
Invariant Low-Dimensional Subspaces in Gradient Descent for Learning Deep Matrix Factorizations

Can Yaras

University of Michigan
cjyaras@umich.edu

Peng Wang

University of Michigan
pengwa@umich.edu

Wei Hu

University of Michigan
vvh@umich.edu

Zhihui Zhu

Ohio State University
zhu.3440@osu.edu

Laura Balzano

University of Michigan
girasole@umich.edu

Qing Qu

University of Michigan
qingqu@umich.edu

Abstract

An extensively studied phenomenon of the past few years in training deep networks is the implicit bias of gradient descent towards parsimonious solutions. In this work, we further investigate this phenomenon by narrowing our focus to deep matrix factorization, where we reveal surprising low-dimensional structures in the learning dynamics when the target matrix is low-rank. Specifically, we show that the evolution of gradient descent starting from arbitrary orthogonal initialization only affects a minimal portion of singular vector spaces across all weight matrices. In other words, the learning process happens only within a small invariant subspace of each weight matrix, despite the fact that all parameters are updated throughout training. From this, we provide rigorous justification for low-rank training in a specific, yet practical setting. In particular, we demonstrate that we can construct compressed factorizations that are equivalent to full-width, deep factorizations throughout training for solving low-rank matrix completion problems efficiently.

1 Introduction

In recent years, deep learning has demonstrated remarkable success across a wide range of applications [1]. Many recent works attempt to explain the exceptional generalization capabilities of deep networks by studying the implicit bias of gradient-based methods, showing that deep networks trained with such algorithms tend to learn simple functions [2–6]. Similarly, it has been shown that gradient descent induces max-margin [6, 7] or low-rank solutions [8–12] in deep networks, to name a few.

In another vein, recent work has explored the increasingly important problem of training deep networks more efficiently via *low-rank training* [13–17], where the number of trainable parameters is effectively reduced by replacing the original network weights with low-rank factorizations. While such methods have shown promising empirical results for reducing training time, theoretical justifications for these approaches remain deficient. The aforementioned works on implicit bias characterize low-rank structure in the limit of gradient descent – they do not address whether the trajectory of the original overparameterized network (along with its generalization/convergence properties) is achievable via low-rank factorization *from initialization* throughout training, which is what low-rank training necessitates.

Contributions. In this work, we draw theoretical connections between the implicit bias of gradient descent in deep networks and the practice of low-rank training. Utilizing deep matrix factorizations as a testbed (commonly assumed for analyzing the complex optimization dynamics of deep networks [10, 18–20]) we demonstrate that for low-rank data, all weight matrices are only updated within

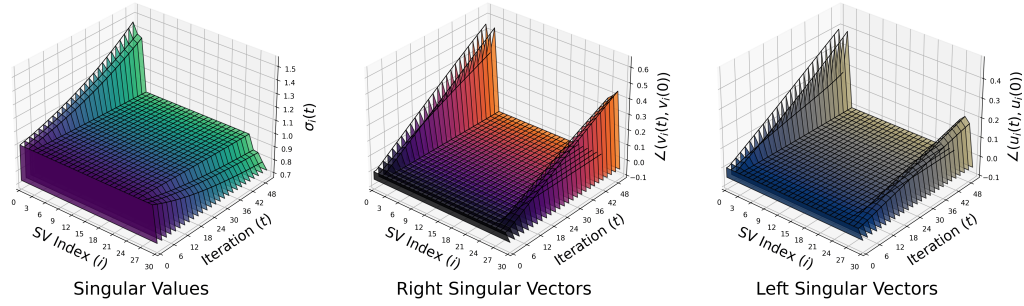


Figure 1: **Evolution of SVD of weight matrices.** We visualize the SVD dynamics of the first layer weight matrix of an $L = 3$ layer deep matrix factorization with $d = 30, r = 3, \sigma_l = 1$ throughout GD without weight decay. *Left:* Magnitude of the i -th singular value $\sigma_i(t)$ at iteration t . *Middle:* Angle $\angle(\mathbf{v}_i(t), \mathbf{v}_i(0))$ between the i -th right singular vector at iteration t and initialization. *Right:* Angle $\angle(\mathbf{u}_i(t), \mathbf{u}_i(0))$ between the i -th left singular vector at iteration t and initialization.

a low-dimensional subspace that is *invariant throughout training*, which can be determined from their arbitrary orthogonal initialization. To our knowledge, this is the first work identifying such invariant structures in gradient descent dynamics from random initialization. From this, we show that we can construct a compressed, low-rank factorization that is nearly equivalent to the original overparameterized network, thereby providing some rigorous foundations for low-rank training. Although deep matrix factorizations are mostly of theoretical interest, they are adopted for low-rank matrix sensing problems - therefore, we demonstrate that our theory can be applied (with slight modifications) towards accelerating practical low-rank deep matrix completion problems.

2 Analysis

Setup. We study the training dynamics of L -layer deep matrix factorizations $f(\Theta)$ given by

$$f(\Theta) := \mathbf{W}_L \mathbf{W}_{L-1} \cdots \mathbf{W}_2 \mathbf{W}_1$$

where $\Theta = (\mathbf{W}_l)_{l=1}^L$ are the parameters or weights with $\mathbf{W}_l \in \mathbf{R}^{d_l \times d_{l-1}}$ for $l \in [L]$. For a given target matrix $\Phi \in \mathbf{R}^{d \times d}$, we learn parameters Θ with $d_0 = d_1 = \cdots = d_L = d$ by minimizing the square loss

$$\ell(\Theta) = \frac{1}{2} \|f(\Theta) - \Phi\|_F^2 \quad (1)$$

via gradient descent (GD) from scaled *orthogonal* initialization, i.e., we initialize parameters $\Theta(0)$ such that all singular values of $\mathbf{W}_l(0)$ are equal to some $\sigma_l > 0$ for each $l \in [L]$. Then, we update all weights for $t = 0, 1, 2, \dots$ as

$$\mathbf{W}_l(t+1) = (1 - \eta\lambda)\mathbf{W}_l(t) - \eta \nabla_{\mathbf{W}_l} \ell(\Theta(t)), \quad l \in [L] \quad (2)$$

where $\lambda \geq 0$ is an optional weight decay parameter and $\eta > 0$ is the learning rate.

Main Result. Under the setting described above, we show learning only occurs within an invariant low-dimensional subspace of the weight matrices, provided that the target matrix Φ is low-rank.

Theorem 1. Suppose $\Phi \in \mathbf{R}^{d \times d}$ is at most rank r where $m := d - 2r > 0$. Then there exist orthogonal matrices $(\mathbf{U}_l)_{l=1}^L \subset \mathbf{R}^{d \times d}$ and $(\mathbf{V}_l)_{l=1}^L \subset \mathbf{R}^{d \times d}$ satisfying $\mathbf{V}_{l+1} = \mathbf{U}_l$ for $l \in [L-1]$, such that $\mathbf{W}_l(t)$ admits the decomposition

$$\mathbf{W}_l(t) = \mathbf{U}_l \begin{bmatrix} \widetilde{\mathbf{W}}_l(t) & \mathbf{0} \\ \mathbf{0} & \rho_l(t) \mathbf{I}_m \end{bmatrix} \mathbf{V}_l^\top \quad (3)$$

for all $l \in [L]$ and $t \geq 0$, where $\widetilde{\mathbf{W}}_l(t) \in \mathbf{R}^{2r \times 2r}$ with $\mathbf{W}_l(0) = \sigma_l \mathbf{I}_{2r}$, and

$$\rho_l(t) = \rho_l(t-1) \cdot (1 - \eta\lambda - \eta \cdot \prod_{k \neq l} \rho_k^2(t-1)) \quad (4)$$

for all $l \in [L]$ and $t \geq 1$ with $\rho_l(0) = \sigma_l$.

