

Constrained Exploration via Reflected Replica Exchange Stochastic Gradient Langevin Dynamics

Haoyang Zheng^{1*} Hengrong Du^{2*} Qi Feng^{3*} Wei Deng⁴ Guang Lin¹

Abstract

Replica exchange stochastic gradient Langevin dynamics (reSGLD) (Deng et al., 2020a) is an effective sampler for non-convex learning in large-scale datasets. However, the simulation may encounter stagnation issues when the high-temperature chain delves too deeply into the distribution tails. To tackle this issue, we propose reflected reSGLD (r2SGLD): an algorithm tailored for constrained non-convex exploration by utilizing reflection steps within a bounded domain. Theoretically, we observe that reducing the diameter of the domain enhances mixing rates, exhibiting a *quadratic* behavior. Empirically, we test its performance through extensive experiments, including identifying dynamical systems with physical constraints, simulations of constrained multi-modal distributions, and image classification tasks. The theoretical and empirical findings highlight the crucial role of constrained exploration in improving the simulation efficiency.

1. Introduction

Stochastic gradient Langevin dynamics (SGLD) (Welling & Teh, 2011) is the go-to Markov Chain Monte Carlo (MCMC) method in big data. The sampler smoothly transitions from stochastic optimization to sampling as the step size decreases and the injected noise enables exploration for posterior sampling. However, the simulation efficiency is severely affected by the pathological curvature of the energy landscape. To address these challenges, the Quasi-Newton Langevin dynamics (Ahn et al., 2012; Şimşekli et al., 2016; Li et al., 2019) proposes to adapt to the varying curvature for

more efficient exploration of the state space; Hamiltonian Monte Carlo (HMC) introduces auxiliary momentum variables and simulates the Hamiltonian dynamics to explore the periodic orbit (Neal, 2012; Campbell et al., 2021; Zou & Gu, 2021; Wang & Wibisono, 2022); Higher-order numerical techniques offer more efficient simulation by maintaining the stability using a larger stepsize (Chen et al., 2015; Li et al., 2019).

Although the above samplers effectively address pathological curvatures, they are still susceptible to getting trapped in local regions during non-convex sampling. To circumvent this local trap phenomenon, importance samplers (Berg & Neuhaus, 1992; Wang & Landau, 2001; Deng et al., 2020b; 2022) inspired by statistical physics encourage the particles to explore the high energy tails of the distribution; simulated tempering (Marinari & Parisi, 1992; Lee et al., 2018) proposes to escape local optima by adjusting temperatures to explore various energy levels of the distribution, which offer viable solutions to these challenges. Building upon these, replica exchange Langevin dynamics (reLD) (Swendsen & Wang, 1986; Earl & Deem, 2005) and the big data extensions via reSGLD (Deng et al., 2020a) utilize multiple diffusion processes at diverse temperatures and incorporate swapping during training, thereby enabling high-temperature processes to act as a bridge across various local modes. Its accelerated convergence has been both theoretically quantified (Dupuis et al., 2012; Dong & Tong, 2022) and empirically validated (Deng et al., 2020a).

The naïve reSGLD algorithm, while effectively navigating non-convex landscapes, still faces over-exploration issues in high-temperature chains, which is often out of our interest in the optimization perspective. This is partly attributed to non-Arrhenius behavior (Zheng et al., 2007; Sindhikara et al., 2010), where increasing temperatures do not linearly translate to improved exploration efficiency.

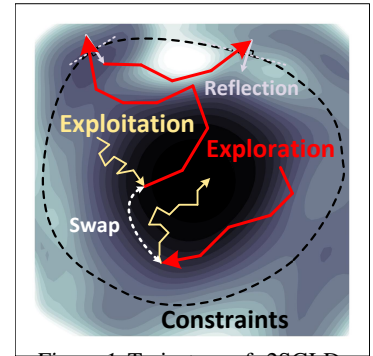


Figure 1. Trajectory of r2SGLD.

^{*}Equal contribution, order determined by coin flip. ¹Purdue University, West Lafayette, IN ²Vanderbilt University, Nashville, TN ³Florida State University, Tallahassee, FL ⁴Machine Learning Research, Morgan Stanley, New York, NY. Correspondence to: Wei Deng <weideng056@gmail.com>, Guang Lin <guanglin@purdue.edu>.

Specifically, over-exploration can result in either exploding or oscillating losses in deep learning training. This phenomenon can deteriorate the model’s stability and optimization performance, and lead to poor predictions. Another interesting interpretation is from the perspective of Thompson sampling approximated by SGLD (Mazumdar et al., 2020; Zheng et al., 2024), where excessive exploration of distribution tails often results in underestimating the optimal arm. This subsequently impedes achieving optimal regret bounds (Jin et al., 2021).

To address this, constrained sampling techniques in MCMC play a pivotal role for different purposes in various forms, such as sampling on explicitly defined manifolds (Lee & Vempala, 2018; Wang et al., 2020), implicitly defined manifolds (Zappa et al., 2018; Lelièvre et al., 2019; Kook et al., 2022; Zhang et al., 2022; Lelièvre et al., 2023), and sampling with moment constraints (Liu et al., 2021).

The proposed algorithm is further guided by constrained gradient Langevin dynamics with explicit form boundaries. A significant contribution to this field was first made by studying the convergence properties of Langevin Monte Carlo in bounded domains (Bubeck et al., 2018), which uncovered a polynomial sample time for log-concave distributions. This research was further extended to non-convex settings later (Lamperski, 2021). Subsequently, Hamiltonian Monte Carlo techniques were employed to advance the field by exploring constrained sampling in the context of ill-conditioned and non-smooth distributions (Kook et al., 2022; Noble et al., 2023). Other notable works include proximal Langevin dynamics (Brosse et al., 2017) for log-concave distributions constrained to convex sets, reflected Schrödinger bridge for constrained generation (Deng et al., 2024), and mirrored Langevin dynamics (Hsieh et al., 2018; Ahn & Chewi, 2021), which focuses on convex settings.

Based on the previous work, we propose an adapted sampler for constrained exploration—reflected replica exchange SGLD (r2SGLD)—to specifically address over-exploration in high-temperature chains and improve the mixing rates in challenging non-convex tasks. The r2SGLD algorithm employs parallel Langevin dynamics with swaps at varying temperatures, which achieves exploration through high-temperature chains and exploitation through low-temperature chains (Figure 1). Chains wandering outside the constrained domain are brought back by a reflection operation, which involves the construction of a tangent line at the nearest boundary for symmetrical reflection. By carefully constraining the exploration domain, we aim to harness the benefits of the high-temperature chain without succumbing to the detriments of over-exploration. A swapping mechanism between chains with different temperatures is subsequently used to balance between exploration and exploitation. We summarize the main contribution to this work as follows:

(1) We contribute to the theoretical landscape by proving that r2SGLD outperforms the naïve reSGLD. For the continuous-time analysis, we provide a quantitative mixing rate, marked by a refined decay rate of $\tilde{O}(1/\text{diam}(\Omega)^2)$, where $\text{diam}(\Omega)$ denotes the domain diameter.

(2) This work introduces the novel use of constrained gradient Langevin dynamics in the identification of dynamical systems, which marks the first known application in this domain and broadens the methodological toolkit for dynamical system analysis.

(3) Extensive testing of r2SGLD against diverse baselines in both bounded 2D multi-modal distribution simulation and large-scale deep learning tasks further confirmed its superior efficacy.

Remark 1.1. It should be noted that compared with the naïve reflected gradient Langevin dynamics, studying the quantitative value of the spectral gap in r2SGLD remains a fundamental challenge. For example, Bakry et al. (2008) delves into a crude lower bound for the spectral gap concerning vanilla Langevin diffusion, where lower values imply faster convergence, particularly on sub-Gaussian distributions. Although these distributions are weaker than log-concavity, they still preclude the emergence of significant non-convexity.

The extension to Gaussian mixture distributions has posed a significant challenge. Dong & Tong (2020) investigates the considerable enhancement of the spectral gap with replica exchange Langevin diffusion compared to vanilla Langevin diffusion. We anticipate a similar acceleration would persist in bounded domains. However, given the methodological nature of our work, we recognize this topic and extensions to general non-convexity as beyond the current scope and defer it to future investigations.

2. Methodology

This section introduces the reflected reLD (r2LD), its infinitesimal generator, and the proposed r2SGLD algorithm. Our discussion centers on their contribution to efficient exploration in constrained non-convex sampling, which offers insights into its underlying mechanisms.

2.1. Reflected Replica Exchange Langevin Diffusion

The fundamental principle of the r2SGLD algorithm lies in the r2LD concept, which involves running multiple Langevin diffusions simultaneously, each at a different temperature. The system dynamics are described by the following stochastic differential equations:

$$\begin{aligned} d\beta_t^{(1)} &= -\nabla U(\beta_t^{(1)})dt + \sqrt{2\tau_1}d\mathbf{W}_t^{(1)} + \nu(\beta_t^{(1)})L^{(1)}(dt), \\ d\beta_t^{(2)} &= -\nabla U(\beta_t^{(2)})dt + \sqrt{2\tau_2}d\mathbf{W}_t^{(2)} + \nu(\beta_t^{(2)})L^{(2)}(dt), \end{aligned} \quad (1)$$

where $\beta_t^{(1)}, \beta_t^{(2)} \in \mathbb{R}^d$ represent the parameter sets at respective temperatures τ_1 and τ_2 . The term $\nabla U(\beta_t)$ denotes the gradient of the potential, and $\mathbf{W}_t^{(1)}$ and $\mathbf{W}_t^{(2)}$ are independent standard Wiener processes. The function $\nu(\beta_t)$ represents the inner unit normal vector at a point on the boundary $\partial\Omega$, while $L^{(1)}$ and $L^{(2)}$ denote independent local times with reference to $\partial\Omega$.

Our focus is on restricting the process to a compact domain $\Omega \subset \mathbb{R}^d$ with a boundary $\partial\Omega$. We ensure reflected boundary conditions through the terms $\nu(\beta^{(1)})L^{(1)}(dt)$ and $\nu(\beta^{(2)})L^{(2)}(dt)$. We further demonstrate that the r2LD, as specified in (1), converges to an invariant distribution $\pi(\cdot, \cdot)$:

$$d\pi(x_1, x_2) = \frac{1}{Z} \underbrace{e^{-\frac{U(x_1)}{\tau_1} - \frac{U(x_2)}{\tau_2}}}_{p(x_1, x_2)} dx_1 dx_2, \quad (2)$$

where Z is a normalizing constant:

$$Z = \int_{\Omega \times \Omega} p(x_1, x_2) dx_1 dx_2.$$

The density of this distribution, $p(\cdot, \cdot)$, applies only within the domain $\Omega \times \Omega$ and the density is zero outside this domain.

The swap between chains results in positional changes from $(\beta_t^{(1)}, \beta_t^{(2)})$ to $(\beta_{t+dt}^{(1)}, \beta_{t+dt}^{(2)})$, at a rate determined by $r(1 \wedge S(\beta_t^{(1)}, \beta_t^{(2)}))dt$, where $r \geq 0$ denotes the swapping intensity to regulate the frequency of swapping configurations between the two processes. The swap function $S(\cdot, \cdot)$ is given as follows:

$$S(\beta_t^{(1)}, \beta_t^{(2)}) := e^{\left(\frac{1}{\tau_1} - \frac{1}{\tau_2}\right)(U(\beta_t^{(1)}) - U(\beta_t^{(2)}))}. \quad (3)$$

Similar to the Langevin diffusion in Chen et al. (2019), r2LD behaves as a reversible Markov jump process, which converges to the same invariant distribution (2).

2.2. The Infinitesimal Generator

The process $\beta_t = (\beta_t^{(1)}, \beta_t^{(2)})$ follows a Markov diffusion process with infinitesimal generator $\mathcal{L}$ to encompass both the r2LD (1) and swapping dynamics (3). For (x_1, x_2) in the interior of $\Omega \times \Omega$, the generator $\mathcal{L}$ applies to a smooth function f and is structured into distinct components: $\mathcal{L}^{(1)}f$ and $\mathcal{L}^{(2)}f$ for the dynamics of x_1 and x_2 , and $\mathcal{L}^{(s)}f$ for swap dynamics:

$$\begin{aligned} \mathcal{L}f = & \underbrace{-\langle \nabla_{x_1} f(x_1, x_2), \nabla U(x_1) \rangle + \tau_1 \Delta_{x_1} f(x_1, x_2)}_{\mathcal{L}^{(1)}f} \\ & \underbrace{-\langle \nabla_{x_2} f(x_1, x_2), \nabla U(x_2) \rangle + \tau_2 \Delta_{x_2} f(x_1, x_2)}_{\mathcal{L}^{(2)}f} \\ & \underbrace{+ r S(x_1, x_2) \cdot (f(x_2, x_1) - f(x_1, x_2))}_{\mathcal{L}^{(s)}f}, \end{aligned} \quad (4)$$

where $\nabla_{x_i}, \Delta_{x_i}$ denote the gradient and Laplacian operator with respect to x_i , $i = 1, 2$, and the entire formulation is subject to the Neumann boundary conditions (Menaldi & Robin, 1985; Kang & Ramanan, 2014), also (Wang, 2014) (Theorem 3.1.3):

$$\begin{cases} \nabla_{x_1} f(x_1, x_2) \cdot \nu(x_1) = 0 & (x_1, x_2) \in \partial\Omega \times \Omega, \\ \nabla_{x_2} f(x_1, x_2) \cdot \nu(x_2) = 0 & (x_1, x_2) \in \Omega \times \partial\Omega. \end{cases} \quad (5)$$

For the invariant measure π that is associated with a smooth potential function $U(\cdot)$ as in (2), the corresponding Dirichlet form that appears in the Poincare inequality (10) and Logarithmic Sobolev (Log-Sobolev) inequality (14) is given as follows:

$$\mathcal{E}(f) = \int \left(\tau_1 \|\nabla_{x_1} f\|^2 + \tau_2 \|\nabla_{x_2} f\|^2 \right) d\pi(x_1, x_2). \quad (6)$$

Furthermore, as established in Theorem 3.3 (Chen et al., 2019), the Dirichlet form $\mathcal{E}_S$ linked to the infinitesimal generator $\mathcal{L}^{(s)}$ includes an additional term representing acceleration:

$$\mathcal{E}_S(f) = \mathcal{E}(f) + \underbrace{\frac{r}{2} \int S(x_1, x_2) \cdot (f(x_2, x_1) - f(x_1, x_2))^2 d\pi(x_1, x_2)}_{\text{acceleration term}},$$

which induces positive acceleration under mild conditions, and it is vital for rapid convergence in the 2-Wasserstein ($\mathcal{W}_2$) distance. Notably, the acceleration's effectiveness depends on the swap function (3).

2.3. The Proposed Algorithm

r2SGLD is the practical interpretation of the continuous process described by the infinitesimal generator in (1). Following Algorithm 1, it starts with updating the parameters $\tilde{\beta}^{(1)}$ and $\tilde{\beta}^{(2)}$ for the diffusion process using stochastic gradients and noises, followed by reflection operations $\mathcal{R}(\cdot)$ to reflect out-of-boundary samples back to the bounded domain Ω . The gradient,

$$\nabla \tilde{U}(\tilde{\beta}_k^{(i)}) = \frac{N}{|\mathbb{B}_k|} \sum_{\mathcal{D}_j \in \mathbb{B}_k} \nabla U(\mathcal{D}_j | \tilde{\beta}_k^{(i)}), \quad i = 1, 2,$$

is estimated by a mini-batch data $\mathbb{B}_k$, and $\{\mathcal{D}_j\}_{j=1}^N$ is the dataset. Next, the algorithm determines whether to swap the chains by comparing a uniformly generated random number u against the corrected swapping intensity $\tilde{S}$:

$$\tilde{S}(\tilde{\beta}_k^{(1)}, \tilde{\beta}_k^{(2)}) = e^{\left(\frac{1}{\tau_1} - \frac{1}{\tau_2}\right)(\tilde{U}(\tilde{\beta}_k^{(1)}) - \tilde{U}(\tilde{\beta}_k^{(2)}) - \left(\frac{1}{\tau_1} - \frac{1}{\tau_2}\right) \frac{\tilde{\sigma}^2}{\mathcal{C}})}, \quad (7)$$

where $\tilde{\sigma}^2$ approximates the variance of $\tilde{U}(\tilde{\beta}_{k+1}^{(1)}) - \tilde{U}(\tilde{\beta}_{k+1}^{(2)})$ and $\mathcal{C}$ acts as an adjustment to balance acceleration and bias. The algorithm proceeds iteratively until a specified number of iterations is reached. The parameter sets $\{\tilde{\beta}_k^{(i)}\}_{k=1}^{K+1}$ are produced as outputs for analytical purposes.

Algorithm 1 The r2SGLD Algorithm.

Input Initial parameters $\tilde{\beta}_1^{(1)}, \tilde{\beta}_1^{(2)}$.

Input Number of iterations K .

Input Temperatures τ_1, τ_2 .

Input Correction factor $\mathcal{C}$.

Input Learn rate η .

for $k = 1, 2, \dots, K$ **do**
Sampling Step

$$\tilde{\beta}_{k+1}^{(1)} = \mathcal{R}\left(\tilde{\beta}_k^{(1)} - \eta \nabla \tilde{U}(\tilde{\beta}_k^{(1)}) + \sqrt{2\eta\tau_1}\xi_k^{(1)}\right),$$

$$\tilde{\beta}_{k+1}^{(2)} = \mathcal{R}\left(\tilde{\beta}_k^{(2)} - \eta \nabla \tilde{U}(\tilde{\beta}_k^{(2)}) + \sqrt{2\eta\tau_2}\xi_k^{(2)}\right).$$

Swapping Step

 Generate a uniform random number $u \in [0, 1]$.

 Compute $\tilde{S}$ follows (7).

if $u < \tilde{S}$ **then**

 Swap $\tilde{\beta}_{k+1}^{(1)}$ and $\tilde{\beta}_{k+1}^{(2)}$.

end if
end for
Output Parameters $\{\tilde{\beta}_k^{(1)}\}_{k=1}^{K+1}$.

3. Theoretical Analysis

In this section, we outline the convergence analysis of continuous-time r2LD under the χ^2 -divergence and $\mathcal{W}_2$ -distance, which extends the previous analysis (Chen et al., 2019; Deng et al., 2020a) within a constrained domain. We further highlight the acceleration benefits achieved through parameter swapping and discuss how the convergence rate is influenced by the diameter of Ω . Lastly, our analysis includes an evaluation of discretization error in the 1-Wasserstein ($\mathcal{W}_1$) distance. Throughout the analysis, we adopt the following assumptions.

Assumption A1. Ω is a compact domain with a boundary $\partial\Omega$ whose *second fundamental form* is bounded below by some constant $\kappa \leq 0$.

