} \text{rk}_*^\alpha \log \left(\frac{d \log^{1.5} d}{r_s} \right)$, the following results hold:

(a) **Quadratic variation bounds.** We have:

$$\|Quad_{t,t}^{(\ell)}(X)\|_2 \leq \frac{C \log d \|\Lambda\|_F \text{rk}^{p_\ell} \text{rk}_*^\alpha}{\kappa_d^2 r_s^{3/2}} \begin{cases} C_{ub} \eta d, & X = \frac{\eta/2}{\sqrt{r_s}} A \\ C_{ub}^2 \eta^3 d^2, & X = \frac{\eta^2/16}{r_s} B \\ \eta^3 d^2, & X = \frac{\eta^2/16}{r_s} C \\ C_{ub} \eta^5 d^3, & X = \frac{\eta^3/32}{r_s^{3/2}} D \\ \eta^7 d^2, & X = \frac{\eta^4/256}{r_s^2} F. \end{cases}$$

(b) **Operator norm bounds.** For $L \geq 8\sqrt{2}e$, there exists $C > 0$ such that

$$\left\| \Gamma_\ell T_t^{-\frac{1}{2}} X_j T_t^{-\frac{1}{2}} \Gamma_\ell \right\|_2 \leq r_{j,t}^{(\ell)}(X) := \begin{cases} \frac{\eta L^2}{2\sqrt{r_s}} \sqrt{\text{Tr}(\Gamma_\ell^2 T_t^{-1}) \text{Tr}(\Gamma_\ell^2 T_t^{-\frac{1}{2}} G_{j-1} T_t^{-\frac{1}{2}})}, & X = \frac{\eta/2}{\sqrt{r_s}} A \\ \frac{\eta^2 L^4}{16 r_s} d \text{Tr}(\Gamma_\ell^2 T_t^{-\frac{1}{2}} G_{j-1} T_t^{-\frac{1}{2}}), & X = \frac{\eta^2/16}{r_s} B \\ \frac{\eta^2 L^4}{16} \text{Tr}(\Gamma_\ell^2 T_t^{-1}), & X = \frac{\eta^2/16}{r_s} C \\ \frac{\eta^3 L^6}{32\sqrt{r_s}} d \sqrt{\text{Tr}(\Gamma_\ell^2 T_t^{-1}) \text{Tr}(\Gamma_\ell^2 T_t^{-\frac{1}{2}} G_{j-1} T_t^{-\frac{1}{2}})}, & X = \frac{\eta^3/32}{r_s^{3/2}} D \\ \frac{\eta^4 L^8}{256} d \text{Tr}(\Gamma_\ell^2 T_t^{-1}), & X = \frac{\eta^4/256}{r_s^2} F \end{cases}$$

$$\leq \frac{C}{\kappa_d} \frac{\text{rk}^{p_\ell}}{r_s} \begin{cases} L^2 \sqrt{C_{ub} \eta} \sqrt{d}, & X = \frac{\eta/2}{\sqrt{r_s}} A \\ L^4 C_{ub} \eta^2 d, & X = \frac{\eta^2/16}{r_s} B \\ L^4 \eta^2 d, & X = \frac{\eta^2/16}{r_s} C \\ L^6 \sqrt{C_{ub} \eta} d^{3/2}, & X = \frac{\eta^3/32}{r_s^{3/2}} D \\ L^8 \eta^4 d^2, & X = \frac{\eta^4/256}{r_s^2} F. \end{cases}$$

Proof. Quadratic variation bounds. We will use the variance bounds given in Proposition 18. For $X = \frac{\eta/2}{\sqrt{r_s}} A$, we have

$$\begin{aligned} Quad_{t,t}^{(\ell)}\left(\frac{\eta/2}{\sqrt{r_s}} A\right) &\preceq \frac{C \eta \|\Lambda\|_F}{\sqrt{r_s}} \left(\text{Tr}(\Gamma_\ell^2 T_t^{-\frac{1}{2}} S_t T_t^{-\frac{1}{2}}) \Gamma_\ell T_t^{-1} \Gamma_\ell + \text{Tr}(\Gamma_\ell^2 T_t^{-1}) \Gamma_\ell T_t^{-\frac{1}{2}} S_t T_t^{-\frac{1}{2}} \Gamma_\ell \right) \\ &\preceq \frac{C \eta \|\Lambda\|_F}{\sqrt{r_s}} \frac{C_{ub} \text{rk}^{p_\ell} \text{rk}_*^\alpha \log d}{\kappa_d^2} \frac{d}{r_s} I_{\text{rk}}. \end{aligned}$$

For $X = \frac{\eta^2/16}{r_s} B$, we have

$$\begin{aligned} Quad_{t,t}^{(\ell)}\left(\frac{\eta^2/16}{r_s} B\right) &\preceq \frac{C C_{ub} \eta^3 d^2 \|\Lambda\|_F}{r_s^{3/2}} \sup_{j \leq t} \left(\text{Tr}(\Gamma_\ell^2 T_t^{-\frac{1}{2}} G_{j-1} T_t^{-\frac{1}{2}}) \right) \Gamma_\ell T_t^{-\frac{1}{2}} S_t T_t^{-\frac{1}{2}} \Gamma_\ell \\ &\preceq \frac{C \eta^3 d^2 \|\Lambda\|_F}{r_s^{3/2}} \frac{C_{ub}^2 \text{rk}^{p_\ell} \text{rk}_*^\alpha \log d}{\kappa_d^2} I_{\text{rk}}. \end{aligned}$$

For $X = \frac{\eta^2/16}{r_s} C$, we have

$$Quad_{t,t}^{(\ell)}\left(\frac{\eta^2/16}{r_s} C\right) \preceq C \eta^4 t \text{Tr}(\Gamma_\ell^2 T_t^{-1}) \Gamma_\ell T_t^{-1} \Gamma_\ell \preceq \frac{C \eta^3 d^2 \|\Lambda\|_F}{r_s^{3/2}} \frac{\text{rk}^{p_\ell} \text{rk}_*^\alpha \log d}{\kappa_d^2} I_{\text{rk}}.$$

For $X = \frac{\eta^3/32}{r_s^{3/2}} D$, we have

$$\begin{aligned} Quad_{t,t}^{(\ell)}\left(\frac{\eta^3/32}{r_s^{3/2}} D\right) &\preceq \frac{C \eta^5 d^2 \|\Lambda\|_F}{\sqrt{r_s}} \left(\text{Tr}(\Gamma_\ell^2 T_t^{-\frac{1}{2}} S_t T_t^{-\frac{1}{2}}) \Gamma_\ell T_t^{-1} \Gamma_\ell + \text{Tr}(\Gamma_\ell^2 T_t^{-1}) \Gamma_\ell T_t^{-\frac{1}{2}} S_t T_t^{-\frac{1}{2}} \Gamma_\ell \right) \\ &\preceq \frac{C \eta^5 d^2 \|\Lambda\|_F}{\sqrt{r_s}} \frac{C_{ub} \text{rk}^{p_\ell} \text{rk}_*^\alpha \log d}{\kappa_d^2} \frac{d}{r_s} I_{\text{rk}}. \end{aligned}$$

For $\mathbf{X} = \frac{\eta^4/256}{r_s^2} \mathbf{F}$, we have

$$\text{Quad}_{t,t}^{(\ell)}\left(\frac{\eta^4/256}{r_s^2} \mathbf{F}\right) \preceq C \eta^8 d^2 t \text{Tr}(\Gamma_\ell^2 \mathbf{T}_t^{-1}) \Gamma_\ell \mathbf{T}_t^{-1} \Gamma_\ell \preceq \frac{C \eta^7 d^2 \|\mathbf{\Lambda}\|_{\mathbf{F}} \text{rk}^{p_\ell} \text{rk}_*^\alpha \log d}{r_s^{3/2} \kappa_d^2} \mathbf{I}_{\text{rk}}.$$

Operator Norm Bounds. We will use the events defined in Proposition 18. For $\mathbf{X} = \frac{\eta/2}{\sqrt{r_s}} \mathbf{A}$, we have

$$r_{j,t}^{(\ell)}\left(\frac{\eta/2}{\sqrt{r_s}} \mathbf{A}\right) = \frac{\eta/2}{\sqrt{r_s}} L^2 \sqrt{\text{Tr}(\Gamma_\ell^2 \mathbf{T}_t^{-1}) \text{Tr}(\Gamma_\ell^2 \mathbf{T}_t^{-\frac{1}{2}} \mathbf{G}_{j-1} \mathbf{T}_t^{-\frac{1}{2}})} \leq \frac{CL^2}{\kappa_d} \frac{\sqrt{C_{\text{ub}}} \eta \sqrt{d} \text{rk}^{p_\ell}}{r_s}.$$

For $\mathbf{X} = \frac{\eta^2/16}{r_s} \mathbf{B}$, we have

$$r_{j,t}^{(\ell)}\left(\frac{\eta^2/16}{r_s} \mathbf{B}\right) = \frac{\eta^2}{16r_s} L^4 d \text{Tr}(\Gamma_\ell^2 \mathbf{T}_t^{-\frac{1}{2}} \mathbf{G}_{j-1} \mathbf{T}_t^{-\frac{1}{2}}) \leq \frac{CL^4}{\kappa_d} \frac{C_{\text{ub}} \eta^2 d \text{rk}^{p_\ell}}{r_s}.$$

For $\mathbf{X} = \frac{\eta^2/16}{r_s} \mathbf{C}$, we have

$$r_{j,t}^{(\ell)}\left(\frac{\eta^2/16}{r_s} \mathbf{C}\right) = \frac{\eta^2}{16r_s} L^4 \text{Tr}(\Gamma_\ell^2 \mathbf{T}_t^{-1}) \leq \frac{CL^4}{\kappa_d} \frac{\eta^2 d \text{rk}^{p_\ell}}{r_s}.$$

For $\mathbf{X} = \frac{\eta^3/32}{r_s^{3/2}} \mathbf{D}$, we have

$$r_{j,t}^{(\ell)}\left(\frac{\eta^3/32}{r_s^{3/2}} \mathbf{D}\right) = \frac{\eta^3}{32\sqrt{r_s}} L^6 d \sqrt{\text{Tr}(\Gamma_\ell^2 \mathbf{T}_t^{-\frac{1}{2}} \mathbf{G}_{j-1} \mathbf{T}_t^{-\frac{1}{2}}) \text{Tr}(\Gamma_\ell^2 \mathbf{T}_t^{-1})} \leq \frac{CL^6}{\kappa_d} \frac{\sqrt{C_{\text{ub}}} \eta^3 d^{3/2} \text{rk}^{p_\ell}}{r_s}.$$

For $\mathbf{X} = \frac{\eta^4/256}{r_s^2} \mathbf{F}$, we have

$$r_{j,t}^{(\ell)}\left(\frac{\eta^4/256}{r_s^2} \mathbf{F}\right) = \frac{\eta^4}{256} L^8 d \text{Tr}(\Gamma_\ell^2 \mathbf{T}_t^{-1}) \leq \frac{CL^8}{\kappa_d} \frac{\eta^4 d^2 \text{rk}^{p_\ell}}{r_s}.$$

□

Proposition 19. Let $\{\mathbf{Y}_t, t = 1, 2, \dots\}$ be a symmetric-matrix martingale with difference sequence $\{\mathbf{X}_t := \mathbf{Y}_{t+1} - \mathbf{Y}_t, t = 1, 2, \dots\}$, whose values are symmetric matrices with dimension $r \leq d$. Let $\{\mathbf{T}_t, t = 1, 2, \dots\}$ be a deterministic sequence, whose values are positive semi-definite matrices with the same dimensionality. Assume that the difference sequence is uniformly bounded in the sense that for a predictable triangular sequence $\{r_{j,t}\}_{j \leq t}$, we have

$$\lambda_{\max}(\mathbf{T}_t^{-\frac{1}{2}} \mathbf{X}_j \mathbf{T}_t^{-\frac{1}{2}}) \leq r_{j,t} \text{ for } j = 1, 2, \dots, t.$$

Define the predictable uniform bound and quadratic variation process of the martingale:

$$R_{k,t} := \max_{j \leq k} r_{j,t} \text{ and } \text{Quad}_{k,t}(\mathbf{X}) := \sum_{j=1}^k \mathbb{E}_{j-1} \left[\left(\mathbf{T}_t^{-\frac{1}{2}} \mathbf{X}_j \mathbf{T}_t^{-\frac{1}{2}} \right)^2 \right] \text{ for } k \leq t = 1, 2, \dots.$$

Let $\mathcal{T} \leq d^3$ be a bounded stopping time. Then, for any deterministic $\sigma^2, \tilde{L} > 0$

$$\begin{aligned} \mathbb{P} \left[\exists t \leq \mathcal{T}, \mathbf{Y}_t \not\preceq u \mathbf{T}_t \text{ and } \max_{t \leq \mathcal{T}} \|\text{Quad}_{t,t}(\mathbf{X})\|_2 \leq \sigma^2 \text{ and } \max_{t \leq \mathcal{T}} R_{t,t} \leq \tilde{L} \right] \\ \leq d^4 \cdot \exp \left(\frac{-u^2/2}{\sigma^2 + \tilde{L}u/3} \right). \end{aligned}$$

Proof. We have

$$\begin{aligned} \mathcal{E}_{\text{target}} &\equiv \left\{ \exists t \leq \mathcal{T}, \mathbf{Y}_t \not\preceq u \mathbf{T}_t \text{ and } \max_{t \leq \mathcal{T}} \|\text{Quad}_{t,t}(\mathbf{X})\|_2 \leq \sigma^2 \text{ and } \max_{t \leq \mathcal{T}} R_{t,t} \leq \tilde{L} \right\} \\ &\subseteq \bigcup_{t=0}^{\mathcal{T}} \left\{ \exists k \leq t, \mathbf{Y}_k \not\preceq u \mathbf{T}_t \text{ and } \|\text{Quad}_{t,t}(\mathbf{X})\|_2 \leq \sigma^2 \text{ and } R_{t,t} \leq \tilde{L} \right\}. \end{aligned}$$

Therefore, we have

$$\mathbb{P}[\mathcal{E}_{\text{target}}] \leq \sum_{n=1}^{d^3} \mathbb{P} \left[\exists k \leq t, \mathbf{Y}_k \not\leq u \mathbf{T}_t \text{ and } \|\text{Quad}_{t,t}(\mathbf{X})\|_2 \leq \sigma^2 \text{ and } R_{t,t} \leq \tilde{L} \right]. \quad (\text{F.39})$$

In the following, we will bound the each term in the right hands-side of (F.39). By [Tro10, Lemma 6.7], we have for $k = 1, \dots, t$ and $\theta > 0$,

$$\mathbb{1}_{R_{k,t} \leq \tilde{L}} \mathbb{E}_{k-1} \left[e^{\frac{\theta}{\tilde{L}} \mathbf{T}_t^{-\frac{1}{2}} \mathbf{X}_k \mathbf{T}_t^{-\frac{1}{2}}} \right] \preceq \mathbb{1}_{R_{k,t} \leq \tilde{L}} \exp \left(\frac{e^\theta - \theta - 1}{\tilde{L}^2} \mathbb{E}_{k-1} \left[(\mathbf{T}_t^{-\frac{1}{2}} \mathbf{X}_k \mathbf{T}_t^{-\frac{1}{2}})^2 \right] \right). \quad (\text{F.40})$$

For notational convenience call $g(\theta) := e^\theta - \theta - 1$. We define a super martingale such that for $0 < k \leq t$,

$$S_k := \text{Tr} \left(\exp \left(\frac{\theta}{\tilde{L}} \mathbf{T}_t^{-\frac{1}{2}} \mathbf{Y}_k \mathbf{T}_t^{-\frac{1}{2}} - \frac{g(\theta)}{\tilde{L}^2} \text{Quad}_{k,t}(\mathbf{X}) \right) \right) \mathbb{1}_{R_{k,t} \leq \tilde{L}},$$

with initial values $R_{0,t} = 0, \mathbf{Y}_0 = \text{Quad}_{0,t} = 0$, and thus, $S_0 = r$. Note that by (F.40) and [Tro10, Corollary 3.3], we can show that $\mathbb{E}_{k-1} S_k \leq S_{k-1}$. We define a stopping time and an event

$$\begin{aligned} \mathcal{T}_{\text{hit}} &:= \{k \geq 0 \mid \lambda_{\max}(\mathbf{T}_t^{-\frac{1}{2}} \mathbf{Y}_k \mathbf{T}_t^{-\frac{1}{2}}) \geq u\} \wedge t, \\ \mathcal{E}_{\text{hit}} &:= \{\mathcal{T}_{\text{hit}} \leq t\} \cap \{\|\text{Quad}_{t,t}(\mathbf{X})\|_2 \leq \sigma^2 \text{ and } R_{t,t} \leq \tilde{L}\}. \end{aligned}$$

