
Fundamental Limits of Two-layer Autoencoders, and Achieving Them with Gradient Methods

Aleksandr Shevchenko^{*1} Kevin Kögler^{*1} Hamed Hassani² Marco Mondelli¹

Abstract

Autoencoders are a popular model in many branches of machine learning and lossy data compression. However, their fundamental limits, the performance of gradient methods and the features learnt during optimization remain poorly understood, even in the two-layer setting. In fact, earlier work has considered either linear autoencoders or specific training regimes (leading to vanishing or diverging compression rates). Our paper addresses this gap by focusing on non-linear two-layer autoencoders trained in the challenging proportional regime in which the input dimension scales linearly with the size of the representation. Our results characterize the minimizers of the population risk, and show that such minimizers are achieved by gradient methods; their structure is also unveiled, thus leading to a concise description of the features obtained via training. For the special case of a sign activation function, our analysis establishes the fundamental limits for the lossy compression of Gaussian sources via (shallow) autoencoders. Finally, while the results are proved for Gaussian data, numerical simulations on standard datasets display the universality of the theoretical predictions.

1. Introduction

Autoencoders represent a key building block in many branches of machine learning (Kingma & Welling, 2014; Rezende et al., 2014), including generative modeling (Bengio et al., 2013) and representation learning (Tschannen et al., 2018). Prompted by the fact that autoencoders learn

succinct representations, neural autoencoding techniques have achieved remarkable success in lossy data compression, even outperforming classical methods, such as jpeg (Ballé et al., 2017; Theis et al., 2017; Agustsson et al., 2017). However, despite the large body of empirical work on neural autoencoders and compressors, basic theoretical questions remain poorly understood even in the shallow case:

What are the fundamental performance limits of autoencoders? Can we achieve such limits with gradient methods? What features does the optimization procedure learn?

Prior work has focused either on linear autoencoders (Baldi & Hornik, 1989; Kunin et al., 2019; Gidel et al., 2019), on the severely under-parameterized setting in which the input dimension is much larger than the number of neurons (Refinetti & Goldt, 2022), or on specific training regimes (lazy training (Nguyen et al., 2021) and mean-field regime with a polynomial number of neurons (Nguyen, 2021)), see Section 2. In contrast, in this paper we consider *non-linear* autoencoders trained in the *challenging proportional regime*, in which the number of inputs to compress scales linearly with the size of the representation. More specifically, we consider the prototypical model of a two-layer autoencoder

$$\hat{x}(x) := \hat{x}(x, A, B) = A\sigma(Bx). \quad (1)$$

Here, $x \in \mathbb{R}^d$ is the input to compress, $\hat{x} \in \mathbb{R}^n$ the reconstruction, $B \in \mathbb{R}^{n \times d}$ the encoding matrix, and $A \in \mathbb{R}^{d \times n}$ the decoding matrix; the activation $\sigma : \mathbb{R} \rightarrow \mathbb{R}$ is applied *element-wise*. We aim at minimizing the population risk

$$\mathcal{R}(A, B) := d^{-1} \mathbb{E}_x \|x - \hat{x}(x)\|_2^2, \quad (2)$$

where the expectation is taken over the distribution of the input x . Our focus is on Gaussian input data, i.e., $x \sim \mathcal{N}(\mathbf{0}, \Sigma)$. When σ is the sign function, the encoder $\sigma(Bx)$ can be interpreted as a *compressor*, namely, it compresses the d -dimensional input signal into n bits. The problem (2) of compressing a Gaussian source with quadratic distortion has been studied in exquisite detail in the information theory literature (Cover & Thomas, 2006), and the optimal performance for general encoder/decoder pairs is known via the *rate-distortion* formalism which characterizes the lowest achievable distortion in terms of the rate $r = n/d$.

^{*}Equal contribution ¹ISTA, Klosterneuburg, Austria

²Department of Electrical and Systems Engineering, University of Pennsylvania, USA. Correspondence to: Aleksandr Shevchenko <alex.shevchenko@ist.ac.at>, Kevin Kögler <kevin.koegler@ist.ac.at>.

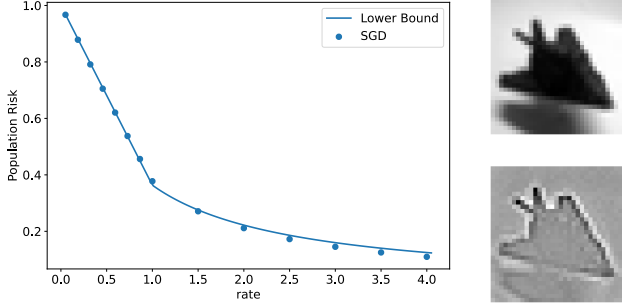


Figure 1. Compression ($\sigma \equiv \text{sign}$) of the CIFAR-10 “airplane” class with a two-layer autoencoder. The data is *whitened* so that $\Sigma = I$: on top, an example of a grayscale image; on the bottom, the corresponding whitening. The blue dots are the population risk obtained via SGD, and they agree well with the solid line corresponding to the lower bounds of Theorem 4.2 and Proposition 4.3.

Here, we focus on encoders and decoders that form the two-layer autoencoder (1): we study the fundamental limits of this learning problem, as well as the performance achieved by commonly used gradient descent methods.

Main contributions. Taken all together, our results show that, for two-layer autoencoders, gradient descent methods achieve a global minimizer of the population risk: this is rigorously proved in the isotropic case ($\Sigma = I$) and corroborated by numerical simulations for a general covariance Σ . Furthermore, we unveil the structure of said minimizer: for $\Sigma = I$, the optimal decoder has unit singular values; for general covariance, the spectrum of the decoder exhibits the same block structure as Σ , and it can be explicitly obtained from Σ via a water-filling criterion; in all cases, weight-tying is optimal, i.e., A is proportional to $B^\top$. Specifically, our technical results can be summarized as follows.

- Section 4.1 characterizes the minimizers of the risk (2) for isotropic data: Theorem 4.2 provides a tight lower bound, which is achieved by the set (7) of weight-tied *orthogonal* matrices, when the compression rate $r = n/d \leq 1$; for $r > 1$, Propositions 4.3 and 4.4 give a lower bound, which is approached (as $d \rightarrow \infty$) by the set (12) of weight-tied *rotationally invariant* matrices.
- Section 4.2 shows that the above minimizers are reached by gradient descent methods for $r \leq 1$: Theorem 4.5 shows linear convergence of *gradient flow* for general initializations, under a weight-tying condition; Theorem 4.6 considers a Gaussian initialization and proves global convergence of the *projected gradient descent* algorithm, in which the encoder matrix B is optimized via a gradient method and the decoder matrix A is obtained directly via linear regression.
- Section 5 focuses on data with general covariance $\Sigma \neq I$. We observe that experimentally weight-tying is optimal

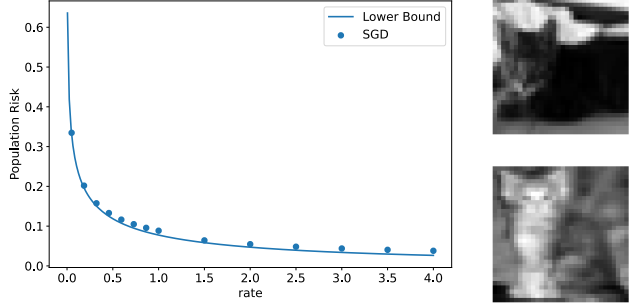


Figure 2. Compression ($\sigma \equiv \text{sign}$) of the CIFAR-10 “cat” class with a two-layer autoencoder. The data is *not whitened* ($\Sigma \neq I$). The blue dots are the SGD population risk, and they are close to the lower bound of Theorem 5.2.

and then derive the corresponding lower bound (see Theorem 5.2), which is also asymptotically achieved (as $d \rightarrow \infty$) by rotationally invariant matrices with a carefully designed spectrum (depending on Σ), see Proposition 5.3.

When $\sigma \equiv \text{sign}$, our analysis characterizes the fundamental limits of the lossy compression of a Gaussian source via two-layer autoencoders. Remarkably, if we restrict to a certain class of *linear encoders* for compression, two-layer autoencoders achieve optimal performance (Tulino et al., 2013), which can be generally obtained via a message passing decoding algorithm (Rangan et al., 2019). However, for *general encoder/decoder pairs*, shallow autoencoders fail to meet the information-theoretic bound given by the rate-distortion curve, see Section 6.

Going beyond the Gaussian assumption on the data, we provide numerical validation to our theoretical predictions on standard datasets, both in the isotropic case (Figure 1) and for general covariance (Figure 2). Additional numerical results – together with the details of the experimental setting – are in Appendix I.

Proof techniques. The lower bound on the population risk of Theorem 4.2 comes from a sequence of relaxations of the objective function, which eventually allows to apply a trace inequality. For $r \geq 1$, Proposition 4.3 crucially exploits an inequality for the Hadamard product of PSD matrices (Khare, 2021), and the asymptotic achievability of Proposition 4.4 takes advantage of concentration-of-measure tools for orthogonal matrices. The key quantity in the analysis of gradient methods is the encoder Gram matrix at iteration t , i.e., $B(t)B(t)^\top$. In particular, for gradient flow (Theorem 4.5), due to the weight-tying condition, tracking $\log \det B(t)B(t)^\top$ leads to a quantitative convergence result. However, when the weights are not tied, this quantity does not appear to increase along the optimization trajectory. Thus, for projected gradient descent (Theorem 4.6), the idea is to decompose $B(t)B(t)^\top$ into (i) its value at the optimum (given by the identity), (ii) the contribution due

to the spectrum evolution (keeping the eigenbasis fixed), and (iii) the change in the eigenbasis. Via a sequence of careful approximations, we are able to show that the term (iii) vanishes. Hence, we can study explicitly the evolution of the spectrum and obtain the desired convergence.

2. Related Work

Theory of autoencoders. A popular line of work has focused on two-layer *linear* autoencoders: Oftadeh et al. (2020) analyze the loss landscape; Kunin et al. (2019) show that the minimizers of the regularized loss recover the principal components of the data and, notably, the corresponding autoencoder is weight-tied; Bao et al. (2020) prove that stochastic gradient descent – after a slight perturbation – escapes the saddles and eventually converges; Gidel et al. (2019) characterize the time-steps at which the network learns different sets of features. Rangamani et al. (2018); Nguyen et al. (2019) prove local convergence for weight-tied two-layer ReLU autoencoders. Nguyen et al. (2021) focus on the lazy training regime (Chizat et al., 2019; Jacot et al., 2018) and bound the over-parameterization needed for global convergence. Radhakrishnan et al. (2020) show that over-parameterized autoencoders learn solutions that are contractive around the training examples. The latent spaces of autoencoders are studied in (Jain et al., 2021), where it is shown that such latent spaces can be aligned by stretching along the left singular vectors of the data. More closely related to our work, Nguyen (2021) and (Refinetti & Goldt, 2022) track the gradient dynamics of *non-linear* two-layer autoencoders via the mean-field PDE and a system of ODEs, respectively. However, these analyses are restricted to diverging and vanishing rates: Nguyen (2021) considers weight-tied autoencoders with polynomially many neurons in the input dimension (so that $r \rightarrow \infty$); Refinetti & Goldt (2022) consider the other extreme regime in which the input dimension diverges (so that $r \rightarrow 0$).

Neural compression. In recent years, compressors based on neural networks have outperformed traditional schemes on real-world data in terms of minimizing distortion and producing visually pleasing reconstructions at reasonable complexity (Ballé et al., 2017; Theis et al., 2017; Agustsson et al., 2017; Ballé et al., 2021). These methods typically use an autoencoder architecture with quantization of the latent variables, which is trained over samples drawn from the source. More recently, other architectures such as attention or diffusion-based models have been incorporated into neural compressors (Cheng et al., 2020; Liu et al., 2019; Yang & Mandt, 2022; Theis et al., 2022), and improvements have been observed. We refer to Yang et al. (2022) for a detailed review on this topic. Given the remarkable success of neural compressors, it is imperative to understand the fundamental limits of compression using neural archi-

tectures. In this regard, Wagner & Ballé (2021) consider a highly-structured and low-dimensional random process, dubbed the *sawbridge*, and show numerically that the rate-distortion function is achieved by a compressor based on deep neural networks trained via stochastic gradient descent. In contrast, our work considers Gaussian sources, which are high-dimensional in nature, and provides the fundamental limits of compression for two-layer autoencoders. Our results also imply that two-layer autoencoders cannot achieve the rate-distortion limit on Gaussian data, see Section 6.

Additional related works on rate-distortion formalism and non-linear inverse problems are discussed in Appendix A.

3. Preliminaries

Notations. We use plain symbols for real numbers (e.g., a, b), bold symbols for vectors (e.g., $\mathbf{a}, \mathbf{b}$), and capitalized bold symbols for matrices (e.g., $\mathbf{A}, \mathbf{B}$). We let $[n] = \{1, \dots, n\}$, $\mathbf{I}$ be the identity matrix and $\mathbf{1}$ the column vector containing ones. Given a matrix $\mathbf{A}$, we denote its operator norm by $\|\mathbf{A}\|_{op}$ and its Frobenius norm by $\|\mathbf{A}\|_F$. Given two matrices $\mathbf{A}$ and $\mathbf{B}$ of the same shape, we denote their element-wise (Hadamard/Schur) product by $\mathbf{A} \circ \mathbf{B}$ and the k -th element-wise power by $\mathbf{A}^{\circ k}$. We write $L^2(\mathbb{R}, \mu)$ for the space of L^2 integrable functions on $\mathbb{R}$ w.r.t. the standard Gaussian measure μ and $h_k(x)$ for the k -th normalized Hermite polynomial (see e.g. O’Donnell (2014)).

Setup. We consider the two-layer autoencoder (1) and aim at minimizing the population risk (2) for a given rate $r = n/d$. In particular, we provide tight lower bounds on the minimum of the population risk computed on Gaussian input data with covariance Σ , i.e.,

$$\hat{\mathcal{R}}(r) := \min_{\mathbf{A}, \mathbf{B}} \mathcal{R}(\mathbf{A}, \mathbf{B}). \quad (3)$$

In the isotropic case ($\Sigma = \mathbf{I}$), our results hold for any odd activation $\sigma \in L^2(\mathbb{R}, \mu)$ after restricting the rows of the encoding matrix $\mathbf{B}$ to have unit norm. We remark that, when $\sigma(x) = \text{sign}(x)$, the restriction is unnecessary since the activation is homogeneous.¹ We also note that restricting the norms of the rows of $\mathbf{B}$ prevents the model from entering the “linear” regime. In fact, when $\|\mathbf{B}\|_F \approx 0$, by linearizing the activation around zero, (1) reduces to the linear model $\hat{\mathbf{x}}(\mathbf{x}) \approx \mathbf{A}\mathbf{B}\mathbf{x}$, which exhibits a PCA-like behaviour. For general covariance Σ , we consider odd homogeneous activations, which includes the sign function and monomials of arbitrary odd degree.

Any function $\sigma \in L^2(\mathbb{R}, \mu)$ can be expanded in terms of Hermite polynomials. This allows to perform Fourier analysis in the Gaussian space $L^2(\mathbb{R}, \mu)$, and it provides a natural tool because of the Gaussian assumption on the data. In

¹We say that a function σ is homogeneous if there exists an integer k s.t. $\sigma(\alpha x) = \alpha^k \sigma(x)$ for all $\alpha \neq 0$.

particular, for odd σ , only odd Hermite polynomials occur, i.e.,

$$\sigma(x) = \sum_{\ell=0}^{\infty} c_{2\ell+1} h_{2\ell+1}(x), \quad (4)$$

where $\{c_\ell\}_{\ell \in \mathbb{N}}$ denote the Hermite coefficients of σ . We also consider the following auxiliary quantity

$$\tilde{\mathcal{R}}(r) := \min_{\mathbf{A}, \|\mathbf{B}\mathbf{D}\|_{i,:} = 1} \mathcal{R}(\mathbf{A}, \mathbf{B}), \quad (5)$$

that defines a minimum of the population risk for the autoencoder (1) with a certain norm constraint on the encoder weights $\mathbf{B}$. Here, $\mathbf{D}$ contains the square roots of the eigenvalues of Σ (i.e., $\Sigma = \mathbf{U}\mathbf{D}^2\mathbf{U}^\top$ for an orthogonal matrix $\mathbf{U}$), and $(\mathbf{B}\mathbf{D})_{i,:}$ stands for the i -th row of the matrix $\mathbf{B}\mathbf{D}$. A few remarks about the restricted population risk (5) are in order. First of all, if σ is homogeneous, the minimum of the restricted population risk (5) and of the unconstrained one (3) coincide (see Lemma 4.1 and Lemma 5.1). Thus, in this case, the analysis of $\tilde{\mathcal{R}}(r)$ directly provides results on the quantity of interest, i.e., $\hat{\mathcal{R}}(r)$. The technical advantage of analysing (5) over (3) comes from fact that the expectation with respect to the Gaussian inputs, which arises in the constrained objective, can be explicitly computed via the reproducing property of Hermite polynomials (see, e.g., O'Donnell (2014)). To exploit this reproducing property, it is crucial that the inner products $\langle \mathbf{B}_{i,:}, \mathbf{x} \rangle$ have the same scale, which is ensured by picking $\|(\mathbf{B}\mathbf{D})_{i,:}\|_2 = 1$. The sole dependence of the constraint on the spectrum $\mathbf{D}$ (and, not on a particular choice of $\mathbf{U}$) stems from the rotational invariance of the isotropic Gaussian distribution.

4. Main Results

In this section, we consider isotropic Gaussian data, i.e., $\Sigma = \mathbf{D} = \mathbf{I}$. First, we derive a closed form expression for the population risk in Lemma 4.1. Then, in Theorem 4.2 we give a lower bound on the population risk for $r \leq 1$ and provide a complete characterization of the autoencoder parameters $(\mathbf{A}, \mathbf{B})$ achieving it. Surprisingly, the minimizer exhibits a *weight-tying* structure and the corresponding matrices are *rotationally invariant*. Later, in Proposition 4.3 we derive an analogous lower bound for $r > 1$. While it is hard to characterize the minimizer structure explicitly for a finite input dimension d (and $r > 1$), we provide a sequence $\{(\mathbf{A}_d, \mathbf{B}_d)\}_{d \in \mathbb{N}}$ that meets the lower bound in the high-dimensional limit ($d \rightarrow \infty$), see Proposition 4.4. Notably, the elements of this sequence share the key features (weight-tying, rotational invariance) of the minimizers for $r \leq 1$. In Section 4.2 we describe gradient methods that provably achieve the optimal value of the population risk. Specifically, we consider gradient flow under a weight-tying constraint and projected (on the sphere) gradient descent with Gaussian initialization. The corresponding results are stated in Theorem 4.5 and Theorem 4.6.

We start by expanding σ in a Hermite series to obtain a closed-form expression for the population risk.

Lemma 4.1. *Consider any odd $\sigma \in L^2(\mathbb{R}, \mu)$ and its Hermite expansion given by (4). Then, $\tilde{\mathcal{R}}(r)$ is equivalent to*

$$\min_{\mathbf{A}, \|\mathbf{B}_{i,:}\|_2=1} \frac{1}{d} \left(\text{Tr} \left[\mathbf{A}^\top \mathbf{A} f(\mathbf{B}\mathbf{B}^\top) \right] - 2c_1 \cdot \text{Tr}[\mathbf{B}\mathbf{A}] \right) + 1, \quad (6)$$

where $f(x) := \sum_{\ell=0}^{\infty} (c_{2\ell+1})^2 x^{2\ell+1}$ is applied element-wise. In particular, if $\sigma(x) = \text{sign}(x)$, then $f(x) = c_1^2 \cdot \arcsin(x)$ and $c_1 = \sqrt{2/\pi}$. Moreover, for any homogeneous σ , we have that $\hat{\mathcal{R}}(r) = \tilde{\mathcal{R}}(r)$.

The proof of the lemma above is contained in Appendix B. Note that, if $c_1 = 0$, it is easy to see that the minimum of $\tilde{\mathcal{R}}(r)$ equals 1 and it is attained when $\mathbf{A}^\top \mathbf{A}$ is the zero-matrix. Furthermore, if $\sum_{\ell=1}^{\infty} (c_{2\ell+1})^2 = 0$, then $\sigma(x) = c_1^2 x$ and we fall back into the simpler case of a linear autoencoder (Baldi & Hornik, 1989; Kunin et al., 2019; Gidel et al., 2019). Thus, for the rest of the section, we will assume that $c_1 \neq 0$ and $\sum_{\ell=1}^{\infty} (c_{2\ell+1})^2 \neq 0$.

4.1. Fundamental Limits: Lower Bound on Risk

We begin by providing a tight lower bound for $r \leq 1$, which is *uniquely* achieved on the set of *weight-tied* orthogonal matrices $\mathcal{H}_{n,d}$ defined as

$$\mathcal{H}_{n,d} := \left\{ \tilde{\mathbf{A}}, \tilde{\mathbf{B}}^\top \in \mathbb{R}^{d \times n} : \tilde{\mathbf{A}} = \frac{c_1}{f(1)} \cdot \tilde{\mathbf{B}}^\top, \tilde{\mathbf{B}}\tilde{\mathbf{B}}^\top = \mathbf{I} \right\}. \quad (7)$$

Theorem 4.2. *Consider any odd $\sigma \in L^2(\mathbb{R}, \mu)$ and fix $r \leq 1$. Then, the following holds*

$$\tilde{\mathcal{R}}(r) \geq \text{LB}_{r \leq 1}(\mathbf{I}) := 1 - \frac{c_1^2}{f(1)} \cdot r,$$

and equality is achieved iff $(\mathbf{A}, \mathbf{B}) \in \mathcal{H}_{n,d}$.

We note that the minimizers $\mathcal{H}_{n,d}$ of $\tilde{\mathcal{R}}(r)$ do not directly correspond to the minimizers of the unconstrained population risk $\hat{\mathcal{R}}(r)$, since in general $\tilde{\mathcal{R}}(r) \neq \hat{\mathcal{R}}(r)$. However, if σ is homogeneous, the “inverse” mapping can be readily obtained. For instance, when $\sigma(x) = \text{sign}(x)$, rescaling the norms of the rows of $\mathbf{B}$ does not affect the compression, i.e., $\text{sign}(\mathbf{B}\mathbf{x}) = \text{sign}(\mathbf{S}\mathbf{B}\mathbf{x})$ for any diagonal $\mathbf{S}$ with positive entries. Hence, to obtain a minimizer, it suffices that the rows of $\mathbf{B}$ form any set of orthogonal (not necessarily normalized) vectors. In contrast, note that $\mathbf{A}$ is still defined with respect to the row-normalized version of $\mathbf{B}$. Similar arguments hold for homogeneous activations.

We also note that the weight-tying structure (7) observed in the minimizers of the population risk is related to the early representation learning literature (Vincent et al., 2008; Hinton & Salakhutdinov, 2006).

We now provide a proof sketch for Theorem 4.2 and defer the full argument to Appendix C.1.

Proof sketch of Theorem 4.2. Using the series expansion of $f(\cdot)$, we can write

$$\begin{aligned} & \text{Tr} \left[\mathbf{A}^\top \mathbf{A} f(\mathbf{B}\mathbf{B}^\top) \right] - 2c_1 \cdot \text{Tr} [\mathbf{B}\mathbf{A}] \\ &= \sum_{\ell=0}^{\infty} c_{2\ell+1}^2 \left(\text{Tr} \left[\mathbf{A}^\top \mathbf{A} (\mathbf{B}\mathbf{B}^\top)^{\circ 2\ell+1} \right] - 2 \frac{c_1}{f(1)} \text{Tr} [\mathbf{B}\mathbf{A}] \right). \end{aligned}$$

Thus, the minimization problem in Lemma 4.1 can be reduced to analysing each Hadamard power individually:

$$\min_{\mathbf{A}, \|\mathbf{B}_{i,:}\|_2=1} \text{Tr} \left[\mathbf{A}^\top \mathbf{A} (\mathbf{B}\mathbf{B}^\top)^{\circ \ell} \right] - \frac{2c_1}{f(1)} \cdot \text{Tr} [\mathbf{B}\mathbf{A}]. \quad (8)$$

The crux of the argument is to provide a suitable sequence of relaxations of (8). The first relaxation gives

$$\text{Tr} \left[(\mathbf{A}^\top \mathbf{A} \circ \mathbf{Q})(\mathbf{B}\mathbf{B}^\top \circ \mathbf{Q}) \right] - \frac{2c_1}{f(1)} \cdot \text{Tr} [\mathbf{B}\mathbf{A}], \quad (9)$$

where $\mathbf{Q}$ is any PSD matrix with unit diagonal. Using the properties of the SVD of $\mathbf{Q}$, (9) can be further relaxed to

$$\sum_{i,j=1}^n \text{Tr} \left[\mathbf{A}_j \mathbf{A}_j^\top \mathbf{B}_j \mathbf{B}_j^\top \right] - \frac{2c_1}{f(1)} \cdot \sum_{i=1}^n \text{Tr} [\mathbf{B}_i \mathbf{A}_i], \quad (10)$$

where now $\mathbf{A}_i, \mathbf{B}_i^\top \in \mathbb{R}^{d \times n}$ are arbitrary matrices. The key observation is that

$$\sum_{i=1}^n \left\| \frac{c_1}{f(1)} \cdot \sqrt{\mathbf{X}}^{-1} \mathbf{A}_i^\top - \sqrt{\mathbf{X}} \mathbf{B}_i \right\|_F^2 = (10) + \frac{c_1^2}{(f(1))^2} \cdot n,$$

with $\mathbf{X} = \sum_{i=1}^n \mathbf{A}_i^\top \mathbf{A}_i$. As each relaxation lower bounds (8) and the Frobenius norm is positive, this argument leads to the lower bound on $\tilde{\mathcal{R}}(r)$. The fact that the lower bound is met for any $(\mathbf{A}, \mathbf{B}) \in \mathcal{H}_{n,d}$ can be verified via a direct calculation. The uniqueness follows by taking the intersection of the minimizers of (8) for different values of ℓ . $\square$

Next, we move to the case $r > 1$.

Proposition 4.3. *Consider any odd $\sigma \in L^2(\mathbb{R}, \mu)$ and fix $r > 1$, then the following holds:*

$$\tilde{\mathcal{R}}(r) \geq \text{LB}_{r>1}(\mathbf{I}) := 1 - \frac{r}{r + \left(\frac{f(1)}{c_1^2} - 1 \right)}.$$

The key difference with the proof of the lower bound in Theorem 4.2 is that the term $\text{Tr} [\mathbf{A}^\top \mathbf{A} \mathbf{B} \mathbf{B}^\top]$ requires a tighter estimate. This is due to the fact that the matrix $\mathbf{B}\mathbf{B}^\top$ is no longer full-rank when $r > 1$. We obtain the desired tighter bound by exploiting the following result by (Khare, 2021):

$$\mathbf{A}^\top \mathbf{A} \circ \mathbf{B}\mathbf{B}^\top \succeq \frac{1}{d} \cdot \text{Diag}(\mathbf{B}\mathbf{A}) \text{Diag}(\mathbf{B}\mathbf{A})^\top, \quad (11)$$

where $\text{Diag}(\mathbf{B}\mathbf{A})$ stands for the vector containing the diagonal entries of $\mathbf{B}\mathbf{A}$. The full argument is contained in Appendix C.2.1.

As for $r \leq 1$, the bound is met (here, in the limit $d \rightarrow \infty$) by considering weight-tied matrices:

$$\hat{\mathbf{B}}^\top = \sqrt{r} \cdot [\mathbf{I}_d, \mathbf{0}_{d,n-d}] \mathbf{U}^\top, \quad \mathbf{b}_i = \frac{\hat{\mathbf{b}}_i}{\|\hat{\mathbf{b}}_i\|_2}, \quad \mathbf{A} = \beta \mathbf{B}^\top, \quad (12)$$

where $\beta = \frac{c_1}{c_1^2 r + f(1) - c_1^2}$ and $\mathbf{U}$ is uniformly sampled from the group of rotation matrices. The idea behind the choice (12) is that, as $d \rightarrow \infty$, $(\mathbf{B}\mathbf{B}^\top)^{\circ 2\ell}$ for $\ell \geq 2$ is close to the identity matrix, and (11) is attained exactly. The formal statement is provided below and proved in Appendix C.2.2.

Proposition 4.4. *Consider any odd $\sigma \in L^2(\mathbb{R}, \mu)$ and fix $r > 1$. Let $\mathbf{A}, \mathbf{B}$ be defined as in (12). Then, for any $\epsilon > 0$ the following holds*

$$|\mathcal{R}(\mathbf{A}, \mathbf{B}) - \text{LB}_{r>1}(\mathbf{I})| \leq C d^{-\frac{1}{2} + \epsilon},$$

with probability $1 - c/d^2$. Here, the constants c, C depend only on r and ϵ .

We note that all the arguments of this section directly apply to Gaussian data with a covariance matrix of the form $\Sigma = \begin{bmatrix} \sigma^2 \mathbf{I}_{d-k} & \mathbf{0}_{(d-k) \times k} \\ \mathbf{0}_{k \times (d-k)} & \mathbf{0}_{k \times k} \end{bmatrix}$. For the details, see Appendix F.

4.2. Gradient Methods Achieve the Lower Bound

In this section, we discuss the achievability of the lower bound obtained in the previous section via gradient methods. We study two procedures which find the minimizer of the population risk $\mathcal{R}(\mathbf{A}, \mathbf{B})$ under the constraint $\|\mathbf{B}_{i,:}\|_2 = 1$. Namely, we analyse (i) *weight-tied gradient flow* on the sphere and (ii) its discrete version (with finite step size) *without* weight-tying, i.e., *projected gradient descent*.

The optimization objective in Lemma 4.1 is equivalent (up to a scaling independent of $(\mathbf{A}, \mathbf{B})$) to

$$\min_{\mathbf{A}, \|\mathbf{B}_{i,:}\|_2=1} \text{Tr} \left[\mathbf{A}^\top \mathbf{A} \cdot f(\mathbf{B}\mathbf{B}^\top) \right] - 2 \cdot \text{Tr} [\mathbf{B}\mathbf{A}], \quad (13)$$

where we have rescaled the function f by $1/c_1^2$. This follows from the fact that the multiplicative factor c_1 can be pushed inside $\mathbf{A}$. Note that such scaling does not affect the properties of gradient-based algorithms (modulo a constant change in their speed). Hence, without loss of generality, we will state and prove all our results for the problem (13).

Weight-tied gradient flow. We start by considering the weight-tied setting

$$\mathbf{A} = \beta \mathbf{B}^\top, \quad \beta \in \mathbb{R}. \quad (14)$$

This is motivated by the fact that the lower bounds on the population risk are approached by weight-tied matrices (see

Theorem 4.2 and Proposition 4.4). We keep the presentation brief and informal, and the formal setup is deferred to Appendix D. Note that for all $\mathbf{A}, \mathbf{B}$ under the weight-tying constraint (14), the optimal β^* in (13) can be found exactly. Thus, to optimize (13), we perform a (Riemannian) gradient flow on $\mathbf{B}$ with rows constrained to the unit sphere, where at each time t we pick the optimal $\beta^*(t)$:

$$\frac{\partial \mathbf{b}_i(t)}{\partial t} = -\mathbf{J}_i(t) \nabla_{\mathbf{b}_i} \Psi(\beta^*(t), \mathbf{B}(t)), \quad (15)$$

where $\mathbf{b}_i(t)$ stands for the i -th row of $\mathbf{B}(t)$ and $\mathbf{J}_i(t)$ projects the gradient $\nabla_{\mathbf{b}_i} \Psi(\beta(t), \mathbf{B}(t))$ of (13) under the weight-tying constraint on the unit sphere (see (64)-(65) in Appendix D for the exact expressions). This ensures that $\|\mathbf{b}_i(t)\|_2 = 1$ along the gradient flow trajectory.

Theorem 4.5 (Informal). *For any $r \leq 1$, the gradient flow (15) initialized with full rank unit-row norm $\mathbf{B}$ converges to a global minimizer of (13), given by $\mathbf{B}\mathbf{B}^\top = \mathbf{I}$.*

Proof sketch of Theorem 4.5. It can be readily shown that the $\mathbf{B}$'s for which $\mathbf{B}\mathbf{B}^\top = \mathbf{I}$ are the unique minimizers of (13) under the weight-tying constraint and they satisfy the stationary point condition of the gradient flow (15). However, if $\mathbf{B}$ becomes not full-rank, such subspaces are never escaped by the gradient flow (15) (see Lemma D.2). Hence, the procedure would fail to converge to the global minimizer that has full-rank. We show that, under the full-rank initialization, this does not happen by lower bounding the time derivative of $\log \det(\mathbf{B}(t)\mathbf{B}(t)^\top)$ (see Lemma D.3), which vanishes uniquely at $\mathbf{B}\mathbf{B}^\top = \mathbf{I}$. This ensures that the solution of (15) will not saturate to a low-rank subspace. $\square$

In Appendix D, we also provide a quantitative bound on the convergence time (see Lemma D.4). We remark that Theorem 4.5 holds for any d and for all full-rank initializations.

Projected gradient descent. We now move to the setting where the encoder and decoder weights are not weight-tied. In this case, we consider the commonly used Gaussian initialization and prove a result for sufficiently large d . The Gaussian initialization allows us to relax the requirement on f : we only need $c_2 = 0$, as opposed to the previous assumption that $c_{2\ell} = 0$ for any $\ell \in \mathbb{N}$ (see the statement of Lemma 4.1). Specifically, we consider the following algorithm to minimize (13):

$$\begin{aligned} \mathbf{A}(t) &= \mathbf{B}(t)^\top (f(\mathbf{B}(t)\mathbf{B}(t)^\top))^{-1} \\ \mathbf{B}'(t) &:= \mathbf{B}(t) - \eta \nabla_{\mathbf{B}(t)} \Psi(\mathbf{B}(t), \mathbf{B}(t+1)) := \text{proj}(\mathbf{B}'(t)), \end{aligned} \quad (16)$$

where $\mathbf{A}(t)$ is the optimal matrix for a fixed $\mathbf{B}(t)$ and $\nabla_{\mathbf{B}(t)}$ (see (86) in Appendix E) is the projected gradient of the objective (13) with respect to $\mathbf{B}(t)$. Furthermore, $\text{proj}(\mathbf{B}'(t))$ rescales all the rows to have unit norm. It will become apparent from the proof of Theorem 4.6 that the inversion in the definition of $\mathbf{A}(t)$ is indeed well defined. We remark

that (16) can be viewed as the discrete counterpart of the Riemannian gradient flow for the weight-tied case (with the optimal $\mathbf{A}(t)$ in place of the weight-tying), where the application of $\text{proj}(\cdot)$ keeps the rows of $\mathbf{B}(t)$ of unit norm. In the related literature, this procedure is often referred to as Riemannian gradient descent (see, e.g., Absil et al. (2009)). Alternatively, (16) may be viewed as coordinate descent (Wright, 2015) on the objective (13), where the step in $\mathbf{A}$ is performed exactly.

Our main result is that the projected gradient descent (16) converges to the global optimum of (13) for large enough d (with high probability). We give a sketch of the argument and defer the complete proof to Appendix E.

Theorem 4.6. *Consider the projected gradient descent (16) applied to the objective (13) for any f of the form $f(x) = x + \sum_{\ell=3} c_\ell^2 x^\ell$, where $\sum_{\ell=3} c_\ell^2 < \infty$. Initialize the algorithm with $\mathbf{B}(0)$ equal to a row-normalized Gaussian, i.e., $\mathbf{B}'_{i,j}(0) \sim \mathcal{N}(0, 1/d)$, $\mathbf{B}(0) = \text{proj}(\mathbf{B}'(0))$. Let the step size η be $\Theta(1/\sqrt{d})$. Then, for any $r < 1$ and sufficiently large d , with probability at least $1 - Ce^{-cd}$, we have that $\mathbf{B}(t)\mathbf{B}(t)^\top$ converges to $\mathbf{I}$, which is the unique global optimum of (13). Moreover, defining $t = T/\eta$, we have the following bound on the rate of convergence*

$$\|\mathbf{B}(t)\mathbf{B}(t)^\top - \mathbf{I}\|_{op} \leq C(1 - c)^T,$$

where $C > 0$ and $c \in (0, 1]$ are universal constants depending only on r and f .

Proof sketch of Theorem 4.6. Let $\mathbf{B}(0)\mathbf{B}(0)^\top = \mathbf{U}\mathbf{\Lambda}(0)\mathbf{U}^\top$ be the singular value decomposition (SVD) of the encoder Gram matrix. Then, the idea is to decompose $\mathbf{B}(t)\mathbf{B}(t)^\top$ at each step of the projected gradient descent dynamics as

$$\mathbf{B}(t)\mathbf{B}(t)^\top = \mathbf{I} + \mathbf{Z}(t) + \mathbf{X}(t), \quad (17)$$

where $\mathbf{Z}(t) = \mathbf{U}(\mathbf{\Lambda}(t) - \mathbf{I})\mathbf{U}^\top$. Here, $\mathbf{I}$ is the global optimum towards which we want to converge; $\mathbf{Z}(t)$ captures the evolution of the eigenvalues while keeping the eigenbasis fixed, as $\mathbf{U}$ comes from the SVD at initialization; and $\mathbf{X}(t)$ is the remaining error term capturing the change in the eigenbasis. The update on $\mathbf{\Lambda}(t)$ is given by $\mathbf{\Lambda}(t+1) = g(\mathbf{\Lambda}(t))$, where $g: \mathbb{R}^n \rightarrow \mathbb{R}^n$ admits an explicit expression. Hence, in light of this explicit expression, if we had $\mathbf{X}(t) \equiv 0$, then the desired convergence would follow from the analysis of the recursion for $\mathbf{\Lambda}(t)$ (see Lemma G.3).

The main technical difficulty lies in carefully controlling the error term $\mathbf{X}(t)$. In particular, we will show that $\mathbf{X}(t)$ decays for large enough d , which means that dynamics (17) is well approximated by $\mathbf{I} + \mathbf{Z}(t)$. The proof can be broken down in four steps. In the *first step*, we compute the leading order term of $\nabla_{\mathbf{B}(t)}$ (see Lemma E.2 and E.3). This simplifies the formula for $\nabla_{\mathbf{B}(t)}$, which can then be

expressed as an explicit nonlinear function of $\mathbf{Z}(t)$ and $\mathbf{X}(t)$. In the *second step*, we perform a Taylor expansion of $\nabla_{\mathbf{B}(t)}$, seen as a matrix-valued function in $\mathbf{Z}(t)$ and $\mathbf{X}(t)$ (see Lemma E.4). The intuition for this expansion comes from the fact that $\mathbf{X}(t)$ is a small quantity, and also $\|\mathbf{Z}(t)\|_{op} \rightarrow 0$ as $t \rightarrow \infty$. In the *third step*, we show that the norm of $\nabla_{\mathbf{B}(t)}$ vanishes sufficiently fast (see Lemma E.5), which implies that the projection step $\mathbf{B}(t+1) := \text{proj}(\mathbf{B}'(t))$ has a negligible effect (see Lemma E.6). After doing these estimates, we finally obtain an explicit recursion for $\mathbf{X}(t)$. In the *fourth step*, we analyse this recursion (see Lemma E.7): first, we show that the error does not amplify too strongly (as in Gronwall's inequality); then, armed with this worst-case estimate, we can prove an exponential decay for $\mathbf{X}(t)$, which suffices to conclude the argument. $\square$

Further discussions on the results in this section can be found in Appendix F.