We defer the proof of Theorem 1 to Appendix A.1. In the following, we discuss several implications of our result and its relationship to previous work.

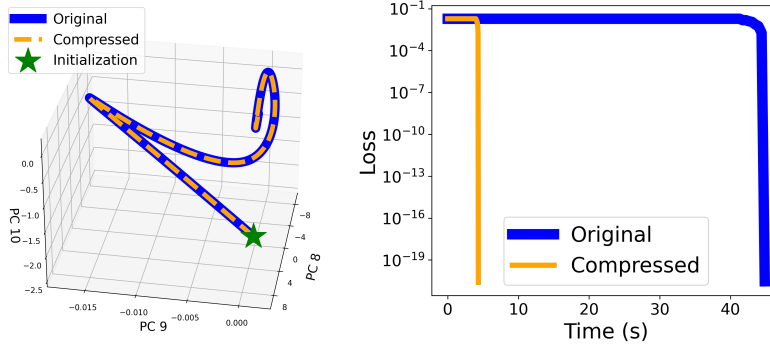


Figure 2: **Network compression for deep matrix factorization.** Comparison of trajectories for optimizing the original problem (1) vs. the compressed problem (6) with $L = 3$, $d = 1000$, $\hat{r} = r = 5$, and $\sigma_l = 10^{-3}$. *Left:* Principal components of end-to-end GD trajectories. *Right:* Training loss vs. wall-time comparison.

- **SVD dynamics of weight matrices.** The decomposition (3) implies that $\mathbf{W}_l(t)$ has m identical singular values that follow the updates given in (4), whose corresponding singular vectors are stationary from initialization throughout GD – this is portrayed in Figure 1. By this, we can decompose the total space $\mathbb{R}^d$ into two invariant singular subspaces: a $2r$ -dimensional space within which learning takes place, and an m -dimensional space corresponding to repeated singular values.
- **Low-rank bias.** From (4), we see that the GD trajectory either remains or tends towards a rank of at most $2r$ when we employ implicit or explicit regularization respectively. Indeed, if we use small initialization $\sigma_l \approx 0$, then the fact that ρ_l is a decreasing sequence implies that $\mathbf{W}_l(t)$ can be no more than rank $2r$ throughout the entire trajectory; whereas if $\lambda > 0$, then $\rho_l(t) \rightarrow 0$ as $t \rightarrow \infty$, which forces $\mathbf{W}_l(t)$ towards a solution of rank at most $2r$, regardless of initialization.
- **Comparison to prior arts.** In contrast to existing work that demonstrates the tendency of GD to find low nuclear-norm solutions [9, 11], our result directly shows that GD tends to find low-rank solutions. Moreover, unlike previous work on implicit bias [11, 21–23], we carefully examine the effect of weight decay, which is commonly employed during the training of deep networks. We note that our analysis is distinct from that of [18, 19], where continuous time dynamics are studied with the special (separable) setting $\mathbf{W}_{L:1}(0) = \mathbf{U}\mathbf{V}^\top$ with $\Phi = \mathbf{U}\Sigma\mathbf{V}^\top$. In contrast, our result applies to discrete time GD and holds for initialization that is agnostic to the target matrix. We also note that our result is unrelated to balanced initialization (as in [24]), since the σ_l can be arbitrarily different from one another.

Compressed Deep Matrix Factorization. We now show that, as a consequence of Theorem 1, we can run gradient descent on dramatically fewer parameters to achieve a near *identical* end-to-end trajectory to the original (full-width) factorization. More specifically, given an initialization $\Theta(0)$ of the original parameters and an upper bound on the rank $\hat{r} \geq r$ such that $d - 2\hat{r} > 0$, we define the *compressed* factorization

$$\hat{f}(\hat{\Theta}, \mathbf{U}_{L,1}, \mathbf{V}_{1,1}) := \mathbf{U}_{L,1} \hat{\mathbf{W}}_L \hat{\mathbf{W}}_{L-1} \cdots \hat{\mathbf{W}}_1 \mathbf{V}_{1,1}^\top \quad (5)$$

where $\hat{\Theta} = (\hat{\mathbf{W}}_l)_{l=1}^L$ are compressed weights with $\hat{\mathbf{W}}_l \in \mathbf{R}^{2\hat{r} \times 2\hat{r}}$ and $\mathbf{U}_{L,1}, \mathbf{V}_{1,1} \in \mathbf{R}^{d \times 2\hat{r}}$ are the first $2\hat{r}$ columns of $\mathbf{U}_L, \mathbf{V}_1 \in \mathbf{R}^{d \times d}$ respectively from Theorem 1 (depends on $\Theta(0)$ and Φ). Then, initializing $\hat{\Theta}(0)$ such that $\hat{\mathbf{W}}_l(0) = \mathbf{U}_{l,1}^\top \mathbf{W}_l(0) \mathbf{V}_{l,1}$ for all $l \in [L]$ and running gradient descent on the loss

$$\hat{\ell}(\hat{\Theta}) = \frac{1}{2} \|\hat{f}(\hat{\Theta}, \mathbf{U}_{L,1}, \mathbf{V}_{1,1}) - \Phi\|_F^2 \quad (6)$$

gives an almost equivalent network in the following sense.

Proposition 1. *For $\hat{r} \geq r$ such that $\hat{m} := d - 2\hat{r} > 0$, running gradient descent on the compressed weights $\hat{\Theta}$ as described above for the loss (6) satisfies*

$$\left\| f(\Theta(t)) - \hat{f}(\hat{\Theta}(t), \mathbf{U}_{L,1}, \mathbf{V}_{1,1}) \right\|_F^2 \leq \hat{m} \cdot \prod_{l=1}^L \sigma_l^2$$

for all iterates $t = 0, 1, 2, \dots$

We defer the proof of Proposition 1 to Appendix A.2. When we start from small initialization ($\sigma_l \approx 0$), Proposition 1 demonstrates that we only need to optimize $4L \cdot \hat{r}^2$ many parameters as opposed to the original $L \cdot d^2$ number of parameters to achieve an almost identical end-to-end trajectory, see Figure 2 (left). Since it is often the case that $r \leq \hat{r} \ll d$, this results in an order of magnitude reduction in time to reach an optimal solution compared to the original network, see Figure 2 (right). In the next section, we demonstrate how this idea can be leveraged (with slight modification) to accelerate a more practical problem.

3 Application: Accelerating Deep Low-Rank Matrix Completion

Problem Setup. We consider the low-rank matrix completion problem [25–27] with ground-truth $\Phi \in \mathbb{R}^{d \times d}$ with rank $r \ll d$, where the goal is to recover Φ from only a few number of observations encoded by a mask $\Omega \in \{0, 1\}^{d \times d}$. Adopting a deep matrix factorization approach [11], we minimize the objective

$$\ell_{\text{mc}}(\Theta) = \frac{1}{2} \|\Omega \odot (f(\Theta) - \Phi)\|_F^2 \quad (7)$$

which simplifies to (1) when $\Omega = \mathbf{1}_d \mathbf{1}_d^\top$ in the full observation case. In practice, the true rank r is not known – instead, we assume to have an upper bound $\hat{r}$ of the same order as r , i.e., $r \leq \hat{r} \ll d$.