We have

$$\begin{aligned} \mathbb{1}_{\mathcal{E}_{\text{hit}}} S_{\mathcal{T}_{\text{hit}}} &\geq \mathbb{1}_{\mathcal{E}_{\text{hit}}} \exp \left(\frac{\theta}{\tilde{L}} u - \frac{g(\theta)}{\tilde{L}^2} \sigma^2 \right) \stackrel{(a)}{\Rightarrow} r \geq \mathbb{P}[\mathcal{E}_{\text{hit}}] \exp \left(\frac{\theta}{\tilde{L}} u - \frac{g(\theta)}{\tilde{L}^2} \sigma^2 \right) \\ &\Rightarrow r \inf_{\theta > 0} \exp \left(-\theta \frac{u}{\tilde{L}} + g(\theta) \frac{\sigma^2}{\tilde{L}^2} \right) \geq \mathbb{P}[\mathcal{E}_{\text{hit}}], \end{aligned}$$

where we use Doob's optional sampling theorem in (a). Since the infimum is attained at $\theta > 0$ and the convex conjugate of $g(\theta)$ is $g^*(\theta) = (\theta + 1) \log(\theta + 1) - \theta$, we have

$$\mathbb{P}[\mathcal{E}_{\text{hit}}] \leq r \cdot \exp \left(-\frac{\sigma^2}{\tilde{L}^2} g^* \left(\frac{u \tilde{L}}{\sigma^2} \right) \right) \leq r \cdot \exp \left(\frac{-u^2/2}{\sigma^2 + \tilde{L}u/3} \right),$$

where we used $g^*(\theta) \geq \frac{u^2/2}{1+u/3}$ in the last step. By $r \leq d$ and (F.39), we have the statement. $\square$

Proposition 20. Let $\mathbb{P}_0$ denote the conditional probability conditioned on $\mathbf{W}_0$. We consider $r_u = \lceil \log^{2.5} d \rceil$, and

$$\begin{aligned} \alpha \in [0, 0.5) : \quad \frac{r_s}{r} &\rightarrow \varphi, \quad \eta \asymp \frac{1}{dr^\alpha \log^{20}(1+d/r_s)}, \quad \kappa_d = \frac{1}{\log^{3.5} d}, \quad C_{ub} = 12 \left(1 + \frac{1}{\sqrt{\varphi}} \right)^2 \\ \alpha > 0.5 : \quad r_s &\asymp 1, \quad \eta \asymp \frac{1}{dr_u^{4\alpha+3} \log^{18} d}, \quad \kappa_d = \frac{1}{r_u \log^{2.5} d}, \quad C_{ub} = 2^\alpha 30 r_u. \end{aligned}$$

For $\alpha \in [0, 0.5)$, we define $\mathcal{T} := \mathcal{T}_{bad} \wedge \frac{1}{2\eta} \left(r_s (1 - \log^{-\frac{1}{2}} d) \wedge r \right)^\alpha \log \left(\frac{d \log^{1.5} d}{r_s} \right)$. We have for $d \geq \Omega_{\alpha, \varphi, \beta}(1)$

$$\begin{aligned} \mathbb{P}_0 \left[\sup_{t \leq \mathcal{T}} \|\mathbf{T}_t^{-\frac{1}{2}} \underline{\nu}_t \mathbf{T}_t^{-\frac{1}{2}}\|_2 \vee r^{\frac{\alpha}{2}} \|\mathbf{\Lambda}^{\frac{1}{2}} \mathbf{T}_t^{-\frac{1}{2}} \underline{\nu}_t \mathbf{T}_t^{-\frac{1}{2}} \mathbf{\Lambda}^{\frac{1}{2}}\|_2 \geq \kappa_d r^{-\frac{\alpha}{2}} \text{ and } \mathcal{G}_{init} \right] \\ \leq 20d^4 \exp(-\log^2 d). \end{aligned}$$

For $\alpha > 0.5$, we set $\mathcal{T} := \mathcal{T}_{bad} \wedge \frac{1}{2\eta} r_s^\alpha \log \left(\frac{d \log^{1.5} d}{r_s} \right)$. We have for $d \geq \Omega_{\alpha, r_s}(1)$

$$\begin{aligned} \mathbb{P}_0 \left[\sup_{t \leq \mathcal{T}} \|\mathbf{T}_t^{-\frac{1}{2}} \underline{\nu}_t \mathbf{T}_t^{-\frac{1}{2}}\|_2 \vee r_u^{\frac{\alpha}{2}} \|\mathbf{\Lambda}_{11}^{\frac{1}{2}} \mathbf{T}_t^{-\frac{1}{2}} \underline{\nu}_t \mathbf{T}_t^{-\frac{1}{2}} \mathbf{\Lambda}_{11}^{\frac{1}{2}}\|_2 \geq \kappa_d r_u^{-\frac{\alpha}{2}} \text{ and } \mathcal{G}_{init} \right] \\ \leq 20d^4 \exp(-\log^2 d). \end{aligned}$$

Proof. For notational convenience, we introduce $\mathcal{X} := \left\{ \frac{\eta/2}{\sqrt{r_s}} \mathbf{A}, \frac{\eta^2/16}{r_s} \mathbf{B}, \frac{\eta^2/16}{r_s} \mathbf{C}, \frac{\eta^3/32}{r_s^{3/2}} \mathbf{D}, \frac{\eta^4/256}{r_s^2} \mathbf{F} \right\}$.

For both cases, we will set the clip threshold to $L = \log^2 d$. We introduce the notation $R_t^{(\ell)} := \max_{\mathbf{X} \in \mathcal{X}} \max_{j \leq t} r_{j,t}^{(\ell)}(\mathbf{X})$ and $\|\text{Quad}_t^{(\ell)}\|_2 := \max_{\mathbf{X} \in \mathcal{X}} \|\text{Quad}_{t,t}^{(\ell)}(\mathbf{X})\|_2$ for $\ell = 1, 2$.

For $\alpha \in [0, 0.5)$, we can write for all $\mathbf{X} \in \mathcal{X}$,

$$\begin{aligned} & \mathbb{P}_0 \left[\sup_{t \leq \mathcal{T}} \left\| \mathbf{T}_t^{-\frac{1}{2}} \underline{\nu}_t \mathbf{T}_t^{-\frac{1}{2}} \right\|_2 \vee r^{\frac{\alpha}{2}} \left\| \mathbf{\Lambda}^{\frac{1}{2}} \mathbf{T}_t^{-\frac{1}{2}} \underline{\nu}_t \mathbf{T}_t^{-\frac{1}{2}} \mathbf{\Lambda}^{\frac{1}{2}} \right\|_2 \geq \kappa_d r^{-\frac{\alpha}{2}} \text{ and } \mathcal{G}_{\text{init}} \right] \\ & \leq \sum_{\mathbf{X} \in \mathcal{X}} \mathbb{P}_0 \left[\sup_{t \leq \mathcal{T}} \left\| \mathbf{T}_t^{-\frac{1}{2}} \left(\sum_{j \leq t} \mathbf{X}_j \right) \mathbf{T}_t^{-\frac{1}{2}} \right\|_2 \geq \frac{\kappa_d r^{-\frac{\alpha}{2}}}{10} \text{ and } \mathcal{G}_{\text{init}} \right] \\ & \quad + \sum_{\mathbf{X} \in \mathcal{X}} \mathbb{P}_0 \left[\sup_{t \leq \mathcal{T}} \left\| \mathbf{\Lambda}^{\frac{1}{2}} \mathbf{T}_t^{-\frac{1}{2}} \left(\sum_{j \leq t} \mathbf{X}_j \right) \mathbf{T}_t^{-\frac{1}{2}} \mathbf{\Lambda}^{\frac{1}{2}} \right\|_2 \geq \frac{\kappa_d r^{-\alpha}}{10} \text{ and } \mathcal{G}_{\text{init}} \right]. \end{aligned}$$

By Propositions 14 and 16 and Corollary 4, $\mathcal{G}_{\text{init}}$ implies the events

$$\begin{aligned} \mathcal{E}_{\text{ht},1} & \equiv \left\{ \max_{t \leq \mathcal{T}} \|\text{Quad}_t^{(1)}\|_2 \leq \frac{O_{\alpha,\beta,\varphi}(r^{-\alpha})}{\log^{12} d} \text{ and } \max_{t \leq \mathcal{T}} R_t^{(1)} \leq \frac{O_{\alpha,\beta,\varphi}(r^{-\alpha})}{\sqrt{d} \log^{12.5} d} \right\} \\ \mathcal{E}_{\text{ht},2} & \equiv \left\{ \max_{t \leq \mathcal{T}} \|\text{Quad}_t^{(2)}\|_2 \leq \frac{O_{\alpha,\beta,\varphi}(r^{-2\alpha})}{\log^{12} d} \text{ and } \max_{t \leq \mathcal{T}} R_t^{(2)} \leq \frac{O_{\alpha,\beta,\varphi}(r^{-2\alpha})}{\sqrt{d} \log^{12.5} d} \right\}. \end{aligned}$$

Therefore, by using Proposition 19

$$\begin{aligned} & \mathbb{P}_0 \left[\sup_{t \leq \mathcal{T}} \left\| \mathbf{T}_t^{-\frac{1}{2}} \left(\sum_{j \leq t} \mathbf{X}_j \right) \mathbf{T}_t^{-\frac{1}{2}} \right\|_2 \geq \frac{\kappa_d r^{-\frac{\alpha}{2}}}{10} \text{ and } \mathcal{G}_{\text{init}} \right] \\ & \leq \mathbb{P}_0 \left[\sup_{t \leq \mathcal{T}} \left\| \mathbf{T}_t^{-\frac{1}{2}} \left(\sum_{j \leq t} \mathbf{X}_j \right) \mathbf{T}_t^{-\frac{1}{2}} \right\|_2 \geq \frac{\kappa_d r^{-\frac{\alpha}{2}}}{10} \text{ and } \mathcal{E}_{\text{ht},1} \right] \\ & \leq 2d^4 \exp(-\log^2 d). \end{aligned}$$

Similarly,

$$\begin{aligned} & \mathbb{P}_0 \left[\sup_{t \leq \mathcal{T}} \left\| \mathbf{\Lambda}^{\frac{1}{2}} \mathbf{T}_t^{-\frac{1}{2}} \left(\sum_{j \leq t} \mathbf{X}_j \right) \mathbf{T}_t^{-\frac{1}{2}} \mathbf{\Lambda}^{\frac{1}{2}} \right\|_2 \geq \frac{\kappa_d r^{-\alpha}}{10} \text{ and } \mathcal{G}_{\text{init}} \right] \\ & \leq \mathbb{P}_0 \left[\sup_{t \leq \mathcal{T}} \left\| \mathbf{\Lambda}^{\frac{1}{2}} \mathbf{T}_t^{-\frac{1}{2}} \left(\sum_{j \leq t} \mathbf{X}_j \right) \mathbf{T}_t^{-\frac{1}{2}} \mathbf{\Lambda}^{\frac{1}{2}} \right\|_2 \geq \frac{\kappa_d r^{-\alpha}}{10} \text{ and } \mathcal{E}_{\text{ht},2} \right] \\ & \leq 2d^4 \exp(-\log^2 d). \end{aligned}$$

For $\alpha > 0.5$, we can write for all $\mathbf{X} \in \mathcal{X}$,

$$\begin{aligned} & \mathbb{P}_0 \left[\sup_{t \leq \mathcal{T}} \left\| \mathbf{T}_t^{-\frac{1}{2}} \underline{\nu}_t \mathbf{T}_t^{-\frac{1}{2}} \right\|_2 \vee r_u^{\frac{\alpha}{2}} \left\| \mathbf{\Lambda}_{11}^{\frac{1}{2}} \mathbf{T}_t^{-\frac{1}{2}} \underline{\nu}_t \mathbf{T}_t^{-\frac{1}{2}} \mathbf{\Lambda}_{11}^{\frac{1}{2}} \right\|_2 \geq \kappa_d r_u^{-\frac{\alpha}{2}} \text{ and } \mathcal{G}_{\text{init}} \right] \\ & \leq \sum_{\mathbf{X} \in \mathcal{X}} \mathbb{P}_0 \left[\sup_{t \leq \mathcal{T}} \left\| \mathbf{T}_t^{-\frac{1}{2}} \left(\sum_{j \leq t} \mathbf{X}_j \right) \mathbf{T}_t^{-\frac{1}{2}} \right\|_2 \geq \frac{\kappa_d r_u^{-\frac{\alpha}{2}}}{10} \text{ and } \mathcal{G}_{\text{init}} \right] \\ & \quad + \sum_{\mathbf{X} \in \mathcal{X}} \mathbb{P}_0 \left[\sup_{t \leq \mathcal{T}} \left\| \mathbf{\Lambda}_{11}^{\frac{1}{2}} \mathbf{T}_t^{-\frac{1}{2}} \left(\sum_{j \leq t} \mathbf{X}_j \right) \mathbf{T}_t^{-\frac{1}{2}} \mathbf{\Lambda}_{11}^{\frac{1}{2}} \right\|_2 \geq \frac{\kappa_d r_u^{-\alpha}}{10} \text{ and } \mathcal{G}_{\text{init}} \right] \end{aligned}$$

By Propositions 14 and 16 and Corollary 4, $\mathcal{G}_{\text{init}}$ implies the events

$$\begin{aligned} \mathcal{E}_{\text{lt},1} & \equiv \left\{ \max_{t \leq \mathcal{T}} \|\text{Quad}_t^{(1)}\|_2 \leq \frac{O_{\alpha,r_s}(r_u^{-4\alpha})}{\log^{12} d} \text{ and } \max_{t \leq \mathcal{T}} R_t^{(1)} \leq \frac{O_{r_s}(r_u^{-4\alpha-1})}{\sqrt{d} \log^{11.5} d} \right\} \\ \mathcal{E}_{\text{lt},2} & \equiv \left\{ \max_{t \leq \mathcal{T}} \|\text{Quad}_t^{(2)}\|_2 \leq \frac{O_{\alpha,r_s}(r_u^{-4\alpha-(\alpha \wedge 1)})}{\log^{12} d \log^{-2} r_u} \text{ and } \max_{t \leq \mathcal{T}} R_t^{(2)} \leq \frac{O_{r_s}(r_u^{-4\alpha-1} r_u^{-(\alpha \wedge 1)})}{\sqrt{d} \log^{11.5} d \log^{-2} r_u} \right\}. \end{aligned}$$

Therefore, by using Proposition 19

$$\begin{aligned} \mathbb{P}_0 \left[\sup_{t \leq \mathcal{T}} \left\| \mathbf{T}_t^{-\frac{1}{2}} \left(\sum_{j \leq t} \mathbf{x}_j \right) \mathbf{T}_t^{-\frac{1}{2}} \right\|_2 \geq \frac{\kappa_d r_u^{-\frac{\alpha}{2}}}{10} \text{ and } \mathcal{G}_{\text{init}} \right] \\ \leq \mathbb{P}_0 \left[\sup_{n \leq \mathcal{T}} \left\| \mathbf{T}_t^{-\frac{1}{2}} \left(\sum_{j \leq t} \mathbf{x}_j \right) \mathbf{T}_t^{-\frac{1}{2}} \right\|_2 \geq \frac{\kappa_d r_u^{-\frac{\alpha}{2}}}{10} \text{ and } \mathcal{E}_{\text{lt},1} \right] \\ \leq 2d^4 \exp(-\log^2 d). \end{aligned}$$

Similarly,

$$\begin{aligned} \mathbb{P}_0 \left[\sup_{t \leq \mathcal{T}} \left\| \mathbf{\Lambda}_{11}^{\frac{1}{2}} \mathbf{T}_t^{-\frac{1}{2}} \left(\sum_{j \leq t} \mathbf{x}_j \right) \mathbf{T}_t^{-\frac{1}{2}} \mathbf{\Lambda}_{11}^{\frac{1}{2}} \right\|_2 \geq \frac{\kappa_d r_u^{-\alpha}}{10} \text{ and } \mathcal{G}_{\text{init}} \right] \\ \leq \mathbb{P}_0 \left[\sup_{t \leq \mathcal{T}} \left\| \mathbf{\Lambda}_{11}^{\frac{1}{2}} \mathbf{T}_t^{-\frac{1}{2}} \left(\sum_{j \leq t} \mathbf{x}_j \right) \mathbf{T}_t^{-\frac{1}{2}} \mathbf{\Lambda}_{11}^{\frac{1}{2}} \right\|_2 \geq \frac{\kappa_d r_u^{-\alpha}}{10} \text{ and } \mathcal{E}_{\text{lt},2} \right] \\ \leq 2d^4 \exp(-\log^2 d). \end{aligned}$$

□

Corollary 5. Consider $\text{rk}_\star = \{r_\star = \lfloor r_s(1 - \log^{-1/2} d) \wedge r \rfloor, r_{u_\star} = r_s\}$ and the parameters in Proposition 20. We have

$$\mathbb{P}_0 \left[\mathcal{T}_{\text{bad}} \geq \frac{1}{2\eta} \text{rk}_\star \log \left(\frac{d \log^{1.5} d}{r_s} \right) \text{ and } \mathcal{G}_{\text{init}} \right] \geq 1 - 20d^4 \exp(-\log^2 d).$$

