

PRINCIPAL COMPONENT ANALYSIS BY OPTIMIZATION OF SYMMETRIC FUNCTIONS HAS NO SPURIOUS LOCAL OPTIMA*

ARMIN EFTEKHARI[†] AND RAPHAEL A. HAUSER[‡]

Abstract. Principal component analysis (PCA) finds the best linear representation of data and is an indispensable tool in many learning and inference tasks. Classically, principal components of a dataset are interpreted as the directions that preserve most of its “energy,” an interpretation that is theoretically underpinned by the celebrated Eckart–Young–Mirsky theorem. This paper introduces many other ways of performing PCA, with various geometric interpretations, and proves that the corresponding family of nonconvex programs has no spurious local optima, while possessing only strict saddle points. These programs therefore loosely behave like convex problems and can be efficiently solved to global optimality, for example, with certain variants of the stochastic gradient descent. Beyond providing new geometric interpretations and enhancing our theoretical understanding of PCA, our findings might pave the way for entirely new approaches to structured dimensionality reduction, such as sparse PCA and nonnegative matrix factorization. More specifically, we study an unconstrained formulation of PCA using determinant optimization that might provide an elegant alternative to the deflating scheme commonly used in sparse PCA.

Key words. principal component analysis, nonconvex optimization, saddle point property

AMS subject classifications. 90C26, 15A23

DOI. 10.1137/18M1188495

1. Introduction. Let $A \in \mathbb{R}^{m \times n}$ be a data matrix, with rows corresponding to m different data vectors and columns corresponding to n different features. Successful dimensionality reduction is at the heart of classification, regression, and other learning tasks that often suffer from the “curse of dimensionality,” where having a small number of training samples in relation to the data dimension (namely, $m \ll n$) typically leads to overfitting [19].

To reduce the dimension of data from n to $p \leq n$, consider a matrix X with orthonormal columns. Then the rows of $AX \in \mathbb{R}^{m \times p}$ correspond to the data vectors (namely, the rows of A) projected onto the column span of X , which we denote by $\text{range}(X)$. In particular, the new data matrix AX has reduced dimension p , while the number m of projected data vectors is unchanged; see Figure 1. Principal component analysis (PCA) is one of the oldest dimensionality reduction techniques that can be traced back to the work of Pearson [35] and Hotelling [23], motivated by the observation that often data lives near a lower-dimensional subspace of $\mathbb{R}^n$; see Figure 2. PCA identifies this subspace by finding a suitable matrix X that retains in AX as much as possible of the energy of A , and the optimal X is called the *loading matrix*. The columns of the loading matrix also reveal the hidden correlations between different features by identifying groups of variables that occur with jointly positive or jointly negative weights, for example, in gene expression data [2].

*Received by the editors May 21, 2018; accepted for publication (in revised form) November 26, 2019; published electronically February 6, 2020.

<https://doi.org/10.1137/18M1188495>

Funding: The work of the authors was supported by EPSRC grant EP/N510129/1. The work of the first author was also supported by the Alan Turing Institute through Turing Seed Funding grant SF019.

[†]EPFL, Lausanne, NW1 2DB, Switzerland (armin.eftekhari@gmail.com).

[‡]Mathematical Institute, Oxford University, Oxford, Oxfordshire, OX1 3LB, UK (hauser@maths.ox.ac.uk).



FIG. 1. This figure illustrates the simple and powerful concept of linear dimensionality reduction. The left panel shows a data matrix A , with rows corresponding to m different data vectors and columns corresponding to n different features. For a matrix $X \in \mathbb{R}^{p \times r}$, the right panel shows the projected data matrix AX , containing again m data vectors (rows) but with only $p \leq n$ features (columns).

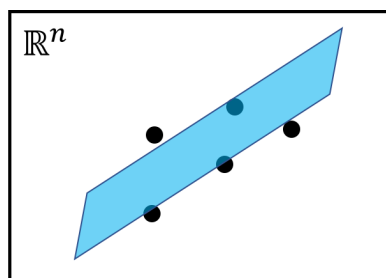


FIG. 2. With each dot corresponding to a data vector with n features, PCA finds a linear subspace (in blue) that best represents the data vectors by capturing most of the energy of the dataset.

PCA is also the building block of other dimensionality reduction techniques such as sparse PCA [28], kernel PCA [36, 37], multidimensional scaling [5], and nonnegative matrix factorization (NMF) [17]. For example, sparse PCA aims to find the important features of data by requiring the loading matrix to be sparse, namely, to have very few nonzero entries. Sparse PCA is useful, for instance, in studying gene expression data, where we are interested in singling out a small number of genes that are responsible for a certain trait or disease [1]. NMF, on the other hand, requires both AX and X to have nonnegative entries, which is valuable in recommender systems, for instance, where the data matrix A containing, say, film ratings is nonnegative and one would expect the same from the projected data matrix AX .

More formally, assume throughout this paper that the data matrix $A \in \mathbb{R}^{m \times n}$ is mean-centered. That is, $\sum_{i=1}^m a_i = 0$, where $a_i \in \mathbb{R}^n$ is the i th row of A , namely, the i th data vector. For $p \leq n$, let $\mathbb{R}_p^{n \times p}$ be the space of full-rank $n \times p$ matrices and consider the *trace inflation function*

$$(1) \quad f_{\text{tr}} : \mathbb{R}_p^{n \times p} \rightarrow \mathbb{R} \quad X \mapsto \frac{\|AX\|_F^2}{\|X\|_F^2} = \frac{\text{tr}(X^* A^* A X)}{\text{tr}(X^* X)}$$

and the program

$$(2) \quad \arg \max \{f_{\text{tr}}(X) : X \in \text{St}(n, p)\}.$$

Above, $\|\cdot\|_F$ and $\text{tr}(\cdot)$ return the Frobenius norm and trace of a matrix, respectively, and A^* is the transpose of matrix A . With $p \leq n$, $\text{St}(n, p)$ above denotes the Stiefel manifold, the set of all $n \times p$ matrices with orthonormal columns. Note that when

$X \in \text{St}(n, p)$, the denominator in the definition of $f_{\text{tr}}(X)$ in (1) is constant, but we include this term to highlight the structural similarity with other inflation functions that we will study later.

It is a consequence of the celebrated Eckart–Young–Mirsky (EYM) theorem that a Stiefel matrix $X \in \text{St}(n, p)$ is a global maximizer of program (2) if and only if it consists of p leading right singular vectors of A , namely, the right singular vectors of A corresponding to its p largest singular values [10, 32]. In other words, program (2) performs PCA on the data matrix A : the loading matrix, namely, the global maximizer of program (2), is a p -leading right singular factor $V_p \in \mathbb{R}^{n \times p}$ of A , and the projected data matrix $AV_p \in \mathbb{R}^{n \times p}$ contains the first p principal components of A .

Note also that program (2) is nonconvex because $\text{St}(n, p) \subset \mathbb{R}^{n \times p}$ is a non-convex set. Even though nonconvex, program (2) behaves like a convex problem in the sense that any local maximizer of program (2) is also a global maximizer. Indeed, it is also a consequence of the EYM theorem that program (2) does not have any spurious local maximizers. Moreover, all saddle points of this program are *strict*, namely, have an ascent direction. Therefore the nonconvex program (2) can be efficiently solved (to global optimality) using (certain variants of) the stochastic gradient descent; see, for instance, [25, 33, 39]. Even more efficiently, solving program (2) or equivalently computing the loading matrix and the principal components of A can be done in $O(\max(m, n)p^2)$ operations using fast algorithms for singular value decomposition (SVD); see, for example, Algorithm 8.6.1 in [18].

Motivation. Our motivation for this work was the following simple observation. The interpretation of PCA as a dimensionality reduction tool suggests that it should suffice to find a matrix $X \in \mathbb{R}_p^{n \times p}$ whose columns span the optimal subspace, which corresponds to p leading right singular vectors of A . That is, one would expect f_{tr} in program (2) to be a function on the Grassmannian $\text{Gr}(n, p)$, the set of all p -dimensional subspaces of $\mathbb{R}^n$. In other words, one would like f_{tr} to be invariant under an arbitrary change of basis in its argument.

That is of course not the case, as a quick inspection of (1) reveals. Generally, we have $f_{\text{tr}}(X\Theta) = f_{\text{tr}}(X)$, only when $\Theta \in \text{Orth}(p)$, namely, when $\Theta \in \mathbb{R}^{p \times p}$ itself is an orthonormal matrix. Program (2) is thus inherently constrained to work with Stiefel matrices, a requirement that is not particularly onerous in the case of PCA but becomes a conceptual nuisance when considering structured dimensionality reduction, such as sparse PCA or NMF. Indeed, enforcing sparsity or nonnegativity in the columns of X in conjunction with orthogonality for the columns of X tends to be very restrictive and is perhaps a questionable objective in the first place.

Contributions. Motivated by the above observation, this paper introduces many other ways of performing PCA, with various geometric interpretations, and proves that the corresponding family of nonconvex programs have no spurious local optima, while possessing only strict saddle points. These new programs therefore loosely behave like convex problems and can be solved to global optimality in polynomial time with, for example, the variants of stochastic gradient ascent in [25, 33]. More specifically, replacing tr in f_{tr} with any elementary symmetric polynomial yields an equivalent formulation for PCA; see the family of problems in (14) and the even larger family of problems in (17).

Program (2) above is indeed a member of this large family. Another notable member of this family is program (6) below, which is effectively unconstrained and consequently does *not* require X to have orthonormal columns. This observation is of particular importance in practice. As we show in section 3, this unconstrained formulation of PCA in program (6) potentially allows for an elegant approach to

structured PCA, in which we wish to impose additional structure on the loading matrix, such as sparsity or nonnegativity.

Let us add that it is known already that program (6) is equivalent to PCA [21]; see [34] for an application to optimal design and [24] for an example in the context of independent component analysis. Of course, this equivalence does not guarantee that program (6), like program (2), can also be solved in polynomial time. In this sense, our contribution is that the nonconvex program (6) has no spurious local optima, has only strict saddle points, and can therefore be solved efficiently by certain variants of the stochastic gradient descent. Moreover, the introduction of the rest of this large family of equivalent formulations of PCA and their analysis in this work is the other novel aspect of this work.

Organization. The rest of this paper is organized as follows. To present this work in an increasing order of complexity, we first introduce in section 2 the unconstrained formulation of PCA, namely, program (6), and discuss in section 3 its potential application in structured dimensionality reduction. In section 4, we then present programs (14) and (17), a large family of equivalent formulations of PCA, of which both programs (2) and (6) are members. The claim that all these programs are indeed equivalent to PCA and can be efficiently solved is proven in sections 5 and 6 and the appendices.

2. PCA by determinant optimization. In analogy to f_{tr} in (1), let us define the *volume inflation function* by

$$(3) \quad \begin{aligned} f_{\det} : \mathbb{R}_p^{n \times p} &\rightarrow \mathbb{R}, \\ X &\mapsto \frac{\det(X^* A^* A X)}{\det(X^* X)}, \end{aligned}$$

where $\det$ stands for determinant and, in analogy to program (2), consider the program

$$(4) \quad \arg \max \{f_{\det}(X) : X \in \text{St}(n, p)\}.$$

Observe that programs (2) and (4) coincide for $p = 1$, namely, when we seek the leading principal component of the matrix A , in which case $X^* A^* A X$ and $X^* X$ are both positive scalars. Unlike f_{tr} , note that $f_{\det}$ is invariant under an arbitrary change of basis. Indeed, for arbitrary $X \in \mathbb{R}_p^{n \times p}$ and $\Theta \in \text{GL}(p)$, we have that

$$(5) \quad \begin{aligned} f_{\det}(X\Theta) &= \frac{\det(\Theta)^2 \det(X^* A^* A X)}{\det(\Theta)^2 \det(X^* X)} \\ &= f_{\det}(X), \end{aligned}$$

where $\text{GL}(p)$ is the general linear group, the set of all invertible $p \times p$ matrices. That is, $f_{\det}$ is naturally defined on the Grassmannian $\text{Gr}(n, p)$ and consequently Program (4) is equivalent to the program

$$(6) \quad \arg \max \{f_{\det}(X) : X \in \mathbb{R}_p^{n \times p}\}.$$

Because $f_{\det}$ is invariant under any change of basis by (5), program (6) inherently constitutes an optimization over the Grassmannian $\text{Gr}(n, p)$. Moreover, it is important that program (6) is effectively unconstrained because $\mathbb{R}_p^{n \times p}$ is an open subset of $\mathbb{R}^{n \times p}$ with nonempty interior. To summarize, the drawback of program (2) in section 1 which served as the motivation of this work is overcome by program (6), because it

is an unconstrained optimization program that involves an objective function defined naturally on the Grassmannian.

