

Accelerated Stochastic Mirror Descent: From Continuous-time Dynamics to Discrete-time Algorithms

Pan Xu*

Department of Computer Science
University of Virginia

Tianhao Wang*

School of Mathematical Sciences
Univ. of Science & Technology of China

Quanquan Gu

Department of Computer Science
University of Virginia

Abstract

We present a new framework to analyze accelerated stochastic mirror descent through the lens of continuous-time stochastic dynamic systems. It enables us to design new algorithms, and perform a unified and simple analysis of the convergence rates of these algorithms. More specifically, under this framework, we provide a Lyapunov function based analysis for the continuous-time stochastic dynamics, as well as several new discrete-time algorithms derived from the continuous-time dynamics. We show that for general convex objective functions, the derived discrete-time algorithms attain the optimal convergence rate. Empirical experiments corroborate our theory.

1 INTRODUCTION

We consider the constrained optimization problem

$$\min_{\mathbf{x} \in \mathcal{X}} f(\mathbf{x}), \quad (1.1)$$

where $f : \mathbb{R}^d \rightarrow \mathbb{R}$ is a convex function and $\mathcal{X}$ is a closed convex subset of $\mathbb{R}^d$. Denote by $\mathbf{x}^*$ the minimizer of (1.1), i.e., $f(\mathbf{x}^*) = \min_{\mathbf{x} \in \mathcal{X}} f(\mathbf{x})$. Projected gradient descent (Luenberger et al., 1984) can be used to solve (1.1) if $\mathcal{X}$ is a simple constraint set. When the computation of the gradient of f is expensive, or it is not directly accessible, stochastic gradient $\nabla \tilde{f}(\mathbf{x}, \xi)$ can be used, where ξ is some random variable and one often assumes the stochastic gradient is an unbiased estimator of the full gradient, i.e.,

$\mathbb{E}[\nabla \tilde{f}(\mathbf{x}, \xi) | \mathbf{x}] = \nabla f(\mathbf{x})$. The projected stochastic gradient descent (SGD) takes the following update formula

$$\mathbf{x}_{k+1} = \Pi_{\mathcal{X}}(\mathbf{x}_k - \eta_k \nabla \tilde{f}(\mathbf{x}_k, \xi_k)), \quad (1.2)$$

where $\eta_k > 0$ is the step size and $\Pi_{\mathcal{X}}$ denotes the Euclidean projection onto $\mathcal{X}$. When the constraint set $\mathcal{X}$ is endowed with a Bregman divergence (Bregman, 1967) that is induced by a continuously differentiable and strongly convex distance generating function $h : \mathcal{X} \rightarrow \mathbb{R}$, the projected SGD algorithm can be generalized to the stochastic mirror descent (SMD) (Nemirovskii et al., 1983):

$$\begin{cases} \mathbf{y}_{k+1} = \nabla h(\mathbf{x}_k) - \eta_k \nabla \tilde{f}(\mathbf{x}_k, \xi_k), \\ \mathbf{x}_{k+1} = \nabla h^*(\mathbf{y}_{k+1}), \end{cases} \quad (1.3)$$

where h^* is the conjugate function of h . It is easy to show that SGD is a special case of SMD when choosing $h(\cdot) = 1/2 \|\cdot\|_2^2$. It is well known that the expected objective function value after k iterations of SMD (i.e., $\mathbb{E}[f(\mathbf{x}_k)]$) converges to the minimal value $f(\mathbf{x}^*)$ at an optimal rate of $O(G/\sqrt{k})$, where G is the Lipschitz constant of f , and k is the number of iterations.

In order to accelerate first-order stochastic optimization algorithms such as SGD and SMD, various momentum techniques have been proposed (Polyak, 1964; Lan, 2012; Nesterov, 2013). In particular, Lan (2012) proposed an accelerated stochastic approximation (AC-SA) algorithm to accelerate SMD. For L -smooth and general convex function, AC-SA achieves the optimal convergence rate $O(L/k^2 + \sigma/\sqrt{k})$ in terms of expected function value gap, where σ^2 is the variance of stochastic gradient. Compared with the non-accelerated SMD (1.3), the acceleration part comes from $O(L/k^2)$ and when the variance of the stochastic gradient vanishes, i.e., $\sigma = 0$, the convergence rate of AC-SA reduces to $O(L/k^2)$, which is the optimal convergence rate in the deterministic setting. Despite the optimal convergence rate, the update formulas in these algorithms are often elusive and lack of intuitive interpretations.

*Equal Contribution

On the other hand, recent years have witnessed the emergence of a line of research which attempts to interpret the stochastic mirror descent from the perspective of continuous-time dynamics. To the best of our knowledge, Raginsky and Bouvrie (2012) is the first work that interprets stochastic mirror descent as the discretization of the following Itô stochastic differential equation (SDE) (Øksendal, 2003)

$$\begin{cases} d\mathbf{Y}_t = -\nabla f(\mathbf{X}_t)dt + \sigma d\mathbf{B}_t, \\ \mathbf{X}_t = \nabla h^*(\mathbf{Y}_t), \end{cases} \quad (1.4)$$

where $-\nabla f(\mathbf{X}_t)dt$ is called the drift term, and σ is called the diffusion term, which corresponds to the variance of the stochastic gradient. Later Mertikopoulos and Staudigl (2016) extended the above first-order SDE in (1.4) by replacing the constant diffusion term σ with a general matrix $\sigma(\mathbf{X}_t, t)$ and proved almost-sure convergence of the function value $f(\mathbf{X}_t)$ along the solution trajectories of SDE to the minimal function value $f(\mathbf{x}^*)$. Very recently, Krichene and Bartlett (2017) proposed the following second-order SDE to interpret accelerated stochastic mirror descent

$$\begin{cases} d\mathbf{Y}_t = -\eta_t(\nabla f(\mathbf{X}_t)dt + \sigma(\mathbf{X}_t, t)d\mathbf{B}_t), \\ d\mathbf{X}_t = a_t(\nabla h^*(\mathbf{Y}_t/s_t) - \mathbf{X}_t)dt, \end{cases} \quad (1.5)$$

where η_t, a_t and s_t are scaling functions of t . They proved $O(1/t^{1/2-q})$ convergence rate of expected function value along the solution trajectories of the SDE to the minimal function value, under the assumption that $\|\sigma(\mathbf{x}, t)\| \leq t^q$ ($q < 1/2$). However, there is still a gap between the continuous-time dynamics and discrete-time algorithms for accelerated stochastic mirror descent. In particular, it remains unclear whether the continuous-time SDE can be used to design and analyze new discrete-time algorithms.

In this paper, we aim to bridge this gap by proposing a new SDE-based interpretation for accelerated stochastic mirror descent, and derive new discrete-time algorithms from this SDE. We provide a unified analysis based on Lyapunov function, which connects both continuous-time dynamics and discrete-time algorithms derived thereof. We also use numerical experiments to backup our theory.

Our Contributions: We summarize the key contributions of our work as follows.

- We propose a new stochastic differential equation-based interpretation for accelerated stochastic mirror descent, motivated by Lagrangian mechanics. We take a Lyapunov function approach to prove that the convergence rate of the solution trajectories of the SDE matches the optimal rate of accelerated stochastic mirror descent for general convex function.
- We derive several new accelerated algorithms of SMD via discretizing the proposed SDE using various Euler discretization schemes, and provide an analogous Lyapunov function-based analysis for the new algorithms, which largely resembles the proof in the continuous-time dynamics. We show that these new algorithms also achieve the optimal convergence rate of accelerated SMD for general convex and smooth functions.

It is worth highlighting that under our framework, the discrete-time algorithms are closely connected with the continuous-time dynamics in the sense that they are nearly equivalent when the discretization step is sufficiently small.

