

Implicit Bias and Loss of Plasticity in Matrix Completion: Depth Promotes Low-Rank Solutions

Baekrok Shin

KAIST AI, Seoul, Republic of Korea

BR.SHIN@KAIST.AC.KR

Chulhee Yun

KAIST AI, Seoul, Republic of Korea

CHULHEE.YUN@KAIST.AC.KR

Abstract

We study matrix completion via deep matrix factorization (a.k.a. deep linear neural networks) as a simplified testbed to examine how network depth influences training dynamics. Despite the simplicity and importance of the problem, prior theory largely focuses on shallow (depth-2) models and does not fully explain the implicit low-rank bias observed in deeper networks. We identify *coupled dynamics* as a key mechanism behind this bias and show that it intensifies with increasing depth. Focusing on gradient flow under diagonal observations, we prove: (a) networks of depth ≥ 3 exhibit coupling unless initialized diagonally, and (b) convergence to rank-1 occurs if and only if the dynamics is coupled—resolving an open question by Menon [24] for a family of initializations. We also revisit the *loss of plasticity* phenomenon in matrix completion [18], where pre-training on few observations and resuming with more degrades performance. We show that deep models avoid plasticity loss due to their low-rank bias, whereas depth-2 networks pre-trained under decoupled dynamics fail to converge to low-rank, even when resumed training (with additional data) satisfies the coupling condition—shedding light on the mechanism behind this phenomenon.

1. Introduction

Matrix completion, a task with practical applications in areas like recommender systems and image restoration, provides a key framework for investigating implicit biases, particularly the tendency towards low-rank solutions. The goal of the matrix completion task is to recover a low-rank ground truth matrix $\mathbf{W}^*$ using only a subset of its entries. A common strategy for matrix completion involves matrix factorization, which can also be viewed as linear neural networks. These networks reparameterize the target matrix $\mathbf{X}$ as a product of factors, $\mathbf{X} = \mathbf{W}_L \mathbf{W}_{L-1} \cdots \mathbf{W}_1$, and train these factors $\mathbf{W}_i$ by minimizing the mean squared error on the observed entries via gradient descent. The observed entries constitute the training set, while the unobserved entries act as the test set.

The problem of predicting $\mathbf{W}^*$ is underdetermined, as infinitely many completions are possible. Nevertheless, both theory and experiments indicate that training even a simple two-layer factorization ($L = 2$) with gradient descent, without explicit rank constraints, typically yields a low-rank solution under reasonable assumptions [5, 23, 27]. Bai et al. [5] recently formalize this phenomenon using the concept of *data connectivity*. They demonstrate that if the observed entries form a connected bipartite graph (meaning any observed entry can be reached from any other via shared rows or columns), a depth-2 factorization initialized at an infinitesimally small scale converges to a low-rank solution. Conversely, the network may converge to a higher-rank matrix if the observations are disconnected (see Definition 8 and Figure 1(a)).

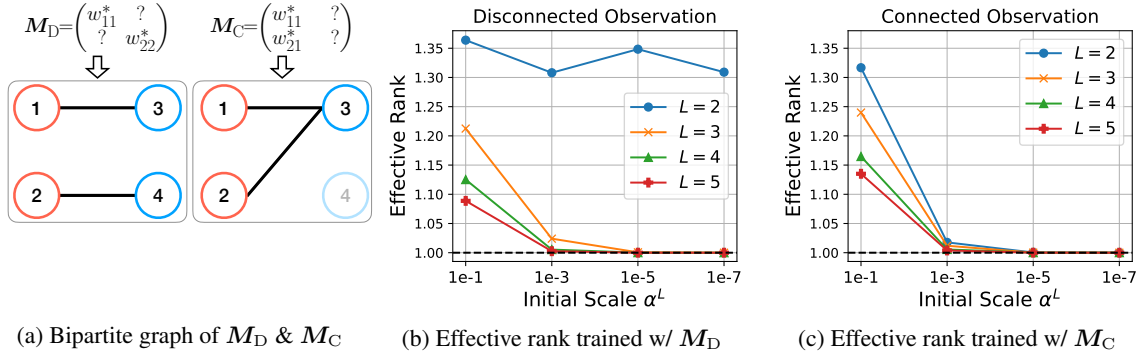


Figure 1: **(a)** Examples of bipartite graphs corresponding to observation patterns of M_D (disconnected) and M_C (connected). **(b-c)** Training results showing effective rank (cf. [28]) for completing rank-1 matrices M_D and M_C , respectively. The rank-1 ground truth matrices were generated via $uv^\top$, where $u, v \in \mathbb{R}^2$ with entries sampled i.i.d. from a standard normal distribution. We initialized each layer’s entries by sampling from a Gaussian distribution with mean zero and std α .

However, the situation changes significantly for deeper ($L \geq 3$) networks, as empirically demonstrated in Figure 1. Consider the task of completing the 2×2 matrix

$$M_D = \begin{pmatrix} w_{11}^* & ? \\ ? & w_{22}^* \end{pmatrix} \quad (1)$$

where only the diagonal entries are observed. This observation pattern forms a disconnected graph as illustrated in Figure 1(a). Consistent with the theory for disconnected graphs, $L = 2$ models fail to find a low-rank solution, empirically converging to rank-2 regardless of initialization scale. In contrast, deeper models ($L \geq 3$) with small initialization tend to converge to a rank-1 solution, as shown in Figure 1(b). This specific example highlights that the implicit low-rank bias appears to be strengthened by depth, in a way that *cannot be explained solely by the data connectivity* developed for $L = 2$ models. Furthermore, considering connected cases as well, Figure 1(c) demonstrates that this strong low-rank bias is generally robust, tending to strengthen further as depth increases.

However, a theoretical understanding of this depth-induced bias remains elusive, largely due to the complex, coupled dynamics during training. Indeed, Menon [24] notes that even for a simple case like (1) with $w_{11}^* = w_{22}^* = 1$, proving that gradient descent with a deep factorization converges to a low-rank solution is still an open problem. Motivated by this gap in understanding, we theoretically analyze such settings, including the example (1).

Investigating the implicit low-rank bias in matrix completion can also shed light on the phenomenon of “*loss of plasticity*”, a challenge widely observed in general neural network training [1, 4, 6, 29]. Discussions on the loss of plasticity phenomenon are provided in Appendix A. To summarize, here are the main research questions that we address throughout the paper:

- What is the fundamental difference between deep ($L \geq 3$) and shallow ($L = 2$) factorizations regarding their implicit low-rank bias, particularly for disconnected observations?
- Can we theoretically establish that deeper models (i.e., with larger $L \geq 3$) exhibit a stronger implicit bias toward low-rank solutions?
- What is the underlying cause of the loss of plasticity and how does depth interplay with it?

2. Problem Setting

We consider the problem of estimating a ground truth matrix $\mathbf{W}^* \in \mathbb{R}^{d \times d}$ based on observations of its entries $\{w_{ij}^*\}_{(i,j) \in \Omega}$, where $\Omega \subseteq [d] \times [d]$ is the set of observed indices. We model the estimate as a linear network $\mathbf{W}_{L:1} \triangleq \mathbf{W}_L \mathbf{W}_{L-1} \cdots \mathbf{W}_1$, where $\mathbf{W}_l \in \mathbb{R}^{d_l \times d_{l-1}}$ with $d_0 = d_L = d$. We denote the (i, j) -th entry of the matrix $\mathbf{W}_{L:1}$ as w_{ij} . The factor matrices $\{\mathbf{W}_l\}_{l=1}^L$ are trained by minimizing an objective function ϕ , defined as the squared loss ℓ over the observed entries in Ω :

$$\phi(\mathbf{W}_1, \dots, \mathbf{W}_L; \Omega) \triangleq \ell(\mathbf{W}_{L:1}; \Omega) = \frac{1}{2} \sum_{(i,j) \in \Omega} (w_{ij} - w_{ij}^*)^2. \quad (2)$$

We study the overparameterized regime where the intermediate dimensions satisfy $d_l \geq d$ for all $l \in [L-1]$, imposing no explicit rank constraints on the product model $\mathbf{W}_{L:1}$. Consistent with prior works, our analysis focuses on *gradient flow* dynamics (gradient descent with an infinitesimal step size) for a given objective function ϕ . The dynamics for each layer $\mathbf{W}_l(t)$ evolve according to:

$$\dot{\mathbf{W}}_l(t) \triangleq \frac{d}{dt} \mathbf{W}_l(t) = -\frac{\partial}{\partial \mathbf{W}_l(t)} \phi(\mathbf{W}_1(t), \mathbf{W}_2(t), \dots, \mathbf{W}_L(t); \Omega), \quad l \in [L], \quad t \geq 0. \quad (3)$$

For depth-2 networks ($L = 2$), the product of factor matrices $\mathbf{A} \in \mathbb{R}^{d \times d_1}$ (representing $\mathbf{W}_2$) and $\mathbf{B} \in \mathbb{R}^{d_1 \times d}$ (representing $\mathbf{W}_1$), we denote $\mathbf{W}_{A,B} \triangleq \mathbf{A}\mathbf{B}$.

3. Implicit Bias of Depth via Coupled Training Dynamics

This section extends Bai et al. [5]'s connectivity argument to general depth factorizations. First, using 2×2 matrix examples, we demonstrate how *coupled training dynamics* explain the role of observation connectivity in depth-2 models. Based on these insights, we hypothesize that deep networks exhibit an intrinsic low-rank bias by maintaining highly coupled training dynamics, regardless of observation patterns. This hypothesis is corroborated by the diagonal observation results.

3.1. Warm-up: Coupled Dynamics vs. Decoupled Dynamics in Depth-2 Networks

We focus on the simple 2×2 matrix completion of $\mathbf{M}_D$ and $\mathbf{M}_C$, using depth-2 models $\mathbf{W}_{A,B}(t) = \mathbf{A}(t)\mathbf{B}(t)$. For brevity, let $\mathbf{a}_i(t) \in \mathbb{R}^{d_1}$ be the transpose of the i -th row of $\mathbf{A}(t)$, and let $\mathbf{b}_j(t) \in \mathbb{R}^{d_1}$ be the j -th column of $\mathbf{B}(t)$. Our aim is to see how training dynamics affect the *alignment* of the rows of $\mathbf{A}(t)$ or the columns of $\mathbf{B}(t)$, as such alignment leads to a rank-1 product matrix $\mathbf{W}_{A,B}(t)$.

Decoupled Dynamics. In the $\mathbf{M}_D$ case (disconnected observations w_{11}^*, w_{22}^*), the gradient flow using the objective defined in (2), results in independent dynamics for the pairs $(\mathbf{a}_1, \mathbf{b}_1)$ and $(\mathbf{a}_2, \mathbf{b}_2)$:

$$\dot{\mathbf{a}}_i(t) = \left(w_{ii}^* - \mathbf{a}_i(t)^\top \mathbf{b}_i(t) \right) \mathbf{b}_i(t), \quad \dot{\mathbf{b}}_i(t) = \left(w_{ii}^* - \mathbf{a}_i(t)^\top \mathbf{b}_i(t) \right) \mathbf{a}_i(t) \quad \text{for } i = 1, 2.$$

Note that while the dynamics couple $\mathbf{a}_1(t)$ with $\mathbf{b}_1(t)$ and $\mathbf{a}_2(t)$ with $\mathbf{b}_2(t)$ within each pair, the two pairs $(\mathbf{a}_1, \mathbf{b}_1)$ and $(\mathbf{a}_2, \mathbf{b}_2)$ are decoupled. This decoupling means the overall system's dynamics separate into two independent systems. Consequently, there is no good reason to align vectors from different pairs, typically leading to high-rank solutions with generic initializations (Figure 1(b)).

Coupled Dynamics. In contrast, for the M_C case (connected observations w_{11}^*, w_{21}^*), the gradient flow on the objective (2) yields coupled dynamics that do not decompose into independent pairs:

$$\begin{aligned}\dot{\mathbf{a}}_1(t) &= \left(w_{11}^* - \mathbf{a}_1(t)^\top \mathbf{b}_1(t)\right) \mathbf{b}_1(t), & \dot{\mathbf{a}}_2(t) &= \left(w_{21}^* - \mathbf{a}_2(t)^\top \mathbf{b}_1(t)\right) \mathbf{b}_1(t), \\ \dot{\mathbf{b}}_1(t) &= \left(w_{11}^* - \mathbf{a}_1(t)^\top \mathbf{b}_1(t)\right) \mathbf{a}_1(t) + \left(w_{21}^* - \mathbf{a}_2(t)^\top \mathbf{b}_1(t)\right) \mathbf{a}_2(t).\end{aligned}\tag{4}$$

The following theorem demonstrates that sufficiently small initial norms in $\mathbf{A}(0)$ also result in the alignment of $\mathbf{a}_1(t)$ and $\mathbf{a}_2(t)$ with $\mathbf{b}_1(t)$.

Theorem 1 *For the product model $\mathbf{W}_{\mathbf{A},\mathbf{B}}(t) = \mathbf{A}(t)\mathbf{B}(t) \in \mathbb{R}^{2 \times 2}$, we consider the gradient flow dynamics (4), where the observations are $w_{11}^* (\neq 0)$ and $w_{21}^* (\neq 0)$. We assume convergence to the zero-loss solution (i.e., $w_{11}(\infty) = w_{11}^*, w_{21}(\infty) = w_{21}^*$). Defining $\mathbf{u}^* = \frac{\mathbf{b}_1(\infty)}{\|\mathbf{b}_1(\infty)\|_2}$ and the orthogonal component $\mathbf{a}_{i\perp}(\infty) = \mathbf{a}_i(\infty) - (\mathbf{a}_i(\infty)^\top \mathbf{u}^*) \mathbf{u}^*$, we have:*

$$\frac{\|\mathbf{a}_{i\perp}(\infty)\|_2^2}{\|\mathbf{a}_i(\infty)\|_2^2} \leq \frac{\|\mathbf{A}(0)\|_F^2 \left(\sqrt{\|\mathbf{b}_1(0)\|_2^4 + 4w_{11}^{*2} + 4w_{21}^{*2}} + \|\mathbf{b}_1(0)\|_2^2 \right)}{2w_{i1}^{*2}}, \text{ for } i = 1, 2.$$

The preceding theorem shows that small initial norms for $\mathbf{A}(0)$ lead to the alignment of $\mathbf{a}_1(\infty)$ and $\mathbf{a}_2(\infty)$ with $\mathbf{b}_1(\infty)$, implying a near rank-1 product matrix $\mathbf{W}_{\mathbf{A},\mathbf{B}}(\infty)$. This suggests that for depth-2 networks, coupled training dynamics (resulting from connected observations) facilitate the emergence of low-rank solutions under such small initialization, in contrast to the decoupled dynamics of disconnected observations, where no such bias exists, regardless of initialization scale.

3.2. Coupled Dynamics in Deep Networks Induce Implicit Bias Towards Low-rank

Section 3.1 illustrated the importance of coupled training dynamics, driven by data connectivity, for achieving low-rank solutions in simple two-layer factorizations ($L = 2$). Building on this understanding, we now extend our analysis to deep networks ($L \geq 3$). For illustrative purposes, consider a depth-3 network $\mathbf{W}_{3:1}$. An arbitrary observed entry w_{ij} from this matrix is given by:

$$w_{ij} = \sum_{k=1}^{d_2} \sum_{l=1}^{d_1} (\mathbf{W}_3)_{ik} (\mathbf{W}_2)_{kl} (\mathbf{W}_1)_{lj}.\tag{5}$$

Crucially, because all elements of the intermediate matrix $\mathbf{W}_2$ contribute to the computation of w_{ij} , independent of (i, j) , gradients of different observed entries will propagate through and update these shared elements in $\mathbf{W}_2$. This inherently couples their training dynamics, a structural feature distinct from the depth-2 case, where coupling is primarily determined by the observation pattern. Such inherent coupling, in turn, implies a potential intrinsic bias towards low-rank solutions for deep models. To formalize this notion, we introduce the following definition of coupled dynamics.

Definition 2 (Coupled/Decoupled Dynamics) *Consider the matrix completion setup with the model $\mathbf{W}_{L:1}(t) = \mathbf{W}_L(t) \cdots \mathbf{W}_1(t) \in \mathbb{R}^{d \times d}$. Let $\boldsymbol{\theta}(t)$ be the vector of all trainable parameters evolving according to the gradient flow dynamics (defined in (3)). Let $w_{ij}(t) \triangleq (\mathbf{W}_{L:1}(t))_{ij}$ be the model prediction for an observed index pair $(i, j) \in \Omega \subseteq [d] \times [d]$. The gradient flow dynamics are **decoupled** if there exists a partition of Ω into non-empty, disjoint subsets $\Omega_1, \dots, \Omega_K$ ($K \geq 2$) such that $\bigcup_{k=1}^K \Omega_k = \Omega$ and the following condition holds for any $(i, j) \in \Omega_k$ and $(x, y) \in \Omega_l$ with $k \neq l$:*

$$\langle \nabla_{\boldsymbol{\theta}(t)} w_{ij}(t), \nabla_{\boldsymbol{\theta}(t)} w_{xy}(t) \rangle = 0, \quad \forall t \geq 0.\tag{6}$$

The gradient flow dynamics are **coupled** if they are not decoupled.

For depth-2 matrices, it is straightforward to verify (based on Definition 8 and 2) that the training dynamics are coupled if and only if the observation graph is connected. Furthermore, Definition 2 indicates that for deep networks ($L \geq 3$), under random Gaussian initialization, the training dynamics are coupled with probability 1, irrespective of the observation pattern.

To gain insight into *how coupled dynamics induce low-rank bias as depth increases*, we further investigate the diagonal observation setting. As highlighted in the 2×2 example (cf. Figure 1(b)), this setting reveals a stark difference between shallow and deep networks despite being a disconnected observation pattern. To investigate this further, we now turn to the general $d \times d$ case.

Specifically, we consider a $d \times d$ ground truth matrix $\mathbf{W}^*$ with positive and identical diagonal observations $w_{ii}^* = w^* > 0$ for $\Omega_{\text{diag}}^{(d)} \triangleq \{(i, i) \mid i \in [d]\}$. We factorize the model with depth- L : $\mathbf{W}_{L:1}(t) = \mathbf{W}_L(t)\mathbf{W}_{L-1}(t) \cdots \mathbf{W}_1(t)$ where $\mathbf{W}_l \in \mathbb{R}^{d \times d}$ for all $l \in [L]$.

To study how dynamic coupling affects the low-rank bias, we consider a family of initializations where, for parameters $\alpha > 0$ and $m > 1$, each factor matrix $\mathbf{W}_l(0)$ is initialized as follows:

$$\mathbf{W}_l(0) = \begin{pmatrix} \alpha & \alpha/m & \cdots & \alpha/m \\ \alpha/m & \alpha & \cdots & \alpha/m \\ \vdots & \vdots & \ddots & \vdots \\ \alpha/m & \alpha/m & \cdots & \alpha \end{pmatrix} \in \mathbb{R}^{d \times d}, \quad \forall l \in [L]. \quad (7)$$

Using this initialization scheme with diagonal observations, the following proposition specifies how parameters m and network depth L determine if training dynamics are coupled or decoupled:

Proposition 3 *Consider a depth- L model, where each factor $\mathbf{W}_l(0) \in \mathbb{R}^{d \times d}$ is initialized with (7) trained with diagonal observations, $\Omega_{\text{diag}}^{(d)}$. Then, according to Definition 2, the following hold:*

- For depth $L = 2$, the training dynamics are **decoupled** for all $m > 1$.
- For depth $L \geq 3$:
 - The training dynamics are **coupled** if $1 < m < \infty$.
 - The training dynamics are **decoupled** if $m = \infty$ (i.e., initialization with $\alpha \mathbf{I}_d$).

Assuming convergence to a zero-loss solution, our objective is to determine the rank of solutions found by gradient flow depending on the coupling of dynamics. The theorem below presents an equation of each singular value of the converged matrix $\mathbf{W}_{L:1}(\infty)$, for all $L \geq 2$.

Theorem 4 *Consider the product matrix $\mathbf{W}_{L:1}$, whose factor matrices $\mathbf{W}_l \in \mathbb{R}^{d \times d}$ are initialized according to (7). Assuming convergence to a zero-loss solution under the diagonal observation $\Omega_{\text{diag}}^{(d)}$, let $(\sigma_1)^L \geq \cdots \geq (\sigma_d)^L \geq 0$ denote the singular values of the converged matrix $\mathbf{W}_{L:1}(\infty)$. Then, for all parameter values $\alpha > 0$, $m > 1$, $d \geq 2$, and $L \geq 2$, the following holds:*

- If $L = 2$ (**decoupled dynamics**): The singular values are explicitly given by

$$\sigma_1 = (m + d - 1) \sqrt{\frac{w^*}{m^2 + d - 1}}, \quad \sigma_r = (m - 1) \sqrt{\frac{w^*}{m^2 + d - 1}} \quad \text{for } r = 2, \dots, d.$$

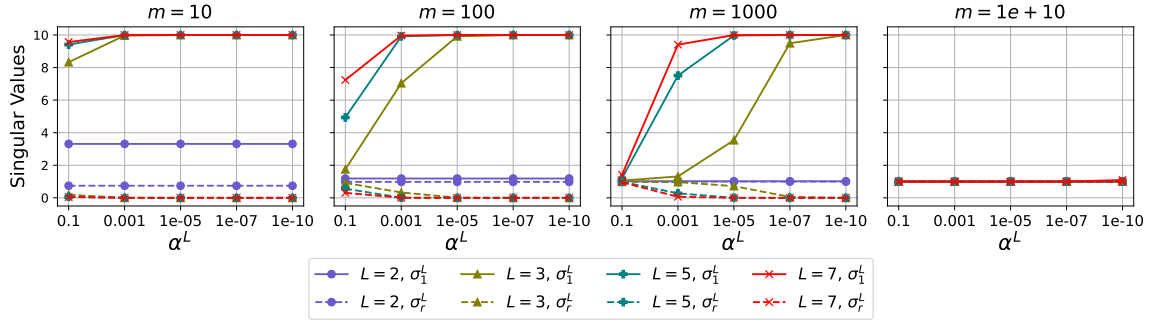


Figure 2: Singular values $(\sigma_i)^L$ (numerically obtained from Theorem 4) against initialization scale α^L , for the diagonal observation task. Solid lines represent the largest singular value $(\sigma_1)^L$; dashed lines denote the other (identical) singular values $(\sigma_r)^L$ for $r \geq 2$. For finite m , these results illustrate that both greater depth L and a smaller initial scale α enhance the low-rank bias, in contrast to the $L = 2$ case. Conversely, a very large m (e.g., $m = 10^{10}$), approximating an $\alpha \mathbf{I}_d$ (rank- d) initialization, leads to decoupled dynamics and a full-rank solution, independent of both L and α .

- If $L \geq 3$ and $1 < m < \infty$ (**coupled dynamics**): The singular values satisfy the implicit equations:

$$(\sigma_1)^{2-L} - \left(\frac{w^* d - (\sigma_1)^L}{d-1} \right)^{\frac{2-L}{L}} = C_{\alpha, m, L, d}, \quad (8)$$

$$\left(w^* d - (d-1)(\sigma_r)^L \right)^{\frac{2-L}{L}} - (\sigma_r)^{2-L} = C_{\alpha, m, L, d}, \quad \text{for } r = 2, \dots, d, \quad (9)$$

where $C_{\alpha, m, L, d} \triangleq \left(\frac{\alpha}{m} \right)^{2-L} ((m+d-1)^{2-L} - (m-1)^{2-L})$.

- If $L \geq 3$ and $m = \infty$ (**decoupled dynamics**): The singular values converge to:

$$(\sigma_i)^L = w^*, \quad \text{for } i = 1, 2, \dots, d.$$

The preceding theorem details converged singular values of $\mathbf{W}_{L:1}(\infty)$ for our initialization scheme (7), with outcomes dependent on the nature of the training dynamics. For decoupled dynamics—specifically, when $L = 2$ (for sufficiently large $m > 1$), or when $L \geq 3$ and $m = \infty$ —all singular values approach w^* independently of the scale α , implying convergence to a full-rank solution. In contrast, for coupled dynamics ($L \geq 3$ with finite m), the outcome is α -dependent. The analytical intractability of this coupled regime motivates a numerical study.

To numerically investigate this, we solve the implicit equations (8) and (9) that determine singular values $(\sigma_i)^L$ for the coupled $L \geq 3$, finite m case. Setting $w^* = 1$ and $d = 10$, we examine how network depth (L) and initialization parameters (α, m) influence the singular value distribution. The results (Figure 2) confirm that these coupled dynamics in models with $L \geq 3$ and finite m indeed induce a low-rank bias, contrasting with the full-rank outcomes of the $L = 2$ (decoupled) case. Moreover, this bias becomes more pronounced as network depth increases, evidenced by a wider gap between $(\sigma_1)^L$ and $(\sigma_r)^L$ for $r \geq 2$.

Remark. Our analysis of low-rank bias for a specific family of deterministic initializations resolves the challenging open problem (1) highlighted by Menon [24]. Experiments in Appendix E further demonstrate that our proposed deterministic initialization exhibits qualitative trends similar to Gaussian initialization. We therefore argue that our results provide foundational insights into low-rank bias applicable to more general random initializations.

Acknowledgements

This work was supported by a National Research Foundation of Korea (NRF) grant (No. RS-2024-00421203) funded by the Korean government (MSIT), and an Institute for Information & communications Technology Planning & Evaluation (IITP) grant (No. RS-2022-II220184, Development and Study of AI Technologies to Inexpensively Conform to Evolving Policy on Ethics) funded by the Korean government (MSIT).

References

- [1] Alessandro Achille, Matteo Rovere, and Stefano Soatto. Critical learning periods in deep networks. In *International Conference on Learning Representations*, 2018.
- [2] Maksym Andriushchenko, Dara Bahri, Hossein Mobahi, and Nicolas Flammarion. Sharpness-aware minimization leads to low-rank features. *Advances in Neural Information Processing Systems*, 36:47032–47051, 2023.
- [3] Sanjeev Arora, Nadav Cohen, Wei Hu, and Yuping Luo. Implicit regularization in deep matrix factorization. *Advances in Neural Information Processing Systems*, 32, 2019.
- [4] Jordan Ash and Ryan P Adams. On warm-starting neural network training. *Advances in neural information processing systems*, 33:3884–3894, 2020.
- [5] Zhiwei Bai, Jiajie Zhao, and Yaoyu Zhang. Connectivity shapes implicit regularization in matrix factorization models for matrix completion. In *The Thirty-eighth Annual Conference on Neural Information Processing Systems*, 2024.
- [6] Tudor Berariu, Wojciech Czarnecki, Soham De, Jorg Bornschein, Samuel Smith, Razvan Pascanu, and Claudia Clopath. A study on the plasticity of neural networks. *arXiv preprint arXiv:2106.00042*, 2021.
- [7] Yihong Chen, Kelly Marchisio, Roberta Raileanu, David Adelani, Pontus Lars Erik Saito Stenertorp, Sebastian Riedel, and Mikel Artetxe. Improving language plasticity via pretraining with active forgetting. *Advances in Neural Information Processing Systems*, 36:31543–31557, 2023.
- [8] Lenaic Chizat, Edouard Oyallon, and Francis Bach. On lazy training in differentiable programming. *Advances in neural information processing systems*, 32, 2019.
- [9] Shibhansh Dohare, Richard S Sutton, and A Rupam Mahmood. Continual backprop: Stochastic gradient descent with persistent randomness. *arXiv preprint arXiv:2108.06325*, 2021.
- [10] Spencer Frei, Gal Vardi, Peter Bartlett, Nathan Srebro, and Wei Hu. Implicit bias in leaky reLU networks trained on high-dimensional data. In *The Eleventh International Conference on Learning Representations*, 2023. URL <https://openreview.net/forum?id=JpbLyEI5EwW>.
- [11] Daniel Gissin, Shai Shalev-Shwartz, and Amit Daniely. The implicit bias of depth: How incremental learning drives generalization. In *International Conference on Learning Representations*, 2020. URL <https://openreview.net/forum?id=H1lj0nNFwB>.

- [12] Suriya Gunasekar, Blake E Woodworth, Srinadh Bhojanapalli, Behnam Neyshabur, and Nati Srebro. Implicit regularization in matrix factorization. *Advances in neural information processing systems*, 30, 2017.
- [13] Minyoung Huh, Hossein Mobahi, Richard Zhang, Brian Cheung, Pulkit Agrawal, and Phillip Isola. The low-rank simplicity bias in deep networks. *arXiv preprint arXiv:2103.10427*, 2021.
- [14] Maximilian Igl, Gregory Farquhar, Jelena Luketina, Wendelin Boehmer, and Shimon Whiteson. Transient non-stationarity and generalisation in deep reinforcement learning. *arXiv preprint arXiv:2006.05826*, 2020.
- [15] Arthur Jacot. Implicit bias of large depth networks: a notion of rank for nonlinear functions. *arXiv preprint arXiv:2209.15055*, 2022.
- [16] Ziwei Ji and Matus Telgarsky. Gradient descent aligns the layers of deep linear networks. In *International Conference on Learning Representations*, 2019. URL <https://openreview.net/forum?id=HJflg30qKX>.
- [17] Ziwei Ji and Matus Telgarsky. The implicit bias of gradient descent on nonseparable data. In *Conference on learning theory*, pages 1772–1798. PMLR, 2019.
- [18] Michael Kleinman, Alessandro Achille, and Stefano Soatto. Critical learning periods emerge even in deep linear networks. In *The Twelfth International Conference on Learning Representations*, 2024. URL <https://openreview.net/forum?id=Aq35gl2c1k>.
- [19] Yiwen Kou, Zixiang Chen, and Quanquan Gu. Implicit bias of gradient descent for two-layer relu and leaky relu networks on nearly-orthogonal data. *Advances in Neural Information Processing Systems*, 36:30167–30221, 2023.
- [20] Saurabh Kumar, Henrik Marklund, and Benjamin Van Roy. Maintaining plasticity in continual learning via regenerative regularization. *arXiv preprint arXiv:2308.11958*, 2023.
- [21] Zhiyuan Li, Yuping Luo, and Kaifeng Lyu. Towards resolving the implicit bias of gradient descent for matrix factorization: Greedy low-rank learning. In *International Conference on Learning Representations*, 2021. URL <https://openreview.net/forum?id=AHOs7Sm5H7R>.
- [22] Clare Lyle, Zeyu Zheng, Evgenii Nikishin, Bernardo Avila Pires, Razvan Pascanu, and Will Dabney. Understanding plasticity in neural networks. In *International Conference on Machine Learning*, pages 23190–23211. PMLR, 2023.
- [23] Jianhao Ma and Salar Fattahi. Convergence of gradient descent with small initialization for unregularized matrix completion. In *The Thirty Seventh Annual Conference on Learning Theory*, pages 3683–3742. PMLR, 2024.
- [24] Govind Menon. The geometry of the deep linear network. *arXiv preprint arXiv:2411.09004*, 2024.
- [25] Evgenii Nikishin, Max Schwarzer, Pierluca D’Oro, Pierre-Luc Bacon, and Aaron Courville. The primacy bias in deep reinforcement learning. In *International conference on machine learning*, pages 16828–16847. PMLR, 2022.