Proof. By using the first items in Proposition 15 and 16, and Lemma 6, $\mathcal{G}_{\text{init}}$ implies that

$$\mathcal{T}_{\text{bad}} \geq \mathcal{T}_{\text{noise}} \wedge \frac{1}{2\eta} \text{rk}_\star \log \left(\frac{d \log^{1.5} d}{r_s} \right).$$

On the other hand, within the (negation) of the events given in Proposition 20, we have

$$\mathcal{T}_{\text{noise}} > \mathcal{T}_{\text{bad}} \wedge \frac{1}{2\eta} \text{rk}_\star \log \left(\frac{d \log^{1.5} d}{r_s} \right).$$

Therefore, the statement follows. □

F.7 Stability near minima

In this section, we will establish that given (SGD) is near global minimum it will stay near global minimum. For the statement, we (re)introduce the block matrix notation: $\text{rk}_\star = \{r_\star = \lfloor r_s(1 - \log^{-1/8} d) \wedge r \rfloor, r_{u_\star} = r_s\}$, we have

$$\mathbf{G}_t = \begin{bmatrix} \mathbf{G}_{t,11} & \mathbf{G}_{t,12} \\ \mathbf{G}_{t,12}^\top & \mathbf{G}_{t,22} \end{bmatrix} \quad \boldsymbol{\nu}_t = \begin{bmatrix} \boldsymbol{\nu}_{t,11} & \boldsymbol{\nu}_{t,12} \\ \boldsymbol{\nu}_{t,12}^\top & \boldsymbol{\nu}_{t,22} \end{bmatrix} \quad \mathbf{\Lambda} = \begin{bmatrix} \mathbf{\Lambda}_{11} & 0 \\ 0 & \mathbf{\Lambda}_{22} \end{bmatrix}, \quad \mathbf{\Lambda}_{\ell_j} = \begin{bmatrix} \mathbf{\Lambda}_{\ell_j,11} & 0 \\ 0 & \mathbf{\Lambda}_{\ell_j,22} \end{bmatrix},$$

where $\mathbf{G}_{t,11}, \boldsymbol{\nu}_{t,11}, \mathbf{\Lambda}_{11}, \mathbf{\Lambda}_{\ell_j,11} \in \mathbb{R}^{\text{rk}_\star \times \text{rk}_\star}$ and $\mathbf{\Lambda}_{\ell_j}$ is introduced (F.2). We define the following iterations:

- Given $\mathbf{G}_0 = \mathbf{I}_{\text{rk}_\star} - \frac{1}{\log d}$ diagonal and $\mathbf{V}_t = 2\mathbf{\Lambda}_{\ell_2,11}^{\frac{1}{2}} \mathbf{G}_t \mathbf{\Lambda}_{\ell_2,11}^{\frac{1}{2}} - \mathbf{\Lambda}_{\ell_1,11}$, we define

$$\begin{aligned} \mathbf{V}_{t+1} &= \mathbf{V}_t \left(\mathbf{I}_{\text{rk}_\star} + \frac{\eta}{1 - 1.1\eta} \mathbf{V}_t \right)^{-1} \\ &\quad + \frac{\eta}{1 - 1.1\eta} \left(\mathbf{\Lambda}_{\ell_1,11}^2 - \frac{8.1}{\text{rk}_\star^\alpha \log d} \mathbf{\Lambda}_{\ell_2,11} - \frac{O(1)}{\log^2 d} \mathbf{\Lambda}_{\ell_2,11}^2 \right). \end{aligned}$$

- For $\boldsymbol{\nu}_0 = 0$, $\boldsymbol{\nu}_{t+1} = \boldsymbol{\nu}_t + \frac{\eta/2}{\sqrt{r_s}} \boldsymbol{\nu}_{t+1,11}$.

- We define a sequence of events $\{\mathcal{E}_t\}_{t \geq 0}$

$$\mathcal{E}_t := \left\{ \frac{-\text{rk}_*^{-\frac{\alpha}{2}}}{\log^2 d} \mathbf{I}_{\text{rk}_*} \preceq \underline{\nu}_t \preceq \frac{\text{rk}_*^{-\frac{\alpha}{2}}}{\log^2 d} \mathbf{I}_{\text{rk}_*} \right\} \cap \left\{ \frac{-\text{rk}_*^{-\alpha}}{\log^4 d} \mathbf{I}_{\text{rk}_*} \preceq \Lambda_{11}^{\frac{1}{2}} \underline{\nu}_t \Lambda_{11}^{\frac{1}{2}} \preceq \frac{\text{rk}_*^{-\alpha}}{\log^4 d} \mathbf{I}_{\text{rk}_*} \right\},$$

We define the stopping times

$$\mathcal{T}_{\text{noise}}(\omega) := \inf \{t \geq 0 \mid \omega \notin \mathcal{E}_t\} \wedge d^3, \quad \mathcal{T}_{\text{bounded}} := \inf \left\{ t \geq 0 \mid \mathbf{G}_t \not\preceq \mathbf{I}_{\text{rk}_*} - \frac{2}{\log d} \mathbf{I}_{\text{rk}_*} \right\}.$$

and $\mathcal{T}_{\text{stable}} := \mathcal{T}_{\text{noise}} \wedge \mathcal{T}_{\text{bounded}}$.

We have the following statement:

Proposition 21. Consider the parameters in Proposition 20. (SGD) guarantees that if $\mathbf{G}_{0,11} \succeq \mathbf{I}_{\text{rk}_*} - \frac{1}{\log d}$, we have $\mathbf{G}_{t,11} \succeq \mathbf{I}_{\text{rk}_*} - \frac{2}{\log d}$ for $t \leq \frac{\text{rk}_*^\alpha \log^2 d}{\eta}$ with probability $1 - d^4 \exp(-\log^2 d)$.

Proof. We define $\underline{\zeta}_t := 2\Lambda_{\ell_2,11}^{\frac{1}{2}} \underline{\nu}_t \Lambda_{\ell_2,11}^{\frac{1}{2}}$. We make the following observations:

- Since $\mathbf{G}_{t,11}^2 + \mathbf{G}_{t,12} \mathbf{G}_{t,12}^\top \preceq \mathbf{I}_{\text{rk}_*}$ for $t \leq \mathcal{T}_{\text{bounded}}$, we have

$$\mathbf{G}_{t,12} \mathbf{G}_{t,12}^\top \preceq \frac{4}{\log d} \mathbf{I}_{\text{rk}_*}$$

Therefore, by using (F.8), we have for $t \leq \mathcal{T}_{\text{bounded}}$

$$\begin{aligned} \mathbf{G}_{t+1,11} &\succeq \mathbf{G}_{t,11} + \eta \left(\Lambda_{\ell_1,11} \mathbf{G}_{t,11} + \mathbf{G}_{t,11} \Lambda_{\ell_1,11} - 2\mathbf{G}_{t,11} \Lambda_{\ell_2,11} \mathbf{G}_{t,11} \right) - \frac{4\eta}{\text{rk}_*^\alpha \log d} \mathbf{I}_{\text{rk}_*} \\ &\quad - C\eta^2 \|\Lambda\|_{\text{F}}^2 r_s \mathbf{I}_{\text{rk}_*} + \frac{\eta/2}{\sqrt{r_s}} \underline{\nu}_{t+1,11} \end{aligned}$$

Then, if we define $\mathbf{V}_t^- := 2\Lambda_{\ell_2,11}^{\frac{1}{2}} \mathbf{G}_{t,11} \Lambda_{\ell_2,11}^{\frac{1}{2}} - \Lambda_{\ell_1,11}$, we have

$$\begin{aligned} \mathbf{V}_{t+1}^- &\succeq \mathbf{V}_t^- - \eta(\mathbf{V}_{t+1}^-)^2 + \eta \Lambda_{\ell_1,11}^2 - \left(\frac{8\eta}{\text{rk}_*^\alpha \log d} + C\eta^2 \|\Lambda\|_{\text{F}}^2 r_s \right) \Lambda_{\ell_2,11} \\ &\quad + \frac{\eta}{\sqrt{r_s}} \Lambda_{\ell_2,11}^{\frac{1}{2}} \underline{\nu}_{t+1,11} \Lambda_{\ell_2,11}^{\frac{1}{2}} \\ &\succeq \mathbf{V}_t^- \left(\mathbf{I}_r + \frac{\eta}{1-1.1\eta} \mathbf{V}_t^- \right)^{-1} + \eta \left(\Lambda_{\ell_1,11}^2 - \frac{8.1}{\text{rk}_*^\alpha \log d} \Lambda_{\ell_2,11} \right) + \Lambda_{\ell_2,11}^{\frac{1}{2}} \underline{\nu}_{t+1,11} \Lambda_{\ell_2,11}^{\frac{1}{2}} \end{aligned}$$

- To derive an upper-bound for $\underline{\mathbf{G}}_t$, assuming $\underline{\mathbf{G}}_t \preceq 1.1\mathbf{I}_{\text{rk}_*}$, we have

$$\begin{aligned} \underline{\mathbf{V}}_{t+1} &= \underline{\mathbf{V}}_t - \frac{\eta}{1-1.1\eta} \underline{\mathbf{V}}_t^2 + \frac{\eta^2}{(1-1.1\eta)^2} \underline{\mathbf{V}}_t^3 \left(\mathbf{I}_{\text{rk}_*} + \frac{\eta}{1-1.1\eta} \underline{\mathbf{V}}_t \right)^{-1} \\ &\quad + \frac{\eta}{1-1.1\eta} \left(\Lambda_{\ell_1,11}^2 - \frac{8.1}{\text{rk}_*^\alpha \log d} \Lambda_{\ell_2,11} - \frac{O(1)}{\log^2 d} \Lambda_{\ell_2,11}^2 \right) \\ &\preceq \underline{\mathbf{V}}_t - \frac{\eta}{1-1.1\eta} \underline{\mathbf{V}}_t^2 + \frac{\eta}{1-1.1\eta} \Lambda_{\ell_1,11}^2. \end{aligned}$$

Then, by Proposition 34, we have $\underline{\mathbf{G}}_{t+1} \preceq 1.1\mathbf{I}_{\text{rk}_*}$. Since the bound holds for $t = 0$, it holds for all $t \in \mathbb{N}$.

- To derive a lower-bound, we first observe that by monotonicity $\underline{\mathbf{V}}_0 \succ 0$. Therefore, assuming $\underline{\mathbf{V}}_t \succ \underline{\mathbf{V}}_0$

$$\begin{aligned} \underline{\mathbf{V}}_{t+1} &\succeq \underline{\mathbf{V}}_t - \frac{\eta/2}{1-1.1\eta} (2\Lambda_{\ell_2,11}^{\frac{1}{2}} \underline{\mathbf{G}}_t \Lambda_{\ell_2,11}^{\frac{1}{2}} - \Lambda_{\ell_1,11})^2 \\ &\quad + \frac{\eta/2}{1-1.1\eta} \left(\Lambda_{\ell_1,11}^2 - \frac{8.1}{\text{rk}_*^\alpha \log d} \Lambda_{\ell_2,11} - \frac{O(1)}{\log^2 d} \Lambda_{\ell_2,11}^2 \right) \end{aligned}$$

$$\begin{aligned} &\succeq \underline{V}_t - \frac{\eta/2}{1-1.1\eta} \left(2\Lambda_{\ell_2,11}^{\frac{1}{2}} \underline{G}_t \Lambda_{\ell_2,11}^{\frac{1}{2}} - \sqrt{\left(\Lambda_{\ell_1,11}^2 - \frac{8.1}{\text{rk}_*^\alpha \log d} \Lambda_{\ell_2,11} - \frac{O(1)}{\log^2 d} \Lambda_{\ell_2,11}^2 \right)} \right)^2 \\ &\quad + \frac{\eta/2}{1-1.1\eta} \left(\Lambda_{\ell_1,11}^2 - \frac{8.1}{\text{rk}_*^\alpha \log d} \Lambda_{\ell_2,11} - \frac{O(1)}{\log^2 d} \Lambda_{\ell_2,11}^2 \right). \end{aligned}$$

Then, by Proposition 34, we have $\underline{V}_t \succeq \underline{V}_0$.

We start our proof by showing that $\mathbf{V}_t^- \succeq \underline{V}_t + \underline{\zeta}_t$ for $t \leq \mathcal{T}_{\text{stable}}$. Assuming the statement holds for $t \in \mathbb{N}$, we have

$$\begin{aligned} \mathbf{V}_{t+1}^- &\succeq (\underline{V}_t + \underline{\zeta}_t) \left(\mathbf{I}_r + \frac{\eta}{1-1.1\eta} (\underline{V}_t + \underline{\zeta}_t) \right)^{-1} + \eta \left(\Lambda_{\ell_1,11}^2 - \frac{8.1}{\text{rk}_*^\alpha \log d} \Lambda_{\ell_2,11} \right) \\ &\quad + \Lambda_{\ell_2,11}^{\frac{1}{2}} \nu_{t+1,11} \Lambda_{\ell_2,11}^{\frac{1}{2}} \\ &= \underline{V}_t \left(\mathbf{I}_r + \frac{\eta}{1-1.1\eta} \underline{V}_t \right)^{-1} + \eta \left(\Lambda_{\ell_1,11}^2 - \frac{8.1}{\text{rk}_*^\alpha \log d} \Lambda_{\ell_2,11} \right) \\ &\quad + \left(\mathbf{I}_r + \frac{\eta}{1-1.1\eta} \underline{V}_t \right)^{-1} \underline{\zeta}_t \left(\mathbf{I}_r + \frac{\eta}{1-1.1\eta} (\underline{V}_t + \underline{\zeta}_t) \right)^{-1} + \Lambda_{\ell_2,11}^{\frac{1}{2}} \nu_{t+1,11} \Lambda_{\ell_2,11}^{\frac{1}{2}} \end{aligned}$$

We have for $t \leq \mathcal{T}_{\text{stable}}$

$$\begin{aligned} &\left(\mathbf{I}_r + \frac{\eta}{1-1.1\eta} \underline{V}_t \right)^{-1} \underline{\zeta}_t \left(\mathbf{I}_r + \frac{\eta}{1-1.1\eta} (\underline{V}_t + \underline{\zeta}_t) \right)^{-1} \\ &= \underline{\zeta}_t - \frac{\eta}{1-1.1\eta} \underline{V}_t \underline{\zeta}_t - \frac{\eta}{1-1.1\eta} \underline{\zeta}_t \underline{V}_t - \frac{\eta}{1-1.1\eta} \underline{\zeta}_t^2 \\ &\quad - \frac{\eta^2}{(1-1.1\eta)^2} \underline{V}_t \left(\mathbf{I}_r + \frac{\eta}{1-1.1\eta} \underline{V}_t \right)^{-1} \underline{\zeta}_t \left(\mathbf{I}_r + \frac{\eta}{1-1.1\eta} (\underline{V}_t + \underline{\zeta}_t) \right)^{-1} (\underline{V}_t + \underline{\zeta}_t) \\ &\succeq \underline{\zeta}_t - \frac{\eta}{1-1.1\eta} \frac{1}{\log^2 d} \underline{V}_t^2 - \frac{\eta}{1-1.1\eta} (1 + \log^2 d) \underline{\zeta}_t^2 - \frac{\eta^2}{(1-1.1\eta)^2} \frac{O(\text{rk}_*^{-\alpha})}{\log^4 d} \mathbf{I}_{\text{rk}_*} \\ &\succeq \underline{\zeta}_t - \frac{\eta}{1-1.1\eta} \frac{O(1)}{\log^2 d} \Lambda_{\ell_2,11}^2. \end{aligned}$$

Since $\mathbf{V}_0^- = \underline{V}_0 + \underline{\zeta}_0$, the claim follows. Then, by the third item above, we have for $t \leq \mathcal{T}_{\text{stable}}$

$$\mathbf{G}_t \succeq \underline{\mathbf{G}}_0 + \underline{\nu}_t \succeq \mathbf{I}_{\text{rk}_*} - \frac{1}{\log d} \mathbf{I}_{\text{rk}_*} - \frac{\text{rk}_*^{-\frac{\alpha}{2}}}{\log^2 d} \mathbf{I}_{\text{rk}_*} \succeq \mathbf{I}_{\text{rk}_*} - \frac{2}{\log d} \mathbf{I}_{\text{rk}_*} \Rightarrow \mathcal{T}_{\text{noise}} \leq \mathcal{T}_{\text{bounded}}.$$

In the following, we will bound $\mathcal{T}_{\text{noise}}$. We have

$$\mathbb{E}_t [\underline{\nu}_{t+1}^2] \stackrel{(a)}{\preceq} \underline{\nu}_t^2 + O(\eta^2) \mathbf{I}_{\text{rk}_*} \Rightarrow \mathbb{E} [\underline{\nu}_t^2] \preceq O(\eta^2 t) \mathbf{I}_{\text{rk}_*}$$