Remark 3.1. For any two tangent vectors v_1, v_2 at $x \in \partial\Omega$, the *second fundamental form* of $\partial\Omega$ is defined by $\mathbb{I}(v_1, v_2)(x) = -\langle \nabla_{v_1} \nu(x), v_2 \rangle$, where $\nu(x)$ is the normal vector field, $\nabla_{v_1} \nu$ denotes the directional derivative of $\nu(x)$ along v_1 . Note that if $\mathbb{I}(v_1, v_2)(x) \geq \kappa$ for some $\kappa \leq 0$, the second fundamental form of the boundary $\partial\Omega$ is bounded below by κ . It is also worth mentioning that if $\kappa = 0$, Ω is convex.

Assumption A2. The function $U \in C^2(\Omega)$. Since Ω is compact, there exists an $L > 0$ such that for all $x, y \in \Omega$,

$$\|\nabla U(x) - \nabla U(y)\| \leq L\|x - y\|. \quad (8)$$

Unless specified otherwise, the theoretical results in this work are based on Assumptions A1 and A2.

3.1. Convergence in χ^2 -Divergence

Our analysis begins by identifying the invariant distribution of r2LD, which delves into the basic properties of the generator $\mathcal{L}$ subject to the boundary condition (5):

Lemma 3.2. $\{\beta_t\}_{t \geq 0}$ is reversible and its invariant distribution π is given by (2).

We define the χ^2 -divergence as

$$\chi^2(\mu_t \| \pi) = \int_{\Omega \times \Omega} \left(\frac{d\mu_t}{d\pi} - 1 \right)^2 d\pi, \quad (9)$$

where $d\mu_t/d\pi$ is the Radon–Nikodym derivative between μ_t and π . The convergence rate of μ_t in χ^2 -divergence is related to the Poincaré inequality:

Lemma 3.3. For any probability measure $\mu \ll \pi$ where the Radon–Nikodym derivative $d\mu/d\pi$ satisfies (5), the following inequality holds:

$$\chi^2(\mu \| \pi) \leq C_P \mathcal{E} \left(\frac{d\mu}{d\pi} \right), \quad (10)$$

where the best constant C_P is called *Poincaré constant*.

Here is our primary convergence result:

Theorem 3.4. Given any initial measure μ_0 for which $d\mu_0/d\pi$ satisfies (5), the χ^2 -divergence to the invariant distribution π decays exponentially according to:

$$\chi^2(\mu_t \| \pi) \leq \chi^2(\mu_0 \| \pi) \exp(-2t(1 + \eta_S)C_P^{-1}), \quad (11)$$

where $\eta_S := \inf_{t > 0} \frac{\mathcal{E}_S \left(\frac{d\mu_t}{d\pi} \right)}{\mathcal{E} \left(\frac{d\mu_t}{d\pi} \right)} - 1$ is the acceleration effect in χ^2 -divergence.

In addition to the observed acceleration effect, our findings highlight the significance of the Poincaré constant C_P in convergence, as shown in (11), where its lower bound directly influences the convergence rate. In the literature, C_P is often referred to the *first Neumann eigenvalue* of $\mathcal{L}$ or *spectral gap*, denoted by $\lambda_1(\mathcal{L})$. For a non-convex domain that satisfies Assumption A1, it has been shown in Wang (2014) (Corollary 3.5.2) that

Lemma 3.5. The Poincaré constant C_P , defined in Lemma 3.3, satisfies the inequality:

$$C_P = \lambda_1(\mathcal{L}) \geq C_1 \left(\frac{\pi^2}{\text{diam}(\Omega)^2} + C_2 \right), \quad (12)$$

where C_1 and C_2 are constants depending on U and κ , in particular, when Ω is convex, we have $C_1 = 1$ and $C_2 = 0$, i.e.,

$$\lambda_1(\mathcal{L}) \geq \frac{\pi^2}{\text{diam}(\Omega)^2}.$$

Remark 3.6. For a non-convex domain Ω whose second fundamental form is bounded below by κ , applying a conformal change of the Euclidean metric (as outlined in Wang (2007), Lemma 2.1) allows for a simplification to a convex domain case, which results in a less explicit bound specified in (12). However, it should be noted that $C_P = \mathcal{O}(1/\text{diam}(\Omega)^2)$ still holds true when $\text{diam}(\Omega) \ll 1$.

3.2. Convergence in 2-Wasserstein Distance

We introduce the p -Wasserstein distance between two Borel probability measures μ and π on $\mathbb{R}^{2d}$

$$\mathcal{W}_p(\mu, \pi) := \inf_{\gamma \in \Gamma(\mu, \pi)} \left(\int_{\mathbb{R}^{2d} \times \mathbb{R}^{2d}} \|x - y\|^p d\gamma(x, y) \right)^{1/p},$$

where $\Gamma(\mu, \pi)$ denotes all joint coupled distributions γ with marginal distributions μ and π . Furthermore, we define the Kullback–Leibler (KL) divergence as

$$D(\mu \| \pi) := \int_{\Omega \times \Omega} \left(\frac{d\mu}{d\pi} \right) \ln \left(\frac{d\mu}{d\pi} \right) d\pi. \quad (13)$$

The convergence rate under $\mathcal{W}_2$ -distance is dependent on the Log-Sobolev inequality which controls the entropy throughout the Fisher information. The proof of the following lemma is referred to Wang (2014) (Corollary 3.5.3).

Lemma 3.7. *For any probability measure $\mu \ll \pi$ and satisfying $d\mu/d\pi \geq 0$ and (5), the following Logarithmic Sobolev (Log-Sobolev) inequality holds:*

$$D(\mu \| \pi) \leq C_{\text{LS}} \mathcal{E} \left(\sqrt{\frac{d\mu}{d\pi}} \right). \quad (14)$$

where the best constant C_{LS} is called Log-Sobolev constant.

Theorem 3.8. *Given any initial measure μ_0 for which $d\mu_0/d\pi \geq 0$ and satisfying (5), the 2-Wasserstein distance between μ_t and π satisfies the following accelerated exponential decay estimate:*

$$\mathcal{W}_2(\mu_t, \pi) \leq \sqrt{2C_{\text{LS}} D(\mu_0 \| \pi)} \exp(-t(1 + \delta_S)C_{\text{LS}}^{-1}), \quad (15)$$

where $\delta_S := \inf_{t>0} \frac{\mathcal{E}_S \left(\sqrt{\frac{d\mu_t}{d\pi}} \right)}{\mathcal{E} \left(\sqrt{\frac{d\mu_t}{d\pi}} \right)} - 1$ is the acceleration effect in $\mathcal{W}_2$ -distance.

Now we apply the estimate for the Log-Sobolev constant for non-convex domain (Wang, 2007) (Corollary 3.5.3):

Lemma 3.9. *The Log-Sobolev constant, defined in Lemma 3.7, satisfies the following inequality:*

$$C_{\text{LS}} \geq \frac{C}{\text{diam}(\Omega)^2}, \quad (16)$$

where C is a constant depending on κ and U . In particular, when Ω is convex, we have

$$C_{\text{LS}} \geq \frac{\sqrt{1 + 4\pi^2} - 1}{2 \cdot \text{diam}(\Omega)^2}. \quad (17)$$

Remark 3.10. For non-convex domain Ω , as highlighted in Remark 3.6, the bound presented in (16) is less explicit.

For detailed derivations, readers may refer to Appendix B.

3.3. Discretization Analysis

Due to the lack of higher-order local time estimates in the reflection term, our study on discrete-time dynamics of r2SGLD reveals its unique aspects compared to naïve reSGLD. Following the techniques in Tanaka (1979); Bubeck et al. (2018), we successfully derive the upper bound for the discretization error within the $\mathcal{W}_1$ -distance framework.

Theorem 3.11 (Discretization error). *Assume that the domain Ω is convex, and Assumptions A1, A2 hold true, then*

$$\mathcal{W}_1(\mu_T, \tilde{\mu}_T) \leq \tilde{\mathcal{O}} \left(\eta^{1/4} + \sqrt{\max_k \mathbb{E}[\|\phi_k\|^2]} + \sqrt{\eta^{-1/2} \max_k \sqrt{\mathbb{E}[\|\psi_k\|^2]}} \right).$$

where $\tilde{\mu}_T$ denotes the distribution of $\tilde{\beta}_T^\eta$, which is the continuous-time interpolation for r2SGLD, $\phi_k := \nabla \tilde{U} - \nabla U$ is the noise in the stochastic gradient, and $\psi_k := \tilde{S} - S$ is the noise in the stochastic swapping rate.

Interested readers can refer to Appendix C for details.

4. Experiments

To validate the r2SGLD algorithm*, we begin by applying it to dynamical system identification (Section 4.1), where the physical constraints are inherent in the dynamical system. Subsequently, its effectiveness in multi-mode distribution simulation is detailed in Section 4.2. Lastly, Section 4.3 illustrates the algorithm’s performance in deep learning tasks.

4.1. Identifying the Lorenz Systems

The Lorenz system is a system of ordinary differential equations that underscores that chaotic systems can be completely deterministic and yet still be inherently unpredictable over long periods of time:

$$\begin{aligned} \dot{x}(t) &= \sigma(y - x), \\ \dot{y}(t) &= x(\rho - z) - y, \\ \dot{z}(t) &= xy - \beta z, \end{aligned} \quad (18)$$

where the system dynamics is determined by three unknown parameters: σ (the Prandtl number), ρ (the Rayleigh number), and β (the aspect ratio)[†]. A prevalent approach to

*Code is available at github.com/haoyangzheng1996/r2SGLD

[†]Following traditional conventions in the Lorenz system, we use β to differentiate it from the parameters β as defined in (1).

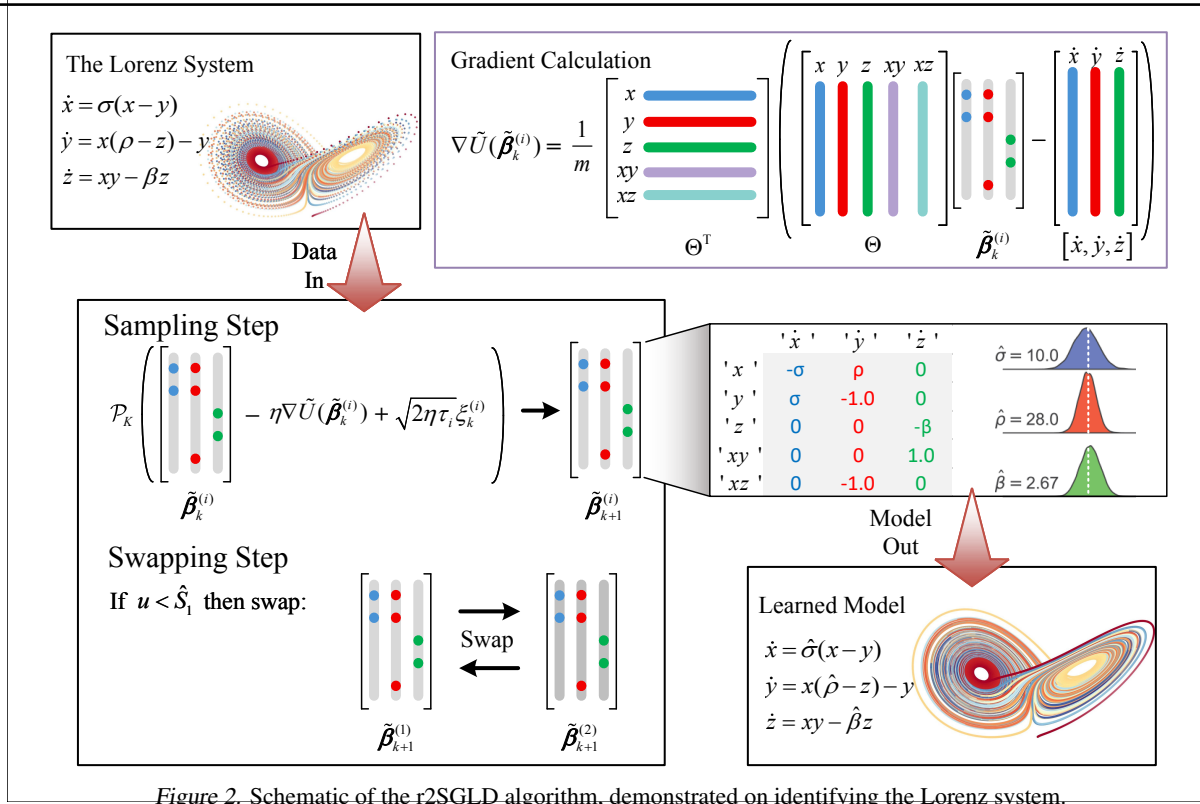


Figure 2. Schematic of the r2SGLD algorithm, demonstrated on identifying the Lorenz system.

learning the Lorenz system involves the Sparse Identification of Nonlinear Dynamics (Brunton et al., 2016; de Silva et al., 2020; Kaptanoglu et al., 2022). This method employs thresholding least squares using the position and its time derivative for $x(t)$, $y(t)$, and $z(t)$ at specific times:

$$\dot{\mathbf{X}} = \Theta(\mathbf{X})\beta, \quad (19)$$

where $\dot{\mathbf{X}}$ is the concatenation of multiple state velocities $[\dot{x}(t_m), \dot{y}(t_m), \dot{z}(t_m)]$ sampled at several times $t_1, t_2, \dots, t_m, \dots$. $\Theta(\mathbf{X})$ is the candidate basis from the sampled states, and β is a sparse matrix to determine which candidate basis is active. In our work, this method is tailored by simplifying the sparse coefficient matrix, which targets the learning of the unknown parameters σ , ρ , and β . Another critical aspect is that our approach relies solely on gradient information of the potentials according to state velocities, rather than having access to the sampled states:

$$\nabla \tilde{U}(\tilde{\beta}_k) = \frac{1}{m} \Theta^\top (\Theta \tilde{\beta}_k - \dot{\mathbf{X}}), \quad (20)$$

where the matrix Θ from $\mathcal{R}^{m \times 5}$ here encompasses five candidate basis functions associated with $x(t_m)$, $y(t_m)$, $z(t_m)$, $x(t_m)y(t_m)$, and $x(t_m)z(t_m)$. Additionally, the matrix $\tilde{\beta} \in \mathcal{R}^{5 \times 3}$, an approximation of β containing parameters σ, ρ, β , interacts with $\dot{\mathbf{X}}$ from $\mathcal{R}^{m \times 3}$, representing m sampled state velocities from t_1 to t_m .

A general framework to identify the Lorenz system with r2SGLD algorithm is shown in Figure 2. The data for this

work is generated using a six-stage, fifth-order Runge-Kutta method (Atkinson, 1991). This method is applied to compute states from time 0 to 100 at intervals of 0.01, followed by calculating the velocity of $\dot{\mathbf{X}}$ using elementary matrix operations. Once the sampling and swapping steps are complete in each iteration, the Runge-Kutta method is used with the latest sampled model parameters, which aims to derive both the approximate states matrix $\tilde{\mathbf{X}}$ and the candidate basis $\Theta(\tilde{\mathbf{X}})$ for the relevant time stamp. Upon convergence of the empirical distribution to its stationary counterpart or upon satisfying certain stopping conditions, the sampling and swapping steps are halted. The outputs are the empirical posterior modes of the unknown parameters, which are used to recover the target dynamical systems.

We apply dual-chain reSGLD and r2SGLD for sampling within the parameter space to estimate the posterior mode of unknown parameters ($\hat{\sigma}$, $\hat{\rho}$, and $\hat{\beta}$). For r2SGLD, its sampling process adheres to physical constraints requiring the positivity of parameters ($\sigma, \rho, \beta > 0$). To ensure global stability within the system, two additional constraints are implemented: one mandates a high rate of viscous dissipation to inhibit amplification of minor perturbations in the system's velocity field ($\sigma > 1 + \beta$), and the other requires strong thermal forcing to counterbalance the stabilizing effects of viscous dissipation ($\rho > 1$). To streamline the algorithm's performance, we enforce accurate values for model parameters associated with the less critical candidate basis, which directs the algorithm's focus primarily toward

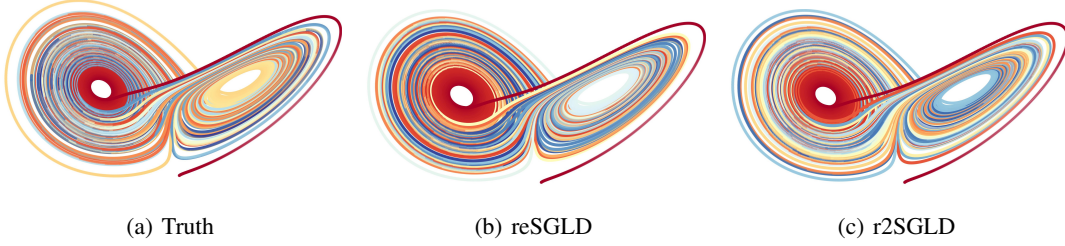


Figure 3. Simulation of the Lorenz system based on the empirical posterior modes of model parameters.

identifying the essential model parameters.

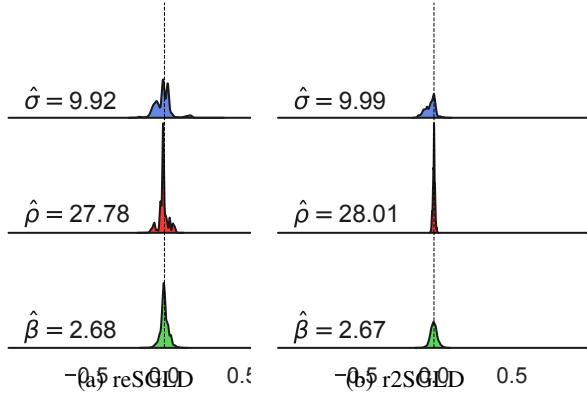


Figure 4. Posterior distributions of the identified model parameters.

Upon adhering to the specified conditions, we evaluate the effectiveness of reflection in reSGLD by contrasting the empirical posteriors with reflection (Figure 4(b)) against those without reflection (Figure 4(a)). The presented figures illustrate estimated values of the empirical posterior modes of $\hat{\sigma}$, $\hat{\rho}$, and $\hat{\beta}$, as generated by the reSGLD and r2SGLD methods. Notably, the posterior distributions displayed here are normalized to facilitate a more straightforward and clear comparison. The result reveals that the posterior modes derived from r2SGLD align more closely with the true parameters ($\sigma = 10.0$, $\rho = 28.0$, $\beta = \frac{8}{3}$) compared to those from reSGLD. Notably, the empirical posterior $\hat{\sigma}$ from reSGLD exhibits dual peaks, whereas r2SGLD’s posterior demonstrates a singular peak, which indicates a more reasonable result. For simulating the dynamics of the Lorenz system, the learned posterior modes are utilized as model parameters, with the simulations presented in Figure 3. While the simulations from reSGLD and r2SGLD are closely matched, r2SGLD demonstrates slightly improved results over extended simulation periods. This implies that r2SGLD yields more reliable outcomes. Further details and results, including baseline comparisons and exploration of identifying the Lotka-Volterra model, are available in A.1 and A.2.

