

NOISE-FREE SAMPLING ALGORITHMS VIA REGULARIZED WASSERSTEIN PROXIMALS

HONG YE TAN, STANLEY OSHER, WUCHEN LI

ABSTRACT. We consider the problem of sampling from a distribution governed by a potential function. This work proposes an explicit score-based MCMC method that is deterministic, resulting in a deterministic evolution for particles rather than a stochastic differential equation evolution. The score term is given in closed form by a regularized Wasserstein proximal, using a kernel convolution that is approximated by sampling. We demonstrate fast convergence on various problems and show improved dimensional dependence of mixing time bounds for the case of Gaussian distributions compared to the unadjusted Langevin algorithm (ULA) and the Metropolis-adjusted Langevin algorithm (MALA). We additionally derive closed form expressions for the distributions at each iterate for quadratic potential functions, characterizing the variance reduction. Empirical results demonstrate that the particles behave in an organized manner, lying on level set contours of the potential. Moreover, the posterior mean estimator of the proposed method is shown to be closer to the maximum a-posteriori estimator compared to ULA and MALA, in the context of Bayesian logistic regression.

1. INTRODUCTION

Sampling from an unknown distribution is a fundamental task in data science. Notable applications include maximum likelihood estimation and uncertainty quantification (Laumont et al., 2022), Bayesian neural networks training (MacKay, 1995), global optimization (Dai et al., 2021), and generative modelling (Batzolis et al., 2021; Hyvärinen & Dayan, 2005; Song et al., 2020). In general, the problem can be formulated as sampling from a Gibbs distribution, with a density of the form

$$\rho(x) \sim \exp(-V(x)),$$

where V is a known bounded $\mathcal{C}^1$ potential function, satisfying appropriate growth conditions such that ρ is a well defined density function. One popular way to do this is using Markov chain Monte Carlo (MCMC) algorithms (Andrieu et al., 2003; Brooks et al., 2011). MCMC algorithms work by first constructing a Markov chain whose stationary distribution is equal or close to the target distribution. By using ergodic theory, the Markov chains can be shown to converge in distribution from a tractable initial distribution to the intractable stationary distribution. Hence, to sample from the target distribution, one needs only evaluate the Markov chain for a suitably large number of iterations.

There are three main paradigms for MCMC: zeroth order methods, first order methods, and score based methods. Some examples of zeroth order methods include the Metropolized random walk and hit-and-run algorithms, which do not use the gradient of the potential ∇V (Mengersen & Tweedie, 1996; Bélisle et al., 1993). First order methods utilize the gradient of our potential V as well as randomness to converge in distribution to the target distribution. Two of the most popular first order methods are the unadjusted Langevin algorithm (ULA) and the Metropolis-adjusted Langevin algorithm (MALA) (Parisi, 1981; Durmus & Moulines, 2019; Rossky et al., 1978; Brooks et al., 2011). These two algorithms were subsequently extended using modifications including acceleration (Wang & Li, 2022), proximal steps (Pereyra, 2016), Riemannian metrics (Patterson & Teh, 2013), Hamiltonians (Betancourt, 2017), and projections (Wang & Li, 2022). ULA and MALA consider discretizing an SDE that corresponds to the Fokker-Planck equation. Many common zeroth and first order sampling methods, including ULA and MALA, rely on randomness that is independent of the samples to guarantee ergodicity of the Markov chains, typically modelled using white Gaussian noise. This randomness generates sufficient diffusion, which is then used show convergence (Mattingly et al., 2002; Meyn & Tweedie, 1994).

While diffusion can be achieved using random noises, we instead consider the third paradigm of achieving diffusion using the score of the density $\nabla \log \rho(x)$. Score based methods reformulate the Fokker-Planck equation into an ODE instead of an SDE, with the ODE depending on the gradient of the log-likelihood (the score) of the density (Maoutsa et al., 2020; Song et al., 2020; Del Moral, 2013). Some recent applications of score based diffusion include conditional generative modelling, utilizing the backwards Kolmogorov equation to diffuse from noise to natural images (Song et al., 2020; Batzolis et al., 2021). However, the score is not available, as it depends on the target density. Various methods have been proposed to approximate the score, including kernel density estimation (Carrillo et al., 2019; Terrell & Scott, 1992; Kim & Scott, 2012; Wand & Jones, 1994), adaptive kernel methods Van Kerm (2003); Botev et al. (2010), and neural ODEs (Bond-Taylor et al., 2021; Chen et al., 2018; Nijkamp et al., 2022). These approaches are generally non-parameteric, without making a priori assumptions on the target distribution. However, such approximations face common problems such as choice of kernel, mode collapse and sensitivity to hyper-parameters (Srivastava et al., 2017; Li et al., 2023a; Gramacki, 2018). We propose an alternative formulation of score approximation using the approximate Wasserstein proximal of the empirical measure, with a principled method of choosing hyper-parameters, that produces samples from a modified density that is close to the target density.

Utilizing Liouville's equation, we consider a score ODE to be solved in the particle space, whose density evolves according to the Fokker-Planck equation. The Jordan-Kinderlehrer-Otto (JKO) scheme considers a discretization of the Fokker-Planck ODE using proximal mappings in the Wasserstein space (Jordan et al., 1998). The target ODE is of the form

$$\frac{dX}{dt} = -\nabla V(X) - \beta \nabla \log \rho(t, X),$$

where $\rho(t)$ is the density of X_t at time t . The JKO scheme discretizes the ODE using Wasserstein proximal operators of the form

$$\rho_{k+1} = \arg \min_{\rho \in \mathcal{P}_2} \int_{\mathbb{R}^d} (\beta \rho \log \rho + V \rho) dx + \frac{1}{2h} \mathcal{W}(\rho_k, \rho)^2,$$

where k is the iteration of the update, $h > 0$ is the stepsize, $\mathcal{P}_2$ is the space of probability densities over $\mathbb{R}^d$ with bounded second moments, and $\mathcal{W}$ is the Wasserstein-2 distance between probability measures. The JKO scheme is the proximal iteration for the free energy functional with the Wasserstein-2 metric. However, the proximal map of the density is generally intractable and requires solving an equivalently difficult problem to our sampling problem. A recent work has considered using a regularized proximal term, formulated in terms of a set of coupled forward- and backward-heat equations (Li et al., 2023b). The score of the regularized Wasserstein proximal term has a *closed-form* solution based on convolutions with heat kernels. Motivated by this, we propose to utilize the closed-form solution for deterministic sampling.

In this work, we propose a deterministic sampling method based on the score flow. We then demonstrate stable convergence as well as convergence in the case of Gaussian densities, where we demonstrate a better dimension dependence bound due to the closed form solution in this case. Our proposed method is then compared with the unadjusted Langevin algorithm (ULA) as well as the Metropolis-adjusted Langevin algorithm (MALA), which are both stochastic methods. In the rest of this section, we introduce the Fokker-Planck equation, as well as the associated SDE and score ODE.

1.1. Definitions. We begin with some preliminary definitions, including the Wasserstein distance metric between probability measures, as well as the Fokker-Planck equation.

Definition 1. For two probability density functions μ, η on $\mathbb{R}^d$ with finite second moment, the Wasserstein-2 distance between μ and η is

$$\mathcal{W}(\mu, \eta) := \left(\inf_{\gamma \in \Gamma(\mu, \eta)} \iint_{\mathbb{R}^d \times \mathbb{R}^d} \|x - y\|^2 \gamma(x, y) dx dy \right)^{1/2},$$

where the norm is the Euclidean norm, and the infimum is taken over all couplings between μ, η , i.e. γ is a joint probability measure on $\mathbb{R}^d \times \mathbb{R}^d$ with

$$\int_{\mathbb{R}^d} \gamma(x, y) dy = \mu(x), \quad \int_{\mathbb{R}^d} \gamma(x, y) dx = \eta(y).$$

Let ρ_0 be a probability density function with finite second moment, and $V \in C^1(\mathbb{R}^d)$ be a bounded potential function. For a scalar $T > 0$, the Wasserstein proximal of ρ_0 is defined as

$$\rho_T = \text{WProx}_{TV}(\rho_0) := \arg \min_{q \in \mathcal{P}_2(\mathbb{R}^d)} \int_{\mathbb{R}^d} V(x)q(x)dx + \frac{\mathcal{W}(\rho_0, q)^2}{2T}, \quad (1)$$

where $\mathcal{W}(\rho_0, q)$ is the Wasserstein-2 distance between ρ_0 and q , and $\mathcal{P}_2$ is the set of probability density functions q with finite second moment.

The Wasserstein proximal does not admit an easily computable solution, and thus we consider an approximation to the Wasserstein proximal. Li et al. (2023b) consider an optimal control formulation based on the Benamou-Brenier formula (Benamou & Brenier, 2000). This reformulates the variational problem into a coupled ODE system. The regularized Wasserstein proximal operator is thus defined by the solution of the regularized PDEs

$$\begin{cases} \partial_t \rho(t, x) + \nabla_x \cdot (\rho(t, x) \nabla_x \Phi(t, x)) = \beta \Delta_x \rho(t, x) \end{cases} \quad (2a)$$

$$\begin{cases} \partial_t \Phi(t, x) + \frac{1}{2} \|\nabla_x \Phi(t, x)\|^2 = -\beta \Delta_x \Phi(t, x) \end{cases} \quad (2b)$$

$$\begin{cases} \rho(0, x) = \rho_0(x), \quad \Phi(T, x) = -V(x). \end{cases} \quad (2c)$$

These coupled ODEs arose from adding regularizing Lagrangian terms $\beta \Delta_x$ to the ODEs given by the Benamou-Brenier formula. Here, Φ is a Kantorovich dual variable that has boundary condition $-V$ at time T . $\rho(T, x)$ is called the *regularized Wasserstein proximal*. Using Hopf-Cole type transformations, Li et al. (2023b) show the following closed-form integral representation for the regularized Wasserstein proximal

$$\rho(T, x) = \int_{\mathbb{R}^d} K(x, y) \rho_0(y) dy, \quad (3)$$

$$K(x, y) = \frac{\exp(-\frac{1}{2\beta}(V(x) + \frac{\|x-y\|^2}{2T}))}{\int_{\mathbb{R}^d} \exp(-\frac{1}{2\beta}(V(z) + \frac{\|z-y\|^2}{2T})) dz}. \quad (4)$$

Observe that the normalizing constant in the kernel is given by a convolution between the potential V and a heat kernel. We note that the integral formulation can be extended to $\rho(t, x)$ and $\Phi(t, x)$ for more general time $t \in [0, T]$, again given by a convolution with a heat kernel.

We are interested in the solution of the Fokker-Planck equation

$$\frac{\partial \rho}{\partial t} = \nabla \cdot (\nabla V(x) \rho) + \beta \Delta \rho, \quad \rho(x, 0) = \rho_0(x). \quad (5)$$

We have the following relations between the Fokker-Planck equation and SDEs. More details can be found in Jordan et al. (1998) and in references therein. The solution $\rho(t, x)$ of the Fokker-Planck equation is equal to the density at time t of the SDE

$$dX(t) = -\nabla V(X(t))dt + \sqrt{2\beta}dW(t), \quad X(0) = X_0, \quad (6)$$

where X_0 is a random variable with density ρ_0 . Under appropriate growth conditions of V (such that the Gibbs measure is finite), the steady state of the Fokker-Planck equation is

$$\rho_\infty(x) \sim \exp(-\beta^{-1}V(x)). \quad (7)$$

Moreover, the Fokker-Planck equation can be viewed as a Wasserstein gradient flow on the free energy (Otto, 2001). Thus, this steady state is the minimizer of the free energy functional $\mathcal{E}$ over probability densities

$$\mathcal{E}(\rho) = \int_{\mathbb{R}^d} \beta \rho \log \rho + V \rho dx. \quad (8)$$

1.2. Score Based Diffusion. Instead of using a random particle formulation arising from a discretization of the SDE in Equation (6), we can use a deterministic version, given knowledge of the density ρ (which is intractable in practice). We now introduce the score-based model, where particles are updated according to the gradient of the potential, and the score function $\nabla_x \log \rho(t, x)$. This formulation arises from Liouville's equation, which states the following (Liouville, 1838; Kardar, 2007; Tolman, 1979).