By clipping strategy we used with $L = \log^2 d$ in (F.38), and defining $\Gamma_1 := \mathbf{I}_{\text{rk}_*}$, $\Gamma_2 := \Lambda_{\ell_1,11}^{\frac{1}{2}}$, and

$$\text{Quad}_{k,t}^{(\ell)}(\mathbf{X}) := \sum_{j=1}^k \mathbb{E}_{j-1} \left[\left(\Gamma_\ell \mathbf{T}_t^{-\frac{1}{2}} \mathbf{X}_j \mathbf{T}_t^{-\frac{1}{2}} \Gamma_\ell \right)^2 \right], \ell \in \{1, 2\},$$

we can show that the following events hold: For any $T \in \mathbb{N}$,

$$\begin{aligned} \widehat{\mathcal{E}}_{\text{ht},1} &\equiv \left\{ \max_{t \leq T} \|\text{Quad}_{t,t}^{(1)}\|_2 \leq O(\eta^2 T) \quad \text{and} \quad \max_{t \leq T} R_{t,t}^{(1)} \leq O(\eta \text{rk}_*^{\frac{1}{2}} \log^2 d) \right\} \\ \widehat{\mathcal{E}}_{\text{ht},2} &\equiv \left\{ \max_{t \leq T} \|\text{Quad}_{t,t}^{(2)}\|_2 \leq O_\alpha(\eta^2 T \text{rk}_*^{p_2-1}) \quad \text{and} \quad \max_{t \leq T} R_{t,t}^{(2)} \leq O_\alpha(\eta \text{rk}_*^{p_2-\frac{1}{2}} \log^2 d) \right\}. \end{aligned}$$

where p_2 is defined in Corollary 4. By using Proposition 19, we can show that with probability $d^4 \exp(-\log^2 d)$, $\mathcal{T}_{\text{noise}} \geq \frac{\text{rk}_*^\alpha \log^2 d}{\eta}$. $\square$

G Auxiliary Statements

G.1 Matrix bounds

Proposition 22. For $A, B \in \mathbb{R}^{d \times r}$, we have

$$-A^\top A - B^\top B \preceq A^\top B + B^\top A \preceq A^\top A + B^\top B.$$

If $r = d$, then $(A + A^\top)^2 \preceq 2A^\top A + 2AA^\top$. Moreover, if $A_1, \dots, A_k$ are symmetric matrices,

$$\left(\sum_{i=1}^k A_i \right)^2 \preceq k \sum_{i=1}^k A_i^2.$$

Proof. We have

$$(A - B)^\top (A - B) \succeq 0 \Rightarrow A^\top A + B^\top B \succeq A^\top B + B^\top A.$$

By using $A \leftarrow -A$, we obtain the left inequality too. For the second inequality, we have

$$\begin{aligned} (A + A^\top)^2 &= A^\top A + AA^\top + AA + A^\top A^\top \\ (A - A^\top)^\top (A - A^\top) &= A^\top A + AA^\top - AA - A^\top A^\top \end{aligned}$$

Therefore, $(A + A^\top)^2 \preceq 2(A^\top A + AA^\top)$. For the last statement,

$$\left(\sum_{i=1}^k A_i \right)^2 = \sum_{i=1}^k A_i^2 + \sum_{i=1}^k \sum_{j=i+1}^k A_i A_j + \sum_{i=1}^k \sum_{j=i+1}^k A_j A_i \preceq k \sum_{i=1}^k A_i^2,$$

where we use the first statement in the last inequality. $\square$

Proposition 23. Consider a symmetric square matrix with block partition

$$M = \begin{bmatrix} A & B \\ B^\top & C \end{bmatrix}.$$

If A is invertible, then $M \succ 0$ if and only if $A \succ 0$ and $C - B^\top A^{-1} B \succ 0$.

Proof. If A is invertible, we have

$$\begin{bmatrix} A & B \\ B^\top & C \end{bmatrix} = \begin{bmatrix} I & 0 \\ B^\top A^{-1} & I \end{bmatrix} \begin{bmatrix} A & 0 \\ 0 & C - B^\top A^{-1} B \end{bmatrix} \begin{bmatrix} I & A^{-1} B \\ 0 & I \end{bmatrix}.$$

Note that

$$\begin{bmatrix} I & 0 \\ B^\top A^{-1} & I \end{bmatrix}^{-1} = \begin{bmatrix} I & 0 \\ -B^\top A^{-1} & I \end{bmatrix}.$$

Therefore, the statement follows. $\square$

Proposition 24. Let $r_u < r$ and $Z \in \mathbb{R}^{r \times r_s}$ such that

$$Z = \begin{bmatrix} Z_1 \\ Z_2 \end{bmatrix}, \text{ where } Z_1 \in \mathbb{R}^{r_u \times r_s}, Z_2 \in \mathbb{R}^{r-r_u \times r_s}.$$

For any $0 \leq \varepsilon < 1$

$$\begin{aligned} ZZ^\top &\succeq \varepsilon \begin{bmatrix} Z_1 Z_1^\top & 0 \\ 0 & 0 \end{bmatrix} + (1 - \varepsilon) \begin{bmatrix} Z_1 Z_1^\top - Z_1 Z_2^\top (Z_2 Z_2^\top)^+ Z_2 Z_1^\top & 0 \\ 0 & 0 \end{bmatrix} \\ &\quad - \frac{\varepsilon}{1 - \varepsilon} \begin{bmatrix} 0 & 0 \\ 0 & Z_2 Z_2^\top \end{bmatrix}, \end{aligned}$$

where $A \rightarrow A^+$ denotes the pseudo inverse operator.

Proof. We will denote $\mathbf{x} \in \mathbb{R}^r$ as

$$\mathbf{x} = \begin{bmatrix} \mathbf{x}_1 \\ \mathbf{x}_2 \end{bmatrix} \quad \text{where } \mathbf{x}_1 \in \mathbb{R}^{r_u}, \mathbf{x}_2 \in \mathbb{R}^{r-r_u}.$$

We have

$$\begin{aligned} \mathbf{x}^\top \mathbf{Z} \mathbf{Z}^\top \mathbf{x} &= \left(\mathbf{x}_1^\top \mathbf{Z}_1 \mathbf{Z}_1^\top \mathbf{x}_1 + 2\mathbf{x}_1^\top \mathbf{Z}_1 \mathbf{Z}_2^\top \mathbf{x}_2 + \frac{1}{1-\varepsilon} \mathbf{x}_2^\top \mathbf{Z}_2 \mathbf{Z}_2^\top \mathbf{x}_2 \right) - \frac{\varepsilon}{1-\varepsilon} \mathbf{x}_2^\top \mathbf{Z}_2 \mathbf{Z}_2^\top \mathbf{x}_2 \\ &\stackrel{(a)}{\geq} \left(\mathbf{x}_1^\top \mathbf{Z}_1 \mathbf{Z}_1^\top \mathbf{x}_1 - (1-\varepsilon) \mathbf{x}_1^\top \mathbf{Z}_1 \mathbf{Z}_2^\top (\mathbf{Z}_2 \mathbf{Z}_2^\top)^+ \mathbf{Z}_2 \mathbf{Z}_1^\top \mathbf{x}_1 \right) - \frac{\varepsilon}{1-\varepsilon} \mathbf{x}_2^\top \mathbf{Z}_2 \mathbf{Z}_2^\top \mathbf{x}_2, \end{aligned}$$

where we minimized the first term in the first line over $\mathbf{x}_2$ in (a). Since (a) holds for all $\mathbf{x}$, the statement follows, $\square$

Proposition 25. Let $\mathbf{A} \in \mathbb{R}^{r \times r}$ be a symmetric matrix. For $\mathbf{S} \succ -\mathbf{A}$, $\mathbf{S} \rightarrow -(\mathbf{S} + \mathbf{A})^{-1}$ is monotone.

Proof. Let $\mathbf{S}_1 \succ \mathbf{S}_2 \succ -\mathbf{A}$. We have

$$-(\mathbf{S}_1 + \mathbf{A})^{-1} + (\mathbf{S}_2 + \mathbf{A})^{-1} = (\mathbf{S}_2 + \mathbf{A})^{-1} ((\mathbf{S}_1 - \mathbf{S}_2)^{-1} + (\mathbf{S}_2 + \mathbf{A})^{-1})^{-1} (\mathbf{S}_2 + \mathbf{A})^{-1} \succ 0. \quad (\text{G.1})$$

For $\mathbf{S}_1 \succeq \mathbf{S}_2$, we can use $\mathbf{S}_1 + \varepsilon \mathbf{I}_r$ in (G.1) and take $\varepsilon \downarrow 0$ $\square$

G.1.1 Additional bounds for continuous-time analysis

Proposition 26. For a symmetric positive definite $\mathbf{D}_1$, $\mathbf{D}_2$, and $C > 0$, we have

$$\begin{aligned} \mathbf{D}_1 \left(\mathbf{D}_1 + \mathbf{Z}_1 (C\mathbf{I}_{r_s} + \mathbf{Z}_2^\top \mathbf{D}_2^{-1} \mathbf{Z}_2)^{-1} \mathbf{Z}_1^\top \right)^{-1} \mathbf{Z}_1 \mathbf{Z}_2^\top (\mathbf{Z}_2 \mathbf{Z}_2^\top + C\mathbf{D}_2)^{-1} \mathbf{D}_2 \\ = \mathbf{Z}_1 (C\mathbf{I}_{r_s} + \mathbf{Z}_2^\top \mathbf{D}_2^{-1} \mathbf{Z}_2 + \mathbf{Z}_1^\top \mathbf{D}_1^{-1} \mathbf{Z}_1)^{-1} \mathbf{Z}_2^\top. \end{aligned}$$

Proof. We have

$$\begin{aligned} &(\mathbf{D}_1 + \mathbf{Z}_1 (C\mathbf{I}_{r_s} + \mathbf{Z}_2^\top \mathbf{D}_2^{-1} \mathbf{Z}_2)^{-1} \mathbf{Z}_1^\top)^{-1} \\ &= \mathbf{D}_1^{-1} - \mathbf{D}_1^{-1} \mathbf{Z}_1 (C\mathbf{I}_{r_s} + \mathbf{Z}_2^\top \mathbf{D}_2^{-1} \mathbf{Z}_2 + \mathbf{Z}_1^\top \mathbf{D}_1^{-1} \mathbf{Z}_1)^{-1} \mathbf{Z}_1^\top \mathbf{D}_1^{-1}. \end{aligned}$$

Therefore,

$$\begin{aligned} \mathbf{D}_1 (\mathbf{D}_1 + \mathbf{Z}_1 (\mathbf{I}_{r_s} + \mathbf{Z}_2^\top \mathbf{D}_2^{-1} \mathbf{Z}_2)^{-1} \mathbf{Z}_1^\top)^{-1} \mathbf{Z}_1 \\ = \mathbf{Z}_1 \left(\mathbf{I}_{r_s} - (C\mathbf{I}_{r_s} + \mathbf{Z}_2^\top \mathbf{D}_2^{-1} \mathbf{Z}_2 + \mathbf{Z}_1^\top \mathbf{D}_1^{-1} \mathbf{Z}_1)^{-1} \mathbf{Z}_1^\top \mathbf{D}_1^{-1} \mathbf{Z}_1 \right) \\ = \mathbf{Z}_1 (C\mathbf{I}_{r_s} + \mathbf{Z}_2^\top \mathbf{D}_2^{-1} \mathbf{Z}_2 + \mathbf{Z}_1^\top \mathbf{D}_1^{-1} \mathbf{Z}_1)^{-1} (C\mathbf{I}_{r_s} + \mathbf{Z}_2^\top \mathbf{D}_2^{-1} \mathbf{Z}_2) \end{aligned}$$

Then,

$$\begin{aligned} \mathbf{Z}_1 (C\mathbf{I}_{r_s} + \mathbf{Z}_2^\top \mathbf{D}_2^{-1} \mathbf{Z}_2 + \mathbf{Z}_1^\top \mathbf{D}_1^{-1} \mathbf{Z}_1)^{-1} (C\mathbf{I}_{r_s} + \mathbf{Z}_2^\top \mathbf{D}_2^{-1} \mathbf{Z}_2) \mathbf{Z}_2^\top (\mathbf{Z}_2 \mathbf{Z}_2^\top + C\mathbf{D}_2)^{-1} \mathbf{D}_2 \\ = \mathbf{Z}_1 (C\mathbf{I}_{r_s} + \mathbf{Z}_2^\top \mathbf{D}_2^{-1} \mathbf{Z}_2 + \mathbf{Z}_1^\top \mathbf{D}_1^{-1} \mathbf{Z}_1)^{-1} \mathbf{Z}_2^\top. \end{aligned}$$

$\square$

Proposition 27. For some diagonal positive definite $\mathbf{A} := \text{diag}(\{a_j\}_{j=1}^{r_u})$ and $\mathbf{B} := \text{diag}(\{b_j\}_{j=1}^{d-r_u})$, we let

$$\mathbf{D}_1 := \frac{\mathbf{A} \exp(-t\mathbf{A})}{\mathbf{I}_{r_u} - \exp(-t\mathbf{A})}, \quad \mathbf{D}_2 := \frac{\mathbf{B} \exp(-t\mathbf{B})}{\mathbf{I}_{d-r_u} - \exp(-t\mathbf{B})},$$

For some $\mathbf{Z}_1 \in \mathbb{R}^{r_u \times r_s}$, $\mathbf{Z}_2 \in \mathbb{R}^{(d-r_u) \times r_s}$, and $C > 0$, we define

$$\mathbf{M} := \exp(0.5t\mathbf{A}) \mathbf{Z}_1 (C\mathbf{I}_{r_s} + \mathbf{Z}_2^\top \mathbf{D}_2^{-1} \mathbf{Z}_2 + \mathbf{Z}_1^\top \mathbf{D}_1^{-1} \mathbf{Z}_1)^{-1} \mathbf{Z}_2^\top \exp(0.5t\mathbf{B}).$$

We have

$$\|M\|_F^2 \leq \tilde{C} \sum_{i=1}^{r_u \wedge r_s} \left(\lambda_{\max}(\mathbf{Z}_1 \mathbf{Z}_1^\top) \exp(t(a_i + b_i)) \wedge \left(C + \frac{\lambda_{\min}(\mathbf{Z}_2^\top \mathbf{Z}_2)}{\lambda_{\max}(\mathbf{D}_2)} \right) \frac{a_i \exp(t(a_i + b_i))}{\exp(ta_i) - 1} \right)$$

where

$$\tilde{C} = \frac{\lambda_{\max}(\mathbf{Z}_2^\top \mathbf{Z}_2)}{\left(C + \frac{\lambda_{\min}(\mathbf{Z}_2^\top \mathbf{Z}_2)}{\lambda_{\max}(\mathbf{D}_2)} \right)^2}$$

Proof. For convenience, we will use

$$\begin{aligned} \tilde{D}_1 &:= \frac{\mathbf{A}}{\mathbf{I}_{r_u} - \exp(-t\mathbf{A})}, & \tilde{D}_2 &= \frac{\mathbf{B}}{\mathbf{I}_{d-r_u} - \exp(-t\mathbf{B})}, \\ \tilde{Z}_1 &:= \exp(0.5t\mathbf{A})\mathbf{Z}_1, & \tilde{Z}_2 &:= \exp(0.5t\mathbf{B})\mathbf{Z}_2. \end{aligned}$$

We let

$$\begin{aligned} M_1 &:= \tilde{Z}_1 \left(C\mathbf{I}_{r_s} + \tilde{Z}_2^\top \tilde{D}_2^{-1} \tilde{Z}_2 + \tilde{Z}_1^\top \tilde{D}_1^{-1} \tilde{Z}_1 \right)^{-\frac{1}{2}}, \\ M_2 &:= \left(C\mathbf{I}_{r_s} + \tilde{Z}_2^\top \tilde{D}_2^{-1} \tilde{Z}_2 + \tilde{Z}_1^\top \tilde{D}_1^{-1} \tilde{Z}_1 \right)^{-\frac{1}{2}} \tilde{Z}_2^\top \end{aligned}$$

We observe that

$$\|M\|_F^2 = \text{Tr}(M_1^\top M_1 M_2 M_2^\top) \leq \sum_{i=1}^{r_u \wedge r_s} \lambda_i(M_1 M_1^\top) \lambda_i(M_2^\top M_2)$$

where we used that $\text{rank}(M_1 M_1^\top) \leq r_u \wedge r_s$ and Von Neumann's trace inequality in the last part. We have

$$\begin{aligned} M_2^\top M_2 &\preceq \exp(0.5t\mathbf{B})\mathbf{Z}_2^\top \left(C\mathbf{I}_{r_s} + \frac{1}{\lambda_{\max}(\mathbf{D}_2)} \mathbf{Z}_2^\top \mathbf{Z}_2 \right)^{-1} \mathbf{Z}_2 \exp(0.5t\mathbf{B}) \\ &\preceq \frac{\lambda_{\max}(\mathbf{Z}_2^\top \mathbf{Z}_2)}{C + \frac{\lambda_{\max}(\mathbf{Z}_2^\top \mathbf{Z}_2)}{\lambda_{\max}(\mathbf{D}_2)}} \exp(t\mathbf{B}) \end{aligned}$$