A key observation of this paper is that program (6) appears to be a good model for dimensionality reduction. Indeed, note that $X^*A^*AX \in \mathbb{R}^{p \times p}$ is the sample covariance matrix of the projected data AX . Consider the normal distribution $\mathcal{N}(0, X^*A^*AX)$ with zero mean and covariance matrix X^*A^*AX , which has ellipsoidal level sets of the form

$$(7) \quad \{z \in \mathbb{R}^p : z^*X^*A^*AXz = c\}$$

for arbitrary $c \geq 0$. Let B_c be the bounding box of this level set and note that the volume of B_c is $c^p \sqrt{\det(X^*A^*AX)}$. We can therefore interpret program (6) as maximizing the volume of this bounding box. That is, program (6), loosely speaking, finds the directions that maximize the volume of the projected dataset.

In contrast, program (2) maximizes the energy of the projected data. That is, program (2) maximizes the diameter of the above bounding box, namely, $c\sqrt{\text{tr}(X^*A^*AX)}$, rather than its volume; see Figure 3. It is perhaps peculiar that $\text{tr}(X^*A^*AX)$ is commonly referred to as the “total variance” of the dataset, for this quantity does not play any role in the normalizing constant of the normal distribution $\mathcal{N}(0, X^*A^*AX)$, whereas $\det(X^*A^*AX)$ does, in direct generalization of the role the variance plays in the one-dimensional case. At any rate, we see that programs (2) and (6) are both sensible approaches for linear dimensionality reduction but that their geometric justifications are very different. Somewhat surprisingly, we find that program (6) also

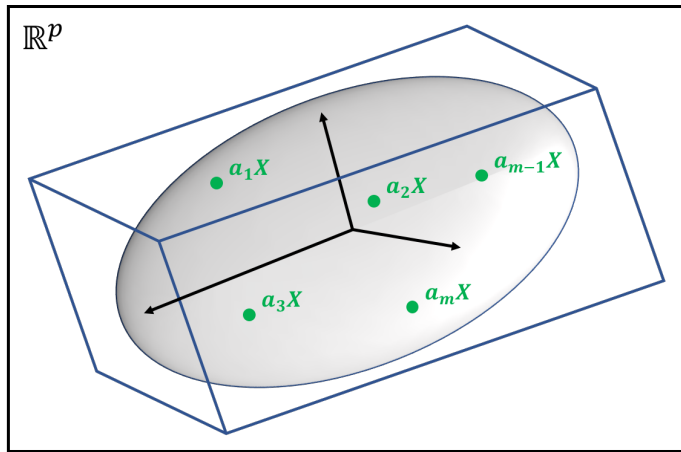


FIG. 3. This figure illustrates the geometric intuition underlying this paper. Suppose that $a_1, \dots, a_m \in \mathbb{R}^n$ are the rows of the data matrix $A \in \mathbb{R}^{m \times n}$, each representing a data vector. Then $a_1X, \dots, a_mX \in \mathbb{R}^p$ are the projected data vectors, with reduced dimension of $p \leq n$. It is easy to see that the sample covariance matrix of these projected data vectors is $X^*A^*AX \in \mathbb{R}^{p \times p}$, with ellipsoidal level sets, one of which and its bounding box are displayed above. Then program (2) maximizes the diameter of this bounding box, which is proportional to $\sqrt{\text{tr}(X^*A^*AX)}$. In contrast, program (6) maximizes the volume of this box, which is proportional to $\sqrt{\det(X^*A^*AX)}$. Remarkably, both programs (2) and (6) perform PCA of data matrix A ; see sections 1 and 2. As discussed in section 3, program (6) is of particular importance in practice as it gives an elegant solution to the problem of structured linear dimensionality reduction. More generally, we also show that maximizing the sum of the volume squares of all q -dimensional facets of this bounding box is equivalent to PCA of matrix A for any $1 \leq q \leq p$; see section 4. In particular, programs (2) and (6) are special cases with $q = 1$ and $q = p$, respectively.

performs PCA of A and has no spurious local optima, exactly like program (2). The next result is proven in section 5.

THEOREM 2.1 (determinant). *The following statements hold true:*

- (i) $\tilde{X} \in \mathbb{R}^{n \times p}$ is a global maximizer of program (6) if and only if there exists a p -leading right singular factor $V_p \in \mathbb{R}^{n \times p}$ of A such that $\text{range}(\tilde{X}) = \text{range}(V_p)$.
- (ii) Program (6) does not have any spurious local optima, namely, any local maximum or minimum of program (6) is also a global maximum or minimum, respectively, and all other stationary points are strict saddle points. Moreover, if $\sigma_p(A) > \sigma_{p+1}(A)$, at any such strict saddle point X_s , there exists an ascent direction $\Delta \in \mathbb{R}^{n \times p}$ such that

$$(8) \quad \nabla^2 f_{\det}(X_s)[\Delta, \Delta] \geq f_{\det}(X_s) \left(\frac{\sigma_p^2(A)}{\sigma_{p+1}^2(A)} - 1 \right) \|\Delta\|_F^2,$$

where the bilinear operator $\nabla^2 f_{\det}(X_s)$ is the Hessian of $f_{\det}$ at X_s . Above, $\sigma_p(A)$ is the p th largest singular value of A .

In words, part (i) of Theorem 2.1 states that program (6) performs PCA on the data matrix A , and therefore programs (2) and (6) are equivalent in this sense. Note that program (6) provides a different geometric interpretation of PCA based on maximizing the “volume” of projected data rather than its “diameter,” which was the case in program (2). Even though we present a new proof for the characterization of the global maximizers of program (6) in part (i) of Theorem 2.1, this result can also be proved using interlacing properties of singular values (see Corollary 3.2 in [22]) or via the Cauchy–Binet formula [29].

The main contribution of Theorem 2.1 is its part (ii) about the global landscape of the objective function $f_{\det}$, stating that the nonconvex program (6) behaves like a convex problem in the sense that any local maximizer (minimizer) of program (6) is also a global maximizer (minimizer). Moreover, saddle points of program (6) are strict. In this way too, programs (2) and (6) are similar; see section 1. Note that part (ii) of Theorem 2.1 is crucial in the design of new dimensionality reduction algorithms: The instability of all stationary points except the global optima and the strictness of all saddle points establish that, for example, stochastic gradient ascent converges to the correct solution in polynomial time [25, 33]. That is, the nonconvex program (6) can be efficiently solved to global optimality. However, as discussed in section 1, computationally efficient algorithms for PCA are already available and application of, say, stochastic gradient ascent to program (6) is not intended to replace those algorithms. Instead, as discussed in section 3, the unconstrained program (6) potentially opens up a radically new approach to structured PCA.

We remark that Theorem 2.1 is in line with a recent trend in computational sciences to understand the geometry and performance of nonconvex programs and algorithms [31, 39, 12, 7, 6, 3, 4, 14, 26, 38, 15, 13, 16]. While the available results do not apply to our problem, the underlying phenomena are closely related. Perhaps the closest result to our work is [14], stating that the (nonconvex) matrix completion program has no spurious local optima when given access to randomly observed matrix entries. This result in a sense extends the EYM theorem [10, 32] to partially observed matrices.

From a computational perspective, we may consider the program

$$(9) \quad \arg \max \{ \log(f_{\det}(X)) : X \in \mathbb{R}_p^{n \times p} \},$$

which is equivalent to program (6) but has better numerical stability. As a numerical example, we generated generic $U, V \in \text{Orth}(100)$ and random matrix $A \in \mathbb{R}^{100 \times 100}$ with SVD $A = U\Sigma V^*$. The singular values of A , namely, the entries of the diagonal matrix $\Sigma \in \mathbb{R}^{100 \times 100}$, were selected according to the power law. To be specific, we took $\sigma_i = i^{-1}$ to generate Figure 4 and $\sigma_i = i^{-2}$ to generate Figure 5, for every $i \in [100] = \{1, 2, \dots, 100\}$. For $p = 5$, we let $V_p \in \mathbb{R}^{n \times p}$ denote the first p columns of V and, by Theorem 2.1, the unique maximizer of programs (6), and (9). (Note that V_p is also the unique maximizer of program (2) by the EYM theorem.) In order to find V_p , we then applied gradient ascent to program (9) with fixed step size of $\rho = 5$ and random initialization, producing a sequence of estimates $\{X_l\}_l \subset \mathbb{R}^{n \times p}$. We also recorded the error $\|X_l X_l^\dagger - V_p V_p^*\|$ in the l th iteration, namely, the sine of the principal angle between $\text{range}(X_l)$ and $\text{range}(V_p)$, which is plotted in Figures 4 and 5. As predicted by Theorem 2.1, the error vanishes in both examples as the algorithm progresses. We also refer the interested reader to [20] for a comparison between programs (2) and (9), as well as LAPACK's implementation of Lanczos' method [18] for performing PCA. While the numerical results presented in [20] are encouraging, a more comprehensive study is required to investigate the competitiveness of program (9) for PCA, as an alternative to more mainstream approaches [18].

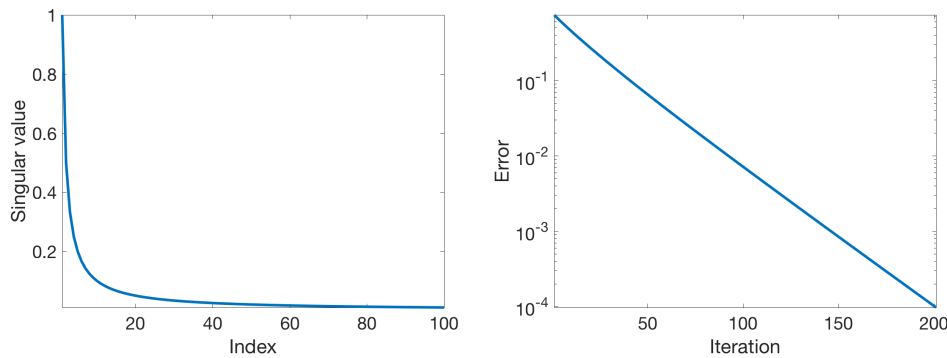


FIG. 4. The left panel shows the spectrum $\{\sigma_i\}_{i=1}^{100}$ of a randomly generated matrix $A \in \mathbb{R}^{100 \times 100}$ with $\sigma_i = i^{-1}$, and the right panel shows the progression of the gradient ascent algorithm with fixed step size, applied to program (9); see section 2 for details.