Proposition 1 (Kubo, 1963). *For an evolution under a density $\rho(t, x)$ given by a Hamiltonian $\mathcal{K}$,*

$$\frac{dX}{dt} = \mathcal{K}(t, X, \rho),$$

the distribution function is constant along the trajectories. In particular, ρ satisfies

$$\frac{\partial \rho}{\partial t} + \nabla \cdot (\rho \mathcal{K}(t, x, \rho)) = 0.$$

We can use Liouville's equation to derive an ODE for the Fokker-Planck dynamics (5). Taking $\mathcal{K}(t, x, \rho) = -\nabla V(x) - \beta \frac{\nabla \rho(t, x)}{\rho(t, x)}$ with $\nabla \log \rho = \frac{\nabla \rho}{\rho}$, we obtain the following ODE, with density $\rho(t, x)$ at time t evolving as in the Fokker-Planck equation

$$\frac{dX}{dt} = -\nabla V(X) - \beta \nabla \log \rho(t, X). \quad (9)$$

If we instead consider the regularized Fokker-Planck equation (2a), this approximates the Fokker-Planck dynamics (5). Applying Liouville's equation with (2a) and $\mathcal{K}(t, x, \rho) = \nabla \Phi(t, x) - \beta \nabla \log \rho$ for time $t \in [0, T]$, we obtain the following particle evolution ODE, whose density at time t is equal to $\rho(t, x)$:

$$\frac{dX}{dt} = \nabla \Phi(t, X) - \beta \nabla \log \rho(t, X). \quad (10)$$

The main difference between this regularized formulation and the non-regularized Fokker-Planck is that V is replaced with the dual variable Φ , and this evolution is only valid for $t \in [0, T]$. Both terms of Equation (10) are problematic. Firstly, we do not have a closed form for $\Phi(t, x)$ for $t > T$ (though an integral formulation is available for $t < T$), and we only have the boundary condition $\Phi(T, x) = -V(x)$. Secondly, the score is not available. In the next section, we propose using the backwards Euler discretization method, utilizing only the boundary information for Φ and ρ , and thus only requiring V and the regularized Wasserstein proximal $\rho_T = \rho(T, x)$.

2. APPROXIMATING THE SCORE

In this section, we present the derivation and formulation of the proposed backwards regularized Wasserstein proximal (BRWP) scheme. Mixing time analysis is then given for the case where the target density is Gaussian, with closed-form updates for the mean and covariance. We characterize the discretization bias and demonstrate the convergence of the distribution to the regularized Wasserstein proximal of the target Gaussian distribution.

Our main goal is to approximately solve the ODE (9) numerically for particles X , using approximations given by (10). In this fashion, we are able to sample particles according to a distribution that evolves approximately according to the corresponding Fokker-Planck equation. The general idea is to consider the regularized Wasserstein proximal map as an approximation to the JKO scheme at each time step. We will demonstrate that the backwards Euler discretization of this approximate scheme is particularly amenable to computation. To begin, we consider the following four approximation steps.

Time approximation using the Wasserstein proximal. For a small time T , the approximate Wasserstein proximal dynamics (2) approximates the Fokker-Planck dynamics (5), where ρ_0 is replaced with $\rho(t, x)$. We thus approximate the Fokker-Planck dynamics by partitioning time into $[0, T]$, $[T, 2T]$, $[2T, 3T]$, ... for $t \geq 0$, and approximating each $[kT, (k+1)T]$ using Equation (2), and $\rho_{k,0}(x) = \rho(kT, x)$. We thus approximate the ODE (9) with (10) on each time partition. To compute this approximation, we can use the following techniques.

Backwards discretization in time. Since we only have particles at each iteration, analytic formulations of the densities $\rho_{k,0}$ and thus $\rho_{k,T}$ are unavailable. Instead of using kernel approximation

methods or otherwise to approximate the score at time $t = kT$, we instead compute *exactly* the score at time $t + T = (k + 1)T$, conditional on $\rho_{k,0}$ being a sum of Dirac masses at the locations of the corresponding particles. This allows for implicit time steps of the Fokker-Planck equation, assuming knowledge of $\rho_{k,T}$. We can compute $\rho_{k,T}$ when $\rho_{k,0}$ is given by an empirical distribution as follows.

Computing using the kernel formulation. For backwards Euler discretization, we need to know $\rho_{k,T}$ and $\Phi(T, \cdot)$ as evolved using Equation (2). $\rho_{k,T}$ is given in Equation (3) using a kernel convolution on $\rho_{k,0}$, and $\Phi(T, x) = -V(x)$ as defined in Equation (2a). We note that while it is possible to perform a forward discretization on the Φ term by computing $\Phi(0, x)$, it is not possible on the ρ term, as the score of a mixture of Dirac masses is undefined. Therefore, we apply a backwards Euler discretization of Equation (10).

Convolution as sampling. Observe the denominator in the convolution kernel given by Equation (4) takes the form of a Gaussian expectation. More precisely, this normalizing constant is given by a convolution of $\exp(-V)$ with a quadratic term. Using this trick similarly to Osher et al. (2023), we can compute the denominator of $K(x, y)$ by sampling from $z \sim \mathcal{N}(y, 2T\beta)$. Moreover, noting the normalizing constants for the Gaussians cancel out, this form means that we can compute integrals using Gaussian expectations. These can be computed using Monte Carlo integration for distributions f as follows.

$$\int_{\mathbb{R}^d} f(y) K(x, y) dy = \frac{\mathbb{E}_{y \sim \mathcal{N}(x, 2T\beta)} \left[f(y) \exp\left(-\frac{V(x)}{2\beta}\right) \right]}{\mathbb{E}_{z \sim \mathcal{N}(x, 2T\beta)} \left[\exp\left(-\frac{V(z)}{2\beta}\right) \right]}. \quad (11)$$

By combining these four approximation steps together, we obtain one step of the regularized Wasserstein proximal ODE Equation (10), discretized using the backwards Euler scheme. One discrete iteration with step-size $\eta > 0$ can be written as

$$\begin{aligned} X_{k+1} &= X_k + \eta [\nabla \Phi(T, X_k) - \beta \nabla \log \rho_{k,0}(T, X_k)] \\ &= X_k - \eta \nabla V(X_k) - \eta \beta \nabla \log \rho_{k,T}(X_k). \end{aligned} \quad (12)$$

To turn this into a discrete update scheme, we consider at each step setting $\rho_{k,0}$ to be the empirical distribution of X_k , rather than $\rho_{k,0} = \rho_{k-1,T}$. If we have N realizations of X_k given by $\{\mathbf{x}_{k,i}\}_{i=1}^N$, we approximate $\rho_{k,0}$ using the empirical distribution,

$$\rho_{k,0} = \frac{1}{N} \sum_{i=1}^N \delta_{\mathbf{x}_{k,i}}.$$

Noting that $\nabla \log \rho_{k,T}(x) = \nabla \rho_{k,T}(x) / \rho_{k,T}(x)$, and using the closed-form expression $\rho_{k,T}(x) = \int \rho_{k,0}(y) K(x, y) dy$, we have the following expression for $\rho_{k,T}$ and the gradient $\nabla \rho_{k,T}$ at a point $\mathbf{x}_i$, temporarily dropping the k subscript:

$$\rho_{k,T}(\mathbf{x}_i) = \frac{1}{N} \sum_{j=1}^N K(\mathbf{x}_i, \mathbf{x}_j) = \frac{1}{N} \sum_{j=1}^N \frac{\exp\left[-\frac{1}{2\beta} \left(V(\mathbf{x}_i) + \frac{\|\mathbf{x}_i - \mathbf{x}_j\|^2}{2T}\right)\right]}{\mathcal{Z}(\mathbf{x}_j)}, \quad (13a)$$

$$\nabla \rho_{k,T}(\mathbf{x}_i) = \frac{1}{N} \sum_{j=1}^N \frac{\left(-\frac{1}{2\beta} \left(\nabla V(\mathbf{x}_i) + \frac{\mathbf{x}_i - \mathbf{x}_j}{T}\right)\right) \exp\left[-\frac{1}{2\beta} \left(V(\mathbf{x}_i) + \frac{\|\mathbf{x}_i - \mathbf{x}_j\|^2}{2T}\right)\right]}{\mathcal{Z}(\mathbf{x}_j)}, \quad (13b)$$

$$\mathcal{Z}(\mathbf{x}_j) := \mathbb{E}_{z \sim \mathcal{N}(\mathbf{x}_j, 2T\beta)} \left[\exp\left(-\frac{V(z)}{2\beta}\right) \right]. \quad (13c)$$

This algorithm can be appropriately vectorized for parallelization. Indeed, as a kernel method, we need to compute the squared distances between all pairs of samples. This computational burden can be lessened by instead subsampling from the current samples to further approximate $\rho_{k,T}$. The Gaussian expectations can be done using Monte Carlo integration. The algorithm, consisting of Equations (12) and (13), is detailed in full in Algorithm 1. Note that the loops can be vectorized to improve run-time by replacing the intermediate variables $\mathcal{Z}, \mathcal{E}, \mathcal{V}$ with appropriately sized tensors.

A heuristic interpretation of the algorithm can be obtained by considering the score function as a weighted search direction. Indeed, $\log \rho_{k,T}$ is computed as a weighted sum of differences of

Algorithm 1 Backwards regularized Wasserstein proximal (BRWP) scheme

Input: Potential V , samples $(\mathbf{x}_{0,i})_{i=1}^N \sim \mu_0^{\otimes N}$, step-size $\eta > 0$, regularization parameters $T, \beta > 0$, Monte-Carlo sample count P

Output: Sequence of samples $(\mathbf{x}_{k,i})_{i=1}^N$ for $k = 1, 2, \dots$

```

1: for  $k \in \mathbb{N}$  do
2:   for  $i = 1, \dots, N$  do
3:     Sample  $(\mathbf{z}_{k,i,p})_{p=1}^P \sim \mathcal{N}(\mathbf{x}_{k,i}, 2\beta TI)$  ▷ Sample for the expectation
4:      $\mathcal{Z}_{k,i} = \frac{1}{P} \sum_{p=1}^P \exp\left(-\frac{V(\mathbf{z}_{k,i,p})}{2\beta}\right)$  ▷ Approximate  $\mathcal{Z}(\mathbf{x}_i)$  from (13c)
5:   end for
6:   for  $i, j = 1, \dots, N$  do
7:      $\mathcal{E}_{k,i,j} = \exp\left[-\frac{1}{2\beta}\left(V(\mathbf{x}_{k,i}) + \frac{\|\mathbf{x}_{k,i} - \mathbf{x}_{k,j}\|^2}{2T}\right)\right]$  ▷ Compute the numerator of (13a)
8:      $\mathcal{V}_{k,i,j} = -\frac{1}{2\beta}\left(\nabla V(\mathbf{x}_{k,i}) + \frac{\mathbf{x}_{k,i} - \mathbf{x}_{k,j}}{T}\right)$  ▷ Compute the numerator of (13b)
9:   end for
10:  for  $i = 1, \dots, N$  do
11:     $\nabla \log \rho_{k,T}(\mathbf{x}_{k,i}) = (\sum_j \mathcal{V}_{k,i,j} \mathcal{E}_{k,i,j} / \mathcal{Z}_{k,j}) / (\sum_j \mathcal{E}_{k,i,j} / \mathcal{Z}_{k,j})$  ▷ Compute the score
12:     $\mathbf{x}_{k+1,i} = \mathbf{x}_{k,i} - \eta \nabla V(\mathbf{x}_{k,i}) - \eta \beta \nabla \log \rho_{k,T}(\mathbf{x}_{k,i})$  ▷ Perform the update (12)
13:  end for
14: end for

```

$\mathcal{V}_{k,i,j}$, which contains a $-(\mathbf{x}_{k,i} - \mathbf{x}_{k,j})$ term in its expression. Considering the update Step 12 in Algorithm 1, the sample particle $\mathbf{x}_{k,i}$ is repelled away from a weighted sum of all the particles. This is the mechanism through which this method achieves diffusion.

