
Alternating Mirror Descent for Constrained Min-Max Games

Andre Wibisono

Department of Computer Science
Yale University
andre.wibisono@yale.edu

Molei Tao

Department of Mathematics
Georgia Institute of Technology
mtao@gatech.edu

Georgios Piliouras

Engineering Systems and Design
Singapore University of Technology and Design
georgios@sutd.edu.sg

Abstract

In this paper we study two-player bilinear zero-sum games with constrained strategy spaces. An instance of natural occurrences of such constraints is when mixed strategies are used, which correspond to a probability simplex constraint. We propose and analyze the alternating mirror descent algorithm, in which each player takes turns to take action following the mirror descent algorithm for constrained optimization. We interpret alternating mirror descent as an alternating discretization of a skew-gradient flow in the dual space, and use tools from convex optimization and modified energy function to establish an $O(K^{-2/3})$ bound on its average regret after K iterations. This quantitatively verifies the algorithm's better behavior than the simultaneous version of mirror descent algorithm, which is known to diverge and yields an $O(K^{-1/2})$ average regret bound. In the special case of an unconstrained setting, our results recover the behavior of alternating gradient descent algorithm for zero-sum games which was studied in [2].

1 Introduction

Multi-agent systems are ubiquitous in many applications and have gained increasing importance in practice, from the classical problems in economics and game theory, to the modern applications in machine learning, in particular via the emergence of learning strategies that can be formulated as min-max games, such as the generative adversarial networks (GANs), robust optimization, reinforcement learning, and many others [11, 16, 23, 27]. In multi-agent systems, interaction between the agents can exhibit nontrivial global behavior, even when each agent individually follows a simple action such as a greedy, optimization-driven strategy [7, 8, 1]. This emergence of complex behavior has been recognized as one source of difficulties in understanding and controlling the global behavior of multi-agent game dynamics [24, 15].

Even in the basic setting of unconstrained two-player zero-sum game with bilinear payoffs, this emergence of non-trivial behavior already presents some difficulties. It is now well-known that if each player follows a classical greedy strategy, such as gradient descent, and if they make their actions *simultaneously*, then their joint trajectories diverge away from equilibrium and lead to increasing regret; however, the average iterates still converge to the equilibrium and yield a vanishing average regret with decreasing step size [6, 3, 7]. This behavior challenges our intuition and is in marked contrast to the case of single-agent optimization, in which greedy strategies are guaranteed to converge. This leads to variations of the greedy strategy to correct the diverging behavior and help

the trajectories to converge to equilibrium, for example via optimistic or extragradient versions of gradient descent [10, 17], which can be seen as approximations of the proximal (implicit) gradient descent [19].

Another variation of the basic greedy strategy is when each player follows gradient descent, but they make their actions in an *alternating* fashion (i.e. one at a time, rather than simultaneously), as studied in [2]. This technique is particularly useful for machine learning applications where the state of the system can be very large with several billions of parameters as one does not need extra memory to store intermediate variables, required both by simultaneous updates as well as extra-gradient methods. The resulting *alternating gradient descent* algorithm has a markedly different behavior than in the simultaneous case. As shown in [2], the trajectories of alternating gradient descent turn out to *cycle* (stay in a bounded orbit), rather than converging to or diverging away from equilibrium. This behavior mimics the ideal setting of continuous-time dynamics, which is the limit as the step size goes to 0, in which case the orbit of the two players cycles around the equilibrium and exactly preserves the “energy function”, which in this case is defined to be the distance to the equilibrium point, achieving average regret bound $O(T^{-1})$ after (continuous) time T [18]. In discrete time, alternating gradient descent does not exactly conserve the energy function; instead, it conserves a *modified energy function*, which is a perturbation of the true energy function with a correction term which is proportional to the step size. This implies that alternating gradient descent has a constant regret for any step size, and thus yields an average regret bound of $O(K^{-1})$ after K iterations, which matches the continuous-time behavior; see [2] for more details.

The results above raise the question of what we can say in the *constrained* setting, when each agent can choose an action from a constrained set. An instance of natural occurrences of such constraints is when mixed strategies are used, which corresponds to a probability simplex constraint. One popular strategy, inspired by optimization techniques, is that each agent now plays the *mirror descent* algorithm for constrained minimization of their own objective function. Mirror descent is closely related to the Follow the Regularized Leader (FTRL) algorithm in online learning, which has a Hamiltonian structure in continuous time [4, 13]. In the idealized continuous-time setting, if both players follow the continuous-time version of mirror descent dynamics, then their trajectories cycle around the equilibrium, and conserve an “energy function” in the dual space, which is now defined to be the sum of the dual functions of the regularizers; this leads to a constant regret, and yields an $O(T^{-1})$ average regret bound after (continuous) time T [18]. In this harder constrained setting, one would hope to weave a single thread connecting in an intuitive manner the behavior of continuous dynamics with multiple distinct discretizations. In discrete time, if the two players follow the simultaneous version of mirror descent, then their trajectories diverge from equilibrium, and yields a $O(K^{-1/2})$ bound on the average regret. If both players follow the *proximal* (implicit) mirror descent algorithm, then as we show their trajectories converge to the equilibrium, resulting in a simple $O(K^{-1})$ average regret bound, which matches the continuous-time behavior. While the $O(K^{-1})$ average regret bound is known to extend to general games [25] under clairvoyant multiplicative weights updates, which is a closely related algorithm to proximal mirror descent, such methods require coordination among the players to solve the implicit update. Nevertheless, there are simpler variations such as the optimistic or extra-gradient, which can get $\tilde{O}(K^{-1})$ convergence rates in general games [9], where $\tilde{O}$ hides polylogarithmic dependence.

In this paper, we propose and study the *alternating mirror descent* algorithm for two-player zero-sum constrained bilinear game, in which the two players take turns to make their moves following the mirror descent algorithm. We show that the total regret of the two players can be expressed in terms of a *modified energy function*, which generalizes the modified energy function in the unconstrained setting (see Theorem 4.5). Recall in the unconstrained setting, the modified energy function is exactly preserved [2]. In the constrained setting, we prove a bound on the growth of the modified energy under third-order smoothness assumption on the energy function (see Theorem 4.4). This yields an $O(K^{-2/3})$ bound on the average regret of the alternating mirror descent algorithm, which improves on the classical $O(K^{-1/2})$ average regret bound of the simultaneous mirror descent algorithm.

As an analysis tool, we study alternating mirror descent algorithm as an alternating discretization of a *skew-gradient flow*. In continuous time, we can study the continuous-time mirror descent dynamics in the dual space; this dynamics corresponds to the skew-gradient flow of the energy function, which conserves the energy and explains the cycling behavior, as we explain in Section 4. In discrete time, since the energy function is convex, the forward discretization of the skew-gradient

flow is proved to increase energy; this corresponds to the diverging behavior of the simultaneous mirror descent algorithm; see Section B.1. In contrast, the backward (implicit) discretization of the skew-gradient flow is proved to decrease energy; this corresponds to the converging behavior of the proximal/clairvoyant learning; see Section B.2. Finally, as in the unconstrained case, the alternating discretization of the flow follows the continuous-time dynamics more closely, leading to improved bounds.

Another reason that we expect the alternating mirror descent algorithm to work well is a connection to symplectic integrators. In the special case when the payoff matrix in the bilinear game is the identity matrix,¹ the continuous-time mirror descent dynamics becomes a *Hamiltonian flow*, i.e. it has a symplectic structure, and the energy function becomes the Hamiltonian function that generates the flow (hence conserved). In discrete time, the alternating mirror descent algorithm corresponds to the *symplectic Euler* discretization in the dual space, which also conserves the symplectic structure. A remarkable feature of symplectic integrators is that they exhibit good energy conservation property until exponentially long time [5, 12]. Recently, symplectic integrators as well as connections between algorithms and continuous-time dynamics, have been shown to be highly relevant in the optimization settings for design of accelerated methods [14, 30, 28, 20]. Specifically, by preserving specific continuous symmetries of the flow, symplectic integrators stabilize the dynamics and allow for larger step sizes for fixed accuracy in the long run. This intuition gives a geometric interpretation to “acceleration”. Our novel connection between game theory, alternating mirror descent dynamics and symplectic integrators opens the door for further cross-fertilization between these areas.

2 Preliminaries

In this paper we study the two-player zero-sum game with bilinear payoff:

$$\min_{p \in \mathcal{P}} \max_{q \in \mathcal{Q}} p^\top A q.$$

Here $\mathcal{P} \subseteq \mathbb{R}^m$ and $\mathcal{Q} \subseteq \mathbb{R}^n$ are closed convex sets which represent the domains of the actions that each player can make, and $A \in \mathbb{R}^{m \times n}$ is an arbitrary payoff matrix. When the first player plays an action $p \in \mathcal{P}$ and the second player plays $q \in \mathcal{Q}$, the first player receives loss $p^\top A q = \sum_{i=1}^m \sum_{j=1}^n p_i q_j A_{ij}$, which they want to minimize; while the second player receives loss $-p^\top A q$, which they want to minimize (equivalently, the second player wants to maximize their utility $p^\top A q$).

An objective of the game is to reach the *Nash equilibrium*, which is a pair $(p^*, q^*) \in \mathcal{P} \times \mathcal{Q}$ that satisfies, for all $(p, q) \in \mathcal{P} \times \mathcal{Q}$:

$$(p^*)^\top A q \leq (p^*)^\top A q^* \leq p^\top A q^*.$$

By von Neumann’s min-max theorem [29], a Nash equilibrium always exists (but is not necessarily unique) as long as $\mathcal{P}$ and $\mathcal{Q}$ are compact, which is the case here.

One way to measure convergence to equilibrium is via the *duality gap* $\text{dg}: \mathcal{P} \times \mathcal{Q} \rightarrow \mathbb{R}$ given by

$$\text{dg}(p, q) = \max_{\tilde{q} \in \mathcal{Q}} p^\top A \tilde{q} - \min_{\tilde{p} \in \mathcal{P}} \tilde{p}^\top A q.$$

One can check that $\text{dg}(p, q) \geq 0$ for all $(p, q) \in \mathcal{P} \times \mathcal{Q}$, and moreover, $\text{dg}(p^*, q^*) = 0$ if and only if (p^*, q^*) is a Nash equilibrium.

In the execution of the zero-sum game, each player follows some algorithm in discrete time (or some dynamics in continuous time). Depending on the precise specification of what algorithms they play and how they play it, the iterates of the actions (p_k, q_k) may not converge to the Nash equilibrium. However, we expect the *average iterates* $(\bar{p}_K, \bar{q}_K)$ to converge to the Nash equilibrium, as measured by the duality gap: $\text{dg}(\bar{p}_K, \bar{q}_K) \rightarrow 0$ as $K \rightarrow \infty$. Furthermore, typically the duality gap of the average iterates are related to the *average regret* of the two players. Therefore, if both players follow a *no-regret* algorithm (so the average regret vanishes asymptotically), then their average iterates converge to Nash equilibrium. Thus, we are interested in bounding the rate at which the average regret of the two players converges to 0. See Section 3.2 for more details.

¹When the payoff matrix is arbitrary, the dynamics has a non-canonical Hamiltonian structure.

Unconstrained case. A special case is the *unconstrained* setting when $\mathcal{P} = \mathbb{R}^m$ and $\mathcal{Q} = \mathbb{R}^n$. Here the equilibrium is at $(0, 0)$. A natural strategy is for each player to follow gradient descent to minimize their own loss functions. As explained in the introduction, the behaviors of the iterates can vary depending on how the two players take the actions: If they both follow simultaneous gradient descent, then the iterates diverge; if they follow simultaneous proximal gradient descent, then the iterates converge; if they follow continuous-time gradient flow, then the iterates cycle; finally, if they follow alternating gradient descent, then the iterates cycle and conserve a modified energy function.

Constrained case. The main question we address in this paper is whether the different behaviors we observe in the unconstrained setting above carry over to the *constrained* setting, when $\mathcal{P} \subsetneq \mathbb{R}^m$ or $\mathcal{Q} \subsetneq \mathbb{R}^n$ (or both) are proper subsets of the Euclidean space. As an example, we may consider $\mathcal{P} = \Delta_m$ and $\mathcal{Q} = \Delta_n$, where $\Delta_m = \{x \in \mathbb{R}^m : x_i \geq 0, \sum_{i=1}^m x_i = 1\}$ is the probability simplex in $\mathbb{R}^m$, and similarly Δ_n is the probability simplex in $\mathbb{R}^n$; these correspond to the setting where each player chooses a random action among a set of discrete choices. However, our analyses and results hold for general constraint sets.

In this constrained setting, a natural strategy is for each player to follow a constrained greedy method, instead of using gradient descent. Following techniques in optimization, we consider the case when each player follows the *mirror descent* algorithm [21] to minimize their own loss functions over their constrained domains.

Mirror descent set-up. We assume that on the domain $\mathcal{P} \subseteq \mathbb{R}^m$ we have a strictly convex regularizer function $\phi: \mathcal{P} \rightarrow \mathbb{R}$ which is a *Legendre function* [26], which means ϕ is continuously differentiable, $\|\nabla \phi(p)\| \rightarrow \infty$ as p approaches the boundary $\partial \mathcal{P}$ of $\mathcal{P}$, and $\nabla \phi$ is a bijection from $\mathcal{P}$ to the range $\nabla \phi(\mathcal{P}) \subseteq \mathbb{R}^m$. Here the gradient $\nabla \phi(p) \in \mathbb{R}^m$ is the vector of partial derivatives.

Let $D_\phi: \mathcal{P} \times \mathcal{P} \rightarrow \mathbb{R}$ be the *Bregman divergence* of ϕ , defined by

$$D_\phi(p, \tilde{p}) = \phi(p) - \phi(\tilde{p}) - \langle \nabla \phi(\tilde{p}), p - \tilde{p} \rangle$$

where $\langle u, v \rangle = u^\top v$ is the ℓ_2 -inner product. Since ϕ is strictly convex, we have that $D_\phi(p, \tilde{p}) \geq 0$, and $D_\phi(p, \tilde{p}) = 0$ if and only if $p = \tilde{p}$. Note that Bregman divergence is not necessarily symmetric: in general, $D_\phi(p, \tilde{p}) \neq D_\phi(\tilde{p}, p)$. The idea of mirror descent [21] is to use the Bregman divergence $D_\phi(p, \tilde{p})$ to measure “distance” (albeit asymmetric) from $\tilde{p}$ to p on $\mathcal{P}$. For example, in the unconstrained setting when $\mathcal{P} = \mathbb{R}^m$, if we choose $\phi(p) = \frac{1}{2} \|p\|_2^2$ where $\|p\|_2^2 = p^\top p$ is the (squared) ℓ_2 -norm, then the Bregman divergence recovers the standard ℓ_2 -distance: $D_\phi(p, \tilde{p}) = \frac{1}{2} \|p - \tilde{p}\|_2^2$, which is symmetric. On the other hand, if $\mathcal{P} = \Delta_m \subset \mathbb{R}^m$ and we choose $\phi(p) = \sum_{i=1}^m p_i \log p_i$ to be negative entropy, then the Bregman divergence is the relative entropy or the Kullback-Leibler divergence: $D_\phi(p, \tilde{p}) = \sum_{i=1}^m p_i \log \frac{p_i}{\tilde{p}_i}$, which is not symmetric.

Similarly, we assume that on the domain $\mathcal{Q} \subseteq \mathbb{R}^n$ we have a strictly convex regularizer function $\psi: \mathcal{Q} \rightarrow \mathbb{R}$ which is a Legendre function, so ψ is continuously differentiable, $\|\nabla \psi(q)\| \rightarrow \infty$ as $q \rightarrow \partial \mathcal{Q}$, and $\nabla \psi$ is a bijection from $\mathcal{Q}$ to the range $\nabla \psi(\mathcal{Q}) \subseteq \mathbb{R}^n$. Let $D_\psi: \mathcal{Q} \times \mathcal{Q} \rightarrow \mathbb{R}$ be the Bregman divergence of ψ , defined by $D_\psi(q, \tilde{q}) = \psi(q) - \psi(\tilde{q}) - \langle \nabla \psi(\tilde{q}), q - \tilde{q} \rangle$.

