

Complete Dictionary Learning via ℓ^4 -Norm Maximization over the Orthogonal Group

Yuexiang Zhai[†]

YSZ@BERKELEY.EDU

Zitong Yang[†]

ZITONG@BERKELEY.EDU

Zhenyu Liao[‡]

LIAOZHENYU2004@GMAIL.COM

John Wright[◊]

JW2966@COLUMBIA.EDU

Yi Ma[†]

YIMA@EECS.BERKELEY.EDU

[†]*Department of Electrical Engineering and Computer Science
University of California, Berkeley, CA 94720-1770*

[‡]*Kuaishou Technology.*

[◊]*Department of Electrical Engineering
Columbia University, New York, NY, 10027*

Editor: Julien Mairal

Abstract

This paper considers the fundamental problem of learning a complete (orthogonal) dictionary from samples of sparsely generated signals. Most existing methods solve the dictionary (and sparse representations) based on heuristic algorithms, usually without theoretical guarantees for either optimality or complexity. The recent ℓ^1 -minimization based methods do provide such guarantees but the associated algorithms recover the dictionary one column at a time. In this work, we propose a new formulation that *maximizes* the ℓ^4 -norm over the orthogonal group, to learn the entire dictionary. We prove that under a random data model, with nearly minimum sample complexity, the global optima of the ℓ^4 -norm are very close to signed permutations of the ground truth. Inspired by this observation, we give a conceptually simple and yet effective algorithm based on “*matching, stretching, and projection*” (MSP). The algorithm provably converges locally and cost per iteration is merely an SVD. In addition to strong theoretical guarantees, experiments show that the new algorithm is significantly more efficient and effective than existing methods, including KSVD and ℓ^1 -based methods. Preliminary experimental results on mixed real imagery data clearly demonstrate advantages of so learned dictionary over classic PCA bases.

Keywords: sparse dictionary learning, ℓ^4 -norm maximization, orthogonal group, measure concentration, fixed point algorithm

1. Introduction and Overview

1.1. Motivation

One of the most fundamental problems in signal processing or data analysis is that given an observed signal $\mathbf{y}$, either continuous or discrete, we would like to find a transform $\mathcal{F}$ such that after applying the transform, the resulting signal $\mathbf{x} = \mathcal{F}(\mathbf{y})$ becomes much more compact, sparse, or compressible. We believe such a compact representation $\mathbf{x}$ can help

reveal intrinsic structures of the observed signal $\mathbf{y}$ and is also more amenable to storage, processing, and transmission.

For computational purposes, the transforms considered are typically orthogonal linear transforms so that both $\mathcal{F}$ and $\mathcal{F}^{-1}$ are easy to represent and compute: In this case, we have $\mathbf{y} = \mathbf{D}\mathbf{x}$ or $\mathbf{x} = \mathbf{D}^*\mathbf{y}$ for some orthogonal matrix (or linear operator) $\mathbf{D}$. Examples include the classical Fourier transform (Oppenheim, 1999; Vetterli et al., 2014) or various wavelets (Vetterli and Kovacevic, 1995). Conventionally, the best transform to use is typically by “*design*”: By assuming the signals of interest have certain physical or mathematical properties (e.g. band-limited, piece-wise smooth, or scale-invariant), one may derive or design the optimal transforms associated with different classes of functions or signals. This classical approach has found its deep mathematical roots in functional analysis (Kreyszig, 1978) and harmonic analysis (Stanton and Weinstein, 1981; Katznelson, 2004) and has seen great empirical successes in digital signal processing (Oppenheim, 1999).

Nevertheless, in the modern big data era, both science and engineering are inundated with tremendous high-dimensional data, such as images, audios, languages, and genetics etc. Many of such data may or may not belong to the classes of functions or signals for which we know the optimal transforms. Assumptions about their intrinsic (low-dim) structures are not clear enough for us to derive any new transform either. Therefore, we are compelled to change our practice and ask whether we can “*learn*” an optimal transform (if exists) directly from the observed data. That is, if we observe many, say p , samples $\mathbf{y}_i \in \mathbb{R}^n$ from a model:

$$\mathbf{y}_i = \mathbf{D}_o \mathbf{x}_i, \quad i = 1, \dots, p$$

where $\mathbf{x}_i \in \mathbb{R}^m$ is presumably much more compact or sparse, can we both learn $\mathbf{D}_o$ and recover the associated $\mathbf{x}_i$ from such $\mathbf{y}_i$? Here $\mathbf{D}_o$ is called the ground truth “dictionary” and the problem is known as “Dictionary Learning.” Dictionary learning is a fundamental problem in data science since finding a sparse representation of data appears in different applications such as computational neural science (Olshausen and Field, 1996, 1997), machine learning (Argyriou et al., 2008; Ranzato et al., 2007), and computer vision (Elad and Aharon, 2006; Yang et al., 2010; Mairal et al., 2014). Note that here we know neither the dictionary $\mathbf{D}_o$ nor the hidden state $\mathbf{x}$. So in machine learning, dictionary learning belongs to the category of *unsupervised learning* problems, and a fundamental one that is.

In this work, we consider the problem of *learning a complete dictionary*¹ from sparsely generated sample signals. More precisely, an n -dimensional sample $\mathbf{y} \in \mathbb{R}^n$ is assumed to be a sparse superposition of columns of a non-singular complete dictionary $\mathbf{D}_o \in \mathbb{R}^{n \times n}$: $\mathbf{y} = \mathbf{D}_o \mathbf{x}$, where $\mathbf{x} \in \mathbb{R}^n$ is a sparse (coefficient) vector. A typical statistical model for the sparse coefficient is that entries of $\mathbf{x}$ are i.i.d. Bernoulli-Gaussian $\{x_i\} \sim_{iid} \text{BG}(\theta)$ ² (Spielman et al., 2012; Sun et al., 2015; Bai et al., 2018).

Suppose we are given a collection of sample signals $\mathbf{Y} = [\mathbf{y}_1, \mathbf{y}_2, \dots, \mathbf{y}_p] \in \mathbb{R}^{n \times p}$, each of which is generated as $\mathbf{y}_i = \mathbf{D}_o \mathbf{x}_i$ for a nonsingular matrix $\mathbf{D}_o$. Write $\mathbf{X}_o = [\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_p] \in \mathbb{R}^{n \times p}$. In this notation, we have:

$$\mathbf{Y} = \mathbf{D}_o \mathbf{X}_o. \tag{1}$$

-
1. In this work, we only consider the case the dictionary is complete. The more general setting in which the dictionary $\mathbf{D}_o$ is over-complete, $m > n$, is beyond the scope of this paper.
 2. I.e., each entry x_i is a product of independent Bernoulli and standard normal random variables: $x_i = \Omega_i V_i$, where $\Omega_i \sim_{iid} \text{Ber}(\theta)$ and $V_i \sim_{iid} \mathcal{N}(0, 1)$.

Dictionary learning is the problem of recovering both the dictionary $\mathbf{D}_o$ and the sparse coefficients $\mathbf{X}_o$, given only the samples $\mathbf{Y}$. Equivalently, we wish to factorize $\mathbf{Y}$ as $\mathbf{Y} = \mathbf{D}\mathbf{X}$, where $\mathbf{D}$ is an estimate of the true dictionary $\mathbf{D}_o$ and $\mathbf{X}$ is the sparsest possible.

Under the Bernoulli-Gaussian assumption, the problem of learning an arbitrary complete dictionary can be reduced to that of learning an *orthogonal* dictionary: As shown in Sun et al. (2015), when the objective is smooth, the problem can be converted to the orthogonal case through a preconditioning:

$$\mathbf{Y} \leftarrow \left(\frac{1}{p\theta} \mathbf{Y}\mathbf{Y}^* \right)^{-\frac{1}{2}} \mathbf{Y}.$$

So without loss of generality, we can assume that $\mathbf{D}_o$ is an orthogonal matrix: $\mathbf{D}_o \in \mathcal{O}(n; \mathbb{R})$.

Because $\mathbf{Y}$ is sparsely generated, the optimal estimate $\mathbf{D}_\star$ should make the associated coefficients $\mathbf{X}_\star$ maximally sparse. In other words, ℓ^0 -norm, the number of non-zero entries, of $\mathbf{X}_\star$ should be as small as possible, therefore, one may formulate the following optimization program to find $\mathbf{X}_\star$:

$$\min_{\mathbf{X}, \mathbf{D}} \|\mathbf{X}\|_0, \quad \text{subject to} \quad \mathbf{Y} = \mathbf{D}\mathbf{X}, \quad \mathbf{D} \in \mathcal{O}(n; \mathbb{R}). \quad (2)$$

Under fairly mild conditions, the global minimizer of the ℓ^0 -norm recovers the true dictionary $\mathbf{D}_o$ Spielman et al. (2012). However, global minimization of the ℓ^0 -norm is a challenging NP-hard problem (Donoho, 2006; Candes and Tao, 2005; Natarajan, 1995). Traditionally, one resorts to local heuristics such as orthogonal matching pursuit, as in the KSVD algorithm (Aharon et al., 2006; Rubinstein et al., 2010).³ This approach has been widely practiced but does not give any strong guarantees for optimality of the algorithm nor correctness of the solution, as we will see through experiments compared with our method.

Since ℓ^1 -norm minimization promotes sparsity (Candès, 2014) and the ℓ^1 -norm is convex and continuous, one may reformulate dictionary learning as an ℓ^1 -minimization problem:

$$\min_{\mathbf{X}, \mathbf{D}} \|\mathbf{X}\|_1, \quad \text{subject to} \quad \mathbf{Y} = \mathbf{D}\mathbf{X}, \quad \mathbf{D} \in \mathcal{O}(n; \mathbb{R}). \quad (3)$$

This reformulation (3) remains a nonsmooth optimization with nonconvex constraints, which is in general still NP-hard (Murty and Kabadi, 1987). Nevertheless, many heuristic algorithms (Mairal et al., 2008, 2009, 2012) have attempted to reformulate optimization (3) as an unconstrained optimization problem with ℓ^1 -regularization by removing the constraint $\mathbf{Y} = \mathbf{D}\mathbf{X}$.

Although the ℓ^0 - or ℓ^1 -minimization has been widely practiced in dictionary learning, rigorous justification for optimality and correctness is only provided recently. Geng and Wright (2014) is the first to show the local optimality of the ℓ^1 -minimization. Spielman et al. (2012) further proves that a complete (square and invertible) $\mathbf{D}$ can be recovered from $\mathbf{Y}$, when each column of $\mathbf{X}$ contains no more than $O(\sqrt{n})$ of nonzero entries. Subsequent works (Agarwal et al., 2013, 2014; Arora et al., 2014, 2015) have provided provable algorithms for overcomplete dictionary learning, under the assumption that each column of $\mathbf{X}$ has $\tilde{O}(\sqrt{n})$ nonzero entries.⁴

3. See also (Ravishankar and Bresler, 2015) algorithms for learning orthogonal sparsifying transformations.

4. $\tilde{O}$ suppresses logarithm factors.

In this work, we mainly focus on *complete dictionary learning*, which implies $\text{row}(\mathbf{Y}) = \text{row}(\mathbf{X})$. Spielman et al. (2012) has proposed to find the sparsest vector $\mathbf{d}^*\mathbf{Y}$ in $\text{row}(\mathbf{Y})$ one by one via solving the following optimization:

$$\min_{\mathbf{d} \in \mathbb{R}^n} \|\mathbf{d}^*\mathbf{Y}\|_1, \quad \text{subject to} \quad \mathbf{d} \neq \mathbf{0} \quad (4)$$

n times, instead of solving the hard nonconvex optimization (3) directly. (4) is easier to solve, since it can be further reduced to a linear programming. Sun et al. (2015) has proposed a new formulation with spherical constraint that finds one column of the dictionary via solving:

$$\min_{\mathbf{d} \in \mathbb{R}^n} \|\mathbf{d}^*\mathbf{Y}\|_1, \quad \text{subject to} \quad \|\mathbf{d}\|_2 = 1. \quad (5)$$

For escaping saddle points, Sun et al. (2015) has proposed a provably correct second-order Riemannian Trust Region method (Absil et al., 2009) to solve (5) and improved the sparsity level of $\mathbf{X}$ to constant.⁵ However, as addressed by Gilboa et al. (2018), the computational complexity of a second-order algorithm is high, not to mention one needs to solve (5) n times! To mitigate the computation complexity, Gilboa et al. (2018) suggests that the first-order gradient descent method with random initialization has the same performance as a second-order one, and Bai et al. (2018) shows that a randomly initialized first-order projected subgradient descent is able to solve (5). But these approaches fail to overcome the main cause for the high complexity – one needs to break down the complete dictionary learning problem (3) into solving n optimization programs like (4) (or (5)).

Besides the ℓ^1 -minimization based dictionary learning framework, works based on the sum-of-square (SoS) SDP hierarchy (Barak et al., 2015; Ma et al., 2016; Schramm and Steurer, 2017) also provide guarantee for exact recovery of (overcomplete) dictionary learning problem in polynomial time, under some specific statistical assumption of the data model. But one still needs to solve the SoS based SDP Problem n times to recover a dictionary with n components,⁶ let alone the high computational complexity for solving a high-dimensional SDP programming each time (Qu et al., 2014; Bai et al., 2018).

1.2. Our Approach and Connection to Prior Works

In this paper, we show that one can actually efficiently learn a complete (orthogonal) dictionary holistically via solving one ℓ^4 -norm optimization over the entire orthogonal group $\text{O}(n; \mathbb{R})$:

$$\max_{\mathbf{D}} \|\mathbf{D}^*\mathbf{Y}\|_4^4, \quad \text{subject to} \quad \mathbf{D} \in \text{O}(n; \mathbb{R}), \quad (6)$$

where the ℓ^4 -norm of a matrix means the sum of 4th powers of all entries: $\forall \mathbf{A} \in \mathbb{R}^{n \times m}$, $\|\mathbf{A}\|_4^4 = \sum_{i,j} a_{i,j}^4$. The intuition for (6) comes from:

$$\max_{\mathbf{X}, \mathbf{D}} \|\mathbf{X}\|_4^4, \quad \text{subject to} \quad \mathbf{Y} = \mathbf{D}\mathbf{X}, \mathbf{D} \in \text{O}(n; \mathbb{R}), \quad (7)$$

where maximizing the ℓ^4 -norm of $\mathbf{X}$ (over a sphere) promotes “spikiness” or “sparsity” of $\mathbf{X}$ (Zhang et al., 2018). It is easy to see this as the sparsest points on a unit ℓ^2 -sphere

5. Each column of $\mathbf{X}$ can contain $O(n)$ non zero entries.

6. See Theorem 1.1 in Schramm and Steurer (2017)

(points $(0, 1)$, $(0, -1)$, $(1, 0)$, and $(-1, 0)$) have the smallest ℓ^1 -norm and largest ℓ^4 -norm, as shown in Figure 1. Since the columns of orthogonal matrices have unit norm, the constraint $\mathbf{D} \in \mathcal{O}(n; \mathbb{R})$ can be viewed as simultaneously enforcing orthonormal constraints on n vectors on the unit ℓ^2 -sphere $\mathbb{S}^1$. Moreover, comparing to the ℓ^1 -norm, the ℓ^4 -norm objective is everywhere smooth, so we expect it is amenable to better optimization.

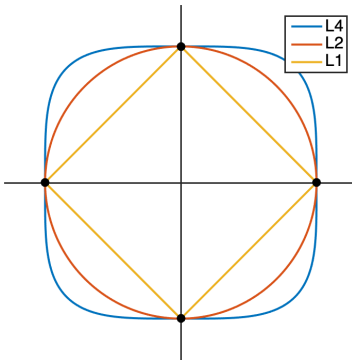


Figure 1: Unit ℓ^1 -, ℓ^2 -, and ℓ^4 - spheres in $\mathbb{R}^2$, a similar picture can be found in Figure 1 of Li and Bresler (2018).

1.2.1. SPHERICAL HARMONIC ANALYSIS

The property of the ℓ^4 -norm has long been realized and used in seeking (orthonormal) functions with similar properties since 1970's, if not any earlier. For instance for spherical harmonics, the Stanton-Weinstein Theorem (Stanton and Weinstein, 1981; Lu, 1987) have shown that “among all the ℓ^2 -normalized spherical harmonics of a given degree, the ℓ^4 -norm is locally maximized by the ‘highest-weight’ function,” which among all the eigenfunctions of the Laplacian on the sphere, is the “most concentrated” in measure (see Theorem 1 of Stanton and Weinstein (1981)). In quantum mechanics, such functions represent trajectories that are the most probable and most closely approximate classical trajectories.

1.2.2. INDEPENDENT COMPONENT ANALYSIS

We should note that 4th-order statistical cumulants have been widely used in blind source separation or independent component analysis (ICA) since the 1990's, see Hyvärinen (1997); Hyvärinen and Oja (1997) and references therein. So if $\mathbf{x}$ are n independent components, by finding extrema of the so-called *kurtosis*: $\text{kurt}(\mathbf{d}^* \mathbf{y}) \doteq \mathbb{E}[(\mathbf{d}^* \mathbf{y})^4] - 3\mathbb{E}[(\mathbf{d}^* \mathbf{y})^2]^2$, one can identify one independent (non-Gaussian) component x_i at a time. Algorithm wise, this is similar to using the ℓ^1 -minimization (5) to identify one column $\mathbf{d}_i$ at a time for $\mathbf{D}$. Fast fixed-point like algorithms have been developed for this purpose (Hyvärinen, 1997; Hyvärinen and Oja, 1997). If $\mathbf{x}$ are indeed i.i.d. Bernoulli-Gaussian, with $\|\mathbf{d}\|_2^2 = 1$, the second term in $\text{kurt}(\mathbf{d}^* \mathbf{y})$ would become a constant and the objective of ICA coincides with maximizing the sparsity-promoting ℓ^4 -norm of a vector over a sphere.

1.2.3. SUM OF SQUARES

The use of ℓ^4 -norm can also be justified from the perspective of sum of squares (SoS). The works of Barak et al. (2015); Ma et al. (2016); Schramm and Steurer (2017) show that in theory, when $\mathbf{x}$ is sufficiently sparse, one can utilize properties of higher order sum of squares polynomials (such as the fourth order polynomials) to correctly recover $\mathbf{D}$. Although Schramm and Steurer (2017) has improved proposed faster tensor decomposition based on SoS method, again the algorithm only recovers one column $\mathbf{d}_i$ at a time.

1.2.4. BLIND DECONVOLUTION

Recent works in blind deconvolution (Zhang et al., 2018; Li and Bresler, 2018) have also explored the sparsity promoting property of the ℓ^4 -norm in their objective function. Zhang et al. (2018); Li and Bresler (2018) have shown that, for any filter on a unit sphere, all local maxima of the ℓ^4 -objective are close to the inverse of the ground truth filter (up to the intrinsic sign and shift ambiguity). Moreover, the global geometry of the ℓ^4 -norm over the sphere is good – all saddle points have negative curvatures. Such nice global geometry guarantees a randomly initialized first-order Riemannian gradient descent algorithm to escape saddle points and find the ground truth, for both single channel (Zhang et al., 2018) and multi-channel (Li and Bresler, 2018; Qu et al., 2019) tasks.

1.3. Main Results

We shall first note that there is an intrinsic “signed permutation” ambiguity in all dictionary learning formulation (2), (3), and (7). For any input data matrix $\mathbf{Y}$, suppose $\mathbf{D}_\star \in \mathcal{O}(n; \mathbb{R})$ is the optimal orthogonal dictionary that gives the sparsest coefficient matrix $\mathbf{X}_\star$ satisfying $\mathbf{Y} = \mathbf{D}_\star \mathbf{X}_\star$. Then for any matrix $\mathbf{P}$ in the signed permutation group $\mathcal{SP}(n)$, the group of orthogonal matrices that only contain 0, ± 1 , we have:

$$\mathbf{Y} = \mathbf{D}_\star \mathbf{X}_\star = \mathbf{D}_\star \mathbf{P} \mathbf{P}^* \mathbf{X}_\star,$$

where $\mathbf{P}^* \mathbf{X}_\star$ is equally sparse as $\mathbf{X}_\star$ and $\mathbf{D}_\star \mathbf{P} \in \mathcal{O}(n; \mathbb{R})$. So we can only expect to recover the correct dictionary (and sparse coefficient matrix) *up to an arbitrary signed permutation*. Therefore, we say the ground truth dictionary $\mathbf{D}_o$ is successfully recovered, if any *signed permuted* version $\mathbf{D}_o \mathbf{P}$ is found.⁷ Unlike approaches (Spielman et al., 2012; Sun et al., 2015; Bai et al., 2018) that solve one column at a time for $\mathbf{D}$, we here attempt recover the entire dictionary $\mathbf{D}$ from solving the problem (6). The signed permutation ambiguity would create numerous equivalent global maximizers, which poses a serious challenge to analysis and optimization.

In this paper, we adopt the Bernoulli-Gaussian model as in prior works (Spielman et al., 2012; Sun et al., 2015; Bai et al., 2018). We assume our observation matrix $\mathbf{Y} \in \mathbb{R}^{n \times p}$ is produced by the product of a ground truth orthogonal dictionary $\mathbf{D}_o$, and a Bernoulli-Gaussian matrix $\mathbf{X}_o \in \mathbb{R}^{n \times p}$:

$$\mathbf{Y} = \mathbf{D}_o \mathbf{X}_o, \quad \mathbf{D}_o \in \mathcal{O}(n; \mathbb{R}), \quad \{\mathbf{X}_o\}_{i,j} \sim_{iid} \text{BG}(\theta). \quad (8)$$

7. In mathematical terms, we are looking for a solution in the quotient space between the orthogonal group and the signed permutation group: $\mathcal{O}(n; \mathbb{R})/\mathcal{SP}(n)$.

The Bernoulli-Gaussian model can be considered as a prototype for dictionary learning because one may adjust θ to control the sparsity level of the ground truth $\mathbf{X}_o$. With the Bernoulli-Gaussian assumption, we now state our main result.

1.3.1. CORRECTNESS OF THE PROPOSED OBJECTIVE FUNCTION

Theorem 1 (Correctness of Global Optima, informal version of Theorem 7) $\forall \theta \in (0, 1)$, let $\mathbf{X}_o \in \mathbb{R}^{n \times p}$, $x_{i,j} \sim_{iid} BG(\theta)$, $\mathbf{D}_o \in \mathcal{O}(n; \mathbb{R})$ an arbitrary orthogonal matrix, and $\mathbf{Y} = \mathbf{D}_o \mathbf{X}_o$. Suppose $\hat{\mathbf{A}}_\star$ is a global maximizer of the optimization problem:

$$\max_{\mathbf{A}} \|\mathbf{A}\mathbf{Y}\|_4^4, \quad \text{subject to} \quad \mathbf{A} \in \mathcal{O}(n; \mathbb{R}), \quad (9)$$

then for any $\varepsilon \in [0, 1]$, there exists a signed permutation matrix $\mathbf{P} \in \mathcal{SP}(n)$, such that

$$\frac{1}{n} \left\| \hat{\mathbf{A}}_\star^* - \mathbf{D}_o \mathbf{P} \right\|_F^2 \leq C\varepsilon, \quad (10)$$

holds with high probability, when p is large enough, and C is a constant depends on θ .

Theorem 1 is obtained through the following line of reasoning: 1) $\forall \mathbf{A} \in \mathcal{O}(n; \mathbb{R})$, we can view $\frac{1}{p} \|\mathbf{A}\mathbf{Y}\|_4^4 = \frac{1}{p} \|\mathbf{A}\mathbf{D}_o \mathbf{X}_o\|_4^4 = \frac{1}{p} \sum_{j=1}^p \|\mathbf{A}\mathbf{D}_o \mathbf{x}_j\|_4^4$ as the mean of p i.i.d. random variables, so it will concentrate to its expectation $\mathbb{E}_{\mathbf{x}_j} \|\mathbf{A}\mathbf{D}_o \mathbf{x}_j\|_4^4$; 2) the value of $\mathbb{E}_{\mathbf{x}_j} \|\mathbf{A}\mathbf{D}_o \mathbf{x}_j\|_4^4$ is largely characterized by $\|\mathbf{A}\mathbf{D}_o\|_4^4$, a deterministic function (with respect to $\mathbf{A}$) on the orthogonal group, whose global optima $\mathbf{A}_\star$ satisfy $\mathbf{A}_\star = \mathbf{P}^* \mathbf{D}_o^*$, $\forall \mathbf{P} \in \mathcal{SP}(n)$. Therefore, when p is large enough, maximizing $\|\mathbf{A}\mathbf{Y}\|_4^4$ over the orthogonal group is equivalent to (with high probability) maximizing the deterministic objective $\|\mathbf{A}\mathbf{D}_o\|_4^4$ over the orthogonal group, which yields the desire result. Formal statements and proofs are given in Section 2 and the Appendices.

1.3.2. A FAST OPTIMIZATION ALGORITHM

Unlike almost all previous algorithms that find the correct dictionary one column $\mathbf{d}_i$ at a time, in Section 3 we introduce a novel *matching, stretching, and projection* (MSP) algorithm that solves the program (9) directly for the entire $\mathbf{D} \in \mathcal{O}(n; \mathbb{R})$. The MSP algorithm directly computes (transpose of) the optimal dictionary $\mathbf{A}_\star$ as the “fixed point” to the following iteration:

$$\mathbf{A}_{t+1} = \mathcal{P}_{\mathcal{O}(n; \mathbb{R})} [(\mathbf{A}_t \mathbf{Y})^{\circ 3} \mathbf{Y}^*], \quad (11)$$

where $\mathcal{P}_{\mathcal{O}(n; \mathbb{R})}(\cdot)$ is the projection onto the orthogonal group $\mathcal{O}(n; \mathbb{R})$, which can be easily calculated from SVD (see Lemma 9). Meanwhile, the MSP algorithm also efficiently maximizes the deterministic objective $\|\mathbf{A}\mathbf{D}_o\|_4^4$ over $\mathcal{O}(n; \mathbb{R})$ by the following iteration:

$$\mathbf{A}_{t+1} = \mathcal{P}_{\mathcal{O}(n; \mathbb{R})} [(\mathbf{A}_t \mathbf{D}_o)^{\circ 3} \mathbf{D}_o^*]. \quad (12)$$

Statistical analysis for proving Theorem 1 suggests that estimates from random samples converge to their expectation. For the deterministic objective $\|\mathbf{A}\mathbf{D}_o\|_4^4$, we further show that the proposed algorithm converges fast with a *cubic rate* around each global maximizer.

Essentially, the update of the MSP algorithm (11), (12) are performing projected gradient ascend with *infinite* step size with respect to objective function $\|\mathbf{A}\mathbf{Y}\|_4^4$, $\|\mathbf{A}\mathbf{D}_o\|_4^4$, over

$O(n; \mathbb{R})$ respectively. The MSP algorithm is similar to the FastICA algorithm (Hyvärinen and Oja, 1997) for independent component analysis (ICA) and such similarity is extensively discussed in a follow-up work (Zhai et al., 2019).

Theorem 2 (Cubic Convergence Rate, informal version of Theorem 16) *Given an orthogonal matrix $\mathbf{A} \in O(n; \mathbb{R})$, if $\|\mathbf{A} - \mathbf{I}\|_F^2 = \varepsilon$ for a small $\varepsilon < 0.579$, and let $\mathbf{A}'$ denote the result after one iteration of our proposed MSP algorithm (12), then we have $\|\mathbf{A}' - \mathbf{I}\|_F^2 \leq O(\varepsilon^3)$.*

In Theorem 2, showing the local convergence of to identity matrix a general orthogonal matrix $\mathbf{A}$ to $\mathbf{I}$ suffices to characterize the local convergence of $\mathbf{A}\mathbf{D}_o$ to a signed permutation matrix $\mathbf{P}$, because of the signed permutation symmetry in the ℓ^4 norm and the orthogonal invariant of SVD. More details are provided in Section 2 and Section 3.

Although Theorem 1 characterizes the properties of the randomized objective $\|\mathbf{A}\mathbf{Y}\|_4^4$ while Theorem 2 describes a deterministic result, they are highly related with each other. We use iteration (11) to maximize the randomized objective $\|\mathbf{A}\mathbf{Y}\|_4^4$ of Theorem 1 and Theorem 2 characterizes the local convergence of iteration (12). $(\mathbf{A}\mathbf{Y})^{\odot 3}\mathbf{Y}^*$ in (11) concentrates to its expectation when there is enough samples and the expectation $\mathbb{E}[(\mathbf{A}\mathbf{Y})^{\odot 3}\mathbf{Y}^*]$ highly depends on $(\mathbf{A}\mathbf{D}_o)^{\odot 3}\mathbf{D}_o$ of (12). Such algorithmic relationship between (11) and (12) is discussed in Proposition 10 and Proposition 11 of Section 3.

As our algorithm is very efficient and scalable, we can test it over very large range of dimensions and settings. Extensive simulations suggest that the MSP algorithm converges globally to the correct solution under broad conditions. We give a global convergence proof for the case $n = 2$ (on $O(2; \mathbb{R})$) and conjecture that similar results hold for arbitrary n (under mild conditions). Extensive experiments show that the algorithm is far more efficient than existing heuristic algorithms and (Riemannian) gradient or subgradient based algorithms. With this efficient algorithm, we characterize empirically the range of success for the program (9), which goes well beyond any existing theoretical guarantees (Sun et al., 2015; Bai et al., 2018) for the complete dictionary case.

1.3.3. OBSERVATIONS AND IMPLICATIONS

Notice that at first sight, the optimization problem associated with dictionary learning is highly nonconvex, with nontrivial orthogonal constraints, and with numerous critical points and ambiguities. Hence, understandably, many recent approaches focus on introducing regularization to the objective function so as to relax the constraints (Aharon et al., 2006; Mairal et al., 2008, 2009, 2012) or analyzing and utilizing local information and design heuristic or gradient descent schemes for such objective functions (with or without regularization) (Agarwal et al., 2014; Arora et al., 2015; Ge et al., 2015; Ma et al., 2017; Li and Liang, 2018; Du et al., 2018; Allen-Zhu et al., 2018; Davis et al., 2018).

However, in our work, we have observed a surprising phenomenon that is rather contrary to conventional views: The discrete signed permutation symmetry $\mathbf{SP}(n)$ associated with the orthogonal group $O(n; \mathbb{R})$ are beneficial. In our work, they both play an important role in making the algorithm efficient and effective, instead of being nuisances or difficulties to be dealt with. As we will see, such discrete symmetry of the orthogonal group makes the global landscape of the objective function amenable to global convergence and enables a fixed-point

type algorithm that converges at a super-linear rate to a correct solution in the quotient space $O(n; \mathbb{R})/SP(n)$. Similar phenomena that symmetry facilitates global convergence of non-convex programs have been also been observed and reported in recent works (Ge et al., 2015; Sun et al., 2015; Zhang et al., 2018; Li and Bresler, 2018; Chi et al., 2018; Kuo et al., 2019). This is a direction that deserves better and deeper study in the future which encourages significant confluence of geometry, algebra, and statistics.

1.4. Notations

We use a bold uppercase and a bold lowercase letter to denote a matrix and a vector, respectively: $\mathbf{X} \in \mathbb{R}^{n \times p}$, $\mathbf{x} \in \mathbb{R}^n$. Moreover, for a matrix $\mathbf{X} \in \mathbb{R}^{n \times p}$, we use $\mathbf{x}_j \in \mathbb{R}^n, \forall j \in [p]$ to denote its j^{th} column vector as default. We use $\mathbf{X}^*$ or $\mathbf{x}^*$ to denote the (conjugate) transpose of a matrix or a vector, respectively. We reserve lower-case letter for scalar: $x \in \mathbb{R}$. We use $\|\mathbf{X}\|_4$ to denote the element-wise ℓ^4 -norm of a matrix $\mathbf{X}$ ($\|\mathbf{X}\|_4^4 = \sum_{i,j} x_{i,j}^4$). We use $\mathbf{D}_o$ to denote the ground truth orthogonal dictionary, and $\mathbf{A}$ is an estimate of $\mathbf{D}_o^*$ from solving (9). $\theta \in (0, 1)$ is sparsity level of the ground truth Bernoulli-Gaussian sparse coefficient $\mathbf{X}_o$: $\mathbf{X}_o \sim \text{BG}(\theta)$. We use $\circ$ to denote the Hadamard product: $\forall \mathbf{A}, \mathbf{B} \in \mathbb{R}^{n \times m}$, $\{\mathbf{A} \circ \mathbf{B}\}_{i,j} = a_{i,j} b_{i,j}$, and $\{\mathbf{A}^{\circ r}\}_{i,j} = a_{i,j}^r$ is the element-wise r^{th} power of $\mathbf{A}$.

Given an input data matrix $\mathbf{Y}$ randomly generated from $\mathbf{Y} = \mathbf{D}_o \mathbf{X}_o$, $\mathbf{X}_o \sim_{iid} \text{BG}(\theta)$, for any orthogonal matrix $\mathbf{A} \in O(n; \mathbb{R})$, we define $\hat{f} : O(n; \mathbb{R}) \times \mathbb{R}^{n \times p} \mapsto \mathbb{R}$ as the 4th power of ℓ^4 -norm of $\mathbf{A}\mathbf{Y}$:

$$\hat{f}(\mathbf{A}, \mathbf{Y}) \doteq \|\mathbf{A}\mathbf{Y}\|_4^4. \quad (13)$$

We define $f : O(n; \mathbb{R}) \mapsto \mathbb{R}$ as the expectation of $\hat{f}$ over $\mathbf{X}_o$:

$$f(\mathbf{A}) \doteq \mathbb{E}_{\mathbf{X}_o}[\hat{f}(\mathbf{A}, \mathbf{Y})] = \mathbb{E}_{\mathbf{X}_o}[\|\mathbf{A}\mathbf{Y}\|_4^4]. \quad (14)$$

For any orthogonal matrix $\mathbf{W} \in O(n; \mathbb{R})$, we define $g : O(n; \mathbb{R}) \mapsto \mathbb{R}$ as 4th power of its ℓ^4 -norm:

$$g(\mathbf{W}) \doteq \|\mathbf{W}\|_4^4. \quad (15)$$

1.5. Organization of this Paper

Rest of the paper is organized as follows. In Section 2, we characterize the global maximizers of (9) statistically via measure concentration. In Section 3, we describe the proposed MSP algorithm. We characterize fixed points of the algorithm and show its convergence results in Section 4. All related proofs can be found in the appendices. Finally, in Section 5, we conduct extensive experiments to show effectiveness and efficiency of our method, by comparing with the state of the art.

2. Key Analysis and Main Result

2.1. Expectation and Concentration of the ℓ^4 -Objective

In this section, we statistically justify that one can recover the ground truth dictionary $\mathbf{D}_o$ by solving

$$\max_{\mathbf{A}} \hat{f}(\mathbf{A}, \mathbf{Y}) = \|\mathbf{A}\mathbf{Y}\|_4^4, \quad \text{subject to} \quad \mathbf{A} \in O(n; \mathbb{R}). \quad (16)$$

Notice that the solution to the above ℓ^4 -norm optimization problem:

$$\hat{\mathbf{A}}_\star = \arg \max_{\mathbf{A} \in \mathbf{O}(n; \mathbb{R})} \hat{f}(\mathbf{A}, \mathbf{Y})$$

is a random variable that depends on the random samples $\mathbf{Y}$. We need to characterize how “close” an estimate $\hat{\mathbf{A}}_\star$ is to the ground truth $\mathbf{D}_o$. A key technique is to show that the random objective function actually concentrates on its expectation (a deterministic function) as the number of observations p increases. We first calculate $f(\mathbf{A})$, the expectation of $\hat{f}(\mathbf{A}, \mathbf{Y})$ and provide its concentration bound in Lemma 3 and Lemma 4 respectively.