2.1. Closed Form Gaussian Evolution. We begin our analysis with the simple case where V is quadratic. Moreover, we find closed forms for the distribution at iteration k , given that the initial distribution is also Gaussian. Consider first the Ornstein-Uhlenbeck process without drift in one dimension, which is a special case of the Fokker-Planck equation. The governing SDE for a constant $a > 0$ is as follows, where W is a Wiener process (Karatzas & Shreve, 1991; Gardiner et al., 1985):

$$dX = -aXdt + \sqrt{2\beta}dW. \quad (14)$$

This can be seen as taking the potential to be $V(x) = ax^2/2$. The true solution for initialization X_0 is given by

$$X_t = X_0 e^{-at} + \frac{\sqrt{2\beta}}{\sqrt{2a}} W_{1-e^{-2at}}. \quad (15)$$

If X_0 is initially normally distributed with mean μ_0 , variance σ_0^2 , then the distribution at X_t will also be normally distributed, with means μ_t and variance σ_t^2 given by

$$\mu_t = \mu_0 e^{-at}, \quad \sigma_t^2 = \sigma_0^2 e^{-2at} + \frac{\beta}{a}(1 - e^{-2at}).$$

The steady state of the flow is Gaussian with mean and variance

$$\mu_\infty = 0, \quad \sigma_\infty^2 = \frac{\beta}{a}.$$

To discretize this flow, we consider two competing methods, ULA and MALA. We can compute the analytic solutions with quadratic potential $V = ax^2/2$ and Gaussian distributed initializations $X_0 \sim \mathcal{N}(\mu_0, \sigma_0^2)$.

ULA. For a step-size $\eta < a^{-1}$, ULA consists of an explicit Euler-Maruyama discretization of Equation (14):

$$X_{k+1} = (1 - a\eta)X_k + \sqrt{2\beta\eta}Z_k,$$

where $(Z_k)_{k=0}^\infty$ are i.i.d standard Gaussians. Therefore, X_t are also Gaussian, with mean and variance satisfying the recurrence relations

$$\mu_{k+1} = (1 - a\eta)\mu_k, \quad \sigma_{k+1}^2 = (1 - a\eta)^2 \sigma_k^2 + 2\beta\eta.$$

Solving the recurrence relations gives the closed form solutions

$$\mu_k = (1 - a\eta)^k \mu_0, \quad \sigma_k^2 = (1 - a\eta)^{2k} \sigma_0^2 + 2\beta\eta \sum_{j=0}^{k-1} (1 - a\eta)^{2j}.$$

Observe that the variance is biased due to the explicit discretization (Wibisono, 2018):

$$\lim_{k \rightarrow \infty} \sigma_k^2 = \frac{2\beta}{(2 - a\eta)a} > \frac{\beta}{a} = \sigma_\infty^2.$$

MALA. The Metropolis-adjusted Langevin algorithm introduces an additional Metropolis-Hastings acceptance step after ULA (Dwivedi et al., 2018; Roberts & Tweedie, 1996). The MALA update is as follows in the case where $\beta = 1$.

$$\begin{aligned} \tilde{X}_{k+1} &= (1 - \eta \nabla V) X_k + \sqrt{2\eta} Z_k; \\ \alpha_{k+1} &= \min \left\{ 1, \frac{\exp \left(-V(\tilde{X}_{k+1}) - \|X_k - \tilde{X}_{k+1} + \eta \nabla V(\tilde{X}_{k+1})\|^2 / 4\eta \right)}{\exp \left(-V(X_k) - \|\tilde{X}_{k+1} - X_k + \eta \nabla V(X_k)\|^2 / 4\eta \right)} \right\}; \\ X_{k+1} &= \begin{cases} \tilde{X}_{k+1}, & \text{with probability } \alpha_{k+1}; \\ X_k, & \text{with probability } 1 - \alpha_{k+1}. \end{cases} \end{aligned}$$

In the case that $\beta \neq 1$, we can perform a change of variables by considering step-size $\beta\eta$ and potential V/β . Then the MALA scheme will have modified acceptance probabilities of the form

$$\begin{aligned} \tilde{X}_{k+1} &= (1 - \eta \nabla V) X_k + \sqrt{2\beta\eta} Z_k; \\ \alpha_{k+1} &= \min \left\{ 1, \frac{\exp \left(-V(\tilde{X}_{k+1})/\beta - \|X_k - \tilde{X}_{k+1} + \eta \nabla V(\tilde{X}_{k+1})\|^2 / (4\beta\eta) \right)}{\exp \left(-V(X_k)/\beta - \|\tilde{X}_{k+1} - X_k + \eta \nabla V(X_k)\|^2 / (4\beta\eta) \right)} \right\}; \\ X_{k+1} &= \begin{cases} \tilde{X}_{k+1}, & \text{with probability } \alpha_{k+1}; \\ X_k, & \text{with probability } 1 - \alpha_{k+1}. \end{cases} \end{aligned}$$

We note that X_k does not follow a Gaussian distribution due to this acceptance step. MALA is unbiased, and converges in distribution to the target Gaussian $\mathcal{N}(\mu_\infty, \sigma_\infty^2)$.

BRWP. Assuming $X_k \sim \mathcal{N}(\mu_k, \sigma_k^2)$, we can compute the closed form of $\rho_{k,T}(x)$ with initial condition $\rho_{k,0} \sim \mathcal{N}(\mu_k, \sigma_k^2)$ using the kernel formulation. A full derivation can be found in Appendix A. The approximate Wasserstein proximal $\rho_{k,T} \sim \mathcal{N}(\tilde{\mu}_{k+1}, \tilde{\sigma}_{k+1}^2)$ is Gaussian, with mean and variance

$$\tilde{\mu}_{k+1} = \frac{\mu_k}{1 + aT}, \quad \tilde{\sigma}_{k+1}^2 = \frac{\sigma_k^2}{(1 + aT)^2} + \frac{2\beta T}{1 + aT}. \quad (16)$$

Applying the discrete backwards iteration given in Equation (12) with this closed form for $\rho_{k,T}$, we have

$$\begin{aligned} X_{k+1} &= (1 - a\eta) X_k - \eta\beta \nabla \log \rho_{k,T}(X_k) \\ &= (1 - a\eta) X_k + \eta\beta \frac{X_k - \tilde{\mu}_{k+1}}{\tilde{\sigma}_{k+1}^2} \\ &= (1 - a\eta + \frac{\eta\beta}{\tilde{\sigma}_{k+1}^2}) X_k - \frac{\eta\beta \tilde{\mu}_{k+1}}{\tilde{\sigma}_{k+1}^2} \\ &= \left(1 - a\eta + \frac{\eta\beta(1 + aT)^2}{\sigma_k^2 + 2\beta T(1 + aT)} \right) X_k - \frac{\eta\beta \mu_k(1 + aT)}{\sigma_k^2 + 2\beta T(1 + aT)}. \end{aligned}$$

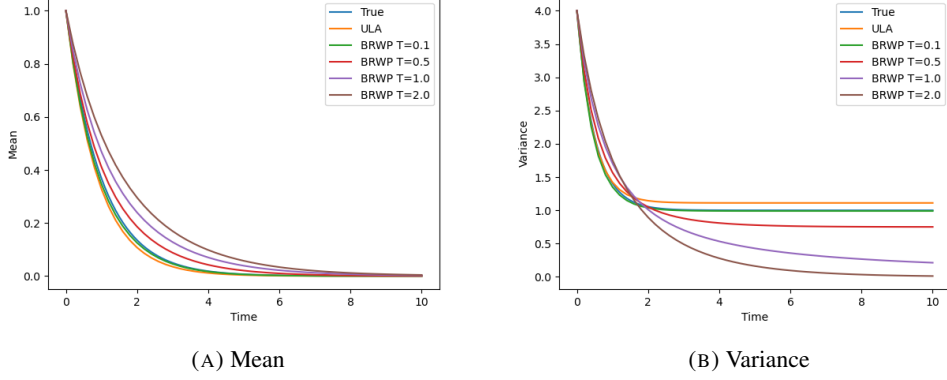


FIGURE 1. Evolution of the analytic solutions for ULA and kernel formula, with initialization $\mathcal{N}(0, 4)$ and target distribution $\mathcal{N}(0, 1)$. The parameters of the OU flow are $a = \beta = 1$, discretized with step-size $\eta = 0.25$. The larger T is, the smaller the stationary variance is. We observe that when $T \geq 1$, the stationary distribution under the Wasserstein proximal flow is degenerate.

Therefore, X_{k+1} is Gaussian, with mean and variance satisfying the recurrence relations

$$\begin{aligned} \mu_{k+1} &= \left(1 - a\eta + \frac{\eta\beta(1+aT)^2}{\sigma_k^2 + 2\beta T(1+aT)}\right) \mu_k - \frac{\eta\beta\mu_k(1+aT)}{\sigma_k^2 + 2\beta T(1+aT)} \\ &= \left(1 - a\eta + \frac{\eta\beta aT(1+aT)}{\sigma_k^2 + 2\beta T(1+aT)}\right) \mu_k, \end{aligned} \quad (17a)$$

$$\sigma_{k+1}^2 = \left(1 - a\eta + \frac{\eta\beta(1+aT)^2}{\sigma_k^2 + 2\beta T(1+aT)}\right)^2 \sigma_k^2. \quad (17b)$$

We can compute the steady states of Equation (17) by setting the front term in Equation (17b) to 1. We can do this by assuming $a\eta < 2$. This results in

$$\begin{aligned} \left(1 - a\eta + \frac{\eta\beta(1+aT)^2}{\sigma_\infty^2 + 2\beta T(1+aT)}\right) &= 1 \\ \Rightarrow \sigma_\infty^2 &= \frac{\beta}{a}(1 - a^2T^2) \text{ if } aT < 1, \sigma_\infty^2 = 0 \text{ otherwise.} \end{aligned}$$

We observe that the bias is different to ULA due to the backwards discretization. Indeed, the bias of ULA results in a variance that is larger than the target variance. On the other hand, the bias for BRWP results in a variance that is smaller than the target variance, and moreover does not depend on the step-size η .