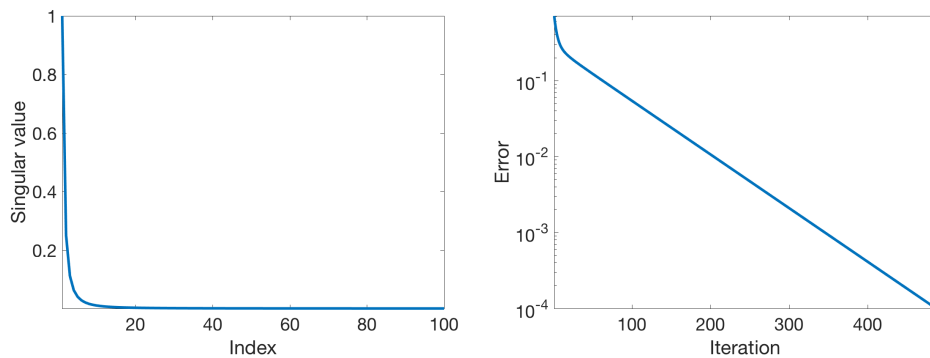


FIG. 5. The left panel shows the spectrum $\{\sigma_i\}_{i=1}^{100}$ of a randomly generated matrix $A \in \mathbb{R}^{100 \times 100}$ with $\sigma_i = i^{-2}$, and the right panel shows the progression of the gradient ascent algorithm with fixed step size, applied to program (9); see section 2 for details.

It might also be helpful to highlight the following practical consideration. Let $\tilde{X} \in \mathbb{R}^{n \times p}$ denote a maximizer of program (6) or (9). Given $\tilde{X}$, a few extra steps are required to compute the complete SVD of A , which we now list. Let $\hat{X} \in \mathbb{R}^{n \times p}$ be an orthonormal basis for $\tilde{X}$, which can be computed in $O(np^2)$ operations by SVD. Then computing the SVD of $A\hat{X} = \hat{U}\hat{\Sigma}\hat{V}^*$ can be performed in merely $O(mp^2)$ operations and yields the diagonal coefficients of $\hat{\Sigma}$ as the p leading singular values of A , as well as $\hat{U}$ and $\hat{X}\hat{V}$ as the corresponding p leading left and right singular vectors of A , respectively.

3. Structured PCA. As we will see in this section, the unconstrained formulation of PCA in program (6) might be of particular interest in practice, in contrast to program (2), which is restricted to the Stiefel manifold. Indeed, the determinant formulation of PCA in program (6) might allow for a more elegant approach to structured PCA, in which we wish to impose additional structure on the loading matrix, such as sparsity or nonnegativity.

For the purposes of this brief and informal discussion, let us focus on sparse PCA, the problem of finding a small number of features that best describe the data matrix $A \in \mathbb{R}^{m \times n}$. As one application, when working with gene expression data, we are interested in a small number of features (genes) which are responsible for certain traits or diseases [27, 9]. The “dual” of sparse PCA can also be interpreted as data clustering.

Loosely speaking, sparse PCA is the problem of finding a *sparse*¹ matrix $X \in \mathbb{R}^{n \times p}$ that retains, in the projected data $AX \in \mathbb{R}^{m \times p}$, as much as possible of the energy of A . More formally, sparse PCA might be formulated as a natural generalization of program (2), namely,

$$(10) \quad \arg \max \{f_{\text{tr}}(X) : X \in \text{St}(n, p) \text{ and } \|X\|_0 \leq k\},$$

where $\|X\|_0$ is the number of nonzero entries of X , and the typically small integer k is the *sparsity level*. Note that program (10) forces X to have orthonormal columns and few nonzero entries, which tends to be restrictive and is also a somewhat questionable objective in the first place. With a few exceptions, particularly [28], this problem is often addressed by *deflating* A , namely, finding the sparse principal components sequentially, that is, one by one. Indeed, note that the Stiefel constraint from program (10) is redundant when $p = 1$, namely, when X is a column vector. One could therefore find the leading sparse principal component of A , say, $\tilde{x}_1 \in \mathbb{R}^n$, by solving program (10) with $p = 1$, remove its contribution from A by forming $A_1 = A - A\tilde{x}_1\tilde{x}_1^*$, and then solve program (10) with A_1 in place of A to find the second sparse principal component $\tilde{x}_2 \in \mathbb{R}^n$, and so on [8]. However, deflating A is believed to be inherently problematic when the problem is ill-posed [28].

The determinant formulation of PCA in program (9) might provide an elegant alternative to program (10). Recall that the feasible set of program (9) is an open subset of $\mathbb{R}^{n \times p}$ with nonempty interior, and thus program (9) is effectively unconstrained. We can therefore formulate sparse PCA by imposing a sparsity constraint on program (9), namely,

$$(11) \quad \arg \max \{\log(f_{\text{det}}(X)) : X \in \mathbb{R}_p^{n \times p} \text{ and } \|X\|_0 \leq k\},$$

which requires X to be full-rank and sparse, relaxing the far more restrictive requirement of being Stiefel and sparse in program (2). Note that removing the full-rank

¹A sparse matrix has a small number of nonzero entries.

requirement in program (11) is impossible as that would mean A has fewer than p principal components and therefore the problem is ill-defined.

Similar ideas might be applied to NMF, in which X and AX are both required to be nonnegative. More generally, the unconstrained nature of program (9) might provide an entirely new approach to many structured dimensionality reduction problems, a research direction that remains to be explored in the future. In particular, what is the global geometry of program (11)? What is the precise relationship between programs (10) and (11)?

4. Generalization to positive symmetric polynomials. So far, we have seen that maximizing the trace objective function in program (2) and maximizing the determinant objective function in program (6) are equivalent, and both provide the leading principal components of the data matrix A . Moreover, both nonconvex programs can be solved to global optimality efficiently; see the discussion after Theorem 2.1, for example. Indeed, these claims for program (2) follow from the EYM theorem [10, 32] and the claims for program (6) follow from Theorem 2.1; see sections 1 and 2.

Note that both $\text{tr}(X^*A^*AX)$ and $\det(X^*A^*AX)$ are *elementary symmetric polynomials*, namely, both are coefficients of the *characteristic polynomial* of X^*A^*AX . More specifically, let $X^*A^*AX = W_X \text{diag}(\lambda_X)W_X^*$ be the eigendecomposition of X^*A^*AX , where $W_X \in \text{Orth}(p)$ is an orthonormal matrix and the vector $\lambda_X \in \mathbb{R}^p$ contains the eigenvalues of X^*A^*AX . Here, $\text{diag}(\lambda_X) \in \mathbb{R}^{p \times p}$ is the diagonal matrix formed by the vector λ_X . Then the characteristic polynomial associated with X^*A^*AX takes $t \in \mathbb{R}$ to

$$\begin{aligned}
 \det(I_p + tX^*A^*AX) &= \det(W_X(I_p + t \text{diag}(\lambda_X))W_X^*) \quad (W_X \in \text{Orth}(p)) \\
 &= \det(W_X) \cdot \det(I_p + t \text{diag}(\lambda_X)) \cdot \det(W_X^*) \\
 &= \det \begin{bmatrix} 1+t \cdot \lambda_{X,1} & 0 & \cdots & 0 \\ 0 & 1+t \cdot \lambda_{X,2} & \cdots & 0 \\ \vdots & \vdots & \ddots & \\ 0 & 0 & \cdots & 1+t \cdot \lambda_{X,p} \end{bmatrix} \quad (\det(W_X) = 1) \\
 &= \prod_{i=1}^p (1 + t \cdot \lambda_{X,i}) \\
 (12) \quad &=: \sum_{q=0}^p s_q(X^*A^*AX) \cdot t^q,
 \end{aligned}$$

where the q th elementary symmetric polynomial $s_q : \text{Sym}(p) \rightarrow \mathbb{R}$ is the coefficient of t^q above, namely,

$$\begin{aligned}
 s_q &: \text{Sym}(p) \rightarrow \mathbb{R}, \\
 (13) \quad B &\mapsto \sum_{1 \leq i_1 < \cdots < i_q \leq p} \prod_{j=1}^q \lambda_{B,i_j},
 \end{aligned}$$

with the convention that $s_0(B) = 1$. Above, $\{\lambda_{B,i}\}_{i=1}^p$ are the eigenvalues of $B \in \text{Sym}(p)$, and $\text{Sym}(p)$ denotes the set of symmetric $p \times p$ matrices over the real numbers. We also remark that elementary symmetric polynomials are *spectral* functions in that they only depend on the eigenvalues of the input matrix. As mentioned earlier, $\text{tr}(X^*A^*AX) = s_1(X^*A^*AX)$ and $\det(X^*A^*AX) = s_p(X^*A^*AX)$.

In analogy to trace and determinant objective functions (1), (3), let us define

$$(14) \quad \begin{aligned} f_{s_q} : \mathbb{R}^{n \times p} &\rightarrow \mathbb{R}, \\ X &\mapsto \frac{s_q(X^* A^* A X)}{s_q(X^* X)} \end{aligned}$$

for every $q \in [p] := \{1, \dots, p\}$. In particular, $f_{s_1} = f_{\text{tr}}$ in (1) and $f_{s_p} = f_{\text{det}}$ in (3) are two special cases. Last, in analogy to programs (2) and (4), consider the program

$$(15) \quad \arg \max \{f_{s_q}(X) : X \in \text{St}(n, p)\}.$$

Again note that programs (2) and (4) are special cases of program (15) for $q = 1$ and $q = p$, respectively. Revisiting the geometric interpretation discussed in section 2, we may also verify that $f_{s_q}(X)$ is proportional to the sum of volumes squared of all q -dimensional facets of the bounding box B_c ; see right after (7) and also Figure 3. In this sense, program (15) finds the projected data AX that maximizes this geometric attribute.

Generalizing the EYM theorem for program (2) and Theorem 2.1 for program (4), the following result states that program (15) performs PCA of A and has no spurious local optima for every $q \in [p]$; see section 6 for the proof.

THEOREM 4.1 (elementary symmetric polynomials). *For every $q \in [p]$, the following statements hold true:*

- (i) $\tilde{X} \in \mathbb{R}^{n \times p}$ is a global maximizer of program (15) if and only if there exists a p -leading right singular factor $V_p \in \mathbb{R}^{n \times p}$ of A such that $\text{range}(\tilde{X}) = \text{range}(V_p)$.
- (ii) Program (15) does not have any spurious local optima, namely, any local maximum or minimum of program (15) is also a global maximum, respectively, minimum, and all other stationary points are strict saddle points.

In words, Theorem 4.1 introduces a family of equivalent formulations for PCA, namely, program (4.1) for every $q \in [p]$. This family includes PCA by trace optimization (program (2)) and PCA by determinant optimization (program (6)). In fact, maximizing *any* conic combination of elementary symmetric polynomials also performs PCA. To be specific, for nonnegative (but not all zero) coefficients $\{w_q\}_{q=0}^p$, consider the symmetric function

$$(16) \quad \begin{aligned} g_w : \text{Sym}(p) &\rightarrow \mathbb{R}, \\ B &\mapsto \sum_{q=0}^p w_q \cdot s_q(B), \end{aligned}$$

where the elementary symmetric polynomial s_q was defined in (13). Also define

$$(17) \quad \begin{aligned} f_{g_w} : \mathbb{R}^{n \times p} &\rightarrow \mathbb{R}, \\ X &\mapsto \frac{g_w(X^* A^* A X)}{g_w(X^* X)}, \end{aligned}$$

and consider the program

$$(18) \quad \arg \max \{f_{g_w}(X) : X \in \text{St}(n, p)\}.$$

The following result is an immediate consequence of Theorem 4.1 and states that program (18) for *any* positive symmetric function performs PCA, thus providing a broad class of equivalent formulations for PCA.

COROLLARY 4.2 (positive symmetric functions). *Let $\{w_q\}_{q=0}^p$ be a set of non-negative coefficients with at least one $w_q \neq 0$, and let f_{g_w} be the function defined in (17). Then the following statements hold true:*

- (i) $\tilde{X} \in \mathbb{R}^{n \times p}$ is a global maximizer of Program (18) if and only if there exists a p -leading right singular factor $V_p \in \mathbb{R}^{n \times p}$ of A such that $\text{range}(\tilde{X}) = \text{range}(V_p)$.
- (ii) Program (18) does not have any spurious local maximizers, namely, any local maximizer of program (18) is also a global maximizer, and all other stationary points are strict saddle points.

