

First-Order Algorithms for Min-Max Optimization in Geodesic Metric Spaces

Michael I. Jordan $\diamond, \dagger$ Tianyi Lin $\diamond$ Emmanouil V. Vlatakis-Gkaragkounis $\diamond$

Department of Electrical Engineering and Computer Sciences $\diamond$

Department of Statistics $\dagger$

University of California, Berkeley

September 29, 2022

Abstract

From optimal transport to robust dimensionality reduction, a plethora of machine learning applications can be cast into the min-max optimization problems over Riemannian manifolds. Though many min-max algorithms have been analyzed in the Euclidean setting, it has proved elusive to translate these results to the Riemannian case. [Zhang et al. \[2022\]](#) have recently shown that geodesic convex concave Riemannian problems always admit saddle-point solutions. Inspired by this result, we study whether a performance gap between Riemannian and optimal Euclidean space convex-concave algorithms is necessary. We answer this question in the negative—we prove that the Riemannian corrected extragradient (RCEG) method achieves last-iterate convergence at a linear rate in the geodesically strongly-convex-concave case, matching the Euclidean result. Our results also extend to the stochastic or non-smooth case where RCEG and Riemannian gradient ascent descent (RGDA) achieve near-optimal convergence rates up to factors depending on curvature of the manifold.

1 Introduction

Constrained optimization problems arise throughout machine learning, in classical settings such as dimension reduction [[Boumal and Absil, 2011](#)], dictionary learning [[Sun et al., 2016a,b](#)], and deep neural networks [[Huang et al., 2018](#)], but also in emerging problems involving decision-making and multi-agent interactions. While simple convex constraints (such as norm constraints) can be easily incorporated in standard optimization formulations, notably (proximal) gradient descent [[Raskutti and Mukherjee, 2015](#), [Giannou et al., 2021b,a](#), [Antonakopoulos et al., 2020](#), [Vlatakis-Gkaragkounis et al., 2020](#)], in a range of other applications such as matrix recovery [[Fornasier et al., 2011](#), [Candes et al., 2008](#)], low-rank matrix factorization [[Han et al., 2021](#)] and generative adversarial nets [[Goodfellow et al., 2014](#)], the constraints are fundamentally nonconvex and are often treated via special heuristics.

Thus, a general goal is to design algorithms that systematically take account of special geometric structure of the feasible set [[Mei et al., 2021](#), [Lojasiewicz, 1963](#), [Polyak, 1963](#)]. A long line of work in the machine learning (ML) community has focused on understanding the geometric properties of commonly used constraints and how they affect optimization; [see, e.g., [Ge et al., 2015](#), [Anandkumar and Ge, 2016](#), [Sra and Hosseini, 2016](#), [Jin et al., 2017](#), [Ge et al., 2017](#), [Du et al., 2017](#), [Reddi et al., 2018](#), [Criscitiello and Boumal, 2019](#), [Jin et al., 2021](#)]. A prominent aspect of this agenda has been the re-expression of these constraints through the lens of Riemannian manifolds. This has given rise to new algorithms [[Sra and Hosseini, 2015](#), [Hosseini and Sra, 2015](#)] with a wide range of ML applications, including online principal component analysis (PCA), the computation of Mahalanobis distance from

noisy measurements [Bonnabel, 2013], consensus distributed algorithms for aggregation in ad-hoc wireless networks [Tron et al., 2012] and maximum likelihood estimation for certain non-Gaussian (heavy- or light-tailed) distributions [Wiesel, 2012].

Going beyond simple minimization problems, the robustification of many ML tasks can be formulated as min-max optimization problems. Well-known examples in this domain include adversarial machine learning [Kumar et al., 2017, Chen et al., 2018], optimal transport [Lin et al., 2020a], and online learning [Mertikopoulos and Sandholm, 2018, Bomze et al., 2019, Antonakopoulos et al., 2020]. Similar to their minimization counterparts, non-convex constraints have been widely applicable to the min-max optimization as well [Heusel et al., 2017, Daskalakis and Panageas, 2018, Balduzzi et al., 2018, Mertikopoulos et al., 2019, Jin et al., 2020]. Recently there has been significant effort in proving tighter results either under more structured assumptions [Thekumprampil et al., 2019, Nouiehed et al., 2019, Lu et al., 2020, Azizian et al., 2020, Diakonikolas, 2020, Golowich et al., 2020, Lin et al., 2020c,b, Liu et al., 2021, Ostrovskii et al., 2021, Kong and Monteiro, 2021], and/or obtaining last-iterate convergence guarantees [Daskalakis and Panageas, 2018, 2019, Mertikopoulos et al., 2019, Adolphs et al., 2019, Liang and Stokes, 2019, Gidel et al., 2019, Mazumdar et al., 2020, Liu et al., 2020, Mokhtari et al., 2020, Lin et al., 2020c, Hamedani and Aybat, 2021, Abernethy et al., 2021, Cai et al., 2022] for computing min-max solutions in convex-concave settings. Nonetheless, the analysis of the iteration complexity in the general *non-convex non-concave* setting is still in its infancy [Vlatakis-Gkaragkounis et al., 2019, 2021]. In response, the optimization community has recently studied how to extend standard min-max optimization algorithms such as gradient descent ascent (GDA) and extragradient (EG) to the Riemannian setting. In mathematical terms, given two Riemannian manifolds $\mathcal{M}, \mathcal{N}$ and a function $f : \mathcal{M} \times \mathcal{N} \rightarrow \mathbb{R}$, the Riemannian min-max optimization (RMMO) problem becomes

$$\min_{x \in \mathcal{M}} \max_{y \in \mathcal{N}} f(x, y).$$

The change of geometry from Euclidean to Riemannian poses several difficulties. Indeed, a fundamental stumbling block has been that this problem may not even have theoretically meaningful solutions. In contrast with minimization where an optimal solution in a bounded domain is always guaranteed [Fearnley et al., 2021], existence of such saddle points necessitates typically the application of topological fixed point theorems [Brouwer, 1911, Kakutani, 1941], KKM Theory [Knaster et al., 1929]). For the case of convex-concave f with compact sets $\mathcal{X}$ and $\mathcal{Y}$, Sion [1958] generalized the celebrated theorem [Neumann, 1928] and guaranteed that a solution (x^*, y^*) with the following property exists

$$\min_{x \in \mathcal{X}} f(x, y^*) = f(x^*, y^*) = \max_{y \in \mathcal{Y}} f(x^*, y).$$

However, at the core of the proof of this result is an ingenious application of Helly’s lemma [Helly, 1923] for the sublevel sets of f , and, until the work of Ivanov [2014], it has been unclear how to formulate an analogous lemma for the Riemannian geometry. As a result, until recently have extensions of the min-max theorem been established, and only for restricted manifold families [Komiya, 1988, Kristály, 2014, Park, 2019].

Zhang et al. [2022] was the first to establish a min-max theorem for a flurry of Riemannian manifolds equipped with unique geodesics. Notice that this family is not a mathematical artifact since it encompasses many practical applications of RMMO, including Hadamard and Stiefel ones used in PCA [Lee et al., 2022]. Intuitively, the unique geodesic between two points of a manifold is the analogue of the a linear segment between two points in convex set: For any two points $x_1, x_2 \in \mathcal{X}$, their connecting geodesic is the unique shortest path contained in $\mathcal{X}$ that connects them.

Even when the RMMO is well defined, transferring the guarantees of traditional min-max optimization algorithms like Gradient Ascent Descent (GDA) and Extra-Gradient (EG) to the Riemannian case is non-trivial. Intuitively speaking, in the Euclidean realm the main leitmotif of the last-iterate analyses the aforementioned algorithms is a proof that $\delta_t = \|x_t - x^*\|^2$ is decreasing over time. To achieve this, typically the proof correlates δ_t and δ_{t-1} via a “square expansion,” namely:

$$\underbrace{\|x_{t-1} - x^*\|^2}_{\alpha^2} = \underbrace{\|x_t - x^*\|^2}_{\beta^2} + \underbrace{\|x_{t-1} - x_t\|^2}_{\gamma^2} - \underbrace{2\langle x_t - x^*, x_{t-1} - x_t \rangle}_{2\beta\gamma\cos(\hat{A})}. \quad (1.1)$$

Notice, however that the above expression relies strongly on properties of Euclidean geometry (and the flatness of the corresponding line), namely that the the lines connecting the three points x_t , x_{t-1} and x^* form a triangle; indeed, it is the generalization of the Pythagorean theorem, known also as the law of cosines, for the induced triangle $(ABC) := \{(x_t, x_{t-1}, x^*)\}$. In a uniquely geodesic manifold such triangle may not belong to the manifold as discussed above. As a result, the difference of distances to the equilibrium using the geodesic paths $d_{\mathcal{M}}^2(x_t, x^*) - d_{\mathcal{M}}^2(x_{t-1}, x^*)$ generally cannot be given in a closed form. The manifold’s curvature controls how close these paths are to forming a Euclidean triangle. In fact, the phenomenon of *distance distortion*, as it is typically called, was hypothesised by Zhang et al. [2022, Section 4.2] to be the cause of exponential slowdowns when applying EG to RMMO problems when compared to their Euclidean counterparts.

Multiple attempts have been made to bypass this hurdle. Huang et al. [2020] analyzed the Riemannian GDA (RGDA) for the non-convex non-concave setting. However, they do not present any last-iterate convergence results and, even in the average/best iterate setting, they only derive sub-optimal rates for the geodesic convex-concave setting due to the lack of the machinery that convex analysis and optimization offers they derive sub-optimal rates for the geodesic convex-concave case, which is the problem of our interest. The analysis of Han et al. [2022] for Riemannian Hamiltonian Method (RHM), matches the rate of second-order methods in the Euclidean case. Although theoretically faster in terms of iterations, second-order methods are not preferred in practice since evaluating second order derivatives for optimization problems of thousands to millions of parameters quickly becomes prohibitive. Finally, Zhang et al. [2022] leveraged the standard averaging output trick in EG to derive a sublinear convergence rate of $O(1/\epsilon)$ for the general geodesically convex-concave Riemannian framework. In addition, they conjectured that the use of a different method could close the exponential gap for the geodesically strongly-convex-strongly-convex scenario and its Euclidean counterpart.

Given this background, a crucial question underlying the potential for successful application of first-order algorithms to Riemannian settings is the following:

Is a performance gap necessary between Riemannian and Euclidean optimal convex-concave algorithms in terms of accuracy and the condition number?

1.1 Our Contributions

Our aim in this paper is to provide an extensive analysis of the Riemannian counterparts of Euclidean optimal first-order methods adapted to the manifold-constrained setting. For the case of the smooth objectives, we consider the *Riemannian corrected extragradient* (RCEG) method while for non-smooth cases, we analyze the textbook *Riemannian gradient descent ascent* (RGDA) method. Our main results are summarized in the following table.

Alg: *RCEG*. Smooth setting with ℓ -Lipschitz Gradient (cf. Assumption 2.1, 3.1 and 3.2)

Perf. Measure	Setting	Complexity	Theorem
Last-Iterate	Det. GSCSC	$O\left(\kappa(\sqrt{\tau_0} + \frac{1}{\xi_0})\log(\frac{1}{\epsilon})\right)$	Thm. 3.1
Last-Iterate	Stoc. GSCSC	$O\left(\kappa(\sqrt{\tau_0} + \frac{1}{\xi_0})\log(\frac{1}{\epsilon}) + \frac{\sigma^2\bar{\xi}_0}{\mu^2\epsilon}\log(\frac{1}{\epsilon})\right)$	Thm. 3.2
Avg-Iterate	Det. GCC	$O\left(\frac{\ell\sqrt{\tau_0}}{\epsilon}\right)$	[Zhang et al., 2022, Thm.1]
Avg-Iterate	Stoc. GCC	$O\left(\frac{\ell\sqrt{\tau_0}}{\epsilon} + \frac{\sigma^2\bar{\xi}_0}{\epsilon^2}\right)$	Thm. 3.3

Alg: *RGDA*. Nonsmooth setting with L -Lipschitz Function (cf. Assumption D.1 and D.2)

Last-Iterate	Det. GSCSC	$O\left(\frac{L^2\bar{\xi}_0}{\mu^2\epsilon}\right)$	Thm. D.1
Last-Iterate	Stoc. GSCSC	$O\left(\frac{(L^2+\sigma^2)\bar{\xi}_0}{\mu^2\epsilon}\right)$	Thm. D.3
Avg-Iterate	Det. GCC	$O\left(\frac{L^2\bar{\xi}_0}{\epsilon^2}\right)$	Thm. D.2
Avg-Iterate	Stoc. GCC	$O\left(\frac{(L^2+\sigma^2)\bar{\xi}_0}{\epsilon^2}\right)$	Thm. D.4

For the definition of the acronyms, Det and Stoc stand for deterministic and stochastic, respectively. GSCSC and GCC stand for geodesically strongly-convex-strongly-concave (cf. Assumption 3.1 or Assumption D.1) and geodesically convex-concave (cf. Assumption 3.2 or Assumption D.2). Here $\epsilon \in (0, 1)$ is the accuracy, L, ℓ the Lipschitzness of the objective and its gradient, $\kappa = \ell/\mu$ is the condition number of the function, where μ is the strong convexity parameter, $(\tau_0, \xi_0, \bar{\xi}_0)$ are curvature parameters (cf. Assumption 2.1), and σ^2 is the variance of a Riemannian gradient estimator.

Our first main contribution is the derivation of a linear convergence rate for RCEG, answering the open conjecture of Zhang et al. [2022] about the performance gap of single-loop extragradient methods. Indeed, while a direct comparison between $d_{\mathcal{M}}^2(x_t, x^*)$ and $d_{\mathcal{M}}^2(x_{t-1}, x^*)$ is infeasible, we are able to establish a relationship between the iterates via appeal to the duality gap function and obtain a contraction in terms of $d_{\mathcal{M}}^2(x_t, x^*)$. In other words, the effect of Riemannian distance distortion is quantitative (the contraction ratio will depend on it) rather than qualitative (the geometric contraction still remains under a proper choice of constant stepsize). More specifically, we use $d_{\mathcal{M}}^2(x_t, x^*) + d_{\mathcal{N}}^2(y_t, y^*)$ and $d_{\mathcal{M}}^2(x_{t+1}, x^*) + d_{\mathcal{N}}^2(y_{t+1}, y^*)$ to bound a gap function defined by $f(\hat{x}_t, y^*) - f(x^*, \hat{y}_t)$. Since the objective function is geodesically strongly-convex-strongly-concave, we have $f(\hat{x}_t, y^*) - f(x^*, \hat{y}_t)$ is lower bounded by $\frac{\mu}{2}(d_{\mathcal{M}}(\hat{x}_t, x^*)^2 + d_{\mathcal{N}}(\hat{y}_t, y^*)^2)$. Then, using the relationship between (x_t, y_t) and $(\hat{x}_t, \hat{y}_t)$, we conclude the desired results in Theorem 3.1. Notably, our approach is not affected by the nonlinear geometry of the manifold.

Secondly, we endeavor to give a systematic analysis of aspects of the objective function, including its smoothness, its convexity and oracle access. As we shall see, similar to the Euclidean case, better finite-time convergence guarantees are connected with a geodesic smoothness condition. For the sake of completeness, in the paper's supplement we present the performance of Riemannian GDA for the full spectrum of stochasticity for the non-smooth case. More specifically, for the stochastic setting, the key ingredient to get the optimal convergence rate is to carefully select the step size such that the noise of the gradient estimator will not affect the final convergence rate significantly. As a highlight, such technique has been used for analyzing stochastic RCEG in the Euclidean setting [Kotsalis et al., 2022] and our analysis can be seen as the extension to the Riemannian setting. For the nonsmooth setting, the analysis is relatively simpler compared to smooth settings but we still need to deal with the issue caused by the nonlinear geometry of manifolds and the interplay between the distortion of Riemannian metrics, the gap function and the bounds of Lipschitzness of our bi-objective. Interestingly, the rates

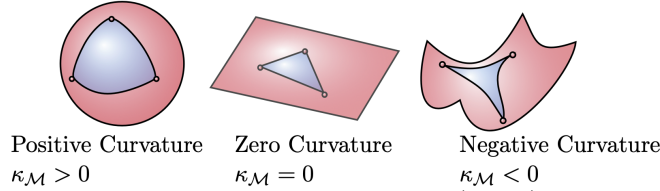
we derive are near optimal in terms of accuracy and condition number of the objective, and analogous to their Euclidean counterparts.

2 Preliminaries and Technical Background

We present the basic setup and optimality conditions for Riemannian min-max optimization. Indeed, we focus on some of key concepts that we need from Riemannian geometry, deferring a fuller presentation, including motivating examples and further discussion of related work, to Appendix A-C.

Riemannian geometry. An n -dimensional manifold $\mathcal{M}$ is a topological space where any point has a neighborhood that is homeomorphic to the n -dimensional Euclidean space. For each $x \in \mathcal{M}$, each tangent vector is tangent to all parametrized curves passing through x and the tangent space $T_x\mathcal{M}$ of a manifold $\mathcal{M}$ at this point is defined as the set of all tangent vectors. A Riemannian manifold $\mathcal{M}$ is a smooth manifold that is endowed with a smooth (“Riemannian”) metric $\langle \cdot, \cdot \rangle_x$ on the tangent space $T_x\mathcal{M}$ for each point $x \in \mathcal{M}$. The inner metric induces a norm $\| \cdot \|_x$ on the tangent spaces.

A geodesic can be seen as the generalization of an Euclidean linear segment and is modeled as a smooth curve (map), $\gamma : [0, 1] \mapsto \mathcal{M}$, which is locally a distance minimizer. Additionally, because of the non-flatness of a manifold a different relation between the angles and the lengths of an arbitrary geodesic triangle is induced. This distortion can be quantified via the *sectional curvature* parameter $\kappa_{\mathcal{M}}$ thanks to Toponogov’s theorem [Cheeger and Ebin, 1975, Burago et al., 1992]. A constructive consequence of



this definition are the trigonometric comparison inequalities (TCIs) that will be essential in our proofs; see Alimisis et al. [2020, Corollary 2.1] and Zhang and Sra [2016, Lemma 5] for detailed derivations. Assuming bounded sectional curvature, TCIs provide a tool for bounding Riemannian “inner products” that are more troublesome than classical Euclidean inner products.

The following proposition summarizes the TCIs that we will need; note that if $\kappa_{\min} = \kappa_{\max} = 0$ (i.e., Euclidean spaces), then the proposition reduces to the law of cosines.

Proposition 2.1 *Suppose that $\mathcal{M}$ is a Riemannian manifold and let Δ be a geodesic triangle in $\mathcal{M}$ with the side length a, b, c and let A be the angle between b and c . Then, we have*

1. *If $\kappa_{\mathcal{M}}$ that is upper bounded by $\kappa_{\max} > 0$ and the diameter of $\mathcal{M}$ is bounded by $\frac{\pi}{\sqrt{\kappa_{\max}}}$, then*

$$a^2 \geq \underline{\xi}(\kappa_{\max}, c) \cdot b^2 + c^2 - 2bc \cos(A),$$

where $\underline{\xi}(\kappa, c) := 1$ for $\kappa \leq 0$ and $\underline{\xi}(\kappa, c) := c\sqrt{\kappa} \cot(c\sqrt{\kappa}) < 1$ for $\kappa > 0$.

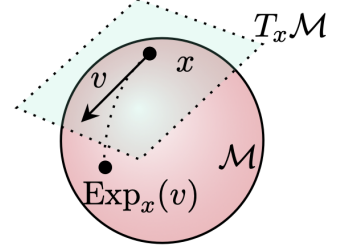
2. *If $\kappa_{\mathcal{M}}$ is lower bounded by $\kappa_{\min}$, then*

$$a^2 \leq \bar{\xi}(\kappa_{\min}, c) \cdot b^2 + c^2 - 2bc \cos(A),$$

where $\bar{\xi}(\kappa, c) := c\sqrt{-\kappa} \coth(c\sqrt{-\kappa}) > 1$ if $\kappa < 0$ and $\bar{\xi}(\kappa, c) := 1$ if $\kappa \geq 0$.

Also, in contrast to the Euclidean case, x and $v = \text{grad}_x f(x)$ do not lie in the same space, since $\mathcal{M}$ and $T_x \mathcal{M}$ respectively are distinct entities. The interplay between these dual spaces typically is carried out via the *exponential maps*. An exponential map at a point $x \in \mathcal{M}$ is a mapping from the tangent space $T_x \mathcal{M}$ to $\mathcal{M}$. In particular, $y := \text{Exp}_x(v) \in \mathcal{M}$ is defined such that there exists a geodesic $\gamma : [0, 1] \mapsto \mathcal{M}$ satisfying $\gamma(0) = x$, $\gamma(1) = y$ and $\gamma'(0) = v$. The inverse map exists since the manifold has a unique geodesic between any two points, which we denote as $\text{Exp}_x^{-1} : \mathcal{M} \mapsto T_x \mathcal{M}$. Accordingly, we have $d_{\mathcal{M}}(x, y) = \|\text{Exp}_x^{-1}(y)\|_x$ is the Riemannian distance induced by the exponential map.

Finally, in contrast again to Euclidean spaces, we cannot compare the tangent vectors at different points $x, y \in \mathcal{M}$ since these vectors lie in different tangent spaces. To resolve this issue, it suffices to define a transport mapping that moves a tangent vector along the geodesics and also preserves the length and Riemannian metric $\langle \cdot, \cdot \rangle_x$; indeed, we can define a parallel transport $\Gamma_x^y : T_x \mathcal{M} \mapsto T_y \mathcal{M}$ such that the inner product between any $u, v \in T_x \mathcal{M}$ is preserved; i.e., $\langle u, v \rangle_x = \langle \Gamma_x^y(u), \Gamma_x^y(v) \rangle_y$.



Riemannian min-max optimization and function classes. We let $\mathcal{M}$ and $\mathcal{N}$ be Riemannian manifolds with unique geodesic and bounded sectional curvature and assume that the function $f : \mathcal{M} \times \mathcal{N} \mapsto \mathbb{R}$ is defined on the product of these manifolds. The regularity conditions that we impose on the function f are as follows.