The remainder of this paper is organized as follows: in Section 2 we review related work. We introduce the preliminaries in Section 3, and present the new continuous-time dynamics for accelerated SMD in Section 4. We show how we can discretize the continuous-time dynamics to invent new discrete-time algorithms in Section 5. We provide numerical experiments in Section 6 to validate our theory and conclude the paper with Section 7.

Notation We use upper case letters $\mathbf{X}_t$ to denote continuous-time curve vector, where $t \geq 0$ is the time index. $\dot{\mathbf{X}}_t$ with an over-dot denotes the derivative of $\mathbf{X}_t$ with time t . Lower case letters $\mathbf{x}_k$ denote the trajectory of a discrete-time algorithm, where $k = 0, 1, \dots$ is the index of iteration number. For all $\mathbf{x} \in \mathbb{R}^d$, we fix a general norm $\|\mathbf{x}\|$ and its dual norm is given by $\|\mathbf{x}\|_* = \sup_{\|\mathbf{y}\| \leq 1} \langle \mathbf{x}, \mathbf{y} \rangle$.

2 RELATED WORK

There is a vast literature of deterministic optimization methods for convex optimization problems. The most widely used first-order deterministic optimization algorithms include gradient descent (Polyak, 1963), mirror descent (Nemirovskii et al., 1983), Nesterov’s accelerated gradient method (Nesterov, 1983, 2005, 2013) and accelerated mirror descent (Nemirovski et al., 2009). While Nesterov’s accelerated gradient method attains the optimal convergence rate for first-order black box model, it falls short of an intuitive interpretation. Recently, there is a series of work, which interprets Nesterov’s accelerated gradient methods from different perspectives, including ordinary differential equation (ODE) (Su et al., 2014), control theory (Lessard et al., 2016; Hu and Lessard, 2017), linear coupling (Allen-Zhu and Orecchia, 2014), geometric interpretation (Bubeck et al., 2015) and game theory (Lan and Zhou, 2015), to name a few. Along these studies, the ODE-based interpretation (Su et al., 2014) is perhaps the most simple and elegant ap-

proach, which models the continuous-time limit of the accelerated methods using ODE and analyzes the stability of the resulting ODE by constructing a Lyapunov function (Su et al., 2014). Moreover, it has been extended to analyze the accelerated mirror descent (Krichene et al., 2015; Wibisono et al., 2016; Wilson et al., 2016; Diakonikolas and Orecchia, 2017). In particular, Wibisono and Wilson (2015); Wibisono et al. (2016) derived a class of ODE-based continuous dynamics for a family of accelerated optimization methods from a novel Lagrangian functional called Bregman Lagrangian. On the other hand, the ODE-based continuous-time dynamics are also able to provide guidance on designing new algorithms by carefully discretizing the ODEs. For example, Krichene et al. (2015); Wilson et al. (2016) rediscovered several discrete-time algorithms and derived a few new algorithms via different Euler discretization schemes for the ODE. Wilson et al. (2016) also found that not all discretization schemes lead to practical algorithms and/or achieve accelerated rates.

Due to its success in large-scale machine learning, stochastic first-order optimization method has also drawn a lot of research interest. Instead of using the gradient, stochastic first-order optimization methods use the stochastic gradient. Representative stochastic first-order optimization methods include stochastic gradient descent (SGD) (Robbins and Monro, 1951), stochastic mirror descent (SMD) (Nemirovski et al., 2009), and their accelerated variants (Hu et al., 2009; Lan, 2012; Ghadimi and Lan, 2012; Chen et al., 2012). Analogous to the ODE-based interpretation of deterministic optimization, stochastic optimization has an SDE-based interpretation. More specifically, Raginsky and Bouvrie (2012) studied the continuous-time dynamics of stochastic mirror descent using Itô's stochastic differential equations, where they showed that the solutions of the stochastic dynamics do not converge to the global minimizer due to the adverse effects brought by Brownian motion. Mertikopoulos and Staudigl (2016) extended the SDE in Raginsky and Bouvrie (2012) for stochastic mirror descent and proved almost-sure convergence of the solution trajectories for SMD. However, continuous-time dynamics for accelerated stochastic mirror descent remain under-studied. To the best of our knowledge, Krichene and Bartlett (2017) is the only existing work that proposes a second-order stochastic dynamics for accelerated stochastic mirror descent and proves convergence rates of the function values along solution trajectories of SDE both in terms of almost-sure convergence and in expectation. Unlike the ODE-based interpretation for deterministic optimization, based on which quite a few discrete-time algorithms are rediscovered or invented, it is unclear whether we can derive new

accelerated SMD algorithms from its SDE-based interpretation. This motivates our work.

3 PRELIMINARIES

3.1 Bregman Divergence

Mirror descent is based on Bregman divergence (Bregman, 1967). In detail, if the domain $\mathcal{X}$ is endowed with a convex and differentiable function $h : \mathcal{X} \rightarrow \mathbb{R}$, the Bregman divergence (Bregman, 1967) is defined as

$$D_h(\mathbf{x}, \mathbf{x}') = h(\mathbf{x}) - h(\mathbf{x}') - \langle \nabla h(\mathbf{x}'), \mathbf{x} - \mathbf{x}' \rangle, \quad (3.1)$$

for all $\mathbf{x}, \mathbf{x}' \in \mathcal{X}$. We always have $D_h(\mathbf{x}, \mathbf{x}') \geq 0$ due to the convexity of h . Following the assumptions in Ghadimi and Lan (2012), we assume that $D_h(\cdot, \cdot)$ grows quadratically, without loss of generality, with parameter 1, i.e.,

$$D_h(\mathbf{x}, \mathbf{x}') \leq \frac{1}{2} \|\mathbf{x} - \mathbf{x}'\|^2, \text{ for all } \mathbf{x}, \mathbf{x}' \in \mathcal{X}. \quad (3.2)$$

A plain example of Bregman divergence is the Euclidean distance: let $h(\mathbf{x}) = 1/2 \|\mathbf{x}\|_2^2$, then $D(\mathbf{x}, \mathbf{x}') = 1/2 \|\mathbf{x} - \mathbf{x}'\|_2^2$ for all $\mathbf{x}, \mathbf{x}' \in \mathcal{X}$. Another example is the Kullback-Leibler (KL) divergence: let $h(\mathbf{x}) = \sum_{i=1}^d x_i \log x_i$ be the negative entropy, then $D(\mathbf{x}, \mathbf{x}') = \sum_{i=1}^d x_i \log(x_i/x'_i)$ for all $\mathbf{x}, \mathbf{x}' \in \mathcal{X}$, which is the Kullback-Leibler divergence. Here $\mathcal{X} = \{\mathbf{x} \in \mathbb{R}_+^d : \sum_{i=1}^d x_i = 1\}$ is the unit simplex in $\mathbb{R}^d$.

Recall the SMD update (1.3), the iteration is mapped back to the primal space $\mathcal{X}$ via the mirror mapping ∇h^* after one step of gradient descent in the dual space. The mirror mapping is the gradient of the convex conjugate function h^* , which is defined as

$$h^*(\mathbf{y}) = \sup_{\mathbf{x} \in \mathcal{X}} \langle \mathbf{y}, \mathbf{x} \rangle - h(\mathbf{x}), \quad \text{for } \mathbf{y} \in E^*,$$

where E^* is the dual space of E and $\mathcal{X} \subset E$.

3.2 Assumptions and Propositions

To ease the presentation, we lay down assumptions and some propositions of Bregman divergence that will be used in our analysis.

Assumption 3.1. f is continuously differentiable and the gradient ∇f is Lipschitz continuous with parameter L_f , that is, for any $\mathbf{x}, \mathbf{y} \in \mathcal{X}$

$$\|\nabla f(\mathbf{x}) - \nabla f(\mathbf{y})\|_* \leq L_f \|\mathbf{x} - \mathbf{y}\|. \quad (3.3)$$

This is also called L_f -smoothness of f .