Proof. Without loss of generality, we can assume that $w_0 = 0$. Indeed, since $s_0 = 1$ is constant by definition, setting $w_0 = 0$ does not change the optima and stationary points of program (18). Note that $X^*X = I_p$ for every X feasible to programs (15) and (18). Therefore program (18) has the same optima and stationary points as

$$(19) \quad \arg \max \{g_w(X^*A^*AX) : X \in \text{St}(n, p)\},$$

which, after recalling (16), has in turn the same optima and stationary points as

$$(20) \quad \arg \max \left\{ \sum_{q=0}^p w_q \cdot \frac{s_q(X^*A^*AX)}{s_q(X^*X)} : X \in \text{St}(n, p) \right\}.$$

Recall Theorem 4.1 about program (15) for every $q \in [p]$. Because the coefficients $\{w_q\}_q$ are nonnegative by assumption, the claims in Theorem 4.1 extend to program (20) and in turn to program (19) and then to program (18). This completes the proof of Corollary 4.2. $\square$

In conclusion, this paper introduces a large family of equivalent interpretations of PCA which can all be solved to global optimality in polynomial time. One member of this family is an unconstrained formulation of PCA that might lead in the future to developing new algorithms and techniques for structured PCA.

5. Proof of Theorem 2.1. We first begin with a change of variables. Let $A = U\Sigma V^*$ be the SVD of A , where $U \in \mathbb{R}^{m \times m}$ and $V \in \mathbb{R}^{n \times n}$ are orthonormal matrices, and the diagonal matrix $\Sigma \in \mathbb{R}^{m \times n}$ is formed by the singular values of A in nonincreasing order, denoted by $\sigma_1 \geq \sigma_2 \geq \dots$. Let us set $\Gamma := \Sigma^*\Sigma \in \mathbb{R}^{n \times n}$ for short and note that

$$(21) \quad \Gamma = \text{diag}(\gamma), \quad \text{where } \gamma = [\sigma_1^2 \ \dots \ \sigma_r^2 \ 0 \ \dots]^* \in \mathbb{R}^n,$$

where $\text{diag}(\gamma)$ shapes the vector γ into a diagonal matrix, and $r = \text{rank}(A)$ is the rank of A , namely, the number of positive singular values of A . Under the change of variables from X to $Y = V^*X$, Program (6) is equivalent to

$$(22) \quad \arg \max \left\{ \frac{\det(Y^*\Gamma Y)}{\det(Y^*Y)} : Y \in \mathbb{R}_p^{n \times p} \right\}.$$

Without loss of generality, we will therefore assume that $A^*A = \Gamma$ in (3), namely, we henceforth set

$$(23) \quad f_{\det}(X) = \frac{\det(X^*\Gamma X)}{\det(X^*X)}.$$

We will prove Theorem 2.1 by studying the stationary points of program (6). This program is unconstrained and therefore a stationary point $X_s \in \mathbb{R}_p^{n \times p}$ of program

(6) is characterized by $\nabla f_{\det}(X_s) = 0$. In light of the invariance in (5), we can also assume without loss of generality that $X_s \in \text{St}(n, p)$. A stationary point $X_s \in \text{St}(n, p)$ of Program (6) thus satisfies

$$(24) \quad \nabla f_{\det}(X_s) = 2\Gamma X_s \nabla \det(X_s^* \Gamma X_s) - 2\det(X_s^* \Gamma X_s) \cdot X_s \nabla \det(X_s^* X_s) = 0.$$

Note that the determinant is a *spectral function*, namely, it only depends on the eigenvalues of its input matrix. More precisely, for a matrix $Z \in \mathbb{R}^{p \times p}$, it holds that

$$(25) \quad \det(Z) = e_p(\lambda_Z) := \prod_{i=1}^p \lambda_{Z,i},$$

where $\lambda_{Z,i}$ is the i th eigenvalue of Z and the vector $\lambda_Z = [\lambda_{Z,1}, \dots, \lambda_{Z,p}]$ contains the eigenvalues of Z . The derivative of a spectral function is well-known. To be specific, let $Z = W_Z \cdot \text{diag}(\lambda_Z) \cdot W_Z^*$ be the eigendecomposition of Z , where $W_Z \in \text{Orth}(p)$ is an orthonormal matrix and, as before, $\text{diag}(\lambda_Z) \in \mathbb{R}^{p \times p}$ is the diagonal matrix formed by the vector λ_Z . For a spectral function $\phi : \text{Sym}(p) \rightarrow \mathbb{R}$, there exists a *symmetric function* $\psi : \mathbb{R}^p \rightarrow \mathbb{R}$ such that

$$(26) \quad \phi(Z) = \phi(\text{diag}(\lambda_Z)) = \psi(\lambda_Z),$$

where we recall that a symmetric function is a function that remains invariant after changing the order of its arguments. We then have from [30] that

$$(27) \quad \nabla \phi(Z) = W_Z \cdot \text{diag}(\nabla \psi(\lambda_Z)) \cdot W_Z^*.$$

In our case, with $\phi = e_p$ and when $Z \in \text{GL}(p)$, namely, when Z is nonsingular, we have that

$$\begin{aligned} \nabla \det(Z) &= W_Z \cdot \text{diag}(\nabla e_p(\lambda_Z)) \cdot W_Z^* && \text{(see (27))} \\ &= W_Z \cdot \text{diag}\left(e_p(\lambda_Z) \begin{bmatrix} \frac{1}{\lambda_{Z,1}} & \cdots & \frac{1}{\lambda_{Z,p}} \end{bmatrix}^*\right) \cdot W_Z^* && \text{(see (25))} \\ (28) \quad &= \det(Z) W_Z \cdot \text{diag}\left(\begin{bmatrix} \frac{1}{\lambda_{Z,1}} & \cdots & \frac{1}{\lambda_{Z,p}} \end{bmatrix}^*\right) \cdot W_Z^* && \text{(see (25)).} \end{aligned}$$

Returning to the proof and recalling that $X_s \in \text{St}(n, p)$, (28) allows us to write that

$$(29) \quad \nabla \det(X_s^* X_s) = I_p.$$

Likewise, after recalling that $X_s^* \Gamma X_s \in \text{GL}(p)$ by assumption, it follows from (28) that

$$(30) \quad \nabla \det(X_s^* \Gamma X_s) = \det(X_s^* \Gamma X_s) \cdot (X_s^* \Gamma X_s)^{-1}.$$

Substituting (29), (30), we find that (24) holds if and only if $\Gamma X_s = X_s X_s^* \Gamma X_s$, and since $X_s X_s^*$ is the projection into $\text{range}(X_s)$, this is true if and only if $\text{range}(\Gamma X_s) \subseteq \text{range}(X_s)$. From this it follows that $X_s \in \text{St}(n, p)$ such that $X_s^* \Gamma X_s \in \text{GL}(p)$ is a stationary point of program (6), if and only if

$$(31) \quad \text{range}(X_s) = \text{range}(\Gamma X_s).$$

The following result, proved in Appendix A, characterizes the stationary points of program (6). That is, the next result characterizes matrices $X_s \in \text{St}(n, p)$ that satisfy (31).

LEMMA 5.1. For a singular value σ_j of A , let n_j denote the multiplicity of σ_j in A . Suppose that $X_s \in \text{St}(n, p)$ satisfies $X_s^* \Gamma X_s \in \text{GL}(p)$ and $\text{range}(X_s) = \text{range}(\Gamma X_s)$. Then there exists an index set $J \subset [n]$ such that

1. $\{\sigma_i\}_{i \in J}$ are distinct and positive, and
2. for every $i \notin J$, the rows of X_s corresponding to σ_i are zero, and
3. for every $i \in J$, the rows of X_s corresponding to σ_i span a $\dim(\sigma_i)$ -dimensional subspace of $\mathbb{R}^p$, and
4. for every distinct pair $\{i, i'\} \subset J$, the corresponding subspaces are orthogonal, and finally
5. $\sum_{i \in J} \dim(\sigma_i) = p$.

Above, for every $i \in J$, we set

$$\dim(\sigma_i) = \max \left[p - \sum_{i' \neq i} n_{i'}, 1 \right].$$

Based on the characterization of stationary points in Lemma 5.1, we next calculate the Hessian of $f_{\det}$ at a stationary point of program (6), which will later help us determine the stability of these stationary points. The following result is in fact more general; see Appendix B for the proof.

LEMMA 5.2. Consider a differentiable spectral function $\phi : \text{Sym}(p) \rightarrow \mathbb{R}$ and the associated symmetric function $\psi : \mathbb{R}^p \rightarrow \mathbb{R}$; see (26). Consider also the function

$$(32) \quad f_\phi : \mathbb{R}_p^{n \times p} \rightarrow \mathbb{R}, \quad X \mapsto \frac{\phi(X^* \Gamma X)}{\phi(X^* X)}.$$

Consider last $X_s \in \text{St}(n, p)$ such that $\text{range}(X_s) = \text{range}(\Gamma X_s)$ and the corresponding index set $J = \{i_1, i_2, \dots\} \subseteq [n]$ in Lemma 5.1. Then we can assume without loss of generality that X_s is block-diagonal. Under this assumption, it holds that

$$(33) \quad X_s^* \Gamma X_s = \text{diag}(\tilde{\gamma}_1),$$

where

$$(34) \quad \tilde{\gamma}_1 := \begin{bmatrix} \overbrace{\sigma_{i_1}^2 \cdots \sigma_{i_1}^2}^{\dim(\sigma_{i_1})} & \overbrace{\sigma_{i_2}^2 \cdots \sigma_{i_2}^2}^{\dim(\sigma_{i_2})} & \cdots \end{bmatrix}^* \in \mathbb{R}^p.$$

Moreover, let $K \subset [n]$ denote the index set corresponding to the nonzero rows of X_s . Then it holds that

$$(35) \quad \nabla^2 f_\phi(X_s)[\Delta, \Delta] = \frac{1}{\phi(I_p)} \sum_{i \in K^C} \sum_{j=1}^p (\sigma_i^2 \partial_j \psi(\tilde{\gamma}_1) - f_\phi(X_s) \partial_1 \psi(1_p)) \Delta_{i,j}^2$$

for every $\Delta \in \mathbb{R}^{n \times p}$ that is zero on the rows indexed by K . Above, the bilinear operator $\nabla^2 f_\phi(X_s) : \mathbb{R}^{n \times p} \times \mathbb{R}^{n \times p} \rightarrow \mathbb{R}$ is the Hessian of f_ϕ at X_s . Also, K^C is the complement of the set K , $\partial_i \psi(\tilde{\gamma}_1)$ is the i th entry of the gradient vector $\nabla \psi(\tilde{\gamma}_1) \in \mathbb{R}^p$, and $1_p \in \mathbb{R}^p$ is the vector of all ones.