Compressed Deep Matrix Completion. In the setting described above, it is advantageous to *over-parameterize* along both the depth and width of the factorization, particularly for accelerating GD convergence to well-generalizing solutions – see Appendix B for a more detailed discussion alongside evidence. Nonetheless, the advantages of over-parameterization are hindered by the fact that depth and width incur much higher *per-iteration* costs – for an L -layer factorization of (full)-width d , we require $O(L \cdot d^3)$ multiplications to evaluate gradients, where d is often very large. However, using ideas from the previous section, we can effectively reduce the computation to $O(\hat{r}^2 \cdot (L\hat{r} + d))$ multiplications via a compressed factorization that emulates the trajectory of a (full)-width d network, thereby enjoying accelerated GD convergence with heavily reduced per-iteration computational cost.

In the full observation case ($\Omega = \mathbf{1}_d \mathbf{1}_d^\top$), we have already seen via Proposition 1 that the compressed factorization (5) with small initialization stays close to the trajectory of the full-width factorization. However, applying this directly to the projection $\Omega \odot \Phi$ will result in the compressed factorization’s trajectory diverging from that of the original – see the orange trace in Figure 3. Intuitively, this is because the factors $U_{L,1}, V_{1,1}$ are initialized from incomplete measurement of Φ – instead, we optimize the modified objective

$$\hat{\ell}_{\text{mc}}(\hat{\Theta}, U_{L,1}, V_{1,1}) = \frac{1}{2} \|\Omega \odot (\hat{f}(\hat{\Theta}, U_{L,1}, V_{1,1}) - \Phi)\|_F^2 \quad (8)$$

where $\hat{\Theta}$ are updated with learning rate η while the $U_{L,1}, V_{1,1}$ factors are updated with a *discrepant* learning rate $\gamma\eta$ where $\gamma > 0$ is small. While this results in an additional $2d\hat{r}$ parameters to be tracked, the trajectory of this compressed factorization will ultimately align with that of the original while converging roughly $5\times$ faster w.r.t. wall-time, as demonstrated in Figure 3. Moreover, the accelerated convergence induced by the full-width trajectory results in the compressed factorization being $3\times$ faster than randomly initialized factorizations of similar width – see Appendix C for more details.

4 Conclusion

This paper offers novel insights into simple structures in gradient descent for learning deep matrix factorizations, from which we derive some rigorous justification for the practice of low-rank training. Through this work, we hope to inspire more principled approaches to designing efficient and effective deep models by exploiting low-dimensional aspects of their training dynamics. Moreover, we plan to apply this result to study progressive neural collapse [28, 29] and hierarchical feature learning [30].

Acknowledgement

CY and QQ acknowledge support from U-M START & PODS grants, NSF CAREER CCF-2143904, NSF CCF-2212066, and NSF CCF-2212326. QQ and PW also acknowledge support from ONR

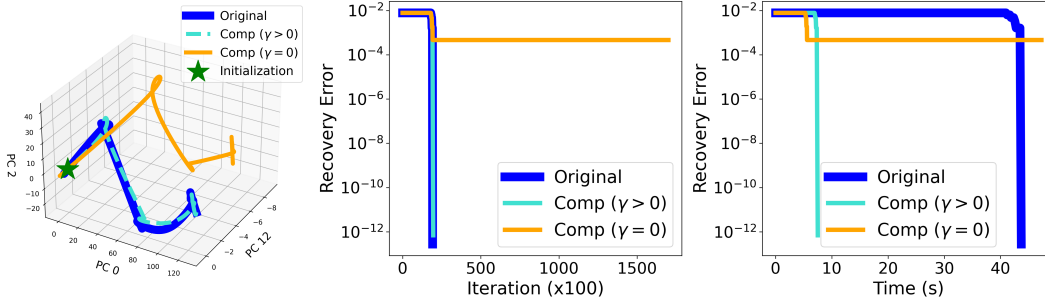


Figure 3: **Network compression for deep matrix completion.** Comparison of trajectories for optimizing the original problem (7) vs. the compressed problem (8) with γ discrepant updates ($\gamma = 0.01$) and ablating γ ($\gamma = 0$) with $L = 3$, $d = 1000$, $r = 5$, $\sigma_l = 10^{-3}$ and 20% of entries observed. *Left*: Principal components of end-to-end trajectories of each factorization. *Middle*: Recovery error vs. iteration comparison. *Right*: Recovery error vs wall-time comparison.

N00014-22-1-2529, NSF IIS 2312842, an AWS AI Award, and a gift grant from KLA. PW and LB acknowledge support from DoE award DE-SC0022186, ARO YIP W911NF1910027, and NSF CAREER CCF-1845076. ZZ acknowledges support from NSF grant CCF-2240708. WH acknowledges support from the Google Research Scholar Program.

References

- [1] Yann LeCun, Yoshua Bengio, and Geoffrey Hinton. Deep learning. *nature*, 521(7553):436–444, 2015.
- [2] Harshay Shah, Kaustav Tamuly, Aditi Raghunathan, Prateek Jain, and Praneeth Netrapalli. The pitfalls of simplicity bias in neural networks. *Advances in Neural Information Processing Systems*, 33:9573–9585, 2020.
- [3] Guillermo Valle-Perez, Chico Q. Camargo, and Ard A. Louis. Deep learning generalizes because the parameter-function map is biased towards simple functions. In *International Conference on Learning Representations*, 2019.
- [4] Suriya Gunasekar, Jason D Lee, Daniel Soudry, and Nati Srebro. Implicit bias of gradient descent on linear convolutional networks. *Advances in neural information processing systems*, 31, 2018.
- [5] Ziwei Ji and Matus Telgarsky. Gradient descent aligns the layers of deep linear networks. *arXiv preprint arXiv:1810.02032*, 2018.
- [6] Daniel Kunin, Atsushi Yamamura, Chao Ma, and Surya Ganguli. The asymmetric maximum margin bias of quasi-homogeneous neural networks. *arXiv preprint arXiv:2210.03820*, 2022.
- [7] Daniel Soudry, Elad Hoffer, Mor Shpigel Nacson, Suriya Gunasekar, and Nathan Srebro. The implicit bias of gradient descent on separable data. *The Journal of Machine Learning Research*, 19(1):2822–2878, 2018.
- [8] Minyoung Huh, Hossein Mobahi, Richard Zhang, Brian Cheung, Pulkit Agrawal, and Phillip Isola. The low-rank simplicity bias in deep networks. *Transactions on Machine Learning Research*, 2023.
- [9] 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.
- [10] Gauthier Gidel, Francis Bach, and Simon Lacoste-Julien. Implicit regularization of discrete gradient dynamics in linear neural networks. *Advances in Neural Information Processing Systems*, 32, 2019.
- [11] Sanjeev Arora, Nadav Cohen, Wei Hu, and Yuping Luo. Implicit regularization in deep matrix factorization. *Advances in Neural Information Processing Systems*, 32, 2019.
- [12] Zhiyuan Li, Yuping Luo, and Kaifeng Lyu. Towards resolving the implicit bias of gradient descent for matrix factorization: Greedy low-rank learning, 2020.
- [13] Samuel Horvath, Stefanos Laskaridis, Shashank Rajput, and Hongyi Wang. Maestro: Uncovering low-rank structures via trainable decomposition, 2023.
- [14] Jiawei Zhao, Yifei Zhang, Beidi Chen, Florian Schäfer, and Anima Anandkumar. Inrank: Incremental low-rank learning, 2023.
- [15] Hongyi Wang, Saurabh Agarwal, Pongsakorn U-chupala, Yoshiki Tanaka, Eric P. Xing, and Dimitris Papailiopoulos. Cuttlefish: Low-rank model training without all the tuning, 2023.
- [16] Edward J Hu, Yelong Shen, Phillip Wallis, Zeyuan Allen-Zhu, Yuanzhi Li, Shean Wang, Lu Wang, and Weizhu Chen. Lora: Low-rank adaptation of large language models. *arXiv preprint arXiv:2106.09685*, 2021.