Assumption 3.2. h is μ_h -strongly convex with some constant $\mu_h > 0$, that is, for any $\mathbf{x}, \mathbf{y} \in \mathcal{X}$

$$h(\mathbf{x}) \geq h(\mathbf{y}) + \langle \nabla h(\mathbf{y}), \mathbf{x} - \mathbf{y} \rangle + \frac{\mu_h}{2} \|\mathbf{x} - \mathbf{y}\|^2.$$

Assumption 3.2 is commonly made in the analysis of mirror descent algorithms (Lan, 2012; Ghadimi and Lan, 2012) and its continuous-time dynamics (Wilson et al., 2016; Krichene and Bartlett, 2017)

Assumption 3.3. $\mathcal{X}$ is convex and compact. There is a constant $M_{h,\mathcal{X}} > 0$ such that

$$M_{h,\mathcal{X}} = \sup_{\mathbf{x}, \mathbf{x}' \in \mathcal{X}} D_h(\mathbf{x}, \mathbf{x}').$$

This is a natural assumption imposed in constrained optimization problems (Lan, 2012).

When the distance generating function h is strongly convex, the following two standard results on conjugate functions hold.

Proposition 3.4. Suppose Assumption 3.2 holds for h . Then its conjugate function h^* is $1/\mu_h$ -smooth and

$$\nabla h^*(\mathbf{y}) = \operatorname{argmax}_{\mathbf{x} \in \mathcal{X}} \langle \mathbf{y}, \mathbf{x} \rangle - h(\mathbf{x}), \quad \text{for } \mathbf{y} \in E^*,$$

Proposition 3.5. Suppose Assumption 3.2 holds for h . For all $\mathbf{x}, \mathbf{x}' \in \mathcal{X}$, we have

$$\nabla h^*(\nabla h(\mathbf{x})) = \mathbf{x}, \quad D_h(\mathbf{x}, \mathbf{x}') = D_{h^*}(\nabla h(\mathbf{x}'), \nabla h(\mathbf{x})).$$

We refer interested readers to Banerjee et al. (2005) for detailed discussions of these properties.

4 THE PROPOSED CONTINUOUS-TIME DYNAMICS

In this section, we present a continuous-time dynamics for accelerated stochastic mirror descent, and analyze its convergence. We start with deriving the dynamics from Bregman Lagrangian (Wibisono et al., 2016).

4.1 Continuous-time Dynamics for Accelerated Stochastic Mirror Descent

We borrow ideas from Wibisono et al. (2016) by viewing the optimizing process as a physical process. For the mechanical system associated with optimization problem (1.1), we use $\mathbf{X}_t$ and $\dot{\mathbf{X}}_t$ to denote the position and velocity respectively. We define the Bregman Lagrangian (Wibisono and Wilson, 2015) as a weighted sum of kinetic Lyapunov $D_h(\mathbf{X}_t + e^{-\alpha_t} \dot{\mathbf{X}}_t, \mathbf{X}_t)$ and potential Lyapunov $f(\mathbf{X}_t)$ as follows

$$\begin{aligned} \mathcal{L}(\mathbf{X}_t, \dot{\mathbf{X}}_t, t) \\ = e^{\alpha_t + \gamma_t} (D_h(\mathbf{X}_t + e^{-\alpha_t} \dot{\mathbf{X}}_t, \mathbf{X}_t) - e^{\beta_t} f(\mathbf{X}_t)), \end{aligned} \quad (4.1)$$

where $\alpha_t, \beta_t, \gamma_t$ are arbitrary scaling functions that are continuously differentiable with respect to t . Consider an action functional $J(\mathbf{X}) = \int_{\mathbb{T}} \mathcal{L}(\mathbf{X}_t, \dot{\mathbf{X}}_t, t) dt$ which is defined on curves $\{\mathbf{X}_t\}_{t \in \mathbb{R}_+}$. By Hamilton's principle (or principle of least action), minimizing the action

functional $J(\mathbf{X})$ requires that the curve $\mathbf{X}_t$ satisfies the following Euler-Lagrange equation

$$\frac{d}{dt} \left\{ \frac{\partial \mathcal{L}}{\partial \dot{\mathbf{X}}_t}(\mathbf{X}_t, \dot{\mathbf{X}}_t, t) \right\} = \frac{\partial \mathcal{L}}{\partial \mathbf{X}_t}(\mathbf{X}_t, \dot{\mathbf{X}}_t, t). \quad (4.2)$$

To simplify the notation, we adopt the following ideal scaling conditions suggested by Wibisono et al. (2016), which are required for the stability of ODE (4.2):

$$\dot{\beta}_t = e^{\alpha_t}, \quad \dot{\gamma}_t = e^{\alpha_t}. \quad (4.3)$$

Submitting the Bregman Lagrangian $\mathcal{L}$ in (4.1) into the Euler-Lagrange equation (4.2), we obtain the following condition for $\mathbf{X}_t$ that minimizes $J(\mathbf{X})$

$$d\nabla h(\mathbf{X}_t + 1/\dot{\beta}_t \dot{\mathbf{X}}_t) = -\dot{\beta}_t e^{\beta_t} \nabla f(\mathbf{X}_t) dt. \quad (4.4)$$

It is worth noting that (4.4) has been used to describe the continuous-time dynamics of many optimization problems in Wibisono et al. (2016). However, in stochastic optimization, one uses the stochastic gradient rather than gradient. To account for this, we add some random noise to the gradient of f . More specifically, to adapt the dynamics (4.4) to stochastic optimization, we add a Brownian motion term after the gradient $\nabla f(\mathbf{X}_t)$ to form the following Itô stochastic differential equation (SDE) (Øksendal, 2003)

$$\begin{aligned} d\nabla h(\mathbf{X}_t + 1/\dot{\beta}_t \dot{\mathbf{X}}_t) \\ = -\dot{\beta}_t e^{\beta_t} (\nabla f(\mathbf{X}_t) dt + \sqrt{\delta} \sigma(\mathbf{X}_t, t) dB_t), \end{aligned} \quad (4.5)$$

where $B_t \in \mathbb{R}^d$ is the standard Brownian motion, $\sigma(\mathbf{X}_t, t) \in \mathbb{R}^{d \times d}$ and $\delta > 0$ is a constant. The term $\sqrt{\delta} \sigma(\mathbf{X}_t, t)$ is called the diffusion coefficient that accounts for the variance of the stochastic gradient.

Note that (4.5) is a second-order SDE. We define a variable $\mathbf{Y}_t$ in the dual space E^* as

$$\mathbf{Y}_t = \nabla h(\mathbf{X}_t + 1/\dot{\beta}_t \dot{\mathbf{X}}_t). \quad (4.6)$$

Based on (4.5) and (4.6), and using Proposition 3.5, we obtain the following continuous-time dynamics for accelerated stochastic mirror descent

$$d\mathbf{X}_t = \dot{\beta}_t (\nabla h^*(\mathbf{Y}_t) - \mathbf{X}_t) dt, \quad (4.7a)$$

$$d\mathbf{Y}_t = -\dot{\beta}_t e^{\beta_t} (\nabla f(\mathbf{X}_t) dt + \sqrt{\delta} \sigma(\mathbf{X}_t, t) dB_t). \quad (4.7b)$$

The solution trajectories of (4.7) do not converge to the minimizer of f due to the stochastic noise. So we modify it by introducing a shrinkage parameter $s_t > 0$ to decrease the adverse effect of the stochastic noise.

$$d\mathbf{X}_t = \dot{\beta}_t (\nabla h^*(\mathbf{Y}_t) - \mathbf{X}_t) dt, \quad (4.8a)$$

$$d\mathbf{Y}_t = -\frac{\dot{\beta}_t e^{\beta_t}}{s_t} (\nabla f(\mathbf{X}_t) dt + \sqrt{\delta} \sigma(\mathbf{X}_t, t) dB_t), \quad (4.8b)$$

where $B_t \in \mathbb{R}^d$ is the standard Brownian motion, $\beta_t, s_t > 0$ are scaling functions of time t and $\delta > 0$ is a

constant. The idea of introducing shrinkage parameter s_t to offset the disadvantage brought by stochastic gradient is similar to that in Mertikopoulos and Staudigl (2016); Krichene and Bartlett (2017). This amount to using time-decaying step size in the discrete-time stochastic mirror descent algorithm to ensure its convergence.