In particular, when $\phi = \det$ and, consequently, $\psi = e_p$ in Lemma 5.2, let us simplify the expression for the Hessian in (35) by noting that

$$\phi(I_p) = \det(I_p) = 1,$$

$$\begin{aligned}
\partial_j \psi(\tilde{\gamma}_1) &= \partial_j e_p(\tilde{\gamma}_1) \\
&= \prod_{i \neq j} \tilde{\gamma}_{1,i} \\
&= \frac{\prod_{i=1}^p \tilde{\gamma}_{1,i}}{\tilde{\gamma}_{1,j}} \\
&= \frac{\det(\text{diag}(\tilde{\gamma}_1))}{\tilde{\gamma}_{1,j}} \\
&= \frac{\det(X_s^* \Gamma X_s)}{\tilde{\gamma}_{1,j} \det(X_s^* X_s)} \quad (\Gamma = \text{diag}(\tilde{\gamma}_1) \text{ and } X_s \in \text{St}(n, p)) \\
&= \frac{f_{\det}(X_s)}{\tilde{\gamma}_{1,j}} \quad (\text{see (32)}).
\end{aligned}$$

$$(36) \quad \partial_1 \psi(1_p) = \partial_1 e_p(1_p) = 1.$$

For any $\Delta \in \mathbb{R}^{n \times p}$ that is zero on the rows indexed by K , substituting the above values back into (35) in Lemma 5.2 with $\phi = \det$ yields that

$$(37) \quad \nabla^2 f_{\det}(X_s)[\Delta, \Delta] = f_{\det}(X_s) \sum_{i \in K^C} \sum_{j=1}^p \left(\frac{\sigma_i^2}{\tilde{\gamma}_{1,j}} - 1 \right) \Delta_{i,j}^2.$$

If $\{\sigma_i\}_{i \in J}$ are *not* the unique numbers in the p leading singular values of A , then there exists $\Delta \in \mathbb{R}^{n \times p}$ such that $\nabla^2 f_{\det}(X_s)[\Delta, \Delta] > 0$, namely, Δ is an ascent direction at X_s . Indeed, let σ_{i_0} with $i_0 \in [n]$ be one of the p leading singular values of A , not listed in $\{\sigma_i\}_{i \in J}$. Then it holds that

$$(38) \quad \sigma_{i_0}^2 > \min_{i \in J} \sigma_i^2 = \min_{j \in [p]} \tilde{\gamma}_{1,j} =: \tilde{\gamma}_{1,j_0} \quad (\text{see (34)})$$

and,² moreover,

$$(39) \quad \frac{\sigma_{i_0}^2}{\tilde{\gamma}_{1,j_0}} \geq \frac{\sigma_p^2}{\sigma_{p+1}^2}.$$

Let $\Delta \in \mathbb{R}^{n \times p}$ be such that Δ_{i_0,j_0} is its only nonzero entry and note that this choice of Δ is indeed zero on the rows indexed by K . With this choice of Δ in (37), we find that

$$(40) \quad \nabla^2 f_{\det}(X_s)[\Delta, \Delta] = f_{\det}(X_s) \left(\frac{\sigma_{i_0}^2}{\tilde{\gamma}_{1,j_0}} - 1 \right) \Delta_{i_0,j_0}^2 > 0 \quad (\text{see (38)}).$$

That is, if $\{\sigma_i\}_{i \in J}$ are *not* the unique numbers in the p leading singular values of A , then there exists an ascent direction at X_s . Moreover, if there is a nontrivial spectral gap $\sigma_p > \sigma_{p+1}$, then it also holds that

$$\begin{aligned}
\nabla^2 f_{\det}(X_s)[\Delta, \Delta] &= f_{\det}(X_s) \left(\frac{\sigma_{i_0}^2}{\tilde{\gamma}_{1,j_0}} - 1 \right) \Delta_{i_0,j_0}^2 \\
&\geq f_{\det}(X_s) \left(\frac{\sigma_p^2}{\sigma_{p+1}^2} - 1 \right) \Delta_{i_0,j_0}^2 \quad (\text{see (39)}) \\
(41) \quad &= f_{\det}(X_s) \left(\frac{\sigma_p^2}{\sigma_{p+1}^2} - 1 \right) \|\Delta\|_F^2,
\end{aligned}$$

²In (38), the index j_0 might not be uniquely defined.

where the last line uses the fact that Δ_{i_0, j_0} is the only nonzero entry of Δ above. Likewise, we can establish that if $\{\sigma_i\}_{i \in J}$ are *not* the unique numbers in the p trailing singular values of A , then there exists a descent direction at X_s . We conclude that if $\{\sigma_i\}_{i \in J}$ are neither the unique numbers in the p leading nor the p trailing singular values of A , then X_s is a strict saddle point (because it has both an ascent and a descent direction).

On the other hand, if $\{\sigma_i\}_{i \in J}$ are the unique numbers in the p leading singular values of A , then all corresponding stationary points take the same objective value $f_{\det}$, which must (globally) maximise the (continuous) objective $f_{\det}$ on the compact set $\text{St}(n, p)$. That is, every such stationary point X_s is in fact a global maximizer of program (6). Likewise, if $\{\sigma_i\}_{i \in J}$ are the unique numbers in the p trailing singular values of A , then all corresponding stationary points are global minimizers. This completes the proof of Theorem 2.1.

The advantage of Lemma 5.2 above is that it gives an explicit expression for the Hessian of f_ϕ , which will be used to prove Theorem 4.1. For the sake of completeness, however, let us show how Lemma 5.2 can be replaced with a simpler argument here, described next. In light of Lemma 5.1 and, if necessary, after a change of basis in (22), we can without loss of generality assume that a stationary point X_s of Program (6) is of the form

$$(42) \quad X_s = \begin{bmatrix} c_{i_1} & \cdots & c_{i_p} \end{bmatrix} \in \text{St}(n, p),$$

where $\sigma_{i_1} \geq \cdots \geq \sigma_{i_p}$, and $c_i \in \mathbb{R}^n$ is the i th canonical vector that takes one at index i and zero elsewhere. If X_s does not correspond to p leading singular values of A , then there exists $i_0 < i_1$ such that $\sigma_{i_0} > \min_{i_j} \sigma_{i_j}$. To simplify the presentation below, let us assume that in fact $\sigma_{i_0} > \sigma_{i_1}$. Now consider the trajectory $\theta \rightarrow X(\theta)$ specified as

$$(43) \quad X(\theta) = \begin{bmatrix} c_{i_0} & c_{i_1} & \cdots & c_{i_p} \end{bmatrix} \begin{bmatrix} \cos \theta & \sin \theta & & \\ -\sin \theta & \cos \theta & & \\ & & \mathbf{I}_{p-1} & \\ & & & \end{bmatrix} \begin{bmatrix} 0 \\ \mathbf{I}_p \end{bmatrix} \in \text{St}(n, p),$$

where the empty blocks in the square matrix above are filled with zeros. It is easy to verify that

$$(44) \quad \frac{d^2 f_{\det}}{d\theta^2}(0) = \frac{2(\sigma_{i_0}^2 - \sigma_{i_1}^2)}{\sigma_{i_1}^2} > 0.$$

That is, there exists an ascent direction at any stationary point that does not correspond to p leading singular values of A . Likewise, one can verify that there exists a descent direction at any stationary point that does not correspond to p trailing singular values of A , and now the rest of the proof of Theorem 2.1 follows as before.

6. Proof of Theorem 4.1. The proof strategy is similar to that of Theorem 2.1 but with some technical subtleties. Without loss of generality, we assume again that $A^*A = \Gamma$ in (14), namely, we assume henceforth that

$$(45) \quad f_{s_q}(X) = \frac{s_q(X^* \Gamma X)}{s_q(X^* X)}.$$

As with Theorem 2.1, we will prove Theorem 4.1 by studying the stationary points of program (15), which we rewrite in the equivalent form

$$(46) \quad \max_{X \in \mathbb{R}^{n \times p}} \min_{\Lambda \in \mathbb{R}^{p \times p}} f_{s_q}(X) + \langle X^* X - \mathbf{I}_p, \Lambda \rangle.$$

Therefore, $X_s \in \text{St}(n, p)$ is a stationary point of programs (15) and (46) if and only if there exists $\Lambda_s \in \mathbb{R}^{p \times p}$ such that

$$(47) \quad \nabla f_{s_q}(X_s) + X_s(\Lambda_s + \Lambda_s^*) = 0,$$

namely, when $\nabla f_{s_q}(X_s)$ belongs to the normal space to the Stiefel manifold at X_s [11]. Without loss of generality, let us assume that $\Lambda_s = \Lambda_s^*$, so that the above condition simplifies to

$$(48) \quad \nabla f_{s_q}(X_s) + 2X_s\Lambda_s = 0.$$

In particular, since $X_s \in \text{St}(n, p)$, we can multiply both sides above by X_s^* and solve for Λ_s above to obtain that

$$(49) \quad \Lambda_s = -\frac{1}{2}X_s^*\nabla f_{s_q}(X_s).$$

Next, from (45), it follows that

$$(50) \quad \nabla f_{s_q}(X_s) = \frac{2\Gamma X_s \nabla s_q(X_s^* \Gamma X_s)}{s_q(X_s^* X_s)} - \frac{2s_q(X_s^* \Gamma X_s) \cdot X_s \nabla s_q(X_s^* X_s)}{s_q(X_s^* X_s)^2}.$$

Let us examine the above expression more carefully. For $Z \in \mathbb{R}^{p \times p}$ with eigen-decomposition $Z = U_Z \text{diag}(\lambda_Z) U_Z^*$, note that the symmetric function corresponding to $\phi = s_q$ is

$$(51) \quad \psi(\lambda_Z) = e_q(\lambda_Z) := \sum_{1 \leq i_1 < \dots < i_q \leq p} \prod_{j=1}^q \lambda_{Z, i_j}.$$

For a nonsingular matrix $Z \in \text{GL}(p)$, it is then not difficult to verify that

$$(52) \quad \nabla e_q(\lambda_Z) = \begin{bmatrix} e_{q-1}(\lambda_Z^1) & \dots & e_{q-1}(\lambda_Z^p) \end{bmatrix}^*,$$

where $\lambda_Z^i \in \mathbb{R}^{p-1}$ is formed from $\lambda_Z \in \mathbb{R}^p$ by removing its i th entry, namely, $\lambda_{Z, i}$. Using (27), we immediately find that

$$(53) \quad \nabla s_q(Z) = W_Z \text{diag} \left(\begin{bmatrix} e_{q-1}(\lambda_Z^1) & \dots & e_{q-1}(\lambda_Z^p) \end{bmatrix}^* \right) W_Z^*.$$

Recalling that $X_s \in \text{St}(n, p)$ and using (53), we calculate the gradients involved in (50) as

$$(54) \quad \begin{aligned} \nabla s_q(X_s^* X_s) &= \binom{p-1}{q-1} X_s^* X_s \quad (\text{see (51), (53)}) \\ &= \binom{p-1}{q-1} I_p, \quad (X_s \in \text{St}(n, p)), \end{aligned}$$

$$(55) \quad \begin{aligned} \nabla s_q(X_s^* \Gamma X_s) &= X_s^* \text{diag} \left(\begin{bmatrix} e_{q-1}(\tilde{\gamma}_1^1) & \dots & e_{q-1}(\tilde{\gamma}_1^p) \end{bmatrix}^* \right) X_s \quad (\text{see (33), (53)}) \\ &= \text{diag} \left(\begin{bmatrix} e_{q-1}(\tilde{\gamma}_1^1) & \dots & e_{q-1}(\tilde{\gamma}_1^p) \end{bmatrix}^* \right) \quad (\text{see Lemma 5.2}) \\ &=: \text{diag}(\tilde{\gamma}_1), \end{aligned}$$

where $\tilde{\gamma}_1^i \in \mathbb{R}^{p-1}$ is formed from $\tilde{\gamma}_1$ by removing its i th entry; see (21). By substituting (54), (55) back into (50), we conclude that as in the proof of Theorem 2.1 $X_s \in \text{St}(n, p)$ such that $X_s^* \Gamma X_s \in \text{GL}(p)$ is a stationary point of Program (15) if and only if

$$(56) \quad \text{range}(X_s) = \text{range}(\Gamma X_s),$$

which is identical to (31) in the proof of Theorem 2.1, and consequently Lemmas 5.1 and 5.2 therein apply here too. Moreover, note that

$$(57) \quad \begin{aligned} s_q(X_s^* X_s) &= s_q(I_p) \quad (X_s \in \text{St}(n, p)) \\ &= e_q(1_p) = \binom{p}{q} \quad (\text{see (51)}), \end{aligned}$$

$$(58) \quad \begin{aligned} s_q(X_s^* \Gamma X_s) &= s_q(\text{diag}(\tilde{\gamma}_1)) \quad (\text{see Lemma 5.2}) \\ &= e_q(\tilde{\gamma}_1) \quad (\text{see (51)}). \end{aligned}$$