Definition 2.1 A function $f : \mathcal{M} \times \mathcal{N} \mapsto \mathbb{R}$ is geodesically L -Lipschitz if for $\forall x, x' \in \mathcal{M}$ and $\forall y, y' \in \mathcal{N}$, the following statement holds true: $|f(x, y) - f(x', y')| \leq L(d_{\mathcal{M}}(x, x') + d_{\mathcal{N}}(y, y'))$. Additionally, if function f is also differentiable, it is called geodesically ℓ -smooth if for $\forall x, x' \in \mathcal{M}$ and $\forall y, y' \in \mathcal{N}$, the following statement holds true,

$$\begin{aligned} \|\text{grad}_x f(x, y) - \Gamma_{x'}^x \text{grad}_x f(x', y')\| &\leq \ell(d_{\mathcal{M}}(x, x') + d_{\mathcal{N}}(y, y')), \\ \|\text{grad}_y f(x, y) - \Gamma_{y'}^y \text{grad}_y f(x', y')\| &\leq \ell(d_{\mathcal{M}}(x, x') + d_{\mathcal{N}}(y, y')), \end{aligned}$$

where $(\text{grad}_x f(x', y'), \text{grad}_y f(x', y')) \in T_{x'} \mathcal{M} \times T_{y'} \mathcal{N}$ is the Riemannian gradient of f at (x', y') , $\Gamma_{x'}^x$ is the parallel transport of $\mathcal{M}$ from x' to x , and $\Gamma_{y'}^y$ is the parallel transport of $\mathcal{N}$ from y' to y .

Definition 2.2 A function $f : \mathcal{M} \times \mathcal{N} \mapsto \mathbb{R}$ is geodesically strongly-convex-strongly-concave with the modulus $\mu > 0$ if the following statement holds true,

$$\begin{aligned} f(x', y) &\geq f(x, y) + \langle \text{subgrad}_x f(x, y), \text{Exp}_x^{-1}(x') \rangle_x + \frac{\mu}{2}(d_{\mathcal{M}}(x, x'))^2, & \text{for each } y \in \mathcal{N}, \\ f(x, y') &\leq f(x, y) + \langle \text{subgrad}_y f(x, y), \text{Exp}_y^{-1}(y') \rangle_y - \frac{\mu}{2}(d_{\mathcal{N}}(y, y'))^2, & \text{for each } x \in \mathcal{M}. \end{aligned}$$

where $(\text{subgrad}_x f(x', y'), \text{subgrad}_y f(x', y')) \in T_{x'} \mathcal{M} \times T_{y'} \mathcal{N}$ is a Riemannian subgradient of f at a point (x', y') . A function f is geodesically convex-concave if the above holds true with $\mu = 0$.

Following standard conventions in Riemannian optimization [Zhang and Sra, 2016, Alimisis et al., 2020, Zhang et al., 2022], we make the following assumptions on the manifolds and objective functions:¹

Assumption 2.1 The objective function $f : \mathcal{M} \times \mathcal{N} \mapsto \mathbb{R}$ and manifolds $\mathcal{M}$ and $\mathcal{N}$ satisfy

1. The diameter of the domain $\{(x, y) \in \mathcal{M} \times \mathcal{N} : -\infty < f(x, y) < +\infty\}$ is bounded by $D > 0$.

¹In particular, our assumed upper and lower bounds $\kappa_{\min}, \kappa_{\max}$ guarantee that TCIs in Proposition 2.1 can be used in our analysis for proving finite-time convergence.

2. $\mathcal{M}, \mathcal{N}$ admit unique geodesic paths for any $(x, y), (x', y') \in \mathcal{M} \times \mathcal{N}$.
3. The sectional curvatures of $\mathcal{M}$ and $\mathcal{N}$ are both bounded in the range $[\kappa_{\min}, \kappa_{\max}]$ with $\kappa_{\min} \leq 0$. If $\kappa_{\max} > 0$, we assume that the diameter of manifolds is bounded by $\frac{\pi}{\sqrt{\kappa_{\max}}}$.

Under these conditions, Zhang et al. [2022] proved an analog of Sion’s minimax theorem [Sion, 1958] in geodesic metric spaces. Formally, we have

$$\max_{y \in \mathcal{N}} \min_{x \in \mathcal{M}} f(x, y) = \min_{x \in \mathcal{M}} \max_{y \in \mathcal{N}} f(x, y),$$

which guarantees that there exists at least one global saddle point $(x^*, y^*) \in \mathcal{M} \times \mathcal{N}$ such that $\min_{x \in \mathcal{M}} f(x, y^*) = f(x^*, y^*) = \max_{y \in \mathcal{N}} f(x^*, y)$. Note that the unicity of geodesics assumption is algorithm-independent and is imposed for guaranteeing that a saddle-point solution always exist. Even though this rules out many manifolds of interest, there are still many manifolds that satisfy such conditions. More specifically, the Hadamard manifold (manifolds with non-positive curvature, $\kappa_{\max} = 0$) has a unique geodesic between any two points. This also becomes a common regularity condition in Riemannian optimization [Zhang and Sra, 2016, Alimisis et al., 2020]. For any point $(\hat{x}, \hat{y}) \in \mathcal{M} \times \mathcal{N}$, the duality gap $f(\hat{x}, y^*) - f(x^*, \hat{y})$ thus gives an optimality criterion.

Definition 2.3 A point $(\hat{x}, \hat{y}) \in \mathcal{M} \times \mathcal{N}$ is an ϵ -saddle point of a geodesically convex-concave function $f(\cdot, \cdot)$ if $f(\hat{x}, y^*) - f(x^*, \hat{y}) \leq \epsilon$ where $(x^*, y^*) \in \mathcal{M} \times \mathcal{N}$ is a global saddle point.

In the setting where f is geodesically strongly-convex-strongly-concave with $\mu > 0$, it is not difficult to verify the uniqueness of a global saddle point $(x^*, y^*) \in \mathcal{M} \times \mathcal{N}$. Then, we can consider the distance gap $(d(\hat{x}, x^*))^2 + (d(\hat{y}, y^*))^2$ as an optimality criterion for any point $(\hat{x}, \hat{y}) \in \mathcal{M} \times \mathcal{N}$.

Definition 2.4 A point $(\hat{x}, \hat{y}) \in \mathcal{M} \times \mathcal{N}$ is an ϵ -saddle point of a geodesically strongly-convex-strongly-concave function $f(\cdot, \cdot)$ if $(d(\hat{x}, x^*))^2 + (d(\hat{y}, y^*))^2 \leq \epsilon$, where $(x^*, y^*) \in \mathcal{M} \times \mathcal{N}$ is a global saddle point. If f is also geodesically ℓ -smooth, we denote $\kappa = \frac{\ell}{\mu}$ as the condition number.

Given the above definitions, we can ask whether it is possible to find an ϵ -saddle point efficiently or not. In this context, Zhang et al. [2022] have answered this question in the affirmative for the setting where f is geodesically ℓ -smooth and geodesically convex-concave; indeed, they derive the convergence rate of Riemannian corrected extragradient (RCEG) method in terms of time-average iterates and also conjecture that RCEG does not guarantee convergence at a linear rate in terms of last iterates when f is geodesically ℓ -smooth and geodesically strongly-convex-strongly-concave, due to the existence of distance distortion; see Zhang et al. [2022, Section 4.2]. Surprisingly, we show in Section 3 that RCEG with constant stepsize can achieve last-iterate convergence at a linear rate. Moreover, we establish the optimal convergence rates of stochastic RCEG for certain choices of stepsize for both geodesically convex-concave and geodesically strongly-convex-strongly-concave settings.

3 Riemannian Corrected Extragradient Method

In this section, we revisit the scheme of Riemannian corrected extragradient (RCEG) method proposed by Zhang et al. [2022] and extend it to a stochastic algorithm that we refer to as *stochastic RCEG*. We present our main results on an optimal last-iterate convergence guarantee for the geodesically strongly-convex-strongly-concave setting (both deterministic and stochastic) and a time-average convergence guarantee for the geodesically convex-concave setting (stochastic). This complements the time-average convergence guarantee for geodesically convex-concave setting (deterministic) [Zhang et al., 2022, Theorem 4.1] and resolves an open problem posted in Zhang et al. [2022, Section 4.2].

Algorithm 1 RCEG

Input: initial points (x_0, y_0) and stepsizes $\eta > 0$.
for $t = 0, 1, 2, \dots, T - 1$ **do**
 Query $(g_x^t, g_y^t) \leftarrow (\text{grad}_x f(x_t, y_t), \text{grad}_y f(x_t, y_t))$, the
 Riemannian gradient of f at a point (x_t, y_t)
 $\hat{x}_t \leftarrow \text{Exp}_{x_t}(-\eta \cdot g_x^t)$.
 $\hat{y}_t \leftarrow \text{Exp}_{y_t}(\eta \cdot g_y^t)$.
 Query $(\hat{g}_x^t, \hat{g}_y^t) \leftarrow (\text{grad}_x f(\hat{x}_t, \hat{y}_t), \text{grad}_y f(\hat{x}_t, \hat{y}_t))$, the
 Riemannian gradient of f at a point $(\hat{x}_t, \hat{y}_t)$
 $x_{t+1} \leftarrow \text{Exp}_{\hat{x}_t}(-\eta \cdot \hat{g}_x^t + \text{Exp}_{\hat{x}_t}^{-1}(x_t))$.
 $y_{t+1} \leftarrow \text{Exp}_{\hat{y}_t}(\eta \cdot \hat{g}_y^t + \text{Exp}_{\hat{y}_t}^{-1}(y_t))$.
end for

Algorithm 2 SRCEG

Input: initial points (x_0, y_0) and stepsizes $\eta > 0$.
for $t = 0, 1, 2, \dots, T - 1$ **do**
 Query (g_x^t, g_y^t) as a **noisy** estimator of Riemannian
 gradient of f at a point (x_t, y_t) .
 $\hat{x}_t \leftarrow \text{Exp}_{x_t}(-\eta \cdot g_x^t)$.
 $\hat{y}_t \leftarrow \text{Exp}_{y_t}(\eta \cdot g_y^t)$.
 Query $(\hat{g}_x^t, \hat{g}_y^t)$ as a **noisy** estimator of Riemannian
 gradient of f at a point $(\hat{x}_t, \hat{y}_t)$.
 $x_{t+1} \leftarrow \text{Exp}_{\hat{x}_t}(-\eta \cdot \hat{g}_x^t + \text{Exp}_{\hat{x}_t}^{-1}(x_t))$.
 $y_{t+1} \leftarrow \text{Exp}_{\hat{y}_t}(\eta \cdot \hat{g}_y^t + \text{Exp}_{\hat{y}_t}^{-1}(y_t))$.
end for

3.1 Algorithmic scheme

The recently proposed *Riemannian corrected extragradient* (RCEG) method [Zhang et al., 2022] is a natural extension of the celebrated extragradient (EG) method to the Riemannian setting. Its scheme resembles that of EG in Euclidean spaces but employs a simple modification in the extrapolation step to accommodate the nonlinear geometry of Riemannian manifolds. Let us provide some intuition how such modifications work.

We start with a basic version of EG as follows, where $\mathcal{M}$ and $\mathcal{N}$ are classically restricted to be convex constraint sets in Euclidean spaces:

$$\begin{aligned} \hat{x}_t &\leftarrow \text{proj}_{\mathcal{M}}(x_t - \eta \cdot \nabla_x f(x_t, y_t)), & \hat{y}_t &\leftarrow \text{proj}_{\mathcal{N}}(y_t + \eta \cdot \nabla_y f(x_t, y_t)), \\ x_{t+1} &\leftarrow \text{proj}_{\mathcal{M}}(x_t - \eta \cdot \nabla_x f(\hat{x}_t, \hat{y}_t)), & y_{t+1} &\leftarrow \text{proj}_{\mathcal{N}}(y_t + \eta \cdot \nabla_y f(\hat{x}_t, \hat{y}_t)). \end{aligned} \quad (3.1)$$

Turning to the setting where $\mathcal{M}$ and $\mathcal{N}$ are Riemannian manifolds, the rather straightforward way to do the generalization is to replace the projection operator by the corresponding exponential map and the gradient by the corresponding Riemannian gradient. For the first line of Eq. (3.1), this approach works and leads to the following updates:

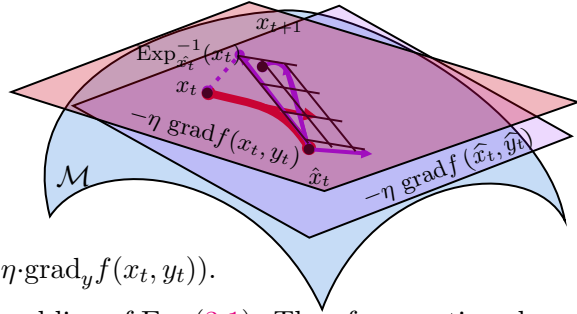
$$\hat{x}_t \leftarrow \text{Exp}_{x_t}(-\eta \cdot \text{grad}_x f(x_t, y_t)), \quad \hat{y}_t \leftarrow \text{Exp}_{y_t}(\eta \cdot \text{grad}_y f(x_t, y_t)).$$

However, we encounter some issues for the second line of Eq. (3.1): The aforementioned approach leads to some problematic updates, $x_{t+1} \leftarrow \text{Exp}_{\hat{x}_t}(-\eta \cdot \text{grad}_x f(\hat{x}_t, \hat{y}_t))$ and $y_{t+1} \leftarrow \text{Exp}_{\hat{y}_t}(\eta \cdot \text{grad}_y f(\hat{x}_t, \hat{y}_t))$; indeed, the exponential maps $\text{Exp}_{x_t}(\cdot)$ and $\text{Exp}_{y_t}(\cdot)$ are defined from $T_{x_t}\mathcal{M}$ to $\mathcal{M}$ and from $T_{y_t}\mathcal{N}$ to $\mathcal{N}$ respectively. However, we have $-\text{grad}_x f(\hat{x}_t, \hat{y}_t) \in T_{\hat{x}_t}\mathcal{M}$ and $\text{grad}_y f(\hat{x}_t, \hat{y}_t) \in T_{\hat{y}_t}\mathcal{N}$. This motivates us to reformulate the second line of Eq. (3.1) as follows:

$$x_{t+1} \leftarrow \text{proj}_{\mathcal{M}}(\hat{x}_t - \eta \cdot \nabla_x f(\hat{x}_t, \hat{y}_t) + (x_t - \hat{x}_t)), \quad y_{t+1} \leftarrow \text{proj}_{\mathcal{N}}(\hat{y}_t + \eta \cdot \nabla_y f(\hat{x}_t, \hat{y}_t) + (y_t - \hat{y}_t)).$$

In the general setting of Riemannian manifolds, the terms $x_t - \hat{x}_t$ and $y_t - \hat{y}_t$ become $\text{Exp}_{\hat{x}_t}^{-1}(x_t) \in T_{\hat{x}_t}\mathcal{M}$ and $\text{Exp}_{\hat{y}_t}^{-1}(y_t) \in T_{\hat{y}_t}\mathcal{N}$. This observation yields the following updates:

$$x_{t+1} \leftarrow \text{Exp}_{\hat{x}_t}(-\eta \cdot \text{grad}_x f(\hat{x}_t, \hat{y}_t) + \text{Exp}_{\hat{x}_t}^{-1}(x_t)), \quad \hat{y}_t \leftarrow \text{Exp}_{\hat{y}_t}(\eta \cdot \text{grad}_y f(\hat{x}_t, \hat{y}_t) + \text{Exp}_{\hat{y}_t}^{-1}(y_t)).$$



We summarize the resulting RCEG method in Algorithm 1 and present the stochastic extension with noisy estimators of Riemannian gradients of f in Algorithm 2.

3.2 Main results

We present our main results on global convergence for Algorithms 1 and 2. To simplify the presentation, we treat separately the following two cases:

Assumption 3.1 *The objective function f is geodesically ℓ -smooth and geodesically strongly-convex-strongly-concave with $\mu > 0$.*

Assumption 3.2 *The objective function f is geodesically ℓ -smooth and geodesically convex-concave.*

Letting $(x^*, y^*) \in \mathcal{M} \times \mathcal{N}$ be a global saddle point of f (which exists under either Assumption 3.1 or 3.2), we let $D_0 = (d_{\mathcal{M}}(x_0, x^*))^2 + (d_{\mathcal{N}}(y_0, y^*))^2 > 0$ and $\kappa = \ell/\mu$ for geodesically strongly-convex-strongly-concave setting. For simplicity of presentation, we also define a ratio $\tau(\cdot, \cdot)$ that measures how non-flatness changes in the spaces: $\tau([\kappa_{\min}, \kappa_{\max}], c) = \frac{\bar{\xi}(\kappa_{\min}, c)}{\underline{\xi}(\kappa_{\max}, c)} \geq 1$. We summarize our results for Algorithm 1 in the following theorem.

Theorem 3.1 *Given Assumptions 2.1 and 3.1, and letting $\eta = \min\{1/(2\ell\sqrt{\tau_0}), \underline{\xi}_0/(2\mu)\}$, there exists some $T > 0$ such that the output of Algorithm 1 satisfies that $(d(x_T, x^*))^2 + (d(y_T, y^*))^2 \leq \epsilon$ (i.e., an ϵ -saddle point of f in Definition 2.4) and the total number of Riemannian gradient evaluations is bounded by*

$$O\left(\left(\kappa\sqrt{\tau_0} + \frac{1}{\underline{\xi}_0}\right) \log\left(\frac{D_0}{\epsilon}\right)\right),$$

where $\tau_0 = \tau([\kappa_{\min}, \kappa_{\max}], D) \geq 1$ measures how non-flatness changes in $\mathcal{M}$ and $\mathcal{N}$ and $\underline{\xi}_0 = \underline{\xi}(\kappa_{\max}, D) \leq 1$ is properly defined in Proposition 2.1.

Remark 3.1 *Theorem 3.1 illustrates the last-iterate convergence of Algorithm 1 for solving geodesically strongly-convex-strongly-concave problems, thereby resolving an open problem delineated by Zhang et al. [2022]. Further, the dependence on κ and $1/\epsilon$ cannot be improved since it matches the lower bound established for min-max optimization problems in Euclidean spaces [Zhang et al., 2021]. However, we believe that the dependence on τ_0 and $\underline{\xi}_0$ is not tight, and it is of interest to either improve the rate or establish a lower bound for general Riemannian min-max optimization.*

Remark 3.2 *The current theoretical analysis covers local geodesic strong-convex-strong-concave settings. The key ingredient is how to define the local region; indeed, if we say the set of $\{(x, y) : d_{\mathcal{M}}(x, x^*) \leq \delta, d_{\mathcal{N}}(y_t, y^*) \leq \delta\}$ is a local region where the function is geodesic strong-convex-strong-concave. Then, the set of $\{(x, y) : (d_{\mathcal{M}}(x, x^*)^2 + d_{\mathcal{N}}(y_t, y^*)^2) \leq \delta^2\}$ must be contained in the above local region and the objective function is also geodesic strong-convex-strong-concave. If $(x_0, y_0) \in \{(x, y) : (d_{\mathcal{M}}(x, x^*)^2 + d_{\mathcal{N}}(y_t, y^*)^2) \leq \delta^2\}$, our theoretical analysis guarantees the last-iterate linear convergence rate. Such argument and definition of local region were standard for min-max optimization in the Euclidean setting; see Liang and Stokes [2019, Assumption 2.1]. For an important optimization problem that is globally geodesically strongly-convex-strongly-concave, we refer to Appendix B where Robust matrix Karcher mean problem is indeed the desired one.*

In the scheme of SRECG, we highlight that (g_x^t, g_y^t) and $(\hat{g}_x^t, \hat{g}_y^t)$ are noisy estimators of Riemannian gradients of f at (x_t, y_t) and $(\hat{x}_t, \hat{y}_t)$. It is necessary to impose the conditions such that these estimators are unbiased and has bounded variance. By abuse of notation, we assume that

$$\begin{aligned} g_x^t &= \text{grad}_x f(x_t, y_t) + \xi_x^t, & g_y^t &= \text{grad}_y f(x_t, y_t) + \xi_y^t, \\ \hat{g}_x^t &= \text{grad}_x f(\hat{x}_t, \hat{y}_t) + \hat{\xi}_x^t, & \hat{g}_y^t &= \text{grad}_y f(\hat{x}_t, \hat{y}_t) + \hat{\xi}_y^t. \end{aligned} \quad (3.2)$$

where the noises (ξ_x^t, ξ_y^t) and $(\hat{\xi}_x^t, \hat{\xi}_y^t)$ are independent and satisfy that

$$\begin{aligned} \mathbb{E}[\xi_x^t] &= 0, & \mathbb{E}[\xi_y^t] &= 0, & \mathbb{E}[\|\xi_x^t\|^2 + \|\xi_y^t\|^2] &\leq \sigma^2, \\ \mathbb{E}[\hat{\xi}_x^t] &= 0, & \mathbb{E}[\hat{\xi}_y^t] &= 0, & \mathbb{E}[\|\hat{\xi}_x^t\|^2 + \|\hat{\xi}_y^t\|^2] &\leq \sigma^2. \end{aligned} \quad (3.3)$$

We are ready to summarize our results for Algorithm 2 in the following theorems.

Theorem 3.2 *Given Assumptions 2.1 and 3.1, letting Eq. (3.2) and Eq. (3.3) hold with $\sigma > 0$ and letting $\eta > 0$ satisfy $\eta = \min\{\frac{1}{24\ell\sqrt{\tau_0}}, \frac{\xi_0}{2\mu}, \frac{2(\log(T) + \log(\mu^2 D_0 \sigma^{-2}))}{\mu T}\}$, there exists some $T > 0$ such that the output of Algorithm 2 satisfies that $\mathbb{E}[(d(x_T, x^*))^2 + (d(y_T, y^*))^2] \leq \epsilon$ and the total number of noisy Riemannian gradient evaluations is bounded by*

$$O\left(\left(\kappa\sqrt{\tau_0} + \frac{1}{\xi_0}\right) \log\left(\frac{D_0}{\epsilon}\right) + \frac{\sigma^2 \bar{\xi}_0}{\mu^2 \epsilon} \log\left(\frac{1}{\epsilon}\right)\right),$$

where $\tau_0 = \tau([\kappa_{\min}, \kappa_{\max}], D) \geq 1$ measures how non-flatness changes in $\mathcal{M}$ and $\mathcal{N}$ and $\xi_0 = \xi(\kappa_{\max}, D) \leq 1$ is properly defined in Proposition 2.1.

Theorem 3.3 *Given Assumptions 2.1 and 3.2 and assume that Eq. (3.2) and Eq. (3.3) hold with $\sigma > 0$ and let $\eta > 0$ satisfies that $\eta = \min\{\frac{1}{4\ell\sqrt{\tau_0}}, \frac{1}{\sigma}\sqrt{\frac{D_0}{\xi_0 T}}\}$, there exists some $T > 0$ such that the output of Algorithm 2 satisfies that $\mathbb{E}[f(\bar{x}_T, y^*) - f(x^*, \bar{y}_T)] \leq \epsilon$ and the total number of noisy Riemannian gradient evaluations is bounded by*

$$O\left(\frac{\ell D_0 \sqrt{\tau_0}}{\epsilon} + \frac{\sigma^2 \bar{\xi}_0}{\epsilon^2}\right),$$

where $\tau_0 = \tau([\kappa_{\min}, \kappa_{\max}], D)$ measures how non-flatness changes in $\mathcal{M}$ and $\mathcal{N}$ and $\bar{\xi}_0 = \bar{\xi}(\kappa_{\min}, D) \geq 1$ is properly defined in Proposition 2.1. The time-average iterates $(\bar{x}_T, \bar{y}_T) \in \mathcal{M} \times \mathcal{N}$ can be computed by $(\bar{x}_0, \bar{y}_0) = (0, 0)$ and the inductive formula: $\bar{x}_{t+1} = \text{Exp}_{\bar{x}_t}(\frac{1}{t+1} \cdot \text{Exp}_{\bar{x}_t}^{-1}(\hat{x}_t))$ and $\bar{y}_{t+1} = \text{Exp}_{\bar{y}_t}(\frac{1}{t+1} \cdot \text{Exp}_{\bar{y}_t}^{-1}(\hat{y}_t))$ for all $t = 0, 1, \dots, T-1$.