4.2. Constrained Multi-modal Simulations

In this study, we examine the enhanced sample efficiency of r2SGLD through simulating bounded multi-modal distributions. Our proposed algorithm (with two chains), along-

side standard baselines such as (reflected) SGLD (SGLD and R-SGLD), (reflected) cyclical SGLD (cycSGLD and R-cycSGLD) (Zhang et al., 2020), and reSGLD, employs gradient information to generate samples and approximate posterior distributions. This comparative analysis aims to demonstrate the acceleration effect offered by the r2SGLD.

Our study focuses on a multi-modal density characterized by a 2D mixture of 25 Gaussian distributions within a flower-shaped boundary, as depicted in Figure 5(a). The boundary is mathematically defined as:

$$r = \sin(2\pi pt) + m, \quad x = r \cos(2\pi t), \quad y = r \sin(2\pi t),$$

where r denotes the radius, p the number of petals, m the extent of radial displacement, and x, y are the horizontal and vertical coordinates, respectively. For the target distribution, we set the boundary with parameters $p = 5$ and $m = 3$.

The empirical posterior distributions generated by constrained sampling methods (R-SGLD, R-cycSGLD, and r2SGLD) are exhibited in Figure 5. In this context, R-SGLD (Figure 5(b)) exhibits the least effective performance, which notably fails to quantify the weights of each mode accurately. The performance is moderately improved with R-cycSGLD (Figure 5(c)), but it remains improvement spaces. In contrast, the results given by r2SGLD (Figure 5(d)) show commendable performance.

Additionally, we consider a penalty-based approach as a baseline for simulating constrained multi-modal distributions. We incorporated an L2 regularization term in the objective function to penalize samples that exit the designated region, defined as:

$$\nabla \tilde{U}_{\text{reg}}(\tilde{\beta}_k) = \nabla \tilde{U}(\tilde{\beta}_k) + \xi \tilde{\beta}_k$$

where ξ represents the shrinkage coefficient. We alternate between the unpenalized objective $\tilde{U}(\tilde{\beta}_k)$ within the region and the penalized $\tilde{U}_{\text{reg}}(\tilde{\beta}_k)$ when out-of-bounds. This approach is evaluated using three penalized variants: penalized SGLD (P-SGLD), penalized cyclic SGLD (P-cycSGLD), and penalized reSGLD (P-reSGLD). Figure 5 (e) presents a comparative analysis of various sampling methods, including SGLD, cycSGLD, reSGLD, P-SGLD, P-cycSGLD, P-reSGLD, R-SGLD, R-cycSGLD, and r2SGLD, focusing on their KL divergence under the same computational cost.

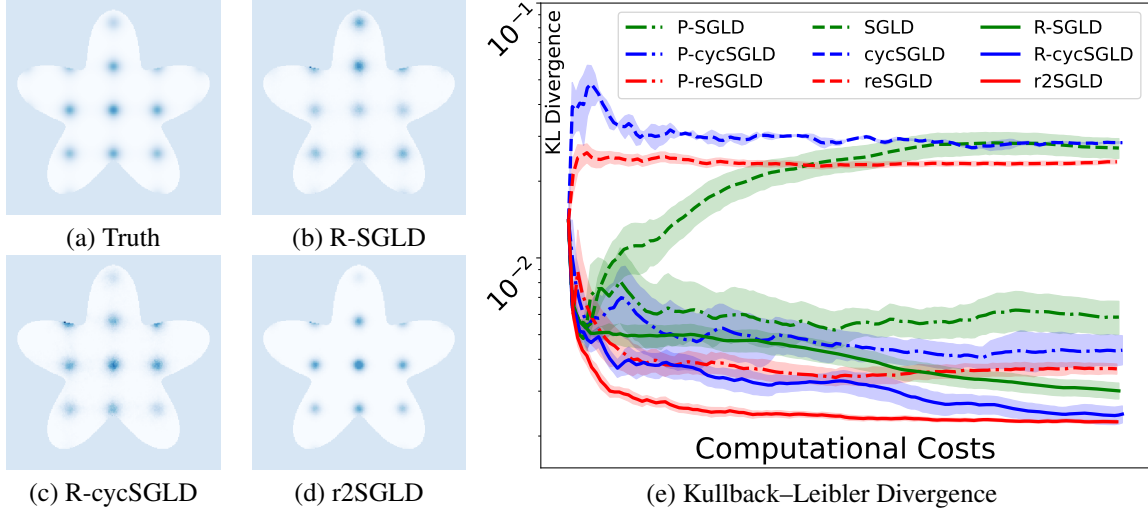


Figure 5. Empirical behavior on multi-mode distributions with flower-shaped boundaries.

This analysis incorporates 95% confidence intervals derived from 10 runs under different initial conditions.

From the result, the implementation of a reflection operation markedly reduces KL divergence by improving sample efficiency through the redirection of out-of-bound samples. While the L2 regularization helps maintain samples within the boundary, it confirms that the reflection-based algorithm outperforms these penalized variants. The comparative inefficiency of the penalized approaches stems from their acceptance of some out-of-bound samples, which compromises the sample efficiency. This further highlights the superiority of the r2SGLD algorithm in avoiding over-exploration and improving sample efficiency. Among the reflection-based sampling methods, R-SGLD reports the least favorable KL divergence, possibly due to its tendency to remain in local modes. R-cycSGLD outperforms R-SGLD owing to its adaptive learning rate, which facilitates both exploration and exploitation. The most effective performance is observed in r2SGLD, where the dual-chain structure allows for more efficient exploration.

To empirically substantiate our theoretical assertion that reducing the domain’s diameter enhances mixing rates with quadratic dependency, we conducted further experiments with r2SGLD. These experiments explored the impact of domain diameters, set between 1.5 and 3.0, on mixing rates within convex bounded domains versus a non-convex distribution with 25 Gaussian modes. For detailed experimental setups and further discussions, please see Section A.3.

4.3. Non-convex Optimization for Image Classifications

We further extend the testing to CIFAR100 benchmarks, which utilize 20 and 56-layer residual networks (ResNet20 and ResNet56, respectively) for training and testing. Our performance evaluation considers modified baseline algo-

rithms with momentum terms, which is crucial for intricate tasks like image classification. This integration of gradient and curvature information enables a more effective exploration of the parameter space to improve model performance. Therefore, our algorithmic baseline includes stochastic gradient with momentum (SGDM), its reflected variant (R-SGDM), (reflected) stochastic gradient Hamiltonian Monte Carlo (SGHMC and R-SGHMC), (reflected) cyclical SGHMC (cycSGHMC and R-cycSGHMC) featuring a cyclical learning rate. To highlight the performance of the proposed algorithm, we include (reflected) replica exchange SGHMC (reSGHMC and r2SGHMC) for the test.

Our study on CIFAR 100 with the proposed algorithm first unveiled a notable advantage: employing reflection operation enables larger learning rates[‡], which is fundamental for detailed model exploration. Algorithms without reflected projection often exhibit failure (test accuracy under 10% or infinite training loss), particularly at initial learning rates over 10.0[§]. Conversely, setting reflection operation between $[-4.0, 4.0]$ not only stabilizes the model but also improves test accuracy to over 70%. This highlights the critical role of reflected projection in ensuring successful learning in complex tasks. Moreover, it indicates the potential for state-of-the-art results given sufficient computational resources for model exploration.

Although a large learning rate has the potential to offer better

[‡]In empirical Deep Neural Network (DNN) training using SGD, a larger learning rate not only increases discretization errors but also effectively raises the temperature of the resulting Markov chain (Mandt et al., 2017). We specify that “large learning rates” are primarily employed during the initial training stages to explore the landscape, and are gradually decayed to implement simulated annealing for enhanced global optimization.

[§]The initial learning rate for CIFAR 100 is scaled to $2e-4$ when accounting for the training data size of 50,000.

Table 1. UNCERTAINTY ESTIMATION COMPARISON BETWEEN ALGORITHMS WITH REFLECTION AND WITHOUT REFLECTION.

METHODS	METRICS (RESNET20)		
	ACC (%) $\uparrow$	NLL $\downarrow$	BRIER (%) $\downarrow$
SGDM	72.13 $\pm$ 0.60	9667 $\pm$ 108	2.78 $\pm$ 0.05
SGHMC	72.47 $\pm$ 0.45	9543 $\pm$ 157	2.75 $\pm$ 0.05
cycSGHMC	73.49 $\pm$ 0.17	8913 $\pm$ 76	2.65 $\pm$ 0.02
reSGHMC	75.01 $\pm$ 0.14	8552 $\pm$ 69	2.50 $\pm$ 0.01
R-SGDM	72.43 $\pm$ 0.35	9626 $\pm$ 94	2.75 $\pm$ 0.03
R-SGHMC	72.85 $\pm$ 0.51	9501 $\pm$ 167	2.73 $\pm$ 0.05
R-cycSGHMC	73.77 $\pm$ 0.22	8953 $\pm$ 52	2.62 $\pm$ 0.02
r2SGHMC	75.38 $\pm$ 0.17	8489 $\pm$ 66	2.46 $\pm$ 0.02

METHODS	METRICS (RESNET56)		
	ACC (%) $\uparrow$	NLL $\downarrow$	BRIER (%) $\downarrow$
SGDM	74.40 $\pm$ 0.71	9724 $\pm$ 169	3.59 $\pm$ 0.23
SGHMC	74.22 $\pm$ 0.66	9723 $\pm$ 214	3.23 $\pm$ 0.21
cycSGHMC	77.98 $\pm$ 0.61	8303 $\pm$ 161	3.19 $\pm$ 0.20
reSGHMC	78.87 $\pm$ 0.44	7406 $\pm$ 130	2.94 $\pm$ 0.06
R-SGDM	74.70 $\pm$ 0.68	9507 $\pm$ 106	3.53 $\pm$ 0.18
R-SGHMC	75.10 $\pm$ 0.55	9232 $\pm$ 158	3.36 $\pm$ 0.23
R-cycSGHMC	78.41 $\pm$ 0.67	7711 $\pm$ 144	3.12 $\pm$ 0.11
r2SGHMC	79.39 $\pm$ 0.30	7155 $\pm$ 91	2.89 $\pm$ 0.02

model performance, we do not directly use a learning rate of more than 10.0 due to computational resource limitations. Instead, we investigate the optimal learning rate across 1,000 epochs for our proposed algorithms and baselines (see A.4 for details). In our study, cycSGHMC and R-cycSGHMC are implemented with a triple-cycle cosine learning rate schedule, while reSGHMC and r2SGHMC employ four chains per model to facilitate exhaustive exploration. It should be noted that incorporating advanced methodologies for chain swapping is essential in this task, and we have strategically adopted more sophisticated schemes for chain swap in both the reSGLD and r2SGLD frameworks. This adaptation is crucial as we engage with multiple chains instead of the conventional two. For an in-depth exploration of these enhanced chain swap mechanisms, readers are encouraged to refer to A.4. The batch size is 2,048 across all methods. We repeat experiments for each algorithm ten times to record the mean and two standard deviations for metrics such as Bayesian model averaging (BMA), negative log-likelihoods (NLL), and Brier scores (Brier).

Upon comparative evaluation of various algorithms shown in Table 1, we observe distinct performance patterns. Algorithms with reflection, particularly r2SGHMC, consistently outperform their non-reflective counterparts across all metrics. Specifically, r2SGHMC achieves the highest accuracy (BMA), NLL, and Brier score, both in ResNet20 and ResNet56 models. While standard SGDM and SGHMC show moderate performance, their reflected versions (R-SGDM and R-SGHMC) exhibit improvements, which indicates the effectiveness of reflection in model optimization. Notably, cycSGHMC and R-cycSGHMC demonstrate significant performance boosts, attributed to their cyclical learning rate schedules. Our proposed method, r2SGHMC,

characterized by its sophisticated chain-swapping mechanism, has proven to be highly effective, which makes it well-suited for complex deep-learning tasks.

5. Conclusion and Discussion

By avoiding unnatural samples and enabling efficient exploration within bounded domains, the r2SGLD algorithm marks a significant advancement in the non-convex exploration of reSGLD, particularly in high-temperature chains. This improvement is substantiated through both theoretical analysis and empirical validation.

Our theoretical investigation comprises two parts: analyzing convergence in continuous scenarios and evaluating the discretization errors of our algorithm. For continuous-time diffusion, we demonstrate its convergence in both χ^2 -divergence and $\mathcal{W}_2$ -distance under non-convex bounded domains. This analysis reveals a quantitative correlation between the constraint radius and the convergence rate, where a smaller radius facilitates accelerating the mixing rate, with rates decaying as quadratic. For discrete-time dynamics, we further provide analysis of the discretization error measured by the $\mathcal{W}_1$ -distance under convex bounded domains.

Experimentally, we establish the versatility of our algorithm in diverse applications. In the realm of dynamical systems, our method pioneers the use of constrained sampling for parameter identification. We offer robust posterior distributions to demonstrate the importance of reflection operations. In sampling from constrained multi-modal distributions, r2SGLD outperforms other methods, as evidenced both in empirical distribution representation and rapid KL divergence convergence. In non-convex optimization tasks within deep learning, r2SGLD successfully handles large learning rates, a task where naïve reSGLD struggles, due to the absence of a reflection mechanism. Despite computational limitations necessitating a reduced learning rate, our algorithm still outperforms others across various metrics. It is noteworthy that these tests were conducted within 1,000 epochs. We anticipate even better results with higher learning rates and with sufficient training epochs.

Acknowledgements

We thank the anonymous reviewers for their helpful comments. Q. Feng is partially supported by the National Science Foundation (DMS-2306769). G. Lin acknowledges the support of the National Science Foundation (DMS-2053746, DMS-2134209, ECCS-2328241, and OAC-2311848), the U.S. Department of Energy (DOE) Office of Science Advanced Scientific Computing Research (DE-SC0023161), and DOE-Fusion Energy Science (DE-SC0024583).

Impact Statement

This paper presents work whose goal is to advance the field of Machine Learning. There are many potential societal consequences of our work, none of which we feel must be specifically highlighted here.

References

- Ahn, K. and Chewi, S. Efficient Constrained Sampling via the Mirror-Langevin Algorithm. *Advances in Neural Information Processing Systems (NeurIPS)*, 34:28405–28418, 2021.
- Ahn, S., Korattikara, A., and Welling, M. Bayesian Posterior Sampling via Stochastic Gradient Fisher Scoring. In *Proc. of the International Conference on Machine Learning (ICML)*, 2012.
- Atkinson, K. *An Introduction to Numerical Analysis*. John Wiley & sons, 1991.
- Bakry, D., F. Barthe, P. C., and Guillin, A. A Simple Proof of the Poincaré Inequality for A Large Class of Probability Measures. *Electron. Comm. Probab.*, 13:60–66, 2008.
- Bakry, D., Gentil, I., and Ledoux, M. Analysis and Geometry of Markov Diffusion Operators. *Springer*, 2014.
- Berg, B. A. and Neuhaus, T. Multicanonical Ensemble: A New Approach to Simulate First-Order Phase Transitions. *Physical Review Letters*, 68(1):9, 1992.
- Brenner, P., Sweet, C. R., VonHandorf, D., and Izaguirre, J. A. Accelerating the Replica Exchange Method through an Efficient All-pairs Exchange. *The Journal of Chemical Physics*, 126:074103, 2007.
- Brosse, N., Durmus, A., Moulines, É., and Pereyra, M. Sampling from a Log-Concave Distribution with Compact Support with Proximal Langevin Monte Carlo. In *Proc. of Conference on Learning Theory (COLT)*, pp. 319–342. PMLR, 2017.
- Brunton, S. L., Proctor, J. L., and Kutz, J. N. Discovering Governing Equations from Data by Sparse Identification of Nonlinear Dynamical Systems. *Proceedings of the National Academy of Sciences*, 113(15):3932–3937, 2016.
- Bubeck, S., Eldan, R., and Lehec, J. Sampling from a Log-Concave Distribution with Projected Langevin Monte Carlo. *Discrete Comput Geom*, pp. 757–783, 2018.
- Campbell, A., Chen, W., Stimper, V., Hernandez-Lobato, J. M., and Zhang, Y. A Gradient Based Strategy for Hamiltonian Monte Carlo Hyperparameter Optimization. In *Proc. of the International Conference on Machine Learning (ICML)*, pp. 1238–1248. PMLR, 2021.
- Cattiaux, P., Guillin, A., and Wu, L.-M. A Note on Talagrand’s Transportation Inequality and Logarithmic Sobolev Inequality. *Probability Theory and Related Fields*, 148:285–334, 2010.
- Chen, C., Ding, N., and Carin, L. On the Convergence of Stochastic Gradient MCMC Algorithms with High-order Integrators. In *Advances in Neural Information Processing Systems (NeurIPS)*, pp. 2278–2286, 2015.
- Chen, Y., Chen, J., Dong, J., Peng, J., and Wang, Z. Accelerating Nonconvex Learning via Replica Exchange Langevin Diffusion. In *Proc. of the International Conference on Learning Representation (ICLR)*, 2019.
- Şimşekli, U., Badeau, R., Cemgil, A. T., and Richard, G. Stochastic Quasi-Newton Langevin Monte Carlo. In *Proc. of the International Conference on Machine Learning (ICML)*, pp. 642–651, 2016.
- de Silva, B., Champion, K., Quade, M., Loiseau, J.-C., Kutz, J., and Brunton, S. Pysindy: A Python Package for the Sparse Identification of Nonlinear Dynamical Systems from Data. *Journal of Open Source Software*, 5(49):2104, 2020.
- Deng, W., Feng, Q., Gao, L., Liang, F., and Lin, G. Non-Convex Learning via Replica Exchange Stochastic Gradient MCMC. In *Proc. of the International Conference on Machine Learning (ICML)*, 2020a.
- Deng, W., Lin, G., and Liang, F. A Contour Stochastic Gradient Langevin Dynamics Algorithm for Simulations of Multi-modal Distributions. In *Advances in Neural Information Processing Systems (NeurIPS)*, 2020b.
- Deng, W., Liang, S., Hao, B., Lin, G., and Liang, F. Interacting contour stochastic gradient langevin dynamics. In *Proc. of the International Conference on Learning Representation (ICLR)*, 2022.
- Deng, W., Zhang, Q., Feng, Q., Liang, F., and Lin, G. Non-reversible Parallel Tempering for Deep Posterior Approximation. In *Proc. of the National Conference on Artificial Intelligence (AAAI)*, volume 37, pp. 7332–7339, 2023.
- Deng, W., Chen, Y., Yang, N., Du, H., Feng, Q., and Chen, R. T. Q. Reflected Schrödinger Bridge for Constrained Generative Modeling. In *Proc. of the Conference on Uncertainty in Artificial Intelligence (UAI)*, 2024.
- Dong, J. and Tong, X. T. Spectral Gap of Replica Exchange Langevin Diffusion on Mixture Distributions. *ArXiv 2006.16193v2*, July 2020.
- Dong, J. and Tong, X. T. Spectral Gap of Replica Exchange Langevin Diffusion on Mixture Distributions. *Stochastic Processes and their Applications*, 151:451–489, 2022.