- [26] Sangyeon Park, Isaac Han, Seungwon Oh, and Kyung-Joong Kim. Activation by interval-wise dropout: A simple way to prevent neural networks from plasticity loss. *arXiv preprint arXiv:2502.01342*, 2025.
- [27] Noam Razin and Nadav Cohen. Implicit regularization in deep learning may not be explainable by norms. *Advances in neural information processing systems*, 33:21174–21187, 2020.
- [28] Olivier Roy and Martin Vetterli. The effective rank: A measure of effective dimensionality. In *2007 15th European signal processing conference*, pages 606–610. IEEE, 2007.
- [29] Baekrok Shin, Junsoo Oh, Hanseul Cho, and Chulhee Yun. Dash: Warm-starting neural network training in stationary settings without loss of plasticity. *Advances in Neural Information Processing Systems*, 37:43300–43340, 2024.
- [30] Daniel Soudry, Elad Hoffer, Mor Shpigel Nacson, Suriya Gunasekar, and Nathan Srebro. The implicit bias of gradient descent on separable data. *Journal of Machine Learning Research*, 19(70):1–57, 2018.
- [31] Dominik Stöger and Mahdi Soltanolkotabi. Small random initialization is akin to spectral learning: Optimization and generalization guarantees for overparameterized low-rank matrix reconstruction. *Advances in Neural Information Processing Systems*, 34:23831–23843, 2021.
- [32] Matus Telgarsky. Deep learning theory lecture notes. *Lecture Notes v0. 0-e7150f2d (alpha)*, Univ. Illinois Urbana-Champaign, Champaign, IL, USA, 2021.
- [33] Nadav Timor, Gal Vardi, and Ohad Shamir. Implicit regularization towards rank minimization in relu networks. In *International Conference on Algorithmic Learning Theory*, pages 1429–1459. PMLR, 2023.
- [34] Blake Woodworth, Suriya Gunasekar, Jason D Lee, Edward Moroshko, Pedro Savarese, Itay Golan, Daniel Soudry, and Nathan Srebro. Kernel and rich regimes in overparametrized models. In *Conference on Learning Theory*, pages 3635–3673. PMLR, 2020.
- [35] Chulhee Yun, Shankar Krishnan, and Hossein Mobahi. A unifying view on implicit bias in training linear neural networks. In *International Conference on Learning Representations*, 2021. URL <https://openreview.net/forum?id=ZsZM-4iMQkH>.
- [36] Chenyang Zhang, Difan Zou, and Yuan Cao. The implicit bias of adam on separable data. In *The Thirty-eighth Annual Conference on Neural Information Processing Systems*, 2024. URL <https://openreview.net/forum?id=xRQxan3WkM>.

Contents

1	Introduction	1
2	Problem Setting	3
3	Implicit Bias of Depth via Coupled Training Dynamics	3
3.1	Warm-up: Coupled Dynamics vs. Decoupled Dynamics in Depth-2 Networks . . .	3
3.2	Coupled Dynamics in Deep Networks Induce Implicit Bias Towards Low-rank . . .	4
A	Understanding The Loss of Plasticity via Depth-2 Matrix Completion	11
A.1	Pre-training with Diagonal Observations	12
A.2	Post-training: 2 by 2 Matrix Example	12
A.3	Post-training: d by d Matrix under Lazy Training Regime	13
B	Conclusion	14
C	Further Related Works	15
C.1	Implicit Regularization in Neural Networks	15
C.2	Loss of Plasticity	15
D	Coupled and Decoupled Training Dynamics	17
D.1	Coupled Dynamics	17
D.2	Decoupled Dynamics	18
E	Additional Experiments	20
E.1	Implicit Bias Experiments	20
E.2	Loss of Plasticity Experiments	21
F	Proof for Section 3	25
F.1	Proof for Theorem 1	25
F.2	Proof for Proposition 3	28
F.3	Proof for Theorem 4	30
G	Proof for Section A	39
G.1	General Form and Proof of Proposition 5	39
G.2	Proof of Theorem 6	44
G.3	Formal Statement and Proof of Theorem 7	57
H	Useful Lemmas	64

Appendix A. Understanding The Loss of Plasticity via Depth-2 Matrix Completion

Investigating the implicit low-rank bias in matrix completion can also shed light on the phenomenon of “*loss of plasticity*”, a challenge widely observed in general neural network training [1, 4, 6, 29]. The term loss of plasticity describes the tendency of neural networks, particularly after initial training, to lose their adaptability to new information, hindering their generalization capabilities. A recent work by Kleinman et al. [18] empirically report the emergence of this phenomenon in matrix completion: models pre-trained on limited observations struggle to adapt when training continues on augmented observations. Notably, they observe that loss of plasticity is further intensified with increasing network depth, a conclusion they reached by measuring a “relative reconstruction loss” when compared to models trained from scratch on the augmented dataset.

However, our findings (Figure 3) offer a more nuanced perspective. We observed that even when pre-trained with a sparser set of observations, deeper models increasingly favor low-rank solutions as their depth increases. This aligns with our argument (Section 3.2) that they inherently achieve low-rank solutions even from limited, disconnected initial data. Consequently, for these deeper models, further training on augmented data (the post-training stage) does not lead to noticeably higher rank compared to training equivalent models from scratch on the augmented observations. Therefore, while their performance might exhibit a relative degradation compared to models trained from scratch, their absolute solution quality can still surpass that of shallower models. Based on our observations, we conclude that the low-rank bias of deep models helps them avoid the loss of plasticity, while the loss is more pronounced in depth-2 models. To theoretically understand the underlying cause of this phenomenon itself, we henceforth focus our analysis on depth-2 models.

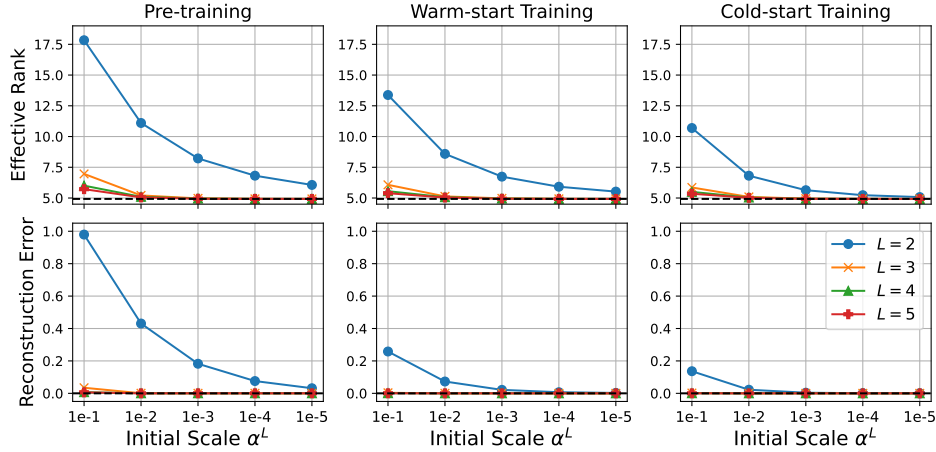


Figure 3: Experiments use a 100×100 rank-5 ground-truth matrix. pre-training utilizes 2000 randomly sampled entries (Ω_{pre} ; $|\Omega_{\text{pre}}| = 2000$), while post-training adds 1000 more, forming Ω_{post} ($\Omega_{\text{pre}} \subset \Omega_{\text{post}}$; $|\Omega_{\text{post}}| = 3000$). The top row of panels displays effective rank, and the bottom row shows reconstruction error, both measured at convergence. The leftmost panels depict training on Ω_{pre} , and the rightmost on Ω_{post} , both starting from random Gaussian initialization. The middle panels show warm-start training on Ω_{post} , initialized from converged pre-trained models with Ω_{pre} .

In Section A.1, models are pre-trained using only diagonal observations, with the set of diagonal indices $\Omega_{\text{diag}}^{(d)}$. We then examine 2×2 (Section A.2) and $d \times d$ (Section A.3) cases. For the 2×2

case, the pre-train observation set is $\Omega_{\text{pre}}^{(2)} \triangleq \Omega_{\text{diag}}^{(2)}$. The post-train set, $\Omega_{\text{post}}^{(2)}$, then incorporates an additional off-diagonal observation to ensure connectivity. Similarly, for the $d \times d$ case, the pre-train set is $\Omega_{\text{pre}}^{(d)} \triangleq \Omega_{\text{diag}}^{(d)}$. Its post-train set, $\Omega_{\text{post}}^{(d)}$, includes additional observations (see Section A.3).

A.1. Pre-training with Diagonal Observations

To clearly observe loss of plasticity in a setting consistent with Section 3.2, we pre-train using only diagonal entries, yielding a disconnected pattern. We consider decoupled-to-coupled scenarios, where additional data is introduced to induce coupled training dynamics. For depth-2 models, they correspond to a disconnected-to-connected observation pattern. For the pre-training, closed-form solutions that depend *solely* on the network’s initialization can be found in the following proposition:

Proposition 5 *Consider a ground truth matrix $\mathbf{W}^* \in \mathbb{R}^{d \times d}$ with diagonal observations $\Omega_{\text{diag}}^{(d)}$. The model is factorized as $\mathbf{W}_{\mathbf{A},\mathbf{B}}(t) = \mathbf{A}(t)\mathbf{B}(t)$, where $\mathbf{A}(t), \mathbf{B}(t) \in \mathbb{R}^{d \times d}$. For each observation $(i, i) \in \Omega_{\text{diag}}^{(d)}$, define the constants P_i and Q_i based on the initial values:*

$$P_i \triangleq \sum_{k=1}^d a_{i,k}(0)b_{k,i}(0) \quad \text{and} \quad Q_i \triangleq \sum_{k=1}^d (a_{i,k}(0)^2 + b_{k,i}(0)^2).$$

Furthermore, for each diagonal observation, let the parameter $\bar{r}_i$ be determined from the ground truth entry $w_{i,i}^*$ and the constants defined above, $\bar{r}_i \triangleq \frac{1}{2} \log \left(\frac{P_i + \frac{Q_i}{2}}{w_{i,i}^* + \sqrt{w_{i,i}^{*2} - P_i^2 + \left(\frac{Q_i}{2}\right)^2}} \right)$. Then, assuming convergence to a zero-loss solution of the loss $\ell(\mathbf{W}_{\mathbf{A},\mathbf{B}}; \Omega_{\text{diag}}^{(d)})$, any entry $a_{p,q}(\infty)$ of the converged matrix $\mathbf{A}(\infty)$ and any entry $b_{p,q}(\infty)$ of the converged matrix $\mathbf{B}(\infty)$ (for any $p, q \in [d]$) are given by:

$$\begin{aligned} a_{p,q}(\infty) &= a_{p,q}(0) \cosh(\bar{r}_p) - b_{q,p}(0) \sinh(\bar{r}_p), \\ b_{p,q}(\infty) &= b_{p,q}(0) \cosh(\bar{r}_q) - a_{q,p}(0) \sinh(\bar{r}_q). \end{aligned}$$

Remark. The above proposition applies to *arbitrary* initializations with *distinct* w_{ii}^* values, which is goes beyond Theorem 4. While the above analysis focuses on diagonal observation cases, it can be generalized to any fully disconnected case (i.e., a single observation per row and column). This yields distinct solutions for various types of observation sets, as detailed in Appendix G.1. We analyze the scenario where training resumes from a state obtained through pre-training. Let the pre-training phase conclude at a sufficiently large timestep T_1 . For simplicity, we assume that the solution $\mathbf{W}_{\mathbf{A},\mathbf{B}}(T_1)$ has perfectly converged with respect to the pre-training objective, neglecting any residual error due to the finite duration of this phase. Our subsequent analysis demonstrates that, starting from $\mathbf{W}_{\mathbf{A},\mathbf{B}}(T_1)$, the model $\mathbf{W}_{\mathbf{A},\mathbf{B}}(t)$ cannot converge to a low-rank solution.

A.2. Post-training: 2 by 2 Matrix Example

We aim to analyze scenarios where training is resumed under coupled dynamics, building upon solutions obtained from an initial decoupled pre-training phase (Proposition 5). To this end, we first define the specific pre-training setup for an illustrative 2×2 case: We observe diagonal entries ($\Omega_{\text{pre}}^{(2)}$), which are identical and positive, i.e., $w^* \triangleq w_{11}^* = w_{22}^* > 0$. To make loss of plasticity

particularly pronounced during the pre-training, we initialize the model with $\alpha \mathbf{I}_2$ (for $\alpha > 0$), which is the $m \rightarrow \infty$ limit of our initialization scheme in (7). Then, from Proposition 5, it follows that:

$$\mathbf{A}(T_1) = \mathbf{B}(T_1) = \begin{pmatrix} \sqrt{w^*} & 0 \\ 0 & \sqrt{w^*} \end{pmatrix}. \quad (10)$$

For the subsequent post-training phase, an additional off-diagonal observation is introduced to establish connectivity. Without loss of generality, we assume $w_{12}^* > 0$ is revealed, while the diagonal entries w_{11}^* and w_{22}^* from the pre-training phase remain observed. Thus, the updated set of observed entries becomes $\Omega_{\text{post}}^{(2)} = \{(1, 1), (1, 2), (2, 2)\}$. The ground-truth matrix is assumed to be rank-1, ensuring the setting is non-trivial, and the task is thus to predict the remaining entry $w_{21}^* = w^{*2}/w_{12}^* > 0$. The following theorem, however, reveals a contrasting outcome for this entry.

Theorem 6 *Let $\mathbf{A}(T_1), \mathbf{B}(T_1)$ be the factor matrices obtained from the pre-training phase, as specified by (10). Then, running gradient flow during the subsequent post-training phase (for $t \geq T_1$), starting from $\mathbf{A}(T_1)$ and $\mathbf{B}(T_1)$, results in exponential decay of the loss:*

$$\ell(\mathbf{W}_{\mathbf{A}, \mathbf{B}}(t); \Omega_{\text{post}}^{(2)}) \leq \frac{1}{2} w_{12}^{*2} e^{-2w^*(t-T_1)}.$$

Consequently, a lower bound for the stable rank of the converged matrix $\mathbf{W}_{\mathbf{A}, \mathbf{B}}(\infty)$ is given by:

$$\frac{\|\mathbf{W}_{\mathbf{A}, \mathbf{B}}(\infty)\|_F^2}{\|\mathbf{W}_{\mathbf{A}, \mathbf{B}}(\infty)\|_2^2} \geq 1 + \exp\left(-8 \frac{w_{12}^*}{w^*}\right).$$

Furthermore, for all $t > T_1$, $w_{21}(t)$ of the evolving matrix $\mathbf{W}_{\mathbf{A}, \mathbf{B}}(t)$ satisfies $w_{21}(t) < 0$.

The theorem indicates that the loss decreases exponentially fast, particularly when starting from high-norm solutions (at a rate governed by w^*). Therefore, since the model converged to high-rank solutions during pre-training, its singular values remain largely unchanged from this initial state, as long as w_{12}^* has a small magnitude compared to w^* . Furthermore, the unobserved entry $w_{21}(t)$ converges to a negative value, which contradicts the positive w_{21}^* expected for the true rank-1 solution.

A.3. Post-training: d by d Matrix under Lazy Training Regime

We attribute Theorem 6 primarily to the model’s “lazy training” [8] as high-norm initializations lead to faster loss decay, causing the model to converge to a nearby global minimum that may not be a low-rank solution. Drawing on this concept, we extend the preceding analysis of loss of plasticity to the more general case of $d \times d$ ground-truth matrices. The following theorem states that when the model is initialized with a sufficiently small loss, resulting from warm-starting that perfectly fits all previously observed data, the model exhibits lazy training. This, in turn, prevents further learning that would reduce the rank and instead steers the model towards a nearby minimum.

Theorem 7 (informal) *For factor matrices $\mathbf{A}, \mathbf{B} \in \mathbb{R}^{d \times d}$, suppose $\mathbf{A}$ and $\mathbf{B}$ are balanced at $t = 0$, i.e., $\mathbf{A}(0)^\top \mathbf{A}(0) = \mathbf{B}(0) \mathbf{B}(0)^\top$. Let $f(\mathbf{A}, \mathbf{B})$ be the function that maps $(\mathbf{A}, \mathbf{B})$ to the vector of model predictions for a given set of observed entries $\Omega_{\text{post}}^{(d)}$. We then define $\sigma_{\max}$ and $\sigma_{\min}$ as the*

maximum and minimum singular values, respectively, of the Jacobian of the function f evaluated at the pre-trained state. If the loss at time T_1 satisfies:

$$\ell(\mathbf{W}_{\mathbf{A},\mathbf{B}}(T_1); \Omega_{\text{post}}^{(d)}) \leq \frac{\sigma_{\min}^6}{1152d\sigma_{\max}^2},$$

this results in exponential decay of the loss:

$$\ell(\mathbf{W}_{\mathbf{A},\mathbf{B}}(t); \Omega_{\text{post}}^{(d)}) \leq \ell(\mathbf{W}_{\mathbf{A},\mathbf{B}}(T_1); \Omega_{\text{post}}^{(d)}) \exp\left(-\frac{1}{2}\sigma_{\min}^2 t\right).$$

Consequently, the stable rank of $\mathbf{A}(t)$ (which is equal to that of $\mathbf{B}(t)$) remains bounded below by

$$\frac{\|\mathbf{A}(t)\|_F^2}{\|\mathbf{A}(t)\|_2^2} \geq \left(\frac{\sqrt{2}\|\mathbf{A}(T_1)\|_F - \frac{\sigma_{\min}}{4\sqrt{2}d}}{\sqrt{2}\|\mathbf{A}(T_1)\|_2 + \frac{\sigma_{\min}}{4\sqrt{2}d}} \right)^2.$$

The above theorem states that if a model has little remaining to learn (achieved via warm-starting), it undergoes lazy training, leading to rapid loss convergence while its stable rank remains largely unchanged from the initial state. Thus, once a model has converged to a high-rank state, it struggles to recover a low-rank structure even when new observations are introduced to form connectivity. Formal statement of Theorem 7 is provided in Appendix G.

As an illustrative example, consider the simple warm-starting scenario from Section A.2: the model is first pre-trained using only diagonal observations of ground truth matrix $\mathbf{W}^*$, $w^* = w_{11}^* = w_{22}^*$, after which off-diagonal entry w_{12}^* is introduced for subsequent training.

Example. Consider a warm-starting scenario where $\mathbf{A}$ and $\mathbf{B}$ are initialized as (10). When observing $\Omega_{\text{post}}^{(2)} = \{(1, 1), (1, 2), (2, 2)\}$, the loss at time T_1 is $\ell(\mathbf{W}_{\mathbf{A},\mathbf{B}}(T_1); \Omega_{\text{post}}^{(2)}) = \frac{1}{2}w_{12}^{*2}$. The two singular values of the Jacobian matrix are both $\sqrt{2w^*}$. If we choose $w_{12}^* \leq \frac{w^*}{12\sqrt{2}}$, Theorem 7 ensures that at every time $t \geq T_1$, the stable rank of $\mathbf{A}(t)$ stays uniformly bounded below:

$$\|\mathbf{A}(t)\|_F^2 / \|\mathbf{A}(t)\|_2^2 > 1.31.$$

Appendix B. Conclusion

We demonstrate that in matrix completion, deeper networks ($L \geq 3$) inherently exhibit a stronger low-rank bias than shallow networks. This tendency is primarily attributed to their coupled training dynamics, which operate irrespective of observation patterns. To enable a tractable theoretical analysis of this phenomenon, we consider gradient flow starting at a family of deterministic initializations. We show in the diagonal observation setting that coupled training dynamics lead to a low-rank bias, and this bias is further intensified as network depth increases. Furthermore, our theoretical analysis of warm-starting scenarios details the loss of plasticity phenomenon, revealing how suboptimal initial states can hinder convergence to low-rank solutions. These findings contribute to a more precise understanding of the implicit bias of depth and loss of plasticity in matrix completion.

Appendix C. Further Related Works

C.1. Implicit Regularization in Neural Networks

Recent studies have shown that deep neural networks have implicit bias or regularization towards certain solutions among the many global minima [2, 3, 5, 10–13, 15–17, 19, 21, 30, 31, 33–36].

Among these, Arora et al. [3], Gissin et al. [11], Li et al. [21] study the implicit bias of depth towards low-rank solutions. Specifically, Gissin et al. [11] and Li et al. [21] examine this bias in deep linear models in relation to the initialization scale. They report that as model depth increases, the dependence on initialization can become weaker, and incremental learning can emerge. However, their analyses consider a matrix factorization task, which they frame as matrix completion with full observations. Therefore, in their setting, convergence to a low-rank solution is guaranteed if the model converges to zero-loss, which does not hold in our matrix completion task settings.

While Arora et al. [3] investigate matrix completion in deep linear networks, offering insights from derived singular value dynamics, they cannot fully track these dynamics to prove low-rank convergence as network depth increases. Their analysis is primarily restricted to the regime where $t \geq t_0$, after which singular vectors are assumed to have stabilized. For $t \geq t_0$, they find that one singular value can be expressed as a function of another, involving a *constant term* that emerges from the state at t_0 (which can be the dominant component). Based on this, they demonstrate that the gap between these singular values widens with increasing depth. In contrast, our Theorem 4, by precisely tracking the converged values of singular values, rigorously establishes their ultimate behavior and the resulting low-rank bias.

Bai et al. [5] introduce the connectivity argument in matrix completion tasks for depth-2 matrices. For an incomplete matrix M , connectivity is characterized by its set of observed indices $\Omega \subseteq [d] \times [d]$ and the corresponding observation matrix P (where $P_{ij} = 1$ if $(i, j) \in \Omega$, and 0 otherwise). The formal definition is as follows:

Definition 8 (Connectivity from Bai et al. [5]) *An incomplete matrix M is **connected** if the bipartite graph $\mathcal{G}_M$, constructed from its observation matrix P using the adjacency matrix $\begin{bmatrix} \mathbf{0} & P^\top \\ P & \mathbf{0} \end{bmatrix}$, is connected after removing isolated vertices. Otherwise, M is **disconnected**.*

They prove that if the observations construct a connected bipartite graph, the model can converge to a low-rank solution when the initialization scale is infinitesimally small, subject to certain technical assumptions. Conversely, if the observations form a disconnected graph, the model generally cannot converge to a low-rank solution. However, a special case occurs if this disconnected graph is composed of complete bipartite components: here, the model converges to the minimum nuclear norm solution, again under specific technical assumptions. This characterization of implicit bias does not readily generalize to matrices with deeper matrices, as depicted in Figure 1.

C.2. Loss of Plasticity

Loss of plasticity describes a widely observed phenomenon where a model’s ability to adapt to new information diminishes over time [1, 4, 9, 25, 29]. This is frequently observed in scenarios with gradually changing datasets, such as those encountered in reinforcement learning [14, 22, 25] or continual learning [7, 9, 20, 26], where the model may struggle to adapt to new environments.

Although loss of plasticity is more extensively studied in non-stationary environments, a similar phenomenon can be observed in stationary settings where a model processes a dataset that grows incrementally from a fixed data distribution [4, 6, 29]. In such stationary scenarios, a model is typically first trained to convergence on an initial, independently and identically distributed (i.i.d.) subset of data (e.g., from CIFAR-10 or CIFAR-100). Subsequently, this converged model serves as a initialization for continued training on an expanded dataset, which incorporates additional samples from the original distribution. Perhaps surprisingly, these warm-started models struggle to generalize to the newly introduced samples, often exhibiting lower test accuracy compared to models trained from scratch on the combined dataset.

While this phenomenon is problematic in many real-world applications where new data is continuously added, theoretical studies on it remain scarce. Shin et al. [29], for instance, offer a theoretical explanation using an artificial framework. Within this framework, they demonstrate that such behavior occurs because warm-started models often complete training by memorizing data-dependent noise, which is not useful for generalization. However, the analytical framework they employ is considered artificial and limited in its ability to accurately characterize the optimization processes of typical deep learning models.

Recently, Kleinman et al. [18] observed loss of plasticity in deep linear networks, identifying “critical learning periods”: an initial phase of effective learning followed by a significantly reduced capacity to learn later. They employ a matrix completion framework to further observe this behavior. When observations from matrix completion tasks are treated as training samples in neural network training, they observed that a model initially trained on a sparse set of observations and subsequently retrained (i.e., warm-started) on an expanded dataset typically exhibits a larger performance gap (in terms of reconstruction error) compared to a model trained from scratch on the entire expanded dataset. However, their work does not offer theoretical guarantees to account for these observations. Motivated by this, in Section A, we attempt to explain this behavior within the specific context of depth-2 matrix completion settings.

Appendix D. Coupled and Decoupled Training Dynamics

This section discusses coupled/decoupled training dynamics defined in Definition 2, illustrated with specific examples.

D.1. Coupled Dynamics

For shallow ($L = 2$) matrices, coupled dynamics typically correspond to connected observations under generic initialization, in accordance with Definitions 8 and 2 (the specific case of initialization such as zero matrices, which leads to decoupled dynamics, will be further detailed in a later subsection). We illustrate this principle with an example where the observed entries form the first column of a 2×2 matrix.

Consider a 2×2 matrix, denoted M_C , which is to be completed using its first column as observations:

$$M_C \triangleq \begin{bmatrix} w_{11}^* & ? \\ w_{21}^* & ? \end{bmatrix}.$$

The corresponding observation pattern matrix P_C is:

$$P_C = \begin{bmatrix} 1 & 0 \\ 1 & 0 \end{bmatrix}.$$

The associated adjacency matrix $\mathcal{A}_C$ for the bipartite graph is constructed as:

$$\mathcal{A}_C = \begin{bmatrix} \mathbf{0}_{2,2} & P_C^\top \\ P_C & \mathbf{0}_{2,2} \end{bmatrix} = \begin{bmatrix} 0 & 0 & 1 & 1 \\ 0 & 0 & 0 & 0 \\ 1 & 0 & 0 & 0 \\ 1 & 0 & 0 & 0 \end{bmatrix},$$

which forms a connected graph as illustrated in Figure 1(a). This setup leads to coupled training dynamics under non-zero initialization. The coupling arises because parameters used to construct w_{11} and w_{21} overlap. Specifically, elements from the first column of matrix B (i.e., b_{11}, b_{21}) are common to the computation of both w_{11} and w_{21} . This shared dependency links the dynamics. The below illustration highlights these shared (teal) and distinct (red/blue) parameters involved in forming the observed entries w_{11} and w_{21} :

$$\begin{aligned} \begin{bmatrix} w_{11} & w_{12} \\ w_{21} & w_{22} \end{bmatrix} &= \begin{bmatrix} a_{11} & a_{12} \\ a_{21} & a_{22} \end{bmatrix} \begin{bmatrix} b_{11} & b_{12} \\ b_{21} & b_{22} \end{bmatrix} \\ w_{11} &= a_{11}b_{11} + a_{12}b_{21} \\ w_{21} &= a_{21}b_{11} + a_{22}b_{21} \end{aligned}$$

The shared use of b_{11} and b_{21} in reconstructing both observed entries is what couples their learning dynamics.

For deeper matrices ($L \geq 3$), training dynamics are typically coupled, irrespective of the observation pattern. Consider, for instance, predicting entries from the disconnected matrix M_D where only diagonal elements are observed:

$$M_D \triangleq \begin{bmatrix} w_{11}^* & ? \\ ? & w_{22}^* \end{bmatrix}.$$

Even with such observations, for $L \geq 3$, coupling arises because parameters in intermediate layers are involved in computing multiple observed entries. This is illustrated in the following depth-3 example ($\mathbf{W} = \mathbf{W}_1 \mathbf{W}_2 \mathbf{W}_3$). Elements of the intermediate matrix $\mathbf{W}_2$ (colored teal) contribute to both the computation of w_{11} and w_{22} :

$$\begin{bmatrix} w_{11} & w_{12} \\ w_{21} & w_{22} \end{bmatrix} = \begin{bmatrix} (w_1)_{11} & (w_1)_{12} \\ (w_1)_{21} & (w_1)_{22} \end{bmatrix} \begin{bmatrix} (w_2)_{11} & (w_2)_{12} \\ (w_2)_{21} & (w_2)_{22} \end{bmatrix} \begin{bmatrix} (w_3)_{11} & (w_3)_{12} \\ (w_3)_{21} & (w_3)_{22} \end{bmatrix}.$$

Specifically, the observed entries are formed as:

$$\begin{aligned} w_{11} &= \left((w_1)_{11}(w_2)_{11} + (w_1)_{12}(w_2)_{21} \right) (w_3)_{11} \\ &\quad + \left((w_1)_{11}(w_2)_{12} + (w_1)_{12}(w_2)_{22} \right) (w_3)_{21}, \\ w_{22} &= \left((w_1)_{21}(w_2)_{11} + (w_1)_{22}(w_2)_{21} \right) (w_3)_{12} \\ &\quad + \left((w_1)_{21}(w_2)_{12} + (w_1)_{22}(w_2)_{22} \right) (w_3)_{22}. \end{aligned}$$

The shared involvement of all elements from $\mathbf{W}_2$ (the teal matrix) in forming both w_{11} and w_{22} leads to coupled dynamics, provided these elements are non-zero. (Conversely, if some elements were to become zero, this could potentially lead to decoupled dynamics, as illustrated in subsequent subsection.)

D.2. Decoupled Dynamics

For shallow ($L = 2$) matrices, decoupled dynamics correspond to disconnected observations. Therefore, to examine this disconnected case, we consider a 2×2 incomplete matrix example, denoted $\mathbf{M}_D$, which is to be completed using only its diagonal entries as observations:

$$\mathbf{M}_D \triangleq \begin{bmatrix} w_{11}^* & ? \\ ? & w_{22}^* \end{bmatrix}.$$

Then the observation matrix $\mathbf{P}_D$ can be constructed as:

$$\mathbf{P}_D = \begin{bmatrix} 1 & 0 \\ 0 & 1 \end{bmatrix},$$

and the adjacency matrix $\mathcal{A}_D$ can be constructed as:

$$\mathcal{A}_D = \begin{bmatrix} \mathbf{0}_{2,2} & \mathbf{P}_D^\top \\ \mathbf{P}_D & \mathbf{0}_{2,2} \end{bmatrix} = \begin{bmatrix} 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 1 \\ 1 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 \end{bmatrix},$$

which forms the disconnected graph as illustrated in Figure 1(a). This setup inherently leads to decoupled training dynamics. The decoupling can be visually understood by examining how distinct sets of elements in the factor matrices $\mathbf{A}$ and $\mathbf{B}$ contribute to the observed entries w_{11} and w_{22} . Specifically, as illustrated below, red-colored entries are exclusively involved in predicting w_{11} , while

blue-colored entries are exclusively involved in predicting w_{22} . These two sets of entries are disjoint, confirming the decoupled nature of the dynamics:

$$\begin{bmatrix} w_{11} & w_{12} \\ w_{21} & w_{22} \end{bmatrix} = \begin{bmatrix} a_{11} & a_{12} \\ a_{21} & a_{22} \end{bmatrix} \begin{bmatrix} b_{11} & b_{12} \\ b_{21} & b_{22} \end{bmatrix},$$

$$\begin{aligned} w_{11} &= a_{11}b_{11} + a_{12}b_{21}, \\ w_{22} &= a_{21}b_{12} + a_{22}b_{22}. \end{aligned}$$

For deep matrices, decoupled training dynamics are observed in at least two key scenarios. First, as detailed in Appendix F.2.3, an $\alpha \mathbf{I}_d$ initialization combined with diagonal-only observations leads to decoupled dynamics for any depth-factorized matrix.