- [17] Albert Gural, Phillip Nadeau, Mehul Tikekar, and Boris Murmann. Low-rank training of deep neural networks for emerging memory technology, 2020.
- [18] Andrew M Saxe, James L McClelland, and Surya Ganguli. Exact solutions to the nonlinear dynamics of learning in deep linear neural networks. *arXiv preprint arXiv:1312.6120*, 2013.
- [19] Andrew M Saxe, James L McClelland, and Surya Ganguli. A mathematical theory of semantic development in deep neural networks. *Proceedings of the National Academy of Sciences*, 116(23):11537–11546, 2019.
- [20] Sanjeev Arora, Nadav Cohen, and Elad Hazan. On the optimization of deep networks: Implicit acceleration by overparameterization. In *International Conference on Machine Learning*, pages 244–253. PMLR, 2018.
- [21] Hancheng Min, Salma Tarmoun, René Vidal, and Enrique Mallada. Convergence and implicit bias of gradient flow on overparametrized linear networks. *arXiv preprint arXiv:2105.06351*, 2022.
- [22] Daniel Gissin, Shai Shalev-Shwartz, and Amit Daniely. The implicit bias of depth: How incremental learning drives generalization. *arXiv preprint arXiv:1909.12051*, 2019.
- [23] Gal Vardi and Ohad Shamir. Implicit regularization in relu networks with the square loss. In *Conference on Learning Theory*, pages 4224–4258. PMLR, 2021.
- [24] Sanjeev Arora, Nadav Cohen, Noah Golowich, and Wei Hu. A convergence analysis of gradient descent for deep linear neural networks. *arXiv preprint arXiv:1810.02281*, 2018.
- [25] Emmanuel Candes and Benjamin Recht. Exact matrix completion via convex optimization. *Communications of the ACM*, 55(6):111–119, 2012.
- [26] Emmanuel J Candès and Terence Tao. The power of convex relaxation: Near-optimal matrix completion. *IEEE Transactions on Information Theory*, 56(5):2053–2080, 2010.
- [27] Mark A Davenport and Justin Romberg. An overview of low-rank matrix recovery from incomplete observations. *IEEE Journal of Selected Topics in Signal Processing*, 10(4):608–622, 2016.
- [28] Hangfeng He and Weijie J Su. A law of data separation in deep learning. *Proceedings of the National Academy of Sciences*, 120(36):e2221704120, 2023.
- [29] Can Yaras, Peng Wang, Zhihui Zhu, Laura Balzano, and Qing Qu. Neural collapse with normalized features: A geometric analysis over the riemannian manifold. *arXiv preprint arXiv:2209.09211*, 2022.
- [30] Peng Wang, Xiao Li, Can Yaras, Zhihui Zhu, Laura Balzano, Wei Hu, and Qing Qu. Understanding deep representation learning via layerwise feature compression and discrimination. *arXiv preprint arXiv:2311.02960*, 2023.
- [31] Samuel Burer and Renato DC Monteiro. A nonlinear programming algorithm for solving semidefinite programs via low-rank factorization. *Mathematical Programming*, 95(2):329–357, 2003.
- [32] Prateek Jain, Praneeth Netrapalli, and Sujay Sanghavi. Low-rank matrix completion using alternating minimization. In *Proceedings of the forty-fifth annual ACM symposium on Theory of computing*, pages 665–674, 2013.
- [33] Qinqing Zheng and John Lafferty. Convergence analysis for rectangular matrix completion using burer-monteiro factorization and gradient descent. *arXiv preprint arXiv:1605.07051*, 2016.
- [34] Ruoyu Sun and Zhi-Quan Luo. Guaranteed matrix completion via non-convex factorization. *IEEE Transactions on Information Theory*, 62(11):6535–6579, 2016.
- [35] Rong Ge, Jason D Lee, and Tengyu Ma. Matrix completion has no spurious local minimum. *Advances in neural information processing systems*, 29, 2016.
- [36] Srinadh Bhojanapalli, Behnam Neyshabur, and Nati Srebro. Global optimality of local search for low rank matrix recovery. *Advances in Neural Information Processing Systems*, 29, 2016.
- [37] Rong Ge, Chi Jin, and Yi Zheng. No spurious local minima in nonconvex low rank problems: A unified geometric analysis. In *International Conference on Machine Learning*, pages 1233–1242. PMLR, 2017.
- [38] Qiuwei Li, Zhihui Zhu, and Gongguo Tang. The non-convex geometry of low-rank matrix optimization. *Information and Inference: A Journal of the IMA*, 8(1):51–96, 2019.
- [39] Yuejie Chi, Yue M Lu, and Yuxin Chen. Nonconvex optimization meets low-rank matrix factorization: An overview. *IEEE Transactions on Signal Processing*, 67(20):5239–5269, 2019.
- [40] Yuanzhi Li, Tengyu Ma, and Hongyang Zhang. Algorithmic regularization in over-parameterized matrix sensing and neural networks with quadratic activations. In *Conference On Learning Theory*, pages 2–47. PMLR, 2018.
- [41] Mahdi Soltanolkotabi, Dominik Stöger, and Changzhi Xie. Implicit balancing and regularization: Generalization and convergence guarantees for overparameterized asymmetric matrix sensing. *arXiv preprint arXiv:2303.14244*, 2023.

Appendix

Notation. Given any $L \in \mathbb{N}$, we use $[L]$ to denote the index set $\{1, \dots, L\}$. We use $I_n \in \mathbb{R}^n$ to denote the identity matrix of size n , and $\mathbf{1}_n$ to denote a vector of length n with all entries equal to 1. We denote by $\|\mathbf{A}\|_F^2$ the squared *Frobenius* norm of matrix $\mathbf{A}$, i.e., the sum of squares of all entries of $\mathbf{A}$. For convenience, whenever $j > i$ we adopt the abbreviations $\mathbf{W}_{j:i} = \mathbf{W}_j \cdots \mathbf{W}_i$ and $\mathbf{W}_{j:i}^\top = \mathbf{W}_i^\top \cdots \mathbf{W}_j^\top$, whereas both are identity if $j < i$.

A Proofs in Section 2

Substituting the analytic form of the gradient into (2), we have the update rule

$$\mathbf{W}_l(t+1) = (1 - \eta\lambda)\mathbf{W}_l(t) - \eta\mathbf{W}_{L:l+1}^\top(t)\mathbf{E}(t)\mathbf{W}_{l-1:1}^\top(t), \quad l \in [L] \quad (9)$$

for $t = 0, 1, 2, \dots$, where $\mathbf{E}(t) = f(\boldsymbol{\Theta}(t)) - \Phi$.

We first establish the following Lemma 1 – the claim in Theorem 1 then follows in a relatively straightforward manner. We note that all statements quantified by i in this section implicitly hold for all $i \in [m]$ (as defined in Theorem 1) for the sake of notational brevity.