3 Algorithm: Alternating Mirror Descent

We consider the **Alternating Mirror Descent (AMD)** algorithm where each player follows the mirror descent algorithm to minimize their own loss function, and they perform the updates in an *alternating* fashion (one at a time).

Concretely, the AMD algorithm starts from an arbitrary initial position $(p_0, q_0) \in \mathcal{P} \times \mathcal{Q}$. At each iteration $k \geq 0$, suppose the players are at position $(p_k, q_k) \in \mathcal{P} \times \mathcal{Q}$. Then in the next iteration $k + 1$, they update their position to:

$$p_{k+1} = \arg \min_{p \in \mathcal{P}} \left\{ p^\top A q_k + \frac{1}{\eta} D_\phi(p, p_k) \right\} \quad (1a)$$

$$q_{k+1} = \arg \min_{q \in \mathcal{Q}} \left\{ -p_{k+1}^\top A q + \frac{1}{\eta} D_\psi(q, q_k) \right\} \quad (1b)$$

where $\eta > 0$ is step size. Observe the first player (the p player) makes the update first using the current value q_k ; then the second player (the q player) updates using the new value p_{k+1} . We assume both players can solve the update equations (1), which is the case because they can choose the convex regularizers ϕ, ψ . Observe that we can write the optimality condition for AMD (1) as:

$$\nabla\phi(p_{k+1}) = \nabla\phi(p_k) - \eta A q_k \quad (2a)$$

$$\nabla\psi(q_{k+1}) = \nabla\psi(q_k) + \eta A^\top p_{k+1}. \quad (2b)$$

In Section 4.1, we interpret AMD as an alternating discretization of a skew-gradient flow in dual space.

3.1 Continuous-Time Motivation: Skew-Gradient Flow

One way to derive the alternating mirror descent algorithm (1) is as a discretization of what we call the *skew-gradient flow* in continuous time. Concretely, let us define the dual functions (convex conjugate) of the convex regularizers $\phi^*: \mathbb{R}^m \rightarrow \mathbb{R}$ and $\psi^*: \mathbb{R}^n \rightarrow \mathbb{R}$ given by:

$$\phi^*(x) = \sup_{p \in \mathcal{P}} \langle p, x \rangle - \phi(p)$$

$$\psi^*(y) = \sup_{q \in \mathcal{Q}} \langle q, y \rangle - \psi(q).$$

We call $\mathbb{R}^m$ and $\mathbb{R}^n$ (the domains of ϕ^* and ψ^*) as the *dual space*. Recall that the gradient of the dual function is the inverse map of the original gradient: $\nabla\phi^* = (\nabla\phi)^{-1}$ and $\nabla\psi^* = (\nabla\psi)^{-1}$.

Given $(p_k, q_k) \in \mathcal{P} \times \mathcal{Q}$, we define the **dual variables** $(x_k, y_k) \in \mathbb{R}^{m+n}$ by:

$$\begin{aligned} x_k = \nabla\phi(p_k) &\Leftrightarrow p_k = \nabla\phi^*(x_k) \\ y_k = \nabla\psi(q_k) &\Leftrightarrow q_k = \nabla\psi^*(y_k). \end{aligned}$$

Suppose $(p_k, q_k) \in \mathcal{P} \times \mathcal{Q}$ evolves following AMD (1), so they satisfy the optimality condition (2). Then the dual variables $(x_k, y_k) \in \mathbb{R}^{m+n}$ evolve by the following algorithm:

$$x_{k+1} = x_k - \eta A \nabla\psi^*(y_k) \quad (3a)$$

$$y_{k+1} = y_k + \eta A^\top \nabla\phi^*(x_{k+1}). \quad (3b)$$

We refer to AMD update in the dual space (3) above as the *alternating method*.

Continuous-time limit. Observe that as the step size $\eta \rightarrow 0$, the update equation (3) for AMD in the dual space recovers the following continuous-time dynamics for $(X_t, Y_t) \in \mathbb{R}^{m+n}$:

$$\dot{X}_t = -A \nabla\psi^*(Y_t) \quad (4a)$$

$$\dot{Y}_t = A^\top \nabla\phi^*(X_t). \quad (4b)$$

Here $\dot{X}_t = \frac{d}{dt} X_t$ is the time derivative. We call the dynamics (4) as the *skew-gradient flow*, generated by the *energy function*:

$$H(x, y) := \phi^*(x) + \psi^*(y). \quad (5)$$

A particular feature of the skew-gradient flow (4) is that it preserves the energy function over time:

$$H(X_t, Y_t) = H(X_0, Y_0)$$

for all $t \geq 0$; see Section 4 for further detail.

AMD in dual space as alternating discretization of skew-gradient flow. We can view the update equation (2) as performing an alternating discretization of the skew-gradient flow (4). Since the energy function is separable, these updates are explicit, and correspond to performing alternating mirror descent in the dual space (3). See Section 4.1 for more detail.

3.2 Regret Analysis of Alternating Mirror Descent

To measure the performance of the algorithm, we can analyze the *regret* of each player, which is the gap between their observed losses and the best loss they could have achieved in hindsight, using a fixed (static) action. We define the regret of the alternating mirror descent algorithm (1) as follows.

Regret definition. From iteration k to iteration $k + 1$ of the algorithm, there are two half steps that happen: The first player updates from p_k to p_{k+1} (while the second player is at q_k), then the second player updates from q_k to q_{k+1} (while the first player is at p_{k+1}). Thus, the first player observes q_k twice: once when the first player is at p_k , and once when they are at p_{k+1} . Therefore, we define the regret of the first player after K iterations, with respect to a static action $p \in \mathcal{P}$, to be:

$$R_{1,K}(p) := \sum_{k=0}^{K-1} \left(\frac{p_k + p_{k+1}}{2} \right)^\top A q_k - \sum_{k=0}^{K-1} p^\top A q_k. \quad (6)$$

Similarly, the second player observes p_{k+1} twice: once when the second player is at q_k , and once when they are at q_{k+1} . Therefore, we define the regret of the second player after K iterations, with respect to a static action $q \in \mathcal{Q}$, to be:

$$R_{2,K}(q) := \sum_{k=0}^{K-1} p_{k+1}^\top A q - \sum_{k=0}^{K-1} p_{k+1}^\top A \left(\frac{q_k + q_{k+1}}{2} \right). \quad (7)$$

Note that in the unconstrained case, this recovers the regret definition of the alternating gradient descent algorithm of [2] (up to a factor of $\frac{1}{2}$ for normalization).

We define the cumulative regret of both players after K iterations, with respect to static actions $(p, q) \in \mathcal{P} \times \mathcal{Q}$, to be:

$$R_K(p, q) := R_{1,K}(p) + R_{2,K}(q).$$

Finally, we define the **total regret** of both players after K iterations to be the best cumulative regret in hindsight:

$$R_K := \sup_{(p,q) \in \mathcal{P} \times \mathcal{Q}} R_K(p, q).$$

Regret and duality gap. We define the average iterates of the players after K iterations to be:

$$\bar{p}_K = \frac{1}{K} \sum_{k=0}^{K-1} p_{k+1} \quad \text{and} \quad \bar{q}_K = \frac{1}{K} \sum_{k=0}^{K-1} q_k.$$

Note that we shift the index of the first player by one, since the second player moves after the first player. Then we observe that the total regret is related to the duality gap of the average iterates. We provide the proof of Lemma 3.1 in Section C.1.1.

Lemma 3.1. *Under the above definitions, for any $K \geq 1$,*

$$\text{dg}(\bar{p}_K, \bar{q}_K) = \frac{1}{K} R_K - \frac{1}{2K} (p_0^\top A q_0 - p_K^\top A q_K). \quad (8)$$

If we are in the constrained case with bounded domains, then the last term in (8) is bounded, so the behavior of the duality gap is controlled by the growth of the total regret R_K .

Bound on regret. If both players follow the alternating mirror descent algorithm with smooth convex regularizers, then we can bound the total regret R_K as in Theorem 3.2 below. The proof uses the interpretation of alternating mirror descent as an alternating discretization of the skew-gradient flow, and a relation between the total regret and the change in the modified energy function, as we explain in Section 4.1.3. Smoothness of the regularizers allow us to control the discretization error, which translates to a bound on the regret. We provide the proof of Theorem 3.2 in Section C.1.2.

Here $\|\cdot\|$ is an arbitrary norm. We define the operator norm of $\nabla^3 \phi(x)$ (which is a 3-tensor) by $\|\nabla^3 \phi(x)\|_{\text{op}} = \sup_{\|v\|=1} |\nabla^3 \phi(x)[v, v, v]|$. Let $\alpha_{\max}$ be the maximum singular value of A .

Theorem 3.2. *Assume $\mathcal{P}$ and $\mathcal{Q}$ are bounded, so $\|p\| \leq M$ and $\|q\| \leq M$ for all $p \in \mathcal{P}$, $q \in \mathcal{Q}$, for some $0 < M < \infty$. Assume the dual functions ϕ^* and ψ^* are M -smooth of order 3, which means $\|\nabla^3 \phi^*(\nabla \phi(p))\|_{\text{op}} \leq M$ and $\|\nabla^3 \psi^*(\nabla \psi(q))\|_{\text{op}} \leq M$ for all $p \in \mathcal{P}$, $q \in \mathcal{Q}$. Suppose both players follow alternating mirror descent (1) with step size $\eta > 0$ from $(p_0, q_0) \in \mathcal{P} \times \mathcal{Q}$. Then the total regret at iteration K with respect to any strategy $(p, q) \in \mathcal{P} \times \mathcal{Q}$ is bounded by:*

$$R_K(p, q) \leq \frac{D_\phi(p, p_0) + D_\psi(q, q_0)}{\eta} + \frac{4 \alpha_{\max}^3 M^4}{3} \eta^2 K.$$

In particular, suppose we start from (p_0, q_0) with $\sup_{p \in \mathcal{P}} D_\phi(p, p_0) < \infty$ and $\sup_{q \in \mathcal{Q}} D_\psi(q, q_0) < \infty$.

Given a horizon K , with step size $\eta = \Theta(K^{-1/3})$, the total regret is

$$R_K = O(K^{1/3}).$$

In the Euclidean case, Theorem 3.2 recovers [2, Theorem 2]. Theorem 3.2 also holds when $\mathcal{P}, \mathcal{Q}$ are probability simplices with entropy regularizers, if we start from p_0, q_0 with full support.

By Lemma 3.1 and Theorem 3.2, we have the following corollary (proof is in Section C.1.3).

Corollary 3.3. *Assume the same assumption as in Theorem 3.2. If both players follow the alternating mirror descent (1) with step size $\eta = \Theta(K^{-1/3})$, then the duality gap of the average iterates decays as $\text{dg}(\bar{p}_K, \bar{q}_K) = O(K^{-2/3})$.*

Compare the result above with the simultaneous mirror descent algorithm for min-max games, which has the classical $O(K^{-1/2})$ convergence in the duality gap of the average iterates, thus highlighting the advantage of using the alternating mirror descent algorithm (at the cost of a third-order smoothness assumption for the analysis).

4 Skew-Gradient Flow and Discretization

In this section we discuss the skew-gradient flow and its discretization. We apply this to analyze the alternating mirror descent algorithm in the dual space $\mathbb{R}^m \times \mathbb{R}^n$.

Recall in the dual space, we have the dual functions of the convex regularizers $f := \phi^*: \mathbb{R}^m \rightarrow \mathbb{R}$ and $g := \psi^*: \mathbb{R}^n \rightarrow \mathbb{R}$. We define the **energy function** $H: \mathbb{R}^{m+n} \rightarrow \mathbb{R}$ by:

$$H(x, y) = f(x) + g(y)$$

for all $z = (x, y) \in \mathbb{R}^m \times \mathbb{R}^n$. Note $f = \phi^*$ and $g = \psi^*$ are convex, being convex conjugate functions. Furthermore, since we assume ϕ and ψ are strictly convex, f and g are differentiable. Therefore, H is a differentiable convex function.

Let $J \in \mathbb{R}^{(m+n) \times (m+n)}$ be the skew-symmetric matrix:

$$J = \begin{pmatrix} 0 & A \\ -A^\top & 0 \end{pmatrix}$$

where recall $A \in \mathbb{R}^{m \times n}$ is the payoff matrix for the min-max game.

We consider the **skew-gradient flow** generated by the energy function H and the skew-symmetric matrix J , which is the solution to the differential equation:

$$\dot{Z}_t = -J \nabla H(Z_t) \tag{9}$$

starting from an arbitrary $Z_0 \in \mathbb{R}^{m+n}$. If we write $Z_t = (X_t, Y_t)$, then the components follow the skew-gradient flow dynamics as in (4):

$$\begin{aligned} \dot{X}_t &= -A \nabla g(Y_t) \\ \dot{Y}_t &= A^\top \nabla f(X_t). \end{aligned}$$

A feature of the skew-gradient flow is that since the velocity is orthogonal to the gradient of the energy function, the skew-gradient flow preserves the energy function. (In contrast, recall the usual gradient flow $\dot{Z}_t = -\nabla H(Z_t)$ decreases the energy function.)

Lemma 4.1. *Along the skew-gradient flow (9), $H(Z_t) = H(Z_0)$ for all $t \geq 0$.*

Proof. We can check that the energy function $H(Z_t)$ has zero time derivative along the flow (9):

$$\frac{d}{dt} H(Z_t) = \langle \nabla H(Z_t), \dot{Z}_t \rangle = -\langle \nabla H(Z_t), J \nabla H(Z_t) \rangle = 0 \tag{10}$$

where the last equality above holds because $J = -J^\top$, so it defines a zero quadratic form. $\square$

Some of our results below hold for general skew-symmetric matrix J (not necessarily in block structure). For simplicity, we focus on the separable case above for the min-max game application.

In continuous time, the skew-gradient flow dynamics (9) preserves the energy function. In discrete time, the behavior of the energy function can vary depending on the discretization method used.

Forward Discretization. A forward discretization of skew-gradient flow corresponds to the two players performing the simultaneous mirror descent algorithm for constrained min-max game. Since the energy function is convex, it is monotonically increasing along the forward method, and it is increasing exponentially fast if H is strongly convex. See Section B.1 for details.

Backward Discretization. A backward discretization of skew-gradient flow corresponds to the two players performing the simultaneous proximal mirror descent algorithm for constrained min-max game. Since the energy function is convex, it is monotonically decreasing along the backward method, and it is decreasing exponentially fast if H is strongly convex. See Section B.2 for details.

4.1 Alternating Discretization of Skew-Gradient Flow

The alternating discretization of the skew-gradient flow performs the updates one component at a time, using the forward method for the first component (x) and the backward method for the second component (y), resulting in the update equation:

$$\begin{aligned} x_{k+1} &= x_k - \eta A \nabla_y H(x_{k+1}, y_k) \\ y_{k+1} &= y_k + \eta A^\top \nabla_x H(x_{k+1}, y_k). \end{aligned}$$

Since the energy function $H(x, y) = f(x) + g(y)$ is separable, this yields an explicit update equation:

$$x_{k+1} = x_k - \eta A \nabla g(y_k) \tag{11a}$$

$$y_{k+1} = y_k + \eta A^\top \nabla f(x_{k+1}). \tag{11b}$$

For the min-max game application, this corresponds to the alternating mirror descent update in the dual space (3).

In contrast to either the forward or the backward discretization, this alternating discretization tracks the continuous-time dynamics more closely, and thus we can bound the deviation in the energy function better. To explain our result, we introduce the notion of modified energy function.

4.1.1 Modified Energy Function

Given a step size $\eta > 0$, we define the **modified energy** function $H_\eta: \mathbb{R}^{m+n} \rightarrow \mathbb{R}$ by:

$$H_\eta(z) = H(z) - \frac{\eta}{2} \langle \nabla f(x), A \nabla g(y) \rangle \tag{12a}$$

$$= f(x) + g(y) - \frac{\eta}{2} \langle \nabla f(x), A \nabla g(y) \rangle \tag{12b}$$

for $z = (x, y) \in \mathbb{R}^{m+n}$. Note that when f and g are quadratic functions, this recovers the definition of the modified energy function in [2].