To illustrate this for a deeper case, we revisit the $\mathbf{M}_D$ observation pattern in a depth-3 context. Lemma 9 in Appendix F.2.3 states that with such an initialization and observing only diagonal entries, all off-diagonal elements of the factor matrices $\mathbf{W}_l(t)$ remain zero throughout training. Consequently, the factor matrices $\mathbf{W}_1, \mathbf{W}_2, \mathbf{W}_3$ are diagonal. The product matrix $\mathbf{W}_{L:1}(t)$ is thus formed as:

$$\begin{bmatrix} w_{11} & w_{12} \\ w_{21} & w_{22} \end{bmatrix} = \begin{bmatrix} (w_1)_{11} & 0 \\ 0 & (w_1)_{22} \end{bmatrix} \begin{bmatrix} (w_2)_{11} & 0 \\ 0 & (w_2)_{22} \end{bmatrix} \begin{bmatrix} (w_3)_{11} & 0 \\ 0 & (w_3)_{22} \end{bmatrix}.$$

The observed entries are therefore computed as products of the respective diagonal elements:

$$\begin{aligned} w_{11} &= (w_1)_{11}(w_2)_{11}(w_3)_{11}, \\ w_{22} &= (w_1)_{22}(w_2)_{22}(w_3)_{22}. \end{aligned}$$

Since w_{11} depends only on the set of parameters $\{(\mathbf{W}_k)_{11}\}_{k=1}^3$ and w_{22} depends only on the entirely disjoint set of parameters $\{(\mathbf{W}_k)_{22}\}_{k=1}^3$, their training dynamics are decoupled.

Second, the training dynamics are also decoupled when the factor matrices are initialized as $d \times d$ zero matrices ($\mathbf{0}_d$). The reasoning is as follows: For any given set of observation indices $\Omega \subset [d] \times [d]$, the gradient flow dynamics for an (i, j) -th entry of a factor matrix $\mathbf{W}_l(t)$ (denoted $(w_l(t))_{ij}$) are given by:

$$\begin{aligned} \frac{d(w_l(t))_{ij}}{dt} &= -\frac{\partial \phi}{\partial (w_l(t))_{ij}} \\ &= -\sum_{(p,q) \in \Omega} (w_{pq}(t) - w_{pq}^*) \frac{\partial w_{pq}(t)}{\partial (w_l(t))_{ij}}. \end{aligned}$$

Here, the derivative of an element $w_{pq}(t)$ of the full product matrix $\mathbf{W}_{L:1}(t)$ with respect to $(w_l(t))_{ij}$ is:

$$\frac{\partial w_{pq}(t)}{\partial (w_l(t))_{ij}} = (\mathbf{W}_L(t) \mathbf{W}_{L-1}(t) \cdots \mathbf{W}_{l+1}(t))_{pi} (\mathbf{W}_{l-1}(t) \mathbf{W}_{l-2}(t) \cdots \mathbf{W}_1(t))_{jq},$$

where the first term is the (p, i) -th element of the product $\mathbf{W}_L(t) \cdots \mathbf{W}_{l+1}(t)$, and the second term is the (j, q) -th element of the product $\mathbf{W}_{l-1}(t) \cdots \mathbf{W}_1(t)$. If all factor matrices $\mathbf{W}_k(0)$ are initialized as zero matrices, then $w_{pq}(0) = 0$. Furthermore, the matrix products forming $\frac{\partial w_{pq}(0)}{\partial (w_l(0))_{ij}}$ are also zero. Consequently, $\frac{d(w_l(t))_{ij}}{dt}|_{t=0} = 0$. Since all entries $(w_l(0))_{ij}$ start at zero and their initial time derivatives are zero, they remain zero throughout training. Thus, all $\mathbf{W}_l(t) = \mathbf{0}_d$ for $t \geq 0$, leading to trivially decoupled dynamics.

Appendix E. Additional Experiments

This section provides additional experiments omitted from the main text.

E.1. Implicit Bias Experiments

In Figure 1, we conducted experiments with a 2×2 rank-1 ground truth matrix featuring specific connected/disconnected examples. To generalize these observations, we extended our experiments to a 3×3 rank-1 ground truth matrix, considering all possible connected/disconnected observation patterns. After accounting for symmetries to eliminate duplicates, this results in a total of 23 unique observation patterns, which are categorized into 17 connected and 6 disconnected cases.

For each of these 23 observation patterns, the 3×3 rank-1 ground truth matrix was generated using constituent vectors whose entries were sampled from a standard normal distribution. Each factor matrix was then initialized by sampling its entries from a Gaussian distribution with a mean of zero and a standard deviation of α . We performed 10 independent trials for each pattern.

Figure 4 illustrates that, consistent with the findings in Figure 1, a significant discrepancy exists between the behavior of depth-2 matrices and that of deeper matrices. This discrepancy becomes notably more pronounced for the disconnected observation patterns.

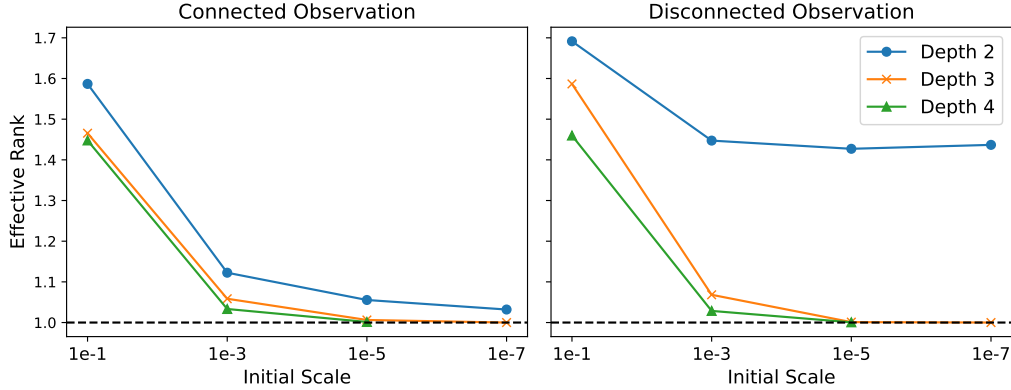


Figure 4: The left panel shows the averaged effective rank of all possible connected patterns as a function of the initial scale α^L . The right panel displays the corresponding data for all possible disconnected patterns.

In the context of Theorem 4, we further test our hypothesis: coupled dynamics can induce a low-rank bias, while decoupled dynamics cannot. We examine this hypothesis under various conditions by varying the ground truth value w^* and the dimension d . The results presented in Figure 5 (for $w^* = 1, d = 3$), Figure 6 (for $w^* = 10, d = 10$), and Figure 7 (for $w^* = 0.1, d = 10$) support this hypothesis.

Furthermore, we conducted experiments using gradient descent with a learning rate chosen to be sufficiently small, to validate our derived equations. For the results presented in Figure 8, we replicated the experimental setup of Figure 5 (however, the $\alpha = 10^{-10}$ case was excluded due to prohibitive computation time). The observed trends align well with those shown in Figure 5.

To validate that our initialization scheme (7) can achieve comparable outcomes to Gaussian initialization while offering more control, we conducted experiments on a 3×3 matrix completion

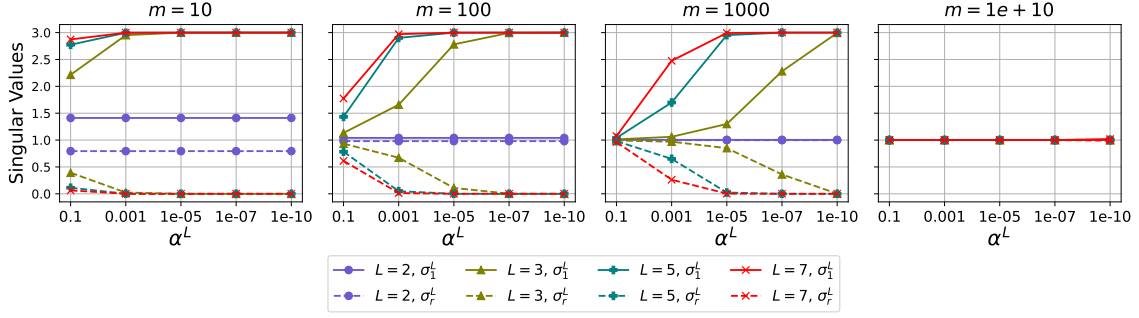


Figure 5: Numerical conditions identical to those in Figure 2, except with ground truth value $w^* = 1$ and dimension $d = 3$.

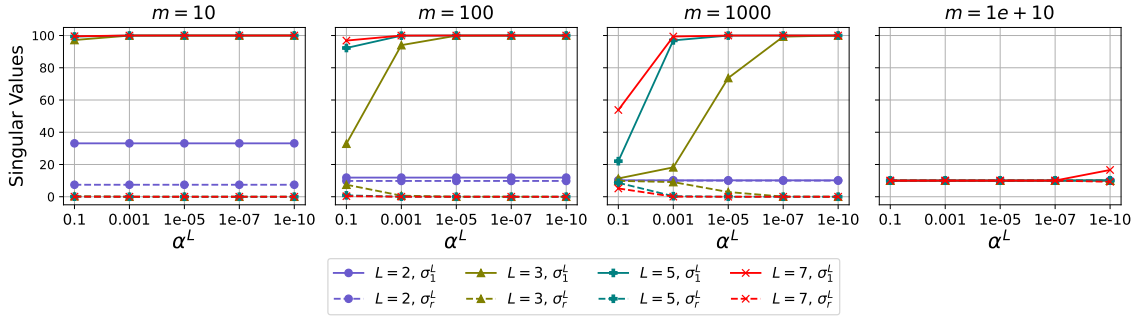


Figure 6: Numerical conditions identical to those in Figure 2, except with ground truth value $w^* = 10$ and dimension $d = 10$.

task with diagonal observations (i.e., $w_{11}^* = w_{22}^* = w_{33}^* = 1$). While our scheme allows initial rank properties to be adjusted via the parameter m , Gaussian initialization’s inherent randomness precludes such direct control. Therefore, for comparison with Gaussian initialization, we ran one thousand trials (seeds) and sorted the converged solutions by their rank.

A comparison of the results presented in Figure 9 indicates that the behavioral trends can appear similar, particularly because distinct low-rank inducing effects are often subtle and difficult to capture definitively in the depth-2 case. For deeper networks ($L \geq 3$), however, a clearer tendency to converge towards lower-rank solutions is typically observed as depth increases.

E.2. Loss of Plasticity Experiments

Section A.2 discussed a scenario where pre-training employs diagonal entries, after which an off-diagonal term (specifically, w_{12}^*) is introduced to restore connectivity, leading to coupled dynamics. Theorem 6 establishes that, in this situation, the model indeed does not converge to a low-rank solution. To empirically validate this theoretical finding, we conducted experiments using the family of initializations (7) tailored to this specific scenario, with results detailed in Figures 10 and 11. These experiments utilized a depth-2 model to reconstruct the ground-truth matrix, with an initialization scale set to $\alpha = 10^{-35}$. Notably, if the initialization scale α is set significantly lower, as the dynamics are coupled, a cold-started model can converge to solutions exhibiting a more pronounced low-rank structure.

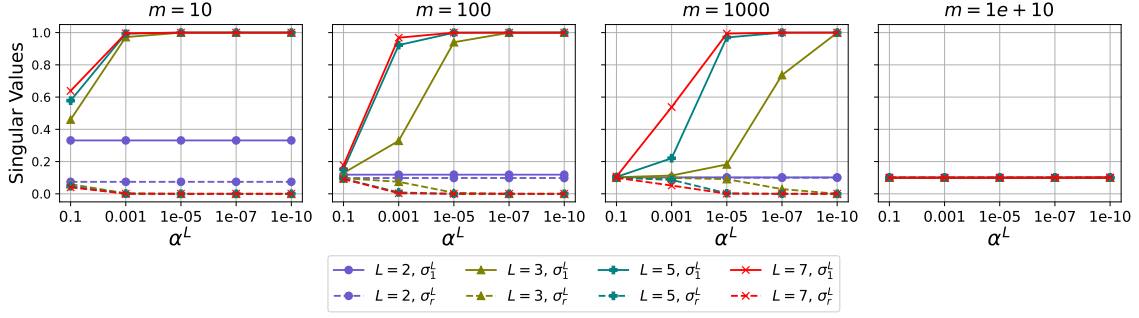


Figure 7: Numerical conditions identical to those in Figure 2, except with ground truth value $w^* = 0.1$ and dimension $d = 10$.

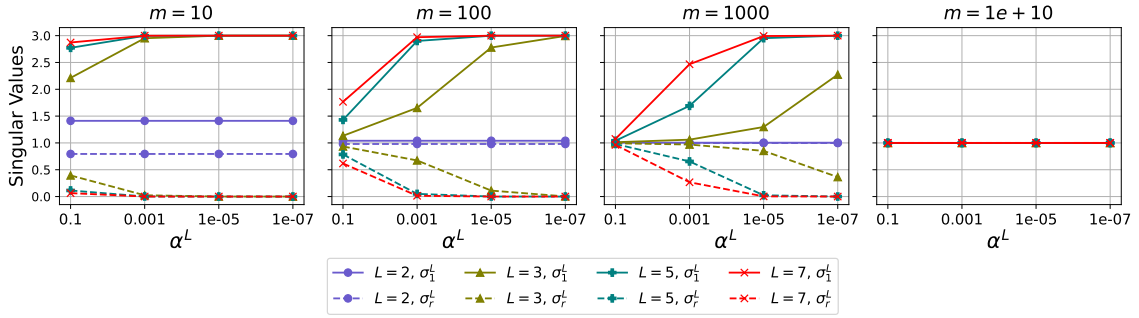


Figure 8: Gradient descent experiments conducted under conditions identical to those in Figure 5.

For the case presented in Figure 10, where $w^* = 1, w_{12}^* = 0.1$, following Theorem 6, the theoretical lower bound on the stable rank for a warm-started model initialized diagonally ($m = \infty$) is approximately 1.45, while the empirically observed stable rank is approximately 1.8. Even in scenarios where substantial new information must be learned (e.g., by setting w_{12}^* to a large value), loss of plasticity is empirically observed, primarily manifesting as high test error (i.e., a significant gap between the target w_{21}^* and the converged w_{21}). While Theorem 6’s analysis via stable rank does not fully explain an accompanying low-rank bias (a point consistent with Figure 11), the theorem does predict that w_{21} converges to a negative value, which implies a large test loss.

Furthermore, we performed additional experiments with different diagonal entry values to investigate whether this argument extends to other scenarios (results shown in Figure 12), although specific theoretical guarantees have not been established for these broader cases. We observe that even in these varied settings, both the effective rank and the stable rank of a warm-started model substantially exceed one, whereas cold-started models can converge to lower-rank solutions.

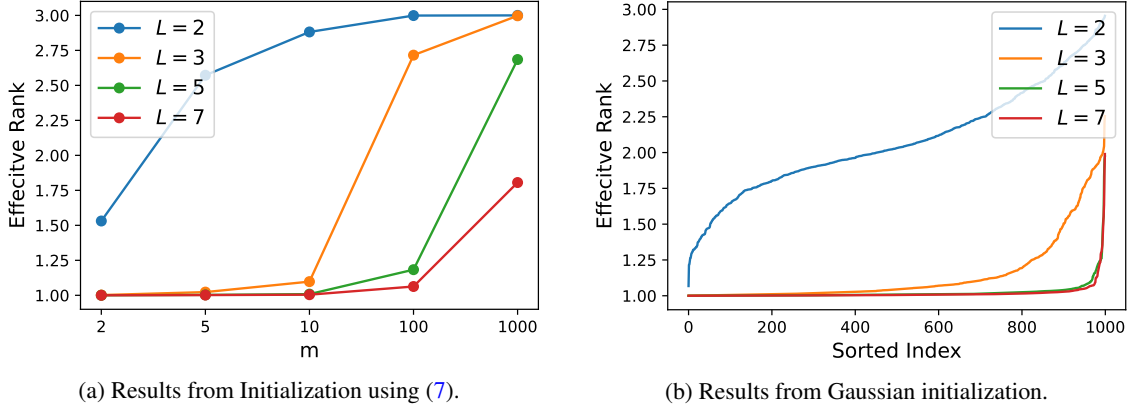


Figure 9: **(a)** Effective rank for the initialization scheme in (7). The x-axis denotes the parameter m , which controls the initial rank characteristics of the model, while the y-axis represents the corresponding effective rank after convergence. **(b)** Effective rank distributions for Gaussian initialization. The results are from 1000 independent trials, sorted by their converged effective rank. The x-axis denotes the sorted trial index (from lowest to highest converged rank), and the y-axis represents the corresponding effective rank after convergence.

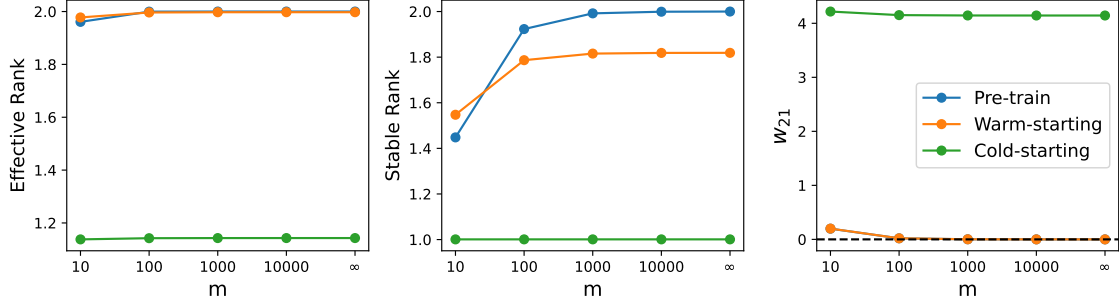


Figure 10: Experimental results for a 2×2 rank-1 ground-truth matrix $\mathbf{W}^*$ with $w_{11}^* = w_{22}^* = 1$ and $w_{12}^* = 0.5$ (implying $w_{21}^* = 2$ for rank-1 structure). Models, initialized according to (7), are first pre-trained on diagonal entries. After achieving zero-loss convergence in pre-training, the off-diagonal element w_{12}^* is introduced, and models are subsequently trained on combined diagonal and off-diagonal observations. The plots display: (Left and Middle) effective rank under different settings; (Right) converged value of $w_{21}(\infty)$. Key observations: (1) Warm-starting with a model that converged to a high-rank solution during pre-training tends to maintain this high rank, even when presented with the same subsequent observations as a cold-started model. (2) In the theoretically analyzed $m = \infty$ case, $w_{21}(\infty) < 0$ is observed, which correlates with the highest effective rank.

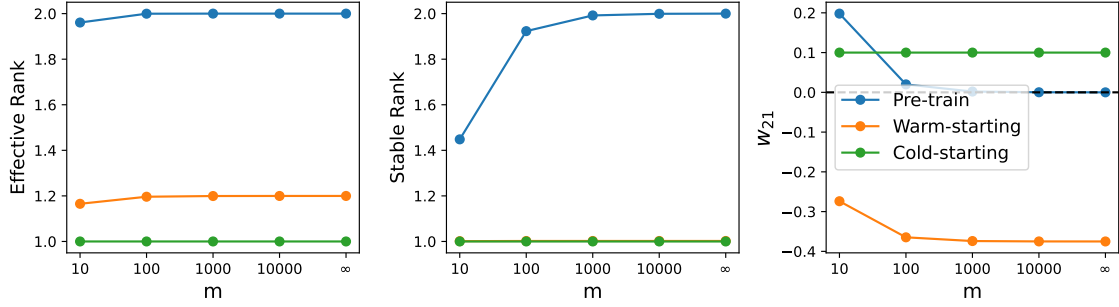


Figure 11: Experimental conditions identical to those in Figure 10, except with ground truth value $w_{12}^* = 10$. The model have to predict w_{21}^* as 0.1

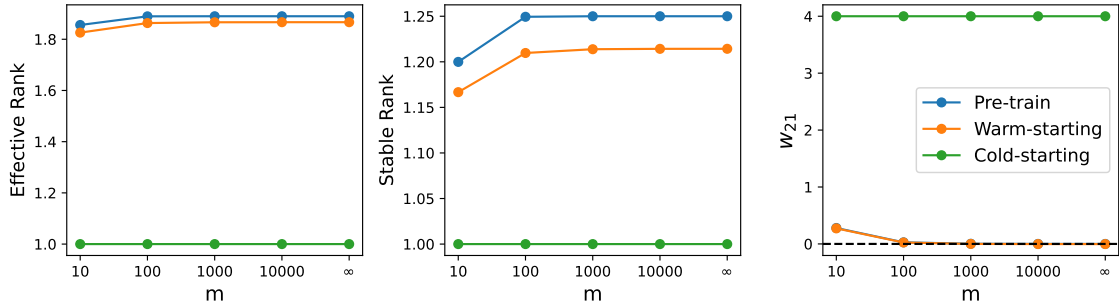


Figure 12: Experimental conditions identical to those in Figure 10, except with ground truth value $w_{11}^* = 1$, $w_{22}^* = 2$, and $w_{12}^* = 0.5$. The model have to predict w_{21}^* as 4.

Appendix F. Proof for Section 3

In this and the following section, we prove the Propositions and Theorems presented in the main text. We begin with the proof of Theorem 1.

F.1. Proof for Theorem 1

When convergence is guaranteed, we can define reference vector $\mathbf{u}^* \triangleq \frac{\mathbf{b}_1(\infty)}{\|\mathbf{b}_1(\infty)\|} \in \mathbb{R}^d$, which is entirely determined by their initial values. We decompose $\mathbf{a}_1(t)$, $\mathbf{a}_2(t)$, and $\mathbf{b}_1(t)$ into two components: one parallel to $\mathbf{u}^*$ and one perpendicular to $\mathbf{u}^*$:

$$\mathbf{a}_1(t) = \mathbf{a}_{1\parallel}(t) + \mathbf{a}_{1\perp}(t), \quad \mathbf{a}_2(t) = \mathbf{a}_{2\parallel}(t) + \mathbf{a}_{2\perp}(t), \quad \mathbf{b}_1(t) = \mathbf{b}_{1\parallel}(t) + \mathbf{b}_{1\perp}(t).$$

For any vector $\mathbf{u} \in \mathbb{R}^d$, the parallel component is defined as $\mathbf{u}_{\parallel} = (\mathbf{u}^{*\top} \mathbf{u}) \mathbf{u}^*$, and the perpendicular component as $\mathbf{u}_{\perp} = \mathbf{u} - \mathbf{u}_{\parallel}$.

We introduce notation to quantify the alignment of each vector with $\mathbf{u}^*$:

$$\alpha_{\mathbf{a}_1}(t) = \mathbf{u}^{*\top} \mathbf{a}_1(t), \quad \alpha_{\mathbf{a}_2}(t) = \mathbf{u}^{*\top} \mathbf{a}_2(t), \quad \alpha_{\mathbf{b}_1}(t) = \mathbf{u}^{*\top} \mathbf{b}_1(t). \quad (11)$$

Additionally, we define notation to measure the magnitude of the perpendicular components:

$$\beta_{\mathbf{a}_1}(t) = \|\mathbf{a}_{1\perp}(t)\|_2^2, \quad \beta_{\mathbf{a}_2}(t) = \|\mathbf{a}_{2\perp}(t)\|_2^2, \quad \beta_{\mathbf{b}_1}(t) = \|\mathbf{b}_{1\perp}(t)\|_2^2. \quad (12)$$

Then, using equation (4), time evolution of each component in equation (11) can be written as:

$$\begin{aligned} \dot{\alpha}_{\mathbf{a}_1}(t) &= \mathbf{u}^{*\top} \dot{\mathbf{a}}_1(t) \\ &= \underbrace{(w_{11}^* - \mathbf{a}_1^\top(t) \mathbf{b}_1(t))}_{\triangleq r_1(t)} \mathbf{u}^{*\top} \mathbf{b}_1(t) \\ &= r_1(t) \alpha_{\mathbf{b}_1}(t). \end{aligned} \quad (13)$$

Likewise, for $\alpha_{\mathbf{a}_2}(t)$, we derive:

$$\begin{aligned} \dot{\alpha}_{\mathbf{a}_2}(t) &= \mathbf{u}^{*\top} \dot{\mathbf{a}}_2(t) \\ &= \underbrace{(w_{21}^* - \mathbf{a}_2^\top(t) \mathbf{b}_1(t))}_{\triangleq r_2(t)} \mathbf{u}^{*\top} \mathbf{b}_1(t) \\ &= r_2(t) \alpha_{\mathbf{b}_1}(t). \end{aligned} \quad (14)$$

Finally, for $\alpha_{\mathbf{b}_1}(t)$, we have:

$$\begin{aligned} \dot{\alpha}_{\mathbf{b}_1}(t) &= \mathbf{u}^{*\top} \dot{\mathbf{b}}_1(t) \\ &= (w_{11}^* - \mathbf{a}_1^\top(t) \mathbf{b}_1(t)) \mathbf{u}^{*\top} \mathbf{a}_1(t) + (w_{21}^* - \mathbf{a}_2^\top(t) \mathbf{b}_1(t)) \mathbf{u}^{*\top} \mathbf{a}_2(t) \\ &= r_1(t) \alpha_{\mathbf{a}_1}(t) + r_2(t) \alpha_{\mathbf{a}_2}(t). \end{aligned} \quad (15)$$

Also, for the perpendicular components, their time evolution can be derived as:

$$\begin{aligned} \dot{\beta}_{\mathbf{a}_1}(t) &= 2\mathbf{a}_{1\perp}(t) \cdot \dot{\mathbf{a}}_{1\perp}(t) \\ &= 2\mathbf{a}_{1\perp}(t) \cdot \frac{d}{dt} \left(\mathbf{a}_1(t) - \left(\mathbf{u}^{*\top} \mathbf{a}_1(t) \right) \mathbf{u}^* \right) \\ &= 2\mathbf{a}_{1\perp}(t) \cdot \left(r_1(t) \mathbf{b}_1(t) - r_1(t) \left(\mathbf{u}^{*\top} \mathbf{b}_1(t) \right) \mathbf{u}^* \right). \end{aligned}$$

Noting that $\mathbf{a}_{1\perp}(t)$ is perpendicular to $\mathbf{u}^*$, the second term in the parenthesis can be considered zero. Thus, we have

$$\dot{\beta}_{\mathbf{a}_1}(t) = 2r_1(t)\mathbf{a}_{1\perp}(t)^\top \mathbf{b}_{1\perp}(t).$$

Likewise, for $\beta_{\mathbf{a}_2}(t)$ and $\beta_{\mathbf{b}_1}(t)$, we can derive their time derivative as:

$$\dot{\beta}_{\mathbf{a}_2}(t) = 2r_2(t)\mathbf{a}_{2\perp}(t)^\top \mathbf{b}_{1\perp}(t), \quad \dot{\beta}_{\mathbf{b}_1}(t) = \dot{\beta}_{\mathbf{a}_1}(t) + \dot{\beta}_{\mathbf{a}_2}(t).$$

Note that by the definition of $\mathbf{u}^*$, we have $\beta_{\mathbf{b}_1}(\infty) = 0$. Integrating the identity $\dot{\beta}_{\mathbf{b}_1}(t) = \dot{\beta}_{\mathbf{a}_1}(t) + \dot{\beta}_{\mathbf{a}_2}(t)$ from $t = 0$ to ∞ gives:

$$\beta_{\mathbf{a}_1}(\infty) + \beta_{\mathbf{a}_2}(\infty) = \underbrace{\beta_{\mathbf{a}_1}(0) + \beta_{\mathbf{a}_2}(0) - \beta_{\mathbf{b}_1}(0)}_{\triangleq \beta_0 \geq 0},$$

which depends solely on the initial values $\mathbf{a}_1(0)$, $\mathbf{a}_2(0)$, and $\mathbf{b}_1(0)$. This equation shows that if the initial value β_0 is small, the solution will eventually align with $\mathbf{u}^*$. However, since we do not know $\mathbf{u}^*$ in advance, one natural way to ensure small perpendicular components is to initialize the entire norms of $\mathbf{a}_1(0)$, $\mathbf{a}_2(0)$ to be sufficiently small.