A.1 Proof of Theorem 1

Lemma 1. *Under the setting of Theorem 1, there exist orthonormal sets $\{\mathbf{u}_i^{(l)}\}_{i=1}^m \subset \mathbb{R}^d$ and $\{\mathbf{v}_i^{(l)}\}_{i=1}^m \subset \mathbb{R}^d$ for $l \in [L]$ satisfying $\mathbf{v}_i^{(l+1)} = \mathbf{u}_i^{(l)}$ for all $l \in [L-1]$ such that the following hold for all $t \geq 0$:*

$$\begin{aligned} \mathcal{A}(t) : \mathbf{W}_l(t)\mathbf{v}_i^{(l)} &= \rho_l(t)\mathbf{u}_i^{(l)} & \forall l \in [L], \\ \mathcal{B}(t) : \mathbf{W}_l^\top(t)\mathbf{u}_i^{(l)} &= \rho_l(t)\mathbf{v}_i^{(l)} & \forall l \in [L], \\ \mathcal{C}(t) : \Phi^\top \mathbf{W}_{L:l+1}(t)\mathbf{u}_i^{(l)} &= \mathbf{0} & \forall l \in [L], \\ \mathcal{D}(t) : \Phi \mathbf{W}_{l-1:1}^\top(t)\mathbf{v}_i^{(l)} &= \mathbf{0} & \forall l \in [L], \end{aligned}$$

where $\rho_l(t) = \rho_l(t-1) \cdot (1 - \eta\lambda - \eta \cdot \prod_{k \neq l} \rho_k(t-1)^2)$ for all $t \geq 1$ with $\rho_l(0) = \sigma_l > 0$.

Proof. Define $\Psi := \mathbf{W}_{L:2}^\top(0)\Phi$. Since the rank of Φ is at most r , we have that the rank of $\Psi \in \mathbb{R}^{d \times d}$ is at most r , which implies that $\dim \mathcal{N}(\Psi) = \dim \mathcal{N}(\Psi^\top) \geq d - r$. We define the subspace

$$\mathcal{S} := \mathcal{N}(\Psi) \cap \mathcal{N}(\Psi^\top \mathbf{W}_1(0)) \subset \mathbb{R}^d.$$

Since $\mathbf{W}_1(0) \in \mathbb{R}^{d \times d}$ is nonsingular, we have

$$\dim(\mathcal{S}) \geq 2(d - r) - d = m.$$

Let $\{\mathbf{v}_i^{(1)}\}_{i=1}^m$ denote an orthonormal set contained in $\mathcal{S}$ and set $\mathbf{u}_i^{(1)} := \mathbf{W}_1(0)\mathbf{v}_i^{(1)}/\sigma_1$, where $\sigma_1 > 0$ is the scale of $\mathbf{W}_1(0)$ – since $\mathbf{W}_1(0)/\sigma_1$ is orthogonal, $\{\mathbf{u}_i^{(1)}\}_{i=1}^m$ is also an orthonormal set. Then we trivially have $\mathbf{W}_1(0)\mathbf{v}_i^{(1)} = \sigma_1\mathbf{u}_i^{(1)}$, which implies $\mathbf{W}_1^\top(0)\mathbf{u}_i^{(1)} = \sigma_1\mathbf{v}_i^{(1)}$. It follows from $\mathbf{v}_i^{(1)} \in \mathcal{S}$ that $\Psi\mathbf{v}_i^{(1)} = \mathbf{0}$ and $\Psi^\top \mathbf{W}_1(0)\mathbf{v}_i^{(1)} = \mathbf{0}$, which is equivalent to $\mathbf{W}_{L:2}^\top(0)\Phi\mathbf{v}_i^{(1)} = \mathbf{0}$ and $\Phi^\top \mathbf{W}_{L:2}(0)\mathbf{W}_1(0)\mathbf{v}_i^{(1)} = \sigma_1\Phi^\top \mathbf{W}_{L:2}(0)\mathbf{u}_i^{(1)} = \mathbf{0}$ respectively. Since $\mathbf{W}_{L:2}^\top(0)$ is full column rank, we further have that $\Phi\mathbf{v}_i^{(1)} = \mathbf{0}$.

Now let $\mathcal{E}(l)$ be the event that we have orthonormal sets $\{\mathbf{u}_i^{(l)}\}_{i=1}^m$ and $\{\mathbf{v}_i^{(l)}\}_{i=1}^m$ satisfying $\mathbf{W}_l(0)\mathbf{v}_i^{(l)} = \sigma_l\mathbf{u}_i^{(l)}$, $\mathbf{W}_l^\top(0)\mathbf{u}_i^{(l)} = \sigma_l\mathbf{v}_i^{(l)}$, $\Phi^\top \mathbf{W}_{L:l+1}(0)\mathbf{u}_i^{(l)} = \mathbf{0}$, and $\Phi \mathbf{W}_{l-1:1}^\top(0)\mathbf{v}_i^{(l)} = \mathbf{0}$. From the above arguments, we have that $\mathcal{E}(1)$ holds – now suppose $\mathcal{E}(k)$ holds for some $1 \leq k < L$.

Set $\mathbf{v}_i^{(k+1)} := \mathbf{u}_i^{(k)}$ and $\mathbf{u}_i^{(k+1)} := \mathbf{W}_{k+1}(0)\mathbf{v}_i^{(k+1)}/\sigma_{k+1}$. This implies that $\mathbf{W}_{k+1}(0)\mathbf{v}_i^{(k+1)} = \sigma_{k+1}\mathbf{u}_i^{(k+1)}$ and $\mathbf{W}_{k+1}^\top(0)\mathbf{u}_i^{(k+1)} = \sigma_{k+1}\mathbf{v}_i^{(k+1)}$. Moreover, we have

$$\begin{aligned}\Phi^\top \mathbf{W}_{L:(k+1)+1}(0)\mathbf{u}_i^{(k+1)} &= \Phi^\top \mathbf{W}_{L:k+1}(0)\mathbf{W}_{k+1}^\top(0)\mathbf{u}_i^{(k+1)}/\sigma_{k+1}^2 \\ &= \Phi^\top \mathbf{W}_{L:k+1}(0)\mathbf{v}_i^{(k+1)}/\sigma_{k+1} \\ &= \Phi^\top \mathbf{W}_{L:k+1}(0)\mathbf{u}_i^{(k)}/\sigma_{k+1} = \mathbf{0},\end{aligned}$$

where the first two equalities follow from orthogonality and $\mathbf{u}_i^{(k+1)} = \mathbf{W}_{k+1}(0)\mathbf{v}_i^{(k+1)}/\sigma_{k+1}$, and the last equality is due to $\mathbf{v}_i^{(k+1)} = \mathbf{u}_i^{(k)}$. Similarly, we have

$$\begin{aligned}\Phi \mathbf{W}_{(k+1)-1:1}^\top(0)\mathbf{v}_i^{(k+1)} &= \Phi \mathbf{W}_{k-1:1}^\top(0)\mathbf{W}_k^\top(0)\mathbf{v}_i^{(k+1)} \\ &= \Phi \mathbf{W}_{k-1:1}^\top(0)\mathbf{W}_k^\top(0)\mathbf{u}_i^{(k)} \\ &= \sigma_k \Phi \mathbf{W}_{k-1:1}^\top(0)\mathbf{v}_i^{(k)} = \mathbf{0},\end{aligned}$$

where the second equality follows from $\mathbf{v}_i^{(k+1)} = \mathbf{u}_i^{(k)}$ and the third equality is due to $\mathbf{W}_k^\top(0)\mathbf{u}_i^{(k)} = \sigma_k\mathbf{v}_i^{(k)}$. Therefore $\mathcal{E}(k+1)$ holds, so we have $\mathcal{E}(l)$ for all $l \in [L]$. As a result, we have shown the base cases $\mathcal{A}(0)$, $\mathcal{B}(0)$, $\mathcal{C}(0)$, and $\mathcal{D}(0)$.

