

# ENHANCING DISTRIBUTIONAL ROBUSTNESS IN REGULARIZED PRINCIPAL COMPONENT ANALYSIS BY WASSERSTEIN DISTANCES\*

LEI WANG<sup>†</sup>, XIN LIU<sup>‡</sup>, AND XIAOJUN CHEN<sup>§</sup>

**Abstract.** We consider the distributionally robust optimization (DRO) model of regularized principal component analysis (PCA) to account for uncertainty in the underlying probability distribution. The resulting formulation leads to a constrained min-max optimization problem, where the ambiguity set captures the distributional uncertainty by the type-2 Wasserstein distance. We prove that the inner maximization problem admits a closed-form optimal value. This explicit characterization equivalently reformulates the original DRO model into a minimization problem on the Stiefel manifold. To solve the resulting problem, we devise an efficient smoothing manifold proximal gradient algorithm. Our analysis establishes Riemannian gradient consistency and global convergence of our algorithm to a stationary point of the reformulated minimization problem. We also provide the iteration complexity  $O(\epsilon^{-3})$  of our algorithm to achieve an  $\epsilon$ -approximate stationary point. Finally, numerical experiments are conducted to validate the effectiveness and scalability of our algorithm, as well as to highlight the necessity and rationality of adopting the DRO model for PCA.

**Key words.** distributionally robust optimization, principal component analysis, Riemannian manifold, smoothing approximation, Wasserstein distance

**MSC codes.** 65K05, 90C17, 90C26, 90C30, 90C47

**1. Introduction.** Let  $\xi \in \mathbb{R}^d$  be a  $d$ -dimensional random vector governed by a probability distribution  $\mathbb{P}_*$ . In this paper, we consider the following distributionally robust optimization (DRO) model of principal component analysis (PCA),

$$(1.1) \quad \min_{X \in \mathcal{O}^{d,r}} \sup_{\mathbb{P} \in \mathcal{P}} \mathbb{E}_{\mathbb{P}} \left[ \left\| (I_d - XX^\top) (\xi - \mathbb{E}_{\mathbb{P}}[\xi]) \right\|_F^2 \right] + s(X).$$

Here, the feasible set  $\mathcal{O}^{d,r} := \{X \in \mathbb{R}^{d \times r} \mid X^\top X = I_r\}$ , commonly referred to as the Stiefel manifold [1, 7, 50], consists of all the  $d \times r$  column-orthonormal matrices with  $1 \leq r < d$ . The ambiguity set  $\mathcal{P}$  represents a collection of distributions that could plausibly contain  $\mathbb{P}_*$  with high confidence. And  $s : \mathbb{R}^{d \times r} \rightarrow \mathbb{R}$  is a convex and Lipschitz continuous function acting as a regularizer to promote certain desired structures of solutions in  $\mathcal{O}^{d,r}$ , such as sparsity [10, 48] or nonnegativity [13, 52].

In most practical scenarios, the underlying distribution  $\mathbb{P}_*$  is unknown and can not be captured precisely, leaving us without the essential information required to solve the PCA problem exactly. Although the technique of sample average approximation provides a practical model, its solutions often suffer from poor out-of-sample performances when the sample size is limited [20]. This dilemma motivates us to

---

\*Submitted to the editors DATE.

**Funding:** This work is supported by National Key R&D Program of China (2023YFA1009300), RGC grant JLFS/P-501/24 for the CAS AMSS-PolyU Joint Laboratory in Applied Mathematics, CAS-Croucher Funding Scheme for Joint Laboratories, Hong Kong Research Grant Council project PolyU15300024, and National Natural Science Foundation of China (12125108, 12021001, 12288201).

<sup>†</sup>Department of Applied Mathematics, The Hong Kong Polytechnic University, Hong Kong, China (lei2wang@polyu.edu.hk).

<sup>‡</sup>State Key Laboratory of Mathematical Sciences, Academy of Mathematics and Systems Science, Chinese Academy of Sciences, and University of Chinese Academy of Sciences, Beijing, China (liuxin@lsec.cc.ac.cn).

<sup>§</sup>Department of Applied Mathematics, The Hong Kong Polytechnic University, Hong Kong, China (maxjchen@polyu.edu.hk).

investigate the DRO model (1.1) of PCA, which minimizes the worst case of the objective function across all distributions in  $\mathcal{P}$ .

In the realm of DRO [25, 36], there is a variety of ambiguity sets available, including those based on moment constraints [14, 17, 57], divergences [46, 56], and Wasserstein distances [20, 22], among others. For a thorough and insightful exposition of these concepts, we refer interested readers to a recent survey [36], which offers an in-depth exploration of DRO problems. Recently, Wasserstein DRO, where the discrepancy between probability measures is dictated by the Wasserstein distance, has garnered tremendous attentions across various domains [35, 36]. In a similar vein, we explore the use of the Wasserstein distance in constructing the ambiguity set  $\mathcal{P}$  in problem (1.1). Let  $\mathcal{Q}_p$  be the space of all probability distributions  $\mathbb{P}$  supported on  $\mathbb{R}^d$  with finite  $p$ -th moments, namely,  $\mathbb{E}_{\mathbb{P}}(\|\xi\|^p) = \int_{\mathbb{R}^d} \|\xi\|^p \mathbb{P}(d\xi) < \infty$ . Below is the definition of the Wasserstein distance on  $\mathcal{Q}_p$ .

**DEFINITION 1.1** ([34]). *The type- $p$  Wasserstein distance  $W_p : \mathcal{Q}_p \times \mathcal{Q}_p \rightarrow \mathbb{R}_+$  between two probability distributions  $\mathbb{P}_1 \in \mathcal{Q}_p$  and  $\mathbb{P}_2 \in \mathcal{Q}_p$  is defined as*

$$W_p(\mathbb{P}_1, \mathbb{P}_2) := \inf_{\mathbb{Q} \in \mathcal{J}(\mathbb{P}_1, \mathbb{P}_2)} \left( \mathbb{E}_{(\xi_1, \xi_2) \sim \mathbb{Q}} [\|\xi_1 - \xi_2\|^p] \right)^{1/p},$$

where  $\|\cdot\|_p$  represents the  $\ell_p$  norm on  $\mathbb{R}^d$ , and  $\mathcal{J}(\mathbb{P}_1, \mathbb{P}_2)$  is the set containing all the joint distributions of  $\xi_1$  and  $\xi_2$  with marginals  $\mathbb{P}_1$  and  $\mathbb{P}_2$ , respectively.

Throughout this paper, we are primarily interested in the case where  $p \in [1, 2]$ , which is of significant importance both theoretically and practically [21]. The corresponding Wasserstein DRO model of PCA can be formulated as

$$\begin{aligned} \min_{X \in \mathcal{O}^{d,r}} \sup_{\mathbb{P} \in \mathcal{Q}_p} \mathbb{E}_{\mathbb{P}} \left[ \left\| (I_d - XX^\top) (\xi - \mathbb{E}_{\mathbb{P}}[\xi]) \right\|_F^2 \right] + s(X) \\ \text{s. t. } \mathbb{P} \in \mathcal{B}_p(\mathbb{P}_o, \rho) := \{\mathbb{P} \in \mathcal{Q}_p \mid W_p(\mathbb{P}, \mathbb{P}_o) \leq \rho\}, \end{aligned}$$

where  $\mathbb{P}_o \in \mathcal{Q}_p$  is a nominal distribution conceived as being close to  $\mathbb{P}_*$  and  $\rho \geq 0$  is a radius of the Wasserstein ball. It is noteworthy that the performance guarantees of the ambiguity set  $\mathcal{B}_p(\mathbb{P}_o, \rho)$  defined in (P<sub>mM</sub>) can be inherited from the theoretical results in Wasserstein DRO [4, 5, 20, 21].

By exploiting the structure of the inner maximization problem, we find that the DRO model (P<sub>mM</sub>) is well-posed only when  $p = 2$ , as its optimal value becomes positive infinity for  $1 \leq p < 2$ . Moreover, the optimal value of the inner maximization problem can be computed in an explicit form for the case  $p = 2$ . This leads us to an equivalent reformulation of the DRO model (P<sub>mM</sub>) with  $p = 2$  as follows,

$$\min_{X \in \mathcal{O}^{d,r}} \text{tr} \left( (I_d - XX^\top) \Sigma_o \right) + s(X) + 2\rho \left\| (I_d - XX^\top) \Sigma_o^{1/2} \right\|_F$$

where  $\Sigma_o = \mathbb{E}_{\mathbb{P}_o}[(\xi - \mathbb{E}_{\mathbb{P}_o}[\xi])(\xi - \mathbb{E}_{\mathbb{P}_o}[\xi])^\top] \in \mathbb{R}^{d \times d}$  is the covariance matrix of  $\xi$  under the nominal distribution  $\mathbb{P}_o$ . It is evident that problem (P<sub>m</sub>) is significantly easier to solve than problem (P<sub>mM</sub>) in its original min-max form. Therefore, we turn our attention to developing efficient algorithms for tackling problem (P<sub>m</sub>).

There have been substantial advancements in nonsmooth optimization on Riemannian manifolds in recent decades, with the majority of efforts directed toward locally Lipschitz continuous functions. During this period, a diverse range of algorithms has emerged, including subgradient-oriented approaches [29, 30, 38], proximal point methods [3, 9], primal-dual frameworks [37, 59], proximal gradient algorithms

[10, 31, 54], smoothing approaches [2, 49, 51], and so on. However, due to the presence of two possibly nonsmooth terms, none of the existing algorithms can be deployed to solve problem  $(P_m)$  efficiently. In particular, the last term in problem  $(P_m)$  poses a formidable challenge, as its proximal operator lacks a closed-form solution and entails costly computations. Since it can be represented by a composition of a smooth mapping and a convex function, we may resort to the proximal linear algorithm [54] for handling this issue. However, this approach requires calculating the Jacobian of  $(I_d - XX^\top)\Sigma_o^{1/2}$ , which involves the square root of  $\Sigma_o$ . And the presence of another nonsmooth term renders the subproblem particularly difficult to tackle. The Riemannian subgradient method proposed in [38] is capable of directly solving problem  $(P_m)$ . Nevertheless, owing to its reliance solely on subgradient information, it suffers from slow convergence with an iteration complexity of  $O(\epsilon^{-4})$ . Consequently, solving problem  $(P_m)$  remains a challenging endeavor.

Within the scope of this work, proximal gradient methods share the closest theoretical and methodological connection with the present study. This class of algorithms is initially proposed in [10] on the Stiefel manifold, laying the foundation for a plethora of subsequent advancements [11]. Later on, Huang and Wei [31] extend this framework to general Riemannian manifolds. The crux of these approaches involves solving a proximal gradient subproblem on the tangent space. Another related line of work is the dynamic smoothing algorithm introduced by Beck and Rosset [2], which aims to solve a class of composite optimization problems based on Moreau envelopes. Although this smoothing-based perspective is related, our problem structure and algorithmic treatment are different.

In this paper, we present a comprehensive analysis of the DRO model  $(P_{mM})$  for PCA. The main contributions are summarized as follows.

- Our investigation reveals that, the optimal value of the inner maximization problem in  $(P_{mM})$  diverges to infinity when  $1 \leq p < 2$ , whereas for  $p = 2$ , it admits a closed-form expression. Accordingly, we shift our attention to the case  $p = 2$  in the subsequent analysis. This explicit characterization facilitates an equivalent reformulation of the original DRO model  $(P_{mM})$  into a minimization form  $(P_m)$ . A particularly intriguing discovery is that, for the classical PCA problem without regularizers, its distributionally robust counterpart is equivalent to the nominal model, uncovering an unexpected equivalence in this specific context.
- We design a novel smoothing method to solve a class of composite optimization problems on the Stiefel manifold. Our theoretical analysis elucidates the Riemannian gradient consistency for the smoothing function and establishes the global convergence of our algorithm to a stationary point. Furthermore, we provide the iteration complexity of  $O(\epsilon^{-3})$  required to achieve an  $\epsilon$ -approximate stationary point.
- Last but not least, preliminary experimental results present the numerical performance of the proposed algorithm, underscoring its practical viability. Moreover, these results validate the necessity and rationality of adopting the DRO model for PCA, demonstrating its superiority in addressing distributional uncertainty and enhancing solution robustness on three real-world datasets.

The rest of this paper proceeds as follows. Section 2 draws into some preliminaries of Riemannian optimization. In Section 3, we make a profound study on the DRO model of PCA and derive its equivalent reformulation. Section 4 discusses the stationarity condition and develops the smoothing algorithm. Its convergence analysis and iteration complexity are provided in Section 5. Numerical results are presented in Section 6. Finally, concluding remarks are given in Section 7.

**2. Preliminaries.** In this section, we introduce the notations and concepts used throughout this paper.

**2.1. Basic Notations.** We use  $\mathbb{R}$  and  $\mathbb{N}$  to denote the sets of real and natural numbers, respectively. And the notations  $\mathbb{R}_+$  and  $\mathbb{R}_{++}$  represent the sets of nonnegative and positive real numbers, respectively. The Euclidean inner product of two matrices  $Y_1, Y_2$  with the same size is defined as  $\langle Y_1, Y_2 \rangle = \text{tr}(Y_1^\top Y_2)$ , where  $\text{tr}(B)$  stands for the trace of a square matrix  $B$ . And the notation  $I_r \in \mathbb{R}^{r \times r}$  represents the  $r \times r$  identity matrix. The Frobenius norm of a given matrix  $C$  is denoted by  $\|C\|_F$ . We denote by  $\mathbb{S}_+^d$  and  $\mathbb{S}_{++}^d$  the spaces of symmetric positive semidefinite matrices and symmetric positive definite matrices in  $\mathbb{R}^{d \times d}$ , respectively. The notation  $B^{1/2}$  stands for the square root of a symmetric positive semidefinite matrix  $B$ .

**2.2. Riemannian Gradients and Clarke Subdifferentials.** Let  $\mathcal{M}$  be a complete submanifold embedded in  $\mathbb{R}^{d \times r}$ . For each point  $X \in \mathcal{M}$ , the tangent space to  $\mathcal{M}$  at  $X$  is referred to as  $\mathcal{T}_X \mathcal{M}$ . In this paper, we consider the Riemannian metric  $\langle \cdot, \cdot \rangle_X$  on  $\mathcal{T}_X \mathcal{M}$  that is induced from the Euclidean inner product  $\langle \cdot, \cdot \rangle$ , i.e.,  $\langle V_1, V_2 \rangle_X = \langle V_1, V_2 \rangle = \text{tr}(V_1^\top V_2)$  for any  $V_1, V_2 \in \mathcal{T}_X \mathcal{M}$ . The tangent bundle of  $\mathcal{M}$  is denoted by  $\mathcal{TM} = \{(X, V) \mid X \in \mathcal{M}, V \in \mathcal{T}_X \mathcal{M}\}$ , that is, the disjoint union of the tangent spaces of  $\mathcal{M}$ . Additionally, we use the notation  $\text{Proj}_{\mathcal{T}_X \mathcal{M}}(\cdot)$  to represent the orthogonal projection operator onto  $\mathcal{T}_X \mathcal{M}$ . In the context of the Stiefel manifold, the tangent space at  $X \in \mathcal{O}^{d,r}$  can be described as  $\mathcal{T}_X \mathcal{O}^{d,r} = \{D \in \mathbb{R}^{d \times r} \mid X^\top D + D^\top X = 0\}$ , and its orthogonal projection operator is given by  $\text{Proj}_{\mathcal{T}_X \mathcal{O}^{d,r}}(V) = V - X(X^\top V + V^\top X)/2$  for any  $V \in \mathbb{R}^{d \times r}$ .

For a smooth function  $f$ , the Riemannian gradient at  $X \in \mathcal{M}$ , denoted by  $\text{grad } f(X)$ , is defined as the unique element of  $\mathcal{T}_X \mathcal{M}$  satisfying

$$\langle \text{grad } f(X), V \rangle = Df(X)[V], \quad \forall V \in \mathcal{T}_X \mathcal{M},$$

where  $Df(X)[V]$  is the directional derivative of  $f$  along the direction  $V$  at the point  $X$ . Since  $f$  is defined on an embedded submanifold in the Euclidean space, its Riemannian gradient can be computed by  $\text{grad } f(X) = \text{Proj}_{\mathcal{T}_X \mathcal{M}}(\nabla f(X))$ .

For a locally Lipschitz continuous function on the manifold, the Riemannian Clarke subdifferential has been intensively studied and used in the literature [58, 30, 29], which is a natural extension of the Clarke subdifferential [16, 44] in the Euclidean space. Throughout this paper, we adopt the following definition for the Riemannian Clarke subdifferential.

**DEFINITION 2.1** ([30]). *Suppose that  $f : \mathcal{M} \rightarrow \mathbb{R}$  is a locally Lipschitz continuous function. Let  $\Omega_R(f) = \{X \in \mathcal{M} \mid f \text{ is differentiable at } X\}$ . Then the Riemannian Clarke subdifferential of  $f$  at  $X \in \mathcal{M}$  is defined as*

$$\partial_R f(X) = \text{conv} \{D \in \mathcal{T}_X \mathcal{M} \mid \text{grad } f(X_t) \rightarrow D, \Omega_R(f) \ni X_t \rightarrow X\}.$$

**2.3. Retractions.** In contrast to the Euclidean setting, the point  $X + V$  does not lie in the manifold in general for  $X \in \mathcal{M}$  and  $V \in \mathcal{T}_X \mathcal{M}$ , due to the absence of a linear structure in  $\mathcal{M}$ . The interplay between  $\mathcal{M}$  and  $\mathcal{T}_X \mathcal{M}$  is typically carried out via the exponential mappings, which are usually computationally intensive to evaluate in practice. As an alternative, the concept of retraction, a first-order approximation of the exponential mapping, is proposed in the literature [1, 7] to alleviate the heavy computational burden.