To develop a more rigorous understanding, we analyze the parallel components. Under the assumption of convergence, we have:

$$\mathbf{a}_1(\infty)^\top \mathbf{b}_1(\infty) = w_{11}^*, \quad \mathbf{a}_2(\infty)^\top \mathbf{b}_1(\infty) = w_{21}^*.$$

Decomposing $\mathbf{a}_1(\infty)$ and $\mathbf{a}_2(\infty)$ leads to:

$$\begin{aligned} \mathbf{a}_1(\infty)^\top \mathbf{b}_1(\infty) &= \left(\mathbf{a}_{1\perp}(\infty) + \mathbf{u}^{*\top} \mathbf{a}_1(\infty) \mathbf{u}^* \right)^\top \mathbf{b}_1(\infty) \\ &= \alpha_{\mathbf{a}_1}(\infty) \alpha_{\mathbf{b}_1}(\infty) = w_{11}^*, \end{aligned} \tag{16}$$

$$\begin{aligned} \mathbf{a}_2(\infty)^\top \mathbf{b}_1(\infty) &= \left(\mathbf{a}_{2\perp}(\infty) + \mathbf{u}^{*\top} \mathbf{a}_2(\infty) \mathbf{u}^* \right)^\top \mathbf{b}_1(\infty) \\ &= \alpha_{\mathbf{a}_2}(\infty) \alpha_{\mathbf{b}_1}(\infty) = w_{21}^*. \end{aligned} \tag{17}$$

Using equations (13)–(15), and noting that

$$\frac{d}{dt} \alpha_{\mathbf{b}_1}^2(t) = \frac{d}{dt} (\alpha_{\mathbf{a}_1}^2(t) + \alpha_{\mathbf{a}_2}^2(t)),$$

we can integrate the above equation both sides over time from 0 to ∞ to obtain:

$$\alpha_{\mathbf{a}_1}^2(\infty) + \alpha_{\mathbf{a}_2}^2(\infty) = \alpha_{\mathbf{b}_1}^2(\infty) + \underbrace{\alpha_{\mathbf{a}_1}^2(0) + \alpha_{\mathbf{a}_2}^2(0) - \alpha_{\mathbf{b}_1}^2(0)}_{\triangleq \alpha_0}. \tag{18}$$

By solving equations (16), (17), and (18), we can obtain closed-form solution of $\alpha_{\mathbf{a}_1}(\infty)$, $\alpha_{\mathbf{a}_2}(\infty)$, and $\alpha_{\mathbf{b}_1}(\infty)$ as follows:

$$\alpha_{\mathbf{a}_1}^2(\infty) = \frac{2w_{11}^{*2}}{\sqrt{\alpha_0^2 + 4w_{11}^{*2} + 4w_{21}^{*2} - \alpha_0}}, \quad \alpha_{\mathbf{a}_2}^2(\infty) = \frac{2w_{21}^{*2}}{\sqrt{\alpha_0^2 + 4w_{11}^{*2} + 4w_{21}^{*2} - \alpha_0}}, \tag{19}$$

$$\alpha_{\mathbf{b}_1}^2(\infty) = \frac{\sqrt{\alpha_0^2 + 4w_{11}^{*2} + 4w_{21}^{*2} - \alpha_0}}{2}. \tag{20}$$

Thus, we can upper bound the proportion of the perpendicular component of $\mathbf{a}_1(\infty)$ and $\mathbf{a}_2(\infty)$ relative to its total magnitude as follows:

$$\begin{aligned}\frac{\|\mathbf{a}_{1\perp}(\infty)\|^2}{\|\mathbf{a}_1(\infty)\|^2} &= \frac{\beta_{\mathbf{a}_1}(\infty)}{\alpha_{\mathbf{a}_1}^2(\infty) + \beta_{\mathbf{a}_1}(\infty)} \leq \frac{\beta_0 \left(\sqrt{\alpha_0^2 + 4w_{11}^{*2} + 4w_{21}^{*2}} - \alpha_0 \right)}{2w_{11}^{*2}}, \\ \frac{\|\mathbf{a}_{2\perp}(\infty)\|^2}{\|\mathbf{a}_2(\infty)\|^2} &= \frac{\beta_{\mathbf{a}_2}(\infty)}{\alpha_{\mathbf{a}_2}^2(\infty) + \beta_{\mathbf{a}_2}(\infty)} \leq \frac{\beta_0 \left(\sqrt{\alpha_0^2 + 4w_{11}^{*2} + 4w_{21}^{*2}} - \alpha_0 \right)}{2w_{21}^{*2}}.\end{aligned}$$

To further refine these bounds, we analyze the terms β_0 and $S(\alpha_0) \triangleq \sqrt{\alpha_0^2 + 4w_{11}^{*2} + 4w_{21}^{*2}} - \alpha_0$. By the definition of β_0 , it is upper bounded by $\|\mathbf{a}_1(0)\|^2 + \|\mathbf{a}_2(0)\|^2$, which equals with $\|\mathbf{A}(0)\|_F^2$. Also, by the definition of α_0 , we have:

$$-\|\mathbf{b}_1(0)\|_2^2 \leq \alpha_0 \leq \|\mathbf{a}_1(0)\|_2^2 + \|\mathbf{a}_2(0)\|_2^2 = \|\mathbf{A}(0)\|_F^2.$$

Noting that the function $f(x) = \sqrt{x^2 + C} - x$ (where $C > 0$) is non-negative and monotonically decreasing for all $x \in \mathbb{R}$, we can upper bound $S(\alpha_0)$ using the lower bound of α_0 :

$$\begin{aligned}S(\alpha_0) &\leq S(-\|\mathbf{b}_1(0)\|_2^2) \\ &= \sqrt{(-\|\mathbf{b}_1(0)\|_2^2)^2 + 4(w_{11}^{*2} + w_{21}^{*2})} - (-\|\mathbf{b}_1(0)\|_2^2) \\ &= \sqrt{\|\mathbf{b}_1(0)\|_2^4 + 4(w_{11}^{*2} + w_{21}^{*2})} + \|\mathbf{b}_1(0)\|_2^2.\end{aligned}$$

Substituting these bounds for β_0 and $S(\alpha_0)$ into the inequality $\frac{\|\mathbf{a}_{1\perp}(\infty)\|^2}{\|\mathbf{a}_1(\infty)\|_2^2} \leq \frac{\beta_0 S(\alpha_0)}{2w_{11}^{*2}}$, we obtain the final upper bound for the proportion of the perpendicular component of $\mathbf{a}_1(\infty)$:

$$\frac{\|\mathbf{a}_{1\perp}(\infty)\|^2}{\|\mathbf{a}_1(\infty)\|_2^2} \leq \frac{\|\mathbf{A}(0)\|_F^2 \left(\sqrt{\|\mathbf{b}_1(0)\|_2^4 + 4(w_{11}^{*2} + w_{21}^{*2})} + \|\mathbf{b}_1(0)\|_2^2 \right)}{2w_{11}^{*2}}.$$

A similar bound applies to $\frac{\|\mathbf{a}_{2\perp}(\infty)\|^2}{\|\mathbf{a}_2(\infty)\|_2^2}$:

$$\frac{\|\mathbf{a}_{2\perp}(\infty)\|^2}{\|\mathbf{a}_2(\infty)\|_2^2} \leq \frac{\|\mathbf{A}(0)\|_F^2 \left(\sqrt{\|\mathbf{b}_1(0)\|_2^4 + 4(w_{11}^{*2} + w_{21}^{*2})} + \|\mathbf{b}_1(0)\|_2^2 \right)}{2w_{21}^{*2}}.$$

F.2. Proof for Proposition 3

According to the definition of coupled/decoupled dynamics presented in Definition 2, for the family of initializations defined in (7) along with the diagonal observations ($\Omega_{\text{diag}}^{(d)}$), we divide the cases to ensure that all possible scenarios for this family of initializations are covered.

F.2.1. CASE FOR $L = 2$

First, we consider the depth-2 ($L = 2$) case. Each diagonal observation, $w_{ii}(t)$, is the inner product of the i -th row of $\mathbf{A}(t)$ and the i -th column of $\mathbf{B}(t)$. Then, when we take the gradient $\nabla_{\theta(t)} w_{ii}(t)$, where $\theta(t)$ represents the concatenation of $\mathbf{A}(t)$ and $\mathbf{B}(t)$, this gradient has non-zero components only corresponding to the i -th row of $\mathbf{A}(t)$ and the i -th column of $\mathbf{B}(t)$; all other components are zero for all $t \geq 0$. Therefore, for any $j \neq i$, the inner product $\langle \nabla_{\theta(t)} w_{ii}(t), \nabla_{\theta(t)} w_{jj}(t) \rangle$ must be zero. This means that there exists a partition of $\Omega_{\text{diag}}^{(d)}$ into disjoint subsets $\Omega_1, \dots, \Omega_d$, where each $\Omega_i = \{(i, i)\}$. Therefore, for any initialization, the training dynamics are **decoupled**.

F.2.2. CASE FOR $L \geq 3$ AND $1 < m < \infty$

For the deeper matrix case ($L \geq 3$), we first note that each diagonal observation $w_{ii}(t)$ can be expressed as:

$$w_{ii}(t) = \sum_{i_{L-1}=1}^d \cdots \sum_{i_1=1}^d (\mathbf{W}_L(t))_{i, i_{L-1}} (\mathbf{W}_{L-1}(t))_{i_{L-1}, i_{L-2}} \cdots (\mathbf{W}_1(t))_{i_1, i}.$$

Now, let us consider m satisfying $1 < m < \infty$, under which every entry of each weight matrix $\mathbf{W}_l(0)$ (for $l = 1, \dots, L$) is initialized as a positive value. Given that $w_{ii}(0)$ is a sum of products of these positive entries, its gradient with respect to the parameters $\theta(0)$ (which includes the entries of $\mathbf{W}_l(0)$), $\nabla_{\theta(0)} w_{ii}(0)$, will also consist of components that are sums of products of positive values. Therefore, it is asserted that each relevant component of $\nabla_{\theta(0)} w_{ii}(0)$ is positive at initialization. Consequently, for any $j \neq i$, since both $\nabla_{\theta(0)} w_{ii}(0)$ and $\nabla_{\theta(0)} w_{jj}(0)$ have all their corresponding components positive, their inner product $\langle \nabla_{\theta(0)} w_{ii}(0), \nabla_{\theta(0)} w_{jj}(0) \rangle$ will be non-zero (specifically, positive). This non-zero inner product signifies **coupled dynamics**.

F.2.3. CASE FOR $L \geq 3$ AND $m = \infty$

Next, we examine the $m = \infty$ case, which corresponds to initializing each factor matrix $\mathbf{W}_l(0)$ as a scaled identity, i.e., $\mathbf{W}_l(0) = \alpha_l \mathbf{I}_d$. The following lemma states that under this initialization, and for dynamics driven by diagonal observations (from $\Omega_{\text{diag}}^{(d)}$), all off-diagonal elements of each $\mathbf{W}_l(t)$ remain zero for all $t \geq 0$.

Lemma 9 *For a set of L matrices $\mathbf{W}_1(t), \dots, \mathbf{W}_L(t) \in \mathbb{R}^{d \times d}$, let $\mathbf{W}_{L:1}(t) = \mathbf{W}_L(t) \cdots \mathbf{W}_1(t)$. Following gradient flow dynamics in (3), if each factor matrix $\mathbf{W}_l(0)$ is initialized as a diagonal matrix (e.g., $\mathbf{W}_l(0) = \alpha_l \mathbf{I}_d$ for scalars α_l), then all off-diagonal elements of each matrix $\mathbf{W}_l(t)$ remain zero for all $t \geq 0$.*

Proof For a given diagonal observation indices $\Omega_{\text{diag}}^{(d)}$, if we consider the gradient flow dynamics for an (i, j) -th entry of the factor matrix $\mathbf{W}_l(t)$ ($\triangleq (w_l(t))_{ij}$), we have:

$$\begin{aligned} \frac{d(w_l(t))_{ij}}{dt} &= -\frac{\partial \phi}{\partial (w_l(t))_{ij}} \\ &= -\sum_{p=1}^d (w_{pp}(t) - w_{pp}^*) \frac{\partial w_{pp}(t)}{\partial (w_l(t))_{ij}}, \end{aligned}$$

Here, the derivative of a diagonal element $w_{pp}(t)$ with respect to $(w_l(t))_{ij}$ is:

$$\frac{\partial w_{pp}(t)}{\partial (w_l(t))_{ij}} = (\mathbf{W}_L(t) \mathbf{W}_{L-1}(t) \cdots \mathbf{W}_{l+1}(t))_{pi} (\mathbf{W}_{l-1}(t) \mathbf{W}_{l-2}(t) \cdots \mathbf{W}_1(t))_{jp},$$

where the first term is (p, i) -th element of the product $\mathbf{W}_L(t) \mathbf{W}_{L-1}(t) \cdots \mathbf{W}_{l+1}(t)$, and the second term is (j, p) -th element of the product $\mathbf{W}_{l-1}(t) \mathbf{W}_{l-2}(t) \cdots \mathbf{W}_1(t)$. We want to show that if all $\mathbf{W}_l(t)$ are diagonal, then $\frac{d(w_l(t))_{ij}}{dt} = 0$ for any off-diagonal element $(w_l(t))_{ij}$ (i.e., $i \neq j$).

Assume at a given time t that all factor matrices $\mathbf{W}_l(t)$ are diagonal. Then, the product $P(t) \triangleq \prod_{k=l+1}^L \mathbf{W}_k(t)$ is diagonal. Similarly, the product $S(t) \triangleq \prod_{k=1}^{l-1} \mathbf{W}_k(t)$ is diagonal. For $\frac{\partial w_{pp}(t)}{\partial (w_l(t))_{ij}}$ to be non-zero (given all $\mathbf{W}_l(t)$ are diagonal), both $(P(t))_{pi}$ and $(S(t))_{jp}$ must be non-zero. This requires $p = i$ and $j = p$, which implies $i = j$.

However, we are considering an off-diagonal element $(w_l(t))_{ij}$, for which $i \neq j$. This means that if all $\mathbf{W}_l(t)$ are diagonal, then for any p :

$$\frac{\partial w_{pp}}{\partial (w_l(t))_{ij}} = 0, \quad \text{if } i \neq j$$

Substituting this into the dynamic equation for $(w_l(t))_{ij}$:

$$\frac{d(w_l(t))_{ij}}{dt} = -\sum_{p=1}^d (w_{pp}(t) - w_{pp}^*) \cdot 0 = 0, \quad \text{if } i \neq j$$

Initially, $\mathbf{W}_l(0)$ are diagonal, so all off-diagonal elements $(w_l(t))_{ij}$ are zero for $i \neq j$. Since their time derivatives are zero when they are zero (i.e., when the matrices are diagonal), these off-diagonal elements remain zero for all $t \geq 0$. $\blacksquare$

With Lemma 9, the factor matrices $\mathbf{W}_l(t)$ remain diagonal, so $w_{ii}(t) = (\mathbf{W}_L(t))_{ii} \cdots (\mathbf{W}_1(t))_{ii}$. This structure leads to decoupled dynamics because each $w_{ii}(t)$ depends exclusively on the set of parameters $\{(\mathbf{W}_k(t))_{ii}\}_{k=1}^L$, while $w_{jj}(t)$ (for $j \neq i$) depends on the distinct set $\{(\mathbf{W}_k(t))_{jj}\}_{k=1}^L$. Consequently, for any $j \neq i$, their respective gradients $\nabla_{\theta(t)} w_{ii}(t)$ and $\nabla_{\theta(t)} w_{jj}(t)$ are orthogonal, meaning their inner product is zero:

$$\langle \nabla_{\theta(t)} w_{ii}(t), \nabla_{\theta(t)} w_{jj}(t) \rangle = 0.$$

This orthogonality implies that the learning for each diagonal entry is independent, allowing a conceptual partition of $\Omega_{\text{diag}}^{(d)}$ into disjoint subsets $\Omega_i = \{(i, i)\}$. Therefore, under this specific diagonal initialization (the $m = \infty$ case), the training dynamics are **decoupled**.

F.3. Proof for Theorem 4

Before presenting the proof of Theorem 4, we first restate the problem setting. The model is defined as $\mathbf{W}_{L:1}(t) = \mathbf{W}_L(t)\mathbf{W}_{L-1}(t) \cdots \mathbf{W}_1(t)$, where each factor matrix $\mathbf{W}_l(t) \in \mathbb{R}^{d \times d}$ is subject to diagonal observations $\Omega_{\text{diag}}^{(d)} = \{(i, i)\}_{i=1}^d$, and follows the gradient flow described in (3). We also assume that all diagonal entries are equal, i.e., $w^* \triangleq w_{11}^* = w_{22}^* \cdots = w_{dd}^*$. To simplify notation, we use $\ell(\mathbf{W}_{L:1}(t))$ in place of $\ell(\mathbf{W}_{L:1}(t); \Omega_{\text{diag}})$ when the context is clear. The explicit gradient flow dynamics for each factor matrix is then given by:

$$\dot{\mathbf{W}}_l(t) = - \prod_{i=l+1}^L \mathbf{W}_i(t)^\top \cdot \nabla \ell(\mathbf{W}_{L:1}(t)) \cdot \prod_{i=1}^{l-1} \mathbf{W}_i(t)^\top, \quad (21)$$

where $\nabla \ell(\mathbf{W}_{L:1}(t)) = \text{diag}(r_1(t), r_2(t), \dots, r_d(t))$. Here, the residual term is defined as $r_i(t) \triangleq w_{ii}(t) - w^*$. To begin, we first present the preliminary lemma required for the following result.

Lemma 10 *Let $\mathbf{I}_n$ denote the $n \times n$ identity matrix and $\mathbf{J}_n \triangleq \mathbb{1}_n \mathbb{1}_n^\top$ denote the $n \times n$ matrix with all entries equal to 1. Then the set*

$$\mathcal{S} = \{a\mathbf{I}_n + b\mathbf{J}_n \mid a, b \in \mathbb{R}\}$$

is closed under scalar multiplication, addition, and matrix multiplication. Also, any two matrices $\mathbf{A}, \mathbf{B} \in \mathcal{S}$ commute.

Proof Let

$$\mathbf{A} = a\mathbf{I}_n + b\mathbf{J}_n \quad \text{and} \quad \mathbf{B} = c\mathbf{I}_n + d\mathbf{J}_n,$$

with $a, b, c, d \in \mathbb{R}$, and let $\lambda \in \mathbb{R}$ be an arbitrary scalar.

Scalar Multiplication:

$$\lambda \mathbf{A} = \lambda(a\mathbf{I}_n + b\mathbf{J}_n) = (\lambda a)\mathbf{I}_n + (\lambda b)\mathbf{J}_n.$$

Since $\lambda a, \lambda b \in \mathbb{R}$, it follows that $\lambda \mathbf{A} \in \mathcal{S}$.

Addition:

$$\mathbf{A} + \mathbf{B} = (a\mathbf{I}_n + b\mathbf{J}_n) + (c\mathbf{I}_n + d\mathbf{J}_n) = (a + c)\mathbf{I}_n + (b + d)\mathbf{J}_n.$$

Since $a + c, b + d \in \mathbb{R}$, we have $\mathbf{A} + \mathbf{B} \in \mathcal{S}$.

Matrix Multiplication:

$$\mathbf{AB} = (a\mathbf{I}_n + b\mathbf{J}_n)(c\mathbf{I}_n + d\mathbf{J}_n).$$

Using the distributive property and the facts that

$$\mathbf{I}_n \mathbf{J}_n = \mathbf{J}_n \mathbf{I}_n = \mathbf{J}_n \quad \text{and} \quad \mathbf{J}_n^2 = n\mathbf{J}_n,$$

we expand:

$$\begin{aligned} \mathbf{AB} &= ac \mathbf{I}_n \mathbf{I}_n + ad \mathbf{I}_n \mathbf{J}_n + bc \mathbf{J}_n \mathbf{I}_n + bd \mathbf{J}_n^2 \\ &= ac \mathbf{I}_n + ad \mathbf{J}_n + bc \mathbf{J}_n + bd (n\mathbf{J}_n) \\ &= ac \mathbf{I}_n + (ad + bc + nbd) \mathbf{J}_n. \end{aligned}$$

Thus, $\mathbf{AB}$ is of the form $\alpha \mathbf{I}_n + \beta \mathbf{J}_n$ with $\alpha = ac$ and $\beta = ad + bc + nbd$, and hence $\mathbf{AB} \in \mathcal{S}$.

Commutativity: By the same procedure as above,

$$\begin{aligned} \mathbf{AB} &= (a\mathbf{I}_n + b\mathbf{J}_n)(c\mathbf{I}_n + d\mathbf{J}_n) \\ &= ac\mathbf{I}_n + (ad + bc + nbd)\mathbf{J}_n \\ &= ca\mathbf{I}_n + (cb + da + ndb)\mathbf{J}_n \\ &= \mathbf{BA}, \end{aligned}$$

which completes the proof. $\blacksquare$

F.3.1. CASE FOR $L = 2$ & $L \geq 3$ AND $1 < m < \infty$

We will first examine two main scenarios: the depth-2 ($L = 2$) case and deeper networks ($L \geq 3$) where $1 < m < \infty$. The $m = \infty$ case will be considered separately in the later subsection, as its initialization with $\alpha \mathbf{I}_d$ warrants distinct treatment.

We now proceed to prove the following auxiliary results, which are used in the proof of Lemma 12. Based on Lemmas 11–13, we will show that all diagonal entries across all layers are identical, and likewise, all off-diagonal entries across layers are also equal.

Lemma 11 *Suppose we have a ground truth matrix $\mathbf{W}^* \in \mathbb{R}^{d \times d}$ whose diagonal entries are the same that we are observing, i.e., $w^* \triangleq w_{11}^* = w_{22}^* = \dots = w_{dd}^*$ and $\Omega_{\text{diag}}^{(d)} = \{(i, i)\}_{i=1}^d$. We factorize a solution matrix at time t as a product of L matrices,*

$$\mathbf{W}_{L:1}(t) = \mathbf{W}_L(t) \mathbf{W}_{L-1}(t) \dots \mathbf{W}_1(t), \quad \mathbf{W}_l(t) \in \mathbb{R}^{d \times d} \quad \text{for all } l \in [L].$$

Suppose that for all $l \in [L]$ and $0 \leq m \leq k$, the following holds:

$$\mathbf{W}_l^{(m)}(t) = x^{(m)} \mathbf{I}_d + y^{(m)} (\mathbf{J}_d - \mathbf{I}_d),$$

for some scalars $x^{(m)}, y^{(m)} \in \mathbb{R}$ where we denote $\mathbf{A}^{(k)}(t)$ as k -th derivative with respect to t of a matrix $\mathbf{A}(t)$. Then, the k -th derivative of the product $\mathbf{W}_{L:1}(t)$ satisfies

$$w_{11}^{(k)}(t) = w_{22}^{(k)}(t) = \dots = w_{dd}^{(k)}(t).$$

Proof Let us denote the m -th derivative of each layer matrix by

$$\mathbf{A}^{(m)} \triangleq \mathbf{W}_l^{(m)}(t).$$

Then, the k -th time derivative of the product $\mathbf{W}_{L:1}(t)$ is given by the Leibniz rule:

$$\frac{d^k}{dt^k} \mathbf{W}_{L:1}(t) = \sum_{k_1 + \dots + k_L = k} \binom{k}{k_1, \dots, k_L} \mathbf{A}^{(k_L)} \mathbf{A}^{(k_{L-1})} \dots \mathbf{A}^{(k_1)}.$$

By the assumption, each $\mathbf{A}^{(m)}$ lies in the span of $\{\mathbf{I}_d, \mathbf{J}_d\}$, and since this span is closed under matrix multiplication and scalar multiplication (by Lemma 10), each term in the sum lies in the same span. Hence, the entire sum $\mathbf{W}^{(k)}(t)$ also lies in $\text{span}\{\mathbf{I}_d, \mathbf{J}_d\}$, which implies that all diagonal entries of $\mathbf{W}^{(k)}(t)$ are equal. $\blacksquare$

Lemma 12 *Under the setting of Lemma 11 where each factor matrix $\mathbf{W}_l(0)$ is initialized according to (7), the following identities hold for all $k \in \mathbb{N} \cup \{0\}$ under the gradient flow dynamics defined in (3):*

$$\begin{aligned} \left(\mathbf{W}_{l_1}^{(k)}(0)\right)_{ii} &= \left(\mathbf{W}_{l_2}^{(k)}(0)\right)_{jj}, \quad i, j \in [d], l_1, l_2 \in [L], \\ \left(\mathbf{W}_{l_1}^{(k)}(0)\right)_{i_1 j_1} &= \left(\mathbf{W}_{l_2}^{(k)}(0)\right)_{i_2 j_2}, \quad i_1 \neq j_1, i_2 \neq j_2 \in [d], l_1, l_2 \in [L]. \end{aligned}$$

Proof For the base case, when $k = 0$, these identities immediately follow from our initialization assumptions. Now, suppose the induction hypothesis holds for all orders $m < k$ (with $k \geq 1$), which means we have:

$$\begin{aligned} \left(\mathbf{W}_{l_1}^{(m)}(0)\right)_{ii} &= \left(\mathbf{W}_{l_2}^{(m)}(0)\right)_{jj}, \quad i, j \in [d], l_1, l_2 \in [L], \\ \left(\mathbf{W}_{l_1}^{(m)}(0)\right)_{i_1 j_1} &= \left(\mathbf{W}_{l_2}^{(m)}(0)\right)_{i_2 j_2}, \quad i_1 \neq j_1, i_2 \neq j_2 \in [d], l_1, l_2 \in [L]. \end{aligned} \quad (22)$$

By applying the Leibniz rule to (21), the k -th derivative of $\mathbf{W}_l(t)$ is given by:

$$\mathbf{W}_l^{(k)}(t) = - \sum_{i_1, \dots, i_L} \binom{k-1}{i_1, \dots, i_L} \prod_{r=l+1}^L \mathbf{W}_r^{(i_r)}(t)^\top \cdot \nabla \ell(\mathbf{W}_{L:1}(t))^{(i_l)} \cdot \prod_{r=1}^{l-1} \mathbf{W}_r^{(i_r)}(t)^\top, \quad (23)$$

with $\sum_l i_l = k-1$ where each $i_l \geq 0$. Given our induction assumption in equation (22) for all $m < k$, let $x^{(m)}(0)$ denote the m -th derivative of the diagonal entries and $y^{(m)}(0)$ the m -th derivative of the off-diagonal entries at initialization. Note that at initialization, by Lemma 11, under the assumption that $\mathbf{W}_l^{(m)}(0)$ lies in the span of $\{\mathbf{I}_d, \mathbf{J}_d\}$ leads to $w_{11}^{(m)}(0) = w_{22}^{(m)}(0) \cdots = w_{dd}^{(m)}(0)$. Therefore, we know $\nabla \ell(\mathbf{W}_{L:1}(0))^{(i_l)} = r^{(i_l)}(0) \mathbf{I}_d$ for all $i_l < k$, where $r^{(i_l)}(0) \triangleq r_{11}^{(i_l)}(0) = \cdots = r_{dd}^{(i_l)}(0)$. Thus, at initialization, since equation (23) consists of terms involving $x^{(m)}(0)$ and $y^{(m)}(0)$ for all $m < k$, we can rewrite the above expression at $t = 0$ in terms of these derivatives as follows:

$$\begin{aligned} \mathbf{W}_l^{(k)}(0) &= - \sum_{i_1, \dots, i_L} \binom{k-1}{i_1, \dots, i_L} r^{(i_l)}(0) \prod_{r \in [L] \setminus \{l\}} \mathbf{W}_r^{(i_r)}(0) \\ &= - \sum_{i_1, \dots, i_L} \binom{k-1}{i_1, \dots, i_L} r^{(i_l)}(0) \prod_{r \in [L] \setminus \{l\}} (a_r \mathbf{I}_d + b_r \mathbf{J}_d), \end{aligned}$$

where constants a_r and b_r are composed of $x^{(r)}(0)$ and $y^{(r)}(0)$. Then, by Lemma 10, $\mathbf{W}_l^{(k)}(0)$ can be expressed in terms of only two values—one for the diagonal entries and one for the off-diagonal entries:

$$\mathbf{W}_l^{(k)}(0) = \alpha \mathbf{I}_d + \beta \mathbf{J}_d, \quad \alpha, \beta \in \mathbb{R},$$

thus concluding the proof. ■

Lemma 13 Under the setting of Lemma 12, the symmetries are preserved for all time $t \geq 0$:

$$\begin{aligned} (\mathbf{W}_{l_1}(t))_{ii} &= (\mathbf{W}_{l_2}(t))_{jj} \quad \text{for all } i, j \in [d], l_1, l_2 \in [L], \\ (\mathbf{W}_{l_1}(t))_{i_1 j_1} &= (\mathbf{W}_{l_2}(t))_{i_2 j_2} \quad \text{for all } i_1 \neq j_1, i_2 \neq j_2 \in [d], l_1, l_2 \in [L]. \end{aligned}$$

Proof By applying Lemma 37 to the result of Lemma 12, we can conclude that the symmetries are preserved for all timestep $t \geq 0$. $\blacksquare$

By the above lemmas, if the initialization follows the scheme in (7), then all diagonal entries of the all layers are identical, and all off-diagonal entries are also identical. Under this condition, the gradient flow dynamics can be easily described by the following lemma.

Lemma 14 Under the same conditions as in Lemma 12, if the diagonal entries of each layer are identical at timestep t (denoted by $x(t)$), and if the off-diagonal entries of each layer are identical at timestep t (denoted by $y(t)$), then the time derivative of $x(t)$ and $y(t)$ are given as:

$$\begin{aligned} \dot{x}(t) &= -\frac{(x(t) + (d-1)y(t))^{L-1} + (d-1)(x(t) - y(t))^{L-1}}{d} r(t), \\ \dot{y}(t) &= -\frac{(x(t) + (d-1)y(t))^{L-1} - (x(t) - y(t))^{L-1}}{d} r(t). \end{aligned}$$

Proof For $l \in [L]$ the gradient flow dynamics of $\mathbf{W}_l$ are written as:

$$\dot{\mathbf{W}}_l(t) = -\prod_{i=l+1}^L \mathbf{W}_i(t)^\top \cdot \nabla \ell(\mathbf{W}_{L:1}(t)) \cdot \prod_{i=1}^{l-1} \mathbf{W}_i(t)^\top, \quad (24)$$

where $\nabla \ell(\mathbf{W}_{L:1}(t)) = \text{diag}(r(t), \dots, r(t))$. Since $\mathbf{W}_l(t)$ is comprised of $x(t)$ in diagonal entries and $y(t)$ in off-diagonal entries, the above dynamics can be rewritten as follows:

$$\begin{aligned} \dot{\mathbf{W}}_l(t) &= -r(t) [\mathbf{W}_l(t)]^{L-l} \cdot \mathbf{I}_d \cdot [\mathbf{W}_l(t)]^{l-1} \\ &= -r(t) [\mathbf{W}_l(t)]^{L-1}. \end{aligned} \quad (25)$$

If we rewrite $\mathbf{W}_l(t) = (x(t) - y(t))\mathbf{I}_d + y(t)\mathbf{J}_d$, its eigenvalues are derived as:

$$\begin{aligned} \lambda_1 &= x(t) + (d-1)y(t) \quad \text{for the eigenvector } \mathbb{1}, \\ \lambda_2 &= x(t) - y(t) \quad \text{for any eigenvector orthogonal to } \mathbb{1} \text{ (multiplicity } d-1). \end{aligned}$$

Here, we denote $\lambda_i \triangleq \lambda_i(\mathbf{W}_{L:1}(t))$, unless otherwise specified. Then, we can decompose $\mathbf{W}_l(t)$ with projection matrix $\mathbf{P}_\parallel = \frac{1}{d}\mathbf{J}_d$ and $\mathbf{P}_\perp = \mathbf{I}_d - \frac{1}{d}\mathbf{J}_d$ as follows:

$$\mathbf{W}_l(t) = \lambda_1 \mathbf{P}_\parallel + \lambda_2 \mathbf{P}_\perp.$$

Therefore, if we take $(L-1)$ -th power of $\mathbf{W}_l(t)$, we can derive:

$$\begin{aligned} [\mathbf{W}_l(t)]^{L-1} &= \lambda_1^{L-1} \mathbf{P}_\parallel + \lambda_2^{L-1} \mathbf{P}_\perp \\ &= (x(t) + (d-1)y(t))^{L-1} \cdot \frac{1}{d} \mathbf{J}_d + (x(t) - y(t))^{L-1} \left(\mathbf{I}_d - \frac{1}{d} \mathbf{J}_d \right) \\ &= (x(t) - y(t))^{L-1} \mathbf{I}_d + \frac{(x(t) + (d-1)y(t))^{L-1} - (x(t) - y(t))^{L-1}}{d} \mathbf{J}_d. \end{aligned}$$

Recalling that $\mathbf{I}_d$ has 1 on the diagonal and 0 off-diagonal, and $\mathbf{J}_d$ has 1 in every entry, the entries of $[\mathbf{W}_l(t)]^{L-1}$ are:

$$\begin{aligned} ([\mathbf{W}_l(t)]^{L-1})_{ii} &= (x(t) - y(t))^{L-1} + \frac{(x(t) + (d-1)y(t))^{L-1} - (x(t) - y(t))^{L-1}}{d} \\ &= \frac{(x(t) + (d-1)y(t))^{L-1} + (d-1)(x(t) - y(t))^{L-1}}{d}, \quad \forall i \in [d], \end{aligned} \quad (26)$$

$$([\mathbf{W}_l(t)]^{L-1})_{ij} = \frac{(x(t) + (d-1)y(t))^{L-1} - (x(t) - y(t))^{L-1}}{d}, \quad \forall i \neq j \in [d]. \quad (27)$$

This concludes the proof by substituting the above equations to equation (25). $\blacksquare$

Under the gradient flow dynamics of the diagonal entry $x(t)$ and $y(t)$, we derive the dynamics of singular value of $\mathbf{W}_l(t)$.