We can also revisit (49) to obtain that

$$(59) \quad \begin{aligned} \Lambda_s &= -\frac{1}{2} X_s^* \nabla f_{s_q}(X_s) \quad (\text{see (49)}) \\ &= -\frac{\text{diag}(\tilde{\gamma}_1) \text{diag}(\tilde{\gamma}_1)}{\binom{p}{q}} + \frac{e_q(\tilde{\gamma}_1) \binom{p-1}{q-1} I_p}{\binom{p}{q}^2} \quad (\text{see (54), (55), (57), (58)}) \\ &= -\frac{\binom{p-1}{q-1}}{\binom{p}{q}} \left(\frac{\text{diag}(\tilde{\gamma}_1 \tilde{\gamma}_1)}{\binom{p-1}{q-1}} - \frac{e_q(\tilde{\gamma}_1) I_p}{\binom{p}{q}} \right). \end{aligned}$$

In particular, when $\phi = s_q$ (and consequently $\psi = e_q$), we next simplify the expression for Hessian in (35) by noting that

$$\phi(I_p) = s_q(I_p) = e_q(1_p) = \binom{p}{q} \quad (\text{see (51)}),$$

$$\partial_j \psi(\tilde{\gamma}_1) = \partial_j e_q(\tilde{\gamma}_1) = e_{q-1}(\tilde{\gamma}_1^j), \quad j \in [p],$$

where $\tilde{\gamma}_1^j \in \mathbb{R}^{p-1}$ is formed from $\tilde{\gamma}_1 \in \mathbb{R}^p$ by removing its j th entry; see (34). Moreover,

$$\partial_1 \psi(1_p) = \partial_1 e_q(1_p) = e_{q-1}(1_{p-1}) = \binom{p-1}{q-1} \quad (\text{see (51)}),$$

$$(60) \quad \begin{aligned} f_{s_q}(X_s) &= \frac{s_q(X_s^* \Gamma X_s)}{s_q(X_s^* X_s)} \quad (\text{see (32)}) \\ &= \frac{s_q(\text{diag}(\tilde{\gamma}_1))}{s_q(I_p)} \quad ((33) \text{ and } X_s \in \text{St}(n, p)) \\ &= \frac{e_q(\tilde{\gamma}_1)}{e_q(1_p)} \\ &= \frac{e_q(\tilde{\gamma}_1)}{\binom{p}{q}} \quad (\text{see (51)}), \end{aligned}$$

$$(61) \quad \begin{aligned} \nabla^2 f_{s_q}(X_s)[\Delta, \Delta] &= \frac{1}{\phi(I_p)} \sum_{i \in K^C} \sum_{j=1}^p (\sigma_i^2 \partial_j e_q(\tilde{\gamma}_1) - f_{s_q}(X_s) \partial_1 e_q(1_p)) \Delta_{i,j}^2 \quad (\text{see (35)}) \\ &= \frac{1}{\binom{p}{q}} \sum_{i \in K^C} \sum_{j=1}^p \left(\sigma_i^2 e_{q-1}(\tilde{\gamma}_1^j) - \frac{e_q(\tilde{\gamma}_1) \binom{p-1}{q-1}}{\binom{p}{q}} \right) \Delta_{i,j}^2 \\ &= \frac{\binom{p-1}{q-1}}{\binom{p}{q}} \sum_{i \in K^C} \sum_{j=1}^p \left(\frac{\sigma_i^2 e_{q-1}(\tilde{\gamma}_1^j)}{\binom{p-1}{q-1}} - \frac{e_q(\tilde{\gamma}_1)}{\binom{p}{q}} \right) \Delta_{i,j}^2. \end{aligned}$$

In light of (46), (48), let us record for future reference that $\Delta \in \mathbb{R}^{n \times p}$ is an ascent direction at X_s if

$$(62) \quad \nabla^2 f_{s_q}(X_s)[\Delta, \Delta] + \langle \Delta^* \Delta, \Lambda_s \rangle \geq 0$$

and

$$(63) \quad X_s^\top \Delta + \Delta^\top X_s = 0.$$

By definition in (34), $\{\sigma_i^2\}_{i \in J}$ are the distinct numbers appearing in $\tilde{\gamma}_1 \in \mathbb{R}^p$. Suppose now that $\{\sigma_i\}_{i \in J}$ are *not* the unique numbers in the p leading singular values of A . Therefore there exist $i_0 \notin K$ and $j_0 \in [p]$ such that

$$(64) \quad \sigma_{i_0}^2 > \min_{i \in J} \sigma_i^2 = \min_{j \in [p]} \tilde{\gamma}_{1,j} =: \tilde{\gamma}_{1,j_0} \quad (\text{see (34)}),$$

and,³ moreover,

$$(65) \quad \frac{\sigma_{i_0}^2}{\tilde{\gamma}_{1,j_0}} \geq \frac{\sigma_p^2}{\sigma_{p+1}^2}.$$

Let us set $\Delta \in \mathbb{R}^{n \times p}$ such that Δ_{i_0,j_0} is its only nonzero entry and note that (63) holds because $i_0 \notin K$. For this choice of Δ , we find that

$$\begin{aligned} (66) \quad & \nabla^2 f_{s_q}(X_s)[\Delta, \Delta] + \langle \Delta^* \Delta, \Lambda_s \rangle \\ &= \frac{\binom{p-1}{q-1}}{\binom{p}{q}} \left(\frac{\sigma_{i_0}^2 e_{q-1}(\tilde{\gamma}_1^j)}{\binom{p-1}{q-1}} - \frac{e_q(\tilde{\gamma}_1)}{\binom{p}{q}} \right) \Delta_{i_0,j_0}^2 + \Lambda_{s,j_0,j_0}^2 \Delta_{i_0,j_0}^2 \quad (\text{see (61)}) \\ &= \frac{\binom{p-1}{q-1}}{\binom{p}{q}} \left(\frac{\sigma_{i_0}^2 e_{q-1}(\tilde{\gamma}_1^{j_0})}{\binom{p-1}{q-1}} - \frac{e_q(\tilde{\gamma}_1)}{\binom{p}{q}} \right) \Delta_{i_0,j_0}^2 \\ &\quad - \frac{\binom{p-1}{q-1}}{\binom{p}{q}} \left(\frac{\tilde{\gamma}_{1,j_0} \hat{\gamma}_{1,j_0}}{\binom{p-1}{q-1}} - \frac{e_q(\tilde{\gamma}_1)}{\binom{p}{q}} \right) \Delta_{i_0,j_0}^2 \quad (\text{see (59)}) \\ &= \frac{\sigma_{i_0}^2 e_{q-1}(\tilde{\gamma}_1^{j_0}) - \tilde{\gamma}_{1,j_0} \hat{\gamma}_{1,j_0}}{\binom{p}{q}} \Delta_{i_0,j_0}^2 \\ &= \frac{(\sigma_{i_0}^2 - \tilde{\gamma}_{1,j_0}) e_{q-1}(\tilde{\gamma}_1^j)}{\binom{p}{q}} \Delta_{i_0,j_0}^2 \quad (\text{see (55)}) \\ &> 0 \quad (\text{see (64)}). \end{aligned}$$

That is, if $\{\sigma_i\}_{i \in J}$ are *not* the unique members in the p leading singular values of A , then there exists an ascent direction at X_s . The rest of the proof of Theorem 4.1 is now the same as that of Theorem 2.1.

Appendix A. Proof of Lemma 5.1. Consider $X_s \in \text{St}(n, p)$ such that $X_s^* \Gamma X_s \in \text{GL}(p)$ and

$$(67) \quad \text{range}(X_s) = \text{range}(\Gamma X_s).$$

Each row of X_s naturally corresponds to a singular value of A , namely, the i th row corresponds to σ_i , where σ_i^2 is the i th diagonal entry of Γ . Let $J \subset [n]$ be the index

³In (64), the index j_0 might not be uniquely defined.

set such that $\Sigma_J := \{\sigma_i\}_{i \in J}$ is the set of *distinct* singular values corresponding to the nonzero rows of X_s .⁴ For future reference, let us record that Σ_J contains only positive singular values, namely,

$$(68) \quad \Sigma_J \subset \mathbb{R}_+.$$

Indeed, if $0 \in \Sigma_J$, namely, if $\sigma_n = 0$, then the rows of X_s corresponding to σ_n are zero too thanks to (67) and consequently $0 \notin \Sigma_J$, which leads to a contradiction. Let us now set

$$(69) \quad \dim(\sigma_i) := \max \left[p - \sum_{j \in J, j \neq i} n_j, 1 \right], \quad i \in J,$$

for short, where n_i is the multiplicity of σ_i . Fix $i_0 \in J$. Consider $\mathcal{K}_{i_0}$, the collection of all index sets $K \subset [n]$ of size p such that

$$(70) \quad \sigma_{i_0} \in \text{unique}(\{\sigma_i\}_{i \in K}) \subseteq \Sigma_J,$$

where $\text{unique}(\{\sigma_i\}_{i \in K})$ returns the distinct members of the set $\{\sigma_i\}_{i \in K}$. In words, every index set $K \in \mathcal{K}_{i_0}$ contains σ_{i_0} and $p-1$ other (not necessarily distinct) singular values of A corresponding to nonzero rows of X_s . Consider an arbitrary $K \in \mathcal{K}_{i_0}$. It follows from (69) that

$$(71) \quad \{\sigma_i\}_{i \in K} \text{ contains at least } \dim(\sigma_{i_0}) \text{ copies of } \sigma_{i_0}.$$

On the other hand, (67) implies that there exists $B \in \mathbb{R}^{p \times p}$ such that

$$(72) \quad X_s B = \Gamma X_s.$$

By multiplying both sides above by X_s^* and using the fact that $X_s \in \text{St}(n, p)$, we infer from (72) that

$$(73) \quad B = X_s^* \Gamma X_s \in \text{GL}(p),$$

where the invertibility of B follows from the assumption of Lemma 5.1. In addition, (72) means that each row of X_s is an eigenvector of B . By restricting (72) to the index set K , we find that

$$(74) \quad X_s[K, :] \cdot B = \Gamma[K, K] \cdot X_s[K, :] \text{ is the eigendecomposition of } B,$$

because $K \in \mathcal{K}_{i_0}$ is a set of size p by the definition of $\mathcal{K}_{i_0}$ earlier. Above, we used the MATLAB matrix notation. For example, $X_s[K, :] \in \mathbb{R}^{p \times p}$ above is the row-submatrix of X_s corresponding to the rows indexed by K . It follows from (74) that

⁴Throughout, we treat $\{\sigma_i\}_{i \in J}$ and similar items as sequences (rather than sets) to allow for repetitions.

(75) σ_{i_0} is an eigenvalue of B with the multiplicity of at least $\dim(\sigma_{i_0})$,

because, by (71), $\{\sigma_i\}_{i \in K}$ contains at least $\dim(\sigma_{i_0})$ copies of σ_{i_0} . In fact, there exists an index set $K_0 \in \mathcal{K}_{i_0}$ such that $\{\sigma_i\}_{i \in K_0}$ contains *exactly* $\dim(\sigma_{i_0})$ copies of σ_{i_0} .⁵ It follows that

(76) σ_{i_0} is an eigenvalue of B with the multiplicity of exactly $\dim(\sigma_{i_0})$.

By (73), B is full-rank and it follows from (76) that the corresponding eigenvectors of B span a $\dim(\sigma_{i_0})$ -dimensional subspace of $\mathbb{R}^p$, namely, the geometric multiplicity of σ_{i_0} is $\dim(\sigma_{i_0})$. Since every row of X_s is an eigenvector of B by (74), it follows that the

(77) rows of X_s that correspond to σ_{i_0} span a $\dim(\sigma_{i_0})$ -dimensional subspace of $\mathbb{R}^p$.