DEFINITION 2.2 ([1]). A retraction on a manifold  $\mathcal{M}$  is a smooth mapping  $\mathfrak{R} : \mathcal{T}\mathcal{M} \rightarrow \mathcal{M}$ , and the restriction of  $\mathfrak{R}$  to  $\mathcal{T}_X\mathcal{M}$ , denoted by  $\mathfrak{R}_X$ , satisfies the following two properties for any  $X \in \mathcal{M}$ .

- (i) It holds that  $\mathfrak{R}_X(0_X) = X$ , where  $0_X$  is the zero vector in  $\mathcal{T}_X\mathcal{M}$ .
- (ii) The differential of  $\mathfrak{R}_X$  at  $0_X$  is an identity map on  $\mathcal{T}_X\mathcal{M}$ .

By leveraging the retraction  $\mathfrak{R}_X(V)$ , we can obtain a point by moving away from  $X \in \mathcal{M}$  along the direction  $V \in \mathcal{T}_X\mathcal{M}$ , while remaining on the manifold. In this way, it defines an update rule to preserve the feasibility. Following the proof of Lemma 2.7 in [8], we know that the retraction satisfies the following properties.

LEMMA 2.3 ([8]). Let  $\mathcal{M}$  be a compact embedded submanifold of a Euclidean space. Then there exist two constants  $M_1 > 0$  and  $M_2 > 0$  such that the following two relationships hold,

$$\begin{cases} \|\mathfrak{R}_X(V) - X\|_F \leq M_1 \|V\|_F, \\ \|\mathfrak{R}_X(V) - (X + V)\|_F \leq M_2 \|V\|_F^2, \end{cases}$$

for any  $X \in \mathcal{M}$  and  $V \in \mathcal{T}_X\mathcal{M}$ .

There are various practical realizations of retractions on the Stiefel manifold, such as QR factorization, polar decomposition and Cayley transformation. We refer interested readers to [1, 7, 55] for more details.

**3. Model Analysis.** The DRO model ( $P_{\text{mM}}$ ) of PCA is investigated in this section. We focus on the inner maximization problem in ( $P_{\text{mM}}$ ) as follows,

$$\begin{aligned} \varphi(X) &:= \sup_{P \in \mathcal{B}_P(P_\circ, \rho)} \mathbb{E}_P \left[ \left\| (I_d - XX^\top) (\xi - \mathbb{E}_P[\xi]) \right\|_F^2 \right] \\ &= \sup_{P \in \mathcal{B}_P(P_\circ, \rho)} \left\{ \mathbb{E}_P \left[ \text{tr} \left( (I_d - XX^\top) \left( \xi \xi^\top - \mathbb{E}_P[\xi] \mathbb{E}_P[\xi]^\top \right) \right) \right] \right\}. \end{aligned}$$

The objective function of (3.1) is quadratic in the underlying distribution  $P$  rather than linear, and hence, existing reformulations of Wasserstein DRO problems [6, 15, 20, 22, 60] are not applicable. To navigate this challenge, we introduce an auxiliary variable  $\eta \in \mathbb{R}^d$  and propose the following splitting formulation of (3.1),

$$\begin{aligned} \sup_{\eta \in \mathcal{M}_P(P_\circ, \rho)} \sup_{P \in \mathcal{B}_P} \mathbb{E}_P \left[ \text{tr} \left( (I_d - XX^\top) \xi \xi^\top \right) \right] - \text{tr} \left( (I_d - XX^\top) \eta \eta^\top \right) \\ \text{s. t. } W_p(P, P_\circ) \leq \rho, \quad \mathbb{E}_P[\xi] = \eta, \end{aligned}$$

where  $\mathcal{M}_P(P_\circ, \rho) := \{\eta = \mathbb{E}_P[\xi] \mid W_p(P, P_\circ) \leq \rho\}$  collects all mean vectors of distributions in the ambiguity set  $\mathcal{B}_P(P_\circ, \rho)$ . The constraint  $\eta \in \mathcal{M}_P(P_\circ, \rho)$  is imposed to avoid an empty feasible region and to ensure the well-posedness of problem (3.2). For  $\tau \geq 0$ ,  $\eta \in \mathcal{M}_P(P_\circ, \rho)$ , and  $X \in \mathcal{O}^{d,r}$ , we define

$$\psi(\tau \mid \eta, X) := \sup_{P \in \mathcal{B}_P} \left\{ \mathbb{E}_P[\omega_X(\xi)] \mid W_p(P, P_\circ) \leq \tau, \mathbb{E}_P[\xi] = \eta \right\},$$

where  $\omega_X(\xi) := \text{tr} \left( (I_d - XX^\top) \xi \xi^\top \right)$ . Then it holds that

$$\varphi(X) = \sup_{\eta \in \mathcal{M}_P(P_\circ, \rho)} \left\{ \psi(\rho^p \mid \eta, X) - \text{tr} \left( (I_d - XX^\top) \eta \eta^\top \right) \right\}.$$

Based on the preceding constructions, the objective function of the outer minimization problem in the Wasserstein DRO model ( $P_{\text{mM}}$ ) can be expressed as  $\varphi(X) + s(X)$ .

**3.1. Dual Representation.** In this subsection, we aim to derive the dual representation of  $\psi(\tau \mid \eta, X)$  defined in (3.3). This part of analysis follows the idea of Zhang et al. [60] based on the Legendre transform [43]. We generalize existing results by handling an additional equality constraint  $\mathbb{E}_{\mathbb{P}}[\xi] = \eta$ . Although our focus is primarily on PCA, the techniques we propose can be naturally extended to more general settings.

The following lemma reveals some useful properties of the function  $\psi(\cdot \mid \eta, X)$  for fixed  $\eta \in \mathcal{M}_p(\mathbb{P}_o, \rho)$  and  $X \in \mathcal{O}^{d,r}$ .

**LEMMA 3.1.** *For fixed  $\eta \in \mathcal{M}_p(\mathbb{P}_o, \rho)$  and  $X \in \mathcal{O}^{d,r}$ , the function  $\psi(\cdot \mid \eta, X)$  is bounded from below, monotonically increasing, and concave on  $\mathbb{R}_+$ .*

*Proof.* The proof follows along the same lines as that of [60, Lemma 1] and is therefore omitted for brevity.  $\square$

For a function  $h : \mathbb{R} \rightarrow \mathbb{R} \cup \{+\infty\}$ , we denote by  $h^\diamond : \mathbb{R} \rightarrow \mathbb{R} \cup \{+\infty\}$  its Legendre transform  $h^\diamond(\lambda) := \sup_{\tau \in \mathbb{R}} \{\lambda\tau - h(\tau)\}$ . Then the dual representation of  $\psi(\tau \mid \eta, X)$  can be constructed by resorting to the Legendre transform.

**THEOREM 3.2.** *For any  $p \in [1, 2]$  and  $\tau > 0$ , it holds that*

$$(3.5) \quad \psi(\tau \mid \eta, X) = \inf_{\lambda \geq 0, \varsigma \in \mathbb{R}^d} \left\{ \lambda\tau + \varsigma^\top \eta + \mathbb{E}_{\xi_o \sim \mathbb{P}_o} \left[ \sup_{\xi \in \mathbb{R}^d} \bar{\omega}_X(\xi, \xi_o) \right] \right\},$$

where  $\bar{\omega}_X(\xi, \xi_o) := \omega_X(\xi) - \lambda \|\xi - \xi_o\|_p^p - \varsigma^\top \xi$ .

*Proof.* For notational simplicity, we use  $\psi(\cdot)$  in place of  $\psi(\cdot \mid \eta, X)$  throughout this proof. Let  $\lambda \in \mathbb{R}$ . If  $\lambda < 0$ , we have  $(-\psi)^\diamond(-\lambda) = \sup_{\tau \geq 0} \{-\lambda\tau + \psi(\tau)\} \geq \sup_{\tau \geq 0} \{-\lambda\tau + \psi(0)\} = +\infty$ . Then our focus is on the case where  $\lambda \geq 0$ . Taking the Legendre transform of  $-\psi$  leads to that

$$\begin{aligned} (-\psi)^\diamond(-\lambda) &= \sup_{\tau \geq 0} \{-\lambda\tau + \psi(\tau)\} \\ &= \sup_{\tau \geq 0} \sup_{\mathbb{P} \in \mathcal{Q}_p} \left\{ \mathbb{E}_{\mathbb{P}}[\omega_X(\xi)] - \lambda\tau \mid \mathbb{W}_p^p(\mathbb{P}, \mathbb{P}_o) \leq \tau, \mathbb{E}_{\mathbb{P}}[\xi] = \eta \right\} \\ &= \sup_{\mathbb{P} \in \mathcal{Q}_p} \sup_{\tau \geq 0} \left\{ \mathbb{E}_{\mathbb{P}}[\omega_X(\xi)] - \lambda\tau \mid \mathbb{W}_p^p(\mathbb{P}, \mathbb{P}_o) \leq \tau, \mathbb{E}_{\mathbb{P}}[\xi] = \eta \right\} \\ &= \sup_{\mathbb{P} \in \mathcal{Q}_p} \left\{ \mathbb{E}_{\mathbb{P}}[\omega_X(\xi)] - \lambda \mathbb{W}_p^p(\mathbb{P}, \mathbb{P}_o) \mid \mathbb{E}_{\mathbb{P}}[\xi] = \eta \right\}. \end{aligned}$$

According to Definition 1.1, it follows that

$$\begin{aligned} (3.6) \quad (-\psi)^\diamond(-\lambda) &= \sup_{\mathbb{P} \in \mathcal{Q}_p} \left\{ \mathbb{E}_{\mathbb{P}}[\omega_X(\xi)] - \lambda \inf_{\mathbb{Q} \in \mathcal{J}(\mathbb{P}, \mathbb{P}_o)} \mathbb{E}_{(\xi, \xi_o) \sim \mathbb{Q}} \left[ \|\xi - \xi_o\|_p^p \right] \mid \mathbb{E}_{\mathbb{P}}[\xi] = \eta \right\} \\ &= \sup_{\mathbb{P} \in \mathcal{Q}_p, \mathbb{Q} \in \mathcal{J}(\mathbb{P}, \mathbb{P}_o)} \left\{ \mathbb{E}_{\mathbb{P}}[\omega_X(\xi)] - \lambda \mathbb{E}_{(\xi, \xi_o) \sim \mathbb{Q}} \left[ \|\xi - \xi_o\|_p^p \right] \mid \mathbb{E}_{\mathbb{P}}[\xi] = \eta \right\} \\ &= \sup_{\mathbb{Q} \in \mathcal{J}(\mathbb{P}_o)} \left\{ \mathbb{E}_{(\xi, \xi_o) \sim \mathbb{Q}} \left[ \omega_X(\xi) - \lambda \|\xi - \xi_o\|_p^p \right] \mid \mathbb{E}_{(\xi, \xi_o) \sim \mathbb{Q}}[\xi] = \eta \right\}, \end{aligned}$$

where  $\mathcal{J}(\mathbb{P}_o)$  stands for the set containing all the joint distributions of  $\xi$  and  $\xi_o$  with second marginal  $\mathbb{P}_o$ . It is clear that both the objective function and the constraint of the last problem in (3.6) are linear with respect to  $\mathbb{Q}$ . Following the discussions in [45, Section 3] and [57, Section 2], we interpret it as a special instance of the problem of

232 moments. As a direct consequence of [57, Proposition 2.1], the Slater condition holds,  
 233 and hence, the strong duality [45, Proposition 3.4] prevails. Hence, we can proceed  
 234 to show that

$$\begin{aligned}
 (-\psi)^\diamond(-\lambda) &= \sup_{\mathbf{Q} \in \mathcal{J}(\mathbb{P}_o)} \left\{ \mathbb{E}_{(\xi, \xi_o) \sim \mathbf{Q}} \left[ \omega_X(\xi) - \lambda \|\xi - \xi_o\|_p^p \right] \mid \mathbb{E}_{(\xi, \xi_o) \sim \mathbf{Q}} [\xi] = \eta \right\} \\
 235 \quad &= \inf_{\varsigma \in \mathbb{R}^d} \left\{ \varsigma^\top \eta + \sup_{\mathbf{Q} \in \mathcal{J}(\mathbb{P}_o)} \mathbb{E}_{(\xi, \xi_o) \sim \mathbf{Q}} \left[ \omega_X(\xi) - \lambda \|\xi - \xi_o\|_p^p - \varsigma^\top \xi \right] \right\} \\
 &= \inf_{\varsigma \in \mathbb{R}^d} \left\{ \varsigma^\top \eta + \sup_{\mathbf{Q} \in \mathcal{J}(\mathbb{P}_o)} \mathbb{E}_{(\xi, \xi_o) \sim \mathbf{Q}} [\bar{\omega}_X(\xi, \xi_o)] \right\},
 \end{aligned}$$

236 where  $\varsigma \in \mathbb{R}^d$  is the Lagrangian multiplier associated with the equality constraint  
 237  $\mathbb{E}_{(\xi, \xi_o) \sim \mathbf{Q}} [\xi] = \eta$ . Lemma 3.1 illustrates that  $\psi$  is bounded from below, monotonically  
 238 increasing, and concave on  $\mathbb{R}_+$ . Hence, either  $\psi(\tau) < +\infty$  for all  $\tau \geq 0$  or  $\psi(\tau) = +\infty$   
 239 for all  $\tau > 0$ . In the former case, by invoking the result of [43, Theorem 12.2], we can  
 240 obtain that

$$\begin{aligned}
 \psi(\tau) &= -(-\psi)^\diamond(\tau) = -\sup_{\lambda \in \mathbb{R}} \{-\lambda\tau - (-\psi)^\diamond(-\lambda)\} = \inf_{\lambda \geq 0} \{\lambda\tau + (-\psi)^\diamond(-\lambda)\} \\
 241 \quad (3.7) \quad &= \inf_{\lambda \geq 0, \varsigma \in \mathbb{R}^d} \left\{ \lambda\tau + \varsigma^\top \eta + \sup_{\mathbf{Q} \in \mathcal{J}(\mathbb{P}_o)} \mathbb{E}_{(\xi, \xi_o) \sim \mathbf{Q}} [\bar{\omega}_X(\xi, \xi_o)] \right\},
 \end{aligned}$$

242 for all  $\tau > 0$ . In the latter case, it holds that  $(-\psi)^\diamond(-\lambda) = +\infty$  for all  $\lambda \geq 0$ , which  
 243 indicates that the relationship (3.7) is also valid. According to [60, Proposition 2],  
 244 the function  $\bar{\omega}_X$  satisfies the interchangeability principle. Then it follows that

$$245 \quad \psi(\tau) = \inf_{\lambda \geq 0, \varsigma \in \mathbb{R}^d} \left\{ \lambda\tau + \varsigma^\top \eta + \mathbb{E}_{\xi_o \sim \mathbb{P}_o} \left[ \sup_{\xi \in \mathbb{R}^d} \bar{\omega}_X(\xi, \xi_o) \right] \right\}.$$

246 The proof is completed.  $\square$

247 Theorem 3.2 also implies a remarkable result that the optimal value of the inner  
 248 maximization problem in the DRO model ( $\mathbb{P}_{\text{mM}}$ ) is always infinite for any  $p \in [1, 2)$   
 249 and  $\rho > 0$ .

250 *Remark 3.3.* Suppose that  $p \in [1, 2)$  and  $\rho > 0$ . Then we have

$$251 \quad \sup_{\xi \in \mathbb{R}^d} \bar{\omega}_X(\xi, \xi_o) = \sup_{\xi \in \mathbb{R}^d} \left\{ \text{tr}((I_d - XX^\top) \xi \xi^\top) - \lambda \|\xi - \xi_o\|_p^p - \varsigma^\top \xi \right\} = +\infty,$$

252 which, together with Theorem 3.2, implies that  $\psi(\rho^p \mid \eta, X) = +\infty$ . From the  
 253 relationship (3.4), it can be deduced that  $\varphi(X) = +\infty$  for all  $X \in \mathcal{O}^{d,r}$ .

254 **3.2. Equivalent Reformulation for  $p = 2$ .** This subsection is devoted to  
 255 deriving the equivalent reformulation ( $\mathbb{P}_{\text{m}}$ ) of the DRO model ( $\mathbb{P}_{\text{mM}}$ ) for the specific  
 256 situation where  $p = 2$ . To this end, we establish that the optimal value  $\varphi(X)$  of  
 257 problem (3.1) admits a closed-form formulation.

258 The following lemma first shows that the supremum  $\sup \{\bar{\omega}_X(\xi, \xi_o) \mid \xi \in \mathbb{R}^d\}$  in  
 259 the dual representation (3.5) of  $\psi(\tau \mid \eta, X)$  can be explicitly computed.