5. Extension to General Covariance

In this section, we consider a Gaussian source with general covariance structure, i.e., $\Sigma = \mathbf{U}\mathbf{D}^2\mathbf{U}^\top$. Without loss of generality, the matrix $\mathbf{D}$ can be written as

$$\mathbf{D} = \text{Diag}(\underbrace{[D_1, \dots, D_1]}_{\times k_1} | \dots | \underbrace{[D_K, \dots, D_K]}_{\times k_K}), \quad (18)$$

where $\sum_{i=1}^K k_i = d$, $k_i \geq 1$ and $D_i > D_{i+1} \geq 0$. We start by deriving a closed-form expression for the population risk, similar to Lemma 4.1. Its proof is given in Appendix B.

Lemma 5.1. *Let $\sigma \in L^2(\mathbb{R}, \mu)$ be an odd homogeneous activation, then $\tilde{\mathcal{R}}(r)$ is equal to the minimum of*

$$\frac{1}{d} \left(\text{Tr} [\mathbf{A}^\top \mathbf{A} f(\mathbf{B}\mathbf{B}^\top)] - 2c_1 \cdot \text{Tr} [\mathbf{B}\mathbf{D}\mathbf{A}] + \text{Tr} [\mathbf{D}^2] \right) \quad (19)$$

under the constraint $\|\mathbf{B}_{i,:}\|_2 = 1$. Moreover, $\hat{\mathcal{R}}(r) = \tilde{\mathcal{R}}(r)$.

Lemma 5.1 can be extended to any odd $\sigma \in L^2(\mathbb{R}, \mu)$ at the cost of losing the equivalence between $\hat{\mathcal{R}}(r)$ and $\tilde{\mathcal{R}}(r)$.

We restrict the theoretical analysis to proving a lower bound on (19) in the weight-tied setting (14). This lower bound is achieved via the choice of $\mathbf{A}, \mathbf{B}$ in Proposition 5.3, and we give numerical evidence (see Figure 6 in Appendix I) that gradient descent saturates the bound *without* the weight-tying constraint. Thus, we expect our lower bound to hold also for general (not necessarily weight-tied) matrices. The lower bound is given by the minimum of

$$\frac{1}{d} \left(\frac{g(1)}{n} \left(\sum_{i=1}^K \beta_i \right)^2 + \sum_{i=1}^K \left(c_1^2 \frac{\beta_i^2}{s_i} - 2c_1 D_i \beta_i + D_i^2 \right) \right) \quad (20)$$

over all $\beta_i \geq 0$ and

$$0 \leq s_i \leq \min\{k_i, n\}, \quad 1 \leq \sum_{i=1}^K s_i \leq \min\{d, n\}. \quad (21)$$

Here $g(x) := f(x) - c_1^2 x$, and we use the convention that $\frac{0^2}{0} = 0$ and $\frac{c}{0} = +\infty$ for $c > 0$. We can also explicitly characterize the optimal s_i, β_i . The optimal s_i are obtained via a *water-filling criterion*:

$$\mathbf{s} = [k_1, \dots, k_{\text{id}(n)-1}, \text{res}(n), 0, \dots, 0], \quad (22)$$

where $\mathbf{s} = [s_1, \dots, s_K]$, $\text{id}(n)$ denotes the first position at which $\min\{n, d\} - \sum_{i=1}^{\text{id}(n)} k_i < 0$, and $\text{res}(n) := \min\{n, d\} - \sum_{i=1}^{\text{id}(n)-1} k_i$. The β_i can also be expressed explicitly in terms of f, s_i, D_i . This is summarized in the following theorem.

Theorem 5.2. *Consider the objective (19) under the weight-tied constraint (14). Then,*

$$(19) \geq \text{LB}(\mathbf{D}) := \min_{s_i, \beta_i} (20), \quad (23)$$

where $\beta_i \geq 0$ and the s_i satisfy (21). Furthermore, the minimizers of (20) are the s_i obtained via the water-filling criterion (22) and

$$\beta_i = \begin{cases} \frac{s_i}{c_1} \cdot \left(\frac{\frac{g(1)}{c_1^2 n} \sum_{j=1}^{M^*} s_j \Delta_j + D_1}{\frac{g(1)}{c_1^2 n} \sum_{j=1}^{M^*} s_j + 1} - \Delta_i \right), & i \leq M^*, \\ 0, & i > M^*, \end{cases} \quad (24)$$

where $\Delta_j = D_1 - D_j$ and M^* is smallest index such that

$$\frac{g(1)}{c_1^2 n} \sum_{j=1}^{M^*+1} s_j (D_{M^*+1} - D_j) + D_{M^*+1} \leq 0.$$

If no such index exists, then $M^* = K$.

We give a high-level overview of the proof below, and the complete argument is provided in Appendix H.

Proof sketch of Theorem 5.2. We first show that (23) holds. Consider the following block decomposition of $\mathbf{B}$ having the same block structure as $\mathbf{D}$:

$$\mathbf{B} = [\Gamma_1 \mathbf{B}_1 | \dots | \Gamma_K \mathbf{B}_K], \quad (25)$$

where $\mathbf{B}_j \in \mathbb{R}^{n \times k_j}$ with $\|(\mathbf{B}_j)_{i,:}\|_2 = 1$ and $\{\Gamma_j\}_{j=1}^K$ are diagonal matrices with $\sum_{j=1}^K \Gamma_j^2 = \mathbf{I}$. Each $\mathbf{B}_i$ plays a similar role to the $\mathbf{B}$ in the isotropic case. The crucial bound for this step comes from Theorem A in Khare (2021):

$$(\Gamma_i \mathbf{B}_i \mathbf{B}_i^\top \Gamma_i)^{\circ 2} \succeq \frac{1}{s_i} \cdot \text{Diag}(\Gamma_i^2) \text{Diag}(\Gamma_i^2)^\top,$$

where $s_i = \text{rank}(\mathbf{B}_i \mathbf{B}_i^\top)$. Now, ignoring the (PSD) cross-terms for $i \neq j$ we can proceed as in the proof of Proposition 4.3 to arrive at (20). It then remains to minimize (20), which is done using tools from convex analysis. $\square$

Asymptotic achievability. We show that the lower bound in Theorem 5.2 can be asymptotically (i.e., as $d \rightarrow \infty$)

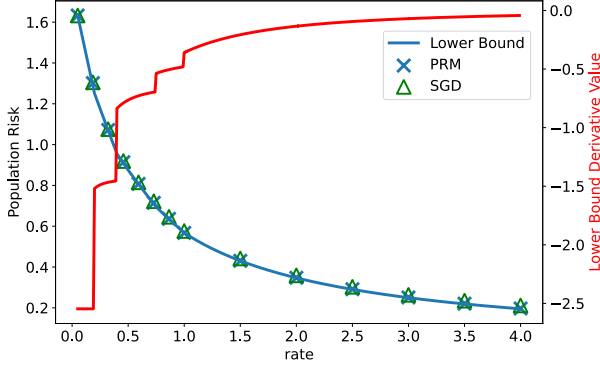


Figure 3. Performance comparison for the compression ($\sigma \equiv \text{sign}$) of a non-isotropic Gaussian source for $\mathbf{k} = (20, 20, 35, 25)$ and $(D_1, D_2, D_3, D_4) = (2, 1.5, 1, 0.8)$.

achieved by using the block form (25), after carefully picking $\mathbf{B}_i$ for each block. Specifically, first we generate a matrix $\mathbf{U} \in \mathbb{R}^{n \times n}$ which is sampled uniformly from the group of orthogonal matrices. Next, we choose each $\mathbf{B}_i$ such that $\hat{\mathbf{B}}_i \hat{\mathbf{B}}_i^\top = \frac{n}{k_i} \mathbf{U} \mathbf{D}_i \mathbf{U}^\top$, where $\mathbf{D}_i$ is a diagonal matrix with

$$(\mathbf{D}_i)_{v,v} = \begin{cases} 1, & \text{if } \sum_{j=1}^{i-1} k_j < v \leq \sum_{j=1}^i k_j, \\ 0, & \text{otherwise,} \end{cases}$$

and the rows of $\mathbf{B}_i$ are given by $\mathbf{b}_i = \frac{\hat{\mathbf{b}}_i}{\|\hat{\mathbf{b}}_i\|_2}$. Furthermore, we pick $\mathbf{\Gamma}_i^2 = \frac{\gamma_i}{n} \mathbf{I}$ and $\mathbf{A} = \beta \mathbf{B}^\top$. The scalings γ_i and β are chosen such that $\beta_i := \beta \gamma_i$ are the minimizers of (20) for s_i as in (21). This is formalized in the following proposition.

Proposition 5.3. *Assume $\mathbf{A}, \mathbf{B}$ are constructed as described above and fix $r > 0$. Also assume that, for all i , $\frac{k_i}{n}$ converges to a strictly positive number as $d \rightarrow \infty$. Then, for any $\epsilon > 0$, with probability $1 - \frac{c}{d^2}$,*

$$|\mathcal{R}(\mathbf{A}, \mathbf{B}) - \text{LB}(\mathbf{D})| \leq C d^{-\frac{1}{2} + \epsilon},$$

where $\text{LB}(\mathbf{D})$ is defined in (23), and the constants c, C only depend on r, ϵ and $\lim_{d \rightarrow \infty} \frac{k_i}{n}$.

The proof of this lemma is similar to that of Proposition 4.4, and it is provided in Appendix H.

Taken together, Proposition 5.3 and Theorem 5.2 show that the optimal $\mathbf{B}$ exhibits the block structure (25), which agrees with the block structure (18) of the data covariance. The individual blocks are orthogonal in the sense that $\mathbf{B}_i^\top \mathbf{\Gamma}_i \mathbf{\Gamma}_j \mathbf{B}_j = \mathbf{0}$. Furthermore, we expect each block to have the same form as the minimizers in the isotropic case, up to some scaling. Such a structure is also confirmed by the numerical experiments: for instance, it is observed in the setting considered for Figure 6 in Appendix I.

Interpretation of the water-filling solution. To provide an intuitive illustration of the property of solutions in Proposition 5.3, consider the case of $K = 2$ and $k_1 = k_2$. In

this case, the eigenvalues of $\mathbf{B} \mathbf{B}^\top$ will take only the two values λ_1, λ_2 corresponding to the two blocks. For rate $r \leq k_1/d = 1/2$, we have $\lambda_2 = 0$ since *water-filling* implies that only the block corresponding to D_1 is utilized by $\mathbf{B}$. For rate $r > 1/2$, we have $\lambda_2 > 0$, as soon as

$$\left(1 + \frac{2c_1^2}{f(1) - c_1^2} r\right) D_2 > D_1.$$

The inequality can be obtained by evaluating explicitly the condition stated in Theorem 5.2 below Equation (24) in this special case. Furthermore, in the limit $r \rightarrow \infty$, we have that $\frac{\lambda_1}{\lambda_2} \rightarrow \frac{D_1}{D_2}$. This means that, for sufficiently large rates, the weight given by the encoder to each block is proportional to the corresponding eigenvalue of the data.

The *water-filling* behaviour can be observed in a setting with four blocks in Figure 3. Namely, at each of the rates $\{k_1/d, (k_1+k_2)/d, (k_1+k_2+k_3)/d, (k_1+k_2+k_3+k_4)/d\}$ that correspond to the earlier blocks being utilized to their full capacity, the derivative of the lower bound experiences a “jump” at which the next λ_i becomes positive.

6. Discussion

Population vs. empirical loss. All our results hold for the optimization of the population loss. Extending them to the empirical loss is an interesting direction for future research. One possible way forward is to exploit progress towards relating the landscape of empirical and population losses, see e.g. (Mei et al., 2018). We remark that, in the simulations of gradient descent, we always use the tempered straight-through estimator of the sign activation (see Appendix I for details). Thus, another promising direction is to show that, in the low-temperature regime (i.e., when the differentiable approximation of the sign becomes almost perfect), the gradient-based scheme converges to the minimizer of the population risk.

Optimality of two-layer autoencoders. This paper characterizes the minimizers of the expected ℓ_2 error incurred by two-layer autoencoders, and it shows that the minimum error is achieved, under certain conditions, by gradient-based algorithms. Thus, for the special case in which $\sigma \equiv \text{sign}$, a natural question is to what degree the model (1) is suitable for data compression. Let us fix the encoder to be a rotationally invariant matrix, i.e., $\mathbf{B} = \mathbf{U} \mathbf{\Lambda} \mathbf{V}^\top$ with $\mathbf{U}, \mathbf{V}$ independent and distributed according to the Haar measure and $\mathbf{\Lambda}$ having bounded entries. Then, the information-theoretically optimal reconstruction error can be computed via the replica method from statistical mechanics (Tulino et al., 2013) and, in a number of scenarios, it coincides with the error of a Vector Approximate Message Passing (VAMP) algorithm (Rangan et al., 2019; Schniter et al., 2016). Furthermore, it is also possible to optimize the spectrum $\mathbf{\Lambda}$ to minimize the error, which leads to the singular values of $\mathbf{B}$ being all

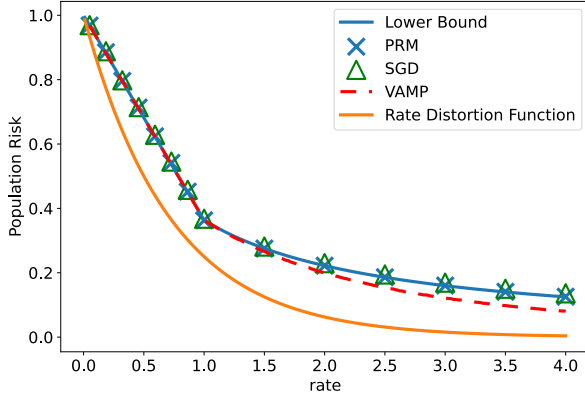


Figure 4. Performance comparison for the compression ($\sigma \equiv \text{sign}$) of an isotropic Gaussian source.

1 (Ma et al., 2021).² Surprisingly, for a compression rate $r \leq 1$, the optimal error found in (Ma et al., 2021) coincides with the minimizer of the population loss given by Theorem 4.2. Hence, two-layer autoencoders are optimal compressors under two conditions: (i) $r \leq 1$, and (ii) fixed encoder given by a rotationally invariant matrix. Both conditions are sufficient and also necessary. For $r > 1$, VAMP outperforms the two-layer autoencoder. Moreover, for a general encoder/decoder pair, the information-theoretically optimal reconstruction error is given by the rate-distortion function, which outperforms two-layer autoencoders for all $r > 0$. This picture is summarized in Figure 4: the blue curve represents the lower bound of Theorem 4.2 (for $r \leq 1$) and Proposition 4.3 (for $r > 1$), which is met by either running GD on the population risk (blue crosses) or SGD on samples taken from a isotropic Gaussian (green triangles) when $d = 100$;³ this lower bound meets the performance of VAMP (red curve) if and only if $r \leq 1$; finally, the rate distortion function (orange curve) provides the best performance for all $r > 0$.

Universality of Gaussian predictions. Figure 4 (and also Figure 6 in Appendix I) show that gradient descent achieves the minimum of the population risk for the compression of Gaussian sources. Going beyond Gaussian inputs, to real-world datasets, Figures 1-2 (as well as those in Appendix I) show an excellent agreement between our predictions (using the empirical covariance of the data) and the performance of autoencoders trained on standard datasets (CIFAR-10, MNIST). As such, this agreement provides a clear indication of the universality of our predictions. In this regard, a flurry of recent research (see e.g. (Hastie et al., 2022; Hu & Lu, 2022; Loureiro et al., 2021; Goldt et al., 2022; Dudeja et al.,

²More specifically, Ma et al. (2021) consider an expectation propagation (EP) algorithm (Minka, 2001; Oppor et al., 2005; Fletcher et al., 2016; He et al., 2017), which has been related to various forms of approximate message passing (Ma & Ping, 2017; Rangan et al., 2019).

³For further details on the experimental setup, see Appendix I.

2022; Montanari & Saeed, 2022; Wang et al., 2022)) has proved that the Gaussian predictions actually hold in a much wider range of models. While none of the existing works exactly fits the setting considered in this paper, this gives yet another indication that our predictions should remain true more generally. The rigorous characterization of this universality is left for future work.

The choice of the activation function. The sign activation function constitutes an important special case of our analysis. However, our results hold for a broader class of activations. In particular, under the restriction that the rows of the encoder B lie on the unit sphere, all the results apply for any odd activation. The reason to fix the norm of the rows of the encoder is to prevent the network from entering the linear regime (e.g., by scaling $B \rightarrow \epsilon B$ and $A \rightarrow \frac{1}{\epsilon} A$). In fact, in the linear regime, perfect recovery can be achieved and this case has been well studied, see e.g. (Baldi & Hornik, 1989; Kunin et al., 2019; Gidel et al., 2019). We also note that, if the activation function is homogeneous, the restriction on the norm of the rows of B can be lifted, as the norm can be scaled out. Extending our analysis to activation functions that are not odd (e.g., ReLU) is an exciting avenue for future research. To achieve this goal, we expect that novel ideas will be needed, since our current approach relies on the fact that the Hermite expansion of the activation function (4) has only odd monomials.

Acknowledgements

Aleksandr Shevchenko, Kevin Kögler and Marco Mondelli are supported by the 2019 Lopez-Loreta Prize. Hamed Hassani acknowledges the support by the NSF CIF award (1910056) and the NSF Institute for CORE Emerging Methods in Data Science (EnCORE).

References

- Absil, P.-A., Mahony, R., and Sepulchre, R. Optimization algorithms on matrix manifolds. Princeton University Press, 2009.
- Agustsson, E., Mentzer, F., Tschannen, M., Cavigelli, L., Timofte, R., Benini, L., and Gool, L. Soft-to-hard vector quantization for end-to-end learning compressible representations. In *NeurIPS*, 2017.
- Arimoto, S. An algorithm for computing the capacity of arbitrary discrete memoryless channels. *IEEE Transactions on Information Theory*, 18(1):14–20, 1972. doi: 10.1109/TIT.1972.1054753.
- Baldi, P. and Hornik, K. Neural networks and principal component analysis: Learning from examples without local minima. *Neural networks*, 2(1):53–58, 1989.

- Ballé, J., Laparra, V., and Simoncelli, E. P. End-to-end optimized image compression. In *International Conference on Learning Representations*, 2017.
- Ballé, J., Chou, P. A., Minnen, D., Singh, S., Johnston, N., Agustsson, E., Hwang, S. J., and Toderici, G. Nonlinear transform coding. *IEEE Trans. on Special Topics in Signal Processing*, 15, 2021. URL <https://arxiv.org/pdf/2007.03034>.
- Bao, X., Lucas, J., Sachdeva, S., and Grosse, R. B. Regularized linear autoencoders recover the principal components, eventually. In *NeurIPS*, 2020.
- Bengio, Y., Yao, L., Alain, G., and Vincent, P. Generalized denoising auto-encoders as generative models. In *NeurIPS*, 2013.
- Blahut, R. Computation of channel capacity and rate-distortion functions. *IEEE Transactions on Information Theory*, 18(4):460–473, 1972. doi: 10.1109/TIT.1972.1054855.
- Boufounos, P. T. and Baraniuk, R. G. 1-bit compressive sensing. In *2008 42nd Annual Conference on Information Sciences and Systems*, pp. 16–21. IEEE, 2008.
- Boyd, S., Boyd, S. P., and Vandenberghe, L. *Convex optimization*. Cambridge University Press, 2004.
- Candes, E. J., Strohmer, T., and Voroninski, V. Phaselift: Exact and stable signal recovery from magnitude measurements via convex programming. *Communications on Pure and Applied Mathematics*, 66(8):1241–1274, 2013.
- Candes, E. J., Li, X., and Soltanolkotabi, M. Phase retrieval via wirtinger flow: Theory and algorithms. *IEEE Transactions on Information Theory*, 61(4):1985–2007, 2015.
- Cheng, Z., Sun, H., Takeuchi, M., and Katto, J. Learned image compression with discretized gaussian mixture likelihoods and attention modules. In *Conference on Computer Vision and Pattern Recognition*, 2020.
- Chizat, L., Oyallon, E., and Bach, F. On lazy training in differentiable programming. In *NeurIPS*, 2019.
- Ciliberti, S., Mézard, M., and Zecchina, R. Message-passing algorithms for non-linear nodes and data compression. *ComplexUs*, 3(1-3):58–65, 2006.
- Cover, T. M. and Thomas, J. A. *Elements of Information Theory (Wiley Series in Telecommunications and Signal Processing)*. Wiley-Interscience, USA, 2006. ISBN 0471241954.
- del Pino, G. E. and Galaz, H. Statistical applications of the inverse gram matrix: A revisitation. *Brazilian Journal of Probability and Statistics*, pp. 177–196, 1995.
- Dudeja, R., Sen, S., and Lu, Y. M. Spectral universality of regularized linear regression with nearly deterministic sensing matrices. *arXiv preprint arXiv:2208.02753*, 2022.
- Fletcher, A., Sahraee-Ardakan, M., Rangan, S., and Schniter, P. Expectation consistent approximate inference: Generalizations and convergence. In *2016 IEEE International Symposium on Information Theory (ISIT)*, pp. 190–194. IEEE, 2016.
- Fletcher, A. K., Rangan, S., and Schniter, P. Inference in deep networks in high dimensions. In *2018 IEEE International Symposium on Information Theory (ISIT)*, pp. 1884–1888. IEEE, 2018.
- Gidel, G., Bach, F., and Lacoste-Julien, S. Implicit regularization of discrete gradient dynamics in linear neural networks. In *NeurIPS*, 2019.
- Goldt, S., Loureiro, B., Reeves, G., Krzakala, F., Mézard, M., and Zdeborová, L. The gaussian equivalence of generative models for learning with shallow neural networks. In *Mathematical and Scientific Machine Learning*, pp. 426–471. PMLR, 2022.
- Gray, R. Vector quantization. *IEEE Assp Magazine*, 1(2): 4–29, 1984.
- Hastie, T., Montanari, A., Rosset, S., and Tibshirani, R. J. Surprises in high-dimensional ridgeless least squares interpolation. *The Annals of Statistics*, 50(2):949–986, 2022.
- He, H., Wen, C.-K., and Jin, S. Generalized expectation consistent signal recovery for nonlinear measurements. In *2017 IEEE International Symposium on Information Theory (ISIT)*, pp. 2333–2337. IEEE, 2017.
- Hinton, G. E. and Salakhutdinov, R. R. Reducing the dimensionality of data with neural networks. *Science*, 313 (5786):504–507, 2006.
- Hu, H. and Lu, Y. M. Universality laws for high-dimensional learning with random features. *IEEE Transactions on Information Theory*, 2022.
- Jacot, A., Gabriel, F., and Hongler, C. Neural tangent kernel: Convergence and generalization in neural networks. In *NeurIPS*, 2018.
- Jain, S., Radhakrishnan, A., and Uhler, C. A mechanism for producing aligned latent spaces with autoencoders. *arXiv preprint arXiv:2106.15456*, 2021.

- Khare, A. Sharp nonzero lower bounds for the schur product theorem. *Proceedings of the American Mathematical Society*, 149(12):5049–5063, 2021.
- Kingma, D. P. and Welling, M. Auto-encoding variational bayes. In *International Conference on Learning Representations*, 2014.
- Korada, S. B. and Urbanke, R. L. Polar codes are optimal for lossy source coding. *IEEE Transactions on Information Theory*, 56(4):1751–1768, 2010.
- Kunin, D., Bloom, J., Goeva, A., and Seed, C. Loss landscapes of regularized linear autoencoders. In *International Conference on Machine Learning*, 2019.
- Lei, E., Hassani, H., and Bidokhti, S. S. Neural estimation of the rate-distortion function for massive datasets. In *2022 IEEE International Symposium on Information Theory (ISIT)*, pp. 608–613. IEEE, 2022.
- Liu, H., Chen, T., Guo, P., Shen, Q., Cao, X., Wang, Y., and Ma, Z. Non-local attention optimized deep image compression. *arXiv preprint arXiv:1904.09757*, 2019.
- Loureiro, B., Gerbelot, C., Cui, H., Goldt, S., Krzakala, F., Mezard, M., and Zdeborová, L. Learning curves of generic features maps for realistic datasets with a teacher-student model. In *NeurIPS*, 2021.
- Ma, J. and Ping, L. Orthogonal amp. *IEEE Access*, 5: 2020–2033, 2017.
- Ma, J., Xu, J., and Maleki, A. Analysis of sensing spectral for signal recovery under a generalized linear model. In *NeurIPS*, 2021.
- Matsumoto, N. and Mazumdar, A. Binary iterative hard thresholding converges with optimal number of measurements for 1-bit compressed sensing. *arXiv preprint arXiv:2207.03427*, 2022.
- Meckes, E. S. *The random matrix theory of the classical compact groups*, volume 218. Cambridge University Press, 2019.
- Mei, S., Bai, Y., and Montanari, A. The landscape of empirical risk for nonconvex losses. *The Annals of Statistics*, 46(6A):2747–2774, 2018.
- Minka, T. P. Expectation propagation for approximate bayesian inference. In *Proceedings of the Seventeenth conference on Uncertainty in artificial intelligence*, pp. 362–369, 2001.
- Montanari, A. and Saeed, B. N. Universality of empirical risk minimization. In *Conference on Learning Theory*, 2022.
- Nguyen, P.-M. Analysis of feature learning in weight-tied autoencoders via the mean field lens. *arXiv preprint arXiv:2102.08373*, 2021.
- Nguyen, T. V., Wong, R. K., and Hegde, C. On the dynamics of gradient descent for autoencoders. In *International Conference on Artificial Intelligence and Statistics*, 2019.
- Nguyen, T. V., Wong, R. K., and Hegde, C. Benefits of jointly training autoencoders: An improved neural tangent kernel analysis. *IEEE Transactions on Information Theory*, 67(7):4669–4692, 2021.
- O’Donnell, R. *Analysis of boolean functions*. Cambridge University Press, 2014.
- Oftadeh, R., Shen, J., Wang, Z., and Shell, D. Eliminating the invariance on the loss landscape of linear autoencoders. In *International Conference on Machine Learning*, 2020.
- Oppor, M., Winther, O., and Jordan, M. J. Expectation consistent approximate inference. *Journal of Machine Learning Research*, 6(12), 2005.
- Radhakrishnan, A., Belkin, M., and Uhler, C. Overparameterized neural networks implement associative memory. *Proceedings of the National Academy of Sciences*, 117 (44):27162–27170, 2020.
- Rangamani, A., Mukherjee, A., Basu, A., Arora, A., Ganapathi, T., Chin, S., and Tran, T. D. Sparse coding and autoencoders. In *2018 IEEE International Symposium on Information Theory (ISIT)*, pp. 36–40. IEEE, 2018.
- Rangan, S., Schniter, P., and Fletcher, A. K. Vector approximate message passing. *IEEE Transactions on Information Theory*, 65(10):6664–6684, 2019.
- Refinetti, M. and Goldt, S. The dynamics of representation learning in shallow, non-linear autoencoders. In *International Conference on Machine Learning*, 2022.
- Rezende, D. J., Mohamed, S., and Wierstra, D. Stochastic backpropagation and approximate inference in deep generative models. In *International Conference on Machine Learning*, 2014.
- Santambrogio, F. {Euclidean, metric, and Wasserstein} gradient flows: an overview. *Bulletin of Mathematical Sciences*, 7(1):87–154, 2017.
- Schniter, P., Rangan, S., and Fletcher, A. K. Vector approximate message passing for the generalized linear model. In *2016 50th Asilomar Conference on Signals, Systems and Computers*, pp. 1525–1529. IEEE, 2016.

- Shannon, C. E. A mathematical theory of communication. *The Bell System Technical Journal*, 27(3):379–423, 1948. doi: 10.1002/j.1538-7305.1948.tb01338.x.
- Shannon, C. E. Coding theorems for a discrete source with a fidelity criterion. *1959 IRE National Convention Record*, pp. 142–163, 1959.
- Theis, L., Shi, W., Cunningham, A., and Huszár, F. Lossy image compression with compressive autoencoders. In *International Conference on Learning Representations*, 2017.
- Theis, L., Salimans, T., Hoffman, M. D., and Mentzer, F. Lossy compression with gaussian diffusion. *arXiv preprint arXiv:2206.08889*, 2022.
- Tschannen, M., Bachem, O., and Lucic, M. Recent advances in autoencoder-based representation learning. *arXiv preprint arXiv:1812.05069*, 2018.
- Tulino, A. M., Caire, G., Verdú, S., and Shamai, S. Support recovery with sparsely sampled free random matrices. *IEEE Transactions on Information Theory*, 59(7):4243–4271, 2013.
- Vershynin, R. *High-dimensional probability: An introduction with applications in data science*, volume 47. Cambridge University Press, 2018.
- Vincent, P., Larochelle, H., Bengio, Y., and Manzagol, P.-A. Extracting and composing robust features with denoising autoencoders. In *International Conference on Machine Learning*, 2008.
- Visick, G. A quantitative version of the observation that the hadamard product is a principal submatrix of the kronecker product. *Linear Algebra and Its Applications*, 304(1-3):45–68, 2000.
- Wagner, A. B. and Ballé, J. Neural networks optimally compress the sawbridge. *2021 Data Compression Conference (DCC)*, pp. 143–152, 2021.
- Wainwright, M. J., Maneva, E., and Martinian, E. Lossy source compression using low-density generator matrix codes: Analysis and algorithms. *IEEE Transactions on Information theory*, 56(3):1351–1368, 2010.
- Wang, T., Zhong, X., and Fan, Z. Universality of approximate message passing algorithms and tensor networks. *arXiv preprint arXiv:2206.13037*, 2022.
- Wright, S. J. Coordinate descent algorithms. *Mathematical Programming*, 151(1):3–34, 2015.
- Yang, R. and Mandt, S. Lossy image compression with conditional diffusion models. *arXiv preprint arXiv:2209.06950*, 2022.
- Yang, Y. and Mandt, S. Towards empirical sandwich bounds on the rate-distortion function. In *International Conference on Learning Representations*, 2021.
- Yang, Y., Mandt, S., and Theis, L. An introduction to neural data compression. *arXiv preprint arXiv:2202.06533*, 2022.
- Yin, P., Lyu, J., Zhang, S., Osher, S., Qi, Y., and Xin, J. Understanding straight-through estimator in training activation quantized neural nets. *arXiv preprint arXiv:1903.05662*, 2019.

A. Additional Related Works

Rate-distortion formalism. Lossy compression of stationary sources is a classical problem in information theory, and several approaches have been proposed, including vector quantization (Gray, 1984), or the usage of powerful channel codes (Korada & Urbanke, 2010; Ciliberti et al., 2006; Wainwright et al., 2010). The rate-distortion function characterizes the optimal trade-off between error and size of the representation for the compression of an i.i.d. source (Shannon, 1948; 1959; Cover & Thomas, 2006). However, computing the rate-distortion function is by itself a challenging task. The Blahut-Arimoto scheme (Blahut, 1972; Arimoto, 1972) provides a systematic approach, but it suffers from the issue of scalability (Lei et al., 2022). Consequently, to compute the rate-distortion of empirical datasets, approximate methods based on generative modeling have been proposed (Yang & Mandt, 2021; Lei et al., 2022).

Non-linear inverse problems. The task of estimating a signal $\mathbf{x}$ from non-linear measurements $\mathbf{y} = \sigma(\mathbf{B}\mathbf{x})$ has appeared in many areas, such as 1-bit compressed sensing where $\sigma(z) = \text{sign}(z)$ (Boufounos & Baraniuk, 2008), or phase retrieval where $\sigma(z) = |z|$ (Candes et al., 2013; 2015). While the focus of these problems is different from ours (e.g., compressed sensing has often an additional sparsity assumption), the ideas and proof techniques developed in this paper might be beneficial to characterize the fundamental limits and the performance of gradient-based methods for general inverse reconstruction tasks, see e.g. (Ma et al., 2021; Matsumoto & Mazumdar, 2022).

B. Closed Forms for the Population Risk

For the proofs of Lemmas 4.1 and 5.1 in the current section, we assume that the rows of $\mathbf{B}$ have non-zero norm, hence, in particular, they may be chosen to have unit norm. In the end of the section, we elaborate on why this assumption holds true.

Let us also mention that we call σ odd in L^2 sense. For this particular case, it means that $\sigma(x) = \sigma(-x)$ for $x \neq 0$ and $|\sigma(0)| < C$, where C is some universal constant. This concern is purely technical, since the main application of our results is 1-bit compression. Namely, we do not set $\sigma(0) = \text{sgn}(0) = 0$. In fact, this would mean that the compressed sequence can take values in $\{-1, 0, 1\}$, which would not result in 1-bit compression, but rather in $\log_2(3)$ -bits compression. It is safe to ignore this technicality and intuitively assume that $\sigma(0) = 0$.

Proof of Lemma 4.1. Opening up the two-norm gives

$$\mathbb{E}\|\mathbf{x} - \mathbf{A}\sigma(\mathbf{B}\mathbf{x})\|_2^2 = \mathbb{E}\|\mathbf{x}\|_2^2 + \mathbb{E}\|\mathbf{A}\sigma(\mathbf{B}\mathbf{x})\|_2^2 - 2\mathbb{E}\langle \mathbf{x}, \mathbf{A}\sigma(\mathbf{B}\mathbf{x}) \rangle. \quad (26)$$

Since $\mathbf{x} \sim \mathcal{N}(\mathbf{0}, \mathbf{I})$, we get

$$\mathbb{E}\|\mathbf{x}\|_2^2 = d. \quad (27)$$

Let $\mathbf{B}^\top = [\mathbf{b}_1, \dots, \mathbf{b}_n] \in \mathbb{R}^{d \times n}$ and $\mathbf{A} = [\mathbf{a}_1, \dots, \mathbf{a}_n] \in \mathbb{R}^{d \times n}$, with $\|\mathbf{b}_i\|_2 = \|\mathbf{B}_{i,:}\| = 1$. Rewriting the second term in (26) gives

$$\mathbb{E}\|\mathbf{A}\sigma(\mathbf{B}\mathbf{x})\|_2^2 = \sum_{i,j=1}^n \langle \mathbf{a}_i, \mathbf{a}_j \rangle \cdot \mathbb{E}[\sigma(\langle \mathbf{b}_i, \mathbf{x} \rangle) \cdot \sigma(\langle \mathbf{b}_j, \mathbf{x} \rangle)]. \quad (28)$$

Using the reproducing property of Hermite coefficients (see, e.g., Chapter 11 in (O'Donnell, 2014)), since the random variables $\langle \mathbf{b}_i, \mathbf{x} \rangle$ and $\langle \mathbf{b}_j, \mathbf{x} \rangle$ are $\langle \mathbf{b}_i, \mathbf{b}_j \rangle$ -correlated, we have

$$\mathbb{E}[h_{2\ell+1}(\langle \mathbf{b}_i, \mathbf{x} \rangle) \cdot h_{2\ell+1}(\langle \mathbf{b}_j, \mathbf{x} \rangle)] = \langle \mathbf{b}_i, \mathbf{b}_j \rangle^{2\ell+1}, \quad \mathbb{E}[h_{2\ell+1}(\langle \mathbf{b}_i, \mathbf{x} \rangle) \cdot h_{2k+1}(\langle \mathbf{b}_j, \mathbf{x} \rangle)] = 0,$$

for $k \neq \ell$. This gives that

$$\mathbb{E}[\sigma(\langle \mathbf{b}_i, \mathbf{x} \rangle) \cdot \sigma(\langle \mathbf{b}_j, \mathbf{x} \rangle)] = \sum_{\ell=0}^{\infty} (c_{2\ell+1})^2 \langle \mathbf{b}_i, \mathbf{b}_j \rangle^{2\ell+1} = f(\langle \mathbf{b}_i, \mathbf{b}_j \rangle),$$

and, hence, using (28) we arrive to

$$\mathbb{E}\|\mathbf{A}\sigma(\mathbf{B}\mathbf{x})\|_2^2 = \sum_{i,j=1}^n \langle \mathbf{a}_i, \mathbf{a}_j \rangle \cdot f(\langle \mathbf{b}_i, \mathbf{b}_j \rangle) = \text{Tr} \left[\mathbf{A}^\top \mathbf{A} \cdot f(\mathbf{B}\mathbf{B}^\top) \right]. \quad (29)$$

Rearranging the last term in (26) gives

$$\mathbb{E}\langle \mathbf{x}, \mathbf{A}\sigma(\mathbf{B}\mathbf{x}) \rangle = \sum_{i=1}^d \sum_{j=1}^n a_j^i \cdot \mathbb{E}[x_i \sigma(\langle \mathbf{b}_j, \mathbf{x} \rangle)], \quad (30)$$

where a_j^i stands for the i -th coordinate of the vector $\mathbf{a}_j$ and x_i stands for the i -th coordinate of the vector $\mathbf{x}$. Let us now compute the inner expected value for each pair (i, j) . Notice that the random variables $\langle \mathbf{b}_j, \mathbf{x} \rangle$ and x_i are jointly Gaussian with zero mean and covariance matrix $\tilde{\Sigma} \in \mathbb{R}^{2 \times 2}$:

$$\tilde{\Sigma}_{21} = \tilde{\Sigma}_{12} = \mathbb{E}x_i \langle \mathbf{b}_j, \mathbf{x} \rangle = \mathbb{E}b_j^i x_i^2 = b_j^i, \quad \tilde{\Sigma}_{11} = \mathbb{E}\langle \mathbf{b}_j, \mathbf{x} \rangle^2 = \|\mathbf{b}_j\|_2^2 = 1, \quad \tilde{\Sigma}_{22} = \mathbb{E}x_i^2 = 1.$$

Hence, the random vectors $(\langle \mathbf{b}_j, \mathbf{x} \rangle, x_i)$ and

$$\left(y_1, b_j^i \cdot y_1 + \sqrt{1 - (b_j^i)^2} \cdot y_2 \right), \quad \text{with } (y_1, y_2) \sim \mathcal{N}(0, \mathbf{I})$$

are identically distributed. In this view, we obtain

$$\begin{aligned} \mathbb{E}[x_i \sigma(\langle \mathbf{b}_j, \mathbf{x} \rangle)] &= \mathbb{E}\left[\left(b_j^i \cdot y_1 + \sqrt{1 - (b_j^i)^2} \cdot y_2\right) \sigma(y_1)\right] \\ &= b_j^i \cdot \mathbb{E}[y_1 \sigma(y_1)] + \sqrt{1 - (b_j^i)^2} \cdot \mathbb{E}[y_2] \cdot \mathbb{E}[\sigma(y_1)] = c_1 \cdot b_j^i, \end{aligned} \quad (31)$$

where we applied the reproducing property to conclude that $\mathbb{E}[y_1 \sigma(y_1)] = c_1$. Consequently, by combining (30) and (31), we get that

$$\mathbb{E}\langle \mathbf{x}, \mathbf{A}\sigma(\mathbf{B}\mathbf{x}) \rangle = c_1 \cdot \sum_{i=1}^d \sum_{j=1}^n a_j^i b_j^i = c_1 \cdot \text{Tr}[\mathbf{B}\mathbf{A}]. \quad (32)$$

By combining (26), (27), (29) and (32), we obtain the desired expression for $\tilde{R}(r)$.

Assume now that σ is homogeneous. Then, in (28) and (30), the norm of $\mathbf{b}_i$ can be pushed into the corresponding $\mathbf{a}_i$ and, hence, we obtain

$$\min_{\mathbf{A}, \mathbf{B}} \mathbb{E}\|\mathbf{x} - \mathbf{A}\sigma(\mathbf{B}\mathbf{x})\|_2^2 = \min_{\mathbf{A}, \|\mathbf{B}_i\|_2=1} \mathbb{E}\|\mathbf{x} - \mathbf{A}\sigma(\mathbf{B}\mathbf{x})\|_2^2,$$

which proves that $\hat{\mathcal{R}}(r) = \tilde{\mathcal{R}}(r)$.