Lemma 3 (Properties of $f(\mathbf{A})$) $\forall \theta \in (0, 1)$, let $\mathbf{X}_o \in \mathbb{R}^{n \times p}$, $x_{i,j} \sim_{iid} BG(\theta)$, $\mathbf{D}_o \in \mathbf{O}(n; \mathbb{R})$ is an orthogonal matrix, and $\mathbf{Y} = \mathbf{D}_o \mathbf{X}_o$. Then, $\forall \mathbf{A} \in \mathbf{O}(n; \mathbb{R})$, we have

$$\frac{1}{3p\theta} f(\mathbf{A}) = (1 - \theta)g(\mathbf{A}\mathbf{D}_o) + \theta n. \quad (17)$$

Proof See A.1 ■

Lemma 4 (Concentration Bound of ℓ^4 -Norm) $\forall \theta \in (0, 1)$, if $\mathbf{X} \in \mathbb{R}^{n \times p}$, $x_{i,j} \sim_{iid} BG(\theta)$, $\forall \delta > 0$, the following inequality holds:

$$\mathbb{P} \left(\sup_{\mathbf{W} \in \mathbf{O}(n; \mathbb{R})} \frac{1}{np} \left| \|\mathbf{W}\mathbf{X}\|_4^4 - \mathbb{E} \|\mathbf{W}\mathbf{X}\|_4^4 \right| \geq \delta \right) < \frac{1}{p}, \quad (18)$$

when $p = \Omega(\theta n^2 \ln n / \delta^2)$.

Proof See A.2. ■

Lemma 4 implies that, for any orthogonal transformation $\mathbf{W}\mathbf{X}$, $\mathbf{W} \in \mathbf{O}(n; \mathbb{R})$ of a Bernoulli-Gaussian matrix $\mathbf{X} \in \mathbb{R}^{n \times p}$, $\frac{1}{np} \|\mathbf{W}\mathbf{X}\|_4^4$ concentrates onto its expectation as long as the sample size p is large enough – in the order $\Omega(\theta n^2 \ln n / \delta^2)$. By our definition, $\hat{f}(\mathbf{A}, \mathbf{Y}) = \|\mathbf{A}\mathbf{Y}\|_4^4 = \|\mathbf{W}\mathbf{X}_o\|_4^4$ ($\mathbf{W} = \mathbf{A}\mathbf{D}_o$ is an orthogonal matrix) satisfies the concentration inequality in Lemma 4. Therefore, including designing optimization algorithms, $f(\mathbf{A})$ can be considered as a good proxy to the original objective $\hat{f}(\mathbf{A}, \mathbf{Y})$ and we can consider (16) as maximizing its expectation:

$$\max_{\mathbf{A}} f(\mathbf{A}) = \mathbb{E} \|\mathbf{A}\mathbf{Y}\|_4^4 \quad \text{subject to } \mathbf{A} \in \mathbf{O}(n; \mathbb{R}). \quad (19)$$

The function $f(\mathbf{A})$ is a deterministic function and its geometric property would help us understand the landscape of the original objective. Moreover, Lemma 3 states that the global maximizers of the mean $f(\mathbf{A})$ are exactly the global maximizers of $g(\mathbf{A}\mathbf{D}_o) = \|\mathbf{A}\mathbf{D}_o\|_4^4$. To understand how such a function can be effectively optimized, we need to study the extrema of ℓ^4 -norm $g(\cdot)$ over the orthogonal group $\mathbf{O}(n; \mathbb{R})$.

2.2. Property of the ℓ^4 -Norm over $O(n; \mathbb{R})$

Lemma 5 (Extrema of ℓ^4 -Norm over Orthogonal Group) *For any orthogonal matrix $\mathbf{A} \in O(n; \mathbb{R})$, $g(\mathbf{A}) = \|\mathbf{A}\|_4^4 \in [1, n]$ and $g(\mathbf{A})$ reaches maximum if and only if $\mathbf{A} \in \text{SP}(n)$.*

Proof See A.3. ■

This lemma implies that if $\mathbf{A}_\star$ is a global maximizer of $g(\mathbf{A}\mathbf{D}_o)$, i.e. $\|\mathbf{A}_\star\mathbf{D}_o\|_4^4 = n$, then it differs from $\mathbf{D}_o^*$ by a signed permutation. However, this lemma does not say the function may or may not have other local minima or maxima. Nevertheless, our experiments in Section 5.2 will show that even if such critical points exist, they are unlikely to be stable (or attractive). As a direct corollary to Lemma 3 and Lemma 5, we know that $\forall \mathbf{A} \in O(n; \mathbb{R}), \theta \in (0, 1)$, the maximum value of $f(\mathbf{A})$ satisfies:

$$\frac{1}{3p\theta} f(\mathbf{A}) \leq n, \quad (20)$$

and the equality holds if and only if $\mathbf{A}\mathbf{D}_o \in \text{SP}(n)$. Although we know that in the stochastic setting, (16) concentrates to the following ℓ^4 -norm maximization over $O(n; \mathbb{R})$:

$$\max_{\mathbf{A}} g(\mathbf{A}\mathbf{D}_o) = \|\mathbf{A}\mathbf{D}_o\|_4^4, \quad \text{subject to } \mathbf{A} \in O(n; \mathbb{R}), \quad (21)$$

we cannot hope that the estimate from $\hat{\mathbf{A}}_\star = \arg \max_{\mathbf{A} \in O(n; \mathbb{R})} \hat{f}(\mathbf{A}, \mathbf{Y})$ would achieve the maximal value of $g(\mathbf{A}\mathbf{D}_o)$ precisely. But if the value is close to the maximal, how close would $\hat{\mathbf{A}}_\star$ be to a global maximizer? The following lemma shows that when the ℓ^4 -norm of an orthogonal matrix $\mathbf{A}$ is close to the maximum value n , it is also close to a signed permutation matrix in Frobenius norm.

Lemma 6 (Approximate Maxima of ℓ^4 -Norm over Orthogonal Group) *Suppose $\mathbf{W}$ is an orthogonal matrix: $\mathbf{W} \in O(n; \mathbb{R})$. $\forall \varepsilon \in [0, 1]$, if $\frac{1}{n} \|\mathbf{W}\|_4^4 \geq 1 - \varepsilon$, then $\exists \mathbf{P} \in \text{SP}(n)$, such that*

$$\frac{1}{n} \|\mathbf{W} - \mathbf{P}\|_F^2 \leq 2\varepsilon. \quad (22)$$

Proof See A.4. ■

This result is useful whenever we evaluate how close a solution given by an algorithm is to the optimal solution, in terms of value of the objective function. With all the above results, we are now ready to characterize in what sense a global maximizer of $\hat{\mathbf{A}}_\star = \arg \max_{\mathbf{A} \in O(n; \mathbb{R})} \hat{f}(\mathbf{A}, \mathbf{Y})$ gives a “correct” estimate of the ground truth dictionary $\mathbf{D}_o$.

2.3. Main Statistical Result

Theorem 7 (Correctness of the Global Optima) $\forall \theta \in (0, 1)$, let $\mathbf{X}_o \in \mathbb{R}^{n \times p}$, $x_{i,j} \sim_{iid} BG(\theta)$, $\mathbf{D}_o \in O(n; \mathbb{R})$ is any orthogonal matrix, and $\mathbf{Y} = \mathbf{D}_o \mathbf{X}_o$. Suppose $\hat{\mathbf{A}}_\star$ is a global maximizer of the optimization problem:

$$\max_{\mathbf{A}} \hat{f}(\mathbf{A}, \mathbf{Y}) = \|\mathbf{A}\mathbf{Y}\|_4^4, \quad \text{subject to } \mathbf{A} \in O(n; \mathbb{R}),$$

then for any $\varepsilon \in [0, 1]$, there exists a signed permutation matrix $\mathbf{P} \in \text{SP}(n)$, such that

$$\frac{1}{n} \left\| \hat{\mathbf{A}}_\star - \mathbf{D}_o \mathbf{P} \right\|_F^2 \leq C\varepsilon, \quad (23)$$

with probability at least $1 - \frac{1}{p}$, when $p = \Omega(\theta n^2 \ln n / \varepsilon^2)$, for a constant $C > \frac{4}{3\theta(1-\theta)}$.

Proof See A.5. ■

Theorem 7 states that with high probability, the global optimal solution to the ℓ^4 -norm optimization (16) is close to the true solution (up to a signed permutation) as the sample size p is of the order $\Omega(\theta n^2 \ln n / \varepsilon^2)$, where ε is the desired accuracy. This result quantifies the sample size needed to achieve certain accuracy of recovery and it matches the intuition that more observations lead to better recovering result. Moreover, this result corresponds to our phase transition curves in Figure 8 of Section 5, and to the best of our knowledge, a sample complexity of $\Omega(n^2 \ln n)$ is currently the best result for dictionary learning task.

Remark 8 (Maximizing ℓ^{2k} -Norm) *If one were to choose maximizing ℓ^{2k} -norm to promoting sparsity, similar analysis of concentration bounds would reveal that for the same error bound, it requires much larger number p of samples for the (random) objective function $\hat{f}(\mathbf{A}, \mathbf{Y})$ to concentrate on its (deterministic) expectation $f(\mathbf{A})$. We will discuss the choice of k in more details in Section 4.4. Experiments in Figure 7 also corroborate with the findings.*

3. Algorithm: Matching, Stretching, and Projection (MSP)

In this section, we introduce an algorithm, based on a simple iterative *matching, stretching, and projection* (MSP) process, which efficiently solves the two related programs (16) and (21).

3.1. Algorithmic Challenges and Related Optimization Methods

Although (16) is everywhere smooth, the associated optimization is non-trivial in several ways. First, one needs to deal with the signed permutation ambiguity. The problem has exponentially many global maximizers. Furthermore, we are maximizing a convex function (or minimizing a concave function) over a constraint set. So conventional methods such as augmented Lagrangian (Bertsekas, 1997) barely works. This is because the Lagrangian:

$$\mathcal{L}(\mathbf{A}, \mathbf{\Lambda}) \doteq -\|\mathbf{A}\mathbf{Y}\|_4^4 + \langle \mathbf{A}\mathbf{A} - \mathbf{I}, \mathbf{\Lambda} \rangle$$

will go to negative infinity due to the concavity of the objective function $-\|\mathbf{A}\mathbf{Y}\|_4^4$. Notice that all of its global maximizers are on the constraint set, an ideal iterative algorithm should converge to a solution that *exactly lies on constraint set* $\mathbf{O}(n; \mathbb{R})$.

Another natural way to optimize (16) (or (21)) is to apply Riemannian gradient (or projected gradient) type methods (Edelman et al., 1998; Absil et al., 2009) on $\mathbf{O}(n; \mathbb{R})$. One can take small gradient steps to ensure convergence and such methods converge at best with a linear rate (say, if the objective function is strongly convex). Nevertheless, by better utilizing the special global geometry of the parameter space (the orthogonal group), we can choose an arbitrary large step size and the process converges much more rapidly (with a superlinear rate).

We next introduce a very effective and efficient algorithm to solve problems (16) and (21), that is based on simple iterative matching, stretching and projection (MSP) operators.

Mathematically, our algorithm for solving (16) or (21) are essentially doing projected gradient ascent on objective $\|\mathbf{A}\mathbf{D}_o\|_4^4$ or $\|\mathbf{A}\mathbf{Y}\|_4^4$ respectively, but with *infinite* step size. Such infinite step size gradient method is able to find global maximum of our proposed objective (16) and (21) because it exploits the coincidence between the signed permutation symmetry in our formulation and the symmetry of $\text{SP}(n)$ over $\text{O}(n; \mathbb{R})$.

3.2. ℓ^4 -Norm Maximization over the Orthogonal Group

Since the objective function of dictionary learning (16) concentrates to the ℓ^4 -norm maximization problem (21) w.h.p., we first introduce our algorithm for the simpler (deterministic) case:

$$\max_{\mathbf{A}} g(\mathbf{A}\mathbf{D}_o) = \|\mathbf{A}\mathbf{D}_o\|_4^4, \quad \text{subject to } \mathbf{A} \in \text{O}(n; \mathbb{R}). \quad (24)$$

In this setting, we only have information of the values of $g(\mathbf{A}\mathbf{D}_o)$ and $\nabla_{\mathbf{A}} g(\mathbf{A}\mathbf{D}_o)$. We want to update $\mathbf{A}$ to recover $\mathbf{D}_o^*$ based on these information. We use the following lemma to enforce the orthogonal constraint.

Lemma 9 (Projection onto the Orthogonal Group) $\forall \mathbf{A} \in \mathbb{R}^{n \times n}$, the orthogonal matrix which is closest to $\mathbf{A}$ in Frobenius norm is the following:

$$\mathcal{P}_{\text{O}(n; \mathbb{R})}(\mathbf{A}) = \arg \min_{\mathbf{M} \in \text{O}(n; \mathbb{R})} \|\mathbf{A} - \mathbf{M}\|_F^2 = \mathbf{U}\mathbf{V}^*, \quad (25)$$

where $\mathbf{U}\mathbf{\Sigma}\mathbf{V}^* = \text{SVD}(\mathbf{A})$.

Proof See B.1. ■

The MSP algorithm that maximizing (21) is outlined as Algorithm 1.

Algorithm 1 MSP for ℓ^4 -Maximization over $\text{O}(n; \mathbb{R})$

- 1: **Initialize** $\mathbf{A}_0 \in \text{O}(n, \mathbb{R})$ ▷ Initialize $\mathbf{A}_0$ for iteration
 - 2: **for** $t = 0, 1, \dots, T - 1$ **do**
 - 3: $\partial \mathbf{A}_t \doteq 4(\mathbf{A}_t \mathbf{D}_o)^{\circ 3} \mathbf{D}_o^*$ ▷ $\nabla_{\mathbf{A}} \|\mathbf{A}\mathbf{D}_o\|_4^4 = 4(\mathbf{A}\mathbf{D}_o)^{\circ 3} \mathbf{D}_o^*$
 - 4: $\mathbf{U}\mathbf{\Sigma}\mathbf{V}^* = \text{SVD}(\partial \mathbf{A}_t)$
 - 5: $\mathbf{A}_{t+1} = \mathbf{U}\mathbf{V}^*$ ▷ Project $\mathbf{A}_{t+1}$ onto orthogonal group
 - 6: **end for**
 - 7: **Output** $\mathbf{A}_T$
-

In each iteration of Algorithm 1, we use $\|\mathbf{A}\mathbf{D}_o\|_4^4/n$ to evaluate how “close” $\mathbf{A}\mathbf{D}_o$ is to a signed permutation matrix, since Lemma 5 shows the global maximum of $\|\mathbf{A}\mathbf{D}_o\|_4^4$ is n and the maximal evaluation is therefore normalized to 1. In Step 3 of the MSP algorithm, the calculation of $\partial \mathbf{A}_t = 4(\mathbf{A}_t \mathbf{D}_o)^{\circ 3} \mathbf{D}_o^*$ does not require knowledge of $\mathbf{D}_o$. It is merely the gradient of the objective function:

$$\nabla_{\mathbf{A}} g(\mathbf{A}\mathbf{D}_o) = \nabla_{\mathbf{A}} \|\mathbf{A}\mathbf{D}_o\|_4^4 = 4(\mathbf{A}\mathbf{D}_o)^{\circ 3} \mathbf{D}_o^*.$$

As the name of the algorithm suggests, each iteration actually performs a “*matching, stretching, and projection*” operation: It first matches the current estimate $\mathbf{A}_t$ to the true

$\mathbf{D}_o$. Then the element-wise cubic function $(\cdot)^{\circ 3}$ stretches all entries of $\mathbf{A}_t \mathbf{D}_o$ by promoting the large ones and suppressing the small ones. $\partial \mathbf{A}_t$ is the correlation between so “sparsified” pattern and the original basis $\mathbf{D}_o^*$, which is then projected back onto the closest orthogonal matrix $\mathbf{A}_{t+1}$ in Frobenius distance.

Repeating this “matching, stretching, and projection” process, $\mathbf{A}_t \mathbf{D}_o$ is increasingly sparsified while ensuring the orthogonality of each $\mathbf{A}_t$. Ideally the process will stop when $\mathbf{A}_t \mathbf{D}_o$ becomes the sparsest, that is, a signed permutation matrix. The iterative MSP algorithm utilizes the global geometry of the orthogonal group and acts more like the *power iteration* method or the fixed point algorithm (Hyvärinen and Oja, 1997). It is easy to see that for any fixed $\mathbf{D}_o$, the optimal $\mathbf{A}_\star$ is the “fixed point” to the following equation:

$$\mathbf{A}_\star = \mathcal{P}_{\mathbf{O}(n; \mathbb{R})} [(\mathbf{A}_\star \mathbf{D}_o)^{\circ 3} \mathbf{D}_o^*], \quad (26)$$

where $\mathcal{P}_{\mathbf{O}(n; \mathbb{R})}$ is the projection onto the orthogonal group $\mathbf{O}(n; \mathbb{R})$. The proposed “matching, stretching, and projection” algorithm is precisely to compute the fixed point of this equation in the most natural way! Our analysis (Theorem 16) will show that this scheme converges extremely well and actually achieves a *super-linear* local convergence rate.

Example 1 (One Run of Algorithm 1) *To help better visualize how well the algorithm works, we consider a special case when $\mathbf{D}_o = \mathbf{I}$. The problem reduces to finding a matrix with the maximum ℓ^4 -norm over the orthogonal group:*

$$\max_{\mathbf{A}} g(\mathbf{A}) \doteq \|\mathbf{A}\|_4^4, \quad \text{subject to} \quad \mathbf{A} \in \mathbf{O}(n; \mathbb{R}). \quad (27)$$

We randomly initialize the MSP algorithm with an orthogonal matrix $\mathbf{A}_0 \in \mathbb{R}^{3 \times 3}$, and the sequences below show how the quickly the MSP algorithm quickly converges to a signed permutation matrix:

$$\begin{array}{ll} \mathbf{A}_0 = \begin{pmatrix} -0.8249 & 0.3820 & -0.4168 \\ -0.5240 & -0.2398 & 0.8173 \\ -0.2122 & -0.8925 & -0.3979 \end{pmatrix} & \xrightarrow{\text{stretching}} \mathbf{A}_0^{\circ 3} = \begin{pmatrix} -0.5613 & 0.0557 & -0.0724 \\ -0.1439 & -0.0138 & 0.5459 \\ -0.0096 & -0.7109 & -0.0630 \end{pmatrix} \\ \xrightarrow{\text{projection}} \mathbf{A}_1 = \begin{pmatrix} -0.9795 & 0.0621 & -0.1917 \\ -0.1953 & -0.0594 & 0.9789 \\ -0.0494 & -0.9963 & -0.0703 \end{pmatrix} & \xrightarrow{\text{stretching}} \mathbf{A}_1^{\circ 3} = \begin{pmatrix} -0.9397 & 0.0002 & -0.0070 \\ -0.0075 & -0.0002 & 0.9381 \\ -0.0001 & -0.9889 & -0.0003 \end{pmatrix} \\ \xrightarrow{\text{projection}} \mathbf{A}_2 = \begin{pmatrix} -1.0000 & 0.0002 & -0.0077 \\ -0.0077 & -0.0003 & 1.000 \\ -0.0002 & -1.0000 & -0.0003 \end{pmatrix} & \xrightarrow{\text{stretching}} \mathbf{A}_2^{\circ 3} = \begin{pmatrix} -0.9999 & 0.0000 & -0.0000 \\ -0.0000 & -0.0000 & 0.9999 \\ -0.0000 & -1.0000 & -0.0000 \end{pmatrix} \\ \xrightarrow{\text{projection}} \mathbf{A}_3 = \begin{pmatrix} -1 & 0 & 0 \\ 0 & 0 & 1 \\ 0 & -1 & 0 \end{pmatrix} & \xrightarrow{\text{output}} \mathbf{A}_3^{\circ 3} = \begin{pmatrix} -1 & 0 & 0 \\ 0 & 0 & 1 \\ 0 & -1 & 0 \end{pmatrix}. \end{array}$$

3.3. ℓ^4 -Norm Maximization for Dictionary Learning

Having understood how to optimize the deterministic case for the expectation, we now consider the original dictionary learning problem (16):

$$\max_{\mathbf{A}} \hat{f}(\mathbf{A}, \mathbf{Y}) = \|\mathbf{A}\mathbf{Y}\|_4^4, \quad \text{subject to} \quad \mathbf{A} \in \mathbf{O}(n; \mathbb{R}).$$

This naturally leads to a similar MSP algorithm, outlined as Algorithm 2.

Algorithm 2 MSP for ℓ^4 -Maximization based Dictionary Learning

```

1: Initialize  $\mathbf{A}_0 \in \mathcal{O}(n, \mathbb{R})$  ▷ Initialize  $\mathbf{A}_0$  for iteration
2: for  $t = 0, 1, \dots, T - 1$  do
3:    $\partial \mathbf{A}_t \doteq 4(\mathbf{A}_t \mathbf{Y})^{\circ 3} \mathbf{Y}^*$  ▷  $\nabla_{\mathbf{A}} \|\mathbf{A} \mathbf{Y}\|_4^4 = 4(\mathbf{A} \mathbf{Y})^{\circ 3} \mathbf{Y}^*$ 
4:    $\mathbf{U} \Sigma \mathbf{V}^* = \text{SVD}(\partial \mathbf{A}_t)$ 
5:    $\mathbf{A}_{t+1} = \mathbf{U} \mathbf{V}^*$  ▷ Project  $\mathbf{A}_{t+1}$  onto orthogonal group
6: end for
7: Output  $\mathbf{A}_T, \|\mathbf{A}_T \mathbf{Y}\|_4^4 / 3np\theta$ 
    
```

Note that in the output we also normalize $\|\mathbf{A} \mathbf{Y}\|_4^4$ by dividing the maximum of its expectation: $3np\theta$ so that the maximal output value would be around 1. Similar to Algorithm 1, we also use $\|\mathbf{A} \mathbf{D}_o\|_4^4 / n$ to evaluate how “close” $\mathbf{A} \mathbf{D}_o$ is to a signed permutation matrix in each iteration.

The same intuition of “matching, stretching, and projection” for the deterministic case naturally carries over here. In Step 3, the estimate $\mathbf{A}_t$ is matched with the observation $\mathbf{Y}$. The cubic function $(\cdot)^{\circ 3}$ re-scales the results and promotes entry-wise spikiness of $\mathbf{X}_t = \mathbf{A}_t \mathbf{Y}$ accordingly. Again, here $\partial \mathbf{A}_t = 4(\mathbf{A}_t \mathbf{Y})^{\circ 3} \mathbf{Y}^*$ is the gradient $\nabla_{\mathbf{A}} \hat{f}(\mathbf{A}, \mathbf{Y})$ of the objective function. However, the algorithm is not performing gradient ascent: The matrix $(\mathbf{A}_t \mathbf{Y})^{\circ 3} \mathbf{Y}^*$ is actually the sample covariance of the following two random vectors: $(\mathbf{A}_t \mathbf{y})^{\circ 3}$ and $\mathbf{y}$. The subsequent projection of this sample covariance matrix onto the orthogonal group $\mathcal{O}(n; \mathbb{R})$ normalizes the scale of the operator $\mathbf{A}_{t+1}$, hence normalize the covariance of $\mathbf{A}_{t+1} \mathbf{Y}$ for the next iteration.

Similar to the deterministic case, for any given sparsely generated data matrix $\mathbf{Y}$, the optimal dictionary $\mathbf{A}_\star$ is the “fixed point” to the following equation:

$$\mathbf{A}_\star = \mathcal{P}_{\mathcal{O}(n; \mathbb{R})}[(\mathbf{A}_\star \mathbf{Y})^{\circ 3} \mathbf{Y}^*], \quad (28)$$

where $\mathcal{P}_{\mathcal{O}(n; \mathbb{R})}$ is the projection onto the orthogonal group $\mathcal{O}(n; \mathbb{R})$. The iterative “matching, stretching, and projection” scheme is precisely to compute the fixed point of this equation in the most natural way!

Although the data and the objective function are random here, Proposition 10 below clarifies the relationship between this expectation and the deterministic gradient $\nabla_{\mathbf{A}} g(\mathbf{A} \mathbf{D}_o)$ and Proposition 11 further shows that $\nabla_{\mathbf{A}} \hat{f}(\mathbf{A}, \mathbf{Y})$ concentrates to its expectation when p increases.

Proposition 10 (Expectation of $\nabla_{\mathbf{A}} \hat{f}(\mathbf{A}, \mathbf{Y})$) *Let $\mathbf{X} \in \mathbb{R}^{n \times p}$, $x_{i,j} \sim_{iid} BG(\theta)$, $\mathbf{D}_o \in \mathcal{O}(n; \mathbb{R})$ is any orthogonal matrix, and $\mathbf{Y} = \mathbf{D}_o \mathbf{X}_o$. The expectation of $\nabla_{\mathbf{A}} \hat{f}(\mathbf{A}, \mathbf{Y})$ satisfies this property:*

$$\mathbb{E}_{\mathbf{X}_o}[\nabla_{\mathbf{A}} \hat{f}(\mathbf{A}, \mathbf{Y})] = 3p\theta(1 - \theta) \nabla_{\mathbf{A}} g(\mathbf{A} \mathbf{D}_o) + 12p\theta^2 \mathbf{A}. \quad (29)$$

Proof See B.2. ■

This proposition indicates that the expected stochastic gradient $\nabla_{\mathbf{A}} \hat{f}(\mathbf{A}, \mathbf{Y})$ agrees well with the gradient of the deterministic objective $g(\mathbf{A} \mathbf{D}_o)$, except for a bias that is linear in the current estimate $\mathbf{A}$. This suggests a possible improvement for the MSP algorithm in the stochastic case: If we have knowledge about the θ (or can estimate it online), we could

subtract a bias term $\alpha \mathbf{A}$ ($\alpha \in (0, 12p\theta^2]$) from $\partial \mathbf{A}_t$ in Step 3 of Algorithm 2. One can verify experimentally that this indeed helps further accelerate the convergence of the algorithm.

Proposition 11 (Concentration Bound of $\frac{1}{np} \nabla \hat{f}(\cdot, \cdot)$) *If $\mathbf{X}_o \in \mathbb{R}^{n \times p}$, $x_{i,j} \sim_{iid} BG(\theta)$, for any $\mathbf{A}, \mathbf{D}_o \in \mathcal{O}(n; \mathbb{R})$, and $\mathbf{Y} = \mathbf{D}_o \mathbf{X}_o$, the following inequality holds*

$$\mathbb{P} \left(\sup_{\mathbf{A} \in \mathcal{O}(n; \mathbb{R})} \frac{1}{4np} \left\| \nabla_{\mathbf{A}} \hat{f}(\mathbf{A}, \mathbf{Y}) - \mathbb{E}[\nabla_{\mathbf{A}} \hat{f}(\mathbf{A}, \mathbf{Y})] \right\|_F \geq \delta \right) < \frac{1}{p}, \quad (30)$$

when $p = \Omega(\theta n^2 \ln n / \delta^2)$.

Proof See B.3. ■

4. Analysis of the MSP Algorithm

In this section, we provide convergence analysis of the proposed MSP Algorithm 1 that maximizes $g(\mathbf{A} \mathbf{D}_o)$ over the orthogonal group $\mathcal{O}(n; \mathbb{R})$. Each iteration of Algorithm 1 performs the following iteration:

$$\mathbf{A}_{t+1} = \mathcal{P}_{\mathcal{O}(n; \mathbb{R})}[(\mathbf{A}_t \mathbf{D}_o)^{\circ 3} \mathbf{D}_o^*], \quad (31)$$

notice that both $\mathbf{A}$ and $\mathbf{D}_o$ are orthogonal matrices, we can further reduce (31) to

$$\mathbf{A}_{t+1} \mathbf{D}_o = \mathcal{P}_{\mathcal{O}(n; \mathbb{R})}[(\mathbf{A}_t \mathbf{D}_o)^{\circ 3}]. \quad (32)$$

If we view $\mathbf{A}_t \mathbf{D}_o$ as another orthogonal matrix $\mathbf{W}_t$, the convergence analysis reduces to prove that the following iteration

$$\mathbf{W}_{t+1} = \mathcal{P}_{\mathcal{O}(n; \mathbb{R})}[(\mathbf{W}_t)^{\circ 3}] \quad (33)$$

will converge to a signed permutation matrix $\mathbf{W}_\infty$. Hence, we conclude that the MSP algorithm 1 for maximizing $g(\mathbf{A} \mathbf{D}_o)$ is invariant of orthogonal rotation. So without loss of generality, we only need to provide convergence analysis for the case $\mathbf{D}_o = \mathbf{I}$ (or slightly abuse the notation a bit by changing $\mathbf{W}_t$ in (33) into $\mathbf{A}_t$).

When $\mathbf{D}_o = \mathbf{I}$, we wish to show the MSP algorithm converges to a signed permutation matrix for the optimization problem:

$$\max_{\mathbf{A}} g(\mathbf{A}) = \|\mathbf{A}\|_4^4, \quad \text{subject to} \quad \mathbf{A} \in \mathcal{O}(n; \mathbb{R}), \quad (34)$$

starting from any randomly initialized $\mathbf{A}_0$ on $\mathcal{O}(n; \mathbb{R})$ with probability 1. For this purpose, we first introduce some basic properties of the space $\mathcal{O}(n; \mathbb{R})$ and our objective function $g(\cdot)$.

4.1. Properties of the Orthogonal Group

The orthogonal group is a special type of Stiefel manifold (Absil et al., 2009) with tangent space

$$T_{\mathbf{W}} \mathcal{O}(n; \mathbb{R}) \doteq \{\mathbf{Z} \mid \mathbf{Z}^* \mathbf{W} + \mathbf{W}^* \mathbf{Z} = \mathbf{0}\}, \quad (35)$$

and the projection operation $\mathcal{P}_{T_{\mathbf{W}}\mathcal{O}(n;\mathbb{R})} : \mathbb{R}^{n \times n} \rightarrow T_{\mathbf{W}}\mathcal{O}(n;\mathbb{R})$ onto tangent space of $\mathcal{O}(n;\mathbb{R})$ is defined as:

$$\mathcal{P}_{T_{\mathbf{W}}\mathcal{O}(n;\mathbb{R})}(\mathbf{Z}) \doteq \mathbf{Z} - \frac{1}{2}\mathbf{W}(\mathbf{W}^*\mathbf{Z} + \mathbf{Z}^*\mathbf{W}) = \frac{1}{2}(\mathbf{Z} - \mathbf{W}\mathbf{Z}^*\mathbf{W}). \quad (36)$$

We use $\nabla_{\mathbf{W}}g(\mathbf{W})$ to denote the gradient of $g(\mathbf{W})$ w.r.t. $\mathbf{W}$ in $\mathbb{R}^{n \times n}$, and $\text{grad}g(\mathbf{W})$ to denote the Riemannian gradient of $g(\mathbf{W})$ w.r.t. $\mathbf{W}$ on $T_{\mathbf{W}}\mathcal{O}(n;\mathbb{R})$. Thus, we can formulate the Riemannian gradient of $\mathbf{W}$ on $T_{\mathbf{W}}\mathcal{O}(n;\mathbb{R})$ as following:

$$\text{grad}g(\mathbf{W}) = \mathcal{P}_{T_{\mathbf{W}}\mathcal{O}(n;\mathbb{R})}(\nabla_{\mathbf{W}}g(\mathbf{W})). \quad (37)$$

The following proposition introduces the critical points of $g(\mathbf{W})$ on $T_{\mathbf{W}}\mathcal{O}(n;\mathbb{R})$.

Proposition 12 *The critical points of $g(\mathbf{W})$ on manifold $\mathcal{O}(n;\mathbb{R})$ satisfies the following condition:*

$$(\mathbf{W}^{\circ 3})^*\mathbf{W} = \mathbf{W}^*\mathbf{W}^{\circ 3}. \quad (38)$$

Proof See C.1. ■

Therefore, $\forall \mathbf{W} \in \mathbb{R}^{n \times n}$ we can write critical points condition of ℓ^4 -norm over $\mathcal{O}(n;\mathbb{R})$ as this following equations:

$$\begin{cases} (\mathbf{W}^{\circ 3})^*\mathbf{W} = \mathbf{W}^*\mathbf{W}^{\circ 3}, \\ \mathbf{W}^*\mathbf{W} = \mathbf{I}. \end{cases} \quad (39)$$

Since the orthogonal group $\mathcal{O}(n;\mathbb{R})$ is a continuous manifold in $\mathbb{R}^{n \times n}$ (Absil et al., 2009; Hall, 2015), this indicates that critical points of $g(\mathbf{W}) = \|\mathbf{W}\|_4^4$ over the orthogonal group $\mathbf{W} \in \mathcal{O}(n;\mathbb{R})$ has measure 0.

Proposition 13 *All global maximizers of ℓ^4 -norm over the orthogonal group are isolated (nondegenerate) critical points.*

Proof See C.2. ■

We conjecture that the all local maximizers of the ℓ^4 -norm over $\mathcal{O}(n;\mathbb{R})$ are global maximizers, as we will discuss in the next subsection.

4.2. Relation between MSP and Projected Gradient Ascent (PGA)

Although the MSP algorithm over orthogonal group (Algorithm 1) and for dictionary learning (Algorithm 2) has the same optimization procedure – they all performs *infinite* step size projected gradient ascent w.r.t. the objective function $\|\mathbf{A}\mathbf{D}_o\|_4^4$ and $\|\mathbf{A}\mathbf{Y}\|_4^4$ respectively, we shall state their intrinsic difference clearly.

In Proposition 10 and Proposition 11, we can see that each iteration of the MSP algorithm for dictionary learning (Algorithm 2) concentrates onto one step of PGA w.r.t. the objective $\|\mathbf{A}\mathbf{D}_o\|_4^4$ with *fixed step size* $\frac{1-\theta}{4\theta}$, while the MSP algorithm over the orthogonal group performs PGA with *infinite* step size. We test PGA Algorithm 3 with different step size α , as shown in Table 1 for detail. We observe that:

- PGA with *arbitrary* fixed step size find *global maximizers* of (34);

- PGA converges *faster* with *larger* step size.

This experimental phenomenon supports the efficiency of the MSP algorithm – it is indeed faster than any first order gradient method.