4.2 Convergence Rate of the Continuous-time Dynamics

In this subsection, we analyze the convergence of our proposed continuous-time dynamics (4.8) for accelerated stochastic mirror descent. The following theorem spells out its convergence rate.

Theorem 4.1. Suppose f is convex and Assumptions 3.2 and 3.3 hold. If the diffusion coefficient in (4.8) satisfies $\|\sigma(\mathbf{X}_t, t)^\top \sigma(\mathbf{X}_t, t)\| \leq t^{2q}$ for some $q < 1/2$, then it holds that

$$\begin{aligned} & \mathbb{E}[f(\mathbf{X}_t) - f(\mathbf{x}^*)] \\ & \leq \frac{\mathcal{E}_0 + (s_t - s_0)M_{h,\mathcal{X}}}{e^{\beta_t}} \\ & \quad + \frac{\mathbb{E}\left[\int_0^t \frac{\delta \dot{\beta}_r^2 e^{2\beta_r}}{s_r} \text{tr}(\sigma_r^\top \nabla^2 h^*(\mathbf{Y}_r) \sigma_r) dr\right]}{2e^{\beta_t}}, \end{aligned} \quad (4.9)$$

where $\sigma_r = \sigma(\mathbf{X}_r, r)$ and $\mathcal{E}_0 = e^{\beta_0}(f(\mathbf{X}_0) - f(\mathbf{x}^*)) + s_0 D_{h^*}(\mathbf{Y}_0, \nabla h(\mathbf{x}^*))$.

Note that the convergence of the continuous-time dynamics (4.8) does not require the smoothness of f .

Remark 4.2. If we choose $\beta_t = p \log t$ for some constant $p > 0$ and $s_t = t^\alpha$ for some $\alpha \in \mathbb{R}$ and note that $\|\sigma(\mathbf{X}_t, t)^\top \sigma(\mathbf{X}_t, t)\| \leq t^{2q}$, we have the following convergence rate

$$\mathbb{E}[f(\mathbf{X}_t) - f(\mathbf{x}^*)] = O\left(\frac{1 + t^\alpha + t^{2p-\alpha+2q-1}}{t^p}\right).$$

If we further choose $\alpha = p + q - 1/2$ and $p = 2$, then we obtain

$$\mathbb{E}[f(\mathbf{X}_t) - f(\mathbf{x}^*)] = O\left(\frac{1}{t^2} + \frac{1}{t^{1/2-q}}\right),$$

which matches the optimal convergence rate for accelerated stochastic mirror descent algorithms (Lan, 2012), where the variance of stochastic gradient is bounded by a constant i.e., $q = 0$. In contrast, the convergence rate of the continuous dynamics (1.5) proposed by Krichene and Bartlett (2017) for accelerated stochastic mirror descent is $O(1/t^{1/2})$, which only matches the optimal rate up to the dominating term.

Remark 4.3. When the variance of stochastic gradient vanishes, i.e., $\sigma(\mathbf{X}, t) = 0$, and if we choose the shrinkage parameter $s_t = 1$ for all $t \in \mathbb{R}_+$, the convergence rate in (4.9) becomes $O(1/t^2)$, which matches

the optimal convergence rate of deterministic accelerated mirror descent algorithms for general convex functions (Nesterov, 1983).

Here we provide a proof of Theorem 4.1 via constructing the following Lyapunov function for the stochastic dynamics system (4.8)

$$\mathcal{E}_t = e^{\beta_t}(f(\mathbf{X}_t) - f(\mathbf{x}^*)) + s_t D_{h^*}(\mathbf{Y}_t, \nabla h(\mathbf{x}^*)). \quad (4.10)$$

Proof of Theorem 4.1. Denote $\sigma_t = \sigma(\mathbf{X}_t, t)$, and applying Itô's Lemma to the Lyapunov function yields

$$\begin{aligned} d\mathcal{E}_t &= \frac{\partial \mathcal{E}_t}{\partial t} dt + \left\langle \frac{\partial \mathcal{E}_t}{\partial \mathbf{X}_t}, d\mathbf{X}_t \right\rangle + \left\langle \frac{\partial \mathcal{E}_t}{\partial \mathbf{Y}_t}, d\mathbf{Y}_t \right\rangle \\ &\quad + \frac{\delta \dot{\beta}_t^2 e^{2\beta_t}}{2s_t^2} \text{tr}\left(\sigma_t^\top \frac{\partial^2 \mathcal{E}_t}{\partial \mathbf{Y}_t^2} \sigma_t\right) dt. \end{aligned} \quad (4.11)$$

By simple calculus, we have

$$\begin{aligned} \frac{\partial \mathcal{E}_t}{\partial t} &= \dot{\beta}_t e^{\beta_t}(f(\mathbf{X}_t) - f(\mathbf{x}^*)) + \dot{s}_t D_{h^*}(\mathbf{Y}_t, \nabla h(\mathbf{x}^*)), \\ \frac{\partial \mathcal{E}_t}{\partial \mathbf{X}_t} &= e^{\beta_t} \nabla f(\mathbf{X}_t), \quad \frac{\partial \mathcal{E}_t}{\partial \mathbf{Y}_t} = s_t (\nabla h^*(\mathbf{Y}_t) - \mathbf{x}^*). \end{aligned}$$

Submitting the above calculations and dynamics (4.8) into (4.11) yields

$$\begin{aligned} d\mathcal{E}_t &= [\dot{\beta}_t e^{\beta_t}(f(\mathbf{X}_t) - f(\mathbf{x}^*)) + \dot{s}_t D_{h^*}(\mathbf{Y}_t, \nabla h(\mathbf{x}^*))] dt \\ &\quad + \dot{\beta}_t e^{\beta_t} \langle \nabla h^*(\mathbf{Y}_t) - \mathbf{x}^*, \nabla f(\mathbf{X}_t) \rangle dt \\ &\quad - \dot{\beta}_t e^{\beta_t} \langle \nabla h^*(\mathbf{Y}_t) - \mathbf{x}^*, \nabla f(\mathbf{X}_t) \rangle dt + \sigma_t d\mathbf{B}_t \\ &\quad + \frac{1}{2} \frac{\delta \dot{\beta}_t^2 e^{2\beta_t}}{s_t^2} \text{tr}\left(\sigma_t^\top \frac{\partial^2 \mathcal{E}_t}{\partial \mathbf{Y}_t^2} \sigma_t\right) dt \\ &= \dot{\beta}_t e^{\beta_t} [f(\mathbf{X}_t) - f(\mathbf{x}^*) + \langle \mathbf{x}^* - \mathbf{X}_t, \nabla f(\mathbf{X}_t) \rangle] dt \\ &\quad + \dot{s}_t D_{h^*}(\mathbf{Y}_t, \nabla h(\mathbf{x}^*)) dt \\ &\quad - \dot{\beta}_t e^{\beta_t} \langle \nabla h^*(\mathbf{Y}_t) - \mathbf{x}^*, \sigma_t d\mathbf{B}_t \rangle \\ &\quad + \mathbb{E}\left[\frac{\delta \dot{\beta}_t^2 e^{2\beta_t}}{2s_t} \text{tr}(\sigma_t^\top \nabla^2 h^*(\mathbf{Y}_t) \sigma_t)\right] dt \\ &\leq \dot{s}_t M_{h,\mathcal{X}} + \frac{1}{2s_t} \delta \dot{\beta}_t^2 e^{2\beta_t} \text{tr}(\sigma_t^\top \nabla^2 h^*(\mathbf{Y}_t) \sigma_t) dt \\ &\quad - \dot{\beta}_t e^{\beta_t} \langle \nabla h^*(\mathbf{Y}_t) - \mathbf{x}^*, \sigma_t d\mathbf{B}_t \rangle, \end{aligned}$$