Now we proceed by induction on $t \geq 0$. Suppose that $\mathcal{A}(t)$, $\mathcal{B}(t)$, $\mathcal{C}(t)$, and $\mathcal{D}(t)$ hold for some $t \geq 0$. First, we show $\mathcal{A}(t+1)$ and $\mathcal{B}(t+1)$. We have

$$\begin{aligned}\mathbf{W}_l(t+1)\mathbf{v}_i^{(l)} &= [(1-\eta\lambda)\mathbf{W}_l(t) - \eta\mathbf{W}_{L:l+1}^\top(t)\mathbf{E}(t)\mathbf{W}_{l-1:1}^\top(t)]\mathbf{v}_i^{(l)} \\ &= [(1-\eta\lambda)\mathbf{W}_l(t) - \eta\mathbf{W}_{L:l+1}^\top(t)(\mathbf{W}_{L:1}(t) - \Phi)\mathbf{W}_{l-1:1}^\top(t)]\mathbf{v}_i^{(l)} \\ &= (1-\eta\lambda)\mathbf{W}_l(t)\mathbf{v}_i^{(l)} - \eta\mathbf{W}_{L:l+1}^\top(t)\mathbf{W}_{L:1}(t)\mathbf{W}_{l-1:1}^\top(t)\mathbf{v}_i^{(l)} \\ &= (1-\eta\lambda)\mathbf{W}_l(t)\mathbf{v}_i^{(l)} - \eta \cdot \left(\prod_{k \neq l} \rho_k^2(t)\right)\mathbf{W}_l(t)\mathbf{v}_i^{(l)} \\ &= \rho_l(t) \cdot (1-\eta\lambda - \eta \cdot \prod_{k \neq l} \rho_k^2(t))\mathbf{u}_i^{(l)} = \rho_l(t+1)\mathbf{u}_i^{(l)}\end{aligned}$$

for all $l \in [L]$, where the first equality follows from (9), the second equality follows from definition of $\mathbf{E}(t)$, the third equality follows from $\mathcal{D}(t)$, and the fourth equality follows from $\mathcal{A}(t)$ and $\mathcal{B}(t)$ applied repeatedly along with $\mathbf{v}_i^{(l+1)} = \mathbf{u}_i^{(l)}$ for all $l \in [L-1]$, proving $\mathcal{A}(t+1)$. Similarly, we have

$$\begin{aligned}\mathbf{W}_l^\top(t+1)\mathbf{u}_i^{(l)} &= [(1-\eta\lambda)\mathbf{W}_l^\top(t) - \eta\mathbf{W}_{l-1:1}(t)\mathbf{E}^\top(t)\mathbf{W}_{L:l+1}(t)]\mathbf{u}_i^{(l)} \\ &= [(1-\eta\lambda)\mathbf{W}_l^\top(t) - \eta\mathbf{W}_{l-1:1}(t)(\mathbf{W}_{L:1}^\top(t) - \Phi^\top)\mathbf{W}_{L:l+1}(t)]\mathbf{u}_i^{(l)} \\ &= (1-\eta\lambda)\mathbf{W}_l^\top(t)\mathbf{u}_i^{(l)} - \eta\mathbf{W}_{l-1:1}(t)\mathbf{W}_{L:1}^\top(t)\mathbf{W}_{L:l+1}(t)\mathbf{u}_i^{(l)} \\ &= (1-\eta\lambda)\mathbf{W}_l^\top(t)\mathbf{u}_i^{(l)} - \eta \cdot \left(\prod_{k \neq l} \rho_k^2(t)\right)\mathbf{W}_l^\top(t)\mathbf{u}_i^{(l)} \\ &= \rho_l(t) \cdot (1-\eta\lambda - \eta \cdot \prod_{k \neq l} \rho_k^2(t))\mathbf{v}_i^{(l)} = \rho_l(t+1)\mathbf{v}_i^{(l)}\end{aligned}$$

for all $l \in [L]$, where the third equality follows from $\mathcal{C}(t)$, and the fourth equality follows from $\mathcal{A}(t)$ and $\mathcal{B}(t)$ applied repeatedly along with $\mathbf{v}_i^{(l+1)} = \mathbf{u}_i^{(l)}$ for all $l \in [L-1]$, proving $\mathcal{B}(t+1)$. Now, we show $\mathcal{C}(t+1)$. For any $k \in [L-1]$, it follows from $\mathbf{v}_i^{(k+1)} = \mathbf{u}_i^{(k)}$ and $\mathcal{A}(t+1)$ that

$$\mathbf{W}_{k+1}(t+1)\mathbf{u}_i^{(k)} = \mathbf{W}_{k+1}(t+1)\mathbf{v}_i^{(k+1)} = \rho_{k+1}(t+1)\mathbf{u}_i^{(k+1)}.$$

Repeatedly applying the above equality for $k = l, l+1, \dots, L-1$, we obtain

$$\Phi^\top \mathbf{W}_{L:l+1}(t)\mathbf{u}_i^{(l)} = \left[\prod_{k=l}^{L-1} \rho_{k+1}(t) \right] \cdot \Phi^\top \mathbf{u}_i^{(L)} = \mathbf{0}$$

which follows from $\mathcal{C}(t)$, proving $\mathcal{C}(t+1)$. Finally, we show $\mathcal{D}(t+1)$. For any $k \in \{2, \dots, L\}$, it follows from $\mathbf{v}_i^{(k)} = \mathbf{u}_i^{(k-1)}$ and $\mathcal{B}(t+1)$ that

$$\mathbf{W}_{k-1}^\top(t+1)\mathbf{v}_i^{(k)} = \mathbf{W}_{k-1}^\top(t+1)\mathbf{u}_i^{(k-1)} = \rho_{k-1}(t+1)\mathbf{v}_i^{(k-1)}.$$

Repeatedly applying the above equality for $k = l, l-1, \dots, 2$, we obtain

$$\Phi \mathbf{W}_{l-1:1}^\top(t)\mathbf{v}_i^{(l)} = \left[\prod_{k=2}^l \rho_{k-1}(t) \right] \cdot \Phi \mathbf{v}_i^{(1)} = \mathbf{0}$$

which follows from $\mathcal{D}(t)$. Thus we have proven $\mathcal{D}(t+1)$, concluding the proof. $\square$

Proof of Theorem 1. By $\mathcal{A}(t)$ and $\mathcal{B}(t)$ of Lemma 1, there exists orthonormal matrices $\{\mathbf{U}_{l,2}\}_{l=1}^L \subset \mathbf{R}^{d \times m}$ and $\{\mathbf{V}_{l,2}\}_{l=1}^L \subset \mathbf{R}^{d \times m}$ for $l \in [L]$ satisfying $\mathbf{U}_{l+1,2} = \mathbf{V}_{l,2}$ for all $l \in [L-1]$ as well as

$$\mathbf{W}_l(t)\mathbf{V}_{l,2} = \rho_l(t)\mathbf{U}_{l,2} \quad \text{and} \quad \mathbf{W}_l(t)^\top \mathbf{U}_{l,2} = \rho_l(t)\mathbf{V}_{l,2} \quad (10)$$

for all $l \in [L]$ and $t \geq 0$, where $\rho_l(t)$ satisfies (4) for $t \geq 1$ with $\rho_l(0) = \sigma_l$. First, complete $\mathbf{V}_{l,2}$ to an orthonormal basis for $\mathbf{R}^d$ as $\mathbf{V}_l = [\mathbf{V}_{l,1} \ \mathbf{V}_{l,2}]$. Then for each $l \in [L-1]$, set $\mathbf{U}_l = [\mathbf{U}_{l,1} \ \mathbf{U}_{l,2}]$ where $\mathbf{U}_{l,1} = \mathbf{W}_l(0)\mathbf{V}_{l,1}/\sigma_l$ and $\mathbf{V}_{l+1} = [\mathbf{V}_{l+1,1} \ \mathbf{V}_{l+1,2}]$ where $\mathbf{V}_{l+1,1} = \mathbf{U}_{l,1}$, and finally set $\mathbf{U}_L = [\mathbf{U}_{L,1} \ \mathbf{U}_{L,2}]$ where $\mathbf{U}_{L,1} = \mathbf{W}_L(0)\mathbf{V}_{L,1}/\sigma_L$. We note that $\mathbf{V}_{l+1} = \mathbf{U}_l$ for each $l \in [L-1]$. Then we have