On the other hand,

$$\begin{aligned} M_1 M_1^\top &\preceq \tilde{Z}_1 \left(\left(C + \frac{\lambda_{\min}(\mathbf{Z}_2^\top \mathbf{Z}_2)}{\lambda_{\max}(\mathbf{D}_2)} \right) \mathbf{I}_{r_s} + \tilde{Z}_1^\top \tilde{D}_1^{-1} \tilde{Z}_1 \right)^{-1} \tilde{Z}_1^\top \\ &= \frac{1}{C + \frac{\lambda_{\min}(\mathbf{Z}_2^\top \mathbf{Z}_2)}{\lambda_{\max}(\mathbf{D}_2)}} \tilde{Z}_1 \left(\mathbf{I}_{r_s} + \tilde{Z}_1^\top \left(\left(C + \frac{\lambda_{\min}(\mathbf{Z}_2^\top \mathbf{Z}_2)}{\lambda_{\max}(\mathbf{D}_2)} \right) \tilde{D}_1 \right)^{-1} \tilde{Z}_1 \right)^{-1} \tilde{Z}_1^\top \\ &= \frac{1}{C + \frac{\lambda_{\min}(\mathbf{Z}_2^\top \mathbf{Z}_2)}{\lambda_{\max}(\mathbf{D}_2)}} \tilde{Z}_1 \tilde{Z}_1^\top \left(\left(C + \frac{\lambda_{\min}(\mathbf{Z}_2^\top \mathbf{Z}_2)}{\lambda_{\max}(\mathbf{D}_2)} \right) \tilde{D}_1 + \tilde{Z}_1 \tilde{Z}_1^\top \right)^{-1} \left(\left(C + \frac{\lambda_{\min}(\mathbf{Z}_2^\top \mathbf{Z}_2)}{\lambda_{\max}(\mathbf{D}_2)} \right) \tilde{D}_1 \right) \end{aligned}$$

We have the following at the same time:

$$\begin{aligned} &\bullet \tilde{Z}_1 \tilde{Z}_1^\top \left(\left(C + \frac{\lambda_{\min}(\mathbf{Z}_2^\top \mathbf{Z}_2)}{\lambda_{\max}(\mathbf{D}_2)} \right) \tilde{D}_1 + \tilde{Z}_1 \tilde{Z}_1^\top \right)^{-1} \left(C + \frac{\lambda_{\min}(\mathbf{Z}_2^\top \mathbf{Z}_2)}{\lambda_{\max}(\mathbf{D}_2)} \right) \tilde{D}_1 \preceq \left(C + \frac{\lambda_{\min}(\mathbf{Z}_2^\top \mathbf{Z}_2)}{\lambda_{\max}(\mathbf{D}_2)} \right) \tilde{D}_1 \\ &\bullet \tilde{Z}_1 \tilde{Z}_1^\top \left(\left(C + \frac{\lambda_{\min}(\mathbf{Z}_2^\top \mathbf{Z}_2)}{\lambda_{\max}(\mathbf{D}_2)} \right) \tilde{D}_1 + \tilde{Z}_1 \tilde{Z}_1^\top \right)^{-1} \left(C + \frac{\lambda_{\min}(\mathbf{Z}_2^\top \mathbf{Z}_2)}{\lambda_{\max}(\mathbf{D}_2)} \right) \tilde{D}_1 \preceq \lambda_{\max}(\mathbf{Z}_1 \mathbf{Z}_1^\top) \exp(t\mathbf{A}) \end{aligned}$$

Therefore, for $i \leq r \wedge r_s$, we have

$$\lambda_i(M_1 M_1^\top) \leq \frac{1}{C + \frac{\lambda_{\min}(\mathbf{Z}_2^\top \mathbf{Z}_2)}{\lambda_{\max}(\mathbf{D}_2)}} \left(\lambda_{\max}(\mathbf{Z}_1 \mathbf{Z}_1^\top) \exp(ta_i) \wedge \left(C + \frac{\lambda_{\min}(\mathbf{Z}_2^\top \mathbf{Z}_2)}{\lambda_{\max}(\mathbf{D}_2)} \right) \frac{a_i \exp(ta_i)}{\exp(ta_i) - 1} \right)$$

Therefore,

$$\|M\|_F^2 \leq \tilde{C} \sum_{i=1}^{r_u \wedge r_s} \left(\lambda_{\max}(\mathbf{Z}_1 \mathbf{Z}_1^\top) \exp(t(a_i + b_i)) \wedge \left(C + \frac{\lambda_{\min}(\mathbf{Z}_2^\top \mathbf{Z}_2)}{\lambda_{\max}(\mathbf{D}_2)} \right) \frac{a_i \exp(t(a_i + b_i))}{\exp(ta_i) - 1} \right).$$

□

G.1.2 Additional bounds for discrete-time analysis

Proposition 28. For some positive definite diagonal matrices $D_0, D_1 \in \mathbb{R}^{r \times r}$ and symmetric matrices $G, \nu \in \mathbb{R}^{r \times r}$, we let

$$V := 2D_0^{\frac{1}{2}}GD_0^{\frac{1}{2}} - D_1 \text{ and } \zeta := D_0^{\frac{1}{2}}\nu D_0^{\frac{1}{2}} \text{ and } \dot{V} := V + \zeta,$$

where

- $\|G\|_2 \leq L_G$ and $\|\nu\|_2 \leq L_\nu$ and $\|D_0\|_2 \leq L_0$.
- $\|D_0^{-1}D_1\|_2 \leq L_{1/0}$ and $\|D_0D_1^{-1}\|_2 \leq L_{0/1}$.
- For notational convenience, let $L_F := 2L_G + L_{1/0}$ and $L_{\dot{F}} := 2(L_G + L_\nu) + L_{1/0}$.

For $0 \leq \eta < \frac{1}{L_{\dot{F}}L_0}$, we have that $(I_r + \eta V)$ and $(I_r + \eta \dot{V})$ are invertible and the following bounds holds:

$$\begin{aligned} -C_1D_1 &\preceq V^2(I_r + \eta V)^{-1}\zeta\dot{V}(I_r + \eta \dot{V})^{-1} \preceq C_1D_1, & -C_2D_1 &\preceq V\zeta\dot{V}^2(I_r + \eta \dot{V})^{-1} \preceq C_2D_1 \\ & & & \text{(G.2)} \\ -C_3D_1 &\preceq V^3(I_r + \eta V)^{-1}\zeta \preceq C_3D_1, & -C_4D_1 &\preceq \zeta\dot{V}^3(I_r + \eta \dot{V})^{-1} \preceq C_4D_1, \\ & & & \text{(G.3)} \end{aligned}$$

where

$$\begin{aligned} C_1 &= \frac{L_\nu L_{0/1} L_F^2 L_{\dot{F}} L_0^3}{(1 - \eta L_F L_0)(1 - \eta L_{\dot{F}} L_0)}, & C_2 &= \frac{L_\nu L_{0/1} L_F L_{\dot{F}}^2 L_0^3}{1 - \eta L_{\dot{F}} L_0} \\ C_3 &= \frac{L_\nu L_{0/1} L_F^3 L_0^3}{1 - \eta L_F L_0}, & C_4 &= \frac{L_\nu L_{0/1} L_{\dot{F}}^3 L_0^3}{1 - \eta L_{\dot{F}} L_0}. \end{aligned}$$

Proof. Note that $\|V\|_2 \vee \|\dot{V}\|_2 \leq L_{\dot{F}}L_0$, therefore, if $0 \leq \eta < \frac{1}{L_{\dot{F}}L_0}$, $(I_r + \eta V)$ and $(I_r + \eta \dot{V})$ are invertible. For the following, we introduce the notation

$$\dot{G} := G + \nu \text{ and } F = 2G + D_0^{-1}D_1 \text{ and } \dot{F} = 2\dot{G} + D_0^{-1}D_1.$$

Note that we have $\|F\|_2 \leq L_F$ and $\|\dot{F}\|_2 \leq L_{\dot{F}}$. For the left part of (G.2), we write

$$\begin{aligned} D_0^{-\frac{1}{2}}V^2(I_r + \eta V)^{-1}\zeta\dot{V}(I_r + \eta \dot{V})^{-1}D_0^{-\frac{1}{2}} \\ = FD_0FD_0(I_r + \eta FD_0)^{-1}\nu D_0\dot{F}(I_r + \eta D_0\dot{F})^{-1}. \end{aligned}$$

Therefore, we have

$$\begin{aligned} &\left\| FD_0FD_0(I_r + \eta FD_0)^{-1}\nu D_0\dot{F}(I_r + \eta D_0\dot{F})^{-1} \right\|_2 \\ &\leq \frac{\|F\|_2^2\|\dot{F}\|_2\|D_0\|_2^3\|\nu\|_2}{(1 - \eta\|F\|_2\|D_0\|_2)(1 - \eta\|\dot{F}\|_2\|D_0\|_2)} \\ &\leq \frac{L_F^2L_{\dot{F}}L_0^3L_\nu}{(1 - \eta L_F L_0)(1 - \eta L_{\dot{F}} L_0)}. \end{aligned}$$

Therefore, we have the bound. For the right part of (G.2), we write

$$D_0^{-\frac{1}{2}}V\zeta\dot{V}^2(I_r + \eta \dot{V})^{-1}D_0^{-\frac{1}{2}} = FD_0\nu D_0\dot{F}D_0\dot{F}(I_r + \eta D_0\dot{F})^{-1}$$

Therefore, we have

$$\left\| FD_0\nu D_0\dot{F}D_0\dot{F}(I_r + \eta D_0\dot{F})^{-1} \right\|_2 \leq \frac{\|F\|_2\|\dot{F}\|_2^2\|D_0\|_2^3\|\nu\|_2}{1 - \eta\|\dot{F}\|_2\|D_0\|_2} \leq \frac{L_F L_{\dot{F}}^2 L_0^3 L_\nu}{1 - \eta L_{\dot{F}} L_0},$$

which gives us the bound. For the left part of (G.3), we write

$$D_0^{-\frac{1}{2}} V^3 (I_r + \eta V)^{-1} \zeta D_0^{-\frac{1}{2}} = (F D_0)^3 (I_r + \eta F D_0)^{-1} \nu$$

Therefore, we have

$$\left\| (F D_0)^3 (I_r + \eta F D_0)^{-1} \nu \right\|_2 \leq \frac{\|F\|_2^3 \|D_0\|_2^3 \|\nu\|_2}{1 - \eta \|F\|_2 \|D_0\|_2} \leq \frac{L_\nu L_F^3 L_0^3}{1 - \eta L_F L_0},$$

which gives us the bound. The the right part of (G.3) can be derived similarly. $\square$

Proposition 29. Let $V, \dot{V} \in \mathbb{R}^{r \times r}$ be symmetric matrices such that $\dot{V} = V + \zeta$. We have

$$\dot{V} (I_r + \eta \dot{V})^{-1} - V (I_r + \eta V)^{-1} = (I_r + \eta V)^{-1} \zeta (I_r + \eta \dot{V})^{-1}.$$

Moreover, given that

$$M := \zeta - \eta V \zeta - \eta \zeta V - \eta \zeta^2 + \eta^2 \zeta V \zeta + \eta^2 \zeta^3 + \eta^2 V \zeta V$$

under the conditions of Proposition 28, we have for any $\kappa_d > 0$,

$$\begin{aligned} -\frac{2}{\kappa_d^2} \eta^2 \zeta^2 - \eta^2 \kappa_d^2 V^4 - \eta^2 \kappa_d^2 V \zeta^2 V - C \eta^3 D_1 &\preceq (I_r + \eta V)^{-1} \zeta (I_r + \eta \dot{V})^{-1} - M \\ &\preceq \frac{2}{\kappa_d^2} \eta^2 \zeta^2 + \eta^2 \kappa_d^2 V^4 + \eta^2 \kappa_d^2 V \zeta^2 V + C \eta^3 D_1 \end{aligned}$$

where $C = C_1 + C_2 + C_3 + C_4$, i.e., the sum of the constants given in Proposition 28.

Proof. We write

$$\begin{aligned} \dot{V} (I_r + \eta \dot{V})^{-1} - (I_r + \eta V)^{-1} V \\ &= (I_r + \eta V)^{-1} \left((I_r + \eta V)(V + \zeta) - V(I_r + \eta \dot{V}) \right) (I_r + \eta \dot{V})^{-1} \\ &= (I_r + \eta V)^{-1} \zeta (I_r + \eta \dot{V})^{-1}. \end{aligned}$$

For the second part, we write

$$\begin{aligned} (I_r + \eta V)^{-1} \zeta (I_r + \eta \dot{V})^{-1} \\ &= (I_r - \eta V (I_r + \eta V)^{-1}) \zeta (I_r - \eta \dot{V} (I_r + \eta \dot{V})^{-1}) \\ &= \zeta - \eta V (I_r - \eta V + \eta^2 V^2 (I_r + \eta V)^{-1}) \zeta \\ &\quad - \eta \zeta \dot{V} (I_r - \eta \dot{V} + \eta^2 \dot{V}^2 (I_r + \eta V)^{-1}) + \eta^2 V (I_r + \eta V)^{-1} \zeta \dot{V} (I_r + \eta \dot{V})^{-1} \\ &= \zeta - \eta V \zeta - \eta \zeta V - \eta \zeta^2 + \eta^2 V^2 \zeta + \eta^2 \zeta \dot{V}^2 - \eta^3 V^3 (I_r + \eta V)^{-1} \zeta \\ &\quad - \eta^3 \zeta \dot{V}^3 (I_r + \eta \dot{V})^{-1} + \eta^2 V (I_r + \eta V)^{-1} \zeta \dot{V} (I_r + \eta \dot{V})^{-1} \end{aligned}$$

We have

$$\eta^2 V^2 \zeta + \eta^2 \zeta \dot{V}^2 = \eta^2 V^2 \zeta + \eta^2 \zeta (V + \zeta)^2 = \underbrace{\eta^2 V^2 \zeta + \eta^2 \zeta V^2 + \eta^2 \zeta^2 V}_{:=M_1} + \eta^2 \zeta V \zeta + \eta^2 \zeta^3.$$

Moreover,

$$\begin{aligned} \eta^2 V (I_r + \eta V)^{-1} \zeta \dot{V} (I_r + \eta \dot{V})^{-1} \\ &= \eta^2 V \zeta \dot{V} (I_r + \eta \dot{V})^{-1} - \eta^3 V^2 (I_r + \eta V)^{-1} \zeta \dot{V} (I_r + \eta \dot{V})^{-1} \\ &= \eta^2 V \zeta \dot{V} - \eta^3 V \zeta \dot{V}^2 (I_r + \eta \dot{V})^{-1} - \eta^3 V^2 (I_r + \eta V)^{-1} \zeta \dot{V} (I_r + \eta \dot{V})^{-1} \end{aligned}$$

$$= \eta^2 \mathbf{V} \zeta \mathbf{V} + \underbrace{\eta^2 \mathbf{V} \zeta^2}_{:=M_2} - \eta^3 \mathbf{V} \zeta \dot{\mathbf{V}}^2 \left(\mathbf{I}_r + \eta \dot{\mathbf{V}} \right)^{-1} - \eta^3 \mathbf{V}^2 (\mathbf{I}_r + \eta \mathbf{V})^{-1} \zeta \dot{\mathbf{V}} \left(\mathbf{I}_r + \eta \dot{\mathbf{V}} \right)^{-1}$$

By Proposition 22, we have

$$-2\eta^2 \zeta^2 - \eta^2 \mathbf{V}^4 - \eta^2 \mathbf{V} \zeta^2 \mathbf{V} \preceq M_1 + M_2 \preceq 2\eta^2 \zeta^2 + \eta^2 \mathbf{V}^4 + \eta^2 \mathbf{V} \zeta^2 \mathbf{V}.$$

Therefore by Proposition 28, we have

$$\begin{aligned} -\frac{2}{\kappa_d^2} \eta^2 \zeta^2 - \eta^2 \kappa_d^2 \mathbf{V}^4 - \eta^2 \kappa_d^2 \mathbf{V} \zeta^2 \mathbf{V} - C \eta^3 \mathbf{D}_1 &\preceq (\mathbf{I}_r + \eta \mathbf{V})^{-1} \zeta \left(\mathbf{I}_r + \eta \dot{\mathbf{V}} \right)^{-1} - M \\ &\preceq \frac{2}{\kappa_d^2} \eta^2 \zeta^2 + \eta^2 \kappa_d^2 \mathbf{V}^4 + \eta^2 \kappa_d^2 \mathbf{V} \zeta^2 \mathbf{V} + C \eta^3 \mathbf{D}_1. \end{aligned}$$

□

Proposition 30. By using the notation in Proposition 28, we consider

$$\eta < \frac{1}{L_F L_0} \quad \text{and} \quad 0 < \varepsilon < \frac{0.5/\eta}{L_F L_0} - 1$$