Remark 3.3 *Theorem 3.2 presents the last-iterate convergence rate of Algorithm 2 for solving geodesically strongly-convex-strongly-concave problems while Theorem 3.3 gives the time-average convergence rate when the function f is only assumed to be geodesically convex-concave. Note that we carefully choose the stepsizes such that our upper bounds match the lower bounds established for stochastic min-max optimization problems in Euclidean spaces [Juditsky et al., 2011, Fallah et al., 2020, Kotsalis et al., 2022], in terms of the dependence on κ , $1/\epsilon$ and σ^2 , up to log factors.*

Discussions: The last-iterate linear convergence rate in terms of Riemannian metrics is only limited to geodesically strongly convex-concave cases but other results, e.g., the average-iterate sublinear convergence rate, are derived under more mild conditions. This is consistent with classical results in the Euclidean setting where geodesic convexity reduces to convexity; indeed, the last-iterate linear convergence rate in terms of squared Euclidean norm is only obtained for strongly convex-concave cases. As such, our setting is not restrictive. Moreover, [Zhang et al. \[2022\]](#) showed that the existence of a global saddle point is only guaranteed under the geodesically convex-concave assumption. For geodesically nonconvex-concave or geodesically nonconvex-nonconcave cases, a global saddle point might not exist and new optimality notions are required before algorithmic design. This question remains open in the Euclidean setting and is beyond the scope of this paper. However, we remark that an interesting class of robustification problems are nonconvex-nonconcave min-max problems in the Euclidean setting can be geodesically convex-concave in the Riemannian setting; see [Appendix B](#).

4 Experiments

We present numerical experiments on the task of robust principal component analysis (RPCA) for symmetric positive definite (SPD) matrices. In particular, we compare the performance of [Algorithm 1](#) and [2](#) with different outputs, i.e., the last iterate (x_T, y_T) versus the time-average iterate $(\bar{x}_T, \bar{y}_T)$ (see the precise definition in [Theorem 3.3](#)). Note that our implementations of both algorithms are based on the MANOPT package [\[Boumal et al., 2014\]](#). All the experiments were implemented in MATLAB R2021b on a workstation with a 2.6 GHz Intel Core i7 and 16GB of memory. Due to space limitations, some additional experimental results are deferred to [Appendix G](#).

Experimental setup. The problem of RPCA [\[Candès et al., 2011, Harandi et al., 2017\]](#) can be formulated as the Riemannian min-max optimization problem with an SPD manifold and a sphere manifold. Formally, we have

$$\max_{M \in \mathcal{M}_{\text{PSD}}^d} \min_{x \in \mathcal{S}^d} \left\{ -x^\top M x - \frac{\alpha}{n} \sum_{i=1}^n d(M, M_i) \right\}. \quad (4.1)$$

In this formulation, $\alpha > 0$ denotes the penalty parameter, $\{M_i\}_{i \in [n]}$ is a sequence of given data SPD matrices, $\mathcal{M}_{\text{PSD}}^d = \{M \in \mathbb{R}^{d \times d} : M \succ 0, M = M^\top\}$ denotes the SPD manifold, $\mathcal{S}^d = \{x \in \mathbb{R}^d : \|x\| = 1\}$ denotes the sphere manifold and $d(\cdot, \cdot) : \mathcal{M}_{\text{PSD}}^d \times \mathcal{M}_{\text{PSD}}^d \mapsto \mathbb{R}$ is the Riemannian distance induced by the exponential map on the SPD manifold $\mathcal{M}_{\text{PSD}}^d$. As demonstrated by [Zhang et al. \[2022\]](#), the problem of RPCA is nonconvex-nonconcave from a Euclidean perspective but is *locally geodesically strongly-convex-strongly-concave* and satisfies most of the assumptions that we make in this paper. In particular, the SPD manifold is complete with sectional curvature in $[-\frac{1}{2}, 1]$ [\[Criscitiello and Boumal, 2022\]](#) and the sphere manifold is complete with sectional curvature of 1. Other reasons why we use such example are: (i) it is a classical one in ML; (ii) [Zhang et al. \[2022\]](#) also uses this example and observes the linear convergence behavior; (iii) the numerical results show that the unicity of geodesics assumption may not be necessary in practice; and (iv) this is an application where both min and max sides are done on Riemannian manifolds.

Following the previous works of [Zhang et al. \[2022\]](#) and [Han et al. \[2022\]](#), we generate a sequence of data matrices M_i satisfying that their eigenvalues are in the range of $[0.2, 4.5]$. In our experiment, we fix $\alpha = 1.0$ and also vary the problem dimension $d \in \{25, 50, 100\}$. The evaluation metric is set as

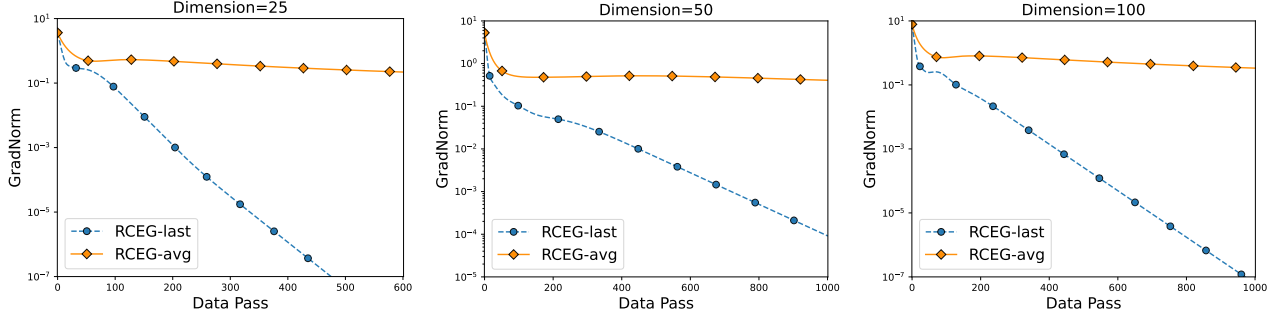


Figure 1: Comparison of last iterate (RCEG-last) and time-average iterate (RCEG-avg) for solving the RPCA problem in Eq. (4.1) with different problem dimensions $d \in \{25, 50, 100\}$. The horizontal axis represents the number of data passes and the vertical axis represents gradient norm.

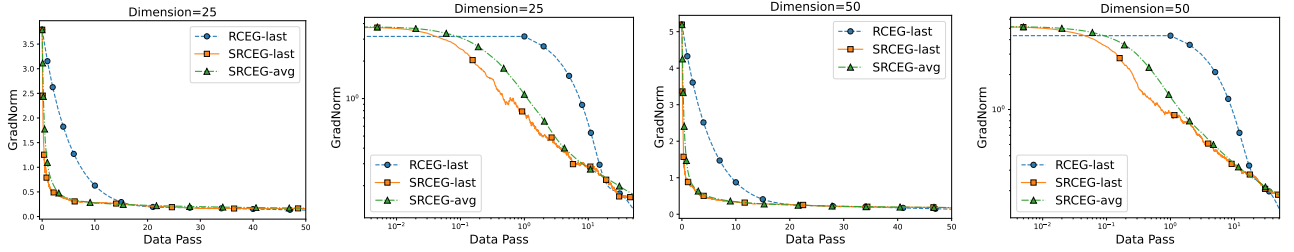


Figure 2: Comparison of RCEG and SRCEG for solving the RPCA problem in Eq. (4.1) with different problem dimensions $d \in \{25, 50\}$. The horizontal axis is the number of data passes and the vertical axis is gradient norm. We set $n = 40$ and $n = 200$ in Figure 1 and 2. For RCEG, we set $\eta = \frac{1}{2\ell}$ where $\ell > 0$ is selected via grid search. For SRCEG, we set $\eta_t = \min\{\frac{1}{2\ell}, \frac{a}{t}\}$ where $\ell, a > 0$ are selected via grid search. Additional results on the effect of stepsize are summarized in Appendix G.

Experimental results. Figure 1 summarizes the effects of different outputs for RCEG; indeed, RCEG-last and RCEG-avg refer to Algorithm 1 with last iterate and time-average iterate respectively. It is clear that the last iterate of RCEG consistently exhibits linear convergence to an optimal solution in all the settings, verifying our theoretical results in Theorem 3.1. In contrast, the average iterate of RCEG converges much slower than the last iterate of RCEG. The possible reason is that the problem of RPCA is *only* locally geodesically strongly-convex-strongly-concave and averaging with the iterates generated during early stage will significantly slow down the convergence of RCEG.

Figure 2 presents the comparison between SRCEG (with either last iterate or time-average iterate) and RCEG with last-iterate; here, SRCEG-last and SRCEG-avg refer to Algorithm 2 with last iterate and time-average iterate respectively. We observe that SRCEG with either last iterate or average iterate converge faster than RCEG at the early stage and all of them finally converge to an optimal solution. This demonstrates the effectiveness and efficiency of SRCEG in practice. It is also worth mentioning that the difference between last-iterate convergence and time-average-iterate convergence is not as significant as in the deterministic setting. This is possibly because the technique of averaging help cancels the negative effect of imperfect information [Kingma and Ba, 2015, Yazıcı et al., 2019].

5 Conclusions

Inspired broadly by the structure of the complex competition that arises in many applications of robust optimization in ML, we focus on the problem of min-max optimization in the pure Riemannian setting (where both min and max player are constrained in a smooth manifold). Answering the open question of Zhang et al. [2022] for the geodesically (strongly) convex-concave case, we showed that the Riemannian correction technique for EG matches the linear last-iterate complexity of their Euclidean counterparts in terms of accuracy and conditional number of objective for both deterministic and stochastic case. Additionally, we provide near-optimal guarantees for both smooth and non-smooth min-max optimization via Riemannian EG and GDA for the simple convex-concave case.

As a consequence of this work numerous open problems emerge; one immediate open question for future work is to explore whether the dependence on the curvature constant is also tight. Additionally, another generalization of interest would be to consider the performance of RCEG in the case of Riemannian Monotone Variational inequalities (RMVI) and examine the generalization of Zhang et al. [2022] existence proof. Finally, there has been recent work in proving last-iterate convergence in the convex-concave setting via Sum-Of-Squares techniques [Cai et al., 2022]. It would be interesting to examine how one could leverage this machinery in a non-Euclidean but geodesic-metric-friendly framework.

Acknowledgments

This work was supported in part by the Mathematical Data Science program of the Office of Naval Research under grant number N00014-18-1-2764 and by the Vannevar Bush Faculty Fellowship program under grant number N00014-21-1-2941. The work of Michael I. Jordan is also partially supported by NSF Grant IIS-1901252. Emmanouil V. Vlatakis-Gkaragkounis is grateful for financial support by the Google-Simons Fellowship, Pancretan Association of America and Simons Collaboration on Algorithms and Geometry. This project was completed while he was a visiting research fellow at the Simons Institute for the Theory of Computing. Additionally, he would like to acknowledge the following series of NSF-CCF grants under the numbers 1763970/2107187/1563155/1814873.

References

- J. Abernethy, K. A. Lai, and A. Wibisono. Last-iterate convergence rates for min-max optimization: Convergence of Hamiltonian gradient descent and consensus optimization. In *ALT*, pages 3–47. PMLR, 2021. (Cited on page 2.)
- P-A. Absil and S. Hosseini. A collection of nonsmooth Riemannian optimization problems. In *Nonsmooth Optimization and Its Applications*, pages 1–15. Springer, 2019. (Cited on page 23.)
- P-A. Absil, R. Mahony, and R. Sepulchre. *Optimization Algorithms on Matrix Manifolds*. Princeton University Press, 2009. (Cited on pages 22 and 24.)
- L. Adolphs, H. Daneshmand, A. Lucchi, and T. Hofmann. Local saddle point optimization: A curvature exploitation approach. In *AISTATS*, pages 486–495. PMLR, 2019. (Cited on page 2.)
- F. Alimisis, A. Orvieto, G. Bécigneul, and A. Lucchi. A continuous-time perspective for modeling

- acceleration in Riemannian optimization. In *AISTATS*, pages 1297–1307. PMLR, 2020. (Cited on pages 5, 6, 7, and 27.)
- A. Anandkumar and R. Ge. Efficient approaches for escaping higher order saddle points in non-convex optimization. In *COLT*, pages 81–102. PMLR, 2016. (Cited on page 1.)
- K. Antonakopoulos, E. V. Belmega, and P. Mertikopoulos. Online and stochastic optimization beyond Lipschitz continuity: A Riemannian approach. In *ICLR*, 2020. URL <https://openreview.net/forum?id=rkxZyaNtwB>. (Cited on pages 1 and 2.)
- W. Azizian, I. Mitliagkas, S. Lacoste-Julien, and G. Gidel. A tight and unified analysis of gradient-based methods for a whole spectrum of differentiable games. In *AISTATS*, pages 2863–2873. PMLR, 2020. (Cited on page 2.)
- M. Bacak. *Convex Analysis and Optimization in Hadamard Spaces*, volume 22. Walter de Gruyter GmbH & Co KG, 2014. (Cited on page 26.)
- D. Balduzzi, S. Racaniere, J. Martens, J. Foerster, K. Tuyls, and T. Graepel. The mechanics of N-player differentiable games. In *ICML*, pages 354–363. PMLR, 2018. (Cited on page 2.)
- H. H. Bauschke and P. L. Combettes. *Convex Analysis and Monotone Operator Theory in Hilbert Spaces*, volume 408. Springer, 2011. (Cited on page 23.)
- G. Becigneul and O-E. Ganea. Riemannian adaptive optimization methods. In *ICLR*, 2019. URL <https://openreview.net/forum?id=r1eiqi09K7>. (Cited on page 23.)
- A. Ben-Tal, L. EL Ghaoui, and A. Nemirovski. *Robust Optimization*, volume 28. Princeton University Press, 2009. (Cited on page 24.)
- G. C. Bento, O. P. Ferreira, and J. G. Melo. Iteration-complexity of gradient, subgradient and proximal point methods on Riemannian manifolds. *Journal of Optimization Theory and Applications*, 173(2): 548–562, 2017. (Cited on pages 23 and 24.)
- R. Bergmann and R. Herzog. Intrinsic formulation of KKT conditions and constraint qualifications on smooth manifolds. *SIAM Journal on Optimization*, 29(4):2423–2444, 2019. (Cited on page 24.)
- I. M. Bomze, P. Mertikopoulos, W. Schachinger, and M. Staudigl. Hessian barrier algorithms for linearly constrained optimization problems. *SIAM Journal on Optimization*, 29(3):2100–2127, 2019. (Cited on page 2.)
- S. Bonnabel. Stochastic gradient descent on Riemannian manifolds. *IEEE Transactions on Automatic Control*, 58(9):2217–2229, 2013. (Cited on pages 2 and 23.)
- N. Boumal and P-A. Absil. RTRMC: A Riemannian trust-region method for low-rank matrix completion. In *NIPS*, pages 406–414, 2011. (Cited on pages 1 and 25.)
- N. Boumal, B. Mishra, P-A. Absil, and R. Sepulchre. Manopt, a Matlab toolbox for optimization on manifolds. *Journal of Machine Learning Research*, 15(1):1455–1459, 2014. (Cited on page 11.)
- N. Boumal, P-A. Absil, and C. Cartis. Global rates of convergence for nonconvex optimization on manifolds. *IMA Journal of Numerical Analysis*, 39(1):1–33, 2019. (Cited on page 22.)

- L. E. J. Brouwer. Über abbildung von mannigfaltigkeiten. *Mathematische Annalen*, 71(1):97–115, 1911. (Cited on page 2.)
- D. Burago, I. D. Burago, Y. Burago, S. Ivanov, S. V. Ivanov, and S. A. Ivanov. *A Course in Metric Geometry*, volume 33. American Mathematical Soc., 2001. (Cited on pages 22 and 26.)
- Y. Burago, M. Gromov, and G. Perel'man. A. D. Alexandrov spaces with curvature bounded below. *Russian Mathematical Surveys*, 47(2):1, 1992. (Cited on page 5.)
- Y. Cai, A. Oikonomou, and W. Zheng. Tight last-iterate convergence of the extragradient method for constrained monotone variational inequalities. *ArXiv Preprint: 2204.09228*, 2022. (Cited on pages 2 and 13.)
- E. J. Candes, M. B. Wakin, and S. P. Boyd. Enhancing sparsity by reweighted ℓ_1 minimization. *Journal of Fourier Analysis and Applications*, 14(5):877–905, 2008. (Cited on page 1.)
- E. J. Candès, X. Li, Y. Ma, and J. Wright. Robust principal component analysis? *Journal of the ACM (JACM)*, 58(3):1–37, 2011. (Cited on page 11.)
- J. Cheeger and D. G. Ebin. *Comparison Theorems in Riemannian Geometry*, volume 9. North-Holland Amsterdam, 1975. (Cited on page 5.)
- N. Chen, A. Klushyn, R. Kurlle, X. Jiang, J. Bayer, and P. Smagt. Metrics for deep generative models. In *AISTATS*, pages 1540–1550. PMLR, 2018. (Cited on page 2.)
- S. Chen, S. Ma, A. M-C. So, and T. Zhang. Proximal gradient method for nonsmooth optimization over the Stiefel manifold. *SIAM Journal on Optimization*, 30(1):210–239, 2020. (Cited on page 23.)
- K. L. Chung. On a stochastic approximation method. *The Annals of Mathematical Statistics*, pages 463–483, 1954. (Cited on page 38.)
- C. Criscitiello and N. Boumal. Efficiently escaping saddle points on manifolds. In *NeurIPS*, pages 5987–5997, 2019. (Cited on pages 1 and 23.)
- C. Criscitiello and N. Boumal. An accelerated first-order method for non-convex optimization on manifolds. *Foundations of Computational Mathematics*, pages 1–77, 2022. (Cited on page 11.)
- C. Daskalakis and I. Panageas. The limit points of (optimistic) gradient descent in min-max optimization. In *NIPS*, pages 9256–9266, 2018. (Cited on page 2.)
- C. Daskalakis and I. Panageas. Last-iterate convergence: Zero-sum games and constrained min-max optimization. In *ITCS*, 2019. (Cited on page 2.)
- J. Diakonikolas. Halpern iteration for near-optimal and parameter-free monotone inclusion and strong solutions to variational inequalities. In *COLT*, pages 1428–1451. PMLR, 2020. (Cited on page 2.)
- S. S. Du, C. Jin, J. D. Lee, M. I. Jordan, B. Póczos, and A. Singh. Gradient descent can take exponential time to escape saddle points. In *NIPS*, pages 1067–1077, 2017. (Cited on page 1.)
- F. Facchinei and J-S. Pang. *Finite-Dimensional Variational Inequalities and Complementarity Problems*. Springer Science & Business Media, 2007. (Cited on page 23.)

- A. Fallah, A. Ozdaglar, and S. Pattathil. An optimal multistage stochastic gradient method for minimax problems. In *CDC*, pages 3573–3579. IEEE, 2020. (Cited on page 10.)
- J. Fearnley, P. W. Goldberg, A. Hollender, and R. Savani. The complexity of gradient descent: $\text{CLS} = \text{PPAD} \cap \text{PLS}$. In *STOC*, pages 46–59, 2021. (Cited on page 2.)
- O. P. Ferreira and P. R. Oliveira. Subgradient algorithm on Riemannian manifolds. *Journal of Optimization Theory and Applications*, 97(1):93–104, 1998. (Cited on page 23.)
- O. P. Ferreira and P. R. Oliveira. Proximal point algorithm on Riemannian manifolds. *Optimization*, 51(2):257–270, 2002. (Cited on page 23.)
- O. P. Ferreira, L. R. Pérez, and S. Z. Németh. Singularities of monotone vector fields and an extragradient-type algorithm. *Journal of Global Optimization*, 31(1):133–151, 2005. (Cited on page 24.)
- P. T. Fletcher and S. Joshi. Riemannian geometry for the statistical analysis of diffusion tensor data. *Signal Processing*, 87(2):250–262, 2007. (Cited on page 24.)
- M. Fornasier, H. Rauhut, and R. Ward. Low-rank matrix recovery via iteratively reweighted least squares minimization. *SIAM Journal on Optimization*, 21(4):1614–1640, 2011. (Cited on page 1.)
- B. Gao, X. Liu, X. Chen, and Y. Yuan. A new first-order algorithmic framework for optimization problems with orthogonality constraints. *SIAM Journal on Optimization*, 28(1):302–332, 2018. (Cited on page 22.)
- R. Ge, F. Huang, C. Jin, and Y. Yuan. Escaping from saddle points—online stochastic gradient for tensor decomposition. In *COLT*, pages 797–842. PMLR, 2015. (Cited on page 1.)
- R. Ge, C. Jin, and Y. Zheng. No spurious local minima in nonconvex low rank problems: A unified geometric analysis. In *ICML*, pages 1233–1242. PMLR, 2017. (Cited on page 1.)
- A. Giannou, E. V. Vlatakis-Gkaragkounis, and P. Mertikopoulos. On the rate of convergence of regularized learning in games: From bandits and uncertainty to optimism and beyond. In *NeurIPS*, pages 22655–22666, 2021a. (Cited on page 1.)
- A. Giannou, E. V. Vlatakis-Gkaragkounis, and P. Mertikopoulos. Survival of the strictest: Stable and unstable equilibria under regularized learning with partial information. In *COLT*, pages 2147–2148. PMLR, 2021b. (Cited on page 1.)
- G. Gidel, R. A. Hemmat, M. Pezeshki, R. Le Priol, G. Huang, S. Lacoste-Julien, and I. Mitliagkas. Negative momentum for improved game dynamics. In *AISTATS*, pages 1802–1811. PMLR, 2019. (Cited on page 2.)
- N. Golowich, S. Pattathil, C. Daskalakis, and A. Ozdaglar. Last iterate is slower than averaged iterate in smooth convex-concave saddle point problems. In *COLT*, pages 1758–1784. PMLR, 2020. (Cited on page 2.)
- I. Goodfellow, J. Pouget-Abadie, M. Mirza, B. Xu, D. Warde-Farley, S. Ozair, A. Courville, and Y. Bengio. Generative adversarial networks. In *NIPS*, pages 2672–2680, 2014. (Cited on page 1.)