Algorithm 3 PGA for ℓ^4 -Maximization over $\mathcal{O}(n; \mathbb{R})$

```

1: Initialize  $\mathbf{A}_0 \in \mathcal{O}(n; \mathbb{R})$ , step size  $\alpha > 0$  ▷ Initialize  $\mathbf{A}_0$ , and  $\alpha$  for iteration
2: for  $t = 0, 1, \dots, T - 1$  do
3:    $\partial \mathbf{A}_t \doteq 4\mathbf{A}^{\odot 3}$  ▷  $\nabla_{\mathbf{A}} \|\mathbf{A}\|_4^4 = 4(\mathbf{A})^{\odot 3}$ 
4:    $\mathbf{U}\mathbf{\Sigma}\mathbf{V}^* = \text{SVD}(\mathbf{A} + \alpha\partial \mathbf{A}_t)$ 
5:    $\mathbf{A}_{t+1} = \mathbf{U}\mathbf{V}^*$  ▷ Project  $\mathbf{A}_{t+1}$  onto  $\mathcal{O}(n; \mathbb{R})$ 
6: end for
7: Output  $\mathbf{A}_T$ 
    
```

	Iterations			
	$\alpha = 1$	$\alpha = 10$	$\alpha = 100$	$\alpha = +\infty$
$\mathcal{O}(5; \mathbb{R})$	13	5	4	4
$\mathcal{O}(25; \mathbb{R})$	23	7	5	5
$\mathcal{O}(50; \mathbb{R})$	35	10	8	6
$\mathcal{O}(100; \mathbb{R})$	63	12	9	9
$\mathcal{O}(200; \mathbb{R})$	70	14	11	9

Table 1: Number of iteration for PGA (Algorithm 3) on orthogonal group of different dimension n to reach global maximizers, with the same initialization for each dimension n and different step size α . We directly apply the MSP Algorithm 1 when $\alpha = +\infty$.

4.3. Convergence Analysis

We first introduce the properties of the critical points of the ℓ^4 -objective (34) in Proposition 14 and Proposition 15 will show that PGA with *any fixed step size* α (even $\alpha = +\infty$) finds a critical points of (34).

Proposition 14 (Fixed Point of the MSP Algorithm) *Given $\mathbf{W} \in \mathcal{O}(n; \mathbb{R})$, $\mathbf{W}$ is a fix point of the MSP Algorithm 1 if and only if $\mathbf{W}$ is a critical point of the ℓ^4 -norm over $\mathcal{O}(n; \mathbb{R})$.*

Proof See C.3. ■

Proposition 15 (Convergence of PGA with Arbitrary Step Size) *Iterative PGA algorithm 3 with any fixed step size $\alpha > 0$ (α can be $+\infty$ and PGA is equivalent to MSP Algorithm 1 when $\alpha = +\infty$) finds a saddle point of optimization problem (34)*

$$\max_{\mathbf{A} \in \mathcal{O}(n; \mathbb{R})} \|\mathbf{A}\|_4^4.$$

Proof See C.4. ■

In addition to Proposition 15, the intuition for *larger gradient converges faster* is due to the landscape of function $h(\mathbf{A})$. As we can see in (117), when α decreases, the landscape of $h(\mathbf{A})$ approaches becomes “flat”, hence the convergence rate decreases. On the contrary, when $\alpha \rightarrow +\infty$, $h(\mathbf{A})$ preserves the “sharp” curvature of $\|\mathbf{A}\|_4^4$, which yields the optimal convergence rate.

Although the function $g(\mathbf{A}) = \|\mathbf{A}\|_4^4$ may have many critical points, the signed permutation group $\mathbf{SP}(n)$ are the only global maximizers. As recent work has shown (Sun et al., 2015), such discrete symmetry helps regulate the global landscape of the objective function and makes it amenable to global optimization. Indeed, we have observed through our extensive experiments that, under broad conditions, the proposed MSP algorithm always converges to the globally optimal solution (set), at a super-linear convergence rate.

We only give a local result on the convergence of the MSP algorithm in this paper.⁸ That is, when the initial orthogonal matrix $\mathbf{A}$ is “close” enough to a signed permutation matrix, the MSP algorithm converges to that signed permutation at a very fast rate. It is easy to verify the algorithm is permutation invariant. Hence w.l.o.g., we may assume the target signed permutation is the identity $\mathbf{I}$.

Theorem 16 (Cubic Convergence Rate around Global Maximizers) *Given an orthogonal matrix $\mathbf{A} \in \mathbf{O}(n; \mathbb{R})$, let $\mathbf{A}'$ denote the output of the MSP Algorithm 1 after one iteration: $\mathbf{A}' = \mathbf{U}\mathbf{V}^*$, where $\mathbf{U}\mathbf{\Sigma}\mathbf{V}^* = \text{SVD}(\mathbf{A}^{\circ 3})$. If $\|\mathbf{A} - \mathbf{I}\|_F^2 = \varepsilon$, for $\varepsilon < 0.579$, then we have $\|\mathbf{A}' - \mathbf{I}\|_F^2 < \|\mathbf{A} - \mathbf{I}\|_F^2$ and $\|\mathbf{A}' - \mathbf{I}\|_F^2 < O(\varepsilon^3)$.*

Proof See C.5. ■

Theorem 16 shows that the MSP Algorithm 1 achieves cubic convergence rate locally, which is much faster than any gradient descent methods. Our experiments in Section 5 confirm this super-linear convergence rate for the MSP algorithms, at least in the deterministic case.

The above theorem only proves local convergence. As shown in section 5.2, We have observed in experiments that the algorithm actually converges *globally* under very broad conditions. We can see why this could be true in general from the special case when $n = 2$. Note that $\mathbf{O}(n, \mathbb{R})$ is a disjoint manifold with two continuous parts $\mathbf{SO}(n, \mathbb{R})$ and its reflection (Hall, 2015), for convenience, we only consider $\mathbf{SO}(n, \mathbb{R})$, which is a “half” of $\mathbf{O}(n, \mathbb{R})$, with determinant $\det(\mathbf{A}) = 1$. The next lemma shows the global convergence of our MSP algorithm in $\mathbf{SO}(2; \mathbb{R})$.

8. We leave the study of ensuring global optimality and convergence to future work.

Proposition 17 (Global Convergence of the MSP Algorithm on $\text{SO}(2; \mathbb{R})$) When $n = 2$, if we parameterize our $\mathbf{A}_t \in \text{SO}(2, \mathbb{R})$ as the following.⁹

$$\mathbf{A}_t = \begin{pmatrix} \cos \theta_t & -\sin \theta_t \\ \sin \theta_t & \cos \theta_t \end{pmatrix}, \quad \forall \theta_t \in \left[-\frac{\pi}{2}, \frac{\pi}{2}\right], \quad (40)$$

then θ_t and θ_{t+1} satisfies the following relation

$$\theta_{t+1} = \tan^{-1}(\tan^3 \theta_t). \quad (41)$$

Proof See C.6. ■

The previous proposition 17 indicates that the MSP algorithm is essentially conducting the following iteration:

$$\theta_{t+1} = \tan^{-1}(\tan^3 \theta_t) \quad \text{or} \quad \tan \theta_{t+1} = \tan^3 \theta_t. \quad (42)$$

See Figure 2 for a plot of this function. The iteration will converge superlinearly to the following values $\theta_\star = \pm\pi/2, 0, \pm\pi$, which all correspond to signed permutation matrices in $\text{SP}(2)$. The fixed points of this iteration are $\theta = k\pi/4$, where $k \in \{-4, -3, -2, -1, 0, 1, 2, 3, 4\}$. Among all these fixed points, the unstable ones are those with odd k , which correspond to Hadamard matrices. This phenomenon also justifies our experiments: Hadamard matrices are fixed point of the MSP iteration, but they are unstable and can be avoided with small random perturbation.

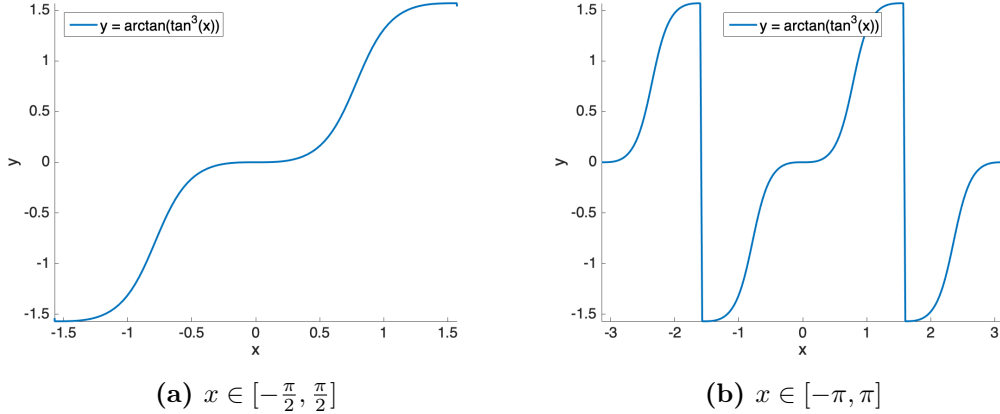


Figure 2: Function $y = \tan^{-1}(\tan^3 x)$, for different domain of x

9. The result of this lemma shows an update on the $\tan x$ function, which is a periodic function with period π , so we set $\theta \in [-\frac{\pi}{2}, \frac{\pi}{2}]$ to avoid the periodic ambiguity, one can easily generalize this result to the case where $\theta \in [-\pi, -\frac{\pi}{2}) \cup (\frac{\pi}{2}, \pi]$, since the signs of 1 won't affect our claim for pursuing a “signed permutation” matrix.

4.4. Generalized ℓ^{2k} -Norm Maximization

One may notice that we can also promote sparsity over the orthogonal group by maximizing the ℓ^{2k} -norm, $\forall k \geq 2$ on the orthogonal group:¹⁰

$$\max_{\mathbf{A}} \|\mathbf{A}\|_{2k}^{2k}, \quad \text{subject to } \mathbf{A} \in \mathbf{O}(n; \mathbb{R}). \quad (43)$$

In fact, the resulting algorithm would have a higher rate of convergence for the deterministic case, as the stretching with the power $(\cdot)^{\circ 2k-1}$ sparsifies the matrix more significantly with a larger k . See Section 5.3 for experimental verification. The following example provides one run of MSP algorithm when $2k = 10$.

Example 2 (One Run of MSP Algorithm with $2k = 10$) *This example shows how a random orthogonal matrix $\mathbf{A}_0 \in \mathbf{O}(n; \mathbb{R})$ evolves into a signed permutation matrix, using ℓ^{10} -norm:*

$$\begin{array}{llll} \mathbf{A}_0 = \begin{pmatrix} -0.6142 & 0.3943 & 0.6836 \\ -0.2039 & 0.7575 & -0.6201 \\ 0.7623 & 0.5203 & 0.3849 \end{pmatrix} & \xrightarrow{\text{stretching}} & \mathbf{A}_0^{\circ 9} = \begin{pmatrix} -0.0124 & 0.0002 & 0.0326 \\ -0.0000 & 0.0821 & -0.0136 \\ 0.0870 & 0.0028 & 0.0002 \end{pmatrix} \\ \xrightarrow{\text{projection}} \mathbf{A}_1 = \begin{pmatrix} -0.2657 & 0.0908 & 0.9598 \\ -0.0150 & 0.9950 & -0.0983 \\ 0.9639 & 0.0406 & 0.2630 \end{pmatrix} & \xrightarrow{\text{stretching}} & \mathbf{A}_1^{\circ 9} = \begin{pmatrix} -0.0000 & 0.0000 & 0.6910 \\ -0.0000 & 0.9562 & -0.0000 \\ 0.7185 & 0.0000 & 0.0000 \end{pmatrix} \\ \xrightarrow{\text{projection}} \mathbf{A}_2 = \begin{pmatrix} 0.0000 & -0.0000 & 1.0000 \\ 0.0000 & 1.0000 & -0.0000 \\ 1.0000 & -0.0000 & 0.0000 \end{pmatrix} & \xrightarrow{\text{output}} & \mathbf{A}_2 = \begin{pmatrix} 0 & 0 & 1 \\ 0 & 1 & 0 \\ 1 & 0 & 0 \end{pmatrix}. \end{array}$$

Comparing with Example 1, where ℓ^4 -norm MSP algorithm takes 6 iterations, the ℓ^{10} -norm MSP finds a signed permutation matrix in only 2 iterations.

In fact, one can show that if $2k \rightarrow \infty$, the corresponding MSP algorithm converges with *only one* iteration for the deterministic case!

So why don't we choose ℓ^∞ -norm instead? Notice that the above behavior is for the expectation of the objective function we actually care here. For dictionary learning, the object function $\hat{f}(\mathbf{A}, \mathbf{Y})$ depends on the random samples $\mathbf{Y}$. Both $\hat{f}(\mathbf{A}, \mathbf{Y})$ and $\nabla_{\mathbf{A}} \hat{f}(\mathbf{A}, \mathbf{Y})$ concentrate on their expectations only when p is large enough. For the same error threshold, if we choose larger k , then we will need much larger sample size p for $\hat{f}(\mathbf{A}, \mathbf{Y})$, $(\mathbf{A}\mathbf{Y})^{\circ 2k-1} \mathbf{Y}^*$ to concentrate on their expectation $f(\mathbf{A})$, $\mathbb{E}[(\mathbf{A}\mathbf{Y})^{\circ 2k-1} \mathbf{Y}^*]$ respectively. As we will see in Section 5.3 from experiments, the choice of $2k = 4$ seems to be the best in terms of balancing these two contending factors of convergence rate and sample size.

5. Experimental Verification

5.1. ℓ^4 -Norm Maximization over the Orthogonal Group

First, to have some basic ideas about how fast and smoothly the proposed MSP algorithm maximizes ℓ^4 -norm. Figure 3 shows one run of the MSP Algorithm 1 that maximizes $\|\mathbf{A}\|_4^4/n$. It reaches global maxima in less than 10 iterations for $n = 50$ and $n = 100$.

Next we run 100 trials of the algorithm to maximize $\|\mathbf{A}\|_4^4/n$ on $\mathbf{O}(n; \mathbb{R})$ with random initialization. Figure 4 shows that all 100 trials of the MSP algorithm reach the global maximum 1. According to Lemma 5, this indicates that all 100 trials (in different dimension $n = 50, 100$) converge to sign permutation matrices in $\mathbf{SP}(n)$.

10. Note that when $2k = 2$, $\|\mathbf{A}\|_2^2 = n$ is a constant, so ℓ^{2k} -norm only promotes sparsity when $2k \geq 4$.

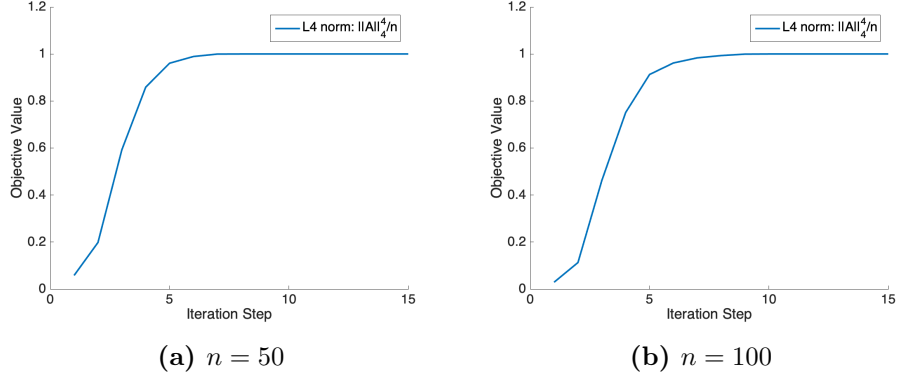


Figure 3: Normalized $\|\mathbf{A}\|_4^4$ in one run of MSP Algorithm 1 for maximizing $\|\mathbf{A}\|_4^4$ on $\mathcal{O}(n; \mathbb{R})$

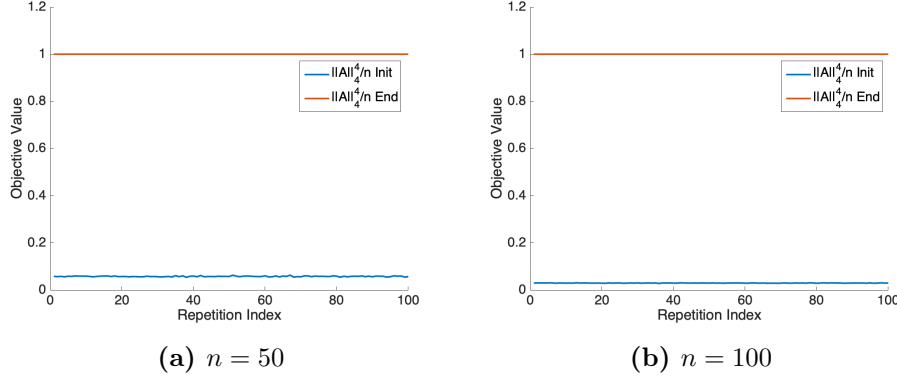


Figure 4: Normalized $\|\mathbf{A}\|_4^4/n$ for 100 trials of the MSP Algorithm 1 on $\mathcal{O}(n; \mathbb{R})$

5.2. Dictionary Learning via MSP

Figure 5(a) presents one trial of the proposed MSP Algorithm 2 for dictionary learning with $\theta = 0.3$, $n = 50$, and $p = 20,000$. The result corroborates with statements in Lemma 4 and Lemma 3: Maximizing $\hat{f}(\mathbf{A}, \mathbf{Y})$ is largely equivalent to optimizing $g(\mathbf{A}\mathbf{D}_o)$, and both values reach global maximum at the same time. Meanwhile, this result also shows our MSP algorithm is able to find the global maximum at ease, since $g(\mathbf{A}\mathbf{D}_o)$ reaches its maximal value 1 (with minor errors) by maximizing $\hat{f}(\mathbf{A}, \mathbf{Y})$. In Figure 5(b), we test the MSP Algorithm 2 in higher dimension $n = 100, p = 40,000, \theta = 0.3$. In both cases, our algorithm is surprisingly efficient: It only takes around 20 iterations to recover the ground truth dictionary.

In Figure 6, we run the MSP Algorithm 2 for 100 random trials with the settings $n = 50, p = 20,000, \theta = 0.3$ and $n = 100, p = 40,000, \theta = 0.3$. Among all 100 trials, $g(\mathbf{A}\mathbf{D}_o)$ achieve the global maximal value (within statistical errors) via optimizing $\hat{f}(\mathbf{A}, \mathbf{Y})$ in less than 30 iterations. This experiment seems to support a conjecture: Within conditions of this experiment, the MSP algorithm recovers the globally optimal dictionary.

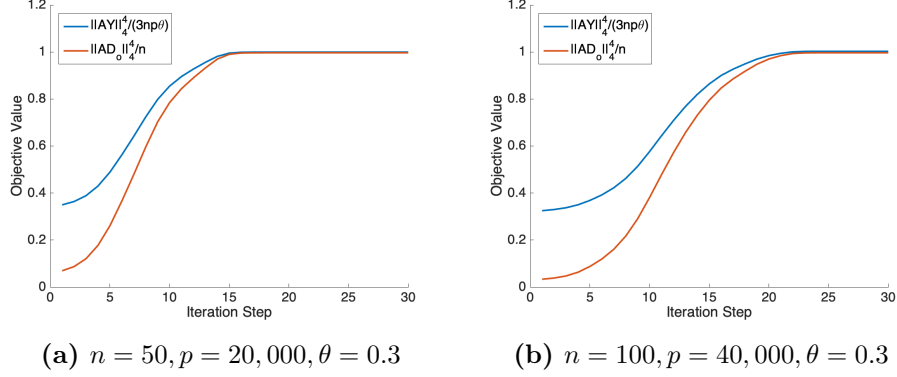


Figure 5: Normalized objective value $\|\mathbf{AD}_o\|_4^4/n$ and $\|\mathbf{AY}\|_4^4/3np\theta$ for individual trials of the MSP Algorithm 2, with different parameters n, p, θ . According to Lemma 5, $g(\mathbf{AD}_o)/n$ reaches 1 indicates successful recovery for $\mathbf{D}_o$. This experiment shows the MSP algorithm finds global maxima of $\hat{f}(\mathbf{A}, \mathbf{Y})$ thus recovers the correct dictionary $\mathbf{D}_o$.

5.3. ℓ^{2k} -Norm Maximization over the Orthogonal Group

In Figure 7, we conduct experiments to support the choice of ℓ^4 -norm. Figure 7(a) shows that for the deterministic case, the MSP Algorithm 1 finds signed permutation matrices faster with higher order ℓ^{2k} -norm. But Figure 7(b) indicates that as the order $2k$ increases, much more samples are needed by Algorithm 2 to achieve the same estimation error: p grows drastically as k increases. Hence, among all these sparsity-promoting norms (ℓ^{2k} -norm), the ℓ^4 -norm strikes a good balance between sample size and convergence rate.¹¹

5.4. Phase Transition Plots for Working Ranges of the MSP Algorithm

Encouraged by previous experiments, we conduct more extensive experiments of the MSP Algorithm 2 in broader settings to find its working range: 1) Figure 8(a) shows the result of varying the sparsity level θ (from 0 to 1) and sample size p (from 500 to 50,000) with a fixed dimension $n = 50$; 2) Figure 8(b) and (c) show results of changing dimension n (from 10 to 1,000) and sample size p (from 1,000 to 100,000) at a fixed sparsity level $\theta = 0.5$. Notice that all figures demonstrate a clear phase transition for the working range.

It is somewhat surprising to see in Figure 8(a) that the MSP algorithm is able to recover the dictionary correctly up to the sparsity level of $\theta \approx 0.6$ if p is large enough, which almost doubles the best existing theoretical guarantee given in Bai et al. (2018); Sun et al. (2015). We shall note that the inconsistency between the non linear phase transition curve in Figure 8(a) and the sample complexity $p = \Omega(\theta n^2 \ln n / \varepsilon^2)$ is due to the constant $C > \frac{4}{3\theta(1-\theta)}$ in Theorem 7.

Figure 8(b) and (c) show the working range for varying n, p with a fixed $\theta = 0.5$. Figure 8(b) is for a smaller range of n (from 10 to 100) and Figure 8(c) for a larger range of n (from

11. Later works (Shen et al., 2020; Xue et al., 2020) have shown that ℓ^3 -norm maximization also has the same sparsifying effect.

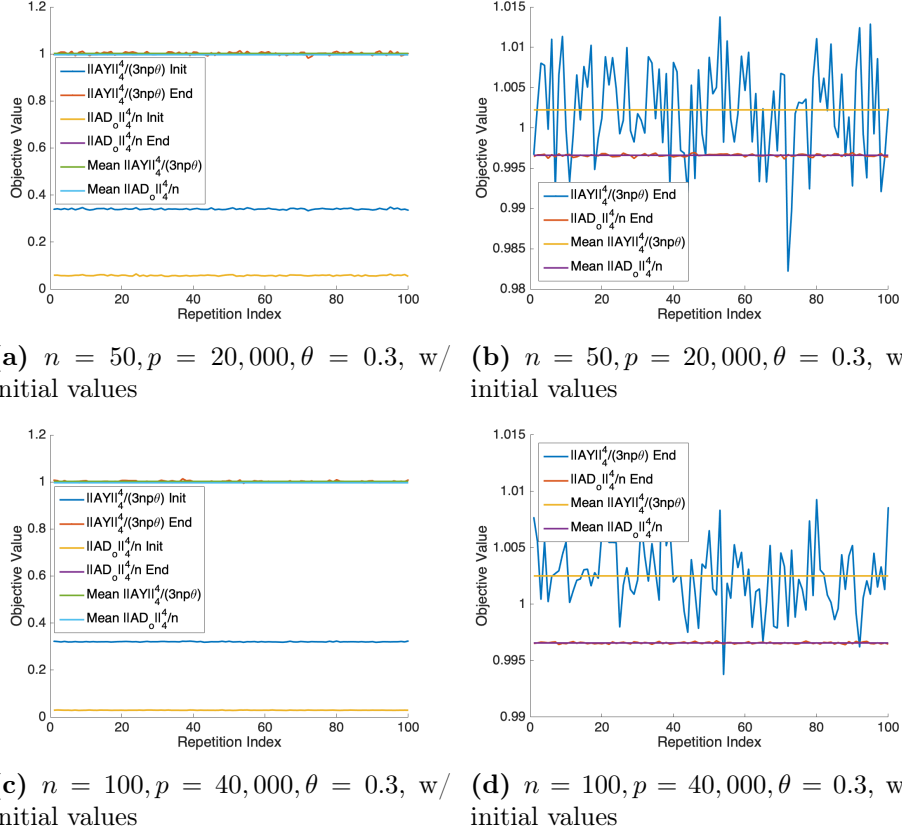
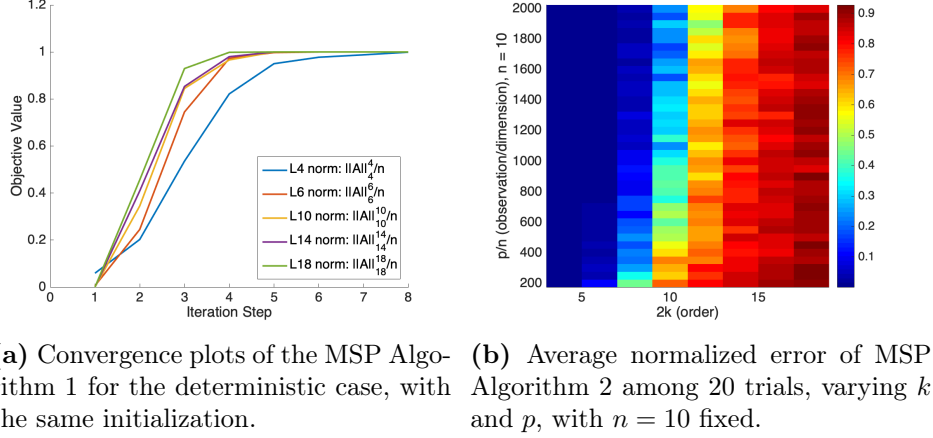
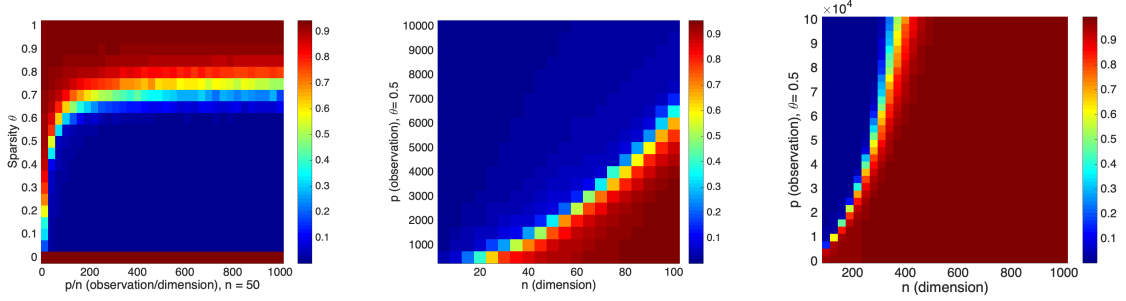


Figure 6: Normalized initial and final objective values of $\|AD_o\|_4^4/n$ and $\|AY\|_4^4/3np\theta$ for 100 trials of the MSP Algorithm 2, with $n = 50$ and $n = 100$. Both objectives converge to 1 (with minor errors) for all 100 trials.

100 to 1,000). Figures (b) and (c) imply that the required sample size p for the algorithm to succeed seems to be quadratic in the dimension n : $p = \Omega(n^2)$, which corresponds to our statistical results ($p = \Omega(\theta n^2 \ln n)$) in Theorem 7 and Proposition 11. This empirical bound is significantly better than the best theoretical bounds given in Bai et al. (2018); Sun et al. (2015) for the sample size required to ensure success. Similar observations have been reported in Schramm and Steurer (2017); Bai et al. (2018).

5.5. Comparison with Prior Arts

Table 2 compares the MSP method with the KSVD (Aharon et al., 2006), the SPAMS dictionary learning package Jenatton et al. (2010), and the latest subgradient method (Bai et al., 2018) for different choices of n, p under the same sparsity level $\theta = 0.3$. As one may see, our algorithm is significantly faster than the other algorithms in all trials. Further more, our algorithm has the potential for large scale experiments: It only takes 374.2 seconds to learn a 400×400 dictionaries from 160,000 samples. While the previous algorithms either fails to


 Figure 7: Use different ℓ^{2k} -norms for Algorithm 1 and Algorithm 2.

 Figure 8: Phase transition of average normalized error $|1 - \|\mathbf{AD}_o\|_4^4/n|$ of 10 random trials for the MSP Algorithm 2 in different settings: (a): Varying θ , p with fixed n ; (b), (c): Varying n , p with fixed θ . Red area indicates large error and blue area small error.

find the correct dictionary or is barely applicable. Within statistical errors, our algorithm gives slightly smaller values for $\|\mathbf{AD}_o\|_4^4/n$ in some trials. But the subgradient method (Bai et al., 2018) uses information of the ground truth dictionary $\mathbf{D}_o$ in their stopping criteria. Our MSP algorithm removes this dependency with only mild loss in accuracy. We shall note that traditional ℓ^1 -minimization based dictionary learning is more accurate than the ℓ^4 -maximization based method, since ℓ^1 -minimization will decrease small entries to 0 while the ℓ^4 -maximization tolerate small noise. Meanwhile, the numerical implementation also affects the experimental performance – the error of SPAMS is consistently better than the Subgradient method in terms of both accuracy and speed.

Although the convex relaxation based on sum-of-square SDP hierarchy (Barak et al., 2015; Ma et al., 2016; Schramm and Steurer, 2017) has theoretical guarantee for correctness under some specific statistical model, we do not include experimental comparison with the

sum-of-square method only because the SoS method incurs high computational complexity in solving large scale SDP programmings or tensor decomposition problems. In fact, on the same computing devices, sum-of-square methods cannot even solve the smallest dictionary learning problem ($n = 25$) in Table 2 due to memory and computation complexity.

			KSVD		SPAMS		Subgradient		MSP (Ours)	
n	$p (\times 10^4)$	θ	Error	Time	Error	Time	Error	Time	Error	Time
25	1	0.3	12.16%	64.5s	0.02%	5.53s	0.25%	9.2s	0.35%	0.2s (15 iter)
50	2	0.3	7.79%	165.4s	0.02%	23.4s	0.27%	107.0s	0.34%	1.3s (20 iter)
100	4	0.3	8.34%	569.4s	0.02%	365.9s	0.27%	2807.9s	0.35%	7.2s (25 iter)
200	8	0.3	8.46%	1064.0s	0.02%	5420.8s	N/A	>12h	0.35%	56.0s (40 iter)
400	16	0.3	11.94%	4646.3s	N/A	> 12h	N/A	>12h	0.35%	494.4s (60 iter)

Table 2: Comparison experiments with KSVD (Aharon et al., 2006), SPAMS Jenatton et al. (2010), and Subgradient method (Bai et al., 2018) in different trials of dictionary learning: (a) $n = 25, p = 1 \times 10^4, \theta = 0.3$; (b) $n = 50, p = 2 \times 10^4, \theta = 0.3$; (c) $n = 100, p = 4 \times 10^4, \theta = 0.3$; (d) $n = 200, p = 4 \times 10^4, \theta = 0.3$; (e) $n = 400, p = 16 \times 10^4, \theta = 0.3$. Recovery error is measured as $|1 - \|\mathbf{A}\mathbf{D}_o\|_4^4/n|$, since Lemma 5 shows that a perfect recovery gives $\|\mathbf{A}\mathbf{D}_o\|_4^4/n = 1$. All experiments are averaged among 5 trials and conducted on a 2.7 GHz Intel Core i5 processor (CPU of a 13-inch Mac Pro 2015).

5.6. Learning Sparsifying Dictionaries of Real Images

Indeed, the theoretical results of the MSP algorithm relies heavily on the Bernoulli-Gaussian assumption of the ground truth sparse code $\mathbf{x}_i$, that is, all $\mathbf{x}_i$ has the same element-wise variance. In order to demonstrate that the MSP algorithm can be applied to broader application scenarios beyond the Bernoulli-Gaussian setting, we test our algorithm on the MNIST dataset of hand-written digits (LeCun et al., 1998), whose element-wise (pixel-wise) variance are barely equal. In our implementation, we vectorize each 28×28 image j in the MNIST dataset into a vector $\mathbf{y}_i \in \mathbb{R}^{784}$ and directly apply our MSP algorithm 2 to the whole data set $\mathbf{Y} = [\mathbf{y}_1, \mathbf{y}_2, \dots, \mathbf{y}_{50,000}]$.¹² Figure 9(a) shows some bases learned from these images which obviously capture shapes of the digits. Figure 9(b) shows some less significant base vectors in the space of $\mathbb{R}^{784}$ that are orthogonal to the above.

We also compare our results with bases learned from Principal Component Analysis (PCA). As shown in Figure 10(a), the top bases learned from PCA are blurred with shapes from different digits mixed together, unlike those bases learned from the MSP in Figure 9(a) that clearly capture features of multiple different digits.

12. The training set of MNIST contains 50,000 images of size 28×28 .

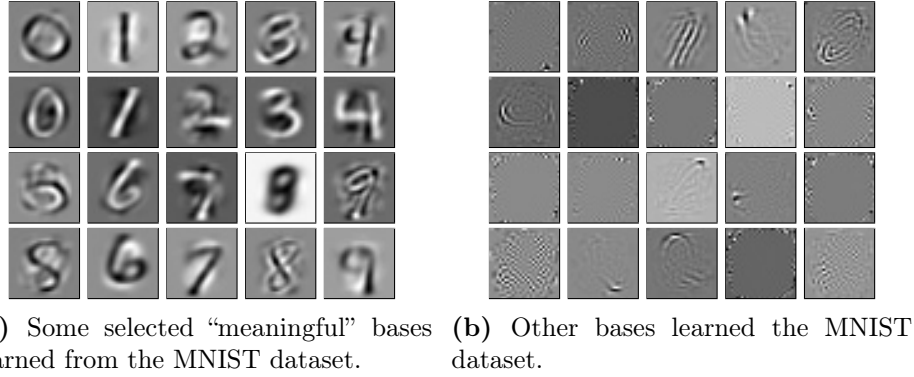


Figure 9: Learned dictionary through the MSP algorithm 2 on the MNIST dataset

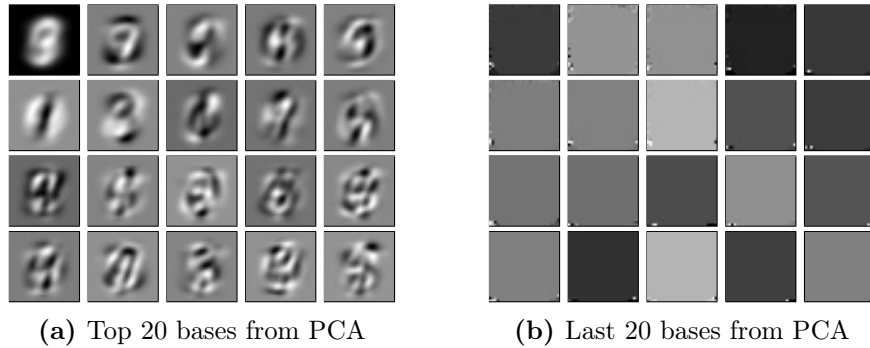


Figure 10: Learned Bases by PCA on MNIST dataset

We test the reconstruction results using the learned dictionary from the MSP algorithm and compare the reconstruction results using bases learned from PCA. In Figure 11, we compare the original MNIST data (LeCun et al., 1998), the reconstruction results from the MSP algorithm and PCA, using the top 1, 2, 3, 4, 5, 10, and 25 bases, respectively. As we can see in Figure 11, the reconstruction results for both the MSP and PCA improve, when the number of used bases increases. As shown in the case when only 2 bases are used, the MSP algorithm is able to represent clear shape for most digits already (see the three top left images), while PCA only provides mostly blurred results. When top 5 bases are used, the MSP algorithm almost capture all information of the input images, while PCA still gives rather blurred reconstruction. Both methods are able to capture salient information when the number of used bases are increased above 10. This is rather reasonable as there are about 10 different digits in this dataset and our method is precisely expected to have advantages when the number of bases is below 10.