Since the choice of $i_0 \in J$ was arbitrary above, we find for every $i \in J$ that

(78) the rows of X_s corresponding to σ_i span a $\dim(\sigma_i)$ -dimensional subspace of $\mathbb{R}^p$, denoted by $S_i \in \text{Gr}(p, \dim(\sigma_i))$.

Because B is symmetric by its definition in (73), these subspaces are orthogonal to one another, namely,

(79) $S_i \perp S_j, \quad i \neq j \text{ and } i, j \in J.$

On the other hand, note that

$$\begin{aligned} B &= X_s^* \Sigma X_s \quad (\text{see (72)}) \\ &= \sum_{i \in J} \sigma_i \left(\sum_{\sigma_j = \sigma_i} X_s[j, :]^* X_s[j, :] \right) + \sum_{\sigma_j \notin \Sigma_J} \sigma_j \cdot X_s[j, :]^* X_s[j, :] \\ &=: \sum_{i \in J} \sigma_i Y_i + \sum_{\sigma_j \notin \Sigma_J} \sigma_j \cdot X_s[j, :]^* X_s[j, :] \\ &= \sum_{i \in J} \sigma_i Y_i, \end{aligned}$$

where the last line above follows from the definition of J , namely, any singular value $\sigma_i \notin \Sigma_J$ corresponds to a zero row of X_s . For every $i \in J$, note that

(80) $\text{range}(Y_i) = S_i$

by definition of S_i in (78). It therefore follows from (79) that $\{Y_i\}_{i \in J}$ are pairwise orthogonal matrices, namely,

(81) $Y_i^* Y_j = 0, \quad i \neq j \text{ and } i, j \in J.$

Therefore, $B = \sum_{i \in J} \sigma_i Y_i$ is the eigendecomposition of B and, because B is full-rank by (73), we find that

⁵Indeed, if $\dim(\sigma_{i_0}) > 1$, such an index set K_0 would include n_i copies of singular value σ_i for every $\sigma_i \neq \sigma_{i_0}$ with $i \in K$. The construction is similar if $\dim(\sigma_{i_0}) = 1$.

$$\begin{aligned}
p &= \sum_{i \in J, \sigma_i \neq 0} \dim(\text{span}(Y_i)) \\
&= \sum_{i \in J, \sigma_i \neq 0} \dim(S_i) \quad (\text{see (80)}) \\
&= \sum_{i \in J, \sigma_i \neq 0} \dim(\sigma_i) \quad (\text{see (78)}) \\
(82) \quad &= \sum_{i \in J} \dim(\sigma_i) \quad (\text{see (68)}).
\end{aligned}$$

This completes the proof of Lemma 5.1.

Appendix B. Proof of Lemma 5.2. Suppose that $X_s \in \text{St}(n, p)$ satisfies

$$(83) \quad \text{range}(X_s) = \text{range}(\Gamma X_s),$$

which implies that

$$(84) \quad X_s B = \Gamma X_s, \quad \text{where } B = X_s^* \Gamma X_s.$$

Let $J \subseteq [n]$ be the corresponding index set prescribed in Lemma 5.1, and recall that $\{\sigma_i\}_{i \in J}$ are distinct by item 1 in Lemma 5.1. Consider an index set $K \supseteq J$ such that $\{\sigma_i\}_{i \in K}$ contains all available copies of the singular values listed in $\{\sigma_i\}_{i \in J}$. Let k denote the size of K . For convenience, we define

$$X_{s,1} := X_s [K, :] \in \mathbb{R}^{k \times p}, \quad X_{s,2} := X_s [K^C, :] \in \mathbb{R}^{(n-k) \times p},$$

$$(85) \quad \Gamma_1 := \Gamma [K, K] \in \mathbb{R}^{k \times k}, \quad \Gamma_2 := \Gamma [K^C, K^C] \in \mathbb{R}^{(n-k) \times (n-k)},$$

where we used the MATLAB matrix notation above. For example, $X_s [K, :]$ is the restriction of X_s to the rows indexed in K . Also, K^C is the complement of index set K with respect to $[n]$. In particular, item 2 in Lemma 5.1 immediately implies that

$$(86) \quad X_{s,2} = 0,$$

and, consequently,

$$\begin{aligned}
X_{s,1}^* X_{s,1} &= X_{s,1}^* X_{s,1} + X_{s,2}^* X_{s,2} \\
&= X_s^* X_s \\
(87) \quad &= I_p \quad (X_s \in \text{St}(n, p)).
\end{aligned}$$

That is,

$$(88) \quad X_{s,1} \in \text{St}(k, p).$$

Note that (83) holds also after a change of basis from X_s to $X_s \Theta$ for invertible $\Theta \in \text{GL}(p)$. Therefore, thanks to (83) and item 4 in Lemma 5.1, we can assume without loss of generality that the supports of rows and also columns of $X_{s,1}$ are disjoint. More specifically, with the enumeration $J = \{i_1, i_2, \dots\}$, we assume without loss of generality that

$$(89) \quad X_{s,1} = \begin{bmatrix} X_{s,1,1} & 0 & \cdots \\ 0 & X_{s,1,2} & \cdots \\ \vdots & \vdots & \ddots \end{bmatrix} \in \mathbb{R}^{k \times p},$$

where the rows of the block $X_{s,1,1}$ correspond to the singular value σ_{i_1} and has $\dim(\sigma_{i_1})$ columns, the block $X_{s,1,2}$ corresponds to σ_{i_2} and so on. In particular, (88) implies that

$$(90) \quad X_{s,1,1}^* X_{s,1,1} = I_{\dim(\sigma_{i_1})}, \quad X_{s,1,2}^* X_{s,1,2} = I_{\dim(\sigma_{i_2})}, \quad \cdots,$$

namely, $X_{s,1,1}$ has orthonormal columns, and so do $X_{s,1,2}$ and the rest of the diagonal blocks of $X_{s,1}$. Another necessary ingredient in our analysis below is the observation that

$$\begin{aligned} X_s^* \Gamma X_s &= X_{s,1}^* \Gamma_1 X_{s,1} + X_{s,2}^* \Gamma_2 X_{s,2} \\ &= X_{s,1}^* \Gamma_1 X_{s,1} \quad (\text{see (86)}) \\ &= \begin{bmatrix} \sigma_{i_1}^2 \cdot X_{s,1,1}^* X_{s,1,1} & 0 & \cdots \\ 0 & \sigma_{i_2}^2 \cdot X_{s,1,2}^* X_{s,1,2} & \cdots \\ \vdots & \vdots & \ddots \end{bmatrix} \\ &= \begin{bmatrix} \sigma_{i_1}^2 \cdot I_{\dim(\sigma_{i_1})} & 0 & \cdots \\ 0 & \sigma_{i_2}^2 \cdot I_{\dim(\sigma_{i_2})} & \cdots \\ \vdots & \vdots & \ddots \end{bmatrix} \quad (\text{see (90)}) \\ &=: \tilde{\Gamma}_1 \in \mathbb{R}^{p \times p} \\ (91) \quad &=: \text{diag}(\tilde{\gamma}_1), \end{aligned}$$

namely, the diagonal matrix $\tilde{\Gamma}_1$ contains $\dim(\sigma_{i_1})$ copies of $\sigma_{i_1}^2$, $\dim(\sigma_{i_2})$ copies of $\sigma_{i_2}^2$, and so on. To compute the Hessian of f_ϕ , we make a small perturbation to its argument. To be specific, consider $\Delta \in \mathbb{R}^{n \times p}$ that is supported only on the rows indexed by K^C and let

$$(92) \quad \Delta_2 := \Delta [K^C, :] \in \mathbb{R}^{(n-k) \times p}$$

be the nonzero block of Δ . Note in particular that

$$(93) \quad X_s^* \Delta = 0,$$

because by construction X_s and Δ are supported on the rows indexed by K and K^C , respectively; see (86). Let $h_\Gamma(X) = \phi(X^* \Gamma X)$ for short and note that

$$\begin{aligned} h_\Gamma(X_s + \Delta) &= \phi((X_s + \Delta)^* \Gamma (X_s + \Delta)) \\ &= \phi(X_s^* \Gamma X_s + X_s^* \Gamma \Delta + \Delta^* \Gamma X_s + \Delta^* \Gamma \Delta) \\ &= \phi(X_s^* \Gamma X_s + B X_s^* \Delta + \Delta^* X_s B + \Delta^* \Gamma \Delta) \quad (\text{see (84)}) \\ &= \phi(X_s^* \Gamma X_s + \Delta^* \Gamma \Delta) \quad (\text{see (93)}) \\ &= \phi(\tilde{\Gamma}_1 + \Delta_2^* \Gamma_2 \Delta_2) \quad (\text{see (91), (92), (85)}) \\ &= \phi(\tilde{\Gamma}_1) + \langle \nabla \phi(\tilde{\Gamma}_1), \Delta_2^* \Gamma_2 \Delta_2 \rangle + O(\|\Delta_2\|^3) \quad (\text{Taylor expansion}) \\ (94) \quad &= \phi(X_s^* \Gamma X_s) + \langle \nabla \phi(\tilde{\Gamma}_1), \Delta_2^* \Gamma_2 \Delta_2 \rangle + O(\|\Delta\|^3) \quad (\text{see (91), (92)}), \end{aligned}$$

where we used the standard Big- O notation above. Recall from (91) that $\tilde{\Gamma}_1 = \text{diag}(\tilde{\gamma}_1)$. Because ϕ is by assumption a spectral function with the corresponding symmetric function ψ , (27) implies that

$$(95) \quad \nabla \phi(\tilde{\Gamma}_1) = \text{diag}(\nabla \psi(\tilde{\gamma}_1)) = \text{diag} \left(\begin{bmatrix} \partial_1 \psi(\tilde{\gamma}_{1,1}) & \cdots & \partial_p \psi(\tilde{\gamma}_{1,p}) \end{bmatrix}^* \right),$$

which allows us to rewrite the last line above as

$$(96) \quad \begin{aligned} h_\Gamma(X_s + \Delta) &= \phi(X_s^* \Gamma X_s) + \langle \text{diag}(\nabla \psi(\tilde{\gamma}_1)), \Delta_2^* \Gamma_2 \Delta_2 \rangle + O(\|\Delta\|^3) \\ &= \phi(X_s^* \Gamma X_s) + \sum_{i \in K^C} \sum_{j=1}^p \sigma_i^2 \cdot \partial_j \psi(\tilde{\gamma}_1) \cdot \Delta[i, j]^2 + O(\|\Delta\|^3), \end{aligned}$$

where $\Delta[i, j]$ is the $[i, j]$ th entry of Δ . Let $1_p \in \mathbb{R}^p$ be the vector of all ones. After setting $h_I(X) = \phi(X^* X)$ and after replacing Γ with I_n above, we find that

$$(97) \quad \begin{aligned} h_I(X_s + \Delta) &= \phi(X_s^* X_s) + \sum_{i \in K^C} \sum_{j=1}^p \partial_j \psi(1_p) \cdot \Delta[i, j]^2 + O(\|\Delta\|^3) \\ &= \phi(I_p) + \partial_1 \psi(1_p) \sum_{i \in K^C} \sum_{j=1}^p \Delta[i, j]^2 + O(\|\Delta\|^3), \end{aligned}$$

where in the last line above we used the fact that $X_s \in \text{St}(n, p)$ and that ψ is a symmetric function, hence $\partial_j \psi(1_p) = \partial_1 \psi(1_p)$ for every $j \in [p]$. Since $f_\phi = h_\Gamma/h_I$ by definition, (96), (97) imply that