- E. Y. Hamedani and N. S. Aybat. A primal-dual algorithm with line search for general convex-concave saddle point problems. *SIAM Journal on Optimization*, 31(2):1299–1329, 2021. (Cited on page 2.)
- A. Han, B. Mishra, P. K. Jawanpuria, and J. Gao. On Riemannian optimization over positive definite matrices with the Bures-Wasserstein geometry. In *NeurIPS*, pages 8940–8953, 2021. (Cited on page 1.)
- A. Han, B. Mishra, P. Jawanpuria, P. Kumar, and J. Gao. Riemannian Hamiltonian methods for min-max optimization on manifolds. *ArXiv Preprint: 2204.11418*, 2022. (Cited on pages 3, 11, and 24.)
- M. Harandi, M. Salzmann, and R. Hartley. Dimensionality reduction on SPD manifolds: The emergence of geometry-aware methods. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, 40(1):48–62, 2017. (Cited on page 11.)
- E. D. Helly. Über mengen konvexer körper mit gemeinschaftlichen punkte. *Jahresbericht der Deutschen Mathematiker-Vereinigung*, 32:175–176, 1923. (Cited on page 2.)
- M. Heusel, H. Ramsauer, T. Unterthiner, B. Nessler, and S. Hochreiter. GANs trained by a two time-scale update rule converge to a local nash equilibrium. In *NIPS*, pages 6629–6640, 2017. (Cited on page 2.)
- R. Hosseini and S. Sra. Matrix manifold optimization for Gaussian mixtures. In *NIPS*, pages 910–918, 2015. (Cited on page 1.)
- J. Hu, A. Milzarek, Z. Wen, and Y. Yuan. Adaptive quadratically regularized Newton method for Riemannian optimization. *SIAM Journal on Matrix Analysis and Applications*, 39(3):1181–1207, 2018. (Cited on page 22.)
- J. Hu, B. Jiang, L. Lin, Z. Wen, and Y. Yuan. Structured quasi-Newton methods for optimization with orthogonality constraints. *SIAM Journal on Scientific Computing*, 41(4):A2239–A2269, 2019. (Cited on page 22.)
- J. Hu, X. Liu, Z-W. Wen, and Y-X. Yuan. A brief introduction to manifold optimization. *Journal of the Operations Research Society of China*, 8(2):199–248, 2020. (Cited on page 24.)
- F. Huang, S. Gao, and H. Huang. Gradient descent ascent for min-max problems on Riemannian manifolds. *ArXiv Preprint: 2010.06097*, 2020. (Cited on pages 3 and 24.)
- L. Huang, X. Liu, B. Lang, A. Yu, Y. Wang, and B. Li. Orthogonal weight normalization: solution to optimization over multiple dependent stiefel manifolds in deep neural networks. In *AAAI*, pages 3271–3278, 2018. (Cited on pages 1 and 25.)
- S. Ivanov. On Helly’s theorem in geodesic spaces. *Electronic Research Announcements*, 21:109, 2014. (Cited on page 2.)
- P. Jawanpuria and B. Mishra. A unified framework for structured low-rank matrix learning. In *ICML*, pages 2254–2263. PMLR, 2018. (Cited on page 25.)
- C. Jin, R. Ge, P. Netrapalli, S. M. Kakade, and M. I. Jordan. How to escape saddle points efficiently. In *ICML*, pages 1724–1732. PMLR, 2017. (Cited on page 1.)

- C. Jin, P. Netrapalli, and M. I. Jordan. What is local optimality in nonconvex-nonconcave minimax optimization? In *ICML*, pages 4880–4889. PMLR, 2020. (Cited on page 2.)
- C. Jin, P. Netrapalli, R. Ge, S. M. Kakade, and M. I. Jordan. On nonconvex optimization for machine learning: Gradients, stochasticity, and saddle points. *Journal of the ACM (JACM)*, 68(2):1–29, 2021. (Cited on page 1.)
- A. Juditsky, A. Nemirovski, and C. Tauvel. Solving variational inequalities with stochastic mirror-prox algorithm. *Stochastic Systems*, 1(1):17–58, 2011. (Cited on page 10.)
- S. Kakutani. A generalization of Brouwer’s fixed point theorem. *Duke Mathematical Journal*, 8(3): 457–459, 1941. (Cited on page 2.)
- H. Kasai and B. Mishra. Inexact trust-region algorithms on Riemannian manifolds. In *NeurIPS*, pages 4249–4260, 2018. (Cited on page 22.)
- H. Kasai, P. Javanpuria, and B. Mishra. Riemannian adaptive stochastic gradient algorithms on matrix manifolds. In *ICML*, pages 3262–3271, 2019. (Cited on page 23.)
- D. P. Kingma and J. Ba. Adam: A method for stochastic optimization. In *ICLR*, 2015. URL <https://openreview.net/forum?id=8gmWwjFyLj>. (Cited on page 12.)
- B. Knaster, C. Kuratowski, and S. Mazurkiewicz. Ein beweis des fixpunktsatzes für n-dimensionale simplexe. *Fundamenta Mathematicae*, 14(1):132–137, 1929. (Cited on page 2.)
- H. Komiya. Elementary proof for Sion’s minimax theorem. *Kodai Mathematical Journal*, 11(1):5–7, 1988. (Cited on pages 2 and 24.)
- W. Kong and R. D. C. Monteiro. An accelerated inexact proximal point method for solving nonconvex-concave min-max problems. *SIAM Journal on Optimization*, 31(4):2558–2585, 2021. (Cited on page 2.)
- G. M. Korpelevich. The extragradient method for finding saddle points and other problems. *Matecon*, 12:747–756, 1976. (Cited on page 23.)
- G. Kotsalis, G. Lan, and T. Li. Simple and optimal methods for stochastic variational inequalities, I: operator extrapolation. *SIAM Journal on Optimization*, 32(3):2041–2073, 2022. (Cited on pages 4 and 10.)
- A. Kristály. Nash-type equilibria on Riemannian manifolds: A variational approach. *Journal de Mathématiques Pures et Appliquées*, 101(5):660–688, 2014. (Cited on pages 2 and 24.)
- A. Kumar, P. Sattigeri, and P. T. Fletcher. Semi-supervised learning with GANs: Manifold invariance with improved inference. In *NIPS*, pages 5540–5550, 2017. (Cited on page 2.)
- J. Lee. *Introduction to Smooth Manifolds*, volume 218. Springer Science & Business Media, 2012. (Cited on page 22.)
- J. Lee, G. Kim, M. Olfat, M. Hasegawa-Johnson, and C. D. Yoo. Fast and efficient MMD-based fair PCA via optimization over Stiefel manifold. In *AAAI*, pages 7363–7371, 2022. (Cited on page 2.)

- C. Li, G. López, and V. Martín-Márquez. Monotone vector fields and the proximal point algorithm on Hadamard manifolds. *Journal of the London Mathematical Society*, 79(3):663–683, 2009. (Cited on page 24.)
- X. Li, S. Chen, Z. Deng, Q. Qu, Z. Zhu, and A. Man-Cho So. Weakly convex optimization over Stiefel manifold using Riemannian subgradient-type methods. *SIAM Journal on Optimization*, 31(3):1605–1634, 2021. (Cited on page 23.)
- T. Liang and J. Stokes. Interaction matters: A note on non-asymptotic local convergence of generative adversarial networks. In *AISTATS*, pages 907–915. PMLR, 2019. (Cited on pages 2 and 9.)
- T. Lin, C. Fan, N. Ho, M. Cuturi, and M. I. Jordan. Projection robust Wasserstein distance and Riemannian optimization. In *NeurIPS*, pages 9383–9397, 2020a. (Cited on pages 2 and 24.)
- T. Lin, C. Jin, and M. I. Jordan. On gradient descent ascent for nonconvex-concave minimax problems. In *ICML*, pages 6083–6093. PMLR, 2020b. (Cited on page 2.)
- T. Lin, C. Jin, and M. I. Jordan. Near-optimal algorithms for minimax optimization. In *COLT*, pages 2738–2779. PMLR, 2020c. (Cited on page 2.)
- T. Lin, Z. Zheng, E. Chen, M. Cuturi, and M. I. Jordan. On projection robust optimal transport: Sample complexity and model misspecification. In *AISTATS*, pages 262–270. PMLR, 2021. (Cited on page 25.)
- H. Liu, A. M-C. So, and W. Wu. Quadratic optimization with orthogonality constraint: explicit Łojasiewicz exponent and linear convergence of retraction-based line-search and stochastic variance-reduced gradient methods. *Mathematical Programming*, 178(1-2):215–262, 2019. (Cited on page 22.)
- M. Liu, Y. Mroueh, J. Ross, W. Zhang, X. Cui, P. Das, and T. Yang. Towards better understanding of adaptive gradient algorithms in generative adversarial nets. In *ICLR*, 2020. URL <https://openreview.net/forum?id=SJxImOVtwH>. (Cited on page 2.)
- M. Liu, H. Rafique, Q. Lin, and T. Yang. First-order convergence theory for weakly-convex-weakly-concave min-max problems. *Journal of Machine Learning Research*, 22(169):1–34, 2021. (Cited on page 2.)
- S. Łojasiewicz. Une propriété topologique des sous-ensembles analytiques réels. *Les équations aux dérivées partielles*, 117:87–89, 1963. (Cited on page 1.)
- S. Lu, I. Tsaknakis, M. Hong, and Y. Chen. Hybrid block successive approximation for one-sided non-convex min-max problems: Algorithms and applications. *IEEE Transactions on Signal Processing*, 68:3676–3691, 2020. (Cited on page 2.)
- B. Martinet. Régularisation d’inéquations variationnelles par approximations successives. *rev. française informat. Recherche Opérationnelle*, 4:154–158, 1970. (Cited on page 23.)
- E. Mazumdar, L. J. Ratliff, and S. S. Sastry. On gradient-based learning in continuous games. *SIAM Journal on Mathematics of Data Science*, 2(1):103–131, 2020. (Cited on page 2.)
- J. Mei, Y. Gao, B. Dai, C. Szepesvari, and D. Schuurmans. Leveraging non-uniformity in first-order non-convex optimization. In *ICML*, pages 7555–7564. PMLR, 2021. (Cited on page 1.)

- P. Mertikopoulos and W. H. Sandholm. Riemannian game dynamics. *Journal of Economic Theory*, 177:315–364, 2018. (Cited on page 2.)
- P. Mertikopoulos, B. Lecouat, H. Zenati, C-S. Foo, V. Chandrasekhar, and G. Piliouras. Optimistic mirror descent in saddle-point problems: Going the extra(-gradient) mile. In *ICLR*, 2019. URL <https://openreview.net/forum?id=Bkg8jjC9KQ>. (Cited on page 2.)
- A. Mokhtari, A. Ozdaglar, and S. Pattathil. A unified analysis of extra-gradient and optimistic gradient methods for saddle point problems: Proximal point approach. In *AISTATS*, pages 1497–1507. PMLR, 2020. (Cited on page 2.)
- A. Nemirovski. Prox-method with rate of convergence $o(1/t)$ for variational inequalities with Lipschitz continuous monotone operators and smooth convex-concave saddle point problems. *SIAM Journal on Optimization*, 15(1):229–251, 2004. (Cited on page 23.)
- A. Nemirovski, A. Juditsky, G. Lan, and A. Shapiro. Robust stochastic approximation approach to stochastic programming. *SIAM Journal on Optimization*, 19(4):1574–1609, 2009. (Cited on page 23.)
- J. V. Neumann. Zur theorie der gesellschaftsspiele. *Mathematische Annalen*, 100(1):295–320, 1928. (Cited on page 2.)
- M. Nouiehed, M. Sanjabi, T. Huang, J. D. Lee, and M. Razaviyayn. Solving a class of non-convex min-max games using iterative first order methods. In *NeurIPS*, pages 14934–14942, 2019. (Cited on page 2.)
- D. M. Ostrovskii, A. Lowy, and M. Razaviyayn. Efficient search of first-order Nash equilibria in nonconvex-concave smooth min-max problems. *SIAM Journal on Optimization*, 31(4):2508–2538, 2021. (Cited on page 2.)
- S. Park. Riemannian manifolds are KKM spaces. *Advances in the Theory of Nonlinear Analysis and its Application*, 3(2):64–73, 2019. (Cited on pages 2 and 24.)
- X. Pennec, P. Fillard, and N. Ayache. A Riemannian framework for tensor computing. *International Journal of Computer Vision*, 66(1):41–66, 2006. (Cited on page 24.)
- P. Petersen. *Riemannian Geometry*, volume 171. Springer, 2006. (Cited on page 27.)
- B. T. Polyak. Gradient methods for minimizing functionals. *Zhurnal Vychislitel’noi Matematiki i Matematicheskoi Fiziki*, 3(4):643–653, 1963. (Cited on page 1.)
- B. T. Polyak and A. B. Juditsky. Acceleration of stochastic approximation by averaging. *SIAM Journal on Control and Optimization*, 30(4):838–855, 1992. (Cited on page 23.)
- G. Raskutti and S. Mukherjee. The information geometry of mirror descent. *IEEE Transactions on Information Theory*, 61(3):1451–1457, 2015. (Cited on page 1.)
- S. Reddi, M. Zaheer, S. Sra, B. Póczos, F. Bach, R. Salakhutdinov, and A. Smola. A generic approach for escaping saddle points. In *AISTATS*, pages 1233–1242. PMLR, 2018. (Cited on page 1.)
- R. T. Rockafellar. Monotone operators and the proximal point algorithm. *SIAM Journal on Control and Optimization*, 14(5):877–898, 1976. (Cited on page 23.)

- D. Ruppert. Efficient estimations from a slowly convergent Robbins-Monro process. Technical report, Cornell University Operations Research and Industrial Engineering, 1988. (Cited on page 23.)
- M. Sion. On general minimax theorems. *Pacific Journal of Mathematics*, 8(1):171–176, 1958. (Cited on pages 2, 7, and 23.)
- S. Sra and R. Hosseini. Conic geometric optimization on the manifold of positive definite matrices. *SIAM Journal on Optimization*, 25(1):713–739, 2015. (Cited on page 1.)
- S. Sra and R. Hosseini. Geometric optimization in machine learning. In *Algorithmic Advances in Riemannian Geometry and Applications*, pages 73–91. Springer, 2016. (Cited on page 1.)
- J. Sun, Q. Qu, and J. Wright. Complete dictionary recovery over the sphere I: Overview and the geometric picture. *IEEE Transactions on Information Theory*, 63(2):853–884, 2016a. (Cited on pages 1 and 25.)
- J. Sun, Q. Qu, and J. Wright. Complete dictionary recovery over the sphere II: Recovery by Riemannian trust-region method. *IEEE Transactions on Information Theory*, 63(2):885–914, 2016b. (Cited on pages 1 and 25.)
- Y. Sun, N. Flammarion, and M. Fazel. Escaping from saddle points on Riemannian manifolds. In *NeurIPS*, pages 7274–7284, 2019. (Cited on page 23.)
- K. K. Thekumprampil, P. Jain, P. Netrapalli, and S. Oh. Efficient algorithms for smooth minimax optimization. In *NeurIPS*, pages 12680–12691, 2019. (Cited on page 2.)
- N. Tripuraneni, N. Flammarion, F. Bach, and M. I. Jordan. Averaging stochastic gradient descent on Riemannian manifolds. In *COLT*, pages 650–687, 2018. (Cited on page 23.)
- R. Tron, B. Afsari, and R. Vidal. Riemannian consensus for manifolds with bounded curvature. *IEEE Transactions on Automatic Control*, 58(4):921–934, 2012. (Cited on page 2.)
- B. Vandereycken. Low-rank matrix completion by Riemannian optimization. *SIAM Journal on Optimization*, 23(2):1214–1236, 2013. (Cited on page 25.)
- E. V. Vlatakis-Gkaragkounis, L. Flokas, and G. Piliouras. Poincaré recurrence, cycles and spurious equilibria in gradient-descent-ascent for non-convex non-concave zero-sum games. In *NeurIPS*, pages 10450–10461, 2019. (Cited on page 2.)
- E. V. Vlatakis-Gkaragkounis, L. Flokas, T. Lianas, P. Mertikopoulos, and G. Piliouras. No-regret learning and mixed Nash equilibria: They do not mix. In *NeurIPS*, pages 1380–1391, 2020. (Cited on page 1.)
- E. V. Vlatakis-Gkaragkounis, L. Flokas, and G. Piliouras. Solving min-max optimization with hidden structure via gradient descent ascent. In *NeurIPS*, pages 2373–2386, 2021. (Cited on page 2.)
- J. H. Wang, G. López, V. Martín-Márquez, and C. Li. Monotone and accretive vector fields on Riemannian manifolds. *Journal of Optimization Theory and Applications*, 146(3):691–708, 2010. (Cited on page 24.)

- Z. Wen and W. Yin. A feasible method for optimization with orthogonality constraints. *Mathematical Programming*, 142(1-2):397–434, 2013. (Cited on page 22.)
- A. Wiesel. Geodesic convexity and covariance estimation. *IEEE Transactions on Signal Processing*, 60(12):6182–6189, 2012. (Cited on page 2.)
- Y. Yazıcı, C-S. Foo, S. Winkler, K-H. Yap, G. Piliouras, and V. Chandrasekhar. The unusual effectiveness of averaging in GAN training. In *ICLR*, 2019. URL https://openreview.net/forum?id=SJgw_sRqFQ. (Cited on page 12.)
- H. Zhang and S. Sra. First-order methods for geodesically convex optimization. In *COLT*, pages 1617–1638. PMLR, 2016. (Cited on pages 5, 6, 7, 23, and 27.)
- H. Zhang, S. J. Reddi, and S. Sra. Riemannian SVRG: Fast stochastic optimization on Riemannian manifolds. In *NeurIPS*, pages 4592–4600, 2016. (Cited on page 23.)
- J. Zhang, S. Ma, and S. Zhang. Primal-dual optimization algorithms over Riemannian manifolds: An iteration complexity analysis. *Mathematical Programming*, 184(1):445–490, 2020. (Cited on page 22.)
- J. Zhang, M. Hong, and S. Zhang. On lower iteration complexity bounds for the convex concave saddle point problems. *Mathematical Programming*, pages 1–35, 2021. (Cited on page 9.)
- P. Zhang, J. Zhang, and S. Sra. Minimax in geodesic metric spaces: Sion’s theorem and algorithms. *ArXiv Preprint: 2202.06950*, 2022. (Cited on pages 1, 2, 3, 4, 6, 7, 8, 9, 11, 13, 23, 24, 30, 35, and 38.)

A Related Work

The literature for the geometric properties of Riemannian Manifolds is immense and hence we cannot hope to survey them here; for an appetizer, we refer the reader to [Burago et al. \[2001\]](#) and [Lee \[2012\]](#) and references therein. On the other hand, as stated, it is not until recently that the long-run non-asymptotic behavior of optimization algorithms in Riemannian manifolds (even the smooth ones) has encountered a lot of interest. For concision, we have deferred here a detailed exposition of the rest of recent results to Appendix A of the paper’s supplement. Additionally, in Appendix B we also give a bunch of motivating examples which can be solved by Riemannian min-max optimization.

Minimization on Riemannian manifolds. Many application problems can be formulated as the minimization or maximization of a smooth function over Riemannian manifold and has triggered a line of research on the extension of the classical first-order and second-order methods to Riemannian setting with asymptotic convergence to first-order stationary points in general [[Absil et al., 2009](#)]. Recent years have witnessed the renewed interests on nonasymptotic convergence analysis of solution methods. In particular, [Boumal et al. \[2019\]](#) proved the global sublinear convergence results for Riemannian gradient descent method and Riemannian trust region method, and further demonstrated that the Riemannian trust region method converges to a second-order stationary point in polynomial time; see also similar results in some other works [[Kasai and Mishra, 2018](#), [Hu et al., 2018, 2019](#)]. We are also aware of recent works on problem-specific methods [[Wen and Yin, 2013](#), [Gao et al., 2018](#), [Liu et al., 2019](#)] and primal-dual methods [[Zhang et al., 2020](#)].

Compared to the smooth counterpart, Riemannian nonsmooth optimization is harder and relatively less explored [Absil and Hosseini, 2019]. A few existing works focus on optimizing geodesically convex functions over Riemannian manifold with subgradient methods [Ferreira and Oliveira, 1998, Zhang and Sra, 2016, Bento et al., 2017]. In particular, Ferreira and Oliveira [1998] provided the first asymptotic convergence result while Zhang and Sra [2016] and [Bento et al., 2017] proved a nonasymptotic global convergence rate of $O(\epsilon^{-2})$ for Riemannian subgradient methods. Further, Ferreira and Oliveira [2002] assumed that the proximal mapping over Riemannian manifold is computationally tractable and proved the global sublinear convergence of Riemannian proximal point method. Focusing on optimization over Stiefel manifold, Chen et al. [2020] studied the composite objective function and proposed Riemannian proximal gradient method which only needs to compute the proximal mapping of nonsmooth component function over the tangent space of Stiefel manifold. Li et al. [2021] consider optimizing a weakly convex function over Stiefel manifold and proposed Riemannian subgradient methods that drive a near-optimal stationarity measure below ϵ within the number of iterations bounded by $O(\epsilon^{-4})$.

There are some results on stochastic optimization over Riemannian manifold. In particular, Bonnabel [2013] proved the first asymptotic convergence result for Riemannian stochastic gradient descent, which is extended by a line of subsequent works [Zhang et al., 2016, Tripuraneni et al., 2018, Becigneul and Ganea, 2019, Kasai et al., 2019]. If the Riemannian Hessian is not positive definite, some recent works have suggested frameworks to escape saddle points [Sun et al., 2019, Criscitiello and Boumal, 2019].

Min-Max optimization in Euclidean spaces. Focusing on solving specifically min-max problems, the algorithms under euclidean geometry have a very rich history in optimization that goes back at least to the original proximal point algorithms [Martinet, 1970, Rockafellar, 1976] for variational inequality (VI) problems; At a high level, if the objective function is Lipschitz and strictly convex-concave, the simple forward-backward schemes are known to converge – and if combined with a Polyak–Ruppert averaging scheme [Ruppert, 1988, Polyak and Juditsky, 1992, Nemirovski et al., 2009], they achieve an $O(1/\epsilon^2)$ complexity² without the caveat of strictness [Bauschke and Combettes, 2011]. If, in addition, the objective admits Lipschitz continuous gradients, then the extragradient (EG) algorithm [Korpelevich, 1976] achieves trajectory convergence without strict monotonicity requirements, while the time-average iterate converges at $O(1/\epsilon)$ steps [Nemirovski, 2004]. Finally, if the problem is strongly convex-concave, forward-backward methods computes an ϵ -saddle point at $O(1/\epsilon)$ steps; and if the operator is also Lipschitz continuous, classical results in operator theory show that simple forward-backward methods suffice to achieve a linear convergence rate [Facchinei and Pang, 2007, Bauschke and Combettes, 2011].

