
Bures-Wasserstein Barycenters and Low-Rank Matrix Recovery

Tyler Maunu
Brandeis University

Thibaut Le Gouic
École Centrale de Marseille

Philippe Rigollet
Massachusetts Institute of Technology

Abstract

We revisit the problem of recovering a low-rank positive semidefinite matrix from rank-one projections using tools from optimal transport. More specifically, we show that a variational formulation of this problem is equivalent to computing a Wasserstein barycenter. In turn, this new perspective enables the development of new geometric first-order methods with strong convergence guarantees in Bures-Wasserstein distance. Experiments on simulated data demonstrate the advantages of our new methodology over existing methods.

1 INTRODUCTION

Recovering a low-rank matrix is a fundamental primitive across many settings, such as matrix completion [Fazel, 2002, Candès and Recht, 2009, Candès and Tao, 2010], phase retrieval [Candès et al., 2015], principal component analysis [Pearson, 1901, Hotelling, 1933], robust subspace recovery [Lerman and Maunu, 2018], and robust principal component analysis [Chandrasekaran et al., 2009, Candès et al., 2011, Xu et al., 2010]. This line of work can be understood as a generalization of the classical compressed sensing question [Donoho, 2006, Candès et al., 2006], where the goal is the recovery of a sparse vector. Indeed, the sparse vector recovery problem can be cast as a low-rank recovery problem over diagonal matrices. In all of these settings, the assumption of a low-rank structure is essential for efficient estimation and optimization in high-dimensional settings.

While the above applications all aim at recovering a low-rank matrix S , the observational—a.k.a sensing—mechanism that governs access to S comes in many declinations. For the purpose of applications, it is often sufficient to focus on linear measurements of the form $\langle S, A \rangle$

for some given sensing matrix A . This setup covers a wide variety of applications ranging from covariance sketching [Chen et al., 2015] and low-rank matrix completion [Candès and Plan, 2010, Recht et al., 2010] to phase retrieval [Fienup, 1978, Candès et al., 2013, 2015] and quantum state tomography [Gross et al., 2010]. Due to the flexibility of this formulation, new solutions to this problem can have many practical implications.

In this paper, we focus on a specific instantiation of this problem, where the measurement matrix $A = xx^\top$ is rank-one and positive semidefinite (PSD), so that $\langle xx^\top, S \rangle = x^\top Sx$. This important case of the low-rank matrix recovery problem has received significant attention over the past few decade [Cai and Zhang, 2015, Chen et al., 2015, Sanghavi et al., 2017, Li et al., 2019].

Assume that we observe a sample

$$y_i = x_i^\top S x_i, i = 1, \dots, n \quad (1.1)$$

where $S \in \mathbb{R}^{d \times d}$ is an unknown rank r PSD matrix and $x_1, \dots, x_n$ are i.i.d from some distribution. Our goal is to recover or estimate S from the pairs $(y_i, x_i), i = 1, \dots, n$. Throughout, we denote by $\mathbb{S}_+^d$ the set of $d \times d$ PSD matrices and $\mathbb{S}_{++}^d$ is the set of positive definite (PD) matrices.

Finding a low-rank matrix S subject to constraints (1.1) is a semidefinite program (SDP) that can be implemented in polynomial time using general-purpose solvers. Furthermore, the specific structure of this SDP may be leveraged to derive faster algorithms. Such solutions include the Burer-Monteiro approach to solving semidefinite programs [Burer and Monteiro, 2003] or nonconvex gradient descent methods for low-rank programs [Sanghavi et al., 2017]. Often, these approaches result in *nonconvex* optimization programs for which theoretical results are limited.

In this work, we take a principled approach to solving this problem by eliciting convexity using a specific geometry on the the space of PSD matrices. More precisely, we employ the Bures-Wasserstein (hereafter BW) geometry, which comes independently from optimal transport and quantum information theory [Bures, 1969, Bhatia et al., 2019]. This geometry allows us to solve the original problem by computing a *BW barycenter* [Agueh and Carlier, 2011, Álvarez-Esteban et al., 2016, Chewi et al., 2020,

[Altschuler et al., 2021]. In turn, we employ geodesic gradient descent on the BW manifold to compute said barycenter. We propose both full gradient and stochastic gradient based methods that are guaranteed to efficiently recover a low-rank matrix. These methods have low computational cost (per iteration complexity of $O(ndr)$ for gradient descent and $O(dr)$ for stochastic gradient descent), have minimal parameter tuning, are easily implemented, and show excellent practical performance. We demonstrate an example application of phase retrieval in Figure 1. In this set-up, BWgradient descent (BWGD) recovers the image faster than Wirtinger Flow (WF) [Candès et al., 2015], and BWGD needs no parameter tuning.



Figure 1: Recovered images after 140 iterations of Wirtinger Flow [Candès et al., 2015] and BWGD (our method). BWGD recovers a much sharper image within the same number of iterations; note that the iteration complexity of both methods is the same.

Main contributions. The main results of this paper are:

1. We prove that a barycenter of a certain distribution of rank-one Gaussians exactly recovers the target low-rank matrix $\mathbf{S}$.
2. With this connection, we give novel geodesic gradient descent and stochastic geodesic gradient descent algorithms for solving the low-rank PSD matrix recovery problem using existing first-order algorithms for computing BW barycenters.
3. Existing first-order algorithms for computing BW barycenters are only guaranteed to work for full rank distributions. Since our method considers barycenters of rank-one PSD matrices, we develop new theory and give a guarantee of local linear convergence in BW distance for the gradient descent method. We also discuss initialization of our method.
4. We demonstrate the competitive edge of our algorithms in a few experimental settings.

Related Work. Many methods have been proposed to solve variants of the matrix recovery problem. Original ideas for this problem trace back to linear systems theory, low-rank matrix completion, low-dimensional Euclidean embeddings, and image compression [Recht et al., 2010].

We focus here on rank-one projections of positive semidefinite matrices as in (1.1). This setup is either specifically considered or a special case of a large number of works including [Candès et al., 2015, Cai and Zhang, 2015, Chen et al., 2015, Zhong et al., 2015, Wang and Giannakis, 2016, Wang et al., 2017, Sanghavi et al., 2017, Li et al., 2019]. These methods can be clustered into two families. The first one aims at minimizing convex relaxations of an energy functional that often based on the nuclear norm [Cai and Zhang, 2015, Chen et al., 2015]. Such convex programs can be solved via standard solvers. Another family of methods directly implement the low-rank constraint into a nonconvex constraint [Li et al., 2019] thus manipulating candidate matrices with smaller representations and thereby boosting computational efficiency; see Chi et al. [2019] for an overview of such algorithms. The special case where $\mathbf{S}$ has rank one corresponds to the classical phase retrieval problem and has received much attention with dedicated algorithms [Fienup, 1978, Candès et al., 2015, Wang and Giannakis, 2016, Wang et al., 2017, Chi et al., 2019].

Note that the methods introduced in the present paper can be readily extended to the *covariance recovery* problem framed in Cai and Zhang [2015], and that is similar to estimation in a random effects model. Under this model, rather than (1.1), we observe $y_i = \langle \mathbf{x}_i, \mathbf{w}_i \rangle + \epsilon_i$, where $\mathbf{w}_i \sim N(\mathbf{0}, \mathbf{S})$, where $N(\mathbf{0}, \mathbf{S})$ is the centered Gaussian distribution on $\mathbb{R}^d$ with PSD covariance matrix $\mathbf{S}$. The goal here is to recover the covariance matrix $\mathbf{S}$ of the weight vectors from observations of $(y_i, \mathbf{x}_i)$, $i = 1, \dots, n$. This problem has natural connections to mixture of regressions models as well [De Veaux, 1989, Yi et al., 2014, Zhong et al., 2016, Sedghi et al., 2016].

In terms of the complexity of various methods, solving the semidefinite program using off-the-shelf solvers takes $O(nd^2 + d^3)$ complexity. Gradient descent on PSD matrices (see the initialization phase in [Tu et al., 2016]) can solve this with per-iteration complexity $O(nd^2)$, and one can prove linear convergence under certain assumptions. The most directly comparable methods are nonconvex gradient descent [Li et al., 2019], which have per-iteration complexity $O(ndr)$ and utilize a Burer-Monteiro factorization [Burer and Monteiro, 2003]. As we will see, our method also has per-iteration complexity of $O(ndr)$.

Finally, we mention several works dedicated to the computation of the nonconvex Wasserstein barycenter problem [Agueh and Carlier, 2011, Álvarez-Esteban et al., 2016, Zemel and Panaretos, 2019, Chewi et al., 2020,

[Altschuler et al., 2021]. No theoretical study in these works allows for rank deficiency.

Notation. Bold capital letters denote matrices while bold lower-case letters denote vectors. The Hilbert-Schmidt inner product is $\langle \cdot, \cdot \rangle$, and $\langle \cdot, \cdot \rangle_{\gamma_{\Sigma}} = \langle \cdot, \cdot \rangle_{\Sigma}$ is the Riemannian metric associated to the (fixed-rank) BW manifold at Σ . Their corresponding norms are written as $\| \cdot \|$ and $\| \cdot \|_{\Sigma}$, respectively. The orthogonal projection onto the column span of Σ is $P_{\Sigma} = P_{\text{Sp}(\Sigma)}$. Similarly, $P_{\Sigma}^{\perp}$ is the projection onto null-space of Σ (orthogonal complement of $\text{Sp}(\Sigma)$). The Dirac distribution at a point x denoted by δ_x .

Outline. We begin by outlining our approach to matrix recovery in Section 2. We then give the main theoretical results for our approach in Section 3. After this, we give experiments demonstrating these advantages in Section 4, and discuss the limitations of our work in Section 5.

2 THE BURES-WASSERSTEIN BARYCENTER APPROACH

Recall that we aim at find the rank r matrix $S \in \mathbb{S}_+^d$ given the observations (1.1). We will use the notation $y_i = \mathcal{A}(S)_i = \langle x_i x_i^{\top}, S \rangle$, $i = 1, \dots, n$. To recover a low-rank matrix S from these measurements, most past work has focused on some form of energy minimization. For example, some works have looked at convex nuclear norm minimization methods. Cai and Zhang [2015] and Chen et al. [2015] concurrently developed a nuclear norm minimization procedure that solves

$$\min_{\Sigma \in \mathbb{S}_+^d, \mathcal{A}(\Sigma)=\mathbf{y}} \text{Tr}(\Sigma), \quad (2.1)$$

where $\mathbf{y} = [y_1, \dots, y_n]^{\top}$ and $\mathcal{A}(\Sigma) = [\mathcal{A}(\Sigma)_1, \dots, \mathcal{A}(\Sigma)_n]^{\top}$. One can directly solve the semidefinite program using standard convex optimization packages. A Lagrangian formulation of (2.1) yields the energy $\frac{1}{2} \|\mathcal{A}(\Sigma) - \mathbf{y}\|_2^2 + \lambda \text{Tr}(\Sigma)$, which can also be minimized using a variety of methods.

In the following, we lay out our approach to the low-rank matrix recovery problem, which focuses on a new energy minimization procedure. In Section 2.1 we discuss the common nonconvex approaches to matrix recovery and outline our novel optimization program. Then, in Section 2.2, we discuss the BW barycenter problem, and show how it recovers solutions to the energy minimization we propose. After this, in Section 2.3 we outline the first-order algorithms for computing BW barycenters. We finish in Section 2.4 by discussing a regularization procedure that allows one to estimate higher rank proxies, from which it is possible to recover S .

2.1 Nonconvex Approaches for Matrix Recovery

Suppose that we know an upper bound for the rank of the underlying matrix S . We could utilize this information in a nonconvex optimization program such as

$$\min_{\Sigma \in \mathbb{S}_+^d, \text{rank}(\Sigma) \leq r} \frac{1}{2n} \|\mathcal{A}(\Sigma) - \mathbf{y}\|_2^2. \quad (2.2)$$

Without the rank restriction (i.e., $r = d$) this problem is in fact convex. For any fixed $r \leq d$, we can parameterize the rank r matrices in $\mathbb{S}_+^d$ by $UU^{\top}$, for $U \in \mathbb{R}^{d \times r}$, which is now commonly referred to as Burer-Monteiro factorization [Burer and Monteiro, 2003]. We thus define the set of PSD matrices of rank at most r using this factorization: $\mathbb{S}_+^{d,r} := \{\Sigma \in \mathbb{S}_+^d : \Sigma = UU^{\top}, U \in \mathbb{R}^{d \times r}\}$. With this parametrization, the matrix recovery problem in (2.2) is equivalent to

$$\min_{U \in \mathbb{R}^{d \times r}} \frac{1}{2} \|\mathcal{A}(UU^{\top}) - \mathbf{y}\|_2^2. \quad (2.3)$$

While past work has focused on these least squares formulations, there have not been many modifications of this energy. We propose the following modifications to the energies (2.2) and (2.3):

$$\min_{\Sigma \in \mathbb{S}_+^d, \text{rank}(\Sigma) \leq r} \frac{1}{2n} \|\sqrt{\mathcal{A}(\Sigma)} - \sqrt{\mathbf{y}}\|_2^2 \quad (2.4)$$

$$= \min_{U \in \mathbb{R}^{d \times r}} \frac{1}{2n} \|\sqrt{\mathcal{A}(UU^{\top})} - \sqrt{\mathbf{y}}\|_2^2, \quad (2.5)$$

where the square root is taken componentwise. As we demonstrate in the following sections, this problem has a natural solution as a BW barycenter.

2.2 The Bures-Wasserstein Barycenter Problem

To explain the connection of (2.4) to BW barycenters, we will first explain how BW space arises from the perspective of optimal transport [Villani, 2009]. Let $\mathcal{P}_2(\mathbb{R}^d)$ be the set of all measures on $\mathbb{R}^d$ with finite second moment. The 2-Wasserstein distance between measures μ and $\nu \in \mathcal{P}_2(\mathbb{R}^d)$ is defined by

$$W_2^2(\mu, \nu) = \inf_{\pi \in \Pi(\mu, \nu)} \mathbb{E}_{(x, y) \sim \pi} \|x - y\|_2^2, \quad (2.6)$$

where $\Pi(\mu, \nu)$ denotes the set of all couplings between μ and ν (i.e., the set of all joint distributions on $\mathbb{R}^d \times \mathbb{R}^d$ with marginals μ and ν). The 2-Wasserstein distance defines a metric over $\mathcal{P}_2(\mathbb{R}^d)$, and the resulting geodesic metric space is referred to as 2-Wasserstein space.

Let $\mathcal{N}(\mathbb{R}^d)$ denote the set of Gaussian distributions on $\mathbb{R}^d$, and $\mathcal{N}_0(\mathbb{R}^d)$ be the set of centered Gaussian distributions. Both are geodesically weakly convex subsets of 2-Wasserstein space, meaning there always exist 2-Wasserstein geodesics between points in these sets that are

contained within these sets. Letting $N(\mathbf{0}, \Sigma) \in \mathcal{N}_0(\mathbb{R}^d)$ denote the Gaussian distribution on $\mathbb{R}^d$ with mean zero and covariance matrix $\Sigma \in \mathbb{S}_+^d$, the 2-Wasserstein distance between $N(\mathbf{0}, \Sigma_0)$ and $N(\mathbf{0}, \Sigma_1) \in \mathcal{N}_0(\mathbb{R}^d)$ has the explicit form

$$W_2^2(N(\mathbf{0}, \Sigma_0), N(\mathbf{0}, \Sigma_1)) = \text{Tr}[\Sigma_0 + \Sigma_1 - 2(\Sigma_0^{1/2}\Sigma_1\Sigma_0^{1/2})^{1/2}]. \quad (2.7)$$

Notice that this is purely a function of the covariance matrices, and so the Wasserstein distance induces a distance metric on PSD matrices called the *Bures-Wasserstein distance* [Bhatia et al., 2019]. To refer to this distance over PSD matrices rather than the Gaussian distributions, we will write

$$d_{\text{BW}}(\Sigma_0, \Sigma_1) = W_2(N(\mathbf{0}, \Sigma_0), N(\mathbf{0}, \Sigma_1)). \quad (2.8)$$

More than just giving the set of PSD matrices a distance metric, this identification endows $\mathbb{S}_+^d$ with a natural Riemannian structure that it inherits from $(\mathcal{N}_0(\mathbb{R}^d), W_2)$.

The barycenter problem seeks to generalize the notion of averages to non-Euclidean spaces. In the 2-Wasserstein barycenter problem, one seeks a solution to

$$\min_{b \in \mathcal{P}_2(\mathbb{R}^d)} \frac{1}{2} \mathbb{E}_{\mu \sim Q} W_2^2(\mu, b), \quad (2.9)$$

where Q is a distribution over $\mathcal{P}_2(\mathbb{R}^d)$ with finite second moment, which we write as $Q \in \mathcal{P}_2(\mathcal{P}_2(\mathbb{R}^d))$. When Q is supported on Gaussians, the minimum is achieved on Gaussians [Knott and Smith, 1994, Agueh and Carlier, 2011, Álvarez-Esteban et al., 2016]. For $\mathcal{N}_0(\mathbb{R}^d)$, due to the identification in (2.7), this is equivalent to the Fréchet mean of PSD matrices on the BW manifold. Without loss of generality, if Q is supported on mean zero Gaussians, we will therefore think of Q as a distribution over PSD matrices.

We finally arrive at the connection between low-rank PSD matrix recovery and BW barycenters. The following proposition connects the barycenter problem (2.9) when $Q = Q_n = \frac{1}{n} \sum_{i=1}^n \delta_{N(\mathbf{0}, y_i \mathbf{x}_i \mathbf{x}_i^\top)}$ to our new low-rank matrix recovery program (2.4), provided that $\frac{1}{n} \sum_{i=1}^n \mathbf{x}_i \mathbf{x}_i^\top = \mathbf{I}$. This proposition indicates that we can recover the matrix S by solving a Wasserstein barycenter problem.