LEMMA 3.4. Suppose that  $p = 2$  and  $\rho > 0$ . If  $\lambda > 1$ , it holds that

$$\mathbb{E}_{\xi_o \sim \mathbb{P}_o} \left[ \sup_{\xi \in \mathbb{R}^d} \bar{\omega}_X(\xi, \xi_o) \right] = \frac{\lambda}{\lambda - 1} \text{tr} \left( (I_d - XX^\top) \mathbb{E}_{\mathbb{P}_o} [\xi_o \xi_o^\top] \right) + \theta_X(\lambda, \varsigma) - \varsigma^\top \eta,$$

where the function  $\theta_X$  is defined as

$$\theta_X(\lambda, \varsigma) = \frac{1}{4\lambda(\lambda - 1)} \varsigma^\top (\lambda I_d - XX^\top) \varsigma + \varsigma^\top \left( \eta - \frac{1}{\lambda - 1} (\lambda I_d - XX^\top) \mathbb{E}_{\mathbb{P}_o} [\xi_o] \right).$$

Moreover, if  $\lambda \in [0, 1]$ , we have

$$\mathbb{E}_{\xi_o \sim \mathbb{P}_o} \left[ \sup_{\xi \in \mathbb{R}^d} \bar{\omega}_X(\xi, \xi_o) \right] = +\infty.$$

*Proof.* Straightforward calculations yield that

$$\begin{aligned} \bar{\omega}_X(\xi, \xi_o) &= \omega_X(\xi) - \lambda \|\xi - \xi_o\|_2^2 - \varsigma^\top \xi \\ &= -\xi^\top ((\lambda - 1) I_d + XX^\top) \xi + (2\lambda \xi_o - \varsigma)^\top \xi - \lambda \xi_o^\top \xi_o, \end{aligned}$$

which is a quadratic function with respect to  $\xi \in \mathbb{R}^d$  for fixed  $\xi_o \in \mathbb{R}^d$ . We move on to investigate the following three cases.

**Case I:**  $\lambda \in [0, 1)$ . Since  $r < d$ , the matrix  $(\lambda - 1)I_d + XX^\top$  has at least one negative eigenvalue  $\lambda - 1 < 0$  associated with an eigenvector  $z \in \mathbb{R}^d$  with  $z^\top z = 1$ . Then for any  $t \in \mathbb{R}$ , we have

$$\bar{\omega}_X(tz, \xi_o) = t^2(1 - \lambda) + t(2\lambda \xi_o - \varsigma)^\top z - \lambda \xi_o^\top \xi_o.$$

As a result of  $\lambda \in [0, 1)$ , it holds that  $\sup_{t \in \mathbb{R}} \bar{\omega}_X(tz, \xi_o) = +\infty$ , which further implies that

$$\sup_{\xi \in \mathbb{R}^d} \bar{\omega}_X(\xi, \xi_o) = +\infty.$$

**Case II:**  $\lambda = 1$ . It is clear that there exists  $z \in \mathbb{R}^d$  with  $z^\top z = 1$  such that  $X^\top z = 0$  as  $r < d$ . Thus, it holds for all  $t \in \mathbb{R}$  that

$$\bar{\omega}_X(tz, \xi_o) = t(2\xi_o - \varsigma)^\top z - \xi_o^\top \xi_o.$$

By simple manipulations, we arrive at

$$\sup_{t \in \mathbb{R}} \bar{\omega}_X(tz, \xi_o) = \begin{cases} -\xi_o^\top \xi_o, & \text{if } (2\xi_o - \varsigma)^\top z = 0, \\ +\infty, & \text{otherwise.} \end{cases}$$

Consequently, it can be readily verified that

$$\mathbb{E}_{\xi_o \sim \mathbb{P}_o} \left[ \sup_{\xi \in \mathbb{R}^d} \bar{\omega}_X(\xi, \xi_o) \right] \geq \mathbb{E}_{\xi_o \sim \mathbb{P}_o} \left[ \sup_{t \in \mathbb{R}} \bar{\omega}_X(tz, \xi_o) \right] = +\infty.$$

**Case III:**  $\lambda > 1$ . In this case, the matrix  $(\lambda - 1)I_d + XX^\top$  is positive definite. Then  $\bar{\omega}_X(\xi, \xi_o)$  is strictly concave with respect to  $\xi \in \mathbb{R}^d$  for fixed  $\xi_o \in \mathbb{R}^d$ . Hence, we can



proceed to show that

$$\sup_{\xi \in \mathbb{R}^d} \bar{\omega}_X(\xi, \xi_\circ) = \left( \lambda \xi_\circ - \frac{1}{2} \varsigma \right)^\top \left( (\lambda - 1) I_d + X X^\top \right)^{-1} \left( \lambda \xi_\circ - \frac{1}{2} \varsigma \right) - \lambda \xi_\circ^\top \xi_\circ,$$

where the supremum is attained at  $\xi = ((\lambda - 1) I_d + X X^\top)^{-1} (\lambda \xi_\circ - \varsigma/2)$ . Moreover, according to the Sherman–Morrison–Woodbury formula [28, page 329], we have

$$((\lambda - 1) I_d + X X^\top)^{-1} = \frac{1}{\lambda - 1} I_d - \frac{1}{\lambda(\lambda - 1)} X X^\top.$$

Then a straightforward verification reveals that

$$\begin{aligned} \sup_{\xi \in \mathbb{R}^d} \bar{\omega}_X(\xi, \xi_\circ) &= \frac{\lambda}{\lambda - 1} \text{tr} \left( (I_d - X X^\top) \xi_\circ \xi_\circ^\top \right) - \frac{1}{\lambda - 1} \varsigma^\top (\lambda I_d - X X^\top) \xi_\circ \\ &\quad + \frac{1}{4\lambda(\lambda - 1)} \varsigma^\top (\lambda I_d - X X^\top) \varsigma, \end{aligned}$$

which completes the proof.  $\square$

Leveraging the result of the previous lemma, we proceed to prove that  $\psi(\rho^2 \mid \eta, X)$  can be expressed in an explicit form. Recall that  $\Sigma_\circ = \mathbb{E}_{\mathbb{P}_\circ}[(\xi - \mathbb{E}_{\mathbb{P}_\circ}[\xi])(\xi - \mathbb{E}_{\mathbb{P}_\circ}[\xi])^\top]$  is the covariance matrix of  $\xi$  under the nominal distribution  $\mathbb{P}_\circ$ .

**LEMMA 3.5.** *Suppose that  $p = 2$  and  $\rho > 0$ . Then, for any  $\eta \in \mathcal{M}_2(\mathbb{P}_\circ, \rho)$  and  $X \in \mathcal{O}^{d,r}$ , it holds that*

$$\begin{aligned} \psi(\rho^2 \mid \eta, X) &= \left( \left( \text{tr} \left( (I_d - X X^\top) \Sigma_\circ \right) \right)^{1/2} + \left( \rho^2 - \|\eta - \mathbb{E}_{\mathbb{P}_\circ}[\xi_\circ]\|_2^2 \right)^{1/2} \right)^2 \\ &\quad + \text{tr} \left( (I_d - X X^\top) \eta \eta^\top \right). \end{aligned}$$

*Proof.* Based on Lemma 3.4, we can restrict our discussion to the case where  $\lambda > 1$ . Moreover, it follows from Theorem 3.2 that

$$\psi(\rho^2 \mid \eta, X) = \inf_{\lambda > 1} \left\{ \lambda \rho^2 + \frac{\lambda}{\lambda - 1} \text{tr} \left( (I_d - X X^\top) \mathbb{E}_{\mathbb{P}_\circ}[\xi_\circ \xi_\circ^\top] \right) + \inf_{\varsigma \in \mathbb{R}^d} \theta_X(\lambda, \varsigma) \right\}.$$

It is clear that  $\theta_X(\lambda, \varsigma)$  is a quadratic function with respect to  $\varsigma \in \mathbb{R}^d$  for fixed  $\lambda > 1$ . Since  $\lambda I_d - X X^\top$  is positive definite, it holds that

$$\begin{aligned} \inf_{\varsigma \in \mathbb{R}^d} \theta_X(\lambda, \varsigma) &= -\frac{\lambda}{\lambda - 1} \text{tr} \left( (I_d - X X^\top) \mathbb{E}_{\mathbb{P}_\circ}[\xi_\circ] \mathbb{E}_{\mathbb{P}_\circ}[\xi_\circ]^\top \right) \\ &\quad + \text{tr} \left( (I_d - X X^\top) \eta \eta^\top \right) - \lambda \|\eta - \mathbb{E}_{\mathbb{P}_\circ}[\xi_\circ]\|_2^2, \end{aligned}$$

where the infimum is attained at  $\varsigma = 2\lambda \mathbb{E}_{\mathbb{P}_\circ}[\xi_\circ] - 2((\lambda - 1) I_d + X X^\top) \eta$ . By simple calculations, we can obtain that

$$\begin{aligned} &\lambda \rho^2 + \frac{\lambda}{\lambda - 1} \text{tr} \left( (I_d - X X^\top) \mathbb{E}_{\mathbb{P}_\circ}[\xi_\circ \xi_\circ^\top] \right) + \inf_{\varsigma \in \mathbb{R}^d} \theta_X(\lambda, \varsigma) \\ &= (\lambda - 1) \left( \rho^2 - \|\eta - \mathbb{E}_{\mathbb{P}_\circ}[\xi_\circ]\|_2^2 \right) + \frac{1}{\lambda - 1} \text{tr} \left( (I_d - X X^\top) \Sigma_\circ \right) \\ &\quad + \text{tr} \left( (I_d - X X^\top) \eta \eta^\top \right) + \rho^2 - \|\eta - \mathbb{E}_{\mathbb{P}_\circ}[\xi_\circ]\|_2^2 + \text{tr} \left( (I_d - X X^\top) \Sigma_\circ \right). \end{aligned}$$

For any  $\eta \in \mathcal{M}_2(\mathbb{P}_\circ, \rho)$ , there exists  $\mathbb{P} \in \mathcal{B}_2(\mathbb{P}_\circ, \rho)$  such that  $\eta = \mathbb{E}_{\mathbb{P}}[\xi]$ . It follows from [23, Theorem 2.1] that  $\|\eta - \mathbb{E}_{\mathbb{P}_\circ}[\xi_\circ]\|_2 \leq \mathbb{W}_2(\mathbb{P}, \mathbb{P}_\circ) \leq \rho$ . Then it can be readily verified that

$$\begin{aligned} \psi(\rho^2 \mid \eta, X) &= \inf_{\lambda > 1} \left\{ (\lambda - 1) \left( \rho^2 - \|\eta - \mathbb{E}_{\mathbb{P}_\circ}[\xi_\circ]\|_2^2 \right) + \frac{1}{\lambda - 1} \text{tr} \left( (I_d - XX^\top) \Sigma_\circ \right) \right\} \\ &\quad + \text{tr} \left( (I_d - XX^\top) \eta \eta^\top \right) + \rho^2 - \|\eta - \mathbb{E}_{\mathbb{P}_\circ}[\xi_\circ]\|_2^2 + \text{tr} \left( (I_d - XX^\top) \Sigma_\circ \right) \\ &= \left( \left( \text{tr} \left( (I_d - XX^\top) \Sigma_\circ \right) \right)^{1/2} + \left( \rho^2 - \|\eta - \mathbb{E}_{\mathbb{P}_\circ}[\xi_\circ]\|_2^2 \right)^{1/2} \right)^2 \\ &\quad + \text{tr} \left( (I_d - XX^\top) \eta \eta^\top \right). \end{aligned}$$

We complete the proof.  $\square$

We are now in a position to derive an explicit expression for  $\varphi(X)$  based on the relationship (3.4), as established in the following theorem.

**THEOREM 3.6.** *For any  $X \in \mathcal{O}^{d,r}$ ,  $\mathbb{P}_\circ \in \mathcal{Q}_2$ , and  $\rho \geq 0$ , the optimal value of problem (3.1) with  $p = 2$  has the following explicit expression,*

$$(3.8) \quad \varphi(X) = \left( \left( \text{tr} \left( (I_d - XX^\top) \Sigma_\circ \right) \right)^{1/2} + \rho \right)^2.$$

*Proof.* We first consider the case where  $\rho = 0$ . Then the feasible region  $\mathcal{B}_2(\mathbb{P}_\circ, 0)$  of problem (3.1) collapses to the singleton  $\{\mathbb{P}_\circ\}$ . As a result, we have

$$\varphi(X) = \text{tr} \left( (I_d - XX^\top) \Sigma_\circ \right),$$

which indicates that the relationship (3.8) holds for  $\rho = 0$ .

Next, our focus is on the case where  $\rho > 0$ . As a direct consequence of Lemma 3.5, we can proceed to show that

$$\begin{aligned} \varphi(X) &= \sup_{\eta \in \mathcal{M}_2(\mathbb{P}_\circ, \rho)} \left\{ \psi(\rho^2 \mid \eta, X) - \text{tr} \left( (I_d - XX^\top) \eta \eta^\top \right) \right\} \\ &= \sup_{\eta \in \mathcal{M}_2(\mathbb{P}_\circ, \rho)} \left( \left( \text{tr} \left( (I_d - XX^\top) \Sigma_\circ \right) \right)^{1/2} + \left( \rho^2 - \|\eta - \mathbb{E}_{\mathbb{P}_\circ}[\xi_\circ]\|_2^2 \right)^{1/2} \right)^2 \\ &= \left( \left( \text{tr} \left( (I_d - XX^\top) \Sigma_\circ \right) \right)^{1/2} + \rho \right)^2, \end{aligned}$$

where the supremum is attained at  $\eta = \mathbb{E}_{\mathbb{P}_\circ}[\xi_\circ]$ . The proof is completed.  $\square$

By expanding the square term in (3.8), it can be obtained that

$$\varphi(X) = \text{tr} \left( (I_d - XX^\top) \Sigma_\circ \right) + 2\rho \left( \text{tr} \left( (I_d - XX^\top) \Sigma_\circ \right) \right)^{1/2} + \rho^2.$$

For any  $X \in \mathcal{O}^{d,r}$ , we have

$$\begin{aligned} \left( \text{tr} \left( (I_d - XX^\top) \Sigma_\circ \right) \right)^{1/2} &= \left( \text{tr} \left( \Sigma_\circ^{1/2} (I_d - XX^\top) (I_d - XX^\top) \Sigma_\circ^{1/2} \right) \right)^{1/2} \\ &= \left\| (I_d - XX^\top) \Sigma_\circ^{1/2} \right\|_{\text{F}}. \end{aligned}$$

Based on this representation and Theorem 3.6, the DRO model ( $\text{P}_{\text{mM}}$ ) with  $p = 2$  can be equivalently reformulated as problem ( $\text{P}_{\text{m}}$ ).

*Remark 3.7.* The equivalent reformulation we derive for the case  $p = 2$  in Theorem 3.6 is similar to that presented in [47, Proposition 3.3] with a moment-based ambiguity set. The distinction of the present work is threefold. First, our ambiguity set is defined directly through the type-2 Wasserstein distance over probability measures, and the reformulation is derived from the original min-max problem without reducing the ambiguity description to moment information. Second, the nominal distribution  $\mathbb{P}_o$  is allowed to be a general distribution in  $\mathcal{Q}_2$ , rather than being restricted to a discrete empirical distribution in [47]. Third, the splitting argument in (3.2) handles the nonlinear dependence of the objective on  $\mathbb{P}$  and also clarifies why the model becomes ill-posed for  $p < 2$ . Finally, we note that the proposed splitting technique can be seamlessly generalized to broader settings beyond PCA.

We end this section by demonstrating that, the solutions of classical PCA without regularizers inherently possess robustness against data perturbations, as characterized by the type-2 Wasserstein distance.

**COROLLARY 3.8.** *Suppose that  $p = 2$  and  $s(X) = 0$  for all  $X \in \mathcal{O}^{d,r}$ . Then for any  $\Sigma_o \in \mathbb{S}_+^d$  and  $\rho \geq 0$ , the global minimizers of problem  $(P_{\text{mM}})$  coincide with those of the nominal model  $\min_{X \in \mathcal{O}^{d,r}} \text{tr}((I_d - XX^\top) \Sigma_o)$  for PCA.*

*Proof.* According to Theorem 3.6, we know that the DRO model  $(P_{\text{mM}})$  is equivalent to the problem of minimizing  $\varphi$  over  $\mathcal{O}^{d,r}$ . Let  $\kappa : \mathbb{R}_+ \rightarrow \mathbb{R}_+$  be a continuous function defined by  $\kappa(t) = (\sqrt{t} + \rho)^2$ . Then  $\kappa$  is monotonically increasing on  $\mathbb{R}_+$  for all  $\rho \geq 0$ . Consequently, the nonnegativity of  $\text{tr}((I_d - XX^\top) \Sigma_o)$  results in that

$$\arg \min_{X \in \mathcal{O}^{d,r}} \varphi(X) = \arg \min_{X \in \mathcal{O}^{d,r}} \kappa(\text{tr}((I_d - XX^\top) \Sigma_o)) = \arg \min_{X \in \mathcal{O}^{d,r}} \text{tr}((I_d - XX^\top) \Sigma_o),$$

which completes the proof.  $\square$

**4. Algorithm Design.** The purpose of this section is to devise an efficient algorithm to solve the equivalent reformulation  $(P_{\text{m}})$  of the DRO model  $(P_{\text{mM}})$  with  $p = 2$  based on the smoothing approximation technique. We consider a broader class of nonsmooth optimization problems of the following form,

$$(4.1) \quad \min_{X \in \mathcal{O}^{d,r}} f(X) := u(X) + s(X) + w(X),$$

where  $u$ ,  $s$ , and  $w$  are appropriate functions satisfying the following conditions.