Recall the Bregman divergence $D_H(z', z) = H(z') - H(z) - \langle \nabla H(z), z' - z \rangle$ is not necessarily symmetric. We define the **Bregman commutator** $C_H: \mathbb{R}^{m+n} \times \mathbb{R}^{m+n} \rightarrow \mathbb{R}$ as a measure of the non-commutativity of Bregman divergence:

$$\begin{aligned} C_H(z', z) &= \frac{1}{2} (D_H(z', z) - D_H(z, z')) \\ &= H(z') - H(z) - \frac{1}{2} \langle \nabla H(z') + \nabla H(z), z' - z \rangle. \end{aligned}$$

Observe that when $H(z) = \frac{1}{2} z^\top M z$ is a quadratic function (for any positive definite matrix M), then $D_H(z', z) = \frac{1}{2} (z' - z)^\top M (z' - z) = D_H(z, z')$ is symmetric, and thus $C_H(z', z) = 0$. Conversely, if the Bregman commutator vanishes: $C_H(z', z) = 0$, then H must be a quadratic function; see Lemma A.1 in Section A.2.1. We can also bound the Bregman commutator under third-order smoothness; see Lemma A.2 in Section A.2.1.

We show that along the alternating method, the modified energy changes by precisely the Bregman commutator; we provide the proof in Section B.3.2.

Lemma 4.2. *Let $z_k = (x_k, y_k)$ evolve following the alternating method (11). Then for any $k \geq 0$,*

$$H_\eta(z_{k+1}) = H_\eta(z_k) + C_H(z_{k+1}, z_k). \tag{13}$$

As a corollary, if H is quadratic—in which case the Bregman commutator vanishes—then the modified energy function is conserved along the alternating method. This recovers the main result of [2] for the unconstrained min-max game. We provide the proof of Corollary 4.3 in Section B.3.1.

Corollary 4.3. *Suppose H is a quadratic function. Along the alternating method (3), for any $k \geq 0$:*

$$H_\eta(z_k) = H_\eta(z_0).$$

4.1.2 Bound under Third-Order Smoothness

Under smoothness assumptions on H , we can deduce the following bound on the modified energy function. Let $\|\cdot\|$ be an arbitrary norm in $\mathbb{R}^{m+n}$ with dual norm $\|\cdot\|_*$. Recall we say H is L -Lipschitz if $|H(z') - H(z)| \leq L\|z' - z\|$ for all $z, z' \in \mathbb{R}^{m+n}$; equivalently, $\|\nabla H(z)\|_* \leq L$. We say H is M -smooth of order 3 if H is three-times differentiable, and $\|\nabla^3 H(z)\|_{\text{op}} \leq M$ for all $z \in \mathbb{R}^{m+n}$, where $\|\cdot\|_{\text{op}}$ is the operator norm. Then we have the following bound. We provide the proof of Theorem 4.4 in Section B.3.3. We note the following result holds without assuming convexity of H .

Theorem 4.4. *Assume that H is L_1 -Lipschitz and L_3 -smooth of order 3. Let $\alpha_{\max}$ be the maximum singular value of A . Along the alternating method (11), for any $\eta \geq 0$ and at any iteration $k \geq 0$:*

$$|H_\eta(z_k) - H_\eta(z_0)| \leq \frac{1}{12} \alpha_{\max}^3 L_3 L_1^3 \eta^3 k.$$

As an example, the function $H(z) = \log \sum_{i=1}^m e^{z_i} + \log \sum_{j=1}^n e^{z_{m+j}}$ satisfies the assumptions of Theorem 4.4 with respect to $\|\cdot\|_\infty$ -norm; see Example A.3 in Section A.2.1. In our min-max game application, this function arises from using the entropy regularizer on the probability simplex.

4.1.3 Regret of Alternating Mirror Descent in terms of Modified Energy

We now consider the constrained min-max game application where $H(x, y) = f(x) + g(y)$ with $f = \phi^*$ and $g = \psi^*$. Given $(p, q) \in \mathcal{P} \times \mathcal{Q}$, let $z = (x, y) \in \mathbb{R}^{m+n}$ be the dual variables, where $x = \nabla \phi(p)$ and $y = \nabla \psi(q)$. Then we can write the cumulative regret of alternating mirror descent in terms of the difference of the modified energy. We provide the proof of Theorem 4.5 in Section B.3.4.

Theorem 4.5. *Let (p_k, q_k) evolve following the alternating mirror descent algorithm (1) with any step size $\eta > 0$, and let $z_k = (x_k, y_k)$ be the dual variables. We can write the cumulative regret of the two players with respect to any $(p, q) \in \mathcal{P} \times \mathcal{Q}$ as:*

$$R_K(p, q) = \frac{1}{\eta} (D_H(z_0, z) - D_H(z_K, z) + H_\eta(z_K) - H_\eta(z_0)).$$

Since H is convex, we can further bound $D_H(z_K, z) \geq 0$, so from the above we get:

$$R_K(p, q) \leq \frac{1}{\eta} (D_H(z_0, z) + H_\eta(z_K) - H_\eta(z_0)).$$

The first term in the bound above is a constant that only depends on the initial point. If we assume the domains are bounded, then this first term is also bounded. Thus, to control the cumulative regret, we need to bound the increase in the modified energy, which we can do via Theorem 4.4 under third-order smoothness. Completing this step yields the proof of Theorem 3.2; see Section C.1.2.

5 Discussion

In this paper we study the alternating mirror descent algorithm for constrained min-max games, and showed that it achieves a better regret bound than the classical simultaneous mirror descent. Our results extend the findings of [2] from the unconstrained to the constrained setting. We have utilized an interpretation of alternating mirror descent as an alternating discretization of the skew-gradient flow in the dual space, and linked the total regret of the players with the growth of the modified energy function. Our analysis highlights the connections between min-max games and Hamiltonian structures, which helps pave the way toward a closer interplay between classical numerical methods and modern algorithmic questions motivated by machine learning applications.

Our results leave many interesting open questions. First, our bound in Theorem 4.4 grows with the number of iterations, and does not yield constant regret as in the unconstrained case. In the special case of identity payoff matrix, in which the skew-gradient flow dynamics has a Hamiltonian and symplectic structure, we conjecture that for sufficiently nice energy functions, the iterates of the

alternating method stay uniformly bounded, as in the unconstrained setting. In the Appendix, we provide some empirical evidence supporting this conjecture. Resolving this question requires making concrete classical bounds from numerical methods, which may be of independent interest.

Second, we have focused on the simple case of two-player bilinear game, which already presents non-trivial behavior. It would be interesting to understand the more general setting of multi-player games with general payoffs, in which the notion of “alternating” play can take many possible forms. It would also be interesting to study the asynchronous setting in which the players make their moves in decentralized or randomized fashions. This will help us understand and control the global behavior of multi-agent systems and how they emerge from simple local or individual strategies.

Acknowledgements

We thank the reviewers for helpful comments and discussions on the paper. This research/project is supported in part by the National Research Foundation, Singapore and DSO National Laboratories under its AI Singapore Program (AISG Award No: AISG2-RP-2020-016), NRF 2018 Fellowship NRF-NRFF2018-07, NRF2019-NRF-ANR095 ALIAS grant, grant PIESGP-AI-2020-01, AME Programmatic Fund (Grant No.A20H6b0151) from the Agency for Science, Technology and Research (A*STAR) and Provost’s Chair Professorship grant RGEPPV2101.

References

- [1] Gabriel P Andrade, Rafael Frongillo, and Georgios Piliouras. Learning in matrix games can be arbitrarily complex. In *Conference on Learning Theory*, pages 159–185. PMLR, 2021.
- [2] James P Bailey, Gauthier Gidel, and Georgios Piliouras. Finite regret and cycles with fixed step-size via alternating gradient descent-ascent. In *Conference on Learning Theory*, pages 391–407. PMLR, 2020.
- [3] James P. Bailey and Georgios Piliouras. Multiplicative weights update in zero-sum games. In *ACM Conference on Economics and Computation*, 2018.
- [4] James P. Bailey and Georgios Piliouras. Multi-agent learning in network zero-sum games is a Hamiltonian system. In *Proceedings of the 18th International Conference on Autonomous Agents and MultiAgent Systems, AAMAS ’19*, page 233–241, Richland, SC, 2019. International Foundation for Autonomous Agents and Multiagent Systems.
- [5] Giancarlo Benettin and Antonio Giorgilli. On the hamiltonian interpolation of near-to-the identity symplectic mappings with application to symplectic integration algorithms. *Journal of Statistical Physics*, 74(5):1117–1143, 1994.
- [6] Nikolo Cesa-Bianchi and Gabor Lugosi. *Prediction, Learning, and Games*. Cambridge University Press, 2006.
- [7] Yun Kuen Cheung and Georgios Piliouras. Vortices instead of equilibria in minmax optimization: Chaos and butterfly effects of online learning in zero-sum games. In *Conference on Learning Theory*, pages 807–834. PMLR, 2019.
- [8] Thiparat Chotibut, Fryderyk Falniowski, Michał Misiurewicz, and Georgios Piliouras. The route to chaos in routing games: When is price of anarchy too optimistic? *Advances in Neural Information Processing Systems*, 33:766–777, 2020.
- [9] Constantinos Daskalakis, Maxwell Fishelson, and Noah Golowich. Near-optimal no-regret learning in general games. *Advances in Neural Information Processing Systems*, 34, 2021.
- [10] Constantinos Daskalakis, Andrew Ilyas, Vasilis Syrgkanis, and Haoyang Zeng. Training GANs with optimism. In *ICLR*, 2018.
- [11] Ian Goodfellow, Jean Pouget-Abadie, Mehdi Mirza, Bing Xu, David Warde-Farley, Sherjil Ozair, Aaron Courville, and Yoshua Bengio. Generative adversarial nets. In *Advances in neural information processing systems*, pages 2672–2680, 2014.

- [12] E. Hairer, C. Lubich, and G. Wanner. *Geometric Numerical Integration: Structure-Preserving Algorithms for Ordinary Differential Equations*. Springer, Berlin Heidelberg New York, second edition, 2006.
- [13] J. Hofbauer. Evolutionary dynamics for bimatrix games: A hamiltonian system? *J. of Math. Biology*, 34:675–688, 1996.
- [14] Michael I Jordan. Dynamical, symplectic and stochastic perspectives on gradient-based optimization. In *Proceedings of the International Congress of Mathematicians: Rio de Janeiro 2018*, pages 523–549. World Scientific, 2018.
- [15] Alistair Letcher. On the impossibility of global convergence in multi-loss optimization. *ICLR*, 2021.
- [16] Aleksander Madry, Aleksandar Makelov, Ludwig Schmidt, Dimitris Tsipras, and Adrian Vladu. Towards deep learning models resistant to adversarial attacks. In *6th International Conference on Learning Representations, ICLR 2018, Vancouver, BC, Canada, April 30 - May 3, 2018, Conference Track Proceedings*. OpenReview.net, 2018.
- [17] Panayotis Mertikopoulos, Bruno Lecouat, Houssam Zenati, Chuan-Sheng Foo, Vijay Chandrasekhar, and Georgios Piliouras. Optimistic mirror descent in saddle-point problems: Going the extra(-gradient) mile. In *ICLR*, 2019.
- [18] Panayotis Mertikopoulos, Christos Papadimitriou, and Georgios Piliouras. Cycles in adversarial regularized learning. In *Proceedings of the Twenty-Ninth Annual ACM-SIAM Symposium on Discrete Algorithms*, pages 2703–2717. SIAM, 2018.
- [19] Aryan Mokhtari, Asuman Ozdaglar, and Sarath Pattathil. A unified analysis of extra-gradient and optimistic gradient methods for saddle point problems: Proximal point approach. In *International Conference on Artificial Intelligence and Statistics*, pages 1497–1507. PMLR, 2020.
- [20] Michael Muehlebach and Michael I Jordan. Optimization with momentum: Dynamical, control-theoretic, and symplectic perspectives. *Journal of Machine Learning Research*, 22(73):1–50, 2021.
- [21] A.S. Nemirovskii and D.B. Yudin. *Problem Complexity and Method Efficiency in Optimization*. A Wiley-Interscience publication. Wiley, 1983.
- [22] Yurii Nesterov. *Lectures on convex optimization*, volume 137. Springer, 2018.
- [23] Shayegan Omidshafiei, Christos Papadimitriou, Georgios Piliouras, Karl Tuyls, Mark Rowland, Jean-Baptiste Lespiau, Wojciech M Czarnecki, Marc Lanctot, Julien Perolat, and Remi Munos. α -rank: Multi-agent evaluation by evolution. *Scientific reports*, 9(1):1–29, 2019.
- [24] Christos Papadimitriou and Georgios Piliouras. Game dynamics as the meaning of a game. *ACM SIGecom Exchanges*, 16(2):53–63, 2019.
- [25] Georgios Piliouras, Ryann Sim, and Stratis Skoulakis. Beyond time-average convergence: Near-optimal uncoupled online learning via Clairvoyant Multiplicative Weights Update. *arXiv preprint arXiv:2111.14737*, 2021.
- [26] R. Tyrrell Rockafellar. *Convex Analysis*. Princeton Landmarks in Mathematics and Physics. Princeton University Press, 1970.
- [27] Julian Schrittwieser, Ioannis Antonoglou, Thomas Hubert, Karen Simonyan, Laurent Sifre, Simon Schmitt, Arthur Guez, Edward Lockhart, Demis Hassabis, Thore Graepel, et al. Mastering atari, go, chess and shogi by planning with a learned model. *Nature*, 588(7839):604–609, 2020.
- [28] Molei Tao and Tomoki Ohsawa. Variational optimization on lie groups, with examples of leading (generalized) eigenvalue problems. In *International Conference on Artificial Intelligence and Statistics*, pages 4269–4280. PMLR, 2020.
- [29] John von Neumann. Zur theorie der gesellschaftsspiele. *Mathematische Annalen*, 100:295–300, 1928.

- [30] Andre Wibisono, Ashia C Wilson, and Michael I Jordan. A variational perspective on accelerated methods in optimization. *Proceedings of the National Academy of Sciences*, 113(47):E7351–E7358, 2016.

A Helper Lemmas

A.1 Properties of Bregman Divergence

Let $f: \mathbb{R}^m \rightarrow \mathbb{R}$ be a differentiable function. Recall the *Bregman divergence* of f is given by:

$$D_f(x', x) = f(x') - f(x) - \langle \nabla f(x), x' - x \rangle$$

for all $x, x' \in \mathbb{R}^m$. The Bregman divergence is in general not symmetric: $D_f(x', x) \neq D_f(x, x')$.

We recall the three-point identity for Bregman divergence:

$$D_f(x, y) - D_f(x, z) - D_f(z, y) = \langle \nabla f(z) - \nabla f(y), x - z \rangle.$$

Recall we say f is α -strongly convex with respect to a norm $\|\cdot\|$ if

$$D_f(x', x) \geq \frac{\alpha}{2} \|x' - x\|^2.$$

For example, if $f(x) = \frac{1}{2}x^\top Mx$ is a quadratic function defined by a symmetric positive definite matrix $M \in \mathbb{R}^{m \times m}$, then f is strongly convex in the ℓ_2 -norm with strong convexity constant equal to the smallest eigenvalue of M . If $f(x) = \sum_{i=1}^m x_i \log x_i$ is the negative entropy function defined on the simplex $\Delta_m \subset \mathbb{R}^m$, then f is 1-strongly convex in the ℓ_1 -norm.

Recall we say f is α -gradient dominated with respect to the dual norm $\|\cdot\|_*$ if

$$\|\nabla f(x)\|_*^2 \geq 2\alpha(f(x) - \min f)$$

for all $x \in \mathbb{R}^m$, where $\min f = \min_x f(x)$ is the minimum value of f .

We recall that strong convexity implies gradient domination with the same constant.

A.2 Properties of Bregman commutator

The *Bregman commutator* of f is the function $C_f: \mathbb{R}^m \times \mathbb{R}^m \rightarrow \mathbb{R}$ that measures the failure of the Bregman divergence to be symmetric:

$$\begin{aligned} C_f(x', x) &= \frac{1}{2}(D_f(x', x) - D_f(x, x')) \\ &= f(x') - f(x) - \frac{1}{2}\langle \nabla f(x') + \nabla f(x), x' - x \rangle. \end{aligned}$$

By definition, the Bregman commutator is antisymmetric:

$$C_f(x, x') = -C_f(x', x).$$

For example, if f is a quadratic function, then the Bregman divergence is also a quadratic function, which is symmetric, and hence the Bregman commutator vanishes: $C_f(x, x') = 0$ for all $x, x' \in \mathbb{R}^m$. The converse also holds: If the Bregman commutator vanishes, then the function must be a quadratic.