2.2. Multi-dimensional Gaussian. With some care, we can extend the analysis of our previous section to the multi-dimensional case, and again obtain closed form expressions for the mean and variance at iteration k . Suppose now that we are working in $\mathbb{R}^d$, and our V takes the form for a zero-mean Gaussian with (symmetric positive-definite) covariance Σ^{-1}

$$V(x) = \frac{1}{2}x^\top \Sigma^{-1}x. \quad (18)$$

As before, we can obtain a closed form for the approximate Wasserstein proximal, and the derivation can be found in Appendix A. We have that $\rho_{k,T} \sim \mathcal{N}(\tilde{\mu}_{k+1}, \tilde{\Sigma}_{k+1})$, with mean and covariance

$$\tilde{\Sigma}_{k+1}^{-1} = (2\beta T(I + T\Sigma^{-1})^{-1} + (I + T\Sigma^{-1})^{-1}\Sigma_k(I + T\Sigma^{-1})^{-1})^{-1}, \quad (19a)$$

$$\tilde{\mu}_{k+1} = (I + T\Sigma^{-1})^{-1}\mu_k. \quad (19b)$$

Applying the discrete backwards iteration Equation (12),

$$\begin{aligned} X_{k+1} &= X_k - \eta \nabla V(X_k) - \eta \beta \nabla \log \rho_{k,T}(X_k) \\ &= (I - \eta \Sigma^{-1})X_k + \eta \beta \tilde{\Sigma}_{k+1}^{-1}(X_k - \tilde{\mu}_{k+1}) \\ &= (I - \eta \Sigma^{-1} + \eta \beta \tilde{\Sigma}_{k+1}^{-1})X_k - \eta \beta \tilde{\Sigma}_{k+1}^{-1} \tilde{\mu}_{k+1}. \end{aligned}$$

Since X_k is Gaussian and affine transformations of Gaussian distributions are Gaussian, we can obtain the following recurrence relations for the parameters of X_{k+1} .

Proposition 2. X_{k+1} is Gaussian with mean μ_{k+1} and covariance Σ_{k+1} given by

$$\begin{aligned} \mu_{k+1} &= (I - \eta \Sigma^{-1})\mu_k + (\eta \beta \tilde{\Sigma}_{k+1}^{-1})(\mu_k - \tilde{\mu}_{k+1}) \\ &= \left(I - \eta \Sigma^{-1} + \eta \beta (2\beta T I + \Sigma_k(I + T \Sigma^{-1})^{-1})^{-1} (T \Sigma^{-1}) \right) \mu_k, \end{aligned} \quad (20a)$$

$$\Sigma_{k+1} = (I - \eta \Sigma^{-1} + \eta \beta \tilde{\Sigma}_{k+1}^{-1})\Sigma_k(I - \eta \Sigma^{-1} + \eta \beta \tilde{\Sigma}_{k+1}^{-1})^\top. \quad (20b)$$

We obtain the covariance of the stationary distribution by setting $I - \eta \Sigma^{-1} + \eta \beta \tilde{\Sigma}_{k+1}^{-1} = I$, yielding

$$\Sigma_\infty = \beta(I - T \Sigma^{-1})\Sigma(I + T \Sigma^{-1}). \quad (21)$$

2.2.1. Mixing Time: Gaussian. We first consider the case where the mean of the initialization is zero, and further that the covariance of the initialization commutes with the covariance of the target distribution. Observe that if Σ commutes with Σ_k , then Σ commutes with $\tilde{\Sigma}_{k+1}$ and hence with Σ_{k+1} . Therefore, without loss of generality, we can work in a simultaneously diagonal basis for Σ , and all Σ_k and $\tilde{\Sigma}_k$. We show linear convergence of the eigenvalues of the covariance matrix to those of the stationary distribution, and give the rate of convergence in terms of T .

We aim to bound the mixing time of the Gaussians under the approximate Wasserstein iterations, defined as follows. This is a measure of how quickly a sequence of distributions converges to a target distribution.

Definition 2 (Mixing time). *The total variation between two probability distributions over a measurable space $(\Omega, \mathcal{F})$ is*

$$\text{TV}(\mathcal{P}, \mathcal{Q}) = \sup_{A \in \mathcal{F}} |\mathcal{P}(A) - \mathcal{Q}(A)|.$$

For a operator $\mathcal{T}_p$ on the space of probability distributions, assume that the chain $\mathcal{T}_p^k(\mu_0) \rightarrow \Pi$ as $k \rightarrow \infty$ for some probability distribution Π . The δ -mixing time with $\delta \in (0, 1)$ and initial distribution μ_0 is

$$t_{\text{mix}}(\delta; \mu_0) = \min \{k \mid \text{TV}(\mathcal{T}_p^k(\mu_0), \Pi) \leq \delta\}.$$

We have the following theorem upper-bounding the total variation between two Gaussians with the same mean. This means that we can control the total variation between two Gaussians with the difference between the covariance matrices.

Theorem 1 (Devroye et al., 2018, Thm. 1.1). *Let $\mu \in \mathbb{R}^d$, Σ_1, Σ_2 be two positive-definite $d \times d$ covariance matrices, and $\lambda_1, \dots, \lambda_d$ denote the eigenvalues of $\Sigma_1^{-1}\Sigma_2 - I$. Then the total variation satisfies*

$$\text{TV}(\mathcal{N}(\mu, \Sigma_1), \mathcal{N}(\mu, \Sigma_2)) \leq \frac{3}{2} \min \left\{ 1, \sqrt{\sum_{i=1}^d \lambda_i^2} \right\}. \quad (22)$$

We firstly assume that $\mu_0 = \mathbf{0}$. Observe that this means that $\mu_k = \mathbf{0}$ for all k . Suppose further that Σ_0 commutes with Σ , for example if $\Sigma_0 = c_0 I$ for some $c_0 > 0$. Under this assumption, we can simultaneously diagonalize Σ_0 and Σ (and indeed, Σ_k as well for all k). Without loss of generality, let us work in an orthonormal eigenbasis, so that $\Sigma = \text{diag}(\xi^{(1)}, \dots, \xi^{(d)})$ and $\Sigma_k = \text{diag}(\tau_k^{(1)}, \dots, \tau_k^{(d)})$ are all diagonal. The following theorem states that the covariance under the BRWP iterations converge linearly to the stationary distribution.

Theorem 2 (Mixing time for multi-dimensional Gaussians with same initialization mean). *Consider the regularized Wasserstein proximal scheme applied to the zero-mean Ornstein-Uhlenbeck process in d dimensions*

$$dX = -\nabla V(X)dt + \sqrt{2\beta}dW, \quad V(x) = \frac{1}{2}x^\top \Sigma^{-1}x,$$

where $\Sigma = \text{diag}(\xi^{(1)}, \dots, \xi^{(d)})$ is positive definite. Let $T > 0, \eta > 0$ be such that $T < \min\{\xi^{(i)} \mid i = 1, \dots, d\}$, and $\eta\xi^{(i)-1} \leq 1/\Delta(\xi^{(i)}, T)$ for all i , where

$$\Delta(\xi^{(i)}, T) = \frac{1}{2} \left(\sqrt{\frac{\xi^{(i)} + T}{2T}} + 1 \right).$$

The stationary distribution Π of the discrete scheme Equation (12) is given by

$$\Pi \sim \mathcal{N}(0, \Sigma_\infty), \quad \Sigma_\infty = \text{diag}(\tau_\infty^{(i)} \mid i = 1, \dots, d) \quad (23)$$

$$\tau_\infty^{(i)} = \beta\xi^{(i)}(1 - T^2\xi^{(i)-2}), \quad i = 1, \dots, d. \quad (24)$$

Suppose the chain is initialized with

$$X_0 \sim \mathcal{N}(0, \Sigma_0), \quad \Sigma_0 = \text{diag}(\tau_0^{(i)} \mid i = 1, \dots, d),$$

with $\tau_0^{(i)} > 0$ for $i = 1, \dots, d$. If $(X_k)_{k \geq 0}$ evolves under the BRWP scheme Equation (12), we have the following closed-form for the distributions of X_k :

$$X_k \sim \mathcal{N}(0, \Sigma_k), \quad \Sigma_k = \text{diag}(\tau_k^{(i)} \mid i = 1, \dots, d),$$

$$\tau_{k+1}^{(i)} = \left(1 - \eta\xi^{(i)-1} + \frac{\eta\beta(1 + T\xi^{(i)-1})^2}{\tau_k^{(i)} + 2\beta T(1 + T\xi^{(i)-1})} \right)^2 \tau_k^{(i)}, \quad i = 1, \dots, d. \quad (25)$$

In particular, the eigenvalues of $\Sigma_k \Sigma_\infty^{-1} - I$ are given by

$$\lambda_k^{(i)} = [\tau_k^{(i)} - \tau_\infty^{(i)}] / \tau_\infty^{(i)}, \quad (26)$$

which converge linearly to 0 with rate of convergence $[1 - \eta\xi^{(i)-1}(\xi - T)/(\xi + T)] \in (0, 1)$. Moreover, the total variation satisfies

$$\text{TV}(\mu(X_k), \Pi_T) \leq \frac{3}{2} \min \left\{ 1, \sqrt{\sum_{i=1}^d (\lambda_k^{(i)})^2} \right\} \leq \frac{3}{2} C \sqrt{d} c^k, \quad (27)$$

where $C = C(\Sigma_0, \Sigma, T) > 0$ is the root mean squared of the initial eigenvalues of $\Sigma_0 \Sigma_T^{-1} - I$, and $c = c(\Sigma_0, \Sigma, T, \beta, \eta) \in (0, 1)$ is the largest rate of convergence. Therefore the mixing time satisfies

$$t_{\text{mix}}(\delta, \mu(X_0)) = \mathcal{O}(\log(C\sqrt{d}/\delta)/\log(c)). \quad (28)$$

Remark 1. The conditions mean that we want T to be small to reduce the asymptotic bias, but also sufficiently large so that we can take a large step-size η .

Sketch proof. Without loss of generality, we can assume all covariance matrices are diagonal. Using the closed-form update Equation (17b) for the variance, we derive a recurrence relation for the eigenvalues of the covariance matrices. This recurrence relation converges to the stationary distribution linearly, provided that the step-size is chosen to be sufficiently small. A full proof can be found in Appendix B. $\square$

For fixed $T < \min_i \xi^{(i)}$, let us compute c in terms of the condition number $\kappa = L/m$ for initializations $\Sigma_0 = L^{-1}(1 - L^{-2}T^2)^{-1}I$, where m and L are the smallest and largest eigenvalues of $\nabla^2 V = \Sigma^{-1}$ respectively. Note that $L(1 - L^{-2}T^2) = \lambda_{\max}(\Sigma_\infty)$ is the Lipschitz constant of the log-Hessian of the stationary density. Without loss of generality, let the eigenvalues of Σ be sorted in descending order, so that $L^{-1} = \xi^{(d)}, m^{-1} = \xi^{(1)}$. In this case, we have that $\omega^{(i)}(\gamma_k^{(i)}) \geq 0$ for each $i = 1, \dots, d$ and all $k \geq 0$. Thus

$$\delta^{(i)} = \min(\omega^{(i)}(0), \omega^{(i)}(\gamma_0^{(i)})) \geq \min \left(\frac{2(\xi^{(i)} - T)}{\xi^{(i)} + T}, 1 \right) \geq \min \left(\frac{2(\xi^{(d)} - T)}{\xi^{(d)} + T}, 1 \right).$$

Let $\eta = \min_i \{\xi^{(i)} / \Delta^{(i)}\} = \xi^{(d)} / \Delta^{(d)}$ be the maximum allowed step-size. Then,

$$\begin{aligned} c &= \max_i \left\{ 1 - \eta \xi^{(i)-1} \delta^{(i)} \right\} \\ &= \max_i \left\{ 1 - \eta \xi^{(i)-1} \delta^{(i)} \right\} \\ &\leq 1 - (\xi^{(d)} / \Delta^{(d)}) \xi^{(1)-1} \min \left(\frac{2(\xi^{(d)} - T)}{\xi^{(d)} + T}, 1 \right) \\ &= 1 - \kappa^{-1} \min \left(\frac{2(\xi^{(d)} - T)}{\xi^{(d)} + T}, 1 \right) / \left(\frac{1}{2} (\sqrt{(\xi^{(d)} + T)/(2T)} + 1) \right). \end{aligned} \quad (29)$$

The choice of T here is not particularly important as long as it is smaller than $\xi^{(d)}$. Taking $T = \xi^{(d)}/3$, we get

$$c \leq 1 - 1 / \left(\frac{\kappa}{2} (\sqrt{3\kappa/2 + 1/2} + 1) \right),$$

so that $-1/\log c = \mathcal{O}(\kappa^{3/2})$. Further note that for this choice of T , $C = 3(\kappa - 1)/2$. We get the following.