$$(98) \quad \begin{aligned} f_\phi(X_s + \Delta) &= \frac{h_\Gamma(X_s + \Delta)}{h_I(X_s + \Delta)} \\ &= \frac{\phi(X_s^* \Gamma X_s) + \sum_{i \in K^C} \sum_{j=1}^p \sigma_i^2 \cdot \partial_j \psi(\tilde{\gamma}_1) \cdot \Delta[i, j]^2 + O(\|\Delta\|^3)}{\phi(I_p) + \partial_1 \psi(1_p) \sum_{i \in K^C} \sum_{j=1}^p \Delta[i, j]^2 + O(\|\Delta\|^3)} \quad (\text{see (96), (97)}) \\ &= \left(\phi(X_s^* \Gamma X_s) + \sum_{i \in K^C} \sum_{j=1}^p \sigma_i^2 \cdot \partial_j \psi(\tilde{\gamma}_1) \cdot \Delta[i, j]^2 + O(\|\Delta\|^3) \right) \\ &\quad \cdot \frac{1}{\phi(I_p)} \left(1 - \frac{\partial_1 \psi(1_p)}{\phi(I_p)} \sum_{i \in K^C} \sum_{j=1}^p \Delta[i, j]^2 + O(\|\Delta\|^3) \right) \quad \left(\frac{1}{1+a} = 1 - a + O(a^2) \right) \\ &= \frac{\phi(X_s^* \Gamma X_s)}{\phi(I_p)} - \frac{\phi(X_s^* \Gamma X_s)}{\phi(I_p)} \cdot \frac{\partial_1 \psi(1_p)}{\phi(I_p)} \sum_{i \in K^C} \sum_{j=1}^p \Delta[i, j]^2 \\ &\quad + \frac{1}{\phi(I_p)} \sum_{i \in K^C} \sum_{j=1}^p \sigma_i^2 \cdot \partial_j \psi(\tilde{\gamma}_1) \cdot \Delta[i, j]^2 + O(\|\Delta\|^3) \\ &= f_\phi(X_s) - f_\phi(X_s) \frac{\partial_1 \psi(1_p)}{\phi(I_p)} \sum_{i \in K^C} \sum_{j=1}^p \Delta[i, j]^2 \\ &\quad + \frac{1}{\phi(I_p)} \sum_{i \in K^C} \sum_{j=1}^p \sigma_i^2 \cdot \partial_j \psi(\tilde{\gamma}_1) \cdot \Delta[i, j]^2 + O(\|\Delta\|^3) \quad (\text{see (32)}), \end{aligned}$$

and, consequently,

$$(99) \quad \nabla^2 f_\phi(X_s)[\Delta, \Delta] = \frac{1}{\phi(I_p)} \sum_{i \in K^C} \sum_{j=1}^p (\sigma_i^2 \partial_j \psi(\tilde{\gamma}_1) - f_\phi(X_s) \partial_1 \psi(1_p)) \Delta_{i,j}^2,$$

for our particular choice of Δ that satisfies $\Delta[K, :] = 0$. Here, the bilinear operator $\nabla^2 f_\phi(X_s) : \mathbb{R}^{n \times p} \times \mathbb{R}^{n \times p} \rightarrow \mathbb{R}$ is the Hessian of f_ϕ at X_s . This completes the proof of Lemma 5.2.

Acknowledgments. AE would like to thank Stephen Becker and David Bortz for pointing out the connection to D-optimality in optimal design and Mike Davis for the connection to independent component analysis.

REFERENCES

- [1] O. ALTER, P. O. BROWN, AND D. BOTSTEIN, *Singular value decomposition for genome-wide expression data processing and modeling*, Proc. Natl. Acad. Sci. USA, 97 (2000), pp. 10101–10106.
- [2] O. ALTER AND G. GOLUB, *Singular value decomposition of genome-scale mRNA lengths distribution reveals asymmetry in RNA gel electrophoresis band broadening*, Proc. Natl. Acad. Sci. USA, 103 (2006), pp. 11828–11833, <https://doi.org/10.1073/pnas.0604756103>.
- [3] S. BHOJANAPALLI, A. KYRILLIDIS, AND S. SANGHAVI, *Dropping convexity for faster semi-definite optimization*, in Proceedings of the Conference on Learning Theory, 2016, pp. 530–582.
- [4] S. BHOJANAPALLI, B. NEYSHABUR, AND N. SREBRO, *Global optimality of local search for low rank matrix recovery*, in Proceedings of Advances in Neural Information Processing Systems, 2016, pp. 3873–3881.
- [5] I. BORG AND P. GROENEN, *Modern Multidimensional Scaling: Theory and Applications*, Springer Ser. Statist., Springer New York, 2013, <https://books.google.co.uk/books?id=NYDSBwAAQBAJ>.
- [6] N. BOUMAL, V. VORONINSKI, AND A. BANDEIRA, *The non-convex Burer-Monteiro approach works on smooth semidefinite programs*, in Proceedings of Advances in Neural Information Processing Systems, 2016, pp. 2757–2765.
- [7] S. BURER AND R. D. MONTEIRO, *A nonlinear programming algorithm for solving semidefinite programs via low-rank factorization*, Math. Program., 95 (2003), pp. 329–357.
- [8] A. D'ASPREMONT, L. E. GHAOUI, M. I. JORDAN, AND G. R. LANCKRIET, *A direct formulation for sparse PCA using semidefinite programming*, in Proceedings of Advances in Neural Information Processing Systems, 2005, pp. 41–48.
- [9] Y. DESHPANDE AND A. MONTANARI, *Information-theoretically optimal sparse pca*, in Proceedings of the IEEE International Symposium on Information Theory, IEEE, 2014, pp. 2197–2201.
- [10] C. ECKART AND G. YOUNG, *The approximation of one matrix by another of lower rank*, Psychometrika, 1 (1936), pp. 211–218, <https://doi.org/10.1007/BF02288367>.
- [11] A. EDELMAN, T. A. ARIAS, AND S. T. SMITH, *The geometry of algorithms with orthogonality constraints*, SIAM J. Matrix Anal. Appl., 20 (1998), pp. 303–353.
- [12] A. EFTEKHARI, G. ONGIE, L. BALZANO, AND M. B. WAKIN, *Streaming principal component analysis from incomplete data*, J. Mach. Learn. Res., 20 (2019), pp. 1–62.
- [13] R. GE, F. HUANG, C. JIN, AND Y. YUAN, *Escaping from saddle points—online stochastic gradient for tensor decomposition*, in Proceedings of the Conference on Learning Theory, 2015, pp. 797–842.
- [14] R. GE, J. D. LEE, AND T. MA, *Matrix completion has no spurious local minimum*, in Proceedings of Advances in Neural Information Processing Systems, 2016, pp. 2973–2981.
- [15] R. GE, J. D. LEE, AND T. MA, *Learning One-Hidden-Layer Neural Networks with Landscape Design*, <https://arxiv.org/abs/1711.00501>, 2017.
- [16] R. GE AND T. MA, *On the optimization landscape of tensor decompositions*, in Proceedings of Advances in Neural Information Processing Systems, 2017, pp. 3653–3663.
- [17] N. GILLIS, *The why and how of nonnegative matrix factorization*, in Regularization, Optimization, Kernels, and Support Vector Machines, J. A. K. Suykens, M. Signoretto, and A. Argyriou, eds., Chapman & Hall/CRC Press, Boca Raton, FL, 2015, pp. 257–292.
- [18] G. GOLUB AND C. VAN LOAN, *Matrix Computations*, Johns Hopkins University Press, Baltimore, 1996, <https://books.google.ch/books?id=mlOa7wPX6OYC>.

- [19] T. HASTIE, R. TIBSHIRANI, AND J. FRIEDMAN, *The Elements of Statistical Learning: Data Mining, Inference, and Prediction*, Springer Ser. Statist., Springer New York, 2013, <https://books.google.co.uk/books?id=yPfZBwAAQBAJ>.
- [20] R. A. HAUSER, A. EFTEKHARI, AND H. F. MATZINGER, *Pca by determinant optimization has no spurious local optima*, in Proceedings of the 24th ACM SIGKDD International Conference on Knowledge Discovery & Data Mining, ACM, 2018, pp. 1504–1511.
- [21] R. HORN, R. HORN, AND C. JOHNSON, *Matrix Analysis*, Cambridge University Press, Cambridge, UK, 1990, <https://books.google.co.uk/books?id=PIYQN0ypTwEC>.
- [22] R. HORN AND C. JOHNSON, *Topics in Matrix Analysis*, Cambridge University Press, Cambridge, UK, 1994, <https://books.google.co.uk/books?id=ukd0AgAAQBAJ>.
- [23] H. HOTELLING, *Relations between two sets of variates*, Biometrika, (1936), pp. 321–377.
- [24] A. HYVARINEN, J. KARHUNEN, AND E. OJA, *Independent Component Analysis*, Wiley Ser. Adaptive Learning Syst. Signal Process. Learning Commun. Control, Wiley, 2004, New York, <https://books.google.co.uk/books?id=96D0ypDwAkkC>.
- [25] C. JIN, R. GE, P. NETRAPALLI, S. M. KAKADE, AND M. I. JORDAN, *How to escape saddle points efficiently*, in J. Mach. Learn. Res., 70 (2017), pp. 1724–1732.
- [26] C. JIN, S. M. KAKADE, AND P. NETRAPALLI, *Provable efficient online matrix completion via non-convex stochastic gradient descent*, in Proceedings of Advances in Neural Information Processing Systems, 2016, pp. 4520–4528.
- [27] I. M. JOHNSTONE AND A. Y. LU, *On consistency and sparsity for principal components analysis in high dimensions*, J. Amer. Statist. Assoc., 104 (2009), pp. 682–693.
- [28] M. JOURNÉE, Y. NESTEROV, P. RICHTÁRIK, AND R. SEPULCHRE, *Generalized power method for sparse principal component analysis*, J. Mach. Learn. Res., 11 (2010), pp. 517–553.
- [29] S. KARLIN AND Y. RINOTT, *A generalized Cauchy Binet formula and applications to total positivity and majorization*, J. Multivariate Anal., 27 (1988), pp. 284–299.
- [30] A. S. LEWIS AND H. S. SENDOV, *Quadratic expansions of spectral functions*, Linear Algebra Appl., 340 (2002), pp. 97–121.
- [31] Q. LI AND G. TANG, *The Nonconvex Geometry of Low-Rank Matrix Optimizations with General Objective Functions*, <https://arxiv.org/abs/1611.03060v1>, 2016.
- [32] L. MIRSKY, *Symmetric gauge functions and unitarily invariant norms*, Quart. J. Math. Oxford, (1966), pp. 1156–1159.
- [33] A. MOKHTARI, A. OZDAGLAR, AND A. JADBABAIE, *Escaping saddle points in constrained optimization*, in Proceedings of Advances in Neural Information Processing Systems, 2018, pp. 3633–3643.
- [34] *NIST/SEMATECH Engineering Statistics Handbook*, <https://books.google.co.uk/books?id=v-XXjwEACAAJ>, 2002.
- [35] K. PEARSON, *On lines and planes of closest fit to systems of points in space*, Philos. Mag., 2 (1901), pp. 559–572, <https://doi.org/10.1080/14786440109462720>.
- [36] B. SCHÖLKOPF AND A. SMOLA, *Learning with Kernels: Support Vector Machines, Regularization, Optimization, and Beyond*, MIT Press, Cambridge, MA, 2002, <https://books.google.co.uk/books?id=y8ORL3DWt4sC>.
- [37] J. SHAWE-TAYLOR AND N. CRISTIANINI, *Kernel Methods for Pattern Analysis*, Cambridge University Press, Cambridge, UK, 2004, <https://books.google.co.uk/books?id=MX0hAwAAQBAJ>.
- [38] M. SOLTANOLKOTABI, A. JAVANMARD, AND J. D. LEE, *Theoretical Insights into the Optimization Landscape of Over-Parameterized Shallow Neural Networks*, <https://arxiv.org/abs/1707.04926>, 2017.
- [39] J. SUN, Q. QU, AND J. WRIGHT, *When Are Nonconvex Problems Not Scary?*, <https://arxiv.org/abs/1510.06096>, 2015.