Finally, consider the case $\sigma(x) = \text{sign}(x)$. Then, Grothendieck's identity (see, e.g., Lemma 3.6.6 in (Vershynin, 2018)) gives

$$\mathbb{E}\sigma(\langle \mathbf{b}_i, \mathbf{x} \rangle) \sigma(\langle \mathbf{b}_j, \mathbf{x} \rangle) = \frac{2}{\pi} \arcsin(\langle \mathbf{b}_i, \mathbf{b}_j \rangle) \Rightarrow f(x) = \frac{2}{\pi} \arcsin(x).$$

Recalling that the first Hermite coefficient of $\sigma(x) = \text{sign}(x)$ is equal to $\sqrt{\frac{2}{\pi}}$ finishes the proof. $\square$

Proof of Lemma 5.1. The proof of Lemma 5.1 follows from similar arguments as that of Lemma 4.1. Given this, we only explain the key differences. We first show that it is enough to consider $\Sigma = \mathbf{D}^2$. Given the SVD $\Sigma = \mathbf{U}\mathbf{D}^2\mathbf{U}^\top$, we have $\mathbf{x} = \mathbf{U}\mathbf{D}\tilde{\mathbf{x}}$, where $\tilde{\mathbf{x}} \sim \mathcal{N}(\mathbf{0}, \mathbf{I})$. Now, we can push the rotation $\mathbf{U}$ in $\mathbf{A}, \mathbf{B}$:

$$\|\mathbf{x} - \mathbf{A}\sigma(\mathbf{B}\mathbf{x})\|_2 = \left\| \mathbf{D}\tilde{\mathbf{x}} - \mathbf{U}^\top \mathbf{A}\sigma(\mathbf{B}\mathbf{U}\mathbf{D}\tilde{\mathbf{x}}) \right\|_2.$$

Thus, after replacing $\mathbf{A}$ with $\mathbf{U}^\top \mathbf{A}$ and $\mathbf{B}$ with $\mathbf{B}\mathbf{U}$, we may assume that $\mathbf{x} = \mathbf{D}\tilde{\mathbf{x}}$.

We again open up the two-norm

$$\mathbb{E}\|\mathbf{x} - \mathbf{A}\sigma(\mathbf{B}\mathbf{x})\|_2^2 = \mathbb{E}\|\mathbf{x}\|_2^2 + \mathbb{E}\|\mathbf{A}\sigma(\mathbf{B}\mathbf{x})\|_2^2 - 2\mathbb{E}\langle \mathbf{x}, \mathbf{A}\sigma(\mathbf{B}\mathbf{x}) \rangle. \quad (33)$$

For the first term, we clearly have

$$\mathbb{E}\|\mathbf{x}\|_2^2 = \text{Tr}[\mathbf{D}^2].$$

Now, for the second term we write

$$\mathbb{E}\|\mathbf{A}\sigma(\mathbf{B}\mathbf{x})\|_2^2 = \mathbb{E}\|\mathbf{A}\sigma(\mathbf{B}\mathbf{D}\tilde{\mathbf{x}})\|_2^2,$$

where $\tilde{\mathbf{x}} \sim \mathcal{N}(\mathbf{0}, \mathbf{I})$. Thus, as in the proof of Lemma 4.1, we have

$$\mathbb{E}\|\mathbf{A}\sigma(\mathbf{B}\mathbf{D}\tilde{\mathbf{x}})\|_2^2 = \text{Tr} \left[\mathbf{A}^\top \mathbf{A} \cdot f(\mathbf{B}\mathbf{D}^2 \mathbf{B}^\top) \right].$$

Similarly, for the last term we obtain

$$\mathbb{E}\langle \tilde{\mathbf{x}}, \mathbf{A}\sigma(\mathbf{B}\mathbf{x}) \rangle = \mathbb{E}\langle \tilde{\mathbf{x}}, \mathbf{D}\mathbf{A}\sigma(\mathbf{B}\mathbf{D}\tilde{\mathbf{x}}) \rangle = c_1 \text{Tr} [\mathbf{D}\mathbf{A}\mathbf{B}\mathbf{D}].$$

Finally, since σ is homogeneous, by abuse of notation we can replace $\mathbf{B}\mathbf{D}$ by any $\mathbf{B}$ with unit-norm rows. This follows from the fact that, similarly to the proof of Lemma 4.1 (namely, equations (28) and (30)), we have that

$$\begin{aligned} \mathbb{E}\|\mathbf{A}\sigma(\mathbf{B}\mathbf{D}\tilde{\mathbf{x}})\|_2^2 &= \sum_{i,j=1}^n \langle \mathbf{a}_i, \mathbf{a}_j \rangle \cdot \mathbb{E} [\sigma(\langle (\mathbf{B}\mathbf{D})_{i,:}, \tilde{\mathbf{x}} \rangle) \cdot \sigma(\langle (\mathbf{B}\mathbf{D})_{j,:}, \tilde{\mathbf{x}} \rangle)], \\ \mathbb{E}\langle \tilde{\mathbf{x}}, \mathbf{D}\mathbf{A}\sigma(\mathbf{B}\mathbf{D}\tilde{\mathbf{x}}) \rangle &= \sum_{i=1}^d \sum_{j=1}^n a_j^i \cdot \mathbb{E} [(D_{i,i} \cdot \tilde{x}_i) \cdot \sigma(\langle (\mathbf{B}\mathbf{D})_{j,:}, \tilde{\mathbf{x}} \rangle)], \end{aligned} \quad (34)$$

which, by homogeneity, readily gives that the norm of $(\mathbf{B}\mathbf{D})_{i,:}$ can be pushed into the corresponding $\mathbf{a}_i$.

As a result, the statement of Lemma 5.1 readily follows by comparing the terms. $\square$

Rows of $\mathbf{B}$ are non-zero. We show that the assumption holds true by contradiction. Without loss of generality, assume that the first $n' \leq n$ rows of $\mathbf{B}$ are zero vectors. Hence, from (34) we can see that the following holds:

$$\begin{aligned} \mathbb{E}\|\mathbf{A}\sigma(\mathbf{B}\mathbf{D}\tilde{\mathbf{x}})\|_2^2 &= \sum_{i,j=n'+1}^n \langle \mathbf{a}_i, \mathbf{a}_j \rangle \cdot \mathbb{E} [\sigma(\langle (\mathbf{B}\mathbf{D})_{i,:}, \tilde{\mathbf{x}} \rangle) \cdot \sigma(\langle (\mathbf{B}\mathbf{D})_{j,:}, \tilde{\mathbf{x}} \rangle)] \\ &\quad + \sum_{i \leq n' \wedge j > n'} \langle \mathbf{a}_i, \mathbf{a}_j \rangle \cdot \mathbb{E} [\sigma(0) \cdot \sigma(\langle (\mathbf{B}\mathbf{D})_{j,:}, \tilde{\mathbf{x}} \rangle)] \\ &\quad + \sum_{i > n' \wedge j \leq n'} \langle \mathbf{a}_i, \mathbf{a}_j \rangle \cdot \mathbb{E} [\sigma(0) \cdot \sigma(\langle (\mathbf{B}\mathbf{D})_{i,:}, \tilde{\mathbf{x}} \rangle)] + \sum_{i,j \leq n'} \langle \mathbf{a}_i, \mathbf{a}_j \rangle \cdot \sigma(0)^2 \\ &= \sum_{i,j=n'+1}^n \langle \mathbf{a}_i, \mathbf{a}_j \rangle \cdot \mathbb{E} [\sigma(\langle (\mathbf{B}\mathbf{D})_{i,:}, \tilde{\mathbf{x}} \rangle) \cdot \sigma(\langle (\mathbf{B}\mathbf{D})_{j,:}, \tilde{\mathbf{x}} \rangle)] + \sum_{i,j \leq n'} \langle \mathbf{a}_i, \mathbf{a}_j \rangle \cdot \sigma(0)^2 \\ &\geq \sum_{i,j=n'+1}^n \langle \mathbf{a}_i, \mathbf{a}_j \rangle \cdot \mathbb{E} [\sigma(\langle (\mathbf{B}\mathbf{D})_{i,:}, \tilde{\mathbf{x}} \rangle) \cdot \sigma(\langle (\mathbf{B}\mathbf{D})_{j,:}, \tilde{\mathbf{x}} \rangle)], \end{aligned} \quad (35)$$

where in the fourth line we used that for $\tilde{\mathbf{x}} \sim \mathcal{N}(\mathbf{0}, \mathbf{I})$, as σ is odd, the following identity holds:

$$\mathbb{E} [\sigma(\langle (\mathbf{B}\mathbf{D})_{j,:}, \tilde{\mathbf{x}} \rangle)] = 0,$$

and the last inequality follows from the fact that for the Gram matrix $\mathbf{M}$ of the vectors $\{\mathbf{a}_i\}_{i=1}^{n'}$:

$$\sum_{i,j \leq n'} \langle \mathbf{a}_i, \mathbf{a}_j \rangle \cdot \sigma(0)^2 = \sigma(0)^2 \cdot \langle \mathbf{1}, \mathbf{M}\mathbf{1} \rangle \geq 0.$$

Similarly, one can verify that

$$\mathbb{E}\langle \tilde{\mathbf{x}}, \mathbf{D}\mathbf{A}\sigma(\mathbf{B}\mathbf{D}\tilde{\mathbf{x}}) \rangle = \sum_{i=1}^d \sum_{j=n'+1}^n a_j^i \cdot \mathbb{E} [(D_{i,i} \cdot \tilde{x}_i) \cdot \sigma(\langle (\mathbf{B}\mathbf{D})_{j,:}, \tilde{\mathbf{x}} \rangle)]. \quad (36)$$

Combining (35) and (36), and recalling the population risk form in (33), we conclude that

$$\mathcal{R}(\mathbf{A}, \mathbf{B}) \geq \mathcal{R}(\mathbf{A}_{:,n'+1:}, \mathbf{B}_{n'+1:,:}),$$

where $A_{:,n'+1:}$ and $B_{n'+1:,}$ are obtained by removing the zero columns/rows from A and B , respectively. This means that considering a matrix B with zero rows is equivalent to looking at a smaller rate $r' < r$. We show in Theorem 4.2, Proposition 4.3 and Theorem 5.2 that the population risk is monotone in the rate. Thus, having zero rows in B is clearly sub-optimal.

C. Proofs of Lower Bound on Loss (Section 4.1)

C.1. Case $r \leq 1$

C.1.1. LOWER BOUND ON $\tilde{R}(r)$

Lemma C.1. *Let $A = [a_1, \dots, a_n] \in \mathbb{R}^{d \times n}$ and $B^\top = [b_1, \dots, b_n] \in \mathbb{R}^{d \times n}$, with $\|b_i\|_2 = 1$ for $i \in [n]$. Let c_1 and $f(\cdot)$ be defined as per Lemma 4.1. Then, the following bound holds:*

$$\mathcal{L}_l(A, B) := \text{Tr} \left[A^\top A \cdot (BB^\top)^{\circ(2\ell+1)} \right] - \frac{2c_1}{f(1)} \cdot \text{Tr} [BA] \geq -\frac{c_1^2}{(f(1))^2} \cdot n. \quad (37)$$

Proof of Lemma C.1. For any symmetric $P, Q, T \in \mathbb{R}^{n \times n}$, a direct computation readily gives that

$$\text{Tr} [P \cdot (Q \circ T)] = \text{Tr} [(P \circ Q) \cdot T]. \quad (38)$$

Thus, by taking $P = A^\top A$, $Q = (BB^\top)^{\circ\ell}$ and $T = (BB^\top)^{\circ(\ell+1)}$, we obtain

$$\text{Tr} \left[A^\top A \cdot (BB^\top)^{\circ(2\ell+1)} \right] = \text{Tr} \left[(A^\top A \circ (BB^\top)^{\circ\ell}) \cdot (BB^\top \circ (BB^\top)^{\circ\ell}) \right].$$

Note that $BB^\top$ is PSD and, therefore, $(BB^\top)^{\circ\ell}$ is also PSD by Schur product theorem. Furthermore, as the rows of B have unit norm, $(BB^\top)^{\circ\ell}$ has unit diagonal. As a result, if we show that, for any PSD matrix Q with unit diagonal entries,

$$\text{Tr} \left[(A^\top A \circ Q) \cdot (BB^\top \circ Q) \right] - \frac{2c_1}{f(1)} \cdot \text{Tr} [BA] \geq -\frac{c_1^2}{(f(1))^2} \cdot n, \quad (39)$$

then the claim (37) immediately follows.

As Q is a PSD matrix with unit diagonal, it admits the following decomposition

$$Q = \sum_{i=1}^n u_i u_i^\top, \quad D_i = \text{Diag}(u_i), \quad \sum_{i=1}^n D_i^2 = I. \quad (40)$$

In this view, defining

$$A_i = A D_i, \quad B_i = D_i B,$$

we can rewrite the LHS of (39) in a more convenient form for further analysis. In particular, for the second term we deduce the following

$$\text{Tr} [BA] = \text{Tr} [AB] = \text{Tr} \left[A \cdot \left(\sum_{i=1}^n D_i^2 \right) \cdot B \right] = \sum_{i=1}^n \text{Tr} [A \cdot D_i^2 \cdot B] = \sum_{i=1}^n \text{Tr} [(A D_i) \cdot (D_i B)] = \sum_{i=1}^n \text{Tr} [A_i B_i].$$

Let us now rearrange the first term of (39). Notice that

$$(A^\top A \circ Q)_{i,j} = \sum_{k=1}^n \langle a_i, a_j \rangle \cdot u_k^i u_k^j = \sum_{k=1}^n \langle a_i \cdot u_k^i, a_j \cdot u_k^j \rangle = \sum_{k=1}^n ((A D_k)^\top \cdot (A D_k))_{i,j} = \sum_{k=1}^n (A_k^\top A_k)_{i,j}.$$

In the same fashion we get

$$(BB^\top \circ Q)_{i,j} = \sum_{k=1}^n (B_k B_k^\top)_{i,j},$$

from which we deduce that

$$\text{Tr} \left[(\mathbf{A}^\top \mathbf{A} \circ \mathbf{Q}) \cdot (\mathbf{B} \mathbf{B}^\top \circ \mathbf{Q}) \right] = \sum_{i,j=1}^n \text{Tr} \left[\mathbf{A}_i^\top \mathbf{A}_i \mathbf{B}_j \mathbf{B}_j^\top \right].$$

Therefore, the proof of (39) can be obtained by proving that, for *any* matrices $\mathbf{A}_1, \dots, \mathbf{A}_n \in \mathbb{R}^{d \times n}$ and $\mathbf{B}_1, \dots, \mathbf{B}_n \in \mathbb{R}^{n \times d}$,

$$\sum_{i,j=1}^n \text{Tr} \left[\mathbf{A}_i^\top \mathbf{A}_i \mathbf{B}_j \mathbf{B}_j^\top \right] - \frac{2c_1}{f(1)} \cdot \sum_{i=1}^n \text{Tr} [\mathbf{A}_i \mathbf{B}_i] + \frac{c_1^2}{(f(1))^2} \text{Tr} [\mathbf{I}] \geq 0. \quad (41)$$

To show the last claim, let us define the following matrices

$$\mathbf{X} = \sum_{i=1}^n \mathbf{A}_i^\top \mathbf{A}_i, \quad \mathbf{Y} = \sum_{i=1}^n \mathbf{B}_i \mathbf{B}_i^\top, \quad \mathbf{Z} = \sum_{i=1}^n \mathbf{B}_i \mathbf{A}_i,$$

which allows us to rewrite the statement of (41) as

$$\text{Tr} \left[\mathbf{X} \mathbf{Y} - \frac{2c_1}{f(1)} \cdot \mathbf{Z} + \frac{c_1^2}{(f(1))^2} \cdot \mathbf{I} \right] \geq 0. \quad (42)$$

Note that $\mathbf{X}$ is PSD, hence it has a symmetric square root, which we denote by $\sqrt{\mathbf{X}}$. Using the continuity of the quantities involved in the LHS of (42), we can assume without loss of generality that $\mathbf{X}$ is invertible. In fact, the following quantities are continuous: trace, matrix product, matrix transpose. In addition, we can always introduce a small perturbation to $\mathbf{A}_i$'s which makes $\mathbf{X}$ full-rank. Thus, it suffices to show that (42) holds for $\mathbf{A}_i$'s such that $\mathbf{X}$ is invertible.

In this view, for any matrix $\mathbf{T} \in \mathbb{R}^{n \times n}$, we have

$$\begin{aligned} 0 &\leq \sum_{i=1}^n \left\| \frac{c_1}{f(1)} \cdot \mathbf{T} \mathbf{A}_i^\top - \sqrt{\mathbf{X}} \mathbf{B}_i \right\|_F^2 = \sum_{i=1}^n \text{Tr} \left[\left(\frac{c_1}{f(1)} \cdot \mathbf{T} \mathbf{A}_i^\top - \sqrt{\mathbf{X}} \mathbf{B}_i \right) \cdot \left(\frac{c_1}{f(1)} \cdot \mathbf{A}_i \mathbf{T}^\top - \mathbf{B}_i^\top \sqrt{\mathbf{X}} \right) \right] \\ &= \sum_{i=1}^n \text{Tr} \left[\frac{c_1^2}{(f(1))^2} \cdot \mathbf{T} \mathbf{A}_i^\top \mathbf{A}_i \mathbf{T}^\top - \frac{2c_1}{f(1)} \sqrt{\mathbf{X}} \mathbf{B}_i \mathbf{A}_i \mathbf{T}^\top + \mathbf{X} \mathbf{B}_i \mathbf{B}_i^\top \right] \\ &= \text{Tr} \left[\frac{c_1^2}{(f(1))^2} \cdot \mathbf{T} \mathbf{X} \mathbf{T}^\top - \frac{2c_1}{f(1)} \sqrt{\mathbf{X}} \mathbf{Z} \mathbf{T}^\top + \mathbf{X} \mathbf{Y} \right], \end{aligned} \quad (43)$$

where in the second line we used that $\text{Tr} [\mathbf{M}] = \text{Tr} [\mathbf{M}^\top]$ for any $\mathbf{M}$, and $\text{Tr} [\mathbf{M} \mathbf{N}] = \text{Tr} [\mathbf{N} \mathbf{M}]$ for any $\mathbf{M}, \mathbf{N}$.

As $\mathbf{X}$ is invertible, its square root $\sqrt{\mathbf{X}}$ is invertible. As $\mathbf{X}$ is also PSD, its inverse, i.e., $\mathbf{X}^{-1}$, is PSD and, hence, it has a symmetric square root, i.e., $\sqrt{\mathbf{X}^{-1}}$. In this view, we get that

$$\sqrt{\mathbf{X}^{-1}} = (\sqrt{\mathbf{X}})^{-1}.$$

Thus, by picking $\mathbf{T} = (\sqrt{\mathbf{X}})^{-1}$, we obtain

$$\mathbf{T}^\top \mathbf{T} = \mathbf{T}^2 = \mathbf{X}^{-1}, \quad \mathbf{T}^\top \sqrt{\mathbf{X}} = \mathbf{T} \sqrt{\mathbf{X}} = \mathbf{I}.$$

Using these observations, we deduce that the RHS of (43) is equal to the LHS of (42), which concludes the proof. $\square$

C.1.2. MATRICES IN $\mathcal{H}_{n,d}$ ARE THE ONLY MINIMIZERS

Lemma C.2. *Let $\mathbf{A} \in \mathbb{R}^{d \times n}$ and $\mathbf{B}^\top = [\mathbf{b}_1, \dots, \mathbf{b}_n] \in \mathbb{R}^{d \times n}$, with $\|\mathbf{b}_i\|_2 = 1$ for $i \in [n]$. Let c_1 and $f(\cdot)$ be defined as per Lemma 4.1. Then, we have that the set of minimizers of*

$$\text{Tr} \left[\mathbf{A}^\top \mathbf{A} \cdot f(\mathbf{B} \mathbf{B}^\top) \right] - 2c_1 \cdot \text{Tr} [\mathbf{B} \mathbf{A}] \quad (44)$$

coincides with the set $\mathcal{H}_{n,d}$ of weight-tied orthogonal matrices.

Proof of Lemma C.2. A direct computation immediately shows that the lower bound (37) is achieved for all $\ell \in \mathbb{N}$ by matrices $(\mathbf{A}, \mathbf{B})$ that belong to the set $\mathcal{H}_{d,n}$. Define the sets of minimizers of (37) as follows

$$\mathcal{M}_\ell := \arg \min_{\mathbf{A}, \mathbf{B}: \|\mathbf{b}_i\|_2=1} \mathcal{L}_\ell(\mathbf{A}, \mathbf{B}) = \left\{ (\mathbf{A}_B, \mathbf{B}) : \mathbf{A}_B \in \arg \min_{\mathbf{A}} \mathcal{L}_\ell(\mathbf{A}, \mathbf{B}), \mathbf{B} \in \arg \min_{\mathbf{B}: \|\mathbf{b}_i\|_2=1} \mathcal{L}_\ell(\mathbf{A}_B, \mathbf{B}) \right\}.$$

We will now show that

$$\bigcap_{\ell=0}^{\infty} \mathcal{M}_\ell = \mathcal{H}_{n,d}. \quad (45)$$

As the Taylor coefficients of $f(\cdot)$ are non-negative, (45) readily gives that the set of minimizers of (44) coincides with $\mathcal{H}_{n,d}$. Further, recall that $c_1 \neq 0$ and $\sum_{\ell=1}^{\infty} (c_{2\ell+1})^2 \neq 0$ and, hence, (45) is the union of the linear term ($\ell = 0$) and at least one non-linear ($\ell > 0$) term.

We first prove that it is enough to consider the case $r = 1$. Thus, assume that the result holds for $n = d$ and consider now $n < d$. We have that, for any orthogonal matrix $\mathbf{O} \in \mathbb{R}^{d \times d}$,

$$\begin{aligned} \mathbb{E}_{\mathbf{x}} \|\mathbf{x} - \mathbf{A}\sigma(\mathbf{B}\mathbf{x})\|_2^2 &= \mathbb{E}_{\mathbf{x}} \|\mathbf{O}\mathbf{x} - \mathbf{A}\sigma(\mathbf{B}\mathbf{O}\mathbf{x})\|_2^2 \\ &= \mathbb{E}_{\mathbf{x}} \left\| \mathbf{x} - \mathbf{O}^\top \mathbf{A} \sigma(\mathbf{B}\mathbf{O}\mathbf{x}) \right\|_2^2, \end{aligned} \quad (46)$$

where in the first step we have used the rotational invariance of $\mathbf{x}$, and in the second step we have multiplied the argument of the norm by the orthogonal matrix $\mathbf{O}^\top$. Thus, (46) gives that $(\mathbf{A}, \mathbf{B}) \in \mathcal{H}_{n,d}$ if and only if $(\mathbf{O}^\top \mathbf{A}, \mathbf{B}\mathbf{O}) \in \mathcal{H}_{n,d}$.

Let us write the SVD of $\mathbf{B}$ as $\mathbf{U}\mathbf{D}\mathbf{V}^\top$, where $\mathbf{U} \in \mathbb{R}^{n \times n}$, $\mathbf{V} \in \mathbb{R}^{d \times d}$ are orthogonal matrices and $\mathbf{D} \in \mathbb{R}^{n \times d}$ is a (rectangular) diagonal matrix. Thus, by taking $\mathbf{O} = \mathbf{V}$, one can assume that $\mathbf{B}$ has the form $(\mathbf{B}_{1:n,1:n}, \mathbf{0}_{1:n,1:d-n})$, where $\mathbf{B}_{1:n,1:n}$ denotes the left $n \times n$ sub-matrix of $\mathbf{B}$ and $\mathbf{0}_{1:n,1:d-n}$ denotes a $n \times (d-n)$ matrix of 0's. We also write the decompositions $\mathbf{A} = ((\mathbf{A}_{1:n,1:n})^\top, (\mathbf{A}_{n+1:d,1:n})^\top)^\top$ and $\mathbf{x} = (\mathbf{x}_{1:n}, \mathbf{x}_{n+1:d})$, where $\mathbf{A}_{1:n,1:n}$ (resp. $\mathbf{A}_{n+1:d,1:n}$) denotes the top $n \times n$ (resp. bottom $(d-n) \times n$) sub-matrix of $\mathbf{A}$, and $\mathbf{x}_{1:n}$ (resp. $\mathbf{x}_{n+1:d}$) denotes the first n (resp. last $d-n$) components of $\mathbf{x}$. Hence, the objective (2) can be expressed (up to the constant multiplicative factor d^{-1}) as the sum of

$$\mathcal{R}_1(\mathbf{A}, \mathbf{B}) = \mathbb{E} \left[\|\mathbf{x}_{1:n} - \mathbf{A}_{1:n,1:n} \sigma(\mathbf{B}_{1:n,1:n} \mathbf{x}_{1:n})\|^2 \right]$$

and

$$\mathcal{R}_2(\mathbf{A}, \mathbf{B}) = \mathbb{E} \left[\|\mathbf{x}_{n+1:d} - \mathbf{A}_{n+1:d,1:n} \sigma(\mathbf{B}_{1:n,1:n} \mathbf{x}_{1:n})\|^2 \right].$$

As $\mathbf{x}_{n+1:d}$ has zero mean and it is independent from $\mathbf{x}_{1:n}$, we have that

$$\mathcal{R}_2(\mathbf{A}, \mathbf{B}) = d - n + \mathbb{E} \left[\|\mathbf{A}_{n+1:d,1:n} \sigma(\mathbf{B}_{1:n,1:n} \mathbf{x}_{1:n})\|^2 \right],$$

which is minimized by setting $\mathbf{A}_{n+1:d,1:n}$ to $\mathbf{0}$. Note that $\mathcal{R}_1$ depends only on $\mathbf{A}_{1:n,1:n}$, $\mathbf{B}_{1:n,1:n}$ (and not on $\mathbf{A}_{n+1:d,1:n}$), hence its minimizers are $(\mathbf{A}_{1:n,1:n}, \mathbf{B}_{1:n,1:n}) \in \mathcal{H}_{n,n}$ by our assumption on the $r = 1$ case. As a result, by using that $(\mathbf{A}, \mathbf{B}) \in \mathcal{H}_{d,n}$ if and only if $(\mathbf{O}^\top \mathbf{A}, \mathbf{B}\mathbf{O}) \in \mathcal{H}_{d,n}$, we conclude that all the minimizers of the desired objective have the form $\mathbf{O}((\mathbf{A}_{1:n,1:n})^\top, (\mathbf{0}_{1:n-d,1:n})^\top)^\top$ and $(\mathbf{B}_{1:n,1:n}, \mathbf{0}_{1:n,1:d-n})\mathbf{O}^\top$, i.e., they form the set $\mathcal{H}_{n,d}$ defined in (7).

It remains to prove the result for $r = 1$. First, consider $\ell = 0$. In this case, we have

$$\begin{aligned} \mathcal{L}_0(\mathbf{A}, \mathbf{B}) &= \text{Tr} \left[\mathbf{A}^\top \mathbf{A} \mathbf{B} \mathbf{B}^\top \right] - \frac{2c_1}{f(1)} \cdot \text{Tr} [\mathbf{B} \mathbf{A}] \\ &= \text{Tr} \left[\mathbf{B}^\top \mathbf{A}^\top \mathbf{A} \mathbf{B} \right] - \frac{2c_1}{f(1)} \cdot \text{Tr} [\mathbf{A} \mathbf{B}] \\ &= \|\mathbf{A} \mathbf{B}\|_F^2 - \frac{2c_1}{f(1)} \cdot \text{Tr} [\mathbf{A} \mathbf{B}], \end{aligned} \quad (47)$$

where we have used that the trace is invariant under cyclic permutation. Notice that the minimizer of (47) is clearly $\mathbf{A} \mathbf{B} = \frac{c_1}{f(1)} \mathbf{I}_d$.

Consider some $\ell \geq 1$. As $\mathbf{AB} = \frac{c_1}{f(1)} \mathbf{I}_d$ and $\mathbf{A}, \mathbf{B}$ are square matrices, $\mathbf{B}$ is invertible and $\mathbf{A}^\top \mathbf{A} = \frac{c_1^2}{(f(1))^2} \cdot (\mathbf{BB}^\top)^{-1}$. Thus,

$$\begin{aligned} \mathcal{L}_\ell(\mathbf{A}, \mathbf{B}) &= \text{Tr} \left[\mathbf{A}^\top \mathbf{A} (\mathbf{BB}^\top)^{\circ(2\ell+1)} \right] - \frac{2c_1}{f(1)} \cdot \text{Tr} [\mathbf{BA}] \\ &= \frac{c_1^2}{(f(1))^2} \cdot \text{Tr} \left[(\mathbf{BB}^\top)^{-1} (\mathbf{BB}^\top)^{\circ(2\ell+1)} \right] - \frac{2c_1^2}{(f(1))^2} \cdot n. \end{aligned} \quad (48)$$

Let $\mathbf{P} = \mathbf{BB}^\top$. Note that $\mathbf{P}$ is symmetric and, hence, also its inverse is symmetric. Then, by using (38), we have that

$$\text{Tr} \left[\mathbf{P}^{-1} \mathbf{P}^{\circ(2\ell+1)} \right] = \text{Tr} \left[(\mathbf{P}^{-1} \circ \mathbf{P}) \mathbf{P}^{\circ 2\ell} \right]. \quad (49)$$

An application of Theorem 5 in (Visick, 2000) gives that

$$\mathbf{P} \circ \mathbf{P}^{-1} \succeq \mathbf{I}, \quad (50)$$

where $\succeq$ denotes majorization in the PSD sense. We now show that $\mathbf{P} \circ \mathbf{P}^{-1} = \mathbf{I}$. To do so, suppose by contradiction that

$$\mathbf{P} \circ \mathbf{P}^{-1} = \mathbf{I} + \mathbf{R},$$

for some $\mathbf{R} \succeq \mathbf{0}$ such that $\mathbf{R} \neq \mathbf{0}$. Hence,

$$\text{Tr} \left[(\mathbf{P}^{-1} \circ \mathbf{P}) \mathbf{P}^{\circ 2\ell} \right] = \text{Tr} \left[\mathbf{P}^{\circ 2\ell} \right] + \text{Tr} \left[\mathbf{R} \mathbf{P}^{\circ 2\ell} \right] = n + \text{Tr} \left[\mathbf{R} \mathbf{P}^{\circ 2\ell} \right], \quad (51)$$

where in the last equality we use that $\mathbf{P}$ (and, consequently, $\mathbf{P}^{\circ 2\ell}$) has unit diagonal. By the Schur product theorem, $\mathbf{P}^{\circ 2\ell} \succ \mathbf{0}$ and, hence, it admits a square root. Thus, we get

$$\text{Tr} \left[\mathbf{R} \mathbf{P}^{\circ 2\ell} \right] = \text{Tr} \left[\sqrt{\mathbf{P}^{\circ 2\ell}} \cdot \mathbf{R} \cdot \sqrt{\mathbf{P}^{\circ 2\ell}} \right].$$

It is easy to see that the matrix $\sqrt{\mathbf{P}^{\circ 2\ell}} \cdot \mathbf{R} \cdot \sqrt{\mathbf{P}^{\circ 2\ell}}$ is PSD and, thus,

$$\text{Tr} \left[\sqrt{\mathbf{P}^{\circ 2\ell}} \cdot \mathbf{R} \cdot \sqrt{\mathbf{P}^{\circ 2\ell}} \right] \geq 0,$$

where the inequality is strict if and only if the corresponding matrix has only zero eigenvalues. However, for any non-zero $\mathbf{v} \in \mathbb{R}^n$, we have that

$$\mathbf{u}_\mathbf{v} := \sqrt{\mathbf{P}^{\circ 2\ell}} \cdot \mathbf{v} \neq \mathbf{0},$$

since $\sqrt{\mathbf{P}^{\circ 2\ell}}$ is strictly positive definite (as $\mathbf{P}^{\circ 2\ell} \succ \mathbf{0}$) and, thus, it does not have 0 eigenvalues. Hence, if

$$\mathbf{v}^\top \cdot \sqrt{\mathbf{P}^{\circ 2\ell}} \cdot \mathbf{R} \cdot \sqrt{\mathbf{P}^{\circ 2\ell}} \cdot \mathbf{v} = \mathbf{u}_\mathbf{v}^\top \mathbf{R} \mathbf{u}_\mathbf{v} = 0,$$

then $\mathbf{u}_\mathbf{v} \neq \mathbf{0}$ is an eigenvector of $\mathbf{R}$ corresponding to a zero eigenvalue. In this view, if $\sqrt{\mathbf{P}^{\circ 2\ell}} \cdot \mathbf{R} \cdot \sqrt{\mathbf{P}^{\circ 2\ell}}$ has all zero eigenvalues, then all eigenvalues of $\mathbf{R}$ are zero. As $\mathbf{R}$ cannot be the zero matrix, by using (51), we conclude that

$$\text{Tr} \left[(\mathbf{P}^{-1} \circ \mathbf{P}) \mathbf{P}^{\circ 2\ell} \right] > n. \quad (52)$$

By combining (48), (49) and (52), we have that $\mathcal{L}_\ell(\mathbf{A}, \mathbf{B}) > -c_1^2 n / (f(1))^2$, which contradicts with the fact that $(\mathbf{A}, \mathbf{B})$ is a minimizer (since any $(\mathbf{A}', \mathbf{B}') \in \mathcal{H}_{n,d}$ achieves the value of $-c_1^2 n / (f(1))^2$). Therefore, we conclude that $\mathbf{P} \circ \mathbf{P}^{-1} = \mathbf{I}$.

At this point, we show that $\mathbf{P} \circ \mathbf{P}^{-1} = \mathbf{I}$ implies that $\mathbf{P} = \mathbf{I}$. Note that $\mathbf{P}$ is a Gram matrix, and let its basis be $\{\mathbf{b}_1, \dots, \mathbf{b}_n\}$. Define

$$\mathbf{b}'_i = \mathbf{b}_i - \tilde{\mathbf{b}}_i,$$

where $\tilde{\mathbf{b}}_i$ is orthogonal projection of $\mathbf{b}_i$ onto the space spanned by $\{\mathbf{b}_j\}_{j \neq i}^n$. From a well-known result (see, for instance, Theorem 2.1 in (del Pino & Galaz, 1995)) we have that

$$\mathbf{P}_{ii}^{-1} = \frac{1}{\|\mathbf{b}'_i\|_2^2}. \quad (53)$$

Hence, we obtain that

$$\|\mathbf{b}'_i\|_2 \leq \|\mathbf{b}_i\|_2 = 1, \quad (54)$$

where the inequality is sharp only if $\mathbf{b}_i$ is orthogonal to all $\{\mathbf{b}_j\}_{j \neq i}^n$. Then, from (53), we deduce

$$n = \text{Tr}[\mathbf{I}] = \text{Tr}[\mathbf{P} \circ \mathbf{P}^{-1}] = \sum_{i=1}^n \|\mathbf{b}_i\|_2^2 \cdot \frac{1}{\|\mathbf{b}'_i\|_2^2} = \sum_{i=1}^n \frac{1}{\|\mathbf{b}'_i\|_2^2}. \quad (55)$$

By combining (54) and (55), we conclude that $\{\mathbf{b}_i\}_{i \in [n]}$ form an orthonormal basis, and, hence, $\mathbf{P} = \mathbf{I}$. This means that (45) holds for $r = 1$ since

$$(44) = \sum_{\ell=1}^{\infty} (c_{2\ell+1})^2 \cdot \mathcal{L}_{\ell}(\mathbf{A}, \mathbf{B}),$$

which concludes the proof. $\square$

Proof of Theorem 4.2. It follows by combining the results of Lemma C.1 and C.2. $\square$

C.2. Case $r > 1$

C.2.1. LOWER BOUND ON $\tilde{R}(r)$

Proof of Proposition 4.3. An application of Theorem A in (Khare, 2021) gives that

$$\text{Tr}[\mathbf{A}^{\top} \mathbf{A} \mathbf{B} \mathbf{B}^{\top}] = \langle \mathbf{1}, (\mathbf{A}^{\top} \mathbf{A} \circ \mathbf{B} \mathbf{B}^{\top}) \mathbf{1} \rangle \geq \frac{1}{d} \langle \mathbf{1}, (\text{Diag}(\mathbf{B} \mathbf{A}) \text{Diag}(\mathbf{B} \mathbf{A})^{\top}) \mathbf{1} \rangle = \frac{1}{d} (\text{Tr}[\mathbf{B} \mathbf{A}])^2,$$

where $\text{Diag}(\mathbf{B} \mathbf{A}) \in \mathbb{R}^n$ stands for the vector with entries corresponding to the diagonal of the matrix $\mathbf{B} \mathbf{A}$. Hence, we have

$$\text{Tr}[\mathbf{A}^{\top} \mathbf{A} \cdot f(\mathbf{B} \mathbf{B}^{\top})] - 2c_1 \cdot \text{Tr}[\mathbf{B} \mathbf{A}] \geq \frac{c_1^2}{d} (\text{Tr}[\mathbf{B} \mathbf{A}])^2 + \sum_{\ell=1}^{\infty} (c_{2\ell+1})^2 \cdot \text{Tr}[\mathbf{A}^{\top} \mathbf{A} \cdot (\mathbf{B} \mathbf{B}^{\top})^{\circ 2\ell+1}] - 2c_1 \cdot \text{Tr}[\mathbf{B} \mathbf{A}]. \quad (56)$$

Define $\alpha := f(1) - c_1^2$. Then, for any $\beta \in [0, 1]$, we can rewrite the RHS of (56) as

$$\left[\frac{c_1^2}{d} (\text{Tr}[\mathbf{B} \mathbf{A}])^2 - 2(1 - \beta)c_1 \cdot \text{Tr}[\mathbf{B} \mathbf{A}] \right] + \sum_{\ell=1}^{\infty} (c_{2\ell+1})^2 \cdot \left(\text{Tr}[\mathbf{A}^{\top} \mathbf{A} \cdot (\mathbf{B} \mathbf{B}^{\top})^{\circ 2\ell+1}] - \frac{2\beta c_1}{\alpha} \cdot \text{Tr}[\mathbf{B} \mathbf{A}] \right). \quad (57)$$

The first term in (57) is a quadratic polynomial in $\text{Tr}[\mathbf{B} \mathbf{A}]$. Hence, we have that

$$\left[\frac{c_1^2}{d} (\text{Tr}[\mathbf{B} \mathbf{A}])^2 - 2(1 - \beta)c_1 \cdot \text{Tr}[\mathbf{B} \mathbf{A}] \right] \geq -d(1 - \beta)^2. \quad (58)$$

Define $\mathbf{B}_e := [\mathbf{B}, \mathbf{0}_{1:n, 1:n-d}]$ and $\mathbf{A}_e^{\top} := [\mathbf{A}^{\top}, \mathbf{0}_{1:n, 1:n-d}]$. One can readily verify that the traces in the second term of (57) remain unchanged if we replace $\mathbf{A}$ and $\mathbf{B}$ with $\mathbf{A}_e$ and $\mathbf{B}_e$, respectively. Note that $\mathbf{A}_e, \mathbf{B}_e$ are square matrices, hence we can apply Lemma C.1 (which readily generalizes to a different scaling in front of the second trace) to get

$$\sum_{\ell=1}^{\infty} (c_{2\ell+1})^2 \cdot \left(\text{Tr}[\mathbf{A}^{\top} \mathbf{A} \cdot (\mathbf{B} \mathbf{B}^{\top})^{\circ 2\ell+1}] - \frac{2\beta c_1}{\alpha} \cdot \text{Tr}[\mathbf{B} \mathbf{A}] \right) \geq -\sum_{\ell=1}^{\infty} (c_{2\ell+1})^2 \cdot \frac{\beta^2 c_1^2}{\alpha^2} n = -\frac{\beta^2 c_1^2}{\alpha} n. \quad (59)$$

By combining (56), (57), (58) and (59), we obtain that

$$\frac{1}{d} \left(\text{Tr}[\mathbf{A}^{\top} \mathbf{A} \cdot f(\mathbf{B} \mathbf{B}^{\top})] - 2 \cdot \text{Tr}[\mathbf{A} \mathbf{B}] \right) + 1 \geq 1 - (1 - \beta)^2 - \frac{\beta^2 c_1^2}{\alpha} r. \quad (60)$$

By taking $\beta = \alpha/(c_1^2 r + \alpha)$ and re-arranging the RHS of (60), the desired result readily follows. $\square$

C.2.2. ASYMPTOTIC ACHIEVABILITY OF THE LOWER BOUND

Lemma C.3. Let $\mathbf{A}, \mathbf{B}$ be defined as in (12). Then, for any $\epsilon > 0$, we have that, with probability at least $1 - c/d^2$,

$$\left| \left(\text{Tr} \left[\mathbf{A}^\top \mathbf{A} f(\mathbf{B}\mathbf{B}^\top) \right] - 2c_1 \text{Tr} [\mathbf{A}\mathbf{B}] \right) - (\beta^2 c_1^2 r n + \beta^2 \alpha n - 2c_1 \beta n) \right| \leq C n^{\frac{1}{2} + \epsilon}.$$

Thus, choosing $\beta = \frac{c_1}{c_1^2 r + \alpha}$ the loss approaches $1 - \frac{r}{r + \frac{\alpha}{c_1^2}}$, i.e., with the same probability,

$$\left| \left(1 + \frac{1}{d} \left(\text{Tr} \left[\mathbf{A}^\top \mathbf{A} f(\mathbf{B}\mathbf{B}^\top) \right] - 2c_1 \text{Tr} [\mathbf{A}\mathbf{B}] \right) \right) - \left(1 - \frac{r}{r + \frac{\alpha}{c_1^2}} \right) \right| \leq C d^{-\frac{1}{2} + \epsilon}.$$

Here, the constants c, C depend only on r and ϵ .