Lemma A.1. *If $C_f(x', x) = 0$ for all $x', x \in \mathbb{R}^m$, then f is a quadratic function.*

Proof. By assumption, we have $D_f(x', x) = D_f(x, x')$ for all $x, x' \in \mathbb{R}^m$. It suffices to argue that $\tilde{f}(x) := f(x) - f(0) - \langle x, \nabla f(0) \rangle$ is a quadratic function, for then f is also quadratic. Since this is a linear transformation, it does not affect the Bregman divergence, and thus by assumption we have $D_{\tilde{f}}(x', x) = D_{\tilde{f}}(x, x')$ for all $x, x' \in \mathbb{R}^m$. Note that by definition, $\tilde{f}(0) = 0$ and $\nabla \tilde{f}(0) = 0$. Now, the relation $D_{\tilde{f}}(x, 0) = D_{\tilde{f}}(0, x)$ implies that $\tilde{f}(x) = \frac{1}{2}\langle \nabla \tilde{f}(x), x \rangle$. Thus, $D_{\tilde{f}}(x', x) = D_{\tilde{f}}(x, x')$ implies that $\langle \nabla \tilde{f}(x), x' \rangle = \langle \nabla \tilde{f}(x'), x \rangle$ for all $x, x' \in \mathbb{R}^m$. This in turn implies, for any x, x' , and y ,

$$\begin{aligned} \langle \nabla \tilde{f}(x + x'), y \rangle &= \langle \nabla \tilde{f}(y), x + x' \rangle \\ &= \langle \nabla \tilde{f}(y), x \rangle + \langle \nabla \tilde{f}(y), x' \rangle \\ &= \langle \nabla \tilde{f}(x), y \rangle + \langle \nabla \tilde{f}(x'), y \rangle \\ &= \langle \nabla \tilde{f}(x) + \nabla \tilde{f}(x'), y \rangle. \end{aligned}$$

This shows that $\nabla \tilde{f}$ is linear, which means $\tilde{f}$ is quadratic, and hence f is also quadratic. $\square$

A.2.1 Bound on Bregman commutator

Recall f is M -smooth of order 3 if $f: \mathbb{R}^m \rightarrow \mathbb{R}$ is three-times differentiable and for all $x \in \mathbb{R}^m$:

$$\|\nabla^3 f(x)\|_{\text{op}} \leq M$$

where $\|\cdot\|_{\text{op}}$ is the operator norm; concretely, this means for all $x \in \mathbb{R}^m$ and $v \in \mathbb{R}^m$ with $\|v\| = 1$,

$$|\nabla^3 f(x)[v, v, v]| \leq M.$$

Equivalently, $\nabla^2 f$ is M -Lipschitz.

We have the following bound on Bregman commutator under third-order smoothness.

Lemma A.2. Assume f is M -smooth of order 3 with respect to a norm $\|\cdot\|$. Then for all $x', x \in \mathbb{R}^m$:

$$|C_f(x, x')| \leq \frac{M}{12} \|x - x'\|^3. \quad (14)$$

Proof. Recall the Taylor expansion

$$r(1) = r(0) + \dot{r}(0) + \int_0^1 (1-t) \ddot{r}(t) dt$$

for any twice-differentiable function $r: [0, 1] \rightarrow \mathbb{R}$. Applying this with $r(t) = f(x + t(x' - x))$ gives

$$f(x') = f(x) + \langle \nabla f(x), x' - x \rangle + \int_0^1 (1-t) \langle (x' - x), \nabla^2 f(x + t(x' - x)) (x' - x) \rangle dt,$$

or equivalently,

$$D_f(x', x) = \left\langle (x' - x), \left(\int_0^1 (1-t) \nabla^2 f(x + t(x' - x)) dt \right) (x' - x) \right\rangle.$$

Similarly,

$$D_f(x, x') = \left\langle (x' - x), \left(\int_0^1 t \nabla^2 f(x + (1-t)(x' - x)) dt \right) (x' - x) \right\rangle.$$

Combining the two equations above, we obtain

$$\begin{aligned} C_f(x', x) &= \frac{1}{2} (D_f(x', x) - D_f(x, x')) \\ &= \frac{1}{2} \left\langle (x' - x), \left(\int_0^1 (1-2t) \nabla^2 f(x + t(x' - x)) dt \right) (x' - x) \right\rangle. \end{aligned} \quad (15)$$

For $0 \leq t \leq 1$, let $s = t - \frac{1}{2}$ and $s' = -s$. We denote $x_s \equiv x + (s + \frac{1}{2})(x' - x)$. Then we can write the integral above as

$$\begin{aligned} \frac{1}{2} \int_0^1 (1-2t) \nabla^2 f(x_{t-\frac{1}{2}}) dt &= \int_{-\frac{1}{2}}^{\frac{1}{2}} (-s) \nabla^2 f(x_s) ds \\ &= \int_{-\frac{1}{2}}^0 (-s) \nabla^2 f(x_s) ds + \int_0^{\frac{1}{2}} (-s) \nabla^2 f(x_s) ds \\ &= \int_0^{\frac{1}{2}} s' \nabla^2 f(x_{-s'}) ds' + \int_0^{\frac{1}{2}} (-s) \nabla^2 f(x_s) ds \\ &= \int_0^{\frac{1}{2}} s (\nabla^2 f(x_{-s}) - \nabla^2 f(x_s)) ds. \end{aligned} \quad (16)$$

By the mean-value theorem and using $x_{-s} - x_s = -2s(x' - x)$, we can write

$$\begin{aligned} \nabla^2 f(x_{-s}) - \nabla^2 f(x_s) &= \int_0^1 \nabla^3 f(x_s + u(x_{-s} - x_s)) [x_{-s} - x_s] du \\ &= -2s \int_0^1 \nabla^3 f(x_s + u(x_{-s} - x_s)) [x' - x] du. \end{aligned}$$

Plugging this in to (16), we obtain

$$\frac{1}{2} \int_0^1 (1-2t) \nabla^2 f(x_{t-\frac{1}{2}}) dt = -2 \int_0^{\frac{1}{2}} \int_0^1 s^2 \nabla^3 f(x_s + u(x_{-s} - x_s)) [x' - x] du ds.$$

Plugging this in to (15), we then obtain

$$C_f(x', x) = -2 \int_0^{\frac{1}{2}} \int_0^1 s^2 \nabla^3 f(x_s + u(x_{-s} - x_s)) [x' - x, x' - x, x' - x] du ds.$$

Using the third-order smoothness property of f , we can then bound

$$\begin{aligned} |C_f(x', x)| &\leq 2 \int_0^{\frac{1}{2}} \int_0^1 s^2 |\nabla^3 f(x_s + u(x_{-s} - x_s)) [x' - x, x' - x, x' - x]| du ds \\ &\leq 2 \int_0^{\frac{1}{2}} \int_0^1 s^2 M \|x' - x\|^3 du ds \\ &= 2M \|x' - x\|^3 \cdot \frac{1}{3} \left(\frac{1}{2}\right)^3 \\ &= \frac{M \|x' - x\|^3}{12} \end{aligned}$$

as desired. $\square$

Example A.3. Consider $f(x) = \log \sum_{i=1}^m e^{x_i}$. This is the log-partition function of an exponential family distribution p_x over a finite set $[m] := \{1, \dots, m\}$ with $p_x(i) = e^{x_i - f(x)}$. Then $\nabla^3 f(x)$ is the third-order covariance, i.e., for all $v \in \mathbb{R}^m$,

$$\nabla^3 f(x)[v, v, v] = \text{Cov}_{I \sim p_x}^3(v_I) = \mathbb{E}_{I \sim p_x}[(v_I - \bar{v})^3]$$

where $\bar{v} = \mathbb{E}_{I \sim p_x}[v_I] = \sum_{i=1}^m v_i p_x(i)$. Suppose $\|v\|_\infty = \max_i |v_i| = 1$, so $|\bar{v}| \leq 1$, and $|v_i - \bar{v}| \leq 2$. Then

$$|\nabla^3 f(x)[v, v, v]| \leq \mathbb{E}_{I \sim p_x}[|v_I - \bar{v}|^3] \leq \mathbb{E}_{I \sim p_x}[2^3] = 8.$$

This shows that f is 8-smooth of order 3 with respect to the ℓ_∞ -norm $\|\cdot\|_\infty$.

B Discretization of Skew-Gradient Flow

B.1 Forward Discretization of Skew-Gradient Flow

Consider the forward method to discretize the skew-gradient flow (9) with step size $\eta > 0$:

$$z_{k+1} = z_k - \eta J \nabla H(z_k). \quad (17)$$

In terms of the components $z_k = (x_k, y_k)$, this corresponds to the simultaneous forward method:

$$\begin{aligned} x_{k+1} &= x_k - \eta A \nabla g(y_k) \\ y_{k+1} &= y_k + \eta A^\top \nabla f(x_k). \end{aligned}$$

For the min-max game application when $f = \nabla \phi^*$ and $g = \psi^*$, this corresponds to the two players following the simultaneous mirror descent algorithm:

$$\begin{aligned} p_{k+1} &= \arg \min_{p \in \mathcal{P}} \left\{ p^\top A q_k + \frac{1}{\eta} D_\phi(p, p_k) \right\} \\ q_{k+1} &= \arg \min_{q \in \mathcal{Q}} \left\{ -p_k^\top A q + \frac{1}{\eta} D_\psi(q, q_k) \right\}. \end{aligned}$$

We show the following properties of the forward method. Here let $\|\cdot\| = \|\cdot\|_2$ be the ℓ_2 -norm. Recall H is L -smooth if $\nabla H(z)$ is L -Lipschitz:

$$\|\nabla H(z) - \nabla H(z')\|_* \leq L \|z - z'\|.$$

Equivalently, $\|\nabla^2 H(z)\|_{\text{op}} \leq L$ where $\|\cdot\|_{\text{op}}$ is the operator norm, and $\nabla^2 H(z)$ is the Hessian matrix of second derivatives.

Theorem B.1. Let z_k evolve following the forward method (17). We have the following properties:

1. If H is convex, then $H(z_{k+1}) \geq H(z_k)$.
2. Assume $m = n$, and assume $A \in \mathbb{R}^{n \times n}$ is invertible with smallest singular value $\alpha_{\min} > 0$. Assume H is μ -strongly convex, and let $z^* = \arg \min_z H(z)$. Then:

$$H(z_k) - H(z^*) \geq (1 + \eta^2 \alpha_{\min}^2 \mu^2)^k (H(z_0) - H(z^*)).$$

3. Assume H is convex and L_2 -smooth. Let $\alpha_{\max}$ be the maximum singular value of A . Then:

$$H(z_k) - H(z^*) \leq (1 + \eta^2 \alpha_{\max}^2 L^2)^k (H(z_0) - H(z^*))$$

The properties above say that if the energy function is convex, it is monotonically increasing along the forward method; furthermore, it is increasing exponentially fast if H is strongly convex. On the other hand, if H is convex and smooth, then it is increasing at most exponentially fast.

We note that while the assumptions in Theorem B.1 are stated with respect to the ℓ_2 -norm $\|\cdot\|_2$, since all norms in $\mathbb{R}^{m+n}$ are equivalent, the results also hold for any other norm $\|\cdot\|$ (with smoothness and strong convexity constants that may have dimension dependence). Here for simplicity we provide the proof for the ℓ_2 -norm. An explicit calculation on a quadratic example shows that the rates in Theorem B.1 are tight; see Example B.2.

Proof of Theorem B.1.

1. First assume H is convex. By Jensen's inequality,

$$H(z_{k+1}) - H(z_k) \geq \langle \nabla H(z_k), z_{k+1} - z_k \rangle = \eta \langle \nabla H(z_k), J \nabla H(z_k) \rangle = 0.$$

This shows the forward method always increases the energy function: $H(z_{k+1}) \geq H(z_k)$.

2. Assume $m = n$, and $A \in \mathbb{R}^{n \times n}$ has smallest singular value $\alpha_{\min} > 0$. This implies $A^\top A \succeq \alpha_{\min}^2 I_n$ and $AA^\top \succeq \alpha_{\min}^2 I_n$, so $J^\top J \succeq \alpha_{\min}^2 I_{2n}$.

Assume further that H is μ -strongly convex. This implies H is μ -gradient dominated:

$$\|\nabla H(z)\|_2^2 \geq 2\mu(H(z) - H(z^*))$$

for all $z \in \mathbb{R}^{2n}$. Then by the μ -strong convexity of H , along the forward method,

$$\begin{aligned} H(z_{k+1}) - H(z_k) &\geq \langle \nabla H(z_k), z_{k+1} - z_k \rangle + \frac{\mu}{2} \|z_{k+1} - z_k\|_2^2 \\ &= -\eta \langle \nabla H(z_k), J \nabla H(z_k) \rangle + \frac{\eta^2 \mu}{2} \|J \nabla H(z_k)\|_2^2 \\ &= 0 + \frac{\eta^2 \mu}{2} \langle \nabla H(z_k), (J^\top J) \nabla H(z_k) \rangle \\ &\geq \frac{\eta^2 \alpha_{\min}^2 \mu}{2} \|\nabla H(z_k)\|_2^2 \\ &\geq \eta^2 \alpha_{\min}^2 \mu^2 (H(z_k) - H(z^*)) \end{aligned}$$

This implies $H(z_{k+1}) - H(z^*) \geq (1 + \eta^2 \alpha_{\min}^2 \mu^2)(H(z_k) - H(z^*))$. Therefore, the forward method increases the energy function exponentially fast:

$$H(z_k) - H(z^*) \geq (1 + \eta^2 \alpha_{\min}^2 \mu^2)^k (H(z_0) - H(z^*)).$$

3. Since A has maximum singular value $\alpha_{\max}$, we have $A^\top A \preceq \alpha_{\max}^2 I_n$ and $AA^\top \preceq \alpha_{\max}^2 I_m$, so $J^\top J \preceq \alpha_{\max}^2 I_{m+n}$.

Since H is convex and L_2 -smooth, we have for all $z \in \mathbb{R}^{m+n}$ (see e.g. [22]):

$$H(z) \geq H(z^*) + \frac{1}{2L_2} \|\nabla H(z)\|_2^2.$$

Then by the L_2 -smoothness of H , along the forward method,

$$\begin{aligned}
H(z_{k+1}) - H(z_k) &\leq \langle \nabla H(z_k), z_{k+1} - z_k \rangle + \frac{L_2}{2} \|z_{k+1} - z_k\|_2^2 \\
&= -\eta \langle \nabla H(z_k), J \nabla H(z_k) \rangle + \frac{\eta^2 L_2}{2} \|J \nabla H(z_k)\|_2^2 \\
&= 0 + \frac{\eta^2 L_2}{2} \langle \nabla H(z_k), (J^\top J) \nabla H(z_k) \rangle \\
&\leq \frac{\eta^2 \alpha_{\max}^2 L_2}{2} \|\nabla H(z_k)\|_2^2 \\
&\leq \eta^2 \alpha_{\max}^2 L_2^2 (H(z_k) - H(z^*))
\end{aligned}$$

This implies $H(z_{k+1}) - H(z^*) \leq (1 + \eta^2 \alpha_{\max}^2 L_2^2)(H(z_k) - H(z^*))$. Therefore, along the forward method, the energy function increases at most exponentially fast:

$$H(z_k) - H(z^*) \leq (1 + \eta^2 \alpha_{\max}^2 L_2^2)^k (H(z_0) - H(z^*)).$$

□

In Section C.3 we analyze the total regret of the forward method for the min-max game application, and show the duality gap of the average iterate decays at an $O(K^{-1/2})$ rate; see Theorem C.2.

B.1.1 Quadratic example

We can check that in the quadratic case, the rates in Theorem B.1 are tight in terms of dependence on step size η .