Lemma 15 *Under the conditions of Lemma 12, the singular values of $\mathbf{W}_l(t)$ evolve according to:*

$$\dot{\sigma}_i(t) = -\sigma_i^{L-1}(t)r(t), \quad i = 1, 2, \dots, d,$$

Proof By Lemma 13, each factor matrix $\mathbf{W}_l(t)$ is symmetric, having $x(t)$ as its diagonal entries and $y(t)$ as its off-diagonal entries. The distinct eigenvalues of $\mathbf{W}_l(t)$ are $\lambda_1(t) = x(t) + (d-1)y(t)$ and $\lambda_2(t) = x(t) - y(t)$ (where $\lambda_2(t)$ has multiplicity $d-1$). Their time derivatives are calculated by:

$$\dot{\lambda}_i(t) = -\lambda_i^{L-1}(t)r(t),$$

Note that by setting $m > 1$, we have $\lambda_1(0) \geq \lambda_2(0) > 0$. If $L = 2$, the solution of above equation is equal to $\lambda_i(t) = \lambda_i(0) \exp\left(-\int_0^t r(\tau)d\tau\right)$, which means it maintains the positiveness of $\lambda_i(0)$ for all $t \geq 0$. For $L > 2$, its general solution can be written as follows:

$$\lambda_i(t) = \left(\lambda_i(0)^{2-L} + (L-2) \int_0^t r(\tau)d\tau \right)^{\frac{1}{2-L}},$$

due to its positiveness at initialization. Then, $\lambda_i(t)$ stays strictly positive, since it may diverge to $+\infty$, but it never reaches zero or changes the sign. Therefore, due to the symmetry and positive definiteness of $\mathbf{W}_l(t)$, we further conclude that $\lambda_i(t) \equiv \sigma_i(t)$. $\blacksquare$

By above lemma, we can solve ODE and find $\sigma_r(t)$ as follows:

$$\sigma_r(t) = \begin{cases} \sigma_r(0) \exp\left(-\int_0^t r(\tau)d\tau\right), & L = 2, \\ \left(\sigma_r(0)^{2-L} + (L-2) \cdot \int_0^t r(\tau)d\tau \right)^{\frac{1}{2-L}}, & L > 2. \end{cases}$$

Since $\sigma_1(0) = x(0) + (d-1)y(0) = \alpha(1 + \frac{d-1}{m})$ and $\sigma_r(0) = x(0) - y(0) = \alpha(1 - \frac{1}{m})$ for all $i \geq 2$, we can separate above equation as following:

$$\begin{aligned} \sigma_1(t) &= \begin{cases} \alpha(1 + \frac{d-1}{m}) \exp\left(-\int_0^t r(\tau) d\tau\right), & L = 2, \\ \left(\alpha^{2-L} \left(1 + \frac{d-1}{m}\right)^{2-L} + (L-2) \cdot \int_0^t r(\tau) d\tau\right)^{\frac{1}{2-L}}, & L > 2, \end{cases} \\ \sigma_r(t) &= \begin{cases} \alpha(1 - \frac{1}{m}) \exp\left(-\int_0^t r(\tau) d\tau\right), & L = 2, \\ \left(\alpha^{2-L} \left(1 - \frac{1}{m}\right)^{2-L} + (L-2) \cdot \int_0^t r(\tau) d\tau\right)^{\frac{1}{2-L}}, & L > 2. \end{cases}, \quad r = 2, 3, \dots, d. \end{aligned}$$

Then, we can establish a relationship between $\sigma_1(t)$ and $\sigma_r(t)$, thereby identifying an invariant property independent of time t :

- For $L = 2$:

$$\frac{\sigma_1(t)}{\sigma_r(t)} = \frac{m+d-1}{m-1}, \quad (28)$$

- For $L > 2$:

$$\sigma_1^{2-L}(t) - \sigma_r^{2-L}(t) = \alpha^{2-L} \left(\left(1 + \frac{d-1}{m}\right)^{2-L} - \left(1 - \frac{1}{m}\right)^{2-L} \right). \quad (29)$$

Furthermore, we can derive a closed-form solution for the singular values by utilizing the convergence guarantee. From equation (26), the diagonal entries of the solution matrix can be expressed as:

$$w_{ii}(t) = ([\mathbf{W}_l(t)]^L)_{ii} = \frac{(x(t) + (d-1)y(t))^L + (d-1)(x(t) - y(t))^L}{d}, \quad \forall i \in [d].$$

Since $w_{ii}(t)$ converges to a fixed value w^* , and noting that $\sigma_1(t) = x(t) + (d-1)y(t)$ and $\sigma_r(t) = x(t) - y(t)$, we obtain the following convergence equation:

$$w^* = \frac{\sigma_1^L(\infty) + (d-1)\sigma_r^L(\infty)}{d}. \quad (30)$$

Combining Equations (28) and (30), we derive a closed-form solution for the singular values of the depth-2 matrix as $t \rightarrow \infty$:

$$\begin{aligned} \sigma_1(\infty) &= (m+d-1) \sqrt{\frac{w^*}{m^2+d-1}}, \\ \sigma_r(\infty) &= (m-1) \sqrt{\frac{w^*}{m^2+d-1}}, \quad r = 2, 3, \dots, d. \end{aligned}$$

For the case when $L \geq 3$, we cannot obtain an exact analytical solution for $\sigma_r(\infty)$. Instead, we derive implicit equations for both $\sigma_1(\infty)$ and $\sigma_r(\infty)$ that cannot be easily solved without specifying numerical values:

$$\begin{aligned} \sigma_1^{2-L}(\infty) - \left(\frac{w^*d - \sigma_1^L(\infty)}{d-1} \right)^{\frac{2-L}{L}} &= C_{\alpha, m, L, d}, \\ (w^*d - (d-1)\sigma_r^L(\infty))^{\frac{2-L}{L}} - \sigma_r^{2-L}(\infty) &= C_{\alpha, m, L, d}, \quad \text{for } r = 2, \dots, d., \end{aligned}$$

where $C_{\alpha,m,L,d} \triangleq \left(\frac{\alpha}{m}\right)^{2-L} \left((m+d-1)^{2-L} - (m-1)^{2-L}\right)$. If we specify the values of $\alpha > 0, m > 1, d \geq 2, L \geq 3$ and $w^* > 0$ for ground-truth value, we can derive $\sigma_1(\infty)$ and $\sigma_r(\infty)$ of solution matrix of depth- L by substituting the values to above equations.

Furthermore, since $\mathbf{W}_{L:1}(t) = [\mathbf{W}_l(t)]^L$, when we diagonalize the factor matrix $\mathbf{W}_l(t)$ as $\mathbf{W}_l(t) = \mathbf{Q}\mathbf{\Lambda}(t)\mathbf{Q}^\top$, we have:

$$\mathbf{W}_{L:1}(t) = \mathbf{Q}\mathbf{\Lambda}^L(t)\mathbf{Q}^\top,$$

which indicates $\sigma_i(\mathbf{W}_{L:1}(t)) = (\sigma_i(\mathbf{W}_{L:1}(t)))^L$ for all $i \in [d]$. This concludes the proof of Theorem 4.

Remark. The $L \geq 3$ and $m = \infty$ case could arguably fall under the preceding analysis when other parameters are held fixed, as $m = \infty$ implies that all singular values are identical. However, a slight dependency on the specific value of α persists; for instance, tracking the overall result becomes challenging if α approaches zero while $m = \infty$. Therefore, we will restrict the scope of the aforementioned analysis to finite m . Consequently, the $L \geq 3$ and $m = \infty$ case will be analyzed separately in the following subsection.

F.3.2. CASE FOR $L \geq 3$ AND $m = \infty$

We now examine the $m = \infty$ case, which corresponds to an initialization scheme like $\mathbf{W}_l(0) = \alpha \mathbf{I}_d$. By Lemma 9, the factor matrices $\mathbf{W}_l(t)$ remain diagonal for all $t \geq 0$, and thus the diagonal entries of the product matrix are $w_{ii}(t) = (\mathbf{W}_L(t))_{ii}(\mathbf{W}_{L-1}(t))_{ii} \cdots (\mathbf{W}_1(t))_{ii}$. Assuming zero-loss convergence is achieved for any initial choice of $\alpha > 0$, it follows that $w_{ii}(\infty) = w^*$ for all i , and consequently, the overall matrix $\mathbf{W}_{L:1}(\infty)$ is diagonal with entries w^* .

Furthermore, let us consider the implications of Lemmas 11–13. These lemmas hold under a condition $y(t) = 0$, thereby belonging to $\text{span}\{\mathbf{I}_d, \mathbf{J}_d\}$, this leads to the result that each diagonal element of the factor matrices at convergence is $(\mathbf{W}_l(\infty))_{ii} = (w^*)^{1/L}$ for all $i \in [d]$ and $l \in [L]$. This means each layer $\mathbf{W}_l(\infty)$ becomes $(w^*)^{1/L} \mathbf{I}_d$, and thus has identical singular values equal to $(w^*)^{1/L}$ (assuming $w^* \geq 0$). This, in turn, leads to the final claim that for the overall product matrix $\mathbf{W}_{L:1}(\infty)$, its singular values $\sigma_i(\mathbf{W}_{L:1}(\infty))$ satisfy $(\sigma_i(\mathbf{W}_{L:1}(\infty)))^L = w^*$ for all $i \in [d]$.

F.3.3. LOSS CONVERGENCE

Actually, we can further show loss convergence under minor initialization assumption with the below proposition.

Proposition 16 *Under the conditions of Theorem 4, assume the initialization scale α is set such that $w_{ii}(0) \leq w^*$, specifically,*

$$0 < \alpha^L \leq \frac{w^* dm^L}{(m+d-1)^L + (d-1)(m-1)^L}.$$

Denoting $\ell(t) \triangleq \ell(\mathbf{W}_{L:1}(t); \Omega_{\text{diag}}^{(d)})$, loss convergence is guaranteed for any $L \geq 2, m > 1$, and $d \geq 2$:

$$\ell(t) \leq \ell(0) \exp \left(\frac{-2L\alpha^{2L-2} ((d-1)(m-1)^{2L-2} + (m+d-1)^{2L-2})}{dm^{2L-2}} t \right).$$

Proof Recall that the eigenvalues are given by $\lambda_1(t) = x(t) + (d-1)y(t)$ and $\lambda_2(t) = x(t) - y(t)$. From Lemma 14, their time derivatives are

$$\begin{aligned}\dot{\lambda}_1(t) &= -\lambda_1^{L-1}(t)r(t), \\ \dot{\lambda}_2(t) &= -\lambda_2^{L-1}(t)r(t).\end{aligned}$$

The diagonal entries $w_{ii}(t)$ of $\mathbf{W}_{L:1}(t)$ can be expressed in terms of these eigenvalues as

$$\begin{aligned}w_{ii}(t) &= \frac{(x(t) + (d-1)y(t))^L + (d-1)(x(t) - y(t))^L}{d} \\ &= \frac{\lambda_1^L(t) + (d-1)\lambda_2^L(t)}{d}.\end{aligned}$$

Let the residual be defined as $r(t) = w_{ii}(t) - w^*$, where w^* is a constant. The time derivative of $r(t)$ can then be calculated by applying the chain rule to $r(t)$'s definition and subsequently substituting the expressions for $\dot{\lambda}_1(t)$ and $\dot{\lambda}_2(t)$:

$$\begin{aligned}\dot{r}(t) &= \frac{d}{dt} (w_{ii}(t) - w^*) \\ &= \frac{L}{d} \lambda_1^{L-1}(t) \dot{\lambda}_1(t) + \frac{L(d-1)}{d} \lambda_2^{L-1}(t) \dot{\lambda}_2(t) \\ &= \frac{L}{d} \lambda_1^{L-1}(t) \left(-\lambda_1^{L-1}(t)r(t) \right) + \frac{L(d-1)}{d} \lambda_2^{L-1}(t) \left(-\lambda_2^{L-1}(t)r(t) \right) \\ &= - \left(\underbrace{\frac{L}{d} \lambda_1^{2L-2}(t) + \frac{L(d-1)}{d} \lambda_2^{2L-2}(t)}_{\triangleq K(t)} \right) r(t).\end{aligned}\tag{31}$$

This gives the first-order linear ordinary differential equation $\dot{r}(t) = -K(t)r(t)$. The general solution for $r(t)$ is therefore

$$r(t) = r(0) \exp \left(- \int_0^t K(\tau) d\tau \right).\tag{32}$$

Given the initial condition provided, $r(0) \leq 0$, and $r(t)$ maintains the same sign as $r(0)$. Therefore, $r(t) \leq 0$ for all $t \geq 0$. The choice of $m > 1$ ensures that $\lambda_1(0), \lambda_2(0) > 0$. With $r(t) \leq 0$ and $\lambda_i(0) > 0$, and assuming $L \geq 1$, the derivatives of the eigenvalues satisfy

$$\begin{aligned}\dot{\lambda}_1(t) &= \frac{d}{dt} (x(t) + (d-1)y(t)) = -(x(t) + (d-1)y(t))^{L-1} r(t) \geq 0, \\ \dot{\lambda}_2(t) &= \frac{d}{dt} (x(t) - y(t)) = -(x(t) - y(t))^{L-1} r(t) \geq 0,\end{aligned}$$

because $\lambda_i^{L-1}(t) \geq 0$ and $r(t) \leq 0$. These non-negative derivatives imply that $\lambda_1(t)$ and $\lambda_2(t)$ are monotonically non-decreasing. Since they start from positive values, they remain positive for all $t \geq 0$, i.e., $\lambda_i(t) \geq \lambda_i(0) > 0$.

With above lower bound for $\lambda_i(t)$, which is $\lambda_i(0)$, we can upper bound the $K(t)$ as follows:

$$\begin{aligned} K(t) &\geq \frac{L}{d} \lambda_1^{2L-2}(0) + \frac{L(d-1)}{d} \lambda_2^{2L-2}(0) \\ &= \frac{L\alpha^{2L-2} ((d-1)(m-1)^{2L-2} + (m+d-1)^{2L-2})}{dm^{2L-2}}. \end{aligned} \quad (33)$$

By upper bounding the absolute value of (32) with (33), we derive:

$$|r(t)| \leq |r(0)| \exp(-K(0)t),$$

and since $\ell(\mathbf{W}_{L:1}(t)) = \frac{d}{2} r^2(t)$, this leads to loss convergence:

$$\ell(\mathbf{W}_{L:1}(t)) \leq \ell(\mathbf{W}_{L:1}(0)) \exp(-2K(0)t).$$

■

Appendix G. Proof for Section A

In this section, we provide the proofs for the propositions and theorems presented in Section A. First, Subsection G.1 presents the general form of Proposition 5 along with its proof. Next, Subsection G.2 details the proof of Theorem 6, focusing on the 2×2 matrix case. Lastly, Subsection G.2.2 generalizes the core ideas to $d \times d$ matrices and provides the proof of Theorem 6 for this general $d \times d$ setting.

G.1. General Form and Proof of Proposition 5

We first present the general form of Proposition 5. This proposition is applicable to any “fully disconnected case”, a scenario that involves the diagonal entries introduced within this same proposition.

For a $d \times d$ ground truth matrix $\mathbf{W}^*$, the observed entries are given by $\Omega = \{(i_n, j_n)\}_{n=1}^d$. Since we consider fully disconnected case, $i_n \neq i_m, j_n \neq j_m$ for all $n, m \in [d]$. We factorize the solution model at time t as $\mathbf{W}_{\mathbf{A}, \mathbf{B}}(t) = \mathbf{A}(t)\mathbf{B}(t)$, where $\mathbf{W}_{\mathbf{A}, \mathbf{B}}(t), \mathbf{A}(t), \mathbf{B}(t) \in \mathbb{R}^{d \times d}$. We consider the gradient flow dynamics with the loss function defined as in (2).

For a given row index k , since there exists a unique entry $(k, j) \in \Omega$, we denote this unique column index by $j^{(k)}$. Thus, $w_{k, j^{(k)}}^*$ and $w_{k, j^{(k)}}(t)$ refer to the ground truth weight $w_{k, j}^*$ and the time-varying weight $w_{k, j}(t)$ respectively, where $j = j^{(k)}$. Similarly, for a given column index l , since there exists a unique entry $(i, l) \in \Omega$, we denote this unique row index by $i^{(l)}$. Thus $w_{i^{(l)}, l}^*$ and $w_{i^{(l)}, l}(t)$ refer to the ground truth weight $w_{i, l}^*$ and the time-varying weight $w_{i, l}(t)$ respectively, where $i = i^{(l)}$. Defining the residuals as $r_{ij}(t) := w_{ij}^* - w_{ij}(t)$, we adopt this compact notation for residuals as well. Then, we can derive a closed-form solution for *arbitrary initialization* with below proposition.

Proposition 17 *Consider a ground truth matrix $\mathbf{W}^* \in \mathbb{R}^{d \times d}$ and a set of d fully disconnected observations $\Omega = \{(i_n, j_n)\}_{n=1}^d$. The model is factorized as $\mathbf{W}_{\mathbf{A}, \mathbf{B}}(t) = \mathbf{A}(t)\mathbf{B}(t)$, where the factors $\mathbf{A}(t), \mathbf{B}(t) \in \mathbb{R}^{d \times d}$. For each observed pair $(i_n, j_n) \in \Omega$, define the constants P_{i_n, j_n} and Q_{i_n, j_n} based on the initial values $\mathbf{A}(0)$ and $\mathbf{B}(0)$:*

$$P_{i_n, j_n} \triangleq \sum_{k=1}^d a_{i_n, k}(0)b_{k, j_n}(0) \quad \text{and} \quad Q_{i_n, j_n} \triangleq \sum_{k=1}^d (a_{i_n, k}(0)^2 + b_{k, j_n}(0)^2).$$

Furthermore, for each such observed pair (i_n, j_n) , let the parameter $\bar{r}_{i_n, j_n}$ be determined from the ground truth entry w_{i_n, j_n}^* and the constants defined above, as follows:

$$\bar{r}_{i_n, j_n} \triangleq \frac{1}{2} \log \left(\frac{P_{i_n, j_n} + \frac{Q_{i_n, j_n}}{2}}{w_{i_n, j_n}^* + \sqrt{w_{i_n, j_n}^{*2} - P_{i_n, j_n}^2 + \left(\frac{Q_{i_n, j_n}}{2}\right)^2}} \right).$$

Then, assuming convergence to a zero-loss solution (i.e., $w_{i_n, j_n}(\infty) = w_{i_n, j_n}^*$ for all $(i_n, j_n) \in \Omega$), any entry $a_{p, q}(\infty)$ of the converged matrix $\mathbf{A}(\infty)$ and any entry $b_{p, q}(\infty)$ of the converged matrix $\mathbf{B}(\infty)$ (for arbitrary indices $p, q \in [d]$) are explicitly given by:

$$\begin{aligned} a_{p, q}(\infty) &= a_{p, q}(0) \cosh \left(\bar{r}_{p, j^{(p)}} \right) - b_{q, j^{(p)}}(0) \sinh \left(\bar{r}_{p, j^{(p)}} \right), \\ b_{p, q}(\infty) &= b_{p, q}(0) \cosh \left(\bar{r}_{i^{(q)}, q} \right) - a_{i^{(q)}, p}(0) \sinh \left(\bar{r}_{i^{(q)}, q} \right). \end{aligned}$$

Proof We can express their evolution in the following vector form using the vectorized parameter

$$\boldsymbol{\theta}(t) := \begin{bmatrix} \text{vec}(\mathbf{A}(t)) \\ \text{vec}(\mathbf{B}(t)) \end{bmatrix} \in \mathbb{R}^{2d^2 \times 2d^2}.$$

$$\dot{\boldsymbol{\theta}}(t) = - \begin{bmatrix} \mathbf{0}_{d^2, d^2} & \mathbf{R}(t) \\ \mathbf{R}(t)^\top & \mathbf{0}_{d^2, d^2} \end{bmatrix} \boldsymbol{\theta}(t) \quad (34)$$

where $\mathbf{R}(t) \in \mathbb{R}^{d^2 \times d^2}$ is defined as:

$$\mathbf{R}(t) = \begin{bmatrix} r_{1,j^{(1)}}(t) \mathbf{e}_{j^{(1)}}^\top \\ r_{1,j^{(1)}}(t) \mathbf{e}_{j^{(1)}+d}^\top \\ \vdots \\ r_{1,j^{(1)}}(t) \mathbf{e}_{j^{(1)}+(d-1)d}^\top \\ r_{2,j^{(2)}}(t) \mathbf{e}_{j^{(2)}}^\top \\ r_{2,j^{(2)}}(t) \mathbf{e}_{j^{(2)}+d}^\top \\ \vdots \\ r_{d,j^{(d)}}(t) \mathbf{e}_{j^{(d)}+(d-1)d}^\top \end{bmatrix} \quad (35)$$

for $\mathbf{e}_i \in \mathbb{R}^{d^2}$ form the standard basis. Since $\begin{bmatrix} \mathbf{0}_{d^2, d^2} & \mathbf{R}(t) \\ \mathbf{R}(t)^\top & \mathbf{0}_{d^2, d^2} \end{bmatrix}$ commutes with any other t values, the solution is given as:

$$\boldsymbol{\theta}(t) = \exp \left(- \int_0^t \begin{bmatrix} \mathbf{0}_{d^2, d^2} & \mathbf{R}(\tau) \\ \mathbf{R}(\tau)^\top & \mathbf{0}_{d^2, d^2} \end{bmatrix} d\tau \right) \cdot \boldsymbol{\theta}(0) \quad (36)$$

$$= \exp \left(- \begin{bmatrix} \mathbf{0}_{d^2, d^2} & \bar{\mathbf{R}}(t) \\ \bar{\mathbf{R}}(t)^\top & \mathbf{0}_{d^2, d^2} \end{bmatrix} \right) \cdot \boldsymbol{\theta}(0) \quad (37)$$

where

$$\bar{\mathbf{R}}(t) := \int_0^t \mathbf{R}(\tau) d\tau = \begin{bmatrix} \bar{r}_{1,j^{(1)}}(t) \mathbf{e}_{j^{(1)}}^\top \\ \bar{r}_{1,j^{(1)}}(t) \mathbf{e}_{j^{(1)}+d}^\top \\ \vdots \\ \bar{r}_{1,j^{(1)}}(t) \mathbf{e}_{j^{(1)}+(d-1)d}^\top \\ \bar{r}_{2,j^{(2)}}(t) \mathbf{e}_{j^{(2)}}^\top \\ \bar{r}_{2,j^{(2)}}(t) \mathbf{e}_{j^{(2)}+d}^\top \\ \vdots \\ \bar{r}_{d,j^{(d)}}(t) \mathbf{e}_{j^{(d)}+(d-1)d}^\top \end{bmatrix}$$

for $\bar{r}_{i,j}(t) = \int_0^t r_{i,j}(\tau) d\tau$. If we assume convergence, we get:

$$\boldsymbol{\theta}(\infty) = \exp \left(- \begin{bmatrix} \mathbf{0}_{d^2, d^2} & \bar{\mathbf{R}}(\infty) \\ \bar{\mathbf{R}}(\infty)^\top & \mathbf{0}_{d^2, d^2} \end{bmatrix} d\tau \right) \cdot \boldsymbol{\theta}(0) \quad (38)$$

$$= \left(\begin{bmatrix} \mathbf{I}_{d^2} & \mathbf{0}_{d^2, d^2} \\ \mathbf{0}_{d^2, d^2} & \mathbf{I}_{d^2} \end{bmatrix} - \begin{bmatrix} \mathbf{0}_{d^2, d^2} & \bar{\mathbf{R}}(t) \\ \bar{\mathbf{R}}(t)^\top & \mathbf{0}_{d^2, d^2} \end{bmatrix} + \frac{1}{2} \begin{bmatrix} \bar{\mathbf{R}}(t) \bar{\mathbf{R}}(t)^\top & \mathbf{0}_{d^2, d^2} \\ \mathbf{0}_{d^2, d^2} & \bar{\mathbf{R}}(t)^\top \bar{\mathbf{R}}(t) \end{bmatrix} \right. \quad (39)$$

$$\left. - \frac{1}{6} \begin{bmatrix} \mathbf{0}_{d^2, d^2} & \bar{\mathbf{R}}(t) \bar{\mathbf{R}}(t)^\top \bar{\mathbf{R}}(t) \\ \bar{\mathbf{R}}(t)^\top \bar{\mathbf{R}}(t) \bar{\mathbf{R}}(t)^\top & \mathbf{0}_{d^2, d^2} \end{bmatrix} + \frac{1}{24} \begin{bmatrix} (\bar{\mathbf{R}}(t) \bar{\mathbf{R}}(t)^\top)^2 & \mathbf{0}_{d^2, d^2} \\ \mathbf{0}_{d^2, d^2} & (\bar{\mathbf{R}}(t)^\top \bar{\mathbf{R}}(t))^2 \end{bmatrix} \right. \quad (40)$$

$$\left. - \dots \right) \cdot \boldsymbol{\theta}(0), \quad (41)$$

which can be simplified as:

$$\boldsymbol{\theta}(\infty) = \begin{bmatrix} \mathbf{C} & \mathbf{D} \\ \mathbf{E} & \mathbf{F} \end{bmatrix} \boldsymbol{\theta}(0), \quad (42)$$

with $\mathbf{C}$, $\mathbf{D}$, $\mathbf{E}$ and $\mathbf{F}$ are defined as following:

$$\begin{aligned} \mathbf{C} &= \cosh \left(\text{diag} \left(\bar{r}_{1,j(1)}, \dots, \bar{r}_{1,j(1)}, \bar{r}_{2,j(2)}, \dots, \bar{r}_{2,j(2)}, \dots, \bar{r}_{d,j(d)}, \dots, \bar{r}_{d,j(d)} \right) \right), \\ \mathbf{F} &= \cosh \left(\text{diag} \left(\bar{r}_{i(1),1}, \bar{r}_{i(2),2}, \dots, \bar{r}_{i(d),d}, \dots, \bar{r}_{i(1),1}, \bar{r}_{i(2),2}, \dots, \bar{r}_{i(d),d} \right) \right), \\ \mathbf{D} &= -\sinh \left(\begin{bmatrix} \bar{r}_{1,j(1)} \mathbf{e}_{j(1)}^\top, \dots, \bar{r}_{1,j(1)} \mathbf{e}_{j(1)+(d-1)d}^\top, \dots, \bar{r}_{d,j(d)} \mathbf{e}_{j(d)}^\top, \dots, \bar{r}_{d,j(d)} \mathbf{e}_{j(d)+(d-1)d}^\top \end{bmatrix}^\top \right), \\ \mathbf{E} &= -\sinh \left(\begin{bmatrix} \bar{r}_{1,j(1)} \mathbf{e}_{j(1)}, \dots, \bar{r}_{1,j(1)} \mathbf{e}_{j(1)+(d-1)d}, \dots, \bar{r}_{d,j(d)} \mathbf{e}_{j(d)}, \dots, \bar{r}_{d,j(d)} \mathbf{e}_{j(d)+(d-1)d} \end{bmatrix} \right). \end{aligned}$$

Here, for any matrix $\mathbf{P}$, the operations $\cosh(\mathbf{P})$ and $\sinh(\mathbf{P})$ are performed elementwise. For a set of d observed indices Ω , there exists d corresponding unknown variables, $\bar{r}_{i_k, j_k}$. If convergence is guaranteed, the model yields d equations relating these variables to the d ground truth values. This implies that the variables $\bar{r}_{i_k, j_k}$ can be characterized as a closed-form. To characterize more

rigorously, we substitute C , D , E , and F into (42):

$$\theta(\infty) = \begin{bmatrix} a_{1,1}(\infty) \\ a_{1,2}(\infty) \\ \vdots \\ a_{1,d}(\infty) \\ a_{2,1}(\infty) \\ a_{2,2}(\infty) \\ \vdots \\ a_{2,d}(\infty) \\ \vdots \\ a_{d,1}(\infty) \\ \vdots \\ a_{d,d}(\infty) \\ b_{1,1}(\infty) \\ b_{1,2}(\infty) \\ \vdots \\ b_{1,d}(\infty) \\ b_{2,1}(\infty) \\ b_{2,2}(\infty) \\ \vdots \\ b_{2,d}(\infty) \\ \vdots \\ b_{d,1}(\infty) \\ \vdots \\ b_{d,d}(\infty) \end{bmatrix} = \begin{bmatrix} a_{1,1}(0) \cosh(\bar{r}_{1,j(1)}) - b_{1,j(1)}(0) \sinh(\bar{r}_{1,j(1)}) \\ a_{1,2}(0) \cosh(\bar{r}_{1,j(1)}) - b_{2,j(1)}(0) \sinh(\bar{r}_{1,j(1)}) \\ \vdots \\ a_{1,d}(0) \cosh(\bar{r}_{1,j(1)}) - b_{d,j(1)}(0) \sinh(\bar{r}_{1,j(1)}) \\ a_{2,1}(0) \cosh(\bar{r}_{2,j(2)}) - b_{1,j(2)}(0) \sinh(\bar{r}_{2,j(2)}) \\ a_{2,2}(0) \cosh(\bar{r}_{2,j(2)}) - b_{2,j(2)}(0) \sinh(\bar{r}_{2,j(2)}) \\ \vdots \\ a_{2,d}(0) \cosh(\bar{r}_{2,j(2)}) - b_{d,j(2)}(0) \sinh(\bar{r}_{2,j(2)}) \\ \vdots \\ a_{d,1}(0) \cosh(\bar{r}_{d,j(d)}) - b_{1,j(d)}(0) \sinh(\bar{r}_{d,j(d)}) \\ \vdots \\ a_{d,d}(0) \cosh(\bar{r}_{d,j(d)}) - b_{d,j(d)}(0) \sinh(\bar{r}_{d,j(d)}) \\ -a_{i(1),1}(0) \sinh(\bar{r}_{i(1),1}) + b_{1,1}(0) \cosh(\bar{r}_{i(1),1}) \\ -a_{i(2),1}(0) \sinh(\bar{r}_{i(2),2}) + b_{1,2}(0) \cosh(\bar{r}_{i(2),2}) \\ \vdots \\ -a_{i(d),1}(0) \sinh(\bar{r}_{i(d),d}) + b_{1,d}(0) \cosh(\bar{r}_{i(d),d}) \\ -a_{i(1),2}(0) \sinh(\bar{r}_{i(1),1}) + b_{2,1}(0) \cosh(\bar{r}_{i(1),1}) \\ -a_{i(2),2}(0) \sinh(\bar{r}_{i(2),2}) + b_{2,2}(0) \cosh(\bar{r}_{i(2),2}) \\ \vdots \\ -a_{i(d),2}(0) \sinh(\bar{r}_{i(d),d}) + b_{2,d}(0) \cosh(\bar{r}_{i(d),d}) \\ \vdots \\ -a_{i(1),d}(0) \sinh(\bar{r}_{i(1),1}) + b_{d,1}(0) \cosh(\bar{r}_{i(1),1}) \\ \vdots \\ -a_{i(d),d}(0) \sinh(\bar{r}_{i(d),d}) + b_{d,d}(0) \cosh(\bar{r}_{i(d),d}) \end{bmatrix}. \quad (43)$$

Then, assuming convergence, for each observation $(i_n, j_n) \in \Omega$ (for $n = 1, \dots, d$), we obtain the equation:

$$\begin{aligned} w_{i_n, j_n}^* &= w_{i_n, j_n}(\infty) = a_{i_n, 1}(\infty) b_{1, j_n}(\infty) + \dots + a_{i_n, d}(\infty) b_{d, j_n}(\infty) \\ &= \sum_{k=1}^d \left[\left(a_{i_n, k}(0) \cosh(\bar{r}_{i_n, j_n}) - b_{k, j(i_n)}(0) \sinh(\bar{r}_{i_n, j_n}) \right) \right. \\ &\quad \left. \cdot (b_{k, j_n}(0) \cosh(\bar{r}_{i_n, j_n}) - a_{i_n, k}(0) \sinh(\bar{r}_{i_n, j_n})) \right]. \end{aligned}$$

Let $C_n = \cosh(\bar{r}_{i_n, j_n})$ and $S_n = \sinh(\bar{r}_{i_n, j_n})$. Then we can rewrite above equation as:

$$\begin{aligned}
 w_{i_n, j_n}^* &= \sum_{k=1}^d (a_{i_n, k}(0)b_{k, j_n}(0)C_n^2 - a_{i_n, k}(0)^2C_nS_n - b_{k, j_n}(0)^2C_nS_n + a_{i_n, k}(0)b_{k, j_n}(0)S_n^2) \\
 &= \left(\sum_{k=1}^d a_{i_n, k}(0)b_{k, j_n}(0) \right) (C_n^2 + S_n^2) - \left(\sum_{k=1}^d (a_{i_n, k}(0)^2 + b_{k, j_n}(0)^2) \right) C_nS_n \\
 &= P_{i_n, j_n} \cosh(2\bar{r}_{i_n, j_n}) - \frac{Q_{i_n, j_n}}{2} \sinh(2\bar{r}_{i_n, j_n}), \tag{44}
 \end{aligned}$$

where $P_{i_n, j_n} = \sum_{k=1}^d a_{i_n, k}(0)b_{k, j_n}(0)$ and $Q_{i_n, j_n} = \sum_{k=1}^d (a_{i_n, k}(0)^2 + b_{k, j_n}(0)^2)$.