This experiment has shown that the proposed MSP algorithm can be applied to more general setting beyond the Bernoulli-Gaussian setting. Moreover, due to its efficiency, the MSP algorithm can be extend to other large scale visual dataset such as CIFAR10 (Krizhevsky et al., 2014; Zhai et al., 2019).

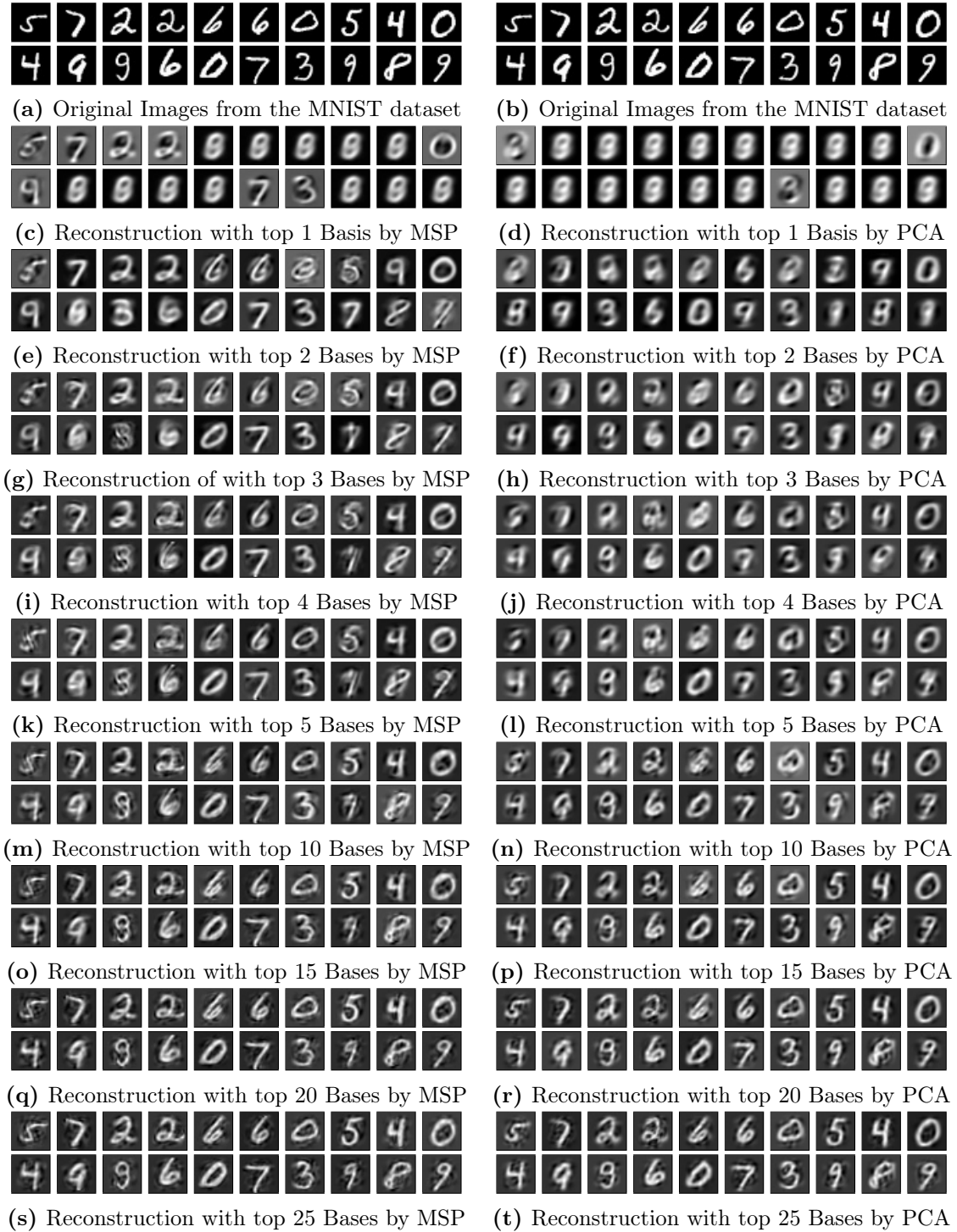


Figure 11: Comparison of compact representation with learned dictionary from the MSP algorithm 2 with the PCA bases for the MNIST image dataset.

6. Conclusions and Discussions

In this paper, we see that a complete dictionary can be effectively and efficiently learned with nearly minimum sample complexity by the proposed simple MSP algorithm. The computational complexity of the algorithm is essentially a few dozens of SVDs, allowing us to learn dictionary for very high-dimensional data.

Regarding sample complexity, the two main measure concentration results of Theorem 7 and Proposition 11 both require a sample complexity of $p = \Omega(\theta n^2 \log n / \varepsilon^2)$ for the global maximizers of the ℓ^4 -objective to be (close to) the correct $n \times n$ dictionary. This bound is consistent with our experiments in section 5.4. Barak et al. (2015); Ma et al. (2016); Schramm and Steurer (2017); Bai et al. (2018) have reported similar empirical evidences that at least $p = \Omega(n^2)$ samples are needed to recover an $n \times n$ dictionary.

Regarding the order of ℓ^{2k} -norm, we adopt ℓ^4 -norm because $2k = 4$ is the minimum even order that promotes sparsity. However, later works (Shen et al., 2020; Xue et al., 2020) show that ℓ^3 -norm maximization also promotes sparsity and the MSP algorithm (PGA with infinite step size) for ℓ^4 -norm maximization naturally generalizes to ℓ^3 -norm maximization.

Regarding optimality, Proposition 15 shows than random initialization of PGA algorithm with any step size finds a critical points of the ℓ^4 -norm over $O(n; \mathbb{R})$, Theorem 16 has shown the local convergence of the proposed MSP Algorithm 1 with a cubic rate for general n around global maximizers, and Proposition 17 has proven its global convergence for $n = 2$. It is natural to conjecture that the signed-permutations would be the only stable maximizers of the ℓ^4 -norm over the entire orthogonal group, hence the proposed algorithm converges globally, just like the experiments have indicated. Nevertheless, a rigorous proof is still elusive at this point.

Although some initial experiments have already suggested that the proposed algorithm works stably with mild noise and real data – showing clear advantages of the so learned dictionary over the classic PCA bases, the proposed algorithm actually generalizes well to data with dense noise, outliers, and sparse corruptions (Zhai et al., 2019). This paper has only addressed the case with a complete (square) dictionary. But there are ample reasons to believe that similar formulations and techniques presented in this paper can be extended to the over-complete case $\mathbf{D} \in \mathbb{R}^{n \times m}$ with $n < m$, at least when $m = O(n)$.

7. Acknowledgement

We would like to thank Yichao Zhou and Dr. Chong You of Berkeley EECS Department for stimulating discussions during preparation of this manuscript. We would like to thank professor Ju Sun of University of Minnesota CSE and Dr. Julien Mairal of Inria for references to existing work and experimental comparison. We would also like to thank Haozhi Qi of Berkeley EECS for help with the experiments. YZ would like to thank professor Yuejie Chi of CMU ECE for meaningful comments on the local convergence result. YZ would also like to thank Ye Xue from HKUST for insightful suggestions and proofreading. Yi likes to thank professor Bernd Sturmfels of Berkeley Math Department for help with analyzing algebraic properties of critical points of the ℓ^4 -norm over the orthogonal group and thank professor Alan Weinstein of Berkeley Math Department for discussions and references on the role of ℓ^4 -norm in spherical harmonic analysis. Yi would like to acknowledge that

this work is partially supported by research grant N00014-20-1-2002 from Office of Naval Research (ONR) and research grant from Tsinghua-Berkeley Shenzhen Institute (TBSI). JW gratefully acknowledges support from NSF grants 1733857, 1838061, 1740833, and 1740391, and thanks Sam Buchanan for discussions related to the geometry of ℓ^4 .

Appendices

A. Proofs of Section 2

A.1. Proof of Lemma 3

Claim 18 $\forall \theta \in (0, 1)$, let $\mathbf{X}_o \in \mathbb{R}^{n \times p}$, $x_{i,j} \sim_{iid} BG(\theta)$, $\mathbf{D}_o \in \mathcal{O}(n; \mathbb{R})$ is any orthogonal matrix, and $\mathbf{Y} = \mathbf{D}_o \mathbf{X}_o$. Then, $\forall \mathbf{A} \in \mathcal{O}(n; \mathbb{R})$, we have

$$\frac{1}{3p\theta} f(\mathbf{A}) = (1 - \theta)g(\mathbf{A}\mathbf{D}_o) + \theta n. \quad (44)$$

Proof For simplicity, since $\mathbf{D}_o \in \mathcal{O}(n; \mathbb{R})$, let $\mathbf{W} = \mathbf{A}\mathbf{D}_o$, we know that $\mathbf{W} \in \mathcal{O}(n; \mathbb{R})$. By the fact that $\mathbf{W} \in \mathcal{O}(n; \mathbb{R})$, we have:

$$\begin{aligned} n &= \sum_{i=1}^n \left(\sum_{k=1}^n w_{i,k}^2 \right)^2 = \underbrace{\sum_{i=1}^n \sum_{k=1}^n w_{i,k}^4}_{g(\mathbf{W})} + 2 \sum_{i=1}^n \sum_{1 \leq k_1 < k_2 \leq n} w_{i,k_1}^2 w_{i,k_2}^2 \\ \Rightarrow \sum_{i=1}^n \sum_{1 \leq k_1 < k_2 \leq n} w_{i,k_1}^2 w_{i,k_2}^2 &= \frac{n - g(\mathbf{W})}{2}. \end{aligned} \quad (45)$$

Next, we calculate $\frac{1}{3p\theta} f(\mathbf{A})$:

$$\begin{aligned} \frac{f(\mathbf{A})}{3p\theta} &= \frac{1}{3p\theta} \mathbb{E}_{\mathbf{X}_o} \hat{f}(\mathbf{A}; \mathbf{Y}) = \frac{1}{3p\theta} \mathbb{E}_{\mathbf{X}_o} \|\mathbf{A}\mathbf{D}_o \mathbf{X}_o\|_4^4 \\ &= \frac{1}{3p\theta} \mathbb{E}_{\mathbf{X}_o} \|\mathbf{W} \mathbf{X}_o\|_4^4 = \frac{1}{3p\theta} \sum_{i=1}^n \sum_{j=1}^p \mathbb{E}_{\mathbf{X}_0} \left(\sum_{k=1}^n w_{i,k} x_{k,j} \right)^4 \\ &= \frac{1}{3p\theta} \sum_{i=1}^n \sum_{j=1}^p \left[3\theta \sum_{k=1}^n w_{i,k}^4 + 6\theta^2 \sum_{1 \leq k_1 < k_2 \leq n} w_{i,k_1}^2 w_{i,k_2}^2 \right] \\ &= \underbrace{\sum_{i=1}^n \sum_{k=1}^n w_{i,k}^4}_{g(\mathbf{W})} + 2\theta \underbrace{\sum_{i=1}^n \sum_{1 \leq k_1 < k_2 \leq n} w_{i,k_1}^2 w_{i,k_2}^2}_{\frac{n - g(\mathbf{W})}{2}} \\ &= (1 - \theta)g(\mathbf{W}) + \theta n = (1 - \theta)g(\mathbf{A}\mathbf{D}_o) + \theta n \quad (\text{By (45)}), \end{aligned} \quad (46)$$

which completes the proof. ■

A.2. Proof of Lemma 4

Claim 19 (Concentration Bound of $\frac{1}{np}\hat{f}(\cdot, \cdot)$) $\forall \theta \in (0, 1)$, if $\mathbf{X} \in \mathbb{R}^{n \times p}$, $x_{i,j} \sim_{iid} BG(\theta)$, for any $\delta > 0$, the following inequality holds

$$\begin{aligned} & \mathbb{P} \left(\sup_{\mathbf{W} \in \mathcal{O}(n; \mathbb{R})} \frac{1}{np} \left| \|\mathbf{W}\mathbf{X}\|_4^4 - \mathbb{E} \|\mathbf{W}\mathbf{X}\|_4^4 \right| \geq \delta \right) \\ & < \exp \left(- \frac{3p\delta^2}{c_1\theta + 8n(\ln p)^4\delta} + n^2 \ln \left(\frac{60np(\ln p)^4}{\delta} \right) \right) \\ & + \exp \left(- \frac{p\delta^2}{c_2\theta} + n^2 \ln \left(\frac{60np(\ln p)^4}{\delta} \right) \right) + 2np\theta \exp \left(- \frac{(\ln p)^2}{2} \right), \end{aligned} \quad (47)$$

for some constants $c_1 > 10^4, c_2 > 3360$. Moreover

$$\begin{aligned} & \exp \left(- \frac{3p\delta^2}{c_1\theta + 8n(\ln p)^4\delta} + n^2 \ln \left(\frac{60np(\ln p)^4}{\delta} \right) \right) \\ & + \exp \left(- \frac{p\delta^2}{c_2\theta} + n^2 \ln \left(\frac{60np(\ln p)^4}{\delta} \right) \right) + 2np\theta \exp \left(- \frac{(\ln p)^2}{2} \right) \leq \frac{1}{p}, \end{aligned} \quad (48)$$

when $p = \Omega(\theta n^2 \ln n / \delta^2)$.

Proof Let $\bar{\mathbf{X}} \in \mathbb{R}^{n \times p}$ denote the truncated $\mathbf{X}$ by bound B

$$\bar{x}_{i,j} = \begin{cases} x_{i,j} & \text{if } |x_{i,j}| \leq B \\ 0 & \text{else} \end{cases}. \quad (49)$$

By Lemma 34, we know that $\|\mathbf{X}\|_\infty \leq B$ happens with probability at least $1 - 2np\theta \exp(-B^2/2)$ and $\bar{\mathbf{X}} = \mathbf{X}$ holds whenever $\|\mathbf{X}\|_\infty \leq B$. So we know that $\bar{\mathbf{X}} \neq \mathbf{X}$ holds with probability at most $2np\theta \exp(-B^2/2)$, and thus:

$$\begin{aligned} & \mathbb{P} \left(\sup_{\mathbf{W} \in \mathcal{O}(n; \mathbb{R})} \frac{1}{np} \left| \|\mathbf{W}\mathbf{X}\|_4^4 - \mathbb{E} \|\mathbf{W}\mathbf{X}\|_4^4 \right| \geq \delta \right) \\ & \leq \mathbb{P} \left(\sup_{\mathbf{W} \in \mathcal{O}(n; \mathbb{R})} \frac{1}{np} \left| \|\mathbf{W}\mathbf{X}\|_4^4 - \mathbb{E} \|\mathbf{W}\mathbf{X}\|_4^4 \right| \geq \delta, \mathbf{X} = \bar{\mathbf{X}} \right) + \mathbb{P}(\mathbf{X} \neq \bar{\mathbf{X}}) \\ & \leq \mathbb{P} \left(\sup_{\mathbf{W} \in \mathcal{O}(n; \mathbb{R})} \frac{1}{np} \left| \|\mathbf{W}\bar{\mathbf{X}}\|_4^4 - \mathbb{E} \|\mathbf{W}\mathbf{X}\|_4^4 \right| \geq \delta \right) + 2np\theta e^{-\frac{B^2}{2}}. \end{aligned} \quad (50)$$

ε -net Covering. For any positive ε satisfying:

$$\varepsilon \leq \frac{\delta}{10npB^4}, \quad (51)$$

by Lemma 35, there exists an ε -nets:

$$\mathcal{S}_\varepsilon = \{\mathbf{W}_1, \mathbf{W}_2, \dots, \mathbf{W}_{|\mathcal{S}_\varepsilon|}\}, \quad (52)$$

which covers $\mathcal{O}(n; \mathbb{R})$

$$\mathcal{O}(n; \mathbb{R}) \subset \bigcup_{l=1}^{|\mathcal{S}_\varepsilon|} \mathbb{B}(\mathbf{W}_l, \varepsilon), \quad (53)$$

in operator norm $\|\cdot\|_2$. Moreover, we have:

$$|\mathcal{S}_\varepsilon| \leq \left(\frac{6}{\varepsilon}\right)^{n^2}. \quad (54)$$

So $\forall \mathbf{W} \in \mathcal{O}(n; \mathbb{R})$, there exists $l \in [\mathcal{S}_\varepsilon]$, such that $\|\mathbf{W} - \mathbf{W}_l\|_2 \leq \varepsilon$. Thus, we have:

$$\begin{aligned} & \mathbb{P}\left(\frac{1}{np} \left| \|\mathbf{W} \bar{\mathbf{X}}\|_4^4 - \mathbb{E} \|\mathbf{W} \mathbf{X}\|_4^4 \right| \geq \delta\right) \\ & \leq \mathbb{P}\left(\frac{1}{np} \left(\left| \|\mathbf{W} \bar{\mathbf{X}}\|_4^4 - \|\mathbf{W}_l \bar{\mathbf{X}}\|_4^4 \right| + \left| \|\mathbf{W}_l \bar{\mathbf{X}}\|_4^4 - \mathbb{E} \|\mathbf{W}_l \mathbf{X}\|_4^4 \right| + \left| \mathbb{E} \|\mathbf{W}_l \mathbf{X}\|_4^4 - \mathbb{E} \|\mathbf{W} \mathbf{X}\|_4^4 \right| \right) \geq \delta\right) \\ & \leq \mathbb{P}\left(\frac{1}{np} \left| \|\mathbf{W}_l \bar{\mathbf{X}}\|_4^4 - \mathbb{E} \|\mathbf{W}_l \mathbf{X}\|_4^4 \right| + [4npB^4 + 12n\theta(1-\theta)] \|\mathbf{W} - \mathbf{W}_l\|_2 \geq \delta\right) \quad (\text{Lemma 38, 39}) \\ & < \mathbb{P}\left(\frac{1}{np} \left| \|\mathbf{W}_l \bar{\mathbf{X}}\|_4^4 - \mathbb{E} \|\mathbf{W}_l \mathbf{X}\|_4^4 \right| + 5npB^4\varepsilon \geq \delta\right) \quad (p, B \text{ are large numbers, } \|\mathbf{W} - \mathbf{W}_l\|_2 \leq \varepsilon) \\ & \leq \mathbb{P}\left(\frac{1}{np} \left| \|\mathbf{W}_l \bar{\mathbf{X}}\|_4^4 - \mathbb{E} \|\mathbf{W}_l \mathbf{X}\|_4^4 \right| + \frac{\delta}{2} \geq \delta\right) \quad (\text{We assume } \varepsilon \leq \frac{\delta}{10npB^4}) \\ & = \mathbb{P}\left(\frac{1}{np} \left| \|\mathbf{W}_l \bar{\mathbf{X}}\|_4^4 - \mathbb{E} \|\mathbf{W}_l \mathbf{X}\|_4^4 \right| \geq \frac{\delta}{2}\right). \end{aligned} \quad (55)$$

Analysis. For random variable $\bar{\mathbf{X}}$, we have:

$$\left| \mathbb{E} \|\mathbf{W}_l \mathbf{X}\|_4^4 - \mathbb{E} \|\mathbf{W}_l \bar{\mathbf{X}}\|_4^4 \right| = \left| \mathbb{E} [\|\mathbf{W}_l \mathbf{X}\|_4^4 \cdot \mathbb{1}_{\{\|\mathbf{X}\|_\infty > B\}}] \right| \leq \sqrt{\mathbb{E} \|\mathbf{W}_l \mathbf{X}\|_4^8} \sqrt{\mathbb{E} \mathbb{1}_{\{\|\mathbf{X}\|_\infty > B\}}}, \quad (56)$$

moreover, if we view $\|\mathbf{W}_l \mathbf{X}\|_4^4 = \sum_{j=1}^p \|\mathbf{W}_l \mathbf{x}_j\|_4^4$ as sum of p independent random variables, we have:

$$\begin{aligned} & \mathbb{E} \|\mathbf{W}_l \mathbf{X}\|_4^8 \\ & = \mathbb{E} \left(\left[\sum_{j=1}^p \|\mathbf{W}_l \mathbf{x}_j\|_4^4 \right]^2 \right) = \mathbb{E} \left(\sum_{j=1}^p \|\mathbf{W}_l \mathbf{x}_j\|_4^8 \right) + 2 \mathbb{E} \left(\sum_{1 \leq j_1 < j_2 \leq p} \|\mathbf{W}_l \mathbf{x}_{j_1}\|_4^4 \|\mathbf{W}_l \mathbf{x}_{j_2}\|_4^4 \right) \\ & = \mathbb{E} \left(\sum_{j=1}^p \|\mathbf{W}_l \mathbf{x}_j\|_4^8 \right) + 2 \sum_{1 \leq j_1 < j_2 \leq p} \mathbb{E} \|\mathbf{W}_l \mathbf{x}_{j_1}\|_4^4 \mathbb{E} \|\mathbf{W}_l \mathbf{x}_{j_2}\|_4^4 \quad (\mathbf{x}_j \text{ are independent}) \\ & = \sum_{j=1}^p \mathbb{E} \|\mathbf{W}_l \mathbf{x}_j\|_4^8 + p(p-1) \left[3\theta(1-\theta) \|\mathbf{W}_l\|_4^4 + 3\theta^2 n \right]^2 \quad (\text{By Lemma 3}) \\ & \leq Cpn^2\theta + 9p(p-1)n^2\theta^2 \quad (\text{By (194) of Lemma 42}) \\ & \leq C_1 p^2 n^2 \theta^2 \end{aligned} \quad (57)$$

for some constants $C_1 > 9$, for sufficiently large n, p . Also, by Lemma 34, we have:

$$\mathbb{E} \mathbb{1}_{\{\|\mathbf{X}\|_\infty > B\}} \leq 2np\theta e^{-\frac{B^2}{2}}. \quad (58)$$

Substitute the results of previous two inequalities (57), (58) into (56), yields

$$\left| \mathbb{E} \|\mathbf{W}_l \mathbf{X}\|_4^4 - \mathbb{E} \|\mathbf{W}_l \bar{\mathbf{X}}\|_4^4 \right| \leq \sqrt{C_1 p n \theta} \cdot \sqrt{2np\theta \exp(-B^2/2)} = C_2 n^{\frac{3}{2}} p^{\frac{3}{2}} \theta^{\frac{3}{2}} e^{-\frac{B^2}{4}}, \quad (59)$$

for some constants $C_2 > 3\sqrt{2}$. Hence, when

$$B > 2\sqrt{\ln\left(\frac{C_3 n^{\frac{1}{2}} p^{\frac{1}{2}} \theta^{\frac{3}{2}}}{\delta}\right)}, \quad (60)$$

for some constants $C_3 > 12\sqrt{2}$, we have:

$$\frac{1}{np} \left| \mathbb{E} \|\mathbf{W}_l \mathbf{X}\|_4^4 - \mathbb{E} \|\mathbf{W}_l \bar{\mathbf{X}}\|_4^4 \right| \leq C_2 n^{\frac{1}{2}} p^{\frac{1}{2}} \theta^{\frac{3}{2}} e^{-\frac{B^2}{4}} < \frac{\delta}{4}. \quad (61)$$

Therefore, combine (55), we have:

$$\begin{aligned} & \mathbb{P}\left(\frac{1}{np} \left| \|\mathbf{W} \bar{\mathbf{X}}\|_4^4 - \mathbb{E} \|\mathbf{W} \mathbf{X}\|_4^4 \right| \geq \delta\right) < \mathbb{P}\left(\frac{1}{np} \left| \|\mathbf{W}_l \bar{\mathbf{X}}\|_4^4 - \mathbb{E} \|\mathbf{W}_l \mathbf{X}\|_4^4 \right| \geq \frac{\delta}{2}\right) \\ & = \mathbb{P}\left(\frac{1}{np} \left| \|\mathbf{W}_l \bar{\mathbf{X}}\|_4^4 - \mathbb{E} \|\mathbf{W}_l \bar{\mathbf{X}}\|_4^4 + \mathbb{E} \|\mathbf{W}_l \bar{\mathbf{X}}\|_4^4 - \mathbb{E} \|\mathbf{W}_l \mathbf{X}\|_4^4 \right| \geq \frac{\delta}{2}\right) \\ & \leq \mathbb{P}\left(\frac{1}{np} \left| \|\mathbf{W}_l \bar{\mathbf{X}}\|_4^4 - \mathbb{E} \|\mathbf{W}_l \bar{\mathbf{X}}\|_4^4 \right| + \frac{1}{np} \left| \mathbb{E} \|\mathbf{W}_l \bar{\mathbf{X}}\|_4^4 - \mathbb{E} \|\mathbf{W}_l \mathbf{X}\|_4^4 \right| \geq \frac{\delta}{2}\right) \\ & \leq \mathbb{P}\left(\frac{1}{np} \left| \|\mathbf{W}_l \bar{\mathbf{X}}\|_4^4 - \mathbb{E} \|\mathbf{W}_l \bar{\mathbf{X}}\|_4^4 \right| \geq \frac{\delta}{4}\right) \quad \text{By (61)} \\ & = \underbrace{\mathbb{P}\left(\frac{1}{np} \left(\|\mathbf{W}_l \bar{\mathbf{X}}\|_4^4 - \mathbb{E} \|\mathbf{W}_l \bar{\mathbf{X}}\|_4^4 \right) \geq \frac{\delta}{4}\right)}_{\Gamma_1} + \underbrace{\mathbb{P}\left(\frac{1}{np} \left(\|\mathbf{W}_l \bar{\mathbf{X}}\|_4^4 - \mathbb{E} \|\mathbf{W}_l \bar{\mathbf{X}}\|_4^4 \right) \leq -\frac{\delta}{4}\right)}_{\Gamma_2}. \end{aligned} \quad (62)$$

Point-wise Bernstein Inequality. Next, we will apply Bernstein's inequality on $\|\mathbf{W}_l \bar{\mathbf{X}}\|_4^4$ to bound its upper (Γ_1) and lower tail (Γ_2). Note that we can view $\|\mathbf{W}_l \bar{\mathbf{X}}\|_4^4$ as sum of p independent variables $\|\mathbf{W}_l \bar{\mathbf{X}}\|_4^4 = \sum_{j=1}^p \|\mathbf{W}_l \bar{\mathbf{x}}_j\|_4^4$, and each of them is bounded by

$$\|\mathbf{W}_l \bar{\mathbf{x}}_j\|_4^4 = \|\mathbf{W}_l \bar{\mathbf{x}}_j\|_2^4 \cdot \left\| \frac{\mathbf{W}_l \bar{\mathbf{x}}_j}{\|\mathbf{W}_l \bar{\mathbf{x}}_j\|_2} \right\|_4^4 \leq \|\mathbf{W}_l \bar{\mathbf{x}}_j\|_2^4 = \|\bar{\mathbf{x}}_j\|_2^4 \leq n^2 B^4. \quad (63)$$

Also, in order to bound $\mathbb{E} \|\mathbf{W}_l \bar{\mathbf{x}}_j\|_4^8$, we consider

$$\mathbb{E} \|\mathbf{W}_l \mathbf{x}_j\|_4^8 - \mathbb{E} \|\mathbf{W}_l \bar{\mathbf{x}}_j\|_4^8 = \mathbb{E}(\|\mathbf{W}_l \mathbf{x}_j\|_4^8 - \|\mathbf{W}_l \bar{\mathbf{x}}_j\|_4^8) = \mathbb{E}(\|\mathbf{W}_l \mathbf{x}_j\|_4^8 \cdot \mathbb{1}_{\{\|\mathbf{X}\|_\infty > B\}}) \geq 0, \quad (64)$$

along with (194) in Lemma 42, this implies

$$\mathbb{E} \|\mathbf{W}_l \bar{\mathbf{x}}_j\|_4^8 \leq \mathbb{E} \|\mathbf{W}_l \mathbf{x}_j\|_4^8 \leq Cn^2\theta, \quad (65)$$

where $C > 105$ is the same constant as (57). Now we apply Bernstein's inequality on Γ_1 :

$$\begin{aligned} \Gamma_1 &= \mathbb{P} \left(\frac{1}{np} \left(\|\mathbf{W}_l \bar{\mathbf{X}}\|_4^4 - \mathbb{E} \|\mathbf{W}_l \bar{\mathbf{X}}\|_4^4 \right) \geq \frac{\delta}{4} \right) \\ &= \mathbb{P} \left(\sum_{j=1}^p \left(\|\mathbf{W}_l \bar{\mathbf{x}}_j\|_4^4 - \mathbb{E} \|\mathbf{W}_l \bar{\mathbf{x}}_j\|_4^4 \right) \geq p \cdot \frac{n\delta}{4} \right) \\ &\leq \exp \left(- \frac{pn^2\delta^2/16}{2 \left(\frac{1}{p} \sum_{j=1}^p \mathbb{E} \|\mathbf{W}_l \bar{\mathbf{x}}_j\|_4^8 + n^2 B^4 \cdot n\delta/12 \right)} \right) \\ &= \exp \left(- \frac{3pn^2\delta^2}{\frac{96}{p} \sum_{j=1}^p \mathbb{E} \|\mathbf{W}_l \bar{\mathbf{x}}_j\|_4^8 + 8n^3 B^4 \delta} \right) \\ &\leq \exp \left(- \frac{3pn^2\delta^2}{\frac{96}{p} \sum_{j=1}^p \mathbb{E} \|\mathbf{W}_l \mathbf{x}_j\|_4^8 + 8n^3 B^4 \delta} \right) \\ &\leq \exp \left(- \frac{3pn^2\delta^2}{c_1 n^2 \theta + 8n^3 B^4 \delta} \right) = \exp \left(- \frac{3p\delta^2}{c_1 \theta + 8n B^4 \delta} \right), \end{aligned} \quad (66)$$

for a constant $c_1 > 10^4$. Next, we apply Bernstein's inequality on Γ_2 , along with result in (57). Note that $\forall j \in [p]$, $\|\mathbf{W}_l \bar{\mathbf{x}}_j\|_4^4$ is lower bounded by 0, hence we have:

$$\begin{aligned} \Gamma_2 &= \mathbb{P} \left(\frac{1}{np} \left(\|\mathbf{W}_l \bar{\mathbf{X}}\|_4^4 - \mathbb{E} \|\mathbf{W}_l \bar{\mathbf{X}}\|_4^4 \right) \leq -\frac{\delta}{4} \right) \\ &= \mathbb{P} \left(\sum_{j=1}^p \left(\|\mathbf{W}_l \bar{\mathbf{x}}_j\|_4^4 - \mathbb{E} \|\mathbf{W}_l \bar{\mathbf{x}}_j\|_4^4 \right) \leq -p \cdot \frac{n\delta}{4} \right) \\ &\leq \exp \left(- \frac{pn^2\delta^2/16}{\frac{2}{p} \sum_{j=1}^p \mathbb{E} \|\mathbf{W}_l \bar{\mathbf{x}}_j\|_4^8} \right) \leq \exp \left(- \frac{pn^2\delta^2/16}{\frac{2}{p} \sum_{j=1}^p \mathbb{E} \|\mathbf{W}_l \mathbf{x}_j\|_4^8} \right) \\ &\leq \exp \left(- \frac{pn^2\delta^2}{32Cn^2\theta} \right) \leq \exp \left(- \frac{p\delta^2}{c_2\theta} \right), \end{aligned} \quad (67)$$

where $C > 105$ is the same constant as (57) and $c_2 > 3360$ is another constant. Replacing (66) and (67) into (62), yields

$$\mathbb{P} \left(\frac{1}{np} \left| \|\mathbf{W} \bar{\mathbf{X}}\|_4^4 - \mathbb{E} \|\mathbf{W} \bar{\mathbf{X}}\|_4^4 \right| \geq \delta \right) \leq \Gamma_1 + \Gamma_2 \leq \exp \left(- \frac{3p\delta^2}{c_1 \theta + 8n B^4 \delta} \right) + \exp \left(- \frac{p\delta^2}{c_2 \theta} \right) \quad (68)$$

for some constants $c_1 > 10^4, c_2 > 3360$.