Example B.2 (Quadratic). Suppose $f(x) = \frac{\beta}{2} \|x\|_2^2$ and $g(y) = \frac{\gamma}{2} \|y\|_2^2$ for some $\beta, \gamma > 0$. Then for $z = (x, y)$, the energy function is $H(z) = \frac{1}{2} z^\top M z$ where $M = \begin{pmatrix} \beta I_m & 0 \\ 0 & \gamma I_n \end{pmatrix}$, where $I_m \in \mathbb{R}^{m \times m}$ and $I_n \in \mathbb{R}^{n \times n}$ are the identity matrices. Here H is μ -strongly convex and L -smooth, where $\mu = \min\{\beta, \gamma\}$ and $L = \max\{\beta, \gamma\}$.

The forward method (17) becomes

$$z_{k+1} = z_k + \eta J M z_k = (I_{m+n} - \eta J M) z_k.$$

Then

$$\begin{aligned}
H(z_{k+1}) &= \frac{1}{2} z_{k+1}^\top M z_{k+1} \\
&= \frac{1}{2} z_k^\top (I_{m+n} - \eta M^\top J^\top) M (I_{m+n} - \eta J M) z_k \\
&= \frac{1}{2} z_k^\top \begin{pmatrix} I_m & -\eta \beta A \\ \eta \gamma A^\top & I_n \end{pmatrix} \begin{pmatrix} \beta I_m & 0 \\ 0 & \gamma I_n \end{pmatrix} \begin{pmatrix} I_m & \eta \gamma A \\ -\eta \beta A^\top & I_n \end{pmatrix} z_k \\
&= \frac{1}{2} z_k^\top \begin{pmatrix} \beta I_m + \eta^2 \beta^2 \gamma A A^\top & 0 \\ 0 & \gamma I_n + \eta^2 \beta \gamma^2 A^\top A \end{pmatrix} z_k.
\end{aligned}$$

If $m \neq n$, then one of $A A^\top$, $A^\top A$ will be low-rank and thus have a zero eigenvalue, so

$$H(z_{k+1}) \geq \frac{1}{2} z_k^\top \begin{pmatrix} \beta I_m & 0 \\ 0 & \gamma I_n \end{pmatrix} z_k = H(z_k).$$

Suppose $m = n$ and A has smallest singular value $\alpha_{\min} > 0$, so $A A^\top \succeq \alpha_{\min}^2 I_n$ and $A^\top A \succeq \alpha_{\min}^2 I_n$. Then

$$\begin{aligned}
H(z_{k+1}) &\geq \frac{1}{2} z_k^\top \begin{pmatrix} \beta I_m + \eta^2 \beta^2 \gamma \alpha_{\min}^2 I_m & 0 \\ 0 & \gamma I_n + \eta^2 \beta \gamma^2 \alpha_{\min}^2 I_n \end{pmatrix} z_k \\
&= \frac{(1 + \eta^2 \alpha_{\min}^2 \beta \gamma)}{2} z_k^\top \begin{pmatrix} \beta I_m & 0 \\ 0 & \gamma I_n \end{pmatrix} z_k \\
&= (1 + \eta^2 \alpha_{\min}^2 \beta \gamma) H(z_k).
\end{aligned}$$

Note in this case the minimizer is $z^* = 0$ with $H(z^*) = 0$. Therefore,

$$H(z_k) - H(z^*) \geq (1 + \eta^2 \alpha_{\min}^2 \beta \gamma)^k (H(z_0) - H(z^*)).$$

Finally, let $\alpha_{\max}$ be the maximum singular value of A . Then

$$\begin{aligned} H(z_{k+1}) &= \frac{1}{2} z_k^\top \begin{pmatrix} \beta I_m + \eta^2 \beta^2 \gamma A A^\top & 0 \\ 0 & \gamma I_n + \eta^2 \beta \gamma^2 A^\top A \end{pmatrix} z_k \\ &\leq \frac{1}{2} z_k^\top \begin{pmatrix} \beta I_m + \eta^2 \beta^2 \gamma \alpha_{\max}^2 I_m & 0 \\ 0 & \gamma I_n + \eta^2 \beta \gamma^2 \alpha_{\max}^2 I_n \end{pmatrix} z_k \\ &= \frac{(1 + \eta^2 \alpha_{\max}^2 \beta \gamma)}{2} z_k^\top \begin{pmatrix} \beta I_m & 0 \\ 0 & \gamma I_n \end{pmatrix} z_k \\ &= (1 + \eta^2 \alpha_{\max}^2 \beta \gamma) H(z_k). \end{aligned}$$

Therefore,

$$H(z_k) - H(z^*) \leq (1 + \eta^2 \alpha_{\max}^2 \beta \gamma)^k (H(z_0) - H(z^*)).$$

B.2 Backward Discretization of Skew-Gradient Flow

Consider the forward method to discretize the skew-gradient flow (9) with step size $\eta > 0$:

$$z_{k+1} = z_k - \eta J \nabla H(z_{k+1}). \quad (18)$$

This is an implicit update, and may require some computation to solve in each iteration. If H is smooth, then the update is well-defined for small enough step size.

In terms of the components $z_k = (x_k, y_k)$, this corresponds to the simultaneous backward method:

$$\begin{aligned} x_{k+1} &= x_k - \eta A \nabla g(y_{k+1}) \\ y_{k+1} &= y_k + \eta A^\top \nabla f(x_{k+1}). \end{aligned}$$

For the min-max game application when $f = \nabla \phi^*$ and $g = \psi^*$, this corresponds to the two players following the simultaneous proximal mirror descent algorithm:

$$\begin{aligned} p_{k+1} &= \arg \min_{p \in \mathcal{P}} \left\{ p^\top A q_{k+1} + \frac{1}{\eta} D_\phi(p, p_k) \right\} \\ q_{k+1} &= \arg \min_{q \in \mathcal{Q}} \left\{ -p_{k+1}^\top A q + \frac{1}{\eta} D_\psi(q, q_k) \right\}. \end{aligned}$$

We show the following properties of the backward method, which are the opposite of the properties of the forward method. Here let $\|\cdot\| = \|\cdot\|_2$ be the ℓ_2 -norm.

Theorem B.3. *Let z_k evolve following the backward method (18). We have the following properties:*

1. *If H is L_2 -smooth and $\eta < \frac{1}{\alpha_{\max} L_2}$, then the backward method is well-defined, i.e. the solution z_{k+1} exists and is unique for any z_k .*
2. *If H is convex, then $H(z_{k+1}) \leq H(z_k)$.*
3. *Assume $m = n$, and assume $A \in \mathbb{R}^{n \times n}$ is invertible with smallest singular value $\alpha_{\min} > 0$. Assume H is μ -strongly convex with respect, and let $z^* = \arg \min_z H(z)$. Then:*

$$H(z_k) - H(z^*) \leq \frac{H(z_0) - H(z^*)}{(1 + \eta^2 \alpha_{\min}^2 \mu^2)^k}.$$

4. *Assume H is convex and L_2 -smooth. Then:*

$$H(z_k) - H(z^*) \geq \frac{H(z_0) - H(z^*)}{(1 + \eta^2 \alpha_{\max}^2 L^2)^k}$$

As in Theorem B.1, the results in Theorem B.3 are stated in terms of the ℓ_2 -norm for simplicity, but also hold for any norm $\|\cdot\|$ in $\mathbb{R}^{m+n}$, at the cost of a potentially dimension-dependence constants. We note that we can also deduce Theorem B.3 by reversing the argument of Theorem B.1, since the backward method is the adjoint of the forward method (i.e. the inverse with negated step size $-\eta$). For completeness, we provide the full proof below. We can also check that in the quadratic case, the rates in Theorem B.3 are tight.

Proof of Theorem B.3.

1. We can write the update (18) as $(I + \eta J \nabla H)(z_{k+1}) = z_k$, where I is the identity operator. The Jacobian matrix of the operator $I + \eta J \nabla H$ is $I_{m+n} + \eta J \nabla^2 H$, which has smallest singular value at least $1 - \eta \alpha_{\max} \|\nabla^2 H\|_{\text{op}} \geq 1 - \eta \alpha_{\max} L_2$. We see that if $\eta < \frac{1}{\alpha_{\max} L_2}$, then the Jacobian matrix is non-singular, thus the operator $I + \eta J \nabla H$ is invertible, and hence $z_{k+1} = (I + \eta J \nabla H)^{-1}(z_k)$ exists and is uniquely defined.
2. Assume H is convex. By Jensen's inequality,
$$H(z_{k+1}) - H(z_k) \leq \langle \nabla H(z_{k+1}), z_{k+1} - z_k \rangle = -\eta \langle \nabla H(z_{k+1}), J \nabla H(z_{k+1}) \rangle = 0.$$
This shows the backward method always decreases the energy function: $H(z_{k+1}) \leq H(z_k)$.
3. Assume $m = n$, and $A \in \mathbb{R}^{n \times n}$ has a positive smallest singular value $\alpha_{\min} > 0$. This implies $A^\top A \succeq \alpha_{\min}^2 I_n$ and $AA^\top \succeq \alpha_{\min}^2 I_n$, so $J^\top J \succeq \alpha_{\min}^2 I_{2n}$.

Assume further that H is μ -strongly convex, which implies H is μ -gradient dominated:

$$\|\nabla H(z)\|_2^2 \geq 2\mu(H(z) - H(z^*))$$

for all $z \in \mathbb{R}^{2n}$. Then by the μ -strong convexity of H , along the backward method,

$$\begin{aligned} H(z_{k+1}) - H(z_k) &\leq \langle \nabla H(z_{k+1}), z_{k+1} - z_k \rangle - \frac{\mu}{2} \|z_{k+1} - z_k\|_2^2 \\ &= -\eta \langle \nabla H(z_{k+1}), J \nabla H(z_{k+1}) \rangle - \frac{\eta^2 \mu}{2} \|J \nabla H(z_{k+1})\|_2^2 \\ &\leq -\frac{\eta^2 \alpha_{\min}^2 \mu}{2} \|\nabla H(z_{k+1})\|_2^2 \\ &\leq -\eta^2 \alpha_{\min}^2 \mu^2 (H(z_{k+1}) - H(z^*)) \end{aligned}$$

This implies

$$H(z_{k+1}) - H(z^*) \leq \frac{H(z_k) - H(z^*)}{1 + \eta^2 \alpha_{\min}^2 \mu^2}.$$

Therefore, the backward method decreases the energy function exponentially fast:

$$H(z_k) - H(z^*) \leq \frac{H(z_0) - H(z^*)}{(1 + \eta^2 \alpha_{\min}^2 \mu^2)^k}.$$

4. Since A has maximum singular value $\alpha_{\max} > 0$, we have $A^\top A \preceq \alpha_{\max}^2 I_n$ and $AA^\top \preceq \alpha_{\max}^2 I_m$, so $J^\top J \preceq \alpha_{\max}^2 I_{m+n}$.

Since H is convex and L -smooth, we have for all $z \in \mathbb{R}^{m+n}$:

$$H(z) \geq H(z^*) + \frac{1}{2L} \|\nabla H(z)\|_2^2.$$

Then by the L_2 -smoothness of H , along the backward method,

$$\begin{aligned} H(z_{k+1}) - H(z_k) &\geq \langle \nabla H(z_{k+1}), z_{k+1} - z_k \rangle - \frac{L_2}{2} \|z_{k+1} - z_k\|_2^2 \\ &= -\eta \langle \nabla H(z_{k+1}), J \nabla H(z_{k+1}) \rangle - \frac{\eta^2 L_2}{2} \|J \nabla H(z_{k+1})\|_2^2 \\ &= 0 - \frac{\eta^2 L_2}{2} \langle \nabla H(z_{k+1}), (J^\top J) \nabla H(z_{k+1}) \rangle \\ &\geq -\frac{\eta^2 \alpha_{\max}^2 L_2}{2} \|\nabla H(z_{k+1})\|_2^2 \\ &\geq -\eta^2 \alpha_{\max}^2 L_2^2 (H(z_{k+1}) - H(z^*)) \end{aligned}$$

This implies

$$H(z_{k+1}) - H(z^*) \geq \frac{H(z_k) - H(z^*)}{1 + \eta^2 \alpha_{\max}^2 L_2^2}.$$

Therefore, along the backward method, the energy function decreases at most exponentially fast:

$$H(z_k) - H(z^*) \geq \frac{H(z_0) - H(z^*)}{(1 + \eta^2 \alpha_{\max}^2 L_2^2)^k}.$$

□

In Section C.4 we analyze the total regret of the backward method for the min-max game application, and show the duality gap of the average iterate decays at an $O(K^{-1})$ rate; see Theorem C.3. This is similar to the performance of the clairvoyant multiplicative weight update [25] which achieves constant regret for a general game.

B.3 Alternating Discretization of Skew-Gradient Flow

B.3.1 Proof of Corollary 4.3

Proof of Corollary 4.3. From Lemma 4.2 and since the Bregman commutator of H vanishes,

$$H_\eta(z_k) = H_\eta(z_0) + \sum_{i=0}^{k-1} C_H(z_{i+1}, z_i) = \tilde{H}_\eta(z_0).$$

□

B.3.2 Proof of Lemma 4.2

Proof of Lemma 4.2. By the definition of Bregman commutator, and using the alternating update, we can write

$$\begin{aligned} C_f(x_{k+1}, x_k) &= f(x_{k+1}) - f(x_k) - \frac{1}{2} \langle \nabla f(x_k) + \nabla f(x_{k+1}), x_{k+1} - x_k \rangle \\ &= f(x_{k+1}) - f(x_k) + \frac{\eta}{2} \langle \nabla f(x_k) + \nabla f(x_{k+1}), A \nabla g(y_k) \rangle. \end{aligned}$$

Similarly, we can also write

$$\begin{aligned} C_g(y_{k+1}, y_k) &= g(y_{k+1}) - g(y_k) - \frac{1}{2} \langle \nabla g(y_k) + \nabla g(y_{k+1}), y_{k+1} - y_k \rangle \\ &= g(y_{k+1}) - g(y_k) - \frac{\eta}{2} \langle \nabla g(y_k) + \nabla g(y_{k+1}), A^\top \nabla f(x_{k+1}) \rangle \\ &= g(y_{k+1}) - g(y_k) - \frac{\eta}{2} \langle \nabla f(x_{k+1}), A(\nabla g(y_k) + \nabla g(y_{k+1})) \rangle. \end{aligned}$$

Since $H(x, y) = f(x) + g(y)$ is separable, we have $D_H(z', z) = D_f(x', x) + D_g(y', y)$ where $z' = (x', y')$ and $z = (x, y)$, and thus $C_H(z', z) = C_f(x', x) + C_g(y', y)$. Adding the two identities above yields

$$\begin{aligned} C_H(z_{k+1}, z_k) &= C_f(x_{k+1}, x_k) + C_g(y_{k+1}, y_k) \\ &= f(x_{k+1}) - f(x_k) + \frac{\eta}{2} \langle \nabla f(x_k) + \nabla f(x_{k+1}), A \nabla g(y_k) \rangle \\ &\quad + g(y_{k+1}) - g(y_k) - \frac{\eta}{2} \langle \nabla f(x_{k+1}), A(\nabla g(y_k) + \nabla g(y_{k+1})) \rangle \\ &= H(z_{k+1}) - H(z_k) + \frac{\eta}{2} \langle \nabla f(x_k), A \nabla g(y_k) \rangle - \frac{\eta}{2} \langle \nabla f(x_{k+1}), A \nabla g(y_{k+1}) \rangle \\ &= H_\eta(z_{k+1}) - H_\eta(z_k) \end{aligned}$$

as desired. □

B.3.3 Proof of Theorem 4.4

Proof of Theorem 4.4. From Lemma 4.2 and using Lemma A.2,

$$|H_\eta(z_k) - H_\eta(z_0)| = \left| \sum_{i=0}^{k-1} C_H(z_{i+1}, z_i) \right| \leq \sum_{i=0}^{k-1} |C_H(z_{i+1}, z_i)| \leq \frac{L_3}{12} \sum_{i=0}^{k-1} \|z_{i+1} - z_i\|^3.$$

Along the alternating method (11), we have

$$z_{i+1} - z_i = \eta \left(\frac{-A \nabla g(y_i)}{A^\top \nabla f(x_{i+1})} \right) = -\eta J \nabla H(z_{i+\frac{1}{2}})$$

where $z_{i+\frac{1}{2}} := \begin{pmatrix} x_{i+1} \\ y_i \end{pmatrix}$. Thus,

$$\|z_{i+1} - z_i\| = \eta \|J \nabla H(z_{i+\frac{1}{2}})\| \leq \eta \alpha_{\max} L_1.$$

Therefore,

$$|H_\eta(z_k) - H_\eta(z_0)| \leq \frac{L_3}{12} \sum_{i=0}^{k-1} (\eta \alpha_{\max} L_1)^3 = \frac{1}{12} \eta^3 \alpha_{\max}^3 L_1^3 L_3 k$$

as desired. $\square$

B.3.4 Proof of Theorem 4.5

Proof of Theorem 4.5. Recall from (6) the regret of the first player is

$$R_{1,K}(p) = \sum_{k=0}^{K-1} \left(\frac{p_k + p_{k+1}}{2} \right)^\top A q_k - \sum_{k=0}^{K-1} p^\top A q_k.$$

Recall also $p_k = \nabla f(x_k)$, $q_k = \nabla g(y_k)$, $p = \nabla f(x)$, and $q = \nabla g(y)$.