- Dupuis, P., Liu, Y., Plattner, N., and Doll, J. D. On the Infinite Swapping Limit for Parallel Tempering. *SIAM J. Multiscale Modeling & Simulation*, 10, 2012.
- Earl, D. J. and Deem, M. W. Parallel Tempering: Theory, Applications, and New Perspectives. *Phys. Chem. Chem. Phys.*, 7:3910–3916, 2005.
- Hindmarsh, A. Odepack, A Systemized Collection of Ode Solvers. *Scientific Computing*, 1983.
- Hsieh, Y.-P., Kavis, A., Rolland, P., and Cevher, V. Mirrored Langevin Dynamics. *Advances in Neural Information Processing Systems (NeurIPS)*, 31, 2018.
- Jin, T., Xu, P., Shi, J., Xiao, X., and Gu, Q. Mots: Minimax Optimal Thompson Sampling. In *International Conference on Machine Learning*, pp. 5074–5083. PMLR, 2021.
- Kang, W. and Ramanan, K. Characterization of Stationary Distributions of Reflected Diffusions. *The Annals of Applied Probability*, 24(4):1329 – 1374, 2014.
- Kaptanoglu, A. A., de Silva, B. M., Fasel, U., Kaheman, K., Goldschmidt, A. J., Callahan, J., Delahunt, C. B., Nicolaou, Z. G., Champion, K., Loiseau, J.-C., Kutz, J. N., and Brunton, S. L. Pysindy: A Comprehensive Python Package for Robust Sparse System Identification. *Journal of Open Source Software*, 7(69):3994, 2022.
- Kook, Y., Lee, Y.-T., Shen, R., and Vempala, S. Sampling with Riemannian Hamiltonian Monte Carlo in a Constrained Space. *Advances in Neural Information Processing Systems (NeurIPS)*, 35:31684–31696, 2022.
- Kufner, A. *Weighted Sobolev spaces*. Wiley, licensed ed., 1985.
- Lamperski, A. Projected Stochastic Gradient Langevin Algorithms for Constrained Sampling and Non-Convex Learning. In *Proc. of Conference on Learning Theory (COLT)*, 2021.
- Lee, H., Risteski, A., and Ge, R. Beyond Log-concavity: Provable Guarantees for Sampling Multi-modal Distributions using Simulated Tempering Langevin Monte Carlo. In *Advances in Neural Information Processing Systems (NeurIPS)*, 2018.
- Lee, Y. T. and Vempala, S. S. Convergence Rate of Riemannian Hamiltonian Monte Carlo and Faster Polytope Volume Computation. In *Proceedings of the 50th Annual ACM SIGACT Symposium on Theory of Computing*, pp. 1115–1121, 2018.
- Lelièvre, T., Rousset, M., and Stoltz, G. Hybrid Monte Carlo Methods for Sampling Probability Measures on Submanifolds. *Numerische Mathematik*, 143:379–421, 2019.
- Lelièvre, T., Stoltz, G., and Zhang, W. Multiple Projection Markov Chain Monte Carlo Algorithms on Submanifolds. *IMA Journal of Numerical Analysis*, 43(2):737–788, 2023.
- Li, X., Wu, D., Mackey, L., and Erdogdu, M. A. Stochastic Runge-Kutta Accelerates Langevin Monte Carlo and Beyond. In *Advances in Neural Information Processing Systems (NeurIPS)*, pp. 7746–7758, 2019.
- Lieberman, G. Regularized distance and its applications. *Pacific journal of Mathematics*, 117(2):329–352, 1985.
- Lingenheil, M., Denschlag, R., Mathias, G., and Tavan, P. Efficiency of Exchange Schemes in Replica Exchange. *Chemical Physics Letters*, 478:80–84, 2009.
- Liu, X., Tong, X., and Liu, Q. Sampling with Trustworthy Constraints: A Variational Gradient Framework. *Advances in Neural Information Processing Systems (NeurIPS)*, 34:23557–23568, 2021.
- Mandt, S., Hoffman, M. D., and Blei, D. M. Stochastic Gradient Descent as Approximate Bayesian Inference. *Journal of Machine Learning Research*, 18:1–35, 2017.
- Marinari, E. and Parisi, G. Simulated Tempering: A New Monte Carlo Scheme. *Europhysics Letters (EPL)*, 19(6): 451–458, 1992.
- Mazumdar, E., Pacchiano, A., Ma, Y., Jordan, M., and Bartlett, P. On Approximate Thompson Sampling with Langevin Algorithms. In *International Conference on Machine Learning*, pp. 6797–6807. PMLR, 2020.
- Menaldi, J.-L. and Robin, M. Reflected Diffusion Processes with Jumps. *The Annals of Probability*, pp. 319–341, 1985.
- Neal, R. M. MCMC using Hamiltonian Dynamics. In *Handbook of Markov Chain Monte Carlo*, volume 54, pp. 113–162. Chapman and Hall/CRC, 2012.
- Noble, M., De Bortoli, V., and Durmus, A. Unbiased Constrained Sampling with Self-Concordant Barrier Hamiltonian Monte Carlo. In *Advances in Neural Information Processing Systems (NeurIPS)*, 2023.
- Okabe, T., Kawata, M., Okamoto, Y., and Mikami, M. Replica Exchange Monte Carlo Method for the Iso-baric–isothermal Ensemble. *Chemical Physics Letters*, 335:435–439, 2001.
- Petzold, L. Automatic Selection of Methods for Solving Stiff and Nonstiff Systems of Ordinary Differential Equations. *SIAM Journal on Scientific and Statistical Computing*, 4(1):136–148, 1983. doi: 10.1137/0904010.

- Qi, R., Wei, G., Ma, B., and Nussinov, R. Replica Exchange Molecular Dynamics: A Practical Application Protocol with Solutions to Common Problems and a Peptide Aggregation and Self-Assembly Example. *Peptide Self-Assembly: Methods and Protocols*, pp. 101–119, 2018.
- Sindhikara, D. J., Emerson, D. J., and Roitberg, A. E. Exchange Often and Properly in Replica Exchange Molecular Dynamics. *Journal of Chemical Theory and Computation*, 6(9):2804–2808, 2010.
- Swendsen, R. H. and Wang, J.-S. Replica Monte Carlo Simulation of Spin-Glasses. *Physical Review Letters*, 57: 2607–2609, 1986.
- Syed, S., Bouchard-Côté, A., Deligiannidis, G., and Doucet, A. Non-reversible Parallel Tempering: A Scalable Highly Parallel Mcmc Scheme. *Journal of the Royal Statistical Society Series B: Statistical Methodology*, 84(2):321–350, 2022.
- Tanaka, H. Stochastic Differential Equations with Reflecting. *Stochastic Processes: Selected Papers of Hiroshi Tanaka*, 9:157, 1979.
- Wang, F. and Landau, D. P. Efficient, Multiple-Range Random Walk Algorithm to Calculate the Density of States. *Physical Review Letters*, 86(10):2050, 2001.
- Wang, F.-Y. Estimates of the First Neumann Eigenvalue and the Log-Sobolev Constant on Non-convex Manifolds. *Mathematische Nachrichten*, 280(12):1431–1439, 2007.
- Wang, F.-Y. *Analysis for Diffusion Processes on Riemannian Manifolds*, volume 18. World Scientific, 2014.
- Wang, J.-K. and Wibisono, A. Accelerating Hamiltonian Monte Carlo via Chebyshev Integration Time. In *Proc. of the International Conference on Learning Representation (ICLR)*, 2022.
- Wang, X., Lei, Q., and Panageas, I. Fast Convergence of Langevin Dynamics on Manifold: Geodesics Meet Log-Sobolev. *Advances in Neural Information Processing Systems (NeurIPS)*, 33:18894–18904, 2020.
- Welling, M. and Teh, Y. W. Bayesian Learning via Stochastic Gradient Langevin Dynamics. In *Proc. of the International Conference on Machine Learning (ICML)*, pp. 681–688, 2011.
- Zappa, E., Holmes-Cerfon, M., and Goodman, J. Monte Carlo on Manifolds: Sampling Densities and Integrating Functions. *Communications on Pure and Applied Mathematics*, 71(12):2609–2647, 2018.
- Zhang, R., Li, C., Zhang, J., Chen, C., and Wilson, A. G. Cyclical Stochastic Gradient MCMC for Bayesian Deep Learning. In *Proc. of the International Conference on Learning Representation (ICLR)*, 2020.
- Zhang, R., Liu, Q., and Tong, X. Sampling in Constrained Domains with Orthogonal-Space Variational Gradient Descent. *Advances in Neural Information Processing Systems (NeurIPS)*, 35:37108–37120, 2022.
- Zheng, H., Deng, W., Moya, C., and Lin, G. Accelerating Approximate Thompson Sampling with Underdamped Langevin Monte Carlo. In *Proceedings of the International Workshop on Artificial Intelligence and Statistics*, 2024.
- Zheng, W., Andrec, M., Gallicchio, E., and Levy, R. M. Simulating Replica Exchange Simulations of Protein Folding with a Kinetic Network Model. *Proceedings of the National Academy of Sciences*, 104(39):15340–15345, 2007.
- Zou, D. and Gu, Q. On the Convergence of Hamiltonian Monte Carlo with Stochastic Gradients. In *Proc. of the International Conference on Machine Learning (ICML)*, pp. 13012–13022. PMLR, 2021.

Supplementary Materials

A. Experimental Details

This section supplements the main text by detailing additional experimental procedures. Our initial focus is on the Lorenz system identification task, including a comparative analysis with four baseline methods. Subsequently, we explore the Lotka-Volterra system identification through constrained sampling, which compares both reflection-based methods and baselines without reflection operations. Furthermore, the implementation specifics of multi-mode sampling within bounded domains are discussed. This includes an examination of hyper-parameter selections and the application of these methods across various domain settings to assess the robustness of our proposed algorithm. An additional experiment is included to explore the relationship between the diameters of constrained domains and the mixing rates. Lastly, we outline hyper-parameters essential for executing large-scale image classification tasks.

A.1. The Lorenz System

For the experiments identifying the Lorenz system, we incorporate baselines such as SGLD, R-SGLD, cycSGLD, and r-cycSGLD. Across all tasks, including the main text (Section 4.1) and baselines, we conduct 60,000 iterations with the first 20,000 iterations as burn-in samples. For SGLD and R-SGLD, the learning rate commences at $5e-6$, decaying at a rate of 0.9999 per iteration after the first 10,000 iterations. For cycSGLD and R-cycSGLD, the initial learning rate is set at $2e-5$, utilizing a 10-cycle cosine learning rate schedule. The learning rates for reSGLD and r2SGLD, set at $2e-5$ and $2e-6$ for the two chains, decaying post 10,000 iterations.

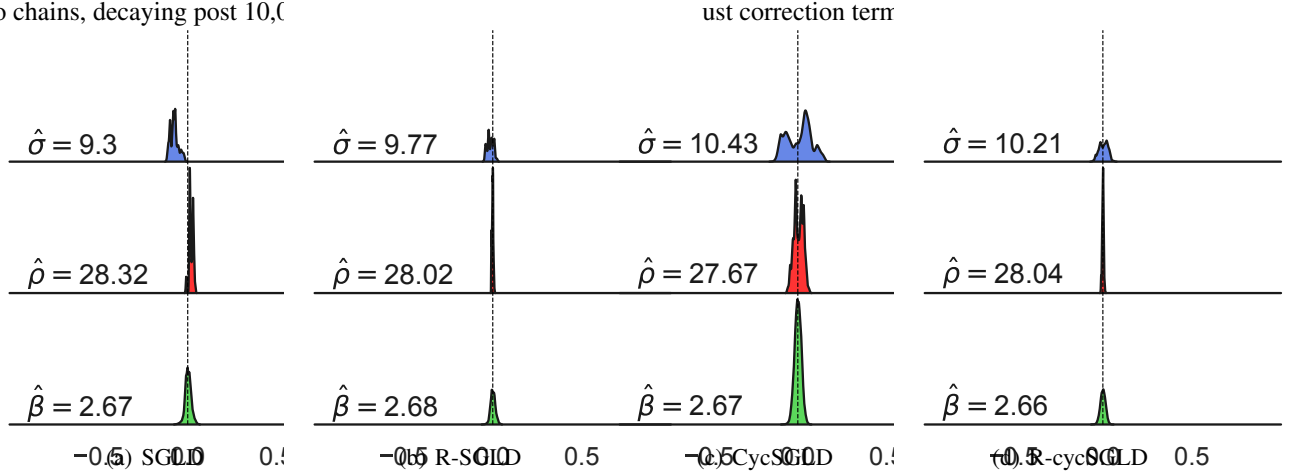


Figure 6. Posterior distributions of the identified model parameters.

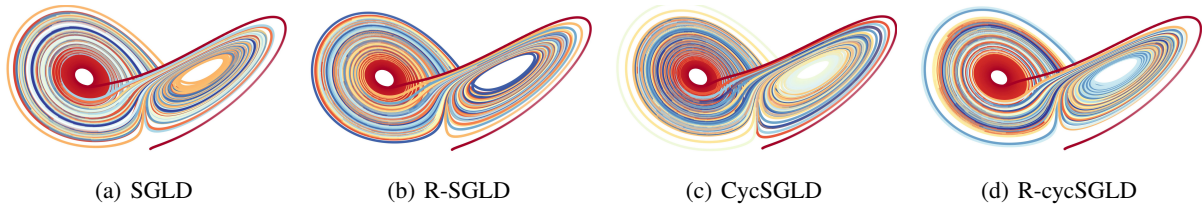


Figure 7. Simulation results of the identified Lorenz system.

As depicted in Figures 4 and 6, sampling with physical constraints and reflection significantly improves the model performance, in terms of the posterior modes and the empirical posteriors. In a comparative analysis of three reflection-based

sampling algorithms, R-SGLD falls behind in performance, while R-cycSGLD better concentrates the parameter posterior near the true parameters. Our proposed algorithm outperforms the others in accuracy. Figure 7 presents the simulation results in the use of the learned parameter according to the baselines. Although all simulations capture the chaotic nature of the Lorenz system, r2SGLD (Figure 3(c)) most closely replicates the true system dynamics (Figure 3(a)).

A.2. Lotka-Volterra Model

The Lotka-Volterra model (the predator-prey model) is a pair of first-order, non-linear, differential equations frequently used to describe the dynamics of biological systems in which two species interact, one as a predator and the other as prey. The equations are as follows:

$$\begin{cases} \dot{x} = \alpha x - \beta xy, \\ \dot{y} = -\delta y + \gamma xy. \end{cases} \quad (21)$$

Here, x and y represent the number of prey and predators, respectively, and $\dot{x}$ and $\dot{y}$ represent the corresponding time derivatives. The model is characterized by four key parameters: **Prey Growth Rate** α represents the rate of growth of the prey population in the absence of predators. It is a positive value, indicating that the prey population grows exponentially when unchallenged. **Predation Rate Coefficient** β governs the rate at which predators destroy prey[¶]. It modulates the interaction between the prey and predators, which signifies the efficiency of predators in consuming prey. **Predator Mortality Rate** δ represents the rate of decline of the predator population in the absence of prey. It signifies the mortality or decay rate of predators when there is no food (prey) available. **Predator Reproduction Rate** γ is tied to the rate at which the predator population increases due to the consumption of prey. It indicates that the growth of the predator population is dependent on the availability of prey.

The Lotka-Volterra model is a simplistic representation and makes several assumptions, such as unlimited food supply for prey and constant rates of predation, which may not always hold true in real ecosystems. However, it provides a foundational framework for understanding the dynamic interactions between predator and prey populations.

In this study, we employ the LSODA solver (Hindmarsh, 1983; Petzold, 1983) to generate data $x(t)$ and $y(t)$, where t ranges from 0 to 100 in increments of 0.01. The framework of identifying the dynamical system aligns with the one established in Section 4.1 on Lorenz system identification, including the sample generation and the simulation of empirical posterior distributions. For the ease of implementation, the corresponding matrices for gradient computations are presented as follows:

$$\dot{\mathbf{X}} = \begin{bmatrix} \dot{x}(t_1) & \dot{y}(t_1) \\ \dot{x}(t_2) & \dot{y}(t_2) \\ \vdots & \vdots \\ \dot{x}(t_m) & \dot{y}(t_m) \\ \vdots & \vdots \end{bmatrix}, \quad \hat{\boldsymbol{\beta}} = \begin{bmatrix} \alpha & 0 \\ 0 & -\delta \\ -\beta & \gamma \end{bmatrix}, \quad \boldsymbol{\Theta}(\mathbf{X}) = \begin{bmatrix} x(t_1) & y(t_1) & x(t_1)y(t_1) \\ x(t_2) & y(t_2) & x(t_2)y(t_2) \\ \vdots & \vdots & \vdots \\ x(t_m) & y(t_m) & x(t_m)y(t_m) \\ \vdots & \vdots & \vdots \end{bmatrix}, \quad (22)$$

where we establish parameter values as follows: $\alpha = 1.0$, $\beta = 0.10$, $\delta = 1.50$, and $\gamma = 0.075$. The chosen batch size is 4,096, and the sampling process involves 60,000 iterations, with the initial 20,000 serving as burn-in samples. For SGLD and R-SGLD, the settings are the same as the ones in identifying the Lorenz system. In the case of cycSGLD and R-cycSGLD, the initial learning rate is 1e-5, implemented with a ten-cycle cosine learning rate schedule. For reSGLD and r2SGLD, learning rates are 1e-5 and 2e-6 respectively for the two chains, following a similar decay pattern post 10,000 iterations at 0.9999.

In Figure 8, we observe the empirical posteriors of model parameters and their posterior modes. These results further confirm the beneficial impact of the reflection operation on parameter estimation. A comparative analysis of different algorithms for constrained sampling tasks reveals that R-SGLD (the blue lines) is observed to be the least effective, whereas R-cycSGLD (the green lines) and r2SGLD (the red lines) demonstrate comparable performance. In Figure 9, the parameters derived from these algorithms are employed to simulate the dynamics of the Lotka-Volterra system over a period from the 100th to the 200th day, as the first 100 days produced overlapping results. The simulations reveal that both R-cycSGLD and r2SGLD provide comparable outcomes, yet r2SGLD shows a marginally closer approximation to the true dynamics (the black lines).

[¶]Following traditional conventions in the Lotka-Volterra model, we use β to differentiate it from the parameters β as defined in (1).