- (i) The function  $u : \mathbb{R}^{d \times r} \rightarrow \mathbb{R}$  is continuously differentiable and its gradient  $\nabla u$  is Lipschitz continuous with the corresponding Lipschitz constant  $L_u > 0$ .
- (ii) The function  $s : \mathbb{R}^{d \times r} \rightarrow \mathbb{R}$  is convex and Lipschitz continuous with the corresponding Lipschitz constant  $L_s > 0$ .
- (iii) The function  $w : \mathbb{R}^{d \times r} \rightarrow \mathbb{R}$  is of the form  $w(X) = 2\rho \|(I_d - XX^\top) \Sigma_o^{1/2}\|_{\text{F}}$  with  $\Sigma_o \in \mathbb{S}_+^d$  and  $\rho > 0$ .

It is evident that model  $(P_{\text{m}})$  is a specific instance of problem (4.1) by identifying  $u(X) = \text{tr}((I_d - XX^\top) \Sigma_o)$ . As a direct consequence of the continuity of  $f$  over the compact manifold  $\mathcal{O}^{d,r}$ , there always exists an optimal solution of problem (4.1).

*Remark 4.1.* The function  $w$  is continuously differentiable at every  $X \in \mathcal{O}^{d,r}$  such that  $\|(I_d - XX^\top) \Sigma_o^{1/2}\|_{\text{F}} > 0$ . In particular, when  $\text{rank}(\Sigma_o) > r$ , this condition holds for any  $X \in \mathcal{O}^{d,r}$ , and hence  $w$  is smooth on the entire feasible set. Possible nonsmoothness occurs at zero-residual points satisfying  $(I_d - XX^\top) \Sigma_o^{1/2} = 0$ , which can arise when  $\text{rank}(\Sigma_o) \leq r$ .

The case of  $\text{rank}(\Sigma_o) \leq r$  is natural from the perspective of data-driven DRO [42]. As discussed in [20], the nominal distribution is often empirical and constructed from limited samples in practice. Therefore, a nominal covariance matrix with rank at most  $r$  does not imply that the true distribution is supported on an  $r$ -dimensional subspace. The Wasserstein ambiguity set is introduced to account for possible distributional perturbations, including mass or variance outside the observed nominal subspace. In this case, the consideration of regularized PCA remains meaningful, as it aims to select a loading matrix  $X$  with desirable structures among possibly many feasible or optimal solutions. Even in the regime where  $\text{rank}(\Sigma_o) > r$ , the gradient of  $w$  contains a factor  $\|(I_d - XX^\top)\Sigma_o^{1/2}\|_{\mathbb{F}}^{-1}$ , which becomes numerically unstable when the denominator is close to zero. The smoothing algorithm developed in this section provides a unified and stable framework that remains applicable without imposing an additional rank condition on  $\Sigma_o$ . The underlying methodology can be seamlessly extended to a wider class of optimization problems, such as the regularized  $\ell_1$ -norm PCA [40] and the DRO model of fair PCA [47].

**4.1. Stationarity Condition.** In this subsection, we establish the stationarity condition for local minimizers of problem (4.1). According to the discussions in [30, 58], a necessary condition for  $f$  to achieve a local minimum at  $X$  on  $\mathcal{O}^{d,r}$  is that

$$0 \in \partial_{\mathbb{R}} f(X).$$

Since  $u$  is smooth,  $s$  is convex, and  $w$  is a composition of a smooth mapping and a convex function, they are all weakly convex and hence regular [18]. As a result, the objective function  $f = u + s + w$  is also regular [16]. Then it follows from [58, Theorem 5.3] that  $\partial_{\mathbb{R}} f(X) = \text{grad } u(X) + \partial_{\mathbb{R}} s(X) + \partial_{\mathbb{R}} w(X)$  for any  $X \in \mathcal{O}^{d,r}$ .

Based on the above discussions, the stationarity condition of the optimization problem (4.1) can be stated as follows.

**DEFINITION 4.2.** *A point  $X_* \in \mathcal{O}^{d,r}$  is called a stationary point of problem (4.1) if the inclusion  $0 \in \text{grad } u(X_*) + \partial_{\mathbb{R}} s(X_*) + \partial_{\mathbb{R}} w(X_*)$  holds.*

When  $w$  is smooth at  $X$ , the term  $\partial_{\mathbb{R}} w(X)$  reduces to the singleton  $\{\text{grad } w(X)\}$ . Thus, the Clarke subdifferential formulation provides a unified stationarity concept.

**4.2. Smoothing Function.** To address the challenges posed by the possibly nonsmooth term  $w$ , we propose to leverage the smoothing approximation technique [12]. Specifically, we construct the smoothing function of  $w$  as follows,

$$(4.2) \quad \tilde{w}(X, \mu) = \begin{cases} w(X), & \text{if } w(X) \geq \mu\rho, \\ \frac{w^2(X)}{2\mu\rho} + \frac{\mu\rho}{2}, & \text{if } w(X) < \mu\rho, \end{cases}$$

where  $\mu > 0$  is a smoothing parameter. Interested readers can refer to [12] for more examples of smoothing functions.

When it is clear from the context, the Euclidean and Riemannian gradients of  $\tilde{w}(X, \mu)$  with respect to  $X$  are simply denoted by  $\nabla \tilde{w}(X, \mu)$  and  $\text{grad } \tilde{w}(X, \mu)$ , respectively. The following proposition reveals that the smoothing function  $\tilde{w}$  enjoys some favorable properties.

**PROPOSITION 4.3.** *The function  $\tilde{w} : \mathbb{R}^{d \times r} \times \mathbb{R}_{++} \rightarrow \mathbb{R}$  constructed in (4.2) satisfies the following conditions.*

(i) *For any  $\mu > 0$ ,  $\tilde{w}(\cdot, \mu)$  is continuously differentiable over  $\mathbb{R}^{d \times r}$ .*

- (ii) For any  $X \in \mathbb{R}^{d \times r}$ , it holds that  $\lim_{X' \rightarrow X, \mu \downarrow 0} \tilde{w}(X', \mu) = w(X)$ .  
 (iii) For any  $X \in \mathbb{R}^{d \times r}$  and  $\mu > 0$ , we have  $w(X) \leq \tilde{w}(X, \mu) \leq w(X) + \mu\rho/2$ .  
 (iv) There exists a constant  $M_w > 0$  such that  $\|\nabla \tilde{w}(X, \mu)\|_F \leq M_w$  for any  $X \in \mathcal{O}^{d,r}$  and  $\mu > 0$ .  
 (v) There exists a constant  $L_w > 0$  such that, for any  $\mu > 0$ ,  $\nabla \tilde{w}(\cdot, \mu)$  is Lipschitz continuous over  $\mathcal{O}^{d,r}$  with the corresponding Lipschitz constant  $L_w \mu^{-1}$ .

*Proof.* The proof is straightforward and is omitted here.  $\square$

We adopt the smoothing function  $\tilde{w}$  given in (4.2) for two key reasons. First, the evaluation of both function values and gradients for  $\tilde{w}$  circumvents the need to compute  $\Sigma_o^{1/2}$ . Second, this particular smoothing function satisfies Riemannian gradient consistency, which will be rigorously demonstrated in the next subsection. This property is crucial for guaranteeing the global convergence of our algorithm.

**4.3. Riemannian Subdifferential.** The algorithm proposed in this paper is based on the smoothing approximation technique. Thus, it is natural that the convergence result is closely tied to the specific smoothing function employed. Below is the definition of the Riemannian subdifferential associated with the smoothing function, which serves as a fundamental concept in our analysis.

**DEFINITION 4.4.** *The Riemannian subdifferential of  $w$  associated with the smoothing function  $\tilde{w}$  at  $X \in \mathcal{O}^{d,r}$  is defined as*

$$\partial_{\mathbb{R}}|_{\tilde{w}}w(X) = \{G \in \mathcal{T}_X \mathcal{O}^{d,r} \mid \text{grad } \tilde{w}(X_t, \mu_t) \rightarrow G, \mathcal{O}^{d,r} \ni X_t \rightarrow X, \mu_t \downarrow 0\}.$$

The following theorem establishes the Riemannian gradient consistency of the smoothing function  $\tilde{w}$ . By bridging two subdifferentials, this property plays a pivotal role in showing the global convergence of the proposed algorithm to a stationary point of problem (4.1).

**THEOREM 4.5.** *For the smoothing function  $\tilde{w}$  constructed in (4.2), it holds that*

$$\partial_{\mathbb{R}}|_{\tilde{w}}w(X) \subseteq \partial_{\mathbb{R}}w(X),$$

for any  $X \in \mathcal{O}^{d,r}$ .

*Proof.* We fix an arbitrary  $X \in \mathcal{O}^{d,r}$ . Let  $G \in \partial_{\mathbb{R}}|_{\tilde{w}}w(X)$ . Then there exist two sequences  $\{X_t\} \subseteq \mathcal{O}^{d,r}$  and  $\{\mu_t\} \subseteq \mathbb{R}_{++}$  with  $X_t \rightarrow X$  and  $\mu_t \downarrow 0$  as  $t \rightarrow \infty$  such that

$$G = \lim_{\mathcal{O}^{d,r} \ni X_t \rightarrow X, \mu_t \downarrow 0} \text{grad } \tilde{w}(X_t, \mu_t).$$

For convenience, we define the index set  $\mathbb{T} := \{t \in \mathbb{N} \mid w(X_t) < \mu_t \rho\}$ .

If  $\mathbb{T}$  is a finite set, there exists  $\bar{t} \in \mathbb{N}$  such that  $w(X_t) \geq \mu_t \rho > 0$  for any  $t \geq \bar{t}$ . Hence, the function  $w$  is continuously differentiable near  $X_t$  and we have

$$\text{grad } \tilde{w}(X_t, \mu_t) = \text{grad } w(X_t),$$

for all  $t \geq \bar{t}$ . Then it can be obtained that

$$G = \lim_{\mathcal{O}^{d,r} \ni X_t \rightarrow X, \mu_t \downarrow 0} \text{grad } \tilde{w}(X_t, \mu_t) = \lim_{\mathcal{O}^{d,r} \ni X_t \rightarrow X} \text{grad } w(X_t),$$

which indicates that  $G \in \partial_{\mathbb{R}}w(X)$ .

Next, we consider the case where  $\mathbb{T}$  is an infinite set. Then it is clear that  $w(X) = 0$ . Thus, the function  $w$  attains the global minimum at  $X$ , which further implies that  $0 \in \partial_{\mathbb{R}} w(X)$ . Let  $\mathbb{T}' := \{t \in \mathbb{N} \mid w(X_t) = 0\}$  be a subset of  $\mathbb{T}$ . If  $\mathbb{T}'$  is also an infinite set, we can obtain that

$$G = \lim_{\mathcal{O}^{d,r} \ni X_t \rightarrow X, \mu_t \downarrow 0, t \in \mathbb{T}'} \text{grad } \tilde{w}(X_t, \mu_t).$$

Straightforward calculations yield that  $\text{grad } \tilde{w}(X_t, \mu_t) = 0 \in \partial_{\mathbb{R}} w(X_t)$  for any  $t \in \mathbb{T}'$ . Hence, the above relationship indicates that  $G = 0 \in \partial_{\mathbb{R}} w(X)$ . Now we assume that  $\mathbb{T}'$  is a finite set. Passing to the subsequence of  $\mathbb{T}$  and relabeling it if necessary, there exists  $\hat{t} \in \mathbb{N}$  such that  $0 < w(X_t) < \mu_t \rho$  for any  $t \geq \hat{t}$ . Moreover, the function  $w$  is continuously differentiable near  $X_t$  and we have

$$(4.3) \quad \text{grad } \tilde{w}(X_t, \mu_t) = \tau_t \text{grad } w(X_t),$$

where  $\tau_t$  is a constant defined by

$$\tau_t = \frac{w(X_t)}{\mu_t \rho} \in (0, 1).$$

A straightforward verification reveals that

$$\|\text{grad } w(X_t)\|_{\text{F}} = 2\rho \frac{\|(I_d - X_t X_t^\top) \Sigma_{\circ} X_t\|_{\text{F}}}{\|(I_d - X_t X_t^\top) \Sigma_{\circ}^{1/2}\|_{\text{F}}} \leq 2\rho \|\Sigma_{\circ}^{1/2}\|_{\text{F}},$$

for any  $X_t \in \mathcal{O}^{d,r}$  satisfying  $w(X_t) \neq 0$ . Hence, the sequence  $\{\text{grad } w(X_t)\}_{t \geq \hat{t}}$  is bounded. By passing to a subsequence if necessary, we may assume without loss of generality that

$$H = \lim_{\mathcal{O}^{d,r} \ni X_t \rightarrow X} \text{grad } w(X_t).$$

According to the definition of Riemannian Clarke subdifferentials, it holds that  $H \in \partial_{\mathbb{R}} w(X)$ . Since  $\tau_t \in (0, 1)$  for any  $t \geq \hat{t}$ , we can assume without loss of generality that  $\lim_{t \rightarrow \infty} \tau_t = \tau$  for a constant  $\tau \in [0, 1]$ . Consequently, it follows from the relationship (4.3) and the fact  $0 \in \partial_{\mathbb{R}} w(X)$  that

$$\begin{aligned} G &= \lim_{\mathcal{O}^{d,r} \ni X_t \rightarrow X, \mu_t \downarrow 0} \text{grad } \tilde{w}(X_t, \mu_t) = \lim_{\mathcal{O}^{d,r} \ni X_t \rightarrow X, \mu_t \downarrow 0} \tau_t \text{grad } w(X_t) \\ &= \tau H = \tau H + (1 - \tau)0 \in \text{conv} \{\partial_{\mathbb{R}} w(X)\} = \partial_{\mathbb{R}} w(X). \end{aligned}$$

The proof is completed.  $\square$

Theorem 4.5 can be regarded as a Riemannian extension of the gradient consistency in Euclidean spaces [12]. It is worth noting, however, that on a Riemannian manifold one must additionally account for the projection onto tangent spaces as well as the use of Riemannian gradients.

**4.4. Algorithm Development.** Based on the smoothing function  $\tilde{w}$ , we can obtain the following approximation of the objective function  $f$  in problem (4.1),

$$\tilde{f}(X, \mu) := \tilde{g}(X, \mu) + s(X),$$

where  $\tilde{g}$  is given by

$$\tilde{g}(X, \mu) := u(X) + \tilde{w}(X, \mu).$$

It is clear that the function  $\tilde{f}(X, \mu)$  exhibits a composite structure. In particular, the first term  $\tilde{g}(X, \mu)$  is smooth for fixed  $\mu > 0$ , whereas the second term  $s(X)$  is possibly nonsmooth. This inherent structure naturally lends itself to the framework of the proximal gradient method on the manifold to minimize  $\tilde{f}(\cdot, \mu)$  over  $\mathcal{O}^{d,r}$ .

Specifically, we intend to solve the following subproblem to find the descent direction  $V_k \in \mathcal{T}_{X_k} \mathcal{O}^{d,r}$  at the  $k$ -th iteration,

$$(4.4) \quad V_k := \arg \min_{V \in \mathcal{T}_{X_k} \mathcal{O}^{d,r}} h_k(V) := \langle \nabla \tilde{g}(X_k, \mu_k), V \rangle + \frac{1}{2\mu_k} \|V\|_{\mathbb{F}}^2 + s(X_k + V),$$

where  $X_k \in \mathcal{O}^{d,r}$  and  $\mu_k > 0$  are the current iterate and smoothing parameter, respectively. The above subproblem involves minimizing a strongly convex function on the tangent space. Although  $V_k$  serves as a descent direction, the updated iterate  $X_k + \alpha_k V_k$ , for an arbitrary stepsize  $\alpha_k > 0$ , does not necessarily remain on  $\mathcal{O}^{d,r}$ . Consequently, we then perform a retraction to bring it back to  $\mathcal{O}^{d,r}$ .