We start by proving the following.

Lemma C.4. Let $\hat{\mathbf{B}}, \mathbf{B}$ be defined as in (12). Then, for any $\epsilon > 0$, we have that, with probability at least $1 - c/d^2$,

$$\max_{i,j} \left| \frac{(\mathbf{B}\mathbf{B}^\top)_{i,j}}{(\hat{\mathbf{B}}\hat{\mathbf{B}}^\top)_{i,j}} - 1 \right| \leq C n^{-\frac{1}{2} + \epsilon}.$$

Here, the constants c, C depend only on r and ϵ .

Proof. If $\mathbf{U} \in \mathbb{R}^{n \times n}$ is sampled uniformly from $\mathbb{SO}(n)$, then it follows from rotational invariance that any fixed row or column is uniformly distributed on the n -dimensional sphere $\mathbb{S}^{n-1}$. Thus, any fixed row of $\mathbf{U}$ is distributed as $\mathbf{g}/\|\mathbf{g}\|_2$, where $\mathbf{g} \sim \mathcal{N}(0, \mathbf{I}/n)$. Now, it follows from the concentration of $\|\mathbf{g}\|_2$ (see e.g. Theorem 3.1.1 in (Vershynin, 2018)) that $\|\|\mathbf{g}\|_2 - 1\|_{\psi_2} \leq C n^{-\frac{1}{2}}$, where $\|\cdot\|_{\psi_2}$ denotes the sub-Gaussian norm. Denote by $\mathbf{g}_d \in \mathbb{R}^d$ the first d components of $\mathbf{g}_d$. Then, by the same reasoning, it holds that $\|\sqrt{r}\|\mathbf{g}_d\|_2 - 1\|_{\psi_2} \leq c d^{-\frac{1}{2}}$. Looking at the definition of $\hat{\mathbf{B}}$, we have that, for any fixed i , the distribution of its rows is given by $\hat{\mathbf{b}}_i \sim \sqrt{r}\mathbf{g}_d/\|\mathbf{g}\|_2$. Furthermore, for any pair of indices i, j , we have that

$$\frac{(\mathbf{B}\mathbf{B}^\top)_{i,j}}{(\hat{\mathbf{B}}\hat{\mathbf{B}}^\top)_{i,j}} = \frac{1}{\|\hat{\mathbf{b}}_i\|_2 \cdot \|\hat{\mathbf{b}}_j\|_2}.$$

Hence,

$$\mathbb{P} \left(\left| \frac{(\mathbf{B}\mathbf{B}^\top)_{i,j}}{(\hat{\mathbf{B}}\hat{\mathbf{B}}^\top)_{i,j}} - 1 \right| \leq n^{-\frac{1}{2} + \epsilon} \right) = \mathbb{P} \left(\left| \frac{1}{\|\hat{\mathbf{b}}_i\|_2 \cdot \|\hat{\mathbf{b}}_j\|_2} - 1 \right| \leq n^{-\frac{1}{2} + \epsilon} \right) \leq C \exp \left(-\frac{d^\epsilon}{C} \right).$$

Now a simple union bound over all rows gives us

$$\mathbb{P} \left(\max_{i,j} \left| \frac{1}{\|\hat{\mathbf{b}}_i\|_2 \cdot \|\hat{\mathbf{b}}_j\|_2} - 1 \right| \leq n^{-\frac{1}{2} + \epsilon} \right) \leq C n \exp \left(-\frac{d^\epsilon}{C} \right) \leq \frac{C}{d^2},$$

which implies the desired result. $\square$

Next, we bound the traces of the terms $\mathbf{B}\mathbf{B}^\top (\mathbf{B}\mathbf{B}^\top)^{\circ(2\ell+1)}$. We start with the case $\ell = 0$.

Lemma C.5. Let $\mathbf{B}$ be defined as in (12). Then, for any $\epsilon > 0$, with probability at least $1 - c/d^2$,

$$\left| \text{Tr} \left[\mathbf{B}\mathbf{B}^\top (\mathbf{B}\mathbf{B}^\top) \right] - r n \right| \leq C d^{\frac{1}{2} + \epsilon}.$$

Here, the constants c, C depend only on r and ϵ .

Proof. Note that

$$\text{Tr} \left[\mathbf{B}\mathbf{B}^\top (\mathbf{B}\mathbf{B}^\top) \right] = \sum_{i,j} \left((\mathbf{B}\mathbf{B}^\top)_{i,j} \right)^2 = \sum_{i,j} \left(\frac{((\mathbf{B}\mathbf{B}^\top)_{i,j})^2}{((\hat{\mathbf{B}}\hat{\mathbf{B}}^\top)_{i,j})^2} - 1 \right) ((\hat{\mathbf{B}}\hat{\mathbf{B}}^\top)_{i,j})^2 + \text{Tr} \left[\hat{\mathbf{B}}\hat{\mathbf{B}}^\top (\hat{\mathbf{B}}\hat{\mathbf{B}}^\top) \right].$$

Thus, an application of Lemma C.4 gives that, with probability at least $1 - c/d^2$,

$$\left| \text{Tr} \left[\mathbf{B} \mathbf{B}^\top (\mathbf{B} \mathbf{B}^\top) \right] - \text{Tr} \left[\hat{\mathbf{B}} \hat{\mathbf{B}}^\top (\hat{\mathbf{B}} \hat{\mathbf{B}}^\top) \right] \right| \leq \text{Tr} \left[\hat{\mathbf{B}} \hat{\mathbf{B}}^\top (\hat{\mathbf{B}} \hat{\mathbf{B}}^\top) \right] \cdot C d^{-\frac{1}{2} + \epsilon}. \quad (61)$$

Since the trace is invariant under cyclic permutation, we readily have that

$$\text{Tr} \left[\hat{\mathbf{B}} \hat{\mathbf{B}}^\top (\hat{\mathbf{B}} \hat{\mathbf{B}}^\top) \right] = rn. \quad (62)$$

By combining (61) and (62), the desired result follows. $\square$

Finally, we consider the higher order terms for $\ell \geq 1$.

Lemma C.6. *Let $\mathbf{B}$ be defined as in (12). Then, for any $\epsilon > 0$, we have that, with probability at least $1 - c/d^2$,*

$$\sup_{\ell \geq 1} \left| \text{Tr} \left[\mathbf{B} \mathbf{B}^\top (\mathbf{B} \mathbf{B}^\top)^{\circ(2\ell+1)} \right] - n \right| \leq C \log^2 n.$$

Here, the constants c, C depend only on r and ϵ .

Proof. We first observe that

$$\text{Tr} \left[\mathbf{B} \mathbf{B}^\top (\mathbf{B} \mathbf{B}^\top)^{\circ(2\ell+1)} \right] = \sum_{i,j} \left((\mathbf{B} \mathbf{B}^\top)_{i,j} \right)^{2\ell+2} = n + \sum_{i \neq j} \left((\mathbf{B} \mathbf{B}^\top)_{i,j} \right)^{2\ell+2}.$$

An application of Lemma C.4 gives that, with probability $1 - c/d^2$,

$$\sup_{\ell \geq 1} \sum_{i \neq j} \left((\mathbf{B} \mathbf{B}^\top)_{i,j} \right)^{2\ell+2} \leq \sup_{\ell \geq 1} \sum_{i \neq j} \left((1 + C d^{-1/2 + \epsilon}) \cdot (\hat{\mathbf{B}} \hat{\mathbf{B}}^\top)_{i,j} \right)^{2\ell+2}. \quad (63)$$

Furthermore, by using the first part of Lemma G.2 with $\mathbf{A} = \hat{\mathbf{B}} \hat{\mathbf{B}}^\top$, we have that, with probability at least $1 - 1/n^2$, the RHS of (63) is lower bounded by

$$\sup_{\ell \geq 1} \sum_{i \neq j} \left((1 + C d^{-1/2 + \epsilon}) \cdot C \sqrt{\frac{\log n}{n}} \right)^{2\ell+2} \leq C \log^2 n,$$

which implies the desired result. $\square$

At this point, we are ready to give the proof of Lemma C.3.

Proof of Lemma C.3. Recall that $\{(c_{2\ell+1})^2\}_{\ell=0}^\infty$ denote the Taylor coefficients of $f(x)$. By using that $\mathbf{A} = \beta \mathbf{B}^\top$, our objective becomes

$$\begin{aligned} & \text{Tr} \left[\mathbf{A}^\top \mathbf{A} f(\mathbf{B} \mathbf{B}^\top) \right] - 2c_1 \text{Tr} [\mathbf{A} \mathbf{B}] = \beta^2 \text{Tr} \left[\mathbf{B} \mathbf{B}^\top f(\mathbf{B} \mathbf{B}^\top) \right] - 2c_1 \beta n \\ &= \beta^2 \sum_{\ell=0}^\infty (c_{2\ell+1})^2 \text{Tr} \left[\mathbf{B} \mathbf{B}^\top (\mathbf{B} \mathbf{B}^\top)^{\circ(2\ell+1)} \right] - 2c_1 \beta n \\ &= \beta^2 c_1^2 rn + \beta^2 \sum_{\ell=1}^\infty (c_{2\ell+1})^2 n - 2c_1 \beta n \\ & \quad + \beta^2 c_1^2 \left(\text{Tr} \left[\mathbf{B} \mathbf{B}^\top (\mathbf{B} \mathbf{B}^\top) \right] - rn \right) + \beta^2 \sum_{\ell=1}^\infty (c_{2\ell+1})^2 \left(\text{Tr} \left[\mathbf{B} \mathbf{B}^\top (\mathbf{B} \mathbf{B}^\top)^{\circ(2\ell+1)} \right] - n \right). \end{aligned}$$

Then, by bounding the last two terms with Lemma C.5 and Lemma C.6, the desired result follows. $\square$

Proof of Proposition 4.4. The proof is a direct application of Lemma C.3. $\square$

D. Global Convergence of Weight-tied Gradient Flow (Theorem 4.5)

We start by giving a formal recap of the weight-tied gradient flow considered in Section 4.2. Under the weight-tying constraint (14), the objective (13) has the following form

$$\begin{aligned}\Psi(\beta, \mathbf{B}) &:= \beta^2 \cdot \text{Tr} \left[\mathbf{B}^\top \mathbf{B} \cdot f(\mathbf{B}\mathbf{B}^\top) \right] - 2\beta n \\ &= \beta^2 \cdot \sum_{i,j=1}^n \langle \mathbf{b}_i, \mathbf{b}_j \rangle \cdot f(\langle \mathbf{b}_i, \mathbf{b}_j \rangle) - 2\beta n,\end{aligned}\tag{64}$$

where $\|\mathbf{b}_i\|_2 = 1$ for all i . Note that the optimal β^* can be found exactly, since (64) is a quadratic polynomial in β . In this view, to optimize (64), we perform a gradient flow on $\{\mathbf{b}_i\}_{i=1}^n$, which are regarded as vectors on the unit sphere, and pick the optimal β^* at each time t . Formally,

$$\begin{aligned}\beta(t) &= \frac{n}{\sum_{i,j=1}^n \langle \mathbf{b}_i, \mathbf{b}_j \rangle \cdot f(\langle \mathbf{b}_i, \mathbf{b}_j \rangle)}, \\ \frac{\partial \mathbf{b}_i(t)}{\partial t} &= -\mathbf{J}_i(t) \nabla_{\mathbf{b}_i} \Psi(\beta(t), \mathbf{B}(t)),\end{aligned}\tag{65}$$

where $\mathbf{J}_i(t) := \mathbf{I} - \mathbf{b}_i(t)\mathbf{b}_i(t)^\top$ projects the gradient $\nabla_{\mathbf{b}_i} \Psi(\beta(t), \mathbf{B}(t))$ onto the tangent space at the point $\mathbf{b}_i(t)$ (see (69) for the closed form expression). This ensures that $\|\mathbf{b}_i(t)\|_2 = 1$ along the gradient flow trajectory. The described procedure can be viewed as Riemannian gradient flow, due to the projection of the gradient $\nabla_{\mathbf{b}_i} \Psi(\beta(t), \mathbf{B}(t))$ on the tangent space of the unit sphere. We now present the formal counterpart of Theorem 4.5.

Theorem D.1. *Fix $r \leq 1$. Let $\mathbf{B}(t)$ be obtained via the gradient flow (65) applied to Ψ defined in (64). Let the initialization $\mathbf{B}(0)$ have unit-norm rows and $\text{rank}(\mathbf{B}(0)) = n$. Then, as $t \rightarrow \infty$, $\mathbf{B}(t)\mathbf{B}(t)^\top$ converges to $\mathbf{I}$, which is the unique global optimum of (64). Moreover, define the residual*

$$\phi(t) = \text{Tr} \left[(\mathbf{B}(t)\mathbf{B}(t)^\top - \mathbf{I}) \cdot f(\mathbf{B}(t)^\top \mathbf{B}(t)) \right] \geq 0,$$

which vanishes at the minimizer, and let T be the first time such that $\phi(T) = \delta$. Then,

$$T \leq -\mathbb{1}\{\phi(0) > nf(1)\} \cdot f(1) \cdot \log \det(\mathbf{B}(0)\mathbf{B}(0)^\top) - \mathbb{1}\{\delta \leq nf(1)\} \cdot \frac{2f^2(1)}{\delta} \cdot \log \det(\mathbf{B}(0)\mathbf{B}(0)^\top).\tag{66}$$

In words, if the residual at initialization is bigger than $nf(1)$, then it takes at most constant time to reach the regime in which the convergence is linear in the precision δ . We also note that by choosing the optimal β^* , the function ϕ can be related to the objective (64) by $\Psi(\beta^*, \mathbf{B}(t)) = -\frac{n}{f(1) + \frac{\phi(t)}{n}}$. Hence, (66) gives a quantitative convergence in terms of the objective function as well.

We now are ready to present the proof of Theorem D.1. Let $\mathbf{B}^\top = [\mathbf{b}_1, \dots, \mathbf{b}_n]$. Recall that, under the weight-tying (14), the objective in (13) can be re-written as

$$\beta^2 \cdot \sum_{i,j=1}^n \langle \mathbf{b}_i, \mathbf{b}_j \rangle \cdot f(\langle \mathbf{b}_i, \mathbf{b}_j \rangle) - 2\beta n.\tag{67}$$

By the definition in Theorem D.1, the residual $\phi(t)$ is given by

$$\phi(t) := \sum_{i \neq j}^n \langle \mathbf{b}_i, \mathbf{b}_j \rangle \cdot f(\langle \mathbf{b}_i, \mathbf{b}_j \rangle).\tag{68}$$

In this view, in accordance with (65), we study the following gradient flow:

$$\begin{cases} \frac{\partial \mathbf{b}_k(t)}{\partial t} = -\beta^2(t) \cdot \left[\mathbf{J}_k(t) \sum_{i \neq j} \mathbf{b}_j(t) \cdot g(\langle \mathbf{b}_k(t), \mathbf{b}_j(t) \rangle) \right], \\ \beta(t) = \frac{n}{nf(1) + \phi(t)}, \\ \|\mathbf{b}_k(0)\|_2 = 1, \end{cases}\tag{69}$$

where $g(x) := x \cdot f'(x) + f(x)$, and we have rescaled the time of the dynamics by a factor 2 to omit the factor 2 in front of $\beta^2(t)$. From here on, we will suppress the time notation when it is clear from the context, for the sake of simplicity. Note that one of the terms is absent in the summation, due to the fact that by definition of the operator $\mathbf{J}_k$:

$$\mathbf{J}_k \mathbf{b}_k = \mathbf{0}.$$

In addition, since $\mathbf{J}_k$ defines the projection of the gradient on the tangent space at the point $\mathbf{b}_k$ of the unit sphere, along the trajectory of the gradient flow (69) we have that $\|\mathbf{b}_k\|_2 = 1$.

The gradient flow (69) is well-defined (i.e., its solution exists and it is unique) when its RHS is Lipschitz continuous (see, for instance, (Santambrogio, 2017)). It suffices to check the Lipschitz continuity of $g(\cdot)$. Note that both $xf'(x)$ and $f(x)$ are Lipschitz continuous on any interval $[-1 + \delta, 1 - \delta]$ for some $\delta > 0$. Hence, the RHS of (69) is Lipschitz continuous, if

$$\max_{i \neq j} |\langle \mathbf{b}_i, \mathbf{b}_j \rangle| \leq 1 - \delta, \quad (70)$$

where δ is bounded away from 0 uniformly in t .

Recall that, by the assumption of Theorem D.1, we have that $\text{rank}(\mathbf{B}(0)\mathbf{B}(0)^\top) = n$, hence $\det(\mathbf{B}(0)\mathbf{B}(0)^\top) \geq \varepsilon_1$ for some $\varepsilon_1 > 0$. Thus, from the result in Lemma D.3, we obtain that

$$\det(\mathbf{B}(t)\mathbf{B}(t)^\top) \geq \varepsilon_1. \quad (71)$$

Let $0 < \lambda_1 < \lambda_2 < \dots < \lambda_n$ denote the eigenvalues of $\mathbf{B}(t)\mathbf{B}(t)^\top$ in increasing order. Then, (71) directly gives that

$$\lambda_1 \prod_{i=2}^n \lambda_i \geq \varepsilon_1 > 0.$$

Since $\mathbf{B}(t)\mathbf{B}(t)^\top$ has unit diagonal, we have that $\sum_{i=1}^n \lambda_i = n$. Hence, the smallest possible value of λ_1 during the gradient flow dynamics can be inferred from

$$\lambda_1 \geq \frac{\varepsilon_1}{\prod_{i=2}^n \lambda_i},$$

by picking the largest possible $\prod_{i=2}^n \lambda_i$ given the constraint $\sum_{i=2}^n \lambda_i \leq n$. This is achieved by taking

$$\lambda_i = \frac{n}{n-1}, \quad \forall i \in \{2, \dots, n\},$$

which gives

$$\prod_{i=2}^n \lambda_i = \left(\frac{n}{n-1} \right)^{n-1} = \left(1 + \frac{1}{n-1} \right)^{n-1} \leq C,$$

where C is a universal constant, since the RHS converges from below to Euler's number as n increases. This proves that λ_1 is bounded away from zero uniformly in t . As a result, we can readily conclude that (70) holds. To see this last claim, consider a vector $\mathbf{v}$ which has 1 on position i and $-\text{sign}\langle \mathbf{b}_i, \mathbf{b}_j \rangle$ on position j . Hence, we have that

$$2\lambda_1 = \lambda_1 \cdot \|\mathbf{v}\|_2^2 \leq \mathbf{v}^\top (\mathbf{B}(t)\mathbf{B}(t)^\top) \mathbf{v} = 2 - 2 \cdot |\langle \mathbf{b}_i, \mathbf{b}_j \rangle| \Rightarrow |\langle \mathbf{b}_i, \mathbf{b}_j \rangle| \leq 1 - \lambda_1.$$

Notice that

$$\phi(t) \leq (n^2 - n)f(1),$$

since $xf(x) \leq f(1)$ for $|x| \leq 1$. Hence, we have that $\beta(t) \geq \frac{1}{nf(1)} > 0$. In this view, along the trajectory of the gradient flow (69), the quantity $\phi(t)$ is *strictly* decreasing until convergence, by the property of gradient flow.

Lemma D.2 (Characterization of stationary points). *Consider the gradient flow (69). Then, the following holds:*

- (A) Any orthogonal set of $\mathbf{b}_i$ is a stationary point and a global minimizer.
- (B) The gradient flow (69) never escapes any subspace spanned by a set of linearly dependent $\mathbf{b}_i$. However, for each such subspace there exists a direction in which (67) can be improved.

Proof of Lemma D.2. Recall that $\beta(t) > 0$ and $\{\mathbf{b}_i\}_{i=1}^n$. Then, the stationary point condition can be expressed as

$$\mathbf{J}_k \sum_{j \neq k} \mathbf{b}_j \cdot g(\langle \mathbf{b}_k, \mathbf{b}_j \rangle) = 0, \quad \forall k \in [n]. \quad (72)$$

Thus, any orthogonal set of vectors is clearly a stationary point by definition of $g(\cdot)$. Moreover, (67) is minimized iff $\mathbf{B}\mathbf{B}^\top = \mathbf{I}$ as $xf(x)$ is an even function since $f(\cdot)$ is odd.

Note that the kernel of the operator $\mathbf{J}_k$ is spanned by the vector $\mathbf{b}_k$. Thus, the condition (72) is equivalent to

$$\sum_{j \neq k} \mathbf{b}_j \cdot g(\langle \mathbf{b}_k, \mathbf{b}_j \rangle) = \gamma_k \cdot \mathbf{b}_k,$$

for some $\gamma_k \in \mathbb{R}$. One can readily verify that $g(x) = 0$ if and only if $x = 0$. Thus, either (i) $\mathbf{b}_k$ is orthogonal to $\mathbf{b}_j$ for all $j \neq k$ and $\gamma_k = 0$, or (ii) $\mathbf{b}_k$ lies in the span of $\{\mathbf{b}_j\}_{j \neq k}$. If condition (i) holds for all $k \in [n]$, then $\{\mathbf{b}_i\}_{i=1}^n$ form an orthogonal set of vectors and we fall in category (A). If condition (ii) holds for some $k \in [n]$, then we fall in category (B).

Now, let us show that, if $\{\mathbf{b}_i\}_{i=1}^n$ spans a sub-space of dimension smaller than n , then there is a direction along which the value of (67) can be improved. Since the $\{\mathbf{b}_i\}_{i=1}^n$ are linearly dependent, there exists $\mathbf{u}$ of unit norm such that

$$\langle \mathbf{u}, \mathbf{b}_j \rangle = 0, \quad \forall j \in [n]. \quad (73)$$

For some $k \in [n]$, consider the perturbation

$$\hat{\mathbf{b}}_k = \frac{1}{\sqrt{1+\lambda^2}} \cdot (\mathbf{b}_k + \lambda \cdot \mathbf{u}),$$

which has unit norm as $\langle \mathbf{b}_k, \mathbf{u} \rangle = 0$. Recall that (67) can be expressed as

$$\beta^2 \left(2 \cdot \sum_{j \neq k} \langle \hat{\mathbf{b}}_k, \mathbf{b}_j \rangle f(\langle \hat{\mathbf{b}}_k, \mathbf{b}_j \rangle) + \frac{\pi}{2} + \sum_{i, j \neq k} \langle \mathbf{b}_i, \mathbf{b}_j \rangle f(\langle \mathbf{b}_i, \mathbf{b}_j \rangle) \right) - 2\beta n. \quad (74)$$

Here, β is chosen to be the minimizer of the quantity (74) having fixed $\{\mathbf{b}_j\}_{j \neq k}$ and $\hat{\mathbf{b}}_k$. Thus, in order to prove that the population risk gets smaller by replacing $\mathbf{b}_k$ with $\hat{\mathbf{b}}_k$ for any $\lambda > 0$, it suffices to show that the following quantity

$$\sum_{j \neq k} \langle \hat{\mathbf{b}}_k, \mathbf{b}_j \rangle f(\langle \hat{\mathbf{b}}_k, \mathbf{b}_j \rangle), \quad (75)$$

is decreasing with λ . This last claim follows from the chain of inequalities below:

$$(75) = \frac{1}{\sqrt{1+\lambda^2}} \sum_{j \neq k} \langle \mathbf{b}_k, \mathbf{b}_j \rangle \cdot f\left(\frac{1}{\sqrt{1+\lambda^2}} \langle \mathbf{b}_k, \mathbf{b}_j \rangle\right) \quad (76)$$

$$= \frac{1}{\sqrt{1+\lambda^2}} \sum_{j \neq k} \langle \mathbf{b}_k, \mathbf{b}_j \rangle \cdot \sum_{\ell=0}^{\infty} \left(\frac{c_{2\ell+1}}{c_1}\right)^2 \cdot \left(\frac{1}{\sqrt{1+\lambda^2}}\right)^{2\ell+1} \cdot \langle \mathbf{b}_k, \mathbf{b}_j \rangle^{2\ell+1} \quad (77)$$

$$\leq \left(\frac{1}{\sqrt{1+\lambda^2}}\right)^2 \sum_{j \neq k} \langle \mathbf{b}_k, \mathbf{b}_j \rangle \cdot \sum_{\ell=0}^{\infty} \left(\frac{c_{2\ell+1}}{c_1}\right)^2 \cdot \langle \mathbf{b}_k, \mathbf{b}_j \rangle^{2\ell+1} \quad (78)$$

$$= \frac{1}{1+\lambda^2} \sum_{j \neq k} \langle \mathbf{b}_k, \mathbf{b}_j \rangle \cdot f(\langle \mathbf{b}_k, \mathbf{b}_j \rangle) < \sum_{j \neq k} \langle \mathbf{b}_k, \mathbf{b}_j \rangle \cdot f(\langle \mathbf{b}_k, \mathbf{b}_j \rangle), \quad (79)$$

where in the second line we substitute the Taylor expansion of $f(\cdot)$, the inequality in the third line uses that the coefficients $\{c_{2\ell+1}^2\}_{\ell=0}^{\infty}$ are all non-negative, and the last inequality follows from the fact that $\lambda > 0$.

Finally, we show that the gradient flow (69) does not escape the degenerate sub-space. If $\dim(\text{span}(\{\mathbf{b}_i\}_{i=1}^n)) < n$, then there exists $\mathbf{u} \in \mathbb{R}^d$ such that (73) holds. By projecting the gradient expression (72) onto $\mathbf{u}$, we have

$$\left\langle \mathbf{u}, \mathbf{J}_k \sum_{j \neq k} \mathbf{b}_j \cdot g(\langle \mathbf{b}_k, \mathbf{b}_j \rangle) \right\rangle = 0.$$

Hence, for any $k \in [n]$, the directional derivative of $\mathbf{b}_k$ in the direction of $\mathbf{u}$ is equal to zero, and the gradient flow does not escape the low-rank sub-space, which concludes the proof. $\square$

In next lemma we show that, if at initialization $\{\mathbf{b}_i\}_{i=1}^n$ spans a sub-space of dimension n , then it will never get stuck in a low-rank sub-space.

Lemma D.3 (Linearly independent $\{\mathbf{b}_i\}_{i=1}^n$ stay linearly independent). *Consider the gradient flow (69) with full rank initialization, i.e., $\text{rank}(\mathbf{B}(0)\mathbf{B}(0)^\top) = n$. Then, the following holds*

$$\frac{\partial}{\partial t} \log \det(\mathbf{B}(t)\mathbf{B}(t)^\top) \geq 2\beta(t)^2 \cdot \phi(t) \geq 0,$$

where $\mathbf{B}(t)^\top = [\mathbf{b}_1(t), \dots, \mathbf{b}_n(t)]$ and $\phi(t)$ is defined in (68). In particular, this implies that $\{\mathbf{b}_i\}_{i=1}^n$ stay full-rank along the gradient flow trajectory.

Proof of Lemma D.3. Applying the chain rule and using that the time derivative of $\mathbf{B}$ is given by the gradient flow (69) implies that

$$\frac{\partial}{\partial t} \log \det(\mathbf{B}\mathbf{B}^\top) = \text{Tr} \left[(\mathbf{B}\mathbf{B}^\top)^{-1} \cdot \left(\frac{\partial \mathbf{B}}{\partial t} \cdot \mathbf{B}^\top + \mathbf{B} \cdot \frac{\partial \mathbf{B}^\top}{\partial t} \right) \right],$$

where

$$\frac{\partial \mathbf{b}_k}{\partial t} = -\beta(t)^2 \cdot \left(\mathbf{J}_k \sum_{j \neq k} \mathbf{b}_j \cdot g(\langle \mathbf{b}_k, \mathbf{b}_j \rangle) \right).$$

Let us compute the quantity

$$\left\langle \frac{\partial \mathbf{b}_k}{\partial t}, \mathbf{b}_\ell \right\rangle = \left(\frac{\partial \mathbf{B}}{\partial t} \cdot \mathbf{B}^\top \right)_{k,\ell}.$$

By definition of $\mathbf{J}_k$, we have that

$$\mathbf{J}_k \sum_{j \neq k} \mathbf{b}_j \cdot g(\langle \mathbf{b}_k, \mathbf{b}_j \rangle) = \sum_{j \neq k} \mathbf{b}_j \cdot g(\langle \mathbf{b}_k, \mathbf{b}_j \rangle) - \sum_{j \neq k} \mathbf{b}_k \cdot \langle \mathbf{b}_k, \mathbf{b}_j \rangle \cdot g(\langle \mathbf{b}_k, \mathbf{b}_j \rangle).$$

Note that

$$\left\langle \sum_{j \neq k} \mathbf{b}_k \cdot \langle \mathbf{b}_k, \mathbf{b}_j \rangle \cdot g(\langle \mathbf{b}_k, \mathbf{b}_j \rangle), \mathbf{b}_\ell \right\rangle = \left[\text{Diag} \left[\mathbf{1}^\top ((\mathbf{B}\mathbf{B}^\top - \mathbf{I}) \circ g(\mathbf{B}\mathbf{B}^\top)) \right] \cdot \mathbf{B}\mathbf{B}^\top \right]_{k,\ell},$$

and that

$$\left\langle \sum_{j \neq k} \mathbf{b}_j \cdot g(\langle \mathbf{b}_k, \mathbf{b}_j \rangle), \mathbf{b}_\ell \right\rangle = \left[g(\mathbf{B}\mathbf{B}^\top) \cdot \mathbf{B}\mathbf{B}^\top \right]_{k,\ell} - g(1) \cdot [\mathbf{B}\mathbf{B}^\top]_{k,\ell}.$$

By combining these last four equations, we conclude that

$$\frac{\partial \mathbf{B}}{\partial t} \cdot \mathbf{B}^\top = -\beta(t)^2 \left(g(\mathbf{B}\mathbf{B}^\top) \cdot \mathbf{B}\mathbf{B}^\top - g(1) \cdot \mathbf{B}\mathbf{B}^\top - \text{Diag} \left[\mathbf{1}^\top ((\mathbf{B}\mathbf{B}^\top - \mathbf{I}) \circ g(\mathbf{B}\mathbf{B}^\top)) \right] \cdot \mathbf{B}\mathbf{B}^\top \right).$$

Furthermore,

$$\mathbf{B} \cdot \frac{\partial \mathbf{B}^\top}{\partial t} = \left(\frac{\partial \mathbf{B}}{\partial t} \cdot \mathbf{B}^\top \right)^\top = -\beta(t)^2 \left(\mathbf{B}\mathbf{B}^\top \cdot g(\mathbf{B}\mathbf{B}^\top) - g(1) \cdot \mathbf{B}\mathbf{B}^\top - \mathbf{B}\mathbf{B}^\top \cdot \text{Diag} \left[\mathbf{1}^\top ((\mathbf{B}\mathbf{B}^\top - \mathbf{I}) \circ g(\mathbf{B}\mathbf{B}^\top)) \right] \right).$$

Hence, by using the cyclic property of the trace, we get that

$$\begin{aligned} \frac{\partial}{\partial t} \log \det(\mathbf{B}\mathbf{B}^\top) &= 2\beta(t)^2 \cdot \text{Tr} \left[\text{Diag} \left[\mathbf{1}^\top ((\mathbf{B}\mathbf{B}^\top - \mathbf{I}) \circ g(\mathbf{B}\mathbf{B}^\top)) \right] \right] - 2\beta(t)^2 \cdot \text{Tr} \left[g(\mathbf{B}\mathbf{B}^\top) - g(1) \cdot \mathbf{I} \right] \\ &= 2\beta(t)^2 \cdot \sum_{i \neq j}^n \langle \mathbf{b}_i, \mathbf{b}_j \rangle \cdot g(\langle \mathbf{b}_i, \mathbf{b}_j \rangle) + 0, \end{aligned}$$

Now, note that

$$xg(x) = x^2 f'(x) + xf(x) \geq 0,$$

since $x^2 f'(x)$ and $xf(x)$ are non-negative functions, which concludes the proof. $\square$

The result of Lemma D.3 gives that $\det(\mathbf{B}\mathbf{B}^\top)$ is non-decreasing. Hence, if $\lambda_{\min}(\mathbf{B}\mathbf{B}^\top) > \delta > 0$ at initialization, then this quantity will be bounded away from zero during the gradient flow dynamics and the gradient flow will not get stuck in a low-rank solution. Therefore, by Lemma D.2, the gradient flow converges to a global minimum, in which the rows of $\mathbf{B}$ are orthogonal vectors with unit norm. The speed at which this happens is characterized by the next lemma.

Lemma D.4 (Rate of convergence). *Consider the gradient flow (69) with full rank initialization, i.e., $\text{rank}(\mathbf{B}(0)\mathbf{B}(0)^\top) = n$. Let T be the time at which $\phi(T)$ hits the value $\delta > 0$. Then, the following holds*

$$T \leq -\det(\mathbf{B}(0)\mathbf{B}(0)^\top) \cdot \left(f(1) \cdot \mathbb{1}\{\phi(0) > n \cdot f(1)\} + \frac{2f^2(1)}{\delta} \cdot \mathbb{1}\{\delta \leq n \cdot f(1)\} \right). \quad (80)$$

Proof of Lemma D.4. For all t , we have that $\text{Tr}[\mathbf{B}(t)\mathbf{B}(t)^\top] = n$, which implies that $\det(\mathbf{B}(t)\mathbf{B}(t)^\top) \leq 1$ and, as a consequence, that $\log \det(\mathbf{B}(t)\mathbf{B}(t)^\top) \leq 0$. From Lemma D.3, we know that

$$\frac{\partial}{\partial t} \log \det(\mathbf{B}(t)\mathbf{B}(t)^\top) \geq 2\beta(t)^2 \cdot \phi(t).$$

In this view, using the exact expression (69) for $\beta(t)$, we get

$$-\log \det(\mathbf{B}(0)\mathbf{B}(0)^\top) \geq \log \det(\mathbf{B}(t)\mathbf{B}(t)^\top) - \log \det(\mathbf{B}(0)\mathbf{B}(0)^\top) \geq \int_0^t \frac{2}{\left(f(1) + \frac{\phi(s)}{n}\right)^2} \cdot \phi(s) ds. \quad (81)$$

Stage 1. Assume that $\phi(0) > n \cdot f(1)$, and let T_1 be such that $\phi(T_1) = n \cdot f(1)$. Recall that the function $\phi(t)$ is decreasing and note that $x/(1+x)^2$ is decreasing for $x \in [1, +\infty)$. In this view, we can lower bound the integrand in the RHS of (81) for all $t \leq T_1$ by

$$\frac{2 \cdot \phi(0)}{\left(f(1) + \frac{\phi(0)}{n}\right)^2} \geq \frac{2(n-1)}{nf(1)} \geq \frac{1}{f(1)}, \quad (82)$$

where the first inequality follows from the definition (68) of $\phi(\cdot)$, which readily implies that $\phi(0) \leq f(1) \cdot n(n-1)$. Hence, by combining (81) with the lower bound (82), we get

$$T_1 \leq -f(1) \cdot \log \det(\mathbf{B}(0)\mathbf{B}(0)^\top).$$

Stage 2. Assume that $\phi(0) \leq n \cdot f(1)$. Let $\delta \in (0, n \cdot f(1)]$ be the desired precision which should be reached during the gradient flow, and let T_2 be such that $\phi(T_2) = \delta$. As $\phi(t)$ is decreasing, we have that

$$\frac{1}{\left(f(1) + \frac{\phi(t)}{n}\right)^2} \geq \frac{1}{\left(f(1) + \frac{\phi(0)}{n}\right)^2} \geq \frac{1}{4f^2(1)}, \quad (83)$$

where in the last step we use that $\phi(0) \leq n \cdot f(1)$. Hence, by combining (81) with the lower bound (83), we get

$$-\log \det(\mathbf{B}(0)\mathbf{B}(0)^\top) \geq \frac{1}{2f^2(1)} \cdot T_2 \delta,$$

which implies that

$$T_2 \leq -\frac{2f^2(1) \cdot \log \det(\mathbf{B}(0)\mathbf{B}(0)^\top)}{\delta}.$$

By combining the results of both stages, the desired result (80) readily follows. $\square$

Proof of Theorem D.1. Theorem D.1 is a compilation of the results presented in current section. $\square$

E. Global Convergence of Projected Gradient Descent (Theorem 4.6)

Recall from statement of Theorem 4.6 that

$$f(x) = x + \sum_{\ell=3}^{\infty} c_{\ell}^2 x^{\ell},$$

with $\sum_{\ell=3}^{\infty} c_{\ell}^2 < \infty$. We also define $\alpha = \sum_{\ell=3}^{\infty} c_{\ell}^2$, and we assume that $\alpha > 0$. In fact, if $\alpha = 0$, then the algorithm trivially converges after one step. We denote by C, c uniform positive constants (depending only on r and α) the value of which might change from term to term. To make the notation lighter we will also but the time t as a subscript (for example $\mathbf{B}(t)$ becomes $\mathbf{B}_t$).