Union Bound. Now, we will give a union bound for

$$\mathbb{P} \left(\sup_{\mathbf{W} \in \mathcal{O}(n; \mathbb{R})} \frac{1}{np} \left| \|\mathbf{W} \bar{\mathbf{X}}\|_4^4 - \mathbb{E} \|\mathbf{W} \bar{\mathbf{X}}\|_4^4 \right| \geq \delta \right). \quad (69)$$

Notice that

$$\begin{aligned}
 & \mathbb{P} \left(\sup_{\mathbf{W} \in \mathcal{O}(n; \mathbb{R})} \frac{1}{np} \left| \|\mathbf{W} \bar{\mathbf{X}}\|_4^4 - \mathbb{E} \|\mathbf{W} \mathbf{X}\|_4^4 \right| \geq \delta \right) \\
 & \leq \sum_{l=1}^{|\mathcal{S}_\varepsilon|} \mathbb{P} \left(\sup_{\mathbf{W} \in \mathbb{B}(\mathbf{W}_l, \varepsilon)} \frac{1}{np} \left| \|\mathbf{W} \bar{\mathbf{X}}\|_4^4 - \mathbb{E} \|\mathbf{W} \mathbf{X}\|_4^4 \right| \geq \delta \right) \quad (\text{By } \varepsilon\text{-covering in (53)}) \\
 & \leq \sum_{l=1}^{|\mathcal{S}_\varepsilon|} \left[\exp \left(-\frac{3p\delta^2}{c_1\theta + 8nB^4\delta} \right) + \exp \left(-\frac{p\delta^2}{c_2\theta} \right) \right] \quad (\text{By (68) when } \varepsilon \leq \frac{\delta}{10npB^4}) \\
 & \leq \left(\frac{6}{\varepsilon} \right)^{n^2} \left[\exp \left(-\frac{3p\delta^2}{c_1\theta + 8nB^4\delta} \right) + \exp \left(-\frac{p\delta^2}{c_2\theta} \right) \right] \quad (\text{By (54)}) \\
 & = \exp \left(n^2 \ln \left(\frac{60npB^4}{\delta} \right) \right) \left[\exp \left(-\frac{3p\delta^2}{c_1\theta + 8nB^4\delta} \right) + \exp \left(-\frac{p\delta^2}{c_2\theta} \right) \right] \quad (\text{Let } \varepsilon = \frac{\delta}{10npB^4}) \\
 & = \exp \left(-\frac{3p\delta^2}{c_1\theta + 8nB^4\delta} + n^2 \ln \left(\frac{60npB^4}{\delta} \right) \right) + \exp \left(-\frac{p\delta^2}{c_2\theta} + n^2 \ln \left(\frac{60npB^4}{\delta} \right) \right).
 \end{aligned} \tag{70}$$

Note that (60) requires a lower bound on B , here we can choose $B = \ln p$, which satisfies (60) when p is large enough (say $p = \Omega(n)$). Combine (50) and substitute $B = \ln p$ into (70), we have:

$$\begin{aligned}
 & \mathbb{P} \left(\sup_{\mathbf{W} \in \mathcal{O}(n; \mathbb{R})} \frac{1}{np} \left| \|\mathbf{W} \mathbf{X}\|_4^4 - \mathbb{E} \|\mathbf{W} \mathbf{X}\|_4^4 \right| \geq \delta \right) \\
 & \leq \mathbb{P} \left(\sup_{\mathbf{W} \in \mathcal{O}(n; \mathbb{R})} \frac{1}{np} \left| \|\mathbf{W} \bar{\mathbf{X}}\|_4^4 - \mathbb{E} \|\mathbf{W} \mathbf{X}\|_4^4 \right| \geq \delta \right) + 2np\theta e^{-\frac{B^2}{2}} \\
 & \leq \exp \left(-\frac{3p\delta^2}{c_1\theta + 8n(\ln p)^4\delta} + n^2 \ln \left(\frac{60np(\ln p)^4}{\delta} \right) \right) \\
 & \quad + \exp \left(-\frac{p\delta^2}{c_2\theta} + n^2 \ln \left(\frac{60np(\ln p)^4}{\delta} \right) \right) + 2np\theta \exp \left(-\frac{(\ln p)^2}{2} \right),
 \end{aligned} \tag{71}$$

for some constants $c_1 > 10^4, c_2 > 3360$, which completes the proof for (47). When $p = \Omega(\theta n^2 \ln n / \delta^2)$, we have:

$$\begin{aligned}
 & \exp \left(-\frac{3p\delta^2}{c_1\theta + 8n(\ln p)^4\delta} + n^2 \ln \left(\frac{60np(\ln p)^4}{\delta} \right) \right) \\
 & + \exp \left(-\frac{p\delta^2}{c_2\theta} + n^2 \ln \left(\frac{60np(\ln p)^4}{\delta} \right) \right) + 2np\theta \exp \left(-\frac{(\ln p)^2}{2} \right) \leq \frac{1}{3p} + \frac{1}{3p} + \frac{1}{3p} = \frac{1}{p},
 \end{aligned} \tag{72}$$

which completes the proof for (48). ■

A.3. Proof of Lemma 5

Claim 20 (Extrema of ℓ^4 -Norm over Orthogonal Group) *For any orthogonal matrix $\mathbf{A} \in \mathcal{O}(n; \mathbb{R})$, $g(\mathbf{A}) = \|\mathbf{A}\|_4^4 \in [1, n]$, $g(\mathbf{A})$ reaches maximum if and only if $\mathbf{A} \in SP(n)$.*

Proof For the maximum of $g(\mathbf{A})$

$$g(\mathbf{A}) = \sum_{i=1}^n \sum_{j=1}^n a_{i,j}^4 \leq \sum_{i=1}^n \sum_{j=1}^n a_{i,j}^4 + 2 \sum_{i=1}^n \sum_{1 \leq j_1 < j_2 \leq n} a_{i,j_1}^2 a_{i,j_2}^2 = \sum_{i=1}^n \left(\sum_{j=1}^n a_{i,j}^2 \right)^2 = n. \quad (73)$$

And when equality holds, we have

$$a_{i,j_1} a_{i,j_2} = 0, \quad \forall i, j_1 \neq j_2 \in [n], \quad (74)$$

which implies $\mathbf{A} \in SP(n)$. ■

A.4. Proof of Lemma 6

Claim 21 (Approximate Maxima of ℓ^4 -Norm over the Orthogonal Group) *Suppose $\mathbf{W}$ is an orthogonal matrix: $\mathbf{W} \in \mathcal{O}(n; \mathbb{R})$. $\forall \varepsilon \in [0, 1]$, if $\frac{1}{n} \|\mathbf{W}\|_4^4 \geq 1 - \varepsilon$, then $\exists \mathbf{P} \in SP(n)$, such that*

$$\frac{1}{n} \|\mathbf{W} - \mathbf{P}\|_F^2 \leq 2\varepsilon. \quad (75)$$

Proof Note that the condition $\frac{1}{n} \|\mathbf{W}\|_4^4 > 1 - \varepsilon$ can be viewed as

$$\frac{1}{n} \sum_{i=1}^n \|\mathbf{w}_i\|_4^4 \geq 1 - \varepsilon, \quad (76)$$

where each column vector $\mathbf{w}_i$ of $\mathbf{W}$ satisfies $\mathbf{w}_i \in \mathbb{S}^{n-1}$. By Lemma 37, we know there exists $j_1, j_2, \dots, j_n \in [n]$, such that

$$\frac{1}{n} \sum_{i=1}^n \|\mathbf{w}_i - s_i \mathbf{e}_{j_i}\|_2^2 \leq 2\varepsilon, \quad (77)$$

where $\mathbf{e}_{j_i}$ are one vector in canonical basis and $s_i \in \{1, -1\}$ indicates the sign of $\mathbf{e}_{j_i}$, $\forall i \in [n]$. Since $\mathbf{W} \in \mathcal{O}(n; \mathbb{R})$, one can easily show that $\forall i_1 \neq i_2, i_1, i_2 \in [n]$, the canonical vector they are corresponding to $\mathbf{e}_{j_{i_1}}, \mathbf{e}_{j_{i_2}}$ are different, that is, $j_{i_1} \neq j_{i_2}$ (otherwise suppose $\exists i_1 \neq i_2$, such that $\mathbf{w}_{i_1}$ and $\mathbf{w}_{i_2}$ correspond to the same canonical vector $\mathbf{e}_j$ in (77), one can easily show that $\mathbf{w}_{i_1}^* \mathbf{w}_{i_2} \neq 0$, which contradicts with the orthogonality of $\mathbf{W}$). Hence, there exists a sign permutation matrix:

$$\mathbf{P} = [\text{sign}(w_{j_1,1}) \mathbf{e}_{j_1}, \text{sign}(w_{j_2,2}) \mathbf{e}_{j_2}, \dots, \text{sign}(w_{j_n,n}) \mathbf{e}_{j_n}], \quad (78)$$

such that

$$\|\mathbf{W} - \mathbf{P}\|_F^2 = \sum_{i=1}^n \|\mathbf{w}_i - \text{sign}(w_{j_i,1}) \mathbf{e}_{j_i}\|_2^2 \leq 2n\varepsilon \implies \frac{1}{n} \|\mathbf{W} - \mathbf{P}\|_F^2 \leq 2\varepsilon, \quad (79)$$

which completes the proof. ■

A.5. Proof of Theorem 7

Claim 22 (Correctness of Global Maxima) $\forall \theta \in (0, 1)$, assume $\mathbf{X}_o = \{x_{i,j}\} \in \mathbb{R}^{n \times p}$ is a Bernoulli-Gaussian matrix, $\mathbf{D}_o \in \mathcal{O}(n; \mathbb{R})$ is any orthogonal matrix, and $\mathbf{Y} = \mathbf{D}_o \mathbf{X}_o$. Suppose $\hat{\mathbf{A}}_\star$ is a global maximizer of the optimization problem

$$\max_{\mathbf{A}} \hat{f}(\mathbf{A}, \mathbf{Y}) = \|\mathbf{A}\mathbf{Y}\|_4^4, \quad \text{subject to } \mathbf{A} \in \mathcal{O}(n; \mathbb{R}),$$

then for any $\varepsilon \in [0, 1]$, there exists a signed permutation matrix $\mathbf{P} \in SP(n)$, such that

$$\frac{1}{n} \left\| \hat{\mathbf{A}}_\star^* - \mathbf{D}_o \mathbf{P} \right\|_F^2 \leq C\varepsilon, \quad (80)$$

with probability at least

$$\begin{aligned} & 1 - \exp \left(-\frac{3p\varepsilon^2}{c_1\theta + 8n(\ln p)^4\varepsilon} + n^2 \ln \left(\frac{60np(\ln p)^4}{\varepsilon} \right) \right) \\ & - \exp \left(-\frac{p\varepsilon^2}{c_2\theta} + n^2 \ln \left(\frac{60np(\ln p)^4}{\varepsilon} \right) \right) - 2np\theta \exp \left(-\frac{(\ln p)^2}{2} \right), \end{aligned} \quad (81)$$

for some constants $c_1 > 10^4, c_2 > 3360, C > \frac{4}{3\theta(1-\theta)}$. Moreover

$$\begin{aligned} & 1 - \exp \left(-\frac{3p\varepsilon^2}{c_1\theta + 8n(\ln p)^4\varepsilon} + n^2 \ln \left(\frac{60np(\ln p)^4}{\varepsilon} \right) \right) \\ & - \exp \left(-\frac{p\varepsilon^2}{c_2\theta} + n^2 \ln \left(\frac{60np(\ln p)^4}{\varepsilon} \right) \right) - 2np\theta \exp \left(-\frac{(\ln p)^2}{2} \right) \geq 1 - \frac{1}{p}, \end{aligned} \quad (82)$$

when $p = \Omega(\theta n^2 \ln n / \varepsilon^2)$.

Proof Suppose $\mathbf{A}_\star$ is the global maximizer of optimization program (19):

$$\max_{\mathbf{A}} f(\mathbf{A}) = \mathbb{E} \|\mathbf{A}\mathbf{Y}\|_4^4 \quad \text{subject to } \mathbf{A} \in \mathcal{O}(n; \mathbb{R}),$$

then by (47) and (48) in Lemma 4, when $p = \Omega(\theta n^2 \ln n / \varepsilon^2)$, we know that with probability at least

$$\begin{aligned} & 1 - \exp \left(-\frac{3p\varepsilon^2}{c_1\theta + 8n(\ln p)^4\varepsilon} + n^2 \ln \left(\frac{60np(\ln p)^4}{\varepsilon} \right) \right) \\ & - \exp \left(-\frac{p\varepsilon^2}{c_2\theta} + n^2 \ln \left(\frac{60np(\ln p)^4}{\varepsilon} \right) \right) - 2np\theta \exp \left(-\frac{(\ln p)^2}{2} \right) \geq 1 - \frac{1}{p}, \end{aligned} \quad (83)$$

we have

$$\frac{1}{np} \left| \hat{f}(\hat{\mathbf{A}}_\star, \mathbf{Y}) - f(\hat{\mathbf{A}}_\star) \right| \leq \varepsilon \quad \text{and} \quad \frac{1}{np} \left| \hat{f}(\mathbf{A}_\star, \mathbf{Y}) - f(\mathbf{A}_\star) \right| \leq \varepsilon. \quad (84)$$

which implies

$$\frac{1}{np} f(\hat{\mathbf{A}}_\star) \leq \frac{1}{np} f(\mathbf{A}_\star) < \frac{1}{np} \hat{f}(\mathbf{A}_\star, \mathbf{Y}) + \varepsilon \leq \frac{1}{np} \hat{f}(\hat{\mathbf{A}}_\star, \mathbf{Y}) + \varepsilon < \frac{1}{np} f(\hat{\mathbf{A}}_\star) + 2\varepsilon. \quad (85)$$

In the above inequality, the first and the third $\leq$ is due to the global optimality of $f(\mathbf{A}_\star)$ and $\hat{f}(\hat{\mathbf{A}}_\star, \mathbf{Y})$, the second and the last $<$ is due to (84). Simplify (85), yields

$$\frac{1}{np}f(\hat{\mathbf{A}}_\star) \in \left(\frac{1}{np}f(\mathbf{A}_\star) - 2\varepsilon, \frac{1}{np}f(\mathbf{A}_\star) \right). \quad (86)$$

From Lemma 3, we have:

$$\frac{1}{3p\theta}f(\mathbf{A}) = (1 - \theta)g(\mathbf{A}\mathbf{D}_o) + \theta n, \quad \forall \mathbf{A} \in \mathcal{O}(n; \mathbb{R}),$$

which implies

$$\begin{aligned} \frac{3\theta(1 - \theta)}{n}g(\hat{\mathbf{A}}_\star\mathbf{D}_o) &\in \left(\frac{3\theta(1 - \theta)}{n}g(\mathbf{A}_\star\mathbf{D}_o) - 2\varepsilon, \frac{3\theta(1 - \theta)}{n}g(\mathbf{A}_\star\mathbf{D}_o) \right] \\ \Rightarrow \frac{1}{n} \|\hat{\mathbf{A}}_\star\mathbf{D}_o\|_4^4 &\in \left(\frac{1}{n} \|\mathbf{A}_\star\mathbf{D}_o\|_4^4 - \frac{2\varepsilon}{3\theta(1 - \theta)}, \frac{1}{n} \|\mathbf{A}_\star\mathbf{D}_o\|_4^4 \right]. \end{aligned} \quad (87)$$

Lemma 3 tells us that $\mathbf{A}_\star\mathbf{D}_o \in \text{SP}(n)$, combining Lemma 5, we know that $\|\mathbf{A}_\star\mathbf{D}_o\|_4^4 = n$. Thus we can further simplify (87) as

$$\frac{1}{n} \|\hat{\mathbf{A}}_\star\mathbf{D}_o\|_4^4 \in \left(1 - \frac{2\varepsilon}{3\theta(1 - \theta)}, 1 \right].$$

Applying Lemma 6 (change ε in Lemma 6 into $2\varepsilon/3\theta(1 - \theta)$), we know that there exists $\mathbf{P} \in \text{SP}(n)$, such that

$$\frac{1}{n} \|\hat{\mathbf{A}}_\star\mathbf{D}_o - \mathbf{P}\|_F^2 \leq \frac{4\varepsilon}{3\theta(1 - \theta)}. \quad (88)$$

By the rotational invariant of Frobenius norm, we have

$$\frac{1}{n} \|\hat{\mathbf{A}}_\star^* - \mathbf{D}_o\mathbf{P}^*\|_F^2 \leq \frac{4\varepsilon}{3\theta(1 - \theta)}, \quad (89)$$

which completes the proof. ■

B. Proofs of Section 3

B.1. Proof of Lemma 9

Claim 23 (Projection onto Orthogonal Group) $\forall \mathbf{A} \in \mathbb{R}^{n \times n}$, the orthogonal matrix which has minimum Frobenius norm with $\mathbf{A}$ is the following

$$\mathcal{P}_{\mathcal{O}(n; \mathbb{R})}(\mathbf{A}) = \arg \min_{\mathbf{M} \in \mathcal{O}(n; \mathbb{R})} \|\mathbf{M} - \mathbf{A}\|_F^2 = \mathbf{U}\mathbf{V}^*, \quad (90)$$

where $\mathbf{U}\mathbf{\Sigma}\mathbf{V}^* = \text{SVD}(\mathbf{A})$.

Proof Notice that

$$\begin{aligned}\|M - A\|_F^2 &= \text{tr}((M - A)(M - A)^*) \\ &= \text{tr}(I - AM^* - MA^* + AA^*) = n - 2\text{tr}(MA^*) + \text{tr}(AA^*).\end{aligned}\quad (91)$$

Since $\text{tr}(AA^*)$ is a constant, we know that

$$\arg \min_{M \in O(n; \mathbb{R})} \|M - A\|_F^2 = \arg \max_{M \in O(n; \mathbb{R})} \text{tr}(AM^*). \quad (92)$$

Let $U\Sigma V^*$ be the SVD of A , then

$$\text{tr}(AM^*) = \text{tr}(U\Sigma V^*M) = \text{tr}(\Sigma V^*M^*U) \leq \sum_{i=1}^n \sigma_i(A) \sigma_i(V^*M^*U) = \sum_{i=1}^n \sigma_i(A), \quad (93)$$

where inequality is obtained through Von Neumann's trace inequality, and the equality holds if and only if V^*M^*U is diagonal matrix (in fact, identity matrix), which implies $V^*M^*U = I \implies M = UV^*$. $\blacksquare$

B.2. Proof of Proposition 10

Claim 24 (Expectation of $\nabla_A \hat{f}(A, Y)$) *Let $X \in \mathbb{R}^{n \times p}$, $x_{i,j} \sim_{iid} BG(\theta)$, $D_o \in O(n; \mathbb{R})$ is any orthogonal matrix, and $Y = D_o X_o$. The expectation of $\nabla_A \hat{f}(A, Y)$ satisfies this property:*

$$\mathbb{E}_{X_o} \nabla_A \hat{f}(A, Y) = 3p\theta(1 - \theta) \nabla_A g(AD_o) + 12p\theta^2 A. \quad (94)$$

Proof For simplicity, since $D_o \in O(n; \mathbb{R})$, let $W = AD_o$, we know that $W \in O(n; \mathbb{R})$. By the fact that $W \in O(n; \mathbb{R})$, we have:

$$\frac{1}{4} \nabla_A \hat{f}(A, Y) = (AY)^{\circ 3} Y^* = (WX_o)^{\circ 3} X_o^* D_o^*, \quad (95)$$

moreover,

$$\begin{aligned}\{(AY)^{\circ 3}\}_{i,j} &= \{(WX_o)^{\circ 3}\}_{i,j} = \left(\sum_{k=1}^n w_{i,k} x_{k,j} \right)^3 \\ &= \sum_{k=1}^n w_{i,k}^3 x_{k,j}^3 + 3 \left(\sum_{1 \leq k_1 < k_2 \leq n} w_{i,k_1}^2 x_{k_1,j}^2 w_{i,k_2} x_{k_2,j} + w_{i,k_1} x_{k_1,j} w_{i,k_2}^2 x_{k_2,j}^2 \right) \\ &\quad + 6 \left(\sum_{1 \leq k_1 < k_2 < k_3 \leq n} w_{i,k_1} x_{k_1,j} w_{i,k_2} x_{k_2,j} w_{i,k_3} x_{k_3,j} \right).\end{aligned}\quad (96)$$

And hence, we know:

$$\begin{aligned}& \{(WX_o)^{\circ 3} X_o^*\}_{i,j'} \\ &= \sum_{j=1}^p \left[x_{j',j} \sum_{k=1}^n w_{i,k}^3 x_{k,j}^3 + 3x_{j',j} \left(\sum_{1 \leq k_1 < k_2 \leq n} w_{i,k_1}^2 x_{k_1,j}^2 w_{i,k_2} x_{k_2,j} + w_{i,k_1} x_{k_1,j} w_{i,k_2}^2 x_{k_2,j}^2 \right) \right. \\ &\quad \left. + 6x_{j',j} \left(\sum_{1 \leq k_1 < k_2 < k_3 \leq n} w_{i,k_1} x_{k_1,j} w_{i,k_2} x_{k_2,j} w_{i,k_3} x_{k_3,j} \right) \right].\end{aligned}\quad (97)$$

Thus,

$$\begin{aligned}\mathbb{E}_{\mathbf{X}_o}\{(\mathbf{W}\mathbf{X}_o)^{\circ 3}\mathbf{X}_o^*\}_{i,j'} &= 3p\theta w_{i,j'}^3 + 3p\theta^2 \sum_{\substack{1 \leq j \leq n \\ j \neq j'}} w_{i,j}^2 w_{i,j'} \\ &= 3p\theta w_{i,j'}^3 + 3p\theta^2(1 - w_{i,j'}^2)w_{i,j'} = 3p\theta(1 - \theta)w_{i,j'}^3 + 3p\theta^2 w_{i,j'},\end{aligned}\tag{98}$$

which implies

$$\begin{aligned}\frac{1}{4p}\mathbb{E}_{\mathbf{X}_o}\nabla_{\mathbf{A}}\hat{f}(\mathbf{A}, \mathbf{Y}) &= \frac{1}{p}\mathbb{E}_{\mathbf{X}_o}(\mathbf{W}\mathbf{X}_o)^{\circ 3}\mathbf{X}_o^*\mathbf{D}_o^* = 3\theta(1 - \theta)\mathbf{W}^{\circ 3}\mathbf{D}_o^* + 3\theta^2\mathbf{W}\mathbf{D}_o^* \\ &= 3\theta(1 - \theta)(\mathbf{A}\mathbf{D}_o)^{\circ 3}\mathbf{D}_o^* + 3\theta^2\mathbf{A} = \frac{3}{4}\theta(1 - \theta)\nabla_{\mathbf{A}}g(\mathbf{A}\mathbf{D}_o) + 3\theta^2\mathbf{A}.\end{aligned}\tag{99}$$

■

B.3. Proof of Proposition 11

Claim 25 (Union Tail Concentration Bound of $\frac{1}{np}\nabla\hat{f}(\cdot, \cdot)$) *If $\mathbf{X}_o \in \mathbb{R}^{n \times p}$, $x_{i,j} \sim iid$ $BG(\theta)$, for any $\mathbf{A}, \mathbf{D}_o \in \mathcal{O}(n; \mathbb{R})$, and $\mathbf{Y} = \mathbf{D}_o\mathbf{X}_o$, the following inequality holds*

$$\begin{aligned}&\mathbb{P}\left(\sup_{\mathbf{A} \in \mathcal{O}(n; \mathbb{R})} \frac{1}{4np} \left\| \nabla_{\mathbf{A}}\hat{f}(\mathbf{A}, \mathbf{Y}) - \mathbb{E}[\nabla_{\mathbf{A}}\hat{f}(\mathbf{A}, \mathbf{Y})] \right\|_F \geq \delta\right) \\ &\leq 2n^2 \exp\left(-\frac{3p\delta^2}{c_1\theta + 8n^{\frac{3}{2}}(\ln p)^4\delta} + n^2 \ln\left(\frac{48np(\ln p)^4}{\delta}\right)\right) + 2np\theta \exp\left(-\frac{(\ln p)^2}{2}\right),\end{aligned}\tag{100}$$

for a constant $c_1 > 1.7 \times 10^4$. Moreover

$$2n^2 \exp\left(-\frac{3p\delta^2}{c_1\theta + 8n^{\frac{3}{2}}(\ln p)^4\delta} + n^2 \ln\left(\frac{48np(\ln p)^4}{\delta}\right)\right) + 2np\theta \exp\left(-\frac{(\ln p)^2}{2}\right) \leq \frac{1}{p}\tag{101}$$

when $p = \Omega(\theta n^2 \ln n / \delta^2)$.

Proof By Proposition 10, we have:

$$\mathbb{E}_{\mathbf{X}_o}\nabla_{\mathbf{A}}\hat{f}(\mathbf{A}, \mathbf{Y}) = 3p\theta(1 - \theta)\nabla_{\mathbf{A}}g(\mathbf{A}\mathbf{D}_o) + 12p\theta^2\mathbf{A},\tag{102}$$

so

$$\begin{aligned}
 & \mathbb{P} \left(\sup_{\mathbf{A} \in \mathcal{O}(n; \mathbb{R})} \frac{1}{4np} \left\| \nabla_{\mathbf{A}} \hat{f}(\mathbf{A}, \mathbf{Y}) - \mathbb{E}[\nabla_{\mathbf{A}} \hat{f}(\mathbf{A}, \mathbf{Y})] \right\|_F \geq \delta \right) \\
 &= \mathbb{P} \left(\sup_{\mathbf{W} \in \mathcal{O}(n; \mathbb{R})} \frac{1}{np} \left\| (\mathbf{W} \mathbf{X}_o)^{\circ 3} \mathbf{X}_o^* \mathbf{D}_o^* - \mathbb{E}[(\mathbf{W} \mathbf{X}_o)^{\circ 3} \mathbf{X}_o^* \mathbf{D}_o^*] \right\|_F \geq \delta \right) \quad (\text{Assume } \mathbf{W} = \mathbf{A} \mathbf{D}_o) \\
 &= \mathbb{P} \left(\sup_{\mathbf{W} \in \mathcal{O}(n; \mathbb{R})} \frac{1}{np} \left\| (\mathbf{W} \mathbf{X}_o)^{\circ 3} \mathbf{X}_o^* - \mathbb{E}[(\mathbf{W} \mathbf{X}_o)^{\circ 3} \mathbf{X}_o^*] \right\|_F \geq \delta \right) \\
 &\leq 2n^2 \exp \left(- \frac{3p\delta^2}{c_1\theta + 8n^{\frac{3}{2}}(\ln p)^4\delta} + n^2 \ln \left(\frac{48np(\ln p)^4}{\delta} \right) \right) \\
 &\quad + 2np\theta \exp \left(- \frac{(\ln p)^2}{2} \right) \quad (\text{By Lemma 44}),
 \end{aligned} \tag{103}$$

for a constant $c_1 > 1.7 \times 10^4$, which completes the proof for (100). When $p = \Omega(\theta n^2 \ln n / \delta^2)$, we have

$$\begin{aligned}
 & 2n^2 \exp \left(- \frac{3p\delta^2}{c_1\theta + 8n^{\frac{3}{2}}(\ln p)^4\delta} + n^2 \ln \left(\frac{48np(\ln p)^4}{\delta} \right) \right) + 2np\theta \exp \left(- \frac{(\ln p)^2}{2} \right) \\
 & \leq \frac{1}{2p} + \frac{1}{2p} = \frac{1}{p},
 \end{aligned} \tag{104}$$

which completes the proof for (101). $\blacksquare$

C. Proofs of Section 4

C.1. Proof of Proposition 12

Claim 26 *The critical points of $g(\mathbf{W})$ on manifold $\mathcal{O}(n; \mathbb{R})$ satisfies the following condition*

$$(\mathbf{W}^{\circ 3})^* \mathbf{W} = \mathbf{W}^* \mathbf{W}^{\circ 3}. \tag{105}$$

Proof Notice that $\nabla_{\mathbf{W}} g(\mathbf{W}) = 4\mathbf{W}^{\circ 3}$, and the critical points of $g(\mathbf{W})$ on $\mathcal{O}(n; \mathbb{R})$ satisfies

$$\text{grad } g(\mathbf{W}) = \mathcal{P}_{T_{\mathbf{W}} \mathcal{O}(n; \mathbb{R})}(\nabla_{\mathbf{A}} g(\mathbf{A})) = \frac{1}{2}(4\mathbf{W}^{\circ 3} - \mathbf{W}(4\mathbf{W}^{\circ 3})^* \mathbf{W}) = \mathbf{0}, \tag{106}$$

which yields

$$(\mathbf{W}^{\circ 3})^* \mathbf{W} = \mathbf{W}^* \mathbf{W}^{\circ 3}. \tag{107}$$

Therefore, $\forall \mathbf{W} \in \mathcal{O}(n; \mathbb{R})$, we can write critical points condition of ℓ^4 -norm over $\mathcal{O}(n; \mathbb{R})$ as the following equations

$$\begin{cases} (\mathbf{W}^{\circ 3})^* \mathbf{W} = \mathbf{W}^* \mathbf{W}^{\circ 3}, \\ \mathbf{W}^* \mathbf{W} = \mathbf{I}. \end{cases} \tag{108}$$

$\blacksquare$

C.2. Proof of Proposition 13

Claim 27 *All global maximizers of ℓ^4 -norm over the orthogonal group are isolated critical points.*

Proof As our objective function is invariant under signed permutation group, without loss of generality, we prove for the identity matrix $\mathbf{I}$. Suppose that they are not isolated. Then, there exists a $\mathbf{W}_0$ such that $(\mathbf{W}_0^{\circ 3})^* \mathbf{W}_0 = \mathbf{W}_0^* \mathbf{W}_0^{\circ 3}$ and $\mathbf{W}_0^* \mathbf{W}_0 = \mathbf{I}$, and in every neighborhood of $\mathbf{W}_0$ there exists some $\mathbf{W}$ that is a critical point. This implies that there exists a path around $\mathbf{W}_0$ such that

$$\mathbf{W}(\cdot) : (-\varepsilon, \varepsilon) \rightarrow \mathbf{O}(n, \mathbb{R}), \quad (\mathbf{W}(t)^{\circ 3})^* \mathbf{W}(t) = \mathbf{W}(t)^* \mathbf{W}(t)^{\circ 3}, \quad \mathbf{W}(0) = \mathbf{W}_0. \quad (109)$$

Then, we expand $\mathbf{W}(t)$ around $t = 0$,

$$\mathbf{W}(t) = \mathbf{W}_0 + t\mathbf{W}_1 + t^2\mathbf{W}_2 + \dots \quad (110)$$

Constraint $(\mathbf{W}(t)^{\circ 3})^* \mathbf{W}(t) = \mathbf{W}(t)^* \mathbf{W}(t)^{\circ 3}$ implies that

$$\begin{aligned} (\mathbf{W}_0 + t\mathbf{W}_1)^* (\mathbf{W}_0 + t\mathbf{W}_1)^{\circ 3} &= [(\mathbf{W}_0 + t\mathbf{W}_1)^{\circ 3}]^* (\mathbf{W}_0 + t\mathbf{W}_1) \\ \implies 3\mathbf{W}_1 + \mathbf{W}_1^* &= 3\mathbf{W}_1^* + \mathbf{W}_1 \implies \mathbf{W}_1 = \mathbf{W}_1^*. \end{aligned} \quad (111)$$

Constraint $\mathbf{W}(t)^* \mathbf{W}(t) = \mathbf{I}$ implies that

$$\left. \frac{d}{dt} [\mathbf{W}(t)^* \mathbf{W}(t)] \right|_{t=0} = \left. \frac{d}{dt} \mathbf{I} \right|_{t=0} = \mathbf{0} \implies \mathbf{W}_1 + \mathbf{W}_1^* = \mathbf{0}. \quad (112)$$

The above computation is equivalent to plugging $\mathbf{W}(t)$ into the constraints, taking the derivative and evaluating at $\mathbf{0}$. Combining (111) and (112), we know $\mathbf{W}_1 = \mathbf{0}$. Since $\mathbf{W}(t)$ is any arbitrary path, this shows that the variety formed by all the critical points does not have a tangent space around $\mathbf{W}_0$. We conclude that the $\mathbf{I}$ is an isolated critical point. In fact, by calculating the Hessian at the maximizers, one can further show that all global maximizers, are nondegenerate critical points. $\blacksquare$

C.3. Proof of Proposition 14

Claim 28 (Fixed Point of the MSP Algorithm) *Given $\mathbf{W} \in \mathbf{O}(n; \mathbb{R})$, $\mathbf{W}$ is a fix point of the MSP algorithm if and only if $\mathbf{W}$ is a critical point of the ℓ^4 -norm over $\mathbf{O}(n; \mathbb{R})$.*

Proof Let $\mathbf{U}\Sigma\mathbf{V}^* = \mathbf{W}^{\circ 3}$ denote the SVD of $\mathbf{W}^{\circ 3}$.

- When $\mathbf{W} \in \mathbf{O}(n; \mathbb{R})$ is a fixed point of MSP algorithm, we have:

$$\mathbf{W} = \mathbf{U}\mathbf{V}^* \implies (\mathbf{W}^{\circ 3})^* \mathbf{W} = \mathbf{V}\Sigma\mathbf{U}^* \mathbf{W} = \mathbf{V}\Sigma\mathbf{V}^*, \quad (113)$$

which implies $(\mathbf{W}^{\circ 3})^* \mathbf{W}$ is symmetric, hence we have:

$$(\mathbf{W}^{\circ 3})^* \mathbf{W} = \mathbf{W}^* \mathbf{W}^{\circ 3}, \quad (114)$$

and by Proposition 12, $\mathbf{W}$ is a critical point of ℓ^4 -norm over the orthogonal group.