Corollary 1. *For initialization $X_0 \sim \mathcal{N}(0, L^{-1}(1 - L^{-2}T^2)^{-1}I)$, where $mI \preceq \nabla^2 V \preceq LI$, and $\kappa = L/m$. Let $T = \xi^{(d)}/3$ and $\eta = \xi^{(1)}/\Delta^{(1)}$. The worst-case mixing time satisfies*

$$t_{\text{mix}}(\delta, \mu(X_0)) = \mathcal{O}(\kappa^{3/2} \log(\kappa\sqrt{d}/\delta)). \quad (30)$$

As a comparison, we have that the mixing times with initialization $\mathcal{N}(0, L^{-1}I)$ for ULA is $\mathcal{O}((d^3 + d \log^2(1/\delta))\kappa^2\delta^{-2})$, and the mixing time for MALA is $\mathcal{O}(d^2\kappa \log(\kappa/\delta))$ (Dalalyan, 2017; Dwivedi et al., 2018). We note that in the analytic case, our mixing time has a small dependence on the dimension, which comes only from translating convergence of the eigenvalues to convergence of the total variation distance.

2.3. Non-commuting Gaussian. We now turn our attention to the case where the initialization Σ_0 does not commute with the target covariance Σ . We stay in the zero mean case. By considering the continuous limit of the BRWP updates, we show that the regularized Wasserstein proximals of the covariances converges to the regularized Wasserstein proximal of the target covariance in terms of Frobenius distance, $\|\tilde{\Sigma}_t^{-1} - \tilde{\Sigma}_\infty^{-1}\|_F^2 \rightarrow 0$. For ease of notation, let us first define $K := I + T\Sigma^{-1}$. Note that K commutes with Σ . We first recall the identities:

$$\tilde{\Sigma}_{k+1} = K^{-1}\Sigma_k K^{-1} + 2\beta T K^{-1}, \quad (31a)$$

$$\Sigma_k = K\tilde{\Sigma}_{k+1}K - 2\beta T K, \quad (31b)$$

$$\Sigma_{k+1} = (I - \eta\Sigma^{-1} + \eta\beta\tilde{\Sigma}_{k+1}^{-1})\Sigma_k(I - \eta\Sigma^{-1} + \eta\beta\tilde{\Sigma}_{k+1}^{-1}). \quad (31c)$$

Moreover, recall that the regularized Wasserstein proximal of the target distribution is $\tilde{\Sigma}_\infty = \beta\Sigma$. Reformulating in terms of $\tilde{\Sigma}$,

$$\begin{aligned} \tilde{\Sigma}_{k+2} &= K^{-1}\Sigma_{k+1}K^{-1} + 2\beta T K^{-1} \\ &= 2\beta T K^{-1} + K^{-1} \left[I - \eta\beta(\tilde{\Sigma}_\infty^{-1} - \tilde{\Sigma}_{k+1}^{-1}) \right] \Sigma_k \left[I - \eta\beta(\tilde{\Sigma}_\infty^{-1} - \tilde{\Sigma}_{k+1}^{-1}) \right] K^{-1} \\ &= 2\beta T K^{-1} + K^{-1}\Sigma_k K^{-1} \\ &\quad - K^{-1}\eta\beta(\tilde{\Sigma}_\infty^{-1} - \tilde{\Sigma}_{k+1}^{-1})\Sigma_k K^{-1} \\ &\quad - K^{-1}\Sigma_k\eta\beta(\tilde{\Sigma}_\infty^{-1} - \tilde{\Sigma}_{k+1}^{-1})K^{-1} \\ &\quad + K^{-1}\eta\beta(\tilde{\Sigma}_\infty^{-1} - \tilde{\Sigma}_{k+1}^{-1})\Sigma_k\eta\beta(\tilde{\Sigma}_\infty^{-1} - \tilde{\Sigma}_{k+1}^{-1})K^{-1} \\ &= \tilde{\Sigma}_{k+1} - \left[\eta\beta K^{-1}(\tilde{\Sigma}_\infty^{-1} - \tilde{\Sigma}_{k+1}^{-1})K \right] \left[K^{-1}\Sigma_k K^{-1} \right] \\ &\quad - \left[K^{-1}\Sigma_k K^{-1} \right] \left[\eta\beta K(\tilde{\Sigma}_\infty^{-1} - \tilde{\Sigma}_{k+1}^{-1})K^{-1} \right] \\ &\quad + \left[\eta\beta K^{-1}(\tilde{\Sigma}_\infty^{-1} - \tilde{\Sigma}_{k+1}^{-1})K \right] \left[K^{-1}\Sigma_k K^{-1} \right] \left[\eta\beta K(\tilde{\Sigma}_\infty^{-1} - \tilde{\Sigma}_{k+1}^{-1})K^{-1} \right]. \end{aligned}$$

We see that this is a discretization of the continuous case by discarding the higher order η terms, defined as follows,

$$d\tilde{\Sigma}_t/dt = -\beta K^{-1} \left[(\tilde{\Sigma}_\infty^{-1} - \tilde{\Sigma}_t^{-1})\Sigma_t + \Sigma_t(\tilde{\Sigma}_\infty^{-1} - \tilde{\Sigma}_t^{-1}) \right] K^{-1}, \quad (32a)$$

$$\Sigma_t = K\tilde{\Sigma}_t K - 2\beta TK. \quad (32b)$$

The above discrete iteration 31 for $\tilde{\Sigma}_k$ is a discretization of the above ODE 32 in the limit as $\eta \rightarrow 0$. We find that the Frobenius norm $\|\tilde{\Sigma}_\infty^{-1} - \tilde{\Sigma}_t^{-1}\|_F^2$ is a Lyapunov function for the ODE formulation of the BRWP scheme.

Proposition 3. *The squared Frobenius norm $\|\tilde{\Sigma}_\infty^{-1} - \tilde{\Sigma}_t^{-1}\|_F^2$ is a Lyapunov function for the continuous limit of the BRWP scheme. Moreover, it converges linearly to zero.*

Sketch proof. The time derivative of the Frobenius norm is given by the trace of the product of a positive definite matrix, and a matrix whose spectrum lies in the positive half line. Using the generalized Hölder's inequality for matrices, we upper bound the time derivative by a negative quantity that is proportional to the squared eigenvalues of $\tilde{\Sigma}_\infty^{-1} - \tilde{\Sigma}_t^{-1}$. A full proof can be found in Appendix C $\square$

In this section, we used the closed form solution for the BRWP scheme to compute the evolution for the Ornstein-Uhlenbeck process. Expressions for the stationary solution and iterations were computed, and linear convergence to the stationary solutions were shown, with better dimension dependence on the mixing time compared to ULA and MALA.

3. EXPERIMENTS

For numerical experiments, we compare our method against ULA and MALA, using the experiments in Dwivedi et al. (2018); Wang & Li (2022). In particular, we consider target densities from an ill-conditioned Gaussian, a Gaussian mixture, a bimodal toy distribution, and Bayesian logistic regression. We will use this to demonstrate convergence to the (approximate) stationary distribution, as well as effectiveness without the requirement of pre-conditioning. Moreover, we will demonstrate the effect of using an ODE to model the particle movement instead of discretizing an SDE, in that the samples do not evolve significantly after some time. We compute ULA and MALA using the algorithms defined in Section 2, and fix $\beta = 1$ for simplicity.

3.1. Ill-Conditioned Gaussian. We first consider the case of a 2-dimensional and 5-dimensional Gaussian, with mean zero and diagonal covariance with eigenvalues evenly spaced from 10 to 1. The corresponding potential $V = x^\top \Sigma^{-1}x/2$ has Lipschitz constant $L = 1$ and strong convexity parameter $m = 0.1$. We consider the step-sizes to be $\eta = 0.1$ for ULA, MALA and the proposed BRWP scheme. For the BRWP scheme, we consider the choices $T = 0.05, 0.1, 0.25, 0.5$. Note that the theory restricts $T < \lambda_{\min}(\Sigma) = 1$, so these choices of T are valid and do not produce degenerate Gaussians for a closed form evolution. The number of Monte-Carlo samples used for computing the normalizing constant was set to $P = 10$. We present three experiments, with dimension d and number of samples N as $(d, N) = (2, 1000), (5, 1000), (5, 200)$, with samples initialized as $\mathcal{N}(0, L^{-1}I) = \mathcal{N}(0, I)$. We present two main findings, that we demonstrate further in following experiments.

Samples are structured. Figure 2 demonstrates the effect of the deterministic sampling. In two dimensions, we observe a clear ellipsoidal structure that is traced out by the iterates, closely matched by the level set contours of the density $\exp(-V)$. This appears to be a consequence of both determinism as well as evolving an empirical approximation to the density at each iteration.

Variance reduction/mode collapse phenomenon. Figure 3 considers a 5-dimensional Gaussian, projected onto the first and last dimensions with target covariance 10 and 1 respectively, using $N = 1000$ and $N = 200$ samples. In the case of sufficiently many samples $N = 1000$, we observe the same structural phenomenon as in Figure 2. However, in the case where $N = 200$, we observe a sample clustering phenomenon. For small values of T , the samples cluster more strongly around the true minimizer of V , which is the origin. For larger values of T , we observe that this clustering phenomenon is weaker, but there is bias due to the approximation as suggested in Section 2.

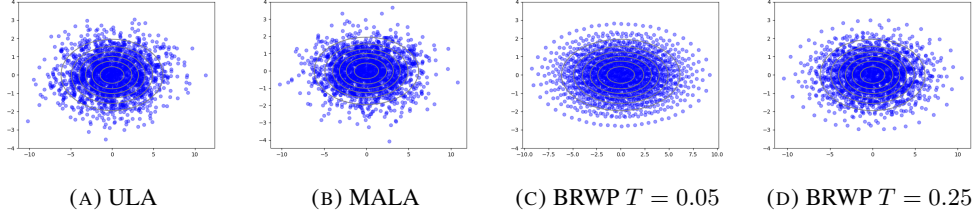


FIGURE 2. 1000 particles after 200 iterations of ULA, MALA and BRWP with $T = 0.05, 0.25$, applied to a two-dimensional Gaussian with condition number $\kappa = 10$. We observe the samples for BRWP organize themselves into rings, clearer for $T = 0.05$ than for $T = 0.25$. This is as opposed to the randomness of ULA and MALA.