From the alternating mirror descent update (3), we have

$$A q_k = A \nabla g(y_k) = -\frac{1}{\eta} (x_{k+1} - x_k).$$

Therefore, the second term in the first player's regret is

$$\sum_{k=0}^{K-1} p^\top A q_k = -\frac{1}{\eta} \sum_{k=0}^{K-1} \nabla f(x)^\top (x_{k+1} - x_k) = -\frac{1}{\eta} \nabla f(x)^\top (x_K - x_0).$$

For the first term in the first player's regret, by the definition of the Bregman commutator, we have

$$\begin{aligned} \sum_{k=0}^{K-1} \left(\frac{p_k + p_{k+1}}{2} \right)^\top A q_k &= -\frac{1}{\eta} \sum_{k=0}^{K-1} \left(\frac{\nabla f(x_{k+1}) + \nabla f(x_k)}{2} \right)^\top (x_{k+1} - x_k) \\ &= -\frac{1}{\eta} \sum_{k=0}^{K-1} (f(x_{k+1}) - f(x_k) - C_f(x_{k+1}, x_k)) \\ &= -\frac{1}{\eta} (f(x_K) - f(x_0)) + \frac{1}{\eta} \sum_{k=0}^{K-1} C_f(x_{k+1}, x_k). \end{aligned}$$

Combining the two terms above, we can write the regret of the first player as

$$\begin{aligned} R_{1,K}(p) &= -\frac{1}{\eta} (f(x_K) - f(x_0) - \nabla f(x)^\top (x_K - x_0)) + \frac{1}{\eta} \sum_{k=0}^{K-1} C_f(x_{k+1}, x_k) \\ &= -\frac{1}{\eta} (D_f(x_K, x) - D_f(x_0, x)) + \frac{1}{\eta} \sum_{k=0}^{K-1} C_f(x_{k+1}, x_k). \end{aligned} \tag{19}$$

Similarly, recall from (7) the regret of the second player is

$$R_{2,K}(q) = \sum_{k=0}^{K-1} p_{k+1}^\top Aq - \sum_{k=0}^{K-1} p_{k+1}^\top A \left(\frac{q_k + q_{k+1}}{2} \right).$$

From the alternating mirror descent update (3), we have

$$A^\top p_{k+1} = A^\top \nabla f(x_{k+1}) = \frac{1}{\eta} (y_{k+1} - y_k).$$

Therefore, the first term in the second player's regret is

$$\sum_{k=0}^{K-1} p_{k+1}^\top Aq = \frac{1}{\eta} \sum_{k=0}^{K-1} (y_{k+1} - y_k)^\top \nabla g(y) = \frac{1}{\eta} (y_K - y_0)^\top \nabla g(y).$$

For the second term in the second player's regret, by the definition of the Bregman commutator, we have

$$\begin{aligned} \sum_{k=0}^{K-1} p_{k+1}^\top A \left(\frac{q_k + q_{k+1}}{2} \right) &= \frac{1}{\eta} \sum_{k=0}^{K-1} (y_{k+1} - y_k)^\top \left(\frac{\nabla g(y_k) + \nabla g(y_{k+1})}{2} \right) \\ &= \frac{1}{\eta} \sum_{k=0}^{K-1} (g(y_{k+1}) - g(y_k) - C_g(y_{k+1}, y_k)) \\ &= \frac{1}{\eta} (g(y_K) - g(y_0)) - \frac{1}{\eta} \sum_{k=0}^{K-1} C_g(y_{k+1}, y_k). \end{aligned}$$

Combining the two terms above, we can write the regret of the second player as

$$\begin{aligned} R_{2,K}(q) &= \frac{1}{\eta} (-g(y_K) + g(y_0) + (y_K - y_0)^\top \nabla g(y)) + \frac{1}{\eta} \sum_{k=0}^{K-1} C_g(y_{k+1}, y_k) \\ &= -\frac{1}{\eta} (D_g(y_K, y) - D_g(y_0, y)) + \frac{1}{\eta} \sum_{k=0}^{K-1} C_g(y_{k+1}, y_k). \end{aligned} \tag{20}$$

Adding (19) and (20), we find that the cumulative regret of both players at iteration K is

$$\begin{aligned} R_K(p, q) &= R_{1,K}(p) + R_{2,K}(q) \\ &= -\frac{1}{\eta} (D_f(x_K, x) - D_f(x_0, x)) + \frac{1}{\eta} \sum_{k=0}^{K-1} C_f(x_{k+1}, x_k) \\ &\quad - \frac{1}{\eta} (D_g(y_K, y) - D_g(y_0, y)) + \frac{1}{\eta} \sum_{k=0}^{K-1} C_g(y_{k+1}, y_k) \\ &= -\frac{1}{\eta} (D_H(z_K, z) - D_H(z_0, z)) + \frac{1}{\eta} \sum_{k=0}^{K-1} C_H(z_{k+1}, z_k) \\ &\stackrel{(13)}{=} -\frac{1}{\eta} (D_H(z_K, z) - D_H(z_0, z)) + \frac{1}{\eta} \sum_{k=0}^{K-1} (H_\eta(z_{k+1}) - H_\eta(z_k)) \\ &= -\frac{1}{\eta} (D_H(z_K, z) - D_H(z_0, z)) + \frac{1}{\eta} (H_\eta(z_K) - H_\eta(z_0)). \end{aligned}$$

In the penultimate line above we have applied (13) from Lemma 4.2 to express the Bregman commutator as a difference of the modified energy. $\square$

C Details for Regret Calculations

C.1 Regret Proofs for Alternating Mirror Descent

C.1.1 Proof of Lemma 3.1

Proof of Lemma 3.1. By the definition of duality gap and total regret of the alternating method,

$$\begin{aligned}
K \cdot \text{dg}(\bar{p}_K, \bar{q}_K) &= K \cdot \max_{(p,q) \in \mathcal{P} \times \mathcal{Q}} (\bar{p}_K^\top Aq - p^\top A\bar{q}_K) \\
&= \max_{(p,q) \in \mathcal{P} \times \mathcal{Q}} \left(\sum_{k=0}^{K-1} p_{k+1}^\top Aq - \sum_{k=0}^{K-1} p^\top Aq_k \right) \\
&= \max_{(p,q) \in \mathcal{P} \times \mathcal{Q}} \left(R_{1,K}(p) + R_{2,K}(q) - \frac{1}{2}(p_0^\top Aq_0 - p_K^\top Aq_K) \right) \\
&= R_K - \frac{1}{2}(p_0^\top Aq_0 - p_K^\top Aq_K).
\end{aligned}$$

□

C.1.2 Proof of Theorem 3.2

Proof of Theorem 3.2. We use the relation between cumulative regret of alternating mirror descent and the modified energy from the skew-gradient flow discretization in Theorem 4.5.

For any reference point $(p, q) \in \mathcal{P} \times \mathcal{Q}$, let $z = (x, y) \in \mathbb{R}^{m+n}$ be the dual variables, where $x = \nabla\phi(p)$ and $y = \nabla\psi(q)$. Recall we define the energy function $H(x, y) = f(x) + g(y)$ where $f = \phi^*$ and $g = \psi^*$. Along the alternating mirror descent (1), let (x_k, y_k) be the dual variables to (p_k, q_k) . By Theorem 4.5, we have

$$\begin{aligned}
R_K(p, q) &= \frac{D_H(z_0, z) - D_H(z_K, z)}{\eta} + \frac{H_\eta(z_K) - H_\eta(z_0)}{\eta} \\
&\leq \frac{D_H(z_0, z)}{\eta} + \frac{H_\eta(z_K) - H_\eta(z_0)}{\eta}.
\end{aligned} \tag{21}$$

The inequality in the last line above follows from $D_H(z_K, z) \geq 0$ since H is convex.

We verify that the assumptions in Theorem 3.2 imply the assumptions in Theorem 4.4 are satisfied. Indeed, H is $2M$ -Lipschitz since its gradient is $\nabla H(x, y) = (\nabla f(x), \nabla g(y)) = (p, q)$, so

$$\|\nabla H(x, y)\| = \|(p, q)\| \leq \|p\| + \|q\| \leq M + M = 2M.$$

Furthermore, H is $2M$ -smooth of order 3 since

$$\nabla^3 H(x, y) = (\nabla^3 f(x), \nabla^3 g(y)) = (\nabla^3 \phi^*(\nabla\phi(p)), \nabla^3 \psi^*(\nabla\psi(q)))$$

so $\|\nabla^3 H(x, y)\|_{\text{op}} \leq \|\nabla^3 \phi^*(\nabla\phi(p))\|_{\text{op}} + \|\nabla^3 \psi^*(\nabla\psi(q))\|_{\text{op}} \leq M + M = 2M$.

Then by Theorem 4.4,

$$|H_\eta(z_K) - H_\eta(z_0)| \leq \frac{\alpha_{\max}^3 (2M) (2M)^3}{12} \eta^3 K = \frac{4}{3} \alpha_{\max}^3 M^4 \eta^3 K.$$

Plugging this to (21) gives

$$R_K(p, q) \leq \frac{D_H(z_0, z)}{\eta} + \frac{4}{3} \alpha_{\max}^3 M^4 \eta^2 K.$$

Since

$$D_H(z_0, z) = D_f(x_0, x) + D_g(y_0, y) = D_\phi(p, p_0) + D_\psi(q, q_0)$$

we also have

$$R_K(p, q) \leq \frac{D_\phi(p, p_0) + D_\psi(q, q_0)}{\eta} + \frac{4}{3} \alpha_{\max}^3 M^4 \eta^2 K. \tag{22}$$

Now suppose we start from (p_0, q_0) such that $\sup_{p \in \mathcal{P}} D_\phi(p, p_0) \leq M_2$ and $\sup_{q \in \mathcal{Q}} D_\psi(q, q_0) \leq M_2$ for some $0 < M_2 < \infty$. Then we can maximize (22) over $(p, q) \in \mathcal{P} \times \mathcal{Q}$ to get the total regret

$$R_K = \sup_{(p,q) \in \mathcal{P} \times \mathcal{Q}} R_K(p, q) \leq \frac{2M_2}{\eta} + \frac{4}{3} \alpha_{\max}^3 M^4 \eta^2 K.$$

With the choice of step size $\eta = \Theta(K^{-1/3})$, we obtain the bound $R_K = O(K^{1/3})$. $\square$

C.1.3 Proof of Corollary 3.3

Proof of Corollary 3.3. By assumption and Theorem 3.2, we have $R_K = O(K^{1/3})$, so $\frac{1}{K} R_K = O(K^{-2/3})$. Since we assume the domains are bounded, we have

$$|p_0^\top A q_0| \leq \alpha_{\max} \|p_0\| \cdot \|q_0\| \leq \alpha_{\max} M^2$$

and similarly, $|p_K^\top A q_K| \leq \alpha_{\max} M^2$. Thus, the last term in (8) is bounded by

$$\frac{1}{K} |p_0^\top A q_0 - p_K^\top A q_K| \leq \frac{2\alpha_{\max} M^2}{K} = O(K^{-1}).$$

Therefore, by Lemma 3.1, $\text{dg}(\bar{p}_K, \bar{q}_K) = O(K^{-2/3}) + O(K^{-1}) = O(K^{-2/3})$. $\square$

C.2 Regret of the Continuous-time Method

Suppose the two players follow the continuous-time variant of mirror descent, namely the natural gradient flow:

$$\begin{aligned} \dot{P}_t &= -\nabla^2 \phi(P_t)^{-1} A Q_t \\ \dot{Q}_t &= \nabla^2 \psi(Q_t)^{-1} A^\top P_t. \end{aligned}$$

This is the continuous-time limit (as the step size $\eta \rightarrow 0$) of the mirror descent method—either the simultaneous, alternating, or proximal version. In terms of the dual variables $X_t = \nabla \phi(P_t)$, $Y_t = \nabla \psi(Q_t)$, they follow the skew-gradient flow (4):

$$\begin{aligned} \dot{X}_t &= -A \nabla g(Y_t) \\ \dot{Y}_t &= A^\top \nabla f(X_t) \end{aligned}$$

where $f = \phi^*$ and $g = \psi^*$. Recall in this case the energy function $H(x, y) = f(x) + g(y)$ is conserved: $H(X_t, Y_t) = H(X_0, Y_0)$, which means the trajectories (X_t, Y_t) cycle and stay in the orbit of constant energy.

For the min-max game application, we show this continuous-time method achieves constant regret.

We define the regret of the first player at (continuous) time T with respect to $\hat{p} \in \mathcal{P}$ to be:

$$R_{1,T}(\hat{p}) := \int_0^T P_t^\top A Q_t dt - \int_0^T \hat{p}^\top A Q_t dt.$$

Similarly, we define the regret of the second player at time T with respect to $\hat{q} \in \mathcal{Q}$ to be:

$$R_{2,T}(\hat{q}) := \int_0^T P_t^\top A \hat{q} dt - \int_0^T P_t^\top A Q_t dt.$$

We define the cumulative regret of both players at time T with respect to $(\hat{p}, \hat{q}) \in \mathcal{P} \times \mathcal{Q}$ to be:

$$R_T(\hat{p}, \hat{q}) := R_{1,T}(\hat{p}) + R_{2,T}(\hat{q}) = \int_0^T P_t^\top A \hat{q} dt - \int_0^T \hat{p}^\top A Q_t dt.$$

Finally, we define the **total regret** at time T to be the maximum cumulative regret of both players:

$$R_T := \max_{(\hat{p}, \hat{q}) \in \mathcal{P} \times \mathcal{Q}} R_T(\hat{p}, \hat{q}).$$

Define the average iterates of the two players at time T :

$$\bar{P}_T = \frac{1}{T} \int_0^T P_t dt \quad \text{and} \quad \bar{Q}_T = \frac{1}{T} \int_0^T Q_t dt$$

We observe that the average total regret is equal to the duality gap of the average iterates:

$$\frac{1}{T} R_T = \max_{q \in \mathcal{Q}} \bar{P}_T^\top A q - \min_{p \in \mathcal{P}} p^\top A \bar{Q}_T = \text{dg}(\bar{p}_T, \bar{q}_T).$$

Then we have the following properties.