- When $\mathbf{W}$ is a critical point of ℓ^4 -norm over orthogonal group, we have $(\mathbf{W}^{\circ 3})^* \mathbf{W} = \mathbf{W}^* \mathbf{W}^{\circ 3}$. So the SVD of $(\mathbf{W}^{\circ 3})^* \mathbf{W}$ should equal to the SVD of $\mathbf{W}^* \mathbf{W}^{\circ 3}$. Also by the rotational invariant of SVD, we know:

$$\text{SVD}((\mathbf{W}^{\circ 3})^* \mathbf{W}) = \mathbf{V} \Sigma \mathbf{U}^* \mathbf{W} = \text{SVD}(\mathbf{W}^* \mathbf{W}^{\circ 3}) = \mathbf{W}^* \mathbf{U} \Sigma \mathbf{V}^*, \quad (115)$$

and by the uniqueness of SVD, we have:

$$\mathbf{V} = \mathbf{W}^* \mathbf{U} \implies \mathbf{W} = \mathbf{U} \mathbf{V}^*. \quad (116)$$

■

C.4. Proof of Proposition 15

Claim 29 (Convergence of PGA with Arbitrary Step Size) *Iterative PGA Algorithm 3 with any fixed step size $\alpha > 0$ (α can be $+\infty$ and PGA is equivalent to MSP Algorithm 1 when $\alpha = +\infty$) finds a saddle point of optimization problem (34)*

$$\max_{\mathbf{A} \in \text{O}(n; \mathbb{R})} \|\mathbf{A}\|_4^4.$$

Proof Consider the following objective function $h(\cdot) : \text{O}(n; \mathbb{R}) \mapsto \mathbb{R}^+$:

$$h(\mathbf{A}) = \begin{cases} \frac{\alpha}{4} \|\mathbf{A}\|_4^4 + \frac{1}{2} \|\mathbf{A}\|_F^2 & \text{when } \alpha < \infty \\ \|\mathbf{A}\|_4^4 & \text{when } \alpha = +\infty \end{cases}. \quad (117)$$

Note that $h(\mathbf{A})$ is convex in both cases ($\alpha < +\infty$ or $\alpha = \infty$). Also note that the Stiefel manifold $\text{O}(n; \mathbb{R})$ is a compact manifold, so by Theorem 1 in Journée et al. (2010), we know that the following iterative update

$$\mathbf{A}_{k+1} \in \arg \max \{h(\mathbf{A}_k) + \langle \partial h(\mathbf{A}_k), \mathbf{W} - \mathbf{A}_k \rangle \mid \mathbf{W} \in \text{O}(n; \mathbb{R})\} \quad (118)$$

will find a saddle point of $h(\mathbf{A})$ with any initialization $\mathbf{A}_0 \in \text{O}(n; \mathbb{R})$, where $\langle \cdot, \cdot \rangle : \text{O}(n; \mathbb{R}) \times \text{O}(n; \mathbb{R}) \mapsto \mathbb{R}$ is defined as

$$\langle \mathbf{W}_1, \mathbf{W}_2 \rangle = \text{tr}(\mathbf{W}_1^* \mathbf{W}_2). \quad (119)$$

Hence, by substituting (119) into (118), yields

$$\mathbf{A}_{k+1} = \begin{cases} \mathcal{P}_{\text{O}(n; \mathbb{R})}(\alpha \mathbf{A}_k^{\circ 3} + \mathbf{A}_k) & \text{when } \alpha < \infty \\ \mathcal{P}_{\text{O}(n; \mathbb{R})}(\mathbf{A}_k^{\circ 3}) & \text{when } \alpha = +\infty \end{cases}, \quad (120)$$

where $\mathcal{P}_{\text{O}(n; \mathbb{R})}(\cdot) : \mathbb{R}^{n \times n} \mapsto \text{O}(n; \mathbb{R})$ is the projection onto $\text{O}(n; \mathbb{R})$ (Absil and Malick, 2012):

$$\mathcal{P}_{\text{O}(n; \mathbb{R})}(\mathbf{A}) = \mathbf{U} \mathbf{V}^*, \quad \text{subject to} \quad \mathbf{U} \Sigma \mathbf{V}^* = \text{SVD}(\mathbf{A}). \quad (121)$$

Moreover, notice that $\|\mathbf{A}\|_F^2 = n$ is a constant, so finding a critical point of $h(\mathbf{A})$ is equivalent to finding a critical point of $\|\mathbf{A}\|_4^4$ over $\text{O}(n; \mathbb{R})$. Hence, we show that projected gradient ascent with any step size $\alpha > 0$ (including $\alpha = +\infty$) finds a critical point of $\|\mathbf{A}\|_4^4$ over $\text{O}(n; \mathbb{R})$.¹³ ■

13. The proof can easily be generalized to any Stiefel manifolds.

C.5. Proof of Theorem 16

Claim 30 (Local Convergence of the MSP Algorithm) *Given an orthogonal matrix $\mathbf{A} \in \mathcal{O}(n; \mathbb{R})$, let $\mathbf{A}'$ denote the output of the MSP Algorithm 1 after one iteration: $\mathbf{A}' = \mathbf{U}\mathbf{V}^*$, where $\mathbf{U}\mathbf{\Sigma}\mathbf{V}^* = \text{SVD}(\mathbf{A}^{\circ 3})$. If $\|\mathbf{A} - \mathbf{I}\|_F^2 = \varepsilon$, for $\varepsilon < 0.579$, then we have $\|\mathbf{A}' - \mathbf{I}\|_F^2 < \|\mathbf{A} - \mathbf{I}\|_F^2$ and $\|\mathbf{A}' - \mathbf{I}\|_F^2 < O(\varepsilon^3)$.*

Proof Let $\mathbf{A} = \mathbf{D} + \mathbf{N}$, where $\mathbf{D}$ is the diagonal part of $\mathbf{A}$ and $\mathbf{N}$ is the off-diagonal part. Therefore, we have:

$$\|\mathbf{N}^{\circ 3}\|_F \leq \|\mathbf{N}\|_F^3 \leq \|\mathbf{A} - \mathbf{I}\|_F^3 = \varepsilon^{3/2}, \quad (122)$$

where the first inequality is achieved through Cauchy-Schwarz inequality and the second inequality holds because $\mathbf{N}$ is the off-diagonal parts of $\mathbf{A} - \mathbf{I}$. We can view $\mathbf{A}^{\circ 3}$ as a $\mathbf{D}^{\circ 3}$ plus a small perturbation $\mathbf{N}^{\circ 3}$ with norm as most $\varepsilon^{3/2}$. By Lemma 45, we have:

$$\|\mathbf{Q}_\mathbf{A} - \mathbf{Q}_\mathbf{D}\|_F \leq \frac{2\|\mathbf{A}^{\circ 3} - \mathbf{D}^{\circ 3}\|_F}{\sigma_n(\mathbf{D}^{\circ 3}) + \sigma_n(\mathbf{A}^{\circ 3})} = \frac{2\|\mathbf{N}^{\circ 3}\|_F}{\sigma_n(\mathbf{D}^{\circ 3}) + \sigma_n(\mathbf{A}^{\circ 3})}, \quad (123)$$

where $\mathbf{U}_\mathbf{A}\mathbf{\Sigma}_\mathbf{A}\mathbf{V}_\mathbf{A}^* = \text{SVD}(\mathbf{A}^{\circ 3})$, $\mathbf{Q}_\mathbf{A} = \mathbf{U}_\mathbf{A}\mathbf{V}_\mathbf{A}^*$ and $\mathbf{U}_\mathbf{D}\mathbf{\Sigma}_\mathbf{D}\mathbf{V}_\mathbf{D}^* = \text{SVD}(\mathbf{D}^{\circ 3})$, $\mathbf{Q}_\mathbf{D} = \mathbf{U}_\mathbf{D}\mathbf{V}_\mathbf{D}^*$. Notice that $\mathbf{D}^{\circ 3}$ is a diagonal matrix, so $\mathbf{Q}_\mathbf{D} = \mathbf{I}$, $\sigma_n(\mathbf{D}^{\circ 3}) = \min_i a_{i,i}^3$. Moreover

$$\varepsilon = \|\mathbf{A} - \mathbf{I}\|_F^2 = \sum_{i=1}^n (a_{i,i} - 1)^2 + \sum_{i \neq j} a_{i,j}^2 = \sum_{i,j} a_{i,j}^2 - 2 \sum_{i=1}^n a_{i,i} + n \iff \sum_{i=1}^n a_{i,i} = n - \frac{\varepsilon}{2}, \quad (124)$$

without loss of generality, we can assume $1 \geq a_{1,1} \geq a_{2,2} \geq \dots \geq a_{n,n} > 0$, so we have:

$$a_{n,n} = n - \frac{\varepsilon}{2} - (a_{1,1} + a_{2,2} + \dots + a_{n-1,n-1}) \geq n - \frac{\varepsilon}{2} - (n-1) = 1 - \frac{\varepsilon}{2}, \quad (125)$$

hence we know that

$$\sigma_n(\mathbf{D}^{\circ 3}) = \min_i a_{i,i}^3 \geq \left(1 - \frac{\varepsilon}{2}\right)^3. \quad (126)$$

Applying Lemma 46 by substituting

$$\mathbf{G} = \mathbf{D}^{\circ 3} = \mathbf{B}\mathbf{M}, \quad \delta\mathbf{G} = \mathbf{A}^{\circ 3} - \mathbf{D}^{\circ 3} = \mathbf{N}^{\circ 3} = \delta\mathbf{B}\mathbf{M}, \quad (127)$$

we know $\mathbf{B} = \mathbf{I}$, $\mathbf{M} = \mathbf{D}^{\circ 3}$, $\delta\mathbf{B} = \mathbf{N}^{\circ 3}(\mathbf{D}^{\circ 3})^{-1}$, and thus:

$$\frac{|\sigma_n(\mathbf{A}^{\circ 3}) - \sigma_n(\mathbf{D}^{\circ 3})|}{\sigma_n(\mathbf{D}^{\circ 3})} \leq \frac{\|\delta\mathbf{B}\|_2}{\sigma_n(\mathbf{B})}, \quad (128)$$

which implies

$$\begin{aligned} |\sigma_n(\mathbf{A}^{\circ 3}) - \sigma_n(\mathbf{D}^{\circ 3})| &\leq \|\mathbf{N}^{\circ 3}(\mathbf{D}^{\circ 3})^{-1}\|_2 \sigma_n(\mathbf{D}^{\circ 3}) \leq \|\mathbf{N}^{\circ 3}\|_2 \|(\mathbf{D}^{\circ 3})^{-1}\|_2 \sigma_n(\mathbf{D}^{\circ 3}) \\ &\leq \|\mathbf{N}^{\circ 3}\|_F \left(1 - \frac{\varepsilon}{2}\right)^{-3} \sigma_n(\mathbf{D}^{\circ 3}) \leq \varepsilon^{3/2} \left(1 - \frac{\varepsilon}{2}\right)^{-3} \sigma_n(\mathbf{D}^{\circ 3}). \end{aligned} \quad (129)$$

Therefore, using (126), we know that:

$$\sigma_n(\mathbf{A}^{\circ 3}) + \sigma_n(\mathbf{D}^{\circ 3}) \geq \left[2 - \varepsilon^{3/2} \left(1 - \frac{\varepsilon}{2}\right)^{-3}\right] \sigma_n(\mathbf{D}^{\circ 3}) \geq 2 \left(1 - \frac{\varepsilon}{2}\right)^3 - \varepsilon^{3/2}. \quad (130)$$

Substituting (122) and (130) into (123), yields

$$\|\mathbf{Q}_A - \mathbf{Q}_D\|_F \leq \frac{2\varepsilon^{3/2}}{2\left(1 - \frac{\varepsilon}{2}\right)^3 - \varepsilon^{3/2}}, \quad (131)$$

hence, we have:

$$\|\mathbf{A}' - \mathbf{I}\|_F^2 = \|\mathbf{Q}_A - \mathbf{Q}_D\|_F^2 \leq O(\varepsilon^3). \quad (132)$$

Moreover, such operation is a contraction whenever:

$$\begin{aligned} \frac{2\varepsilon^{3/2}}{2\left(1 - \frac{\varepsilon}{2}\right)^3 - \varepsilon^{3/2}} &< \|\mathbf{A} - \mathbf{D}\|_F = \varepsilon^{1/2} \\ \iff 2\varepsilon - 2\left(1 - \frac{\varepsilon}{2}\right)^3 - \varepsilon^{3/2} &< 0 \\ \iff \varepsilon &< 0.579, \end{aligned} \quad (133)$$

which completes the proof. ■

C.6. Proof of Proposition 17

Claim 31 (Global Convergence of the MSP Algorithm on $\text{SO}(2; \mathbb{R})$) When $n = 2$, if we denote our $\mathbf{A}_t \in \text{SO}(2, \mathbb{R})$ as following form:

$$\mathbf{A}_t = \begin{pmatrix} \cos \theta_t & -\sin \theta_t \\ \sin \theta_t & \cos \theta_t \end{pmatrix}, \quad \forall \theta_t \in \left[-\frac{\pi}{2}, \frac{\pi}{2}\right], \quad (134)$$

then: 1) $\mathbf{A}_{t+1} \in \text{SO}(n; \mathbb{R})$; 2) if we let

$$\mathbf{A}_{t+1} = \begin{pmatrix} \cos \theta_{t+1} & -\sin \theta_{t+1} \\ \sin \theta_{t+1} & \cos \theta_{t+1} \end{pmatrix}, \quad \forall \theta_{t+1} \in \left[-\frac{\pi}{2}, \frac{\pi}{2}\right], \quad (135)$$

θ_t and θ_{t+1} satisfies the following relation

$$\theta_{t+1} = \tan^{-1}(\tan^3 \theta_t). \quad (136)$$

Proof By the update of the MSP algorithm, we know that

$$\mathbf{A}_{t+1} = \arg \min_{\mathbf{M} \in \text{O}(n; \mathbb{R})} \|\mathbf{M} - \mathbf{A}_t^{\circ 3}\|_F^2. \quad (137)$$

$\forall \mathbf{M} \in \text{O}(n; \mathbb{R})$, we can denote $\mathbf{M}$ as

$$\mathbf{M} = \begin{cases} \begin{pmatrix} \cos \theta & -\sin \theta \\ \sin \theta & \cos \theta \end{pmatrix} & \text{if } \det(\mathbf{M}) = 1 \\ \begin{pmatrix} \cos \theta & \sin \theta \\ \sin \theta & -\cos \theta \end{pmatrix} & \text{if } \det(\mathbf{M}) = -1, \end{cases}, \theta \in (-\pi, \pi] \quad (138)$$

Then if $\det(\mathbf{M}) = 1$, we have:

$$\begin{aligned}\|\mathbf{M} - \mathbf{A}_t^{\circ 3}\|_F^2 &= 2(\cos^3 \theta_t - \cos \theta)^2 + 2(\sin^3 \theta_t - \sin \theta)^2 \\ &= 2\cos^6 \theta_t + 2\sin^6 \theta_t - 4\cos \theta \cos^3 \theta_t - 4\sin \theta \sin^3 \theta_t + 2.\end{aligned}\quad (139)$$

When $\det(\mathbf{M}) = -1$, we have:

$$\begin{aligned}\|\mathbf{M} - \mathbf{A}_t^{\circ 3}\|_F^2 &= (\cos^3 \theta_t - \cos \theta)^2 + (\cos^3 \theta_t + \cos \theta)^2 + (\sin^3 \theta_t - \sin \theta)^2 + (\sin^3 \theta_t + \sin \theta)^2 \\ &= 2\cos^6 \theta_t + 2\sin^6 \theta_t + 2.\end{aligned}\quad (140)$$

Notice that when $\det(\mathbf{A}) = -1$, $\|\mathbf{M} - \mathbf{A}_t^{\circ 3}\|_F^2$ is a constant, so as long as we pick θ in the same quadrant with θ_t , then

$$2\cos^6 \theta_t + 2\sin^6 \theta_t - 4\cos \theta \cos^3 \theta_t - 4\sin \theta \sin^3 \theta_t + 2 \leq 2\cos^6 \theta_t + 2\sin^6 \theta_t + 2, \quad (141)$$

and therefore, we know $\mathbf{A}_{t+1} \in \text{SO}(n; \mathbb{R})$. So, we know that θ should satisfy the first order condition of critical point condition of the following optimization problem:

$$\mathbf{A}_{t+1} = \arg \min_{\mathbf{M} \in \text{SO}(n; \mathbb{R})} \|\mathbf{M} - \mathbf{A}_t^{\circ 3}\|_F^2, \quad (142)$$

which implies

$$\nabla_{\theta} [2(\cos^3 \theta_t - \cos \theta)^2 + 2(\sin^3 \theta_t - \sin \theta)^2] = 0 \implies \sin \theta \cos^3 \theta_t - \cos \theta \sin^3 \theta_t = 0 \quad (143)$$

which implies

$$\begin{cases} \tan \theta = \tan^3 \theta_t, & \text{when } \theta_t \neq \pm\pi/2 \\ \cot \theta = \cot^3 \theta_t, & \text{when } \theta_t \neq 0, \pi. \end{cases} \quad (144)$$

Note that we distinguish the different cases $\theta_t = \pm\pi/2, 0, \pi$ to avoid the division by zero in (143). One can also ignore this by taking the inverse on tangent, which yields

$$\theta_{t+1} = \tan^{-1}(\tan^3 \theta_t), \quad (145)$$

as stated. ■

D. Related Lemmas and Inequalities

D.1. Some Basic Inequalities

Lemma 32 (One-sided Bernstein's Inequality) *Given n random variables $x_1, x_2, \dots, x_n$, if $\forall i \in [n], x_i \leq b$ almost surely, then*

$$\mathbb{P}\left(\sum_{i=1}^n (x_i - \mathbb{E}[x_i]) \geq nt\right) \leq \exp\left(-\frac{nt^2}{2(\frac{1}{n} \sum_{i=1}^n \mathbb{E}[x_i^2] + bt/3)}\right). \quad (146)$$

Proof See Proposition 2.14 in Wainwright (2019). ■

Lemma 33 (Some Useful Norm Matrix Norm Inequalities) *Given two matrix $\mathbf{A}, \mathbf{B}$:*

1. *if $\mathbf{A}, \mathbf{B} \in \mathbb{R}^{n \times m}$, then $\|\mathbf{A} \circ \mathbf{B}\|_F \leq \|\mathbf{A}\|_F \|\mathbf{B}\|_F$*
2. *if $\mathbf{A} \in \mathbb{R}^{n \times r}$, $\mathbf{B} \in \mathbb{R}^{r \times m}$, then $\|\mathbf{AB}\|_F \leq \|\mathbf{A}\|_2 \|\mathbf{B}\|_F$*
3. *if $\mathbf{A} \in \mathbb{R}^{n \times r}$, $\mathbf{B} \in \mathbb{R}^{r \times m}$, then $\|\mathbf{AB}\|_F \leq \|\mathbf{A}\|_F \|\mathbf{B}\|_F$*

Proof

1. $\|\mathbf{A} \circ \mathbf{B}\|_F^2 = \sum_{i,j} (a_{i,j} b_{i,j})^2 \leq \left(\sum_{i,j} a_{i,j}^2 \right) \left(\sum_{i,j} b_{i,j}^2 \right) = \|\mathbf{A}\|_F^2 \|\mathbf{B}\|_F^2.$
2. $\|\mathbf{AB}\|_F^2 = \sum_j \|\mathbf{A} \mathbf{b}_j\|_F^2 \leq \|\mathbf{A}\|_2^2 \sum_j \|\mathbf{b}_j\|_2^2 = \|\mathbf{A}\|_2^2 \|\mathbf{B}\|_F^2.$
3. Let $\mathbf{a}_i, i \in [n]$ be the i^{th} row vector of $\mathbf{A}$ and $\mathbf{b}_j, j \in [m]$ be the j^{th} column vector of $\mathbf{B}$. So

$$\|\mathbf{AB}\|_F^2 = \sum_{i=1}^n \sum_{j=1}^m |\mathbf{a}_i^* \mathbf{b}_j|^2 \leq \left(\sum_{i=1}^n \|\mathbf{a}_i\|_2^2 \right) \left(\sum_{j=1}^m \|\mathbf{b}_j\|_2^2 \right) = \|\mathbf{A}\|_F^2 \|\mathbf{B}\|_F^2. \quad (147)$$

■

D.2. Truncation of Bernoulli-Gaussian Matrix

Lemma 34 (Entry-wise Truncation of a Bernoulli Gaussian Matrix) *Let $\mathbf{X} \in \mathbb{R}^{n \times p}$, where $x_{i,j} \sim_{iid} BG(\theta)$ and let $\|\cdot\|_\infty$ denote the maximum element (in absolute value) of a matrix, then*

$$\mathbb{P}(\|\mathbf{X}\|_\infty \geq t) \leq 2np\theta \exp\left(-\frac{t^2}{2}\right). \quad (148)$$

Proof A Bernoulli Gaussian variable $x_{i,j}, \forall i \in [n], j \in [p]$ satisfies $x_{i,j} = b_{i,j} \cdot g_{i,j}$, where $b_{i,j} \sim_{iid} \text{Ber}(\theta)$, $g_{i,j} \sim_{iid} \mathcal{N}(0, 1)$ and therefore

$$\mathbb{P}(|x_{i,j}| \geq t) = \theta \cdot \mathbb{P}(|g_{i,j}| \geq t) \leq 2\theta \exp\left(-\frac{t^2}{2}\right). \quad (149)$$

By union bound, we have:

$$\mathbb{P}(\|\mathbf{X}\|_\infty \geq t) \leq \sum_{i=1}^n \sum_{j=1}^p \mathbb{P}(|x_{i,j}| \geq t) \leq 2np\theta \exp\left(-\frac{t^2}{2}\right). \quad (150)$$

■

D.3. ε -covering of Stiefel Manifolds

Lemma 35 (ε -Net Covering of Stiefel Manifolds) ¹⁴ *There is a covering ε -net $\mathcal{S}_\varepsilon$ for Stiefel manifold $\mathcal{M} = \{\mathbf{W} \in \mathbb{R}^{n \times r} | \mathbf{W}^* \mathbf{W} = \mathbf{I}\}, (n \geq r)$ in operator norm*

$$\forall \mathbf{W} \in \mathcal{M}, \exists \mathbf{W}' \in \mathcal{S}_\varepsilon \quad \text{subject to} \quad \|\mathbf{W} - \mathbf{W}'\|_2 \leq \varepsilon, \quad (151)$$

of size $|\mathcal{S}_\varepsilon| \leq \left(\frac{6}{\varepsilon}\right)^{nr}$.

Proof Let $\mathcal{S}'_{\varepsilon/2} = \{\mathbf{A}_1, \mathbf{A}_2, \dots, \mathbf{A}_{|\mathcal{S}'_{\varepsilon/2}|}\}$ be an $\varepsilon/2$ -nets for the unit operator norm ball of $n \times r$ matrix $\{\mathbf{A} \in \mathbb{R}^{n \times r} | \|\mathbf{A}\|_2 \leq 1\}$, ε -net covering theorem shows that such construction of $\mathcal{S}'_{\varepsilon/2}$ exists and $|\mathcal{S}'_{\varepsilon/2}| \leq \left(\frac{6}{\varepsilon}\right)^{nr}$. Next, let $\mathcal{S}'$ be the subset of $\mathcal{S}'_{\varepsilon/2}$, which consists elements of $\mathcal{S}'_{\varepsilon/2}$ whose distance are within $\varepsilon/2$ with $\mathcal{M}$:

$$\mathcal{S}' = \{\mathbf{A} \in \mathcal{S}'_{\varepsilon/2} | \exists \mathbf{W} \in \mathcal{M}, \text{subject to } \|\mathbf{W} - \mathbf{A}\|_2 \leq \varepsilon/2\}. \quad (152)$$

Then, $\forall \mathbf{A} \in \mathcal{S}'$, let $\hat{\mathbf{W}}(\mathbf{A})$ be the nearest element of $\mathbf{A}$ in $\mathcal{M}$:

$$\hat{\mathbf{W}}(\mathbf{A}) = \arg \min_{\mathbf{W} \in \mathcal{M}} \|\mathbf{W} - \mathbf{A}\|_2, \quad (153)$$

and let $\mathcal{S}_\varepsilon$ be the set of the nearest element of each $\mathbf{A}$ in $\mathcal{S}'$: $\mathcal{S}_\varepsilon = \{\hat{\mathbf{W}}(\mathbf{A}) | \mathbf{A} \in \mathcal{S}'\}$. Since $\mathcal{S}'_{\varepsilon/2}$ is an $\varepsilon/2$ -nets for $\mathcal{M}$. So $\forall \mathbf{W} \in \mathcal{M}$, there exists $\mathbf{A}_l \in \mathcal{S}'_{\varepsilon/2}$, such that

$$\|\mathbf{W} - \mathbf{A}_l\|_2 \leq \frac{\varepsilon}{2}, \quad (154)$$

so $\mathbf{A}_l \in \mathcal{S}'$, and therefore there exists $\hat{\mathbf{W}}(\mathbf{A}_l) \in \mathcal{S}_\varepsilon$, such that

$$\left\| \hat{\mathbf{W}}(\mathbf{A}_l) - \mathbf{A}_l \right\|_2 \leq \|\mathbf{W} - \mathbf{A}_l\|_2 \leq \frac{\varepsilon}{2}, \quad (155)$$

hence by triangle inequality

$$\left\| \mathbf{W} - \hat{\mathbf{W}}(\mathbf{A}_l) \right\|_2 \leq \left\| \hat{\mathbf{W}}(\mathbf{A}_l) - \mathbf{A}_l \right\|_2 + \|\mathbf{W} - \mathbf{A}_l\|_2 \leq \varepsilon. \quad (156)$$

Thus, $\mathcal{S}_\varepsilon$ is an ε -net for $\mathcal{M}$ and $|\mathcal{S}_\varepsilon| = |\mathcal{S}'| \leq |\mathcal{S}'_{\varepsilon/2}| \leq \left(\frac{6}{\varepsilon}\right)^{nr}$. ■

D.4. Convergence to Maxima of ℓ^4 -Norm over Unit Sphere

Lemma 36 (Single Vector Convergence over Unit Sphere) *Suppose $\mathbf{q}$ is a vector on the unit sphere: $\mathbf{q} \in \mathbb{S}^{n-1}$. $\forall \varepsilon \in [0, 1]$, if $\|\mathbf{q}\|_4^4 \geq 1 - \varepsilon$, then $\exists i \in [n]$, such that*

$$\|\mathbf{q} - \mathbf{e}_i\|_2^2 \leq 2\varepsilon \quad \text{when } q_i > 0, \quad \|\mathbf{q} + \mathbf{e}_i\|_2^2 \leq 2\varepsilon \quad \text{when } q_i < 0, \quad (157)$$

where $\{\mathbf{e}_1, \mathbf{e}_2, \dots, \mathbf{e}_n\}$ is the canonical basis of $\mathbb{R}^n$.

14. A similar result can be found in Lemma 4.5 of Recht et al. (2010).

Proof Let $\mathbf{q} = [q_1, q_2, \dots, q_n]^*$, and without loss of generality, we can assume

$$1 \geq q_1^2 \geq q_2^2 \geq \dots \geq q_n^2 \geq 0. \quad (158)$$

Also, from the assumption

$$q_1^4 + q_2^4 + \dots + q_n^4 \geq 1 - \varepsilon, \quad (159)$$

and along with (158) and $\mathbf{q} \in \mathbb{S}^{n-1}$, we have:

$$q_1^4 + q_2^2(q_2^2 + q_3^2 + \dots + q_n^2) = q_1^4 + q_2^2(1 - q_1^2) \geq 1 - \varepsilon, \quad (160)$$

which implies:

$$q_1^4 + q_1^2(1 - q_1^2) \geq q_1^4 + q_2^2(1 - q_1^2) \geq 1 - \varepsilon \implies q_1^2 \geq 1 - \varepsilon. \quad (161)$$

Hence, we know

$$q_2^2 + \dots + q_n^2 \leq \varepsilon. \quad (162)$$

Moreover, (161) also implies

$$\varepsilon \geq (1 - q_1)(1 + q_1) \implies \begin{cases} 1 - q_1 \leq \varepsilon/(1 + q_1) \leq \varepsilon & \text{when } q_1 > 0, \\ 1 + q_1 \leq \varepsilon/(1 - q_1) \leq \varepsilon & \text{when } q_1 < 0. \end{cases} \quad (163)$$

When $q_1 > 0$, combine (161) and (162), we have

$$\|\mathbf{q} - \mathbf{e}_1\|_2^2 = (q_1 - 1)^2 + q_2^2 + \dots + q_n^2 \leq \varepsilon^2 + \varepsilon \leq 2\varepsilon. \quad (164)$$

And similar result for $\|\mathbf{q} + \mathbf{e}_1\|_2^2 \leq 2\varepsilon$ can be obtained through the same reasoning. $\blacksquare$

D.5. Multiple Vectors Convergence to Maxima of ℓ^4 -Norm over Unit Sphere

Lemma 37 (Multiple Vectors Convergence over Unit Sphere) Suppose $\mathbf{q}_1, \mathbf{q}_2, \dots, \mathbf{q}_k$ are k vectors on the unit sphere: $\mathbf{q}_i \in \mathbb{S}^{n-1}, \forall i \in [k]. \forall \varepsilon \in [0, 1]$, if

$$\frac{1}{k} \sum_{i=1}^k \|\mathbf{q}_i\|_4^4 \geq 1 - \varepsilon, \quad (165)$$

then $\exists j_1, j_2, \dots, j_k \in [n]$, such that

$$\frac{1}{k} \sum_{i=1}^k \|\mathbf{q}_i - s_i \mathbf{e}_{j_i}\|_2^2 \leq 2\varepsilon, \quad (166)$$

where $\mathbf{e}_{j_i}$ are one vector in canonical basis and $s_i \in \{1, -1\}$ indicates the sign of $\mathbf{e}_{j_i}, \forall i \in [k]$.

Proof Note that we can reformulate the condition (165) as

$$\sum_{i=1}^k \|\mathbf{q}_i\|_4^4 \geq k - k\varepsilon. \quad (167)$$

Since $\forall i \in [k]$, $0 \leq \|\mathbf{q}_i\|_4^4 \leq \|\mathbf{q}_i\|_2^4 = 1$, we can assume

$$\|\mathbf{q}_i\|_4^4 = 1 - \alpha_i \varepsilon, \quad \forall i \in [k], \quad (168)$$

where α_i satisfies $\alpha_i \varepsilon \in [0, 1]$, for all $i \in [k]$ and

$$\sum_{i=1}^k \alpha_i \leq k. \quad (169)$$

By Lemma 36, we know there exists $s_i \in \{1, -1\}$ and j_i , such that

$$\|\mathbf{q}_i - s_i \mathbf{e}_{j_i}\|_2^2 \leq 2\alpha_i \varepsilon, \quad \forall i \in [k]. \quad (170)$$

Along with (169), we have:

$$\frac{1}{k} \sum_{i=1}^k \|\mathbf{q}_i - s_i \mathbf{e}_{j_i}\|_2^2 \leq \frac{2}{k} \varepsilon \sum_{i=1}^k \alpha_i \leq 2\varepsilon, \quad (171)$$

which completes the proof. ■

D.6. Related Lipschitz Constants

Lemma 38 (Lipschitz Constant of $\frac{1}{np} \hat{f}(\cdot, \cdot)$ over $\mathcal{O}(n; \mathbb{R})$) *If $\mathbf{X} \in \mathbb{R}^{n \times p}$, $x_{i,j} \sim_{iid} BG(\theta)$, and let $\bar{\mathbf{X}}$ be the truncation of $\mathbf{X}$ by bound B*

$$\bar{x}_{i,j} = \begin{cases} x_{i,j} & \text{if } |x_{i,j}| \leq B \\ 0 & \text{else} \end{cases}, \quad (172)$$

then $\forall \mathbf{W}_1, \mathbf{W}_2 \in \mathcal{O}(n; \mathbb{R})$, we have

$$\frac{1}{np} \left| \|\mathbf{W}_1 \bar{\mathbf{X}}\|_4^4 - \|\mathbf{W}_2 \bar{\mathbf{X}}\|_4^4 \right| \leq L_1 \|\mathbf{W}_1 - \mathbf{W}_2\|_2, \quad (173)$$

for a constant $L_1 \leq 4npB^4$.

Proof Notice that

$$\begin{aligned} \left| \|\mathbf{W}_1 \bar{\mathbf{X}}\|_4^4 - \|\mathbf{W}_2 \bar{\mathbf{X}}\|_4^4 \right| &= \left| \sum_{i,j} [(\mathbf{W}_1 \bar{\mathbf{X}})_{i,j}^4 - (\mathbf{W}_2 \bar{\mathbf{X}})_{i,j}^4] \right| \\ &= \left| \sum_{i,j} \left\{ [(\mathbf{W}_1 \bar{\mathbf{X}})_{i,j}^2 - (\mathbf{W}_2 \bar{\mathbf{X}})_{i,j}^2] [(\mathbf{W}_1 \bar{\mathbf{X}})_{i,j}^2 + (\mathbf{W}_2 \bar{\mathbf{X}})_{i,j}^2] \right\} \right| \\ &\leq \left\{ \sum_{i,j} [(\mathbf{W}_1 \bar{\mathbf{X}})_{i,j}^2 - (\mathbf{W}_2 \bar{\mathbf{X}})_{i,j}^2]^2 \right\}^{1/2} \left\{ \sum_{i,j} [(\mathbf{W}_1 \bar{\mathbf{X}})_{i,j}^2 + (\mathbf{W}_2 \bar{\mathbf{X}})_{i,j}^2]^2 \right\}^{1/2} \\ &= \underbrace{\left\| (\mathbf{W}_1 \bar{\mathbf{X}})^{\circ 2} - (\mathbf{W}_2 \bar{\mathbf{X}})^{\circ 2} \right\|_F}_{\Gamma_1} \underbrace{\left\| (\mathbf{W}_1 \bar{\mathbf{X}})^{\circ 2} + (\mathbf{W}_2 \bar{\mathbf{X}})^{\circ 2} \right\|_F}_{\Gamma_2}, \end{aligned} \quad (174)$$

the only inequality is obtained through Cauchy–Schwarz inequality. For Γ_1 , we have

$$\begin{aligned}
 \Gamma_1 &= \left\| (\mathbf{W}_1 \bar{\mathbf{X}})^{\circ 2} - (\mathbf{W}_2 \bar{\mathbf{X}})^{\circ 2} \right\|_F = \left\| (\mathbf{W}_1 \bar{\mathbf{X}} - \mathbf{W}_2 \bar{\mathbf{X}}) \circ (\mathbf{W}_1 \bar{\mathbf{X}} + \mathbf{W}_2 \bar{\mathbf{X}}) \right\|_F \\
 &\leq \left\| (\mathbf{W}_1 - \mathbf{W}_2) \bar{\mathbf{X}} \right\|_F \left\| (\mathbf{W}_1 + \mathbf{W}_2) \bar{\mathbf{X}} \right\|_F \quad (\text{By inequality 1 in Lemma 33}) \\
 &\leq \|\mathbf{W}_1 - \mathbf{W}_2\|_2 \|\mathbf{W}_1 + \mathbf{W}_2\|_2 \|\bar{\mathbf{X}}\|_F^2 \quad (\text{By inequality 2 in Lemma 33}) \\
 &\leq \|\mathbf{W}_1 - \mathbf{W}_2\|_2 (\|\mathbf{W}_1\|_2 + \|\mathbf{W}_2\|_2) \|\bar{\mathbf{X}}\|_F^2 \\
 &= 2 \|\mathbf{W}_1 - \mathbf{W}_2\|_2 \|\bar{\mathbf{X}}\|_F^2 \quad (\mathbf{W}_1, \mathbf{W}_2 \text{ are orthogonal, } \|\mathbf{W}_1\|_2 = \|\mathbf{W}_2\|_2 = 1) \\
 &\leq 2npB^2 \|\mathbf{W}_1 - \mathbf{W}_2\|_2 \quad (|\bar{x}_{i,j}| \leq B).
 \end{aligned} \tag{175}$$

For Γ_2 , we have

$$\begin{aligned}
 \Gamma_2 &= \left\| (\mathbf{W}_1 \bar{\mathbf{X}})^{\circ 2} + (\mathbf{W}_2 \bar{\mathbf{X}})^{\circ 2} \right\|_F \\
 &\leq \left\| (\mathbf{W}_1 \bar{\mathbf{X}})^{\circ 2} \right\|_F + \left\| (\mathbf{W}_2 \bar{\mathbf{X}})^{\circ 2} \right\|_F \\
 &\leq \|\mathbf{W}_1 \bar{\mathbf{X}}\|_F^2 + \|\mathbf{W}_2 \bar{\mathbf{X}}\|_F^2 \quad (\text{By inequality 1 in Lemma 33}) \\
 &\leq \|\mathbf{W}_1\|_2^2 \|\bar{\mathbf{X}}\|_F^2 + \|\mathbf{W}_2\|_2^2 \|\bar{\mathbf{X}}\|_F^2 \\
 &= 2 \|\bar{\mathbf{X}}\|_F^2 \quad (\mathbf{W}_1, \mathbf{W}_2 \text{ are orthogonal matrices, } \|\mathbf{W}_1\|_2 = \|\mathbf{W}_2\|_2 = 1) \\
 &\leq 2npB^2 \quad (|\bar{x}_{i,j}| \leq B).
 \end{aligned} \tag{176}$$

Thus, we know that

$$\frac{1}{np} \left| \|\mathbf{W}_1 \bar{\mathbf{X}}\|_4^4 - \|\mathbf{W}_2 \bar{\mathbf{X}}\|_4^4 \right| \leq \frac{1}{np} \Gamma_1 \Gamma_2 \leq 4npB^4 \|\mathbf{W}_1 - \mathbf{W}_2\|_2. \tag{177}$$

■

Lemma 39 (Lipschitz Constant of $\frac{1}{np}f(\cdot)$ over $\mathcal{O}(n; \mathbb{R})$) *If $\mathbf{X} \in \mathbb{R}^{n \times p}$, $x_{i,j} \sim_{iid} BG(\theta)$, then $\forall \mathbf{W}_1, \mathbf{W}_2 \in \mathcal{O}(n; \mathbb{R})$, we have*

$$\frac{1}{np} \left| \mathbb{E} \|\mathbf{W}_1 \mathbf{X}\|_4^4 - \mathbb{E} \|\mathbf{W}_2 \mathbf{X}\|_4^4 \right| \leq L_2 \|\mathbf{W}_1 - \mathbf{W}_2\|_2, \tag{178}$$

with a constant $L_2 \leq 12n\theta(1 - \theta)$.