We analyze the following projected gradient descent procedure for minimizing the population risk

$$\sum_{i,j=1}^n \langle \mathbf{a}_i, \mathbf{a}_j \rangle \cdot f \left(\left\langle \frac{\mathbf{b}_i}{\|\mathbf{b}_i\|_2}, \frac{\mathbf{b}_j}{\|\mathbf{b}_j\|_2} \right\rangle \right) - 2 \sum_{i=1}^n \left\langle \mathbf{a}_i, \frac{\mathbf{b}_i}{\|\mathbf{b}_i\|_2} \right\rangle. \quad (84)$$

Given unit-norm initial $\{\mathbf{b}_i\}_{i \in [n]}$, at each step we pick the optimal value of $\mathbf{A}$ given $\mathbf{B}$

$$\mathbf{A}_t = \mathbf{B}_t^{\top} \left(f(\mathbf{B}_t \mathbf{B}_t^{\top}) \right)^{-1}. \quad (85)$$

Then, we update $\mathbf{B}_t$ with a gradient step and a projection on the sphere to keep the unit norm:

$$\mathbf{B}'_t := \mathbf{B}_t - \eta \nabla_{\mathbf{B}_t}, \quad \mathbf{B}_{t+1} := \text{proj}(\mathbf{B}'_t).$$

Here, the operator $\text{proj}(\mathbf{M})$ normalizes the rows of $\mathbf{M}$ to be of unit norm and each row of $\nabla_{\mathbf{B}_t}$ is defined as the corresponding row of the gradient of $\mathbf{B}_t$, i.e.,

$$(\nabla_{\mathbf{B}_t})_{k,:} = \underbrace{-2\mathbf{J}_k \mathbf{a}_k + 2 \sum_{j \neq k} \langle \mathbf{a}_k, \mathbf{a}_j \rangle \mathbf{J}_k \mathbf{b}_j}_{:= \nabla_{\mathbf{B}_t}^1 \text{ (part 1)}} + \underbrace{\sum_{l=3}^{\infty} \ell c_{\ell}^2 \sum_{j \neq k} \langle \mathbf{a}_k, \mathbf{a}_j \rangle \langle \mathbf{b}_k, \mathbf{b}_j \rangle^{l-1} \mathbf{J}_k \mathbf{b}_j}_{:= \nabla_{\mathbf{B}_t}^2 \text{ (part 2)}}, \quad (86)$$

where $\mathbf{J}_k := \mathbf{I} - \mathbf{b}_k \mathbf{b}_k^{\top}$ and we have omitted the iteration number t on $\{\mathbf{a}_j, \mathbf{b}_j\}_{j \in [n]}$ to keep notation light. Note that in (86) the norms $\|\mathbf{b}_i\|_2, \|\mathbf{b}_j\|_2$ no longer appear as the projection step enforces $\|\mathbf{b}_i\|_2 = 1$. At each step of the projected gradient descent dynamics, we decompose $\mathbf{B}_t \mathbf{B}_t^{\top}$ as follows:

$$\mathbf{B}_t \mathbf{B}_t^{\top} = \mathbf{I} + \mathbf{Z}_t + \mathbf{X}_t, \quad (87)$$

where $\mathbf{B}_0 \mathbf{B}_0^{\top} = \mathbf{U} \mathbf{\Lambda}_0 \mathbf{U}^{\top}$, $\mathbf{Z}_t = \mathbf{U}(\mathbf{\Lambda}_t - \mathbf{I})\mathbf{U}^{\top}$ and $\mathbf{\Lambda}_{t+1} = g(\mathbf{\Lambda}_t)$ for some function $g : \mathbb{R}^n \rightarrow \mathbb{R}^n$ which defines the spectrum evolution. Here, $\mathbf{U}$ is an orthogonal matrix that importantly does not depend on t and $\mathbf{\Lambda}_t$ is the diagonal matrix containing the eigenvalues (i.e., $\mathbf{U} \mathbf{\Lambda}_t \mathbf{U}^{\top}$ is the SVD). We also define $\mathbf{X}_t^D := \text{Diag}(\mathbf{X}_t)$ and $\mathbf{X}_t^O := \mathbf{X}_t - \mathbf{X}_t^D$.

For now we will make the following assumptions, which will be proved later in the argument. There exist universal constants $C, C_X > 0$ and $\delta \in (0, 1)$ (depending only on r) such that, with probability at least $1 - Ce^{-cd}$,

$$\begin{aligned} \inf_{t \geq 0} \lambda_{\min}(\mathbf{Z}_t) &\geq -1 + \delta_r, \\ \sup_{t \geq 0} \|\mathbf{Z}_t\|_{op} &\leq C, \\ \sup_{t \geq 0} \|\mathbf{X}_t\|_{op} &\leq C_X \frac{\text{poly}(\log d)}{\sqrt{d}}, \\ \|\mathbf{\Lambda}_t - \mathbf{I}\|_{op} &\leq C e^{-c\eta t}. \end{aligned} \quad (88)$$

Here, $\text{poly}(\log d)$ is used to denote polynomial powers of $\log d$, i.e., $(\log d)^C$ for some universal constant C . In the assumptions (88), we specifically distinguish the constant C_X in the bound on $\|\mathbf{X}_t\|_{op}$ from the others. This important distinction between C and C_X will be apparent later to show that assumptions (88) indeed hold. Note also that, for sufficiently large d , (88) implies that

$$\sup_{t \geq 0} \|\mathbf{X}_t\|_{op} \leq 1. \quad (89)$$

We are now ready to give the proof Theorem 4.6. For the convenience of the reader we restate it here.

Theorem E.1. Consider the projected gradient descent algorithm as described above applied to the objective (13) for any f of the form $f(x) = x + \sum_{\ell=3} c_\ell^2 x^\ell$, where $\sum_{\ell=3} c_\ell^2 < \infty$. Initialize the algorithm with $\mathbf{B}_0$ equal to a row-normalized Gaussian, i.e., $(\mathbf{B}'_0)_{i,j} \sim \mathcal{N}(0, 1/d)$, $(\mathbf{B}_0)_{i,:} = \text{Proj}_{\mathbb{S}^{d-1}}((\mathbf{B}'_0)_{i,:})$. Let the step size η be $\Theta(1/\sqrt{d})$. Then, for any $r < 1$, we have that at any time $t = T/\eta$, with probability at least $1 - Ce^{-cd}$,

$$\left\| \mathbf{B}_t \mathbf{B}_t^\top - \mathbf{I} \right\|_{op} \leq C(1 - c)^T,$$

where $C > 0$ and $c \in (0, 1]$ are universal constants depending only on r and f .

Let $\mathbf{E}^t := \mathbf{E}(\mathbf{X}_t, \mathbf{Z}_t) \in \mathbb{R}^{n \times n}$ be a generic matrix whose operator norm is upper bounded by

$$\left\| \mathbf{E}^t \right\|_{op} \leq C \left(\frac{\text{poly}(\log d)}{\sqrt{d}} \cdot \left\| \mathbf{Z}_t \right\|_{op}^{1/2} + \left\| \mathbf{X}_t \right\|_{op}^2 + \left\| \mathbf{X}_t \right\|_{op} \left\| \mathbf{Z}_t \right\|_{op}^{1/2} \right). \quad (90)$$

We highlight that the constant in front of the upper-bound on the error term $\mathbf{E}^t$ is independent of C_X and t .

Lemma E.2 (Bound for the matrix inverse). Assume that (88) holds. Then, for all sufficiently large n , with probability at least $1 - 1/d^2$, jointly for all $t \geq 0$ and $\ell \geq 3$, the following bounds hold

$$\left\| (\mathbf{B}_t \mathbf{B}_t^\top - \mathbf{I})^{\circ \ell} \right\|_{op} \leq \left\| \mathbf{E}^t \right\|_{op}, \quad (91)$$

$$\left\| (f(\mathbf{B}_t \mathbf{B}_t^\top))^{-1} - (\alpha \mathbf{I} + \mathbf{B}_t \mathbf{B}_t^\top)^{-1} \right\|_{op} \leq \left\| \mathbf{E}^t \right\|_{op}, \quad (92)$$

where α was defined as $\alpha = \sum_{\ell=3}^\infty c_\ell^2$.

Proof of Lemma E.2. Note that, for any square matrices $\mathbf{R}, \mathbf{S} \in \mathbb{R}^{n \times n}$,

$$\left\| \mathbf{R} \circ \mathbf{S} \right\|_{op} \leq \sqrt{n} \left\| \mathbf{S} \right\|_{op} \max_{i,j} |\mathbf{R}_{i,j}|. \quad (93)$$

Thus, for $\ell \geq 3$,

$$\begin{aligned} \left\| (\mathbf{B}_t \mathbf{B}_t^\top - \mathbf{I})^{\circ \ell} \right\|_{op} &\leq \sqrt{n} \left\| (\mathbf{B}_t \mathbf{B}_t^\top - \mathbf{I})^{\circ(\ell-3)} \right\|_{op} \max_{i,j} |(\mathbf{B}_t \mathbf{B}_t^\top - \mathbf{I})^{\circ 3}_{i,j}| \\ &= \sqrt{n} \left\| (\mathbf{B}_t \mathbf{B}_t^\top - \mathbf{I})^{\circ(\ell-3)} \right\|_{op} \max_{i \neq j} |(\mathbf{B}_t \mathbf{B}_t^\top - \mathbf{I})^{\circ 3}_{i,j}| \\ &= \sqrt{n} \left\| (\mathbf{B}_t \mathbf{B}_t^\top - \mathbf{I})^{\circ(\ell-3)} \right\|_{op} \max_{i \neq j} |(\mathbf{Z}_t + \mathbf{X}_t)^{\circ 3}_{i,j}|, \end{aligned} \quad (94)$$

where in the first line we use (93), in the second line we use that $((\mathbf{B}_t \mathbf{B}_t^\top - \mathbf{I})^{\circ 3})_{i,i} = 0$ for $i \in [n]$ and in the third line we use the decomposition (87).

Let us bound the off-diagonal entries of $\mathbf{X}_t$ via (88) and the off-diagonal entries of $\mathbf{Z}_t$ via Lemma G.2. This gives that, with probability at least $1 - 1/d^2$, jointly for all $t \geq 0$,

$$\max_{i \neq j} |(\mathbf{Z}_t + \mathbf{X}_t)^{\circ 3}_{i,j}| \leq (C + C_X)^3 \left(\frac{\text{poly}(\log d)}{d} \right)^{3/2}. \quad (95)$$

We will condition on this event (without explicitly mentioning it every time) for the remainder of the argument. By combining (94) and (95), we have that

$$\begin{aligned} \left\| (\mathbf{B}_t \mathbf{B}_t^\top - \mathbf{I})^{\circ \ell} \right\|_{op} &\leq \sqrt{n} \left[(C + C_X)^3 \left(\frac{\text{poly}(\log d)}{d} \right)^{3/2} \right] \left\| (\mathbf{B}_t \mathbf{B}_t^\top - \mathbf{I})^{\circ(\ell-3)} \right\|_{op} \\ &\leq \left\| (\mathbf{B}_t \mathbf{B}_t^\top - \mathbf{I})^{\circ(\ell-3)} \right\|_{op} \end{aligned} \quad (96)$$

where the last inequality holds for all sufficiently large n . Note that, for any square matrices $\mathbf{R}, \mathbf{S}$, an application of Theorem 1 in (Visick, 2000) gives that

$$\left\| \mathbf{R} \circ \mathbf{S} \right\|_{op} \leq \left\| \mathbf{R} \right\|_{op} \left\| \mathbf{S} \right\|_{op}. \quad (97)$$

Hence,

$$\|(B_t B_t^\top - I)^{\circ \ell}\|_{op} \leq \|(B_t B_t^\top - I)^{\circ(\ell-3)}\|_{op} \|(B_t B_t^\top - I)^{\circ 3}\|_{op}. \quad (98)$$

Now, by using again (97) and the assumptions (88), we have that, for $\ell \in [3]$,

$$\|(B_t B_t^\top - I)^{\circ \ell}\|_{op} \leq C. \quad (99)$$

Thus, by combining (96) and (99), we obtain that $\|(B_t B_t^\top - I)^{\circ(\ell-3)}\|_{op}$ is uniformly bounded in ℓ , which together with (98) gives that

$$\|(B_t B_t^\top - I)^{\circ \ell}\|_{op} \leq C \|(B_t B_t^\top - I)^{\circ 3}\|_{op}. \quad (100)$$

We remark here that C is independent of l and C_X . This means that it suffices to prove the claim (91) for $\ell = 3$.

To do so, define $H := \mathbf{1}\mathbf{1}^\top - I$, hence, since $B_t B_t^\top$ has unit diagonal, we have that

$$\begin{aligned} (B_t B_t^\top - I)^{\circ 3} &= (B_t B_t^\top - I)^{\circ 3} \circ H = (U(\Lambda_t - I)U^\top + X_t^O + X_t^D)^{\circ 3} \circ H = (Z_t \circ H + X_t^O \circ H + X_t^D \circ H)^{\circ 3} \\ &= (Z_t \circ H + X_t^O)^{\circ 3} = (Z_t \circ H)^{\circ 3} + 3(Z_t \circ H)^{\circ 2} \circ X_t^O + 3(Z_t \circ H) \circ (X_t^O)^{\circ 2} + (X_t^O)^{\circ 3}. \end{aligned}$$

Using again (97) and that, by Lemma G.1 for any $R \in \mathbb{R}^{n \times n}$,

$$\|R \circ H\|_{op} = \|R - \text{diag}(R)\|_{op} \leq C \|R\|_{op},$$

we get

$$\begin{aligned} \|(B_t B_t^\top - I)^{\circ 3}\|_{op} &\leq C \left(\|(Z_t \circ H)^{\circ 3}\|_{op} + \|Z_t\|_{op}^2 \|X_t^O\|_{op} + \|Z_t\|_{op} \|X_t^O\|_{op}^2 + \|X_t^O\|_{op}^3 \right) \\ &\leq C \left(\|(Z_t \circ H)^{\circ 3}\|_{op} + \|Z_t\|_{op}^{1/2} \|X_t^O\|_{op} + \|X_t^O\|_{op}^2 \right), \end{aligned} \quad (101)$$

where the second step holds since $\|X_t^O\|_{op} \leq 1$ and $\|Z_t\|_{op} \leq C$ by (88)-(89). Another application of (93) gives that

$$\begin{aligned} \|(Z_t \circ H)^{\circ 3}\|_{op} &= \|(Z_t \circ H)^{\circ 2} \circ Z_t\|_{op} \leq \sqrt{n} \cdot \max_{i \neq j} |(Z_t)_{i,j}|^2 \cdot \|Z_t\|_{op} \\ &\leq C \frac{\log d}{\sqrt{d}} \cdot \|Z_t\|_{op} \leq C \frac{\log d}{\sqrt{d}} \cdot \|Z_t\|_{op}^{1/2}, \end{aligned} \quad (102)$$

where the second passage follows from Lemma G.2 and the last from $\|Z_t\|_{op} \leq C$. By combining (101) and (102), the proof of (91) for $\ell = 3$ is complete.

To prove (92), define the following quantity

$$Y := \sum_{\ell=3}^{\infty} c_\ell^2 (B_t B_t^\top - I)^{\circ \ell}.$$

By definition of $f(\cdot)$ we have that

$$f(B_t B_t^\top) = \alpha I + B_t B_t^\top + Y,$$

which implies that

$$\begin{aligned} (f(B_t B_t^\top))^{-1} &= (\alpha I + B_t B_t^\top + Y)^{-1} \\ &= (I + Y(\alpha I + B_t B_t^\top)^{-1})^{-1} (\alpha I + B_t B_t^\top)^{-1} \\ &= \left(I + \sum_{k=1}^{\infty} (-1)^k (Y(\alpha I + B_t B_t^\top)^{-1})^k \right) (\alpha I + B_t B_t^\top)^{-1}. \end{aligned} \quad (103)$$

By definition (90), we have that $\|E^t\|_{op} \leq 1/2$ under assumptions (88) for sufficiently large d . Hence, by the result (91) we have just proved, $\|(B_t B_t^\top - I)^{\circ \ell}\|_{op} \leq 1/2$, which implies that $\sum_{\ell=3}^{\infty} c_\ell^2 \|(B_t B_t^\top - I)^{\circ \ell}\|_{op} \leq \alpha/2$. Thus, we have

$$\|Y(B_t B_t^\top + \alpha I)^{-1}\|_{op} \leq \|Y\|_{op} \|(B_t B_t^\top + \alpha I)^{-1}\|_{op} \leq \frac{\alpha}{2} \cdot \frac{1}{\alpha} \leq \frac{1}{2}. \quad (104)$$

Therefore, we can conclude that

$$\begin{aligned}
 \|(f(\mathbf{B}_t \mathbf{B}_t^\top))^{-1} - (\alpha \mathbf{I} + \mathbf{B}_t \mathbf{B}_t^\top)^{-1}\|_{op} &\leq \|(\alpha \mathbf{I} + \mathbf{B}_t \mathbf{B}_t^\top)^{-1}\|_{op} \cdot \sum_{k=1}^{\infty} \|\mathbf{Y}(\alpha \mathbf{I} + \mathbf{B}_t \mathbf{B}_t^\top)^{-1}\|_{op}^k \\
 &\leq \frac{1}{\alpha} \cdot \frac{\|\mathbf{Y}(\alpha \mathbf{I} + \mathbf{B}_t \mathbf{B}_t^\top)^{-1}\|_{op}}{1 - \|\mathbf{Y}(\alpha \mathbf{I} + \mathbf{B}_t \mathbf{B}_t^\top)^{-1}\|_{op}} \\
 &\leq \frac{2}{\alpha} \cdot \|\mathbf{Y}\|_{op} \|(\alpha \mathbf{I} + \mathbf{B}_t \mathbf{B}_t^\top)^{-1}\|_{op} \\
 &\leq \frac{2}{\alpha^2} \cdot \|\mathbf{Y}\|_{op},
 \end{aligned} \tag{105}$$

where the third inequality uses (104). By bounding $\|\mathbf{Y}\|_{op}$ via (91), the proof of (92) is complete. $\square$

Lemma E.3 (Bound for the Schur product with $\mathbf{A}^\top \mathbf{A}$). *Assume that (88) holds, and let $\mathbf{A}_t$ be given by (85). Then, we have that, with probability at least $1 - 1/d^2$, jointly for all $t \geq 0$ and $\ell \geq 2$,*

$$\left\| \mathbf{A}_t^\top \mathbf{A}_t \circ (\mathbf{B}_t \mathbf{B}_t^\top - \mathbf{I})^{\circ \ell} \right\|_{op} \leq \|\mathbf{E}^t\|_{op}. \tag{106}$$

Proof of Lemma E.3. We have that

$$\begin{aligned}
 \|\mathbf{A}_t^\top \mathbf{A}_t \circ (\mathbf{B}_t \mathbf{B}_t^\top - \mathbf{I})^{\circ \ell}\|_{op} &\leq \|\mathbf{A}_t^\top \mathbf{A}_t \circ (\mathbf{B}_t \mathbf{B}_t^\top - \mathbf{I})^{\circ 2}\|_{op} \left\| (\mathbf{B}_t \mathbf{B}_t^\top - \mathbf{I})^{\circ(\ell-2)} \right\|_{op} \\
 &\leq C \|\mathbf{A}_t^\top \mathbf{A}_t \circ (\mathbf{B}_t \mathbf{B}_t^\top - \mathbf{I})^{\circ 2}\|_{op},
 \end{aligned} \tag{107}$$

where the first inequality uses (97) and the second inequality uses that $\left\| (\mathbf{B}_t \mathbf{B}_t^\top - \mathbf{I})^{\circ(\ell-2)} \right\|_{op}$ is uniformly bounded in l , which follows from (96) and (99).

Let us now focus on bounding the RHS of (107). An application of Lemma E.2 gives that

$$(f(\mathbf{B}_t \mathbf{B}_t^\top))^{-1} = (\alpha \mathbf{I} + \mathbf{B}_t \mathbf{B}_t^\top)^{-1} + \mathbf{E}_1,$$

where

$$\|\mathbf{E}\|_{op} \leq \|\mathbf{E}^t\|_{op}.$$

Hence, by using (85), we get that

$$\begin{aligned}
 \mathbf{A}_t^\top \mathbf{A}_t &= ((\alpha \mathbf{I} + \mathbf{B}_t \mathbf{B}_t^\top)^{-1} \mathbf{B}_t + \mathbf{E}_1^\top \mathbf{B}_t) (\mathbf{B}_t^\top (\alpha \mathbf{I} + \mathbf{B}_t \mathbf{B}_t^\top)^{-1} + \mathbf{B}_t^\top \mathbf{E}_1) \\
 &= \mathbf{B}_t \mathbf{B}_t^\top (\alpha \mathbf{I} + \mathbf{B}_t \mathbf{B}_t^\top)^{-2} + \mathbf{E}_1^\top \mathbf{B}_t \mathbf{B}_t^\top (\alpha \mathbf{I} + \mathbf{B}_t \mathbf{B}_t^\top)^{-1} + (\alpha \mathbf{I} + \mathbf{B}_t \mathbf{B}_t^\top)^{-1} \mathbf{B}_t \mathbf{B}_t^\top \mathbf{E}_1 + \mathbf{E}_1^\top \mathbf{B}_t \mathbf{B}_t^\top \mathbf{E}_1,
 \end{aligned} \tag{108}$$

where we rearranged the first term in (108) using that $\mathbf{B}_t \mathbf{B}_t^\top$ and $(\alpha \mathbf{I} + \mathbf{B}_t \mathbf{B}_t^\top)^{-1}$ commute. By using the assumptions (88), we have that

$$\|\mathbf{B}_t \mathbf{B}_t^\top\|_{op} \leq C, \quad \|\mathbf{E}_1\|_{op} \leq 1/2, \quad \|(\alpha \mathbf{I} + \mathbf{B}_t \mathbf{B}_t^\top)^{-1}\|_{op} \leq \frac{1}{\alpha}.$$

Hence, we can upper bound the operator norm of the last three terms in (108) as

$$\left\| \mathbf{E}_1^\top \mathbf{B}_t \mathbf{B}_t^\top (\alpha \mathbf{I} + \mathbf{B}_t \mathbf{B}_t^\top)^{-1} + (\alpha \mathbf{I} + \mathbf{B}_t \mathbf{B}_t^\top)^{-1} \mathbf{B}_t \mathbf{B}_t^\top \mathbf{E}_1 + \mathbf{E}_1^\top \mathbf{B}_t \mathbf{B}_t^\top \mathbf{E}_1 \right\|_{op} \leq C \|\mathbf{E}_1\|_{op}. \tag{109}$$

Let us now take a closer look at the first term in (108). Recall that

$$\mathbf{B}_t \mathbf{B}_t^\top = \mathbf{U} \mathbf{\Lambda}_t \mathbf{U}^\top + \mathbf{X}_t.$$

As the operator norm is sub-multiplicative, we have that

$$\|\mathbf{X}_t \cdot (\alpha \mathbf{I} + \mathbf{B}_t \mathbf{B}_t^\top)^{-2}\|_{op} \leq C \|\mathbf{X}_t\|_{op}. \tag{110}$$

Furthermore,

$$\begin{aligned} U\Lambda_t U^\top (\alpha I + U\Lambda_t U^\top + X_t)^{-2} &= U\Lambda_t U^\top ((I + X_t(\alpha I + U\Lambda_t U^\top)^{-1})(\alpha I + U\Lambda_t U^\top))^{-2} \\ &= U\Lambda_t U^\top T_1^{-1} T_2^{-1} T_1^{-1} T_2^{-1}, \end{aligned} \quad (111)$$

where we have defined

$$T_1 = \alpha I + U\Lambda_t U^\top, \quad T_2 = I + X_t(\alpha I + U\Lambda_t U^\top)^{-1}.$$

By expanding T_2^{-1} as in (103)-(105), we get

$$\|T_2^{-1} - I\|_{op} \leq C\|X_t\|_{op},$$

or equivalently

$$T_2^{-1} = I + E_2,$$

with $\|E_2\|_{op} \leq C\|X_t\|_{op}$. In this view, looking at (111) we have

$$U\Lambda_t U^\top T_1^{-1} T_2^{-1} T_1^{-1} T_2^{-1} = U\Lambda_t B U^\top T_1^{-1} (I + E_2) T_1^{-1} (I + E_2).$$

All the terms which involve E_2 can be controlled. We provide the analysis for two terms of different nature, the rest follows from similar arguments. As $\|T_1^{-1}\|_{op} \leq 1/\alpha$ and $\|\Lambda_t\|_{op} \leq C$, we have that

$$\begin{aligned} \|U\Lambda_t U^\top T_1^{-1} E_2 T_1^{-1} E_2\|_{op} &\leq \|T_1^{-1}\|_{op}^2 \|E_2\|_{op}^2 \leq \frac{C}{\alpha^2} \|X_t\|_{op}^2 \leq \frac{C}{\alpha^2} \|X_t\|_{op}, \\ \|U\Lambda_t U^\top T_1^{-1} I T_1^{-1} E_2\|_{op} &\leq \|T_1^{-1}\|_{op}^2 \|E_2\|_{op} \leq \frac{C}{\alpha^2} \|X_t\|_{op}, \end{aligned}$$

where we have also used that $\|X_t\|_{op}$ is bounded via assumptions (88). Furthermore, a simple manipulation gives

$$U\Lambda_t U^\top T_1^{-2} = U\Lambda_t U^\top (\alpha I + U\Lambda_t U^\top)^{-2} = U\Lambda_t (\alpha I + \Lambda_t)^{-2} U^\top = U\phi(\Lambda_t) U^\top,$$

where $\phi(x) = \frac{x}{(\alpha+x)^2}$. As a result,

$$\left\| U\Lambda_t U^\top T_1^{-1} T_2^{-1} T_1^{-1} T_2^{-1} - U\phi(\Lambda_t) U^\top \right\|_{op} \leq C\|X_t\|_{op},$$

which implies that

$$\|B_t B_t^\top (\alpha I + B_t B_t^\top)^{-2} - U\phi(\Lambda_t) U^\top\|_{op} \leq C\|X_t\|_{op}. \quad (112)$$

By combining (108), (109) and (112), we have that

$$\|A_t^\top A_t - U\phi(\Lambda_t) U^\top\|_{op} \leq C(\|X_t\|_{op} + \|E_1\|_{op}). \quad (113)$$

At this point, we are ready to analyze the operator norm of $\|A_t^\top A_t \circ (B_t B_t^\top - I)^{\circ 2}\|_{op}$:

$$\begin{aligned} A_t^\top A_t \circ (B_t B_t^\top - I)^{\circ 2} &= (U\phi(\Lambda_t) U^\top + E_3) \circ (U(\Lambda_t - I) U^\top + X_t)^{\circ 2} \circ H \\ &= (U\phi(\Lambda_t) U^\top + E_3) \circ ((U(\Lambda_t - I) U^\top)^{\circ 2} + X_t^{\circ 2} + 2(U(\Lambda_t - I) U^\top) \circ X_t) \circ H, \end{aligned} \quad (114)$$

where we have defined $H := 11^\top - I$ and $\|E_3\|_{op} \leq C(\|X_t\|_{op} + \|E_1\|_{op})$. We now decompose the quantity into three terms:

$$A_t^\top A_t \circ (B_t B_t^\top - I)^{\circ 2} = S_1 + S_2 + S_3,$$

where

$$\begin{aligned} S_1 &= (U\phi(\Lambda_t) U^\top \circ U(\Lambda_t - I) U^\top \circ H) \circ U(\Lambda_t - I) U^\top, \\ S_2 &= H \circ E_3 \circ ((U(\Lambda_t - I) U^\top)^{\circ 2} + X_t^{\circ 2} + 2(U(\Lambda_t - I) U^\top) \circ X_t), \\ S_3 &= H \circ U\phi(\Lambda_t) U^\top \circ (X_t^{\circ 2} + 2(U(\Lambda_t - I) U^\top) \circ X_t). \end{aligned}$$

We proceed to bound each of these terms separately.

We start with $\mathbf{S}_1$. As $\phi(x)$ is differentiable for $x \geq 0$, the derivative of $\phi(x)$ is bounded for any compact interval $I \subseteq \mathbb{R}_+$. Hence, $\phi(x)$ is locally Lipschitz on I with Lipschitz constant $C_I > 0$, which implies that

$$|\phi(x) - \phi(1)| = \left| \phi(x) - \frac{1}{(1+\alpha)^2} \right| \leq C_I |x - 1|.$$

By assumption (88), we have that $\Lambda_t \succ 0$ and $\|\Lambda_t\|_{op} \leq C$, hence

$$\left\| \mathbf{U}\phi(\Lambda_t)\mathbf{U}^\top - \frac{1}{(1+\alpha)^2}\mathbf{I} \right\|_{op} \leq C_I \cdot \|\mathbf{Z}_t\|_{op}. \quad (115)$$

Hence, an application of Lemma G.2 gives that, with probability at least $1 - 1/d^2$,

$$\sup_{t \geq 0} m \left(\mathbf{U}\phi(\Lambda_t)\mathbf{U}^\top - \frac{1}{(1+\alpha)^2}\mathbf{I} \right) \leq c\sqrt{\frac{\log d}{d}}, \quad (116)$$

where $c > 0$ is a universal constant. Another application of Lemma G.2 also gives that, with the same probability,

$$\sup_{t \geq 0} m \left(\mathbf{U}(\Lambda_t - \mathbf{I})\mathbf{U}^\top \right) \leq c\sqrt{\frac{\log d}{d}}. \quad (117)$$

As a result, we obtain the bound

$$\|\mathbf{S}_1\|_{op} = \|([\mathbf{U}\phi(\Lambda_t)\mathbf{U}^\top - 1/(1+\alpha)^2\mathbf{I}] \circ \mathbf{U}(\Lambda_t - \mathbf{I})\mathbf{U}^\top \circ \mathbf{H}) \circ \mathbf{U}(\Lambda_t - \mathbf{I})\mathbf{U}^\top\|_{op} \leq C \frac{\log d}{\sqrt{d}} \|\mathbf{Z}_t\|_{op}. \quad (118)$$

Here, the first equality is due to the fact that we are taking the Hadamard product with the matrix $\mathbf{H}$ which has 0 on the diagonal, hence we can add multiples of the identity to $\mathbf{U}\phi(\Lambda_t)\mathbf{U}^\top$; and the second inequality uses (93) with $\mathbf{R} = [\mathbf{U}\phi(\Lambda_t)\mathbf{U}^\top - 1/(1+\alpha)^2\mathbf{I}] \circ \mathbf{U}(\Lambda_t - \mathbf{I})\mathbf{U}^\top \circ \mathbf{H}$ and $\mathbf{S} = \mathbf{U}(\Lambda_t - \mathbf{I})\mathbf{U}^\top$ in combination with (116)-(117).

Next, we bound $\|\mathbf{S}_2\|_{op}$. We inspect the terms appearing in the expression for $\mathbf{S}_2$ one by one. First note that we can omit $\mathbf{H}$ in the expression since, by Lemma G.1 for any square matrix $\mathbf{R}$

$$\|\mathbf{R} \circ \mathbf{H}\|_{op} \leq C\|\mathbf{R}\|_{op}. \quad (119)$$

Hence, by using (97), we get

$$\begin{aligned} \|\mathbf{H} \circ \mathbf{E}_3 \circ ((\mathbf{U}(\Lambda_t - \mathbf{I})\mathbf{U}^\top)^{\circ 2})\|_{op} &\leq C\|\mathbf{E}_3\|_{op}\|\mathbf{Z}_t\|_{op}^2 \\ \|\mathbf{H} \circ \mathbf{E}_3 \circ \mathbf{X}_t^{\circ 2}\|_{op} &\leq C\|\mathbf{E}_3\|_{op}\|\mathbf{X}_t\|_{op}^2 \\ \|\mathbf{H} \circ \mathbf{E}_3 \circ 2(\mathbf{U}(\Lambda_t - \mathbf{I})\mathbf{U}^\top) \circ \mathbf{X}_t\|_{op} &\leq C\|\mathbf{E}_3\|_{op}\|\mathbf{X}_t\|_{op}\|\mathbf{Z}_t\|_{op}, \end{aligned}$$

which leads to the bound

$$\|\mathbf{S}_2\|_{op} \leq C\|\mathbf{E}_3\|_{op} (\|\mathbf{X}_t\|_{op}^2 + \|\mathbf{Z}_t\|_{op}^2 + \|\mathbf{X}_t\|_{op}\|\mathbf{Z}_t\|_{op}). \quad (120)$$

Finally, we bound $\|\mathbf{S}_3\|_{op}$. Consider the term

$$\|[\mathbf{H} \circ \mathbf{U}\phi(\Lambda_t)\mathbf{U}^\top \circ 2(\mathbf{U}(\Lambda_t - \mathbf{I})\mathbf{U}^\top)] \circ \mathbf{X}_t\|_{op}.$$

Then, by using (119) and (115), we have

$$\|\mathbf{H} \circ \mathbf{U}\phi(\Lambda_t)\mathbf{U}^\top\|_{op} = \left\| \mathbf{H} \circ [\mathbf{U}\phi(\Lambda_t)\mathbf{U}^\top - \frac{1}{(1+\alpha)^2}\mathbf{I}] \right\|_{op} \leq C \left\| \mathbf{U}\phi(\Lambda_t)\mathbf{U}^\top - \frac{1}{(1+\alpha)^2}\mathbf{I} \right\|_{op} \leq C\|\mathbf{Z}_t\|_{op}. \quad (121)$$

Hence, in conjunction with (97), we get

$$\|\mathbf{H} \circ \mathbf{U}\phi(\Lambda_t)\mathbf{U}^\top \circ 2\mathbf{U}(\Lambda_t - \mathbf{I})\mathbf{U}^\top\|_{op} \leq C \cdot \|\mathbf{Z}_t\|_{op}^2,$$

which invoking (97) one more time gives

$$\|[\mathbf{H} \circ \mathbf{U}\phi(\mathbf{\Lambda}_t)\mathbf{U}^\top \circ 2(\mathbf{U}(\mathbf{\Lambda}_t - \mathbf{I})\mathbf{U}^\top)] \circ \mathbf{X}_t\|_{op} \leq C\|\mathbf{Z}_t\|_{op}^2\|\mathbf{X}_t\|_{op}.$$

Furthermore, by combining (97) and (121), we get

$$\|[\mathbf{H} \circ \mathbf{\Lambda}\phi(\mathbf{\Lambda}_t)\mathbf{\Lambda}^\top] \circ \mathbf{X}_t^{\circ 2}\|_{op} \leq C\|\mathbf{Z}_t\|_{op}\|\mathbf{X}_t\|_{op}^2.$$

Thus,

$$\|\mathbf{S}_3\|_{op} \leq C(\|\mathbf{Z}_t\|_{op}^2\|\mathbf{X}_t\|_{op} + \|\mathbf{Z}_t\|_{op}\|\mathbf{X}_t\|_{op}^2). \quad (122)$$

Recall that, from assumptions (88)-(89), $\|\mathbf{X}_t\|_{op}, \|\mathbf{Z}_t\|_{op} \leq C$. Then, by combining the bounds in (118), (120) and (122), the desired result readily follows. $\square$

By exploiting the above lemmas, we are able to make the following approximation for the gradient.

Lemma E.4 (Gradient approximation). *Assume that (88) holds, and let $\nabla_{\mathbf{B}_t}$ be given by (86). Further define $\gamma = 1 + \alpha$ and $F(x) = \frac{1+x}{(\gamma+x)^2}$. Then, for all sufficiently large n , with probability $1 - 1/d^2$, jointly for all $t \geq 0$,*

$$\left\| \frac{1}{2}\nabla_{\mathbf{B}_t}\mathbf{B}_t^\top + \alpha F(\mathbf{Z}_t) - \alpha \text{Diag}(F(\mathbf{Z}_t))(\mathbf{I} + \mathbf{Z}_t) - \frac{2\alpha}{\gamma^3}\mathbf{X}_t^O - \frac{\alpha}{\gamma^2}\mathbf{X}_t^D \right\|_{op} \leq \|\mathbf{E}^t\|_{op}. \quad (123)$$

Proof of Lemma E.4. We start by showing that, with probability $1 - 1/d^2$, jointly for all $t \geq 0$,

$$\left\| \frac{1}{2}\nabla_{\mathbf{B}_t} + \alpha(\alpha\mathbf{I} + \mathbf{B}_t\mathbf{B}_t^\top)^{-2}\mathbf{B}_t - \alpha \text{Diag}\left((\alpha\mathbf{I} + \mathbf{B}_t\mathbf{B}_t^\top)^{-2}(\mathbf{B}_t\mathbf{B}_t^\top)\right)\mathbf{B}_t \right\|_{op} \leq \|\mathbf{E}^t\|_{op}. \quad (124)$$

Let us first consider the term $\nabla_{\mathbf{B}_t}^1$, which can be equivalently expressed as

$$\nabla_{\mathbf{B}_t}^1 = 2(-\mathbf{A}_t^\top + \text{Diag}(\mathbf{B}_t\mathbf{A}_t)\mathbf{B}_t + \mathbf{T}\mathbf{B}_t - \text{Diag}(\mathbf{T}(\mathbf{B}_t\mathbf{B}_t^\top))\mathbf{B}_t),$$

where $\mathbf{T} = \mathbf{A}_t^\top \mathbf{A}_t - \text{Diag}(\mathbf{A}_t^\top \mathbf{A}_t)$. It is then easy to verify that

$$\frac{1}{2}\nabla_{\mathbf{B}_t}^1 = -\mathbf{A}_t^\top + \mathbf{A}_t^\top \mathbf{A}_t \mathbf{B}_t + \text{Diag}(\mathbf{B}_t\mathbf{A}_t)\mathbf{B}_t - \text{Diag}(\mathbf{A}_t^\top \mathbf{A}_t \mathbf{B}_t \mathbf{B}_t^\top)\mathbf{B}_t. \quad (125)$$

Using Lemma E.2, we get

$$\mathbf{A}_t^\top \mathbf{A}_t = ((\alpha\mathbf{I} + \mathbf{B}_t\mathbf{B}_t^\top)^{-1} + \mathbf{E}_1)\mathbf{B}_t\mathbf{B}_t^\top((\alpha\mathbf{I} + \mathbf{B}_t\mathbf{B}_t^\top)^{-1} + \mathbf{E}_1), \quad (126)$$

where $\|\mathbf{E}_1\|_{op} \leq \|\mathbf{E}^t\|_{op}$. It follows from (88) that $\|\mathbf{B}_t\mathbf{B}_t^\top\|_{op} \leq C$. Hence, using that $\mathbf{B}_t\mathbf{B}_t^\top$ and $(\alpha\mathbf{I} + \mathbf{B}_t\mathbf{B}_t^\top)$ commute in conjunction with $\|(\alpha\mathbf{I} + \mathbf{B}_t\mathbf{B}_t^\top)^{-1}\|_{op} \leq 1/\alpha$ we get

$$\mathbf{A}_t^\top \mathbf{A}_t = \mathbf{B}_t\mathbf{B}_t^\top(\alpha\mathbf{I} + \mathbf{B}_t\mathbf{B}_t^\top)^{-2} + \mathbf{E}_2, \quad (127)$$

where $\|\mathbf{E}_2\|_{op} \leq \|\mathbf{E}^t\|_{op}$. Noting that $\frac{1}{\alpha+x} - \frac{\alpha}{(\alpha+x)^2} = \frac{x}{(\alpha+x)^2}$ and using the spectral theorem for the symmetric matrix $\mathbf{B}_t\mathbf{B}_t^\top$, we can further rewrite (127) as

$$\mathbf{A}_t^\top \mathbf{A}_t = (\alpha\mathbf{I} + \mathbf{B}_t\mathbf{B}_t^\top)^{-1} - \alpha(\alpha\mathbf{I} + \mathbf{B}_t\mathbf{B}_t^\top)^{-2} + \mathbf{E}_2. \quad (128)$$

With similar arguments, by Lemma E.2, we can write

$$\mathbf{B}_t\mathbf{A}_t = \mathbf{B}_t\mathbf{B}_t^\top(\alpha\mathbf{I} + \mathbf{B}_t\mathbf{B}_t^\top)^{-1} + \mathbf{E}_3, \quad (129)$$

where $\|\mathbf{E}_3\|_{op} \leq \|\mathbf{E}^t\|_{op}$. Noting that $1 - \frac{\alpha}{\alpha+x} = \frac{x}{\alpha+x}$, again by the spectral theorem for $\mathbf{B}_t\mathbf{B}_t^\top$, we get

$$\mathbf{B}_t\mathbf{A}_t = \mathbf{I} - \alpha(\alpha\mathbf{I} + \mathbf{B}_t\mathbf{B}_t^\top)^{-1} + \mathbf{E}_3, \quad (130)$$

and, consequently, we obtain

$$\text{Diag}(\mathbf{B}_t \mathbf{A}_t) \mathbf{B}_t = \mathbf{B}_t - \alpha \text{Diag}((\alpha \mathbf{I} + \mathbf{B}_t \mathbf{B}_t^\top)^{-1}) \mathbf{B}_t + \mathbf{E}_4, \quad (131)$$

where $\|\mathbf{E}_4\|_{op} \leq \|\mathbf{E}^t\|_{op}$. Using (128) and $1 - \frac{\alpha}{\alpha+x} = \frac{1}{x+\alpha}$, we get

$$\begin{aligned} \text{Diag}(\mathbf{A}_t^\top \mathbf{A}_t \mathbf{B}_t \mathbf{B}_t^\top) \mathbf{B}_t &= \text{Diag}((\alpha \mathbf{I} + \mathbf{B}_t \mathbf{B}_t^\top)^{-1} \mathbf{B}_t \mathbf{B}_t^\top) \mathbf{B}_t - \alpha \text{Diag}((\alpha \mathbf{I} + \mathbf{B}_t \mathbf{B}_t^\top)^{-2} \mathbf{B}_t \mathbf{B}_t^\top) \mathbf{B}_t + \mathbf{E}_5 \\ &= \mathbf{B}_t - \alpha \text{Diag}((\alpha \mathbf{I} + \mathbf{B}_t \mathbf{B}_t^\top)^{-1}) \mathbf{B}_t - \alpha \text{Diag}((\alpha \mathbf{I} + \mathbf{B}_t \mathbf{B}_t^\top)^{-2} \mathbf{B}_t \mathbf{B}_t^\top) \mathbf{B}_t + \mathbf{E}_5, \end{aligned} \quad (132)$$

where $\|\mathbf{E}_5\|_{op} \leq \|\mathbf{E}^t\|_{op}$.