where the inequality follows from the convexity of f and Assumption 3.3 that $D_{h^*}(\mathbf{Y}_t, \nabla h(\mathbf{x}^*)) \leq M_{h,\mathcal{X}}$. Integrating from 0 to t , we obtain

$$\begin{aligned} \mathcal{E}_t &\leq \mathcal{E}_0 + (s_t - s_0)M_{h,\mathcal{X}} \\ &\quad + \frac{1}{2} \int_0^t \frac{\delta \dot{\beta}_r^2 e^{2\beta_r}}{s_r} \text{tr}(\sigma_r^\top \nabla^2 h^*(\mathbf{Y}_r) \sigma_r) dr \\ &\quad - \int_0^t \dot{\beta}_r e^{\beta_r} \langle \nabla h^*(\mathbf{Y}_r) - \mathbf{x}^*, \sigma_r d\mathbf{B}_r \rangle. \end{aligned}$$

By the definition of $\mathcal{E}_t$ and the fact that $D_h^* \geq 0$, we immediately get

$$\begin{aligned} & e^{\beta t} (f(\mathbf{X}_t) - f(\mathbf{x}^*)) \\ & \leq \mathcal{E}_0 + (s_t - s_0) M_{h, \mathcal{X}} \\ & \quad + \int_0^t \frac{\delta \dot{\beta}_r^2 e^{2\beta_r}}{2s_r} \text{tr}(\sigma_r^\top \nabla^2 h^*(\mathbf{Y}_r) \sigma_r) dr \\ & \quad - \int_0^t \dot{\beta}_r e^{\beta_r} \langle \nabla h^*(\mathbf{Y}_r) - \mathbf{x}^*, \sigma_r d\mathbf{B}_r \rangle. \end{aligned}$$

Taking expectation of both sides and using the martingale property of Itô integral, we obtain

$$\begin{aligned} & \mathbb{E}[f(\mathbf{X}_t) - f(\mathbf{x}^*)] \\ & \leq \frac{\mathcal{E}_0 + (s_t - s_0) M_{h, \mathcal{X}}}{e^{\beta t}} \\ & \quad + \frac{\mathbb{E}[\int_0^t \frac{\delta \dot{\beta}_r^2 e^{2\beta_r}}{s_r} \text{tr}(\sigma_r^\top \nabla^2 h^*(\mathbf{Y}_r) \sigma_r) dr]}{2e^{\beta t}}. \end{aligned}$$

This completes the proof. $\square$

5 THE PROPOSED DISCRETE-TIME ALGORITHMS

Our primary goal in proposing the continuous-time dynamics for accelerated stochastic mirror descent is to design new discrete-time algorithms. In this section, we provide several discretizations of (4.8). We will adapt the Lyapunov function-based analysis in previous section to the discrete-time algorithms and deliver very simple and intuitive proofs for the new algorithms. The high-level idea of this section is inspired by Wilson et al. (2016), which discussed various discretization schemes for the continuous-time dynamics of deterministic optimization.

We use Euler discretization schemes (Kloeden and Platen, 1992) of differential equations and its composition to derive several discrete algorithms of dynamic (4.8). Specifically, note that (4.8) is a system of two first-order differential equations, and thus we can compose explicit and implicit Euler discretizations in four different ways. Let δ be the time step and

$$\mathbf{x}_k = \mathbf{X}_t, \quad \mathbf{x}_{k+1} = \mathbf{X}_{t+\delta}, \quad \mathbf{y}_k = \mathbf{Y}_t, \quad \mathbf{y}_{k+1} = \mathbf{Y}_{t+\delta}.$$

The following approximations

$$(\mathbf{x}_{k+1} - \mathbf{x}_k)/\delta \approx \dot{\mathbf{X}}_t, \quad (\mathbf{y}_{k+1} - \mathbf{y}_k)/\delta \approx \dot{\mathbf{Y}}_t$$

are the explicit (forward) Euler discretization for the time derivatives of $\mathbf{X}_t$ and $\mathbf{Y}_t$, and

$$(\mathbf{x}_k - \mathbf{x}_{k-1})/\delta \approx \dot{\mathbf{X}}_t, \quad (\mathbf{y}_k - \mathbf{y}_{k-1})/\delta \approx \dot{\mathbf{Y}}_t$$

are the implicit (backward) Euler discretization for the time derivatives of $\mathbf{X}_t$ and $\mathbf{Y}_t$. For the scaling parameters, we choose $A_k = e^{\beta t}$, $s_k = s_t$, and the discretizations as follows

$$(A_{k+1} - A_k)/\delta \approx d e^{\beta t} / dt, \quad (A_{k+1} - A_k)/(\delta A_k) \approx \dot{\beta}_t.$$

5.1 Implicit Euler Discretization

Algorithm 1 Accelerated Stochastic Mirror Descent (Implicit)

-
- 1: **Input:** $A_0 = 1/2$, $\mathbf{x}_0 = \mathbf{y}_0 = \mathbf{0}$.
 - 2: **for** $k = 0$ to K **do**
 - 3: $A_{k+1} = (k+1)(k+2)/2$, $\tau_k = (A_{k+1} - A_k)/A_k$,
 $s_k = \sigma k^{3/2} + 1$
 - 4: $\nabla h^*(\mathbf{y}_{k+1}) = \mathbf{x}_{k+1} + \frac{1}{\tau_k}(\mathbf{x}_{k+1} - \mathbf{x}_k)$.
 - 5: $\mathbf{x}_{k+1} = \underset{\mathbf{x} \in \mathcal{X}}{\text{argmin}} \left\{ \tilde{f}(\mathbf{x}, \xi_{k+1}) + \frac{s_k}{A_{k+1}} D_h(\mathbf{u}, \nabla h^*(\mathbf{y}_k)) \right\}$, and $\mathbf{u} = \mathbf{x} + \frac{1}{\tau_k}(\mathbf{x} - \mathbf{x}_k)$.
 - 6: **end for**
-

The first discrete-time algorithm we derive is from implicit Euler discretization of dynamics (4.8). The discretization process is displayed in Algorithm 1. Here we use Borel function $\tilde{f}(\mathbf{x}_k, \xi_k)$ to denote the noisy objective function, where $\{\xi_k\}_{k=0,1,\dots}$ is a sequence of independent random variables. Denote $\Delta(\mathbf{x}_k) = \nabla \tilde{f}(\mathbf{x}_k, \xi_k) - \nabla f(\mathbf{x}_k)$ and we assume that

$$\begin{aligned} \mathbb{E}[\tilde{f}(\mathbf{x}_k, \xi_k) | \mathbf{x}_k] &= f(\mathbf{x}_k), \quad \mathbb{E}[\nabla \tilde{f}(\mathbf{x}_k, \xi_k) | \mathbf{x}_k] = \nabla f(\mathbf{x}_k), \\ \mathbb{E}[\|\Delta(\mathbf{x}_k)\|_*^2 | \mathbf{x}_k] &\leq \sigma^2. \end{aligned}$$

The optimality condition of Algorithm 1 is given by

$$\begin{aligned} & \left\{ \begin{aligned} \nabla h^*(\mathbf{y}_{k+1}) &= \mathbf{x}_{k+1} + \frac{A_k}{A_{k+1} - A_k}(\mathbf{x}_{k+1} - \mathbf{x}_k), \\ \mathbf{y}_{k+1} - \mathbf{y}_k &= -\frac{A_{k+1} - A_k}{s_k} \nabla \tilde{f}(\mathbf{x}_{k+1}, \xi_{k+1}), \end{aligned} \right. \end{aligned} \quad (5.1a) \quad (5.1b)$$

where $\nabla \tilde{f}(\mathbf{x}_{k+1}, \xi_{k+1})$ is the stochastic gradient.