Min-Max optimization on Riemannian manifolds. In the case of nonlinear geometry, the literature has been devoted on two different orthogonal axes: *a)* the existence of saddle point for min-max objective bi-functions and *b)* the design of algorithms for the computation of such points. For the existence of saddle point, a long line of recent work tried to generalize the seminal minima theorem for quasi-convex-quasi-concave problems of Sion [1958]. The crucial bottleneck of this generalization to Riemannian smooth manifolds had been the application of both Knaster–Kuratowski–Mazurkiewicz (KKM) theorem and Helly’s theorem in non-flat spaces. Before Zhang et al. [2022], the existence of

²For the rest of the presentation, we adopt the convention of presenting the *fine-grained complexity* performance measure for computing an $O(\epsilon)$ -close solution instead of the *convergence rate* of a method. Thus a rate of the form $\|\mathbf{x}_t - \mathbf{x}^*\| \leq O(1/t^{1/p})$ typically corresponds to $O(1/\epsilon^p)$ gradient computations and the geometric rate $\|\mathbf{x}_t - \mathbf{x}^*\| \leq O(\exp(-\mu t))$ matches usually up with the $O(\ln(1/\epsilon))$ computational complexity.

saddle points had been identified for the special case of Hadamard manifolds [Komiya, 1988, Kristály, 2014, Bento et al., 2017, Park, 2019].

Similar with the existence results, initially the developed methods referred to the computation of singularities in monotone variational operators typically in hyperbolic Hadamard manifolds with negative curvature [Li et al., 2009]. More recently, Huang et al. [2020] proposed a Riemannian gradient descent ascent method (RGDA), yet the analysis is restricted to $\mathcal{N}$ being a convex subset of the Euclidean space and $f(x, y)$ being strongly concave in y . It is worth mentioning that for the case Hadamard and generally hyperbolic manifolds, extra-gradient style algorithms have been proposed [Wang et al., 2010, Ferreira et al., 2005] in the literature, establishing mainly their asymptotic convergence. However it was not until recent Zhang et al. [2022] that the riemannian correction trick has been analyzed for the case of the extra-gradient algorithm. Bearing in our mind the higher-order methods, Han et al. [2022] has recently proposed the Riemannian Hamiltonian Descent and versions of Newton’s method for geodesic convex geodesic concave functions. Since in this work, we focus only on first-order methods, we don’t compare with the aforementioned Hamiltonian alternative since it incorporates always the extra computational burden of second-derivatives and hessian over a manifold.

B Motivating Examples

We provide some examples of Riemannian min-max optimization to give a sense of their expressivity. Two of the examples are the generic models from the optimization literature [Ben-Tal et al., 2009, Absil et al., 2009, Hu et al., 2020] and the two others are the formulations of application problems arising from machine learning and data analytics [Pennec et al., 2006, Fletcher and Joshi, 2007, Lin et al., 2020a].

Example B.1 (Riemannian optimization with nonlinear constraints) *We can consider a rather straightforward generalization of constrained optimization problem from Euclidean spaces to Riemannian manifolds [Bergmann and Herzog, 2019]. This formulation finds a wide range of real-world applications, e.g., non-negative principle component analysis, weighted max-cut and so on. Letting $\mathcal{M}$ be a finite-dimensional Riemannian manifold with unique geodesic, we focus on the following problem:*

$$\min_{x \in \mathcal{M}} f(x), \quad \text{s.t. } g(x) \leq 0, \quad h(x) = 0,$$

where $g := (g_1, g_2, \dots, g_m) : \mathcal{M} \mapsto \mathbb{R}^m$ and $h := (h_1, h_2, \dots, h_n) : \mathcal{M} \mapsto \mathbb{R}^n$ are two mappings. Then, we can introduce the dual variables λ and μ and reformulate the aforementioned constrained optimization problem as follows,

$$\min_{x \in \mathcal{M}} \max_{(\lambda, \mu) \in \mathbb{R}_+^m \times \mathbb{R}^n} f(x) + \langle \lambda, g(x) \rangle + \langle \mu, h(x) \rangle.$$

Suppose that f and all of g_i and h_i are geodesically convex and smooth, the above problem is a geodesic-convex-Euclidean-concave min-max optimization problem.

Example B.2 (Distributionally robust Riemannian optimization) *Distributionally robust optimization (DRO) is an effective method to deal with the noisy data, adversarial data, and imbalanced data. We consider the problem of DRO over Riemannian manifold; indeed, given a set of data samples $\{\xi_i\}_{i=1}^N$, the problem of DRO over Riemannian manifold $\mathcal{M}$ can be written in the form of*

$$\min_{x \in \mathcal{M}} \max_{\mathbf{p} \in \mathcal{S}} \sum_{i=1}^N p_i \ell(x; \xi_i) - \|\mathbf{p} - \frac{1}{N} \mathbf{1}\|^2,$$

where $\mathbf{p} = (p_1, p_2, \dots, p_N)$ and $\mathcal{S} = \{\mathbf{p} \in \mathbb{R}^N : \sum_{i=1}^N p_i = 1, p_i \geq 0\}$. In general, $\ell(x; \xi_i)$ denotes the loss function over Riemannian manifold $\mathcal{M}$. If ℓ is geodesically convex and smooth, the above problem is a geodesic-convex-Euclidean-concave min-max optimization problem.

Example B.3 (Robust matrix Karcher mean problem) We consider a robust version of classical matrix Karcher mean problem. More specifically, the Karcher mean of N symmetric positive definite matrices $\{A_i\}_{i=1}^N$ is defined as the matrix $X \in \mathcal{M} = \{X \in \mathbb{R}^{n \times n} : X \succ 0, X = X^\top\}$ that minimizes the sum of squared distance induced by the Riemannian metric:

$$d(X, Y) = \|\log(X^{-1/2} Y X^{-1/2})\|_F.$$

The loss function is thus defined by

$$f(X; \{A_i\}_{i=1}^N) = \sum_{i=1}^N (d(X, A_i))^2.$$

which is known to be nonconvex in Euclidean spaces but geodesically strongly convex. Then, the robust version of classical matrix Karcher mean problem is aiming at solving the following problem:

$$\min_{X \in \mathcal{M}} \max_{Y_i \in \mathcal{M}} f(X; \{Y_i\}_{i=1}^N) - \gamma \left(\sum_{i=1}^N (d(Y_i, A_i))^2 \right),$$

where $\gamma > 0$ stands for the trade-off between the computation of Karcher mean over a set of $\{Y_i\}_{i=1}^N$ and the difference between the observed samples $\{A_i\}_{i=1}^N$ and $\{Y_i\}_{i=1}^N$. It is clear that the above problem is a geodesically strongly-convex-strongly-concave min-max optimization problem.

Example B.4 (Projection robust optimal transport problem) We consider the projection robust optimal transport (OT) problem – a robust variant of the OT problem – that achieves superior sample complexity bound [Lin et al., 2021]. Let $\{x_1, x_2, \dots, x_n\} \subseteq \mathbb{R}^d$ and $\{y_1, y_2, \dots, y_n\} \subseteq \mathbb{R}^d$ denote sets of n atoms, and let $(r_1, r_2, \dots, r_n)$ and $(c_1, c_2, \dots, c_n)$ denote weight vectors. We define discrete probability measures $\mu = \sum_{i=1}^n r_i \delta_{x_i}$ and $\nu = \sum_{j=1}^n c_j \delta_{y_j}$. In this setting, the computation of the k -dimensional projection robust OT distance between μ and ν resorts to solving the following problem:

$$\max_{U \in \text{St}(d, k)} \min_{\pi \in \Pi(\mu, \nu)} \sum_{i=1}^n \sum_{j=1}^n \pi_{i,j} \|U^\top x_i - U^\top y_j\|^2,$$

where $\text{St}(d, k) = \{U \in \mathbb{R}^{d \times k} \mid U^\top U = I_k\}$ is a Stiefel manifold and $\Pi(r, c) = \{\pi \in \mathbb{R}_+^{n \times n} \mid \sum_{j=1}^n \pi_{ij} = r_i, \sum_{i=1}^n \pi_{ij} = c_j\}$ is a transportation polytope. It is worth mentioning that the above problem is a geodesically-nonconvex-Euclidean-concave min-max optimization problem with special structures, making the computation of stationary points tractable. While the global convergence guarantee for our algorithm does not apply, the above problem might be locally geodesically-convex-Euclidean-concave such that our algorithm with sufficiently good initialization works here.

In addition to these examples, it is worth mentioning that Riemannian min-max optimization problems contain all general min-max optimization problems in Euclidean spaces and all Riemannian minimization or maximization optimization problems. It is also an abstraction of many machine learning problems, e.g., principle component analysis [Boumal and Absil, 2011], dictionary learning [Sun et al., 2016a,b], deep neural networks (DNNs) [Huang et al., 2018] and low-rank matrix learning [Vandereyken, 2013, Jawanpuria and Mishra, 2018]; indeed, the problem of principle component analysis resorts to optimization problems on Grassmann manifolds for example.

C Metric Geometry

To generalize the first-order methods in Euclidean setting, we introduce several basic concepts in metric geometry [Burago et al., 2001], which are known to include both Euclidean spaces and Riemannian manifolds as special cases. Formally, we have

Definition C.1 (Metric Space) *A metric space (X, d) is a pair of a set X and a distance function $d(\cdot, \cdot)$ satisfying: (i) $d(x, x') \geq 0$ for any $x, x' \in X$; (ii) $d(x, x') = d(x', x)$ for any $x, x' \in X$; and (iii) $d(x, x'') \leq d(x, x') + d(x', x'')$ for any $x, x', x'' \in X$. In other words, the distance function $d(\cdot, \cdot)$ is non-negative, symmetrical and satisfies the triangle inequality.*

A path $\gamma : [0, 1] \mapsto X$ is a continuous mapping from the interval $[0, 1]$ to X and the *length* of γ is defined as $\text{length}(\gamma) := \lim_{n \rightarrow +\infty} \sup_{0=t_0 < \dots < t_n=1} \sum_{i=1}^n d(\gamma(t_{i-1}), \gamma(t_i))$. Note that the triangle inequality implies that $\sup_{0=t_0 < \dots < t_n=1} \sum_{i=1}^n d(\gamma(t_{i-1}), \gamma(t_i))$ is nondecreasing. Then, the length of a path γ is well defined since the limit is either $+\infty$ or a finite scalar. Moreover, for $\forall \epsilon > 0$, there exists $n \in \mathbb{N}$ and the partition $0 = t_0 < \dots < t_n = 1$ of the interval $[0, 1]$ such that $\text{length}(\gamma) \leq \sum_{i=1}^n d(\gamma(t_{i-1}), \gamma(t_i)) + \epsilon$.

Definition C.2 (Length Space) *A metric space (X, d) is a length space if, for any $x, x' \in X$ and $\epsilon > 0$, there exists a path $\gamma : [0, 1] \mapsto X$ connecting x and x' such that $\text{length}(\gamma) \leq d(x, x') + \epsilon$.*

We can see from Definition C.2 that a set of length spaces is strict subclass of metric spaces; indeed, for some $x, x' \in X$, there does not exist a path γ such that its length can be approximated by $d(x, x')$ for some tolerance $\epsilon > 0$. In metric geometry, a *geodesic* is a path which is locally a distance minimizer everywhere. More precisely, a path γ is a geodesic if there is a constant $\nu > 0$ such that for any $t \in [0, 1]$ there is a neighborhood I of $[0, 1]$ such that,

$$d(\gamma(t_1), \gamma(t_2)) = \nu |t_1 - t_2|, \quad \text{for any } t_1, t_2 \in I.$$

Note that the above generalizes the notion of geodesic for Riemannian manifolds. Then, we are ready to introduce the geodesic space and uniquely geodesic space [Bacak, 2014].

Definition C.3 *A metric space (X, d) is a geodesic space if, for any $x, x' \in X$, there exists a geodesic $\gamma : [0, 1] \mapsto X$ connecting x and x' . Furthermore, it is called uniquely geodesic if the geodesic connecting x and x' is unique for any $x, x' \in X$.*

Trigonometric geometry in nonlinear spaces is intrinsically different from Euclidean space. In particular, we remark that the law of cosines in Euclidean space (with $\|\cdot\|$ as ℓ_2 -norm) is crucial for analyzing the convergence property of optimization algorithms, e.g.,

$$\|a\|^2 = \|b\|^2 + \|c\|^2 - 2bc \cos(A),$$

where a, b, c are sides of a *geodesic triangle* in Euclidean space and A is the angle between b and c . However, such nice property does not hold for nonlinear spaces due to the lack of flat geometry, further motivating us to extend the law of cosines under nonlinear trigonometric geometry. That is to say, given a geodesic triangle in X with sides a, b, c where A is the angle between b and c , we hope to establish the relationship between a^2, b^2, c^2 and $2bc \cos(A)$ in nonlinear spaces; see the main context for the comparing inequalities.

Finally, we specify the definition of *section curvature* of Riemannian manifolds and clarify how such quantity affects the trigonometric comparison inequalities. More specifically, the sectional curvature is

defined as the Gauss curvature of a 2-dimensional sub-manifold that are obtained from the image of a two-dimensional subspace of a tangent space after exponential mapping. It is worth mentioning that the above 2-dimensional sub-manifold is locally isometric to a 2-dimensional sphere, a Euclidean plane, and a hyperbolic plane with the same Gauss curvature if its sectional curvature is positive, zero and negative respectively. Then we are ready to summarize the existing trigonometric comparison inequalities for Riemannian manifold with bounded sectional curvatures. Note that the following two propositions are the full version of Proposition 2.1 and will be used in our subsequent proofs.

Proposition C.1 *Suppose that $\mathcal{M}$ is a Riemannian manifold with sectional curvature that is upper bounded by $\kappa_{\max}$ and let Δ be a geodesic triangle in $\mathcal{M}$ with the side length a, b, c and A which is the angle between b and c . If $\kappa_{\max} > 0$, we assume the diameter of $\mathcal{M}$ is bounded by $\frac{\pi}{\sqrt{\kappa_{\max}}}$. Then, we have*

$$a^2 \geq \underline{\xi}(\kappa_{\max}, c) \cdot b^2 + c^2 - 2bc \cos(A),$$

where $\underline{\xi}(\kappa, c) := 1$ for $\kappa \leq 0$ and $\underline{\xi}(\kappa, c) := c\sqrt{\kappa} \cot(c\sqrt{\kappa}) < 1$ for $\kappa > 0$.

Proposition C.2 *Suppose that $\mathcal{M}$ is a Riemannian manifold with sectional curvature that is lower bounded by $\kappa_{\min}$ and let Δ be a geodesic triangle in $\mathcal{M}$ with the side length a, b, c and A which is the angle between b and c . Then, we have*

$$a^2 \leq \bar{\xi}(\kappa_{\min}, c) \cdot b^2 + c^2 - 2bc \cos(A),$$

where $\bar{\xi}(\kappa, c) := c\sqrt{-\kappa} \coth(c\sqrt{-\kappa}) > 1$ if $\kappa < 0$ and $\bar{\xi}(\kappa, c) := 1$ if $\kappa \geq 0$.

Remark C.1 *Proposition C.1 and C.2 are simply the restatement of Alimisis et al. [2020, Corollary 2.1] and Zhang and Sra [2016, Lemma 5]. The former inequality is obtained when the sectional curvature is bounded from above while the latter inequality characterizes the relationship between the trigonometric lengths when the sectional curvature is bounded from below. If $\kappa_{\min} = \kappa_{\max} = 0$ (i.e., Euclidean spaces), we have $\bar{\xi}(\kappa_{\min}, c) = \underline{\xi}(\kappa_{\max}, c) = 1$. The proof is based on Toponogov's theorem and Riccati comparison estimate [Petersen, 2006, Proposition 25] and we refer the interested readers to Zhang and Sra [2016] and Alimisis et al. [2020] for the details.*

D Riemannian Gradient Descent Ascent for Nonsmooth Setting

In this section, we propose and analyze Riemannian gradient descent ascent (RGDA) method for nonsmooth Riemannian min-max optimization and extend it to stochastic RGDA. We present our results on the optimal last-iterate convergence guarantee for geodesically strongly-convex-strongly-concave setting (both deterministic and stochastic) and time-average convergence guarantee for geodesically convex-concave setting (both deterministic and stochastic).

D.1 Algorithmic scheme

Compared to Riemannian corrected extragradient (RCEG) method, our Riemannian gradient descent ascent (RGDA) method is a relatively straightforward generalization of GDA in Euclidean spaces. More specifically, we start with the scheme of GDA as follows (just consider $\mathcal{M}$ and $\mathcal{N}$ as convex constraint sets in Euclidean spaces),

$$x_{t+1} \leftarrow \text{proj}_{\mathcal{M}}(x_t - \eta_t \cdot g_x^t), \quad y_{t+1} \leftarrow \text{proj}_{\mathcal{N}}(y_t + \eta_t \cdot g_y^t). \quad (\text{D.1})$$

Algorithm 3 RGDA

Input: initial points (x_0, y_0) and stepsizes $\eta_t > 0$.
for $t = 0, 1, 2, \dots, T - 1$ **do**
 Query $(g_x^t, g_y^t) \leftarrow (\text{subgrad}_x f(x_t, y_t), \text{subgrad}_y f(x_t, y_t))$
 as Riemannian subgradient of f at a point (x_t, y_t) .
 $x_{t+1} \leftarrow \text{Exp}_{x_t}(-\eta_t \cdot g_x^t)$.
 $y_{t+1} \leftarrow \text{Exp}_{y_t}(\eta_t \cdot g_y^t)$.
end for

Algorithm 4 SRGDA

Input: initial points (x_0, y_0) and stepsizes $\eta_t > 0$.
for $t = 0, 1, 2, \dots, T - 1$ **do**
 Query (g_x^t, g_y^t) as a **noisy** estimator of Riemannian subgradient of f at a point (x_t, y_t) .
 $x_{t+1} \leftarrow \text{Exp}_{x_t}(-\eta_t \cdot g_x^t)$.
 $y_{t+1} \leftarrow \text{Exp}_{y_t}(\eta_t \cdot g_y^t)$.
end for

where $(g_x^t, g_y^t) \in (\partial_x f(x_t, y_t), \partial_y f(x_t, y_t))$ is one subgradient of f . By replacing the projection operator by the corresponding exponential map and the gradient by the corresponding Riemannian gradient, we have

$$x_{t+1} \leftarrow \text{Exp}_{x_t}(-\eta_t \cdot g_x^t), \quad y_{t+1} \leftarrow \text{Exp}_{y_t}(\eta_t \cdot g_y^t).$$

where $(g_x^t, g_y^t) \leftarrow (\text{subgrad}_x f(x_t, y_t), \text{subgrad}_y f(x_t, y_t))$ is one Riemannian subgradient of f . Then, we summarize the resulting scheme of RGDA method in Algorithm 3 and its stochastic extension with noisy estimators of Riemannian gradients of f in Algorithm 4.

D.2 Main results

We present our main results on the global convergence rate estimation for Algorithm 3 and 4 in terms of Riemannian gradient and noisy Riemannian gradient evaluations. The following assumptions are made throughout for geodesically strongly-convex-strongly-concave and geodesically convex-concave settings.

Assumption D.1 *The objective function $f : \mathcal{M} \times \mathcal{N} \mapsto \mathbb{R}$ and manifolds $\mathcal{M}$ and $\mathcal{N}$ satisfy*

1. *f is geodesically L -Lipschitz and geodesically strongly-convex-strongly-concave with $\mu > 0$.*
2. *The diameter of the domain $\{(x, y) \in \mathcal{M} \times \mathcal{N} : -\infty < f(x, y) < +\infty\}$ is bounded by $D > 0$.*
3. *The sectional curvatures of $\mathcal{M}$ and $\mathcal{N}$ are both bounded in the range $[\kappa_{\min}, +\infty)$ with $\kappa_{\min} \leq 0$.*

Assumption D.2 *The objective function $f : \mathcal{M} \times \mathcal{N} \mapsto \mathbb{R}$ and manifolds $\mathcal{M}$ and $\mathcal{N}$ satisfy*

1. *f is geodesically L -Lipschitz and geodesically convex-concave.*
2. *The diameter of the domain $\{(x, y) \in \mathcal{M} \times \mathcal{N} : -\infty < f(x, y) < +\infty\}$ is bounded by $D > 0$.*
3. *The sectional curvatures of $\mathcal{M}$ and $\mathcal{N}$ are both bounded in the range $[\kappa_{\min}, +\infty)$ with $\kappa_{\min} \leq 0$.*

Imposing the geodesically Lipschitzness condition is crucial to achieve finite-time convergence guarantee if we do not assume the geodesically smoothness condition. Note that we only require the lower bound for the sectional curvatures of manifolds and this is weaker than that presented in the main context.

Letting $(x^*, y^*) \in \mathcal{M} \times \mathcal{N}$ be a global saddle point of f (it exists under either Assumption D.1 or D.2), we let $D_0 = (d_{\mathcal{M}}(x_0, x^*))^2 + (d_{\mathcal{N}}(y_0, y^*))^2 > 0$ and summarize our results for Algorithm 3 in the following theorems.

Theorem D.1 *Under Assumption D.1 and let $\eta_t > 0$ satisfies that $\eta_t = \frac{1}{\mu} \min\{1, \frac{2}{t}\}$. There exists some $T > 0$ such that the output of Algorithm 3 satisfies that $(d(x_T, x^*))^2 + (d(y_T, y^*))^2 \leq \epsilon$ and the total number of Riemannian subgradient evaluations is bounded by*

$$O\left(\frac{\bar{\xi}_0 L^2}{\mu^2 \epsilon}\right),$$

where $\bar{\xi}_0 = \bar{\xi}(\kappa_{\min}, D)$ measures the lower bound for the change of non-flatness in $\mathcal{M}$ and $\mathcal{N}$.