By solving (44) with respect to $\bar{r}_{i_n, j_n}$, we can get:

$$\begin{aligned}
 2w_{i_n, j_n}^* &= P_{i_n, j_n} (e^{2\bar{r}_{i_n, j_n}} + e^{-2\bar{r}_{i_n, j_n}}) - \frac{Q_{i_n, j_n}}{2} (e^{2\bar{r}_{i_n, j_n}} - e^{-2\bar{r}_{i_n, j_n}}) \\
 &= e^{2\bar{r}_{i_n, j_n}} \left(P_{i_n, j_n} - \frac{Q_{i_n, j_n}}{2} \right) + e^{-2\bar{r}_{i_n, j_n}} \left(P_{i_n, j_n} + \frac{Q_{i_n, j_n}}{2} \right).
 \end{aligned}$$

Multiply by $e^{2\bar{r}_{i_n, j_n}}$ leads to:

$$2w_{i_n, j_n}^* e^{2\bar{r}_{i_n, j_n}} = e^{4\bar{r}_{i_n, j_n}} \left(P_{i_n, j_n} - \frac{Q_{i_n, j_n}}{2} \right) + P_{i_n, j_n} + \frac{Q_{i_n, j_n}}{2}.$$

Rearrange into a quadratic equation by setting $u = e^{2\bar{r}_{i_n, j_n}}$:

$$\left(P_{i_n, j_n} - \frac{Q_{i_n, j_n}}{2} \right) u^2 - 2w_{i_n, j_n}^* u + P_{i_n, j_n} + \frac{Q_{i_n, j_n}}{2} = 0.$$

By solving the above equation while noting that $P_{i_n, j_n} - \frac{Q_{i_n, j_n}}{2} \leq 0$ by the definition, we can get explicit solutions for $\bar{r}_{i_n, j_n}$:

$$\bar{r}_{i_n, j_n} = \frac{1}{2} \log \left(\frac{P_{i_n, j_n} + \frac{Q_{i_n, j_n}}{2}}{w_{i_n, j_n}^* + \sqrt{w_{i_n, j_n}^{*2} - P_{i_n, j_n}^2 + \left(\frac{Q_{i_n, j_n}}{2} \right)^2}} \right).$$

Note that each $\bar{r}_{i_n, j_n}$ is solely determined by the initial points $\theta(0)$. With $\bar{r}_{i_n, j_n}$ determined for each observed entry, we have closed-form expressions characterizing the model's learned relationship for these observations. Consequently, by (43), we have:

$$\begin{aligned}
 a_{p, q}(\infty) &= a_{p, q}(0) \cosh(\bar{r}_{p, j(p)}) - b_{q, j(p)}(0) \sinh(\bar{r}_{p, j(p)}), \\
 b_{p, q}(\infty) &= b_{p, q}(0) \cosh(\bar{r}_{i(q), q}) - a_{i(q), p}(0) \sinh(\bar{r}_{i(q), q}).
 \end{aligned}$$

■

G.2. Proof of Theorem 6

In this section, we will provide the analysis of 2×2 matrix that starting from pre-trained weights with diagonal observations $w^* \triangleq w_{11}^* = w_{22}^*$, $\mathbf{W}_{A,B}(t)$ cannot converge to a low-rank solution. Let $T_1 > t_1$ be the timestep that concludes the pre-train phase. For the sake of simplicity, we omit the ϵ term introduced in the pre-train phase. Then, we know from Proposition 17, we have:

$$\mathbf{A}(T_1) = \mathbf{B}(T_1) = \begin{pmatrix} \sqrt{w^*} & 0 \\ 0 & \sqrt{w^*} \end{pmatrix}. \quad (45)$$

In the post-train phase, we introduce an additional observation in the off-diagonal entries, specifically w_{12}^* or w_{21}^* . Without loss of generality, we assume $w_{12}^* > 0$ is revealed while other observations remain the same, i.e., $\Omega_{\text{post}} = \{(1, 1), (1, 2), (2, 2)\}$. Note that the gradient of the post-train loss is:

$$\begin{aligned} \nabla \ell(\mathbf{W}_{A,B}) &= \begin{pmatrix} w_{11} - w^* & w_{12} - w_{12}^* \\ 0 & w_{22} - w^* \end{pmatrix} \\ &= \begin{pmatrix} a_{11}b_{11} + a_{12}b_{21} - w^* & a_{11}b_{12} + a_{12}b_{22} - w_{12}^* \\ 0 & a_{21}b_{12} + a_{22}b_{22} - w^* \end{pmatrix}. \end{aligned}$$

For simplicity, we again omit the Ω term in the loss specification. We define the residuals for the relevant matrix elements as $r_{11} := w_{11} - w^*$, $r_{12} := w_{12} - w_{12}^*$, and $r_{22} := w_{22} - w^*$.

We begin by demonstrating a pairwise symmetry between the entries of $\mathbf{A}(t)$ and $\mathbf{B}(t)$, which simplifies subsequent analysis. To this end, we first provide the time derivatives for the elements of $\mathbf{A}(t)$ and $\mathbf{B}(t)$. Given the general gradient flow dynamics $\dot{\mathbf{A}}(t) = -\nabla \ell(\mathbf{W}_{A,B}(t))\mathbf{B}^\top(t)$ and $\dot{\mathbf{B}}(t) = -\mathbf{A}^\top(t)\nabla \ell(\mathbf{W}_{A,B}(t))$, the component-wise updates are as follows. For $\mathbf{A}(t)$:

$$\begin{aligned} \dot{a}_{11}(t) &= b_{11}(t)(w^* - w_{11}(t)) + b_{12}(t)(w_{12}^* - w_{12}(t)), \\ \dot{a}_{12}(t) &= b_{21}(t)(w^* - w_{11}(t)) + b_{22}(t)(w_{12}^* - w_{12}(t)), \\ \dot{a}_{21}(t) &= b_{12}(t)(w^* - w_{22}(t)), \\ \dot{a}_{22}(t) &= b_{22}(t)(w^* - w_{22}(t)), \end{aligned} \quad (46)$$

and for $\mathbf{B}(t)$:

$$\begin{aligned} \dot{b}_{11}(t) &= a_{11}(t)(w^* - w_{11}(t)), \\ \dot{b}_{12}(t) &= a_{11}(t)(w_{12}^* - w_{12}(t)) + a_{21}(t)(w^* - w_{22}(t)), \\ \dot{b}_{21}(t) &= a_{12}(t)(w^* - w_{11}(t)), \\ \dot{b}_{22}(t) &= a_{12}(t)(w_{12}^* - w_{12}(t)) + a_{22}(t)(w^* - w_{22}(t)). \end{aligned} \quad (47)$$

Using the equations above, we first present a result showing that the k -th derivative of each element in $\mathbf{A}(t)$ and $\mathbf{B}(t)$ at initialization exhibits a pairwise symmetry:

Lemma 18 *Let $\mathbf{W}_{A,B}(T_1) = \mathbf{A}(T_1)\mathbf{B}(T_1) \in \mathbb{R}^{2 \times 2}$ be a product matrix, where $\mathbf{A}(T_1)$ and $\mathbf{B}(T_1)$ are matrices that are obtained at the end of the pre-training phase. Suppose the ground truth matrix satisfies $w_{11}^* = w_{22}^*$. Then for every $k \in \mathbb{N} \cup \{0\}$, the following identities hold:*

$$\begin{aligned} a_{11}^{(k)}(T_1) &= b_{22}^{(k)}(T_1), & a_{12}^{(k)}(T_1) &= b_{12}^{(k)}(T_1), \\ a_{21}^{(k)}(T_1) &= b_{21}^{(k)}(T_1), & a_{22}^{(k)}(T_1) &= b_{11}^{(k)}(T_1), \end{aligned} \quad (48)$$

and consequently,

$$w_{11}^{(k)}(T_1) = w_{22}^{(k)}(T_1). \quad (49)$$

Proof We prove the statement by induction on k . When $k = 0$, by the initialization assumption, we have

$$a_{11}(T_1) = b_{22}(T_1), \quad a_{12}(T_1) = b_{12}(T_1), \quad a_{21}(T_1) = b_{21}(T_1), \quad a_{22}(T_1) = b_{11}(T_1),$$

and therefore $w_{11}(T_1) = w_{22}(T_1)$.

Assume that for all orders $m < k$ (with $k \geq 1$) the identities

$$a_{11}^{(m)}(T_1) = b_{22}^{(m)}(T_1), \quad a_{12}^{(m)}(T_1) = b_{12}^{(m)}(T_1), \quad a_{21}^{(m)}(T_1) = b_{21}^{(m)}(T_1), \quad a_{22}^{(m)}(T_1) = b_{11}^{(m)}(T_1),$$

hold, and hence also $w_{11}^{(m)}(T_1) = w_{22}^{(m)}(T_1)$. By the Leibniz rule, each element of the k -th derivative can be written as a finite sum involving derivatives of orders strictly less than k . For $\mathbf{A}(t)$:

$$\begin{aligned} a_{11}^{(k)}(t) &= - \sum_{j=0}^{k-1} \binom{k-1}{j} \left(b_{11}^{(k-1-j)}(t) r_{11}^{(j)}(t) + b_{12}^{(k-1-j)}(t) r_{12}^{(j)}(t) \right), \\ a_{12}^{(k)}(t) &= - \sum_{j=0}^{k-1} \binom{k-1}{j} \left(b_{21}^{(k-1-j)}(t) r_{11}^{(j)}(t) + b_{22}^{(k-1-j)}(t) r_{12}^{(j)}(t) \right), \\ a_{21}^{(k)}(t) &= - \sum_{j=0}^{k-1} \binom{k-1}{j} b_{12}^{(k-1-j)}(t) r_{22}^{(j)}(t), \\ a_{22}^{(k)}(t) &= - \sum_{j=0}^{k-1} \binom{k-1}{j} b_{22}^{(k-1-j)}(t) r_{22}^{(j)}(t), \end{aligned}$$

and for $\mathbf{B}(t)$:

$$\begin{aligned} b_{11}^{(k)}(t) &= - \sum_{j=0}^{k-1} \binom{k-1}{j} a_{11}^{(k-1-j)}(t) r_{11}^{(j)}(t), \\ b_{12}^{(k)}(t) &= - \sum_{j=0}^{k-1} \binom{k-1}{j} \left(a_{11}^{(k-1-j)}(t) r_{12}^{(j)}(t) + a_{21}^{(k-1-j)}(t) r_{22}^{(j)}(t) \right), \\ b_{21}^{(k)}(t) &= - \sum_{j=0}^{k-1} \binom{k-1}{j} a_{12}^{(k-1-j)}(t) r_{11}^{(j)}(t), \\ b_{22}^{(k)}(t) &= - \sum_{j=0}^{k-1} \binom{k-1}{j} \left(a_{12}^{(k-1-j)}(t) r_{12}^{(j)}(t) + a_{22}^{(k-1-j)}(t) r_{22}^{(j)}(t) \right). \end{aligned}$$

By the inductive hypothesis, all derivatives of order less than k satisfy the symmetric relations at $t = T_1$. Inserting these equalities into the expressions with $t = T_1$ above shows that the symmetry is maintained at the k -th order:

$$a_{11}^{(k)}(T_1) = b_{22}^{(k)}(T_1), \quad a_{12}^{(k)}(T_1) = b_{12}^{(k)}(T_1), \quad a_{21}^{(k)}(T_1) = b_{21}^{(k)}(T_1), \quad a_{22}^{(k)}(T_1) = b_{11}^{(k)}(T_1),$$

proving equations (48) and (49). ■

Lemma 19 *Under the setting of Lemma 18, below relationships hold for all $t \geq T_1$:*

$$\begin{aligned} a_{11}(t) &= b_{22}(t), & a_{12}(t) &= b_{12}(t), \\ a_{21}(t) &= b_{21}(t), & a_{22}(t) &= b_{11}(t), \end{aligned} \quad (50)$$

which further leads to $w_{11}(t) = w_{22}(t)$.

Proof By Lemmas 37 and 18, we may conclude that for all $t \geq T_1$, equation (50) holds, and therefore $w_{11}(t) = w_{22}(t)$. ■

By Lemma 19, all entries of $\mathbf{B}(t)$ can be expressed in terms of the entries of $\mathbf{A}(t)$ for all $t \geq T_1$. From this point onward, we will represent $\mathbf{W}_{\mathbf{A},\mathbf{B}}(t)$ solely using the elements of $\mathbf{A}(t)$. We begin by simplifying the time derivative of $\mathbf{A}(t)$ as follows:

$$\begin{aligned} \dot{a}_{11}(t) &= a_{22}(t)(w^* - w_{11}(t)) + a_{12}(t)(w_{12}^* - w_{12}(t)), \\ \dot{a}_{12}(t) &= a_{21}(t)(w^* - w_{11}(t)) + a_{11}(t)(w_{12}^* - w_{12}(t)), \\ \dot{a}_{21}(t) &= a_{12}(t)(w^* - w_{22}(t)), \\ \dot{a}_{22}(t) &= a_{11}(t)(w^* - w_{22}(t)). \end{aligned} \quad (51)$$

Rewriting $\mathbf{W}_{\mathbf{A},\mathbf{B}}(t)$ in terms of the elements of $\mathbf{A}(t)$ yields:

$$\begin{aligned} \mathbf{W}_{\mathbf{A},\mathbf{B}}(t) &= \mathbf{A}(t)\mathbf{B}(t) \\ &= \begin{pmatrix} a_{11}(t) & a_{12}(t) \\ a_{21}(t) & a_{22}(t) \end{pmatrix} \begin{pmatrix} a_{22}(t) & a_{12}(t) \\ a_{21}(t) & a_{11}(t) \end{pmatrix} \\ &= \begin{pmatrix} a_{11}(t)a_{22}(t) + a_{12}(t)a_{21}(t) & 2a_{11}(t)a_{12}(t) \\ 2a_{21}(t)a_{22}(t) & a_{11}(t)a_{22}(t) + a_{12}(t)a_{21}(t) \end{pmatrix}. \end{aligned} \quad (52)$$

We can also simplify the time derivative of $\mathbf{W}_{\mathbf{A},\mathbf{B}}(t)$ as follows:

$$\begin{aligned} \dot{w}_{11}(t) &= (w^* - w_{11}(t)) (a_{11}^2(t) + a_{12}^2(t) + a_{21}^2(t) + a_{22}^2(t)) \\ &\quad + (w_{12}^* - w_{12}(t)) (a_{11}(t)a_{21}(t) + a_{12}(t)a_{22}(t)), \\ \dot{w}_{12}(t) &= 2(w_{12}^* - w_{12}(t)) (a_{11}^2(t) + a_{12}^2(t)) \\ &\quad + 2(w^* - w_{11}(t)) (a_{11}(t)a_{21}(t) + a_{12}(t)a_{22}(t)), \\ \dot{w}_{21}(t) &= 2(w^* - w_{11}(t)) (a_{11}(t)a_{21}(t) + a_{12}(t)a_{22}(t)), \\ \dot{w}_{22}(t) &= \dot{w}_{11}(t). \end{aligned} \quad (53)$$

Using (52), we state the basic conservation law: if the matrices are initialized in a balanced manner, this balancedness is preserved throughout the training process. That is,

$$\mathbf{A}(T_1)^\top \mathbf{A}(T_1) = \mathbf{B}(T_1)\mathbf{B}(T_1)^\top,$$

holds at initialization, this leads to

$$a_{11}^2(t) + a_{21}^2(t) = a_{12}^2(t) + a_{22}^2(t), \quad \forall t \geq T_1. \quad (54)$$

Now, we are going to examine the time derivative of the loss:

$$\begin{aligned}
 \frac{d}{dt}\ell(\mathbf{W}_{\mathbf{A},\mathbf{B}}(t)) &= \left\langle \nabla\ell(\mathbf{W}_{\mathbf{A},\mathbf{B}}(t)), \dot{\mathbf{W}}(t) \right\rangle \\
 &= \left\langle \nabla\ell(\mathbf{W}_{\mathbf{A},\mathbf{B}}(t)), \dot{\mathbf{A}}(t)\mathbf{B}(t) + \mathbf{A}(t)\dot{\mathbf{B}}(t) \right\rangle \\
 &= \text{Tr} \left(\nabla\ell^\top(\mathbf{W}_{\mathbf{A},\mathbf{B}}(t)) \left(\dot{\mathbf{A}}(t)\mathbf{B}(t) + \mathbf{A}(t)\dot{\mathbf{B}}(t) \right) \right) \\
 &= \text{Tr} \left(\nabla\ell^\top(\mathbf{W}_{\mathbf{A},\mathbf{B}}(t)) \dot{\mathbf{A}}(t)\mathbf{B}(t) \right) + \text{Tr} \left(\nabla\ell^\top(\mathbf{W}_{\mathbf{A},\mathbf{B}}(t)) \mathbf{A}(t)\dot{\mathbf{B}}(t) \right) \\
 &= -\text{Tr} \left(\nabla\ell^\top(\mathbf{W}_{\mathbf{A},\mathbf{B}}(t)) \nabla\ell(\mathbf{W}_{\mathbf{A},\mathbf{B}}(t)) \mathbf{B}^\top(t)\mathbf{B}(t) \right) \\
 &\quad - \text{Tr} \left(\nabla\ell^\top(\mathbf{W}_{\mathbf{A},\mathbf{B}}(t)) \mathbf{A}(t) \mathbf{A}^\top(t) \nabla\ell(\mathbf{W}_{\mathbf{A},\mathbf{B}}(t)) \right) \\
 &= -\text{Tr} \left(\underbrace{\nabla\ell(\mathbf{W}_{\mathbf{A},\mathbf{B}}(t)) \mathbf{B}^\top(t)\mathbf{B}(t) \nabla\ell(\mathbf{W}_{\mathbf{A},\mathbf{B}}^\top(t))}_{:=\mathbf{L}_1(t)} \right) \\
 &\quad - \text{Tr} \left(\underbrace{\nabla\ell(\mathbf{W}_{\mathbf{A},\mathbf{B}}^\top(t)) \mathbf{A}(t) \mathbf{A}^\top(t) \nabla\ell(\mathbf{W}_{\mathbf{A},\mathbf{B}}(t))}_{:=\mathbf{L}_2(t)} \right). \tag{55}
 \end{aligned}$$

The third equality follows from the fact that for any two matrices $\mathbf{A}$ and $\mathbf{B}$ of the same size, $\langle \mathbf{A}, \mathbf{B} \rangle = \text{Tr}(\mathbf{A}^\top \mathbf{B})$. The last equation holds due to the cyclic property of the trace. Combining (55) with Lemma 38, we can ensure $\mathbf{L}_1(t)$ and $\mathbf{L}_2(t)$ are both positive semidefinite, which implies the loss is monotonically non-increasing for all $t \geq T_1$.

With Lemma 19 and the monotonicity of the loss, we can guarantee positiveness of a_{11}, a_{22}, w_{11} , and w_{22} after the pre-train phase:

Lemma 20 *For a product matrix $\mathbf{W}_{\mathbf{A},\mathbf{B}}(t) = \mathbf{A}(t)\mathbf{B}(t) \in \mathbb{R}^{2 \times 2}$, if $a_{11}(T_1), a_{22}(T_1), w_{11}(T_1)$, and $w_{22}(T_1)$ have all positive values, following inequalities hold for all $t \geq T_1$:*

$$a_{11}(t), a_{22}(t) > 0, \quad a_{12}(t) \geq 0.$$

Furthermore,

$$w_{11}(t), w_{22}(t) > 0$$

holds for all $t \geq T_1$.

Proof We will prove the inequalities step by step.

Positiveness of $a_{11}(t)$. For the sake of contradiction, assume that there exists a timestep $\tau_1 > T_1$ where $a_{11}(\tau_1) = 0$ holds. From (52) and Lemma 34, we must have $\det(\mathbf{A}(\tau_1)) > 0$, which implies that $a_{12}(\tau_1)a_{21}(\tau_1) < 0$. Given the monotonicity of ℓ , $\mathbf{W}_{\mathbf{A},\mathbf{B}}(t)$ must satisfy:

$$\ell(\mathbf{W}_{\mathbf{A},\mathbf{B}}(t)) \leq \ell(\mathbf{W}_{\mathbf{A},\mathbf{B}}(T_1)). \tag{56}$$

for all $t \geq T_1$. However, $\mathbf{W}_{\mathbf{A},\mathbf{B}}(\tau_1)$ cannot satisfy (56) because $w_{11}(\tau_1), w_{22}(\tau_1) < 0$ and $w_{12}(\tau_1) = 0$ for any $\tau_1 \geq 0$. This contradiction implies that such a τ_1 cannot exist.

Positiveness of $a_{22}(t)$. Similarly, let's assume there exists a time $\tau_2 > T_1$ such that $a_{22}(\tau_2) = 0$ for the first time. We can express $\mathbf{W}_{\mathbf{A},\mathbf{B}}(\tau_2)$ as:

$$\mathbf{W}_{\mathbf{A},\mathbf{B}}(\tau_2) = \begin{pmatrix} a_{12}(\tau_2)a_{21}(\tau_2) & 2a_{11}(\tau_2)a_{12}(\tau_2) \\ 0 & a_{12}(\tau_2)a_{21}(\tau_2) \end{pmatrix}.$$

where the diagonal entries are negative due to the condition $\det(\mathbf{A}(\tau_2)) > 0$. Therefore, the time derivative of a_{22} at timestep τ_2 is positive:

$$\dot{a}_{22}(\tau_2) = a_{11}(\tau_2)(w^* - w_{11}(\tau_2)) > 0.$$

Since $a_{22}(t)$ is increasing at point τ_2 , there exists time $t' < \tau_2$ such that $a_{22}(t') < 0$ (since $a_{22}(t)$ is continuous and differentiable), which is contradictory. Consequently, there cannot exist a τ_2 such that $a_{22}(\tau_2) = 0$.

Positiveness of $\mathbf{a}_{12}(\mathbf{t})$. Given that ℓ is non-decreasing, we can state:

$$\begin{aligned} \ell(\mathbf{W}_{\mathbf{A},\mathbf{B}}(t)) &= \frac{1}{2} [(w^* - w_{11}(t))^2 + (w_{12}^* - w_{12}(t))^2 + (w^* - w_{22}(t))^2] \\ &\leq \ell(\mathbf{W}_{\mathbf{A},\mathbf{B}}(T_1)) = \frac{1}{2} w_{12}^{*2}, \end{aligned}$$

for all $t \geq T_1$. Since $(w^* - w_{11}(t))^2$ and $(w^* - w_{22}(t))^2$ are non-negative, $w_{12}(t)$ must be non-negative for all $t \geq T_1$. From (52), we know $w_{12}(t) = 2a_{11}(t)a_{12}(t)$, which implies $a_{12}(t) \geq 0$ for all $t \geq T_1$ with the above conclusion which states $a_{11}(t) > 0$.

Positiveness of $\mathbf{w}_{11}(\mathbf{t})$, $\mathbf{w}_{22}(\mathbf{t})$. Likewise, assume for the sake of contradiction that there exists a time $\tau_3 \geq T_1$ when $w_{11}(\tau_3) = 0$ is first satisfied. This directly implies that $a_{11}(\tau_3)a_{22}(\tau_3) = -a_{12}(\tau_3)a_{21}(\tau_3)$. Squaring both sides of the equation yields:

$$a_{11}^2(\tau_3)a_{22}^2(\tau_3) = a_{12}^2(\tau_3)a_{21}^2(\tau_3).$$

Subtracting $a_{12}^2(\tau_3)a_{22}^2(\tau_3)$ from both sides:

$$a_{11}^2(\tau_3)a_{22}^2(\tau_3) - a_{12}^2(\tau_3)a_{22}^2(\tau_3) = a_{12}^2(\tau_3)a_{21}^2(\tau_3) - a_{12}^2(\tau_3)a_{22}^2(\tau_3).$$

Factoring:

$$a_{22}^2(\tau_3) (a_{11}^2(\tau_3) - a_{12}^2(\tau_3)) = a_{12}^2(\tau_3) (a_{21}^2(\tau_3) - a_{22}^2(\tau_3)).$$

By the conservation law in (54), we have $a_{11}^2(\tau_3) + a_{21}^2(\tau_3) = a_{12}^2(\tau_3) + a_{22}^2(\tau_3)$, which leads to $a_{11}^2(\tau_3) - a_{12}^2(\tau_3) = a_{22}^2(\tau_3) - a_{21}^2(\tau_3)$. Replacing $a_{11}^2(\tau_3) - a_{12}^2(\tau_3)$ with $-(a_{21}^2(\tau_3) - a_{22}^2(\tau_3))$:

$$-a_{22}^2(\tau_3) (a_{21}^2(\tau_3) - a_{22}^2(\tau_3)) = a_{12}^2(\tau_3) (a_{21}^2(\tau_3) - a_{22}^2(\tau_3)).$$

This gives us:

$$(a_{12}^2(\tau_3) + a_{22}^2(\tau_3)) (a_{21}^2(\tau_3) - a_{22}^2(\tau_3)) = 0.$$

Since $a_{22}(\tau_3) > 0$ from the previous result, we can conclude that $a_{21}(\tau_3) = \pm a_{22}(\tau_3)$. To determine the sign of $a_{21}(\tau_3)$, recall that $\mathbf{W}_{\mathbf{A},\mathbf{B}}(\tau_3)$ is written as:

$$\mathbf{W}_{\mathbf{A},\mathbf{B}}(\tau_3) = \begin{pmatrix} 0 & 2a_{11}(\tau_3)a_{12}(\tau_3) \\ 2a_{21}(\tau_3)a_{22}(\tau_3) & 0 \end{pmatrix}.$$

Since $a_{11}(\tau_3) > 0$, $a_{12}(\tau_3) \geq 0$ from the previous result, $2a_{11}(\tau_3)a_{12}(\tau_3) \geq 0$ holds. Also, given that $\det(\mathbf{W}_{\mathbf{A},\mathbf{B}}(\tau_3)) > 0$, we can determine that $a_{21}(\tau_3)$ is negative, which implies $a_{21}(\tau_3) =$

$-a_{22}(\tau_3)$. Additionally, by the conservation law, we have $a_{11}^2(\tau_3) = a_{12}^2(\tau_3)$, which leads to $a_{11}(\tau_3) = a_{12}(\tau_3) > 0$.

Finally, consider the time derivative of w_{11} at timestep τ_3 , substituting $a_{11}(\tau_3)$ and $a_{21}(\tau_3)$ with $a_{12}(\tau_3)$ and $-a_{22}(\tau_3)$, respectively:

$$\begin{aligned}\dot{w}_{11}(\tau_3) &= (w^* - w_{11}(\tau_3))(a_{11}^2(\tau_3) + a_{12}^2(\tau_3) + a_{21}^2(\tau_3) + a_{22}^2(\tau_3)) \\ &\quad + (w_{12}^* - w_{12}(\tau_3))(a_{11}(\tau_3)a_{21}(\tau_3) + a_{12}(\tau_3)a_{22}(\tau_3)) \\ &= 2w^*(a_{12}^2(\tau_3) + a_{22}^2(\tau_3)) \\ &> 0,\end{aligned}$$

which contradicts our initial assumption. ■

Given that the time derivative in the (53) includes the term $a_{11}(t)a_{21}(t) + a_{12}(t)a_{22}(t)$, we need to verify the sign of $a_{11}a_{21} + a_{12}a_{22}$ in order to proceed with the analysis. Below lemma shows that as long as $w_{12}(t) \leq w_{12}^*$ holds, $a_{11}(t)a_{21}(t) + a_{12}(t)a_{22}(t)$ is always lower bounded by zero.

Lemma 21 *For a product matrix $\mathbf{W}_{\mathbf{A},\mathbf{B}}(t) = \mathbf{A}(t)\mathbf{B}(t) \in \mathbb{R}^{2 \times 2}$, if at any point $t \in [T_1, T_2]$ we have $w_{12}(t) \leq w_{12}^*$, then the following inequality holds throughout the entire interval $[T_1, T_2]$:*

$$a_{11}(t)a_{21}(t) + a_{12}(t)a_{22}(t) \geq 0.$$

Proof We first define $g(t) \triangleq a_{11}(t)a_{21}(t) + a_{12}(t)a_{22}(t)$. Recall that at T_1 , we have $a_{12}(T_1) = a_{21}(T_1) = 0$, which implies $g(T_1) = 0$ as well. Note that by (51), at timestep T_1 , we have

$$\dot{a}_{12}(T_1) = a_{11}(T_1)(w_{12}^* - w_{12}(T_1)) + a_{21}(T_1)(w^* - w_{11}(T_1)) > 0,$$

while other elements remain unchanged. This indicates that $g(t) > 0$ immediately after T_1 . We now show that if $g(\tau) > 0$ for any $\tau \in (T_1, T_2]$, then there is no $\tau' \in [\tau, T_2]$ which satisfies both $g(\tau') = 0$ and $\frac{d}{dt}g(t)\Big|_{t=\tau'} < 0$. This implies that $g(t)$ never becomes negative under the assumption of $w_{12}(t) \leq w_{12}^*$.

Suppose, for the sake of contradiction, that there exists a $\tau' \in [\tau, T_2]$ where $g(\tau') = 0$ and $\frac{d}{dt}g(t)\Big|_{t=\tau'} < 0$. Given $g(\tau') = 0$ and the conservation law in (54), and the inequalities from Lemma 20, we can determine that there exist two combinations of the solution:

1. $a_{11}(\tau') = a_{22}(\tau')$, $a_{12}(\tau') = -a_{21}(\tau')$, $a_{11}(\tau') > a_{12}(\tau')$.
2. $a_{11}(\tau') = a_{22}(\tau')$, $a_{12}(\tau') = a_{21}(\tau') = 0$.