Theorem C.1. *Assume the domains $\mathcal{P}$ and $\mathcal{Q}$ are bounded, so there exists $0 < M < \infty$ such that $D_\phi(p', p) \leq M$ and $D_\psi(q', q) \leq M$ for all $p', p \in \mathcal{P}$, $q', q \in \mathcal{Q}$. If both players play the natural gradient flow dynamics (23), then the total regret is bounded for all $T \geq 0$:*

$$R_T \leq 2M.$$

Therefore, the duality gap of the average iterate decays at an $O(1/T)$ rate:

$$\text{dg}(\bar{P}_T, \bar{Q}_T) \leq \frac{2M}{T}.$$

Proof. Let $\hat{x} = \nabla \phi(\hat{p})$ and $\hat{y} = \nabla \psi(\hat{q})$, so $\hat{p} = \nabla f(\hat{x})$ and $\hat{q} = \nabla g(\hat{y})$. By the definition of the dynamics, we have

$$\begin{aligned} (P_t - \hat{p})^\top A Q_t &= -(\nabla f(X_t) - \nabla f(\hat{x}))^\top \dot{X}_t \\ &= \frac{d}{dt}(-f(X_t) + \langle \nabla f(\hat{x}), X_t - \hat{x} \rangle) \\ &= -\frac{d}{dt} D_f(X_t, \hat{x}). \end{aligned}$$

Therefore, we can write the regret of the first player as

$$\begin{aligned} R_{1,T}(\hat{p}) &= \int_0^T (P_t - \hat{p})^\top A Q_t dt \\ &= D_f(X_0, \hat{x}) - D_f(X_T, \hat{x}) \\ &= D_\phi(\hat{p}, P_0) - D_\phi(\hat{p}, P_T). \end{aligned}$$

Similarly, by the definition of the dynamics, we have

$$\begin{aligned} P_t^\top A(\hat{q} - Q_t) &= \dot{Y}_t^\top (\nabla g(\hat{y}) - \nabla g(Y_t))^\top \\ &= \frac{d}{dt}(-g(Y_t) + \langle \nabla g(\hat{y}), Y_t - \hat{y} \rangle) \\ &= -\frac{d}{dt} D_g(Y_t, \hat{y}). \end{aligned}$$

Therefore, we can write the regret of the second player as

$$\begin{aligned} R_{2,T}(\hat{q}) &= \int_0^T P_t^\top A(\hat{q} - Q_t) dt \\ &= D_g(Y_0, \hat{y}) - D_g(Y_T, \hat{y}) \\ &= D_\psi(\hat{q}, Q_0) - D_\psi(\hat{q}, Q_T). \end{aligned}$$

Combining the two quantities above, and since ϕ, ψ are convex, we find that

$$\begin{aligned} R_T(\hat{p}, \hat{q}) &= D_\phi(\hat{p}, P_0) - D_\phi(\hat{p}, P_T) + D_\psi(\hat{q}, Q_0) - D_\psi(\hat{q}, Q_T) \\ &\leq D_\phi(\hat{p}, P_0) + D_\psi(\hat{q}, Q_0) \\ &\leq 2M. \end{aligned}$$

Therefore the total regret is also bounded:

$$R_T = \max_{(\hat{p}, \hat{q}) \in \mathcal{P} \times \mathcal{Q}} R_T(\hat{p}, \hat{q}) \leq 2M.$$

Thus, the duality gap of the average iterate is bounded by:

$$\text{dg}(\bar{P}_T, \bar{Q}_T) = \frac{1}{T} R_T \leq \frac{2M}{T}.$$

□

C.3 Regret of the Forward Method

Consider the forward method (17) for discretizing the skew-gradient flow (4), which means the two players follow the simultaneous mirror descent algorithm. Recall from Theorem B.1 that the energy function is increasing along the forward method. We show that the total regret of the two players increases at an $O(K^{1/2})$ rate, and thus the average total regret decays as $O(K^{-1/2})$ after K iterations, which recovers the classical rate of simultaneous mirror descent.

We define the regret of the first player at iteration K with respect to $\hat{p} \in \mathcal{P}$ to be:

$$R_{1,K}(\hat{p}) := \sum_{k=0}^{K-1} p_k^\top A q_k - \sum_{k=0}^{K-1} \hat{p}^\top A q_k.$$

Similarly, we define the regret of the second player at iteration K with respect to $\hat{q} \in \mathcal{Q}$ to be:

$$R_{2,K}(\hat{q}) := \sum_{k=0}^{K-1} p_k^\top A \hat{q} - \sum_{k=0}^{K-1} p_k^\top A q_k.$$

We define the cumulative regret of both players at iteration K with respect to $(\hat{p}, \hat{q}) \in \mathcal{P} \times \mathcal{Q}$ to be:

$$R_K(\hat{p}, \hat{q}) := R_{1,K}(\hat{p}) + R_{2,K}(\hat{q}) = \sum_{k=0}^{K-1} p_k^\top A \hat{q} - \sum_{k=0}^{K-1} \hat{p}^\top A q_k.$$

Finally, we define the **total regret** at iteration K to be the maximum cumulative regret of both players:

$$R_K := \max_{(\hat{p}, \hat{q}) \in \mathcal{P} \times \mathcal{Q}} R_K(\hat{p}, \hat{q}).$$

Define the average iterates of the two players at iteration K :

$$\bar{p}_K = \frac{1}{K} \sum_{k=0}^{K-1} p_k \quad \text{and} \quad \bar{q}_K = \frac{1}{K} \sum_{k=0}^{K-1} q_k.$$

We observe that the average regret is equal to the duality gap of the average iterates:

$$\frac{1}{K} R_K = \max_{q \in \mathcal{Q}} \bar{p}_K^\top A q - \min_{p \in \mathcal{P}} p^\top A \bar{q}_K = \text{dg}(\bar{p}_K, \bar{q}_K).$$

Then we have the following properties.

Theorem C.2. Assume the domains $\mathcal{P}$ and $\mathcal{Q}$ are bounded, so there exists $0 < M < \infty$ such that $D_\phi(p', p) \leq M$ and $D_\psi(q', q) \leq M$ for all $p', p \in \mathcal{P}$, $q', q \in \mathcal{Q}$. Assume the energy function $H = f + g = \phi^* + \psi^*$ is L_1 -Lipschitz and L_2 -smooth. Let $\alpha_{\max}$ be the largest singular value of A . If both players play the forward method (17), then the total regret is bounded by:

$$R_K \leq \frac{1}{2} \eta \alpha_{\max}^2 L_1^2 L_2 K + \frac{2M}{\eta}.$$

In particular, choosing $\eta = \Theta(K^{-1/2})$ yields the bound $R_K \leq O(K^{1/2})$, and therefore,

$$\text{dg}(\bar{p}_K, \bar{q}_K) \leq O(K^{-1/2}).$$

Proof. Let $\hat{x} = \nabla \phi(\hat{p})$ and $\hat{y} = \nabla \psi(\hat{q})$, so $\hat{p} = \nabla f(\hat{x})$ and $\hat{q} = \nabla g(\hat{y})$. From the forward method update (17),

$$\begin{aligned} (p_k - \hat{p})^\top A q_k &= -\frac{1}{\eta} (\nabla f(x_k) - \nabla f(\hat{x}))^\top (x_{k+1} - x_k) \\ &= \frac{1}{\eta} (D_f(x_{k+1}, x_k) + D_f(x_k, \hat{x}) - D_f(x_{k+1}, \hat{x})) \\ &= \frac{1}{\eta} (D_\phi(p_k, p_{k+1}) + D_\phi(\hat{p}, p_k) - D_\phi(\hat{p}, p_{k+1})). \end{aligned}$$

Therefore,

$$\begin{aligned} R_{1,K}(\hat{p}) &= \frac{1}{\eta} \sum_{k=0}^{K-1} D_\phi(p_k, p_{k+1}) + \frac{1}{\eta} (D_\phi(\hat{p}, p_0) - D_\phi(\hat{p}, p_K)) \\ &\leq \frac{1}{\eta} \sum_{k=0}^{K-1} D_\phi(p_k, p_{k+1}) + \frac{1}{\eta} D_\phi(\hat{p}, p_0) \end{aligned}$$

where the last inequality follows since ϕ is convex.

Similarly, from the forward method update (17),

$$\begin{aligned} p_k^\top A(\hat{q} - q_k) &= \frac{1}{\eta} (y_{k+1} - y_k)^\top (\nabla g(\hat{y}) - \nabla g(y_k)) \\ &= \frac{1}{\eta} (D_g(y_{k+1}, y_k) + D_g(y_k, \hat{y}) - D_g(y_{k+1}, \hat{y})) \\ &= \frac{1}{\eta} (D_\psi(q_k, q_{k+1}) + D_\psi(\hat{q}, q_k) - D_\psi(\hat{q}, q_{k+1})). \end{aligned}$$

Therefore,

$$\begin{aligned} R_{2,K}(\hat{q}) &= \frac{1}{\eta} \sum_{k=0}^{K-1} D_\psi(q_k, q_{k+1}) + \frac{1}{\eta} (D_\psi(\hat{q}, q_0) - D_\psi(\hat{q}, q_K)) \\ &\leq \frac{1}{\eta} \sum_{k=0}^{K-1} D_\psi(q_k, q_{k+1}) + \frac{1}{\eta} D_\psi(\hat{q}, q_0) \end{aligned}$$

where the last inequality follows since ψ is convex.

Combining the two quantities above, we have

$$\begin{aligned} R_K(\hat{p}, \hat{q}) &\leq \frac{1}{\eta} \sum_{k=0}^{K-1} (D_\phi(p_k, p_{k+1}) + D_\psi(q_k, q_{k+1})) + \frac{1}{\eta} (D_\phi(\hat{p}, p_0) + D_\psi(\hat{q}, q_0)) \\ &= \frac{1}{\eta} \sum_{k=0}^{K-1} D_H(z_{k+1}, z_k) + \frac{1}{\eta} D_H(z_0, \hat{z}). \end{aligned}$$

Since $z_{k+1} = z_k - \eta J \nabla H(z_k)$, and H is L_1 -Lipschitz and L_2 -smooth, we can bound:

$$\begin{aligned} D_H(z_{k+1}, z_k) &\leq \frac{L_2}{2} \|z_{k+1} - z_k\|^2 \\ &= \frac{\eta^2 L_2}{2} \|J \nabla H(z_k)\|^2 \\ &\leq \frac{\eta^2 \alpha_{\max}^2 L_2}{2} \|\nabla H(z_k)\|^2 \\ &\leq \frac{\eta^2 \alpha_{\max}^2 L_2 L_1^2}{2}. \end{aligned}$$

This gives a regret bound:

$$\begin{aligned} R_K(\hat{p}, \hat{q}) &\leq \frac{\eta K \alpha_{\max}^2 L_2 L_1^2}{2} + \frac{1}{\eta} D_H(z_0, \hat{z}) \\ &\leq \frac{\eta K \alpha_{\max}^2 L_2 L_1^2}{2} + \frac{2M}{\eta} \end{aligned}$$

where the last inequality follows since $D_H(z_0, \hat{z}) = D_f(x_0, \hat{x}) + D_g(y_0, \hat{y}) \leq 2M$. Thus, the total regret is also bounded:

$$R_K \leq \frac{\eta K \alpha_{\max}^2 L_2 L_1^2}{2} + \frac{2M}{\eta}.$$

Choosing $\eta = \frac{2\sqrt{M}}{\alpha_{\max} L_1 \sqrt{L_2 K}} = \Theta(K^{-1/2})$ gives

$$R_K \leq 2\alpha_{\max} L_1 \sqrt{L_2 M K} = O(K^{1/2})$$

Therefore, the duality gap of the average iterates is bounded by $\text{dg}(\bar{p}_K, \bar{q}_K) = \frac{1}{K} R_K \leq O(K^{-1/2})$, as desired. $\square$

C.4 Regret of the Backward Method

Consider the backward method (18) for discretizing the skew-gradient flow (4), which means the two players follow the simultaneous proximal mirror descent algorithm. Recall from Theorem B.3 that the energy function is decreasing along the forward method. We show that the total regret of the two players along the backward method is bounded by a constant, and thus the average total regret decays as $O(K^{-1})$ after K iterations. This is similar to the property of the clairvoyant multiplicative weight update [25], which achieves a constant regret bound for general games.

We define the regret of the first player at iteration K with respect to $\hat{p} \in \mathcal{P}$ to be:

$$R_{1,K}(\hat{p}) := \sum_{k=0}^{K-1} p_{k+1}^\top A q_{k+1} - \sum_{k=0}^{K-1} \hat{p}^\top A q_{k+1}.$$

Note the difference with the forward method; here we are shifting the time index by 1. Similarly, We define the regret of the second player (q) at the beginning of iteration K (at time $T = \eta K$) with respect to a static point $\hat{q} \in \mathcal{Q}$ to be:

$$R_{2,K}(\hat{q}) := \sum_{k=0}^{K-1} p_{k+1}^\top A \hat{q} - \sum_{k=0}^{K-1} p_{k+1}^\top A q_{k+1}.$$

We define the cumulative regret of both players at the iteration K with respect to $(\hat{p}, \hat{q}) \in \mathcal{P} \times \mathcal{Q}$ to be:

$$R_K(\hat{p}, \hat{q}) := R_{1,K}(\hat{p}) + R_{2,K}(\hat{q}) = \sum_{k=0}^{K-1} p_{k+1}^\top A \hat{q} - \sum_{k=0}^{K-1} \hat{p}^\top A q_{k+1}.$$

Finally, we define the **total regret** at iteration K to be the maximum cumulative regret of both players:

$$R_K := \max_{(\hat{p}, \hat{q}) \in \mathcal{P} \times \mathcal{Q}} R_K(\hat{p}, \hat{q}).$$

Define the average iterates of the two players at iteration K :

$$\bar{p}_K = \frac{1}{K} \sum_{k=0}^{K-1} p_{k+1} \quad \text{and} \quad \bar{q}_K = \frac{1}{K} \sum_{k=0}^{K-1} q_{k+1}.$$

We observe that the average regret is equal to the duality gap of the average iterates:

$$\frac{1}{K} R_K = \max_{q \in \mathcal{Q}} \bar{p}_K^\top A q - \min_{p \in \mathcal{P}} p^\top A \bar{q}_K = \text{dg}(\bar{p}_K, \bar{q}_K).$$

Then we have the following properties.

Theorem C.3. Assume the domains $\mathcal{P}$ and $\mathcal{Q}$ are bounded, so there exists $0 < M < \infty$ such that $D_\phi(p', p) \leq M$ and $D_\psi(q', q) \leq M$ for all $p', p \in \mathcal{P}$, $q', q \in \mathcal{Q}$. If both players play the backward method (18) for any step size $\eta > 0$, then the total regret is bounded by:

$$R_K \leq \frac{2M}{\eta}.$$

In particular, the duality gap decays as:

$$\text{dg}(\bar{p}_K, \bar{q}_K) \leq \frac{2M}{\eta K}.$$

Proof. Let $\hat{x} = \nabla\phi(\hat{p})$ and $\hat{y} = \nabla\psi(\hat{q})$, so $\hat{p} = \nabla f(\hat{x})$ and $\hat{q} = \nabla g(\hat{y})$. From the backward method update (18),

$$\begin{aligned} (p_{k+1} - \hat{p})^\top A q_{k+1} &= -\frac{1}{\eta} (\nabla f(x_{k+1}) - \nabla f(\hat{x}))^\top (x_{k+1} - x_k) \\ &= \frac{1}{\eta} (D_f(x_k, \hat{x}) - D_f(x_k, x_{k+1}) - D_f(x_{k+1}, \hat{x})) \\ &= \frac{1}{\eta} (D_\phi(\hat{p}, p_k) - D_\phi(p_{k+1}, p_k) - D_\phi(\hat{p}, p_{k+1})). \end{aligned}$$

Therefore,

$$\begin{aligned} R_{1,K}(\hat{p}) &= -\frac{1}{\eta} \sum_{k=0}^{K-1} D_\phi(p_k, p_{k+1}) + \frac{1}{\eta} (D_\phi(\hat{p}, p_0) - D_\phi(\hat{p}, p_K)) \\ &\leq \frac{1}{\eta} D_\phi(\hat{p}, p_0) \end{aligned}$$

where the last inequality follows since ϕ is convex.

Similarly, from the backward method update (18),

$$\begin{aligned} p_{k+1}^\top A(\hat{q} - q_{k+1}) &= \frac{1}{\eta} (y_{k+1} - y_k)^\top (\nabla g(\hat{y}) - \nabla g(y_{k+1})) \\ &= \frac{1}{\eta} (D_g(y_k, \hat{y}) - D_g(y_k, y_{k+1}) - D_g(y_{k+1}, \hat{y})) \\ &= \frac{1}{\eta} (D_\psi(\hat{q}, q_k) - D_\psi(q_{k+1}, q_k) - D_\psi(\hat{q}, q_{k+1})). \end{aligned}$$

Therefore,

$$\begin{aligned} R_{2,K}(\hat{q}) &= -\frac{1}{\eta} \sum_{k=0}^{K-1} D_\psi(q_{k+1}, q_k) + \frac{1}{\eta} (D_\psi(\hat{q}, q_0) - D_\psi(\hat{q}, q_K)) \\ &\leq \frac{1}{\eta} D_\psi(\hat{q}, q_0) \end{aligned}$$

where the last inequality follows since ψ is convex.