Theorem D.2 Under Assumption D.2 and let $\eta_t > 0$ satisfies that $\eta_t = \frac{1}{L} \sqrt{\frac{D_0}{2\xi_0 T}}$. There exists some $T > 0$ such that the output of Algorithm 3 satisfies that $f(\bar{x}_T, y^*) - f(x^*, \bar{y}_T) \leq \epsilon$ and the total number of Riemannian subgradient evaluations is bounded by

$$O\left(\frac{\bar{\xi}_0 L^2 D_0}{\epsilon^2}\right),$$

where $\bar{\xi}_0 = \bar{\xi}(\kappa_{\min}, D)$ measures the lower bound for the change of non-flatness in $\mathcal{M}$ and $\mathcal{N}$, and the time-average iterates $(\bar{x}_T, \bar{y}_T) \in \mathcal{M} \times \mathcal{N}$ can be computed by $(\bar{x}_0, \bar{y}_0) = (0, 0)$ and the inductive formula: $\bar{x}_{t+1} = \text{Exp}_{\bar{x}_t}(\frac{1}{t+1} \cdot \text{Exp}_{\bar{x}_t}^{-1}(x_t))$ and $\bar{y}_{t+1} = \text{Exp}_{\bar{y}_t}(\frac{1}{t+1} \cdot \text{Exp}_{\bar{y}_t}^{-1}(y_t))$ for all $t = 0, 1, \dots, T-1$.

Remark D.1 Theorem D.1 and D.2 establish the last-iterate and time-average rates of convergence of Algorithm 3 for solving Riemannian min-max optimization problems under Assumption D.1 and D.2 respectively. Further, the dependence on L and $1/\epsilon$ can not be improved since it has matched the lower bound established for the nonsmooth min-max optimization problems in Euclidean spaces.

In the scheme of SRGDA, we highlight that (g_x^t, g_y^t) is a noisy estimators of Riemannian subgradient of f at (x_t, y_t) . It is necessary to impose the conditions such that these estimators are unbiased and has bounded variance. By abuse of notation, we assume that

$$g_x^t = \text{subgrad}_x f(x_t, y_t) + \xi_x^t, \quad g_y^t = \text{subgrad}_y f(x_t, y_t) + \xi_y^t, \quad (\text{D.2})$$

where the noises (ξ_x^t, ξ_y^t) satisfy that

$$\mathbb{E}[\xi_x^t] = 0, \quad \mathbb{E}[\xi_y^t] = 0, \quad \mathbb{E}[\|\xi_x^t\|^2 + \|\xi_y^t\|^2] \leq \sigma^2. \quad (\text{D.3})$$

We are ready to summarize our results for Algorithm 4 in the following theorems.

Theorem D.3 Under Assumption D.1 and let Eq. (D.2) and Eq. (D.3) hold with $\sigma > 0$ and let $\eta_t > 0$ satisfies that $\eta_t = \frac{1}{\mu} \min\{1, \frac{2}{t}\}$. There exists some $T > 0$ such that the output of Algorithm 4 satisfies that $\mathbb{E}[(d(x_T, x^*))^2 + (d(y_T, y^*))^2] \leq \epsilon$ and the total number of noisy Riemannian gradient evaluations is bounded by

$$O\left(\frac{\bar{\xi}_0(L^2 + \sigma^2)}{\mu^2 \epsilon}\right),$$

where $\bar{\xi}_0 = \bar{\xi}(\kappa_{\min}, D)$ measures the lower bound for the change of non-flatness in $\mathcal{M}$ and $\mathcal{N}$.

Theorem D.4 Under Assumption D.2 and let Eq. (D.2) and Eq. (D.3) hold with $\sigma > 0$ and let $\eta_t > 0$ satisfies that $\eta_t = \frac{1}{2} \sqrt{\frac{D_0}{\bar{\xi}_0(L^2 + \sigma^2)T}}$. There exists some $T > 0$ such that the output of Algorithm 4 satisfies that $\mathbb{E}[f(\bar{x}_T, y^*) - f(x^*, \bar{y}_T)] \leq \epsilon$ and the total number of noisy Riemannian gradient evaluations is bounded by

$$O\left(\frac{\bar{\xi}_0(L^2 + \sigma^2)D_0}{\epsilon^2}\right),$$

where $\bar{\xi}_0 = \bar{\xi}(\kappa_{\min}, D)$ measures the lower bound for the change of non-flatness in $\mathcal{M}$ and $\mathcal{N}$, and the time-average iterates $(\bar{x}_T, \bar{y}_T) \in \mathcal{M} \times \mathcal{N}$ can be computed by $(\bar{x}_0, \bar{y}_0) = (0, 0)$ and the inductive formula: $\bar{x}_{t+1} = \text{Exp}_{\bar{x}_t}(\frac{1}{t+1} \cdot \text{Exp}_{\bar{x}_t}^{-1}(x_t))$ and $\bar{y}_{t+1} = \text{Exp}_{\bar{y}_t}(\frac{1}{t+1} \cdot \text{Exp}_{\bar{y}_t}^{-1}(y_t))$ for all $t = 0, 1, \dots, T-1$.

Remark D.2 Theorem D.3 and D.4 establish the last-iterate and time-average rates of convergence of Algorithm 4 for solving Riemannian min-max optimization problems under Assumption D.1 and D.2. Moreover, the dependence on L and $1/\epsilon$ can not be improved since it has matched the lower bound established for nonsmooth stochastic min-max optimization problems in Euclidean spaces.

E Missing Proofs for Riemannian Corrected Extragradient Method

In this section, we present some technical lemmas for analyzing the convergence property of Algorithm 1 and 2. We also give the proofs of Theorem 3.1, 3.2 and 3.3.

E.1 Technical lemmas

We provide two technical lemmas for analyzing Algorithm 1 and 2 respectively. Parts of the first lemma were presented in Zhang et al. [2022, Lemma C.1]. For the completeness, we provide the proof details.

Lemma E.1 Under Assumption 3.1 and let $\{(x_t, y_t), (\hat{x}_t, \hat{y}_t)\}_{t=0}^{T-1}$ be generated by Algorithm 1 with the stepsize $\eta > 0$. Then, we have

$$\begin{aligned} 0 \leq & \frac{1}{2} ((d_{\mathcal{M}}(x_t, x^*))^2 - (d_{\mathcal{M}}(x_{t+1}, x^*))^2 + (d_{\mathcal{N}}(y_t, y^*))^2 - (d_{\mathcal{N}}(y_{t+1}, y^*))^2) \\ & + 2\bar{\xi}_0 \eta^2 \ell^2 ((d_{\mathcal{M}}(\hat{x}_t, x_t))^2 + (d_{\mathcal{N}}(\hat{y}_t, y_t))^2 - \frac{1}{2}\bar{\xi}_0 ((d_{\mathcal{M}}(\hat{x}_t, x_t))^2 + (d_{\mathcal{N}}(\hat{y}_t, y_t))^2) \\ & - \frac{\mu\eta}{2} ((d_{\mathcal{M}}(\hat{x}_t, x^*))^2 + (d_{\mathcal{N}}(\hat{y}_t, y^*))^2). \end{aligned}$$

where $(x^*, y^*) \in \mathcal{M} \times \mathcal{N}$ is a global saddle point of f .

Proof. Since f is geodesically ℓ -smooth, we have the Riemannian gradients of f , i.e., $(\text{grad}_x f, \text{grad}_y f)$, are well defined. Since f is geodesically strongly-concave-strongly-concave with the modulus $\mu \geq 0$ (here $\mu = 0$ means that f is geodesically concave-concave), we have

$$\begin{aligned} f(\hat{x}_t, y^*) - f(x^*, \hat{y}_t) &= f(\hat{x}_t, \hat{y}_t) - f(x^*, \hat{y}_t) - (f(\hat{x}_t, \hat{y}_t) - f(\hat{x}_t, y^*)) \\ &\stackrel{\text{Definition 2.2}}{\leq} -\langle \text{grad}_x f(\hat{x}_t, \hat{y}_t), \text{Exp}_{\hat{x}_t}^{-1}(x^*) \rangle + \langle \text{grad}_y f(\hat{x}_t, \hat{y}_t), \text{Exp}_{\hat{y}_t}^{-1}(y^*) \rangle - \frac{\mu}{2} (d_{\mathcal{M}}(\hat{x}_t, x^*))^2 - \frac{\mu}{2} (d_{\mathcal{N}}(\hat{y}_t, y^*))^2. \end{aligned}$$

Since $(x^*, y^*) \in \mathcal{M} \times \mathcal{N}$ is a global saddle point of f , we have $f(\hat{x}_t, y^*) - f(x^*, \hat{y}_t) \geq 0$. Recalling also from the scheme of Algorithm 1 that we have

$$\begin{aligned} x_{t+1} &\leftarrow \text{Exp}_{\hat{x}_t}(-\eta \cdot \text{grad}_x f(\hat{x}_t, \hat{y}_t) + \text{Exp}_{\hat{x}_t}^{-1}(x_t)), \\ y_{t+1} &\leftarrow \text{Exp}_{\hat{y}_t}(\eta \cdot \text{grad}_y f(\hat{x}_t, \hat{y}_t) + \text{Exp}_{\hat{y}_t}^{-1}(y_t)). \end{aligned}$$

By the definition of an exponential map, we have

$$\begin{aligned} \text{Exp}_{\hat{x}_t}^{-1}(x_{t+1}) &= -\eta \cdot \text{grad}_x f(\hat{x}_t, \hat{y}_t) + \text{Exp}_{\hat{x}_t}^{-1}(x_t), \\ \text{Exp}_{\hat{y}_t}^{-1}(y_{t+1}) &= \eta \cdot \text{grad}_y f(\hat{x}_t, \hat{y}_t) + \text{Exp}_{\hat{y}_t}^{-1}(y_t). \end{aligned} \tag{E.1}$$

This implies that

$$\begin{aligned} -\langle \text{grad}_x f(\hat{x}_t, \hat{y}_t), \text{Exp}_{\hat{x}_t}^{-1}(x^*) \rangle &= \frac{1}{\eta} (\langle \text{Exp}_{\hat{x}_t}^{-1}(x_{t+1}), \text{Exp}_{\hat{x}_t}^{-1}(x^*) \rangle - \langle \text{Exp}_{\hat{x}_t}^{-1}(x_t), \text{Exp}_{\hat{x}_t}^{-1}(x^*) \rangle), \\ \langle \text{grad}_y f(\hat{x}_t, \hat{y}_t), \text{Exp}_{\hat{y}_t}^{-1}(y^*) \rangle &= \frac{1}{\eta} (\langle \text{Exp}_{\hat{y}_t}^{-1}(y_{t+1}), \text{Exp}_{\hat{y}_t}^{-1}(y^*) \rangle - \langle \text{Exp}_{\hat{y}_t}^{-1}(y_t), \text{Exp}_{\hat{y}_t}^{-1}(y^*) \rangle). \end{aligned}$$

Putting these pieces together yields that

$$0 \leq \frac{1}{\eta} (\langle \text{Exp}_{\hat{x}_t}^{-1}(x_{t+1}), \text{Exp}_{\hat{x}_t}^{-1}(x^*) \rangle - \langle \text{Exp}_{\hat{x}_t}^{-1}(x_t), \text{Exp}_{\hat{x}_t}^{-1}(x^*) \rangle) - \frac{\mu}{2} (d_{\mathcal{M}}(\hat{x}_t, x^*))^2 \\ + \frac{1}{\eta} (\langle \text{Exp}_{\hat{y}_t}^{-1}(y_{t+1}), \text{Exp}_{\hat{y}_t}^{-1}(y^*) \rangle - \langle \text{Exp}_{\hat{y}_t}^{-1}(y_t), \text{Exp}_{\hat{y}_t}^{-1}(y^*) \rangle) - \frac{\mu}{2} (d_{\mathcal{N}}(\hat{y}_t, y^*))^2.$$

Equivalently, we have

$$0 \leq \langle \text{Exp}_{\hat{x}_t}^{-1}(x_{t+1}), \text{Exp}_{\hat{x}_t}^{-1}(x^*) \rangle - \langle \text{Exp}_{\hat{x}_t}^{-1}(x_t), \text{Exp}_{\hat{x}_t}^{-1}(x^*) \rangle - \frac{\mu\eta}{2} (d_{\mathcal{M}}(\hat{x}_t, x^*))^2 \quad (\text{E.2}) \\ + \langle \text{Exp}_{\hat{y}_t}^{-1}(y_{t+1}), \text{Exp}_{\hat{y}_t}^{-1}(y^*) \rangle - \langle \text{Exp}_{\hat{y}_t}^{-1}(y_t), \text{Exp}_{\hat{y}_t}^{-1}(y^*) \rangle - \frac{\mu\eta}{2} (d_{\mathcal{N}}(\hat{y}_t, y^*))^2.$$

It suffices to bound the terms in the right-hand side of Eq. (E.2) by leveraging the celebrated comparison inequalities on Riemannian manifold with bounded sectional curvature (see Proposition C.1 and C.2). More specifically, we define the constants using $\bar{\xi}(\cdot, \cdot)$ and $\underline{\xi}(\cdot, \cdot)$ from Proposition C.1 and C.2 as follows,

$$\bar{\xi}_0 = \bar{\xi}(\kappa_{\min}, D), \quad \underline{\xi}_0 = \underline{\xi}(\kappa_{\max}, D).$$

By Proposition C.1 and using that $\max\{d_{\mathcal{M}}(\hat{x}_t, x^*), d_{\mathcal{N}}(\hat{y}_t, y^*)\} \leq D$, we have

$$-\langle \text{Exp}_{\hat{x}_t}^{-1}(x_t), \text{Exp}_{\hat{x}_t}^{-1}(x^*) \rangle \leq -\frac{1}{2} \left(\underline{\xi}_0 (d_{\mathcal{M}}(\hat{x}_t, x_t))^2 + (d_{\mathcal{M}}(\hat{x}_t, x^*))^2 - (d_{\mathcal{M}}(x_t, x^*))^2 \right), \quad (\text{E.3}) \\ -\langle \text{Exp}_{\hat{y}_t}^{-1}(y_t), \text{Exp}_{\hat{y}_t}^{-1}(y^*) \rangle \leq -\frac{1}{2} \left(\underline{\xi}_0 (d_{\mathcal{N}}(\hat{y}_t, y_t))^2 + (d_{\mathcal{N}}(\hat{y}_t, y^*))^2 - (d_{\mathcal{N}}(y_t, y^*))^2 \right).$$

By Proposition C.2 and using that $\max\{d_{\mathcal{M}}(\hat{x}_t, x^*), d_{\mathcal{N}}(\hat{y}_t, y^*)\} \leq D$, we have

$$\langle \text{Exp}_{\hat{x}_t}^{-1}(x_{t+1}), \text{Exp}_{\hat{x}_t}^{-1}(x^*) \rangle \leq \frac{1}{2} (\bar{\xi}_0 (d_{\mathcal{M}}(\hat{x}_t, x_{t+1}))^2 + (d_{\mathcal{M}}(\hat{x}_t, x^*))^2 - (d_{\mathcal{M}}(x_{t+1}, x^*))^2).$$

and

$$\langle \text{Exp}_{\hat{y}_t}^{-1}(y_{t+1}), \text{Exp}_{\hat{y}_t}^{-1}(y^*) \rangle \leq \frac{1}{2} (\bar{\xi}_0 (d_{\mathcal{N}}(\hat{y}_t, y_{t+1}))^2 + (d_{\mathcal{N}}(\hat{y}_t, y^*))^2 - (d_{\mathcal{N}}(y_{t+1}, y^*))^2).$$

By the definition of an exponential map and Riemannian metric, we have

$$d_{\mathcal{M}}(\hat{x}_t, x_{t+1}) = \|\text{Exp}_{\hat{x}_t}^{-1}(x_{t+1})\| \stackrel{\text{Eq. (E.1)}}{=} \|\eta \cdot \text{grad}_x f(\hat{x}_t, \hat{y}_t) - \text{Exp}_{\hat{x}_t}^{-1}(x_t)\|, \quad (\text{E.4}) \\ d_{\mathcal{N}}(\hat{y}_t, y_{t+1}) = \|\text{Exp}_{\hat{y}_t}^{-1}(y_{t+1})\| \stackrel{\text{Eq. (E.1)}}{=} \|\eta \cdot \text{grad}_y f(\hat{x}_t, \hat{y}_t) + \text{Exp}_{\hat{y}_t}^{-1}(y_t)\|.$$

Further, we see from the scheme of Algorithm 1 that we have

$$\hat{x}_t \leftarrow \text{Exp}_{x_t}(-\eta \cdot \text{grad}_x f(x_t, y_t)), \\ \hat{y}_t \leftarrow \text{Exp}_{y_t}(\eta \cdot \text{grad}_y f(x_t, y_t)).$$

By the definition of an exponential map, we have

$$\text{Exp}_{\hat{x}_t}^{-1}(\hat{x}_t) = -\eta \cdot \text{grad}_x f(x_t, y_t), \quad \text{Exp}_{\hat{y}_t}^{-1}(\hat{y}_t) = \eta \cdot \text{grad}_y f(x_t, y_t).$$

Using the definition of a parallel transport map and the above equations, we have

$$\text{Exp}_{\hat{x}_t}^{-1}(x_t) = \eta \cdot \Gamma_{x_t}^{\hat{x}_t} \text{grad}_x f(x_t, y_t), \quad \text{Exp}_{\hat{y}_t}^{-1}(y_t) = -\eta \cdot \Gamma_{y_t}^{\hat{y}_t} \text{grad}_y f(x_t, y_t)$$

Since f is geodesically ℓ -smooth, we have

$$\|\text{grad}_x f(\hat{x}_t, \hat{y}_t) - \Gamma_{x_t}^{\hat{x}_t} \text{grad}_x f(x_t, y_t)\| \leq \ell (d_{\mathcal{M}}(\hat{x}_t, x_t) + d_{\mathcal{N}}(\hat{y}_t, y_t)), \\ \|\text{grad}_y f(\hat{x}_t, \hat{y}_t) - \Gamma_{y_t}^{\hat{y}_t} \text{grad}_y f(x_t, y_t)\| \leq \ell (d_{\mathcal{M}}(\hat{x}_t, x_t) + d_{\mathcal{N}}(\hat{y}_t, y_t)).$$

Plugging the above inequalities into Eq. (E.4) yields that

$$\max \{d_{\mathcal{M}}(\hat{x}_t, x_{t+1}), d_{\mathcal{N}}(\hat{y}_t, y_{t+1})\} \leq \eta \ell(d_{\mathcal{M}}(\hat{x}_t, x_t) + d_{\mathcal{N}}(\hat{y}_t, y_t)).$$

Therefore, we have

$$\begin{aligned} \langle \text{Exp}_{\hat{x}_t}^{-1}(x_{t+1}), \text{Exp}_{\hat{x}_t}^{-1}(x^*) \rangle &\leq \frac{1}{2} (2\bar{\xi}_0 \eta^2 \ell^2 ((d_{\mathcal{M}}(\hat{x}_t, x_t))^2 + (d_{\mathcal{N}}(\hat{y}_t, y_t))^2) + (d_{\mathcal{M}}(\hat{x}_t, x^*))^2 - (d_{\mathcal{M}}(x_{t+1}, x^*))^2), \\ \langle \text{Exp}_{\hat{y}_t}^{-1}(y_{t+1}), \text{Exp}_{\hat{y}_t}^{-1}(y^*) \rangle &\leq \frac{1}{2} (2\bar{\xi}_0 \eta^2 \ell^2 ((d_{\mathcal{M}}(\hat{x}_t, x_t))^2 + (d_{\mathcal{N}}(\hat{y}_t, y_t))^2) + (d_{\mathcal{N}}(\hat{y}_t, y^*))^2 - (d_{\mathcal{N}}(y_{t+1}, y^*))^2). \end{aligned}$$

Plugging the above inequalities and Eq. (E.3) into Eq. (E.2) yields the desired inequality. $\square$

The second lemma gives another key inequality that is satisfied by the iterates generated by Algorithm 2.

Lemma E.2 *Under Assumption 3.1 (or Assumption 3.2) and the noisy model (cf. Eq. (3.2) and (3.3)) and let $\{(x_t, y_t), (\hat{x}_t, \hat{y}_t)\}_{t=0}^{T-1}$ be generated by Algorithm 2 with the stepsize $\eta > 0$. Then, we have*

$$\begin{aligned} \mathbb{E}[f(\hat{x}_t, y^*) - f(x^*, \hat{y}_t)] &\leq \frac{1}{2\eta} \mathbb{E}[(d_{\mathcal{M}}(x_t, x^*))^2 - (d_{\mathcal{M}}(x_{t+1}, x^*))^2 + (d_{\mathcal{N}}(y_t, y^*))^2 - (d_{\mathcal{N}}(y_{t+1}, y^*))^2] \\ &\quad + 6\bar{\xi}_0 \eta \ell^2 \mathbb{E}[(d_{\mathcal{M}}(\hat{x}_t, x_t))^2 + (d_{\mathcal{N}}(\hat{y}_t, y_t))^2] - \frac{1}{2\eta} \xi \mathbb{E}[(d_{\mathcal{M}}(\hat{x}_t, x_t))^2 + (d_{\mathcal{N}}(\hat{y}_t, y_t))^2] \\ &\quad - \frac{\mu}{2} \mathbb{E}[(d_{\mathcal{M}}(\hat{x}_t, x^*))^2 + (d_{\mathcal{N}}(\hat{y}_t, y^*))^2] + 3\bar{\xi}_0 \eta \sigma^2, \end{aligned}$$

where $(x^*, y^*) \in \mathcal{M} \times \mathcal{N}$ is a global saddle point of f .