Inspired by the Lyapunov function (4.10) for the continuous-time dynamics, we construct the following Lyapunov function to analyze the convergence of Algorithm 1.

$$\mathcal{E}_k = \mathbb{E}[A_k(f(\mathbf{x}_k) - f(\mathbf{x}^*)) + s_k D_h(\mathbf{x}^*, \nabla h^*(\mathbf{y}_k))]. \quad (5.2)$$

Based on the above Lyapunov function, we can prove the following theorem, which states the convergence rate of Algorithm 1.

Theorem 5.1. Suppose f is convex. Under Assumptions 3.2 and 3.3, if we choose $A_k = k(k+1)/2$, $A_0 = 1/2$ and $s_k = \sigma k^{3/2} + 1$, the expected function value gap at $\mathbf{x}_k$ output by Algorithm 1 is bounded as

$$\begin{aligned} \mathbb{E}[f(\mathbf{x}_k) - f(\mathbf{x}^*)] &\leq \frac{2(\mathcal{E}_0 + M_{h, \mathcal{X}})}{k(k+1)} \\ &\quad + \sigma \left(M_{h, \mathcal{X}} + \frac{1}{3\mu_h} \right) \frac{\sqrt{k+1}}{k}, \end{aligned}$$

where $\mathcal{E}_0 = A_0(f(\mathbf{x}_0) - f(\mathbf{x}^*)) + s_0 D_h(\mathbf{x}^*, \nabla h^*(\mathbf{y}_0))$.

Remark 5.2. Theorem 5.1 suggests that the convergence rate of Algorithm 1 is $O((1 + M_{h,\mathcal{X}})/k^2 + M_{h,\mathcal{X}}\sigma/\sqrt{k})$, which matches the optimal rate of accelerated stochastic optimization for general convex and smooth functions (Lan, 2012). Note that when variance of the stochastic gradient vanishes, we can choose $s_k = 1$ for all k and obtain the optimal convergence rate of accelerated gradient descent for general convex and smooth functions $O(1/k^2)$.

Remark 5.3. The optimal convergence rate of Algorithm 1 does not require the smooth assumption on f , which is aligned with the analysis of continuous-time dynamics (4.8). In fact, Algorithm 1 can be seen as an accelerated proximal point algorithm whose optimal convergence rate is also $O(1/k^2)$ (Güler, 1992).

5.2 Hybrid Euler Discretization

Although Algorithm 1 is succinct and attains the optimal convergence rate, the implicit Euler discretizations of both stochastic processes make it hard to implement in practice due to the requirement of an exact minimization in each iteration (Step 5 in Algorithm 1). Nevertheless, the implicit discretization sheds great light on the connection between continuous-time dynamics and discrete-time algorithms. Now we derive a more practical algorithm combining implicit and explicit Euler discretizations of dynamics (4.8). The discretization process is displayed in Algorithm 2. The

Algorithm 2 Accelerated Stochastic Mirror Descent (ASMD)

- 1: **Input:** $A_0 = s_0 = 1/2, \mathbf{x}_0 = \mathbf{y}_0 = \mathbf{0}$.
 - 2: **for** $k = 0$ to K **do**
 - 3: $A_{k+1} = (k+1)(k+2)/2, \tau_k = (A_{k+1} - A_k)/A_k, s_{k+1} = (k+1)^{3/2}$.
 - 4: $\mathbf{x}_{k+1} = \frac{\tau_k}{\tau_k+1} \nabla h^*(\mathbf{y}_k) + \frac{1}{\tau_k+1} \mathbf{x}_k$.
 - 5: $\mathbf{y}_{k+1} = \mathbf{y}_k - \frac{A_{k+1}-A_k}{s_k} \nabla \tilde{f}(\mathbf{x}_{k+1}, \xi_{k+1})$.
 - 6: **end for**
-

optimality condition of Algorithm 2 is given by

$$\begin{cases} \nabla h^*(\mathbf{y}_k) = \mathbf{x}_{k+1} + \frac{A_k}{A_{k+1} - A_k} (\mathbf{x}_{k+1} - \mathbf{x}_k), & (5.3a) \\ \mathbf{y}_{k+1} - \mathbf{y}_k = -\frac{A_{k+1} - A_k}{s_k} \nabla \tilde{f}(\mathbf{x}_{k+1}, \xi_{k+1}). & (5.3b) \end{cases}$$

Based on the same Lyapunov function (5.2) of Algorithm 1, we can prove the convergence rate of Algorithm 2.

Theorem 5.4. Suppose f is convex. Under Assumptions 3.1, 3.2 and 3.3, if we choose $A_k = k(k+1)/2$, $s_k = k^{3/2}$ and $A_0 = s_0 = 1/2$, the expected function value gap at $\mathbf{x}_k$ output by Algorithm 2 is bounded as

$$\mathbb{E}[f(\mathbf{x}_k) - f(\mathbf{x}^*)] \leq \frac{2\mathcal{E}_0}{k(k+1)} + \frac{C_0\sqrt{k}}{\mu_h^2(k+1)},$$

where $C_0 = 2((4L_f^2 + \mu_h^2)\mathcal{M}_{h,\mathcal{X}} + \mu_h\sigma^2 + 2\|\nabla f(\mathbf{x}^*)\|_*^2)$ and $\mathcal{E}_0 = A_0(f(\mathbf{x}_0) - f(\mathbf{x}^*)) + s_0 D_h(\mathbf{x}^*, \nabla h^*(\mathbf{y}_0))$.

Remark 5.5. Compared with Algorithm 1, Algorithm 2 is a much more practical algorithm and can be easily implemented. The convergence rate of Algorithm 2 is $O(M_{h,\mathcal{X}}/k^2 + (L_f^2 + \sigma^2)/\sqrt{k})$. Although this matches the optimal rate of accelerated stochastic mirror descent (Lan, 2012), yet when the variance of the stochastic gradient vanishes, i.e., $\sigma = 0$, it is not reducible to the optimal convergence rate $O(1/k^2)$ for deterministic optimization.