Proof According to Lemma 3, we have

$$\mathbb{E} \|\mathbf{W} \mathbf{X}\|_4^4 = 3p\theta(1 - \theta) \|\mathbf{W}\|_4^4 + 3\theta^2 np, \quad \forall \mathbf{W} \in \mathcal{O}(n; \mathbb{R}), \tag{179}$$

so

$$\left| \mathbb{E} \|\mathbf{W}_1 \mathbf{X}\|_4^4 - \mathbb{E} \|\mathbf{W}_2 \mathbf{X}\|_4^4 \right| = 3p\theta(1 - \theta) \left| \|\mathbf{W}_1\|_4^4 - \|\mathbf{W}_2\|_4^4 \right|. \tag{180}$$

Notice that

$$\begin{aligned}
 \left| \|\mathbf{W}_1\|_4^4 - \|\mathbf{W}_2\|_4^4 \right| &= \left| \sum_{i,j} \left[(\mathbf{W}_1)_{i,j}^4 - (\mathbf{W}_2)_{i,j}^4 \right] \right| \\
 &= \left| \sum_{i,j} \left\{ \left[(\mathbf{W}_1)_{i,j}^2 - (\mathbf{W}_2)_{i,j}^2 \right] \left[(\mathbf{W}_1)_{i,j}^2 + (\mathbf{W}_2)_{i,j}^2 \right] \right\} \right| \\
 &= \left\{ \sum_{i,j} \left[(\mathbf{W}_1)_{i,j}^2 - (\mathbf{W}_2)_{i,j}^2 \right]^2 \right\}^{1/2} \left\{ \sum_{i,j} \left[(\mathbf{W}_1)_{i,j}^2 + (\mathbf{W}_2)_{i,j}^2 \right]^2 \right\}^{1/2} \quad (\text{Cauchy-Schwarz}) \\
 &= \|\mathbf{W}_1^{\circ 2} - \mathbf{W}_2^{\circ 2}\|_F \|\mathbf{W}_1^{\circ 2} + \mathbf{W}_2^{\circ 2}\|_F = \|(\mathbf{W}_1 - \mathbf{W}_2) \circ (\mathbf{W}_1 + \mathbf{W}_2)\|_F \|\mathbf{W}_1^{\circ 2} + \mathbf{W}_2^{\circ 2}\|_F \\
 &\leq \|\mathbf{W}_1 - \mathbf{W}_2\|_F \|\mathbf{W}_1 + \mathbf{W}_2\|_F (\|\mathbf{W}_1^{\circ 2}\|_F + \|\mathbf{W}_2^{\circ 2}\|_F) \quad (\text{By Lemma 33}) \\
 &\leq n \|\mathbf{W}_1 - \mathbf{W}_2\|_2 \|\mathbf{W}_1 + \mathbf{W}_2\|_2 (\|\mathbf{W}_1\|_F^2 + \|\mathbf{W}_2\|_F^2) \quad (\text{By Lemma 33, } \|\mathbf{W}\|_F \leq \sqrt{n} \|\mathbf{W}\|_2) \\
 &\leq 4n^2 \|\mathbf{W}_1 - \mathbf{W}_2\|_2 \quad (\|\mathbf{W}_1 + \mathbf{W}_2\|_2 \leq \|\mathbf{W}_1\|_2 + \|\mathbf{W}_2\|_2 \text{ and } \|\mathbf{W}_1\|_2 = \|\mathbf{W}_2\|_2 = 1).
 \end{aligned} \tag{181}$$

Hence,

$$\frac{1}{np} \left| \mathbb{E} \|\mathbf{W}_1 \mathbf{X}\|_4^4 - \mathbb{E} \|\mathbf{W}_2 \mathbf{X}\|_4^4 \right| \leq 12n\theta(1-\theta) \|\mathbf{W}_1 - \mathbf{W}_2\|_2, \tag{182}$$

which completes the proof. $\blacksquare$

Lemma 40 (Lipschitz Constant of $\nabla \hat{f}(\cdot, \cdot)$) *If $\mathbf{X} \in \mathbb{R}^{n \times p}$, $x_{i,j} \sim_{iid} BG(\theta)$, and let $\bar{\mathbf{X}}$ be the truncation of $\mathbf{X}$ by bound B*

$$\bar{x}_{i,j} = \begin{cases} x_{i,j} & \text{if } |x_{i,j}| \leq B \\ 0 & \text{else} \end{cases}, \tag{183}$$

then $\forall \mathbf{W}_1, \mathbf{W}_2 \in \mathcal{O}(n; \mathbb{R})$, we have

$$\frac{1}{np} \left\| (\mathbf{W}_1 \bar{\mathbf{X}})^{\circ 3} \bar{\mathbf{X}}^* - (\mathbf{W}_2 \bar{\mathbf{X}})^{\circ 3} \bar{\mathbf{X}}^* \right\|_F \leq L_1 \|\mathbf{W}_1 - \mathbf{W}_2\|_2, \tag{184}$$

with a constant $L_1 \leq 3npB^4$.

Proof Notice that

$$\begin{aligned}
 &\frac{1}{np} \left\| (\mathbf{W}_1 \bar{\mathbf{X}})^{\circ 3} \bar{\mathbf{X}}^* - (\mathbf{W}_2 \bar{\mathbf{X}})^{\circ 3} \bar{\mathbf{X}}^* \right\|_F = \frac{1}{np} \left\| \left[(\mathbf{W}_1 \bar{\mathbf{X}})^{\circ 3} - (\mathbf{W}_2 \bar{\mathbf{X}})^{\circ 3} \right] \bar{\mathbf{X}}^* \right\|_F \\
 &\leq \frac{1}{np} \left\| (\mathbf{W}_1 \bar{\mathbf{X}} - \mathbf{W}_2 \bar{\mathbf{X}}) \circ \left[(\mathbf{W}_1 \bar{\mathbf{X}})^{\circ 2} + (\mathbf{W}_2 \bar{\mathbf{X}})^{\circ 2} + (\mathbf{W}_1 \bar{\mathbf{X}}) \circ (\mathbf{W}_2 \bar{\mathbf{X}}) \right] \right\|_F \|\bar{\mathbf{X}}^*\|_F \quad (185) \\
 &\leq \frac{1}{np} \underbrace{\|\mathbf{W}_1 \bar{\mathbf{X}} - \mathbf{W}_2 \bar{\mathbf{X}}\|_F}_{\Gamma_1} \underbrace{\left\| (\mathbf{W}_1 \bar{\mathbf{X}})^{\circ 2} + (\mathbf{W}_2 \bar{\mathbf{X}})^{\circ 2} + (\mathbf{W}_1 \bar{\mathbf{X}}) \circ (\mathbf{W}_2 \bar{\mathbf{X}}) \right\|_F}_{\Gamma_2} \|\bar{\mathbf{X}}^*\|_F,
 \end{aligned}$$

where the $\leq$ is achieved by inequality 1 in Lemma 33. For Γ_1 , using inequality 2 in Lemma 33, we have

$$\Gamma_1 = \|(\mathbf{W}_1 - \mathbf{W}_2)\bar{\mathbf{X}}\|_F \leq \|\mathbf{W}_1 - \mathbf{W}_2\|_2 \|\bar{\mathbf{X}}\|_F. \quad (186)$$

For Γ_2 , we have

$$\begin{aligned} \Gamma_2 &= \|(\mathbf{W}_1\bar{\mathbf{X}})^{\circ 2} + (\mathbf{W}_2\bar{\mathbf{X}})^{\circ 2} + (\mathbf{W}_1\bar{\mathbf{X}}) \circ (\mathbf{W}_2\bar{\mathbf{X}})\|_F \\ &\leq \left(\|(\mathbf{W}_1\bar{\mathbf{X}})^{\circ 2}\|_F + \|(\mathbf{W}_2\bar{\mathbf{X}})^{\circ 2}\|_F + \|(\mathbf{W}_1\bar{\mathbf{X}}) \circ (\mathbf{W}_2\bar{\mathbf{X}})\|_F \right) \\ &\leq \left(\|\mathbf{W}_1\bar{\mathbf{X}}\|_F^2 + \|\mathbf{W}_2\bar{\mathbf{X}}\|_F^2 + \|\mathbf{W}_1\bar{\mathbf{X}}\|_F \|\mathbf{W}_2\bar{\mathbf{X}}\|_F \right) \quad (\text{By inequality 1 in Lemma 33}) \\ &= 3 \|\bar{\mathbf{X}}\|_F^2 \quad (\|\cdot\|_F \text{ is rotation invariant}). \end{aligned} \quad (187)$$

Hence, we have

$$\begin{aligned} &\frac{1}{np} \|(\mathbf{W}_1\bar{\mathbf{X}})^{\circ 3} \bar{\mathbf{X}}^* - (\mathbf{W}_2\bar{\mathbf{X}})^{\circ 3} \bar{\mathbf{X}}^*\|_F \leq \frac{1}{np} \Gamma_1 \Gamma_2 \|\bar{\mathbf{X}}\|_F \\ &\leq \frac{3}{np} \|\bar{\mathbf{X}}\|_F^4 \|\mathbf{W}_1 - \mathbf{W}_2\|_2 = 3npB^4 \|\mathbf{W}_1 - \mathbf{W}_2\|_2, \end{aligned} \quad (188)$$

which completes the proof. $\blacksquare$

Lemma 41 (Lipschitz Constant of $\mathbb{E}_p^1 \nabla \hat{f}(\cdot, \cdot)$) *If $\mathbf{X} \in \mathbb{R}^{n \times p}$, $x_{i,j} \sim_{iid} BG(\theta)$, then $\forall \mathbf{W}_1, \mathbf{W}_2 \in \mathcal{O}(n; \mathbb{R})$, we have*

$$\frac{1}{np} \|\mathbb{E}[(\mathbf{W}_1\bar{\mathbf{X}}^*)^{\circ 3} \bar{\mathbf{X}}] - \mathbb{E}[(\mathbf{W}_2\bar{\mathbf{X}})^{\circ 3} \bar{\mathbf{X}}^*]\|_F \leq L_2 \|\mathbf{W}_1 - \mathbf{W}_2\|_2, \quad (189)$$

with a constant $L_2 \leq 9\theta(1 - \theta) + \frac{3\theta^2}{\sqrt{n}}$.

Proof By (99) in proof B.2 of Proposition 10, we have

$$\mathbb{E}[(\mathbf{W}\bar{\mathbf{X}})^{\circ 3} \bar{\mathbf{X}}] = 3p\theta(1 - \theta)\mathbf{W}^{\circ 3} + 3p\theta^2\mathbf{W}, \quad \forall \mathbf{W} \in \mathcal{O}(n; \mathbb{R}). \quad (190)$$

Hence, we have:

$$\begin{aligned} &\frac{1}{np} \|\mathbb{E}[(\mathbf{W}_1\bar{\mathbf{X}}^*)^{\circ 3} \bar{\mathbf{X}}] - \mathbb{E}[(\mathbf{W}_2\bar{\mathbf{X}})^{\circ 3} \bar{\mathbf{X}}^*]\|_F \\ &= \frac{1}{n} \|3\theta(1 - \theta)(\mathbf{W}_1^{\circ 3} - \mathbf{W}_2^{\circ 3}) + 3\theta^2(\mathbf{W}_1 - \mathbf{W}_2)\|_F \\ &\leq \frac{3\theta(1 - \theta)}{n} \|\mathbf{W}_1^{\circ 3} - \mathbf{W}_2^{\circ 3}\|_F + \frac{3\theta^2}{n} \|\mathbf{W}_1 - \mathbf{W}_2\|_F \\ &\leq \frac{3\theta(1 - \theta)}{n} \underbrace{\|\mathbf{W}_1^{\circ 3} - \mathbf{W}_2^{\circ 3}\|_F}_{\Gamma} + \frac{3\theta^2}{\sqrt{n}} \|\mathbf{W}_1 - \mathbf{W}_2\|_2. \end{aligned} \quad (191)$$

Note that

$$\begin{aligned}
 \Gamma &= \|\mathbf{W}_1^{\circ 3} - \mathbf{W}_2^{\circ 3}\|_F = \|(\mathbf{W}_1 - \mathbf{W}_2)(\mathbf{W}_1^{\circ 2} + \mathbf{W}_2^{\circ 2} + \mathbf{W}_1 \circ \mathbf{W}_2)\|_F \\
 &\leq \|\mathbf{W}_1 - \mathbf{W}_2\|_2 \|\mathbf{W}_1^{\circ 2} + \mathbf{W}_2^{\circ 2} + \mathbf{W}_1 \circ \mathbf{W}_2\|_F \quad (\text{By inequality 2 in Lemma 33}) \\
 &\leq (\|\mathbf{W}_1\|_F^2 + \|\mathbf{W}_2\|_F^2 + \|\mathbf{W}_1\|_F \|\mathbf{W}_2\|_F) \|\mathbf{W}_1 - \mathbf{W}_2\|_2 = 3n \|\mathbf{W}_1 - \mathbf{W}_2\|_2,
 \end{aligned} \tag{192}$$

therefore, we have:

$$\frac{1}{np} \|\mathbb{E}[(\mathbf{W}_1 \bar{\mathbf{X}}^*)^{\circ 3} \bar{\mathbf{X}}] - \mathbb{E}[(\mathbf{W}_2 \bar{\mathbf{X}})^{\circ 3} \bar{\mathbf{X}}^*]\|_F \leq \left(9\theta(1-\theta) + \frac{3\theta^2}{\sqrt{n}}\right) \|\mathbf{W}_1 - \mathbf{W}_2\|_2, \tag{193}$$

which completes the proof. $\blacksquare$

D.7. High Order Moment Bound of Bernoulli Gaussian Random Variables

Lemma 42 (Second Moment of $\|\cdot\|_4^4$) Assume that $\mathbf{W} \in \mathcal{O}(n; \mathbb{R})$, $\mathbf{x} \in \mathbb{R}^n$, $x_i \sim_{iid} BG(\theta)$, $\forall i \in [n]$, then the second order moment of $\|\mathbf{W}\mathbf{x}\|_4^4$ satisfies

$$\mathbb{E} \|\mathbf{W}\mathbf{x}\|_4^8 \leq Cn^2\theta, \tag{194}$$

for a constant $C > 105$.

Proof Assume $\mathbf{v}$ is a Gaussian vector and the support for Bernoulli-Gaussian vector $\mathbf{x}$ is $\mathcal{S}$, that is, $\forall i \in [n]$,

$$x_i = \begin{cases} v, v \sim \mathcal{N}(0, 1) & \text{if } i \in \mathcal{S}, \\ 0 & \text{otherwise.} \end{cases} \tag{195}$$

Let $\mathbf{P}_{\mathcal{S}} : \mathbb{R}^n \mapsto \mathbb{R}^n$ be the projection onto set $\mathcal{S}$, that is, $\forall \mathbf{q} \in \mathbb{R}^n$

$$(\mathbf{P}_{\mathcal{S}}\mathbf{q})_i = \begin{cases} q_i & \text{if } i \in \mathcal{S}, \\ 0 & \text{otherwise.} \end{cases} \tag{196}$$

Let $\mathbf{w}_i$ denote the i^{th} row vector of $\mathbf{W}$, so

$$\begin{aligned}
 \mathbb{E} \|\mathbf{W}\mathbf{x}\|_4^8 &= \mathbb{E} \left[(\|\mathbf{W}\mathbf{x}\|_4^4)^2 \right] = \mathbb{E} \left[\sum_{i=1}^n \sum_{j=1}^n \langle \mathbf{w}_i, \mathbf{x} \rangle^4 \langle \mathbf{w}_j, \mathbf{x} \rangle^4 \right] \\
 &= \mathbb{E}_{\mathcal{S}} \sum_{i=1}^n \sum_{j=1}^n \mathbb{E} \left[\langle \mathbf{P}_{\mathcal{S}}\mathbf{w}_i, \mathbf{v} \rangle^4 \langle \mathbf{P}_{\mathcal{S}}\mathbf{w}_j, \mathbf{v} \rangle^4 \right] \\
 &\leq \mathbb{E}_{\mathcal{S}} \sum_{i=1}^n \sum_{j=1}^n \left[\mathbb{E} \langle \mathbf{P}_{\mathcal{S}}\mathbf{w}_i, \mathbf{v} \rangle^8 \mathbb{E} \langle \mathbf{P}_{\mathcal{S}}\mathbf{w}_j, \mathbf{v} \rangle^8 \right]^{\frac{1}{2}} \quad (\text{Cauchy-Schwarz}).
 \end{aligned} \tag{197}$$

Since $\mathbf{v} \sim \mathcal{N}(\mathbf{0}, \mathbf{I})$, so

$$\langle \mathbf{P}_{\mathcal{S}}\mathbf{w}_i, \mathbf{v} \rangle = \sum_{k=1}^n (\mathbf{P}_{\mathcal{S}}\mathbf{w}_i)_k v_k \sim \mathcal{N}(0, \|\mathbf{P}_{\mathcal{S}}\mathbf{w}_i\|_2^2). \tag{198}$$

Therefore,

$$\mathbb{E} \langle \mathbf{P}_S \mathbf{w}_i, \mathbf{v} \rangle^8 = 105 \|\mathbf{P}_S \mathbf{w}_i\|_2^8. \quad (199)$$

Hence, combine (197), we have

$$\begin{aligned} \mathbb{E} \|\mathbf{W} \mathbf{x}\|_4^8 &\leq \mathbb{E}_S \sum_{i=1}^n \sum_{j=1}^n \left[\mathbb{E} \langle \mathbf{P}_S \mathbf{w}_i, \mathbf{v} \rangle^8 \mathbb{E} \langle \mathbf{P}_S \mathbf{w}_j, \mathbf{v} \rangle^8 \right]^{\frac{1}{2}} \\ &= 105 \mathbb{E}_S \sum_{i=1}^n \sum_{j=1}^n \|\mathbf{P}_S \mathbf{w}_i\|_2^4 \|\mathbf{P}_S \mathbf{w}_j\|_2^4 \\ &= 105 \sum_{i=1}^n \sum_{j=1}^n \sum_{k_1, k_2, k_3, k_4} \mathbb{E}_S \left[w_{i,k_1}^2 \mathbb{1}_{k_1 \in S} w_{i,k_2}^2 \mathbb{1}_{k_2 \in S} w_{j,k_3}^2 \mathbb{1}_{k_3 \in S} w_{j,k_4}^2 \mathbb{1}_{k_4 \in S} \right]. \end{aligned} \quad (200)$$

Now we discuss these four different cases separately:

- With probability $c_1 \theta^4$ ($c_1 \leq 1$), all $k_1, k_2, k_3, k_4 \in S$, in this case, we have:

$$\begin{aligned} &\sum_{i=1}^n \sum_{j=1}^n \sum_{k_1, k_2, k_3, k_4} \mathbb{E}_S \left[w_{i,k_1}^2 \mathbb{1}_{k_1 \in S} w_{i,k_2}^2 \mathbb{1}_{k_2 \in S} w_{j,k_3}^2 \mathbb{1}_{k_3 \in S} w_{j,k_4}^2 \mathbb{1}_{k_4 \in S} \right] \\ &= \sum_{i=1}^n \sum_{j=1}^n \sum_{k_1, k_2, k_3, k_4} \mathbb{E}_S \left[w_{i,k_1}^2 w_{i,k_2}^2 w_{j,k_3}^2 w_{j,k_4}^2 \right] = \sum_{i=1}^n \sum_{j=1}^n \sum_{k_1, k_2, k_3, k_4} w_{i,k_1}^2 w_{i,k_2}^2 w_{j,k_3}^2 w_{j,k_4}^2 \\ &= \sum_{i=1}^n \sum_{j=1}^n \sum_{k_1, k_2, k_3} w_{i,k_1}^2 w_{i,k_2}^2 w_{j,k_3}^2 = \sum_{i=1}^n \sum_{j=1}^n \sum_{k_1, k_2} w_{i,k_1}^2 w_{i,k_2}^2 = n^2. \end{aligned} \quad (201)$$

- With probability $c_2 \theta^3$ ($c_2 \leq 1$), only three among $k_1, k_2, k_3, k_4 \in S$, in this case, we have:

$$\begin{aligned} &\sum_{i=1}^n \sum_{j=1}^n \sum_{k_1, k_2, k_3, k_4} \mathbb{E}_S \left[w_{i,k_1}^2 \mathbb{1}_{k_1 \in S} w_{i,k_2}^2 \mathbb{1}_{k_2 \in S} w_{j,k_3}^2 \mathbb{1}_{k_3 \in S} w_{j,k_4}^2 \mathbb{1}_{k_4 \in S} \right] \\ &= \sum_{i=1}^n \sum_{j=1}^n \sum_{k_1, k_2, k_3} \left[w_{i,k_1}^2 w_{i,k_2}^2 w_{j,k_3}^4 \right] + \sum_{i=1}^n \sum_{j=1}^n \sum_{k_1, k_2, k_3} \left[w_{i,k_1}^2 w_{i,k_2}^2 w_{j,k_2}^2 w_{j,k_3}^2 \right] = n \|\mathbf{W}\|_4^4 + 1. \end{aligned} \quad (202)$$

- With probability $c_3 \theta^2$ ($c_3 \leq 1$), only two among $k_1, k_2, k_3, k_4 \in S$, in this case, we have:

$$\begin{aligned} &\sum_{i=1}^n \sum_{j=1}^n \sum_{k_1, k_2, k_3, k_4} \mathbb{E}_S \left[w_{i,k_1}^2 \mathbb{1}_{k_1 \in S} w_{i,k_2}^2 \mathbb{1}_{k_2 \in S} w_{j,k_3}^2 \mathbb{1}_{k_3 \in S} w_{j,k_4}^2 \mathbb{1}_{k_4 \in S} \right] \\ &= \sum_{i=1}^n \sum_{j=1}^n \sum_{k_1, k_2} w_{i,k_1}^4 w_{i,k_2}^4 + \sum_{i=1}^n \sum_{j=1}^n \sum_{k_1, k_2} w_{i,k_1}^4 w_{j,k_1}^2 w_{j,k_2}^2 + \sum_{i=1}^n \sum_{j=1}^n \sum_{k_1, k_2} w_{i,k_1}^2 w_{i,k_2}^2 w_{j,k_1}^2 w_{j,k_2}^2 \\ &\leq 3 \sum_{i=1}^n \sum_{j=1}^n \sum_{k_1, k_2} w_{i,k_1}^4 w_{i,k_2}^4 = 3 \|\mathbf{W}\|_4^8. \end{aligned} \quad (203)$$

The only inequality above is achieved by Rearrangement inequality.

- With probability $c_4\theta$ ($c_4 \leq 1$), only one among $k_1, k_2, k_3, k_4 \in \mathcal{S}$, in this case, we have:

$$\begin{aligned} & \sum_{i=1}^n \sum_{j=1}^n \sum_{k_1, k_2, k_3, k_4} \mathbb{E}_{\mathcal{S}} \left[w_{i,k_1}^2 \mathbb{1}_{k_1 \in \mathcal{S}} w_{i,k_2}^2 \mathbb{1}_{k_2 \in \mathcal{S}} w_{j,k_3}^2 \mathbb{1}_{k_3 \in \mathcal{S}} w_{j,k_4}^2 \mathbb{1}_{k_4 \in \mathcal{S}} \right] \\ &= \sum_{i=1}^n \sum_{j=1}^n \sum_{k_1} w_{i,k_1}^4 w_{j,k_1}^4 \leq \sum_{i=1}^n \sum_{j=1}^n \sum_{k_1} w_{i,k_1}^8 = n \|\mathbf{W}\|_8^8. \end{aligned} \quad (204)$$

The only inequality above is achieved by Rearrangement inequality.

Substitute (201), (202), (203), and (204) into (200), yields

$$\begin{aligned} \mathbb{E} \|\mathbf{W}\mathbf{x}\|_4^8 &\leq 105 \sum_{i=1}^n \sum_{j=1}^n \sum_{k_1, k_2, k_3, k_4} \mathbb{E}_{\mathcal{S}} \left[w_{i,k_1}^2 \mathbb{1}_{k_1 \in \mathcal{S}} w_{i,k_2}^2 \mathbb{1}_{k_2 \in \mathcal{S}} w_{j,k_3}^2 \mathbb{1}_{k_3 \in \mathcal{S}} w_{j,k_4}^2 \mathbb{1}_{k_4 \in \mathcal{S}} \right] \\ &\leq C(\theta^4 n^2 + \theta^3 n \|\mathbf{W}\|_4^4 + \theta^3 + 3\theta^2 \|\mathbf{W}\|_4^8 + \theta n \|\mathbf{W}\|_8^8) \leq Cn^2\theta, \end{aligned} \quad (205)$$

for a constant $C > 105$, which completes the proof. $\blacksquare$

Lemma 43 (Second Moment of $\nabla_{\mathbf{A}} \hat{f}(\mathbf{A}, \mathbf{Y})$) Assume that $\mathbf{W} \in \mathcal{O}(n; \mathbb{R}), \mathbf{X} \in \mathbb{R}^{n \times p}$, $x_{i,j} \sim_{iid} BG(\theta)$, $\forall i \in [n], j \in [p]$, then each element of $\{(\mathbf{W}\mathbf{X}_o)^{\circ 3} \mathbf{X}_o^*\}_{i,j'}$ satisfies

1. $\{(\mathbf{W}\mathbf{X}_o)^{\circ 3} \mathbf{X}_o^*\}_{i,j'}$ can be represented as sum of p i.i.d random variables z_j

$$\{(\mathbf{W}\mathbf{X}_o)^{\circ 3} \mathbf{X}_o^*\}_{i,j'} = \sum_{j=1}^p z_j, \quad \forall i, j' \in [n]. \quad (206)$$

2. The second moment of z_j is bounded by $C\theta$

$$\mathbb{E} z_j^2 \leq C, \quad \forall j \in [p]. \quad (207)$$

for a constant $C > 177$.

Proof Assume $\mathbf{v}$ is a Gaussian vector and the support for Bernoulli-Gaussian vector $\mathbf{x}$ is $\mathcal{S}$, that is, $\forall i \in [n]$,

$$x_i = \begin{cases} v, & v \sim \mathcal{N}(0, 1) & \text{if } i \in \mathcal{S}, \\ 0 & & \text{otherwise.} \end{cases} \quad (208)$$

Let $\mathbf{P}_{\mathcal{S}} : \mathbb{R}^n \mapsto \mathbb{R}^n$ be the projection onto set $\mathcal{S}$, that is, $\forall \mathbf{q} \in \mathbb{R}^n$

$$(\mathbf{P}_{\mathcal{S}} \mathbf{q})_i = \begin{cases} q_i & \text{if } i \in \mathcal{S}, \\ 0 & \text{otherwise.} \end{cases} \quad (209)$$

By (95) and (96) in proof B.2 of Proposition 10, we know that

$$\frac{1}{4} \nabla_{\mathbf{A}} \hat{f}(\mathbf{A}, \mathbf{Y}) = (\mathbf{A}\mathbf{Y})^{\circ 3} \mathbf{Y}^* = (\mathbf{W}\mathbf{X}_o)^{\circ 3} \mathbf{X}_o^* \mathbf{D}_o^*, \quad (210)$$

and

$$\{(\mathbf{W}\mathbf{X}_o)^{\circ 3}\mathbf{X}_o^*\}_{i,j'} = \sum_{j=1}^p \left[x_{j',j} \left(\sum_{k=1}^n w_{i,k} x_{k,j} \right)^3 \right] = \sum_{j=1}^p \left[x_{j',j} \langle \mathbf{w}_i, \mathbf{x}_j \rangle^3 \right], \quad (211)$$

where $\mathbf{w}_i$ is the i^{th} row vector of $\mathbf{W}$. Notice that $x_{i,j}$ are independent with each other, if we define p random variables $z_1, z_2, \dots, z_p$ as the following

$$z_j = x_{j',j} \langle \mathbf{w}_i, \mathbf{x}_j \rangle, \forall j \in [p], \quad (212)$$

then we can view $\{(\mathbf{W}\mathbf{X})^{\circ 3}\mathbf{X}^*\}_{i,j'}$ as the mean of p i.i.d. random variables z_j . Notice that

$$\mathbb{E}(z_j^2) = \mathbb{E} \left[x_{j',j}^2 \left(\sum_{k=1}^n w_{i,k} x_{k,j} \right)^6 \right] \leq \left(\mathbb{E} x_{j',j}^4 \right)^{\frac{1}{2}} \left(\mathbb{E} \langle \mathbf{w}_i, \mathbf{x}_j \rangle^{12} \right)^{\frac{1}{2}} = \sqrt{3}\theta^{\frac{1}{2}} \left(\mathbb{E}_{\mathcal{S}} \mathbb{E}_{\mathbf{v}} \langle \mathbf{P}_{\mathcal{S}} \mathbf{w}_i, \mathbf{v} \rangle^{12} \right)^{\frac{1}{2}}, \quad (213)$$

and

$$\langle \mathbf{P}_{\mathcal{S}} \mathbf{w}_i, \mathbf{v} \rangle = \sum_{k=1}^n (\mathbf{P}_{\mathcal{S}} \mathbf{w}_i)_k v_k \sim \mathcal{N}(0, \|\mathbf{P}_{\mathcal{S}} \mathbf{w}_i\|_2^2). \quad (214)$$

So we have

$$\left(\mathbb{E}_{\mathcal{S}} \mathbb{E}_{\mathbf{v}} \langle \mathbf{P}_{\mathcal{S}} \mathbf{w}_i, \mathbf{v} \rangle^{12} \right)^{\frac{1}{2}} = \left(11!! \mathbb{E}_{\mathcal{S}} \|\mathbf{P}_{\mathcal{S}} \mathbf{w}_i\|_2^{12} \right)^{\frac{1}{2}} = \sqrt{11!!} \left(\mathbb{E}_{\mathcal{S}} \|\mathbf{P}_{\mathcal{S}} \mathbf{w}_i\|_2^{12} \right)^{\frac{1}{2}}. \quad (215)$$

Follow the same pipe line (201), (202), (203), and (204) in Lemma 42, one can show that

$$\begin{aligned} \mathbb{E}_{\mathcal{S}} \|\mathbf{P}_{\mathcal{S}} \mathbf{w}_i\|_2^{12} &= \sum_{k_1, k_2, \dots, k_6} w_{i,k_1}^2 \mathbb{1}_{k_1 \in \mathcal{S}} w_{i,k_2}^2 \mathbb{1}_{k_2 \in \mathcal{S}} w_{i,k_3}^2 \mathbb{1}_{k_3 \in \mathcal{S}} w_{i,k_4}^2 \mathbb{1}_{k_4 \in \mathcal{S}} w_{i,k_5}^2 \mathbb{1}_{k_5 \in \mathcal{S}} w_{i,k_6}^2 \mathbb{1}_{k_6 \in \mathcal{S}} \\ &\leq C'(\theta^6 + \theta^5 + \theta^4 + \theta^3 + \theta^2 + \theta) \leq C''\theta, \end{aligned} \quad (216)$$

for some constant $C', C'' > 1$. Therefore, combine (213), (215), and (216), we have

$$\mathbb{E}(z_j^2) \leq \sqrt{3}\theta^{\frac{1}{2}} \left(\mathbb{E}_{\mathcal{S}} \mathbb{E}_{\mathbf{v}} \langle \mathbf{P}_{\mathcal{S}} \mathbf{w}_i, \mathbf{v} \rangle^{12} \right)^{\frac{1}{2}} \leq \sqrt{3 \times 11!!\theta^2} \leq C\theta, \quad (217)$$

for a constant $C > 177$, which completes the proof. $\blacksquare$

D.8. Union Tail Concentration Bound

Lemma 44 (Union Tail Concentration Bound of $(\mathbf{W}\mathbf{X})^{\circ 3}\mathbf{X}^*$) *If $\mathbf{X} \in \mathbb{R}^{n \times p}$, $x_{i,j} \sim_{iid} BG(\theta)$, the following inequality holds*

$$\begin{aligned} &\mathbb{P} \left(\sup_{\mathbf{W} \in O(n; \mathbb{R})} \frac{1}{np} \|(\mathbf{W}\mathbf{X})^{\circ 3}\mathbf{X}^* - \mathbb{E}[(\mathbf{W}\mathbf{X})^{\circ 3}\mathbf{X}^*]\|_F \geq \delta \right) \\ &\leq 2n^2 \exp \left(- \frac{3p\delta^2}{c_1\theta + 8n^{\frac{3}{2}}(\ln p)^4\delta} + n^2 \ln \left(\frac{48np(\ln p)^4}{\delta} \right) \right) + 2np\theta \exp \left(- \frac{(\ln p)^2}{2} \right), \end{aligned} \quad (218)$$

for a constant $c_1 > 1.7 \times 10^4$.