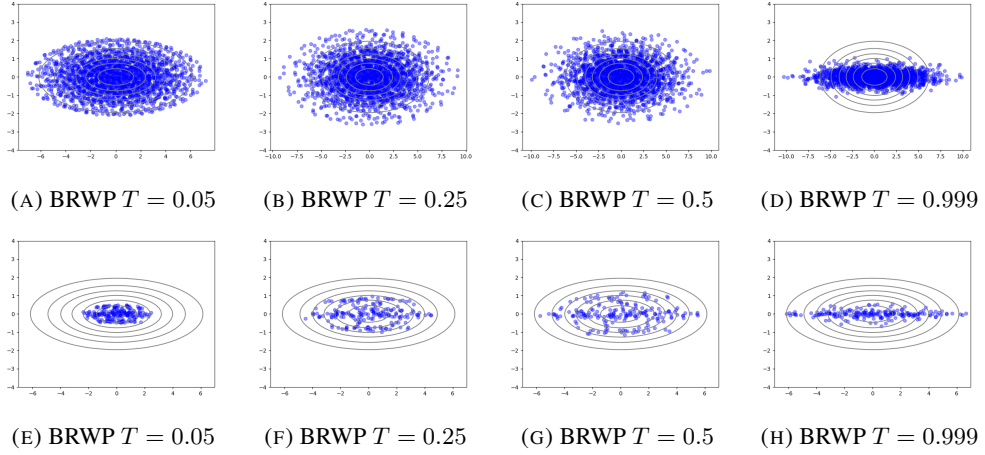


FIGURE 3. Particles after 200 iterations of BRWP with $T = 0.05, 0.25, 0.5, 0.999$, applied to a two-dimensional Gaussian with condition number $\kappa = 10$. Figures (a-d) in the top row have $N = 1000$ samples, while figures (e-h) in the bottom row have $N = 200$ samples. We observe that in higher dimensions, having fewer samples results in partial mode collapse as $T \rightarrow 0$. Moreover, for T close to 1, the variance in the vertical direction of $\lambda(\Sigma) = 1$ is reduced, demonstrating the bias of BRWP as described in Section 2.

The variance reduction suggests that the error incurred by approximating the distribution after each forward iteration by the empirical measure plays an effect in the convergence behavior. In the case where T is small, this can be partially explained by the quadratic term dominating V in the score formulation.

3.2. Gaussian Mixture. To further illustrate the structure phenomenon, we can also use a mixture of Gaussians. Using the experiment setup in (Dwivedi et al., 2018), we consider sampling from the target density, given by a mixture of Gaussians $\mathcal{N}(a, I)$ and $\mathcal{N}(-a, I)$:

$$p(x) = \frac{1}{2(2\pi)^{d/2}} \left(e^{-\|x-a\|_2^2/2} + e^{-\|x+a\|_2^2/2} \right).$$

The corresponding potential is given by

$$V(x) = \frac{1}{2} \|x - a\|_2^2 - \log \left(1 + e^{-2x^\top a} \right), \quad (33a)$$

$$\nabla V(x) = x - a + 2a(1 + e^{2x^\top a})^{-1}. \quad (33b)$$

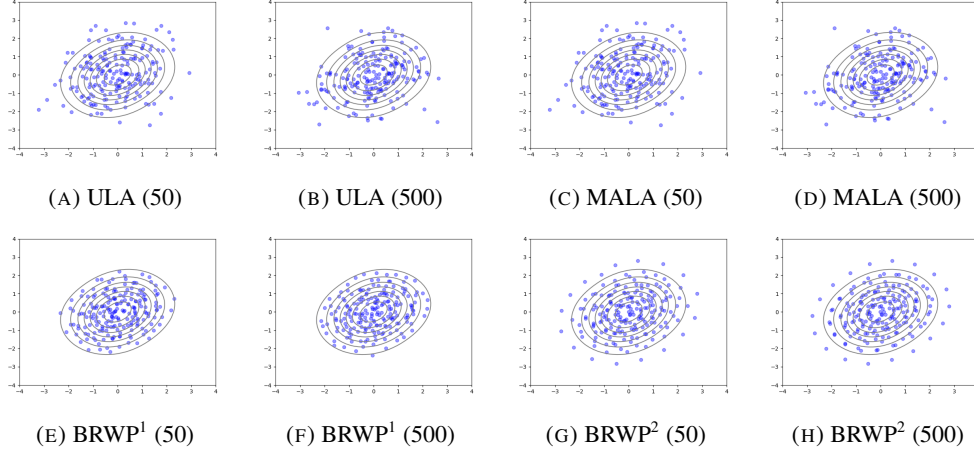


FIGURE 4. Evolution of particles under ULA, MALA, and BRWP for the bimodal distribution, with step-size $\eta = 0.1$. Superscripts indicate different parameters of T , with BRWP^1 having $T = 0.01$ and BRWP^2 having $T = 0.1$. We observe that the iterates of BRWP converge in an organized manner as opposed to the randomness of ULA and MALA, with the lower level of T giving some variance reduction properties.

We consider the same problem parameters as in Dwivedi et al. (2018); Dalalyan (2017), taking dimension $d = 2$ and the parameter $a = (1/2, 1/2)$. This gives strong convexity parameter $m = \frac{1}{2}$ and Lipschitz constant $L = 1$. The initial distribution is chosen as $\mathcal{N}(0, L^{-1}I) = \mathcal{N}(0, I)$, and we initialize 200 particles with this distribution. For consistency, we use the same initialization for each of the compared methods. We compare with BRWP with parameters $T = 0.01, 0.1$, with $P = 25$ Monte Carlo samples for approximating the normalizing constant $\mathcal{Z}$.

We observe in Figure 4 that the samples of BRWP for parameters $T = 0.01$ and $T = 0.1$ both converge to roughly ellipsoidal patterns for this non-Gaussian case, fitting the level sets of the density. Moreover, we observe that the samples themselves exhibit some sort of structure, and do not have random-walk-like movements between the iterations as a result of the deterministic discretization.

3.3. Bimodal Distribution. As a more complicated toy example, we consider the two-dimensional bi-modal distribution as in Wang & Li (2022). This objective function has the form

$$p(x) \propto \exp(-2(\|x\| - 3)^2) [\exp(-2(x_1 - 3)^2) + \exp(-2(x_1 + 3)^2)].$$

This is generated by the potential V with gradient ∇V as follows:

$$V(x) = 2(\|x\| - 3)^2 - 2 \log [\exp(-2(x_1 - 3)^2) + \exp(-2(x_1 + 3)^2)], \quad (34a)$$

$$\nabla V(x) = 4 \frac{(\|x\| - 3)x}{\|x\|} + \frac{4(x_1 - 3) \exp(-2(x_1 - 3)^2) + 4(x_1 + 3) \exp(-2(x_1 + 3)^2)}{\exp(-2(x_1 - 3)^2) + \exp(-2(x_1 + 3)^2)} e_1, \quad (34b)$$

where $e_1 = (1, 0)^\top$ is the first standard coordinate vector. We fix the step-size for ULA and MALA to be $\eta = 0.01$, and regularization parameter $T = 0.01, 0.05, 0.1$ for the BRWP method. The samples are initialized as standard Gaussian $\mathcal{N}(0, I)$, and we use 200 particles for simulation.

In Figure 5, we plot the evolution of ULA, MALA and BRWP with $T = 0.01$ at iteration numbers 10, 50, 100 and 2000. This figure illustrates that the samples of BRWP travel in a structured manner, and indeed stay approximately the same even after many iterations. In contrast, ULA and MALA continue to exhibit random behaviors after reaching the neighborhoods of the modes.

Figure 6 explores the behavior of the compared algorithms in the very large step-size regime, where none of the methods are expected to converge. Taking the step-size $\eta = 0.5$, we have that

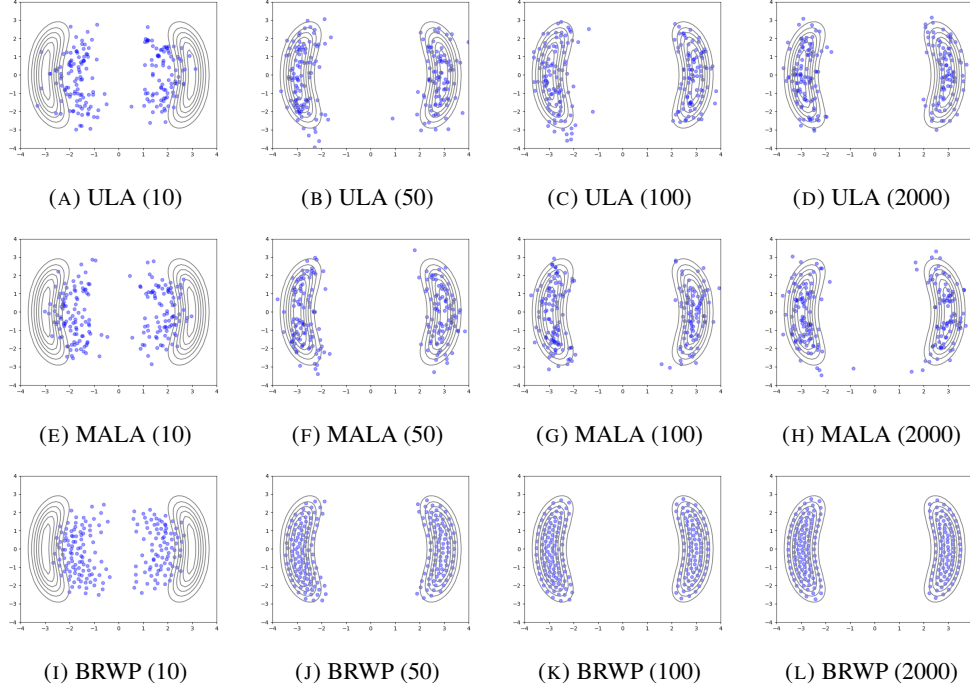


FIGURE 5. Evolution of particles under ULA, MALA and BRWP for the bimodal distribution, with step-size $\eta = 0.01$. The parameter of T was taken to be $T = 0.01$ for BRWP. We observe that the iterates of BRWP converge in a structured manner to fit the distribution, and the iterates stay almost identical from 100 to 200 iterations. In contrast, the SDE based methods ULA and MALA have samples that continue to be random.

ULA diverges, while MALA and BRWP with $T = 0.1$ do not converge to neighborhoods of the modes. However, once we take $T = 0.2$ to be sufficiently large, we again observe a convergent behavior. The iterates converge towards a curve that follows the valleys of V . This suggests that T implicitly performs a variance reduction even in the case where the target distribution is not log-concave.