5.3 Discretization with an Additional Sequence

Algorithm 3 Accelerated Stochastic Mirror Descent (ASMD3)

- 1: **Input:** $A_0 = 1/2, \mathbf{x}_0 = \mathbf{y}_0 = \mathbf{z}_0 = \mathbf{0}$.
 - 2: **for** $k = 0$ to K **do**
 - 3: $A_{k+1} = \mu_h^2(k+1)(k+2)/(4L_f), s_k = \sigma/L_f(k+1)^{3/2} + 1, M_k = L_f(A_{k+1} - A_k)^2/(\mu_h^2 s_k A_{k+1})$.
 - 4: $\mathbf{z}_{k+1} = \frac{A_{k+1}-A_k}{A_{k+1}} \nabla h^*(\mathbf{y}_k) + \frac{A_k}{A_{k+1}} \mathbf{x}_k$.
 - 5: $\mathbf{y}_{k+1} = \mathbf{y}_k - \frac{A_{k+1}-A_k}{s_k} \nabla \tilde{f}(\mathbf{z}_{k+1}, \xi_{k+1})$.
 - 6: $\mathbf{x}_{k+1} = \operatorname{argmin}_{\mathbf{x} \in \mathcal{X}} \left\{ \langle \nabla \tilde{f}(\mathbf{z}_{k+1}, \xi_{k+1}), \mathbf{x} \rangle + \frac{L_f}{M_k} D_h(\mathbf{z}_{k+1}, \mathbf{x}) \right\}$.
 - 7: **end for**
-

As we showed in previous subsection, Algorithm 2 is not able to obtain the optimal rate for deterministic optimization when using the full gradient instead of stochastic gradient. In order to address this issue, we will modify the algorithm by adding an additional sequence in a similar way to linear coupling (Allen-Zhu and Orecchia, 2014). The new algorithm with three sequences is presented in Algorithm 3. Note that the additionally added sequence can be seen as a step of mirror descent. The optimality condition of Algorithm 3 is given by

$$\begin{cases} \mathbf{z}_{k+1} = \frac{A_{k+1} - A_k}{A_{k+1}} \nabla h^*(\mathbf{y}_k) + \frac{A_k}{A_{k+1}} \mathbf{x}_k, & (5.4a) \\ \mathbf{y}_{k+1} - \mathbf{y}_k = -\frac{A_{k+1} - A_k}{s_k} \nabla \tilde{f}(\mathbf{z}_{k+1}, \xi_{k+1}), & (5.4b) \\ \mathbf{x}_{k+1} = \operatorname{argmin}_{\mathbf{x} \in \mathcal{X}} \left\{ \langle \nabla \tilde{f}(\mathbf{z}_{k+1}, \xi_{k+1}), \mathbf{x} \rangle + \frac{L_f}{M_k} D_h(\mathbf{z}_{k+1}, \mathbf{x}) \right\}. & (5.4c) \end{cases}$$

In order to prove the convergence of Algorithm 3 with the additional sequence, we define a new Lyapunov function as $\mathcal{E}_k = \mathbb{E}[A_k(f(\mathbf{x}_k) - f(\mathbf{x}^*)) + s_k D_h(\mathbf{x}^*, \nabla h^*(\mathbf{y}_{k+1}))]$, which is slightly different from that of previous algorithms. Then we can show the following result.

Theorem 5.6. Suppose f is convex. Under Assumptions 3.1, 3.2 and 3.3, if we choose $A_k = \mu_h^2 k(k+1)$

$1)/(4L_f)$, $s_k = \sigma/L_f(k+1)^{3/2} + 1$ and $M_k = L_f(A_{k+1} - A_k)^2/(\mu_h^2 s_k A_{k+1})$, the expected function value gap at $\mathbf{x}_k$ output by Algorithm 3 is bounded as

$$\begin{aligned} & \mathbb{E}[f(\mathbf{x}_k) - f(\mathbf{x}^*)] \\ & \leq \frac{4L_f(\mathcal{E}_0 + M_{h,\mathcal{X}})}{\mu_h^2 k(k+1)} + \frac{\sigma(\mu_h^2 + 12M_{h,\mathcal{X}})\sqrt{k+1}}{3\mu_h^2 k}, \end{aligned}$$

where $\mathcal{E}_0 = A_0(f(\mathbf{x}_0) - f(\mathbf{x}^*)) + s_0 D_h(\mathbf{x}^*, \nabla h^*(\mathbf{y}_0))$.

Remark 5.7. The convergence rate of Algorithm 3 is in the order of $O(L_f(1 + M_{h,\mathcal{X}})/k^2 + \sigma(1 + M_{h,\mathcal{X}})/\sqrt{k})$, which matches the optimal rate of accelerated stochastic mirror descent for general convex and smooth functions (Lan, 2012). More importantly, when the stochastic gradient reduces to the full gradient, namely, $\sigma = 0$, the $\sigma(1 + M_{h,\mathcal{X}})/\sqrt{k}$ term diminishes and the optimal convergence rate for stochastic optimization reduces to the optimal convergence rate $O(1/k^2)$ for deterministic optimization of general convex and smooth functions.

6 EXPERIMENTS

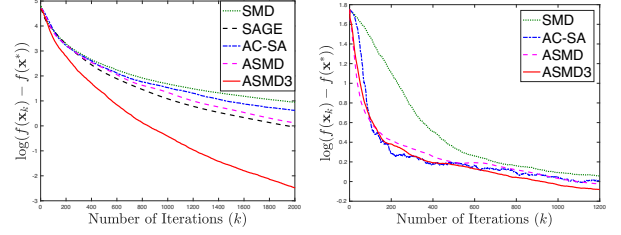
In this section, we conduct numerical experiments to verify the convergence rate of the proposed algorithms derived from the continuous-time dynamics. We compare our Algorithm 2 (**ASMD**) and Algorithm 3 (**ASMD3**) with stochastic mirror descent (**SMD**), accelerated stochastic approximation (**AC-SA**) (Lan, 2012) and stochastic accelerated gradient (**SAGE**) (Hu et al., 2009).

We apply the above optimization algorithms to the linear regression problem on a convex compact subset of $\mathbb{R}^d$.

$$\min_{\mathbf{x} \in \mathcal{X}} f(\mathbf{x}) = \sum_{i=1}^n (\mathbf{A}_{i*} \mathbf{x} - y_i)^2, \quad (6.1)$$

where $\mathbf{y} = (y_1, \dots, y_n)^\top \in \mathbb{R}^n$ and $\mathbf{A}_{i*}$ denotes the i -th row of $\mathbf{A} \in \mathbb{R}^{n \times d}$. We set $n = 100$, $d = 200$, and generated the design matrix $\mathbf{A}$ with entries following $N(0, 1)$. The response vector was generated by $\mathbf{y} = \mathbf{A}\mathbf{u}^* + \boldsymbol{\epsilon}$, where $\boldsymbol{\epsilon} \sim N(\mathbf{0}, \mathbf{I}_{n \times n})$ and $\mathbf{u}^* \sim N(\mathbf{0}, \mathbf{I}_{d \times d})$ were randomly generated. It is easy to verify that the objective function f is convex and L -smooth with $L = \|\mathbf{A}^\top \mathbf{A}\|_2$. We considered two settings of distance generating function h and the constrained set $\mathcal{X}$:

- Setting (1): $h(\mathbf{x}) = 1/2\|\mathbf{x}\|_2^2$ is the squared Euclidean norm, and $\mathcal{X} = \{\mathbf{x} : \|\mathbf{x}\|_2 \leq 2\|\mathbf{u}^*\|_2\}$. Then the mirror mapping $\nabla h^*(\mathbf{x}) = \operatorname{argmin}_{\mathbf{u} \in \mathcal{X}} \|\mathbf{x} - \mathbf{u}\|_2^2$ reduces to projection onto $\mathcal{X}$ and has a closed form solution.
- Setting (2): $h(\mathbf{x}) = \sum_{i=1}^d x_i \log x_i$ is the negative entropy and $\mathcal{X} = \{\mathbf{x} \in \mathbb{R}_+^d : \sum_{i=1}^d x_i = 1\}$ is a simplex. In this case the mirror mapping also has



(a) Setting (1): $h(\mathbf{x}) = 1/2\|\mathbf{x}\|_2^2$, and $\mathcal{X} = \{\mathbf{x} \in \sum_{i=1}^d x_i \log x_i, \text{ and } \mathcal{X} = \mathbb{R}^d : \|\mathbf{x}\|_2 \leq 2\|\mathbf{u}^*\|_2\}$ (b) Setting (2): $h(\mathbf{x}) = 1/2\|\mathbf{x}\|_2^2$, and $\mathcal{X} = \{\mathbf{x} \in \mathbb{R}_+^d : \sum_{i=1}^d x_i = 1\}$

Figure 1: Logarithmic averaged function value gap over 50 repetitions for all methods under different settings.

a closed-form solution $[\nabla h^*(\mathbf{x})]_i = e^{x_i} / \sum_{i=1}^d e^{x_i}$ for all $i = 1, \dots, d$ (Banerjee et al., 2005).