Proof Let $\bar{\mathbf{X}} \in \mathbb{R}^{n \times p}$ denote the truncated $\mathbf{X}$ by bound B

$$\bar{x}_{i,j} = \begin{cases} x_{i,j} & \text{if } |x_{i,j}| \leq B, \\ 0 & \text{else.} \end{cases} \quad (219)$$

Note that $\bar{\mathbf{X}} = \mathbf{X}$ holds whenever $\|\mathbf{X}\|_\infty \leq B$, and by Lemma 34, we know that with probability $\|\mathbf{X}\|_\infty \leq B$ happens with probability at least $1 - 2np\theta \exp(-B^2/2)$. So we know that $\bar{\mathbf{X}} \neq \mathbf{X}$ holds with probability at most $2np\theta \exp(-B^2/2)$, and thus

$$\begin{aligned} & \mathbb{P} \left(\sup_{\mathbf{W} \in \mathcal{O}(n; \mathbb{R})} \frac{1}{np} \|(\mathbf{W}\mathbf{X})^{\circ 3} \mathbf{X}^* - \mathbb{E}[(\mathbf{W}\mathbf{X})^{\circ 3} \mathbf{X}^*]\|_F > \delta \right) \\ & \leq \mathbb{P} \left(\sup_{\mathbf{W} \in \mathcal{O}(n; \mathbb{R})} \frac{1}{np} \|(\mathbf{W}\mathbf{X})^{\circ 3} \mathbf{X}^* - \mathbb{E}[(\mathbf{W}\mathbf{X})^{\circ 3} \mathbf{X}^*]\|_F > \delta, \mathbf{X} = \bar{\mathbf{X}} \right) + \mathbb{P}(\mathbf{X} \neq \bar{\mathbf{X}}) \quad (220) \\ & \leq \mathbb{P} \left(\sup_{\mathbf{W} \in \mathcal{O}(n; \mathbb{R})} \frac{1}{np} \|(\mathbf{W}\bar{\mathbf{X}})^{\circ 3} \bar{\mathbf{X}}^* - \mathbb{E}[(\mathbf{W}\mathbf{X})^{\circ 3} \mathbf{X}^*]\|_F > \delta \right) + 2np\theta e^{-\frac{B^2}{2}} \end{aligned}$$

ε -**net Covering.** For any positive ε satisfy

$$\varepsilon \leq \frac{\delta}{8npB^4}, \quad (221)$$

Lemma 35 shows there exists an ε -nets

$$\mathcal{S}_\varepsilon = \{\mathbf{W}_1, \mathbf{W}_2, \dots, \mathbf{W}_{|\mathcal{S}_\varepsilon|}\}, \quad (222)$$

which covers $\mathcal{O}(n; \mathbb{R})$

$$\mathcal{O}(n; \mathbb{R}) \subset \bigcup_{l=1}^{|\mathcal{S}_\varepsilon|} \mathbb{B}(\mathbf{W}_l, \varepsilon), \quad (223)$$

in operator norm $\|\cdot\|_2$. Moreover, we have

$$|\mathcal{S}_\varepsilon| \leq \left(\frac{6}{\varepsilon}\right)^{n^2}. \quad (224)$$

So $\forall \mathbf{W} \in \mathcal{O}(n; \mathbb{R})$, there exists $l \in [\mathcal{S}_\varepsilon]$, such that $\|\mathbf{W} - \mathbf{W}_l\|_2 \leq \varepsilon$. Thus, we have:

$$\begin{aligned}
 & \mathbb{P}\left(\frac{1}{np} \|(\mathbf{W}\bar{\mathbf{X}})^{\circ 3} \bar{\mathbf{X}}^* - \mathbb{E}[(\mathbf{W}\mathbf{X})^{\circ 3} \mathbf{X}^*]\|_F \leq \delta\right) \\
 & \leq \mathbb{P}\left(\frac{1}{np} \left(\|(\mathbf{W}\bar{\mathbf{X}})^{\circ 3} \bar{\mathbf{X}}^* - (\mathbf{W}_l\bar{\mathbf{X}})^{\circ 3} \bar{\mathbf{X}}^*\|_F + \|(\mathbf{W}_l\bar{\mathbf{X}})^{\circ 3} \bar{\mathbf{X}}^* - \mathbb{E}[(\mathbf{W}_l\mathbf{X})^{\circ 3} \mathbf{X}^*]\|_F\right.\right. \\
 & \quad \left.\left.+ \|\mathbb{E}[(\mathbf{W}_l\mathbf{X})^{\circ 3} \mathbf{X}^*] - \mathbb{E}[(\mathbf{W}\mathbf{X})^{\circ 3} \mathbf{X}^*]\|_F\right) \geq \delta\right) \\
 & \leq \mathbb{P}\left(\frac{1}{np} \|(\mathbf{W}_l\bar{\mathbf{X}})^{\circ 3} \bar{\mathbf{X}}^* - \mathbb{E}[(\mathbf{W}_l\mathbf{X})^{\circ 3} \mathbf{X}^*]\|_F\right. \\
 & \quad \left.+ \left(3npB^4 + 9\theta(1-\theta) + \frac{3\theta^2}{\sqrt{n}}\right) \|\mathbf{W} - \mathbf{W}_l\|_2 \geq \delta\right) \quad (\text{By Lemma 40, 41}) \\
 & \leq \mathbb{P}\left(\frac{1}{np} \|(\mathbf{W}_l\bar{\mathbf{X}})^{\circ 3} \bar{\mathbf{X}}^* - \mathbb{E}[(\mathbf{W}_l\mathbf{X})^{\circ 3} \mathbf{X}^*]\|_F + 4npB^4\varepsilon \geq \delta\right) \quad (p, B \text{ are large, } \|\mathbf{W} - \mathbf{W}_l\|_2 \leq \varepsilon) \\
 & \leq \mathbb{P}\left(\frac{1}{np} \|(\mathbf{W}_l\bar{\mathbf{X}})^{\circ 3} \bar{\mathbf{X}}^* - \mathbb{E}[(\mathbf{W}_l\mathbf{X})^{\circ 3} \mathbf{X}^*]\|_F + \frac{\delta}{2} \geq \delta\right) \quad (\text{We assume } \varepsilon \leq \frac{\delta}{8npB^4}) \\
 & = P\left(\frac{1}{np} \|(\mathbf{W}_l\bar{\mathbf{X}})^{\circ 3} \bar{\mathbf{X}} - \mathbb{E}[(\mathbf{W}_l\mathbf{X})^{\circ 3} \mathbf{X}^*]\|_F \geq \frac{\delta}{2}\right).
 \end{aligned} \tag{225}$$

Analysis. For random variable $\bar{\mathbf{X}}$, we have

$$\begin{aligned}
 & \|\mathbb{E}[(\mathbf{W}_l\mathbf{X})^{\circ 3} \mathbf{X}^*] - \mathbb{E}[(\mathbf{W}_l\bar{\mathbf{X}})^{\circ 3} \bar{\mathbf{X}}^*]\|_F = \left\| \mathbb{E}[(\mathbf{W}_l\mathbf{X})^{\circ 3} \mathbf{X}^*] - \mathbb{E}[(\mathbf{W}_l\mathbf{X})^{\circ 3} \mathbf{X}^* \cdot \mathbb{1}_{\{\|\mathbf{X}\|_\infty \leq B\}}] \right\|_F \\
 & = \left\| \mathbb{E}[(\mathbf{W}_l\mathbf{X})^{\circ 3} \mathbf{X}^* \cdot \mathbb{1}_{\{\|\mathbf{X}\|_\infty > B\}}] \right\|_F = \sqrt{\sum_{\substack{1 \leq i \leq n \\ 1 \leq j' \leq n}} \left(\mathbb{E}\left\{[(\mathbf{W}_l\mathbf{X})^{\circ 3} \mathbf{X}^*]_{i,j'} \cdot \mathbb{1}_{\{\|\mathbf{X}\|_\infty > B\}}\right\} \right)^2} \\
 & \leq \sqrt{\sum_{\substack{1 \leq i \leq n \\ 1 \leq j' \leq n}} \left(\left\{ \mathbb{E}[(\mathbf{W}_l\mathbf{X})^{\circ 3} \mathbf{X}^*]_{i,j'}^2 \right\} \left\{ \mathbb{E} \mathbb{1}_{\{\|\mathbf{X}\|_\infty > B\}} \right\} \right)} = \|\mathbb{E}[(\mathbf{W}_l\mathbf{X})^{\circ 3} \mathbf{X}^*]\|_F \cdot \sqrt{\mathbb{E} \mathbb{1}_{\{\|\mathbf{X}\|_\infty > B\}}}.
 \end{aligned} \tag{226}$$

Note that in (99) of Lemma B.2, we know that

$$\mathbb{E}[(\mathbf{W}_l\mathbf{X})^{\circ 3} \mathbf{X}^*] = 3p\theta(1-\theta)\mathbf{W}_l^{\circ 3} + 3p\theta^2\mathbf{W}_l, \tag{227}$$

hence

$$\begin{aligned}
 \|\mathbb{E}[(\mathbf{W}_l\mathbf{X})^{\circ 3} \mathbf{X}^*]\|_F & = \|3p\theta(1-\theta)\mathbf{W}_l^{\circ 3} + 3p\theta^2\mathbf{W}_l\|_F \leq 3p\theta(1-\theta) \|\mathbf{W}_l^{\circ 3}\|_F + 3p\theta^2 \|\mathbf{W}_l\|_F \\
 & \leq 3p\theta(1-\theta) \|\mathbf{W}_l\|_F^3 + 3p\theta^2 \|\mathbf{W}_l\|_F \quad (\text{By inequality 1 in Lemma 33}) \\
 & = 3p\theta(1-\theta)n^{\frac{3}{2}} + 3p\theta^2n^{\frac{1}{2}} < 4n^{\frac{3}{2}}p\theta \quad (n \text{ is a large number}).
 \end{aligned} \tag{228}$$

Moreover, Lemma 34 shows that

$$\mathbb{E}\mathbb{1}_{\{\|\mathbf{X}\|_\infty > B\}} \leq 2np\theta^{-\frac{B^2}{2}}. \quad (229)$$

Substitute (229) and (228) into (226), yield

$$\frac{1}{np} \left\| \mathbb{E}[(\mathbf{W}_l \mathbf{X})^{\circ 3} \mathbf{X}^*] - \mathbb{E}[(\mathbf{W}_l \bar{\mathbf{X}})^{\circ 3} \bar{\mathbf{X}}^*] \right\|_F \leq 4\sqrt{2}np^{\frac{1}{2}}\theta^{\frac{3}{2}}e^{-\frac{B^2}{4}}. \quad (230)$$

Hence, when

$$B \geq 2\sqrt{\ln\left(\frac{16\sqrt{2}np^{\frac{1}{2}}\theta^{\frac{3}{2}}}{\delta}\right)}, \quad (231)$$

we have

$$\frac{1}{np} \left\| \mathbb{E}[(\mathbf{W}_l \bar{\mathbf{X}})^{\circ 3} \bar{\mathbf{X}}^*] - \mathbb{E}[(\mathbf{W}_l \mathbf{X})^{\circ 3} \mathbf{X}^*] \right\|_F \leq 4\sqrt{2}np^{\frac{1}{2}}\theta^{\frac{3}{2}}e^{-\frac{B^2}{4}} \leq \frac{\delta}{4}. \quad (232)$$

Therefore, combine (225), we have

$$\begin{aligned} & \mathbb{P}\left(\frac{1}{np} \left\| (\mathbf{W} \bar{\mathbf{X}})^{\circ 3} \bar{\mathbf{X}}^* - \mathbb{E}[(\mathbf{W} \mathbf{X})^{\circ 3} \mathbf{X}^*] \right\|_F \leq \delta\right) \\ & \leq \mathbb{P}\left(\frac{1}{np} \left\| (\mathbf{W}_l \bar{\mathbf{X}})^{\circ 3} \bar{\mathbf{X}}^* - \mathbb{E}[(\mathbf{W}_l \mathbf{X})^{\circ 3} \mathbf{X}^*] \right\|_F \geq \frac{\delta}{2}\right) \\ & \leq \mathbb{P}\left(\frac{1}{np} \left\| (\mathbf{W} \bar{\mathbf{X}})^{\circ 3} \bar{\mathbf{X}}^* - \mathbb{E}[(\mathbf{W}_l \bar{\mathbf{X}})^{\circ 3} \bar{\mathbf{X}}^*] \right\|_F + \frac{1}{np} \left\| \mathbb{E}[(\mathbf{W}_l \bar{\mathbf{X}})^{\circ 3} \bar{\mathbf{X}}^*] - \mathbb{E}[(\mathbf{W}_l \mathbf{X})^{\circ 3} \mathbf{X}^*] \right\|_F \geq \frac{\delta}{2}\right) \\ & \leq \mathbb{P}\left(\frac{1}{np} \left\| (\mathbf{W} \bar{\mathbf{X}})^{\circ 3} \bar{\mathbf{X}}^* - \mathbb{E}[(\mathbf{W}_l \bar{\mathbf{X}})^{\circ 3} \bar{\mathbf{X}}^*] \right\|_F \geq \frac{\delta}{4}\right) \quad (\text{By (232)}) \\ & \leq n^2 \mathbb{P}\left(\left| \left[(\mathbf{W}_l \bar{\mathbf{X}})^{\circ 3} \bar{\mathbf{X}}^* \right]_{i,j'} - \mathbb{E} \left[(\mathbf{W}_l \bar{\mathbf{X}})^{\circ 3} \bar{\mathbf{X}}^* \right]_{i,j'} \right| \geq p \cdot \frac{\delta}{4}\right) \quad (\text{By union bound, } \forall i, j' \in [n]). \end{aligned} \quad (233)$$

Point-wise Bernstein's Inequality. Next, we apply Bernstein's inequality on $\left[(\mathbf{W}_l \bar{\mathbf{X}})^{\circ 3} \bar{\mathbf{X}}^* \right]_{i,j'}$. Note that for each $i, j' \in [n]$

$$\left[(\mathbf{W}_l \bar{\mathbf{X}})^{\circ 3} \bar{\mathbf{X}}^* \right]_{i,j'} = \sum_{j=1}^p \left[\bar{x}_{j',j} \left(\sum_{k=1}^n w_{i,k} \bar{x}_{k,j} \right)^3 \right], \quad (234)$$

can be viewed as sum of p independent variables

$$\bar{z}_j = \bar{x}_{j',j} \left(\sum_{k=1}^n w_{i,k} \bar{x}_{k,j} \right)^3 = \bar{x}_{j',j} \langle \mathbf{w}_i, \bar{\mathbf{x}}_j \rangle^3, \quad \forall j \in [p], \quad (235)$$

where $\mathbf{w}_i$ is the i^{th} row vector of $\mathbf{W}$. Note that each $\bar{z}_j$ are bounded by

$$\left| \bar{x}_{j',j} \langle \mathbf{w}_i, \bar{\mathbf{x}}_j \rangle^3 \right| \leq B \left| \langle \mathbf{w}_i, \bar{\mathbf{x}}_j \rangle^3 \right| = B \|\bar{\mathbf{x}}_j\|_2^3 \left| \left\langle \mathbf{w}_i, \frac{\bar{\mathbf{x}}_j}{\|\bar{\mathbf{x}}_j\|_2} \right\rangle \right| \leq B \cdot (nB^2)^{\frac{3}{2}} = n^{\frac{3}{2}} B^4, \quad (236)$$

also for each $\bar{x}_{j',j} \left(\sum_{k=1}^n w_{i,j} \bar{x}_{k,j} \right)^3$, we have

$$\begin{aligned} & \mathbb{E} \left[x_{j',j}^2 \left(\sum_{k=1}^n w_{i,j} x_{k,j} \right)^6 \right] - \mathbb{E} \left[\bar{x}_{j',j}^2 \left(\sum_{k=1}^n w_{i,j} \bar{x}_{k,j} \right)^6 \right] \\ &= \mathbb{E} \left[x_{j',j}^2 \left(\sum_{k=1}^n w_{i,j} x_{k,j} \right)^6 - \bar{x}_{j',j}^2 \left(\sum_{k=1}^n w_{i,j} \bar{x}_{k,j} \right)^6 \right] \\ &= \mathbb{E} \left[x_{j',j}^2 \left(\sum_{k=1}^n w_{i,j} x_{k,j} \right)^6 \cdot \mathbb{1}_{\{\|\mathbf{x}\|_\infty > B\}} \right] \geq 0, \end{aligned} \quad (237)$$

and (217) in Lemma 43 shows that

$$\mathbb{E} \left[\bar{x}_{j',j}^2 \left(\sum_{k=1}^n w_{i,j} \bar{x}_{k,j} \right)^6 \right] \leq \mathbb{E} \left[x_{j',j}^2 \left(\sum_{k=1}^n w_{i,j} x_{k,j} \right)^6 \right] \leq C, \quad (238)$$

for a constant $C > 177$. Hence, by Bernstein's inequality, $\forall i, j' \in [n]$, we have

$$\begin{aligned} & \mathbb{P} \left(\left| \left[(\mathbf{W}_l \bar{\mathbf{X}})^{\circ 3} \bar{\mathbf{X}}^* \right]_{i,j'} - \mathbb{E} \left[(\mathbf{W}_l \bar{\mathbf{X}})^{\circ 3} \bar{\mathbf{X}}^* \right]_{i,j'} \right| \geq p \cdot \frac{\delta}{4} \right) \\ &= \mathbb{P} \left(\left| \sum_{j=1}^p \left[\bar{x}_{j',j} \left(\sum_{k=1}^n w_{i,k} \bar{x}_{k,j} \right)^3 \right] - \sum_{j=1}^p \mathbb{E} \left[\bar{x}_{j',j} \left(\sum_{k=1}^n w_{i,k} \bar{x}_{k,j} \right)^3 \right] \right| \geq p \cdot \frac{\delta}{4} \right) \\ &= \mathbb{P} \left(\sum_{j=1}^p \left[\bar{x}_{j',j} \left(\sum_{k=1}^n w_{i,k} \bar{x}_{k,j} \right)^3 \right] - \sum_{j=1}^p \mathbb{E} \left[\bar{x}_{j',j} \left(\sum_{k=1}^n w_{i,k} \bar{x}_{k,j} \right)^3 \right] \geq p \cdot \frac{\delta}{4} \right) \\ & \quad + \mathbb{P} \left(\sum_{j=1}^p \left[\bar{x}_{j',j} \left(\sum_{k=1}^n w_{i,k} \bar{x}_{k,j} \right)^3 \right] - \sum_{j=1}^p \mathbb{E} \left[\bar{x}_{j',j} \left(\sum_{k=1}^n w_{i,k} \bar{x}_{k,j} \right)^3 \right] \leq -p \cdot \frac{\delta}{4} \right) \\ &\leq 2 \exp \left(- \frac{p\delta^2/16}{2 \left[\frac{1}{p} \sum_{j=1}^p \bar{x}_{j',j}^2 \left(\sum_{k=1}^n w_{i,k} \bar{x}_{k,j} \right)^6 + \frac{n^{\frac{3}{2}} B^4 \delta}{12} \right]} \right) \quad (\text{By (236)}) \\ &\leq 2 \exp \left(- \frac{p\delta^2}{\frac{96}{p} \sum_{j=1}^p C\theta + 8\delta n^{\frac{3}{2}} B^4 \delta} \right) \quad (\text{By (237)}) \\ &= 2 \exp \left(- \frac{3p\delta^2}{c_1 \theta + 8n^{\frac{3}{2}} B^4 \delta} \right), \end{aligned} \quad (239)$$

for a constant $c_1 > 1.7 \times 10^4$. Combine (233), we have

$$\begin{aligned} & \mathbb{P}\left(\frac{1}{np} \|(\mathbf{W}\bar{\mathbf{X}})^{\circ 3} \bar{\mathbf{X}}^* - \mathbb{E}[(\mathbf{W}\mathbf{X})^{\circ 3} \mathbf{X}^*]\|_F \leq \delta\right) \\ & \leq n^2 \mathbb{P}\left(\left|[(\mathbf{W}_l \bar{\mathbf{X}})^{\circ 3} \bar{\mathbf{X}}^*]_{i,j'} - \mathbb{E}[(\mathbf{W}_l \bar{\mathbf{X}})^{\circ 3} \bar{\mathbf{X}}^*]_{i,j'}\right| \geq p \cdot \frac{\delta}{4}\right). \end{aligned} \quad (240)$$

Union Bound. Now, we will give a union bound for

$$\mathbb{P}\left(\sup_{\mathbf{W} \in \mathcal{O}(n;\mathbb{R})} \frac{1}{np} \|(\mathbf{W}\bar{\mathbf{X}})^{\circ 3} \bar{\mathbf{X}}^* - \mathbb{E}[(\mathbf{W}\mathbf{X})^{\circ 3} \mathbf{X}^*]\|_F > \delta\right). \quad (241)$$

Notice that

$$\begin{aligned} & \mathbb{P}\left(\sup_{\mathbf{W} \in \mathcal{O}(n;\mathbb{R})} \frac{1}{np} \|(\mathbf{W}\bar{\mathbf{X}})^{\circ 3} \bar{\mathbf{X}}^* - \mathbb{E}[(\mathbf{W}\mathbf{X})^{\circ 3} \mathbf{X}^*]\|_F > \delta\right) \\ & \leq \sum_{l=1}^{|\mathcal{S}_\varepsilon|} \mathbb{P}\left(\sup_{\mathbf{W} \in \mathbb{B}(\mathbf{W}_l, \varepsilon)} \frac{1}{np} \|(\mathbf{W}\bar{\mathbf{X}})^{\circ 3} \bar{\mathbf{X}}^* - \mathbb{E}[(\mathbf{W}\mathbf{X})^{\circ 3} \mathbf{X}^*]\|_F > \delta\right) \quad (\text{By } \varepsilon\text{-covering in (223)}) \\ & \leq \sum_{l=1}^{|\mathcal{S}_\varepsilon|} n^2 \mathbb{P}\left(\left|[(\mathbf{W}_l \bar{\mathbf{X}})^{\circ 3} \bar{\mathbf{X}}^*]_{i,j'} - \mathbb{E}[(\mathbf{W}_l \bar{\mathbf{X}})^{\circ 3} \bar{\mathbf{X}}^*]_{i,j'}\right| \geq p \cdot \frac{\delta}{4}\right) \quad (\text{By (240)}) \\ & \leq \sum_{l=1}^{|\mathcal{S}_\varepsilon|} \left[2n^2 \exp\left(-\frac{3p\delta^2}{c_1\theta + 8n^{\frac{3}{2}}B^4\delta}\right)\right] \quad (\text{By (239)}) \\ & \leq \left(\frac{6}{\varepsilon}\right) n^2 \left[2n^2 \exp\left(-\frac{3p\delta^2}{c_1\theta + 8n^{\frac{3}{2}}B^4\delta}\right)\right] \quad (\text{By (224)}) \\ & = \exp\left(n^2 \ln\left(\frac{48npB^4}{\delta}\right)\right) \left[2n^2 \exp\left(-\frac{3p\delta^2}{c_1\theta + 8n^{\frac{3}{2}}B^4\delta}\right)\right] \quad (\text{Let } \varepsilon = \frac{\delta}{8npB^4}) \\ & = 2n^2 \exp\left(-\frac{3p\delta^2}{c_1\theta + 8n^{\frac{3}{2}}B^4\delta} + n^2 \ln\left(\frac{48npB^4}{\delta}\right)\right), \end{aligned} \quad (242)$$

for a constant $c_1 > 1.4 \times 10^4$. Note that (231) requires a lower bound on B , here we can choose $B = \ln p$, which satisfies (231) when p is large enough (say $p = \Omega(n)$). Combine (220) and substitute $B = \ln p$, we have

$$\begin{aligned} & \mathbb{P}\left(\sup_{\mathbf{W} \in \mathcal{O}(n;\mathbb{R})} \frac{1}{np} \|(\mathbf{W}\mathbf{X})^{\circ 3} \mathbf{X}^* - \mathbb{E}[(\mathbf{W}\mathbf{X})^{\circ 3} \mathbf{X}^*]\|_F > \delta\right) \\ & \leq \mathbb{P}\left(\sup_{\mathbf{W} \in \mathcal{O}(n;\mathbb{R})} \frac{1}{np} \|(\mathbf{W}\bar{\mathbf{X}})^{\circ 3} \bar{\mathbf{X}}^* - \mathbb{E}[(\mathbf{W}\mathbf{X})^{\circ 3} \mathbf{X}^*]\|_F > \delta\right) + 2np\theta e^{-\frac{B^2}{2}} \\ & \leq 2n^2 \exp\left(-\frac{3p\delta^2}{c_1\theta + 8n^{\frac{3}{2}}(\ln p)^4\delta} + n^2 \ln\left(\frac{48np(\ln p)^4}{\delta}\right)\right) + 2np\theta \exp\left(-\frac{(\ln p)^2}{2}\right), \end{aligned} \quad (243)$$

for a constant $c_1 > 1.7 \times 10^4$, which completes the proof. $\blacksquare$

D.9. Perturbation Bound for Unitary Polar Factor

Lemma 45 (Perturbation Bound for Unitary Polar Factor) *Let $\mathbf{A}_1, \mathbf{A}_2 \in \mathbb{R}^{n \times n}$ be two nonsingular matrices, and let $\mathbf{Q}_1 = \mathbf{U}_1 \mathbf{V}_1^*$, $\mathbf{Q}_2 = \mathbf{U}_2 \mathbf{V}_2^*$, where*

$$\mathbf{U}_1 \boldsymbol{\Sigma}_1 \mathbf{V}_1^* = \text{SVD}(\mathbf{A}_1), \quad \mathbf{U}_2 \boldsymbol{\Sigma}_2 \mathbf{V}_2^* = \text{SVD}(\mathbf{A}_2).$$

Let $\sigma_n(\mathbf{A}_1)$, $\sigma_n(\mathbf{A}_2)$ denote the smallest singular value of $\mathbf{A}_1, \mathbf{A}_2$ respectively. Then, for any unitary invariant norm $\|\cdot\|_\diamond$ we have

$$\|\mathbf{Q}_1 - \mathbf{Q}_2\|_\diamond \leq \frac{2}{\sigma_n(\mathbf{A}_1) + \sigma_n(\mathbf{A}_2)} \|\mathbf{A}_1 - \mathbf{A}_2\|_\diamond. \quad (244)$$

Proof See theorem 1 in Li (1995). $\blacksquare$

D.10. Perturbation Bound for Singular Values

Lemma 46 (Singular Value Perturbation) *Let $\mathbf{G} = \mathbf{B}\mathbf{M}$ be a general full rank matrix, where $\mathbf{M}$ ($M_{i,i}$ equals to the ℓ^2 -norm of the i^{th} column of $\mathbf{G}$) is a chosen diagonal matrix so $\mathbf{B}$ has unit matrix 2 norm ($\|\mathbf{B}\|_2 = 1$). Let $\delta\mathbf{G} = \delta\mathbf{B}\mathbf{M}$ be a perturbation of $\mathbf{G}$ such that $\|\delta\mathbf{B}\|_2 \leq \sigma_{\min}(\mathbf{B})$ and σ_i and σ'_i be the i^{th} singular value of $\mathbf{G}$ and $\mathbf{G} + \delta\mathbf{G}$ respectively. Then*

$$\frac{|\sigma_i - \sigma'_i|}{\sigma_i} \leq \frac{\|\delta\mathbf{B}\|_2}{\sigma_{\min}(\mathbf{B})}. \quad (245)$$

Proof See theorem 2.17 in Demmel and Veselić (1992). $\blacksquare$

References

- P.-A. Absil and J. Malick. Projection-like retractions on matrix manifolds. *SIAM Journal on Optimization*, 22(1):135–158, 2012.
- P.-A. Absil, R. Mahony, and R. Sepulchre. *Optimization algorithms on matrix manifolds*. Princeton University Press, 2009.
- A. Agarwal, A. Anandkumar, and P. Netrapalli. Exact recovery of sparsely used overcomplete dictionaries. *stat*, 1050:8–39, 2013.
- A. Agarwal, A. Anandkumar, P. Jain, P. Netrapalli, and R. Tandon. Learning sparsely used overcomplete dictionaries. In *Conference on Learning Theory*, pages 123–137, 2014.
- M. Aharon, M. Elad, A. Bruckstein, et al. K-svd: An algorithm for designing overcomplete dictionaries for sparse representation. *IEEE Transactions on signal processing*, 54(11):4311, 2006.

- Z. Allen-Zhu, Y. Li, and Z. Song. A convergence theory for deep learning via over-parameterization. *arXiv preprint arXiv:1811.03962*, 2018.
- A. Argyriou, T. Evgeniou, and M. Pontil. Convex multi-task feature learning. *Machine Learning*, 73(3):243–272, 2008.
- S. Arora, R. Ge, and A. Moitra. New algorithms for learning incoherent and overcomplete dictionaries. In *Conference on Learning Theory*, pages 779–806, 2014.
- S. Arora, R. Ge, T. Ma, and A. Moitra. Simple, efficient, and neural algorithms for sparse coding. 2015.
- Y. Bai, Q. Jiang, and J. Sun. Subgradient descent learns orthogonal dictionaries. *arXiv preprint arXiv:1810.10702*, 2018.
- B. Barak, J. Kelner, and D. Steurer. Dictionary learning and tensor decomposition via the sum-of-squares method. In *STOC*, 2015.
- D. P. Bertsekas. Nonlinear programming. *Journal of the Operational Research Society*, 48(3):334–334, 1997.
- E. Candes and T. Tao. Decoding by linear programming. *arXiv preprint math/0502327*, 2005.
- E. J. Candès. Mathematics of sparsity (and a few other things). Citeseer, 2014.
- Y. Chi, Y. M. Lu, and Y. Chen. Nonconvex optimization meets low-rank matrix factorization: An overview. *arXiv preprint arXiv:1809.09573*, 2018.
- D. Davis, D. Drusvyatskiy, S. Kakade, and J. D. Lee. Stochastic subgradient method converges on tame functions. *arXiv preprint arXiv:1804.07795*, 2018.
- J. Demmel and K. Veselić. Jacobi’s method is more accurate than qr. *SIAM Journal on Matrix Analysis and Applications*, 13(4):1204–1245, 1992.
- D. L. Donoho. For most large underdetermined systems of linear equations the minimal ℓ^1 -norm solution is also the sparsest solution. *Communications on Pure and Applied Mathematics: A Journal Issued by the Courant Institute of Mathematical Sciences*, 59(6):797–829, 2006.
- S. S. Du, J. D. Lee, H. Li, L. Wang, and X. Zhai. Gradient descent finds global minima of deep neural networks. *arXiv preprint arXiv:1811.03804*, 2018.
- A. Edelman, T. A. Arias, and S. T. Smith. The geometry of algorithms with orthogonality constraints. *SIAM journal on Matrix Analysis and Applications*, 20(2):303–353, 1998.
- M. Elad and M. Aharon. Image denoising via sparse and redundant representations over learned dictionaries. *IEEE Transactions on Image processing*, 15(12):3736–3745, 2006.
- R. Ge, F. Huang, C. Jin, and Y. Yuan. Escaping from saddle points—online stochastic gradient for tensor decomposition. In *Conference on Learning Theory*, pages 797–842, 2015.

- Q. Geng and J. Wright. On the local correctness of ℓ_1 -minimization for dictionary learning. In *2014 IEEE International Symposium on Information Theory*, pages 3180–3184. IEEE, 2014.
- D. Gilboa, S. Buchanan, and J. Wright. Efficient dictionary learning with gradient descent. *arXiv preprint arXiv:1809.10313*, 2018.
- B. C. Hall. *Lie Groups, Lie Algebras, and Representations: An Elementary Introduction*. Springer, 2 edition, 2015.
- A. Hyvärinen. A family of fixed-point algorithms for independent component analysis. In *IEEE Int. Conf. on Acoustics, Speech and Signal Processing (ICASSP)*, pages 3917–3920, 1997.
- A. Hyvärinen and E. Oja. A fast fixed-point algorithm for independent component analysis. *Neural Computation*, 9:1483–1492, 1997.
- R. Jenatton, J. Mairal, G. Obozinski, and F. R. Bach. Proximal methods for sparse hierarchical dictionary learning. In *ICML*, 2010.
- M. Journée, Y. Nesterov, P. Richtárik, and R. Sepulchre. Generalized power method for sparse principal component analysis. *Journal of Machine Learning Research*, 11(Feb): 517–553, 2010.
- Y. Katznelson. *An introduction to harmonic analysis*. Cambridge University Press, 2004.
- E. Kreyszig. *Introductory functional analysis with applications*, volume 1. wiley New York, 1978.
- A. Krizhevsky, V. Nair, and G. Hinton. The cifar-10 dataset. *online: [http://www. cs. toronto. edu/kriz/cifar. html](http://www.cs.toronto.edu/kriz/cifar.html)*, 55, 2014.
- H.-W. Kuo, Y. Lau, Y. Zhang, and J. Wright. Geometry and symmetry in short-and-sparse deconvolution. *arXiv preprint arXiv:1901.00256*, 2019.
- Y. LeCun, L. Bottou, Y. Bengio, P. Haffner, et al. Gradient-based learning applied to document recognition. *Proceedings of the IEEE*, 86(11):2278–2324, 1998.
- R.-C. Li. New perturbation bounds for the unitary polar factor. *SIAM Journal on Matrix Analysis and Applications*, 16(1):327–332, 1995.
- Y. Li and Y. Bresler. Global geometry of multichannel sparse blind deconvolution on the sphere. In *Advances in Neural Information Processing Systems*, pages 1132–1143, 2018.
- Y. Li and Y. Liang. Learning overparameterized neural networks via stochastic gradient descent on structured data. In *Advances in Neural Information Processing Systems*, pages 8157–8166, 2018.
- J.-H. Lu. A note on a theorem of Stanton-Weinstein on the L_4 -norm of spherical harmonics. In *Mathematical Proceedings of the Cambridge Philosophical Society*, volume 102, page 561, 1987.

- C. Ma, K. Wang, Y. Chi, and Y. Chen. Implicit regularization in nonconvex statistical estimation: Gradient descent converges linearly for phase retrieval, matrix completion and blind deconvolution. *arXiv preprint arXiv:1711.10467*, 2017.
- T. Ma, J. Shi, and D. Steurer. Polynomial-time tensor decompositions with sum-of-squares. In *2016 IEEE 57th Annual Symposium on Foundations of Computer Science (FOCS)*, pages 438–446. IEEE, 2016.
- J. Mairal, F. Bach, J. Ponce, G. Sapiro, and A. Zisserman. Discriminative learned dictionaries for local image analysis. Technical report, MINNESOTA UNIV MINNEAPOLIS INST FOR MATHEMATICS AND ITS APPLICATIONS, 2008.
- J. Mairal, J. Ponce, G. Sapiro, A. Zisserman, and F. R. Bach. Supervised dictionary learning. In *Advances in neural information processing systems*, pages 1033–1040, 2009.
- J. Mairal, F. Bach, and J. Ponce. Task-driven dictionary learning. *IEEE transactions on pattern analysis and machine intelligence*, 34(4):791–804, 2012.
- J. Mairal, F. Bach, and J. Ponce. Sparse modeling for image and vision processing. *Foundations and Trends in Computer Graphics and Vision*, 8(2-3):85–283, 2014. ISSN 1572-2740. doi: 10.1561/06000000058. URL <http://dx.doi.org/10.1561/06000000058>.
- K. G. Murty and S. N. Kabadi. Some np-complete problems in quadratic and nonlinear programming. *Mathematical programming*, 39(2):117–129, 1987.
- B. K. Natarajan. Sparse approximate solutions to linear systems. *SIAM journal on computing*, 24(2):227–234, 1995.
- B. A. Olshausen and D. J. Field. Emergence of simple-cell receptive field properties by learning a sparse code for natural images. *Nature*, 381(6583):607–609, 1996.
- B. A. Olshausen and D. J. Field. Sparse coding with an overcomplete basis set: A strategy employed by v1? *Vision research*, 37(23):3311–3325, 1997.
- A. V. Oppenheim. *Discrete-time signal processing*. Pearson Education India, 1999.
- Q. Qu, J. Sun, and J. Wright. Finding a sparse vector in a subspace: Linear sparsity using alternating directions. In *Advances in Neural Information Processing Systems*, pages 3401–3409, 2014.
- Q. Qu, X. Li, and Z. Zhu. A nonconvex approach for exact and efficient multichannel sparse blind deconvolution. *arXiv preprint arXiv:1908.10776*, 2019.
- M. Ranzato, Y. Boureau, and Y. LeCun. Sparse feature learning for deep belief networks. In *Advances in Neural Information Processing Systems (NIPS 2007)*, volume 20, 2007.
- S. Ravishanker and Y. Bresler. l0 sparsifying transform learning with efficient optimal updates and convergence guarantees. 2015.
- B. Recht, M. Fazel, and P. A. Parrilo. Guaranteed minimum-rank solutions of linear matrix equations via nuclear norm minimization. *SIAM review*, 52(3):471–501, 2010.

- R. Rubinstein, A. M. Bruckstein, and M. Elad. Dictionaries for sparse representation modeling. *Proceedings of the IEEE*, 98(6):1045–1057, 2010.
- T. Schramm and D. Steurer. Fast and robust tensor decomposition with applications to dictionary learning. *arXiv preprint arXiv:1706.08672*, 2017.
- Y. Shen, Y. Xue, J. Zhang, K. B. Letaief, and V. Lau. Complete dictionary learning via ℓ_p -norm maximization. *arXiv preprint arXiv:2002.10043*, 2020.
- D. A. Spielman, H. Wang, and J. Wright. Exact recovery of sparsely-used dictionaries. In *Conference on Learning Theory*, pages 37–1, 2012.
- R. J. Stanton and A. Weinstein. On the L_4 norm of spherical harmonics. In *Mathematical Proceedings of the Cambridge Philosophical Society*, volume 89, page 343, 1981.
- J. Sun, Q. Qu, and J. Wright. Complete dictionary recovery over the sphere. *arXiv preprint arXiv:1504.06785*, 2015.
- M. Vetterli and J. Kovacevic. Wavelets and subband coding. Technical report, Prentice-Hall, 1995.
- M. Vetterli, J. Kovačević, and V. K. Goyal. *Foundations of signal processing*. Cambridge University Press, 2014.
- M. J. Wainwright. *High-dimensional statistics: A non-asymptotic viewpoint*, volume 48. Cambridge University Press, 2019.
- Y. Xue, Y. Shen, V. Lau, J. Zhang, and K. B. Letaief. Blind data detection in massive mimo via ℓ_3 -norm maximization over the stiefel manifold. *arXiv preprint arXiv:2004.12301*, 2020.
- J. Yang, J. Wright, T. S. Huang, and Y. Ma. Image super-resolution via sparse representation. *IEEE transactions on image processing*, 19(11):2861–2873, 2010.
- Y. Zhai, H. Mehta, Z. Zhou, and Y. Ma. Understanding l4-based dictionary learning: Interpretation, stability, and robustness. In *International Conference on Learning Representations*, 2019.
- Y. Zhang, H.-W. Kuo, and J. Wright. Structured local optima in sparse blind deconvolution. *arXiv preprint arXiv:1806.00338*, 2018.