Combining the two quantities above, we have

$$\begin{aligned} R_K(\hat{p}, \hat{q}) &\leq \frac{1}{\eta} (D_\phi(\hat{p}, p_0) + D_\psi(\hat{q}, q_0)) \\ &\leq \frac{2M}{\eta}. \end{aligned}$$

Therefore,

$$\text{dg}(\bar{p}_K, \bar{q}_K) = \frac{1}{K} R_K \leq \frac{2M}{\eta K}.$$

□

D Toward a Better Regret Bound under Hamiltonian Structure

There are still some gaps between our results for alternating mirror descent for constrained min-max games and the results from [2] for alternating gradient descent for unconstrained min-max games. In particular, our bound in Theorem 4.4 still grows with the number of iterations, and does not yield constant regret, unlike in the unconstrained case. We believe that it is possible to close this gap in some particular settings, for example when the system has a Hamiltonian structure.

In the special case when $m = n$ and the payoff matrix is the identity matrix $A = I$, the skew-gradient flow dynamics (9) becomes a *Hamiltonian flow*; this means in addition to conserving the energy

function, the continuous dynamics also preserves the canonical symplectic structure. (When the payoff matrix A is arbitrary, the skew-gradient flow dynamics has a non-canonical Hamiltonian structure, for which some of the properties still apply; for simplicity, here we consider identity payoff matrix.) In this case, the update (2) for alternating mirror descent becomes what is known as the *symplectic Euler* discretization method, which also preserves the symplectic structure. From classical results in numerical analysis (e.g. see [12]), the symplectic Euler method has a good energy conservation property in discrete time. In particular, informally, under some nice assumptions on the energy function and the trajectories of the algorithm, the symplectic Euler method approximately preserves the energy function with bounded $O(\eta)$ error for exponentially many (in terms of step size) number of iterations. Furthermore, using the technique of backward error analysis, there is a modified energy function (expressed in terms of a formal series) which the discrete-time algorithm should conserve. When translated into the min-max game application, this would yield a poly-logarithmic bound on the total regret. However, to apply these powerful tools, rather nontrivial conditions need to be satisfied (e.g., integrability, and Diophantine condition). On the other hand, a key property of the energy function that we have in our context, namely convexity, has not been utilized by these tools. Therefore, we leave as a future work the task of quantifying with concrete bounds the performance of the symplectic Euler method. Here we simply conjecture that for sufficiently nice and *convex* energy functions, the iterates of the alternating method with small enough step size stay uniformly bounded, as in the unconstrained setting.

Conjecture 1. *Assume H is convex, differentiable, and satisfies some mild assumptions on the growth of derivatives. There exists $\eta_{\max} > 0$ such that for all step sizes $0 < \eta < \eta_{\max}$, along the alternating method (3), the modified energy function remains bounded:*

$$\sup_{k \geq 0} |H_\eta(z_k)| < \infty.$$

As a concrete example, we conjecture this property holds for the *alternating multiplicative weight update algorithm*, which is the alternating mirror descent method for simplex constraints with negative entropy regularizer. In this case, the energy function in the dual space is given by

$$H(x, y) = \log \sum_{i=1}^m e^{x_i} + \log \sum_{j=1}^n e^{y_j}.$$

Below, we provide some empirical evidence supporting this conjecture.

D.1 Simulation results in one-dimension

We present some simulation results in one-dimension that supports the conjecture above. For simplicity we take $A = 1$ (otherwise we can rescale both f and g). We provide the code in Section D.2.

For each example, we plot the trajectory of (x_k, y_k) on the left figure, and we plot the energy function and the modified energy function on the right figure.

D.1.1 Quadratic

As a sanity check, let $f(x) = g(x) = \frac{1}{2}x^2$. This corresponds to the unconstrained ($\mathcal{P} = \mathcal{Q} = \mathbb{R}$) min-max game with quadratic regularizer, i.e. the alternating gradient descent method of [2].

In Figure 1 we use initial position $z_0 = (x_0, y_0) = (3, 3)$, step size $\eta = 0.1$, and number of iterations $K = 300$. We see the modified energy function is conserved exactly, and the trajectory is very close to a circle.

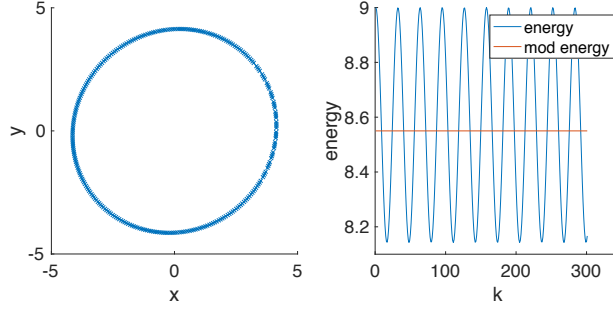


Figure 1: f and g are quadratic with $z_0 = (3, 3)$, $\eta = 0.1$, $K = 300$.

In Figure 2 we use $z_0 = (x_0, y_0) = (3, 3)$, $\eta = 1.1$, and $K = 50$. The modified energy function is still conserved exactly. The trajectory is not a circle anymore (it is an ellipse), but still bounded.

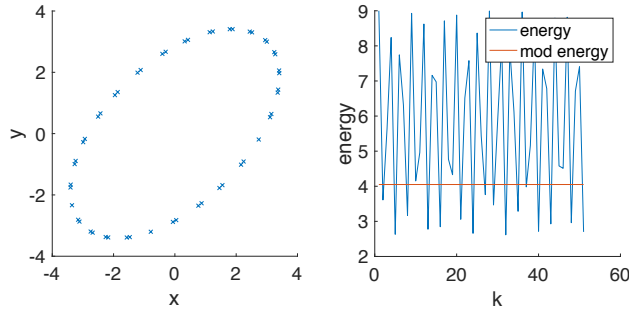


Figure 2: f and g are quadratic with $z_0 = (3, 3)$, $\eta = 1.1$, $K = 50$.

D.1.2 Log-cosh

Let $f(x) = g(x) = \log \cosh(x)$. This corresponds to the min-max game with two-dimensional simplex constraints ($\mathcal{P} = \mathcal{Q} = \Delta_2 \subset \mathbb{R}^2$) with negative entropy regularizers. In the dual space, the dual function becomes $\phi^*(x_1, x_2) = \log(e^{x_1} + e^{x_2}) = \frac{x_1 + x_2}{2} + \log \cosh(\frac{x_1 - x_2}{2}) + \log 2$. So up to a linear function, the dual function becomes $f(x) = \log \cosh x$ with $x = \frac{x_1 - x_2}{2}$.

In Figure 3 we use initial position $z_0 = (x_0, y_0) = (3, 3)$, step size $\eta = 0.1$, and number of iterations $K = 300$. Other initial positions give the same qualitative behavior. We see the modified energy function is almost conserved and the trajectory is almost periodic.

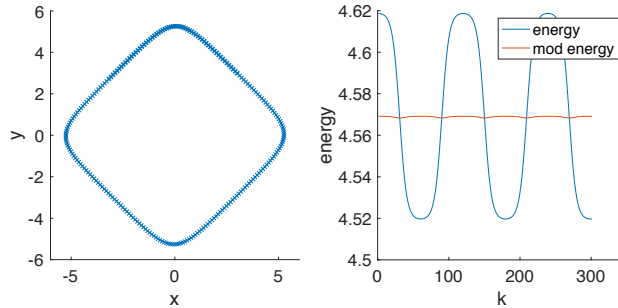


Figure 3: f and g are log cosh with $z_0 = (3, 3)$, $\eta = 0.1$, $K = 300$.

In Figure 4 we use $z_0 = (x_0, y_0) = (3, 3)$, $\eta = 1$, and $K = 200$.

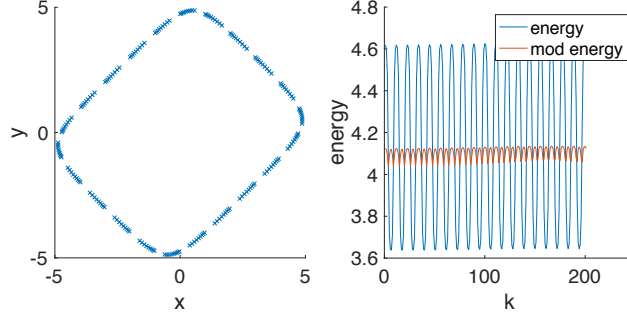


Figure 4: f and g are log cosh with $z_0 = (3, 3)$, $\eta = 1$, $K = 200$.

In Figure 5 we use $z_0 = (x_0, y_0) = (3, 3)$, $\eta = 3$, and $K = 800$.

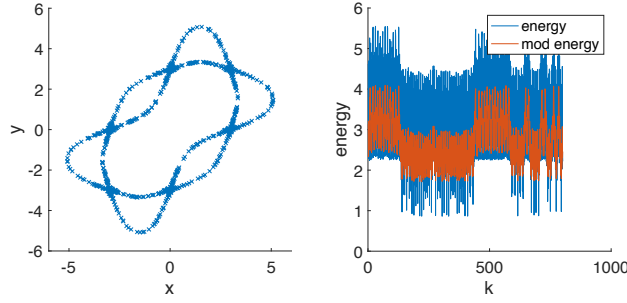


Figure 5: f and g are log cosh with $z_0 = (3, 3)$, $\eta = 3$, $K = 800$.

D.1.3 Quadratic and log cosh

Let $f(x) = \frac{1}{2}x^2$ and $g(x) = \log \cosh(x)$. This corresponds to a min-max game with one-dimensional unconstrained space for the first player with quadratic regularizer, and a two-dimensional simplex constraint for the second player with negative entropy regularizer.

In Figure 6 we use initial position $z_0 = (x_0, y_0) = (3, 3)$, step size $\eta = 0.1$, and number of iterations $K = 300$.

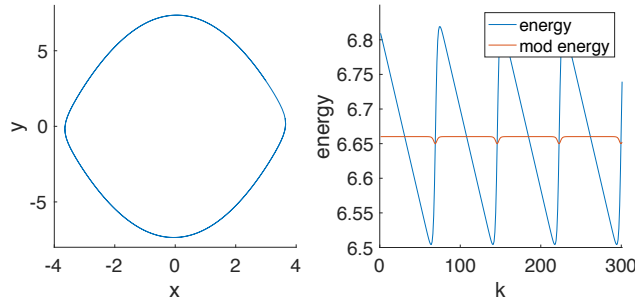


Figure 6: f is quadratic and g is log cosh with $z_0 = (3, 3)$, $\eta = 0.1$, $K = 300$.

In Figure 7 we use $z_0 = (x_0, y_0) = (3, 3)$, $\eta = 1$, and $K = 800$.

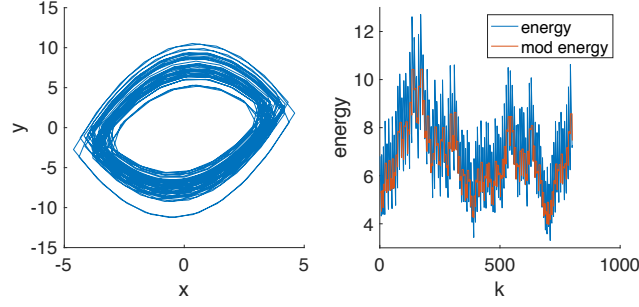


Figure 7: f is quadratic and g is log cosh with $z_0 = (3, 3)$, $\eta = 1$, $K = 800$.

D.1.4 Cubic

Let $f(x) = g(x) = \frac{1}{3}|x|^3$. In Figure 8 we use initial position $z_0 = (x_0, y_0) = (3, 3)$, step size $\eta = 0.1$, and number of iterations $K = 300$.

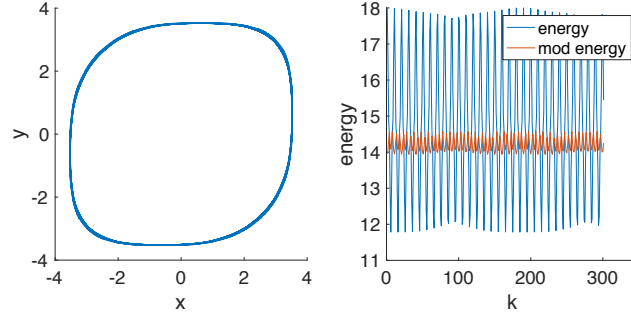


Figure 8: f and g are cubic with $z_0 = (3, 3)$, $\eta = 0.1$, $K = 300$.

In Figure 9 we use $z_0 = (x_0, y_0) = (3, 3)$, $\eta = 0.3$, and $K = 100$.

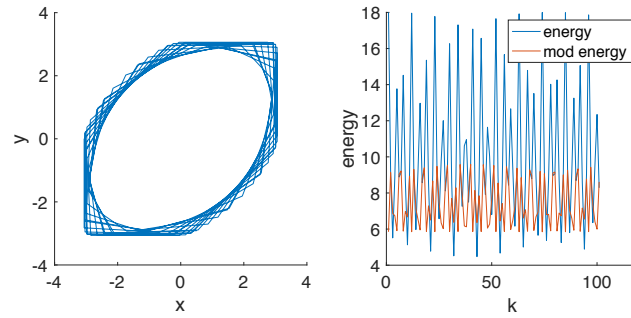


Figure 9: f and g are cubic with $z_0 = (3, 3)$, $\eta = 0.3$, $K = 100$.

D.2 Code for simulation

This is the MATLAB code for the simulation results presented above. We ran the simulations on a personal computer with 2.7 GHz processor and 16 GB of RAM. Each simulation took less than 2 seconds.

```
%% Script to simulate alternating method

% Define functions and their gradients

% Polynomial of degree p
polyp = @(p,x) 1/p.*abs(x).^p;
grad_polyp = @(p,x) x.*abs(x).^(p-2);

% log-cosh (log-sum-exp in 1-dimension)
logcosh = @(x) log(cosh(x));
grad_logcosh = @(x) tanh(x);

% Choose f and g from the list above, or define explicitly

% First option: Quadratic
f = @(x) polyp(2,x);
gradf = @(x) grad_polyp(2,x);
g = @(y) polyp(2,y);
gradg = @(y) grad_polyp(2,y);

% % Second option: logcosh
% f = logcosh;
% gradf = grad_logcosh;
% g = logcosh;
% gradg = grad_logcosh;

% % Third option: quadratic and logcosh
% f = @(x) polyp(2,x);
% gradf = @(x) grad_polyp(2,x);
% g = logcosh;
% gradg = grad_logcosh;

% % Fourth option: cubic
% f = @(x) polyp(3,x);
% gradf = @(x) grad_polyp(3,x);
% g = @(x) polyp(3,x);
% gradg = @(x) grad_polyp(3,x);

% Initial point
x0 = 3;
y0 = 3;

K = 200; % number of rounds
eta = 1; % step size

% Define energy and modified energy
energy = @(x,y) f(x) + g(y);
modenergy = @(x,y) energy(x,y) - eta/2.*gradf(x).*gradg(y);

% Prepare storage
xList = zeros(K+1,1);
yList = zeros(K+1,1);
energyList = zeros(K+1,1);
modenergyList = zeros(K+1,1);
```

```

xList(1) = x0;
yList(1) = y0;
energyList(1) = energy(x0,y0);
modenergyList(1) = modenergy(x0,y0);

% Iterate
for k = 1:K
    curX = xList(k);
    curY = yList(k);
    newX = curX - eta*gradg(curY);
    newY = curY + eta*gradf(newX);
    xList(k+1) = newX;
    yList(k+1) = newY;
    energyList(k+1) = energy(newX,newY);
    modenergyList(k+1) = modenergy(newX,newY);
end

% Plot
fig = figure;

subplot(1,2,1); hold all; set(gca,'FontSize',16);
axis square;
scatter(xList,yList,'x');
xlabel('x');
ylabel('y');

subplot(1,2,2); hold all; set(gca,'FontSize',16);
axis square;
plot(energyList);
plot(modenergyList);
xlabel('k');
ylabel('energy');
legend('energy', 'mod energy');

set(fig,'PaperOrientation','landscape');
set(fig,'PaperUnits','normalized');
set(fig,'PaperPosition',[0 0 1 1]);

```