In practice, subproblem (4.4) can be addressed effectively via a semi-smooth Newton method [10] or through a fixed-point method [39]. Below, we provide a concise exposition of the computational procedure underlying the latter approach. The KKT conditions of subproblem (4.4) can be expressed as

$$(4.5) \quad \begin{cases} 0 \in \nabla \tilde{g}(X_k, \mu_k) - X_k \Lambda + \frac{1}{\mu_k} V + \partial s(X_k + V), \\ 0 = X_k^\top V + V^\top X_k, \end{cases}$$

where  $\Lambda \in \mathbb{R}^{r \times r}$  represents the Lagrangian multiplier associated with the linear constraint  $V \in \mathcal{T}_{X_k} \mathcal{O}^{d,r}$ . For a given  $\Lambda$ , the first relationship of (4.5) allows  $V$  to be characterized explicitly by

$$(4.6) \quad V(\Lambda) := \mathfrak{P}_{\mu_k, s}(X_k - \mu_k (\nabla \tilde{g}(X_k, \mu_k) - X_k \Lambda)) - X_k,$$

where the proximal mapping  $\mathfrak{P}_{\mu, s}$  is defined as

$$\mathfrak{P}_{\mu, s}(Y) := \arg \min_{Z \in \mathbb{R}^{d \times r}} s(Z) + \frac{1}{2\mu} \|Z - Y\|_{\mathbb{F}}^2.$$

Substituting (4.6) into the second relationship of (4.5) reveals that  $\Lambda$  satisfies the nonlinear equation  $E(\Lambda) = 0$  with  $E(\Lambda) := X_k^\top V(\Lambda) + V(\Lambda)^\top X_k$ . Building upon this insight, Liu et al. [39] propose the following fixed-point method for solving  $E(\Lambda) = 0$  with respect to  $\Lambda$ ,

$$\Lambda_{t+1} = \Lambda_t - \frac{1}{\mu_k} E(\Lambda_t).$$

In practice, the stopping criterion can be specified as  $\|E(\Lambda_t)\|_{\mathbb{F}} \leq \varepsilon$  for a prescribed tolerance  $\varepsilon \in (0, 1)$ . He et al. [27] have shown that this fixed-point method—viewed as a special instance of their primal-dual hybrid gradient algorithm—enjoys an  $O(1/t)$  convergence rate in the ergodic sense under mild conditions, which are satisfied in the present setting.

The complete framework of our approach for solving problem (4.1) is outlined in Algorithm 1, which is named *smoothing manifold proximal gradient* and abbreviated as SMPG. It is noteworthy that SMPG involves an Armijo line search procedure in (4.7) to determine the stepsize. As we will show later, this backtracking line search procedure is well-defined and guaranteed to terminate in a finite number of steps.

---

**Algorithm 1:** smoothing manifold proximal gradient (SMPG).

---

```

1 Input:  $X_0 \in \mathcal{O}^{d,r}$ ,  $\mu_0 > 0$ ,  $\bar{\mu} \in [0, \mu_0]$ ,  $\theta \in (0, 1)$ , and  $\beta \in (0, 1)$ .
2 for  $k = 0, 1, 2, \dots$  do
3   Solve subproblem (4.4) to obtain  $V_k$ .
4   if  $\|V_k\|_F \leq \bar{\mu}^2$  and  $\mu_k \leq \bar{\mu}$  then
5     Return  $X_k$ .
6   else
7     Find  $\alpha_k := \beta^{m_k}$  such that  $m_k$  is the smallest integer satisfying
      (4.7)  $\tilde{f}(\Re_{X_k}(\beta^{m_k} V_k), \mu_k) \leq \tilde{f}(X_k, \mu_k) - \frac{\beta^{m_k}}{2\mu_k} \|V_k\|_F^2$ .
8     Update  $X_{k+1} := \Re_{X_k}(\alpha_k V_k)$ .
9     Set
      
$$\mu_{k+1} := \begin{cases} \mu_k, & \text{if } \|V_k\|_F > \mu_k^2, \\ \theta\mu_k, & \text{if } \|V_k\|_F \leq \mu_k^2. \end{cases}$$

10 Output:  $X_k$ .
```

---

*Remark 4.6.* Among the four parameters introduced in SMPG,  $\mu_0$  denotes the initial smoothing parameter,  $\theta$  governs the decay rate of the smoothing parameter,  $\beta$  controls the decay rate of the stepsize in the line search procedure, and  $\bar{\mu}$  serves as the termination tolerance. In practice,  $\bar{\mu}$  specifies the desired accuracy and therefore does not require tuning. Our numerical experiments presented in Section 6.2 indicate that SMPG is not particularly sensitive to the other three parameters.

Algorithm 1 bears a certain resemblance to the ManPG method proposed in [10]. The latter similarly solves a subproblem akin to (4.4) in the tangent space, followed by a retraction to preserve the feasibility. It should be emphasized, however, that in Algorithm 1 we further employ a smoothing approximation to handle the possibly nonsmooth term arising from the DRO formulation, together with an update scheme for the smoothing parameter. These features introduce additional challenges to the convergence analysis of Algorithm 1, which we shall elaborate on in the subsequent section.

**5. Convergence Analysis.** This section delves into a comprehensive convergence analysis of the proposed algorithm. Specifically, we establish that any accumulation point of the sequence generated by Algorithm 1 is a stationary point. And the iteration complexity of Algorithm 1 is provided to attain an approximate stationary point.



**5.1. Descent Property.** In the following lemma, we first prove that  $V_k$ , obtained by solving subproblem (4.4), serves as a descent direction in the tangent space  $\mathcal{T}_{X_k} \mathcal{O}^{d,r}$  at the current iterate  $X_k \in \mathcal{O}^{d,r}$ .

LEMMA 5.1. *Suppose that  $\{X_k\}$  is the sequence generated by Algorithm 1. Then it holds that*

$$h_k(0) - h_k(\alpha V_k) \geq \frac{\alpha(2-\alpha)}{2\mu_k} \|V_k\|_F^2,$$

for any  $\alpha \in [0, 1]$ .

*Proof.* Since  $h_k(V)$  is strongly convex with the modulus  $\mu_k^{-1}$ , we have

$$(5.1) \quad h_k(\bar{V}) \geq h_k(V) + \langle \partial h_k(V), \bar{V} - V \rangle + \frac{1}{2\mu_k} \|\bar{V} - V\|_F^2,$$

for any  $V, \bar{V} \in \mathbb{R}^{d \times r}$ . In particular, for  $V, \bar{V} \in \mathcal{T}_{X_k} \mathcal{O}^{d,r}$ , it holds that

$$\langle \partial h_k(V), \bar{V} - V \rangle = \left\langle \text{Proj}_{\mathcal{T}_{X_k} \mathcal{O}^{d,r}}(\partial h_k(V)), \bar{V} - V \right\rangle.$$

Moreover, it follows from the optimality condition of subproblem (4.4) that  $0 \in \text{Proj}_{\mathcal{T}_{X_k} \mathcal{O}^{d,r}}(\partial h_k(V_k))$ . Taking  $V = V_k$  and  $\bar{V} = 0$  in (5.1) yields that

$$h_k(0) \geq h_k(V_k) + \frac{1}{2\mu_k} \|V_k\|_F^2,$$

which, after a suitable rearrangement, can be equivalently written as

$$s(X_k) \geq \langle \nabla \tilde{g}(X_k, \mu_k), V_k \rangle + \frac{1}{\mu_k} \|V_k\|_F^2 + s(X_k + V_k).$$

According to the convexity of  $s$ , we have

$$\begin{aligned} s(X_k + \alpha V_k) - s(X_k) &= s((1-\alpha)X_k + \alpha(X_k + V_k)) - s(X_k) \\ &\leq \alpha (s(X_k + V_k) - s(X_k)). \end{aligned}$$

Collecting the above two relationships together results in that

$$\begin{aligned} h_k(\alpha V_k) - h_k(0) &= \alpha \langle \nabla \tilde{g}(X_k, \mu_k), V_k \rangle + \frac{\alpha^2}{2\mu_k} \|V_k\|_F^2 + s(X_k + \alpha V_k) - s(X_k) \\ &\leq \alpha \left( \langle \nabla \tilde{g}(X_k, \mu_k), V_k \rangle + \frac{\alpha}{2\mu_k} \|V_k\|_F^2 + s(X_k + V_k) - s(X_k) \right) \\ &\leq \frac{\alpha(\alpha-2)}{2\mu_k} \|V_k\|_F^2, \end{aligned}$$

which completes the proof.  $\square$

Based on Lemma 5.1, we can proceed to show that the line search procedure in Algorithm 1 is well-defined, ensuring that the stepsize  $\alpha_k$  can be determined in a finite number of trials.

LEMMA 5.2. *Let  $\{X_k\}$  be the sequence generated by Algorithm 1. Then there exists a constant  $\bar{\alpha} \in (0, 1]$  such that*

$$\tilde{f}(X_k, \mu_k) - \tilde{f}(\mathfrak{R}_{X_k}(\alpha V_k), \mu_k) \geq \frac{\alpha}{2\mu_k} \|V_k\|_F^2,$$

for any  $\alpha \in (0, \bar{\alpha}]$ .

*Proof.* According to the Lipschitz continuity of  $\nabla \tilde{g}(X, \mu)$  for fixed  $\mu > 0$ , it follows that

$$\begin{aligned} \tilde{g}(\mathfrak{R}_{X_k}(\alpha V_k), \mu_k) &\leq \tilde{g}(X_k, \mu_k) + \langle \nabla \tilde{g}(X_k, \mu_k), \mathfrak{R}_{X_k}(\alpha V_k) - X_k \rangle \\ &\quad + \frac{L_u \mu_k + L_w}{2\mu_k} \|\mathfrak{R}_{X_k}(\alpha V_k) - X_k\|_F^2 \\ &\leq \tilde{g}(X_k, \mu_k) + \langle \nabla \tilde{g}(X_k, \mu_k), \mathfrak{R}_{X_k}(\alpha V_k) - (X_k + \alpha V_k) \rangle \\ &\quad + \alpha \langle \nabla \tilde{g}(X_k, \mu_k), V_k \rangle + \frac{L_u \mu_0 + L_w}{2\mu_k} \|\mathfrak{R}_{X_k}(\alpha V_k) - X_k\|_F^2, \end{aligned}$$

where the last inequality holds due to the fact that  $\mu_k \leq \mu_0$ . Since  $\nabla u$  is continuous over the compact manifold  $\mathcal{O}^{d,r}$ , there exists a constant  $M_u > 0$  such that  $\|\nabla u(X)\|_F \leq M_u$  for any  $X \in \mathcal{O}^{d,r}$ . Hence, it holds that

$$\|\nabla \tilde{g}(X_k, \mu_k)\|_F \leq \|\nabla u(X_k)\|_F + \|\nabla \tilde{w}(X_k, \mu_k)\|_F \leq \frac{(M_u + M_w)\mu_0}{\mu_k}.$$

By invoking Lemma 2.3, we can obtain that

$$\begin{aligned} &\langle \nabla \tilde{g}(X_k, \mu_k), \mathfrak{R}_{X_k}(\alpha V_k) - (X_k + \alpha V_k) \rangle \\ &\leq \|\nabla \tilde{g}(X_k, \mu_k)\|_F \|\mathfrak{R}_{X_k}(\alpha V_k) - (X_k + \alpha V_k)\|_F \\ &\leq \frac{\alpha^2 \mu_0 M_2 (M_u + M_w)}{\mu_k} \|V_k\|_F^2, \end{aligned}$$

and

$$\|\mathfrak{R}_{X_k}(\alpha V_k) - X_k\|_F^2 \leq \alpha^2 M_1^2 \|V_k\|_F^2.$$

Let  $C = 2\mu_0 M_2 (M_u + M_w) + M_1^2 (L_u \mu_0 + L_w) > 0$  be a constant. Then it can be readily verified that

$$\tilde{g}(\mathfrak{R}_{X_k}(\alpha V_k), \mu_k) - \tilde{g}(X_k, \mu_k) \leq \alpha \langle \nabla \tilde{g}(X_k, \mu_k), V_k \rangle + \frac{\alpha^2 C}{2\mu_k} \|V_k\|_F^2.$$

Moreover, according to the Lipschitz continuity of  $s$ , we have

$$\begin{aligned} s(\mathfrak{R}_{X_k}(\alpha V_k)) - s(X_k) &= s(\mathfrak{R}_{X_k}(\alpha V_k)) - s(X_k + \alpha V_k) + s(X_k + \alpha V_k) - s(X_k) \\ &\leq L_s \|\mathfrak{R}_{X_k}(\alpha V_k) - (X_k + \alpha V_k)\|_F + s(X_k + \alpha V_k) - s(X_k) \\ &\leq \frac{\alpha^2 \mu_0 M_2 L_s}{\mu_k} \|V_k\|_F^2 + s(X_k + \alpha V_k) - s(X_k), \end{aligned}$$

where the last inequality results from Lemma 2.3 and the fact that  $\mu_k \leq \mu_0$ . Collecting the above two inequalities together yields that

$$\begin{aligned} &\tilde{f}(\mathfrak{R}_{X_k}(\alpha V_k), \mu_k) - \tilde{f}(X_k, \mu_k) \\ &= \tilde{g}(\mathfrak{R}_{X_k}(\alpha V_k), \mu_k) - \tilde{g}(X_k, \mu_k) + s(\mathfrak{R}_{X_k}(\alpha V_k)) - s(X_k) \\ &\leq \alpha \langle \nabla \tilde{g}(X_k, \mu_k), V_k \rangle + s(X_k + \alpha V_k) - s(X_k) + \frac{\alpha^2 (C + 2\mu_0 M_2 L_s)}{2\mu_k} \|V_k\|_F^2. \end{aligned} \tag{5.2}$$

As a direct consequence of Lemma 5.1, we can proceed to show that

$$\alpha \langle \nabla \tilde{g}(X_k, \mu_k), V_k \rangle + s(X_k + \alpha V_k) - s(X_k) \leq -\frac{\alpha}{\mu_k} \|V_k\|_F^2,$$

which together with the relationship (5.2) implies that

$$\tilde{f}(\mathfrak{R}_{X_k}(\alpha V_k), \mu_k) - \tilde{f}(X_k, \mu_k) \leq -\frac{\alpha}{2\mu_k} (2 - \alpha(C + 2\mu_0 M_2 L_s)) \|V_k\|_F^2.$$

Let  $\bar{\alpha} = \min\{1, 1/(C + 2\mu_0 M_2 L_s)\} \in (0, 1]$ . Then for any  $\alpha \in (0, \bar{\alpha}]$ , we can conclude that

$$\tilde{f}(\mathfrak{R}_{X_k}(\alpha V_k), \mu_k) - \tilde{f}(X_k, \mu_k) \leq -\frac{\alpha}{2\mu_k} \|V_k\|_F^2,$$

as desired. The proof is completed.  $\square$

Lemma 5.2 guarantees that the line search procedure in (4.7) terminates in at most  $\lceil \log_{\beta} \bar{\alpha} \rceil$  steps, which is independent of the smoothing parameter  $\mu_k$ . Here, the notation  $\lceil m \rceil$  represents the smallest integer greater than or equal to  $m \in \mathbb{R}$ .

**5.2. Global Convergence.** The following lemma lays the foundation for the global convergence analysis of Algorithm 1 with  $\bar{\mu} = 0$ , as established in Theorem 5.4.

**LEMMA 5.3.** *Let  $\{X_k\}$  be the sequence generated by Algorithm 1 with  $\bar{\mu} = 0$ . Then  $\mathbb{K} := \{k \in \mathbb{N} \mid \|V_k\|_F \leq \mu_k^2\}$  is an infinite set.*

*Proof.* Suppose, on the contrary, that  $\mathbb{K}$  is a finite set. Then there exists  $\bar{k} \in \mathbb{N}$  such that

$$(5.3) \quad \mu_k = \mu_{\bar{k}} > 0, \text{ and } \|V_k\|_F > \mu_{\bar{k}}^2 > 0,$$

for any  $k \geq \bar{k}$ . Thus, we have  $X_{k+1} = \mathfrak{R}_{X_k}(\alpha_k V_k)$  for all  $k \geq \bar{k}$ , where the stepsize  $\alpha_k$  is obtained by using the line search procedure in (4.7) with  $\mu_k$  fixed as  $\mu_{\bar{k}}$ . From Lemma 5.2, we know that  $\alpha_k \geq \bar{\alpha}\beta$  and

$$\tilde{f}(X_k, \mu_{\bar{k}}) - \tilde{f}(X_{k+1}, \mu_{\bar{k}}) \geq \frac{\alpha_k}{2\mu_{\bar{k}}} \|V_k\|_F^2 \geq \frac{\bar{\alpha}\beta}{2\mu_{\bar{k}}} \|V_k\|_F^2,$$

for any  $k \geq \bar{k}$ . This observation indicates that the sequence  $\{\tilde{f}(X_k, \mu_{\bar{k}})\}_{k \geq \bar{k}}$  is monotonically decreasing. Moreover, as a direct consequence of Proposition 4.3, we can proceed to show that

$$f(X) \leq \tilde{f}(X, \mu) \leq f(X) + \frac{\mu\rho}{2},$$

for any  $X \in \mathcal{O}^{d,r}$  and  $\mu > 0$ . In light of the continuity of  $f$  over the compact manifold  $\mathcal{O}^{d,r}$ , there exist two constants  $\tilde{f}_{\min}$  and  $\tilde{f}_{\max}$  such that

$$(5.4) \quad \tilde{f}_{\min} \leq \tilde{f}(X, \mu) \leq \tilde{f}_{\max},$$

for any  $X \in \mathcal{O}^{d,r}$  and  $\mu \in (0, \mu_0]$ . Hence, the sequence  $\{\tilde{f}(X_k, \mu_{\bar{k}})\}_{k \geq \bar{k}}$  is convergent. Then we can obtain that

$$\lim_{k \rightarrow \infty} \|V_k\|_F^2 \leq \frac{2\mu_{\bar{k}}}{\bar{\alpha}\beta} \lim_{k \rightarrow \infty} (\tilde{f}(X_k, \mu_{\bar{k}}) - \tilde{f}(X_{k+1}, \mu_{\bar{k}})) = 0,$$

which contradicts the second relationship in (5.3). Consequently, we can conclude that  $\mathbb{K}$  is an infinite set. The proof is completed.  $\square$

Now we are in a position to establish the global convergence of Algorithm 1 to a stationary point of problem (4.1) under the setting  $\bar{\mu} = 0$ .