We take the time derivative of $g(t)$ at timestep τ' and substitute the values from (51) as follows:

$$\begin{aligned}\frac{d}{dt}g(t)\Big|_{t=\tau'} &= \dot{a}_{11}(\tau')a_{21}(\tau') + a_{11}(\tau')\dot{a}_{21}(\tau') + \dot{a}_{12}(\tau')a_{22}(\tau') + a_{12}(\tau')\dot{a}_{22}(\tau') \\ &= 2(w^* - w_{11}(\tau'))(a_{11}(\tau')a_{12}(\tau') + a_{21}(\tau')a_{22}(\tau')) \\ &\quad + (w_{12}^* - w_{12}(\tau'))(a_{11}(\tau')a_{22}(\tau') + a_{12}(\tau')a_{21}(\tau')).\end{aligned}\tag{57}$$

For the first case, substituting equations $a_{11}(\tau') = a_{22}(\tau')$ and $a_{12}(\tau') = -a_{21}(\tau')$ to (57) leads to:

$$\left. \frac{d}{dt}g(t) \right|_{t=\tau'} = (w_{12}^* - w_{12}(\tau'))w_{11}(\tau').$$

Since $w_{11}(t) > 0$ for all $t \geq T_1$, if $w_{12}(\tau') \leq w_{12}^*$ holds, then $g(t)$ cannot take negative values at time τ' .

For the second case, substituting equations $a_{11}(\tau') = a_{22}(\tau')$ and $a_{12}(\tau') = a_{21}(\tau') = 0$ to (57) leads to:

$$\left. \frac{d}{dt}g(t) \right|_{t=\tau'} = (w_{12}^* - w_{12}(\tau'))a_{11}^2(\tau'),$$

which is again a non-negative value if $w_{12}(\tau') \leq w_{12}^*$, leading to a contradiction. $\blacksquare$

Lemma 22 *For a product matrix $\mathbf{W}_{A,B}(t) = \mathbf{A}(t)\mathbf{B}(t) \in \mathbb{R}^{2 \times 2}$, the following inequalities holds for all timestep $t \geq T_1$:*

$$\begin{aligned} w_{12}(t) &\leq w_{12}^*, \\ w_{11}(t), w_{22}(t) &\geq w^*, \\ w_{21}(t) &\leq 0. \end{aligned}$$

Proof We will prove this lemma in several steps:

Step 1: $w_{12}(t) \leq w_{12}^*$ for all $t \geq T_1$.

We know $w_{12}(T_1) = 0 \leq w_{12}^*$. Assume, for the sake of contradiction, that there exists a time $t' > T_1$ where t' is the first timestep such that $w_{12}(t') > w_{12}^*$. If this were true, there must exist a time s where $T_1 \leq s < t'$ such that:

$$w_{12}(s) = w_{12}^*, \quad \dot{w}_{12}(s) > 0.$$

For these conditions to be met, $w_{12}(s)$ must satisfy:

$$\dot{w}_{12}(s) = 2(w^* - w_{11}(s))(a_{11}(s)a_{21}(s) + a_{12}(s)a_{22}(s)) > 0. \quad (58)$$

To satisfy (58), there are two possibilities:

$$(w^* - w_{11}(s)) > 0 \quad \text{and} \quad (a_{11}(s)a_{21}(s) + a_{12}(s)a_{22}(s)) > 0, \quad (59)$$

$$\text{or} \quad (w^* - w_{11}(s)) < 0 \quad \text{and} \quad (a_{11}(s)a_{21}(s) + a_{12}(s)a_{22}(s)) < 0. \quad (60)$$

However, neither of these can be true:

1. Equation (60) contradicts Lemma 21, given that $s < t'$.
2. Equation (59) cannot be satisfied because there is no s where $w^* > w_{11}(s)$. If there were, there would be a time s' where $T_1 \leq s' < s$ both satisfying $w_{11}(s') = w^*$, and $\dot{w}_{11}(s') < 0$. But we find:

$$\dot{w}_{11}(s') = (w_{12}^* - w_{12}(s'))(a_{11}(s')a_{21}(s') + a_{12}(s')a_{22}(s')) \geq 0.$$

This is because $w_{12}(s') < w_{12}^*$, and thus $a_{11}(s')a_{21}(s') + a_{12}(s')a_{22}(s') \geq 0$ by Lemma 21. Therefore, our initial assumption must be false, implying that $w_{12}(t) \leq w_{12}^*$ for all $t \geq T_1$.

Step 2: Prove $w_{11}(t) \geq w_{11}^*$ and $w_{22}(t) \geq w_{22}^*$ for all $t \geq T_1$.

Given $w_{12}(t) \leq w_{12}^*$ for all $t \geq T_1$, Lemma 21 implies $a_{11}(t)a_{21}(t) + a_{12}(t)a_{22}(t) \geq 0$ for all $t \geq T_1$. The evolution of w_{11} is given by:

$$\dot{w}_{11}(t) = (w^* - w_{11}(t))(a_{11}^2(t) + a_{12}^2(t) + a_{21}^2(t) + a_{22}^2(t)) + (w_{12}^* - w_{12}(t))(a_{11}(t)a_{21}(t) + a_{12}(t)a_{22}(t)).$$

By above equation, if there exists a time $t' \geq T_1$ where $w_{11}(t') = w^*$, we can conclude $\dot{w}_{11}(t') \geq 0$, and thus $w_{11}(t) \geq w^*$ for all $t \geq T_1$. By Lemma 19, w_{22} has the same value as w_{11} , so $w_{22}(t) \geq w^*$ for all $t \geq T_1$.

Step 3: Prove $w_{21}(t) \leq 0$ for all $t \geq T_1$.

The evolution of w_{21} is given by:

$$\dot{w}_{21}(t) = 2(w^* - w_{11}(t))(a_{11}(t)a_{21}(t) + a_{12}(t)a_{22}(t)).$$

Since $w_{11}(t) \geq w^*$ and $a_{11}(t)a_{21}(t) + a_{12}(t)a_{22}(t) \geq 0$ for all $t \geq T_1$, we can conclude $\dot{w}_{21}(t) \leq 0$ for all $t \geq T_1$. ■

G.2.1. PROOF OF LOSS CONVERGENCE

Recall that the time derivative of the loss function is written as:

$$\frac{d}{dt}\ell(\mathbf{W}_{\mathbf{A},\mathbf{B}}(t)) = -\text{Tr}(\mathbf{L}_1(t)) - \text{Tr}(\mathbf{L}_2(t)),$$

where $\mathbf{L}_1(t)$ and $\mathbf{L}_2(t)$ are defined in (55). To further our analysis, we can expand the time derivative of the loss by calculating the trace of $\mathbf{L}_1(t)$ and $\mathbf{L}_2(t)$. We omit the time index t when clear from context.

$$\begin{aligned} \mathbf{L}_1 &= \begin{pmatrix} r_{11} & r_{12} \\ 0 & r_{22} \end{pmatrix} \begin{pmatrix} a_{21}^2 + a_{22}^2 & a_{11}a_{21} + a_{12}a_{22} \\ a_{11}a_{21} + a_{12}a_{22} & a_{11}^2 + a_{12}^2 \end{pmatrix} \begin{pmatrix} r_{11} & 0 \\ r_{12} & r_{22} \end{pmatrix} \\ &= \begin{pmatrix} r_{11}^2(a_{21}^2 + a_{22}^2) + 2r_{11}r_{12}(a_{11}a_{21} + a_{12}a_{22}) + r_{12}^2(a_{11}^2 + a_{12}^2) & C_1 \\ C_1 & r_{22}^2(a_{11}^2 + a_{12}^2) \end{pmatrix}, \end{aligned}$$

for some time-dependent value C_1 . Following a similar process, we calculate $\mathbf{L}_2$:

$$\begin{aligned} \mathbf{L}_2 &= \begin{pmatrix} r_{11} & 0 \\ r_{12} & r_{22} \end{pmatrix} \begin{pmatrix} a_{11}^2 + a_{12}^2 & a_{11}a_{21} + a_{12}a_{22} \\ a_{11}a_{21} + a_{12}a_{22} & a_{21}^2 + a_{22}^2 \end{pmatrix} \begin{pmatrix} r_{11} & r_{12} \\ 0 & r_{22} \end{pmatrix} \\ &= \begin{pmatrix} r_{11}^2(a_{11}^2 + a_{12}^2) & C_2 \\ C_2 & r_{12}^2(a_{11}^2 + a_{12}^2) + 2r_{12}r_{22}(a_{11}a_{21} + a_{12}a_{22}) + r_{22}^2(a_{21}^2 + a_{22}^2) \end{pmatrix}, \end{aligned}$$

again for the time-dependent value C_2 . With these expressions for $\mathbf{L}_1$ and $\mathbf{L}_2$, we can now rewrite equation (55) in a more explicit form:

$$\begin{aligned} \frac{d}{dt}\ell(\mathbf{W}_{\mathbf{A},\mathbf{B}}(t)) &= -\text{Tr}(\mathbf{L}_1(t)) - \text{Tr}(\mathbf{L}_2(t)) \\ &= -r_{11}^2(t)(a_{11}^2(t) + a_{12}^2(t) + a_{21}^2(t) + a_{22}^2(t)) \\ &\quad - 2r_{12}^2(t)(a_{11}^2(t) + a_{12}^2(t)) \\ &\quad - r_{22}^2(t)(a_{11}^2(t) + a_{12}^2(t) + a_{21}^2(t) + a_{22}^2(t)) \\ &\quad - 2r_{12}(t)r_{22}(t)(a_{11}(t)a_{21}(t) + a_{12}(t)a_{22}(t)) \\ &\quad - 2r_{11}(t)r_{12}(t)(a_{11}(t)a_{21}(t) + a_{12}(t)a_{22}(t)). \end{aligned} \tag{61}$$

Note that the (61) is the non-positive term. Given that $\mathbf{L}_1$ and $\mathbf{L}_2$ are positive semi-definite, we can analyze each diagonal entry separately. This leads us to the following inequalities:

$$\begin{aligned} r_{11}^2(a_{21}^2 + a_{22}^2) + 2r_{11}r_{12}(a_{11}a_{21} + a_{12}a_{22}) + r_{12}^2(a_{11}^2 + a_{12}^2) &\geq 0, \\ r_{12}^2(a_{11}^2 + a_{12}^2) + 2r_{12}r_{22}(a_{11}a_{21} + a_{12}a_{22}) + r_{22}^2(a_{21}^2 + a_{22}^2) &\geq 0. \end{aligned}$$

By rearranging the above inequalities, we obtain:

$$\begin{aligned} -2r_{11}r_{12}(a_{11}a_{21} + a_{12}a_{22}) &\leq r_{11}^2(a_{21}^2 + a_{22}^2) + r_{12}^2(a_{11}^2 + a_{12}^2), \\ -2r_{12}r_{22}(a_{11}a_{21} + a_{12}a_{22}) &\leq r_{12}^2(a_{11}^2 + a_{12}^2) + r_{22}^2(a_{21}^2 + a_{22}^2). \end{aligned}$$

Substituting these inequalities into equation (61), we derive:

$$\frac{d}{dt}\ell(\mathbf{W}_{\mathbf{A},\mathbf{B}}(t)) \leq -r_{11}^2(t)(a_{11}^2(t) + a_{12}^2(t)) - r_{22}^2(t)(a_{11}^2(t) + a_{12}^2(t)). \quad (62)$$

This provides a tighter upper bound on the time derivative of the loss. However, it is still insufficient to guarantee convergence, as the bound does not depend on the term $r_{12}(t)$. As a result, even though the right-hand side converges to zero, this alone does not imply that the loss itself converges.

To further tighten the bound, we leverage the positive semidefiniteness of $\mathbf{L}_1$ and $\mathbf{L}_2$. Specifically, note that for both $\mathbf{Q}\mathbf{K}\mathbf{Q}^\top$ and $\mathbf{Q}^\top\mathbf{K}\mathbf{Q}$ to be positive semi-definite, the only necessary condition is $\mathbf{K} \succcurlyeq 0$. Therefore, we modify $\mathbf{L}_1(t)$ to $\widetilde{\mathbf{L}}_1(t) \triangleq \nabla\ell(\mathbf{W}_{\mathbf{A},\mathbf{B}}(t))(\mathbf{B}^\top(t)\mathbf{B}(t) - \mu(t) \cdot \mathbf{e}_2\mathbf{e}_2^\top) \nabla\ell^\top(\mathbf{W}_{\mathbf{A},\mathbf{B}}(t))$, where $\mu(t)$ is chosen to ensure that the matrix $\mathbf{B}^\top(t)\mathbf{B}(t) - \mu(t) \cdot \mathbf{e}_2\mathbf{e}_2^\top$ remains positive semidefinite. This guarantees that $\widetilde{\mathbf{L}}_1(t) \succcurlyeq 0$. To ensure this condition, $\mu(t)$ must satisfy:

$$\begin{aligned} \left| \mathbf{B}(t)^\top \mathbf{B}(t) - \mu(t) \cdot \mathbf{e}_2\mathbf{e}_2^\top \right| &= \left| \begin{pmatrix} a_{21}^2(t) + a_{22}^2(t) & a_{11}(t)a_{21}(t) + a_{12}(t)a_{22}(t) \\ a_{11}(t)a_{21}(t) + a_{12}(t)a_{22}(t) & a_{11}^2(t) + a_{12}^2(t) - \mu(t) \end{pmatrix} \right| \\ &= - (a_{21}^2(t) + a_{22}^2(t)) \mu(t) + (a_{11}(t)a_{22}(t) - a_{12}(t)a_{21}(t))^2 \\ &\geq 0. \end{aligned}$$

Rearranging this inequality with respect to $\mu(t)$, we get:

$$\begin{aligned} \mu(t) &\leq \frac{(a_{11}(t)a_{22}(t) - a_{12}(t)a_{21}(t))^2}{a_{21}^2(t) + a_{22}^2(t)} \\ &= \frac{\det(\mathbf{B}(t))^2}{a_{21}^2(t) + a_{22}^2(t)}. \end{aligned} \quad (63)$$

Therefore, if we set $\mu(t)$ to satisfy the above inequality, we can guarantee $\widetilde{\mathbf{L}}_1$ to be a positive semidefinite matrix. Now, $\widetilde{\mathbf{L}}_1(t)$ can be calculated as:

$$\begin{aligned} \widetilde{\mathbf{L}}_1 &= \begin{pmatrix} r_{11} & r_{12} \\ 0 & r_{22} \end{pmatrix} \begin{pmatrix} a_{21}^2 + a_{22}^2 & a_{11}a_{21} + a_{12}a_{22} \\ a_{11}a_{21} + a_{12}a_{22} & a_{11}^2 + a_{12}^2 - \mu \end{pmatrix} \begin{pmatrix} r_{11} & 0 \\ r_{12} & r_{22} \end{pmatrix} \\ &= \begin{pmatrix} r_{11}^2(a_{21}^2 + a_{22}^2) + 2r_{11}r_{12}(a_{11}a_{21} + a_{12}a_{22}) + r_{12}^2(a_{11}^2 + a_{12}^2 - \mu) & \tilde{C} \\ \tilde{C} & r_{22}^2(a_{11}^2 + a_{12}^2 - \mu) \end{pmatrix}, \end{aligned}$$

for some $\tilde{C}$. Since the matrix $\mathbf{B}^\top \mathbf{B} - \mu \cdot \mathbf{e}_2 \mathbf{e}_2^\top$ is positive semi-definite, we can ensure $a_{12}^2 + a_{22}^2 - \mu \geq 0$. This leads to the following inequality from $(\tilde{\mathbf{L}}_1)_{11}$:

$$-2r_{11}r_{12}(a_{11}a_{21} + a_{12}a_{22}) \leq r_{11}^2(a_{21}^2 + a_{22}^2) + r_{12}^2(a_{11}^2 + a_{12}^2 - \mu).$$

Finally, substituting this inequality into (61), we arrive at:

$$\frac{d}{dt} \ell(\mathbf{W}_{\mathbf{A}, \mathbf{B}}(t)) \leq -(r_{11}^2(t) + r_{22}^2(t)) (a_{11}^2(t) + a_{12}^2(t)) - r_{12}^2(t) \mu(t). \quad (64)$$

To prove the convergence of the loss, our main remaining goal is to establish a time-invariant lower bound for

$$\min \{a_{11}^2(t) + a_{12}^2(t), \mu(t)\}$$

to apply Grönwall's inequality.

Lemma 23 *For a solution matrix $\mathbf{W}_{\mathbf{A}, \mathbf{B}}(t)$ initialized as $\mathbf{W}_{\mathbf{A}, \mathbf{B}}(T_1)$, which represents the state of the matrix after pre-training up to time T_1 , the inequality*

$$\det(\mathbf{W}_{\mathbf{A}, \mathbf{B}}(t)) \geq w^{*2}$$

holds for all $t \geq T_1$.

Proof Since $w_{12}(t)$ must satisfy $|w_{12}(t) - w_{12}^*| \leq \sqrt{2\ell(\mathbf{W}_{\mathbf{A}, \mathbf{B}}(t))} \leq w_{12}^*$ by the monotonicity of the loss, we can ensure that $w_{12}(t) \geq 0$ for all $t \geq T_1$. Also, by Lemma 22, we have $w_{11}(t), w_{22}(t) \geq w^*$, and $w_{21}(t) \leq 0$ for all $t \geq T_1$. Under these conditions, $\det(\mathbf{W}_{\mathbf{A}, \mathbf{B}}(t))$ can be lower bounded as:

$$\det(\mathbf{W}_{\mathbf{A}, \mathbf{B}}(t)) = w_{11}(t)w_{22}(t) - w_{12}(t)w_{21}(t) \geq w^{*2},$$

for all timesteps $t \geq T_1$. ■

Lemma 24 *For $\mu(t)$ defined to satisfy (63) and the entries in $\mathbf{A}(t)$, the following inequality holds for all timesteps $t \geq T_1$:*

$$\min \{a_{11}^2(t) + a_{12}^2(t), \mu(t)\} \geq w^*.$$

Proof

To prove the lower bound of $a_{11}^2(t) + a_{12}^2(t)$, Our goal is to demonstrate that $a_{11}^2(t) + a_{12}^2(t) \geq w^*$ for all timesteps t after T_1 . By Lemma 24, we have $\|\mathbf{W}_{\mathbf{A}, \mathbf{B}}(t)\|_F \geq \sqrt{2}w^*$, which leads to:

$$\begin{aligned} \sqrt{2}w^* &\leq \|\mathbf{W}_{\mathbf{A}, \mathbf{B}}(t)\|_F \\ &= \sqrt{\sigma_1^2(\mathbf{W}_{\mathbf{A}, \mathbf{B}}(t)) + \sigma_2^2(\mathbf{W}_{\mathbf{A}, \mathbf{B}}(t))}. \end{aligned}$$

By applying Lemma 35, we have:

$$\begin{aligned}
 \sqrt{\sigma_1^2(\mathbf{W}_{\mathbf{A},\mathbf{B}}(t)) + \sigma_2^2(\mathbf{W}_{\mathbf{A},\mathbf{B}}(t))} &= \sqrt{\sigma_1^4(\mathbf{A}(t)) + \sigma_2^4(\mathbf{A}(t))} \\
 &= \sqrt{(\sigma_1^2(\mathbf{A}(t)) + \sigma_2^2(\mathbf{A}(t)))^2 - 2\sigma_1^2(\mathbf{A}(t))\sigma_2^2(\mathbf{A}(t))} \\
 &= \sqrt{\|\mathbf{A}(t)\|_F^4 - 2\det(\mathbf{A}(t))^2}. \tag{65}
 \end{aligned}$$

Rewriting (65) while applying Lemmas 35 and 23 leads to:

$$\begin{aligned}
 \|\mathbf{A}(t)\|_F^4 &\geq 2w^{*2} + 2\det(\mathbf{A}(t))^2 \\
 &= 2w^{*2} + 2\det(\mathbf{W}_{\mathbf{A},\mathbf{B}}(t)) \\
 &\geq 4w^{*2}.
 \end{aligned}$$

Thus, $\mathbf{A}(t)$ have to satisfy $\|\mathbf{A}(t)\|_F^2 \geq 2w^*$ for all timesteps $t \geq T_1$. Now, assume that there exists a time $t' > T_1$ such that $a_{11}^2(t') + a_{12}^2(t') < w^*$. To satisfy inequality $\|\mathbf{A}(t')\|_F^2 \geq 2w^*$, we would need at least $a_{21}^2(t') + a_{22}^2(t') > w^*$ to hold. To verify the value of $a_{21}^2(t') + a_{22}^2(t')$, we take its time derivative using (51):

$$\begin{aligned}
 \frac{d}{dt}(a_{21}^2(t) + a_{22}^2(t)) &= 2a_{21}(t)\dot{a}_{21}(t) + 2a_{22}(t)\dot{a}_{22}(t) \\
 &= -2a_{12}(t)a_{21}(t)r_{22}(t) - 2a_{11}(t)a_{22}(t)r_{22}(t) \\
 &= -2r_{22}(t)(a_{11}(t)a_{22}(t) + a_{12}(t)a_{21}(t)) \\
 &= 2w_{11}(t)(w^* - w_{11}(t)).
 \end{aligned}$$

Since $w_{11}(t) \geq w^*$ holds by Lemma 22 for all $t \geq T_1$, we conclude $a_{21}^2(t) + a_{22}^2(t)$ is monotonically non-increasing from time $t \geq T_1$. Since $a_{12}^2(T_1) + a_{22}^2(T_1)$ is initialized as w^* , this implies that $a_{21}^2(t') + a_{22}^2(t') \leq w^*$. Consequently, there cannot exist a $t' > T_1$ such that $a_{11}^2(t') + a_{12}^2(t') < w^*$ holds, which leads to contradiction.

Next, we are now showing that the term $\frac{\det(\mathbf{B}(t))^2}{a_{21}^2(t) + a_{22}^2(t)}$ is lower bounded by w^* . Therefore, if we set $\mu(t)$ as w^* , we can guarantee the positive semidefiniteness of $\widetilde{\mathbf{L}}_1(t)$.

By applying Lemma 35 and the lower bound of $\det(\mathbf{W}_{\mathbf{A},\mathbf{B}}(t))$ by Lemma 23, we have

$$\frac{\det(\mathbf{B}(t))^2}{a_{21}^2(t) + a_{22}^2(t)} = \frac{\det(\mathbf{W}_{\mathbf{A},\mathbf{B}}(t))}{a_{21}^2(t) + a_{22}^2(t)} \geq \frac{w^{*2}}{a_{21}^2(t) + a_{22}^2(t)}.$$

Also, from the previous result, we have an upper bound on $a_{21}^2(t) + a_{22}^2(t)$, which is $a_{21}^2(t) + a_{22}^2(t) \leq w^*$. Combining these results, the following inequality holds:

$$\frac{\det(\mathbf{W}_{\mathbf{A},\mathbf{B}}(t))}{a_{21}^2(t) + a_{22}^2(t)} \geq w^*.$$

Therefore, if we set $\mu(t)$ to be w^* , $\mu(t)$ can satisfy the positive semidefiniteness condition. By combining the results, we can finally guarantee:

$$\min \{a_{11}^2(t) + a_{12}^2(t), \mu(t)\} \geq w^*.$$

■

Using the results of Lemma 24, we can rewrite (64) as follows:

$$\begin{aligned} \frac{d}{dt} \ell(\mathbf{W}_{A,B}(t)) &\leq - (r_{11}^2(t) + r_{22}^2(t)) (a_{11}^2(t) + a_{12}^2(t)) - r_{12}^2(t) \mu(t) \\ &\leq - (r_{11}^2(t) + r_{12}^2(t) + r_{22}^2(t)) w^* \\ &\leq -2w^* \ell(\mathbf{W}_{A,B}(t)). \end{aligned}$$

Applying Grönwall's inequality to our previous result, we can now demonstrate loss convergence where $t \geq T_1$:

$$\begin{aligned} \ell(\mathbf{W}_{A,B}(t)) &\leq \ell(\mathbf{W}_{A,B}(T_1)) e^{-2w^*(t-T_1)} \\ &= \frac{1}{2} w_{12}^{*2} e^{-2w^*(t-T_1)}. \end{aligned} \tag{66}$$

This inequality allows us to conclude that $\ell(\mathbf{W}_{A,B}(t))$ converges to zero exponentially.

G.2.2. PROOF OF STABLE RANK BOUND

From (66), we know that at convergence, $w_{11}(\infty) = w_{22}(\infty) = w^*$ and $w_{12}(\infty) = w_{12}^*$. Although a closed-form expression for $w_{21}(\infty)$ is unavailable, Lemma 22 shows that $w_{21}(t) \leq 0$ for $t \geq T_1$, which implies $w_{21}(\infty) \leq 0$. This indicates that the test loss remains strictly positive, as the ground-truth value $w_{21}^* = \frac{w_{12}^{*2}}{w^*}$ is assumed to be strictly positive.

In this section, we leverage the fast convergence rate detailed in (66) to establish bounds on the singular values of the converged matrix $\mathbf{W}_{A,B}(\infty)$. Subsequently, these singular value bounds are used to further bound the stable rank of $\mathbf{W}_{A,B}(\infty)$.

Lemma 25 *The singular values of $\mathbf{W}_{A,B}(\infty)$ fulfill:*

$$\begin{aligned} \sigma_1(\mathbf{W}_{A,B}(\infty)) &\leq w^* \cdot \exp\left(2 \frac{w_{12}^*}{w^*}\right) \\ \sigma_2(\mathbf{W}_{A,B}(\infty)) &\geq w^* \cdot \exp\left(-2 \frac{w_{12}^*}{w^*}\right) \end{aligned}$$

Proof We denote the singular values of $\mathbf{W}_{A,B}(t)$ as $\sigma_r(t)$ for simplicity. By Lemma 32, we can get general solution of each singular value $\sigma_r(t)$ by solving linear differential equation:

$$\sigma_r(t) = \sigma_r(s) \cdot \exp\left(-2 \int_{t'=s}^t \langle \nabla \ell(\mathbf{W}_{A,B}(t')), \mathbf{u}_r(t') \mathbf{v}_r^\top(t') \rangle dt'\right), \quad r = 1, 2 \tag{67}$$

where $\mathbf{u}_r(t)$ and $\mathbf{v}_r(t)$ denotes left and right singular vector of corresponding r -th singular value, respectively. Since $\mathbf{u}_r(t)$ and $\mathbf{v}_r(t)$ are both unit vectors, applying Cauchy-Schwartz inequality, we can bound $\langle \nabla \ell(\mathbf{W}_{A,B}(t)), \mathbf{u}_r(t) \mathbf{v}_r^\top(t) \rangle$ by:

$$\begin{aligned} \left| \langle \nabla \ell(\mathbf{W}_{A,B}(t)), \mathbf{u}_r(t) \mathbf{v}_r^\top(t) \rangle \right| &\leq \|\nabla \ell(\mathbf{W}_{A,B}(t))\|_F \cdot \|\mathbf{u}_r(t) \mathbf{v}_r^\top(t)\|_F \\ &= \|\nabla \ell(\mathbf{W}_{A,B}(t))\|_F \\ &= \sqrt{2\ell(\mathbf{W}_{A,B}(t))}. \end{aligned}$$

we can get bound $\sigma_r(t)$ as following:

$$\sigma_r(s) \cdot \exp \left(-2\sqrt{2} \int_{t'=s}^t \sqrt{\ell(\mathbf{W}_{\mathbf{A},\mathbf{B}}(t'))} dt' \right) \leq \sigma_r(t) \leq \sigma_r(s) \cdot \exp \left(2\sqrt{2} \int_{t'=s}^t \sqrt{\ell(\mathbf{W}_{\mathbf{A},\mathbf{B}}(t'))} dt' \right) \quad (68)$$

With the setting above, in the pre-train section, after T_1 timesteps, we prove that $\sigma_1(T_1) = \sigma_2(T_1) = w^*$. Starting from T_1 with pre-trained weights, we can lower bound $\sigma_2(\mathbf{W}_{\mathbf{A},\mathbf{B}}(t))$ with equations (66) and (68) when $t \geq T_1$ as follows:

$$\begin{aligned} \sigma_2(t) &\geq \sigma_2(T_1) \cdot \exp \left(-2\sqrt{2} \int_{t'=T_1}^t \sqrt{\ell(\mathbf{W}_{\mathbf{A},\mathbf{B}}(t'))} dt' \right) \\ &\geq w^* \cdot \exp \left(-2w_{12}^* \int_{t'=T_1}^t e^{-w^*(t'-T_1)} dt' \right) \\ &= w^* \cdot \exp \left(-\frac{2w_{12}^*}{w^*} \left(1 - e^{-w^*(t-T_1)} \right) \right) \end{aligned}$$

and when $t \rightarrow \infty$, $\sigma_2(\infty)$ can be lower bounded by:

$$\sigma_2(\infty) \geq w^* \cdot e^{-2 \cdot \frac{w_{12}^*}{w^*}}$$

In the same way, we can upper bound $\sigma_1(\infty)$ by:

$$\sigma_1(\infty) \leq w^* \cdot e^{2 \cdot \frac{w_{12}^*}{w^*}}$$

■

By Lemma 25, we can now lower bound the stable rank of a matrix $\mathbf{W}_{\mathbf{A},\mathbf{B}}(\infty)$:

$$\begin{aligned} \frac{\|\mathbf{W}_{\mathbf{A},\mathbf{B}}(\infty)\|_F^2}{\|\mathbf{W}_{\mathbf{A},\mathbf{B}}(\infty)\|_2^2} &= \frac{\sigma_1^2(\mathbf{W}_{\mathbf{A},\mathbf{B}}(\infty)) + \sigma_2^2(\mathbf{W}_{\mathbf{A},\mathbf{B}}(\infty))}{\sigma_1^2(\mathbf{W}_{\mathbf{A},\mathbf{B}}(\infty))} \\ &= 1 + \frac{\sigma_2^2(\mathbf{W}_{\mathbf{A},\mathbf{B}}(\infty))}{\sigma_1^2(\mathbf{W}_{\mathbf{A},\mathbf{B}}(\infty))} \\ &\geq 1 + \exp \left(-8 \frac{w_{12}^*}{w^*} \right), \end{aligned}$$

which concludes the proof of Theorem 6.

G.3. Formal Statement and Proof of Theorem 7

We now extend the preceding analysis to the general case involving a ground truth matrix $\mathbf{W}^* \in \mathbb{R}^{d \times d}$. The solution matrix $\mathbf{W}_{\mathbf{A}, \mathbf{B}} \in \mathbb{R}^{d \times d}$ is again factorized as $\mathbf{W}_{\mathbf{A}, \mathbf{B}} = \mathbf{A}\mathbf{B}$, where both $\mathbf{A}, \mathbf{B} \in \mathbb{R}^{d \times d}$. In this section, our detailed presentation and proof of Theorem 7 (from the main text) are structured as follows: we first introduce and prove Theorem 26, which is then followed by its direct consequence, Corollary 27.