With this in mind, we get back to (125). Combining the results of (128), (131) and (132) we get

$$\begin{aligned} \nabla_{\mathbf{B}_t}^1 &= \underbrace{-(\alpha \mathbf{I} + \mathbf{B}_t \mathbf{B}_t^\top)^{-1} \mathbf{B}_t}_{-\mathbf{A}_t^\top} + \underbrace{(\alpha \mathbf{I} + \mathbf{B}_t \mathbf{B}_t^\top)^{-1} \mathbf{B}_t - \alpha(\alpha + \mathbf{B}_t \mathbf{B}_t^\top)^{-2} \mathbf{B}_t}_{\mathbf{A}_t^\top \mathbf{A}_t \mathbf{B}_t} + \underbrace{\mathbf{B}_t - \alpha \text{Diag}((\alpha \mathbf{I} + \mathbf{B}_t \mathbf{B}_t^\top)^{-1}) \mathbf{B}_t}_{\text{Diag}(\mathbf{B}_t \mathbf{A}_t) \mathbf{B}_t} \\ &\quad - \underbrace{\mathbf{B}_t + \alpha \text{Diag}((\alpha \mathbf{I} + \mathbf{B}_t \mathbf{B}_t^\top)^{-1}) \mathbf{B}_t + \alpha \text{Diag}((\alpha \mathbf{I} + \mathbf{B}_t \mathbf{B}_t^\top)^{-2} \mathbf{B}_t \mathbf{B}_t^\top) \mathbf{B}_t}_{-\text{Diag}(\mathbf{A}_t^\top \mathbf{A}_t \mathbf{B}_t \mathbf{B}_t^\top) \mathbf{B}_t} + \mathbf{E}_6 \\ &= -\alpha(\alpha \mathbf{I} + \mathbf{B}_t \mathbf{B}_t^\top)^{-2} \mathbf{B}_t + \alpha \text{Diag}((\alpha \mathbf{I} + \mathbf{B}_t \mathbf{B}_t^\top)^{-2} \mathbf{B}_t \mathbf{B}_t^\top) \mathbf{B}_t + \mathbf{E}_6, \end{aligned} \quad (133)$$

where $\|\mathbf{E}_6\|_{op} \leq \|\mathbf{E}^t\|_{op}$.

Let us now analyze the second part of the gradient which involves terms of the form below for $\ell \geq 3$:

$$\nabla_{\mathbf{B}_t}^{2,k,\ell} := c_\ell^2 \cdot \ell \cdot \sum_{j \neq k} \langle \mathbf{a}_k, \mathbf{a}_j \rangle \langle \mathbf{b}_k, \mathbf{b}_j \rangle^{(\ell-1)} \mathbf{J}_k \mathbf{b}_j.$$

Now, from the fact that

$$\mathbf{J}_k = \mathbf{I} - \mathbf{b}_k \mathbf{b}_k^\top,$$

we can write

$$c_\ell^2 \cdot \ell \cdot \sum_{j \neq k} \langle \mathbf{a}_k, \mathbf{a}_j \rangle \langle \mathbf{b}_k, \mathbf{b}_j \rangle^{(\ell-1)} \mathbf{J}_k \mathbf{b}_j = c_\ell^2 \cdot \ell \cdot \sum_{j \neq k} \langle \mathbf{a}_k, \mathbf{a}_j \rangle \langle \mathbf{b}_k, \mathbf{b}_j \rangle^{(\ell-1)} (\mathbf{b}_j - \langle \mathbf{b}_k, \mathbf{b}_j \rangle \mathbf{b}_k). \quad (134)$$

The second term of the RHS gives the following contribution to the $\mathbf{B}_t$ update

$$\text{Diag}(\mathbf{A}_t^\top \mathbf{A}_t (\mathbf{B}_t \mathbf{B}_t^\top - \mathbf{I})^{\circ \ell}) \mathbf{B}_t.$$

By recalling that $\|\mathbf{A}_t^\top \mathbf{A}_t\|_{op} \leq C$ and $\|\mathbf{B}_t\|_{op} \leq C$, we have

$$\|\text{Diag}(\mathbf{A}_t^\top \mathbf{A}_t (\mathbf{B}_t \mathbf{B}_t^\top - \mathbf{I})^{\circ \ell}) \mathbf{B}_t\|_{op} \leq C \|\mathbf{A}_t^\top \mathbf{A}_t (\mathbf{B}_t \mathbf{B}_t^\top - \mathbf{I})^{\circ \ell}\|_{op} \|\mathbf{B}_t\|_{op} \leq C \|(\mathbf{B}_t \mathbf{B}_t^\top - \mathbf{I})^{\circ \ell}\|_{op}. \quad (135)$$

Now, for $\ell < 5$, we upper bound the RHS of (135) via Lemma E.2, which gives that

$$\|\text{Diag}(\mathbf{A}_t^\top \mathbf{A}_t (\mathbf{B}_t \mathbf{B}_t^\top - \mathbf{I})^{\circ \ell}) \mathbf{B}_t\|_{op} \leq C \|\mathbf{E}^t\|_{op}. \quad (136)$$

Furthermore, if we follow passages analogous to (94)-(95) (the only difference being that we exchange the roles of the Hadamard powers 3 and $\ell - 3$), we have that, with probability at least $1 - 1/d^2$, jointly for all $t \geq 0$ and $\ell \geq 5$,

$$\|\text{Diag}(\mathbf{A}_t^\top \mathbf{A}_t (\mathbf{B}_t \mathbf{B}_t^\top - \mathbf{I})^{\circ \ell}) \mathbf{B}_t\|_{op} \leq C \sqrt{n} \|\mathbf{E}^t\|_{op} \left(\frac{\text{poly}(\log d)}{d} \right)^{(\ell-3)/2} \leq C \|\mathbf{E}^t\|_{op} \left(\frac{\text{poly}(\log d)}{d} \right)^{(\ell-4)/2}, \quad (137)$$

for sufficiently large d .

Define the following quantity:

$$\mathbf{Y} = (\mathbf{A}_t^\top \mathbf{A}_t) \circ (\mathbf{B}_t \mathbf{B}_t^\top - \mathbf{I})^{\circ(\ell-1)}. \quad (138)$$

In this view, the first term in (134) can be written as $\mathbf{Y} \mathbf{B}_t$. For $l < 5$, by Lemma E.3 we have that $\|\mathbf{Y}\|_{op} \leq \|\mathbf{E}^t\|_{op}$, hence $\|\mathbf{Y} \mathbf{B}_t\|_{op} \leq C \|\mathbf{E}^t\|_{op}$ as $\|\mathbf{B}_t\|_{op} \leq C$. Furthermore, with probability at least $1 - 1/d^2$, jointly for all $t \geq 0$ and $\ell \geq 5$, we have

$$\begin{aligned} \|\mathbf{Y} \mathbf{B}_t\|_{op} &\leq C \|\mathbf{Y}\|_{op} = C \sqrt{n} \|(\mathbf{A}_t^\top \mathbf{A}_t) \circ (\mathbf{B}_t \mathbf{B}_t^\top - \mathbf{I})^{\circ 2}\|_{op} \max_{i,j} |(\mathbf{B}_t \mathbf{B}_t^\top - \mathbf{I})_{i,j}|^{\ell-3} \\ &\leq \sqrt{n} \|\mathbf{E}^t\|_{op} \max_{i,j} |(\mathbf{B}_t \mathbf{B}_t^\top - \mathbf{I})_{i,j}|^{\ell-3} \\ &\leq \sqrt{n} \|\mathbf{E}^t\|_{op} \left[(C + C_X)^{\ell-3} \left(\frac{\text{poly}(\log d)}{d} \right)^{(\ell-3)/2} \right] \\ &\leq (C + C_X)^{\ell-3} \|\mathbf{E}^t\|_{op} \left(\frac{\text{poly}(\log d)}{d} \right)^{(\ell-4)/2}. \end{aligned} \quad (139)$$

Here, in the second line we use Lemma E.3; and in the third line we bound the off-diagonal entries of $\mathbf{X}_t$ via (88) and the off-diagonal entries of $\mathbf{Z}_t$ via Lemma G.2. Hence, by combining (137) and (139), we conclude that

$$\|\nabla_{\mathbf{B}_t}^2\|_{op} \leq C \|\mathbf{E}^t\|_{op} + \|\mathbf{E}^t\|_{op} \sum_{\ell=5}^{\infty} (C + C_X)^{\ell-3} c_\ell^2 \ell \left(\frac{\text{poly}(\log d)}{\sqrt{d}} \right)^{\ell-4} \leq C \|\mathbf{E}^t\|_{op}, \quad (140)$$

where we used that the series $\sum_{\ell=5}^{\infty} (C + C_X)^{\ell-3} c_\ell^2 \ell \left(\frac{\text{poly}(\log d)}{\sqrt{d}} \right)^{\ell-4}$ converges to a finite value for all sufficiently large d , since $(C + C_X) \frac{\text{poly}(\log d)}{\sqrt{d}} < 1$. This finishes the proof of (124).

We now further analyse the gradient in (124). Defining $F(x) = \frac{1+x}{(\gamma+x)^2}$, with $\gamma = 1 + \alpha$, we can write

$$\frac{1}{2} \nabla_{\mathbf{B}_t} \mathbf{B}_t^\top = -\alpha F(\mathbf{Z}_t + \mathbf{X}_t) + \alpha \text{Diag}(F(\mathbf{Z}_t + \mathbf{X}_t)) + \alpha \text{Diag}(F(\mathbf{Z}_t + \mathbf{X}_t)) (\mathbf{Z}_t + \mathbf{X}_t) + \mathbf{E}^t. \quad (141)$$

By a slight abuse of notation, we will denote by $F^{(l)}(0)$ the l -th derivative of the unidimensional function $F(x) = \frac{1+x}{(\gamma+x)^2}$ computed at $x = 0$. Here, $F(\mathbf{Z}_t + \mathbf{X}_t)$ is defined by the spectral theorem (note that indeed $\mathbf{Z}_t + \mathbf{X}_t = \mathbf{B}_t \mathbf{B}_t^\top - \mathbf{I}$ is symmetric).

We will now compute the error we incur if in (141) we replace $F(\mathbf{X}_t + \mathbf{Z}_t)$ by $F(\mathbf{Z}_t)$. We first consider the case when $\|\mathbf{Z}_t\|_{op} > \frac{\gamma}{3}$. In this case, we have that

$$\left\| F(\mathbf{Z}_t + \mathbf{X}_t) - F(\mathbf{Z}_t) - F^{(1)}(0) \mathbf{X}_t \right\|_{op} \leq C \|\mathbf{X}_t\|_{op} \leq C \|\mathbf{Z}_t\|_{op} \|\mathbf{X}_t\|_{op}. \quad (142)$$

Here, the second inequality trivially holds since $\|\mathbf{Z}_t\|_{op} > \frac{\gamma}{3}$. To prove the first inequality, let DF be the derivative of the matrix-valued function $F(\mathbf{M}) = (\mathbf{I} + \mathbf{M})(\gamma \mathbf{I} + \mathbf{M})^{-2}$. Then, by evaluating this derivative for $\mathbf{M} = \mathbf{Z}_t$ in the direction of $\mathbf{X}_t$, we obtain

$$DF(\mathbf{Z}_t) \mathbf{X}_t = -(\mathbf{I} + \mathbf{Z}_t)(\gamma \mathbf{I} + \mathbf{Z}_t)^{-1} \mathbf{X}_t (\gamma \mathbf{I} + \mathbf{Z}_t)^{-2} - (\mathbf{I} + \mathbf{Z}_t)(\gamma \mathbf{I} + \mathbf{Z}_t)^{-2} \mathbf{X}_t (\gamma \mathbf{I} + \mathbf{Z}_t)^{-1} + \mathbf{X}_t (\gamma \mathbf{I} + \mathbf{Z}_t)^{-2}. \quad (143)$$

To verify this expression we first note that the derivative of the function $G(\mathbf{M}) = \mathbf{M}^{-1}$ in the direction of $\mathbf{X}$ is given by $DG(\mathbf{M}) \mathbf{X} = -\mathbf{M}^{-1} \mathbf{X} \mathbf{M}^{-1}$. Now, (143) easily follows from the product rule applied to $F(\mathbf{Z}) = (\mathbf{I} + \mathbf{Z})(\gamma \mathbf{I} + \mathbf{Z})^{-1} (\gamma \mathbf{I} + \mathbf{Z})^{-1}$. By the assumptions in (88), we have that $\mathbf{Z}_t, (\gamma \mathbf{I} + \mathbf{Z}_t)^{-1}$ are uniformly bounded, hence the map DF is uniformly bounded as well. This implies that

$$\|F(\mathbf{Z}_t + \mathbf{X}_t) - F(\mathbf{Z}_t)\|_{op} \leq C \|\mathbf{X}_t\|_{op}.$$

As $\|F^{(1)}(0) \mathbf{X}_t\|_{op} \leq C \|\mathbf{X}_t\|_{op}$, we readily obtain (142).

Now we consider the case where $\|\mathbf{Z}_t\|_{op} \leq \frac{\gamma}{3}$. First note that, by (88), $\|\mathbf{X}_t\|_{op} \leq \frac{\gamma}{3}$. Hence,

$$F(\mathbf{Z}_t + \mathbf{X}_t) = \sum_{\ell=0}^{\infty} F^{(\ell)}(0) \frac{(\mathbf{Z}_t + \mathbf{X}_t)^\ell}{\ell!}.$$

The series above converges absolutely since $F^{(\ell)}(0)$ scales as $\frac{\ell!}{\gamma^\ell} \text{poly}(\ell)$. To see this, first we note that, if $h(x) = \frac{1}{(\gamma+x)^2}$, then $h^{(\ell)}(0) = (-1)^\ell (\ell+1)! \frac{1}{\gamma^{\ell+2}}$. Thus, by the product rule, $F^{(\ell)}(0) = (-1)^\ell (\ell+1)! \frac{1}{\gamma^{\ell+2}} + (-1)^{\ell-1} \ell! \frac{1}{\gamma^{\ell+1}}$ which has the desired asymptotic behaviour. Expanding the brackets and applying the triangle inequality yields

$$\left\| F(\mathbf{Z}_t + \mathbf{X}_t) - \sum_{\ell=0}^{\infty} F^{(\ell)}(0) \frac{\mathbf{Z}_t^\ell}{\ell!} - F^{(1)}(0) \mathbf{X}_t \right\|_{op} \leq \sum_{\ell=2}^{\infty} F^{(\ell)}(0) \frac{\|\mathbf{X}_t\|_{op}^\ell}{\ell!} + \sum_{\ell=2}^{\infty} F^{(\ell)}(0) \frac{1}{\ell!} \sum_{i=1}^{\ell-1} \binom{\ell}{i} \|\mathbf{Z}_t\|_{op}^i \|\mathbf{X}_t\|_{op}^{\ell-i}.$$

As $\|\mathbf{Z}_t\|_{op}, \|\mathbf{X}_t\|_{op} \leq \frac{\gamma}{3}$, we have

$$\sum_{\ell=2}^{\infty} F^{(\ell)}(0) \frac{\|\mathbf{X}_t\|_{op}^\ell}{\ell!} \leq \|\mathbf{X}_t\|_{op}^2 \sum_{\ell=2}^{\infty} F^{(\ell)}(0) \left(\frac{\gamma}{3}\right)^{\ell-2} \frac{1}{\ell!} \leq C \|\mathbf{X}_t\|_{op}^2,$$

and

$$\sum_{\ell=2}^{\infty} F^{(\ell)}(0) \frac{1}{\ell!} \sum_{i=1}^{\ell-1} \binom{\ell}{i} \|\mathbf{Z}_t\|_{op}^i \|\mathbf{X}_t\|_{op}^{\ell-i} \leq \sum_{\ell=2}^{\infty} F^{(\ell)}(0) \frac{1}{\ell!} 2^\ell \left(\frac{\gamma}{3}\right)^{\ell-2} \|\mathbf{Z}_t\|_{op} \|\mathbf{X}_t\|_{op} \leq C \|\mathbf{Z}_t\|_{op} \|\mathbf{X}_t\|_{op}.$$

By combining the last three expressions and using that

$$F(\mathbf{Z}_t) = \sum_{\ell=0}^{\infty} F^{(\ell)}(0) \frac{\mathbf{Z}_t^\ell}{\ell!},$$

we obtain

$$\left\| F(\mathbf{X}_t + \mathbf{Z}_t) - F(\mathbf{Z}_t) - F^{(1)}(0) \mathbf{X}_t \right\|_{op} \leq C \left(\|\mathbf{X}_t\|_{op} \|\mathbf{Z}_t\|_{op} + \|\mathbf{X}_t\|_{op}^2 \right). \quad (144)$$

As the map DF is uniformly bounded, we have

$$\|F(\mathbf{Z}_t) - F(0)\mathbf{I}\|_{op} \leq C \|\mathbf{Z}_t\|_{op}. \quad (145)$$

By combining (144), (145) and (141), we obtain

$$\frac{1}{2} \nabla_{\mathbf{B}_t} \mathbf{B}_t^\top = -\alpha F(\mathbf{Z}_t) + \alpha \text{Diag}(F(\mathbf{Z}_t)) (\mathbf{I} + \mathbf{Z}_t) - \alpha F^{(1)}(0) \mathbf{X}_t + \alpha \text{Diag}(\mathbf{X}_t F^{(1)}(0)) + \alpha \mathbf{X}_t F(0) + \mathbf{E}^t. \quad (146)$$

Using that $F(0) = \frac{1}{\gamma^2}$ and $F^{(1)}(0) = \frac{1}{\gamma^2} (1 - \frac{2}{\gamma})$, we finally obtain

$$\frac{1}{2} \nabla_{\mathbf{B}_t} \mathbf{B}_t^\top = -\alpha F(\mathbf{Z}_t) + \alpha \text{Diag}(F(\mathbf{Z}_t)) (\mathbf{I} + \mathbf{Z}_t) + \frac{2\alpha}{\gamma^3} \mathbf{X}_t^O + \frac{\alpha}{\gamma^2} \mathbf{X}_t^D + \mathbf{E}^t, \quad (147)$$

which concludes the proof. $\square$

Now let us return to the update equation of $\mathbf{B}_t \mathbf{B}_t^\top$ during the gradient step

$$\mathbf{B}_t' \mathbf{B}_t'^\top = (\mathbf{B}_t - \eta \nabla_{\mathbf{B}_t}) (\mathbf{B}_t - \eta \nabla_{\mathbf{B}_t})^\top = \mathbf{B}_t \mathbf{B}_t^\top - \eta \cdot \nabla_{\mathbf{B}_t} \mathbf{B}_t^\top - \eta \cdot \mathbf{B}_t (\nabla_{\mathbf{B}_t})^\top + \eta^2 \cdot \nabla_{\mathbf{B}_t} (\nabla_{\mathbf{B}_t})^\top. \quad (148)$$

Note that we can control the terms $\mathbf{B}_t (\nabla_{\mathbf{B}_t})^\top$ and $\nabla_{\mathbf{B}_t} \mathbf{B}_t^\top$ via Lemma E.4. In this view, it remains to argue that the contribution of the term $\eta^2 \cdot \nabla_{\mathbf{B}_t} (\nabla_{\mathbf{B}_t})^\top$ and of the projection step are of order $\eta \|\mathbf{E}^t\|_{op}$. For convenience of the upcoming lemmas we define the following quantity:

$$\tilde{\nabla}_{\mathbf{B}_t} := 2 \left(-\alpha (\alpha \mathbf{I} + \mathbf{B}_t \mathbf{B}_t^\top)^{-2} \mathbf{B}_t + \alpha \text{Diag} \left((\alpha \mathbf{I} + \mathbf{B}_t \mathbf{B}_t^\top)^{-2} (\mathbf{B}_t \mathbf{B}_t^\top) \right) \mathbf{B}_t \right). \quad (149)$$

Lemma E.5. Assume that (88) holds, and let $\nabla_{\mathbf{B}_t}$ be given by (86) with $\eta \leq C/\sqrt{d}$. Then, for all sufficiently large n , with probability $1 - 1/d^2$, jointly for all $t \geq 0$:

$$\eta^2 \|\nabla_{\mathbf{B}_t} (\nabla_{\mathbf{B}_t})^\top\|_{op} \leq \eta \|\mathbf{E}^t\|_{op}.$$

Proof of Lemma E.5. We start by showing that

$$\|\tilde{\nabla}_{\mathbf{B}_t}\|_{op} \leq C(\|\mathbf{X}_t\|_{op} + \|\mathbf{Z}_t\|_{op}). \quad (150)$$

Recall that $\|\mathbf{B}_t\|_{op}, \|(\alpha\mathbf{I} + \mathbf{B}_t\mathbf{B}_t^\top)^{-2}\|_{op} \leq C$. Hence, the following chain of inequalities holds

$$\begin{aligned} \|\tilde{\nabla}_{\mathbf{B}_t}\|_{op} &\leq \|\mathbf{B}_t\|_{op} \cdot \left\| -\alpha(\alpha\mathbf{I} + \mathbf{B}_t\mathbf{B}_t^\top)^{-2} + \alpha \text{Diag} \left((\alpha\mathbf{I} + \mathbf{B}_t\mathbf{B}_t^\top)^{-2} (\mathbf{B}_t\mathbf{B}_t^\top) \right) \right\|_{op} \\ &\leq C \left\| -\alpha(\alpha\mathbf{I} + \mathbf{B}_t\mathbf{B}_t^\top)^{-2} (\mathbf{I} - \mathbf{B}_t\mathbf{B}_t^\top + \mathbf{B}_t\mathbf{B}_t^\top) + \alpha \text{Diag} \left((\alpha\mathbf{I} + \mathbf{B}_t\mathbf{B}_t^\top)^{-2} (\mathbf{B}_t\mathbf{B}_t^\top) \right) \right\|_{op} \\ &\leq C \left(\left\| (\alpha\mathbf{I} + \mathbf{B}_t\mathbf{B}_t^\top)^{-2} (\mathbf{Z}_t + \mathbf{X}_t) \right\|_{op} \right. \\ &\quad \left. + \left\| (\alpha\mathbf{I} + \mathbf{B}_t\mathbf{B}_t^\top)^{-2} \mathbf{B}_t\mathbf{B}_t^\top - \text{Diag} \left((\alpha\mathbf{I} + \mathbf{B}_t\mathbf{B}_t^\top)^{-2} (\mathbf{B}_t\mathbf{B}_t^\top) \right) \right\|_{op} \right) \\ &\leq C(\|\mathbf{X}_t\|_{op} + \|\mathbf{Z}_t\|_{op} + \|F(\mathbf{X}_t + \mathbf{Z}_t) - \text{Diag}(F(\mathbf{X}_t + \mathbf{Z}_t))\|_{op}), \end{aligned} \quad (151)$$

where we recall the definition $F(x) = \frac{1+x}{(\gamma+x)^2}$, with $\gamma = 1 + \alpha$. By combining (144) and (145) (in the proof of Lemma E.4), we have

$$\|F(\mathbf{X}_t + \mathbf{Z}_t) - F(0)\mathbf{I}\|_{op} \leq C(\|\mathbf{X}_t\|_{op} + \|\mathbf{Z}_t\|_{op}),$$

As $\|\text{Diag}(\mathbf{M})\|_{op} \leq C\|\mathbf{M}\|_{op}$ for any matrix $\mathbf{M}$, we also have that

$$\|\text{Diag}(F(\mathbf{X}_t + \mathbf{Z}_t)) - F(0)\mathbf{I}\|_{op} \leq C(\|\mathbf{X}_t\|_{op} + \|\mathbf{Z}_t\|_{op}).$$

Hence,

$$\|F(\mathbf{X}_t + \mathbf{Z}_t) - \text{Diag}(F(\mathbf{X}_t + \mathbf{Z}_t))\|_{op} \leq C(\|\mathbf{X}_t\|_{op} + \|\mathbf{Z}_t\|_{op}),$$

which finishes the proof of (150).

At this point, recall from (124) and (149) that

$$\|\nabla_{\mathbf{B}_t} - \tilde{\nabla}_{\mathbf{B}_t}\|_{op} \leq \|\mathbf{E}^t\|_{op}. \quad (152)$$

Thus,

$$\|\nabla_{\mathbf{B}_t} \nabla_{\mathbf{B}_t}^\top\|_{op} \leq 2 \left\| \tilde{\nabla}_{\mathbf{B}_t} \mathbf{E}^t \right\|_{op} + \left\| \tilde{\nabla}_{\mathbf{B}_t} (\tilde{\nabla}_{\mathbf{B}_t})^\top \right\|_{op} + \|(\mathbf{E}^t)^2\|_{op}.$$

Recalling the previous bound on $\|\tilde{\nabla}_{\mathbf{B}_t}\|_{op}$ in (150) and using the assumptions in (88), we get that

$$\left\| \tilde{\nabla}_{\mathbf{B}_t} \mathbf{E}^t \right\|_{op}, \|\mathbf{E}^t\|_{op}^2 \leq C\|\mathbf{E}^t\|_{op},$$

and

$$\begin{aligned} \eta^2 \|\tilde{\nabla}_{\mathbf{B}_t}\|_{op}^2 &\leq C\eta(\|\mathbf{X}_t\|_{op}^2 + \|\mathbf{X}_t\|_{op} \|\mathbf{Z}_t\|_{op}) + C\eta^2 \|\mathbf{Z}_t\|_{op}^2 \\ &\leq C\eta \left(\frac{1}{\sqrt{d}} \|\mathbf{Z}_t\|_{op} + \|\mathbf{X}_t\|_{op}^2 + \|\mathbf{X}_t\|_{op} \|\mathbf{Z}_t\|_{op} \right) \leq C\eta \|\mathbf{E}^t\|_{op}, \end{aligned} \quad (153)$$

where we have also used that $\eta \leq C/\sqrt{d}$. This concludes the proof. $\square$

The next lemma controls the contribution of the projection step.

Lemma E.6 (Projection step). *Assume that (88) holds and $\eta \leq C/\sqrt{d}$. Then, for all sufficiently large n , with probability $1 - 1/d^2$, jointly for all $t \geq 0$:*

$$\|\text{proj}(\mathbf{B}'_t) - \mathbf{B}'_t\|_{op} \leq \eta \|\mathbf{E}^t\|_{op},$$

which implies that, by differentiability of the bilinear form,

$$\|\text{proj}(\mathbf{B}'_t)\text{proj}(\mathbf{B}'_t)^\top - \mathbf{B}'_t(\mathbf{B}'_t)^\top\|_{op} \leq \eta \|\mathbf{E}^t\|_{op}.$$

Proof of Lemma E.6. Recall that the objective (84) does not depend on the norm of $\{\mathbf{b}_i\}_{i=1}^n$, hence $(\nabla_{\mathbf{B}_t})_{i,:}$ is orthogonal to $(\mathbf{B}_t)_{i,:}$, which implies that

$$\text{proj}_i(\mathbf{B}'_t) = \frac{(\mathbf{B}_t)_{i,:} - \eta(\nabla_{\mathbf{B}_t})_{i,:}}{\sqrt{1 + \eta^2\|(\nabla_{\mathbf{B}_t})_{i,:}\|^2}}.$$

Let us define

$$\mathbf{D}_t := \text{Diag}\left(\frac{1}{\sqrt{1 + \eta^2\|(\nabla_{\mathbf{B}_t})_{1,:}\|^2}}, \dots, \frac{1}{\sqrt{1 + \eta^2\|(\nabla_{\mathbf{B}_t})_{n,:}\|^2}}\right).$$

Then, we obtain the following compact form:

$$\text{proj}(\mathbf{B}'_t) = \mathbf{D}_t(\mathbf{B}_t - \eta\nabla_{\mathbf{B}_t}) = \mathbf{D}_t\mathbf{B}'_t.$$

In this view, it remains to bound $\|\mathbf{D}_t - \mathbf{I}\|_{op}$. In more details, by (150) and (152), we have

$$\|\nabla_{\mathbf{B}_t}\|_{op} \leq \|\tilde{\nabla}_{\mathbf{B}_t}\|_{op} + \|\mathbf{E}^t\|_{op} \leq C(\|\mathbf{X}_t\|_{op} + \|\mathbf{Z}_t\|_{op} + \|\mathbf{E}^t\|_{op}) \leq C',$$

where $C' > 0$ is a universal constant (independent of C_X, n, d). Hence, by recalling that $\|\mathbf{B}_t\|_{op} \leq C$ by assumption (88), we have

$$\|\text{proj}(\mathbf{B}'_t) - \mathbf{B}'_t\|_{op} = \|(\mathbf{D}_t - \mathbf{I})(\mathbf{B}_t - \eta\nabla_{\mathbf{B}_t})\|_{op} \leq C\|\mathbf{D}_t - \mathbf{I}\|_{op}.$$

Note that function $1/\sqrt{1+x}$ is differentiable at 0, hence, we have that for small enough η (which follows from $\eta \leq C/\sqrt{d}$):

$$\left| \frac{1}{\sqrt{1 + \eta^2\|(\nabla_{\mathbf{B}_t})_{i,:}\|^2}} - 1 \right| \leq C\eta^2\|(\nabla_{\mathbf{B}_t})_{i,:}\|^2.$$

In this view, we have

$$\|\mathbf{D}_t - \mathbf{I}\|_{op} \leq C\eta^2\|\nabla_{\mathbf{B}_t}\|_{op}^2 \leq C\eta^2\|\tilde{\nabla}_{\mathbf{B}_t}\|_{op}^2 + C\eta^2\|\tilde{\nabla}_{\mathbf{B}_t}\|_{op}\|\mathbf{E}^t\|_{op} + C\eta^2\|\mathbf{E}^t\|_{op}^2.$$

Inspecting each term one by one and applying (150) in conjunction with $\eta \leq C/\sqrt{d}$ gives that

$$\begin{aligned} \eta^2\|\mathbf{E}^t\|_{op}^2 &\leq C\eta\|\mathbf{E}^t\|_{op}, \\ \eta^2\|\tilde{\nabla}_{\mathbf{B}_t}\|_{op}\|\mathbf{E}^t\|_{op} &\leq C\eta\|\mathbf{E}^t\|_{op}, \\ \eta^2\|\tilde{\nabla}_{\mathbf{B}_t}\|_{op}^2 &\leq C\eta\|\mathbf{E}^t\|_{op}, \end{aligned}$$

where in the last step we have used (153). This concludes the proof. $\square$

In this view, using (148) and Lemmas E.4, E.5 and E.6, we obtain

$$\begin{aligned} \mathbf{I} + \mathbf{Z}_{t+1} + \mathbf{X}_{t+1} &= \mathbf{B}_{t+1}\mathbf{B}_{t+1}^\top = \mathbf{I} + \mathbf{Z}_t + \mathbf{X}_t + 4\eta\alpha F(\mathbf{Z}_t) - 2\eta\alpha \text{Diag}(F(\mathbf{Z}_t))(\mathbf{I} + \mathbf{Z}_t) \\ &\quad - 2\eta\alpha(\mathbf{I} + \mathbf{Z}_t)\text{Diag}(F(\mathbf{Z}_t)) - \frac{8\alpha\eta}{\gamma^3}\mathbf{X}_t^O - \frac{4\alpha\eta}{\gamma^2}\mathbf{X}_t^D + \eta\mathbf{E}^t. \end{aligned} \quad (154)$$

Furthermore, we have that

$$\begin{aligned} \text{Diag}(F(\mathbf{Z}_t))(\mathbf{I} + \mathbf{Z}_t) &= (\text{Diag}(F(\mathbf{Z}_t) - F(0)\mathbf{I}) + F(0)\mathbf{I})(\mathbf{I} + \mathbf{Z}_t) \\ &= \frac{1}{\gamma^2}(\mathbf{I} + \mathbf{Z}_t) + (\text{Diag}(F(\mathbf{Z}_t) - F(0)\mathbf{I}))(\mathbf{I} + \mathbf{Z}_t) \\ &= \frac{1}{\gamma^2}(\mathbf{I} + \mathbf{Z}_t) + \left(\frac{1}{n}\text{Tr}[F(\mathbf{Z}_t) - F(0)\mathbf{I}] + \mathbf{D}'_t\right)(\mathbf{I} + \mathbf{Z}_t), \end{aligned} \quad (155)$$

where $\mathbf{D}'_t$ is a diagonal matrix such that, with probability at least $1 - 1/d^2$, its entries are upper bounded in modulus by $\frac{C \log d}{\sqrt{d}} \|\mathbf{Z}_t\|_{op}^{1/2}$. The last passage follows from Lemma G.2. Note that $\frac{1}{\gamma^2}(\mathbf{I} + \mathbf{Z}_t) = \frac{1}{n} \text{Tr}[F(0)\mathbf{I}]$ and recall that $\|\mathbf{Z}_t\|_{op} \leq C$. Hence, (155) implies that

$$\text{Diag}(F(\mathbf{Z}_t))(\mathbf{I} + \mathbf{Z}_t) = \frac{1}{n} \text{Tr}[F(\mathbf{Z}_t)](\mathbf{I} + \mathbf{Z}_t) + \mathbf{E}^t. \quad (156)$$

Similarly, we have that

$$(\mathbf{I} + \mathbf{Z}_t)\text{Diag}(F(\mathbf{Z}_t)) = \frac{1}{n} \text{Tr}[F(\mathbf{Z}_t)](\mathbf{I} + \mathbf{Z}_t) + \mathbf{E}^t. \quad (157)$$

By combining (156)-(157) with (154) and using that $\mathbf{X}_t = \mathbf{X}_t^O + \mathbf{X}_t^D$, we get

$$\begin{aligned} \mathbf{Z}_{t+1} + \mathbf{X}_{t+1} &= \left(1 - \frac{8\alpha}{\gamma^3}\eta\right) \mathbf{X}_t^O + \left(1 - \frac{4\alpha}{\gamma^2}\eta\right) \mathbf{X}_t^D + \mathbf{Z}_t + 4\eta\alpha F(\mathbf{Z}_t) \\ &\quad - 4\eta\alpha \frac{1}{n} \text{Tr}[F(\mathbf{Z}_t)](\mathbf{I} + \mathbf{Z}_t) + \eta \mathbf{E}^t. \end{aligned} \quad (158)$$

Hence, we can write the following system capturing the dynamics of the spectrum $\mathbf{Z}_t$ and of the errors $(\mathbf{X}_t^O, \mathbf{X}_t^D)$

$$\mathbf{Z}_{t+1} = \mathbf{Z}_t + 4\eta\alpha F(\mathbf{Z}_t) - 4\eta\alpha \frac{1}{n} \text{Tr}[F(\mathbf{Z}_t)](\mathbf{I} + \mathbf{Z}_t), \quad (159)$$

$$\mathbf{X}_{t+1}^D = \left(1 - \frac{4\alpha}{\gamma^2}\eta\right) \mathbf{X}_t^D + \eta \mathbf{E}^t, \quad (160)$$

$$\mathbf{X}_{t+1}^O = \left(1 - \frac{8\alpha}{\gamma^3}\eta\right) \mathbf{X}_t^O + \eta \mathbf{E}^t. \quad (161)$$

Here, the operator norm of $\mathbf{E}^t$ is upper bounded as in (90), where we recall that the constant C is uniformly bounded in t .

In the view of (159), one can readily see that the updates on the spectrum of $\mathbf{Z}_t$ follow the one described in Lemma G.3 and, thus, converges exponentially. This means that the set of assumptions on $\mathbf{Z}_t$ in (88) is satisfied by suitably picking C .

Now it only remains to take care of $\mathbf{X}_t$. If we write $x_t^D = \|\mathbf{X}_t^D\|_{op}$, $x_t^O = \|\mathbf{X}_t^O\|_{op}$, $z_t = \|\mathbf{Z}_t\|_{op}^{1/2}$, then recalling the definition of $\mathbf{E}_t$ in (90), (160), (161) we have that

$$x_{t+1}^D \leq \left(1 - \frac{4\alpha}{\gamma^2}\eta\right) x_t^D + \eta C_D \left(\frac{\text{poly}(\log d)}{\sqrt{d}} \cdot z_t + (x_t^D + x_t^O)^2 + (x_t^D + x_t^O)z_t \right) \quad (162)$$

$$x_{t+1}^O \leq \left(1 - \frac{8\alpha}{\gamma^3}\eta\right) x_t^O + \eta C_O \left(\frac{\text{poly}(\log d)}{\sqrt{d}} \cdot z_t + (x_t^D + x_t^O)^2 + (x_t^D + x_t^O)z_t \right). \quad (163)$$

Since both of these recursive bounds are monotone in x_t^D, x_t^O , we can dominate them as follows. If we recursively define x_t by

$$x_{t+1} = \left(1 - \eta \min \left\{ \frac{4\alpha}{\gamma^2}, \frac{8\alpha}{\gamma^3} \right\} \right) x_t + \eta \max\{C_D, C_O\} \left(\frac{\text{poly}(\log d)}{\sqrt{d}} \cdot z_t + (x_t + x_t)^2 + (x_t + x_t)z_t \right), \quad (164)$$

then by monotonicity $\max\{x_t^D, x_t^O\} \leq x_t$. Thus, we only need to analyse the recursion (164), which we do in the following lemma. Note that the condition $z_t \leq C e^{-c_t \eta}$ required by Lemma E.7 holds by (88).