Then,

$$\begin{aligned} \mathbf{V} (\mathbf{I}_r + \eta \mathbf{V})^{-1} - \varepsilon \eta \mathbf{V}^2 &\succeq \mathbf{V} (\mathbf{I}_r + \eta(1 + \varepsilon) \mathbf{V})^{-1} - 2.5 \varepsilon \eta^2 C \mathbf{D}_1 \\ \mathbf{V} (\mathbf{I}_r + \eta \mathbf{V})^{-1} + \varepsilon \eta \mathbf{V}^2 &\preceq \mathbf{V} (\mathbf{I}_r + \eta(1 - \varepsilon) \mathbf{V})^{-1} + 1.5 \varepsilon \eta^2 C \mathbf{D}_1, \end{aligned}$$

$$\text{where } C = \frac{L_{0/1} L_F^3 L_0^3}{1 - \eta L_F L_0}.$$

Proof. For the lower bound, we have

$$\begin{aligned} &\mathbf{V} (\mathbf{I}_r + \eta \mathbf{V})^{-1} - \varepsilon \eta \mathbf{V}^2 \\ &= \mathbf{V} - (1 + \varepsilon) \eta \mathbf{V}^2 + \eta^2 \mathbf{V}^3 (\mathbf{I}_r + \eta \mathbf{V})^{-1} \\ &= \mathbf{V} - (1 + \varepsilon) \eta \mathbf{V}^2 + (1 + \varepsilon)^2 \eta^2 \mathbf{V}^3 (\mathbf{I}_r + (1 + \varepsilon) \eta \mathbf{V})^{-1} \\ &\quad - (2\varepsilon + \varepsilon^2) \eta^2 \mathbf{V}^3 (\mathbf{I}_r + \eta \mathbf{V})^{-1} + (1 + \varepsilon)^2 \varepsilon \eta^3 \mathbf{V}^4 (\mathbf{I}_r + (1 + \varepsilon) \eta \mathbf{V})^{-1} (\mathbf{I}_r + \eta \mathbf{V})^{-1} \\ &\succeq \mathbf{V} (\mathbf{I}_r + \eta(1 + \varepsilon) \mathbf{V})^{-1} - 2.5 \varepsilon \eta^2 C \mathbf{D}_1, \end{aligned}$$

where we used C_3 with $L_\nu = 1$ in Proposition 28 in the last step. For the upper bound,

$$\begin{aligned} &\mathbf{V} (\mathbf{I}_r + \eta \mathbf{V})^{-1} + \varepsilon \eta \mathbf{V}^2 \\ &= \mathbf{V} - (1 - \varepsilon) \eta \mathbf{V}^2 + \eta^2 \mathbf{V}^3 (\mathbf{I}_r + \eta \mathbf{V})^{-1} \\ &= \mathbf{V} - (1 - \varepsilon) \eta \mathbf{V}^2 + (1 - \varepsilon)^2 \eta^2 \mathbf{V}^3 (\mathbf{I}_r + (1 - \varepsilon) \eta \mathbf{V})^{-1} \\ &\quad + (2\varepsilon - \varepsilon^2) \eta^2 \mathbf{V}^3 (\mathbf{I}_r + \eta \mathbf{V})^{-1} - (1 - \varepsilon)^2 \varepsilon \eta^3 \mathbf{V}^4 (\mathbf{I}_r + (1 - \varepsilon) \eta \mathbf{V})^{-1} (\mathbf{I}_r + \eta \mathbf{V})^{-1} \\ &\preceq \mathbf{V} (\mathbf{I}_r + \eta(1 - \varepsilon) \mathbf{V})^{-1} + 1.5 \varepsilon \eta^2 C \mathbf{D}_1. \end{aligned}$$

□

Lemma 7. For any $\eta \in \mathbb{R}$ and $t \in \mathbb{N}$, we have

$$\begin{bmatrix} \mathbf{I}_r & \eta \mathbf{I}_r \\ \eta \mathbf{\Lambda}^2 & \mathbf{I}_r \end{bmatrix}^t = \begin{bmatrix} \frac{(\mathbf{I}_r + \eta \mathbf{\Lambda})^t + (\mathbf{I}_r - \eta \mathbf{\Lambda})^t}{2} & \mathbf{\Lambda}^{-1} \frac{(\mathbf{I}_r + \eta \mathbf{\Lambda})^t - (\mathbf{I}_r - \eta \mathbf{\Lambda})^t}{2} \\ \mathbf{\Lambda} \frac{(\mathbf{I}_r + \eta \mathbf{\Lambda})^t - (\mathbf{I}_r - \eta \mathbf{\Lambda})^t}{2} & \frac{(\mathbf{I}_r + \eta \mathbf{\Lambda})^t + (\mathbf{I}_r - \eta \mathbf{\Lambda})^t}{2} \end{bmatrix}. \quad (\text{G.4})$$

Proof. We observe that

$$\mathbf{A} := \begin{bmatrix} 0 & \mathbf{I}_r \\ \mathbf{\Lambda}^2 & 0 \end{bmatrix} \Rightarrow (\text{G.4}) = \sum_{k=0}^t \binom{t}{k} \eta^k \mathbf{A}^k.$$

Note that

$$\mathbf{A}^{2k} = \begin{bmatrix} \mathbf{\Lambda}^{2k} & 0 \\ 0 & \mathbf{\Lambda}^{2k} \end{bmatrix} \quad \text{and} \quad \mathbf{A}^{2k+1} = \begin{bmatrix} 0 & \mathbf{\Lambda}^{2k} \\ \mathbf{\Lambda}^{2k+2} & 0 \end{bmatrix}.$$

Therefore,

$$\begin{aligned} \sum_{k=0}^t \binom{t}{k} \eta^k \mathbf{A}^k &= \begin{bmatrix} \sum_{k \text{ even}}^t \binom{t}{k} \eta^k \mathbf{\Lambda}^k & \sum_{k \text{ odd}}^t \binom{t}{k} \eta^k \mathbf{\Lambda}^{k-1} \\ \sum_{k \text{ odd}}^t \binom{t}{k} \eta^k \mathbf{\Lambda}^{k+1} & \sum_{k \text{ even}}^t \binom{t}{k} \eta^k \mathbf{\Lambda}^k \end{bmatrix} \\ &= \begin{bmatrix} \frac{(\mathbf{I}_r + \eta \mathbf{\Lambda})^t + (\mathbf{I}_r - \eta \mathbf{\Lambda})^t}{2} & \mathbf{\Lambda}^{-1} \frac{(\mathbf{I}_r + \eta \mathbf{\Lambda})^t - (\mathbf{I}_r - \eta \mathbf{\Lambda})^t}{2} \\ \mathbf{\Lambda} \frac{(\mathbf{I}_r + \eta \mathbf{\Lambda})^t - (\mathbf{I}_r - \eta \mathbf{\Lambda})^t}{2} & \frac{(\mathbf{I}_r + \eta \mathbf{\Lambda})^t + (\mathbf{I}_r - \eta \mathbf{\Lambda})^t}{2} \end{bmatrix}. \end{aligned}$$

□

G.2 Some moment bounds and concentration inequalities

Lemma 8 (Hypercontractivity). *Let $P_k : \mathbb{R}^d \rightarrow \mathbb{R}$ be a polynomial of degree- k and $\mathbf{x} \sim \mathcal{N}(0, \mathbf{I}_d)$. For $q \geq 2$, we have $\mathbb{E}[P_k(\mathbf{x})^q]^{1/q} \leq (q-1)^{k/2} \mathbb{E}[P_k(\mathbf{x})^2]^{1/2}$.*

Lemma 9. *Let $\mathbf{x} \sim \mathcal{N}(0, \mathbf{I}_d)$ and $\mathbf{S} \in \mathbb{R}^{d \times d}$ be a symmetric matrix. For $u > 0$,*

$$\mathbb{P}[|\mathbf{x}^\top \mathbf{S} \mathbf{x} - \text{Tr}(\mathbf{S})| \geq 2\|\mathbf{S}\|_F u + 2\|\mathbf{S}\|_2 u^2] \leq 2e^{-u^2}.$$

Proof. We note that $\mathbf{x}^\top \mathbf{S} \mathbf{x} - \text{Tr}(\mathbf{S})$ has the same distribution with $\sum_{i=1}^d \lambda_i(\mathbf{S})(Z_i^2 - 1)$, where $Z_i \sim_{iid} \mathcal{N}(0, 1)$. By using the Laurent-Massart lemma [LM00], we have the result. □

Corollary 6. *Let $y = \mathbf{x}^\top \mathbf{S} \mathbf{x} - \text{Tr}(\mathbf{S})$ and $\mathbf{x} \sim \mathcal{N}(0, \mathbf{I}_d)$. For $p \geq 2$, we have $\mathbb{E}[|y|^p]^{\frac{1}{p}} \leq (p-1)\sqrt{2}\|\mathbf{S}\|_F$.*

Proof. By observing that $\mathbb{E}[|y|^2] = 2\|\mathbf{S}\|_F^2$, we have the result. □

Corollary 7. *For $\mathbf{A} \in \mathbb{R}^{d \times r}$, $p \geq 2$ and $\mathbf{x} \sim \mathcal{N}(0, \mathbf{I}_d)$, we have $\mathbb{E}[\|\mathbf{A}^\top \mathbf{x}\|_2^{2p}]^{\frac{1}{p}} \leq \sqrt{3}(p-1)\text{Tr}(\mathbf{A}^\top \mathbf{A})$.*

Proof. By Lemma 8, we have $\mathbb{E}[\|\mathbf{A}^\top \mathbf{x}\|_2^{2p}]^{\frac{1}{p}} \leq (p-1)\mathbb{E}[\|\mathbf{A}^\top \mathbf{x}\|_2^4]^{\frac{1}{2}}$. For $\mathbf{S} = \mathbf{A}\mathbf{A}^\top$, we have

$$\mathbb{E}[\|\mathbf{A}^\top \mathbf{x}\|_2^4] = \mathbb{E}[(\mathbf{x}^\top \mathbf{S} \mathbf{x})^2] = \text{Tr}(\mathbb{E}[(\mathbf{x}^\top \mathbf{S} \mathbf{x}) \mathbf{x} \mathbf{x}^\top] \mathbf{S}).$$

We have

$$\mathbb{E}[(\mathbf{x}^\top \mathbf{S} \mathbf{x}) \mathbf{x} \mathbf{x}^\top] = \text{Tr}(\mathbf{S}) \mathbf{I}_d + 2\mathbf{S} \Rightarrow \mathbb{E}[\|\mathbf{A}^\top \mathbf{x}\|_2^4] = \text{Tr}(\mathbf{S})^2 + 2\|\mathbf{S}\|_F^2 \stackrel{(a)}{\leq} 3\text{Tr}(\mathbf{S})^2,$$

where (a) follows that $\mathbf{S}$ is positive semi-definite. Since $\text{Tr}(\mathbf{S}) = \text{Tr}(\mathbf{A}^\top \mathbf{A})$, we have the statement. □

Proposition 31. *Let $\mathbf{x}_j \sim_{i.i.d} \mathcal{N}(0, \mathbf{I}_r)$, for $j \in [N]$. There exists a constant $c > 0$ such that for $\delta = \frac{u(r + \sqrt{Cr \log d} + C \log d)}{\sqrt{N}}$, we have*

$$\begin{aligned} \mathbb{P} \left[\sup_{\substack{\mathbf{S} \in \mathbb{R}^{r \times r} \\ \|\mathbf{S}\|_F = 1}} \left| \frac{1}{N} \sum_{j=1}^N \frac{1}{2} \text{Tr}(\mathbf{S}(\mathbf{x}_j \mathbf{x}_j^\top - \mathbf{I}_r))^2 - 1 \right| \geq \max\{2\delta, \delta^2\} + 10d^{-C/2} \right] \\ \leq d^2 \exp(-cu^2) + 2Nd^{-C}. \end{aligned}$$

Proof. We observe that

$$\frac{1}{2} \|\mathbf{x}_j \mathbf{x}_j^\top - \mathbf{I}_r\|_F \leq \frac{1}{2} (\|\mathbf{x}_j\|_2^2 + r).$$

By using Lemma 9, we can derive

$$\mathbb{P} \left[\underbrace{\|\mathbf{x}_j\|_2^2}_{=: \mathcal{E}_j} \leq r + 2\sqrt{r} \sqrt{C \log d} + 2C \log d \right] \geq 1 - 2d^{-C}.$$

We have

$$\begin{aligned} \left| \mathbb{E} \left[\frac{1}{2} \text{Tr}(\mathbf{S}(\mathbf{x}_j \mathbf{x}_j^\top - \mathbf{I}_r))^2 \mathbb{1}_{\mathcal{E}_j} \right] - 1 \right| &= \frac{1}{2} \mathbb{E} [\text{Tr}(\mathbf{S}(\mathbf{x}_j \mathbf{x}_j^\top - \mathbf{I}_r))^2 \mathbb{1}_{\mathcal{E}_j^c}] \\ &\leq \frac{1}{2} \mathbb{E} [\text{Tr}(\mathbf{S}(\mathbf{x}_j \mathbf{x}_j^\top - \mathbf{I}_r))^4]^{1/2} \sqrt{2} d^{-C/2} \\ &\leq 9\sqrt{2} d^{-C/2}. \end{aligned}$$

By using [Ver10, Theorem 5.41], for $\delta = \frac{u(r+\sqrt{Cr \log d} + C \log d)}{\sqrt{N}}$, we have

$$\begin{aligned} &\mathbb{P} \left[\sup_{\substack{\mathbf{S} \in \mathbb{R}^{r \times r} \\ \|\mathbf{S}\|_F = 1}} \left| \frac{1}{N} \sum_{j=1}^N \frac{1}{2} \text{Tr}(\mathbf{S}(\mathbf{x}_j \mathbf{x}_j^\top - \mathbf{I}_r))^2 - 1 \right| \geq \max\{2\delta, \delta^2\} + 10d^{-C/2} \right] \\ &\leq \mathbb{P} \left[\sup_{\substack{\mathbf{S} \in \mathbb{R}^{r \times r} \\ \|\mathbf{S}\|_F = 1}} \left| \frac{1}{N} \sum_{j=1}^N \frac{1}{2} \text{Tr}(\mathbf{S}(\mathbf{x}_j \mathbf{x}_j^\top - \mathbf{I}_r))^2 \mathbb{1}_{\mathcal{E}_j} - \mathbb{E} \left[\frac{1}{2} \text{Tr}(\mathbf{S}(\mathbf{x}_j \mathbf{x}_j^\top - \mathbf{I}_r))^2 \mathbb{1}_{\mathcal{E}_j} \right] \right| \geq \max\{2\delta, \delta^2\} \right] \\ &\quad + 2Nd^{-C} \\ &\leq d^2 \exp(-cu^2) + 2Nd^{-C}. \end{aligned}$$

□

Proposition 32. Let $\mathbf{x}_j \sim_{i.i.d} \mathcal{N}(0, \mathbf{I}_d)$, for $j \in [N]$, and $\mathbf{W} \in \mathbb{R}^{d \times r}$ be an orthonormal matrix. For a fixed $\mathbf{S} \in \mathbb{R}^{d \times d}$, $C \geq 16$ and $N \geq Cr \log d$, we have

$$\begin{aligned} \mathbb{P} \left[\left\| \frac{1}{N} \sum_{j=1}^N \frac{1}{2} \text{Tr}(\mathbf{S}(\mathbf{x}_j \mathbf{x}_j^\top - \mathbf{I}_d)) \mathbf{W}^\top (\mathbf{x}_j \mathbf{x}_j^\top - \mathbf{I}_d) \mathbf{W} - \mathbf{W}^\top \mathbf{S} \mathbf{W} \right\|_2 \geq 24e \|\mathbf{S}\|_F \left(\sqrt{\frac{Cr}{N}} + d^{-\frac{C}{2}} \right) \right] \\ \leq 2e^{-\frac{Cr}{8}} + 2Nd^{-C}. \end{aligned}$$

Proof. Without loss of generality, we assume $\|\mathbf{S}\|_F = 1$. By using Lemma 9, we have

$$\mathbb{P} \left[\underbrace{|\text{Tr}(\mathbf{S}(\mathbf{x}_j \mathbf{x}_j^\top - \mathbf{I}_d))|}_{=:\mathcal{E}_j} \leq 4\sqrt{C \log d} \right] \geq 1 - 2d^{-C}.$$

For the following, we fix a $\mathbf{v} \in S^{d-1}$. First, to bound the bias due to clipping, we write:

$$\begin{aligned} &\left| \mathbb{E} \left[\text{Tr}(\mathbf{S}(\mathbf{x}_j \mathbf{x}_j^\top - \mathbf{I}_d)) (\langle \mathbf{v}, \mathbf{x} \rangle^2 - 1) \mathbb{1}_{\mathcal{E}_j^c} \right] \right| \\ &\leq \mathbb{E} \left[\text{Tr}(\mathbf{S}(\mathbf{x}_j \mathbf{x}_j^\top - \mathbf{I}_d))^4 \right]^{\frac{1}{4}} \mathbb{E} \left[(\langle \mathbf{v}, \mathbf{x} \rangle^2 - 1)^4 \right]^{\frac{1}{4}} \sqrt{2} d^{-C/2} \leq 18\sqrt{2} d^{-C/2}. \end{aligned}$$