Proposition 1. *If $\frac{1}{n} \sum_i \mathbf{x}_i \mathbf{x}_i^\top = \mathbf{I}$, then*

$$\begin{aligned} \argmin_{\Sigma \in \mathbb{S}_+^d} \frac{1}{2n} \|\sqrt{\mathcal{A}(\Sigma)} - \sqrt{\mathbf{y}}\|_2^2 \\ = \argmin_{\Sigma \in \mathbb{S}_+^d} \frac{1}{2n} \sum_{i=1}^n d_{\text{BW}}^2(\Sigma, y_i \mathbf{x}_i \mathbf{x}_i^\top). \end{aligned} \quad (2.10)$$

To make this result practical, we cannot assume in general that $\frac{1}{n} \sum_i \mathbf{x}_i \mathbf{x}_i^\top = \mathbf{I}$. If we instead encounter a case where $\frac{1}{n} \sum_{i=1}^n \mathbf{x}_i \mathbf{x}_i^\top = C_n$, where C_n is the PD sample covariance matrix of the vectors $\mathbf{x}_1, \dots, \mathbf{x}_n$, then the

transformation $\mathbf{x}_i \mapsto C_n^{-1/2} \mathbf{x}_i$ outputs vectors with identity covariance (i.e., $\frac{1}{n} \sum_{i=1}^n C_n^{-1/2} \mathbf{x}_i \mathbf{x}_i^\top C_n^{-1/2} = \mathbf{I}$). We are able use this fact to recover the matrix S , as we show in the following proposition.

Proposition 2. *Let $C_n = \frac{1}{n} \sum_{i=1}^n \mathbf{x}_i \mathbf{x}_i^\top \in \mathbb{S}_{++}^d$. Then*

$$\begin{aligned} \argmin_{\Sigma \in \mathbb{S}_+^d} \frac{1}{2n} \|\sqrt{\mathcal{A}(C_n^{-1/2} \Sigma C_n^{-1/2})} - \sqrt{\mathbf{y}}\|_2^2 \\ = \argmin_{\Sigma \in \mathbb{S}_+^d} \frac{1}{2n} \sum_{i=1}^n d_{\text{BW}}^2(\Sigma, y_i C_n^{-1/2} \mathbf{x}_i \mathbf{x}_i^\top C_n^{-1/2}), \end{aligned} \quad (2.11)$$

and $C_n^{1/2} S C_n^{1/2}$ is a solution to both problems.

Notice that one can recover S from $C_n^{1/2} S C_n^{1/2}$ as $S = C_n^{-1/2} C_n^{1/2} S C_n^{1/2} C_n^{-1/2}$. Thus, in the sample setting where C_n is not exactly the identity and assuming we can solve the barycenter problem, we envision a two stage procedure: 1) recover the barycenter Σ_n of $y_1 C_n^{-1/2} \mathbf{x}_1 \mathbf{x}_1^\top C_n^{-1/2}, \dots, y_n C_n^{-1/2} \mathbf{x}_n \mathbf{x}_n^\top C_n^{-1/2}$, and 2) transform the barycenter by $C_n^{-1/2} \Sigma_n C_n^{-1/2}$ to find S .

The whitening step can be efficiently computed since we can use any linear transformation L^{-1} such that $L^{-1} \frac{1}{n} \sum_{i=1}^n \mathbf{x}_i \mathbf{x}_i^\top L^{-1\top} = \mathbf{I}$. For example, with the Cholesky factorization $C_n = L L^\top$, we can solve the equations $L \mathbf{z}_i = \mathbf{x}_i$ for $\mathbf{z}_i$, and these satisfy $\frac{1}{n} \sum_i \mathbf{z}_i \mathbf{z}_i^\top = \mathbf{I}$.

Remark 3. *The assumption that C_n is full rank can be relaxed by using the pseudoinverse of C_n rather than the inverse. However, for our later theorems we assume $n = O(dr)$ i.i.d. samples from a $N(\mathbf{0}, \mathbf{I})$ distribution in order to obtain a restricted isometry property. In this case, the matrix C_n will be full rank with probability 1.*

The connections established by Propositions 1 and 2 enables the development of novel methods for the matrix recovery problem, since we can solve a specific Wasserstein barycenter problem rather than the original matrix recovery problem (2.4). In other words, any methods that solve this Wasserstein barycenter problem could be used to solve the matrix recovery problem. Since barycenters are geometric notions of averages, this naturally leads to the development of novel geometric methods for matrix recovery.

2.3 Algorithms for Barycenters

The primary way to compute BW barycenters involves Riemannian gradient descent [Álvarez-Esteban et al., 2016, Chewi et al., 2020, Altschuler et al., 2021]. Following the results in the last section, we wish to find the BW barycenter of the matrices $\mathbf{X}_i = y_i C_n^{-1/2} \mathbf{x}_i \mathbf{x}_i^\top C_n^{-1/2}$, $i = 1, \dots, n$. In other words, we seek to minimize the energy function $F : \mathbb{S}_+^d \rightarrow \mathbb{R}$ given by

$$F(\Sigma) = \frac{1}{2n} \sum_{i=1}^n d_{\text{BW}}^2(\mathbf{X}_i, \Sigma). \quad (2.12)$$

For ease of notation, we will write F as an expectation over $y_i C_n^{-1/2} \mathbf{x}_i \mathbf{x}_i^\top C_n^{-1/2} \sim Q$, and whether or not Q is a discrete distribution will be made clear from context.

The gradient of F at full rank Σ_0 is

$$\nabla F(\Sigma_0) = \mathbf{I} - \mathbb{E}_{\mathbf{x} \sim Q} \frac{\mathbf{X}}{\sqrt{\text{Tr}(\mathbf{X} \Sigma_0)}} =: \mathbf{I} - \tilde{T}(\Sigma_0),$$

where for convenience we have defined the quantity $\tilde{T}(\Sigma_0)$. At low-rank Σ_0 , this is a subgradient [Clarke, 1990]. We note that this holds for general measures Q where we can differentiate under the integral.

Wasserstein gradient descent uses the gradient to determine a geodesic along which to move. In Wasserstein space, this is the “pushforward” direction in a base measure is transported. For a complete description of Wasserstein gradient descent, see [Álvarez-Esteban et al., 2016, Zemel and Panaretos, 2019, Chewi et al., 2020, Altschuler et al., 2021], and for a more complete discussion of BW geometry, the reader should consult Bhatia et al. [2019]. We have included a discussion of the geodesic structure of BW space in the appendix. For our purposes, we consider BWGD (BWGD) with step size η_k :

$$\Sigma_{k+1} = (\mathbf{I} - \eta_k \nabla F(\Sigma_k)) \Sigma_k (\mathbf{I} - \eta_k \nabla F(\Sigma_k)). \quad (2.13)$$

When $\eta_k = 1$, this corresponds to the fixed point iteration of Álvarez-Esteban et al. [2016] (see also Appendix B). Note that this is easy to extend to the stochastic setting: if we observe a stochastic gradient $\mathbf{G}_k$ rather than $\nabla F(\Sigma_k)$, then the BW stochastic gradient descent (BWSGD) iteration would use $\mathbf{G}_k$ in place of $\nabla F(\Sigma_k)$. For example, one common variant of such a stochastic gradient method uses

$$\mathbf{G}_k = \mathbf{I} - \frac{\mathbf{X}_k}{\text{Tr}(\mathbf{X}_k \Sigma_k)^{1/2}}, \quad (2.14)$$

where $k = 1, \dots, n$. In other words, this variant of stochastic gradient descent passes over each sample one at a time. At each point in time, we take a gradient with respect to that sample alone and move in the minus gradient direction.

To save computational time and to allow efficient computation in the low-rank case, we modify the BWGD iteration (2.13) and the BWSGD iteration to instead operate on factorized matrices (i.e., on the formulation seen in (2.5)). This means that instead of storing the sequence Σ_k for $k \in \mathbb{N}$, we instead store the sequence

$$\mathbf{U}_{k+1} = (\mathbf{I} - \eta_k \nabla F(\mathbf{U}_k \mathbf{U}_k^\top)) \mathbf{U}_k. \quad (2.15)$$

We note that if Σ_k is low-rank in (2.13), then the update in (2.15) is equivalent to (2.13). In particular, it is not hard to show that if $\mathbf{U}_k \mathbf{U}_k^\top = \Sigma_k$, then $\mathbf{U}_{k+1} \mathbf{U}_{k+1}^\top = \Sigma_{k+1}$. In the same way as before, stochastic gradient methods naturally extend to the low-rank setting. This update corresponds to a geodesic gradient descent update over the fixed rank BW manifold [Massart and Absil, 2020].

When $Q = Q_n$ is a discrete distribution, which is the main focus of this paper, the BWGD updates can be computed in $O(ndr)$ time. This follows from the fact that we can rewrite the update in (2.15) as

$$\mathbf{U}_{k+1} = (1 - \eta_k) \mathbf{U}_k + \frac{\eta_k}{n} \sum_{i=1}^n \frac{\mathbf{x}_i \mathbf{x}_i^\top \mathbf{U}_k}{\|\mathbf{U}_k^\top \mathbf{x}_i\|_2}. \quad (2.16)$$

In the case of single-sample streaming BWSGD, which uses the gradient in (2.14), the updates take $O(dr)$ time. We also note that both the BWGD and BWSGD iterations maintain the rank of the updated matrix for all $\eta_k \in [0, 1]$. For $\eta_k = 1$, the rank is maintained for BWGD as long as $\text{rank}(\sum_i \mathbf{x}_i \mathbf{x}_i^\top) = d$.

2.4 Regularization through Perturbed Gradient Descent

As we demonstrate later, our current theorem only applies for $r \geq 3$, despite the fact that we observe good performance in practice for the iteration (2.15) when $r = 1$ and $r = 2$. We comment that a modified version of our method can yield results in the cases of $r = 1$ or $r = 2$.

To demonstrate how this works, consider the observation model $y_i = \langle \mathbf{x}_i \mathbf{x}_i^\top, \mathbf{S} \rangle$ for a rank r matrix $\mathbf{S}$ and $i = 1, \dots, n$. If we choose an arbitrary rank r' matrix Δ , we could instead try to recover the rank at most $r + r'$ matrix $\mathbf{S} + \Delta$ from the observations $\tilde{y}_i = \langle \mathbf{x}_i \mathbf{x}_i^\top, \mathbf{S} + \Delta \rangle = \langle \mathbf{x}_i \mathbf{x}_i^\top, \mathbf{S} \rangle + \langle \mathbf{x}_i \mathbf{x}_i^\top, \Delta \rangle$, which can be computed by computing and adding the factors $\langle \mathbf{x}_i \mathbf{x}_i^\top, \Delta \rangle$ to the observations y_i . If we then find the barycenter of $\tilde{y}_i \mathbf{x}_i \mathbf{x}_i^\top$ (provided that it is unique, which we prove later under conditions on the vectors $\mathbf{x}_i$), then this would recover $\mathbf{S} + \Delta$! Since this is true for any such Δ , the user can pick Δ beforehand and then recover $\mathbf{S}$ from simple subtraction: $\mathbf{S} = (\mathbf{S} + \Delta) - \Delta$. In brief, every rank r recovery problem can be solved by instead first solving the rank $r + r'$ problem to find $\mathbf{S} + \Delta$ and then subtracting off the perturbation factor Δ .

3 THEORETICAL RESULTS

We now discuss the main theoretical results of this paper. We make the following assumption on our measurement model in (1.1).

Assumption 1. *We observe data from the model (1.1) with $\mathbf{x}_i \stackrel{i.i.d.}{\sim} N(\mathbf{0}, \mathbf{I})$. The underlying matrix $\mathbf{S}$ is rank r and satisfies $m \leq \lambda_r(\mathbf{S}) \leq \lambda_1(\mathbf{S}) \leq M$.*

The assumption of Gaussianity can be weakened to sub-Gaussianity, but we state the result for Gaussian vectors since it simplifies some of the later proofs (for example, the proof of Euclidean smoothness in Appendix C.2.3). The sub-Gaussianity is essential for an ℓ_2/ℓ_1 restricted isometry property [Chen et al., 2015, Cai and Zhang, 2015] that we need to hold (see Appendix C.2.1).

Our main result on matrix recovery with gradient descent for BW barycenters is given in the following theorem.

Theorem 4. Suppose that we observe $y_i = \langle \mathbf{x}_i \mathbf{x}_i^\top, \mathbf{S} \rangle$, $i = 1, \dots, n$, where $\mathbf{x}_i$ and $\mathbf{S}$ satisfy Assumption 1. Suppose further that $r = \text{rank}(\mathbf{S}) \geq 5$, and let $\mathbf{S}_n = \mathbf{C}_n^{1/2} \mathbf{S} \mathbf{C}_n^{1/2}$. Then, for constants c_1, c_2 , if $n \gtrsim dr$, with probability at least $1 - \exp(-c_2 n) - 1/n$,

1. $\mathbf{S}_n$ is the unique global minimizer of F over $\mathbb{S}_+^d$.
2. Let $\mathbf{U}_0$ be the initial iterate of BWGD. If $\lambda_1(\mathbf{U}_0 \mathbf{U}_0^\top) \leq M$, and $F(\mathbf{U}_0 \mathbf{U}_0^\top) - F(\mathbf{S}_n) \leq \frac{c_1^5 m^8 (r-2)^2 \beta^{10}}{6^5 M^{15/2} d^3 r^2}$ for an $\mathbf{S}$ dependent constant β , then BWGD with step size 1 satisfies the following bound for $\mathcal{C} = O(c_1 m / M^{3/2})$.

$$d_{\text{BW}}^2(\mathbf{U}_k \mathbf{U}_k^\top, \mathbf{C}_n^{1/2} \mathbf{S} \mathbf{C}_n^{1/2}) \leq (1 - \mathcal{C})^k (F(\mathbf{U}_0 \mathbf{U}_0^\top) - F(\mathbf{S})). \quad (3.1)$$

This is the first result for convergence in Bures-Wasserstein distance in the literature. The constant $\mathcal{C}$ is the local strong geodesic convexity constant seen in (C.50). We note that this amounts to showing that our new energy given by (2.4) has a positive definite Hessian in a neighborhood around $\mathbf{S}$. A couple of remarks are in order to discuss two issues that arise with the analysis of our method: initialization and the rank constraint on $\mathbf{S}$. In practice, we observe convergence to $\mathbf{S}$ from random initialization, but do not currently have a proof of this fact.

Remark 5 (Initialization). We note that the convergence bound in Theorem 4 is local. A few procedures can be used to initialize in the correct neighborhood. First, under the Gaussian assumption, $\mathbb{E}_y(\mathbf{x} \mathbf{x}^\top - \mathbf{I}) = 2\mathbf{S}$, and so one could use a rank r approximation of $\frac{1}{2n} \sum_{i=1}^n y_i (\mathbf{x}_i \mathbf{x}_i^\top - \mathbf{I})$ to initialize the gradient descent. This approximation can be computed in $O(ndr)$ time with the power method. A finite sample approximation result follows from concentration of the fourth moment tensor, see [Diakonikolas et al., 2019, Theorem 4.13] for details. Another path to initialization would be to use gradient descent on full rank PSD matrices in the first stage since, as we show in the Appendix, the function (2.10) is convex over PD matrices. This would allow one to recover a good approximation to $\mathbf{S}$, and take a rank r approximation of it, and then run BWGD from there. The downside of this method is that it takes complexity $O(nd^2)$ to compute, but it would not need concentration of the fourth moment tensor. Also, the gradient descent procedure with overspecified rank converges sublinearly.

Remark 6 (Rank Constraint). The result holds for $\text{rank}(\mathbf{S}) \geq 5$. This is due to a smoothness bound that requires bounds on the expectation and variance of $\mathbb{E} \sqrt{1 + (x_{r+1}^2 + \dots + x_{2r}^2) / (x_1^2 + \dots + x_r^2)}$ for $x_i \stackrel{i.i.d.}{\sim} N(0, 1)$ random variables. In order for $\mathbb{E} 1 / (x_1^2 + \dots + x_r^2)^2$ to exist, we need $r \geq 3$, and in order for $1 / (x_1^2 + \dots + x_r^2)^2$

to have finite variance, we need $r \geq 5$. In practice, and as we show in the experiments, the method still succeeds much of the time for $r = 1, \dots, 5$. To have a theoretically guaranteed method in these settings, one can use the regularized BWGD method of Section 2.4.

We give a sketch of the proof of Theorem 4.

- We first show that the energy F in (2.12) is Euclidean strongly convex over $\mathbb{S}_+^d$ and that $\mathbf{S}_n$ is the unique minimizer with high probability. This result uses the ℓ_2/ℓ_1 -RIP condition of Chen et al. [2015] (or restricted uniform boundedness condition of Cai and Zhang [2015]).
- We then prove that the gradient is locally $1/2$ -Hölder continuous: $\|\nabla F(\mathbf{\Sigma}) - \nabla F(\mathbf{S}_n)\|_F \lesssim \|\mathbf{\Sigma} - \mathbf{S}_n\|_F^{1/2}$ for rank r matrices $\mathbf{\Sigma}$.
- Using this smoothness result, we are able to show that F is locally-BW-geodesically convex around $\mathbf{S}_n$ in $\mathbb{S}_+^{d,r}$.
- Local strong geodesic convexity along with a geodesic smoothness result yields the local linear convergence result via standard optimization arguments (see, for example, Bubeck [2015], Zhang and Sra [2016]).

We also include the following theorem on the convergence of BWSGD. This guarantees a slow rate of convergence for BWSGD with gradient given by (2.14).

Theorem 7. Suppose that we observe $y_i = \langle \mathbf{x}_i \mathbf{x}_i^\top, \mathbf{S} \rangle$, $i = 1, \dots, n$, where $\mathbf{x}_i$ and $\mathbf{S}$ satisfy Assumption 1. Suppose that we run single sample streaming BWSGD, which uses gradient (2.14), for n iterations with step size $1/\sqrt{n}$. Then, for a constant c_1 , if $n \gtrsim dr$, with probability at least $1 - \exp(-c_1 n) - 1/n$,

$$\min_{k=1, \dots, n} \mathbb{E} \|\nabla F(\mathbf{\Sigma}_k)\|_{\mathbf{\Sigma}_k}^2 = O(n^{-1/2}), \quad (3.2)$$

where $\|\cdot\|_{\mathbf{\Sigma}}^2 = \mathbb{E}_{\mathbf{z} \sim N(\mathbf{0}, \mathbf{\Sigma})} \|\cdot\|_{\mathbf{z}}^2$ is the norm induced by the BW Riemannian metric.

This states that the best iterate of the BWSGD sequence outputs an approximate stationary point with respect to the norm induced by the BW Riemannian metric. However, we cannot guarantee that this stationary point is the global minimum. We comment that we do not tend to run into spurious local minima in practice, and future work should go into studying this fact. While we do not have a justification for this, we give a theorem on the $r = 1$ case, where we show that the energy function has no local minima in the asymptotic limit. The proof of this theorem is left to the supplementary material.