Since (6.1) has a finite-sum structure, we simply choose the stochastic gradient by uniformly sampling i from $\{1, 2, \dots, n\}$ with mini batch size 1. We plotted the function value gap $f(\mathbf{x}_k) - f(\mathbf{x}^*)$ for each algorithm in Figure 1, where k is the number of iterations and $\mathbf{x}^*$ is the global minimizer of (6.1) solved by CVX program (Grant and Boyd, 2008). Figure 1(a) and Figure 1(b) show the averaged results over 50 repetitions for Setting (1) and Setting (2) respectively. Note that **SAGE** is not applicable to mirror descent. Experimental results in both settings demonstrate that our algorithms **ASMD** and **ASMD3** achieve comparable convergence rate as AC-SA, and converge much faster than SMD. This is well-aligned with our theory.

7 CONCLUSIONS

In this paper, we bridge the gap between continuous-time dynamics and discrete-time algorithms for accelerated stochastic mirror descent by presenting a variational analysis based on Bregman Lagrangian and Lyapunov functions. We not only propose a new continuous-time dynamics of accelerate stochastic mirror descent, but also derive several new algorithms from the continuous dynamics. Both the continuous-time dynamics and the discrete-time algorithms achieve the optimal convergence rate for general convex and smooth functions in the stochastic optimization setting.

Acknowledgements We would like to thank the anonymous reviewers for their helpful comments. This research was sponsored in part by the National Science Foundation IIS-1618948 and IIS-1652539. The views and conclusions contained in this paper are those of the authors and should not be interpreted as representing any funding agencies.

References

- ALLEN-ZHU, Z. and ORECCHIA, L. (2014). Linear coupling: An ultimate unification of gradient and mirror descent. *arXiv preprint arXiv:1407.1537* .
- BANERJEE, A., MERUGU, S., DHILLON, I. S. and GHOSH, J. (2005). Clustering with bregman divergences. *Journal of machine learning research* **6** 1705–1749.
- BREGMAN, L. M. (1967). The relaxation method of finding the common point of convex sets and its application to the solution of problems in convex programming. *USSR computational mathematics and mathematical physics* **7** 200–217.
- BUBECK, S., LEE, Y. T. and SINGH, M. (2015). A geometric alternative to nesterov’s accelerated gradient descent. *arXiv preprint arXiv:1506.08187* .
- CHEN, G. and TEBoulLE, M. (1993). Convergence analysis of a proximal-like minimization algorithm using bregman functions. *SIAM Journal on Optimization* **3** 538–543.
- CHEN, X., LIN, Q. and PENA, J. (2012). Optimal regularized dual averaging methods for stochastic optimization. In *Advances in Neural Information Processing Systems*.
- CHLEBUS, E. (2009). An approximate formula for a partial sum of the divergent p-series. *Applied Mathematics Letters* **22** 732–737.
- DIKONIKOLAS, J. and ORECCHIA, L. (2017). Accelerated extra-gradient descent: A novel accelerated first-order method. *arXiv preprint arXiv:1706.04680* .
- GHADIMI, S. and LAN, G. (2012). Optimal stochastic approximation algorithms for strongly convex stochastic composite optimization i: A generic algorithmic framework. *SIAM Journal on Optimization* **22** 1469–1492.
- GRANT, M. and BOYD, S. (2008). Cvx: Matlab software for disciplined convex programming.
- GÜLER, O. (1992). New proximal point algorithms for convex minimization. *SIAM Journal on Optimization* **2** 649–664.
- HU, B. and LESSARD, L. (2017). Dissipativity theory for nesterov’s accelerated method. *arXiv preprint arXiv:1706.04381* .
- HU, C., PAN, W. and KWOK, J. T. (2009). Accelerated gradient methods for stochastic optimization and online learning. In *Advances in Neural Information Processing Systems*.
- KLOEDEN, P. E. and PLATEN, E. (1992). Higher-order implicit strong numerical schemes for stochastic differential equations. *Journal of statistical physics* **66** 283–314.
- KRICHE, W. and BARTLETT, P. L. (2017). Acceleration and averaging in stochastic mirror descent dynamics. *arXiv preprint arXiv:1707.06219* .
- KRICHE, W., BAYEN, A. and BARTLETT, P. L. (2015). Accelerated mirror descent in continuous and discrete time. In *Advances in neural information processing systems*.
- LAN, G. (2012). An optimal method for stochastic composite optimization. *Mathematical Programming* **133** 365–397.
- LAN, G. and ZHOU, Y. (2015). An optimal randomized incremental gradient method. *arXiv preprint arXiv:1507.02000* .
- LESSARD, L., RECHT, B. and PACKARD, A. (2016). Analysis and design of optimization algorithms via integral quadratic constraints. *SIAM Journal on Optimization* **26** 57–95.
- LUENBERGER, D. G., YE, Y. ET AL. (1984). *Linear and nonlinear programming*, vol. 2. Springer.
- MERTIKOPOULOS, P. and STAUDIGL, M. (2016). On the convergence of gradient-like flows with noisy gradient input. *arXiv preprint arXiv:1611.06730* .
- NEMIROVSKI, A., JUDITSKY, A., LAN, G. and SHAPIRO, A. (2009). Robust stochastic approximation approach to stochastic programming. *SIAM Journal on optimization* **19** 1574–1609.
- NEMIROVSKII, A., YUDIN, D. B. and DAWSON, E. R. (1983). Problem complexity and method efficiency in optimization .
- NESTEROV, Y. (1983). A method of solving a convex programming problem with convergence rate $O(1/k^2)$. In *Soviet Mathematics Doklady*, vol. 27.
- NESTEROV, Y. (2005). Smooth minimization of non-smooth functions. *Mathematical programming* **103** 127–152.
- NESTEROV, Y. (2013). *Introductory lectures on convex optimization: A basic course*, vol. 87. Springer Science & Business Media.
- ØKSENDAL, B. (2003). Stochastic differential equations. In *Stochastic differential equations*. Springer, 65–84.
- POLYAK, B. T. (1963). Gradient methods for minimizing functionals. *Zhurnal Vychislitel’noi Matematiki i Matematicheskoi Fiziki* **3** 643–653.
- POLYAK, B. T. (1964). Some methods of speeding up the convergence of iteration methods. *USSR Computational Mathematics and Mathematical Physics* **4** 1–17.
- RAGINSKY, M. and BOUVRIE, J. (2012). Continuous-time stochastic mirror descent on a network: Variance reduction, consensus, convergence. In *Decision*

and Control (CDC), 2012 IEEE 51st Annual Conference on. IEEE.

ROBBINS, H. and MONRO, S. (1951). A stochastic approximation method. *The annals of mathematical statistics* 400–407.

SU, W., BOYD, S. and CANDÈS, E. (2014). A differential equation for modeling nesterovs accelerated gradient method: Theory and insights. In *Advances in Neural Information Processing Systems*.

WIBISONO, A. and WILSON, A. C. (2015). On accelerated methods in optimization. *arXiv preprint arXiv:1509.03616* .

WIBISONO, A., WILSON, A. C. and JORDAN, M. I. (2016). A variational perspective on accelerated methods in optimization. *Proceedings of the National Academy of Sciences* 201614734.

WILSON, A. C., RECHT, B. and JORDAN, M. I. (2016). A lyapunov analysis of momentum methods in optimization. *arXiv preprint arXiv:1611.02635* .