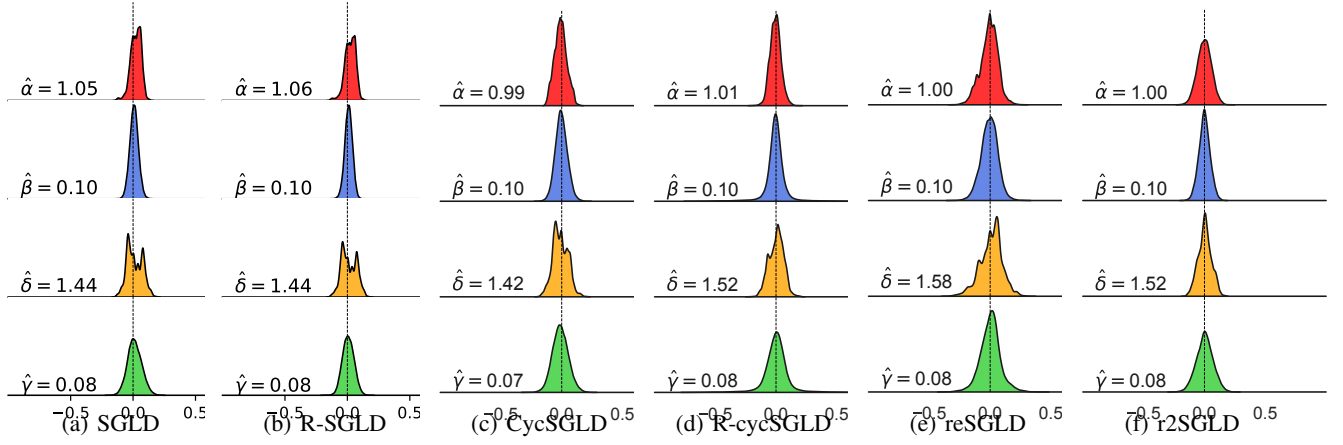
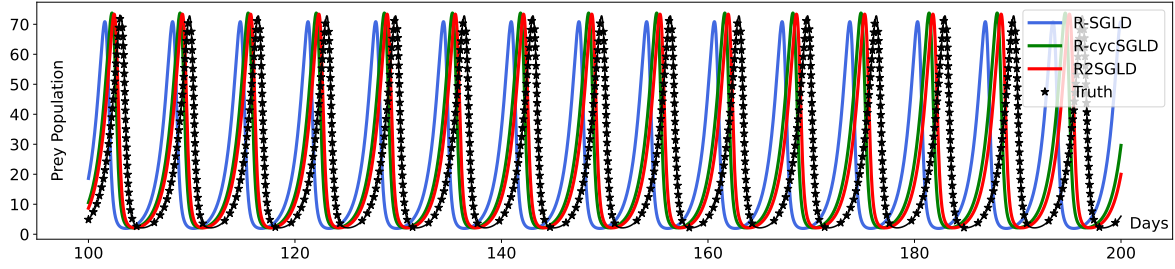
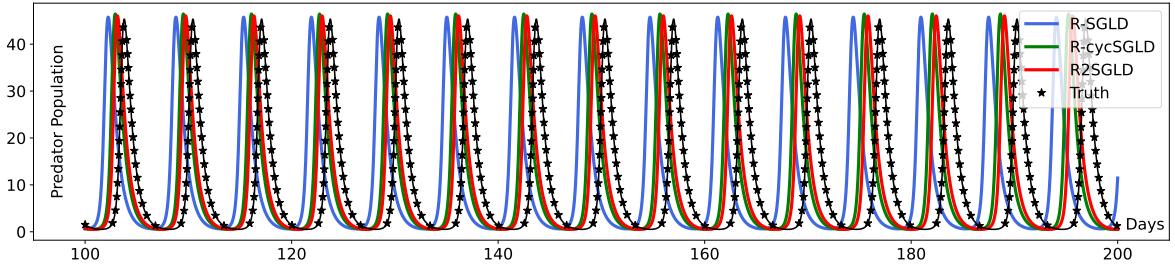


Figure 8. Posterior distribution of model parameters in the Lotka-Volterra system.



(a) Prey simulations.



(b) Predator simulations.

Figure 9. Simulation of the identified Lotka-Volterra system from 100 to 200 days.

A.3. Constrained Sampling from Multi-mode Distributions

This section focuses on constrained sampling from a multi-mode distribution. We study a 2D mixture of 25 Gaussians

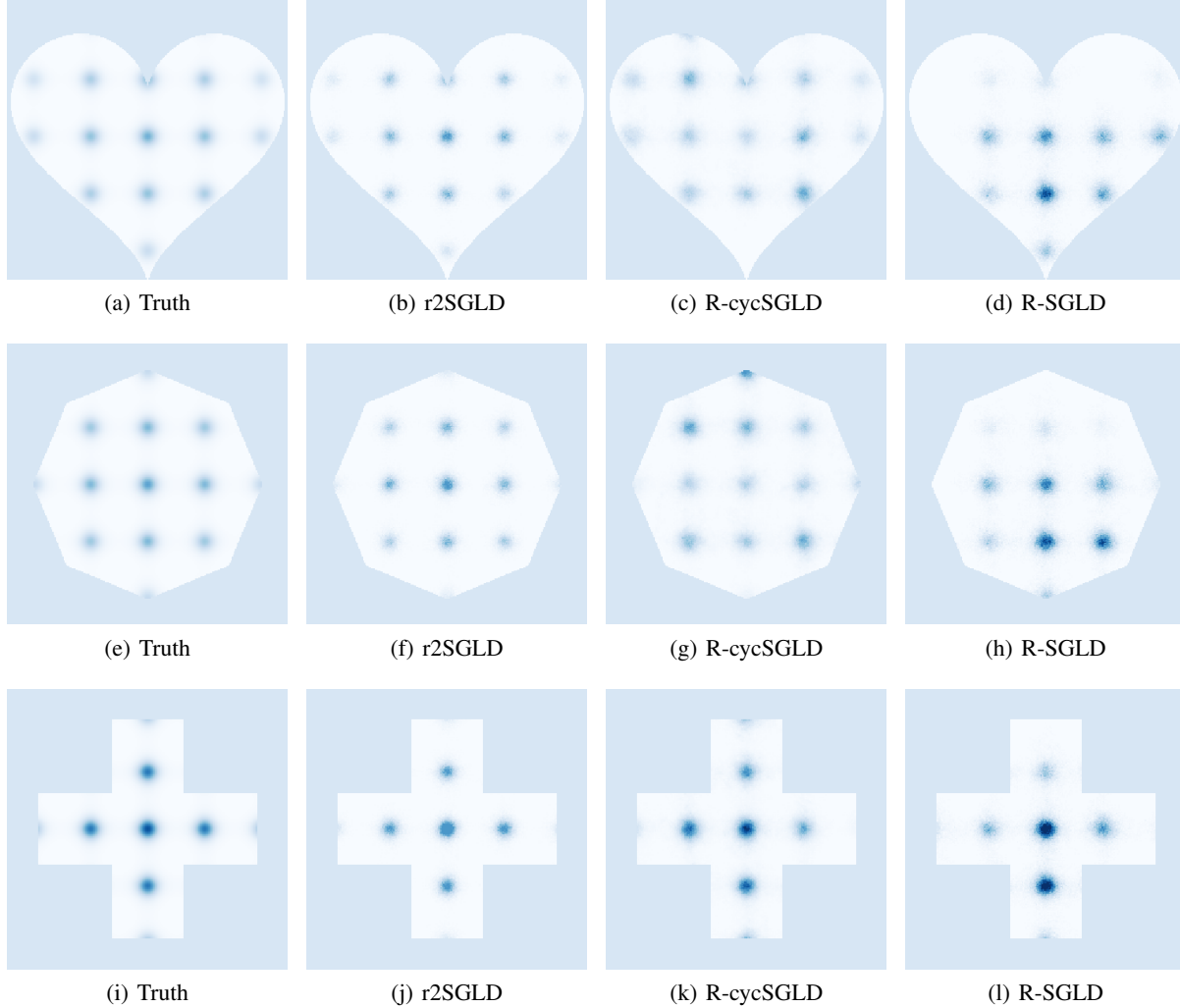


Figure 10. Empirical behavior on multi-mode distributions with diverse boundaries.

within uniquely shaped domains: flower (Figure 5), heart (Figure 10(a)), polygon (Figure 10(e)), and cross (Figure 10(i)). The heart domain in this simulation is defined by the following parametric equations:

$$\begin{aligned} x &= 16 \sin^3(2\pi t), \\ y &= 13 \cos(2\pi t) - 5 \cos(4\pi t) - 2 \cos(6\pi t) - \cos(8\pi t). \end{aligned}$$

Subsequently, the octagon domain is described by the equations:

$$\begin{aligned} r &= ct - \lfloor ct \rfloor, \\ x &= (1 - r) \cdot X_{\lfloor ct \rfloor} + r \cdot X_{\lfloor ct \rfloor + 1}, \\ y &= (1 - r) \cdot Y_{\lfloor ct \rfloor} + r \cdot Y_{\lfloor ct \rfloor + 1}, \end{aligned}$$

where $(X_i, Y_i)_{i=0}^C$ are the edge sequences with $(X_C, Y_C) = (X_0, Y_0)$ and $C = 8$.

The probability densities are set to zero outside these shapes. For this study, we employ R-SGLD (initial rate: $5e-4$), R-cycSGLD (initial rate: $1e-3$ with a 5-cycle cosine schedule), and r2SGLD (two chains at $5e-4$ and $1.5e-3$, aiming for

a 5-20% swap rate). A total of 50,000 samples are generated, with the initial 10,000 as burn-in. The obtained empirical posteriors are shown in Figures 5 and 10.

Upon examining the empirical posteriors, we note the limitations of R-SGLD, where the results often get trapped in local modes and fail to fully capture the target distributions. Conversely, R-cycSGLD demonstrates a more robust ability to capture most modes, albeit with a tendency to overstay in specific modes and underrepresent others, which more samples could improve. In comparison, our proposed algorithm shows significant improvements in accurately characterizing all modes, which provides the most precise empirical behavior in multi-mode distributions across different constrained domains.

Exploration of the Constrained Domain Diameter To empirically validate our theoretical analysis, we conduct additional experiments with r2SGLD, which focuses on how the diameters of the constrained domain affect mixing rates. The target distributions are selected as convex bounded domains (octagons) against a non-convex distribution of 25 Gaussian modes (refer to Figure 10(e)). The domain diameters are set between 1.5 and 3.0 (the target distribution is 5.0×5.0). We employ a dual-chain r2SGLD with learning rates of $2e-5$ and $2e-4$. For each diameter setting, we execute ten runs, where each generates 10,000 samples. The expected KL divergences and their 95% confidence intervals are depicted in Figure 11 for purposes of analysis.

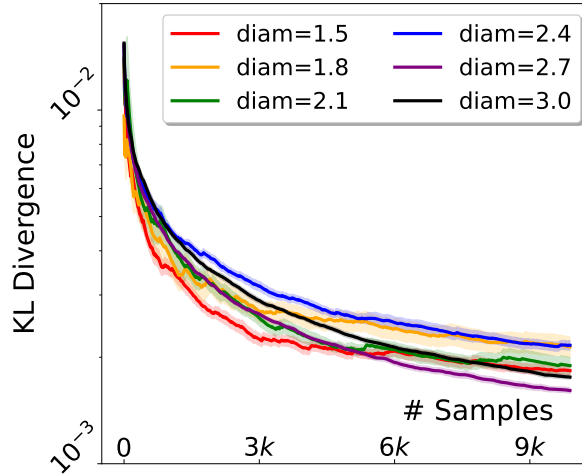


Figure 11. KL divergence with respect to different domain diameter setts.

Illustrated by the corresponding figure, a discernible correlation emerges between the domain diameter and the convergence rate of the algorithm. Notably, the result demonstrates a trend where a reduction in domain diameter results in a more rapid decline in KL divergence. This further leads to an advanced mixing rate. This observation suggests that a smaller domain diameter facilitates the algorithm’s expedited convergence to the target distribution, which also implies that constraining the parameter space informed by prior knowledge can markedly enhance the algorithm’s efficiency in its convergence.

A.4. Image Classification

The following part is dedicated to an empirical evaluation of algorithm settings. Our objective is to the choices of hyper-parameters and demonstrate their effectiveness in the tests.

Learning Rate This study acknowledges that while reflection operations in sampling methods can lead to more comprehensive exploration given unlimited computational resources, our experiments are conducted within the confines of 1,000 epochs. A critical component of these experiments is the determination of suitable initial learning rates to enhance model performance.

In our study, we conduct a systematic exploration of learning rates for the single-chain R-SGHMC algorithm. We initiate this exploration by setting the learning rate at a standard value of 0.1 and incrementally increasing it to identify the point of optimal model performance. The effectiveness of various learning rates is assessed using two key metrics: BMA and NLL. These metrics, depicted in Figure 12, guide our determination of suitable learning rates. Building upon these findings, we

then adjust the learning rates of other baseline models to align with the identified optimal learning rate. Our results indicate that an initial learning rate of approximately 2.0 is effective for both the baseline and the proposed algorithms.

For consistency in our methodology, this initial learning rate (2.0) is applied to SGDM, R-SGDM, SGHMC, R-SGHMC, cycSGHMC, and R-cycSGHMC. For SGDM, R-SGDM, SGHMC, and R-SGHMC, this learning rate is consistently held during the initial 300 epochs to facilitate exploration and subsequently decayed at a rate of 0.990 each epoch. In the case of cycSGHMC and R-cycSGHMC, we opt for a triple-cycle cosine learning rate schedule. For reSGHMC and r2SGHMC, which utilize four chains, the chain with the lowest learning rate begins at 1.0, scaling up to approximately 4.0 for the chain with the highest rate.

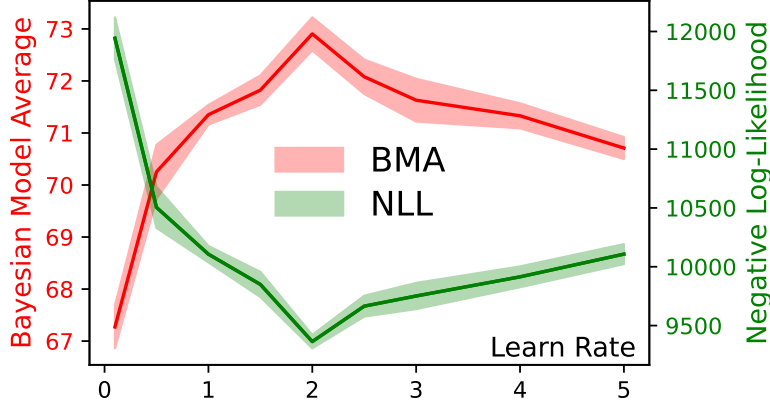


Figure 12. Metric comparisons of CIFAR 100 with different initial learning rates.

Swap Efficiency In the realm of reSGLD and r2SGLD algorithms, the implementation of effective swapping schemes is crucial to enhance exploration efficiency. A fundamental challenge arises due to the minimal overlap between the distributions of the extremal temperatures, $\pi^{(1)}$ and $\pi^{(P)}$ (P denotes the last chain or the number of chains), which hinders acceleration despite a large $\tau^{(P)}$. A viable solution involves introducing intermediate temperature particles ($\tau^{(2)}, \dots, \tau^{(P-1)}$) to facilitate “tunnels” for more efficient chain swapping. Common strategies include all-pairs exchange (APE), adjacent (ADJ) pair swaps, and deterministic even-odd (DEO) schemes.

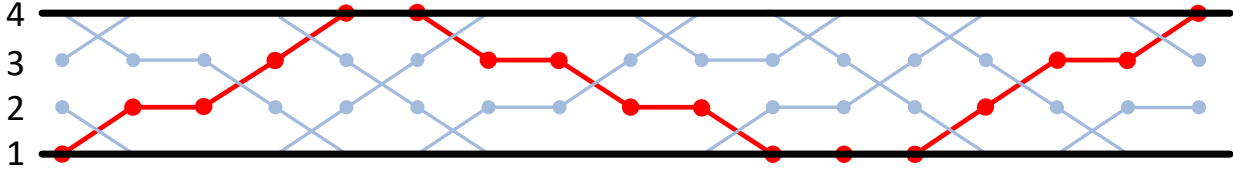


Figure 13. Non-reversible chain swapping with the ODE scheme. The red line demonstrates a single chain’s round-trip path.

The APE scheme attempts swaps between arbitrary chain pairs, offering thorough exploration (Brenner et al., 2007; Lingenheil et al., 2009). However, its cubic swap time complexity ($O(P^3)$) and practical usability constraints limit its widespread adoption. The ADJ scheme, alternatively, swaps adjacent pairs in sequence, from $(1, 2)$ to $(P - 1, P)$ (Qi et al., 2018). While this method is simpler and more intuitive, its sequential nature results in exchange information dependency and is less effective in multi-core or distributed environments, particularly with a larger number of chains. To address these limitations, the deterministic even-odd (DEO) scheme was introduced (Okabe et al., 2001; Syed et al., 2022; Deng et al., 2023), which alternates swaps between even and odd pairs in successive iterations. The DEO scheme, characterized by its non-reversible and asymmetric nature, significantly increases swap rates and reduces round trip times, with its design detailed in Figure 13.

In this research, we adopt the DEO scheme (as detailed in Algorithm 2) for reSGLD and r2SGLD algorithms to optimize path efficiency in the deep learning task. The algorithm is structured to achieve a near-linear path and anticipate three to four round trips within a span of 1,000 epochs for increased swapping efficiency.

Algorithm 2 Reflected Replica Exchange Stochastic Gradient Langevin Dynamics with the DEO scheme.

Input Window size $\mathcal{W} := \left\lceil \frac{\log P + \log \log P}{-\log(1-\mathbb{S})} \right\rceil$.

Input Number of iterations K .

Input Number of chains P .

Input Target swap rate $\mathbb{S}$.

Input Step size γ .

for $k = 1, 2, \dots, K$ **do**

 Sampling with reflection $\tilde{\beta}_{k+1} = \mathcal{P}\left(\tilde{\beta}_k - \eta \nabla \tilde{U}(\tilde{\beta}_k) + \sqrt{2\eta\tau_2}\xi_k\right)$
for $p = 1, 2, \dots, P-1$ **do**

 Update the swap indicator: $\mathcal{S}^{(p)} := 1_{\tilde{U}(\beta_{k+1}^{(p+1)}) + \mathcal{C}_k < \tilde{U}(\beta_{k+1}^{(p)})}$
if $k \bmod \mathcal{W} = 0$ **then**

 Allow swaps $\mathcal{G}^{(p)} := 1$
end if
if $\mathcal{G}^{(p)}$ and $\mathcal{S}^{(p)}$ **then**

 Chain swap $\beta_{k+1}^{(p)}$ and $\beta_{k+1}^{(p+1)}$

 Refuse swaps $\mathcal{G}^{(p)} := 0$
end if
end for

 Update correction $\mathcal{C}_{k+1} = \mathcal{C}_k + \gamma \left(\frac{1}{P-1} \sum_{p=1}^{P-1} 1_{\tilde{U}(\beta_{k+1}^{(p+1)}) + \mathcal{C}_k - \tilde{U}(\beta_{k+1}^{(p)}) < 0} - \mathbb{S} \right)$.

end for
Output Parameters $\{\beta_k^{(1)}\}_{k=1}^K$.