Theorem 8. Consider the observation model with the rank one matrix $\mathbf{S} = \mathbf{v} \mathbf{v}^\top$ and $y = \langle \mathbf{x} \mathbf{x}^\top, \mathbf{S} \rangle$. Then, the only

fixed points of the population version of (2.15) (which corresponds to gradient descent on (2.12) with the sum replaced by an integral) are $\mathbf{v}$ or orthogonal to $\mathbf{v}$. In particular, this implies that population BWGD from any initialization such that $\mathbf{u}_0 \not\perp \mathbf{v}$ converges to $\mathbf{v}$.

For the case of general r , we believe that a similar result holds, although we have not yet been able to show it. Furthermore, we also believe that these results can be extended to high probability results in the finite sample case.

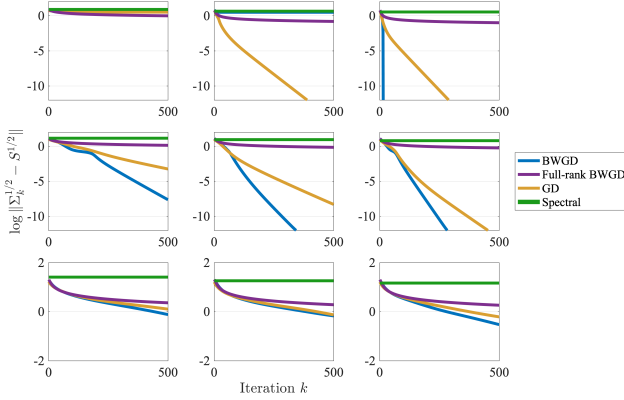


Figure 2: Convergence of BWGD (2.15) compared with full rank BWGD as well as the Euclidean GD [Li et al., 2019]. Here, $d = 32$, and down the columns we use $r = 1, 4, 16$, respectively. Across the rows we vary the number of points, with $n = 3dr$, $n = 10dr$, and $n = 20dr$, respectively. As we can see, for low to moderate ranks, the BWGD method converges much faster than the standard GD method.

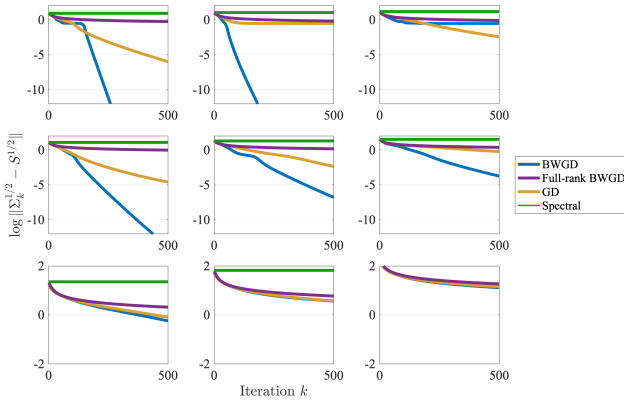


Figure 3: Convergence of BWGD (2.15) compared with full rank BWGD as well as the Euclidean GD [Li et al., 2019]. Here, $d = 32$, and down the columns we use $r = 2, 4, 16$, respectively. Across the rows we vary scale factor α , with $\alpha = 0, 1, 2$, respectively. As we can see, the BWGD method converges much faster than the standard GD method.

4 EXPERIMENTS

Here we present some numerical simulations that demonstrate the effectiveness of our method. All experiments were run in MATLAB on a 2020 Macbook Pro with a 2Ghz Quad-Core Intel Core i5 CPU and 16 GB of RAM.

4.1 Synthetic Experiments

We begin with some experiments on generated datasets to better understand the performance of our method. Since our main results focus on the performance of gradient descent rather than stochastic gradient descent, we focus on its performance here.

The methods we compare are rank r BWGD, full rank BWGD, the nonconvex Euclidean Gradient Descent (GD) of Li et al. [2019], and a spectral method, which takes a low-rank approximation to $\mathbf{S}_n = \frac{1}{2n} \sum_{i=1}^n y_i (\mathbf{x}_i \mathbf{x}_i^\top - \mathbf{I})$ since $\frac{1}{2} \mathbb{E} y^2 (\mathbf{x} \mathbf{x}^\top - \mathbf{I}) = \mathbf{S}$ [Sedghi et al., 2016]. As an error metric, we compute $\|\Sigma_k^{1/2} - \mathbf{S}^{1/2}\|_F = \Theta(d_{\text{BW}}(\Sigma_k, \mathbf{S}))$. We do not compare with other matrix recovery algorithms since they are not guaranteed to work in the symmetric rank one projection setting. Also, we find the comparison with GD to be the most relevant, since BWGD and GD have comparable complexity and are both first-order methods.

In our first experiment, we test the accuracy of the methods over time as we vary the rank of $\mathbf{S}$ and the sample size. Since the spectral method is not iterative, we include it as a horizontal line. Figure 2 displays the results of this experiment. Across the rows we vary the number of points and down the columns we vary the rank of $\mathbf{S}$. The fixed dimension is $d = 32$, the ranks from top to bottom are $r = 1, 4, 16$, and across the rows the number of points are $3dr$, $10dr$, and $20dr$, respectively. For each frame, we generate 20 datasets and run the four methods on them. All methods are run with random initialization, where the entries of $\mathbf{U}_0$ are i.i.d. $N(0, 1)$. For $r = 1$, we see that rank r BWGD succeeds once the number of points is sufficiently large. Furthermore, the convergence when it is successful is extremely fast. For moderate ranks, rank r BWGD converges faster than the previous GD method of Li et al. [2019]. We note that in this figure, there are some conflating factors that affect the convergence. Most notably, the conditioning in this problem (m versus M) gets worse as rank increases since $\mathbf{S} = \mathbf{V} \mathbf{V}^\top$ and the entries of $\mathbf{V}$ are i.i.d. $N(0, 1)$. Better conditioning leads to faster for higher ranks, see Figure 7 in the appendix, where the columns of $\mathbf{V}$ are orthogonal and length $\sqrt{d}$. Also, there is the additional factor of c_1 in the strong convexity constant in Theorem 3, which is the RIP constant. While in theory the constant is universal once the number of points is sufficiently large, one expects some scaling in terms of rank for finite n .

In Figure 3, we examine the performance of the methods under varying conditioning of $\mathbf{S}$. We set $d = 32$ and $n = 5dr$ and vary r as well as the conditioning of the matrix $\mathbf{S}$. Here, $\mathbf{S} = \mathbf{V} \text{diag}(r^\alpha, (r-1)^\alpha, \dots, 1^\alpha) \mathbf{V}^\top$, where the entries of $\mathbf{V}$ are i.i.d. $N(0, 1)$ and α is a scale factor. Figure 3 displays the results on 20 randomly generated datasets per frame. The rows correspond to $r = 2, 4$, and 16 respectively. The columns correspond to $\alpha = 0, 1$ and 2 respectively. Rank r BWGD performs uniformly well throughout.

In Figure 4 we show the dependence on sample size. Here, $d = 64$ and r is varied from 1 to 20. The error of BWGD and GD after 200 iterations is shown for sample sizes of $d, 2d, \dots, 20d$ for each value of r . For each r, n pair, we generate 20 datasets and compute the average error across them. The color indicates the average error value across these datasets. As we see, BWGD performs the best out of these methods. GD with linesearch is also competitive, but is more time consuming.

For the final synthetic experiment in Figure 5, we demonstrate the scalability of BWGD to higher dimensions. We note that it scales much better in terms of actual computational time when compared with the Euclidean GD method of Li et al. [2019].

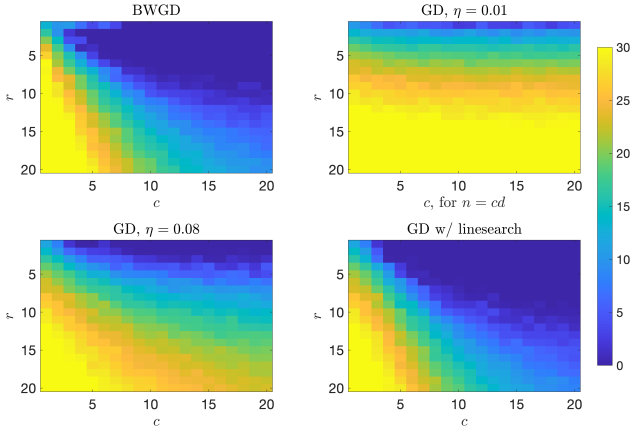


Figure 4: Convergence of BWGD (2.15) with the GD algorithm of [Li et al., 2019]. Here, $d = 64$, and down the columns of each inset we use $r = 1, \dots, 20$, respectively. Across the rows we vary the scale factor c , where $n = cd$. As we can see, the BWGD method converges much faster than the standard GD method with fixed step size, and performs on par with GD with linesearch. We note that the GD method with line search is much more computationally intensive than BWGD, as is displayed in Figure 5.

4.2 Real Data Experiment: Phase Retrieval

Here we give the details of the experiment displayed in Figure 1. We replicate the phase retrieval experiments in Candès et al. [2015]. In this paper, the authors study the

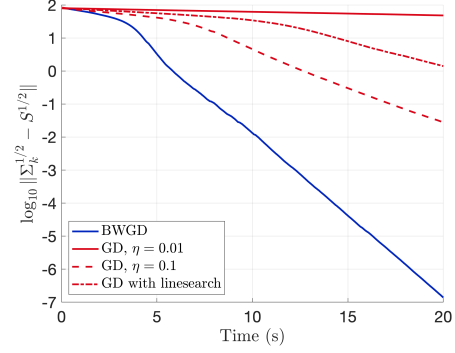


Figure 5: Convergence of BWGD (2.15) and GD [Li et al., 2019]. Here, $d = 512$, $r = 8$, and $n = 3rd$. As we can see, the BWGD method converges much faster than the standard GD method with any choice of step sizes. The per-iteration convergence rate for BWGD is comparable to GD with linesearch, but doing linesearch increases computational burden.

nonconvex Wirtinger Flow algorithm. Since phase retrieval is non equivalent to the recovery problem in (1.1) with a rank one complex $\mathbf{S}$, we can apply our algorithm in this setting.

Here, we use $i = \sqrt{-1}$. We use an image of Denali National Park, which is denoted by the 2-dimensional array $\mathbf{J}$ for each color band. In our simulated acquisition model, for $m = (u, v, \ell)$, we acquire data of the form

$$y_m = \left| \sum_{j=1, k=1}^{j=d_1, k=d_2} \mathbf{J}_{jk} \bar{d}_\ell(j, k) e^{-i2\pi(ju+kv)} \right|^2. \quad (4.1)$$

Here, $d_\ell(j, k) \sim b_1 b_2$, where b_1 is uniform on $\{1, -1, i, -i\}$, and b_2 takes values $\sqrt{2}/2$ with probability $4/5$ and $\sqrt{3}$ with probability $1/5$. The goal is to recover the image $\mathbf{J}$ from these measurements. We note that these measurements can be equivalently written as $y_m = \mathbf{F}_{u,v,\ell}^b \mathbf{J}^b \mathbf{J}^{b\top} \mathbf{F}_{u,v,\ell}^b$, where $\cdot^b$ denotes the vectorization operation and $\mathbf{F}$ is a matrix with entries $\bar{d}_\ell(j, k) e^{-i2\pi(ju+kv)}$. This notation makes clear the connection with the original matrix recovery problem in (1.1). We display the errors versus runtimes in Figure 6. Despite the fact that this example has $r = 1$, BWGD still efficiently recovers the underlying image, even though our current theory only works for $r \geq 5$.

We also give additional phase retrieval experiments on more images in Appendix D.2.

5 LIMITATIONS

There are a few notable limitations for the current work. First of all, the theory does not directly extend to the cases of $r = 1, \dots, 4$, and we must resort instead to the regularized methods discussed in Section 2.4. It is unclear whether

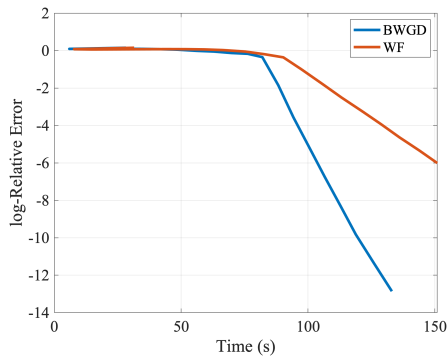


Figure 6: Error vs Runtime of BWGD (our method) and Wirtinger Flow Candès et al. [2015] on the phase retrieval experiment in Figure 1. As we can see, BWGD recovers the underlying image faster. Both methods are initialized with the power method as is done in [Candès et al., 2015].

or not this is a limitation of the methods or the analysis, although experiments indicate that the method works well for small ranks in practice. Second, while our experiments indicate well-behaved energy landscapes, we are only able to prove local convergence results. Third, while we give the first optimization rates in terms of BW distance for this problem, our constants are not optimized. Fourth, we assume that one knows the rank of S in advance. This is not a large issue since our framework allows for one to pick any $r' \geq r$ and still recover r , albeit at a slower rate.

6 CONCLUSION

We have presented a novel connection between the BW barycenter problem and low-rank matrix recovery from rank one measurements. We show that a novel energy minimization problem coincides with the barycenter minimization. This connection allows us to extend algorithms from the barycenter problem to the matrix recovery problem, giving new algorithms for recovering low-rank PSD matrices. Our methods are guaranteed to have local linear convergence in BW distance, which is a stronger guarantee than existing methods.

This work leaves open many unexplored directions. For example, it would be interesting to show that the energy landscape does not exhibit spurious local minima. Beyond this, it would also be interesting to know when and how one can go beyond the RIP assumption, which currently relies on sub-Gaussianity of the sensing vectors. Finally, it would be interesting to connect these ideas to optimization landscapes for neural networks [Zhong et al., 2017, Du et al., 2019, Li et al., 2019].

7 ACKNOWLEDGEMENTS

P. Rigollet is supported by NSF grants IIS-1838071, DMS-2022448, and CCF-2106377.

References

- Martial Agueh and Guillaume Carlier. Barycenters in the Wasserstein space. *SIAM J. Math. Anal.*, 43(2):904–924, 2011.
- Jason Altschuler, Sinho Chewi, Patrik R Gerber, and Austin Stromme. Averaging on the Bures-Wasserstein manifold: dimension-free convergence of gradient descent. *Advances in Neural Information Processing Systems*, 34: 22132–22145, 2021.
- Pedro C Álvarez-Esteban, E Del Barrio, JA Cuesta-Albertos, and C Matrán. A fixed-point approach to barycenters in Wasserstein space. *Journal of Mathematical Analysis and Applications*, 441(2):744–762, 2016.
- Luigi Ambrosio, Nicola Gigli, and Giuseppe Savaré. *Gradient flows in metric spaces and in the space of probability measures*. Lectures in Mathematics ETH Zürich. Birkhäuser Verlag, Basel, second edition, 2008.
- Rajendra Bhatia, Tanvi Jain, and Yongdo Lim. On the Bures-Wasserstein distance between positive definite matrices. *Expo. Math.*, 37(2):165–191, 2019.
- C. L. Blake and C. J. Merz. UCI repository of machine learning databases, 1998. URL <https://archive.ics.uci.edu/ml/datasets/abalone>.
- Sébastien Bubeck. Convex optimization: Algorithms and complexity. *Foundations and Trends® in Machine Learning*, 8(3-4):231–357, 2015.
- Samuel Burer and Renato DC Monteiro. A nonlinear programming algorithm for solving semidefinite programs via low-rank factorization. *Mathematical Programming*, 95(2):329–357, 2003.
- Donald Bures. An extension of Kakutani’s theorem on infinite product measures to the tensor product of semifinite w^* -algebras. *Transactions of the American Mathematical Society*, 135:199–212, 1969.
- T Tony Cai and Anru Zhang. ROP: Matrix recovery via rank-one projections. *The Annals of Statistics*, 43(1): 102–138, 2015.
- Emmanuel J Candès and Yaniv Plan. Matrix completion with noise. *Proceedings of the IEEE*, 98(6):925–936, 2010.
- Emmanuel J Candès and Benjamin Recht. Exact matrix completion via convex optimization. *Foundations of Computational mathematics*, 9(6):717–772, 2009.
- Emmanuel J Candès and Terence Tao. The power of convex relaxation: Near-optimal matrix completion. *IEEE*