Proof. Using the same argument, we have ($\mu = 0$ refers to geodesically convex-concave case)

$$\begin{aligned} f(\hat{x}_t, y^*) - f(x^*, \hat{y}_t) &= f(\hat{x}_t, \hat{y}_t) - f(x^*, \hat{y}_t) - (f(\hat{x}_t, \hat{y}_t) - f(\hat{x}_t, y^*)) \\ &\leq -\langle \text{grad}_x f(\hat{x}_t, \hat{y}_t), \text{Exp}_{\hat{x}_t}^{-1}(x^*) \rangle + \langle \text{grad}_y f(\hat{x}_t, \hat{y}_t), \text{Exp}_{\hat{y}_t}^{-1}(y^*) \rangle - \frac{\mu}{2} (d_{\mathcal{M}}(\hat{x}_t, x^*))^2 - \frac{\mu}{2} (d_{\mathcal{N}}(\hat{y}_t, y^*))^2. \end{aligned}$$

Combining the arguments used in Lemma E.1 and the scheme of Algorithm 2, we have

$$\begin{aligned} -\langle \hat{g}_x^t, \text{Exp}_{\hat{x}_t}^{-1}(x^*) \rangle &= \frac{1}{\eta} (\langle \text{Exp}_{\hat{x}_t}^{-1}(x_{t+1}), \text{Exp}_{\hat{x}_t}^{-1}(x^*) \rangle - \langle \text{Exp}_{\hat{x}_t}^{-1}(x_t), \text{Exp}_{\hat{x}_t}^{-1}(x^*) \rangle), \\ \langle \hat{g}_y^t, \text{Exp}_{\hat{y}_t}^{-1}(y^*) \rangle &= \frac{1}{\eta} (\langle \text{Exp}_{\hat{y}_t}^{-1}(y_{t+1}), \text{Exp}_{\hat{y}_t}^{-1}(y^*) \rangle - \langle \text{Exp}_{\hat{y}_t}^{-1}(y_t), \text{Exp}_{\hat{y}_t}^{-1}(y^*) \rangle). \end{aligned}$$

Putting these pieces together with Eq. (3.2) yields that

$$\begin{aligned} f(\hat{x}_t, y^*) - f(x^*, \hat{y}_t) &\leq \frac{1}{\eta} (\langle \text{Exp}_{\hat{x}_t}^{-1}(x_{t+1}), \text{Exp}_{\hat{x}_t}^{-1}(x^*) \rangle - \langle \text{Exp}_{\hat{x}_t}^{-1}(x_t), \text{Exp}_{\hat{x}_t}^{-1}(x^*) \rangle) \\ &\quad + \frac{1}{\eta} (\langle \text{Exp}_{\hat{y}_t}^{-1}(y_{t+1}), \text{Exp}_{\hat{y}_t}^{-1}(y^*) \rangle - \langle \text{Exp}_{\hat{y}_t}^{-1}(y_t), \text{Exp}_{\hat{y}_t}^{-1}(y^*) \rangle) - \frac{\mu}{2} (d_{\mathcal{M}}(\hat{x}_t, x^*))^2 - \frac{\mu}{2} (d_{\mathcal{N}}(\hat{y}_t, y^*))^2 \\ &\quad + \langle \hat{\xi}_x^t, \text{Exp}_{\hat{x}_t}^{-1}(x^*) \rangle - \langle \hat{\xi}_y^t, \text{Exp}_{\hat{y}_t}^{-1}(y^*) \rangle. \end{aligned} \tag{E.5}$$

By the same argument as used in Lemma E.1, we have

$$\begin{aligned} -\langle \text{Exp}_{\hat{x}_t}^{-1}(x_t), \text{Exp}_{\hat{x}_t}^{-1}(x^*) \rangle &\leq -\frac{1}{2} \left(\bar{\xi}_0 (d_{\mathcal{M}}(\hat{x}_t, x_t))^2 + (d_{\mathcal{M}}(\hat{x}_t, x^*))^2 - (d_{\mathcal{M}}(x_t, x^*))^2 \right), \\ -\langle \text{Exp}_{\hat{y}_t}^{-1}(y_t), \text{Exp}_{\hat{y}_t}^{-1}(y^*) \rangle &\leq -\frac{1}{2} \left(\bar{\xi}_0 (d_{\mathcal{N}}(\hat{y}_t, y_t))^2 + (d_{\mathcal{N}}(\hat{y}_t, y^*))^2 - (d_{\mathcal{N}}(y_t, y^*))^2 \right), \end{aligned} \tag{E.6}$$

and

$$\begin{aligned} \langle \text{Exp}_{\hat{x}_t}^{-1}(x_{t+1}), \text{Exp}_{\hat{x}_t}^{-1}(x^*) \rangle &\leq \frac{1}{2} (\bar{\xi}_0 \eta^2 \|\hat{g}_x^t - \Gamma_{\hat{x}_t}^{\hat{x}_t} g_x^t\|^2 + (d_{\mathcal{M}}(\hat{x}_t, x^*))^2 - (d_{\mathcal{M}}(x_{t+1}, x^*))^2), \\ \langle \text{Exp}_{\hat{y}_t}^{-1}(y_{t+1}), \text{Exp}_{\hat{y}_t}^{-1}(y^*) \rangle &\leq \frac{1}{2} (\bar{\xi}_0 \eta^2 \|\hat{g}_y^t - \Gamma_{\hat{y}_t}^{\hat{y}_t} g_y^t\|^2 + (d_{\mathcal{N}}(\hat{y}_t, y^*))^2 - (d_{\mathcal{N}}(y_{t+1}, y^*))^2). \end{aligned}$$

Since f is geodesically ℓ -smooth and Eq. (3.2) holds, we have

$$\begin{aligned}\|\hat{g}_x^t - \Gamma_{x_t}^{\hat{x}_t} g_x^t\|^2 &\leq 3\|\hat{\xi}_x^t\|^2 + 3\|\xi_x^t\|^2 + 6\ell^2(d_{\mathcal{M}}(\hat{x}_t, x_t))^2 + 6\ell^2(d_{\mathcal{N}}(\hat{y}_t, y_t))^2, \\ \|\hat{g}_y^t - \Gamma_{y_t}^{\hat{y}_t} g_y^t\|^2 &\leq 3\|\hat{\xi}_y^t\|^2 + 3\|\xi_y^t\|^2 + 6\ell^2(d_{\mathcal{M}}(\hat{x}_t, x_t))^2 + 6\ell^2(d_{\mathcal{N}}(\hat{y}_t, y_t))^2.\end{aligned}$$

Therefore, we have

$$\begin{aligned}&\langle \text{Exp}_{\hat{x}_t}^{-1}(x_{t+1}), \text{Exp}_{\hat{x}_t}^{-1}(x^*) \rangle + \langle \text{Exp}_{\hat{y}_t}^{-1}(y_{t+1}), \text{Exp}_{\hat{y}_t}^{-1}(y^*) \rangle \\ &\leq 6\bar{\xi}_0\eta^2\ell^2((d_{\mathcal{M}}(\hat{x}_t, x_t))^2 + (d_{\mathcal{N}}(\hat{y}_t, y_t))^2) + \frac{3}{2}\bar{\xi}_0\eta^2(\|\hat{\xi}_x^t\|^2 + \|\xi_x^t\|^2 + \|\hat{\xi}_y^t\|^2 + \|\xi_y^t\|^2) \\ &\quad + \frac{1}{2}((d_{\mathcal{M}}(\hat{x}_t, x^*))^2 - (d_{\mathcal{M}}(x_{t+1}, x^*))^2 + (d_{\mathcal{N}}(\hat{y}_t, y^*))^2 - (d_{\mathcal{N}}(y_{t+1}, y^*))^2).\end{aligned}$$

Plugging the above inequalities and Eq. (E.6) into Eq. (E.5) yields that

$$\begin{aligned}f(\hat{x}_t, y^*) - f(x^*, \hat{y}_t) &\leq \frac{1}{2\eta}((d_{\mathcal{M}}(x_t, x^*))^2 - (d_{\mathcal{M}}(x_{t+1}, x^*))^2 + (d_{\mathcal{N}}(y_t, y^*))^2 - (d_{\mathcal{N}}(y_{t+1}, y^*))^2) \\ &\quad + 6\bar{\xi}_0\eta\ell^2((d_{\mathcal{M}}(\hat{x}_t, x_t))^2 + (d_{\mathcal{N}}(\hat{y}_t, y_t))^2) + \frac{3}{2}\bar{\xi}_0\eta(\|\hat{\xi}_x^t\|^2 + \|\xi_x^t\|^2 + \|\hat{\xi}_y^t\|^2 + \|\xi_y^t\|^2) \\ &\quad - \frac{1}{2\eta}\xi_0((d_{\mathcal{M}}(\hat{x}_t, x_t))^2 + (d_{\mathcal{N}}(\hat{y}_t, y_t))^2) - \frac{\mu}{2}(d_{\mathcal{M}}(\hat{x}_t, x^*))^2 - \frac{\mu}{2}(d_{\mathcal{N}}(\hat{y}_t, y^*))^2 \\ &\quad + \langle \hat{\xi}_x^t, \text{Exp}_{\hat{x}_t}^{-1}(x^*) \rangle - \langle \hat{\xi}_y^t, \text{Exp}_{\hat{y}_t}^{-1}(y^*) \rangle.\end{aligned}$$

Taking the expectation of both sides and using Eq. (3.3) yields the desired inequality. $\square$

E.2 Proof of Theorem 3.1

Since Riemannian metrics satisfy the triangle inequality, we have

$$(d_{\mathcal{M}}(\hat{x}_t, x^*))^2 + (d_{\mathcal{N}}(\hat{y}_t, y^*))^2 \geq \frac{1}{2}((d_{\mathcal{M}}(x_t, x^*))^2 + (d_{\mathcal{N}}(y_t, y^*))^2) - (d_{\mathcal{M}}(\hat{x}_t, x_t))^2 + (d_{\mathcal{N}}(\hat{y}_t, y_t))^2.$$

Plugging the above inequality into the inequality from Lemma E.1 yields that

$$\begin{aligned}&(d_{\mathcal{M}}(x_{t+1}, x^*))^2 + (d_{\mathcal{N}}(y_{t+1}, y^*))^2 \\ &\leq \left(1 - \frac{\mu\eta}{2}\right)((d_{\mathcal{M}}(x_t, x^*))^2 + (d_{\mathcal{N}}(y_t, y^*))^2) + (4\bar{\xi}_0\eta^2\ell^2 + \mu\eta - \xi_0)((d_{\mathcal{M}}(\hat{x}_t, x_t))^2 + (d_{\mathcal{N}}(\hat{y}_t, y_t))^2).\end{aligned}$$

Since $\eta = \min\{\frac{1}{4\ell\sqrt{\tau_0}}, \frac{\xi_0}{2\mu}\}$, we have $4\bar{\xi}_0\eta^2\ell^2 + \mu\eta - \xi_0 \leq 0$. By the definition, we have $\tau_0 \geq 1$, $\kappa \geq 1$ and $\xi_0 \leq 1$. This implies that

$$1 - \frac{\mu\eta}{2} = 1 - \min\left\{\frac{1}{8\kappa\sqrt{\tau_0}}, \frac{\xi_0}{4}\right\} > 0.$$

Putting these pieces together yields that

$$\begin{aligned}(d_{\mathcal{M}}(x_T, x^*))^2 + (d_{\mathcal{N}}(y_T, y^*))^2 &\leq \left(1 - \min\left\{\frac{1}{8\kappa\sqrt{\tau_0}}, \frac{\xi_0}{4}\right\}\right)^T (d_{\mathcal{M}}(x_0, x^*))^2 + (d_{\mathcal{N}}(y_0, y^*))^2 \\ &\leq \left(1 - \min\left\{\frac{1}{8\kappa\sqrt{\tau_0}}, \frac{\xi_0}{4}\right\}\right)^T D_0.\end{aligned}$$

This completes the proof.

E.3 Proof of Theorem 3.2

Since Riemannian metrics satisfy the triangle inequality, we have

$$(d_{\mathcal{M}}(\hat{x}_t, x^*))^2 + (d_{\mathcal{N}}(\hat{y}_t, y^*))^2 \geq \frac{1}{2}((d_{\mathcal{M}}(x_t, x^*))^2 + (d_{\mathcal{N}}(y_t, y^*))^2) - (d_{\mathcal{M}}(\hat{x}_t, x_t))^2 + (d_{\mathcal{N}}(\hat{y}_t, y_t))^2.$$

Plugging the above inequality into the inequality from Lemma E.2 yields that

$$\begin{aligned} \mathbb{E}[f(\hat{x}_t, y^*) - f(x^*, \hat{y}_t)] &\leq \frac{1}{2\eta} \mathbb{E}[(d_{\mathcal{M}}(x_t, x^*))^2 - (d_{\mathcal{M}}(x_{t+1}, x^*))^2 + (d_{\mathcal{N}}(y_t, y^*))^2 - (d_{\mathcal{N}}(y_{t+1}, y^*))^2] \\ &\quad + (6\bar{\xi}_0\eta\ell^2 + \frac{\mu}{2} - \frac{1}{2\eta}\bar{\xi}_0) \mathbb{E}[(d_{\mathcal{M}}(\hat{x}_t, x_t))^2 + (d_{\mathcal{N}}(\hat{y}_t, y_t))^2] - \frac{\mu}{4} \mathbb{E}[(d_{\mathcal{M}}(\hat{x}_t, x^*))^2 + (d_{\mathcal{N}}(\hat{y}_t, y^*))^2] + 3\bar{\xi}_0\eta\sigma^2. \end{aligned}$$

Since $(x^*, y^*) \in \mathcal{M} \times \mathcal{N}$ is a global saddle point of f , we have $\mathbb{E}[f(\hat{x}_t, y^*) - f(x^*, \hat{y}_t)] \geq 0$. Then, we have

$$\begin{aligned} &\mathbb{E}[(d_{\mathcal{M}}(x_{t+1}, x^*))^2 + (d_{\mathcal{N}}(y_{t+1}, y^*))^2] \\ &\leq (1 - \frac{\mu\eta}{2}) \mathbb{E}[(d_{\mathcal{M}}(x_t, x^*))^2 + (d_{\mathcal{N}}(y_t, y^*))^2] + (12\bar{\xi}_0\eta^2\ell^2 + \mu\eta - \bar{\xi}_0) \mathbb{E}[(d_{\mathcal{M}}(\hat{x}_t, x_t))^2 + (d_{\mathcal{N}}(\hat{y}_t, y_t))^2] \\ &\quad + 6\bar{\xi}_0\eta^2\sigma^2. \end{aligned}$$

Since $\eta \leq \min\{\frac{1}{24\ell\sqrt{\tau_0}}, \frac{\bar{\xi}_0}{2\mu}\}$, we have $12\bar{\xi}_0\eta^2\ell^2 + \mu\eta - \bar{\xi}_0 \leq 0$. This implies that

$$\mathbb{E}[(d_{\mathcal{M}}(x_{t+1}, x^*))^2 + (d_{\mathcal{N}}(y_{t+1}, y^*))^2] \leq (1 - \frac{\mu\eta}{2}) \mathbb{E}[(d_{\mathcal{M}}(x_t, x^*))^2 + (d_{\mathcal{N}}(y_t, y^*))^2] + 6\bar{\xi}_0\eta^2\sigma^2.$$

By the definition, we have $\tau_0 \geq 1$, $\kappa \geq 1$ and $\bar{\xi}_0 \leq 1$. This implies that

$$1 - \frac{\mu\eta}{2} \geq 1 - \min\left\{\frac{1}{48\kappa\sqrt{\tau_0}}, \frac{\bar{\xi}_0}{4}\right\} > 0.$$

By the inductive arguments, we have

$$\begin{aligned} \mathbb{E}[(d_{\mathcal{M}}(x_T, x^*))^2 + (d_{\mathcal{N}}(y_T, y^*))^2] &\leq (1 - \frac{\mu\eta}{2})^T ((d_{\mathcal{M}}(x_0, x^*))^2 + (d_{\mathcal{N}}(y_0, y^*))^2) + 6\bar{\xi}_0\eta^2\sigma^2 \left(\sum_{t=0}^{T-1} (1 - \frac{\mu\eta}{2})^t\right) \\ &\leq (1 - \frac{\mu\eta}{2})^T D_0 + \frac{12\bar{\xi}_0\eta\sigma^2}{\mu}. \end{aligned}$$

Since $\eta = \min\{\frac{1}{24\ell\sqrt{\tau_0}}, \frac{\bar{\xi}_0}{2\mu}, \frac{2(\log(T) + \log(\mu^2 D_0 \sigma^{-2}))}{\mu T}\}$, we have

$$\begin{aligned} (1 - \frac{\mu\eta}{2})^T D_0 &\leq \left(1 - \min\left\{\frac{1}{48\kappa\sqrt{\tau_0}}, \frac{\bar{\xi}_0}{4}\right\}\right)^T D_0 + \left(1 - \frac{\log(\mu^2 D_0 \sigma^{-2} T)}{T}\right)^T D_0 \\ &\stackrel{1+x \leq e^x}{\leq} \left(1 - \min\left\{\frac{1}{48\kappa\sqrt{\tau_0}}, \frac{\bar{\xi}_0}{4}\right\}\right)^T D_0 + \frac{\sigma^2}{\mu^2 T}, \end{aligned}$$

and

$$\frac{12\bar{\xi}_0\eta\sigma^2}{\mu} \leq \frac{24\bar{\xi}_0\sigma^2}{\mu^2 T} \log\left(\frac{\mu^2 D_0 T}{\sigma^2}\right).$$

Putting these pieces together yields that

$$\mathbb{E}[(d_{\mathcal{M}}(x_T, x^*))^2 + (d_{\mathcal{N}}(y_T, y^*))^2] \leq \left(1 - \min\left\{\frac{1}{48\kappa\sqrt{\tau_0}}, \frac{\bar{\xi}_0}{4}\right\}\right)^T D_0 + \frac{\sigma^2}{\mu^2 T} + \frac{24\bar{\xi}_0\sigma^2}{\mu^2 T} \log\left(\frac{\mu^2 D_0 T}{\sigma^2}\right).$$

This completes the proof.

E.4 Proof of Theorem 3.3

By the inductive formulas of $\bar{x}_{t+1} = \text{Exp}_{\bar{x}_t}(\frac{1}{t+1} \cdot \text{Exp}_{\bar{x}_t}^{-1}(\hat{x}_t))$ and $\bar{y}_{t+1} = \text{Exp}_{\bar{y}_t}(\frac{1}{t+1} \cdot \text{Exp}_{\bar{y}_t}^{-1}(\hat{y}_t))$ and using Zhang et al. [2022, Lemma C.2], we have

$$f(\bar{x}_T, y^*) - f(x^*, \bar{y}_T) \leq \frac{1}{T} \left(\sum_{t=0}^{T-1} f(\hat{x}_t, y^*) - f(x^*, \hat{y}_t) \right).$$

Plugging the above inequality into the inequality from Lemma E.2 yields that (recall that $\mu = 0$ in geodesically convex-concave setting here)

$$\begin{aligned} \mathbb{E}[f(\bar{x}_T, y^*) - f(x^*, \bar{y}_T)] &\leq \frac{1}{2\eta T} ((d_{\mathcal{M}}(x_0, x^*))^2 + (d_{\mathcal{N}}(y_0, y^*))^2) \\ &\quad + \frac{1}{T} \left(6\bar{\xi}_0\eta\ell^2 - \frac{1}{2\eta}\bar{\xi}_0 \right) \left(\sum_{t=0}^{T-1} \mathbb{E}[(d_{\mathcal{M}}(\hat{x}_t, x_t))^2 + (d_{\mathcal{N}}(\hat{y}_t, y_t))^2] \right) + 3\bar{\xi}_0\eta\sigma^2. \end{aligned}$$

Since $\eta \leq \frac{1}{4\ell\sqrt{\tau_0}}$, we have $6\bar{\xi}_0\eta\ell^2 - \frac{1}{2\eta}\bar{\xi}_0 \leq 0$. Then, this together with $(d_{\mathcal{M}}(x_0, x^*))^2 + (d_{\mathcal{N}}(y_0, y^*))^2 \leq D_0$ implies that

$$\mathbb{E}[f(\bar{x}_T, y^*) - f(x^*, \bar{y}_T)] \leq \frac{D_0}{2\eta T} + 3\bar{\xi}_0\eta\sigma^2.$$

Since $\eta = \min\{\frac{1}{4\ell\sqrt{\tau_0}}, \frac{1}{\sigma}\sqrt{\frac{D_0}{\bar{\xi}_0 T}}\}$, we have

$$\frac{D_0}{2\eta T} \leq \frac{2\ell D_0\sqrt{\tau_0}}{T} + \frac{\sigma}{2}\sqrt{\frac{\bar{\xi}_0 D_0}{T}},$$

and

$$3\bar{\xi}_0\eta\sigma^2 \leq 3\sigma\sqrt{\frac{\bar{\xi}_0 D_0}{T}}.$$

Putting these pieces together yields that

$$\mathbb{E}[f(\bar{x}_T, y^*) - f(x^*, \bar{y}_T)] \leq \frac{2\ell D_0\sqrt{\tau_0}}{T} + \frac{7\sigma}{2}\sqrt{\frac{\bar{\xi}_0 D_0}{T}}.$$

This completes the proof.

F Missing Proofs for Riemannian Gradient Descent Ascent

In this section, we present some technical lemmas for analyzing the convergence property of Algorithm 3 and 4. We also give the proofs of Theorem D.1, D.2, D.3 and D.4.

F.1 Technical lemmas

We provide two technical lemmas for analyzing Algorithm 3 and 4 respectively. The first lemma gives a key inequality that is satisfied by the iterates generated by Algorithm 3.

Lemma F.1 *Under Assumption D.1 (or Assumption D.2) and let $\{(x_t, y_t)\}_{t=0}^{T-1}$ be generated by Algorithm 3 with the stepsize $\eta_t > 0$. Then, we have*

$$\begin{aligned} f(x_t, y^*) - f(x^*, y_t) &\leq \frac{1}{2\eta_t} ((d_{\mathcal{M}}(x_t, x^*))^2 - (d_{\mathcal{M}}(x_{t+1}, x^*))^2) \\ &\quad + \frac{1}{2\eta_t} ((d_{\mathcal{N}}(y_t, y^*))^2 - (d_{\mathcal{N}}(y_{t+1}, y^*))^2) - \frac{\mu}{2}(d_{\mathcal{M}}(x_t, x^*))^2 - \frac{\mu}{2}(d_{\mathcal{N}}(y_t, y^*))^2 + \bar{\xi}_0\eta_t L^2, \end{aligned}$$

where $(x^*, y^*) \in \mathcal{M} \times \mathcal{N}$ is a global saddle point of f .