The primary difference between Algorithm 2 and Algorithm 1 is the consideration of a multiple-chain variant, which innovates chain swapping management. This is achieved through a carefully designed swap indicator $\mathcal{S}^{(p)}$ that dynamically adjusts based on the relative positions of the chains, thereby ensuring more effective and frequent exchanges. The algorithm incorporates a window-based approach, where swaps are allowed only at specific intervals, determined by the window size $\mathcal{W}$. This strategic regulation of swaps is further enhanced by combining with a correction term $\mathcal{C}_k$, which is updated adaptively:

$$\mathcal{C}_{k+1} = \mathcal{C}_k + \gamma \left(\frac{1}{P-1} \sum_{p=1}^{P-1} 1_{\tilde{U}(\beta_{k+1}^{(p+1)}) + \mathcal{C}_k - \tilde{U}(\beta_{k+1}^{(p)}) < 0} - \mathbb{S} \right).$$

Here, γ denotes the step size, $P = 4$ represents the number of chains involved in the process, and $\mathbb{S}$ is the target swap rate among chains. It should be noted that as the iteration index k trends toward infinity, it is anticipated that both the threshold and adaptive learning rates will converge toward their respective fixed points. This implementation of a uniform $\mathcal{C}$ ensures that swaps occur with a frequency that optimizes the overall convergence towards the desired target distribution. This balance between exploration and exploitation in the r2SGLD algorithm is crucial for tasks involving complex datasets like CIFAR 100, where traditional methods might struggle with efficient parameter space navigation.

B. Continuous Convergence Analysis

Proof of Lemma 3.2 First we show the reversibility of π under the semigroup $\{e^{t\mathcal{L}}\}_{t \geq 0}$:

$$\int_{\Omega \times \Omega} g \mathcal{L} f \, d\pi(x_1, x_2) = \int_{\Omega \times \Omega} f \mathcal{L} g \, d\pi(x_1, x_2), \quad \text{for all } f, g \in \mathcal{D}(\mathcal{L}). \quad (23)$$

Recall in (4) that $\mathcal{L} = \mathcal{L}^{(1)} + \mathcal{L}^{(2)} + \mathcal{L}^{(s)}$. We have

$$\begin{aligned} & \int_{\Omega} \int_{\Omega} g \mathcal{L}^{(1)} f \, d\pi(x_1, x_2) \\ &= \frac{1}{Z} \int_{\Omega} \int_{\Omega} g(x_1, x_2) [-\langle \nabla_{x_1} f(x_1, x_2), \nabla_{x_1} U(x_1) \rangle + \tau_1 \Delta_{x_1} f(x_1, x_2)] e^{-\frac{U(x_1)}{\tau_1}} dx_1 e^{-\frac{U(x_2)}{\tau_2}} dx_2 \\ &= \frac{1}{Z} \int_{\Omega} \int_{\Omega} g(x_1, x_2) \nabla \cdot \left(\tau_1 \nabla_{x_1} f(x_1, x_2) e^{-\frac{U(x_1)}{\tau_1}} \right) dx_1 e^{-\frac{U(x_2)}{\tau_2}} dx_2 \\ &= - \int_{\Omega} \int_{\Omega} \tau_1 \nabla_{x_1} g(x_1, x_2) \cdot \nabla_{x_1} f(x_1, x_2) d\pi(x_1, x_2) \quad (\text{Integration by parts with } f \text{ satisfying (5)}_1) \\ &= \int_{\Omega} \int_{\Omega} f \mathcal{L}^{(1)} g \, d\pi(x_1, x_2) \end{aligned} \quad (24)$$

A similar calculation with f satisfying (5)₂ yields that

$$\int_{\Omega} \int_{\Omega} g \mathcal{L}^{(2)} f \, d\pi(x_1, x_2) = \int_{\Omega} \int_{\Omega} f \mathcal{L}^{(2)} g \, d\pi(x_1, x_2) = - \int_{\Omega} \int_{\Omega} \tau_2 \nabla_{x_2} f(x_1, x_2) \cdot \nabla_{x_2} g(x_1, x_2) d\pi(x_1, x_2). \quad (25)$$

By the same argument as in [Chen et al. \(2019\)](#) (Lemma 3.2), we can also derive

$$\int_{\Omega} \int_{\Omega} g \mathcal{L}^{(s)} f \, d\pi(x_1, x_2) = \int_{\Omega} \int_{\Omega} f \mathcal{L}^{(s)} g \, d\pi(x_1, x_2). \quad (26)$$

Combining (24), (25) and (26), we obtain (23). Now by choosing $g = 1 \in \mathcal{D}(\mathcal{L})$ in (23), since $\mathcal{L}1 = 0$, we obtain

$$\int_{\Omega \times \Omega} \mathcal{L} f \, d\pi = \int_{\Omega \times \Omega} f \mathcal{L} 1 \, d\pi = 0.$$

Thus π is invariant under $\{e^{t\mathcal{L}}\}_{t \geq 0}$. □

Next, we introduce a Lyapunov function subject to the homogeneous Neumann boundary condition, this is an extension of the Lyapunov function introduced in [Cattiaux et al., 2010](#).

Lemma B.1. *Let $\mathcal{L}^{(0)} = \mathcal{L}^{(1)} + \mathcal{L}^{(2)}$, then there exists a Lyapunov function $V(x_1, x_2) \geq 1$ such that satisfies*

$$\begin{cases} \mathcal{L}^{(0)} V(x_1, x_2) \leq -\lambda V(x_1, x_2) + b, & (x_1, x_2) \in \Omega \times \Omega, \\ \nabla_{x_1} V(x_1, x_2) \cdot \nu(x_1) = 0, & (x_1, x_2) \in \partial\Omega \times \Omega, \\ \nabla_{x_2} V(x_1, x_2) \cdot \nu(x_2) = 0, & (x_1, x_2) \in \Omega \times \partial\Omega. \end{cases} \quad (27)$$

for some constants $\lambda > 0, b \geq 0$.

Proof Let $d(x) := \text{dist}(x, \partial\Omega)$ be the distance function between x and the boundary $\partial\Omega$. There exists an infinitely differentiable *regularized distance function* which is equivalent to d ([Kufner, 1985](#); [Lieberman, 1985](#)). More precisely, there exists a positive function $\rho(x) \in C^\infty(\Omega)$ and positive constants $C_k, k = 0, 1, 2, \dots$, such that for all $x \in \Omega$,

$$\frac{1}{C_0} d(x) \leq \rho(x) \leq C_0 d(x),$$

and for $k \geq 1$,

$$|D^k \rho(x)| \leq C_k \rho(x)^{1-k}.$$

Now we defined the Lyapunov function V by ρ :

$$V(x_1, x_2) := \exp[\rho(x_1)^2 + \rho(x_2)^2] \geq 1.$$

We can compute the Hessian $\nabla_{x_1 x_1}^2 V(x_1, x_2)$:

$$\begin{aligned} \nabla_{x_1 x_1}^2 V(x_1, x_2) &= 2\nabla_{x_1} \rho(x_1) \otimes \nabla_{x_1} \rho(x_1) V(x_1, x_2) \\ &\quad + 2\rho(x_1) \nabla_{x_1 x_1}^2 \rho(x_1) V(x_1, x_2) + 4\rho(x_1)^2 \nabla_{x_1} \rho(x_1) \otimes \nabla_{x_1} \rho(x_1) V(x_1, x_2), \end{aligned}$$

where $\otimes$ is the tensor product. Using the estimates for ρ , we obtain the following bound:

$$|\nabla_{x_1 x_1}^2 V(x_1, x_2)| \leq [2C_1^2 + 2C_2 + 4C_0^2 \text{diam}(\Omega)^2] \exp[2C_0^2 \text{diam}(\Omega)^2].$$

Given any $\lambda > 0$, we can take $b = \sup\{\mathcal{L}^{(0)}V(x_1, x_2) + \lambda V(x_1, x_2) : (x_1, x_2) \in \bar{\Omega} \times \bar{\Omega}\} \vee 0 < \infty$. Then (27)₁ is satisfied. Furthermore, we notice that since

$$|\nabla_{x_1} V(x_1, x_2)| \leq 2\rho(x_1) |\nabla_{x_1} \rho(x_1)| V(x_1, x_2) \leq 2C_0 d(x_1) \cdot C_1 \cdot \exp[2C_0^2 \text{diam}(\Omega)^2] \rightarrow 0 \text{ as } x_1 \rightarrow \partial\Omega,$$

thus the homogeneous Neumann boundary condition (27)₂ with respect to x_1 is also true for V . Similarly, we can get (27)₃ also holds. $\square$

Proof of Lemma 3.3 Next, we show the Poincaré inequality using the Lyapunov function V in Lemma B.1. Suppose $f \in C^1(\Omega \times \Omega)$, and u is an arbitrary constant. Multiplying both sides of (27)₁ with $(f - u)^2 / \lambda V$ and integrating over $\Omega \times \Omega$ against $\pi(x_1, x_2)$ gives

$$\int_{\Omega} \int_{\Omega} (f - u)^2 d\pi(x_1, x_2) \leq \int_{\Omega} \int_{\Omega} -\frac{\mathcal{L}^{(0)}V}{\lambda V} (f - u)^2 d\pi(x_1, x_2) + b \int_{\Omega} \int_{\Omega} \frac{(f - u)^2}{\lambda V} d\pi(x_1, x_2). \quad (28)$$

Notice that

$$\begin{aligned} &\int_{\Omega} \int_{\Omega} -\frac{\mathcal{L}^{(1)}V}{\lambda V} (f - u)^2 d\pi(x_1, x_2) \\ &= \int_{\Omega} \int_{\Omega} -\frac{(f - u)^2}{\lambda V} \nabla_{x_1} \cdot \left(\tau_1 \nabla_{x_1} V e^{-\frac{U(x_1)}{\tau_1}} \right) dx_1 e^{-\frac{U(x_2)}{\tau_2}} dx_2 \\ &= \int_{\Omega} \int_{\Omega} \tau_1 \nabla_{x_1} \left[\frac{(f - u)^2}{\lambda V} \right] \cdot \nabla_{x_1} V e^{-\frac{U(x_1)}{\tau_1}} dx_1 e^{-\frac{U(x_2)}{\tau_2}} dx_2 \quad (\text{Integration by parts with } V \text{ satisfying (27)}_2) \\ &= \tau_1 \int_{\Omega} \int_{\Omega} \nabla_{x_1} \left[\frac{(f - u)^2}{\lambda V} \right] \cdot \nabla_{x_1} V d\pi(x_1, x_2) \\ &= \tau_1 \int_{\Omega} \int_{\Omega} \frac{2(f - u)}{\lambda V} \nabla_{x_1} (f - u) \cdot \nabla_{x_1} V - (f - u)^2 \frac{|\nabla_{x_1} V|^2}{\lambda V^2} d\pi(x_1, x_2) \\ &= \frac{\tau_1}{\lambda} \int_{\Omega} \int_{\Omega} \left(|\nabla_{x_1} (f - u)|^2 - \left| \nabla_{x_1} (f - u) - \frac{f - u}{V} \nabla_{x_1} V \right|^2 \right) d\pi(x_1, x_2) \\ &\leq \frac{\tau_1}{\lambda} \int_{\Omega} \int_{\Omega} |\nabla_{x_1} (f - u)|^2 d\pi(x_1, x_2). \end{aligned}$$

Similar computation works for $\mathcal{L}^{(2)}$, we reach that

$$\int_{\Omega} \int_{\Omega} -\frac{\mathcal{L}^{(0)}V}{\lambda V} (f - u)^2 d\pi(x_1, x_2) \leq \frac{1}{\lambda} \int_{\Omega} \int_{\Omega} (\tau_1 |\nabla_{x_1} (f - u)|^2 + \tau_2 |\nabla_{x_2} (f - u)|^2) d\pi(x_1, x_2). \quad (29)$$

On the other hand, we choose u^* such that $\int_{\Omega} \int_{\Omega} (f - u^*) d\pi(x_1, x_2) = 0$, then

$$b \int_{\Omega} \int_{\Omega} \frac{(f - u^*)^2}{\lambda V} d\pi(x_1, x_2) = \frac{b}{\lambda} \int_{\Omega} \int_{\Omega} (f - u^*)^2 d\pi(x_1, x_2) \leq \frac{b}{\lambda} c_P \int_{\Omega} \int_{\Omega} |\nabla f|^2 d\pi(x_1, x_2). \quad (30)$$

Combining (28), (29) and (30) we obtain the Poincaré inequality (10). $\square$

Now we have all the ingredients to show the accelerated exponential convergence of μ_t .

Proof of Theorem 3.4 Using the generator $\mathcal{L}$, we have $d\mu_t/d\pi = e^{t\mathcal{L}}(d\mu_0/d\pi)$. In other words, given that $d\mu_0/d\pi$ satisfies (5), $f_t := d\mu_t/d\pi$ satisfies $\partial_t f_t = \mathcal{L}f_t$ and (5) for all $t > 0$. Then using the integration by parts we can show

$$\begin{aligned} \frac{d}{dt}\chi^2(\mu_t|\pi) &= \frac{d}{dt} \int_{\Omega \times \Omega} \left[e^{t\mathcal{L}} \left(\frac{d\mu_0}{d\pi} \right) \right]^2 d\pi = 2 \int_{\Omega \times \Omega} e^{t\mathcal{L}} \left(\frac{d\mu_0}{d\pi} \right) \frac{d}{dt} \left[e^{t\mathcal{L}} \left(\frac{d\mu_0}{d\pi} \right) \right] d\pi \\ &= 2 \int_{\Omega \times \Omega} e^{t\mathcal{L}} \left(\frac{d\mu_0}{d\pi} \right) \mathcal{L} \left[e^{t\mathcal{L}} \left(\frac{d\mu_0}{d\pi} \right) \right] d\pi \\ &= -2\mathcal{E}_S \left(e^{t\mathcal{L}} \left(\frac{d\mu_0}{d\pi} \right) \right) = -2\mathcal{E}_S \left(\frac{d\mu_t}{d\pi} \right). \end{aligned} \quad (31)$$

By Lemma 3.3, we have following estimate

$$\chi^2(\mu_t|\pi) \leq C_P \mathcal{E} \left(\frac{d\mu_t}{d\pi} \right) \leq C_P(1 + \eta_S)^{-1} \mathcal{E}_S \left(\frac{d\mu_t}{d\pi} \right). \quad (32)$$

Combining (31) and (32) yields

$$\frac{d}{dt}\chi^2(\mu_t|\pi) \leq -2C_P^{-1}(1 + \eta_S)\chi^2(\mu_t|\pi).$$

According to Gronwall's inequality, we conclude that

$$\chi^2(\mu_t|\pi) \leq \chi^2(\mu_0|\pi) \exp\{-2(C_P^{-1}(1 + \eta_S))t\}.$$

□

Proof of Theorem 3.8 Given that $d\mu_0/d\pi \geq 0$ and satisfying (5), we have that f_t satisfies $\partial_t f_t = \mathcal{L}f_t$, $f_t \geq 0$ and (5) for all $t > 0$. Then again we use the integration by parts to derive

$$\begin{aligned} \frac{d}{dt}D(\mu_t|\pi) &= \frac{d}{dt} \int_{\Omega \times \Omega} e^{t\mathcal{L}} \left(\frac{d\mu_0}{d\pi} \right) \ln \left[e^{t\mathcal{L}} \left(\frac{d\mu_0}{d\pi} \right) \right] d\pi \\ &= \int_{\Omega \times \Omega} \left[1 + \ln \left(e^{t\mathcal{L}} \left(\frac{d\mu_0}{d\pi} \right) \right) \right] \mathcal{L} \left[e^{t\mathcal{L}} \left(\frac{d\mu_0}{d\pi} \right) \right] d\pi \\ &= \int_{\Omega \times \Omega} \ln \left(\frac{d\mu_t}{d\pi} \right) \mathcal{L} \left(\frac{d\mu_t}{d\pi} \right) d\pi \leq -2\mathcal{E}_S \left(\sqrt{\frac{d\mu_t}{d\pi}} \right). \end{aligned} \quad (33)$$

Now we apply the Log-Sobolev inequality (14) to get

$$D(\mu_t|\pi) \leq C_{LS} \mathcal{E} \left(\sqrt{\frac{d\mu_t}{d\pi}} \right) \leq C_{LS}(1 + \delta_S)^{-1} \mathcal{E}_S \left(\sqrt{\frac{d\mu_t}{d\pi}} \right). \quad (34)$$

Combining (33) and (34) together, we arrive at

$$\frac{d}{dt}D(\mu_t|\pi) \leq -2C_{LS}^{-1}(1 + \delta_S)D(\mu_t|\pi).$$

Then we apply the Gronwall inequality to get

$$D(\mu_t|\pi) \leq D(\mu_0|\pi) \exp\{-2(C_{LS}^{-1}(1 + \delta_S))t\}.$$

By the Otto–Villani Theorem (Bakry et al., 2014) (Theorem 9.6.1), we have the following quadratic transportation cost inequality

$$\mathcal{W}_2(\mu_t, \pi) \leq \sqrt{2C_{LS}D(\mu_t|\pi)} \leq \sqrt{2C_{LS}D(\mu_0|\pi)} \exp\{-tC_{LS}^{-1}(1 + \delta_S)\}$$