- Transactions on Information Theory*, 56(5):2053–2080, 2010.
- Emmanuel J Candès, Justin Romberg, and Terence Tao. Robust uncertainty principles: Exact signal reconstruction from highly incomplete frequency information. *IEEE Transactions on information theory*, 52(2):489–509, 2006.
- Emmanuel J Candès, Xiaodong Li, Yi Ma, and John Wright. Robust principal component analysis? *Journal of the ACM (JACM)*, 58(3):1–37, 2011.
- Emmanuel J Candès, Thomas Strohmer, and Vladislav Voroninski. Phaselift: Exact and stable signal recovery from magnitude measurements via convex programming. *Communications on Pure and Applied Mathematics*, 66(8):1241–1274, 2013.
- Emmanuel J Candès, Yonina C Eldar, Thomas Strohmer, and Vladislav Voroninski. Phase retrieval via matrix completion. *SIAM review*, 57(2):225–251, 2015.
- Venkat Chandrasekaran, Sujay Sanghavi, Pablo A Parrilo, and Alan S Willsky. Sparse and low-rank matrix decompositions. *IFAC Proceedings Volumes*, 42(10):1493–1498, 2009.
- Yuxin Chen, Yuejie Chi, and Andrea J Goldsmith. Exact and stable covariance estimation from quadratic sampling via convex programming. *IEEE Transactions on Information Theory*, 61(7):4034–4059, 2015.
- Sinho Chewi, Tyler Maunu, Philippe Rigollet, and Austin J Stromme. Gradient descent algorithms for Bures-Wasserstein barycenters. In *Conference on Learning Theory*, pages 1276–1304. PMLR, 2020.
- Yuejie Chi, Yue M Lu, and Yuxin Chen. Nonconvex optimization meets low-rank matrix factorization: An overview. *IEEE Transactions on Signal Processing*, 67(20):5239–5269, 2019.
- Frank H Clarke. *Optimization and nonsmooth analysis*. SIAM, 1990.
- Richard D De Veaux. Mixtures of linear regressions. *Computational Statistics & Data Analysis*, 8(3):227–245, 1989.
- Ilias Diakonikolas, Gautam Kamath, Daniel Kane, Jerry Li, Ankur Moitra, and Alistair Stewart. Robust estimators in high-dimensions without the computational intractability. *SIAM Journal on Computing*, 48(2):742–864, 2019.
- David L Donoho. Compressed sensing. *IEEE Transactions on information theory*, 52(4):1289–1306, 2006.
- Simon Du, Jason Lee, Haochuan Li, Liwei Wang, and Xiyu Zhai. Gradient descent finds global minima of deep neural networks. In *International conference on machine learning*, pages 1675–1685. PMLR, 2019.
- Maryam Fazel. *Matrix rank minimization with applications*. PhD thesis, PhD thesis, Stanford University, 2002.
- James R Fienup. Reconstruction of an object from the modulus of its fourier transform. *Optics letters*, 3(1):27–29, 1978.
- Saeed Ghadimi and Guanghui Lan. Stochastic first-and zeroth-order methods for nonconvex stochastic programming. *SIAM Journal on Optimization*, 23(4):2341–2368, 2013.
- David Gross, Yi-Kai Liu, Steven T Flammia, Stephen Becker, and Jens Eisert. Quantum state tomography via compressed sensing. *Physical review letters*, 105(15):150401, 2010.
- Harold Hotelling. Analysis of a complex of statistical variables into principal components. *Journal of educational psychology*, 24(6):417, 1933.
- Martin Knott and Cyril S Smith. On a generalization of cyclic monotonicity and distances among random vectors. *Linear algebra and its applications*, 199:363–371, 1994.
- Gilad Lerman and Tyler Maunu. An overview of robust subspace recovery. *Proceedings of the IEEE*, 106(8):1380–1410, 2018.
- Yuanxin Li, Cong Ma, Yuxin Chen, and Yuejie Chi. Non-convex matrix factorization from rank-one measurements. In *The 22nd International Conference on Artificial Intelligence and Statistics*, pages 1496–1505. PMLR, 2019.
- Ziwei Liu, Ping Luo, Xiaogang Wang, and Xiaoou Tang. Deep learning face attributes in the wild. In *Proceedings of the IEEE international conference on computer vision*, pages 3730–3738, 2015.
- Estelle Massart and P-A Absil. Quotient geometry with simple geodesics for the manifold of fixed-rank positive-semidefinite matrices. *SIAM Journal on Matrix Analysis and Applications*, 41(1):171–198, 2020.
- Yurii Nesterov. *Introductory lectures on convex optimization: A basic course*, volume 87. Springer Science & Business Media, 2004.
- Karl Pearson. LIII. on lines and planes of closest fit to systems of points in space. *The London, Edinburgh, and Dublin philosophical magazine and journal of science*, 2(11):559–572, 1901.
- Benjamin Recht, Maryam Fazel, and Pablo A Parrilo. Guaranteed minimum-rank solutions of linear matrix equations via nuclear norm minimization. *SIAM review*, 52(3):471–501, 2010.
- Sujay Sanghavi, Rachel Ward, and Chris D White. The local convexity of solving systems of quadratic equations. *Results in Mathematics*, 71(3):569–608, 2017.
- Hanie Sedghi, Majid Janzamin, and Anima Anandkumar. Provable tensor methods for learning mixtures of generalized linear models. In *Artificial Intelligence and Statistics*, pages 1223–1231. PMLR, 2016.

- Stephen Tu, Ross Boczar, Max Simchowitz, Mahdi Soltanolkotabi, and Ben Recht. Low-rank solutions of linear matrix equations via Procrustes flow. In *International Conference on Machine Learning*, pages 964–973. PMLR, 2016.
- Cédric Villani. *Optimal transport: old and new*, volume 338 of *Grundlehren der Mathematischen Wissenschaften [Fundamental Principles of Mathematical Sciences]*. Springer-Verlag, Berlin, 2009.
- Gang Wang and Georgios Giannakis. Solving random systems of quadratic equations via truncated generalized gradient flow. *Advances in Neural Information Processing Systems*, 29, 2016.
- Gang Wang, Georgios Giannakis, Yousef Saad, and Jie Chen. Solving most systems of random quadratic equations. *Advances in Neural Information Processing Systems*, 30, 2017.
- Huan Xu, Constantine Caramanis, and Sujay Sanghavi. Robust pca via outlier pursuit. *Advances in neural information processing systems*, 23, 2010.
- Xinyang Yi, Constantine Caramanis, and Sujay Sanghavi. Alternating minimization for mixed linear regression. In *International Conference on Machine Learning*, pages 613–621. PMLR, 2014.
- Yoav Zemel and Victor M. Panaretos. Fréchet means and Procrustes analysis in Wasserstein space. *Bernoulli*, 25(2):932–976, 2019.
- Hongyi Zhang and Suvrit Sra. First-order methods for geodesically convex optimization. In *29th Annual Conference on Learning Theory*, volume 49 of *Proceedings of Machine Learning Research*, pages 1617–1638, Columbia University, New York, New York, USA, 2016.
- Kai Zhong, Prateek Jain, and Inderjit S Dhillon. Efficient matrix sensing using rank-1 Gaussian measurements. In *International conference on algorithmic learning theory*, pages 3–18. Springer, 2015.
- Kai Zhong, Prateek Jain, and Inderjit S Dhillon. Mixed linear regression with multiple components. In *NIPS*, pages 2190–2198, 2016.
- Kai Zhong, Zhao Song, Prateek Jain, Peter L Bartlett, and Inderjit S Dhillon. Recovery guarantees for one-hidden-layer neural networks. In *International conference on machine learning*, pages 4140–4149. PMLR, 2017.

A GEOMETRIES OF $\mathbb{S}_+^d$

Since the BW barycenter problem is inherently a geometric minimization problem over $\mathbb{S}_+^d$, we will quickly comment on how the choice of geometry effects optimization algorithms over this space.

There are many ways to define geometries over $\mathbb{S}_+^d$. For sake of comparison with previous methods for matrix recovery, we will compare BW space to the Euclidean geometry over PSD matrices. Among other things, the choice of geometry gives us a way of defining distance minimizing paths, or *geodesics*, over $\mathbb{S}_+^d$.

Consider two matrices $\Sigma_0, \Sigma_1 \in \mathbb{S}_+^d$ such that $\text{rank}(\Sigma_0 \Sigma_1) = \text{rank}(\Sigma_0) = \text{rank}(\Sigma_1)$. This is a sufficient but not necessary condition to ensure that our following definition of BW geodesic is well defined since the same map can work for transporting higher rank PSD matrices to lower rank matrices. In any case, under this assumption, there exists a transport map from Σ_0 to Σ_1 given by

$$T = \Sigma_1^{1/2} (\Sigma_1^{1/2} \Sigma_0 \Sigma_1^{1/2})^{-1/2} \Sigma_1^{1/2}, \quad (\text{A.1})$$

where the inverse is actually a pseudoinverse and one can check that $T \Sigma_0 T = \Sigma_1$. In this case, the Euclidean and BW geodesics $\Sigma_t : [0, 1] \rightarrow \mathbb{S}_+^d$ are given by

$$\Sigma_t = (1-t)\Sigma_0 + t\Sigma_1, \quad (\text{EG})$$

$$\Sigma_t = (I + t(T - I))\Sigma_0(I + t(T - I)). \quad (\text{BWG})$$

The first choice of geodesic, the Euclidean Geodesic (EG), corresponds to the distance functional $\|\Sigma_0 - \Sigma_1\|_F$. The second choice of geodesic, the BW Geodesic (BWG), corresponds to the BW distance functional $d_{\text{BW}}(\Sigma_0, \Sigma_1)$.

We note that (BWG) is equivalent to

$$\Sigma_t = (I + t(T - I))\Sigma_0(I + t(T - I)) \quad (\text{A.2})$$

$$= (\Sigma_0^{1/2} + t\Delta)(\Sigma_0^{1/2} + t\Delta)^\top. \quad (\text{A.3})$$

where $\Delta = (T - I)\Sigma_0^{1/2}$.

One of the big differences in these two paths is that while the Euclidean geodesics are linear in t , which is a result of the underlying flatness of the space, the BW geodesic contains terms that are quadratic in t . This points to the fact that this choice of geometry adds *curvature* to the space $\mathbb{S}_+^d$.

If we restrict ourselves to rank r PSD matrices, $\mathbb{S}_+^{d,r}$, the geodesic (BWG) becomes

$$\begin{aligned} \Sigma_t &= (I + t(T - I)U_0U_0^\top(I + t(T - I))^\top) \\ &= ((1-t)U_0 + tU_1)((1-t)U_0 + tU_1)^\top, \end{aligned} \quad (\text{A.4})$$

where we use the fact that $T\Sigma_0T = TU_0U_0^\top T = \Sigma_1 = U_1U_1^\top$. Note that there is an inherent rotational symmetry in the problem, since for any $R \in \mathbb{R}^{r \times r}$ such that $RR^\top = I$ and $UU^\top \in \mathbb{S}_+^{d,r}$, $UU^\top = URR^\top U^\top$.

B FIXED POINTS AND BWGD

Here we briefly show how one can interpret the fixed point equation of [Álvarez-Esteban et al. \[2016\]](#) as gradient descent with step size 1. Consider the case where Σ_k is full rank. The fixed point iteration is given by

$$\Sigma_{k+1} = \Sigma_k^{-1/2} \left(\frac{1}{n} \sum_{i=1}^n (\Sigma_k^{1/2} X_i \Sigma_k^{1/2})^{1/2} \right)^2 \Sigma_k^{-1/2} \quad (\text{B.1})$$

We can write

$$\begin{aligned} &\Sigma_k^{-1/2} \left(\frac{1}{n} \sum_{i=1}^n (\Sigma_k^{1/2} X_i \Sigma_k^{1/2})^{1/2} \right)^2 \Sigma_k^{-1/2} \\ &= \Sigma_k^{-1/2} \left(\frac{1}{n} \sum_{i=1}^n (\Sigma_k^{1/2} X_i \Sigma_k^{1/2})^{1/2} \right) \Sigma_k^{-1/2} \Sigma_k \Sigma_k^{-1/2}. \end{aligned} \quad (\text{B.2})$$

$$\begin{aligned} & \left(\frac{1}{n} \sum_{i=1}^n (\Sigma_k^{1/2} \mathbf{X}_i \Sigma_k^{1/2})^{1/2} \right) \Sigma_k^{-1/2} \\ &= (\mathbf{I} - \nabla F(\Sigma_k)) \Sigma_k (\mathbf{I} - \nabla F(\Sigma_k)). \end{aligned}$$

On the other hand, we note that the fixed point equation in [Álvarez-Esteban et al. \[2016\]](#) does not work for low-rank matrices, since the span in (B.1) is not allowed to change between iterations. Instead, one must use the transport map $\mathbf{X}_i^{-1/2} (\mathbf{X}_i^{1/2} \Sigma_k \mathbf{X}_i^{1/2})^{+1/2} \mathbf{X}_i^{1/2}$, which results in the fixed point iteration

$$\begin{aligned} \Sigma_{k+1} &= \left(\frac{1}{n} \sum_{i=1}^n \mathbf{X}_i^{1/2} (\mathbf{X}_i^{1/2} \Sigma_k \mathbf{X}_i^{1/2})^{+1/2} \mathbf{X}_i^{1/2} \right) \Sigma_k. \\ & \left(\frac{1}{n} \sum_{i=1}^n \mathbf{X}_i^{1/2} (\mathbf{X}_i^{1/2} \Sigma_k \mathbf{X}_i^{1/2})^{+1/2} \mathbf{X}_i^{1/2} \right) \end{aligned} \quad (\text{B.3})$$

and fixed point equation

$$\begin{aligned} \Sigma &= \left(\frac{1}{n} \sum_{i=1}^n \mathbf{X}_i^{1/2} (\mathbf{X}_i^{1/2} \Sigma \mathbf{X}_i^{1/2})^{+1/2} \mathbf{X}_i^{1/2} \right) \Sigma. \\ & \left(\frac{1}{n} \sum_{i=1}^n \mathbf{X}_i^{1/2} (\mathbf{X}_i^{1/2} \Sigma \mathbf{X}_i^{1/2})^{+1/2} \mathbf{X}_i^{1/2} \right). \end{aligned} \quad (\text{B.4})$$

In essence, this says that Σ is an exponential barycenter.

C SUPPLEMENTARY PROOFS

C.1 Proof of Propositions 1 and 2

Proposition 1 *If $\frac{1}{n} \sum_i \mathbf{x}_i \mathbf{x}_i^\top = \mathbf{I}$, then*

$$\begin{aligned} & \operatorname{argmin}_{\Sigma \in \mathbb{S}_+^d} \frac{1}{2n} \|\sqrt{\mathcal{A}(\Sigma)} - \sqrt{\mathbf{y}}\|_2^2 \\ &= \operatorname{argmin}_{\Sigma \in \mathbb{S}_+^d} \frac{1}{2n} \sum_{i=1}^n d_{\text{BW}}^2(\Sigma, y_i \mathbf{x}_i \mathbf{x}_i^\top). \end{aligned}$$

Proof of Proposition 1. First, if $\frac{1}{n} \sum_i \mathbf{x}_i \mathbf{x}_i^\top = \mathbf{I}$, then

$$\frac{1}{n} \sum_i \mathcal{A}(\Sigma)_i = \operatorname{Tr}(\Sigma).$$

Indeed, this follows from the fact that

$$\begin{aligned} \frac{1}{n} \sum_i \mathcal{A}(\Sigma)_i &= \frac{1}{n} \sum_i \mathbf{x}_i^\top \Sigma \mathbf{x}_i \\ &= \operatorname{Tr} \left(\Sigma \frac{1}{n} \sum_i \mathbf{x}_i \mathbf{x}_i^\top \right) \\ &= \operatorname{Tr}(\Sigma). \end{aligned}$$

With this in mind, we expand the square in (2.4)

$$\frac{1}{2n} \|\sqrt{\mathcal{A}(\Sigma)} - \sqrt{\mathbf{y}}\|_2^2 = \quad (\text{C.1})$$

$$\begin{aligned} & \frac{1}{2n} \sum_i \mathcal{A}(\Sigma)_i + \frac{1}{2n} \sum_i y_i - \frac{1}{n} \sum_i \sqrt{y_i} \sqrt{\mathbf{x}_i^\top \Sigma \mathbf{x}_i} \\ &= \frac{1}{2} \text{Tr}(\Sigma) - \frac{1}{n} \sum_i \sqrt{y_i} \sqrt{\mathbf{x}_i^\top \Sigma \mathbf{x}_i} + \frac{1}{2n} \sum_i y_i. \end{aligned}$$

Thus, the minimization in (2.4) is equivalent to the program

$$\min_{\Sigma \in \mathbb{S}_+^d} \text{Tr}(\Sigma) - \frac{2}{n} \sum_{i=1}^n \sqrt{y_i} \sqrt{\mathbf{x}_i^\top \Sigma \mathbf{x}_i} \quad (\text{C.2})$$

On the other hand, it is not hard to show that the BW distance between a matrix $\Sigma \in \mathbb{S}_+^d$ and a rank one matrix $\mathbf{w}\mathbf{w}^\top$ is

$$d_{\text{BW}}^2(\Sigma, \mathbf{w}\mathbf{w}^\top) = \left[\text{Tr}(\Sigma) + \text{Tr}(\mathbf{x}\mathbf{x}^\top) - 2\sqrt{\text{Tr}(\mathbf{x}\mathbf{x}^\top \Sigma)} \right]. \quad (\text{C.3})$$

Letting $\mathbf{w}\mathbf{w}^\top = y_i \mathbf{x}_i \mathbf{x}_i^\top$ and summing over i , we see that (C.2) is equivalent to minimizing (2.9) when

$$Q = \frac{1}{n} \sum_{i=1}^n \delta_{y_i \mathbf{x}_i \mathbf{x}_i^\top}.$$

□

Proposition 2 Let $C_n = \frac{1}{n} \sum_{i=1}^n \mathbf{x}_i \mathbf{x}_i^\top \in \mathbb{S}_{++}^d$. Then

$$\begin{aligned} & \underset{\Sigma \in \mathbb{S}_+^d}{\text{argmin}} \frac{1}{2n} \left\| \sqrt{\mathcal{A}(C_n^{-1/2} \Sigma C_n^{-1/2})} - \sqrt{\mathbf{y}} \right\|^2 \\ &= \underset{\Sigma \in \mathbb{S}_+^d}{\text{argmin}} \frac{1}{2n} \sum_{i=1}^n d_{\text{BW}}^2(\Sigma, y_i C_n^{-1/2} \mathbf{x}_i \mathbf{x}_i^\top C_n^{-1/2}), \end{aligned}$$

and $C_n^{1/2} \mathbf{S} C_n^{1/2}$ is a solution to both problems.

Proof of Proposition 2. Consider the first order condition for the rank one matrices $y_i^2 C_n^{-1/2} \mathbf{x}_i \mathbf{x}_i^\top C_n^{-1/2}$ given by

$$\frac{1}{n} \sum_{i=1}^n \frac{\text{Tr}(\mathbf{x}_i \mathbf{x}_i^\top \mathbf{S})^{1/2} C_n^{-1/2} \mathbf{x}_i \mathbf{x}_i^\top C_n^{-1/2}}{\text{Tr}(C_n^{-1/2} \mathbf{x}_i \mathbf{x}_i^\top C_n^{-1/2} \Sigma_n)^{1/2}} = \mathbf{I}. \quad (\text{C.4})$$

Since $C_n^{-1/2} C_n C_n^{-1/2} = \mathbf{I}$, we see that $C_n^{-1/2} \Sigma_n C_n^{-1/2} = \mathbf{S}$ is a sufficient condition for Σ_n to be a barycenter of

$$Q = \frac{1}{n} \sum_{i=1}^n \delta_{y_i C_n^{-1/2} \mathbf{x}_i \mathbf{x}_i^\top C_n^{-1/2}} \quad (\text{C.5})$$

□

C.2 Proof of Theorem 4

The proof of Theorem 4 proceeds in the following sections. In Section C.2.1, we establish some restricted isometry properties (RIP) that will be used in our proof. Then, Section C.2.2 proves that the function F is Euclidean strongly convex over $\mathbb{S}_+^d$ with high probability. After this, Section C.2.3 shows that F also satisfies a certain Euclidean smoothness over rank r matrices. Section C.2.4 discusses first order optimality conditions for the barycenter, and in particular gives sufficient conditions for a point $\bar{\Sigma}$ to be the minimizer of F . Section C.2.5 gives a descent lemma for the fixed-rank gradient descent method. After this, Section C.2.6 proves local strong convexity of F over fixed-rank PSD matrices in $\mathbb{S}_+^{d,r}$ that are close to $\mathbf{S}$. We finish in Section C.2.7 by putting all of these facts together.