$$\mathbf{U}_{l,1}^\top \mathbf{W}_l(t)\mathbf{V}_{l,2} = \rho_l(t)\mathbf{U}_{l,1}^\top \mathbf{U}_{l,2} = \mathbf{0} \quad (11)$$

for all $l \in [L]$, where the first equality follows from (10). Similarly, we also have

$$\mathbf{U}_{l,2}^\top \mathbf{W}_l(t)\mathbf{V}_{l,1} = \rho_l(t)\mathbf{V}_{l,2}^\top \mathbf{V}_{l,1} = \mathbf{0} \quad (12)$$

for all $l \in [L]$, where the first equality also follows from (10). Therefore, combining (10), (11), and (12) yields

$$\mathbf{U}_l^\top \mathbf{W}_l(t)\mathbf{V}_l = [\mathbf{U}_{l,1} \ \mathbf{U}_{l,2}]^\top \mathbf{W}_l(t) [\mathbf{V}_{l+1,1} \ \mathbf{V}_{l+1,2}] = \begin{bmatrix} \widetilde{\mathbf{W}}_l(t) & \mathbf{0} \\ \mathbf{0} & \rho_l(t)\mathbf{I}_m \end{bmatrix}$$

for all $l \in [L]$, where $\widetilde{\mathbf{W}}_l(0) = \sigma_l \mathbf{I}_{2r}$ by construction of $\mathbf{U}_{l,1}$. This directly implies (3), completing the proof. $\square$

A.2 Proof of Proposition 1

Proof. First, it follows from Theorem 1 that for any $1 \leq i \leq j \leq L$ we have

$$\mathbf{W}_{j:i}(t) = \mathbf{U}_{j,1} \widetilde{\mathbf{W}}_{j:i}(t) \mathbf{V}_{i,1}^\top + \left(\prod_{k=i}^j \rho_k(t) \right) \cdot \mathbf{U}_{j,2} \mathbf{V}_{i,2}^\top \quad (13)$$

for all $t \geq 0$, where $\mathbf{U}_{l,1}, \mathbf{V}_{l,1} \in \mathbf{R}^{d \times 2\hat{r}}$ and $\mathbf{U}_{l,2}, \mathbf{V}_{l,2} \in \mathbf{R}^{d \times \hat{m}}$ are the first $2\hat{r}$ and last $\hat{m}$ columns of $\mathbf{U}_l, \mathbf{V}_l \in \mathbf{R}^{d \times d}$ respectively.

The key claim to be shown here is that $\widehat{\mathbf{W}}_l(t) = \widetilde{\mathbf{W}}_l(t)$ for all $l \in [L]$ and $t \geq 0$. Afterwards, it follows straightforwardly from (13) that

$$\begin{aligned} & \left\| f(\Theta(t)) - \widehat{f}(\widehat{\Theta}(t), \mathbf{U}_{L,1}, \mathbf{V}_{1,1}) \right\|_F^2 \\ &= \left\| \mathbf{U}_{L,1} \widetilde{\mathbf{W}}_{L:1}(t) \mathbf{V}_{1,1}^\top + \left(\prod_{l=1}^L \rho_l(t) \right) \cdot \mathbf{U}_{L,2} \mathbf{V}_{1,2}^\top - \mathbf{U}_{L,1} \widehat{\mathbf{W}}_{L:1}(t) \mathbf{V}_{1,1}^\top \right\|_F^2 \\ &= \left\| \mathbf{U}_{L,1} (\widetilde{\mathbf{W}}_{L:1}(t) - \widehat{\mathbf{W}}_{L:1}(t)) \mathbf{V}_{1,1}^\top + \left(\prod_{l=1}^L \rho_l(t) \right) \cdot \mathbf{U}_{L,2} \mathbf{V}_{1,2}^\top \right\|_F^2 \\ &= \left\| \left(\prod_{l=1}^L \rho_l(t) \right) \cdot \mathbf{U}_{L,2} \mathbf{V}_{1,2}^\top \right\|_F^2 \leq \hat{m} \cdot \prod_{l=1}^L \sigma_l^2. \end{aligned}$$

We proceed by induction. For $t = 0$, we have that

$$\widehat{\mathbf{W}}_l(0) = \mathbf{U}_{l,1}^\top \mathbf{W}_l(0) \mathbf{V}_{l,1} = \widetilde{\mathbf{W}}_l(0)$$

for all $l \in [L]$ by (13) and choice of initialization.

Now suppose $\widehat{\mathbf{W}}_l(t) = \widetilde{\mathbf{W}}_l(t)$ for all $l \in [L]$. Comparing

$$\widehat{\mathbf{W}}_l(t+1) = (1 - \eta\lambda)\widehat{\mathbf{W}}_l(t) - \eta \nabla_{\widehat{\mathbf{W}}_l} \widehat{\ell}(\widehat{\Theta}(t))$$

with

$$\begin{aligned} \widetilde{\mathbf{W}}_l(t+1) &= \mathbf{U}_{l,1}^\top \mathbf{W}_l(t+1) \mathbf{V}_{l,1} \\ &= \mathbf{U}_{l,1}^\top [(1 - \eta\lambda)\mathbf{W}_l(t) - \eta \nabla_{\mathbf{W}_l} \ell(\Theta(t))] \mathbf{V}_{l,1} \\ &= (1 - \eta\lambda)\widetilde{\mathbf{W}}_l(t) - \eta \mathbf{U}_{l,1}^\top \nabla_{\mathbf{W}_l} \ell(\Theta(t)) \mathbf{V}_{l,1} \end{aligned}$$

it suffices to show that

$$\nabla_{\widehat{\mathbf{W}}_l} \widehat{\ell}(\widehat{\Theta}(t)) = \mathbf{U}_{l,1}^\top \nabla_{\mathbf{W}_l} \ell(\Theta(t)) \mathbf{V}_{l,1}, \quad \forall l \in [L] \quad (14)$$

to yield $\widehat{\mathbf{W}}_l(t+1) = \widetilde{\mathbf{W}}_l(t+1)$ for all $l \in [L]$. Computing the right hand side of (14), we have

$$\begin{aligned} \mathbf{U}_{l,1}^\top \nabla_{\mathbf{W}_l} \ell(\Theta(t)) \mathbf{V}_{l,1} &= \mathbf{U}_{l,1}^\top \mathbf{W}_{L:l+1}^\top(t) (\mathbf{W}_{L:1}(t) - \Phi) \mathbf{W}_{l-1:1}^\top(t) \mathbf{V}_{l,1} \\ &= (\mathbf{W}_{L:l+1}(t) \mathbf{U}_{l,1})^\top (\mathbf{W}_{L:1}(t) - \Phi) (\mathbf{V}_{l,1}^\top \mathbf{W}_{l-1:1}(t))^\top \end{aligned}$$

where

$$\mathbf{W}_{L:l+1}(t) \mathbf{U}_{l,1} = \left(\mathbf{U}_{L,1} \widetilde{\mathbf{W}}_{L:l+1}(t) \mathbf{V}_{l+1,1}^\top + \left(\prod_{k=l+1}^L \rho_k(t) \right) \cdot \mathbf{U}_{L,2} \mathbf{V}_{l+1,2}^\top \right) \mathbf{U}_{l,1} = \mathbf{U}_{L,1} \widetilde{\mathbf{W}}_{L:l+1}(t)$$

by (13) and the fact that $\mathbf{U}_l = \mathbf{V}_{l+1}$, and similarly

$$\mathbf{V}_{l,1}^\top \mathbf{W}_{l-1:1}(t) = \mathbf{V}_{l,1}^\top \left(\mathbf{U}_{l-1,1} \widetilde{\mathbf{W}}_{l-1:1}(t) \mathbf{V}_{1,1}^\top + \left(\prod_{k=1}^{l-1} \rho_k(t) \right) \cdot \mathbf{U}_{l-1,2} \mathbf{V}_{1,2}^\top \right) = \widetilde{\mathbf{W}}_{l-1:1}(t) \mathbf{V}_{1,1}^\top.$$