We use the slightly modified loss function:

$$\mathcal{L}(\mathbf{A}, \mathbf{B}) = \frac{1}{2} \sum_{n=1}^N (\langle \mathbf{A}\mathbf{B}, \mathbf{X}_n \rangle - y_n)^2, \quad (69)$$

where the measurement matrix $\mathbf{X}_n = \mathbf{e}_{i_n} \mathbf{e}_{j_n}^\top$ represents a masking matrix, with the n -th observed entry set to one and all other entries set to zero, and $y_n \in \mathbb{R}$ denotes the ground truth value of the n -th observation. Then, by defining $\Theta = \begin{bmatrix} \mathbf{A} \\ \mathbf{B}^\top \end{bmatrix} \in \mathbb{R}^{2d \times d}$ and $\bar{\mathbf{X}}_n = \frac{1}{2} \begin{bmatrix} \mathbf{0} & \mathbf{X}_n \\ \mathbf{X}_n^\top & \mathbf{0} \end{bmatrix} \in \mathbb{R}^{2d \times 2d}$, we can rewrite the (69) as:

$$\begin{aligned} \mathcal{L}(\mathbf{A}, \mathbf{B}) &= \tilde{\mathcal{L}}(\Theta) = \frac{1}{2} \sum_{n=1}^N (\langle \Theta \Theta^\top, \bar{\mathbf{X}}_n \rangle - y_n)^2 \\ &= \frac{1}{2} \|F(\Theta) - \mathbf{y}\|_2^2. \end{aligned} \quad (70)$$

Here, $F(\Theta)$ and $\mathbf{y}$ represent vectors defined as:

$$F(\Theta) \triangleq \begin{bmatrix} \langle \Theta \Theta^\top, \bar{\mathbf{X}}_1 \rangle \\ \langle \Theta \Theta^\top, \bar{\mathbf{X}}_2 \rangle \\ \vdots \\ \langle \Theta \Theta^\top, \bar{\mathbf{X}}_N \rangle \end{bmatrix} \in \mathbb{R}^N, \quad \mathbf{y} \triangleq \begin{bmatrix} y_1 \\ y_2 \\ \vdots \\ y_N \end{bmatrix} \in \mathbb{R}^N. \quad (71)$$

By reparameterizing $\mathbf{A}, \mathbf{B}$ to Θ , and $\mathbf{X}_n$ to $\bar{\mathbf{X}}_n$, we can reduce the parameter matrices into a single matrix Θ while ensuring the symmetry of $\Theta \Theta^\top$. We train the model Θ via gradient flow, where the

loss evolution is given by:

$$\begin{aligned}
 \dot{\tilde{\mathcal{L}}}(\boldsymbol{\Theta}(t)) &= (F(\boldsymbol{\Theta}(t)) - \mathbf{y})^\top \dot{F}(\boldsymbol{\Theta}(t)) \\
 &= (F(\boldsymbol{\Theta}(t)) - \mathbf{y})^\top \begin{bmatrix} \frac{d}{dt} \langle \boldsymbol{\Theta}(t) \boldsymbol{\Theta}(t)^\top, \bar{\mathbf{X}}_1 \rangle \\ \frac{d}{dt} \langle \boldsymbol{\Theta}(t) \boldsymbol{\Theta}(t)^\top, \bar{\mathbf{X}}_2 \rangle \\ \vdots \\ \frac{d}{dt} \langle \boldsymbol{\Theta}(t) \boldsymbol{\Theta}(t)^\top, \bar{\mathbf{X}}_N \rangle \end{bmatrix} \\
 &= 2(F(\boldsymbol{\Theta}(t)) - \mathbf{y})^\top \begin{bmatrix} \langle \bar{\mathbf{X}}_1 \boldsymbol{\Theta}(t), \dot{\boldsymbol{\Theta}}(t) \rangle \\ \langle \bar{\mathbf{X}}_2 \boldsymbol{\Theta}(t), \dot{\boldsymbol{\Theta}}(t) \rangle \\ \vdots \\ \langle \bar{\mathbf{X}}_N \boldsymbol{\Theta}(t), \dot{\boldsymbol{\Theta}}(t) \rangle \end{bmatrix} \\
 &= 2(F(\boldsymbol{\Theta}(t)) - \mathbf{y})^\top \begin{bmatrix} \text{vec}(\bar{\mathbf{X}}_1 \boldsymbol{\Theta}(t))^\top \\ \text{vec}(\bar{\mathbf{X}}_2 \boldsymbol{\Theta}(t))^\top \\ \vdots \\ \text{vec}(\bar{\mathbf{X}}_N \boldsymbol{\Theta}(t))^\top \end{bmatrix} \text{vec}(\dot{\boldsymbol{\Theta}}(t)) \quad (72) \\
 &= (F(\boldsymbol{\Theta}(t)) - \mathbf{y})^\top J(\boldsymbol{\Theta}(t)) \text{vec}(\dot{\boldsymbol{\Theta}}(t)). \quad (73)
 \end{aligned}$$

Here, the Jacobian matrix $J(\boldsymbol{\Theta}(t))$ is defined as:

$$J(\boldsymbol{\Theta}(t)) \triangleq \frac{\partial F(\boldsymbol{\Theta}(t))}{\partial \text{vec}(\boldsymbol{\Theta}(t))} = \begin{bmatrix} \text{vec}(\nabla_{\boldsymbol{\Theta}} \langle \boldsymbol{\Theta}(t) \boldsymbol{\Theta}(t)^\top, \bar{\mathbf{X}}_1 \rangle)^\top \\ \text{vec}(\nabla_{\boldsymbol{\Theta}} \langle \boldsymbol{\Theta}(t) \boldsymbol{\Theta}(t)^\top, \bar{\mathbf{X}}_2 \rangle)^\top \\ \vdots \\ \text{vec}(\nabla_{\boldsymbol{\Theta}} \langle \boldsymbol{\Theta}(t) \boldsymbol{\Theta}(t)^\top, \bar{\mathbf{X}}_N \rangle)^\top \end{bmatrix} = 2 \begin{bmatrix} \text{vec}(\bar{\mathbf{X}}_1 \boldsymbol{\Theta}(t))^\top \\ \text{vec}(\bar{\mathbf{X}}_2 \boldsymbol{\Theta}(t))^\top \\ \vdots \\ \text{vec}(\bar{\mathbf{X}}_N \boldsymbol{\Theta}(t))^\top \end{bmatrix} \in \mathbb{R}^{N \times 2d^2}. \quad (74)$$

With the notations defined above, we state the following theorem:

Theorem 26 *Let the combined weight matrix be*

$$\boldsymbol{\Theta} \triangleq \begin{bmatrix} \mathbf{A} \\ \mathbf{B}^\top \end{bmatrix} \in \mathbb{R}^{2d \times d},$$

and consider the loss function $\tilde{\mathcal{L}}$ defined in (69). Denote

$$\sigma_{\min} \triangleq \sigma_{\min}(J(\boldsymbol{\Theta}(0))), \quad \sigma_{\max} \triangleq \sigma_{\max}(J(\boldsymbol{\Theta}(0))).$$

If the initialization satisfies:

$$\tilde{\mathcal{L}}(\boldsymbol{\Theta}(0)) \leq \frac{\sigma_{\min}^6}{1152d\sigma_{\max}^2},$$

then for every $t \geq 0$ the following hold:

$$\begin{aligned}
 \tilde{\mathcal{L}}(\boldsymbol{\Theta}(t)) &\leq \tilde{\mathcal{L}}(\boldsymbol{\Theta}(0)) \exp\left(-\frac{1}{2}\sigma_{\min}^2 t\right), \\
 \|\boldsymbol{\Theta}(t) - \boldsymbol{\Theta}(0)\|_F &\leq \frac{6\sqrt{2}\sigma_{\max}}{\sigma_{\min}^2} \sqrt{\tilde{\mathcal{L}}(\boldsymbol{\Theta}(0))}.
 \end{aligned}$$

The above theorem tells us that, if the model is initialized with a sufficiently small loss, the model's loss will converge to zero quickly, and the parameters will not move significantly from the initialization. With the above theorem, we can state the following corollary:

Corollary 27 *Suppose $\mathbf{A}$ and $\mathbf{B}$ are initialized as balanced, i.e.:*

$$\mathbf{A}(0)^\top \mathbf{A}(0) = \mathbf{B}(0) \mathbf{B}(0)^\top.$$

Under the conditions of Theorem 26, for every singular index $i \in [d]$ and all $t \geq 0$:

$$\sigma_i(\mathbf{A}(t)) = \sigma_i(\mathbf{B}(t)) \quad \text{and} \quad |\sigma_i(\mathbf{A}(t)) - \sigma_i(\mathbf{A}(0))| \leq \frac{\sigma_{\min}}{4\sqrt{2d}}.$$

Consequently, the stable rank of $\mathbf{A}(t)$ remains bounded below by

$$\frac{\|\mathbf{A}(t)\|_F^2}{\|\mathbf{A}(t)\|_2^2} \geq \left(\frac{\sqrt{2}\|\mathbf{A}(0)\|_F - \frac{\sigma_{\min}}{4\sqrt{2d}}}{\sqrt{2}\|\mathbf{A}(0)\|_2 + \frac{\sigma_{\min}}{4\sqrt{2d}}} \right)^2.$$

G.3.1. PROOF OF THEOREM 26

We begin the proof of the theorem by noting that the Jacobian $J(\cdot)$ is a Lipschitz function, as stated in the following lemma:

Lemma 28 *The Jacobian matrix $J(\mathbf{W})$, as defined in (74), is $\sqrt{d}$ -Lipschitz. Specifically, for any matrices $\mathbf{W}, \mathbf{V} \in \mathbb{R}^{d \times d}$, the following inequality holds:*

$$\|J(\mathbf{W}) - J(\mathbf{V})\| \leq \sqrt{d} \|\text{vec}(\mathbf{W}) - \text{vec}(\mathbf{V})\|, \quad (75)$$

Proof Note that for each n -th observation,

$$\begin{aligned} J_n(\Theta) &= 2\text{vec}(\bar{\mathbf{X}}_n \Theta)^\top \\ &= \text{vec} \left(\begin{pmatrix} 0 & \mathbf{X}_n \\ \mathbf{X}_n^\top & 0 \end{pmatrix} \begin{pmatrix} \mathbf{A} \\ \mathbf{B}^\top \end{pmatrix} \right)^\top \\ &= \text{vec} \left(\begin{pmatrix} \mathbf{X}_n \mathbf{B}^\top \\ \mathbf{X}_n^\top \mathbf{A} \end{pmatrix} \right)^\top \in \mathbb{R}^{2d^2}. \end{aligned}$$

Let $\mathbf{M}_l$ denote the l -th row of a matrix $\mathbf{M}$, and let $\mathbf{M}_{\cdot, l}$ denote its l -th column. We have

$$\begin{aligned} \|J_n(\Theta)\|_F^2 &= \|\mathbf{X}_n^\top \mathbf{A}\|_F^2 + \|\mathbf{X}_n \mathbf{B}^\top\|_F^2 \\ &= \|\mathbf{e}_{j_n} \mathbf{e}_{i_n}^\top \mathbf{A}\|_F^2 + \|\mathbf{e}_{i_n} \mathbf{e}_{j_n}^\top \mathbf{B}^\top\|_F^2 \\ &= \|\mathbf{A}_{i_n}\|_2^2 + \|\mathbf{B}_{\cdot, j_n}\|_2^2. \end{aligned}$$

Now, suppose we observe all entries, i.e., $N = d^2$. Then for any fixed n , $i_n = i_m$ can be satisfied for all $m \in [d]$, meaning each element of $\mathbf{A}$ is observed d times. Similarly, each element of $\mathbf{B}$ is also observed d times.

Therefore, we can upper bound the Frobenius norm of the Jacobian matrix by the Frobenius norm of the Jacobian under full observation:

$$\begin{aligned}\|J(\Theta)\|_F^2 &\leq \sum_{n=1}^{d^2} \left(\|\mathbf{X}_n^\top \mathbf{A}\|_F^2 + \|\mathbf{X}_n \mathbf{B}^\top\|_F^2 \right) \\ &= d \left(\|\mathbf{A}\|_F^2 + \|\mathbf{B}\|_F^2 \right) \\ &= d \|\Theta\|_F^2.\end{aligned}$$

By upper bounding the spectral norm of the difference between two Jacobian matrices and applying the inequality above, we obtain:

$$\begin{aligned}\|J(\mathbf{W}) - J(\mathbf{V})\|^2 &= \|J(\mathbf{W} - \mathbf{V})\|^2 \\ &\leq \|J(\mathbf{W} - \mathbf{V})\|_F^2 \\ &\leq d \|\mathbf{W} - \mathbf{V}\|_F^2,\end{aligned}$$

which concludes the proof. ■

Next, we borrow a lemma from Telgarsky [32], which states that for a Lipschitz function J , if we consider a sufficiently small neighborhood around the initialization $\Theta(0)$, then the singular values of the Jacobian $J(\Theta)$ remain close to those at initialization:

Lemma 29 (Lemma 8.3 in Telgarsky [32]) *If we suppose $\|\text{vec}(\Theta) - \text{vec}(\Theta(0))\| \leq \frac{\sigma_{\min}}{2\sqrt{d}}$, we have the following:*

$$\sigma_{\min}(J(\Theta)) \geq \frac{\sigma_{\min}}{2}, \quad \sigma_{\max}(J(\Theta)) \leq \frac{3\sigma_{\max}}{2},$$

where we denote $\sigma_{\min} \triangleq \sigma_{\min}(J(\Theta(0)))$, and $\sigma_{\max} \triangleq \sigma_{\max}(J(\Theta(0)))$.

For simplicity, we denote θ as the vectorized version of Θ , i.e., $\theta \triangleq \text{vec}(\Theta)$. We define the time step τ , which is the first time step when the trajectory of $\theta(t)$ touches the boundary:

$$\tau \triangleq \inf_{t \geq 0} \left\{ t \mid \|\theta(t) - \theta(0)\| \geq \frac{\sigma_{\min}}{2\sqrt{d}} \right\}.$$

We now demonstrate the convergence of the loss when $t \in [0, \tau]$ using the following lemma.

Lemma 30 *For all $t \in [0, \tau]$, the loss defined in (69) converges as follows:*

$$\tilde{\mathcal{L}}(\Theta(t)) \leq \tilde{\mathcal{L}}(\Theta(0)) \exp \left(-\frac{1}{2} \sigma_{\min}^2 t \right),$$

where we define $\sigma_{\min} \triangleq \sigma_{\min}(J(\Theta(0)))$.

Proof Recall that the time derivative of the loss can be written as follows, according to (73):

$$\begin{aligned}\dot{\tilde{\mathcal{L}}}(\Theta(t)) &= -(F(\Theta(t)) - \mathbf{y})^\top J(\Theta(t)) \dot{\theta}(t) \\ &= -(F(\Theta(t)) - \mathbf{y})^\top J(\Theta(t)) J(\Theta(t))^\top (F(\Theta(t)) - \mathbf{y}),\end{aligned}$$

noting that

$$\dot{\boldsymbol{\theta}}(t) = -\nabla_{\boldsymbol{\theta}(t)} \tilde{\mathcal{L}}(\boldsymbol{\Theta}(t)) = -J(\boldsymbol{\Theta}(t))^\top (F(\boldsymbol{\Theta}(t)) - \mathbf{y}).$$

By Lemma 29, for any $t \in [0, \tau]$, we can upper bound the above term as follows:

$$\begin{aligned} \dot{\tilde{\mathcal{L}}}(\boldsymbol{\Theta}(t)) &\leq -\lambda_{\min} \left(J(\boldsymbol{\Theta}(t)) J(\boldsymbol{\Theta}(t))^\top \right) \|F(\boldsymbol{\Theta}(t)) - \mathbf{y}\|^2 \\ &\leq -\frac{1}{2} \sigma_{\min}^2 \tilde{\mathcal{L}}(\boldsymbol{\Theta}(t)). \end{aligned}$$

Applying Grönwall's inequality gives:

$$\tilde{\mathcal{L}}(\boldsymbol{\Theta}(t)) \leq \tilde{\mathcal{L}}(\boldsymbol{\Theta}(0)) \exp \left(-\frac{1}{2} \sigma_{\min}^2 t \right) \quad \text{for } t \in [0, \tau].$$

■

The above lemma shows that the loss decays rapidly to zero if $\boldsymbol{\theta}(t)$ stays within a small neighborhood around the initialization. We now show that if the loss converges quickly near initialization, then $\boldsymbol{\theta}(t)$ does not move far from its initial value:

Lemma 31 *Let $\sigma_{\min} \triangleq \sigma_{\min}(J(\boldsymbol{\Theta}(0)))$ and $\sigma_{\max} \triangleq \sigma_{\max}(J(\boldsymbol{\Theta}(0)))$. For all $t \in [0, \tau]$, the distance between the weight vector at time t and the initial weight vector is bounded by:*

$$\|\boldsymbol{\theta}(t) - \boldsymbol{\theta}(0)\| \leq \frac{6\sqrt{2}\sigma_{\max}}{\sigma_{\min}^2} \sqrt{\tilde{\mathcal{L}}(\boldsymbol{\Theta}(0))}.$$

Proof We start by evaluating the distance between $\boldsymbol{\theta}(t)$ and $\boldsymbol{\theta}(0)$ using Lemma 29:

$$\begin{aligned} \|\boldsymbol{\theta}(t) - \boldsymbol{\theta}(0)\| &= \left\| \int_0^t \dot{\boldsymbol{\theta}}(s) \, ds \right\| \\ &= \int_0^t \left\| J(\boldsymbol{\Theta}(s))^\top (F(\boldsymbol{\Theta}(s)) - \mathbf{y}) \right\| \, ds \\ &\leq \int_0^t \sigma_{\max}(J(\boldsymbol{\Theta}(s))) \|F(\boldsymbol{\Theta}(s)) - \mathbf{y}\| \, ds \\ &\leq \frac{3}{2} \sigma_{\max} \int_0^t \|F(\boldsymbol{\Theta}(s)) - \mathbf{y}\| \, ds. \end{aligned}$$

By Lemma 30, we know that the objective function $\tilde{\mathcal{L}}(\boldsymbol{\Theta})$ satisfies:

$$\|F(\boldsymbol{\Theta}(t)) - \mathbf{y}\|^2 \leq \|F(\boldsymbol{\Theta}(0)) - \mathbf{y}\|^2 \exp \left(-\frac{1}{2} \sigma_{\min}^2 t \right).$$

Taking the square root of both sides, we obtain:

$$\|F(\boldsymbol{\Theta}(t)) - \mathbf{y}\| \leq \|F(\boldsymbol{\Theta}(0)) - \mathbf{y}\| \exp \left(-\frac{1}{4} \sigma_{\min}^2 t \right).$$

Substituting this into the previous inequality:

$$\begin{aligned}\|\boldsymbol{\theta}(t) - \boldsymbol{\theta}(0)\| &\leq \frac{3}{2}\sigma_{\max}\|F(\boldsymbol{\Theta}(0)) - \mathbf{y}\| \int_0^t \exp\left(-\frac{1}{4}\sigma_{\min}^2 s\right) ds \\ &\leq \frac{6\sigma_{\max}}{\sigma_{\min}^2}\|F(\boldsymbol{\Theta}(0)) - \mathbf{y}\|,\end{aligned}$$

where we used the fact that:

$$\int_0^t \exp(-Cs) ds \leq \frac{1}{C}, \quad \text{for } C > 0.$$

■

By combining Lemmas 30 and 31, we obtain the following results:

$$\tilde{\mathcal{L}}(\boldsymbol{\Theta}(t)) \leq \tilde{\mathcal{L}}(\boldsymbol{\Theta}(0)) \exp\left(-\frac{1}{2}\sigma_{\min}^2 t\right), \quad (76)$$

$$\|\boldsymbol{\theta}(t) - \boldsymbol{\theta}(0)\| \leq \frac{6\sqrt{2}\sigma_{\max}}{\sigma_{\min}^2} \sqrt{\tilde{\mathcal{L}}(\boldsymbol{\Theta}(0))}, \quad (77)$$

which hold for $t \in [0, \tau]$. If we can demonstrate that $\tau = \infty$, the proof is complete.

Actually, if we initialize $\boldsymbol{\Theta}(0)$ to satisfy the condition:

$$\tilde{\mathcal{L}}(\boldsymbol{\Theta}(0)) \leq \frac{\sigma_{\min}^6}{1152d\sigma_{\max}^2},$$

and substitute this condition into (77), we obtain an upper bound for $\|\boldsymbol{\theta}(t) - \boldsymbol{\theta}(0)\|$:

$$\|\boldsymbol{\theta}(t) - \boldsymbol{\theta}(0)\| \leq \frac{6\sqrt{2}\sigma_{\max}}{\sigma_{\min}^2} \frac{\sigma_{\min}^3}{\sqrt{1152d}\sigma_{\max}} = \frac{\sigma_{\min}}{4\sqrt{d}}.$$

Recall the definition of τ , which is the first time when $\boldsymbol{\theta}(t)$ touches the boundary of the small ball around the initialization:

$$\tau \triangleq \inf_{t \geq 0} \left\{ t \mid \|\boldsymbol{\theta}(t) - \boldsymbol{\theta}(0)\| \geq \frac{\sigma_{\min}}{2\sqrt{d}} \right\}.$$

However, with the condition $\tilde{\mathcal{L}}(\boldsymbol{\Theta}(0)) \leq \frac{\sigma_{\min}^6}{1152d\sigma_{\max}^2}$, $\boldsymbol{\theta}(t)$ cannot ever touch the boundary. This is because $\|\boldsymbol{\theta}(t) - \boldsymbol{\theta}(0)\|$ is bounded above by $\frac{\sigma_{\min}}{4\sqrt{d}}$, which is strictly less than $\frac{\sigma_{\min}}{2\sqrt{d}}$. Therefore, the parameter will remain inside the ball indefinitely, meaning $\tau = \infty$. This completes the proof of the theorem.

G.3.2. PROOF OF COROLLARY 27

First, we establish the equality $\sigma_i(\mathbf{A}(t)) = \sigma_i(\mathbf{B}(t))$ for all $i \in [d]$. Corollary 27 assumes that $\mathbf{A}(0)$ and $\mathbf{B}(0)$ are initialized as “balanced”, satisfying $\mathbf{A}(0)^\top \mathbf{A}(0) = \mathbf{B}(0)\mathbf{B}(0)^\top$. By Lemma 35, this balanced condition ensures that the singular values of $\mathbf{A}(t)$ and $\mathbf{B}(t)$ remain identical for all $t \geq 0$:

$$\sigma_i(\mathbf{A}(t)) = \sigma_i(\mathbf{B}(t)).$$

Second, we address the change in the singular values of a combined parameter matrix $\Theta(t)$ (related to $\mathbf{A}(t)$ and $\mathbf{B}(t)$). Theorem 26 states that under a specified condition on the initial loss, $\tilde{\mathcal{L}}(\Theta(0)) \leq \frac{\sigma_{\min}^6}{1152d\sigma_{\max}^2}$, the deviation of $\Theta(t)$ from its initialization $\Theta(0)$ is bounded for all $t \geq 0$ by:

$$\|\Theta(t) - \Theta(0)\|_F \leq \frac{\sigma_{\min}}{4\sqrt{d}}.$$

Let $K = \frac{\sigma_{\min}}{4\sqrt{d}}$. By Weyl's inequality, $|\sigma_i(\mathbf{X}) - \sigma_i(\mathbf{Y})| \leq \|\mathbf{X} - \mathbf{Y}\|_2$, and noting that $\|\cdot\|_2 \leq \|\cdot\|_F$, we have for all $i \in [d]$:

$$\begin{aligned} |\sigma_i(\Theta(t)) - \sigma_i(\Theta(0))| &\leq \|\Theta(t) - \Theta(0)\|_2 \\ &\leq \|\Theta(t) - \Theta(0)\|_F \\ &\leq K. \end{aligned}$$

This inequality allows us to establish bounds for $\|\Theta(t)\|_F$ (using reverse triangle inequality) and its largest singular value $\sigma_1(\Theta(t)) = \|\Theta(t)\|_2$:

$$\begin{aligned} \|\Theta(t)\|_F &\geq \|\Theta(0)\|_F - K, \\ \sigma_1(\Theta(t)) &\leq \sigma_1(\Theta(0)) + K. \end{aligned}$$

This yields the following lower bound on the stable rank of $\Theta(t)$:

$$\frac{\|\Theta(t)\|_F^2}{\|\Theta(t)\|_2^2} \geq \left(\frac{\|\Theta(0)\|_F - K}{\sigma_1(\Theta(0)) + K} \right)^2 = \left(\frac{\|\Theta(0)\|_F - \frac{\sigma_{\min}}{4\sqrt{d}}}{\|\Theta(0)\|_2 + \frac{\sigma_{\min}}{4\sqrt{d}}} \right)^2.$$

Furthermore, the balancedness condition implies $\mathbf{A}(t)^\top \mathbf{A}(t) = \mathbf{B}(t)\mathbf{B}(t)^\top$. By the definition of $\Theta(t)$, $\Theta(t)^\top \Theta(t) = \mathbf{A}(t)^\top \mathbf{A}(t) + \mathbf{B}(t)\mathbf{B}(t)^\top$, this leads to $\Theta(t)^\top \Theta(t) = 2\mathbf{A}(t)^\top \mathbf{A}(t)$. This relationship implies $\sigma_i(\Theta(t)) = \sqrt{2}\sigma_i(\mathbf{A}(t))$ for all i . Substituting this into the bounds for $\Theta(t)$, we have $\|\Theta(0)\|_F = \sqrt{2}\|\mathbf{A}(0)\|_F$ and $\sigma_1(\Theta(0)) = \sqrt{2}\sigma_1(\mathbf{A}(0)) = \sqrt{2}\|\mathbf{A}(0)\|_2$. This leads to the final lower bound on the stable rank of $\mathbf{A}(t)$ (which, by balancedness, is equal to that of $\mathbf{B}(t)$):

$$\frac{\|\mathbf{A}(t)\|_F^2}{\|\mathbf{A}(t)\|_2^2} \geq \left(\frac{\sqrt{2}\|\mathbf{A}(0)\|_F - K}{\sqrt{2}\|\mathbf{A}(0)\|_2 + K} \right)^2 = \left(\frac{\sqrt{2}\|\mathbf{A}(0)\|_F - \frac{\sigma_{\min}}{4\sqrt{d}}}{\sqrt{2}\|\mathbf{A}(0)\|_2 + \frac{\sigma_{\min}}{4\sqrt{d}}} \right)^2.$$

Appendix H. Useful Lemmas

Lemma 32 (Adaptation of Lemma 1 and Theorem 3 in [3]) *For any time t , the product matrix $\mathbf{W}(t) \in \mathbb{R}^{d,d}$ can be decomposed into its singular value decomposition:*

$$\mathbf{W}(t) = \sum_{r=1}^d \sigma_r(t) \mathbf{u}_r(t) \mathbf{v}_r(t)^\top$$

where $\sigma_r(t)$ are the singular values of $\mathbf{W}(t)$, and $\mathbf{u}_r(t)$, $\mathbf{v}_r(t)$ are the corresponding left and right singular vectors, respectively. Moreover, if $\mathbf{A}$, $\mathbf{B}$ are balanced at initialization, i.e.,

$$\mathbf{A}^\top(0)\mathbf{A}(0) = \mathbf{B}(0)\mathbf{B}^\top(0),$$

the time evolution of the singular values $\sigma_r(t)$ is represented as:

$$\dot{\sigma}_r(t) = -2 \cdot \sigma_r(t) \cdot \left\langle \nabla \ell(\mathbf{W}(t)), \mathbf{u}_r(t) \mathbf{v}_r(t)^\top \right\rangle, \quad r = 1, \dots, d \quad (78)$$

Lemma 33 *For any real-valued square matrix $\mathbf{A} \in \mathbb{R}^{d \times d}$, the absolute value of its determinant equals the product of its singular values:*

$$|\det(\mathbf{A})| = \prod_{r=1}^d \sigma_r$$

where σ_r are the singular values of $\mathbf{A}$.

Proof We express $\mathbf{A}$ using SVD: $\mathbf{A} = \mathbf{U}\mathbf{\Sigma}\mathbf{V}^\top$. Applying the determinant to both sides, we get:

$$\begin{aligned} \det(\mathbf{A}) &= \det(\mathbf{U}\mathbf{\Sigma}\mathbf{V}^\top) \\ &= \det(\mathbf{U}) \det(\mathbf{\Sigma}) \det(\mathbf{V}^\top) \end{aligned}$$

Here, $\mathbf{U}$ and $\mathbf{V}$ have orthonormal columns, and $\mathbf{\Sigma}$ is diagonal with singular values along its main diagonal. Since the determinant of an orthonormal matrix is either ± 1 ,

$$|\det(\mathbf{A})| = \det(\mathbf{\Sigma}) = \prod_{r=1}^d \sigma_r.$$

■

Lemma 34 (Determinant of $\mathbf{A}(t)$) *Consider a matrix $\mathbf{A}(t) \in \mathbb{R}^{d,d}$ initialized as $\det(\mathbf{A}(0)) > 0$. Then, $\det(\mathbf{A}(t)) > 0$ for all $t \geq 0$.*

Proof This follows directly from Lemma 32 and 33. Since the singular values are initialized as positive, and their evolution is continuous according to the given differential equation, they cannot become zero or negative. Therefore, $\mathbf{A}(t)$ maintains its sign of the determinant at initialization throughout the optimization process. ■

Lemma 35 (Adaptation of Lemma 8 in [27]) *Consider a product matrix $\mathbf{W}(t) = \mathbf{A}(t)\mathbf{B}(t) \in \mathbb{R}^{d \times d}$, where $\mathbf{A}(t)$ and $\mathbf{B}(t)$ are of equal size and balanced at initialization. Under these conditions, the following equality holds for all $t \geq 0$ and all singular values:*

$$\sigma_r(\mathbf{W}(t)) = \sigma_r(\mathbf{A}(t))^2 = \sigma_r(\mathbf{B}(t))^2$$

where $\sigma_r(\cdot)$ denotes the r -th singular value of the respective matrix where $r \in [d]$. Moreover, if $\det(\mathbf{A}(0))$ and $\det(\mathbf{B}(0))$ are both positive, then by Lemma 34, we can guarantee that for all $t \geq 0$:

$$\det(\mathbf{W}(t)) = \det(\mathbf{A}(t))^2 = \det(\mathbf{B}(t))^2$$

Lemma 36 (Adaptation of Theorem 1 in [3]) *Consider a product matrix $\mathbf{W}(t) = \mathbf{A}(t)\mathbf{B}(t) \in \mathbb{R}^{d \times d}$. We can guarantee $\mathbf{A}(t)$ and $\mathbf{B}(t)$ are analytic functions of t . As a result, $\mathbf{W}(t)$ is also an analytic function of t .*

Lemma 37 (Lemma 10 in Razin and Cohen [27]) *Let $f, g : [0, \infty] \rightarrow \mathbb{R}$ be real analytic functions such that $f^{(k)}(0) = g^{(k)}(0)$ for all $k \in \mathbb{N} \cup \{0\}$. Then, $f(t) = g(t)$ for all $t \geq 0$.*

Lemma 38 (Positive Semidefiniteness of $\mathbf{ABA}^\top$) *For matrices $\mathbf{A}, \mathbf{B} \in \mathbb{R}^{d,d}$, if $\mathbf{B}$ is positive semi-definite, then both $\mathbf{ABA}^\top$ and $\mathbf{A}^\top \mathbf{B} \mathbf{A}$ are positive semi-definite.*

Proof For any vector $\mathbf{x} \in \mathbb{R}^d$:

$$\mathbf{x}^\top \mathbf{ABA}^\top \mathbf{x} = (\mathbf{A}^\top \mathbf{x})^\top \mathbf{B}(\mathbf{A}^\top \mathbf{x}) \geq 0$$

since $\mathbf{B}$ is a positive semi-definite matrix. In the same way, for any vector $\mathbf{x} \in \mathbb{R}^d$ we have:

$$\mathbf{x}^\top \mathbf{A}^\top \mathbf{B} \mathbf{A} \mathbf{x} = (\mathbf{A} \mathbf{x})^\top \mathbf{B}(\mathbf{A} \mathbf{x}) \geq 0$$

which concludes the proof. ■