**THEOREM 5.4.** *Suppose that  $\{X_k\}$  is the sequence generated by Algorithm 1 with  $\bar{\mu} = 0$ . Then the sequence  $\{X_k\}$  has at least one accumulation point. Moreover, any accumulation point is a stationary point of problem (4.1).*

*Proof.* For each  $k \in \mathbb{K}$ , we have  $\mu_{k+1} = \theta\mu_k$  with  $\theta \in (0, 1)$  being a decaying factor. According to Lemma 5.3, the index set  $\mathbb{K}$  is infinite. Then it can be readily verified that

$$(5.5) \quad \lim_{\mathbb{K} \ni k \rightarrow \infty} \frac{1}{\mu_k} \|V_k\|_F \leq \lim_{\mathbb{K} \ni k \rightarrow \infty} \mu_k = 0.$$

Since  $\mathcal{O}^{d,r}$  is a compact manifold, the sequence  $\{X_k\}$  is bounded. Then from the Bolzano-Weierstrass theorem, it can be deduced that the sequence  $\{X_k\}$  has at least one accumulation point. Let  $X_*$  be an accumulation point of  $\{X_k\}$ . The completeness of  $\mathcal{O}^{d,r}$  guarantees that  $X_* \in \mathcal{O}^{d,r}$ . By passing to a subsequence if necessary, we may assume without loss of generality that  $\lim_{\mathbb{K} \ni k \rightarrow \infty} X_k = X_*$ .

Next, by virtue of the optimality condition [10, 58] of subproblem (4.4), there exists  $H_k \in \partial s(X_k + V_k)$  such that

$$(5.6) \quad \text{grad } u(X_k) + \text{grad } \tilde{w}(X_k, \mu_k) + \text{Proj}_{\mathcal{T}_{X_k} \mathcal{O}^{d,r}}(H_k) + \frac{1}{\mu_k} V_k = 0,$$

for any  $k \in \mathbb{K}$ . It is clear that the sequence  $\{\text{grad } u(X_k)\}$  is convergent and

$$\lim_{\mathbb{K} \ni k \rightarrow \infty} \text{grad } u(X_k) = \text{grad } u(X_*).$$

According to the Lipschitz continuity of  $s$ , the sequence  $\{H_k\}$  is bounded [16]. Without loss of generality, we assume that it is also convergent. Then there exists  $H_*$  such that  $\lim_{\mathbb{K} \ni k \rightarrow \infty} H_k = H_*$ . It follows from [43, Theorem 24.4] that  $H_* \in \partial s(X_*)$ . In addition, the boundedness of the sequence  $\{\mu_k^{-1} V_k\}$  results from the fact that it is convergent in (5.5). As a result, the sequence  $\{\text{grad } \tilde{w}(X_k, \mu_k)\}$  is also bounded. Without loss of generality, we can assume that there exists  $G_*$  such that

$$\lim_{\mathbb{K} \ni k \rightarrow \infty} \text{grad } \tilde{w}(X_k, \mu_k) = G_*.$$

According to the definition of the Riemannian subdifferential of  $w$  associated with  $\tilde{w}$ , it holds that  $G_* \in \partial_{\mathbb{R}}|_{\tilde{w}} w(X_*)$ . Then it follows from Theorem 4.5 that  $G_* \in \partial_{\mathbb{R}} w(X_*)$ .

Finally, upon taking  $k \rightarrow \infty$  in the relationship (5.6), we can conclude that

$$0 = \text{grad } u(X_*) + G_* + \text{Proj}_{\mathcal{T}_{X_*} \mathcal{O}^{d,r}}(H_*) \in \text{grad } u(X_*) + \partial_{\mathbb{R}} s(X_*) + \partial_{\mathbb{R}} w(X_*),$$

which indicates that  $X_*$  is a stationary point of problem (4.1). We complete the proof.  $\square$

**5.3. Iteration Complexity.** The final task is to derive the iteration complexity of Algorithm 1, a critical challenge that remains unresolved in existing works. The proof of Theorem 5.4 previously discussed leads to the insight that an accumulation point of  $\{X_k\}$  is a stationary point of (4.1) if the following conditions are satisfied,

$$\begin{cases} \lim_{k \rightarrow \infty} \text{dist} \left( 0, \text{grad } u(X_k) + \text{grad } \tilde{w}(X_k, \mu_k) + \text{Proj}_{\mathcal{T}_{X_k} \mathcal{O}^{d,r}}(\partial s(X_k + V_k)) \right) = 0, \\ \lim_{k \rightarrow \infty} \|V_k\|_F = 0, \quad \lim_{k \rightarrow \infty} \mu_k = 0. \end{cases}$$

This observation motivates us to define the concept of  $\epsilon$ -approximate stationarity for problem (4.1) as follows.

DEFINITION 5.5. *A point  $X \in \mathcal{O}^{d,r}$  is called an  $\epsilon$ -approximate stationary point of problem (4.1) if there exists  $V \in \mathcal{T}_X \mathcal{O}^{d,r}$  with  $\|V\|_F \leq \epsilon$  and  $\mu \in (0, \epsilon]$  such that*

$$\text{dist} \left( 0, \text{grad } u(X) + \text{grad } \tilde{w}(X, \mu) + \text{Proj}_{\mathcal{T}_X \mathcal{O}^{d,r}} (\partial s(X + V)) \right) \leq \epsilon.$$

We show that Algorithm 1 is capable of identifying an  $\epsilon$ -approximate stationary point of problem (4.1) under the setting  $\bar{\mu} = \epsilon$ .

LEMMA 5.6. *For any  $\epsilon \in (0, 1)$  and  $\mu_0 \geq \epsilon$ , Algorithm 1 with  $\bar{\mu} = \epsilon$  will terminate at an  $\epsilon$ -approximate stationary point of problem (4.1).*

*Proof.* We first define the constant

$$\iota_\epsilon = \lceil \log_\theta (\epsilon / \mu_0) \rceil + 1.$$

Let  $\mathbb{k}_i$  be the  $i$ -th smallest element of  $\mathbb{K}$ . Then it holds that

$$\mu_{\mathbb{k}_{\iota_\epsilon}} = \theta^{\iota_\epsilon - 1} \mu_0 \leq \epsilon, \text{ and } \|V_{\mathbb{k}_{\iota_\epsilon}}\|_F \leq \mu_{\mathbb{k}_{\iota_\epsilon}}^2 \leq \epsilon^2,$$

which reveals that Algorithm 1 will terminate at the iterate  $X_{\mathbb{k}_{\iota_\epsilon}}$ .

The next step is to show that  $X_{\mathbb{k}_{\iota_\epsilon}}$  is an  $\epsilon$ -approximate stationary point of problem (4.1). According to the relationship (5.6), there exists  $H_{\mathbb{k}_{\iota_\epsilon}} \in \partial s(X_{\mathbb{k}_{\iota_\epsilon}} + V_{\mathbb{k}_{\iota_\epsilon}})$  such that

$$-\frac{1}{\mu_{\mathbb{k}_{\iota_\epsilon}}} V_{\mathbb{k}_{\iota_\epsilon}} = \text{grad } u(X_{\mathbb{k}_{\iota_\epsilon}}) + \text{grad } \tilde{w}(X_{\mathbb{k}_{\iota_\epsilon}}, \mu_{\mathbb{k}_{\iota_\epsilon}}) + \text{Proj}_{\mathcal{T}_{X_{\mathbb{k}_{\iota_\epsilon}}} \mathcal{O}^{d,r}} (H_{\mathbb{k}_{\iota_\epsilon}}),$$

which implies that

$$\begin{aligned} & \text{dist} \left( 0, \text{grad } u(X_{\mathbb{k}_{\iota_\epsilon}}) + \text{grad } \tilde{w}(X_{\mathbb{k}_{\iota_\epsilon}}, \mu_{\mathbb{k}_{\iota_\epsilon}}) + \text{Proj}_{\mathcal{T}_{X_{\mathbb{k}_{\iota_\epsilon}}} \mathcal{O}^{d,r}} (\partial s(X_{\mathbb{k}_{\iota_\epsilon}} + V_{\mathbb{k}_{\iota_\epsilon}})) \right) \\ & \leq \frac{1}{\mu_{\mathbb{k}_{\iota_\epsilon}}} \|V_{\mathbb{k}_{\iota_\epsilon}}\|_F \leq \mu_{\mathbb{k}_{\iota_\epsilon}} \leq \epsilon. \end{aligned}$$

Therefore, we conclude that  $X_{\mathbb{k}_{\iota_\epsilon}}$  is an  $\epsilon$ -approximate stationary point of problem (4.1), which completes the proof.  $\square$

The iteration complexity of Algorithm 1 is established in the following theorem for finding an  $\epsilon$ -approximate stationary point.

THEOREM 5.7. *For any  $\epsilon \in (0, 1)$  and  $\mu_0 \geq \epsilon$ , Algorithm 1 with  $\bar{\mu} = \epsilon$  will reach an  $\epsilon$ -approximate stationary point of problem (4.1) after at most  $O(\epsilon^{-3})$  iterations.*

*Proof.* According to Lemma 5.6, Algorithm 1 with  $\bar{\mu} = \epsilon$  will terminate at  $X_{\mathbb{k}_{\iota_\epsilon}}$ , which is an  $\epsilon$ -approximate stationary point of problem (4.1). We give an upper bound on the iteration number  $\mathbb{k}_{\iota_\epsilon}$ . For convenience, we denote  $\mathbb{k}_0 = 0$ . The iterations from  $\mathbb{k}_i$  to  $\mathbb{k}_{i+1}$  solve subproblem (4.4) with the smoothing parameter fixed as  $\mu_{\mathbb{k}_{i+1}}$  for  $0 \leq i \leq \iota_\epsilon - 1$ . According to Lemma 5.2, we have

$$\tilde{f}(X_k, \mu_{\mathbb{k}_{i+1}}) - \tilde{f}(X_{k+1}, \mu_{\mathbb{k}_{i+1}}) \geq \frac{\alpha_k}{2\mu_{\mathbb{k}_{i+1}}} \|V_k\|_F^2 \geq \frac{\bar{\alpha}\beta}{2\mu_{\mathbb{k}_{i+1}}} \|V_k\|_F^2,$$

for  $\mathbb{k}_i \leq k \leq \mathbb{k}_{i+1} - 1$ . Then it can be readily verified that

$$\begin{aligned} \sum_{k=\mathbb{k}_i}^{\mathbb{k}_{i+1}-1} \|V_k\|_F^2 &\leq \frac{2\mu_{\mathbb{k}_{i+1}}}{\bar{\alpha}\beta} \sum_{k=\mathbb{k}_i}^{\mathbb{k}_{i+1}-1} \left( \tilde{f}(X_k, \mu_{\mathbb{k}_{i+1}}) - \tilde{f}(X_{k+1}, \mu_{\mathbb{k}_{i+1}}) \right) \\ &= \frac{2\mu_{\mathbb{k}_{i+1}}}{\bar{\alpha}\beta} \left( \tilde{f}(X_{\mathbb{k}_i}, \mu_{\mathbb{k}_{i+1}}) - \tilde{f}(X_{\mathbb{k}_{i+1}}, \mu_{\mathbb{k}_{i+1}}) \right) \leq \frac{2\mu_{\mathbb{k}_{i+1}}}{\bar{\alpha}\beta} \left( \tilde{f}_{\max} - \tilde{f}_{\min} \right), \end{aligned}$$

where the last inequality follows from (5.4). From the definition of  $\mathbb{k}_{i+1}$ , we know that  $\|V_k\|_F > \mu_{\mathbb{k}_{i+1}}^2$  for  $\mathbb{k}_i \leq k \leq \mathbb{k}_{i+1} - 1$ . Hence, it holds that

$$\sum_{k=\mathbb{k}_i}^{\mathbb{k}_{i+1}-1} \|V_k\|_F^2 \geq (\mathbb{k}_{i+1} - \mathbb{k}_i) \mu_{\mathbb{k}_{i+1}}^4,$$

which together with the relationship  $\mu_{\mathbb{k}_{i+1}} = \theta^i \mu_0$  implies that

$$\mathbb{k}_{i+1} - \mathbb{k}_i \leq \frac{2(\tilde{f}_{\max} - \tilde{f}_{\min})}{\bar{\alpha}\beta\mu_{\mathbb{k}_{i+1}}^3} = \frac{2(\tilde{f}_{\max} - \tilde{f}_{\min})}{\bar{\alpha}\beta\mu_0^3\theta^{3i}}.$$

Finally, a straightforward verification reveals that

$$\mathbb{k}_{\iota_\epsilon} = \mathbb{k}_0 + \sum_{i=0}^{\iota_\epsilon-1} (\mathbb{k}_{i+1} - \mathbb{k}_i) \leq \frac{2(\tilde{f}_{\max} - \tilde{f}_{\min})}{\bar{\alpha}\beta\mu_0^3} \sum_{i=0}^{\iota_\epsilon-1} \frac{1}{\theta^{3i}} \leq \frac{2\theta^3(\tilde{f}_{\max} - \tilde{f}_{\min})}{\bar{\alpha}\beta(1 - \theta^3)\epsilon^3}.$$

The proof is completed.  $\square$

Theorem 5.7 demonstrates that SMPG achieves an iteration complexity of  $O(\epsilon^{-3})$ . This result represents a clear improvement over the Riemannian subgradient method proposed in [38], which suffers from an inferior iteration complexity of  $O(\epsilon^{-4})$ .

**6. Numerical Experiments.** Preliminary numerical results are presented in this section to provide additional insights into the performance guarantees of model (P<sub>mM</sub>) and Algorithm 1 (SMPG). All code is implemented in MATLAB R2018b on a workstation with dual Intel Xeon Gold 6242R CPU processors (at 3.10 GHz  $\times 20 \times 2$ ) and 510 GB of RAM under Ubuntu 20.04.

**6.1. Experimental Setting.** In the numerical experiments, we estimate an empirical distribution [20] to serve as the nominal distribution  $\mathbb{P}_\circ$  in problem (P<sub>mM</sub>). Although the true distribution  $\mathbb{P}_*$  remains inherently elusive, it is often partially observable through a finite collection of  $n \in \mathbb{N}$  independent samples [20, 41], such as past realizations of the random vector  $\xi$ . Let the training dataset comprising these samples be denoted by  $\hat{\Xi}_n := \{\hat{\xi}_i\}_{i=1}^n \subseteq \mathbb{R}^d$ . Then we construct the empirical distribution  $\hat{\mathbb{P}}_n := \sum_{i=1}^n \delta_{\hat{\xi}_i}/n$ , where  $\delta_{\hat{\xi}_i}$  represents the Dirac distribution concentrating unit mass at  $\hat{\xi}_i \in \mathbb{R}^d$ . In fact, the empirical distribution  $\hat{\mathbb{P}}_n$  can be interpreted as the uniform distribution over the finite samples in  $\hat{\Xi}_n$ .

Based on the preceding constructions, the sample average approximation (SAA) model of PCA can be expressed as

$$(6.1) \quad \min_{X \in \mathcal{O}^{d,r}} \mathbb{E}_{\hat{\mathbb{P}}_n} \left[ \left\| (I_d - XX^\top) \left( \xi - \mathbb{E}_{\hat{\mathbb{P}}_n}[\xi] \right) \right\|_2^2 \right] + s(X).$$

The above formulation, which does not account for uncertainty in the underlying distribution, has been extensively investigated in the literature [41, 48, 53]. In the

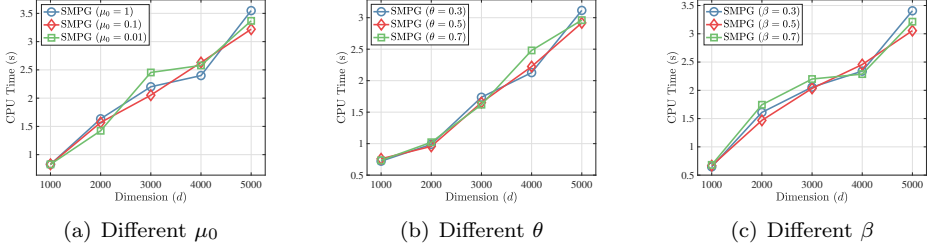


FIG. 1. Numerical comparison of different algorithmic parameters for SMPG.

subsequent experiments, we will engage in a performance comparison between the equivalent reformulation ( $P_m$ ) of the DRO model ( $P_{mM}$ ) and the SAA model (6.1).