□

C. Discretization Analysis

To better formulate the discretized error, we rewrite the **Reflected Replica exchange Langevin diffusion** as below. For any fixed learning rate $\eta > 0$, we define

$$\begin{cases} d\beta_t = -\nabla G(\beta_t)dt + \Sigma(\alpha_t)d\mathbf{W}_t + \nu_t L(dt), \\ \mathbb{P}(\alpha(t) = j | \alpha(t-dt) = l, \beta(\lfloor t/\eta \rfloor \eta) = \beta) = rS(\beta)\eta \mathbf{1}_{\{t=\lfloor t/\eta \rfloor \eta\}} + o(dt), \text{ for } l \neq j, \end{cases} \quad (35)$$

where we denote $\{\beta_t\}_{t \geq 0} := \left\{ \begin{pmatrix} \beta_t^{(1)} \\ \beta_t^{(2)} \end{pmatrix} \right\}_{t \geq 0}$, $\nabla G(\beta) := \begin{pmatrix} \nabla U(\beta^{(1)}) \\ \nabla U(\beta^{(2)}) \end{pmatrix}$, and $\nu_t L(dt) = \begin{pmatrix} \nu_t^{(1)} & 0 \\ 0 & \nu_t^{(2)} \end{pmatrix} \begin{pmatrix} L^{(1)}(dt) \\ L^{(2)}(dt) \end{pmatrix}$,

where the local time $\nu_t L(dt)$ ensures the process β_t stays in the domain $\Omega \times \Omega$. We denote $\mathbf{1}_{t=\lfloor t/\eta \rfloor \eta}$ as the indicator function, i.e. for every $t = i\eta$ with $i \in \mathbb{N}^+$, given $\beta(i\eta) = \beta$, we have $\mathbb{P}(\alpha(t) = j | \alpha(t-dt) = l) = rS(\beta)\eta$, where $S(\beta)$ is defined as $\min\{1, S(\beta^{(1)}, \beta^{(2)})\}$ and $S(\beta^{(1)}, \beta^{(2)})$ is defined in (3). In this case, the Markov Chain $\alpha(t)$ is a constant on the time interval $[\lfloor t/\eta \rfloor \eta, \lfloor t/\eta \rfloor \eta + \eta)$ with some state in the finite-state space $\{0, 1\}$ and the generator matrix Q follows

$$Q = \begin{pmatrix} -rS(\beta)\eta\delta(t - \lfloor t/\eta \rfloor \eta) & rS(\beta)\eta\delta(t - \lfloor t/\eta \rfloor \eta) \\ rS(\beta)\eta\delta(t - \lfloor t/\eta \rfloor \eta) & -rS(\beta)\eta\delta(t - \lfloor t/\eta \rfloor \eta) \end{pmatrix},$$

where $\delta(\cdot)$ is a Dirac delta function. The diffusion matrix $\Sigma(\alpha_t)$ is thus defined as $(\Sigma(0), \Sigma(1)) := \left\{ \begin{pmatrix} \sqrt{2\tau^{(1)}}\mathbf{I}_d & 0 \\ 0 & \sqrt{2\tau^{(2)}}\mathbf{I}_d \end{pmatrix}, \begin{pmatrix} \sqrt{2\tau^{(2)}}\mathbf{I}_d & 0 \\ 0 & \sqrt{2\tau^{(1)}}\mathbf{I}_d \end{pmatrix} \right\}$. The process β_t is the unique solution to the Skorokhod problem for the jump process defined by:

$$\begin{cases} d\mathbf{y}_t = -\nabla G(\beta_t)dt + \Sigma(\alpha_t)d\mathbf{W}_t, \\ \mathbb{P}(\alpha(t) = j | \alpha(t-dt) = l, \beta(\lfloor t/\eta \rfloor \eta) = \beta) = rS(\beta)\eta \mathbf{1}_{\{t=\lfloor t/\eta \rfloor \eta\}} + o(dt), \text{ for } l \neq j. \end{cases} \quad (36)$$

The existence of the unique solution for the Skorokhod problem for $\mathbf{y}$ follows from the original proof in (Tanaka, 1979) (see also (Menaldi & Robin, 1985)) since the process $\mathbf{y}_t$ has continuous paths. The switching happens after the reflection at each time $i\eta < T$. We should think of the process as concatenated on the time interval $[i\eta, (i+1)\eta)$ up to time horizon T . Similarly, we consider the following **Reflected Replica exchange stochastic gradient Langevin diffusion**, for the same learning rate $\eta > 0$ as above, we have

$$\begin{cases} d\tilde{\beta}_t^\eta = -\nabla \tilde{G}(\tilde{\beta}_{\lfloor t/\eta \rfloor \eta}^\eta)dt + \Sigma(\tilde{\alpha}_{\lfloor t/\eta \rfloor \eta})d\mathbf{W}_t + \tilde{\nu}_t \tilde{L}(dt), \\ \mathbb{P}(\tilde{\alpha}(t) = j | \tilde{\alpha}(t-dt) = l, \tilde{\beta}(\lfloor t/\eta \rfloor \eta) = \tilde{\beta}) = r\tilde{S}(\tilde{\beta})\eta \mathbf{1}_{\{t=\lfloor t/\eta \rfloor \eta\}} + o(dt), \text{ for } l \neq j, \end{cases} \quad (37)$$

where $\nabla \tilde{G}(\tilde{\beta}) := \begin{pmatrix} \nabla \tilde{U}(\tilde{\beta}^{(1)}) \\ \nabla \tilde{U}(\tilde{\beta}^{(2)}) \end{pmatrix}$ and $\tilde{S}(\tilde{\beta}) = \min\{1, \tilde{S}(\tilde{\beta}^{(1)}, \tilde{\beta}^{(2)})\}$ and $\tilde{S}(\tilde{\beta}^{(1)}, \tilde{\beta}^{(2)})$ is defined in Section 2.3. Here

$\int_0^t \tilde{\nu}_s \tilde{L}(ds) = \int_0^t \begin{pmatrix} \tilde{\nu}_s^{(1)} & 0 \\ 0 & \tilde{\nu}_s^{(2)} \end{pmatrix} \begin{pmatrix} \tilde{L}^{(1)}(ds) \\ \tilde{L}^{(2)}(ds) \end{pmatrix}$ denotes the bounded variation reflection process ensuring $\tilde{\beta}_t^\eta$ stays inside the domain $\Omega \times \Omega$.

The reflection process $\tilde{\beta}_t^\eta$ is then the unique solution to the Skorokhod problem for the following process,

$$\begin{cases} d\tilde{\mathbf{y}}_t^\eta = -\nabla \tilde{G}(\tilde{\beta}_{\lfloor t/\eta \rfloor \eta}^\eta)dt + \Sigma(\tilde{\alpha}_{\lfloor t/\eta \rfloor \eta})d\mathbf{W}_t, \\ \mathbb{P}(\tilde{\alpha}(t) = j | \tilde{\alpha}(t-dt) = l, \tilde{\beta}(\lfloor t/\eta \rfloor \eta) = \tilde{\beta}) = r\tilde{S}(\tilde{\beta})\eta \mathbf{1}_{\{t=\lfloor t/\eta \rfloor \eta\}} + o(dt), \text{ for } l \neq j, \end{cases} \quad (38)$$

which is the continuous interpolation of the replica exchange stochastic gradient Langevin dynamics.

We then show the following proof of Theorem 3.11.

Proof of Theorem 3.11 According to the definition of the **Reflected Replica exchange Langevin diffusion** $\{\beta_t\}_{t \geq 0}$ and the **Reflected Replica exchange stochastic gradient Langevin diffusion** $\{\tilde{\beta}_t^\eta\}_{t \geq 0}$, we can write the difference of the two

process $\beta_t - \tilde{\beta}_t^\eta$ in the following form,

$$\beta_t - \tilde{\beta}_t^\eta = \mathbf{y}_t - \tilde{\mathbf{y}}_t + \int_0^t \nu_s L(ds) - \int_0^t \tilde{\nu}_s \tilde{L}(ds). \quad (39)$$

Since both $\mathbf{y}_t$ and $\tilde{\mathbf{y}}_t$ are piece wise continuous paths, and $(\beta_t, \int_0^t \nu_s L(ds))$ and $(\tilde{\beta}_t, \int_0^t \tilde{\nu}_s \tilde{L}(ds))$ solve the Skorokhod problems for $\mathbf{y}_t$ and $\tilde{\mathbf{y}}_t$, respectively. Following (Tanaka, 1979) (Lemma 2.2), we have the following estimates,

$$\|\beta_t - \tilde{\beta}_t^\eta\|^2 \leq \underbrace{\|\mathbf{y}_t - \tilde{\mathbf{y}}_t^\eta\|^2 + 2 \int_0^t (\mathbf{y}_t - \tilde{\mathbf{y}}_t^\eta - \mathbf{y}_s + \tilde{\mathbf{y}}_s^\eta)^\top (\nu_s L(ds) - \tilde{\nu}_s \tilde{L}(ds))}_{\mathcal{K}}. \quad (40)$$

Denote $h_{\Omega \times \Omega}$ as the support function of $\Omega \times \Omega$, which is defined by $h_{\Omega \times \Omega}(y) = \sup\{\langle x, y \rangle; x \in \Omega \times \Omega\}$, for $y \in \Omega \times \Omega$. For simplification, we assume that the origin is inside the convex domain Ω . Thus the Euclidean norm $\|\cdot\|$ and the norm $\|\cdot\|_{\Omega \times \Omega}$ are equivalent, up to a constant C , by taking Ω as d -dimensional ball. Following the estimates (Bubeck et al., 2018) (Proposition 1) and the equivalent norm in our setting, for the inner normal vectors ν_s and $\tilde{\nu}_s$, we have

$$\begin{aligned} \sup_{0 \leq t \leq T} \sqrt{\mathcal{K}} &\leq \sup_{0 \leq t \leq T} \left\{ 2C(\|\mathbf{y}_t - \tilde{\mathbf{y}}_t^\eta\| + \sup_{0 \leq s \leq t} \|\mathbf{y}_s - \tilde{\mathbf{y}}_s^\eta\|) \int_0^t h_{\Omega \times \Omega}(-\nu_s) L(ds) \right. \\ &\quad \left. + 2C(\|\mathbf{y}_t - \tilde{\mathbf{y}}_t^\eta\| + \sup_{0 \leq s \leq t} \|\mathbf{y}_s - \tilde{\mathbf{y}}_s^\eta\|) \int_0^t h_{\Omega \times \Omega}(-\tilde{\nu}_s) \tilde{L}(ds) \right\}^{1/2} \\ &\leq 2C \sqrt{\sup_{0 \leq t \leq T} \|\mathbf{y}_t - \tilde{\mathbf{y}}_t^\eta\| \left(\int_0^T h_{\Omega \times \Omega}(-\nu_s) L(ds) + \int_0^T h_{\Omega \times \Omega}(-\tilde{\nu}_s) \tilde{L}(ds) \right)}, \end{aligned} \quad (41)$$

where the last inequality follows from the fact L_t ($\tilde{L}_t$ resp.) is increasing. Taking square root, $\sup_{0 \leq t \leq T}$ and expectation $\mathbb{E}[\cdot]$ on both sides of (40), together with Cauchy-Schwartz inequality, applying (41), we get

$$\begin{aligned} \mathbb{E} \left[\sup_{0 \leq t \leq T} \|\beta_t - \tilde{\beta}_t^\eta\| \right] &\leq \mathbb{E} \left[\sup_{0 \leq t \leq T} \|\mathbf{y}_t - \tilde{\mathbf{y}}_t^\eta\| \right] \cdots \mathcal{I} \\ &\quad + \underbrace{\sqrt{\mathbb{E} \left[4C \sup_{0 \leq t \leq T} \|\mathbf{y}_t - \tilde{\mathbf{y}}_t^\eta\| \right]} \sqrt{\mathbb{E} \left[\left(\int_0^T h_{\Omega \times \Omega}(-\nu_s) L(ds) + \int_0^T h_{\Omega \times \Omega}(-\tilde{\nu}_s) \tilde{L}(ds) \right) \right]}}_{\mathcal{J}}. \end{aligned} \quad (42)$$

Estimates of $\mathcal{I}$: The estimates of the first term $\mathcal{I}$ is similar to (Deng et al., 2020a) (Lemma 1). We follow the similar steps here and highlight the major differences coming from the fact that β_t and $\tilde{\beta}_t^\eta$ are reflected processes in the current setting. Plugging in the dynamics for $\mathbf{y}_t$ and $\tilde{\mathbf{y}}_t$, we first have

$$\mathbf{y}_t - \tilde{\mathbf{y}}_t^\eta = \int_0^t \nabla G(\beta_s) - \nabla \tilde{G}(\tilde{\beta}_{[s/\eta]\eta}^\eta) ds + \int_0^t (\Sigma_s - \tilde{\Sigma}_{[s/\eta]\eta}^\eta) d\mathbf{W}_s. \quad (43)$$

Taking $\sup_{0 \leq t \leq T}$, and expectation of both sides of (43), then applying Burkholder-Davis-Gundy inequality and Cauchy-Schwarz inequality, we have

$$\begin{aligned} \mathbb{E} \left[\sup_{0 \leq t \leq T} \|\mathbf{y}_t - \tilde{\mathbf{y}}_t^\eta\| \right] &\leq \mathbb{E} \left[\int_0^T \|\nabla G(\beta_s) - \nabla \tilde{G}(\tilde{\beta}_{[s/\eta]\eta}^\eta)\| ds + \sup_{0 \leq t \leq T} \left\| \int_0^t (\Sigma_s - \tilde{\Sigma}_{[s/\eta]\eta}^\eta) d\mathbf{W}_s \right\| \right] \\ &\leq \underbrace{\mathbb{E} \left[\int_0^T \|\nabla G(\beta_s) - \nabla \tilde{G}(\tilde{\beta}_{[s/\eta]\eta}^\eta)\| ds \right]}_{\mathcal{I}_1} + C \underbrace{\sqrt{\mathbb{E} \left[\int_0^T \|\Sigma_s - \tilde{\Sigma}_{[s/\eta]\eta}^\eta\|^2 ds \right]}}_{\mathcal{I}_2}. \end{aligned}$$

The estimate of $\mathcal{I}_2$ follows the same as in (Deng et al., 2020a) (Lemma 1, estimates (15), (16), (17)), which gives us

$$\begin{aligned}
 (\mathcal{I}_2)^2 &= C\mathbb{E} \left[\int_0^T \|\Sigma_s - \tilde{\Sigma}_{\lfloor s/\eta \rfloor \eta}\|^2 ds \right] \\
 &\leq C \sum_{j=1}^{2d} \sum_{k=0}^{\lfloor T/\eta \rfloor} \int_{k\eta}^{(k+1)\eta} \mathbb{E} \left[\|\Sigma_s(j) - \tilde{\Sigma}_{k\eta}^\eta(j)\|^2 \right] ds \\
 &\leq C \sum_{j=1}^{2d} \sum_{k=0}^{\lfloor T/\eta \rfloor} \int_{k\eta}^{(k+1)\eta} \mathbb{E} \left[\|\Sigma_s(j) - \Sigma_{k\eta}^\eta(j) + \Sigma_{k\eta}^\eta(j) - \tilde{\Sigma}_{k\eta}^\eta(j)\|^2 \right] ds \\
 &\leq C \sum_{j=1}^{2d} \sum_{k=0}^{\lfloor T/\eta \rfloor} \left[\int_{k\eta}^{(k+1)\eta} \mathbb{E} \left[\|\Sigma_s(j) - \Sigma_{k\eta}^\eta(j)\|^2 \right] ds + \int_{k\eta}^{(k+1)\eta} \mathbb{E} \left[\|\Sigma_{k\eta}^\eta(j) - \tilde{\Sigma}_{k\eta}^\eta(j)\|^2 \right] ds \right] \\
 &\leq Cd(1+T)\tilde{\delta}_2(\tau_1, \tau_2) \left(\eta + \max_k \sqrt{\mathbb{E}[\|\psi_k\|^2]} \right).
 \end{aligned} \tag{44}$$

where $\tilde{\delta}_2(\tau_1, \tau_2) = 4(\sqrt{\tau_2} - \sqrt{\tau_1})^2$, and ψ_k is the noise in the swapping rate. For convenience, the constant C may vary from line to line. We mainly focus on showing the estimates for the term $\mathcal{I}_1$. Applying the inequality $\|a + b + c\|^2 \leq 3(\|a\|^2 + \|b\|^2 + \|c\|^2)$, we get

$$\begin{aligned}
 \mathcal{I}_1 &\leq \underbrace{\mathbb{E} \left[\int_0^T \|\nabla G(\beta_s) - \nabla G(\tilde{\beta}_s^\eta)\| ds \right]}_{\mathcal{I}_{11}} + \underbrace{\left(T\mathbb{E} \left[\int_0^T \|\nabla G(\tilde{\beta}_s^\eta) - \nabla G(\tilde{\beta}_{\lfloor s/\eta \rfloor \eta}^\eta)\|^2 ds \right] \right)^{1/2}}_{\mathcal{I}_{12}} \\
 &\quad + \underbrace{\left(T\mathbb{E} \left[\int_0^T \|\nabla G(\tilde{\beta}_{\lfloor s/\eta \rfloor \eta}^\eta) - \nabla \tilde{G}(\tilde{\beta}_{\lfloor s/\eta \rfloor \eta}^\eta)\|^2 ds \right] \right)^{1/2}}_{\mathcal{I}_{13}} \\
 &\leq \mathcal{I}_{11} + \mathcal{I}_{12} + \mathcal{I}_{13}.
 \end{aligned} \tag{45}$$

By using the smoothness assumption (A2), we first get

$$\mathcal{I}_{11} \leq L\mathbb{E} \left[\int_0^T \|\beta_s - \tilde{\beta}_s^\eta\| ds \right].$$

Following (Deng et al., 2020a) (Lemma 1, estimates (10)), we next get

$$\begin{aligned}
 (\mathcal{I}_{12})^2 &\leq TL^2\mathbb{E} \left[\int_0^T \|\tilde{\beta}_s^\eta - \tilde{\beta}_{\lfloor s/\eta \rfloor \eta}^\eta\|^2 ds \right] \\
 &\leq TL^2 \sum_{k=0}^{\lfloor T/\eta \rfloor} \mathbb{E} \left[\int_{k\eta}^{(k+1)\eta} \|\tilde{\beta}_s^\eta - \tilde{\beta}_{\lfloor s/\eta \rfloor \eta}^\eta\|^2 ds \right] \\
 &\leq TL^2 \sum_{k=0}^{\lfloor T/\eta \rfloor} \int_{k\eta}^{(k+1)\eta} \mathbb{E} \left[\sup_{k\eta \leq s < (k+1)\eta} \|\tilde{\beta}_s^\eta - \tilde{\beta}_{\lfloor s/\eta \rfloor \eta}^\eta\|^2 \right] ds.
 \end{aligned} \tag{46}$$