On the other hand, to bound the moments of the clipped random variable, we have for $p \geq 2$,

$$\begin{aligned} &\mathbb{E} \left[|\text{Tr}(\mathbf{S}(\mathbf{x}_j \mathbf{x}_j^\top - \mathbf{I}_d)) (\langle \mathbf{v}, \mathbf{x} \rangle^2 - 1)|^p \mathbb{1}_{\mathcal{E}_j} \right] \\ &\leq (4\sqrt{C \log d})^{p-2} \mathbb{E} \left[\text{Tr}(\mathbf{S}(\mathbf{x}_j \mathbf{x}_j^\top - \mathbf{I}_d))^2 |(\langle \mathbf{v}, \mathbf{x} \rangle^2 - 1)|^p \right] \leq (12e)^2 (8\sqrt{2}e\sqrt{C \log d})^{p-2} \frac{p!}{2}. \end{aligned}$$

By using ε -cover argument, we can derive

$$\begin{aligned} &\mathbb{P} \left[\left\| \frac{1}{N} \sum_{j=1}^N \frac{1}{2} \text{Tr}(\mathbf{S}(\mathbf{x}_j \mathbf{x}_j^\top - \mathbf{I}_d)) \mathbf{W}^\top (\mathbf{x}_j \mathbf{x}_j^\top - \mathbf{I}_d) \mathbf{W} - \mathbf{W}^\top \mathbf{S} \mathbf{W} \right\|_2 \geq 24eu + 18\sqrt{2} d^{-C/2} \right] \\ &\leq 2 \cdot 9^r \exp \left(\frac{-Nu^2/2}{1 + u\sqrt{C \log d}} \right) + 2Nd^{-C}. \end{aligned}$$

By using $u = \sqrt{Cr/N}$, we have the result. □

Proof. Without loss of generality, we assume $\|\mathbf{S}\|_F = 1$. We have

$$\left\| \sum_{j=1}^N y_j (\mathbf{x}_j \mathbf{x}_j^\top - \mathbf{I}_d) - \mathbf{S} \right\|_F^2 \leq \sup_{\substack{\mathbf{S} \in \mathbb{R}^{r \times r} \\ \|\mathbf{S}\|_F = 1}} \left| \frac{1}{N} \sum_{j=1}^N \frac{1}{2} \text{Tr}(\mathbf{S}(\mathbf{x}_j \mathbf{x}_j^\top - \mathbf{I}_r)) \right|^2.$$

Hence, by considering the event in Proposition 31, we have the statement. $\square$

Proposition 33. Let $X \in \mathbb{R}$ be a random variable such that for some $K, C > 0$, $\mathbb{E}[|X|^p] \leq CK^p p^{pc}$ for some $c > 0$ and $p \geq k$. Then, $\mathbb{P}[|X| \geq Ku] \leq Ce^{-\frac{u^{1/c}}{e}}$ for $u \geq (ke)^c$.

Proof. Use Markov inequality with $p = \frac{u^{1/c}}{e}$. $\square$

G.3 Miscellaneous

Proposition 34. We consider $\eta \leq \frac{1}{10}$. The following statements holds:

- For $0.2 \geq \delta > 0$, let

$$u_{t+1} = u_t + \eta u_t(1 - u_t), \quad 1 + \delta \geq u_0 \geq 0.$$

We have $1 + (\delta \vee \frac{\eta^2}{4}) \geq \sup_t u_t \geq 0$. Moreover, $t^* = \inf\{t : u_t \geq 1\}$, we have $u_{t+1} \geq u_t$ for $t < t^*$ and $u_{t^*} \geq u_t \geq 1$ for $t \geq t^*$.

- For $0.5 > \varepsilon > 0$ and $1.1 > \bar{u}_0 \geq u_0 \geq \underline{u}_0 > 0$, let

$$\bar{u}_{t+1} = \bar{u}_t + \eta(1 + \varepsilon)\bar{u}_t(1 - \bar{u}_t) \quad \text{and} \quad \underline{u}_{t+1} = \underline{u}_t + \eta \underline{u}_t(1 - \underline{u}_t).$$

and

$$u_t + \eta u_t(1 - u_t) \leq u_{t+1} \leq u_t + \eta(1 + \varepsilon)u_t(1 - u_t).$$

We have

$$\frac{1}{2} \left(1 \wedge \underline{u}_0 e^{\frac{\eta t}{1+\eta}} \right) \leq \underline{u}_t \leq u_t \leq \bar{u}_t \leq \left(1.1 \wedge \bar{u}_0 e^{\eta(1+\varepsilon)t} \right).$$

Proof. If $t \geq t^*$, by monotonicity of the update, we have $1 \leq u_{t+1} \leq u_t \leq u_{t^*}$. If $t^* > 0$, then for $t < t^*$, we have $1 \geq u_t \geq 0$ and $u_t(1 - u_t) \geq 0$, and thus, we have $u_{t+1} \geq u_t \geq 0$. Next, we observe that $u_{t^*} \leq 1 + 0.25\eta$ and by monotonicity of the update for $t \geq t^*$, we have $1 \leq u_t \leq u_{t^*}$. Hence, it is sufficient to bound u_{t^*} to bound $\sup_t u_t$. Note that, we have $1 \geq u_{t^*-1} \geq 1 - 0.25\eta$, and thus,

$$\frac{u_{t^*}}{u_{t^*-1}} = 1 + \eta(1 - u_{t^*-1}) \leq 1 + \eta^2 \Rightarrow u_{t^*} \leq 1 + \frac{\eta^2}{4}.$$

For the second item, by monotonicity, we have $0 < \underline{u}_t \leq u_t \leq \bar{u}_t < 1.1$. Moreover, by [AGP24, Lemma A.2], we have for $t < t_u := \inf\{t : \underline{u}_t \geq 0.5\}$

$$\frac{\underline{u}_0 e^{\frac{\eta t}{1+\eta}}}{1 + \underline{u}_0 e^{\frac{\eta t}{1+\eta}}} \leq \underline{u}_t \Rightarrow \frac{\underline{u}_0}{2} e^{\frac{\eta t}{1+\eta}} \leq \underline{u}_0.$$

For $t \geq t_u$, by the first item, we have $\underline{u}_t \geq 0.5$. Therefore, we have

$$\frac{1}{2} \left(\underline{u}_0 e^{\frac{\eta t}{1+\eta}} \wedge 1 \right) \leq \underline{u}_t.$$

On the other hand, for all $t \in \mathbb{N}$, we have $\bar{u}_t \leq \bar{u}_0 e^{\eta(1+\varepsilon)t}$. By the first item, we have $\bar{u}_t \leq \bar{u}_0 e^{\eta(1+\varepsilon)t} \wedge 1.1$. $\square$

Proposition 35. For $t, \lambda > 0$, we have

$$\frac{1}{t \exp(t\lambda)} \leq \frac{\lambda}{\exp(t\lambda) - 1} \leq \frac{1}{t}.$$

Proof. The upper bound follows $\exp(t\lambda) - 1 \geq t\lambda$. For the lower bound,

$$\frac{1}{t} - \frac{\lambda}{\exp(t\lambda) - 1} = \frac{\exp(t\lambda) - t\lambda - 1}{t(\exp(t\lambda) - 1)}. \quad (\text{G.5})$$

We have

$$\exp(t\lambda) - t\lambda - 1 \leq \sum_{k=2}^{\infty} \frac{(t\lambda)^k}{k!} = t\lambda \sum_{k=1}^{\infty} \frac{(t\lambda)^k}{(k+1)!} \leq t\lambda \sum_{k=1}^{\infty} \frac{(t\lambda)^k}{k!} = t\lambda(\exp(t\lambda) - 1).$$

Therefore,

$$(\text{G.5}) \leq \lambda \Rightarrow \frac{1}{t} \leq \frac{\lambda}{\exp(t\lambda) - 1} + \lambda \Rightarrow \frac{1}{t \exp(t\lambda)} \leq \frac{\lambda}{\exp(t\lambda) - 1}.$$

□

Lemma 10. Let $r_s \asymp d^\gamma$, $\gamma \in [0, 1)$, and $\log^{-1} d \ll C_d \ll \log^{10} d$. We define F_d, G_d, H_d as

$$F_d(u) := \left(1 - \frac{1}{1 + \left(\frac{dC_d}{r_s} \frac{1}{u} - 1 \right) \left(\frac{d}{r_s} \right)^{-\frac{1}{u}}} \right)^2, \quad G_d(u) := 1 - \frac{1}{1 + \left(\frac{dC_d}{r_s} - 1 \right) \left(\frac{d}{r_s} \right)^{-\frac{1}{u}}},$$

$$H_d(u) := \left(1 - C_d \left(\frac{d}{r_s} \right)^{\frac{1}{u} - 1} \right)_+,$$

We have

- For any $C > 0$, $\sup_{u \leq \log^C d} |F_d(u)| \leq 1$ for $d \geq \Omega_C(1)$.
- $\sup_u |G_d(u)| \vee |H_d(u)| \leq 1$.
- For any $\delta \in (0, 0.5)$, let $\mathcal{C}_\delta := \{u \geq 0 : |u - 1| < \delta\}$. For any compact $\mathcal{K} \subset (0, \infty] \setminus \mathcal{C}_\delta$, we have $F_d(u) \xrightarrow{d \rightarrow \infty} \mathbb{1}\{u > 1\}$ uniformly on $\mathcal{K}$.
- For any compact $\mathcal{K} \subseteq [0, \infty] \setminus \mathcal{C}_\delta$, we have

$$G_d(u), H_d(u) \xrightarrow{d \rightarrow \infty} \mathbb{1}\{u > 1\}, \quad G_d^2(u) \xrightarrow{d \rightarrow \infty} \mathbb{1}\{u > 1\}$$

all uniformly on $\mathcal{K}$.

Proof. For the first item, if $u \leq \log^C d$, for $d \geq \Omega_C(1)$

$$\frac{dC_d u}{r_s} - 1 \geq \frac{d}{r_s} \frac{t}{\log^{C+1} d} - 1 > 0.$$

Therefore, $|F_d(u)| \leq 1$. For the second item, since $\frac{dC_d}{r_s} > 1$ for $d \geq \Omega(1)$, the item follows.

For the third item, since $E := [0, \infty) \setminus \mathcal{C}_\delta$ is closed in $[0, \infty)$, it suffices to establish the result on small open intervals around each point of E within $[0, \infty)$. Fix $u_0 \in E$ and choose $\epsilon \in (0, \delta/2)$. Since $B(u_0, \epsilon) := (u_0 - \epsilon, u_0 + \epsilon) \cap [0, \infty)$ is convex it can be either in $P_> := \{u : u > 1 + \delta/2\}$ or $P_< := \{u : u < 1 - \delta/2\}$. Without loss of generality let us assume it is in $P_<$. Then,

$$\sup_{u \in B(u_0, \epsilon) \subset P_<} |F_d(u)| \leq 1 - \frac{1}{1 + \left(O_\delta(C_d) - \left(\frac{d}{r_s} \right)^{-1} \right) \left(\frac{d}{r_s} \right)^{-O_\delta(1)}} \rightarrow 0.$$

A similar step can be repeated if $B_\epsilon \subset P_<$.

For the last item, we first observe that uniform convergence of $G_d(u)$ implies the uniform convergence of $G_d^2(u)$. Therefore, we will only prove the first result. Since $E := [0, \infty) \setminus \mathcal{C}_\delta$ is compact, and thus, $P_> \cap E$ and $P_< \cap E$ are also compact, we can directly use these sets. Without loss of generality let us use $P_< \cap E$. Then,

$$\sup_{u \in P_< \cap E} |G_d(u)| \vee |H_d(u)| \leq \left(1 - C_d \frac{d}{r_s}^{O_\delta(1)} \right)_+ \rightarrow 0.$$

A similar step can be repeated if $P_> \cap E$. Therefore, the statement follows. □

Proposition 36. Let $r_u \leq r$ and

$$t \in \begin{cases} (0, \infty), & \alpha \in [0, 0.5) \\ (0, \infty) \setminus \{j^\alpha : j \in \mathbb{N}\}, & \alpha > 0.5, \end{cases} \quad \kappa_{\text{eff}} := \begin{cases} r^\alpha, & \alpha \in [0, 0.5) \\ 1, & \alpha > 0.5. \end{cases}$$

We have

- For $K \in \{G, H\}$ and $t \neq \lim_{d \rightarrow \infty} \frac{1}{\lambda_j \kappa_{\text{eff}}}$, we have

$$K_d\left(\frac{1}{\lambda_j t \kappa_{\text{eff}}}\right) - \mathbb{1}\left\{\frac{1}{\lambda_j} > t \kappa_{\text{eff}}\right\} = o_d(1).$$

- For $K \in \{F, G, H\}$,

$$\frac{1}{\|\Lambda\|_F^2} \sum_{j=1}^{r_u} \lambda_j^2 \left(K_d\left(\frac{1}{\lambda_j t \kappa_{\text{eff}}}\right) - \mathbb{1}\left\{\frac{1}{\lambda_j} > t \kappa_{\text{eff}}\right\} \right) = o_d(1). \quad (\text{G.6})$$

Proof. The first item immediately follows Lemma 10. In the following, we will prove the second item for the heavy and light tailed cases separately.

For $\alpha \in [0, 0.5)$: We define a sequence of measures $\mu_d\{j/r\} \propto j^{-2\alpha}$, $j \leq [r]$. We observe that

- We have $\mu_d \rightarrow \mu$ weakly such that μ is supported on $[0, 1]$ and $\mu([0, \tau]) = \tau^{1-2\alpha}$ for $\tau \in [0, 1]$.
- Moreover, $(\text{G.6}) = \mathbb{E}_{X \sim \mu_d}[(K_d(X^\alpha/t) - \mathbb{1}\{X^\alpha > t\})\mathbb{1}\{X \leq \frac{r_u}{r}\}]$.

By using the $\mathcal{C}_\delta$ definition in Lemma 10:

$$\begin{aligned} & |\mathbb{E}_{X \sim \mu_d}[(K_d(X^\alpha/t) - \mathbb{1}\{X^\alpha > t\})\mathbb{1}\{X \leq \frac{r_u}{r}\}]| \\ & \leq \mathbb{E}_{X \sim \mu_d}[|K_d(X^\alpha/t) - \mathbb{1}\{X^\alpha > t\}|\mathbb{1}\{X^\alpha \in [0, 1] \setminus \mathcal{C}_\delta\}] \\ & + \mathbb{E}_{X \sim \mu_d}[|K_d(X^\alpha/t) - \mathbb{1}\{X^\alpha > t\}|\mathbb{1}\{X^\alpha \in \mathcal{C}_\delta\}] \\ & \stackrel{(a)}{\leq} o_d(1) + \mathbb{P}_{X \sim \mu}[X^\alpha \in \mathcal{C}_\delta], \end{aligned}$$

where we used the second item in Lemma 10 for (a). Since $\mathbb{P}_{X \sim \mu}[X^\alpha \in \mathcal{C}_\delta] \xrightarrow{\delta \rightarrow 0} 0$, we have the first result.

For $\alpha > 0.5$: We define a sequence of measures $\mu_d\{j\} \propto j^{-2\alpha}$, $j \leq [r]$. We observe that

- We have $\mu_d \rightarrow \mu$ weakly such that $\mu\{j\} \propto j^{-2\alpha}$ for $j \in \mathbb{N}$.
- Moreover, $(\text{G.6}) = \mathbb{E}_{X \sim \mu_d}[(K_d(X^\alpha/t) - \mathbb{1}\{X^\alpha > t\})\mathbb{1}\{X \leq r_u\}]$.