C.2.1 RIP Conditions

We discuss here a case an RIP condition that becomes essential for our later proof. As discussed in Cai and Zhang [2015], issues arise in trying to prove a full ℓ_2 RIP for this problem, due to the fact that the fourth moments of $\mathbf{x}$ that show up. Instead, both Cai and Zhang [2015], Chen et al. [2015] prove the following ℓ_2/ℓ_1 RIP condition (also referred to as “Restricted Uniform Boundedness”).

Theorem 9 (Chen et al. [2015] Proposition 1). *Suppose that $\mathbf{x}_1, \dots, \mathbf{x}_n$ are a sample from a sub-Gaussian distribution with $\mathbb{E}\mathbf{x}_i = \mathbf{0}$, $\mathbb{E}x_{ij}^2 = 1$, and $\mathbb{E}x_{ij}^4 > 1$. Then, there are constants $c_1, c_2, c_3 > 0$ such that with probability exceeding $1 - \exp(-c_3 n)$*

$$c_1 \|\Delta\|_F \leq \frac{1}{n} \left\| \begin{bmatrix} \mathbf{x}_1^\top \Delta \mathbf{x}_1 - \mathbf{x}_2^\top \Delta \mathbf{x}_2 \\ \mathbf{x}_3^\top \Delta \mathbf{x}_3 - \mathbf{x}_4^\top \Delta \mathbf{x}_4 \\ \vdots \\ \mathbf{x}_{n-1}^\top \Delta \mathbf{x}_{n-1} - \mathbf{x}_n^\top \Delta \mathbf{x}_n \end{bmatrix} \right\|_1 \leq c_2 \|\Delta\|_F. \quad (\text{C.6})$$

hold simultaneously for all rank r matrices Δ provided that $n \gtrsim dr$.

We have the following corollary to this theorem.

Corollary 10. *Suppose that a random vector $\mathbf{w} = (w_1, \dots, w_d)^\top$ is sub-Gaussian with $\mathbb{E}\mathbf{w} = \mathbf{0}$, $\mathbb{E}w_j^2 = 1$, $\mathbb{E}w_j^4 > 1$. Then, the following population RIP holds:*

$$\mathbb{E}_{\mathbf{w}} \text{Tr}(\mathbf{w}\mathbf{w}^\top \mathbf{A})^2 \geq c_1 \|\mathbf{A}\|_F^2. \quad (\text{C.7})$$

Furthermore, for a sample of n i.i.d. copies of $\mathbf{W}$, $\mathbf{w}_1, \dots, \mathbf{w}_n$, there are constants $c_2, c_3 > 0$ such that with probability exceeding $1 - \exp(-c_3 n)$

$$\frac{1}{n} \sum_{i=1}^n \text{Tr}(\mathbf{w}_i \mathbf{w}_i^\top \mathbf{A})^2 \geq c_2 \|\mathbf{A}\|_F^2 \quad (\text{C.8})$$

hold simultaneously for all rank r matrices Δ provided that $n \gtrsim dr$.

Proof. The first part holds from the argument in Appendix A of Chen et al. [2015].

For the second part, we compute

$$\begin{aligned} \frac{1}{n} \sum_{i=1}^n \text{Tr}(\mathbf{w}_i \mathbf{w}_i^\top \mathbf{A})^2 &= \frac{1}{n} \sum_{i=1}^n \text{Tr}(\mathbf{w}_i \mathbf{w}_i^\top \mathbf{A})^2 \\ &= \frac{1}{n} \sum_{i=1}^{n/2} \text{Tr}(\mathbf{w}_{2i} \mathbf{w}_{2i}^\top \mathbf{A})^2 + \text{Tr}(\mathbf{w}_{2i-1} \mathbf{w}_{2i-1}^\top \mathbf{A})^2 \\ &\geq \frac{1}{2n} \sum_{i=1}^{n/2} [\text{Tr}(\mathbf{w}_{2i} \mathbf{w}_{2i}^\top \mathbf{A}) - \text{Tr}(\mathbf{w}_{2i-1} \mathbf{w}_{2i-1}^\top \mathbf{A})]^2 \\ &= \frac{1}{n^2} \sum_{i=1}^{n/2} [\text{Tr}(\mathbf{w}_{2i} \mathbf{w}_{2i}^\top \mathbf{A}) - \text{Tr}(\mathbf{w}_{2i-1} \mathbf{w}_{2i-1}^\top \mathbf{A})]^2 \\ &\geq \frac{1}{n^2} \left[\sum_{i=1}^{n/2} |\text{Tr}(\mathbf{w}_{2i} \mathbf{w}_{2i}^\top \mathbf{A}) - \text{Tr}(\mathbf{w}_{2i-1} \mathbf{w}_{2i-1}^\top \mathbf{A})| \right]^2 \\ &\geq c_2^2 \|\mathbf{A}\|_F^2. \end{aligned} \quad (\text{C.9})$$

□

C.2.2 Euclidean Strong Convexity

We begin by proving strong convexity along Euclidean geodesics. In particular, this guarantees that $\mathbf{S}$ is the unique minimizer of F over $\mathbb{S}_+^d$ in the population setting and $\mathbf{S}_n$ is the unique rank r global minimizer over $\mathbb{S}_+^{d,r}$.

Population Setting: First, it is easy to see that $\mathbf{S}$ is stationary because

$$\nabla F(\mathbf{S}) = \mathbf{I} - \mathbb{E}_Q \sqrt{\frac{\mathbf{x}^\top \mathbf{S} \mathbf{x}}{\mathbf{x}^\top \mathbf{S} \mathbf{x}}} \mathbf{x} \mathbf{x}^\top = \mathbf{I} - \mathbb{E}_Q \mathbf{x} \mathbf{x}^\top = 0. \quad (\text{C.10})$$

Define the set

$$\mathcal{S}(m, M) = \{\mathbf{\Sigma} \in \mathbb{S}_+^d : m \leq \lambda_r(\mathbf{\Sigma}) \leq \lambda_1(\mathbf{\Sigma}) \leq M\}. \quad (\text{C.11})$$

We have the following lemma over this set.

Lemma 11. *Let $y = \langle \mathbf{x} \mathbf{x}^\top, \mathbf{S} \rangle$, where $\mathbf{x} \sim N(\mathbf{0}, \mathbf{I})$, $\mathbf{S}$ is rank r , and m, M be such that $\mathbf{S} \in \mathcal{S}(m, M)$. Then, the population objective F is Euclidean strongly convex over $\mathcal{S}(0, M)$ with constant $\frac{c_1 \beta^2 d}{6M^{3/2}}$.*

In other words, the function $F(\mathbf{\Sigma}_t)$ is strongly convex when $\mathbf{\Sigma}_t$ is defined by (EG).

Proof. Let $\mathbf{\Sigma}_t$ denote the geodesic between $\mathbf{\Sigma}_0$ and $\mathbf{\Sigma}_1$ given by $(1-t)\mathbf{\Sigma}_0 + t\mathbf{\Sigma}_1$. We compute the derivatives of the BW distance as

$$\partial_t d_{\text{BW}}(\mathbf{\Sigma}_t, \mathbf{x} \mathbf{x}^\top)^2 = \text{Tr}(\mathbf{\Sigma}_1 - \mathbf{\Sigma}_0) \quad (\text{C.12})$$

$$- \frac{\text{Tr}((\mathbf{\Sigma}_1 - \mathbf{\Sigma}_0) \mathbf{x} \mathbf{x}^\top)}{\sqrt{\mathbf{x}^\top \mathbf{\Sigma}_t \mathbf{x}}}, \quad (\text{C.13})$$

$$\partial_t^2 d_{\text{BW}}(\mathbf{\Sigma}_t, \mathbf{x} \mathbf{x}^\top)^2 = \frac{1}{2} \frac{\text{Tr}(\mathbf{x} \mathbf{x}^\top (\mathbf{\Sigma}_1 - \mathbf{\Sigma}_0))^2}{\text{Tr}(\mathbf{x} \mathbf{x}^\top \mathbf{\Sigma}_t)^{3/2}} \quad (\text{C.14})$$

By the definition of M , we have the uniform bound

$$\begin{aligned} \partial_t^2 F(\mathbf{\Sigma}_t)|_{t=s} &= \mathbb{E}_{\mathbf{x} \mathbf{x}^\top \sim Q} \frac{1}{2} \frac{\text{Tr}(\mathbf{x} \mathbf{x}^\top (\mathbf{\Sigma}_1 - \mathbf{\Sigma}_0))^2}{(\text{Tr}(\mathbf{x} \mathbf{x}^\top \mathbf{\Sigma}_s))^{3/2}} \\ &= \mathbb{E}_{\mathbf{x} \mathbf{x}^\top \sim Q} \frac{1}{2 \|\mathbf{x}\|^3} \frac{\text{Tr}(\mathbf{x} \mathbf{x}^\top (\mathbf{\Sigma}_1 - \mathbf{\Sigma}_0))^2}{(\text{Tr}(\mathbf{x} \mathbf{x}^\top \mathbf{\Sigma}_s / \|\mathbf{x}\|^2))^{3/2}} \\ &\geq \frac{1}{2(2M)^{3/2}} \mathbb{E}_{\mathbf{x} \mathbf{x}^\top \sim Q} \frac{1}{\|\mathbf{x}\|^3} \text{Tr}(\mathbf{x} \mathbf{x}^\top (\mathbf{\Sigma}_1 - \mathbf{\Sigma}_0))^2. \end{aligned} \quad (\text{C.15})$$

This then implies that

$$\begin{aligned} \partial_t^2 F(\mathbf{\Sigma}_t)|_{t=s} &\geq \frac{1}{6M^{3/2}} \mathbb{E}_{\mathbf{x} \sim N(\mathbf{0}, \mathbf{I})} \frac{\sqrt{\mathbf{x}^\top \mathbf{S} \mathbf{x}}}{\|\mathbf{x}\|^3} \text{Tr}(\mathbf{x} \mathbf{x}^\top (\mathbf{\Sigma}_1 - \mathbf{\Sigma}_0))^2 \\ &\geq \frac{1}{6M^{3/2}} \mathbb{E}_{\mathbf{x} \sim N(\mathbf{0}, \mathbf{I})} \|\mathbf{x}\|^2 \sqrt{\frac{\mathbf{x}^\top \mathbf{S} \mathbf{x}}{\|\mathbf{x}\|^2}} \text{Tr}\left(\frac{\mathbf{x} \mathbf{x}^\top}{\|\mathbf{x}\|^2} (\mathbf{\Sigma}_1 - \mathbf{\Sigma}_0)\right)^2 \\ &= \frac{d}{6M^{3/2}} \mathbb{E}_{\mathbf{x} \sim N(\mathbf{0}, \mathbf{I})} \sqrt{\frac{\mathbf{x}^\top \mathbf{S} \mathbf{x}}{\|\mathbf{x}\|^2}} \text{Tr}\left(\frac{\mathbf{x} \mathbf{x}^\top}{\|\mathbf{x}\|^2} (\mathbf{\Sigma}_1 - \mathbf{\Sigma}_0)\right)^2. \end{aligned} \quad (\text{C.16})$$

Define the random vector

$$\mathbf{w} := \left(\frac{\mathbf{x}^\top \mathbf{S} \mathbf{x}}{\|\mathbf{x}\|^2} \right)^{1/8} \frac{\mathbf{x}}{\|\mathbf{x}\|}. \quad (\text{C.17})$$

By symmetry, $\mathbf{w}$ is mean zero. Furthermore, $\mathbf{w}$ is bounded and thus sub-Gaussian. Furthermore, it is a simple exercise to show that

$$\mathbf{C}_W := \mathbb{E} \mathbf{w} \mathbf{w}^\top \succeq \beta \mathbf{I}, \quad (\text{C.18})$$

for some dimension and $\mathbf{S}$ -dependent constant β . We can thus bound

$$\begin{aligned} \mathbb{E}_{\mathbf{x}} \text{Tr}(\mathbf{w} \mathbf{w}^\top (\mathbf{\Sigma}_1 - \mathbf{\Sigma}_0))^2 &= \mathbb{E}_W \text{Tr}\left(\mathbf{C}_W^{1/2} \mathbf{C}_W^{-1/2} \mathbf{w} \mathbf{w}^\top \mathbf{C}_W^{-1/2} \mathbf{C}_W^{1/2} (\mathbf{\Sigma}_1 - \mathbf{\Sigma}_0)\right)^2 \end{aligned} \quad (\text{C.19})$$

$$\geq \beta^2 \mathbb{E}_W \text{Tr}(C_W^{-1/2} \mathbf{w} \mathbf{w}^\top C_W^{-1/2} (\boldsymbol{\Sigma}_1 - \boldsymbol{\Sigma}_0))^2.$$

Using the fact that $C_W^{-1/2} \mathbf{w}$ is mean zero sub-Gaussian with identity covariance (and furthermore the 4th moment condition is trivially satisfied), we can now use the RIP condition of Corollary 10 to bound

$$\begin{aligned} \mathbb{E}_{\mathbf{w} \mathbf{w}^\top \sim Q} \text{Tr}(C_W^{-1/2} \mathbf{w} \mathbf{w}^\top C_W^{-1/2} (\boldsymbol{\Sigma}_1 - \boldsymbol{\Sigma}_0))^2 \\ \geq c_1 \|\boldsymbol{\Sigma}_1 - \boldsymbol{\Sigma}_0\|_F^2. \end{aligned} \quad (\text{C.20})$$

Putting this all together,

$$\begin{aligned} \partial_t^2 F(\boldsymbol{\Sigma}_t)|_{t=s} &\geq \frac{c_1 \beta^2 d}{2(2M)^{3/2}} \|\boldsymbol{\Sigma}_1 - \boldsymbol{\Sigma}_0\|_F^2 \\ &\geq \frac{c_1 \beta^2 d}{6M^{3/2}} \|\boldsymbol{\Sigma}_1 - \boldsymbol{\Sigma}_0\|_F^2. \end{aligned} \quad (\text{C.21})$$

□

Sample Setting: The sample setting requires a bit more work. In this case, we observe y_i and $\mathbf{x}_i$, $i = 1, \dots, n$, and to recover the matrix $\mathbf{S}$ we first compute the barycenter of $y_i C_n^{-1/2} \mathbf{x}_i \mathbf{x}_i^\top C_n^{-1/2}$, where C_n is the sample covariance. We denote this whitened discrete distribution by $\tilde{Q}$. In this case, by Proposition 2, the matrix we hope to recover is $\mathbf{S}_n = C_n^{-1/2} \mathbf{S} C_n^{-1/2}$.

Lemma 12. *If $n \gtrsim d^2$, then with probability at least $1 - \exp(-c_2 n)$ for some constant c_2 , F is Euclidean strongly convex over $\mathcal{S}(0, M)$ with constant $\frac{c_3 \beta^2 d}{6M^{3/2}}$.*

Proof. In this case, the second statement of Corollary 10 holds with rank d ,

$$\begin{aligned} \partial_t^2 F(\boldsymbol{\Sigma}_t)|_{t=s} &\geq \\ &\frac{1}{6M^{3/2}} \frac{1}{n} \sum_{i=1}^n \frac{\sqrt{\mathbf{x}_i^\top \mathbf{S} \mathbf{x}_i}}{\|C_n^{-1/2} \mathbf{x}_i\|^3} \text{Tr}(C_n^{-1/2} \mathbf{x}_i \mathbf{x}_i^\top C_n^{-1/2} (\mathbf{S}_n - \boldsymbol{\Sigma}))^2 \\ &\geq \frac{1}{6M^{3/2}} \frac{\lambda_{\min}(C_n)^{3/2}}{\lambda_{\max}(C_n)^2} \frac{1}{n} \sum_{i=1}^n \frac{\sqrt{\mathbf{x}_i^\top \mathbf{S} \mathbf{x}_i}}{\|\mathbf{x}_i\|^3} \text{Tr}(\mathbf{x}_i \mathbf{x}_i^\top (\mathbf{S}_n - \boldsymbol{\Sigma}))^2 \\ &= \frac{d}{6M^{3/2}} \frac{\lambda_{\min}(C_n)^{3/2}}{\lambda_{\max}(C_n)^2} \\ &\quad \frac{1}{n} \sum_{i=1}^n \text{Tr} \left(\frac{\|\mathbf{x}_i\|}{d} \left(\frac{\mathbf{x}_i^\top \mathbf{S} \mathbf{x}_i}{\|\mathbf{x}_i\|^2} \right)^{1/4} \frac{\mathbf{x}_i \mathbf{x}_i^\top}{\|\mathbf{x}_i\|^2} (\mathbf{S}_n - \boldsymbol{\Sigma}) \right)^2. \end{aligned} \quad (\text{C.22})$$

Then, following the same line of reasoning as in the previous proof to obtain strong convexity the sub-Gaussianity of the random vector

$$\mathbf{w}_i = \frac{\|\mathbf{x}_i\|^{1/2}}{\sqrt{d}} \left(\frac{\mathbf{x}_i^\top \mathbf{S} \mathbf{x}_i}{\|\mathbf{x}_i\|^2} \right)^{1/8} \frac{\mathbf{x}_i}{\|\mathbf{x}_i\|}, \quad (\text{C.23})$$

which again a simple exercise shows $\mathbb{E} \mathbf{w}_i \mathbf{w}_i^\top \succeq \beta \mathbf{I}$ and $\mathbb{E} \mathbf{w}_i = 0$. Then, with probability at least $1 - \exp(-c_2 n)$, the discrete distribution is strongly convex with constant $\frac{c_3 \beta^2 d}{6M^{3/2}}$. □

Notice that in both the population and the sample setting with high probability we get a unique barycenter $\mathbf{S}$ (and $\mathbf{S}_n$, respectively).