For $\forall k \in \mathbb{N}$ and $s \in [k\eta, (k+1)\eta]$, since $\tilde{\beta}_{\lfloor s/\eta \rfloor \eta}^\eta$ stays constant on $[k\eta, (k+1)\eta]$, which does not involve reflection, we thus have

$$\tilde{\beta}_s^\eta - \tilde{\beta}_{\lfloor s/\eta \rfloor \eta}^\eta = \tilde{\beta}_s^\eta - \tilde{\beta}_{k\eta}^\eta = -\nabla \tilde{G}(\tilde{\beta}_{k\eta}^\eta) \cdot (s - k\eta) + \tilde{\Sigma}_{k\eta}^\eta \int_{k\eta}^s d\mathbf{W}_r + \int_{k\eta}^s \tilde{\nu}_r \tilde{L}(dr),$$

where $\int_{k\eta}^s \tilde{\nu}_r \tilde{L}(dr)$ ensures the process $\tilde{\beta}_s^\eta$ stays inside the domain $\Omega \times \Omega$. Since there is no swap for the diffusion matrix $\Sigma(\tilde{\alpha}_{\lfloor s/\eta \rfloor \eta})$ for $s \in [k\eta, (k+1)\eta)$, the existence of $\int_{k\eta}^s \tilde{\nu}_r \tilde{L}(dr)$ follows from the Skorohod problem for a drift-diffusion process (Tanaka, 1979). Applying Itô's formula to $\|\tilde{\beta}_s^\eta - \tilde{\beta}_{\lfloor s/\eta \rfloor \eta}^\eta\|^2$, we have

$$\begin{aligned} \|\tilde{\beta}_s^\eta - \tilde{\beta}_{\lfloor s/\eta \rfloor \eta}^\eta\|^2 &= 2 \int_{k\eta}^s \langle \tilde{\beta}_r^\eta - \tilde{\beta}_{\lfloor r/\eta \rfloor \eta}^\eta, \Sigma(\tilde{\alpha}_{\lfloor r/\eta \rfloor \eta}) d\mathbf{W}_r \rangle - 2 \int_{k\eta}^s \langle \tilde{\beta}_r^\eta - \tilde{\beta}_{\lfloor r/\eta \rfloor \eta}^\eta, \nabla \tilde{G}(\tilde{\beta}_{\lfloor r/\eta \rfloor \eta}^\eta) \rangle dr \\ &\quad + 2 \int_{k\eta}^s \langle \tilde{\beta}_r^\eta - \tilde{\beta}_{\lfloor r/\eta \rfloor \eta}^\eta, \tilde{\nu}_r \rangle \tilde{L}(dr) \\ &\leq 2 \int_{k\eta}^s \langle \tilde{\beta}_r^\eta - \tilde{\beta}_{\lfloor r/\eta \rfloor \eta}^\eta, \Sigma(\tilde{\alpha}_{\lfloor r/\eta \rfloor \eta}) d\mathbf{W}_r \rangle - 2 \int_{k\eta}^s \langle \tilde{\beta}_r^\eta - \tilde{\beta}_{\lfloor r/\eta \rfloor \eta}^\eta, \nabla \tilde{G}(\tilde{\beta}_{\lfloor r/\eta \rfloor \eta}^\eta) \rangle dr, \end{aligned}$$

where the last inequality follows from the fact that both $\tilde{\beta}_r^\eta$ and $\tilde{\beta}_{\lfloor r/\eta \rfloor \eta}^\eta$ are inside the domain, $\tilde{\nu}_r$ is a inner normal vector at $\tilde{\beta}_s^\eta$, and $\tilde{L}$ is non-negative measure, thus $\langle \tilde{\beta}_r^\eta - \tilde{\beta}_{\lfloor r/\eta \rfloor \eta}^\eta, \tilde{\nu}_r \rangle \leq 0$, since the domain is convex. Taking the $\sup_{k\eta \leq s \leq (k+1)\eta}$ and the expectation $\mathbb{E}$ on both sides, and applying the Burkholder-Davis-Gundy inequality, Cauchy-Schwarz inequality and the Young inequality (i.e. $2ab \leq \varepsilon a^2 + b^2/\varepsilon$), we have

$$\begin{aligned} \mathbb{E} \left[\sup_{k\eta \leq s \leq (k+1)\eta} \|\tilde{\beta}_s^\eta - \tilde{\beta}_{\lfloor s/\eta \rfloor \eta}^\eta\|^2 \right] &\leq 2\mathbb{E} \left[\left\| \sup_{k\eta \leq s \leq (k+1)\eta} \int_{k\eta}^s \langle \tilde{\beta}_r^\eta - \tilde{\beta}_{\lfloor r/\eta \rfloor \eta}^\eta, \Sigma(\tilde{\alpha}_{k\eta}) d\mathbf{W}_r \rangle \right\|^2 \right] \\ &\quad + 2\mathbb{E} \left[\sup_{k\eta \leq s \leq (k+1)\eta} \int_{k\eta}^s \langle \tilde{\beta}_r^\eta - \tilde{\beta}_{\lfloor r/\eta \rfloor \eta}^\eta, \nabla \tilde{G}(\tilde{\beta}_{k\eta}^\eta) \rangle dr \right] \\ &\leq 2C\mathbb{E} \left[\left(\int_{k\eta}^{(k+1)\eta} \|\langle \tilde{\beta}_r^\eta - \tilde{\beta}_{\lfloor r/\eta \rfloor \eta}^\eta, \Sigma(\tilde{\alpha}_{k\eta}) \rangle\|^2 dr \right)^{1/2} \right] \\ &\quad + \mathbb{E} \left[\sup_{k\eta \leq s \leq (k+1)\eta} \int_{k\eta}^s \|\tilde{\beta}_r^\eta - \tilde{\beta}_{\lfloor r/\eta \rfloor \eta}^\eta\|^2 dr + \eta^2 \|\nabla \tilde{G}(\tilde{\beta}_{k\eta}^\eta)\|^2 \right] \\ &\leq C\mathbb{E} \left[\sup_{k\eta \leq s \leq (k+1)\eta} \|\langle \tilde{\beta}_s^\eta - \tilde{\beta}_{\lfloor s/\eta \rfloor \eta}^\eta\|^2 \varepsilon + \frac{1}{\varepsilon} \tau_2^2 \eta \right] \\ &\quad + \int_{k\eta}^{(k+1)\eta} \mathbb{E} \left[\sup_{k\eta \leq r \leq (k+1)\eta} \|\tilde{\beta}_r^\eta - \tilde{\beta}_{\lfloor r/\eta \rfloor \eta}^\eta\|^2 \right] dr + \mathbb{E}[\eta^2 \|\nabla \tilde{G}(\tilde{\beta}_{k\eta}^\eta)\|^2], \end{aligned}$$

where $0 < \varepsilon < 1$ is a small constant. Pick ε such that $C\varepsilon < 1$, we end up with

$$\begin{aligned} \mathbb{E} \left[\sup_{k\eta \leq s \leq (k+1)\eta} \|\tilde{\beta}_s^\eta - \tilde{\beta}_{\lfloor s/\eta \rfloor \eta}^\eta\|^2 \right] &\leq \frac{1}{1 - C\varepsilon} \left(\frac{1}{\varepsilon} \tau_2^2 \eta + \mathbb{E}[\eta^2 \|\nabla \tilde{G}(\tilde{\beta}_{k\eta}^\eta)\|^2] \right) \\ &\quad + \frac{1}{1 - C\varepsilon} \int_{k\eta}^{(k+1)\eta} \mathbb{E} \left[\sup_{k\eta \leq r \leq (k+1)\eta} \|\tilde{\beta}_r^\eta - \tilde{\beta}_{\lfloor r/\eta \rfloor \eta}^\eta\|^2 \right] dr, \end{aligned}$$

We further have the following estimates,

$$\begin{aligned} \eta^2 \mathbb{E}[\|\nabla \tilde{G}(\tilde{\beta}_{k\eta}^\eta)\|^2] &= \eta^2 \mathbb{E}[\|(\nabla G(\tilde{\beta}_{k\eta}^\eta) + \phi_k)\|^2] \\ &\leq 2\eta^2 \mathbb{E}[\|\nabla G(\tilde{\beta}_{k\eta}^\eta) - \nabla G(\beta^*)\|^2 + \|\phi_k\|^2] \\ &\leq 4C^2 \eta^2 \mathbb{E}[\|\tilde{\beta}_{k\eta}^\eta\|^2 + \|\beta^*\|^2] + 2\eta^2 \mathbb{E}[\|\phi_k\|^2], \end{aligned}$$

where the first inequality follows from the separation of the noise from the stochastic gradient and the choice of stationary point β^* of $G(\cdot)$ with $\nabla G(\beta^*) = 0$, and ϕ_k is the stochastic noise in the gradient at step k . Thus, combining the above two parts with Grownall's inequality, we get

$$\begin{aligned} &\mathbb{E} \left[\sup_{k\eta \leq s \leq (k+1)\eta} \|\tilde{\beta}_s^\eta - \tilde{\beta}_{\lfloor s/\eta \rfloor \eta}^\eta\|^2 \right] \\ &\leq \frac{1}{1 - C\varepsilon} \left(4L^2(1 + \eta)\eta^2 \mathbb{E}[\|\tilde{\beta}_{k\eta}^\eta\|^2 + \|\beta^*\|^2] + 4(1 + \eta)\eta^2 \mathbb{E}[\|\phi_k\|^2] + \frac{1}{\varepsilon} \tau_2^2 \eta(1 + \eta) \right). \end{aligned}$$

Plugging the above estimates back in (46), we get

$$\begin{aligned} (\mathcal{I}_{12})^2 &\leq TL^2(1+T/\eta) \left[4L^2(1+\eta)\eta^3 \sup_{k \geq 0} \mathbb{E}[\|\tilde{\beta}_{k\eta}^\eta\|^2 + \|\beta^*\|^2] + 4\eta^3(1+\eta) \max_{k \geq 0} \mathbb{E}[\|\phi_k\|^2] + \frac{1}{\varepsilon} \tau_2^2 \eta^2 (1+\eta) \right] \\ &\leq \tilde{\delta}_1(d, \tau_2, T, C, \varepsilon) \eta^2 + 4TL^2(1+T)\eta^2 \max_k \mathbb{E}[\|\phi_k\|^2], \end{aligned}$$

where $\tilde{\delta}_1(d, \tau_2, T, C, \varepsilon)$ is a constant depending on d, τ_2, T, C , and ε . Note that the above inequality requires a result on the bounded second moment of $\sup_{k \geq 0} \mathbb{E}[\|\tilde{\beta}_{k\eta}^\eta\|^2]$, which follows from the bounded domain assumption. We are now left to estimate the term $\mathcal{I}_{13}$ and we have

$$\begin{aligned} (\mathcal{I}_{13})^2 &\leq T \sum_{k=0}^{\lfloor T/\eta \rfloor} \mathbb{E} \left[\int_{k\eta}^{(k+1)\eta} \|\nabla G(\tilde{\beta}_{k\eta}^\eta) - \nabla \tilde{G}(\tilde{\beta}_{k\eta}^\eta)\|^2 ds \right] \\ &\leq T(1+T/\eta) \max_k \mathbb{E}[\|\phi_k\|^2] \eta \\ &\leq T(1+T) \max_k \mathbb{E}[\|\phi_k\|^2]. \end{aligned} \tag{47}$$

Combing all the estimates of $\mathcal{I}_{11}, \mathcal{I}_{12}$ and $\mathcal{I}_{13}$, we obtain

$$\begin{aligned} \mathcal{I}_1 &\leq L \underbrace{\int_0^T \mathbb{E} \left[\sup_{0 \leq s \leq T} \|\beta_s - \tilde{\beta}_s^\eta\| \right] ds}_{\mathcal{I}_{11}} + \underbrace{\left(\tilde{\delta}_1(d, \tau_2, T, C, \varepsilon) \eta^2 + 4TL^2(1+T)\eta^2 \max_k \mathbb{E}[\|\phi_k\|^2] \right)^{1/2}}_{\mathcal{I}_{12}} \\ &\quad + \underbrace{\left(T(1+T) \max_k \mathbb{E}[\|\phi_k\|^2] \right)^{1/2}}_{\mathcal{I}_{13}}. \end{aligned} \tag{48}$$

Combining the estimates in $\mathcal{I}_1$ and $\mathcal{I}_2$, we have

$$\begin{aligned} \mathbb{E} \left[\sup_{0 \leq t \leq T} \|\mathbf{y}_t - \tilde{\mathbf{y}}_t^\eta\| \right] &\leq L \underbrace{\int_0^T \mathbb{E} \left[\sup_{0 \leq t \leq T} \|\beta_t - \tilde{\beta}_t^\eta\| \right] dt}_{\mathcal{I}_{11}} \\ &\quad + \underbrace{\left(\tilde{\delta}_1(d, \tau_2, T, C, \varepsilon) \eta^2 + 4TL^2(1+T)\eta^2 \max_k \mathbb{E}[\|\phi_k\|^2] \right)^{1/2}}_{\mathcal{I}_{12}} \\ &\quad + \underbrace{\left(T(1+T) \max_k \mathbb{E}[\|\phi_k\|^2] \right)^{1/2}}_{\mathcal{I}_{13}} + \underbrace{\sqrt{Cd(1+T)\tilde{\delta}_2(\tau_1, \tau_2) \left(\eta + \max_k \sqrt{\mathbb{E}[\|\psi_k\|^2]} \right)}}_{\mathcal{I}_2} \end{aligned} \tag{49}$$

Estimates of $\mathcal{J}$: we first observe that,

$$\begin{aligned} \mathcal{J} &= \sqrt{\mathbb{E} \left[4C \sup_{0 \leq t \leq T} \|\mathbf{y}_t - \tilde{\mathbf{y}}_t^\eta\| \right]} \sqrt{\mathbb{E} \left[\left(\int_0^T h_{\Omega \times \Omega}(\nu_s) L(ds) + \int_0^T h_{\Omega \times \Omega}(\tilde{\nu}_s) \tilde{L}(ds) \right) \right]} \\ &\leq \frac{1}{2} \left(\frac{1}{\eta^{1/4}} \mathbb{E} \left[4C \sup_{0 \leq t \leq T} \|\mathbf{y}_t - \tilde{\mathbf{y}}_t^\eta\| \right] + \eta^{1/4} \mathbb{E} \left[\left(\int_0^T h_{\Omega \times \Omega}(\nu_s) L(ds) + \int_0^T h_{\Omega \times \Omega}(\tilde{\nu}_s) \tilde{L}(ds) \right) \right] \right). \end{aligned}$$

The above inequality follows from Young's inequality $2ab \leq a^2 \frac{1}{\eta^c} + b^2 \eta^c$. The choice of $0 < c < 1/2$ will affect the rate in (54). The first term follows from (49). We are left to estimate the second term. For $s \in [k\eta, (k+1)\eta)$, β_s and $\tilde{\beta}_s$ are reflected diffusion process without swap, by Itô's formula, we have

$$\begin{aligned} 2 \int_{k\eta}^s \langle \tilde{\beta}_r^\eta, -\tilde{\nu}_r \rangle \tilde{L}(dr) &= 2 \int_{k\eta}^s \langle \tilde{\beta}_r^\eta, \Sigma(\tilde{\alpha}_{\lfloor r/\eta \rfloor \eta}) d\mathbf{W}_r \rangle - 2 \int_{k\eta}^s \langle \tilde{\beta}_r^\eta, \nabla \tilde{G}(\tilde{\beta}_{\lfloor r/\eta \rfloor \eta}^\eta) dr \rangle \\ &\quad + \|\tilde{\beta}_{k\eta}^\eta\|^2 - \|\tilde{\beta}_s^\eta\|^2 + 2d(s - k\eta). \end{aligned} \tag{50}$$

Take expectation on both sides, since the domain is convex, and $\tilde{\nu}_s$ is an inner normal vector, similar to (Bubeck et al., 2018) (Lemma 9), for the support function $h_{\Omega \times \Omega}(\cdot)$, we have

$$\begin{aligned} \mathbb{E}[2 \int_{k\eta}^s h_{\Omega \times \Omega}(-\tilde{\nu}_r) \tilde{L}(dr)] &\leq \mathbb{E}[\int_{k\eta}^s \|\langle \tilde{\beta}_r^\eta, \nabla \tilde{G}(\tilde{\beta}_{\lfloor r/\eta \rfloor \eta}^\eta) \rangle\| dr] + 2d(s - k\eta) \\ &\leq \frac{(2d + R\tilde{L})(s - k\eta)}{2} \leq \frac{(2d + R\tilde{L})\eta}{2}, \end{aligned} \quad (51)$$

where R is the upper bound of the diameter of the domain, and $\tilde{L}$ is the upper bound of the gradient $\nabla \tilde{G}$. Summing over the estimates $\int_{k\eta}^{(k+1)\eta}$ for $k = 0, \dots, \lfloor T/\eta \rfloor$, we get the bound,

$$\int_0^T h_{\Omega \times \Omega}(\tilde{\nu}_s) \tilde{L}(ds) \leq \frac{(2d + R\tilde{L})T}{2}. \quad (52)$$

The same estimate also holds true for $\int_0^T h_{\Omega \times \Omega}(\nu_s) L(ds)$, which gives us

$$\mathcal{J} \leq \frac{2C}{\eta^{1/4}} \mathbb{E} \left[\sup_{0 \leq t \leq T} \|\mathbf{y}_t - \tilde{\mathbf{y}}_t^\eta\| \right] + (2d + R\tilde{L})T\eta^{1/4}. \quad (53)$$

Estimates of (42): Plugging the estimates for $\mathcal{I}$ (49) and $\mathcal{J}$ (53) into (42), we have

$$\begin{aligned} \mathbb{E} \left[\sup_{0 \leq t \leq T} \|\beta_t - \tilde{\beta}_t^\eta\| \right] &\leq (1 + \frac{2C}{\eta^{1/4}}) \mathbb{E} \left[\sup_{0 \leq t \leq T} \|\mathbf{y}_t - \tilde{\mathbf{y}}_t^\eta\| \right] + (2d + R\tilde{L})T\eta^{1/4} \\ &\leq \frac{C}{\eta^{1/4}} \mathbb{E} \left[\sup_{0 \leq t \leq T} \|\mathbf{y}_t - \tilde{\mathbf{y}}_t^\eta\| \right] + (2d + R\tilde{L})T\eta^{1/4} \\ &\leq \frac{C}{\eta^{1/4}} \int_0^T \mathbb{E} \left[\sup_{0 \leq t \leq T} \|\beta_t - \tilde{\beta}_t^\eta\| \right] dt + (2d + R\tilde{L})T\eta^{1/4} \\ &\quad + \frac{C}{\eta^{1/4}} \left(\tilde{\delta}_1(d, \tau_2, T, C, \varepsilon)\eta^2 + 4TL^2(1+T)\eta^2 \max_k \mathbb{E}[\|\phi_k\|^2] \right)^{1/2} \\ &\quad + \frac{C}{\eta^{1/4}} \left(T(1+T) \max_k \mathbb{E}[\|\phi_k\|^2] \right)^{1/2} + \frac{C}{\eta^{1/4}} \left(Cd(1+T)\tilde{\delta}_2(\tau_1, \tau_2) \left(\eta + \max_k \sqrt{\mathbb{E}[\|\psi_k\|^2]} \right) \right)^{1/2}. \end{aligned}$$

Here the constant C varies from line to line. Applying Grownall's inequality, we have

$$\mathbb{E} \left[\sup_{0 \leq t \leq T} \|\beta_t - \tilde{\beta}_t^\eta\| \right] \leq \delta_1 \eta^{1/4} + \delta_2 \sqrt{\max_k \mathbb{E}[\|\phi_k\|^2]} + \delta_3 \sqrt{\eta^{-1/2} \max_k \sqrt{\mathbb{E}[\|\psi_k\|^2]}}, \quad (54)$$

where δ_1 is a constant depending on $\tau_1, \tau_2, d, T, C, R, L, \tilde{L}, \varepsilon$; δ_2 depends on T , and C ; δ_3 depends on d, T , and C . By the definition of $\mathcal{W}_1$ distance, we get the following convergence rate in $\mathcal{W}_1$ distance,

$$\mathcal{W}_1(\mathcal{L}(\beta_T), \mathcal{L}(\tilde{\beta}_T^\eta)) \leq \mathbb{E}[\|\beta_T - \tilde{\beta}_T^\eta\|] \leq E \left[\sup_{0 \leq t \leq T} \|\beta_t - \tilde{\beta}_t^\eta\| \right]. \quad (55)$$

□