Lemma E.7 (Error decay). *Let $\{z_t\}_{t=0}^\infty$ be a non-negative exponentially decaying sequence, i.e., $z_t \leq C_z e^{-\eta c_z t}$, and consider a non-negative sequence $\{x_t\}_{t=0}^\infty$ such that at each time-step t the following condition holds for $\eta = \Theta(1/\sqrt{d})$ and sufficiently large d :*

$$x_{t+1} = (1 - \eta c_1)x_t + \eta C_2 \cdot z_t \cdot x_t + \eta C_3 x_t^2 + \eta C_4 \cdot \frac{\text{poly}(\log d)}{\sqrt{d}} \cdot z_t, \quad (165)$$

with $x_0 = 0$. Then, the following holds

$$x_t \leq C \frac{\text{poly}(\log d)}{\sqrt{d}} \cdot T e^{-cT}, \quad (166)$$

where $T = t\eta$.

Proof of Lemma E.7. We proceed in two parts. In the first part, we show that our recursion does not blow up in $t = K/\eta$ steps. In the second part, $z_t \leq C_z \exp(-c_z K)$ will be small, which allows us to deduce (166).

Error does not blow up in finite time. Let $t = K/\eta$ where K is such that $K/\eta \in \mathbb{N}$. We start by analysing the simpler recursion

$$x_{t+1} = (1 - \eta c_1)x_t + \eta C_2 \cdot z_t \cdot x_t + \eta C_4 \cdot \frac{\text{poly}(\log d)}{\sqrt{d}} \cdot z_t.$$

By hypothesis, $z_t \leq C_z$. Hence, we arrive to

$$x_{t+1} = (1 - \eta c_1)x_t + \eta C_2 C_z \cdot x_t + \eta C_4 C_z \frac{\text{poly}(\log d)}{\sqrt{d}}.$$

Writing $C_5 = C_2 C_z - c_1$, unrolling the recursion on the RHS and using $x_0 = 0$ gives

$$\begin{aligned} x_{t+1} &= \eta C_4 C_z \frac{\text{poly}(\log d)}{\sqrt{d}} \sum_{j=0}^t (1 + \eta C_5)^j \\ &\leq \eta C_4 C_z \frac{\text{poly}(\log d)}{\sqrt{d}} \sum_{j=0}^{K/\eta} e^{\eta C_5 j} \\ &= \eta C_4 C_z \frac{\text{poly}(\log d)}{\sqrt{d}} \cdot e^{C_5 K} \sum_{j=0}^{K/\eta} e^{-C_5 \eta(t-j)} \\ &\leq \eta C_4 C_z \frac{\text{poly}(\log d)}{\sqrt{d}} \cdot \frac{e^{C_5 K}}{1 - e^{-\eta C_5}}, \end{aligned}$$

where the inequality holds for $t \leq K/\eta$ and we have used $1 + x \leq e^x$. For small enough η , we have that

$$\frac{\eta}{1 - e^{-\eta C_5}} \leq \frac{2}{C_5},$$

hence, for all $t \leq K/\eta$,

$$x_{t+1} \leq 2 \frac{\text{poly}(\log d)}{\sqrt{d}} \frac{C_4 C_z}{C_5} \exp(C_5 K). \quad (167)$$

Let us now go back to our original recursion (165), which contains the term x_t^2 . We claim that this recursion satisfies a bound like (167). Assume by contradiction that it exceeds the bound

$$x_t \leq 4 \frac{\text{poly}(\log d)}{\sqrt{d}} \frac{C_4 C_z}{C_5} \exp(C_5 K) \quad (168)$$

for the first time at step t' . Then, for all $t < t'$, (168) holds. Noting that $x_t^2 \leq 4 \frac{\text{poly}(\log d)}{\sqrt{d}} \frac{C_4 C_z}{C_5} \exp(C_5 K) x_t$ we define $C'_5 = C_2 C_z + 4 C_3 \frac{\text{poly}(\log d)}{\sqrt{d}} \frac{C_4 C_z}{C_5} \exp(C_5 K) - c_1$. By unrolling the recursion exactly as before, we obtain

$$x_{t+1} \leq 2 \frac{\text{poly}(\log d)}{\sqrt{d}} \frac{C_4 C_z}{C'_5} \exp(C'_5 K) \leq 3 \frac{\text{poly}(\log d)}{\sqrt{d}} \frac{C_4 C_z}{C_5} \exp(C_5 K), \quad (169)$$

for d large enough. Here, the second inequality follows for large d , since it is clear from the definitions that $|C_5 - C'_5|$ vanishes for large d . This shows that we cannot violate (168), thus (169) holds for all $t \leq K/\eta$.

Convergence of errors x_t to zero. We now choose K large enough so that

$$z_t = C_z e^{-\eta c_z t} < \frac{c_1}{2C_2}, \quad \forall t \geq K/\eta.$$

Hence, the term corresponding to $\eta C_2 z_t x_t$ can be pushed inside the $(1 - \eta c_1)x_t$ term. Consequently, we can equivalently study the following dynamics

$$x_{t+1} = (1 - \eta c'_1)x_t + \eta C_3 x_t^2 + \eta C_4 C_z \frac{\text{poly}(\log d)}{\sqrt{d}} e^{-\eta c_z t}, \quad (170)$$

where $c'_1 = c_1/2$. Here, we initialize again at $t = 0$, but now starting at

$$x_0 = C_6 \frac{\text{poly}(\log d)}{\sqrt{d}},$$

where $C_6 = 4 \frac{C_4 C_z}{C_5} \exp(C_5 K)$, corresponding to the bound in (168). Rearranging we have

$$x_{t+1} = x_t + \eta \left(-c'_1 x_t + C_3 x_t^2 + C_4 C_z \frac{\text{poly}(\log d)}{\sqrt{d}} e^{-\eta c_z t} \right). \quad (171)$$

As the last term inside the brackets vanishes when $d \rightarrow \infty$, we have two roots of the polynomial inside the brackets, corresponding to the fixed points of the iteration. The left root r_l scales as

$$r_l \leq C_l \frac{\text{poly}(\log d)}{\sqrt{d}} e^{-\eta c_z t},$$

and the right root r_r as

$$r_r \geq \frac{c'_1}{C_3} - C \frac{\text{poly}(\log d)}{\sqrt{d}} e^{-\eta c_z t}.$$

In addition, it is easy to see that both roots are non-negative.

Next, we prove that $x_t \leq C \frac{\text{poly}(\log d)}{\sqrt{d}}$ for all t . We will show this by contradiction. At initialization we have

$$x_0 = C_6 \frac{\text{poly}(\log d)}{\sqrt{d}}.$$

Choose A, B as follows:

$$A := \max\{C_l, C_6\}, \quad B = C_7 A.$$

We first note that, for small enough η and large enough d , we can choose C_7 such that $x_t \leq A \frac{\text{poly}(\log d)}{\sqrt{d}}$ implies $x_{t+1} \leq B \frac{\text{poly}(\log d)}{\sqrt{d}}$. We now show that $x_t \leq B \frac{\text{poly}(\log d)}{\sqrt{d}}$ for all t . To do so, assume by contradiction that $x_{t+1} > B \frac{\text{poly}(\log d)}{\sqrt{d}}$. Then $x_t \in [A \frac{\text{poly}(\log d)}{\sqrt{d}}, B \frac{\text{poly}(\log d)}{\sqrt{d}}] \subseteq [r_l, r_r]$, thus

$$-c'_1 x_t + C_3 x_t^2 + C_4 C_z \frac{\text{poly}(\log d)}{\sqrt{d}} e^{-\eta c_z t} < 0.$$

Hence, from (171) it follows that

$$x_{t+1} \leq x_t \leq B \frac{\text{poly}(\log d)}{\sqrt{d}},$$

which gives us the desired contradiction.

Thus, for all t ,

$$x_t^2 \leq B \frac{\text{poly}(\log d)}{\sqrt{d}} x_t.$$

This allows us to push the second term in (170) into the first one (for d large enough), which reduces the recursion to

$$x_{t+1} = (1 - \eta c''_1) x_t + \eta C_4 C_z \frac{\text{poly}(\log d)}{\sqrt{d}} e^{-\eta c_z t},$$

where $c_1'' \geq c_1'/2$. By unrolling this last recursion and using $x_0 = C_6 \frac{\text{poly}(\log d)}{\sqrt{d}}$, we have that, for $t \geq 1$,

$$x_t = C_6 \frac{\text{poly}(\log d)}{\sqrt{d}} (1 - \eta c_1'')^t + \eta C_4 C_z \frac{\text{poly}(\log d)}{\sqrt{d}} \sum_{\ell=1}^t (1 - \eta c_1'')^{t-\ell} e^{-\eta c_z \ell} \quad (172)$$

$$\leq C_6 \frac{\text{poly}(\log d)}{\sqrt{d}} \exp(-\eta c_1'' t) + \eta C_4 C_z \frac{\text{poly}(\log d)}{\sqrt{d}} \sum_{\ell=1}^t e^{-\eta(c_z \ell + c_1''(t-\ell))}, \quad (173)$$

where the inequality follows from $1 - x \leq e^{-x}$. Since the term in the exponents of the sum is a linear function in ℓ , its maximum value is attained in the endpoints. Thus,

$$x_t \leq C_6 \frac{\text{poly}(\log d)}{\sqrt{d}} \exp(-\eta c_1'' t) + \eta C_4 C_z \frac{\text{poly}(\log d)}{\sqrt{d}} t \max\{e^{-\eta c_z t}, e^{-\eta c_1'' t}\},$$

which implies (166). $\square$

By Lemma E.7 we know that

$$\|\mathbf{X}_t\|_{op} \leq \frac{C}{\sqrt{d}} \cdot T e^{-cT},$$

where C is independent of C_X by definition. Hence, we can pick C_X such that, for sufficiently large d , the assumptions on $\mathbf{X}_t$ in (88) are satisfied. With this in mind, we can use Lemma G.3 to bound the dynamics involving $\mathbf{Z}_t$ and Lemma E.7 to claim that the error $\mathbf{X}_t$ vanishes at least geometrically fast. This concludes the proof of Theorem E.1.

F. Discussion of Isotropic Gaussian Results

Degenerate isotropic Gaussian data. All the arguments of Section 4.1 directly apply for $\mathbf{x} \sim \mathcal{N}(\mathbf{0}, \sigma^2 \mathbf{I})$, the only differences being the scaling of the term $\text{Tr}[\mathbf{B}\mathbf{A}]$ (which is additionally multiplied by σ) and the constant variance term σ^2 (in place of 1) in (6). Our results can be also easily extended to the case of degenerate isotropic Gaussian data, i.e., $\mathbf{x} \sim \mathcal{N}(\mathbf{0}, \mathbf{\Sigma})$ with $\lambda_i(\mathbf{\Sigma}) = \sigma^2$ for $i \leq d - k$ and $\lambda_i(\mathbf{\Sigma}) = 0$ for $i > d - k$, where $\lambda_i(\mathbf{\Sigma})$ stands for the i -th eigenvalue of $\mathbf{\Sigma}$ in non-increasing order. In fact, by the rotational invariance of the Gaussian distribution, we can assume without loss of generality that $\mathbf{x} = [x_1, \dots, x_{d-k}, 0, \dots, 0]$, where $(x_i) \sim_{\text{i.i.d.}} \mathcal{N}(0, \sigma^2)$. Hence, by considering $\mathbf{A} \in \mathbb{R}^{(d-k) \times n}$ and $\mathbf{B} \in \mathbb{R}^{n \times (d-k)}$ and substituting d with $d - k$ where suitable, analogous results follow.

Scaling of the learning rate. Theorem 4.6 is stated for $\eta = \Theta(1/\sqrt{d})$, as this corresponds to the biggest learning rate for which our argument works (thus requiring the least amount of steps for convergence). The same result can be proved for $\eta = \Theta(d^{-\kappa})$ with $\kappa \geq 1/2$. The only piece of the proof affected by this change is the third part of Lemma G.2 (in particular, the chain of inequalities (187)), which continues to hold as long as η is polynomial in d^{-1} .

Assumptions on compression rate r . We expect an analog of Theorem 4.5 (see Theorem D.1 for the formal statement) to hold for $r > 1$, as long as d is sufficiently large. In fact, for a fixed d , it appears to be difficult to even characterize the global minimizer: the choice (12) approaches the lower bound $\text{LB}_{r>1}(\mathbf{I})$ only as $d \rightarrow \infty$, see Proposition 4.4. We also expect Theorem 4.6 to hold for $r \geq 1$. Here, an additional challenge is that the minimizer has non-zero off-diagonal entries. In combination with the lack of an exact characterization of the minimizer, this leads to an additional error term that would be difficult to control with the current tools. At the same time, the restriction $r < 1$ is likely to be an artifact of the proof as experimentally (see, for instance, Figure 4) the algorithm still converges to the global optimum for $r \geq 1$.

Gaussian initialization in Theorem 4.6. The Gaussian initialization ensures that, with high probability, the off-diagonal entries of $\mathbf{B}(t)\mathbf{B}(t)^\top$ are small. This allows us to approximate higher-order Hadamard powers of $\mathbf{B}(t)\mathbf{B}(t)^\top$ with $\mathbf{I}$. However, in experiments the Gaussian assumption seems to be unnecessary, and we expect the convergence result to hold for all (non-degenerate) initializations.

G. Auxiliary Results

Lemma G.1. *For any $\mathbf{R} \in \mathbb{R}^{n \times n}$ the following holds*

$$\|\mathbf{R} - \text{diag}(\mathbf{R})\|_{op} \leq C \|\mathbf{R}\|_{op}.$$

Proof. By definition of the operator norm we have that

$$\|\mathbf{R}\|_{op} = \sup_{\|\mathbf{x}\|_2=1} \|\mathbf{R}\mathbf{x}\|_2.$$

Note that by Cauchy-Schwarz, the following holds for $\|\mathbf{y}\|_2 = 1$:

$$\langle \mathbf{y}, \mathbf{R}\mathbf{x} \rangle \leq \|\mathbf{R}\mathbf{x}\|_2,$$

and the inequality is met when $\mathbf{y}$ is aligned with $\mathbf{R}\mathbf{x}$. Hence, we get

$$\sup_{\|\mathbf{y}\|_2=1} \langle \mathbf{y}, \mathbf{R}\mathbf{x} \rangle = \|\mathbf{R}\mathbf{x}\|_2,$$

and, thus, the operator norm can be rewritten as

$$\|\mathbf{R}\|_{op} = \sup_{\|\mathbf{x}\|_2=1} \|\mathbf{R}\mathbf{x}\|_2 = \sup_{\|\mathbf{x}\|_2=\|\mathbf{y}\|_2=1} \langle \mathbf{y}, \mathbf{R}\mathbf{x} \rangle.$$

Note also that $\|\text{diag}(\mathbf{R})\|_{op}$ is equal to the maximal diagonal element (in absolute value). Hence, by letting $\mathbf{e}_i$ be the i -th element of the canonical basis, we get

$$\|\text{diag}(\mathbf{R})\|_{op} = \sup_i |\mathbf{R}_{i,i}| \leq \sup_i |\langle \mathbf{e}_i, \mathbf{R}\mathbf{e}_i \rangle| \leq \sup_{\|\mathbf{x}\|_2=\|\mathbf{y}\|_2=1} \langle \mathbf{y}, \mathbf{R}\mathbf{x} \rangle = \|\mathbf{R}\|_{op}.$$

In this view, an application of triangle inequality, i.e.,

$$\|\mathbf{R} - \text{diag}(\mathbf{R})\|_{op} \leq \|\mathbf{R}\|_{op} + \|\text{diag}(\mathbf{R})\|_{op} \leq 2 \|\mathbf{R}\|_{op},$$

finishes the proof. $\square$

Lemma G.2. *Consider the matrix $\mathbf{A}_t = \mathbf{U}\mathbf{\Lambda}_t\mathbf{U}^\top$, where the matrix $\mathbf{U}$ is distributed according to the Haar measure and it is independent from the diagonal matrix $\mathbf{\Lambda}_t$. Further, assume that all the diagonal entries of $\mathbf{\Lambda}_t$ are bounded in absolute value by a constant. Then, the following results hold.*

1. *We have that, with probability at least $1 - 1/d^2$,*

$$\max_{i \neq j} |(\mathbf{A}_t)_{i,j}| \leq c \sqrt{\frac{\log d}{d}}, \quad (174)$$

for some absolute constant $c > 0$.

2. *Let $\mathbf{D}_t = \text{diag}(\mathbf{A}_t)$. Then,*

$$\mathbf{D}_t = \alpha \mathbf{I} + \mathbf{D}'_t,$$

where

$$\alpha = \frac{1}{n} \text{Tr}(\mathbf{\Lambda}_t),$$

and $\mathbf{D}'_t$ is a diagonal matrix such that, with probability at least $1 - 1/d^2$,

$$\max_{i \in [n]} |(\mathbf{D}'_t)_{i,i}| \leq c \frac{\log d}{\sqrt{d}}. \quad (175)$$

3. Assume that, for all $t \in \mathbb{N}$,

$$\|\mathbf{\Lambda}_t\|_{op} \leq C e^{-c\eta t}, \quad (176)$$

where $c, C > 0$ are absolute constants and $\eta = \Theta(1/\sqrt{d})$. Then, with probability at least $1 - 1/d^2$,

$$\sup_{t \geq 0} \max_{i \neq j} |(\mathbf{A}_t)_{i,j}| \leq c \sqrt{\frac{\log d}{d}}, \quad (177)$$

$$\sup_{t \geq 0} \max_{i \in [n]} |(\mathbf{D}'_t)_{i,i}| \leq c \frac{\log d}{\sqrt{d}}. \quad (178)$$

Proof. We start by proving (174). Consider the metric measure space $(\mathbb{SO}(d), \|\cdot\|_F, \mathbb{P})$. Here, $\mathbb{SO}(d)$ denotes the special orthogonal group containing all $d \times d$ orthogonal matrices with determinant 1 (i.e., all rotation matrices), and $\mathbb{P}$ is the uniform probability measure on $\mathbb{SO}(d)$, i.e., the Haar measure. Given a diagonal matrix $\mathbf{\Lambda}_t$ and two indices $i, j \in [d]$, define $f : \mathbb{SO}(d) \rightarrow \mathbb{R}$ as

$$f(\mathbf{M}) = (\mathbf{M}\mathbf{\Lambda}_t\mathbf{M}^\top)_{i,j}. \quad (179)$$

Note that

$$\begin{aligned} |f(\mathbf{M}) - f(\mathbf{M}')| &= |(\mathbf{M}\mathbf{\Lambda}_t\mathbf{M}^\top)_{i,j} - (\mathbf{M}'\mathbf{\Lambda}_t(\mathbf{M}')^\top)_{i,j}| \\ &\leq |(\mathbf{M}\mathbf{\Lambda}_t\mathbf{M}^\top)_{i,j} - (\mathbf{M}'\mathbf{\Lambda}_t\mathbf{M}^\top)_{i,j}| + |(\mathbf{M}'\mathbf{\Lambda}_t\mathbf{M}^\top)_{i,j} - (\mathbf{M}'\mathbf{\Lambda}_t(\mathbf{M}')^\top)_{i,j}| \\ &\leq |((\mathbf{M} - \mathbf{M}')\mathbf{\Lambda}_t\mathbf{M}^\top)_{i,j}| + |(\mathbf{M}'\mathbf{\Lambda}_t(\mathbf{M} - \mathbf{M}')^\top)_{i,j}| \\ &\leq \|(\mathbf{M} - \mathbf{M}')\mathbf{\Lambda}_t\mathbf{M}^\top\|_F + \|\mathbf{M}'\mathbf{\Lambda}_t(\mathbf{M} - \mathbf{M}')^\top\|_F \\ &\leq 2\|\mathbf{M} - \mathbf{M}'\|_F \|\mathbf{\Lambda}_t\|_{op} \|\mathbf{M}\|_{op} \leq 2\|\mathbf{M} - \mathbf{M}'\|_F \|\mathbf{\Lambda}_t\|_{op}, \end{aligned} \quad (180)$$

where in the fourth inequality we use that, for any two matrices $\mathbf{A}$ and $\mathbf{B}$, $\|\mathbf{AB}\|_F \leq \|\mathbf{A}\|_{op} \|\mathbf{B}\|_F$, and in the fifth inequality we use that $\|\mathbf{M}\|_{op} = 1$ as $\mathbf{M} \in \mathbb{SO}(d)$. Hence, f has Lipschitz constant upper bounded by $2\|\mathbf{\Lambda}_t\|_{op}$ and an application of Theorem 5.2.7 of (Vershynin, 2018) gives that

$$\mathbb{P}(|f(\mathbf{U}) - \mathbb{E}[f(\mathbf{U})]| \geq u) \leq 2 \exp\left(-c_1 \frac{du^2}{2\|\mathbf{\Lambda}_t\|_{op}}\right), \quad (181)$$

where c_1 is a universal constant.

Let $\mathbf{u}_i$ denote the i -th row of $\mathbf{U}$. Then,

$$f(\mathbf{U}) = \langle \mathbf{u}_i, \mathbf{\Lambda}_t \mathbf{u}_j \rangle. \quad (182)$$

Suppose that $i \neq j$. Since $\mathbf{U}$ is distributed according to the Haar measure, $\mathbf{u}_i$ is uniform on the unit sphere and $\mathbf{u}_j$ is uniformly distributed on the unit sphere in the orthogonal complement of $\mathbf{u}_i$ (see Section 1.2 of (Meckes, 2019)). Thus, $(\mathbf{u}_i, \mathbf{u}_j)$ has the same distribution as $(-\mathbf{u}_i, \mathbf{u}_j)$, which implies that, whenever $i \neq j$

$$\mathbb{E}[f(\mathbf{U})] = 0. \quad (183)$$

By combining (181)-(183) with a union bound over i, j , we have that

$$\mathbb{P}(\max_{i \neq j} |(\mathbf{U}\mathbf{\Lambda}_t\mathbf{U}^\top)_{i,j}| \geq u) \leq 2d^2 \exp\left(-c_1 \frac{du^2}{2\|\mathbf{\Lambda}_t\|_{op}}\right). \quad (184)$$

As $\|\mathbf{\Lambda}_t\|_{op}$ is upper bounded by a universal constant, the result (174) readily follows.

For the second part, note that

$$(\mathbf{D}_t)_{i,i} = \langle \mathbf{u}_i, \mathbf{\Lambda}_t \mathbf{u}_i \rangle. \quad (185)$$

Furthermore, the following chain of equalities hold

$$\mathbb{E}[(\mathbf{D}_t)_{i,i}] = \frac{1}{n} \sum_{i=1}^n \mathbb{E}[(\mathbf{D}_t)_{i,i}] = \mathbb{E}\left[\frac{1}{n} \sum_{i=1}^n (\mathbf{D}_t)_{i,i}\right] = \frac{1}{n} \text{Tr}(\mathbf{D}_t), \quad (186)$$

where the first equality uses that the u_i 's have the same (marginal) distribution, and the last term does not contain an expectation since $\text{Tr}(\mathbf{D}_t) = \text{Tr}(\mathbf{A}_t) = \sum_{i=1}^d (\mathbf{\Lambda}_t)_{i,i}$, which does not depend on U . Therefore, by using (181) and by performing a union bound over $i \in [n]$, the result (175) follows.

For the third part, by performing a union bound over $t \geq 0$ in (184), we have that (177) holds with probability at least

$$\begin{aligned}
 2 \sum_{t=0}^{\infty} \exp\left(-c_1 \frac{du^2}{2\|\mathbf{\Lambda}_t\|_{op}}\right) &\leq 2 \sum_{t=0}^{\infty} \exp(-c_2 d u^2 e^{C\eta t}) \\
 &\leq 2 \sum_{t=0}^{\infty} \exp(-c_2 d u^2 e^{C\lfloor \eta t \rfloor}) \\
 &\leq 2 \left\lceil \frac{1}{\eta} \right\rceil \sum_{t=0}^{\infty} \exp(-c_2 d u^2 e^{Ct}) \\
 &\leq C\sqrt{d} \sum_{t=0}^{\infty} \exp(-c_2 d u^2 e^{Ct}),
 \end{aligned} \tag{187}$$

where the first inequality follows from (176) and the last one from $\eta = \Theta(1/\sqrt{d})$. Choosing $u = c \frac{\log d}{\sqrt{d}}$ we can get that $b := \exp(-c_2 d u^2) < 1$ and, hence, the following holds

$$\sum_{t=0}^{\infty} \exp(-c_2 d u^2) e^{Ct} \leq \sum_{t=0}^{\infty} \exp(-c_2 d u^2)^{Ct+1} = \frac{b}{1-b^C} \leq \frac{1}{d^3},$$

where the first inequality uses that $e^t \geq 1+t$ and the second inequality follows from the definition of b . This concludes the proof of (177). The proof of (178) uses an analogous union bound on $t \geq 0$. $\square$

Lemma G.3. Let $\lambda^0 = \{\lambda_1^0, \dots, \lambda_n^0\}$ be a set of numbers in $\mathbb{R}$ such that

$$\lambda_{min}^0 := \min_{i \in [n]} \lambda_i^0 \geq \delta > 0, \quad \lambda_{max}^0 := \max_{i \in [n]} \lambda_i^0 \leq M < +\infty, \quad \sum_{j=1}^n \lambda_j^0 = n.$$

Let the values $\{\lambda_i^t\}_{i=1}^n$ be updated according to the equation below

$$\lambda_i^{t+1} = \lambda_i^t + \eta \left(F(\lambda_i^t) - \lambda_i^t \cdot \frac{1}{n} \sum_{j=1}^n F(\lambda_j^t) \right) = G(\lambda_i^t, \lambda^t), \tag{188}$$

where $F(\cdot)$ is defined as per Lemma E.4, $\eta = \Theta(1/\sqrt{d})$ and $\lambda^t := \{\lambda_1^t, \dots, \lambda_n^t\}$. Then, for large enough d , we have

$$|\lambda_i^{t+1} - 1| \leq (1 - c\delta \cdot \eta) |\lambda_i^t - 1|$$

and thus after t iterations

$$|\lambda_i^t - 1| \leq \max\{(M-1), (1-\delta)\} \exp(-c\delta \cdot \eta t),$$

where $c, C > 0$ are constants.

Proof. We first show by induction that $\sum_{i=1}^n \lambda_i^t = n$ holds for all t . In fact,

$$\begin{aligned}
 \sum_{i=1}^n \lambda_i^{t+1} &= \sum_{i=1}^n \lambda_i^t + \eta \left(\sum_{i=1}^n F(\lambda_i^t) - \sum_{i=1}^n \lambda_i^t \cdot \frac{1}{n} \sum_{j=1}^n F(\lambda_j^t) \right) \\
 &= n + \eta \left(\sum_{i=1}^n F(\lambda_i^t) - \sum_{j=1}^n F(\lambda_j^t) \right) = n.
 \end{aligned}$$

Now, we will show the convergence of λ_{min}^t and λ_{max}^t . To do so, we assume that $\lambda_{max}^t \leq M$ and $\lambda_{min}^t \geq \delta$ holds at time step t (we will verify this later). Define the function $g : \mathbb{R} \rightarrow \mathbb{R}$ as

$$g(x) := x + \eta (F(x) - x \cdot C). \quad (189)$$

By taking the derivative, we have that, for sufficiently large d ,

$$g'(x) = 1 + \eta (F'(x) - C) > 0,$$

as $\|F'\|_\infty \leq C$. This implies that $g(\cdot)$ is a monotone increasing function, which gives that

$$\begin{aligned} \max_{i \in [n]} g(\lambda_i^t) &= g(\lambda_{max}^t), \\ \min_{i \in [n]} g(\lambda_i^t) &= g(\lambda_{min}^t). \end{aligned} \quad (190)$$

Note that the updates on λ_i^t in (188) have a common part for all $i \in [n]$, i.e.,

$$\left| \frac{1}{n} \sum_{j=1}^n F(\lambda_j^t) \right| \leq C,$$

where we used that $\|F'\|_\infty \leq C$. In this view, by definition of g and (190), we have

$$\begin{aligned} \lambda_{max}^{t+1} &= G(\lambda_{max}^t, \lambda^t), \\ \lambda_{min}^{t+1} &= G(\lambda_{min}^t, \lambda^t), \end{aligned} \quad (191)$$

which means that the min/max value at the previous step are mapped to the min/max value at the next step of (188). Using that $\frac{1}{n} \sum_{i=1}^n \lambda_i^t = 1$ we can write

$$\begin{aligned} \lambda_i^{t+1} &= \lambda_i^t + \eta \left(\frac{1}{n} \sum_{j=1}^n \lambda_j^t \cdot F(\lambda_i^t) - \lambda_i^t \cdot \frac{1}{n} \sum_{j=1}^n F(\lambda_j^t) \right) \\ &= \lambda_i^t + \eta \left(\frac{1}{n} \sum_{j=1}^n \left[\frac{\lambda_j^t \lambda_i^t}{(\alpha + \lambda_i^t)^2} - \frac{\lambda_i^t \lambda_j^t}{(\alpha + \lambda_j^t)^2} \right] \right) \\ &= \lambda_i^t + \eta \left(\frac{1}{n} \sum_{j=1}^n \lambda_i^t \lambda_j^t \left(\frac{(2\alpha + \lambda_i^t + \lambda_j^t)(\lambda_j^t - \lambda_i^t)}{(\alpha + \lambda_i^t)^2 (\alpha + \lambda_j^t)^2} \right) \right). \end{aligned} \quad (192)$$

Recall that we assumed $\lambda_{max}^t \leq M$ and $\lambda_{min}^t \geq \delta$. In this view, we get the following bound

$$\lambda_{max}^t \lambda_j^t \left(\frac{(2\alpha + \lambda_{max}^t + \lambda_j^t)(\lambda_{max}^t - \lambda_j^t)}{(\alpha + \lambda_{max}^t)^2 (\alpha + \lambda_j^t)^2} \right) \geq (\lambda_{max}^t - \lambda_j^t) \cdot \frac{2\alpha\delta}{(\alpha + M)^4}, \quad (193)$$

which is justified as follows

$$\begin{aligned} \lambda_{max}^t \lambda_j^t \left(\frac{(2\alpha + \lambda_{max}^t + \lambda_j^t)(\lambda_{max}^t - \lambda_j^t)}{(\alpha + \lambda_{max}^t)^2 (\alpha + \lambda_j^t)^2} \right) &= (\lambda_{max}^t - \lambda_j^t) \cdot \left(\frac{(2\alpha + \lambda_{max}^t + \lambda_j^t) \lambda_{max}^t \lambda_j^t}{(\alpha + \lambda_{max}^t)^2 (\alpha + \lambda_j^t)^2} \right) \\ &\geq (\lambda_{max}^t - \lambda_j^t) \cdot \frac{2\alpha \cdot 1 \cdot \delta}{(\alpha + M)^2 (\alpha + M)^2}, \end{aligned}$$

where we used that $\lambda_{max}^t \geq 1$ since $\sum_{i=1}^n \lambda_i^t = n$. Hence, using the previous observation about mapping of extremes in (191) and the observation above, we get from (192) that

$$\lambda_{max}^{t+1} \leq \lambda_{max}^t - \eta \cdot \frac{1}{n} \sum_{j=1}^n \left[(\lambda_{max}^t - \lambda_j^t) \cdot \frac{2\alpha\delta}{(\alpha + M)^4} \right], \quad (194)$$

which leads to

$$\begin{aligned}
 \lambda_{max}^{t+1} - 1 &\leq \lambda_{max}^t - 1 - \eta \cdot \frac{1}{n} \sum_{j=1}^n \left[(\lambda_{max}^t - \lambda_j^t) \cdot \frac{2\alpha\delta}{(\alpha + M)^4} \right] \\
 &= \lambda_{max}^t - 1 - \eta \cdot \left[(\lambda_{max}^t - 1) \cdot \frac{2\alpha\delta}{(\alpha + M)^4} \right] \\
 &= (\lambda_{max}^t - 1) \left(1 - \eta \cdot \frac{2\alpha\delta}{(\alpha + M)^4} \right) = (\lambda_{max}^t - 1)(1 - c\delta \cdot \eta),
 \end{aligned} \tag{195}$$

where we used that $\sum_{j=1}^n \lambda_j^t = n$ in the first equality. Hence, using that $\lambda_{max}^t \geq 1$ as $\sum_{j=1}^n \lambda_j^t = n$ we have

$$|\lambda_{max}^{t+1} - 1| = \lambda_{max}^{t+1} - 1 \leq |\lambda_{max}^t - 1| \cdot (1 - c\delta \cdot \eta). \tag{196}$$

Similarly to the previous bound, we get that

$$\lambda_{min}^t \lambda_j^t \left(\frac{(2\alpha + \lambda_{min}^t + \lambda_j^t)(\lambda_{min}^t - \lambda_j^t)}{(\alpha + \lambda_{min}^t)^2 (\alpha + \lambda_j^t)^2} \right) \leq \lambda_j^t (\lambda_{min}^t - \lambda_j^t) \frac{2\alpha\delta}{(\alpha + M)^4},$$

since $\lambda_{min}^t \leq \lambda_t$. Hence, using the previous observation about mapping of extremes in (191) and the observation above, we deduce from (192) that

$$\begin{aligned}
 \lambda_{min}^{t+1} - 1 &\geq (\lambda_{min}^t - 1) - \eta \cdot \frac{1}{n} \sum_{j=1}^n \left[\lambda_j^t (\lambda_{min}^t - \lambda_j^t) \frac{2\alpha\delta}{(\alpha + M)^4} \right] \\
 &= (\lambda_{min}^t - 1) - \eta \cdot \lambda_{min}^t \cdot \frac{2\alpha\delta}{(\alpha + M)^4} + \eta \cdot \frac{2\alpha\delta}{(\alpha + M)^4} \cdot \frac{1}{n} \sum_{j=1}^n (\lambda_j^t)^2 \\
 &\geq (\lambda_{min}^t - 1) - \eta \cdot (\lambda_{min}^t - 1) \cdot \frac{2\alpha\delta}{(\alpha + M)^4} \\
 &= (\lambda_{min}^t - 1) \cdot (1 - c\delta \cdot \eta),
 \end{aligned} \tag{197}$$

where in the second inequality we used Jensen's inequality for x^2 as $\sum_{j=1}^n \lambda_j^t = n$. Hence, we get the following

$$|\lambda_{min}^{t+1} - 1| = 1 - \lambda_{min}^{t+1} \leq |\lambda_{min}^t - 1| \cdot (1 - c\delta \cdot \eta), \tag{198}$$

since $\lambda_{min}^t \leq 1$ as $\sum_{j=1}^n \lambda_j^t = n$.

In this view, the assumptions $\lambda_{max}^t \leq M$ and $\lambda_{min}^t \geq \delta$ follow from (196) and (198) since the extremes are getting closer to one after each iteration. Recalling that by the assumption on initialization

$$\max_i |\lambda_i^0 - 1| \leq \max\{(M - 1), (1 - \delta)\},$$

the claim follows. $\square$

H. Proofs for General Covariance

Lemma H.1. Assume that $\{\hat{\gamma}_i\}_{i \in [K]}, \{\hat{s}_i\}_{i \in [K]}$ minimize

$$-\frac{\left(\sum_{i=1}^K D_i \gamma_i\right)^2}{\left(g(1) \cdot n + \sum_{i=1}^K \frac{\gamma_i^2}{s_i}\right)}. \tag{199}$$

Then, for any $i < j$, we must have $\hat{s}_i = \min\{\hat{s}_j + \hat{s}_j, k_i\}$.

Proof of Lemma H.1. Since the $\{\hat{\gamma}_i\}_{i \in [K]}, \{\hat{s}_i\}_{i \in [K]}$ are optimal, if we fix two indices $i < j$ the corresponding $\hat{\gamma}_i, \hat{\gamma}_j, \hat{s}_i, \hat{s}_j$ are optimal among all $\gamma_i, \gamma_j, s_i, s_j$ satisfying

$$\begin{cases} 0 < \gamma_i + \gamma_j = \gamma := \hat{\gamma}_i + \hat{\gamma}_j \leq n, \\ 0 < s_i + s_j = s := \hat{s}_i + \hat{s}_j \leq \min\{n, k_i + k_j\}. \end{cases} \quad (200)$$

Thus, we proceed by analysing the solution for two fixed indices under the constraints (200) (keeping all other $\hat{\gamma}_l, \hat{s}_l$ for $l \notin \{i, j\}$ fixed). Note that, for each fixed (γ_i, γ_j) satisfying the constraints (200), the following objective

$$\begin{aligned} \frac{\gamma_i^2}{s_i} + \frac{\gamma_j^2}{s_j} &\rightarrow \min_{s_i, s_j} \\ \text{s.t. } &s_i \leq k_i, s_j \leq k_j, s_i + s_j = s \end{aligned} \quad (201)$$

is equivalent to finding optimal ranks for (199). Importantly, in (201) we consider continuous (s_i, s_j) . This relaxation has the same minimum, since we will show that the optimal s_i, s_j have integer values. We may also assume that $\gamma_j > 0$ as otherwise clearly $s_i = \min\{s, k_i\}$ is optimal.

Since (201) is strictly convex (on the domain given by the constraints), we can find its unique minimizer by finding a solution to the KKT conditions:

$$-\frac{\gamma_i^2}{s_i^2} + (\lambda + \mu_i) = 0, \quad -\frac{\gamma_j^2}{s_j^2} + (\lambda + \mu_j) = 0, \quad \mu_i, \mu_j \geq 0, \quad \mu_i(s_i - k_i) = 0, \quad \mu_j(s_j - k_j) = 0, \quad s = s_i + s_j.$$

If $s_i = k_i$ or $s_j = 0$, then the claim is readily obtained. We will now prove that, if this is not the case, then we can find new $\tilde{s}_i, \tilde{s}_j, \tilde{\gamma}_i, \tilde{\gamma}_j$ which achieve a better value.

We first show that for $s_i < k_i, 0 < s_j < k_j$

$$\frac{\gamma_i^2}{s_i} + \frac{\gamma_j^2}{s_j} = \gamma_i \frac{\gamma}{s} + \gamma_j \frac{\gamma}{s} = \frac{\gamma^2}{s}. \quad (202)$$

Note that, in this case, $\mu_i = \mu_j = 0$, so the first two KKT conditions imply

$$\frac{\gamma_i}{s_i} = \sqrt{\lambda} = \frac{\gamma_j}{s_j}.$$

Thus, we have

$$\frac{\gamma_i}{s_i} = \frac{\gamma_j}{s_j} = \frac{\gamma_i + \gamma_j}{s_i + s_j} = \frac{\gamma}{s}, \quad (203)$$

from which (202) is immediate.