3.4. Bayesian Logistic Regression. We additionally explore the performance for Bayesian logistic regression, in the framework detailed in Dwivedi et al. (2018); Dalalyan (2017). The problem is as follows. Suppose that we have covariates $x \in \mathbb{R}^d$ as well as a binary variable $y \in \{0, 1\}$. The logistic model is for the conditional distribution of y given x for a parameter $\theta \in \mathbb{R}^d$ is

$$\mathbb{P}(y = 1 \mid x, \theta) = \frac{\exp(\theta^\top x)}{1 + \exp(\theta^\top x)}.$$

Given a binary vector $Y \in \{0, 1\}^n$ and a feature matrix $X \in \mathbb{R}^{n \times d}$ with rows $x_i \in \mathbb{R}^d$, suppose we impose a prior density $\theta \sim \mathcal{N}(0, \Sigma_X)$, where $\Sigma_X = \frac{1}{n} X^\top X$ is the sample covariance matrix of X . Then the posterior density of θ is given by

$$p(\theta \mid X, Y) \propto \exp \left\{ Y^\top X \theta - \sum_{i=1}^n \log(1 + \exp(\theta^\top x_i)) - \alpha \|\Sigma_X^{\frac{1}{2}} \theta\|_2^2 \right\}, \quad (35)$$

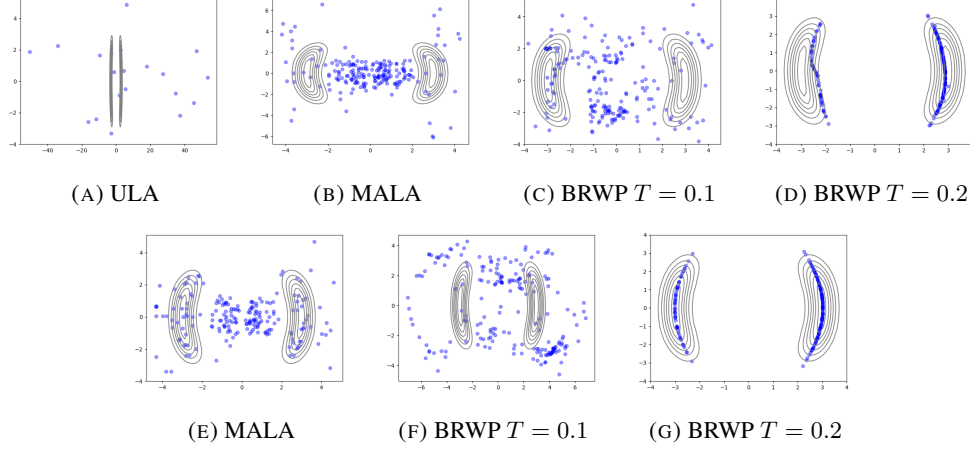


FIGURE 6. Evolution of particles under ULA, MALA and BRWP for the bimodal distribution, with very large step-size $\eta = 0.5$. Figures (a-d) are at iteration 10, while (e-g) are at iteration 100. ULA diverges after 10 iterations. MALA and BRWP with $T = 0.1$ do not converge to the modes. For sufficiently large T , BRWP removes low covariance components and allows for a larger step-size.

where $\alpha > 0$ is a regularization parameter. This can be cast into the problem of sampling from a Gibbs distribution, with potential

$$V(\theta) = -Y^\top X\theta + \sum_{i=1}^n \log(1 + \exp(\theta^\top x_i)) + \alpha \|\Sigma_X^{\frac{1}{2}} \theta\|_2^2,$$

$$\nabla V(\theta) = -X^\top Y + \sum_{i=1}^n \frac{x_i}{1 + \exp(-\theta^\top x_i)} + \alpha \Sigma_X \theta.$$

As in Dwivedi et al. (2018), the eigenvalues of the Hessian are bounded by $L = (0.25n + \alpha)\lambda_{\max}(\Sigma_X)$ and $m = \alpha\lambda_{\min}(\Sigma_X)$. In our experiments, we choose the logistic regression parameters as $\alpha = 0.5, d = 2, n = 50$. We fix the step-size to be $\eta = 0.05$, and run each of the methods for 5000 iterations. We initialize $N = 1000$ samples using the distribution $\mathcal{N}(0, L^{-1}I)$ as in Dwivedi et al. (2018).

For evaluation, we consider the error with respect to the true minimizers of V , denoted by θ^* . This is also known as the maximum a posteriori (MAP) estimate in the Bayesian optimization literature. To compute θ^* , we run gradient descent for 1000 iterations with step-size 1e-3, followed by 1000 iterations with step-size 1e-4, initialized at $\theta = (1, 1)^\top$. For the computed samples, we compute the expected ℓ_1 deviation from θ^* divided by d , as well as the ℓ_1 distance of the sample mean to θ^* divided by d . The metrics are, where $\hat{\theta}_k$ is the empirical distribution and $\bar{\theta}$ is the sample mean at the k -th iteration,

$$\varepsilon_1 = \frac{1}{d} \|\bar{\theta} - \theta^*\|_1, \quad \varepsilon_2 = \frac{1}{d} \mathbb{E} \|\hat{\theta}_k - \theta^*\|_1. \quad (36)$$

These metrics deviate from that of Dwivedi et al. (2018) in the sense that θ^* is chosen to be the minimum of V , instead of the $\theta = (1, 1)$ used to generate the samples. This compensates for the bias generated by the added regularization, and makes it easier to compare the posterior means to the MAP estimate.

Figure 7 plots the error metrics $\varepsilon_1, \varepsilon_2$ as defined in Equation (36). We observe that the metrics for the BRWP scheme for regularization parameters $T = 0.025, 0.05, 0.1, 0.2$ are lower than ULA and MALA. Moreover, we observe that the error metrics converge after around 200 iterations and have significantly less noise across the iterations. From ε_1 being smaller, we have that the posterior mean is closer to the MAP estimate for BRWP. ε_2 being smaller demonstrates again the variance reduction of the scheme, with larger values of T corresponding to less variance. Figures 8 and 9 plot

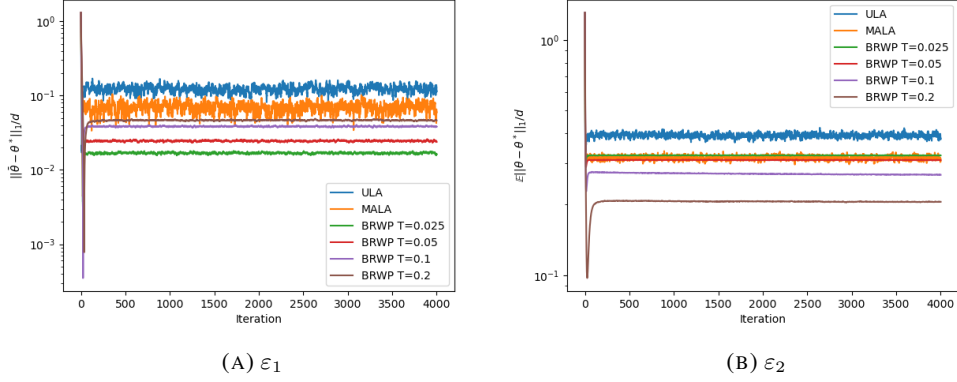


FIGURE 7. Plot of the ε_1 and ε_2 metrics. The regularization parameter is $\alpha = 0.5$, with condition number $\kappa \approx 28.2$. The step-size is $\eta = 0.05$ for all methods. We observe that for small values of T , the sample mean is closer to the true parameter value θ^* ; for larger T , the variance is lower. This demonstrates a bias-variance trade-off of T .

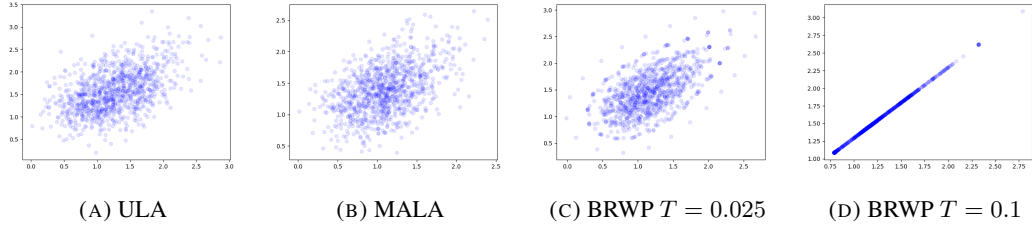


FIGURE 8. Plots of the samples of θ after 4000 iterations, with $N = 1000$ samples. Parameters are $\alpha = 0.5, \eta = 0.05$. For this particular instantiation, we find that $\theta^* \approx (1.16, 1.45)$. We observe that for small T , we have a teardrop shaped structure. For large T , we have mode collapse in one direction.

the samples after 4000 iterations for the various levels of T . We observe that for $T = 0.025, 0.05$, a clear teardrop-shaped structure arises, traced out by the outer samples. For $T = 0.1, 0.2$, the samples appear collinear.

We can additionally interpret ε_2 as an optimization objective, rather than a sampling objective. Using this interpretation, we have that for larger T , the optimization effect on V is larger and dominates the diffusion. For sampling schemes such as ULA and MALA, this would be dictated by the regularization parameter β . For a convex objective V , as $\beta \rightarrow 0$, we have less diffusion effects, and the target density $\exp(-V(x)/\beta)$ converges to the Dirac mass at the minimizer of V .

4. CONCLUSION

This work presents a novel deterministic approach to sampling using the regularized Wasserstein proximal. By approximating the density as a regularized Wasserstein proximal of the empirical distribution, we obtain a particle-based ODE approximation to the Fokker-Planck equation at each time step. Discretizing this approximate ODE using a backwards Euler step gives a deterministic sampling algorithm. We fully characterize the convergence and give closed-form iterations in the case of an Ornstein-Uhlenbeck process with quadratic potential, corresponding to a Gaussian target distribution. Moreover, we observe numerically that the proposed BRWP scheme converges in a visually structured manner by foregoing stochasticity.

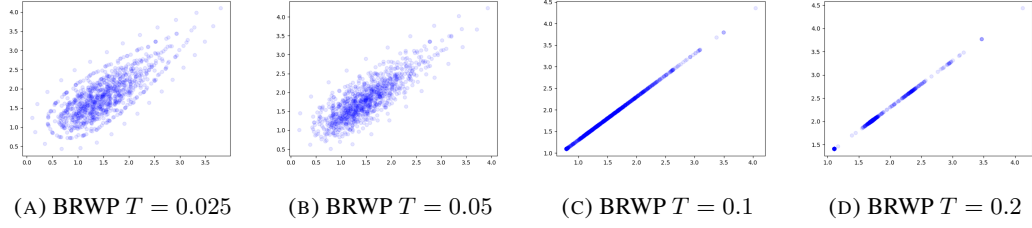


FIGURE 9. Plots of the samples of θ after 4000 iterations, with $N = 1000$ samples. Parameters are $\alpha = 0.1, \eta = 0.05$. For this particular instantiation, we find that $\theta^* \approx (1.32, 1.62)$. We observe variance reduction in the approximately $y = -x$ direction as T increases.

While the empirical results demonstrate the practicality of our scheme as an alternative to non-deterministic sampling algorithms such as ULA and MALA in the case of low-dimensional non-log-concave distributions, there are two main limitations. Firstly, the variance reduction/mode collapse phenomenon increases the number of samples required, and thus the complexity of the BRWP method. Secondly, the analysis is currently limited to the case of Gaussians. To further cement this method as a suitable and provably convergent method for sampling, demonstrating the convergence rate for more general distributions such as log-concave distributions is required. We conjecture that the variance reducing behavior of T can be shown to implicitly reduce or remove deviations in directions of small covariance, which may be useful to remove small noise in data. Various open questions corresponding to the proposed scheme follow.

Convergence rates for log-concave density. The preliminary analysis given is only for Gaussian densities, with empirical results suggesting that the method continues to work. We believe that the closed-form updates of the BRWP scheme can lead to an analytic solution for convergence rates for log-concave target densities.

Discretization and approximation error. We made four approximating steps at the start of Section 2 to construct the BRWP scheme, including approximating the Fokker-Planck equation, ODE discretization using the backwards Euler method, and replacing densities with empirical measures at each iteration. The impact and convergence rate of these approximations with respect to the number of samples or the regularization parameters could be an interesting direction.

Sample scaling in dimension and variance reduction phenomenon. We observed in Section 3.1 that in higher dimensions, the number of samples plays a role in variance reduction, even if the analytic rates are dimension independent. Quantifying or mitigating this effect for either sampling or optimization would be beneficial for high-dimensional applications.