Proof. Since f is geodesically strongly-concave-strongly-concave with the modulus $\mu \geq 0$ (here $\mu = 0$ means that f is geodesically concave-concave), we have

$$\begin{aligned} f(x_t, y^*) - f(x^*, y_t) &= f(x_t, y_t) - f(x^*, y_t) - (f(x_t, y_t) - f(x_t, y^*)) \\ &\leq -\langle \text{subgrad}_x f(x_t, y_t), \text{Exp}_{x_t}^{-1}(x^*) \rangle + \langle \text{subgrad}_y f(x_t, y_t), \text{Exp}_{y_t}^{-1}(y^*) \rangle - \frac{\mu}{2}(d_{\mathcal{M}}(x_t, x^*))^2 - \frac{\mu}{2}(d_{\mathcal{N}}(y_t, y^*))^2. \end{aligned}$$

Recalling also from the scheme of Algorithm 3 that we have

$$\begin{aligned} x_{t+1} &\leftarrow \text{Exp}_{x_t}(-\eta_t \cdot \text{subgrad}_x f(x_t, y_t)), \\ y_{t+1} &\leftarrow \text{Exp}_{y_t}(\eta_t \cdot \text{subgrad}_y f(x_t, y_t)). \end{aligned}$$

By the definition of an exponential map, we have

$$\begin{aligned} \text{Exp}_{x_t}^{-1}(x_{t+1}) &= -\eta_t \cdot \text{subgrad}_x f(x_t, y_t), \\ \text{Exp}_{y_t}^{-1}(y_{t+1}) &= \eta_t \cdot \text{subgrad}_y f(x_t, y_t). \end{aligned} \tag{F.1}$$

This implies that

$$\begin{aligned} -\langle \text{subgrad}_x f(x_t, y_t), \text{Exp}_{x_t}^{-1}(x^*) \rangle &= \frac{1}{\eta_t} \langle \text{Exp}_{x_t}^{-1}(x_{t+1}), \text{Exp}_{x_t}^{-1}(x^*) \rangle, \\ \langle \text{subgrad}_y f(x_t, y_t), \text{Exp}_{y_t}^{-1}(y^*) \rangle &= \frac{1}{\eta_t} \langle \text{Exp}_{y_t}^{-1}(y_{t+1}), \text{Exp}_{y_t}^{-1}(y^*) \rangle. \end{aligned}$$

Putting these pieces together yields that

$$\begin{aligned} f(x_t, y^*) - f(x^*, y_t) &\leq \frac{1}{\eta_t} \langle \text{Exp}_{x_t}^{-1}(x_{t+1}), \text{Exp}_{x_t}^{-1}(x^*) \rangle \\ &\quad + \frac{1}{\eta_t} \langle \text{Exp}_{y_t}^{-1}(y_{t+1}), \text{Exp}_{y_t}^{-1}(y^*) \rangle - \frac{\mu}{2}(d_{\mathcal{M}}(x_t, x^*))^2 - \frac{\mu}{2}(d_{\mathcal{N}}(y_t, y^*))^2. \end{aligned} \tag{F.2}$$

It suffices to bound the terms in the right-hand side of Eq. (F.2) by leveraging the celebrated comparison inequalities on Riemannian manifold with lower bounded sectional curvature (see Proposition C.2). More specifically, we define the constants using $\bar{\xi}(\cdot, \cdot)$ and $\underline{\xi}(\cdot, \cdot)$ from Proposition C.2 as follows,

$$\bar{\xi}_0 = \bar{\xi}(\kappa_{\min}, D).$$

By Proposition C.2 and using that $\max\{d_{\mathcal{M}}(x_t, x^*), d_{\mathcal{N}}(y_t, y^*)\} \leq D$, we have

$$\begin{aligned} \langle \text{Exp}_{x_t}^{-1}(x_{t+1}), \text{Exp}_{x_t}^{-1}(x^*) \rangle &\leq \frac{1}{2} (\bar{\xi}_0 (d_{\mathcal{M}}(x_t, x_{t+1}))^2 + (d_{\mathcal{M}}(x_t, x^*))^2 - (d_{\mathcal{M}}(x_{t+1}, x^*))^2), \\ \langle \text{Exp}_{y_t}^{-1}(y_{t+1}), \text{Exp}_{y_t}^{-1}(y^*) \rangle &\leq \frac{1}{2} (\bar{\xi}_0 (d_{\mathcal{N}}(y_t, y_{t+1}))^2 + (d_{\mathcal{N}}(y_t, y^*))^2 - (d_{\mathcal{N}}(y_{t+1}, y^*))^2). \end{aligned}$$

Since f is geodesically L -Lipschitz, we have

$$\|\text{subgrad}_x f(x_t, y_t)\| \leq L, \quad \|\text{subgrad}_y f(x_t, y_t)\| \leq L.$$

By the definition of an exponential map and Riemannian metric, we have

$$\begin{aligned} d_{\mathcal{M}}(x_t, x_{t+1}) &= \|\text{Exp}_{x_t}^{-1}(x_{t+1})\| \stackrel{\text{Eq. (F.1)}}{=} \|\eta_t \cdot \text{subgrad}_x f(x_t, y_t)\| \leq \eta_t L, \\ d_{\mathcal{N}}(y_t, y_{t+1}) &= \|\text{Exp}_{y_t}^{-1}(y_{t+1})\| \stackrel{\text{Eq. (F.1)}}{=} \|\eta_t \cdot \text{subgrad}_y f(x_t, y_t)\| \leq \eta_t L. \end{aligned}$$

Putting these pieces together yields that

$$\begin{aligned} \langle \text{Exp}_{x_t}^{-1}(x_{t+1}), \text{Exp}_{x_t}^{-1}(x^*) \rangle &\leq \frac{1}{2} (\bar{\xi}_0 \eta_t^2 L^2 + (d_{\mathcal{M}}(x_t, x^*))^2 - (d_{\mathcal{M}}(x_{t+1}, x^*))^2), \\ \langle \text{Exp}_{y_t}^{-1}(y_{t+1}), \text{Exp}_{y_t}^{-1}(y^*) \rangle &\leq \frac{1}{2} (\bar{\xi}_0 \eta_t^2 L^2 + (d_{\mathcal{N}}(y_t, y^*))^2 - (d_{\mathcal{N}}(y_{t+1}, y^*))^2). \end{aligned}$$

Plugging the above inequalities into Eq. (F.2) yields the desired inequality. $\square$

The second lemma gives another key inequality that is satisfied by the iterates generated by Algorithm 4.

Lemma F.2 Under Assumption D.1 (or Assumption D.2) and the noisy model (cf. Eq. (D.2) and (D.3)) and let $\{(x_t, y_t)\}_{t=0}^{T-1}$ be generated by Algorithm 4 with the stepsize $\eta_t > 0$. Then, we have

$$\begin{aligned} \mathbb{E}[f(x_t, y^*) - f(x^*, y_t)] &\leq \frac{1}{2\eta_t} \mathbb{E}[(d_{\mathcal{M}}(x_t, x^*))^2 - (d_{\mathcal{M}}(x_{t+1}, x^*))^2] \\ &\quad + \frac{1}{2\eta_t} \mathbb{E}[(d_{\mathcal{N}}(y_t, y^*))^2 - (d_{\mathcal{N}}(y_{t+1}, y^*))^2] - \frac{\mu}{2} \mathbb{E}[(d_{\mathcal{M}}(x_t, x^*))^2 + (d_{\mathcal{N}}(y_t, y^*))^2] + 2\bar{\xi}_0 \eta_t (L^2 + \sigma^2), \end{aligned}$$

where $(x^*, y^*) \in \mathcal{M} \times \mathcal{N}$ is a global saddle point of f .

Proof. Using the same argument, we have ($\mu = 0$ refers to geodesically convex-concave case)

$$\begin{aligned} f(x_t, y^*) - f(x^*, y_t) &= f(x_t, y_t) - f(x^*, y_t) - (f(x_t, y_t) - f(x_t, y^*)) \\ &\leq -\langle \text{subgrad}_x f(x_t, y_t), \text{Exp}_{x_t}^{-1}(x^*) \rangle + \langle \text{subgrad}_y f(x_t, y_t), \text{Exp}_{y_t}^{-1}(y^*) \rangle - \frac{\mu}{2} (d_{\mathcal{M}}(x_t, x^*))^2 - \frac{\mu}{2} (d_{\mathcal{N}}(y_t, y^*))^2. \end{aligned}$$

Combining the arguments used in Lemma F.1 and the scheme of Algorithm 2, we have

$$\begin{aligned} -\langle g_x^t, \text{Exp}_{x_t}^{-1}(x^*) \rangle &= \frac{1}{\eta_t} \langle \text{Exp}_{x_t}^{-1}(x_{t+1}), \text{Exp}_{x_t}^{-1}(x^*) \rangle, \\ \langle g_y^t, \text{Exp}_{y_t}^{-1}(y^*) \rangle &= \frac{1}{\eta_t} \langle \text{Exp}_{y_t}^{-1}(y_{t+1}), \text{Exp}_{y_t}^{-1}(y^*) \rangle. \end{aligned}$$

Putting these pieces together with Eq. (D.2) yields that

$$\begin{aligned} f(x_t, y^*) - f(x^*, y_t) &\leq \frac{1}{\eta_t} \langle \text{Exp}_{x_t}^{-1}(x_{t+1}), \text{Exp}_{x_t}^{-1}(x^*) \rangle \\ &\quad + \frac{1}{\eta_t} \langle \text{Exp}_{y_t}^{-1}(y_{t+1}), \text{Exp}_{y_t}^{-1}(y^*) \rangle - \frac{\mu}{2} (d_{\mathcal{M}}(x_t, x^*))^2 - \frac{\mu}{2} (d_{\mathcal{N}}(y_t, y^*))^2 + \langle \xi_x^t, \text{Exp}_{x_t}^{-1}(x^*) \rangle - \langle \xi_y^t, \text{Exp}_{y_t}^{-1}(y^*) \rangle. \end{aligned} \tag{F.3}$$

By the same argument as used in Lemma F.1 and Eq. (D.2), we have

$$\begin{aligned} \langle \text{Exp}_{x_t}^{-1}(x_{t+1}), \text{Exp}_{x_t}^{-1}(x^*) \rangle &\leq \frac{1}{2} (\bar{\xi}_0 (d_{\mathcal{M}}(x_t, x_{t+1}))^2 + (d_{\mathcal{M}}(x_t, x^*))^2 - (d_{\mathcal{M}}(x_{t+1}, x^*))^2), \\ \langle \text{Exp}_{y_t}^{-1}(y_{t+1}), \text{Exp}_{y_t}^{-1}(y^*) \rangle &\leq \frac{1}{2} (\bar{\xi}_0 (d_{\mathcal{N}}(y_t, y_{t+1}))^2 + (d_{\mathcal{N}}(y_t, y^*))^2 - (d_{\mathcal{N}}(y_{t+1}, y^*))^2), \end{aligned}$$

and

$$\begin{aligned} d_{\mathcal{M}}(x_t, x_{t+1}) &= \|\text{Exp}_{x_t}^{-1}(x_{t+1})\| = \|\eta_t \cdot g_x^t\| \leq \eta_t (L + \|\xi_x^t\|), \\ d_{\mathcal{N}}(y_t, y_{t+1}) &= \|\text{Exp}_{y_t}^{-1}(y_{t+1})\| = \|\eta_t \cdot g_y^t\| \leq \eta_t (L + \|\xi_y^t\|). \end{aligned}$$

Therefore, we have

$$\begin{aligned} &\langle \text{Exp}_{x_t}^{-1}(x_{t+1}), \text{Exp}_{x_t}^{-1}(x^*) \rangle + \langle \text{Exp}_{y_t}^{-1}(y_{t+1}), \text{Exp}_{y_t}^{-1}(y^*) \rangle \\ &\leq \frac{1}{2} \bar{\xi}_0 \eta_t^2 (4L^2 + 2\|\xi_x^t\|^2 + 2\|\xi_y^t\|^2) + \frac{1}{2} ((d_{\mathcal{M}}(x_t, x^*))^2 - (d_{\mathcal{M}}(x_{t+1}, x^*))^2 + (d_{\mathcal{N}}(y_t, y^*))^2 - (d_{\mathcal{N}}(y_{t+1}, y^*))^2). \end{aligned}$$

Plugging the above inequalities into Eq. (F.3) yields that

$$\begin{aligned} f(x_t, y^*) - f(x^*, y_t) &\leq \frac{1}{2\eta_t} ((d_{\mathcal{M}}(x_t, x^*))^2 - (d_{\mathcal{M}}(x_{t+1}, x^*))^2 + (d_{\mathcal{N}}(y_t, y^*))^2 - (d_{\mathcal{N}}(y_{t+1}, y^*))^2) \\ &\quad + \bar{\xi}_0 \eta_t (2L^2 + \|\xi_x^t\|^2 + \|\xi_y^t\|^2) - \frac{\mu}{2} (d_{\mathcal{M}}(x_t, x^*))^2 - \frac{\mu}{2} (d_{\mathcal{N}}(y_t, y^*))^2 + \langle \xi_x^t, \text{Exp}_{x_t}^{-1}(x^*) \rangle - \langle \xi_y^t, \text{Exp}_{y_t}^{-1}(y^*) \rangle. \end{aligned}$$

Taking the expectation of both sides and using Eq. (D.3) yields the desired inequality. $\square$

F.2 Proof of Theorem D.1

Since $(x^*, y^*) \in \mathcal{M} \times \mathcal{N}$ is a global saddle point of f , we have $f(x_t, y^*) - f(x^*, y_t) \geq 0$. Plugging this inequality into the inequality from Lemma F.1 yields that

$$(d_{\mathcal{M}}(x_{t+1}, x^*))^2 + (d_{\mathcal{N}}(y_{t+1}, y^*))^2 \leq (1 - \mu\eta_t) ((d_{\mathcal{M}}(x_t, x^*))^2 + (d_{\mathcal{N}}(y_t, y^*))^2) + 2\bar{\xi}_0\eta_t^2 L^2.$$

Since $\eta_t = \frac{1}{\mu} \min\{1, \frac{2}{t}\}$, we have

$$(d_{\mathcal{M}}(x_{t+1}, x^*))^2 + (d_{\mathcal{N}}(y_{t+1}, y^*))^2 \leq (1 - \frac{2}{t}) ((d_{\mathcal{M}}(x_t, x^*))^2 + (d_{\mathcal{N}}(y_t, y^*))^2) + \frac{8\bar{\xi}_0 L^2}{\mu^2 t^2}, \quad \text{for all } t \geq 2.$$

Letting $\{b_t\}_{t \geq 1}$ be a nonnegative sequence such that $a_{t+1} \leq (1 - \frac{P}{t})a_t + \frac{Q}{t^2}$ where $P > 1$ and $Q > 0$. Then, Chung [1954] proved that $a_t \leq \frac{Q}{P-1} \frac{1}{t}$. Therefore, we have

$$(d_{\mathcal{M}}(x_t, x^*))^2 + (d_{\mathcal{N}}(y_t, y^*))^2 \leq \frac{8\bar{\xi}_0 L^2}{\mu^2 t}, \quad \text{for all } t \geq 2.$$

This completes the proof.

F.3 Proof of Theorem D.2

By the inductive formulas of $\bar{x}_{t+1} = \text{Exp}_{\bar{x}_t}(\frac{1}{t+1} \cdot \text{Exp}_{\bar{x}_t}^{-1}(x_t))$ and $\bar{y}_{t+1} = \text{Exp}_{\bar{y}_t}(\frac{1}{t+1} \cdot \text{Exp}_{\bar{y}_t}^{-1}(y_t))$ and using Zhang et al. [2022, Lemma C.2], we have

$$f(\bar{x}_T, y^*) - f(x^*, \bar{y}_T) \leq \frac{1}{T} \left(\sum_{t=0}^{T-1} f(x_t, y^*) - f(x^*, y_t) \right).$$

Plugging the above inequality into the inequality from Lemma F.1 yields that (recall that $\mu = 0$ in geodesically convex-concave setting and $\eta_t = \eta = \frac{1}{L} \sqrt{\frac{D_0}{2\bar{\xi}_0 T}}$)

$$f(\bar{x}_T, y^*) - f(x^*, \bar{y}_T) \leq \frac{1}{2\eta T} ((d_{\mathcal{M}}(x_0, x^*))^2 + (d_{\mathcal{N}}(y_0, y^*))^2) + \bar{\xi}_0 \eta L^2.$$

This together with $(d_{\mathcal{M}}(x_0, x^*))^2 + (d_{\mathcal{N}}(y_0, y^*))^2 \leq D_0$ implies that

$$f(\bar{x}_T, y^*) - f(x^*, \bar{y}_T) \leq \frac{D_0}{2\eta T} + \bar{\xi}_0 \eta L^2.$$

Since $\eta = \frac{1}{L} \sqrt{\frac{D_0}{2\bar{\xi}_0 T}}$, we have

$$f(\bar{x}_T, y^*) - f(x^*, \bar{y}_T) \leq L \sqrt{\frac{2\bar{\xi}_0 D_0}{T}}.$$

This completes the proof.

F.4 Proof of Theorem D.3

Since $(x^*, y^*) \in \mathcal{M} \times \mathcal{N}$ is a global saddle point of f , we have $\mathbb{E}[f(x_t, y^*) - f(x^*, y_t)] \geq 0$. Plugging this inequality into the inequality from Lemma F.2 yields that

$$\mathbb{E} [(d_{\mathcal{M}}(x_{t+1}, x^*))^2 + (d_{\mathcal{N}}(y_{t+1}, y^*))^2] \leq (1 - \mu\eta_t) \mathbb{E} [(d_{\mathcal{M}}(x_t, x^*))^2 + (d_{\mathcal{N}}(y_t, y^*))^2] + 4\bar{\xi}_0\eta_t^2 (L^2 + \sigma^2).$$

Since $\eta_t = \frac{1}{\mu} \min\{1, \frac{2}{t}\}$, we have

$$\mathbb{E} [(d_{\mathcal{M}}(x_{t+1}, x^*))^2 + (d_{\mathcal{N}}(y_{t+1}, y^*))^2] \leq (1 - \frac{2}{t}) \mathbb{E} [(d_{\mathcal{M}}(x_t, x^*))^2 + (d_{\mathcal{N}}(y_t, y^*))^2] + \frac{16\bar{\xi}_0(L^2 + \sigma^2)}{\mu^2 t^2}, \quad \text{for all } t \geq 2.$$

Applying the same argument as used in Theorem D.1, we have

$$(d_{\mathcal{M}}(x_t, x^*))^2 + (d_{\mathcal{N}}(y_t, y^*))^2 \leq \frac{16\bar{\xi}_0(L^2 + \sigma^2)}{\mu^2 t}, \quad \text{for all } t \geq 2.$$

This completes the proof.

F.5 Proof of Theorem D.4

Using the same argument, we have

$$f(\bar{x}_T, y^*) - f(x^*, \bar{y}_T) \leq \frac{1}{T} \left(\sum_{t=0}^{T-1} f(x_t, y^*) - f(x^*, y_t) \right).$$

Plugging the above inequality into the inequality from Lemma F.2 yields that (recall that $\mu = 0$ in geodesically convex-concave setting and $\eta_t = \eta = \frac{1}{2} \sqrt{\frac{D_0}{\bar{\xi}_0(L^2 + \sigma^2)T}}$)

$$\mathbb{E}[f(\bar{x}_T, y^*) - f(x^*, \bar{y}_T)] \leq \frac{1}{2\eta T} ((d_{\mathcal{M}}(x_0, x^*))^2 + (d_{\mathcal{N}}(y_0, y^*))^2) + 2\bar{\xi}_0\eta(L^2 + \sigma^2).$$

This together with $(d_{\mathcal{M}}(x_0, x^*))^2 + (d_{\mathcal{N}}(y_0, y^*))^2 \leq D_0$ implies that

$$\mathbb{E}[f(\bar{x}_T, y^*) - f(x^*, \bar{y}_T)] \leq \frac{D_0}{2\eta T} + 2\bar{\xi}_0\eta(L^2 + \sigma^2).$$

Since $\eta = \frac{1}{2} \sqrt{\frac{D_0}{\bar{\xi}_0(L^2 + \sigma^2)T}}$, we have

$$f(\bar{x}_T, y^*) - f(x^*, \bar{y}_T) \leq 2\sqrt{\frac{\bar{\xi}_0(L^2 + \sigma^2)D_0}{T}}.$$

This completes the proof.

G Additional Experimental Results

We present some additional experimental results for the effect of different choices of α as well the effect of different choices of η for RCEG. In our experiment here, we set $n = 40$ consistently.

Figure 3 presents the performance of RCEG when $\alpha = 2.0$. We observe that the results are similar to that summarized in Figure 1. In particular, the last iterate of RCEG consistently achieves the linearly convergence to an optimal solution in all the settings. In contrast, the average iterate of RCEG converges much slower than the last iterate of RCEG. Figure 4 summarizes the effect of different choices of η in RCEG. We observe that setting η as a relatively larger value will speed up the convergence to an optimal solution while all of the choices here lead to the linear convergence. This suggests that the choice of stepsize η in RCEG can be aggressive in practice.

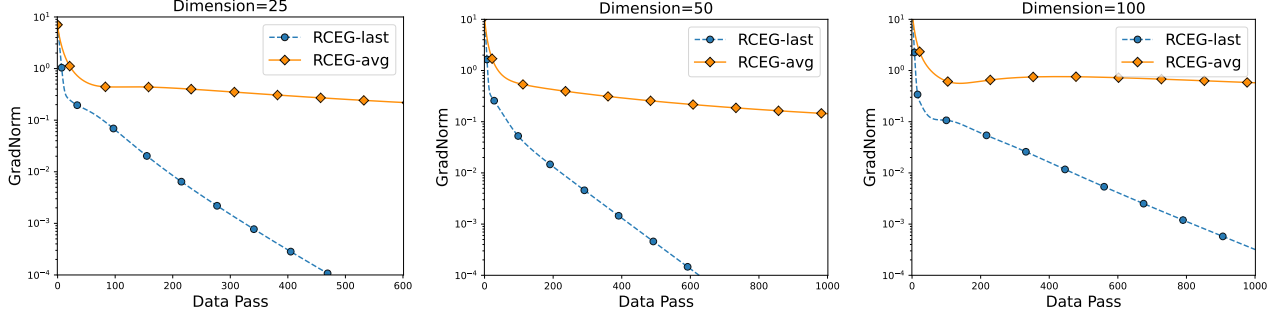


Figure 3: Comparison of last iterate (RCEG-last) and time-average iterate (RCEG-avg) for solving the RPCA problem when $\alpha = 2.0$. The horizontal axis represents the number of data passes and the vertical axis represents gradient norm.

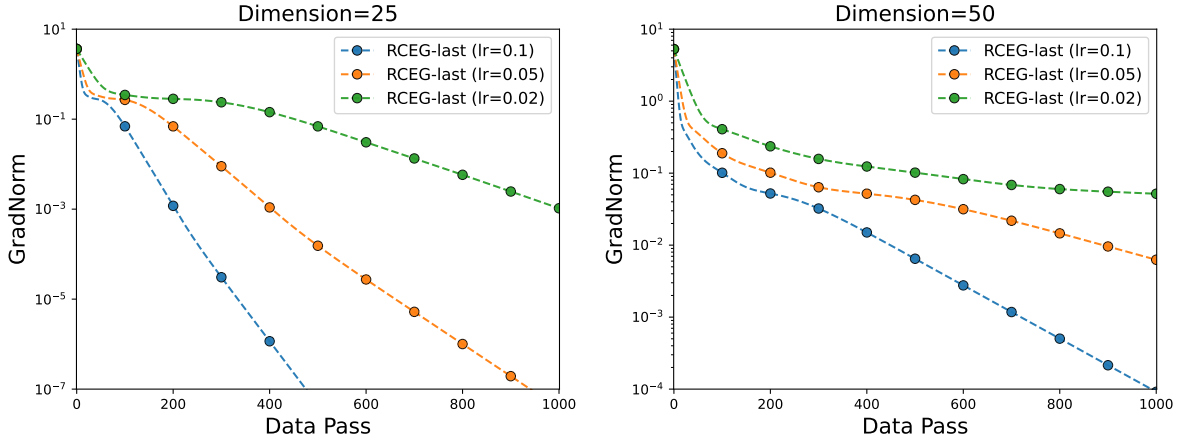


Figure 4: Comparison of different step sizes ($\eta \in \{0.1, 0.05, 0.02\}$) for solving the RPCA problem with different dimensions when $\alpha = 2.0$. The horizontal axis represents the number of data passes and the vertical axis represents gradient norm.