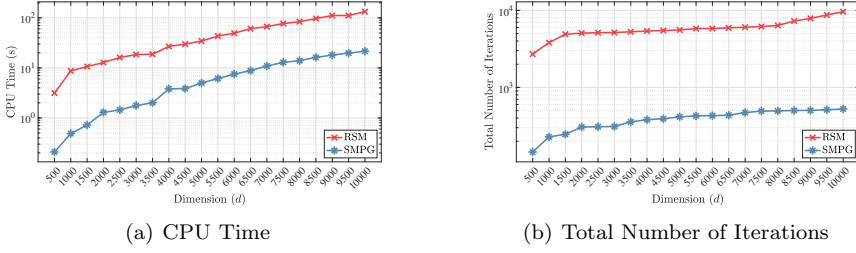
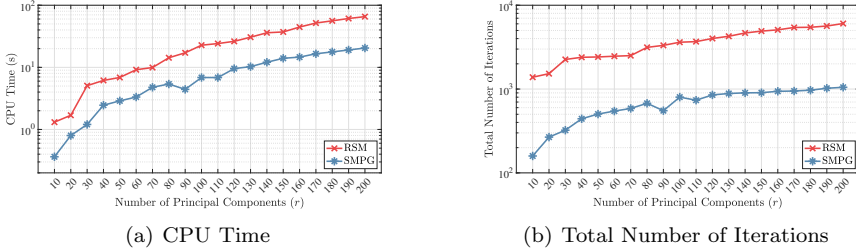
In addition, we focus on the  $\ell_1$ -norm regularizer  $s(X) = \gamma \|X\|_1$  with a parameter  $\gamma > 0$  to control the amount of sparsity, where  $\|X\|_1 := \sum_{i,j} |X_{i,j}|$  with  $X_{i,j}$  being the  $(i, j)$ -th entry of  $X$ . Our empirical experiments reveal that the choice of regularizers does not affect the numerical results dramatically.

**6.2. Implementation Details.** The proposed algorithm SMPG involves three tuning parameters, including the initial smoothing parameter  $\mu_0 > 0$ , the decay rate  $\theta \in (0, 1)$  of the smoothing parameter, and the decay rate  $\beta \in (0, 1)$  of the stepsize in the line search procedure. This subsection investigates the performance of SMPG on problem ( $P_m$ ) under various choices of these parameters.

For our testing, we randomly generate a collection of matrices  $\Sigma_o \in \mathbb{S}_+^d$  with dimension  $d \in \{1000, 2000, 3000, 4000, 5000\}$ . The other parameters in problem ( $P_m$ ) are set to  $r = 10$ ,  $\gamma = 0.1$ , and  $\rho = 1$ . We then evaluate the behavior of SMPG under a broad range of parameter configurations, taking  $\mu_0 \in \{0.01, 0.1, 1\}$ ,  $\theta \in \{0.3, 0.5, 0.7\}$ , and  $\beta \in \{0.3, 0.5, 0.7\}$ . The CPU times consumed by SMPG to attain the termination tolerance  $\bar{\mu} = 10^{-3}$  are recorded and displayed in Figure 1. As evidenced therein, SMPG demonstrates a remarkable degree of stability across different parameter choices, with its performance exhibiting no pronounced sensitivity to the specific values selected. Accordingly, we fix  $\mu_0 = 0.1$ ,  $\theta = 0.5$ , and  $\beta = 0.5$  by default in the algorithm SMPG.

**6.3. Performance of SMPG.** The experiments conducted in this subsection are designed to demonstrate the effectiveness and efficiency of SMPG for solving problem ( $P_m$ ) in comparison with the Riemannian subgradient method (RSM) proposed in [38]. Building on the implications of [38, Theorem 2], we equip RSM with a stepsize scheme of the form  $\sigma k^{-1/2}$  at the  $k$ -th iteration with  $\sigma > 0$ . For each test instance, the parameter  $\sigma$  is tuned via a manual grid search in  $[1, 20]$  with an increment of 1 to identify the value that delivers the best empirical performance. We employ the fixed-point method proposed in [39] to solve the subproblem (4.4) of SMPG. The retraction operator is realized by the polar decomposition in both SMPG and RSM. We set the termination tolerance of SMPG to  $\bar{\mu} = 10^{-3}$ . And RSM is stopped once the function value at the current iterate does not exceed  $1.0001 \times f_{\text{SMPG}}$ , where  $f_{\text{SMPG}}$  denotes the final function value returned by SMPG. The initial point for both algorithms is generated by the leading  $r$  eigenvectors of  $\Sigma_o$  in problem ( $P_m$ ).

We begin by assessing the performance of SMPG and RSM on synthetic datasets. Specifically, the true distribution  $\mathbb{P}_*$  is chosen as the normal distribution  $\mathbb{N}(0, \Sigma_*)$ ,

FIG. 2. Numerical comparison between SMPG and RSM for different values of  $d$  with  $r = d/100$ .FIG. 3. Numerical comparison between SMPG and RSM on MNIST for different values of  $r$ .

where the true covariance matrix  $\Sigma_* \in \mathbb{S}_+^d$  is obtained by projecting a randomly generated matrix onto  $\mathbb{S}_+^d$ . Then the nominal distribution  $\mathbb{P}_o$  is constructed by the empirical distribution  $\hat{\mathbb{P}}_n$  based on  $n = 300$  samples drawn independently and identically from  $\mathcal{N}(0, \Sigma_*)$ . In this test, we vary the dimension  $d$  within  $\{500j \mid j = 1, 2, \dots, 20\}$  while fixing  $r = d/100$ ,  $\gamma = 0.05$ , and  $\rho = 1$  in problem  $(P_m)$ . Figure 2 comprises two subplots that depict CPU times and total numbers of iterations required by two algorithms. It can be observed that SMPG outperforms RSM in terms of computational efficiency. Although each iteration of RSM incurs a relatively low cost, its slow convergence necessitates a large number of iterations to reach a desirable accuracy. Therefore, the total CPU time of RSM is considerably higher than that of SMPG.

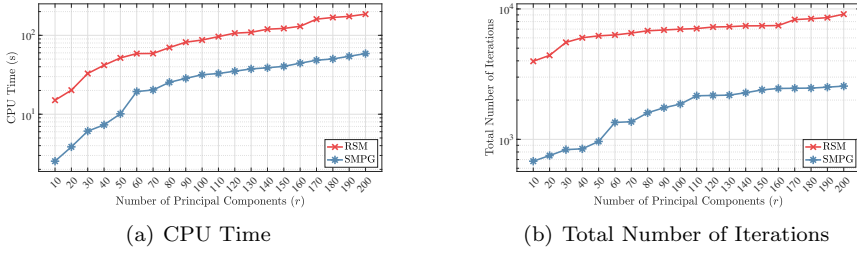
In the subsequent experiments, we then extend our evaluation to two real-world datasets, MNIST<sup>1</sup> and CIFAR-10<sup>2</sup>, to further examine the practical behavior of SMPG and RSM. MNIST contains 60000 samples, each with  $d = 784$  features; and CIFAR-10 consists of 50000 samples with  $d = 3072$  features per sample. The nominal distribution  $\mathbb{P}_o$  is constructed in the same manner as before, based on the empirical distribution  $\hat{\mathbb{P}}_n$  derived from the first  $n = 500$  samples. Other parameters in problem  $(P_m)$  are set to  $\gamma = 0.02$  and  $\rho = 1$  with  $r \in \{10j \mid j = 1, 2, \dots, 20\}$ . The corresponding numerical results are presented in Figure 3 and Figure 4 for MNIST and CIFAR-10, respectively. These evidences clearly show that the computational advantage of SMPG is not confined to synthetic scenarios, whose superiority remains pronounced when applied to real-world datasets.

**6.4. Performance of DRO Model.** The purpose of this subsection is to illustrate the rationality and necessity of adopting the DRO model for PCA. For conve-

<sup>1</sup><https://yann.lecun.com/exdb/mnist/>

<sup>2</sup><https://www.cs.toronto.edu/~kriz/cifar.html>




 FIG. 4. Numerical comparison between SMPG and RSM on CIFAR-10 for different values of  $r$ .

nience, the DRO model ( $P_m$ ) and the SAA model (6.1) are denoted by DRPCA and PCA-SAA, respectively. The performances of DRPCA and PCA-SAA are evaluated on three real-world datasets, including MNIST, CIFAR-10, and WHO-Mortality<sup>3</sup>. The last one includes 6080 samples, each characterized by  $d = 24$  features.

In the first experiment, we extract the first  $n$  samples of each dataset to construct the empirical distribution  $\hat{P}_n$  as the nominal distribution  $P_0$ . Then SMPG and ManPG [10] are deployed to solve the DRPCA model ( $P_m$ ) and the PCA-SAA model (6.1), respectively. For our simulation in this case, we fix  $r = 5$  and  $\gamma = 0.02$ . The radius  $\rho$  in problem ( $P_m$ ) is set to  $5n^{-1/2}$  for each sample size  $n$ . To assess the quality of solutions in DRPCA and PCA-SAA, we calculate the out-of-sample performance [20] for the function value of PCA as follows,

$$(6.2) \quad f_*(X) = \mathbb{E}_{P_*} \left[ \left\| (I_d - XX^\top) (\xi - \mathbb{E}_{P_*} [\xi]) \right\|_2^2 \right],$$

which represents the reconstruction error of PCA with the true distribution  $P_*$  being the empirical distribution generated by all samples in the dataset. Furthermore, the sparsity levels of solutions, quantified by the percentage of zero entries, are recorded as well. When computing the sparsity level, we set an entry to zero when its absolute value is less than  $10^{-5}$  (see [10]). Figure 5 and Figure 6 visualize the out-of-sample performances of DRPCA and PCA-SAA on three datasets, evaluated across sample sizes  $n \in \{100, 200, 300, 400, 500\}$ . As evidenced therein, the solutions of DRPCA consistently demonstrate superior performances compared to those of PCA-SAA in terms of function values. In addition, the solutions obtained from both models exhibit comparable sparsity levels, with their differences remaining within an acceptable margin. These numerical results highlight the rationality and necessity of adopting the DRO model for PCA.

Due to the presence of noise and the variability of practical environments, real-world applications often face the risk of distributional corruption or the emergence of extreme scenarios. We then perform a stress test [19, 26] on DRPCA and PCA-SAA to simulate this situation. The robustness of solutions delivered by both models is examined under a contaminated distribution  $P_\nu = (1-\nu)\hat{P}_n + \nu\bar{P}$ , where  $\bar{P}$  represents a contaminant of  $\hat{P}_n$  and  $\nu \in [0, 1]$  encodes the contamination level. Similar to the setting of [32], we choose a Gaussian contaminant inspired by the extremal-covariance construction in Appendix A. Specifically, we set  $\bar{P} = \mathcal{N}(\zeta, \bar{\Sigma}_{\bar{X}})$  for a solution  $\bar{X}$  of DRPCA, where  $\zeta \in \mathbb{R}^d$  is a  $d$ -dimensional vector of all ones and  $\bar{\Sigma}_{\bar{X}} \in \mathbb{S}_+^d$  is given in Theorem A.1. Next, the function values of PCA under  $P_\nu$  are evaluated for  $\nu \in [0, 1]$ ,

<sup>3</sup><https://www.who.int/data/gho/data/themes/mortality-and-global-health-estimates>

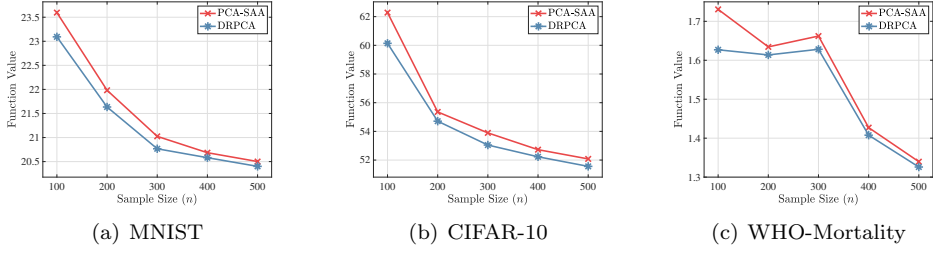


FIG. 5. Numerical comparison of function values between DRPCA and PCA-SAA.

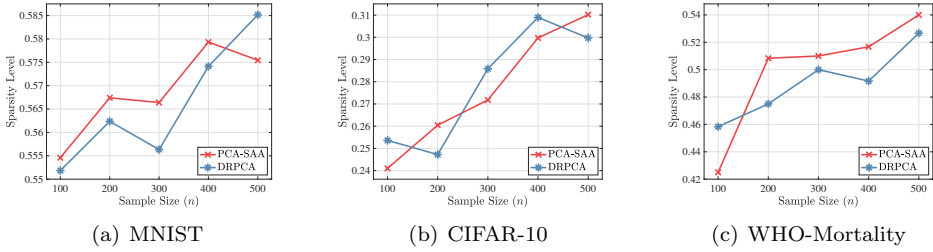


FIG. 6. Numerical comparison of sparsity levels between DRPCA and PCA-SAA.

obtained by substituting  $\mathbb{P}_*$  with  $\mathbb{P}_\nu$  in (6.2). In this experiment, we fix  $r = 20$ ,  $\gamma = 0.1$ ,  $n = 100$ , and  $\rho = 10$ . Figure 7 illustrates the numerical results of DRPCA and PCA-SAA on three datasets under the stress test. Beyond a certain level of contamination, the solutions of DRPCA begin to exhibit better performance than those of PCA-SAA. Moreover, the function values of PCA-SAA change more steeply than those of DRPCA, which suggests the enhanced robustness of DRPCA against contaminated data.

**7. Concluding Remarks.** The DRO model ( $\mathbb{P}_{mM}$ ) of regularized PCA constitutes a constrained min-max optimization problem on a Riemannian manifold. When the ambiguity set is characterized by the type-2 Wasserstein distance, we equivalently reformulate it as a minimization problem ( $\mathbb{P}_m$ ) by providing a closed-form expression for the optimal value of the inner maximization problem in ( $\mathbb{P}_{mM}$ ). However, problem ( $\mathbb{P}_m$ ) can hardly be solved efficiently by existing Riemannian optimization algorithms due to the involvement of two possibly nonsmooth terms in the objective function. To surmount this issue, we develop an efficient algorithm SMPG for problem ( $\mathbb{P}_m$ ), which incorporates the smoothing approximation technique into the proximal gradient method on Riemannian manifolds. We rigorously demonstrate that SMPG achieves global convergence to a stationary point and further establish an iteration complexity. Preliminary numerical results are presented to validate the efficiency of SMPG and the effectiveness of our DRO model, illuminating their potential in addressing the challenges inherent in PCA under distributional uncertainty.

**Acknowledgments.** We sincerely express our gratitude to Prof. Man-Chung Yue and two anonymous reviewers for their insightful and constructive comments, which have substantially improved the quality, clarity, and presentation of this paper.

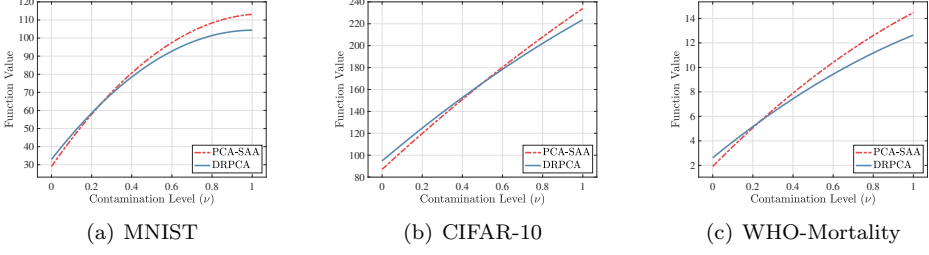


FIG. 7. Numerical comparison of stress tests between DRPCA and PCA-SAA.

**Appendix A. Extremal Distribution.** In this appendix, we focus on the case where  $p = 2$ . For any  $X \in \mathcal{O}^{d,r}$ , we construct a candidate extremal distribution  $\bar{\mathbb{P}}_X$  that achieves the worst-case function value in problem (3.1), namely,

$$(A.1) \quad \phi(X, \bar{\mathbb{P}}_X) = \varphi(X),$$

where  $\phi(X, \mathbb{P}) = \mathbb{E}_{\mathbb{P}}[\|(I_d - XX^\top)(\xi - \mathbb{E}_{\mathbb{P}}[\xi])\|_{\mathbb{F}}^2]$  denotes the objective function of problem (3.1). The closed-form expression of  $\varphi$  is provided in Theorem 3.6. Recall that  $\Sigma_o \in \mathbb{S}_+^d$  is the covariance matrix of  $\xi$  under the nominal distribution  $\mathbb{P}_o$ . For convenience, we define  $\mathcal{S}(\Sigma_o) = \{X \in \mathcal{O}^{d,r} \mid (I_d - XX^\top)\Sigma_o = 0\}$ . Given that  $\Sigma_o \in \mathbb{S}_+^d$  and  $I_d - XX^\top \in \mathbb{S}_+^d$ , it is immediately evident that  $X \in \mathcal{S}(\Sigma_o)$  if and only if  $\text{tr}((I_d - XX^\top)\Sigma_o) = 0$ .

**THEOREM A.1.** *Let  $\bar{\mathbb{P}}_X$  be a probability distribution with the following covariance matrix,*

$$(A.2) \quad \bar{\Sigma}_X = \begin{cases} \bar{L}_X \Sigma_o \bar{L}_X, & \text{if } X \notin \mathcal{S}(\Sigma_o), \\ \Sigma_o + \frac{\rho^2}{d-r} (I_d - XX^\top), & \text{if } X \in \mathcal{S}(\Sigma_o), \end{cases}$$

where  $\bar{L}_X = I_d + \rho (\text{tr}((I_d - XX^\top)\Sigma_o))^{-1/2} (I_d - XX^\top) \in \mathbb{S}_{++}^d$ . Then, for any  $X \in \mathcal{O}^{d,r}$ , the relationship (A.1) holds.