Remark 13. *When $n \gtrsim dr$, we can repeat the argument in Lemma 12 by applying Corollary 10 with the rank set to $2r$ to yield that $\mathbf{S}_n$ is the unique rank r minimizer of F over $\mathcal{S}_+^{d,r}$.*

C.2.3 Local Generalized Euclidean Smoothness

We remind ourselves that

$$\nabla F(\Sigma) = \mathbf{I} - \mathbb{E}_Q \sqrt{\frac{\mathbf{x}^\top \mathbf{S} \mathbf{x}}{\mathbf{x}^\top \Sigma \mathbf{x}}} \mathbf{x} \mathbf{x}^\top, \quad (\text{C.24})$$

In gradient-based optimization, a function is typically called *smooth* if it has a Lipschitz continuous gradient in the Euclidean sense. In our case, this would correspond to proving a bound of the form

$$\|\nabla F(\mathbf{S}) - \nabla F(\Sigma)\|_F \lesssim \|\Sigma - \mathbf{S}\|_F. \quad (\text{C.25})$$

While we are not able to prove a bound of this form, locally we are able to prove that the gradient is $1/2$ -Hölder continuous.

Using this, we have the following lemma.

Lemma 14. *Let $y = \langle \mathbf{x} \mathbf{x}^\top, \mathbf{S} \rangle$, $X \sim N(\mathbf{0}, \mathbf{I})$, $\text{rank}(\Sigma) \geq 3$, and $\lambda_r(\Sigma) \geq m$. Then,*

$$\begin{aligned} \|\nabla F(\mathbf{S}) - \nabla F(\Sigma)\|_F &= \|\nabla F(\Sigma)\|_F \\ &\leq \frac{d\sqrt{r}}{m\sqrt{r-2}} \|\Sigma - \mathbf{S}\|_F^{1/2} \end{aligned} \quad (\text{C.26})$$

Proof. We have that

$$\begin{aligned} &\left\| \mathbb{E}_Q \sqrt{\frac{\mathbf{x}^\top \mathbf{S} \mathbf{x}}{\mathbf{x}^\top \Sigma \mathbf{x}}} \mathbf{x} \mathbf{x}^\top - \mathbb{E}_Q \sqrt{\frac{\mathbf{x}^\top \mathbf{S} \mathbf{x}}{\mathbf{x}^\top \Sigma \mathbf{x}}} \mathbf{x} \mathbf{x}^\top \right\| \leq \\ &\mathbb{E}_Q \left| \frac{\sqrt{\mathbf{x}^\top \Sigma \mathbf{x}} - \sqrt{\mathbf{x}^\top \mathbf{S} \mathbf{x}}}{\sqrt{\mathbf{x}^\top \Sigma \mathbf{x}}} \right| \|\mathbf{x}\|^2 \\ &= d \mathbb{E}_Q \frac{1}{\sqrt{\mathbf{x}^\top \Sigma \mathbf{x}}} \sqrt{|\mathbf{x}^\top (\Sigma - \mathbf{S}) \mathbf{x}|} \\ &\leq d \|\Sigma - \mathbf{S}\|_F^{1/2} \mathbb{E}_Q \frac{\|P_{\Sigma - \mathbf{S}} \mathbf{x}\|}{\sqrt{\mathbf{x}^\top \Sigma \mathbf{x}}} \end{aligned} \quad (\text{C.27})$$

As long as $r \geq 3$ and $\lambda_r(\Sigma) \geq m$, we have that

$$\mathbb{E}_Q \frac{\|P_{\Sigma - \mathbf{S}} \mathbf{x}\|}{\sqrt{\mathbf{x}^\top \Sigma \mathbf{x}}} \leq \frac{1}{m} \sqrt{1 + \frac{r}{r-2}} = \frac{\sqrt{r}}{m\sqrt{r-2}}. \quad (\text{C.28})$$

□

This demonstrates that as $\Sigma \rightarrow \mathbf{S}$ for $\text{rank}(\Sigma) \geq 3$ $\|\nabla F(\Sigma)\|_F \rightarrow 0$.

This can be extended to the finite sample case as well using concentration of the quantity

$$\frac{1}{n} \sum_{i=1}^n \frac{\|P_{\Sigma - \mathbf{S}} \mathbf{x}_i\|_2}{\sqrt{\mathbf{x}_i^\top \Sigma \mathbf{x}_i}}. \quad (\text{C.29})$$

In particular, we have that

$$\frac{1}{n} \sum_{i=1}^n \frac{\|P_{\Sigma - \mathbf{S}} \mathbf{x}_i\|_2}{\sqrt{\mathbf{x}_i^\top \Sigma \mathbf{x}_i}} \leq \sqrt{\frac{1}{n} \sum_{i=1}^n \frac{\|P_{\Sigma - \mathbf{S}} \mathbf{x}_i\|_2^2}{\mathbf{x}_i^\top \Sigma \mathbf{x}_i}} \quad (\text{C.30})$$

and

$$\frac{1}{n} \sum_{i=1}^n \frac{\|P_{\Sigma - \mathbf{S}} \mathbf{x}_i\|_2^2}{\mathbf{x}_i^\top \Sigma \mathbf{x}_i} \leq \frac{1}{m} \left(1 + \frac{1}{n} \sum_{i=1}^n \frac{Y_i}{W_i} \right), \quad (\text{C.31})$$

where Y_i is a χ_r^2 random variable and W_i is a χ_r^2 random variable. By Chebyshevs inequality, denoting $\text{var}(Y_i/W_i) = \sigma^2$, we have

$$\Pr \left(\left| \frac{1}{n} \sum_{i=1}^n \frac{Y_i}{W_i} - \frac{r}{r-2} \right| \geq \frac{\epsilon \sigma^2}{n} \right) \leq \frac{1}{\epsilon^2}. \quad (\text{C.32})$$

Letting $\epsilon = \sqrt{n}$, we arrive at the following result.

Lemma 15. *Let n be such that $\sigma^2/\sqrt{n} \leq r/(2(r-2))$. Then, with probability at least $1 - 1/n$*

$$\|\nabla F(\mathbf{S}_n) - \nabla F(\mathbf{\Sigma})\|_F \leq \frac{d\sqrt{r}}{m\sqrt{r-2}} \|\mathbf{\Sigma} - \mathbf{S}_n\|_F^{1/2} \quad (\text{C.33})$$

C.2.4 First Order Optimality of the Low-Rank Barycenter

Moving on to the BW geometry, we first show that the low-rank barycenter is a first-order stationary point with respect to BW distance. While this follows from the previous results due to the fact that it is a global minimum over $\mathbb{S}_+^d$, we take a different approach here based on the fixed point iteration of [Agueh and Carlier \[2011\]](#). This fixed point iteration forms the basis of our efficient low-rank algorithm.

By [Agueh and Carlier \[2011\]](#), a sufficient condition for $\gamma_{\mathbf{\Sigma}}$ to solve (2.9), is

$$\mathbb{E}_{\mathbf{x}\mathbf{x}^\top \sim Q} \frac{\mathbf{x}\mathbf{x}^\top}{\text{Tr}(\mathbf{x}\mathbf{x}^\top \mathbf{\Sigma})^{1/2}} = \mathbf{I}, \quad (\text{C.34})$$

where $\mathbf{I}$ is the identity matrix in $\mathbb{R}^d$. Notice that this corresponds to the gradient ∇F being equal to $\mathbf{0}$. In the following, we let

$$\tilde{T}(\mathbf{\Sigma}) := \mathbb{E}_Q \frac{\mathbf{x}\mathbf{x}^\top}{\text{Tr}(\mathbf{x}\mathbf{x}^\top \mathbf{\Sigma})^{1/2}} \quad (\text{C.35})$$

Consider the map corresponding to gradient descent with step size 1,

$$\mathbf{\Sigma}_{k+1} = \left(\tilde{T}(\mathbf{\Sigma}_k) \right) \mathbf{\Sigma}_k \left(\tilde{T}(\mathbf{\Sigma}_k) \right). \quad (\text{C.36})$$

The corresponding fixed point equation is

$$\mathbf{\Sigma} = \left(\tilde{T}(\mathbf{\Sigma}) \right) \mathbf{\Sigma} \left(\tilde{T}(\mathbf{\Sigma}) \right). \quad (\text{C.37})$$

[Chewi et al. \[2020\]](#) prove that the operator norm $\|\cdot\|_2$ is convex along generalized geodesics. Therefore, (C.37) maps a compact subset of $\mathbb{S}_+$ to itself, and one can apply the Brouwer fixed point theorem to guarantee a solution. The fixed point satisfies a restricted first order condition given by

$$\mathbf{P}_{\text{Sp}(\mathbf{\Sigma})} \tilde{T}(\mathbf{\Sigma}) \mathbf{P}_{\text{Sp}(\mathbf{\Sigma})} = \mathbf{P}_{\text{Sp}(\mathbf{\Sigma})} = \text{id}_{\text{Sp}(\mathbf{\Sigma})} = \mathbf{P}_{\text{Sp}(\mathbf{\Sigma})}. \quad (\text{C.38})$$

Notice that if $\mathbf{\Sigma}$ is full rank, then this implies (C.34). More generally, we need an extra condition on top of first order optimality to guarantee that $\mathbf{\Sigma}$ is a barycenter.

Proposition 16 (Sufficient Condition for Barycenter). *If $\mathbf{\Sigma}$ satisfies the first order conditions*

$$\mathbf{P}_{\mathbf{\Sigma}} \tilde{T}(\mathbf{\Sigma}) \mathbf{P}_{\mathbf{\Sigma}} = \mathbf{P}_{\mathbf{\Sigma}}, \quad (\text{C.39})$$

$$\tilde{T}(\mathbf{\Sigma}) \preceq \mathbf{I}. \quad (\text{C.40})$$

then $\mathbf{\Sigma}$ is a barycenter of Q . Furthermore, in our observation model where Q is the law of $\sqrt{\mathbf{x}^\top \mathbf{S} \mathbf{x} \mathbf{x}^\top}$, for $\mathbf{x} \sim N(\mathbf{0}, \mathbf{I})$, $\mathbf{S}$ satisfies (C.39) and (C.40).

Proof. The first equation, (C.39), guarantees that $\mathbf{\Sigma}$ is a fixed point satisfying (C.37). On the other hand, (C.40) guarantees that all directional derivatives are positive (and that $\tilde{T}(\mathbf{\Sigma}) \preceq \mathbf{I}$), and thus $\mathbf{\Sigma}$ is a local minimum. To see this, suppose that $\mathbf{\Sigma} = \mathbf{\Sigma}_0$ is a fixed point and $\mathbf{\Sigma}_1$ is another PSD matrix. The directional derivatives are

$$\partial_t F(\mathbf{\Sigma}_t)|_{t=0} = \text{Tr} \left((\mathbf{I} - \tilde{T}(\mathbf{\Sigma})) (\mathbf{\Sigma}_1 - \mathbf{\Sigma}_0) \right) \quad (\text{C.41})$$

$$= \text{Tr} \left((I - \tilde{T}(\Sigma)) \Sigma_1 \right).$$

If $\tilde{T}(\Sigma) \not\preceq I$, then there is a Σ_1 such that this is less than zero, and therefore Σ_0 cannot be a minimum. On the other hand, if $\tilde{T}(\Sigma) \preceq I$, then all directional derivatives are positive. Combined with the Euclidean convexity result in the paper, this proves that Σ would be the global minimum and thus the barycenter.

Finally, in our observation model, we have

$$\tilde{T}(S) = I, \quad (\text{C.42})$$

and so both (C.39) and (C.40) hold. $\square$

Notice alternatively that this also implies that the barycenter is the only stationary point in the set where $\tilde{T}(\Sigma) \succeq I$. In particular, this is because at all points where $\tilde{T}(\Sigma) \neq I$, one can find a direction of decrease.

C.2.5 Smoothness and a Descent Lemma

The barycenter functional is *geodesically smooth*, as is shown in Theorem 7 of Chewi et al. [2020]. This result extends to non absolutely continuous measures by noting that 1) nonnegative curvature extends to measures that are not absolutely continuous with respect to Lebesgue, (see Ambrosio et al. [2008] Theorem 7.3.2), and 2) the characterization of the derivative of the Wasserstein distance extends to cases where the measures are not absolutely continuous (see Ambrosio et al. [2008] Lemma 7.3.6).

With smoothness, we have the following descent lemma over fixed rank BW space. Note that such a descent lemma is standard in the analysis of gradient descent methods (see Nesterov [2004, Theorem 2.1.5]). Here, Σ^+ is the update after one gradient step.

Lemma 17. *For any $\Sigma \in \mathbb{S}_+^{d,r}$ and $\Sigma^+ = \tilde{T}(\Sigma)\Sigma\tilde{T}(\Sigma)$, it holds that*

$$F(\Sigma^+) - F(\Sigma) \leq -\frac{1}{2} \left\| I - \tilde{T}(\Sigma) \right\|_{\gamma_\Sigma}^2. \quad (\text{C.43})$$

C.2.6 Local Geodesic Strong Convexity

We now prove local strong convexity in the population setting. By the previous section, we know that S is the unique barycenter of Q since the variance inequality holds for arbitrarily small m and ρ and arbitrarily large M .

Proposition 18. *Let $y = \langle \mathbf{x}\mathbf{x}^\top, S \rangle$, $X \sim N(\mathbf{0}, I)$, $\text{rank}(S) \geq 3$, and m, M be such that $m < \lambda_r(S) \leq \lambda_1(S) < M$. Then, if*

$$\rho \leq \frac{m^2(r-2)}{d^2r} \frac{c_1^2 \beta^4 m^2}{36M^3} \quad (\text{C.44})$$

F is locally geodesically strongly convex over the the largest BW ball inside the set

$$\mathcal{S}(m, M) \cap \{\Sigma \in \mathbb{S}_+^{d,r} : \|\Sigma - S\|_F \leq \rho\}. \quad (\text{C.45})$$

Proof. Fix $\Sigma_0, \Sigma_1 \in \mathcal{S}(m, M) \cap \{\Sigma \in \mathbb{S}_+^{d,r} : \|\Sigma - S\|_F \leq \rho\}$.

Suppose there exists a transport map T between Σ_0 and Σ_1 , which we are guaranteed for ρ sufficiently small. Then for the BW geodesic $\Sigma_t = ((1-t)I + tT)\Sigma_0((1-t)I + tT)$, we can compute

$$\begin{aligned} \partial_t^2 F(\Sigma_t)|_{t=s} &= \mathbb{E}_{\mathbf{x}\mathbf{x}^\top \sim Q} \text{Tr}[(T - I)\Sigma_0(T - I)] \\ &+ 2 \frac{\text{Tr}(\mathbf{x}\mathbf{x}^\top (T - I)\Sigma_0)^2}{(\text{Tr}(\mathbf{x}\mathbf{x}^\top \Sigma_s))^{3/2}} \\ &- \text{Tr} \left[\frac{\mathbf{x}\mathbf{x}^\top (T - I)\Sigma_0(T - I)}{\sqrt{\text{Tr}(\mathbf{x}\mathbf{x}^\top \Sigma_s)}} \right]. \end{aligned} \quad (\text{C.46})$$

Some manipulation when $s = 0$ yields

$$\partial_t^2 F(\Sigma_t)|_{t=0} = \left\langle (T - I)\Sigma_0(T - I), \nabla F(\Sigma_0) \right\rangle \quad (\text{C.47})$$

$$\begin{aligned}
 & + \mathbb{E}_{\mathbf{x}\mathbf{x}^\top \sim Q} 2 \frac{\text{Tr}(\mathbf{x}\mathbf{x}^\top (\mathbf{T} - \mathbf{I}) \boldsymbol{\Sigma}_0)^2}{(\text{Tr}(\mathbf{x}\mathbf{x}^\top \boldsymbol{\Sigma}_0))^{3/2}} \\
 & \geq -d_{\text{BW}}^2(\boldsymbol{\Sigma}_0, \boldsymbol{\Sigma}_1) \|\nabla F(\boldsymbol{\Sigma}_0)\|_F \\
 & + \mathbb{E}_{\mathbf{x}\mathbf{x}^\top \sim Q} 2 \frac{\text{Tr}(\mathbf{x}\mathbf{x}^\top (\mathbf{T} - \mathbf{I}) \boldsymbol{\Sigma}_0)^2}{(\text{Tr}(\mathbf{x}\mathbf{x}^\top \boldsymbol{\Sigma}_0))^{3/2}}.
 \end{aligned}$$

Following the proof of Lemma 11 last term satisfies the lower bound

$$\begin{aligned}
 & \mathbb{E}_{\mathbf{x}\mathbf{x}^\top \sim Q} 2 \frac{\text{Tr}(\mathbf{x}\mathbf{x}^\top (\mathbf{T} - \mathbf{I}) \boldsymbol{\Sigma}_0)^2}{(\text{Tr}(\mathbf{x}\mathbf{x}^\top \boldsymbol{\Sigma}_0))^{3/2}} \\
 & \geq \frac{c_1 \beta^2}{3M^{3/2}} \|(\mathbf{T} - \mathbf{I}) \boldsymbol{\Sigma}_0\|_F^2 \\
 & \geq \frac{c_1 \beta^2 m}{3M^{3/2}} \|(\mathbf{T} - \mathbf{I}) \boldsymbol{\Sigma}_0^{1/2}\|_F^2 \\
 & = \frac{c_1 \beta^2 m}{3M^{3/2}} d_{\text{BW}}^2(\boldsymbol{\Sigma}_0, \boldsymbol{\Sigma}_1)
 \end{aligned} \tag{C.48}$$

On the other hand, by our Euclidean arguments, $\mathbf{S}$ is the unique rank r point such that $\nabla F = \mathbf{0}$. We can upper bound $\|\nabla F(\boldsymbol{\Sigma})\|_F$ on the set $\{\boldsymbol{\Sigma} \in \mathbb{S}_+^{d,r} : \|\boldsymbol{\Sigma} - \mathbf{S}\|_F \leq \rho\}$ using Lemma 12 by

$$\|\nabla F(\boldsymbol{\Sigma})\|_F \leq \frac{d\sqrt{r}}{m\sqrt{r-2}} \sqrt{\rho}. \tag{C.49}$$

Thus, if

$$\rho \leq \frac{m^2(r-2)}{d^2 r} \frac{c_1^2 \beta^4 m^2}{36M^3},$$

then

$$\partial_t^2 F(\boldsymbol{\Sigma}_t)|_{t=0} \geq \frac{c_1 \beta^2 m}{6M^{3/2}} d_{\text{BW}}^2(\boldsymbol{\Sigma}_0, \boldsymbol{\Sigma}_1). \tag{C.50}$$

for all $\boldsymbol{\Sigma}_0 \in \{\boldsymbol{\Sigma} \in \mathbb{S}_+^{d,r} : \|\boldsymbol{\Sigma} - \mathbf{S}\|_F \leq \rho\}$. We thus have local strong geodesic convexity in the set

$$\mathcal{S}(m, M) \cap \{\boldsymbol{\Sigma} \in \mathbb{S}_+^{d,r} : \|\boldsymbol{\Sigma} - \mathbf{S}\|_F \leq \rho\}. \tag{C.51}$$

□

Remark 19. We note that the same proof extends this to the sample setting with high probability in an analogous way to how Lemma 12 extends Lemma 11 to the sample setting.