For the case $s_j = k_j$ and $s_i < k_i$, we have that $\mu_j \geq \mu_i = 0$, hence

$$\frac{\gamma_i}{s_i} = \sqrt{\lambda + \mu_i} \leq \sqrt{\lambda + \mu_j} = \frac{\gamma_j}{s_j}.$$

From the previous case, we know that without the constraints on k_i, k_j the optimal value in (201) is $\frac{\gamma^2}{s}$. Thus,

$$\frac{\gamma_i^2}{s_i} + \frac{\gamma_j^2}{s_j} \geq \frac{\gamma^2}{s}.$$

Now, for $\epsilon > 0$, define $\tilde{s}_i = s_i + \epsilon, \tilde{s}_j = s_j - \epsilon$. Note that, as $s_i < k_i$ and $s_j > 0$, we can choose ϵ small enough such that $0 < \tilde{s}_i < k_i, 0 < \tilde{s}_j < k_j$. At this point, let us simply choose $\tilde{\gamma}_i, \tilde{\gamma}_j$ such that

$$\frac{\tilde{\gamma}_i}{\tilde{s}_i} = \frac{\tilde{\gamma}_j}{\tilde{s}_j}$$

which as in (202), (203) implies that

$$\frac{\tilde{\gamma}_i^2}{\tilde{s}_i} + \frac{\tilde{\gamma}_j^2}{\tilde{s}_j} = \frac{\gamma^2}{s} \leq \frac{\gamma_i^2}{s_i} + \frac{\gamma_j^2}{s_j}. \quad (204)$$

We also have $\tilde{\gamma}_i > \gamma_i$, as otherwise

$$\frac{\tilde{\gamma}_i}{\tilde{s}_i} < \frac{\gamma_i}{s_i} \leq \frac{\gamma_j}{s_j} < \frac{\tilde{\gamma}_j}{\tilde{s}_j}$$

would be a contradiction. This gives that

$$D_i \gamma_i + D_j \gamma_j < D_i \tilde{\gamma}_i + D_j \tilde{\gamma}_j,$$

which implies that our new choice achieves a lower value for (199), thus giving the desired contradiction. $\square$

Lemma H.2. Assume that f, f_i are differentiable strictly convex functions on $\mathbb{R}$ such that

$$f'_i(0) < f'_j(0) < 0, \quad i < j, \quad \lim_{m_i \rightarrow +\infty} f'_i(m_i) = +\infty, \quad \lim_{m_i \rightarrow -\infty} f'_i(m_i) = -\infty, \quad (205)$$

and

$$f(0) = f'(0) = 0, \quad \lim_{m \rightarrow +\infty} f'(m) = +\infty. \quad (206)$$

Then, the objective given by

$$\min_{m_i \geq 0} f(m) + \sum_{i=1}^K f_i(m_i), \quad m = \sum_{i=1}^K m_i \quad (207)$$

has a unique minimizer. It is uniquely characterised by being of the form $(m_1, \dots, m_M, 0, \dots, 0)$ and satisfying

$$m = \sum_{i=1}^M \left((-f'_i)^{-1} \circ f' \right) (m), \quad m_i = \left((-f'_i)^{-1} \circ f' \right) (m) \geq 0, \quad f'(m) + f'_i(m_i) \geq 0, \quad i \in [M]. \quad (208)$$

Furthermore, it can be obtained via binary search by finding the largest index M , such that the corresponding m_i are all strictly positive.

While the assumptions of this theorem might seem technical, most of them can be relaxed. However, we note that all such assumptions are fulfilled by the setting being studied and relaxing them would come at the cost of the readability of the proof of Lemma H.2.

Proof of Lemma H.2. We start by showing that (207) has a unique minimizer. Recall that f and f_i are strictly convex functions, and, hence, their derivatives f' and f'_i are increasing. From (206), we also obtain that $\lim_{m \rightarrow +\infty} f'(m) = +\infty$. By monotonicity, we have $f'_i(m_i) \geq f'_i(0)$. Therefore,

$$\lim_{m \rightarrow +\infty} f'(m) + \sum_{i=1}^K f'_i(m_i) = +\infty,$$

and thus

$$\lim_{m \rightarrow +\infty} f(m) + \sum_{i=1}^K f_i(m_i) = +\infty.$$

As a consequence, the objective achieves its infimum. Therefore, as $f(m) + \sum_{i=1}^K f_i(m_i)$ is strictly convex, the minimum is unique.

Notice that Slater's condition is satisfied, since the feasible set of (207) has an interior point. Hence, $\{m_i\}_{i=1}^K$ is a unique minimizer of (207) if and only if it satisfies the following KKT conditions (for the “if and only if” statement, see for instance page 244 in Boyd et al. (2004)):

1. Stationary condition: $f'(m) + f'_i(m_i) - \lambda_i = 0$.

2. Primal feasibility: $m_i \geq 0$.
3. Complementary slackness: $\lambda_i m_i = 0$.
4. Dual feasibility: $\lambda_i \geq 0$.

In particular, the uniqueness of the minimizer implies that the KKT conditions have a unique solution. Thus, we only need to show that the m_i found by this procedure satisfy the above equations.

We now show that the active set $\mathcal{A} := \{i : m_i > 0\}$ for the optimal m_i is monotone, meaning that $\mathcal{A} = [M]$ for some $M \leq K$. We prove the statement by contradiction. Assume that there exists $m_i = 0$ and $m_j > 0$ where $i < j$. Recall that f'_j is strictly increasing, which by the ordering condition (205) implies that

$$f'_i(0) + f' \left(\sum_{\ell=1}^K m_\ell \right) < f'_j(m_j) + f' \left(\sum_{\ell=1}^K m_\ell \right).$$

Hence, taking some sufficiently small mass from m_j and redistributing it in m_i will decrease the objective value in (207), which concludes the proof.

Fix $M \leq K$. We now show that the solution of the following system of equations

$$f'(m) + f'_i(m_i) = 0, \quad \forall i \leq M \quad (209)$$

exists and unique. Note that this system comes from the 1. and 3. KKT conditions.

As f'_i is strictly monotone, its inverse exists and, hence, from (209) we get

$$m_i = (-f'_i)^{-1}(f'(m)), \quad (210)$$

which gives

$$m = \sum_{i=1}^M (-f'_i)^{-1}(f'(m)). \quad (211)$$

Let us argue the existence and uniqueness of the solution of equation (211) for a fixed M . Recall that f'_i is increasing and, thus, $-f'_i$ is decreasing. The inverse of a decreasing function is decreasing, hence $(-f'_i)^{-1}$ is decreasing. Recalling that f' is increasing and that the composition of an increasing and a decreasing function is decreasing, it follows that $(-f'_i)^{-1}(f'(m))$ is decreasing. By assumption $f'_i(0) < 0$ and f'_i is increasing such that $\lim_{m_i \rightarrow +\infty} f'_i(m_i) = +\infty$, therefore the value $(-f'_i)^{-1}(0)$ is well-defined and

$$(-f'_i)^{-1}(0) > 0.$$

Thus, we have that

$$g_M(m) = \sum_{i=1}^M (-f'_i)^{-1}(f'(m)) - m$$

is a strictly decreasing function with

$$\lim_{m \rightarrow +\infty} g_M(m) = -\infty, \quad g_M(0) > 0.$$

In this view, the solution of (211) exists and unique.

Next, we elaborate on why (210) is well-defined given the solution of (211). Note that, by our assumptions,

$$\lim_{m_i \rightarrow +\infty} f'_i(m_i) = +\infty, \quad \lim_{m_i \rightarrow -\infty} f'_i(m_i) = -\infty,$$

hence, the same holds for $(-f'_i)^{-1}$, and, thus, due to continuity the quantity

$$(-f'_i)^{-1}(x)$$

is well-defined for any $x \in \mathbb{R}$. Given this, we readily have that the solution of the system (209) exists and unique. Furthermore, this solution can be found using (211) and (210). Note also that (211) and (210) agree with (208).

We now show that the following procedure finds the optimal active set $\mathcal{A}^* = [M^*]$. Let $m_i(M)$, $i \leq M$ be a solution of (209) for fixed value of $M \leq K$, and define $m(M) := \sum_{i=1}^M m_i(M)$. Using (211) and (210) find the smallest M such that the corresponding $m_M(M)$ is non-negative, then $M^* = M - 1$ if $M \geq 1$, otherwise, $m = m_i = 0$, $\forall i \in [K]$. If no such M was found, $M^* = [K]$. To show that the described procedure in fact gives the optimal active set $\mathcal{A}^* = [M^*]$, we need to prove that

1. If $M < M^*$, then $m_i(M) \geq 0$.
2. If $M > M^*$, then $m_M(M) \leq 0$.

Clearly, these two conditions imply that the active set of the minimizer is given by $[M^*]$, and it can be found via binary search.

We start by proving the first property. Note that, by the KKT conditions on the optimizer M^* , we have that

$$m_i(M^*) \geq 0.$$

First assume that $m(M) > m(M^*)$. By monotonicity, it follows from (210) that

$$m_i(M) < m_i(M^*),$$

but

$$m(M^*) = \sum_{i=1}^M m_i(M^*) + \sum_{i=M+1}^{M^*} m_i(M^*) \geq \sum_{i=1}^M m_i(M^*) > \sum_{i=1}^M m_i(M) = m(M),$$

where we have used that $m_i(M^*) \geq 0$, which is a contradiction. Thus, we have that $m(M) \leq m(M^*)$. Again, by (210) and monotonicity,

$$m_i(M) \geq m_i(M^*),$$

and, hence, all $m_i(M)$ are non-negative.

We finally argue the second property. We start by proving a weaker statement, i.e., there exists $i \geq M^* + 1$ such that $m_i(M) < 0$. Assume that $m(M) < m(M^*)$. By (210) and monotonicity

$$m_i(M) > m_i(M^*),$$

hence, the following holds:

$$m(M) = \sum_{i=1}^M m_i(M) = \sum_{i=1}^{M^*} m_i(M) + \sum_{i=M^*+1}^M m_i(M) > \sum_{i=1}^{M^*} m_i(M^*) + \sum_{i=M^*+1}^M m_i(M) = m(M^*) + \sum_{i=M^*+1}^M m_i(M),$$

which since $m(M) < m(M^*)$ implies that $\sum_{i=M^*+1}^M m_i(M)$ is a negative quantity. Thus, there exists $i \geq M^* + 1$ such that $m_i(M) < 0$. Assume now that $m(M) \geq m(M^*)$. Recall that only the minimizer satisfies the KKT conditions, thus

$$f'(m(M^*)) + f'_M(0) \geq 0,$$

which, as f' is increasing, implies that

$$f'(m(M)) + f'_M(0) \geq 0.$$

By construction of $m_M(M)$, we know that

$$f'(m(M)) + f'_M(m_M(M)) = 0,$$

thus, by monotonicity of f'_M we have $m_M(M) \leq 0$.

It remains to show that it suffices to check $m_M(M) \leq 0$ and not an arbitrary $m_i(M)$ for $i \geq M^* + 1$. Assume that $m_i(M) \leq 0$ for some $i \leq M$. Recall that by assumption

$$f'_i(0) < f'_M(0) < 0,$$

and by construction we have

$$f'_i(m_i(M)) = f'_M(m_M(M)) = -f'(m(M)).$$

Since f'_i is a decreasing function, we get that $-f'(m(M)) < f'_i(0)$. Recalling that $f'_i(0) < f'_M(0)$, we get $-f'(m(M)) < f'_M(0)$ and, hence, by monotonicity of f'_M we obtain that $m_M(M) \leq 0$, which concludes the proof. $\square$

Lemma H.3. *The minimizer of (20) can be computed in $\log(K)$ steps via binary search by finding the smallest index M^* such that*

$$\frac{g(1)}{c_1^2 n} \sum_{j=1}^{M^*+1} s_j (D_{M^*+1} - D_j) + D_{M^*+1} \leq 0. \quad (212)$$

Then, the optimal active set has the form $\mathcal{A} = [M^*]$ and corresponding non-zero β_i , for $i \leq M^*$, are computed as

$$\beta_i = \frac{s_i}{c_1} \cdot \left(\frac{\frac{g(1)}{c_1^2 n} \sum_{j \in \mathcal{A}} s_j \Delta_j + D_1}{\frac{g(1)}{c_1^2 n} \sum_{j \in \mathcal{A}} s_j + 1} - \Delta_i \right), \quad (213)$$

where $\Delta_j = D_1 - D_j$.

Proof of Lemma H.3. By rescaling $g(x)$ as $\frac{g(x)}{c_1^2}$ and β_i as $c_1 \beta_i$, we may without loss of generality assume that $c_1 = 1$. From the results of Lemma H.2, by a direct computation, we get that for $\mathcal{A} = [M]$

$$\beta_j(M) = m_j(M) = s_j \cdot \left(\frac{\frac{g(1)}{n} \sum_{i=1}^M s_i \Delta_i + D_1}{\frac{g(1)}{n} \sum_{i=1}^M s_i + 1} - \Delta_j \right), \quad \forall j \leq M,$$

thus, applying the described binary search procedure to find M^* such that $M^* + 1 = \min(\arg \min_M \mathbb{1}[m_M(M) > 0])$ finishes the proof.

We now elaborate on the computations. For the compactness of the notation, we omit the dependence on active set in m_i 's and m . We apply Lemma H.2 with

$$f(x) = \frac{g(1)}{n} \cdot x^2, \quad f_i(x) = \frac{x^2}{s_i} - 2D_i x,$$

which gives

$$f'(x) = \frac{2g(1)}{n} \cdot x, \quad f'_i(x) = \frac{2x}{s_i} - 2D_i.$$

Hence, we obtain that

$$(-f'_i)^{-1}(x) = \frac{s_i \cdot (2D_i - x)}{2},$$

and, thus, by (211) we obtain

$$m = \sum_{i=1}^M (-f'_i)^{-1}(f'(m)) = -f'(m) \cdot \sum_{i=1}^M \frac{s_i}{2} + \sum_{i=1}^M D_i s_i = -\frac{g(1)}{n} \cdot m \cdot \sum_{i=1}^M s_i + \sum_{i=1}^M D_i s_i.$$

In this view, we get

$$m = \frac{\sum_{i=1}^M D_i s_i}{\frac{g(1)}{n} \sum_{i=1}^M s_i + 1},$$

and, hence, since by (210) the following holds

$$m_j = (-f'_j)^{-1}(f'(m)),$$

we get

$$\begin{aligned}
 m_j &= s_j \cdot \frac{2D_j - f'(m)}{2} = s_j \cdot \frac{2D_j \left(\frac{g(1)}{n} \sum_{i=1}^M s_i + 1 \right) - \frac{2g(1)}{n} \cdot \sum_{i=1}^M D_i s_i}{2 \cdot \left(\frac{g(1)}{n} \sum_{i=1}^M s_i + 1 \right)} \\
 &= s_j \cdot \frac{\frac{g(1)}{n} \sum_{i=1}^M D_j s_i + D_j - \frac{g(1)}{n} \sum_{i=1}^M D_i s_i + \frac{g(1)}{n} \sum_{i=1}^M D_1 s_i - \frac{g(1)}{n} \sum_{i=1}^M D_1 s_i - D_1 + D_1}{\frac{g(1)}{n} \sum_{i=1}^M s_i + 1} = \\
 &= s_j \cdot \left(\frac{\frac{g(1)}{n} \sum_{i=1}^M s_i \Delta_i + D_1}{\frac{g(1)}{n} \sum_{i=1}^M s_i + 1} - \Delta_j \right),
 \end{aligned}$$

where $\Delta_j = D_1 - D_j$. It is easy to verify that the condition

$$\frac{g(1)}{n} \sum_{j=1}^{M^*+1} s_j (D_{M^*+1} - D_j) + D_{M^*+1} \leq 0$$

described in the statement of the lemma is equivalent to $\beta_{M^*+1}(M^*+1) = m_{M^*+1}(M^*+1) \leq 0$, which concludes the proof. $\square$

Proof of Theorem 5.2. We start by showing how the lower bound reduces to the objective in (20). Consider the following block decomposition of $\mathbf{B}$ in accordance with $\mathbf{D}$ as in (25)

$$\mathbf{B} = [\mathbf{\Gamma}_1 \mathbf{B}_1 | \cdots | \mathbf{\Gamma}_K \mathbf{B}_K],$$

where $\mathbf{B}_j \in \mathbb{R}^{n \times k_j}$ with $\|(\mathbf{B}_j)_{i,:}\|_2 = 1$ and $\{\mathbf{\Gamma}_j\}_{j=1}^K$ are diagonal matrices.

Since we require $\|\mathbf{B}_{i,:}\|_2 = 1$, the $\mathbf{\Gamma}_i$ must satisfy

$$\sum_{j=1}^K \mathbf{\Gamma}_j^2 = \mathbf{I}. \tag{214}$$

Thus, up to a multiplicative factor $1/d$ and an additive term $\text{Tr}[\mathbf{D}^2]$, the objective (19) can be written as:

$$\beta^2 (\text{Tr}[\mathbf{M} f(\mathbf{M})]) - 2c_1 \beta \cdot \sum_{i=1}^K D_i \cdot \text{Tr}[\mathbf{\Gamma}_i^2], \tag{215}$$

where $\mathbf{M} = \sum_{i=1}^K \mathbf{M}_i := \sum_{i=1}^K \mathbf{\Gamma}_i \mathbf{B}_i \mathbf{B}_i^\top \mathbf{\Gamma}_i$. Recall that $f(x) = c_1^2 x + g(x)$, where g is the sum of odd monomials. Hence, we will be able to lower bound the terms in the first trace of (215) in a similar fashion to Proposition 4.4. Note that

$$\text{Tr}[\mathbf{M}_i^2] = \langle \mathbf{1}, \mathbf{M}_i^{\circ 2} \mathbf{1} \rangle,$$

so applying Theorem A in Khare (2021) gives that

$$(\mathbf{\Gamma}_i \mathbf{B}_i \mathbf{B}_i^\top \mathbf{\Gamma}_i)^{\circ 2} \succeq \frac{1}{s_i} \cdot \text{Diag}(\mathbf{\Gamma}_i^2) \text{Diag}(\mathbf{\Gamma}_i^2)^\top,$$

where $s_i = \text{rank}(\mathbf{B}_i \mathbf{B}_i^\top)$. Thus, we have the bound

$$\text{Tr}[\mathbf{M}_i^2] \geq \frac{1}{s_i} (\text{Tr}[\mathbf{\Gamma}_i^2])^2$$

Since $xg(x) \geq 0$, we can lower bound the rest of the terms with the identity, i.e.,

$$\text{Tr}[\mathbf{M} g(\mathbf{M})] = \langle \mathbf{1}, \mathbf{M} \circ g(\mathbf{M}) \mathbf{1} \rangle \geq g(1) \cdot n$$

as $\text{Diag}(\mathbf{M}) = \mathbf{I}$. Consequently, neglecting the cross-terms $\text{Tr}[\mathbf{M}_i \mathbf{M}_j]$ (as the trace of the product of PSD matrices is non-negative) we arrive at

$$\text{Tr}[\mathbf{M}f(\mathbf{M})] \geq g(1) \cdot n + c_1^2 \cdot \sum_{i=1}^K \frac{1}{s_i} (\text{Tr}[\mathbf{\Gamma}_i^2])^2.$$

Defining $\gamma_i := \text{Tr}[\mathbf{\Gamma}_i^2] \geq 0$, we arrive at the following lower bound on (215):

$$\beta^2 \left(g(1) \cdot n + \sum_{i=1}^K \frac{\gamma_i^2}{s_i} \right) - 2\beta \cdot \sum_{i=1}^K D_i \gamma_i, \quad (216)$$

where, with an abuse of notation, we rescale $g(1) := g(1)/c_1^2$ and $\beta := c_1 \beta$. Now, by choosing $\beta_i := \beta \gamma_i$ and using that $\sum_{i=1}^K \gamma_i = n$ due to (214), the objective (216) is seen to be equivalent to (20). This shows that (19) $\geq \text{LB}(\mathbf{D})$. We now give a brief outline of how one can obtain the optimal s_i and β_i for (20).

For finding the optimal s_i , it is more natural to still consider (216). Due to the block form (25), the s_i have to satisfy the constraints in (21). Note that (216) evaluated at the optimal β is equal to

$$(216) \geq - \frac{\left(\sum_{i=1}^K D_i \gamma_i \right)^2}{\left(g(1) \cdot n + \sum_{i=1}^K \frac{\gamma_i^2}{s_i} \right)}. \quad (217)$$

The optimal s_i for this objective are *water-filled*, i.e.,

$$\begin{cases} \mathbf{s} = [n, 0, \dots, 0], & n \leq k_1, \\ \mathbf{s} = [k_1, k_2, \dots, k_K], & d \leq n, \\ \mathbf{s} = [k_1, \dots, k_{\text{id}(n)-1}, \text{res}(n), 0, \dots, 0] & \text{otherwise,} \end{cases} \quad (218)$$

where $\mathbf{s} = [s_1, \dots, s_K]$ and $\text{id}(n)$ denotes the first position at which

$$\min\{n, d\} - \sum_{i=1}^{\text{id}(n)} k_i < 0,$$

and

$$\text{res}(n) = \min\{n, d\} - \sum_{i=1}^{\text{id}(n)-1} k_i.$$

This follows directly from Lemma H.1. It only remains to show that the optimal β_i can be obtained via (24), which is done in Lemma H.3. This concludes the proof. $\square$

Proof of Proposition 5.3. Except for terms of the form $\text{Tr}[\mathbf{B}_i \mathbf{B}_i^\top \mathbf{B}_j \mathbf{B}_j^\top]$, all the other terms can be estimated as in the proof of Proposition 4.4. The only technical difference is that all the constants now depend on the ratios $\frac{k_i}{n}$.

We will show that, with probability at least $1 - c \exp(-cd^\epsilon)$, for all $i \neq j$,

$$\text{Tr}[\mathbf{B}_i \mathbf{B}_i^\top \mathbf{B}_j \mathbf{B}_j^\top] \leq n^{\frac{1}{2} + \epsilon}. \quad (219)$$

Thus, by a simple union bound, we have that, with probability at least $1 - \frac{c}{d^2}$, this bound holds jointly for all pairs $\mathbf{B}_i, \mathbf{B}_j$. It follows as in the proof of Lemma C.3 that we can write

$$\mathbf{B}_i \mathbf{B}_i^\top = \mathbf{P}_i \mathbf{U} \mathbf{D}_i \mathbf{U}^\top \mathbf{P}_i,$$

where by abuse of notation we pushed the factor $\frac{n}{k_i}$ in $\mathbf{D}_i$ (which will only affect the constants c, C). Here, $\mathbf{P}_i$ is a diagonal matrix such that, for any $\epsilon > 0$, with probability at least $1 - c \exp(-cd^\epsilon)$, we have that

$$\|\mathbf{P}_i - \mathbf{I}\|_{\text{op}} \leq n^{-\frac{1}{2} + \epsilon}.$$

To see this, first observe that $\Theta : (\mathbb{R}^{n \times n})^4 \mapsto \mathbb{R}$ given by $\Theta(X_1, X_2, X_3, X_4) = \text{Tr} [X_1 U D_i U^\top X_2 X_3 U D_j U^\top X_4]$ is differentiable (as it is the composition of the trace function with 4-linear form). Since by construction

$$\text{Tr} [U D_i U^\top U D_j U^\top] = \text{Tr} [\mathbf{0}] = 0,$$

this implies that, with probability at least $1 - \frac{c}{d^2}$,

$$\begin{aligned} 0 &\leq \text{Tr} [B_i B_i^\top B_j B_j^\top] \\ &= \text{Tr} [P_i U D_i U^\top P_i P_j U D_j U^\top P_j] \\ &= \text{Tr} [P_i U D_i U^\top P_i P_j U D_j U^\top P_j] - \text{Tr} [U D_i U^\top U D_j U^\top] \\ &\leq C n n^{-\frac{1}{2} + \epsilon}, \end{aligned}$$

where in the last step we used that the derivative of the trace function is bounded by $n \cdot \|\cdot\|_{op}$. Thus, (219) holds.

By construction, the sum of all the cross terms is of the form

$$\sum_{i \neq j} \text{Tr} [M_i M_j],$$

where $M_i = \Gamma_i B_i B_i^\top \Gamma_i$, $\Gamma_i^2 = \frac{\gamma_i}{n} \mathbf{I}$ and $\sum_{i=1}^K \gamma_i = n$. We have

$$\begin{aligned} \left| \sum_{i \neq j} \text{Tr} [M_i M_j] \right| &= \left| \sum_{i \neq j} \frac{\gamma_i \gamma_j}{n^2} \text{Tr} [B_i B_i^\top B_j B_j^\top] \right| \\ &\leq \sum_{i \neq j} \frac{\gamma_i \gamma_j}{n^2} \left| \text{Tr} [B_i B_i^\top B_j B_j^\top] \right| \\ &\leq C \sum_{i \neq j} \frac{\gamma_i \gamma_j}{n^2} n^{\frac{1}{2} + \epsilon} \\ &\leq C n^{\frac{1}{2} + \epsilon}, \end{aligned}$$

where in the third step we used a union bound on (219) and in the last step we used $\sum_{i=1}^K \frac{\gamma_i}{n} = 1$. $\square$

I. Details of Experiments and Additional Numerical Results

We first describe the training details and the whitening procedure that is used to preprocess natural images for MNIST (Figure 8) and CIFAR-10 (Figures 1, 5 and 7). Next, we give some remarks about the experiments concerning VAMP (Figure 4) and about the discontinuous behaviour of the derivative of the lower bound highlighted in Figure 6. In addition, we present additional numerical experiments which cover extra classes of natural images.

Activation function and weight parameterization. Note that the derivative of the sign activation is zero almost everywhere (except one point, which is the origin). In this view, we cannot use conventional gradient-based algorithms to find the optimal set of parameters for an autoencoder with the sign activation. We tackle this issue by using a straight-through estimator (see, for instance, (Yin et al., 2019)) of the sign activation. During the forward pass the activations of the first layer are computed for $\sigma(x) = \text{sign}(x)$, while during the backward pass $\sigma(x) = \tanh(x/\tau)$ is used. Here, the temperature parameter $\tau > 0$ controls how well the differentiable surrogate $\tanh(x/\tau)$ approximates $\text{sign}(x)$, as

$$\lim_{\tau \rightarrow 0} \tanh(x/\tau) = \text{sign}(x), \quad \forall x \in \mathbb{R} \setminus \{0\}.$$

More precisely, the differentiable approximation becomes more accurate for smaller values of τ . However, we also note that extremely small values of τ might cause numerical issues, since the derivative of the differentiable surrogate diverges at the origin as $\tau \rightarrow 0$. For the numerical experiments, we pick $\tau \in [0.01, 0.2]$, with the exact value depending on the specific setting.

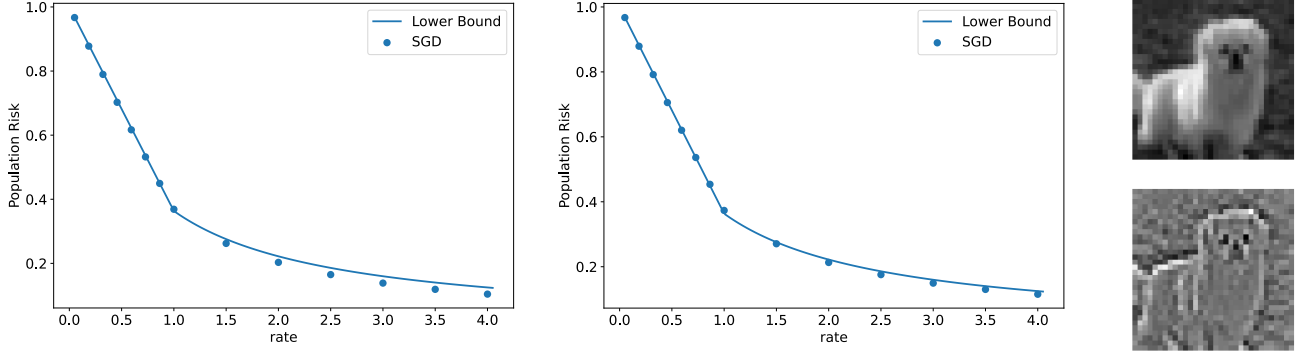


Figure 5. Compression ($\sigma \equiv \text{sign}$) of the CIFAR-10 “dog” class with a two-layer autoencoder. The data is *whitened* so that $\Sigma = I$: on top, an example of a grayscale image; on the bottom, the corresponding whitening. The blue dots are the population risk obtained via SGD, and they agree well with the solid line corresponding to the lower bounds of Theorem 4.2 and Proposition 4.3. Here, the effect of the number of augmentations used per image is shown. For the left plot each image was augmented 10 times, while for the right plot each image was augmented 15 times.

Note that the constraint on the encoder weights $\|B_{i,:}\|_2 = 1$ can be enforced via a simple reparameterization that forces the rows of B to lie on the unit sphere $\mathbb{S}^{d-1}$. More precisely, we use the following classical differentiable reparameterization of $B^\top = [b_1, \dots, b_n]$, where $b_i = \frac{\hat{b}_i}{\|\hat{b}_i\|_2}$, with $\{\hat{b}_i\}_{i=1}^n$ being the trainable parameters. We note that it is not clear a priori whether we need to impose the constraints directly for the straight-through estimator, since during the forward pass we use the norm-agnostic sign function.

Augmentation and whitening. For the experiments on natural images, we augment the data of each class 15 times. This is done to emulate the optimization of the population risk, since the amount of initial data (approximately 5000 samples per class) leads to a gap between empirical and population risks, especially for high rates. The effect of the data augmentation is represented in Figure 5 for a whitened CIFAR-10 class. It can be seen that a mild amount of augmentation, i.e., $\times 10$ and $\times 15$, is already enough for our purposes, and the difference between the two plots is rather small. Notably, this amount of augmentation brings the dataset to the scale of the original data when all classes are considered (around 50000 training examples).

The whitening procedure used in the experiments concerning isotropic data is performed as follows: given the *centered* augmented data $X \in \mathbb{R}^{n_{\text{samples}} \times d}$, we compute its empirical covariance matrix given by

$$\hat{\Sigma} = \frac{1}{n_{\text{samples}} - 1} \cdot \sum_{i=1}^{n_{\text{samples}}} X_{i,:} X_{i,:}^\top,$$

and then we multiply each input by the inverse square root of it, i.e.,

$$\hat{X}_{i,:} = \hat{\Sigma}^{-\frac{1}{2}} X_{i,:}.$$

The resulting whitened images are represented in Figures 1, 5 and 8.

In the experiments concerning non-isotropic data (Figures 2 and 9), we center the data with the empirical mean and divide by a *scalar* empirical variance computed across all the pixels, which is the standard preprocessing procedure widely used for computer vision tasks.

VAMP experiments. For the VAMP experiments, we implement the State Evolution (SE) recursion which exactly characterizes the limiting performance of VAMP as $d \rightarrow \infty$, see (Schniter et al., 2016; Rangan et al., 2019) for an overview. We then plot the fixed point of said SE recursion. A concrete description for VAMP is provided by Algorithm 2 in (Fletcher et al., 2018), which however covers a more general multi-layer setting.

“Jumps” of the lower bound derivative. The derivative switch described in Figure 6 does not necessarily happen precisely at the point when the block is filled. A switch may occur at a later point since, even if $s_i > 0$, the corresponding optimal β_i may be 0. Intuitively, this phenomenon occurs in cases when it is still better to put more mass in the block where the rank is

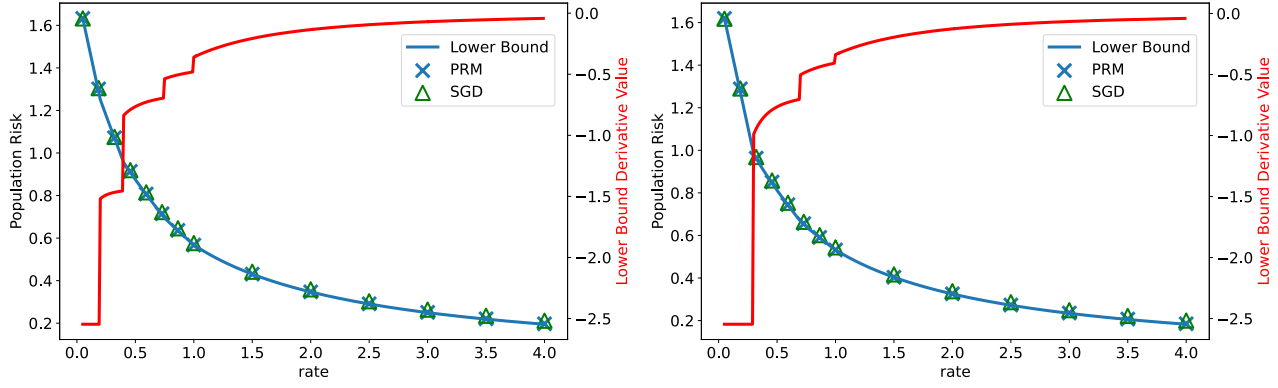


Figure 6. Compression ($\sigma \equiv \text{sign}$) of a non-isotropic Gaussian source, whose covariance matrix is obtained by taking $\mathbf{k} = (20, 20, 35, 25)$ and $(D_1, D_2, D_3, D_4) = (2, 1.5, 1, 0.8)$ for the left plot, and $\mathbf{k} = (30, 40, 30)$ and $(D_1, D_2, D_3) = (2, 1, 0.7)$ for the right plot. The blue crosses (Population Risk Minimizer, PRM) are obtained by optimizing (19) via GD. The green triangles are obtained by training an autoencoder via SGD on Gaussian samples with the given covariance structure. The red solid line plots the derivative of the population risk computed using a finite differences scheme. Note that the derivative jumps when the corresponding blocks are getting filled, although this may not happen in general, see Appendix I. A similar behavior can be observed in the isotropic case at $r = 1$, as there is only one block to fill (see Figure 4).

utilized to the fullest ($s_j = k_j$). This corresponds to the following condition on the derivatives of the objective (20):

$$\frac{\partial(20)}{\partial\beta_i}(0) > \frac{\partial(20)}{\partial\beta_j}(\beta_j^*),$$

where β_i^* stands for the optimal β_i and j denotes the first index at which $\beta_j^* > 0$. This behaviour occurs when the spectrum $\mathbf{D}$ has a large variation in scale, e.g.,

$$\mathbf{D} = [5, 0.02, 0.01].$$

In this case, the last components will be utilized for n significantly larger than k_1 ($n = k_1$ precisely characterizes the point where the rank of the first block of $\mathbf{B}$, i.e., $\mathbf{B}_1$, is the maximum possible). Note that, for this choice of $\mathbf{D}$, the plot of the derivative analogous to Figure 6 will not indicate such prominent “jumps”. In fact, the contribution of the last components to the derivative value is less significant in comparison to the analogous quantity evaluated for the top-most eigenvalues.

Additional experimental data. We also provide additional numerical simulations, similar to those presented in the body of the paper. In particular, we provide more class variations for the natural data experiments (MNIST and CIFAR-10).

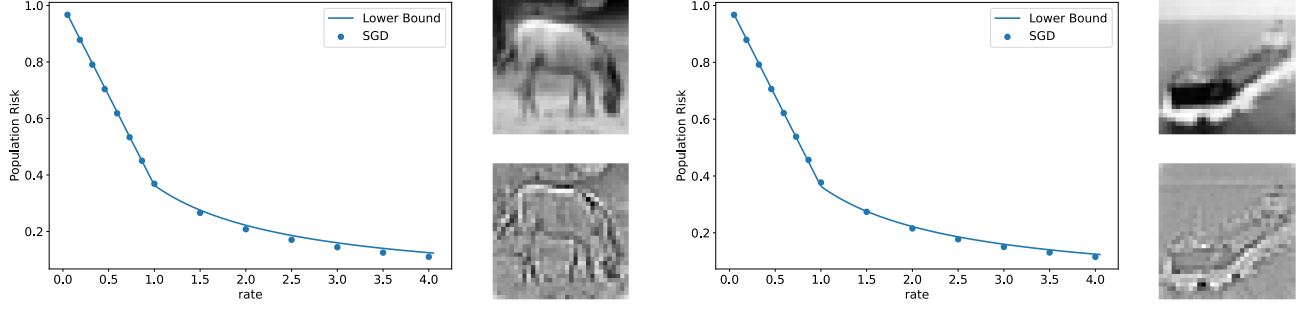


Figure 7. Compression ($\sigma \equiv \text{sign}$) of the CIFAR-10 "horse" class (left) and "ship" class (right) with a two-layer autoencoder. The data is *whitened* so that $\Sigma = \mathbf{I}$: on top, an example of a grayscale image; on the bottom, the corresponding whitening. The blue dots are the population risk obtained via SGD, and they agree well with the solid line corresponding to the lower bounds of Theorem 4.2 and Proposition 4.3. Here, in both cases the amount of augmentations per image is equal to 15.

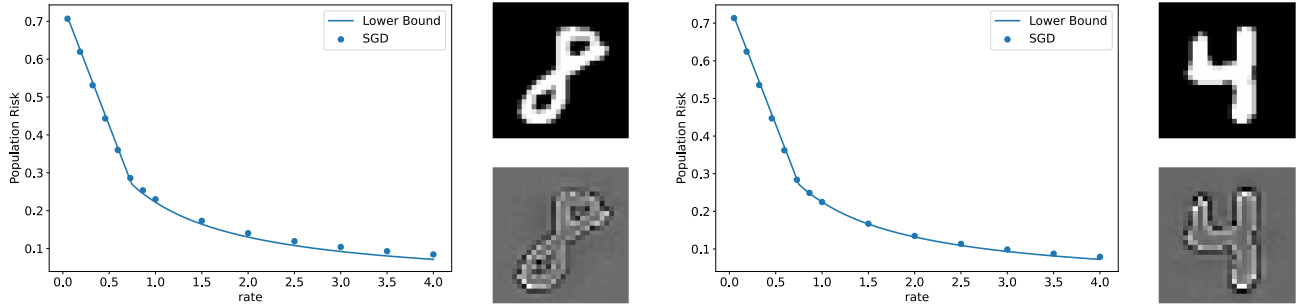


Figure 8. Compression ($\sigma \equiv \text{sign}$) of the MNIST "8" class (left) and "4" class (right) with a two-layer autoencoder. The data is *whitened* so that $\Sigma = \mathbf{I}$: on top, an example of a grayscale image; on the bottom, the corresponding whitening. The blue dots are the population risk obtained via SGD, and they agree well with the solid line corresponding to the lower bounds of Theorem 4.2 and Proposition 4.3. Here, in both cases the amount of augmentations per image is equal to 10.

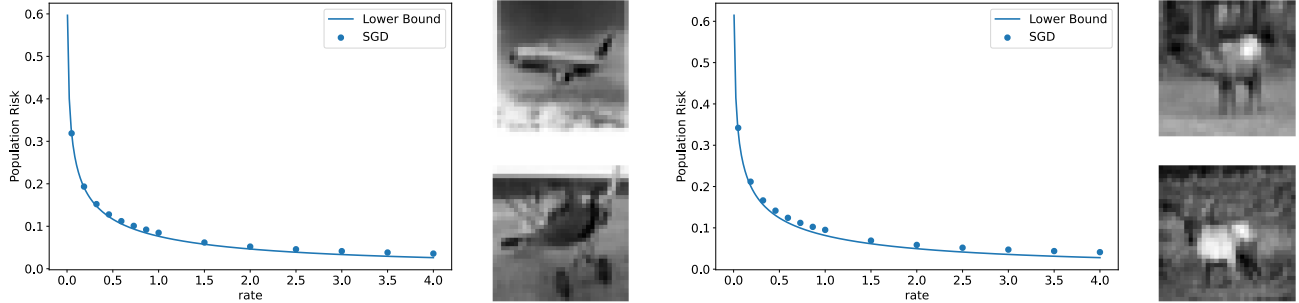


Figure 9. Compression ($\sigma \equiv \text{sign}$) of the CIFAR-10 "airplane" class (left) and "deer" class (right) with a two-layer autoencoder. The data is *not whitened* ($\Sigma \neq \mathbf{I}$). The blue dots are the SGD population risk, and they are close to the lower bound of Theorem 5.2. Here, in both cases the amount of augmentations per image is equal to 15.