*Proof.* By simple manipulations, we can obtain that

$$\phi(X, \mathbb{P}) = \text{tr} \left( (I_d - XX^\top) \mathbb{E}_{\mathbb{P}} \left[ (\xi - \mathbb{E}_{\mathbb{P}}[\xi]) (\xi - \mathbb{E}_{\mathbb{P}}[\xi])^\top \right] \right) = \text{tr}((I_d - XX^\top) \Sigma),$$

where  $\Sigma \in \mathbb{S}_+^d$  is the covariance matrix of  $\xi$  under the distribution  $\mathbb{P}$ . Thus, the objective function of problem (3.1) only depends on the covariance matrix. We investigate the following two cases.

**Case I:**  $X \notin \mathcal{S}(\Sigma_o)$ . It can be readily verified that

$$\text{tr}(\bar{\Sigma}_X) = \text{tr}(\bar{L}_X^2 \Sigma_o) = \text{tr}(\Sigma_o) + 2\rho (\text{tr}((I_d - XX^\top)\Sigma_o))^{1/2} + \rho^2.$$

By virtue of the orthogonality of  $X$ , we have  $\bar{L}_X X = X$ . Hence, it holds that

$$\begin{aligned} \phi(X, \bar{\mathbb{P}}_X) &= \text{tr}((I_d - XX^\top) \bar{\Sigma}_X) = \text{tr}(\bar{\Sigma}_X) - \text{tr}(X^\top \bar{L}_X \Sigma_o \bar{L}_X X) \\ &= \text{tr}(\Sigma_o) + 2\rho (\text{tr}((I_d - XX^\top)\Sigma_o))^{1/2} + \rho^2 - \text{tr}(X^\top \Sigma_o X) \\ &= \left( (\text{tr}((I_d - XX^\top)\Sigma_o))^{1/2} + \rho \right)^2 = \varphi(X). \end{aligned}$$

**Case II:**  $X \in \mathcal{S}(\Sigma_\circ)$ . Straightforward calculations give rise to that

$$\begin{aligned} \phi(X, \bar{\mathbb{P}}_X) &= \text{tr}((I_d - XX^\top) \bar{\Sigma}_X) = \text{tr}((I_d - XX^\top) \Sigma_\circ) + \frac{\rho^2}{d-r} \text{tr}(I_d - XX^\top) \\ &= \rho^2 = \varphi(X). \end{aligned}$$

We complete the proof.  $\square$

*Remark A.2.* The formulation in Theorem A.1 does not imply that the eigenvectors of  $\bar{\Sigma}_X$  are sparse. In general, the sparsity pattern of these eigenvectors depends on the eigenspace of  $\Sigma_\circ$ , the loading matrix  $X$ , and the Wasserstein radius  $\rho$ . No sparsity guarantee should be expected. Sparse PCA seeks to construct a sparse approximation to the principal subspace, balancing explained variance and interpretability. In general, the solution of sparse PCA does not necessarily coincide with that of PCA. From a theoretical perspective, sparsity regularizers are useful to remediate certain inconsistency phenomena inherent in PCA [33].

*Remark A.3.* The pair  $(X, \bar{\Sigma}_X)$  can not be interpreted as a Nash equilibrium in general. Intuitively, the construction of  $\bar{\Sigma}_X$  increases the variance in the residual subspace spanned by  $I_d - XX^\top$ . As a result, the leading eigenspace of  $\bar{\Sigma}_X$  may move away from  $X$ , rather than coincide with it. From the viewpoint of min-max optimization, the outer variable  $X$  in our DRO model ( $\text{P}_{\text{MM}}$ ) is constrained to lie on the nonconvex Stiefel manifold. Moreover, the objective function, with respect to  $X$ , is neither convex when viewed as a function in the Euclidean space, nor geodesically convex when regarded as a function on the Stiefel manifold. Therefore, the standard convex-concave assumptions for the existence of a Nash equilibrium are not available.

Even though Theorem A.1 furnishes a probability distribution  $\bar{\mathbb{P}}_X$  that attains the worst-case function value in problem (3.1), whether this construction falls within the feasible region  $\mathcal{B}_2(\mathbb{P}_\circ, \rho)$  remains, at present, beyond our reach to confirm for general cases. This difficulty stems from the fact that computing the Wasserstein distance between two arbitrary distributions is typically intractable, which is shown to be  $\#P$ -hard in [35]. Fortunately, the type-2 Wasserstein distance between two normal distributions admits a closed-form expression [23, 24]. Leveraging this result, we are able to identify a probability distribution that not only satisfies the relationship (A.1) but is also feasible for problem (3.1), when the nominal distribution is normal.

**COROLLARY A.4.** *Suppose that the nominal distribution  $\mathbb{P}_\circ$  is the normal distribution  $\mathbb{N}(\zeta_\circ, \Sigma_\circ)$  with  $\zeta_\circ \in \mathbb{R}^d$  and  $\Sigma_\circ \in \mathbb{S}_+^d$ . Let  $\bar{\Sigma}_X \in \mathbb{S}_+^d$  be the matrix constructed in (A.2). Then the global maximum of problem (3.1) is attained at  $\bar{\mathbb{P}}_X = \mathbb{N}(\zeta_\circ, \bar{\Sigma}_X)$ . Namely, the relationship (A.1) holds with  $\bar{\mathbb{P}}_X \in \mathcal{B}_2(\mathbb{P}_\circ, \rho)$ .*

*Proof.* This is a direct consequence of Theorem A.1 and [23, Theorem 2.1]. The proof is omitted here.  $\square$

## REFERENCES

- [1] P.-A. ABSIL, R. MAHONY, AND R. SEPULCHRE, *Optimization Algorithms on Matrix Manifolds*, Princeton University Press, Princeton, 2008.
- [2] A. BECK AND I. ROSSET, *A dynamic smoothing technique for a class of nonsmooth optimization problems on manifolds*, SIAM J. Optim., 33 (2023), pp. 1473–1493.
- [3] G. C. BENTO, O. P. FERREIRA, AND J. G. MELO, *Iteration-complexity of gradient, subgradient and proximal point methods on Riemannian manifolds*, J. Optim. Theory Appl., 173 (2017), pp. 548–562.

- [4] J. BLANCHET AND Y. KANG, *Sample out-of-sample inference based on Wasserstein distance*, Oper. Res., 69 (2021), pp. 985–1013.
- [5] J. BLANCHET, Y. KANG, AND K. MURTHY, *Robust Wasserstein profile inference and applications to machine learning*, J. Appl. Probab., 56 (2019), pp. 830–857.
- [6] J. BLANCHET AND K. MURTHY, *Quantifying distributional model risk via optimal transport*, Math. Oper. Res., 44 (2019), pp. 565–600.
- [7] N. BOUMAL, *An Introduction to Optimization on Smooth Manifolds*, Cambridge University Press, 2023.
- [8] N. BOUMAL, P.-A. ABSIL, AND C. CARTIS, *Global rates of convergence for nonconvex optimization on manifolds*, IMA J. Numer. Anal., 39 (2018), pp. 1–33.
- [9] S. CHEN, Z. DENG, S. MA, AND A. M.-C. SO, *Manifold proximal point algorithms for dual principal component pursuit and orthogonal dictionary learning*, IEEE Trans. Signal Process., 69 (2021), pp. 4759–4773.
- [10] S. CHEN, S. MA, A. M.-C. SO, AND T. ZHANG, *Proximal gradient method for nonsmooth optimization over the Stiefel manifold*, SIAM J. Optim., 30 (2020), pp. 210–239.
- [11] S. CHEN, S. MA, A. M.-C. SO, AND T. ZHANG, *Nonsmooth optimization over the Stiefel manifold and beyond: Proximal gradient method and recent variants*, SIAM Rev., 66 (2024), pp. 319–352.
- [12] X. CHEN, *Smoothing methods for nonsmooth, nonconvex minimization*, Math. Program., 134 (2012), pp. 71–99.
- [13] X. CHEN, Y. HE, AND Z. ZHANG, *Tight error bounds for the sign-constrained Stiefel manifold*, SIAM J. Optim., 35 (2025), pp. 302–329.
- [14] X. CHEN, H. SUN, AND H. XU, *Discrete approximation of two-stage stochastic and distributionally robust linear complementarity problems*, Math. Program., 177 (2019), pp. 255–289.
- [15] H. T. M. CHU, M. LIN, AND K.-C. TOH, *Wasserstein distributionally robust optimization and its tractable regularization formulations*, arXiv:2402.03942, (2024).
- [16] F. H. CLARKE, *Optimization and Nonsmooth Analysis*, Society for Industrial and Applied Mathematics, Philadelphia, 1990.
- [17] E. DELAGE AND Y. YE, *Distributionally robust optimization under moment uncertainty with application to data-driven problems*, Oper. Res., 58 (2010), pp. 595–612.
- [18] D. DRUSVYATSKIY AND C. PAQUETTE, *Efficiency of minimizing compositions of convex functions and smooth maps*, Math. Program., 178 (2019), pp. 503–558.
- [19] J. DUPAČOVÁ, *Stress testing via contamination*, in Coping with Uncertainty: Modeling and Policy Issues, Berlin, 2006, Springer, pp. 29–46.
- [20] P. M. ESFAHANI AND D. KUHN, *Data-driven distributionally robust optimization using the Wasserstein metric: Performance guarantees and tractable reformulations*, Math. Program., 171 (2018), pp. 115–166.
- [21] R. GAO, *Finite-sample guarantees for Wasserstein distributionally robust optimization: Breaking the curse of dimensionality*, Oper. Res., 71 (2023), pp. 2291–2306.
- [22] R. GAO AND A. KLEYWEGT, *Distributionally robust stochastic optimization with Wasserstein distance*, Math. Oper. Res., 48 (2023), pp. 603–655.
- [23] M. GELBRICH, *On a formula for the  $L^2$  Wasserstein metric between measures on Euclidean and Hilbert spaces*, Math. Nachr., 147 (1990), pp. 185–203.
- [24] C. R. GIVENS AND R. M. SHORTT, *A class of Wasserstein metrics for probability distributions*, Mich. Math. J., 31 (1984), pp. 231–240.
- [25] J. GOH AND M. SIM, *Distributionally robust optimization and its tractable approximations*, Oper. Res., 58 (2010), pp. 902–917.
- [26] G. A. HANASUSANTO, D. KUHN, S. W. WALLACE, AND S. ZYMLER, *Distributionally robust multi-item newsvendor problems with multimodal demand distributions*, Math. Program., 152 (2015), pp. 1–32.
- [27] B. HE, Y. YOU, AND X. YUAN, *On the convergence of primal-dual hybrid gradient algorithm*, SIAM J. Imaging Sci., 7 (2014), pp. 2526–2537.
- [28] N. J. HIGHAM, *Functions of Matrices: Theory and Computation*, Society for Industrial and Applied Mathematics, Philadelphia, 2008.
- [29] S. HOSSEINI, W. HUANG, AND R. YOUSEFPOUR, *Line search algorithms for locally Lipschitz functions on Riemannian manifolds*, SIAM J. Optim., 28 (2018), pp. 596–619.
- [30] S. HOSSEINI AND A. USCHMAJEW, *A Riemannian gradient sampling algorithm for nonsmooth optimization on manifolds*, SIAM J. Optim., 27 (2017), pp. 173–189.
- [31] W. HUANG AND K. WEI, *Riemannian proximal gradient methods*, Math. Program., 194 (2022), pp. 371–413.
- [32] J. JIANG, H. SUN, AND X. CHEN, *Data-driven distributionally robust multiproduct pricing problems under pure characteristics demand models*, SIAM J. Optim., 34 (2024), pp. 2917–

- 2942.
- [33] I. M. JOHNSTONE AND A. Y. LU, *On consistency and sparsity for principal components analysis in high dimensions*, J. Am. Stat. Assoc., 104 (2009), pp. 682–693.
- [34] L. V. KANTOROVICH AND S. RUBINShteIN, *On a space of totally additive functions*, Vestn. St. Petersburg. Univ.: Math., 13 (1958), pp. 52–59.
- [35] D. KUHN, P. M. ESFAHANI, V. A. NGUYEN, AND S. SHAFIEEZADEH-ABADEH, *Wasserstein distributionally robust optimization: Theory and applications in machine learning*, in Operations Research & Management Science in the Age of Analytics, INFORMS, 2019, pp. 130–166.
- [36] D. KUHN, S. SHAFIEE, AND W. WIESEMANN, *Distributionally robust optimization*, Acta Numer., 34 (2025), pp. 579–804.
- [37] R. LAI AND S. OSHER, *A splitting method for orthogonality constrained problems*, J. Sci. Comput., 58 (2014), pp. 431–449.
- [38] X. LI, S. CHEN, Z. DENG, Q. QU, Z. ZHU, AND A. M.-C. SO, *Weakly convex optimization over Stiefel manifold using Riemannian subgradient-type methods*, SIAM J. Optim., 31 (2021), pp. 1605–1634.
- [39] X. LIU, N. XIAO, AND Y.-X. YUAN, *A penalty-free infeasible approach for a class of nonsmooth optimization problems over the Stiefel manifold*, J. Sci. Comput., 99 (2024), p. 30.
- [40] G.-F. LU, J. ZOU, Y. WANG, AND Z. WANG, *L1-norm-based principal component analysis with adaptive regularization*, Pattern Recognit., 60 (2016), pp. 901–907.
- [41] Z. LU AND Y. ZHANG, *An augmented Lagrangian approach for sparse principal component analysis*, Math. Program., 135 (2012), pp. 149–193.
- [42] V. A. NGUYEN, D. KUHN, AND P. M. ESFAHANI, *Distributionally robust inverse covariance estimation: The Wasserstein shrinkage estimator*, Oper. Res., 70 (2022), pp. 490–515.
- [43] R. T. ROCKAFELLAR, *Convex Analysis*, Princeton University Press, Princeton, 1970.
- [44] R. T. ROCKAFELLAR AND R. J.-B. WETS, *Variational Analysis*, Springer Science & Business Media, 2009.
- [45] A. SHAPIRO, *On duality theory of conic linear problems*, in Semi-Infinite Programming: Recent Advances, M. Á. Goberna and M. A. López, eds., Springer, Boston, 2001, pp. 135–165.
- [46] B. P. VAN PARYS, P. M. ESFAHANI, AND D. KUHN, *From data to decisions: Distributionally robust optimization is optimal*, Manag. Sci., 67 (2021), pp. 3387–3402.
- [47] H. VU, T. TRAN, M.-C. YUE, AND V. A. NGUYEN, *Distributionally robust fair principal components via geodesic descents*, in Proc. 10th Int. Conf. Learn. Represent., 2022, pp. 1–23.
- [48] L. WANG, L. BAO, AND X. LIU, *A decentralized proximal gradient tracking algorithm for composite optimization on Riemannian manifolds*, J. Mach. Learn. Res., 26 (2025), pp. 1–37.
- [49] L. WANG AND X. CHEN, *An adaptive smoothing algorithm for non-Lipschitz optimization on manifolds with complexity guarantees*, arXiv:2604.18325, (2026).
- [50] L. WANG AND X. LIU, *Decentralized optimization over the Stiefel manifold by an approximate augmented Lagrangian function*, IEEE Trans. Signal Process., 70 (2022), pp. 3029–3041.
- [51] L. WANG AND X. LIU, *Smoothing gradient tracking for decentralized optimization over the Stiefel manifold with non-smooth regularizers*, in Proc. 62nd IEEE Conf. Decis. Control., IEEE, 2023, pp. 126–132.
- [52] L. WANG, X. LIU, AND X. CHEN, *A support-set algorithm for optimization problems with nonnegative and orthogonal constraints*, arXiv:2511.03443, (2025).
- [53] L. WANG, X. LIU, AND Y. ZHANG, *Seeking consensus on subspaces in federated principal component analysis*, J. Optim. Theory Appl., 203 (2024), pp. 529–561.
- [54] Z. WANG, B. LIU, S. CHEN, S. MA, L. XUE, AND H. ZHAO, *A manifold proximal linear method for sparse spectral clustering with application to single-cell RNA sequencing data analysis*, INFORMS J. Optim., 4 (2022), pp. 200–214.
- [55] Z. WEN AND W. YIN, *A feasible method for optimization with orthogonality constraints*, Math. Program., 142 (2013), pp. 397–434.
- [56] W. WIESEMANN, D. KUHN, AND B. RUSTEM, *Robust Markov decision processes*, Math. Oper. Res., 38 (2013), pp. 153–183.
- [57] H. XU, Y. LIU, AND H. SUN, *Distributionally robust optimization with matrix moment constraints: Lagrange duality and cutting plane methods*, Math. Program., 169 (2018), pp. 489–529.
- [58] W. H. YANG, L.-H. ZHANG, AND R. SONG, *Optimality conditions for the nonlinear programming problems on Riemannian manifolds*, Pac. J. Optim., 10 (2014), pp. 415–434.
- [59] J. ZHANG, S. MA, AND S. ZHANG, *Primal-dual optimization algorithms over Riemannian manifolds: An iteration complexity analysis*, Math. Program., 184 (2020), pp. 445–490.
- [60] L. ZHANG, J. YANG, AND R. GAO, *A short and general duality proof for Wasserstein distributionally robust optimization*, Oper. Res., 73 (2025), pp. 2146–2155.