We also have that

$$\begin{aligned} \mathbf{U}_{L,1}^\top (\mathbf{W}_{L:1}(t) - \Phi) \mathbf{V}_{1,1} &= \mathbf{U}_{L,1}^\top \left(\mathbf{U}_{L,1} \widetilde{\mathbf{W}}_{L:1}(t) \mathbf{V}_{1,1}^\top + \left(\prod_{k=1}^L \rho_k(t) \right) \cdot \mathbf{U}_{L,2} \mathbf{V}_{1,2}^\top - \Phi \right) \mathbf{V}_{1,1} \\ &= \widetilde{\mathbf{W}}_{L:1}(t) - \mathbf{U}_{L,1}^\top \Phi \mathbf{V}_{1,1} \end{aligned}$$

so putting together the previous four equalities yields

$$\begin{aligned} \mathbf{U}_{l,1}^\top \nabla_{\mathbf{W}_l} \ell(\Theta(t)) \mathbf{V}_{l,1} &= (\mathbf{W}_{L:l+1}(t) \mathbf{U}_{l,1})^\top (\mathbf{W}_{L:1}(t) - \Phi) (\mathbf{V}_{l,1}^\top \mathbf{W}_{l-1:1}(t))^\top \\ &= \widetilde{\mathbf{W}}_{L:l+1}^\top(t) \mathbf{U}_{L,1}^\top (\mathbf{W}_{L:1}(t) - \Phi) \mathbf{V}_{1,1} \widetilde{\mathbf{W}}_{l-1:1}^\top(t) \\ &= \widetilde{\mathbf{W}}_{L:l+1}^\top(t) (\widetilde{\mathbf{W}}_{L:1}(t) - \mathbf{U}_{L,1}^\top \Phi \mathbf{V}_{1,1}) \widetilde{\mathbf{W}}_{l-1:1}^\top(t). \end{aligned}$$

On the other hand, the left hand side of (14) gives

$$\begin{aligned} \nabla_{\widehat{\mathbf{W}}_l} \widehat{\ell}(\widehat{\Theta}(t)) &= \widehat{\mathbf{W}}_{L:l+1}(t)^\top \mathbf{U}_{L,1}^\top (\mathbf{U}_{L,1} \widehat{\mathbf{W}}_{L:1}(t) \mathbf{V}_{1,1}^\top - \Phi) \mathbf{V}_{1,1} \widehat{\mathbf{W}}_{l-1:1}(t)^\top \\ &= \widehat{\mathbf{W}}_{L:l+1}(t)^\top (\widehat{\mathbf{W}}_{L:1}(t) - \mathbf{U}_{L,1}^\top \Phi \mathbf{V}_{1,1}) \widehat{\mathbf{W}}_{l-1:1}(t)^\top \end{aligned}$$

so (14) holds by the fact that $\widehat{\mathbf{W}}_l(t) = \widetilde{\mathbf{W}}_l(t)$ for all $l \in [L]$, completing the proof. $\square$

B Benefits of Over-Parameterization in Deep Matrix Completion

In the setup described in Section 3, we claim that depth and width are beneficial for accelerating GD convergence to well-generalizing solutions, and therefore constructing more computationally efficient factorizations that share the same trajectory is a fruitful endeavour. Below, we give a more detailed explanation of these ideas:

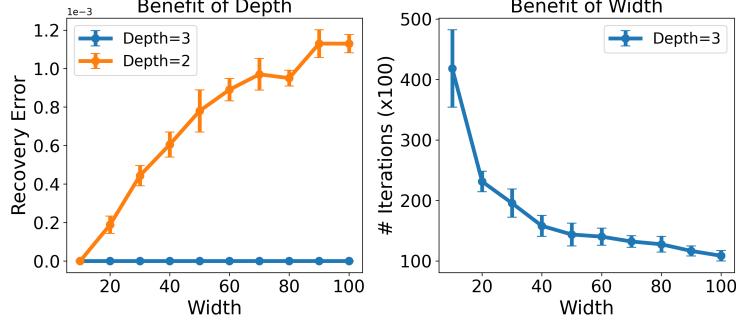


Figure 4: **Benefits of depth & width in overparameterized matrix completion** with $d = 100$, $r = 5$, $\sigma_l = 10^{-3}$ and 30% of entries observed. *Left*: Recovery error. *Right*: Number of GD iterations to converge to 10^{-10} error.

- **Benefits of depth.** When $L = 2$, (7) reduces to Burer-Monteiro factorization [31], whose global optimality and convergence under GD have been widely studied under various settings [9, 32–41]. However, it has been demonstrated [11] that in the over-parameterized regime $\hat{r} > r$, deeper factorizations (starting from small random initialization) continue to generalize well beyond the exact parameterization $\hat{r} = r$ unlike their shallow counterparts, see Figure 4 (left).
- **Benefits of width.** On the other hand, increasing the width $\hat{r}$ of the deep factorization beyond r results in accelerated convergence of GD in terms of iterations, see Figure 4 (right).

C Compressed vs. Narrow Factorizations

We compare the training efficiency of a $2\hat{r}$ -compressed factorization (with trajectory equivalent to a wide factorization of width $d \gg \hat{r}$) versus a narrow factorization with width $2\hat{r}$ under different over-parameterized estimates $\hat{r}$. As depicted in Figure 5 (left), the compressed factorization requires fewer iterations to reach convergence, and the number of iterations necessary is almost unaffected by $\hat{r}$. Consequently, **training compressed factorizations is considerably more time-efficient than training narrow ones of the same size**, provided that $\hat{r}$ is not significantly larger than r .

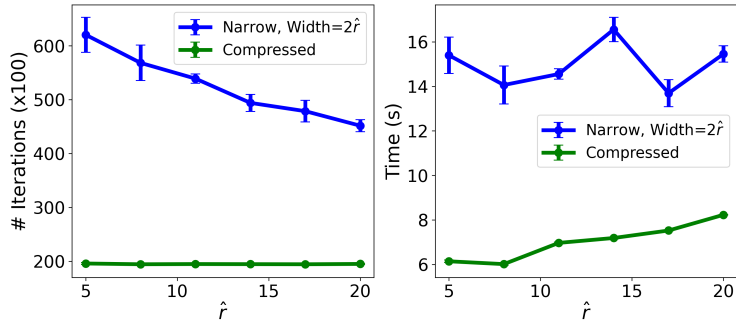


Figure 5: **Efficiency of compressed vs. narrow factorizations** for different overestimated $\hat{r}$ with $L = 3$, $d = 1000$, $r = 5$, $\sigma_l = 10^{-3}$ and 20% of entries observed. *Left*: Number of iterations to converge to 10^{-10} error. *Right*: Time to converge.

The distinction between the compressed and narrow factorizations underscores the benefits of width, as previously demonstrated and discussed in Figure 4 (right), where increasing the width results in faster convergence. However, increasing the width alone also increases the number of parameters. By employing our compression methodology, we can achieve the best of both worlds.