Let $t \in ((j-1)^\alpha, j^\alpha)$ for some $j \in \mathbb{N}$. For small enough $\delta > 0$, we have

$$\begin{aligned} & |\mathbb{E}_{X \sim \mu_d}[(K_d(X^\alpha/t) - \mathbb{1}\{X^\alpha > t\})\mathbb{1}\{X \leq r_u\}]| \\ & = \mathbb{E}_{X \sim \mu_d}[|K_d(X^\alpha/t) - \mathbb{1}\{X^\alpha > t\}|\mathbb{1}\{X \in [0, r_u]\}\mathbb{1}\{X^\alpha \notin \mathcal{C}_\delta\}] \stackrel{(b)}{=} o_d(1), \end{aligned}$$

where we used both items in Lemma 10 for (b). □

Corollary 8. For $1 \geq c_d \gg \log^{-5} d$, we define

$$g_d(\lambda, t) := \frac{-\lambda \exp(-t\lambda)}{1 - \exp(-t\lambda)} + \frac{\lambda^2 \exp(-t\lambda)}{(1 - \exp(-t\lambda))^2} \left(\frac{c_d}{t} \frac{r_s}{d} + \frac{\lambda \exp(-t\lambda)}{1 - \exp(-t\lambda)} \right)^{-1}$$

Let $r_u \leq r$ and

$$\kappa_{\text{eff}} := \begin{cases} r^\alpha, & \alpha \in [0, 0.5) \\ 1, & \alpha > 0.5 \end{cases}, \quad \mathsf{T}_{\text{eff}} := \kappa_{\text{eff}} \log d/r_s.$$

We have

$$\frac{1}{\|\mathbf{\Lambda}\|_{\text{F}}^2} \left(\sum_{j=1}^{r_u} g_d^2(\lambda_j; t\mathsf{T}_{\text{eff}}) - \sum_{j=1}^{r_u} \lambda_j^2 \mathbb{1}\left\{\frac{1}{\lambda_j} > t\kappa_{\text{eff}}\right\} \right) = o_d(1)$$

for any fixed

$$t \in \begin{cases} (0, \infty), & \alpha \in [0, 0.5) \\ (0, \infty) \setminus \{j^\alpha : j \in \mathbb{N}\}, & \alpha > 0.5. \end{cases}$$

Proof. We observe that

$$g_d(\lambda; t) = \lambda \left(1 - \frac{1}{1 - \exp(-t\lambda) + \frac{d}{r_s} \frac{\lambda t}{c_d} \exp(-t\lambda)} \right)$$

Therefore, we have $g_d^2(\lambda; t\mathsf{T}_{\text{eff}}) = \lambda^2 F_d(\frac{1}{\lambda_j t\kappa_{\text{eff}}})$. Then, by Proposition 36

$$\begin{aligned} \frac{1}{\|\mathbf{\Lambda}\|_{\text{F}}^2} \left(\sum_{j=1}^{r_u} g_d^2(\lambda_j; t\mathsf{T}_{\text{eff}}) - \sum_{j=1}^{r_u} \lambda_j^2 \mathbb{1}\left\{\frac{1}{\lambda_j} \geq t\kappa_{\text{eff}}\right\} \right) \\ = \frac{1}{\|\mathbf{\Lambda}\|_{\text{F}}^2} \sum_{j=1}^{r_u} \lambda_j^2 \left(F_d\left(\frac{1}{\lambda_j t\kappa_{\text{eff}}}\right) - \mathbb{1}\left\{\frac{1}{\lambda_j} > t\kappa_{\text{eff}}\right\} \right) = o_d(1). \end{aligned}$$

□

Checklist

1. Claims

Question: Do the main claims made in the abstract and introduction accurately reflect the paper's contributions and scope?

Answer: [\[Yes\]](#)

Justification: : The main claims made in the abstract and introduction are supported by our main theorems, Theorem 1 and 2, and their corollaries, Corollary 1 and 2.

Guidelines:

- The answer NA means that the abstract and introduction do not include the claims made in the paper.
- The abstract and/or introduction should clearly state the claims made, including the contributions made in the paper and important assumptions and limitations. A No or NA answer to this question will not be perceived well by the reviewers.
- The claims made should match theoretical and experimental results, and reflect how much the results can be expected to generalize to other settings.
- It is fine to include aspirational goals as motivation as long as it is clear that these goals are not attained by the paper.

2. Limitations

Question: Does the paper discuss the limitations of the work performed by the authors?

Answer: [\[Yes\]](#)

Justification: Yes, we discuss limitations in the Conclusion part (see Section 6)

Guidelines:

- The answer NA means that the paper has no limitation while the answer No means that the paper has limitations, but those are not discussed in the paper.
- The authors are encouraged to create a separate "Limitations" section in their paper.
- The paper should point out any strong assumptions and how robust the results are to violations of these assumptions (e.g., independence assumptions, noiseless settings, model well-specification, asymptotic approximations only holding locally). The authors should reflect on how these assumptions might be violated in practice and what the implications would be.
- The authors should reflect on the scope of the claims made, e.g., if the approach was only tested on a few datasets or with a few runs. In general, empirical results often depend on implicit assumptions, which should be articulated.
- The authors should reflect on the factors that influence the performance of the approach. For example, a facial recognition algorithm may perform poorly when image resolution is low or images are taken in low lighting. Or a speech-to-text system might not be used reliably to provide closed captions for online lectures because it fails to handle technical jargon.
- The authors should discuss the computational efficiency of the proposed algorithms and how they scale with dataset size.
- If applicable, the authors should discuss possible limitations of their approach to address problems of privacy and fairness.
- While the authors might fear that complete honesty about limitations might be used by reviewers as grounds for rejection, a worse outcome might be that reviewers discover limitations that aren't acknowledged in the paper. The authors should use their best judgment and recognize that individual actions in favor of transparency play an important role in developing norms that preserve the integrity of the community. Reviewers will be specifically instructed to not penalize honesty concerning limitations.

3. Theory assumptions and proofs

Question: For each theoretical result, does the paper provide the full set of assumptions and a complete (and correct) proof?

Answer: [\[Yes\]](#)

Justification: The details, including assumptions, for the theoretical setup are detailed in Section 2. The proofs are presented in appendix (see the Supplementary file)

Guidelines:

- The answer NA means that the paper does not include theoretical results.
- All the theorems, formulas, and proofs in the paper should be numbered and cross-referenced.
- All assumptions should be clearly stated or referenced in the statement of any theorems.
- The proofs can either appear in the main paper or the supplemental material, but if they appear in the supplemental material, the authors are encouraged to provide a short proof sketch to provide intuition.
- Inversely, any informal proof provided in the core of the paper should be complemented by formal proofs provided in appendix or supplemental material.
- Theorems and Lemmas that the proof relies upon should be properly referenced.

4. Experimental result reproducibility

Question: Does the paper fully disclose all the information needed to reproduce the main experimental results of the paper to the extent that it affects the main claims and/or conclusions of the paper (regardless of whether the code and data are provided or not)?

Answer: [Yes]

Justification: Details to reproduce the simulations in Figure 1b are provided in its caption.

Guidelines:

- The answer NA means that the paper does not include experiments.
- If the paper includes experiments, a No answer to this question will not be perceived well by the reviewers: Making the paper reproducible is important, regardless of whether the code and data are provided or not.
- If the contribution is a dataset and/or model, the authors should describe the steps taken to make their results reproducible or verifiable.
- Depending on the contribution, reproducibility can be accomplished in various ways. For example, if the contribution is a novel architecture, describing the architecture fully might suffice, or if the contribution is a specific model and empirical evaluation, it may be necessary to either make it possible for others to replicate the model with the same dataset, or provide access to the model. In general, releasing code and data is often one good way to accomplish this, but reproducibility can also be provided via detailed instructions for how to replicate the results, access to a hosted model (e.g., in the case of a large language model), releasing of a model checkpoint, or other means that are appropriate to the research performed.
- While NeurIPS does not require releasing code, the conference does require all submissions to provide some reasonable avenue for reproducibility, which may depend on the nature of the contribution. For example
 - (a) If the contribution is primarily a new algorithm, the paper should make it clear how to reproduce that algorithm.
 - (b) If the contribution is primarily a new model architecture, the paper should describe the architecture clearly and fully.
 - (c) If the contribution is a new model (e.g., a large language model), then there should either be a way to access this model for reproducing the results or a way to reproduce the model (e.g., with an open-source dataset or instructions for how to construct the dataset).
 - (d) We recognize that reproducibility may be tricky in some cases, in which case authors are welcome to describe the particular way they provide for reproducibility. In the case of closed-source models, it may be that access to the model is limited in some way (e.g., to registered users), but it should be possible for other researchers to have some path to reproducing or verifying the results.

5. Open access to data and code

Question: Does the paper provide open access to the data and code, with sufficient instructions to faithfully reproduce the main experimental results, as described in supplemental material?

Answer: [NA]

Justification: This is a theory paper; toy experiments are conducted on Gaussian data.

Guidelines:

- The answer NA means that paper does not include experiments requiring code.
- Please see the NeurIPS code and data submission guidelines (<https://nips.cc/public/guides/CodeSubmissionPolicy>) for more details.
- While we encourage the release of code and data, we understand that this might not be possible, so “No” is an acceptable answer. Papers cannot be rejected simply for not including code, unless this is central to the contribution (e.g., for a new open-source benchmark).
- The instructions should contain the exact command and environment needed to run to reproduce the results. See the NeurIPS code and data submission guidelines (<https://nips.cc/public/guides/CodeSubmissionPolicy>) for more details.
- The authors should provide instructions on data access and preparation, including how to access the raw data, preprocessed data, intermediate data, and generated data, etc.
- The authors should provide scripts to reproduce all experimental results for the new proposed method and baselines. If only a subset of experiments are reproducible, they should state which ones are omitted from the script and why.
- At submission time, to preserve anonymity, the authors should release anonymized versions (if applicable).
- Providing as much information as possible in supplemental material (appended to the paper) is recommended, but including URLs to data and code is permitted.

6. Experimental setting/details

Question: Does the paper specify all the training and test details (e.g., data splits, hyper-parameters, how they were chosen, type of optimizer, etc.) necessary to understand the results?

Answer: [Yes]

Justification: The details of our toy experiments are provided in the caption of Figure 1.

Guidelines:

- The answer NA means that the paper does not include experiments.
- The experimental setting should be presented in the core of the paper to a level of detail that is necessary to appreciate the results and make sense of them.
- The full details can be provided either with the code, in appendix, or as supplemental material.

7. Experiment statistical significance

Question: Does the paper report error bars suitably and correctly defined or other appropriate information about the statistical significance of the experiments?

Answer: [Yes]

Justification: Error bars are provided in Figure 1b; however, they are not visible.

Guidelines:

- The answer NA means that the paper does not include experiments.
- The authors should answer "Yes" if the results are accompanied by error bars, confidence intervals, or statistical significance tests, at least for the experiments that support the main claims of the paper.
- The factors of variability that the error bars are capturing should be clearly stated (for example, train/test split, initialization, random drawing of some parameter, or overall run with given experimental conditions).
- The method for calculating the error bars should be explained (closed form formula, call to a library function, bootstrap, etc.)
- The assumptions made should be given (e.g., Normally distributed errors).
- It should be clear whether the error bar is the standard deviation or the standard error of the mean.

- It is OK to report 1-sigma error bars, but one should state it. The authors should preferably report a 2-sigma error bar than state that they have a 96% CI, if the hypothesis of Normality of errors is not verified.
- For asymmetric distributions, the authors should be careful not to show in tables or figures symmetric error bars that would yield results that are out of range (e.g. negative error rates).
- If error bars are reported in tables or plots, The authors should explain in the text how they were calculated and reference the corresponding figures or tables in the text.

8. Experiments compute resources

Question: For each experiment, does the paper provide sufficient information on the computer resources (type of compute workers, memory, time of execution) needed to reproduce the experiments?

Answer: [NA]

Justification: This is a theory paper; toy experiments are conducted on Gaussian data.

Guidelines:

- The answer NA means that the paper does not include experiments.
- The paper should indicate the type of compute workers CPU or GPU, internal cluster, or cloud provider, including relevant memory and storage.
- The paper should provide the amount of compute required for each of the individual experimental runs as well as estimate the total compute.
- The paper should disclose whether the full research project required more compute than the experiments reported in the paper (e.g., preliminary or failed experiments that didn't make it into the paper).

9. Code of ethics

Question: Does the research conducted in the paper conform, in every respect, with the NeurIPS Code of Ethics <https://neurips.cc/public/EthicsGuidelines>?

Answer: [Yes]

Justification: Yes, this submission follows the NeurIPS Code of Ethics

Guidelines:

- The answer NA means that the authors have not reviewed the NeurIPS Code of Ethics.
- If the authors answer No, they should explain the special circumstances that require a deviation from the Code of Ethics.
- The authors should make sure to preserve anonymity (e.g., if there is a special consideration due to laws or regulations in their jurisdiction).

10. Broader impacts

Question: Does the paper discuss both potential positive societal impacts and negative societal impacts of the work performed?

Answer: [NA]

Justification: The goal of this paper is to advance the theoretical understanding of training dynamics of two-layer neural networks. There are no direct societal impacts of our work that should be highlighted here

Guidelines:

- The answer NA means that there is no societal impact of the work performed.
- If the authors answer NA or No, they should explain why their work has no societal impact or why the paper does not address societal impact.
- Examples of negative societal impacts include potential malicious or unintended uses (e.g., disinformation, generating fake profiles, surveillance), fairness considerations (e.g., deployment of technologies that could make decisions that unfairly impact specific groups), privacy considerations, and security considerations.

- The conference expects that many papers will be foundational research and not tied to particular applications, let alone deployments. However, if there is a direct path to any negative applications, the authors should point it out. For example, it is legitimate to point out that an improvement in the quality of generative models could be used to generate deepfakes for disinformation. On the other hand, it is not needed to point out that a generic algorithm for optimizing neural networks could enable people to train models that generate Deepfakes faster.
- The authors should consider possible harms that could arise when the technology is being used as intended and functioning correctly, harms that could arise when the technology is being used as intended but gives incorrect results, and harms following from (intentional or unintentional) misuse of the technology.
- If there are negative societal impacts, the authors could also discuss possible mitigation strategies (e.g., gated release of models, providing defenses in addition to attacks, mechanisms for monitoring misuse, mechanisms to monitor how a system learns from feedback over time, improving the efficiency and accessibility of ML).

11. Safeguards

Question: Does the paper describe safeguards that have been put in place for responsible release of data or models that have a high risk for misuse (e.g., pretrained language models, image generators, or scraped datasets)?

Answer: [NA]

Justification: This paper poses no such risks.

Guidelines:

- The answer NA means that the paper poses no such risks.
- Released models that have a high risk for misuse or dual-use should be released with necessary safeguards to allow for controlled use of the model, for example by requiring that users adhere to usage guidelines or restrictions to access the model or implementing safety filters.
- Datasets that have been scraped from the Internet could pose safety risks. The authors should describe how they avoided releasing unsafe images.
- We recognize that providing effective safeguards is challenging, and many papers do not require this, but we encourage authors to take this into account and make a best faith effort.

12. Licenses for existing assets

Question: Are the creators or original owners of assets (e.g., code, data, models), used in the paper, properly credited and are the license and terms of use explicitly mentioned and properly respected?

Answer: [NA]

Justification: This paper does not use existing assets.

Guidelines:

- The answer NA means that the paper does not use existing assets.
- The authors should cite the original paper that produced the code package or dataset.
- The authors should state which version of the asset is used and, if possible, include a URL.
- The name of the license (e.g., CC-BY 4.0) should be included for each asset.
- For scraped data from a particular source (e.g., website), the copyright and terms of service of that source should be provided.
- If assets are released, the license, copyright information, and terms of use in the package should be provided. For popular datasets, paperswithcode.com/datasets has curated licenses for some datasets. Their licensing guide can help determine the license of a dataset.
- For existing datasets that are re-packaged, both the original license and the license of the derived asset (if it has changed) should be provided.

- If this information is not available online, the authors are encouraged to reach out to the asset's creators.

13. **New assets**

Question: Are new assets introduced in the paper well documented and is the documentation provided alongside the assets?

Answer: [NA]

Justification: This paper does not release new assets.

Guidelines:

- The answer NA means that the paper does not release new assets.
- Researchers should communicate the details of the dataset/code/model as part of their submissions via structured templates. This includes details about training, license, limitations, etc.
- The paper should discuss whether and how consent was obtained from people whose asset is used.
- At submission time, remember to anonymize your assets (if applicable). You can either create an anonymized URL or include an anonymized zip file.

14. **Crowdsourcing and research with human subjects**

Question: For crowdsourcing experiments and research with human subjects, does the paper include the full text of instructions given to participants and screenshots, if applicable, as well as details about compensation (if any)?

Answer: [NA]

Justification: This paper does not involve human subjects.

Guidelines:

- The answer NA means that the paper does not involve crowdsourcing nor research with human subjects.
- Including this information in the supplemental material is fine, but if the main contribution of the paper involves human subjects, then as much detail as possible should be included in the main paper.
- According to the NeurIPS Code of Ethics, workers involved in data collection, curation, or other labor should be paid at least the minimum wage in the country of the data collector.

15. **Institutional review board (IRB) approvals or equivalent for research with human subjects**

Question: Does the paper describe potential risks incurred by study participants, whether such risks were disclosed to the subjects, and whether Institutional Review Board (IRB) approvals (or an equivalent approval/review based on the requirements of your country or institution) were obtained?

Answer: [NA]

Justification: This paper does not involve human subjects.

Guidelines:

- The answer NA means that the paper does not involve crowdsourcing nor research with human subjects.
- Depending on the country in which research is conducted, IRB approval (or equivalent) may be required for any human subjects research. If you obtained IRB approval, you should clearly state this in the paper.
- We recognize that the procedures for this may vary significantly between institutions and locations, and we expect authors to adhere to the NeurIPS Code of Ethics and the guidelines for their institution.
- For initial submissions, do not include any information that would break anonymity (if applicable), such as the institution conducting the review.

16. **Declaration of LLM usage**

Question: Does the paper describe the usage of LLMs if it is an important, original, or non-standard component of the core methods in this research? Note that if the LLM is used only for writing, editing, or formatting purposes and does not impact the core methodology, scientific rigorousness, or originality of the research, declaration is not required.

Answer:[NA]

Justification: We do not use LLMs.

Guidelines:

- The answer NA means that the core method development in this research does not involve LLMs as any important, original, or non-standard components.
- Please refer to our LLM policy (<https://neurips.cc/Conferences/2025/LLM>) for what should or should not be described.