C.2.7 Proof of Theorem 4

We first restate the theorem for ease of reference.

Theorem 4 Suppose that we observe $y_i = \langle \mathbf{x}_i \mathbf{x}_i^\top, \mathbf{S} \rangle$, $i = 1, \dots, n$, where $\mathbf{x}_i$ and $\mathbf{S}$ satisfy Assumption 1. Suppose further that $r = \text{rank}(\mathbf{S}) \geq 5$, and let $\mathbf{S}_n = \mathbf{C}_n^{1/2} \mathbf{S} \mathbf{C}_n^{1/2}$. Then, for constants c_1, c_2 , if $n \gtrsim dr$, with probability at least $1 - \exp(-c_2 n)$,

1. $\mathbf{S}_n$ is the unique global minimizer of F over $\mathbb{S}_+^d$.
2. Let $\mathbf{U}_0$ be the initial iterate of BWGD. If $\lambda_1(\mathbf{U}_0 \mathbf{U}_0^\top) \leq M$, and $F(\mathbf{U}_0 \mathbf{U}_0^\top) - F(\mathbf{S}_n) \leq \frac{c_1^5 m^8 (r-2)^2 \beta^{10}}{6^5 M^{15/2} d^3 r^2}$ for an $\mathbf{S}$ dependent constant β , then BWGD with step size 1 satisfies the following bound for $\mathcal{C} = O(c_1 m / M^{3/2})$.

$$\begin{aligned}
 & d_{\text{BW}}^2(\mathbf{U}_k \mathbf{U}_k^\top, \mathbf{C}_n^{1/2} \mathbf{S} \mathbf{C}_n^{1/2}) \\
 & \leq (1 - \mathcal{C})^k (F(\mathbf{U}_0 \mathbf{U}_0^\top) - F(\mathbf{S})).
 \end{aligned}$$

Theorem 4. The first statement follows by the Euclidean strong convexity in Lemma 12.

To prove the second statement, suppose that we initialize such that $m \leq \lambda_r(\mathbf{U}_0 \mathbf{U}_0^\top) \leq \lambda_1(\mathbf{U}_0 \mathbf{U}_0^\top) \leq M$, and

$$\begin{aligned} F(\mathbf{U}_0 \mathbf{U}_0^\top) - F(\mathbf{S}) &\leq \left(\frac{m^2(r-2)}{d^2 r} \frac{c_1^2 \beta^4 m^2}{36 M^3} \right)^2 \cdot \frac{c_1 \beta^2 d}{6 M^{3/2}} \\ &= \frac{c_1^5 m^8 (r-2)^2 \beta^{10}}{6^5 M^{15/2} d^3 r^2}. \end{aligned} \quad (\text{C.52})$$

First, by Euclidean strong convexity, we have

$$\|\mathbf{\Sigma} - \mathbf{S}\|_F \leq \sqrt{(F(\mathbf{\Sigma}) - F(\mathbf{S})) \frac{6 M^{3/2}}{c_1 \beta^2 d}}. \quad (\text{C.53})$$

Therefore, the initialization condition on $F(\mathbf{U}_0 \mathbf{U}_0^\top) - F(\mathbf{S})$ is enough to guarantee that

$$\|\mathbf{U}_0 \mathbf{U}_0^\top - \mathbf{S}\| \leq \frac{m^2(r-2)}{d^2 r} \frac{c_1^2 \beta^4 m^2}{36 M^3}, \quad (\text{C.54})$$

and therefore Proposition 18 holds. Furthermore, by Lemma 17, $F(\mathbf{U}_k \mathbf{U}_k^\top) \leq F(\mathbf{U}_{k-1} \mathbf{U}_{k-1}^\top)$ for all k , and thus

$$\|\mathbf{U}_k \mathbf{U}_k^\top - \mathbf{S}\| \leq \frac{m^2(r-2)}{d^2 r} \frac{2 c_1^2 \beta^4 m^2}{36 M^3} \quad (\text{C.55})$$

for all k , and the iterates remain in the ball of strong geodesic convexity.

Using the strong geodesic convexity of Proposition 18 and Lemma 17, it is then a standard argument to show linear convergence. We can apply, for example, [Zhang and Sra, 2016, Theorem 15] to yield the result. $\square$

C.3 Proof of Theorem 7

We give a proof of this theorem by following Ghadimi and Lan [2013].

Let $T_{\mathbf{\Sigma}_0 \rightarrow \mathbf{\Sigma}_1} = T_{N(\mathbf{0}, \mathbf{\Sigma}_0) \rightarrow N(\mathbf{0}, \mathbf{\Sigma}_1)}$ be the optimal transport map between the measures $N(\mathbf{0}, \mathbf{\Sigma}_0)$ and $N(\mathbf{0}, \mathbf{\Sigma}_1)$. For simplicity, we will write $\gamma_{\mathbf{\Sigma}} = N(\mathbf{0}, \mathbf{\Sigma})$. Consider a observing a sequence of rank one matrices $y_1 \mathbf{x}_1 \mathbf{x}_1^\top, \dots, y_n \mathbf{x}_n \mathbf{x}_n^\top$. Let

$$\begin{aligned} \mathbf{\Sigma}_1 &= \left((1-\eta) \mathbf{I} + \eta \frac{\sqrt{y_1} \mathbf{x}_1 \mathbf{x}_1^\top}{\text{Tr}(\mathbf{x}_1 \mathbf{x}_1^\top \mathbf{\Sigma}_0)^{1/2}} \right) \mathbf{\Sigma}_0 \\ &\quad \cdot \left((1-\eta) \mathbf{I} + \eta \frac{\sqrt{y_1} \mathbf{x}_1 \mathbf{x}_1^\top}{\text{Tr}(\mathbf{x}_1 \mathbf{x}_1^\top \mathbf{\Sigma}_0)^{1/2}} \right). \end{aligned} \quad (\text{C.56})$$

or written in terms of measures,

$$\gamma_{\mathbf{\Sigma}_1} = [(1-\eta) \text{id} + \eta T_{\mathbf{\Sigma}_0 \rightarrow y_1 \mathbf{x}_1 \mathbf{x}_1^\top}]_{\#} \gamma_{\mathbf{\Sigma}_0}. \quad (\text{C.57})$$

Due to nonnegative curvature, we have for any rank one matrix $\mathbf{w} \mathbf{w}^\top$

$$\begin{aligned} W_2^2(\gamma_{\mathbf{\Sigma}_1}, \gamma_{\mathbf{w} \mathbf{w}^\top}) &\leq \|T_{\mathbf{\Sigma}_0 \rightarrow \mathbf{\Sigma}_1} - T_{\mathbf{\Sigma}_0 \rightarrow \mathbf{w} \mathbf{w}^\top}\|_{\mathbf{\Sigma}_0}^2 \\ &= \|(1-\eta) \mathbf{I} + \eta T_{\mathbf{\Sigma}_0 \rightarrow \mathbf{x}_1 \mathbf{x}_1^\top} - T_{\mathbf{\Sigma}_0 \rightarrow \mathbf{w} \mathbf{w}^\top}\|_{\mathbf{\Sigma}_0}^2 \\ &= \|(\mathbf{I} - T_{\mathbf{\Sigma}_0 \rightarrow \mathbf{w} \mathbf{w}^\top}) + \eta(T_{\mathbf{\Sigma}_0 \rightarrow \mathbf{x}_1 \mathbf{x}_1^\top} - \mathbf{I})\|_{\mathbf{\Sigma}_0}^2 \\ &= \|\mathbf{I} - T_{\mathbf{\Sigma}_0 \rightarrow \mathbf{w} \mathbf{w}^\top}\|_{\mathbf{\Sigma}_0}^2 + \eta^2 \|\mathbf{I} - T_{\mathbf{\Sigma}_0 \rightarrow \mathbf{x}_1 \mathbf{x}_1^\top}\|_{\mathbf{\Sigma}_0}^2 \\ &\quad - 2\eta \langle \mathbf{I} - T_{\mathbf{\Sigma}_0 \rightarrow \mathbf{w} \mathbf{w}^\top}, \mathbf{I} - T_{\mathbf{\Sigma}_0 \rightarrow \mathbf{x}_1 \mathbf{x}_1^\top} \rangle. \end{aligned} \quad (\text{C.58})$$

Taking the expectation with respect to $\mathbf{w} \mathbf{w}^\top = y \mathbf{x} \mathbf{x}^\top$, where $y = \mathbf{x}^\top \mathbf{S} \mathbf{x}$,

$$F(\mathbf{\Sigma}_1) - F(\mathbf{\Sigma}_0) \leq -2\eta \langle \nabla F(\mathbf{\Sigma}_0), \mathbf{I} - T_{\mathbf{\Sigma}_0 \rightarrow \mathbf{x}_1 \mathbf{x}_1^\top} \rangle_{\mathbf{\Sigma}_0} \quad (\text{C.59})$$

$$+ \eta^2 d_{\text{BW}}^2(\Sigma_0, y_1 \mathbf{x}_1 \mathbf{x}_1^\top).$$

Summing over iterations, we find

$$\begin{aligned} F(\Sigma_K) - F(\Sigma_0) &\leq \\ &- 2\eta \sum_{k=1}^K \langle \nabla F(\Sigma_{k-1}), \mathbf{I} - T_{\Sigma_{k-1} \rightarrow y_k \mathbf{x}_k \mathbf{x}_k^\top} \rangle_{\Sigma_{k-1}} \\ &+ \eta^2 \sum_{k=1}^K d_{\text{BW}}^2(\Sigma_{k-1}, y_k \mathbf{x}_k \mathbf{x}_k^\top). \end{aligned} \quad (\text{C.60})$$

Taking the expectation with respect to $y_1 \mathbf{x}_1 \mathbf{x}_1^\top, \dots, y_K \mathbf{x}_K \mathbf{x}_K^\top$ conditioned on $\Sigma_1, \dots, \Sigma_{K-1}$, and using $\mathbb{E}[d_{\text{BW}}^2(\Sigma_{k-1}, y_k \mathbf{x}_k \mathbf{x}_k^\top) | \Sigma_{k-1}] \leq \sigma^2$,

$$\mathbb{E}F(\Sigma_K) - F(\Sigma_0) \leq -2\eta \sum_{k=1}^K \|\nabla F(\Sigma_k)\|_{\Sigma_k}^2 + \eta^2 K \sigma^2. \quad (\text{C.61})$$

With this,

$$\sum_{k=1}^K \|\nabla F(\Sigma_k)\|_{\Sigma_k}^2 \leq \mathbb{E} \frac{F(\Sigma_K) - F(\Sigma_0)}{2\eta} + \frac{\eta}{2} K \sigma^2. \quad (\text{C.62})$$

Choosing $\eta = 1/\sqrt{K}$,

$$\frac{1}{K} \sum_{k=1}^K \|\nabla F(\Sigma_k)\|_{\Sigma_k}^2 \leq \mathbb{E} \frac{F(\Sigma_K) - F(\Sigma_0)}{2\sqrt{K}} + \frac{\eta \sigma^2}{2\sqrt{K}}. \quad (\text{C.63})$$

Within the sequence of iterates $\Sigma_0, \dots, \Sigma_K$, at least one element must have gradient bounded by $O(1/\sqrt{K})$.

C.4 Suboptimal Stationary Points

There potentially exist stationary points that are not optimal in the low-rank case. In particular, with the parametrization $\Sigma = \mathbf{U}\mathbf{U}^\top$, these are points such that

$$\mathbb{E} \frac{\sqrt{\mathbf{x}^\top \mathbf{V} \mathbf{V}^\top \mathbf{x}}}{\sqrt{\mathbf{x}^\top \mathbf{U} \mathbf{U}^\top \mathbf{x}}} \mathbf{x} \mathbf{x}^\top \mathbf{U} = \mathbf{U}, \quad (\text{C.64})$$

where $\mathbf{V} \mathbf{V}^\top = \mathbf{S}$.

In the following, we will show that at least in the $r = 1$ case, there are no local minima $\mathbf{u} \mathbf{u}^\top$ that are not orthogonal to $\mathbf{v} \mathbf{v}^\top$.

C.4.1 No Local Minima in 1D Case

Theorem 20. *Consider the observation model with the rank one matrix $\mathbf{S} = \mathbf{v} \mathbf{v}^\top$ and $y = \langle \mathbf{x} \mathbf{x}^\top, \mathbf{S} \rangle$. Then, the only fixed points of the iteration*

$$\mathbb{E} \sqrt{\frac{\mathbf{x}^\top \mathbf{v} \mathbf{v}^\top \mathbf{x}}{\mathbf{x}^\top \mathbf{u} \mathbf{u}^\top \mathbf{x}}} \mathbf{x} \mathbf{x}^\top \mathbf{u} = \mathbf{u} \quad (\text{C.65})$$

are orthogonal to $\mathbf{v}$ or $\mathbf{u} = \mathbf{v}$. Since the points orthogonal to $\mathbf{v}$ are local maxima, population gradient descent converges to $\mathbf{v}$.

Proof. In the 1-dimensional case, the stationary points are

$$\mathbb{E} \frac{|\mathbf{v}^\top \mathbf{x}|}{|\mathbf{u}^\top \mathbf{x}|} \mathbf{x} \mathbf{x}^\top \mathbf{u} = \mathbf{u}. \quad (\text{C.66})$$

Obviously, $\mathbf{u} = \mathbf{v}$ is a stationary point.

On the other hand, suppose that $u \perp v$. Then, we can write

$$\begin{aligned}
 & \mathbb{E} \left[\frac{|v^\top x|}{|u^\top x|} x x^\top u \right] \\
 &= \mathbb{E} \left[\frac{|v^\top x|}{|u^\top x|} (u u^\top / \|u\|^2 + v v^\top / \|v\|^2 + w w^\top) x x^\top u \right] \\
 &= \frac{1}{\|u\|^2} \mathbb{E} \left[\frac{|v^\top x|}{|u^\top x|} u u^\top x x^\top u \right] + \frac{1}{\|v\|^2} \mathbb{E} \left[\frac{|v^\top x|}{|u^\top x|} v v^\top x x^\top u \right] \\
 &= \left(\frac{1}{\|u\|^2} \mathbb{E} |v^\top x| \mathbb{E} |u^\top x| \right) u + \frac{1}{\|v\|^2} \left(\mathbb{E} (v^\top x)^2 \mathbb{E} \frac{x^\top u}{|u^\top x|} \right) v \\
 &= \frac{2}{\pi} \|v\| \frac{u}{\|u\|}.
 \end{aligned}$$

Therefore, for this to be a stationary point, we need

$$\frac{2}{\pi} \frac{\|v\|}{\|u\|} = 1, \text{ or } \|u\| = \frac{2\|v\|}{\pi}. \quad (\text{C.67})$$

Thus, any orthogonal vector with this length is a stationary point.

Finally, we show that there are no other stationary points. For u to be a fixed point, we need to have that

$$\mathbb{E} |v^\top x| \text{sign}(u^\top x) x = u. \quad (\text{C.68})$$

Due to the rotational symmetry of the x 's, we can assume without loss of generality that only the first two coordinates of v , u are nonzero. Furthermore, we can reduce to the two dimensional case, since the coordinates $x_3, \dots, x_d$ do not contribute to the expectation. Finally, we can assume without loss of generality that u and v are rotated so that $u = [a, 0]^\top$. In this way, if $v = (v_1, v_2)$, (C.68) becomes

$$\mathbb{E} |v_1 x_1 + v_2 x_2| \text{sign}(x_1) [x_1, x_2]^\top. \quad (\text{C.69})$$

The first coordinate is obviously positive. If we can show that the second coordinate is nonzero, then u cannot be a fixed point. We assume without loss of generality that $v_2 > 0$.

Thus we consider

$$\mathbb{E} |v_1 x_1 + v_2 x_2| \text{sign}(x_1) x_2 \quad (\text{C.70})$$

for i.i.d. $N(0, 1)$ random variables x_1, x_2 . By symmetry, (x_1, x_2) occurs with the same probability as $(-x_1, -x_2)$, and

$$\begin{aligned}
 & |v_1(-x_1) + v_2(-x_2)| \text{sign}(-x_1)(-x_2) = \\
 & |v_1 x_1 + v_2 x_2| \text{sign}(x_1) x_2.
 \end{aligned} \quad (\text{C.71})$$

Therefore, we can integrate over any half-plane, and so

$$\begin{aligned}
 & \mathbb{E} |v_1 x_1 + v_2 x_2| \text{sign}(x_1) x_2 = \\
 & \mathbb{E}_{(x_1, x_2) | x_1 > 0} |v_1 x_1 + v_2 x_2| x_2.
 \end{aligned} \quad (\text{C.72})$$

For all fixed $x_1 > 0$, it is easy to see that

$$\mathbb{E}_{x_2 | x_1 > 0} |v_1 x_1 + v_2 x_2| x_2 > 0, \quad (\text{C.73})$$

since $|v_1 x_1 + |v_2 x_2|| > |v_1 x_1 - |v_2 x_2||$ when $v_2 > 0$. Therefore the second coordinate cannot be zero, and this means that u is not a fixed point.

Together with the monotonicity of gradient descent, which implies convergence to a fixed point, we conclude that population gradient descent in the 1D case converges to the underlying vector v . $\square$

C.5 Highly Local Recovery in Discrete 1D Case

Suppose that we have a discrete set of sensing vectors $\mathbf{x}_1, \dots, \mathbf{x}_n$ that satisfy the ℓ_2/ℓ_1 -RIP condition, and that $\frac{1}{n} \sum_i \mathbf{x}_i \mathbf{x}_i^\top = \mathbf{I}$. Suppose that $\mathbf{S} = \mathbf{v} \mathbf{v}^\top$. Define the set

$$\mathcal{B} = \{\mathbf{u} : |\mathbf{u}^\top \mathbf{x}_i| > \epsilon \forall i\}. \quad (\text{C.74})$$

This is an open set. Over this set, we can bound

$$\begin{aligned} \|\nabla F(\mathbf{u}) - \nabla F(\mathbf{v})\| &\lesssim \|\mathbf{v} \mathbf{v}^\top - \mathbf{u} \mathbf{u}^\top\|_F^{1/2} \frac{1}{n} \sum_i \frac{\|\mathbf{x}_i\|}{|\mathbf{u}^\top \mathbf{x}_i|} \\ &\leq C \|\mathbf{v} \mathbf{v}^\top - \mathbf{u} \mathbf{u}^\top\|_F^{1/2}, \end{aligned} \quad (\text{C.75})$$

for some C that depends on all parameters. Assume that $\mathbf{v} \in \mathcal{B}$. This implies that there exists a ball around $\mathbf{v}$ that is contained within $\mathcal{B}$. In this ball, we can get the local Euclidean smoothness bound given in Section C.2.3. In turn, this implies local geodesic strong convexity over a small subset of this ball, which implies local linear convergence.

C.6 Lack of Strong Geodesic Convexity

We illustrate here that the functional F , while being locally strongly convex about $\mathbf{S}$ when restricted to rank r matrices, is not locally strongly convex around $\mathbf{S}$ for higher rank matrices.

Let $\mathbf{S}$ be the matrix

$$\mathbf{S} = \begin{bmatrix} a & 0 \\ 0 & 0 \end{bmatrix}, \quad (\text{C.76})$$

Σ_0 be the matrix

$$\Sigma_0 = \begin{bmatrix} a & 0 \\ 0 & b \end{bmatrix}, \quad (\text{C.77})$$

and let

$$\mathbf{T} = \begin{bmatrix} 1 & 0 \\ 0 & 0 \end{bmatrix} \quad (\text{C.78})$$

be the transport map from Σ_0 to $\mathbf{S}$. Let Σ_t be the geodesic from $\mathbf{S}$. Using the first display in (C.47), we find

$$\begin{aligned} \partial_t^2 F(\Sigma_t)|_{t=0} &= \left\langle (\mathbf{T} - \mathbf{I}) \Sigma_0 (\mathbf{T} - \mathbf{I}), \nabla F(\Sigma_0) \right\rangle \\ &\quad + \mathbb{E}_{\mathbf{x} \mathbf{x}^\top \sim Q} 2 \frac{\text{Tr}(\mathbf{x} \mathbf{x}^\top (\mathbf{T} - \mathbf{I}) \Sigma_0)^2}{(\text{Tr}(\mathbf{x} \mathbf{x}^\top \Sigma_0))^{3/2}} \\ &= \left\langle \begin{bmatrix} 0 & 0 \\ 0 & b \end{bmatrix}, \nabla F(\Sigma_0) \right\rangle + \mathbb{E}_{\mathbf{x} \mathbf{x}^\top \sim Q} 2 \frac{\text{Tr}(\mathbf{x} \mathbf{x}^\top (\mathbf{T} - \mathbf{I}) \Sigma_0)^2}{(\text{Tr}(\mathbf{x} \mathbf{x}^\top \Sigma_0))^{3/2}} \end{aligned} \quad (\text{C.79})$$

By a trace inequality,

$$2 \mathbb{E}_Q \frac{(\text{Tr}[\mathbf{x} \mathbf{x}^\top (\mathbf{T} - \mathbf{I}) \Sigma_0^{1/2}])^2}{\text{Tr}[\mathbf{x} \mathbf{x}^\top \Sigma_0]^{3/2}} \lesssim \|(\mathbf{T} - \mathbf{I}) \Sigma_0\|_F^2. \quad (\text{C.80})$$

To have geodesic strong convexity, we need to lower bound $\partial_t^2 F(\Sigma_t)|_{t=0}$ by $c \|(\mathbf{T} - \mathbf{I}) \Sigma_0^{1/2}\|_F^2 = c d_{\text{BW}}(\Sigma_0, \Sigma_1)^2$ for some $c > 0$. On the other hand, using the result of Section C.2.3, $\|\nabla F(\Sigma_0)\|_F \lesssim \|\Sigma_0 - \mathbf{S}\|_F = |b|$, and so

$$\partial_t^2 F(\Sigma_t)|_{t=0} \lesssim b^2 + \|(\mathbf{T} - \mathbf{I}) \Sigma_0\|_F^2. \quad (\text{C.81})$$

We note that

$$\|(\mathbf{T} - \mathbf{I}) \Sigma_0\|_F^2 = b^2.$$

On the other hand,

$$d_{\text{BW}}^2(\Sigma_0, \Sigma_1) = \|(\mathbf{T} - \mathbf{I}) \Sigma_0^{1/2}\|_F^2 = b,$$

There is no $c > 0$ such that

$$b^2 \geq cb \quad (\text{C.82})$$

for all $b > 0$. Therefore, F is not strongly geodesically convex at full-rank Σ_0 that are close the boundary. In particular, if the true barycenter is low-rank, then as the full-rank Σ approach $\bar{\Sigma}$, we cannot expect strong convexity.

D SUPPLEMENTAL EXPERIMENTS

D.1 Convergence of BWGD $M/m = 1$

We first repeat the experiment in Figure 2, except we generate the data using $S = VV^\top$, where V now has orthogonal columns. The results are displayed in Figure 7.

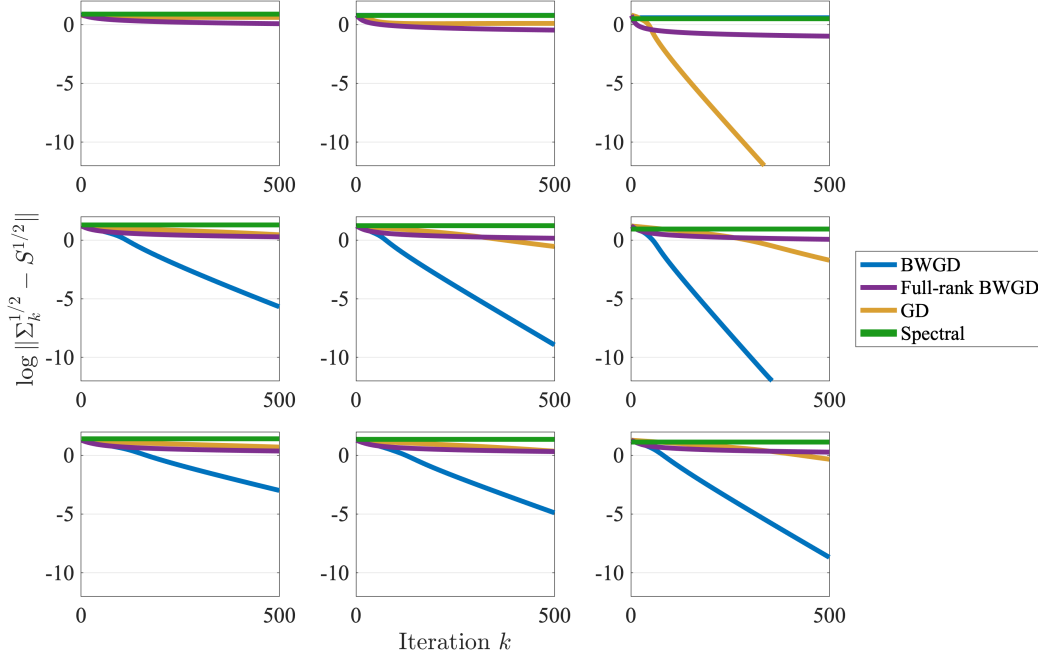


Figure 7: Convergence of BWGD (2.15) with the full gradient and GD [Li et al., 2019]. Here, the model is the same as in Figure 2, except the columns of V are orthogonal and length $\sqrt{d}$. We set $d = 32$, and down the columns we use $r = 1, 4, 16$, respectively. Across the rows we vary the number of points, with $n = 3dr, n = 10dr$, and $n = 20dr$, respectively. As we can see, for low to moderate ranks, the BWGD method converges much faster than the standard GD method.

D.2 Phase Retrieval on CelebA Data

Here, we repeat the experiment of Section 4.2 on 20 images taken from the CelebA dataset Liu et al. [2015]. We give two plots. First, we plot error versus iteration for both BWGD and Wirtinger Flow. Second, we give some examples of the generated images. As we can see, BWGD recovers sharp images much faster than Wirtinger flow across a range of examples.

D.3 Convergence of BWGD Versus BWSGD

In the plot for SGD, we also give lines to show the different rates. Here, it appears that SGD is converging to the true barycenter. Also, it appears to be converging at a faster than anticipated rate. An explanation of this phenomenon will be explored in future work. In the left plot of Figure 9, we set $d = 20$ and $n = d^2$, and we plot the error versus iteration for BWGD for the various ranks. As we can see, the convergence takes longer as the rank increases. In the right plot of Figure 9, we plot the error versus iteration for BWSGD using the single sample gradient of (2.14), where at each iteration we draw a new sample. As we can see, the BWSGD interpolates between two convergence regimes: a slow regime where the rate is $k^{-1/4}$ and a fast regime where the rate is $k^{-1/2}$. The latter rate is typical of cases where there is local strong convexity or a Polyak-Łojasiewicz inequality.

D.4 Abalone Dataset

We also present an experiment with real data. This example is one where we attempt to measure the heterogeneity in a regression dataset. Here, we pick the classical Abalone dataset available from the UCI machine learning repository [Blake

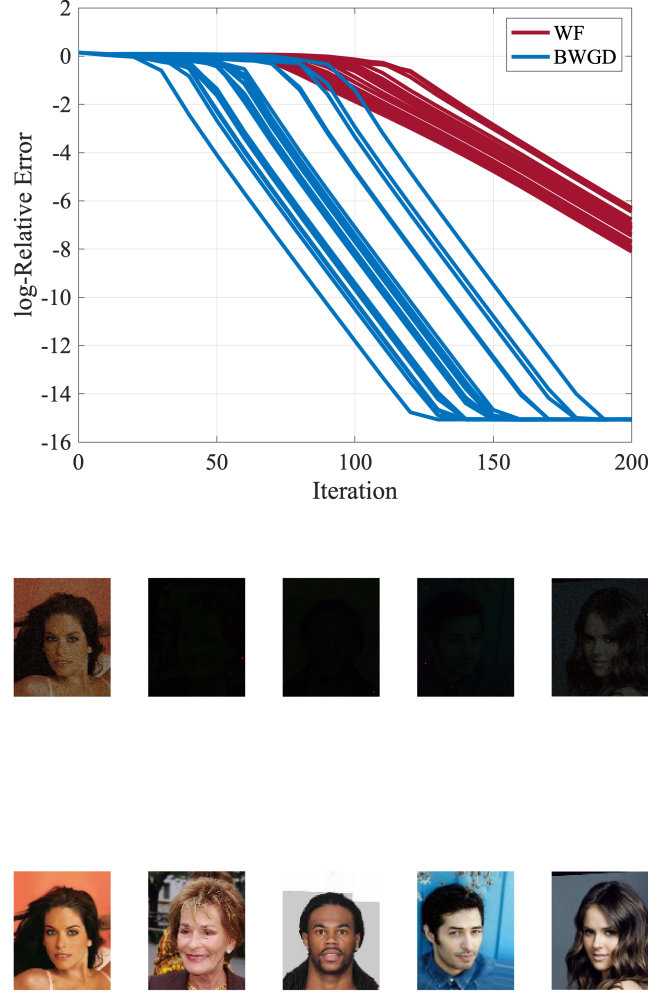


Figure 8: Convergence and example images of BWGD and WF on the CelebA Phase Retrieval experiment.

and Merz, 1998]. In this dataset, one attempts to predict the age of abalone from certain covariates. Linear regression models exhibit poor fit on this data for a variety of reasons. One reason, which we demonstrate here, is the presence of heterogeneity in the measurements.

Assuming the covariance recovery model $y_i = \langle \mathbf{x}_i, \mathbf{w}_i \rangle + \epsilon_i$, where $\mathbf{w}_i \sim N(\mathbf{0}, \mathbf{S})$, we could try to measure heterogeneity in the data by recovering the covariance of the regression vectors, and then plotting how $\mathbf{x}$ relates to y in its top principal space. Here, we recover such a covariance, and then project $\mathbf{x}$ onto the top two principal directions. We then plot y against these two directions. The resulting plots show varying degrees of heterogeneity depending on the method employed. Here, we compare BWGD, GD, PCA on the $\mathbf{x}$'s, and the spectral method, which finds the top principal directions of the matrix

$$S_n = \frac{1}{2n} \sum_{i=1}^n y_i (\mathbf{x}_i \mathbf{x}_i^\top - \mathbf{I}) \quad (\text{D.1})$$

As we can see, the spectral methods completely fails here. The BWGD method gives a similar but qualitatively different result from the Euclidean gradient descent method. In particular, there appear to be two primary directions of variation, which would indicate that this may be a mixture of two different regression components. Both BWGD and GD recover a stretched direction, but the secondary direction (in the PX_1 direction) appears to be more pronounced for BWGD.

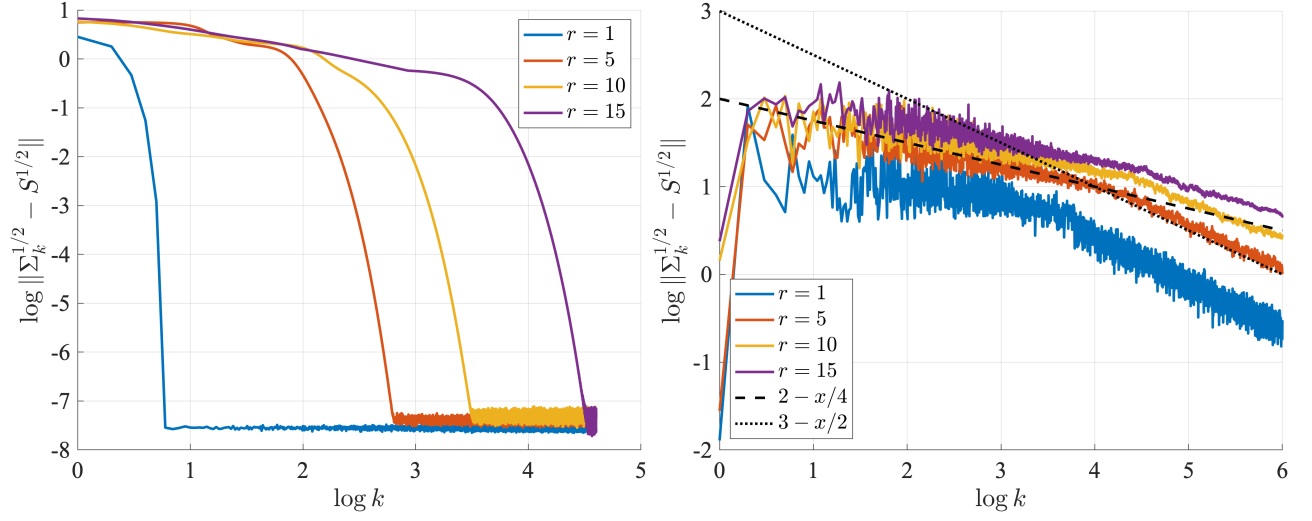


Figure 9: Convergence of BWGD and WF on the CelebA Phase Retrieval experiment.

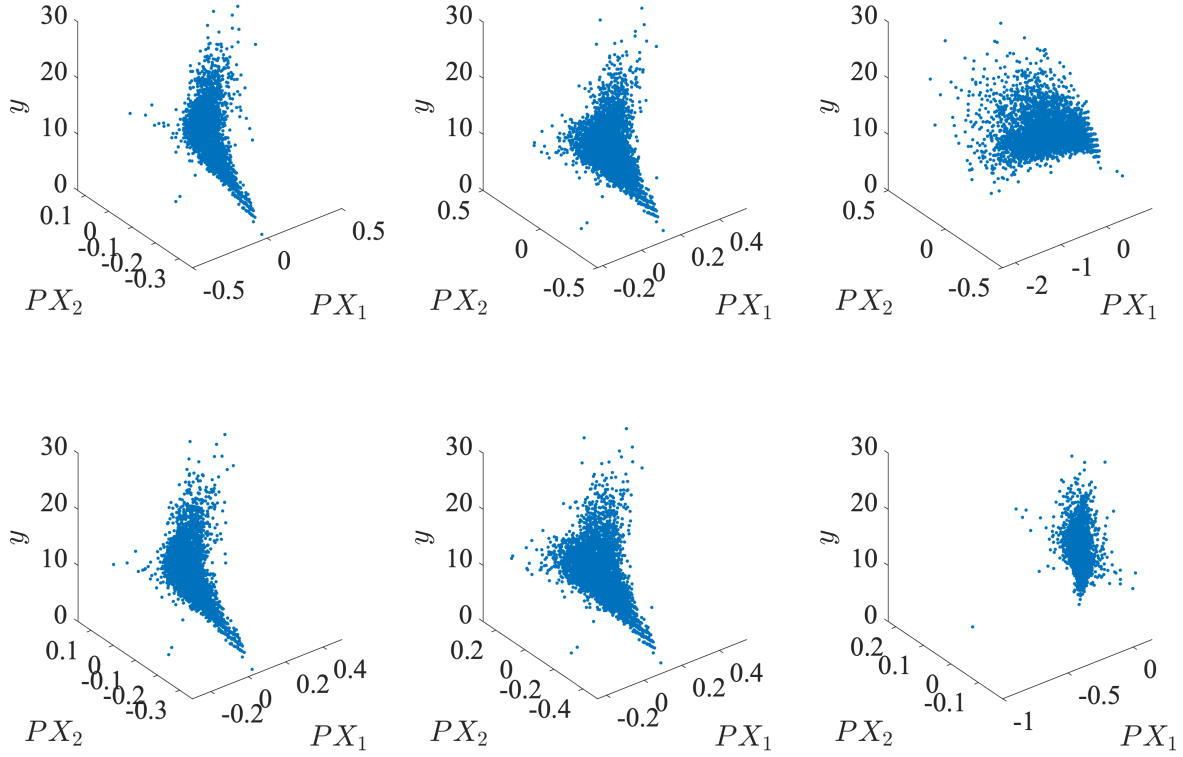


Figure 10: Projections for abalone data. The left column are projections for a fixed rank BW and full-rank BW barycenter, the middle column corresponds to low-rank and full-rank Euclidean gradient descent on (2.2), and the right column corresponds to the top principal subspace of $\mathbf{x}$ and the result of the spectral method.