Structure of the iterates. We observed empirically that the iterates cluster in a visually cohesive manner, with external iterates approximately lying on level sets of the density. However, this is deeply connected with the discretization method, and could prove to be a difficult yet rewarding problem.

ACKNOWLEDGEMENTS

H.Y. Tan acknowledges support from GSK.ai, the Masason Foundation, and the European Union Horizon 2020 research and innovation programme under the Marie Skłodowska-Curie grant agreement No. 777826 NoMADS. S. Osher is supported in part by AFOSR MURI FP 9550-18-1-502 and ONR grants: N00014-20-1-2093 and N00014-20-1-2787. W. Li's work is supported by AFOSR MURI FP 9550-18-1-502, AFOSR YIP award No. FA9550-23-1-0087, NSF DMS-2245097, and NSF RTG: 2038080.

REFERENCES

Christophe Andrieu, Nando De Freitas, Arnaud Doucet, and Michael I Jordan. An introduction to MCMC for machine learning. *Machine learning*, 50:5–43, 2003.

- Georgios Batzolis, Jan Stanczuk, Carola-Bibiane Schönlieb, and Christian Etmann. Conditional image generation with score-based diffusion models. *arXiv preprint arXiv:2111.13606*, 2021.
- Bernhard Baumgartner. An inequality for the trace of matrix products, using absolute values. *arXiv preprint arXiv:1106.6189*, 2011.
- Claude JP Bélisle, H Edwin Romeijn, and Robert L Smith. Hit-and-run algorithms for generating multivariate distributions. *Mathematics of Operations Research*, 18(2):255–266, 1993.
- Jean-David Benamou and Yann Brenier. A computational fluid mechanics solution to the Monge-Kantorovich mass transfer problem. *Numerische Mathematik*, 84(3):375–393, 2000.
- Michael Betancourt. A conceptual introduction to Hamiltonian Monte Carlo. *arXiv preprint arXiv:1701.02434*, 2017.
- Sam Bond-Taylor, Adam Leach, Yang Long, and Chris G Willcocks. Deep generative modelling: A comparative review of vaes, gans, normalizing flows, energy-based and autoregressive models. *IEEE transactions on pattern analysis and machine intelligence*, 2021.
- Zdravko I Botev, Joseph F Grotowski, and Dirk P Kroese. Kernel density estimation via diffusion. *Annals of Statistics*, 38(5):2916–2957, 2010.
- Steve Brooks, Andrew Gelman, Galin Jones, and Xiao-Li Meng. *Handbook of Markov chain Monte Carlo*. CRC press, 2011.
- José Antonio Carrillo, Katy Craig, and Francesco S Patacchini. A blob method for diffusion. *Calculus of Variations and Partial Differential Equations*, 58:1–53, 2019.
- Ricky TQ Chen, Yulia Rubanova, Jesse Bettencourt, and David K Duvenaud. Neural ordinary differential equations. *Advances in neural information processing systems*, 31, 2018.
- Yin Dai, Yuling Jiao, Lican Kang, Xiliang Lu, and Jerry Zhijian Yang. Global optimization via Schrödinger-Föllmer diffusion. *arXiv e-prints*, pp. arXiv–2111, 2021.
- Arnak S Dalalyan. Theoretical guarantees for approximate sampling from smooth and log-concave densities. *Journal of the Royal Statistical Society Series B: Statistical Methodology*, 79(3):651–676, 2017.
- Pierre Del Moral. *Mean field simulation for Monte Carlo integration*. CRC press, 2013.
- Luc Devroye, Abbas Mehrabian, and Tommy Reddad. The total variation distance between high-dimensional Gaussians with the same mean. *arXiv preprint arXiv:1810.08693*, 2018.
- Alain Durmus and Eric Moulines. High-dimensional bayesian inference via the unadjusted Langevin algorithm. 2019.
- Raaz Dwivedi, Yuansi Chen, Martin J Wainwright, and Bin Yu. Log-concave sampling: Metropolis-hastings algorithms are fast! In *Conference on learning theory*, pp. 793–797. PMLR, 2018.
- Crispin W Gardiner et al. *Handbook of stochastic methods*, volume 3. springer Berlin, 1985.
- Artur Gramacki. *Nonparametric kernel density estimation and its computational aspects*, volume 37. Springer, 2018.
- Roger A Horn and Charles R Johnson. *Matrix analysis*. Cambridge university press, 2012.
- Aapo Hyvärinen and Peter Dayan. Estimation of non-normalized statistical models by score matching. *Journal of Machine Learning Research*, 6(4), 2005.
- Richard Jordan, David Kinderlehrer, and Felix Otto. The variational formulation of the Fokker-Planck equation. *SIAM journal on mathematical analysis*, 29(1):1–17, 1998.
- Ioannis Karatzas and Steven Shreve. *Brownian motion and stochastic calculus*, volume 113. Springer Science & Business Media, 1991.
- Mehran Kardar. *Statistical physics of particles*. Cambridge University Press, 2007.
- JooSeuk Kim and Clayton D Scott. Robust kernel density estimation. *The Journal of Machine Learning Research*, 13(1):2529–2565, 2012.
- Ryogo Kubo. Stochastic liouville equations. *Journal of Mathematical Physics*, 4(2):174–183, 1963.
- Rémi Laumont, Valentin De Bortoli, Andrés Almansa, Julie Delon, Alain Durmus, and Marcelo Pereyra. Bayesian imaging using plug & play priors: when langevin meets tweedie. *SIAM Journal on Imaging Sciences*, 15(2):701–737, 2022.
- Wei Li, Wei Liu, Jinlin Chen, Libing Wu, Patrick D Flynn, Wei Ding, and Ping Chen. Reducing mode collapse with Monge-Kantorovich optimal transport for generative adversarial networks. *IEEE Transactions on Cybernetics*, 2023a.

- Wuchen Li, Siting Liu, and Stanley Osher. A kernel formula for regularized Wasserstein proximal operators. *arXiv preprint arXiv:2301.10301*, 2023b.
- Joseph Liouville. Note sur la théorie de la variation des constantes arbitraires. *Journal de mathématiques pures et appliquées*, 3:342–349, 1838.
- David JC MacKay. Bayesian neural networks and density networks. *Nuclear Instruments and Methods in Physics Research Section A: Accelerators, Spectrometers, Detectors and Associated Equipment*, 354(1):73–80, 1995.
- Dimitra Maoutsa, Sebastian Reich, and Manfred Opper. Interacting particle solutions of Fokker–Planck equations through gradient–log–density estimation. *Entropy*, 22(8):802, 2020.
- Jonathan C Mattingly, Andrew M Stuart, and Desmond J Higham. Ergodicity for sdes and approximations: locally Lipschitz vector fields and degenerate noise. *Stochastic processes and their applications*, 101(2):185–232, 2002.
- Kerrie L Mengersen and Richard L Tweedie. Rates of convergence of the Hastings and Metropolis algorithms. *The annals of Statistics*, 24(1):101–121, 1996.
- Sean P Meyn and Robert L Tweedie. Computable bounds for geometric convergence rates of Markov chains. *The Annals of Applied Probability*, pp. 981–1011, 1994.
- Erik Nijkamp, Ruiqi Gao, Pavel Sountsov, Srinivas Vasudevan, Bo Pang, Song-Chun Zhu, and Ying Nian Wu. Mcmc should mix: learning energy-based model with neural transport latent space mcmc. In *International Conference on Learning Representations (ICLR 2022)*, 2022.
- Stanley Osher, Howard Heaton, and Samy Wu Fung. A Hamilton–Jacobi-based proximal operator. *Proceedings of the National Academy of Sciences*, 120(14):e2220469120, 2023.
- Felix Otto. The geometry of dissipative evolution equations: the porous medium equation. 2001.
- Giorgio Parisi. Correlation functions and computer simulations. *Nuclear Physics B*, 180(3):378–384, 1981.
- Sam Patterson and Yee Whye Teh. Stochastic gradient Riemannian Langevin dynamics on the probability simplex. *Advances in neural information processing systems*, 26, 2013.
- Marcelo Pereyra. Proximal Markov chain Monte Carlo algorithms. *Statistics and Computing*, 26: 745–760, 2016.
- Gareth O Roberts and Richard L Tweedie. Exponential convergence of Langevin distributions and their discrete approximations. *Bernoulli*, pp. 341–363, 1996.
- Peter J Rossky, Jimmie D Doll, and Harold L Friedman. Brownian dynamics as smart Monte Carlo simulation. *The Journal of Chemical Physics*, 69(10):4628–4633, 1978.
- Yang Song, Jascha Sohl-Dickstein, Diederik P Kingma, Abhishek Kumar, Stefano Ermon, and Ben Poole. Score-based generative modeling through stochastic differential equations. *arXiv preprint arXiv:2011.13456*, 2020.
- Akash Srivastava, Lazar Valkov, Chris Russell, Michael U Gutmann, and Charles Sutton. Veegan: Reducing mode collapse in gans using implicit variational learning. *Advances in neural information processing systems*, 30, 2017.
- George R Terrell and David W Scott. Variable kernel density estimation. *The Annals of Statistics*, pp. 1236–1265, 1992.
- Richard Chace Tolman. *The principles of statistical mechanics*. Courier Corporation, 1979.
- Philippe Van Kerm. Adaptive kernel density estimation. *The Stata Journal*, 3(2):148–156, 2003.
- Matt P Wand and M Chris Jones. *Kernel smoothing*. CRC press, 1994.
- Yifei Wang and Wuchen Li. Accelerated information gradient flow. *Journal of Scientific Computing*, 90:1–47, 2022.
- Andre Wibisono. Sampling as optimization in the space of measures: The Langevin dynamics as a composite optimization problem. In *Conference on Learning Theory*, pp. 2093–3027. PMLR, 2018.

APPENDIX A. DERIVATION OF UPDATES FOR GAUSSIAN

In this section, we derive the closed form expressions for updating a Gaussian distribution, under the Ornstein-Uhlenbeck process.

We begin with the derivation of Equation (16), which is the approximate Wasserstein proximal of the distribution at iteration k . In the following derivation, we discard constants that do not depend on x and y (but are allowed to depend on μ_k, σ_k^2). We begin with computing the normalization constant of $K(x, y)$, given by the denominator of Equation (4).

$$\begin{aligned}
& \int_{\mathbb{R}^d} \exp \left(-\frac{1}{2\beta} \left(V(z) + \frac{(z-y)^2}{2T} \right) \right) dz \\
&= \int_{\mathbb{R}^d} \exp \left(-\frac{1}{2\beta} \left(\frac{az^2}{2} + \frac{(z-y)^2}{2T} \right) \right) dz \\
&= \int_{\mathbb{R}^d} \exp \left(-\frac{1}{2\beta} \left(\frac{a}{2} + \frac{1}{2T} \right) \left(z - \frac{y/2T}{a/2 + 1/2T} \right)^2 \right) \exp \left(\frac{1}{2\beta} \frac{(y/2T)^2}{a/2 + 1/2T} \right) \exp \left(\frac{-y^2}{4\beta T} \right) dz \\
&\propto \exp \left(\frac{y^2}{2\beta} \left(\frac{1}{2T(1+aT)} - \frac{1}{2T} \right) \right).
\end{aligned}$$

Substituting into the definition of $\rho_{k,T}$,

$$\begin{aligned}
\rho_{k,T}(x) &= \int_{\mathbb{R}} K(x, y) \frac{1}{\sqrt{2\pi}\sigma_k} \exp \left(-\frac{(y-\mu_k)^2}{2\sigma_k^2} \right) dy \\
&\propto \int \exp \left(-\frac{1}{2\beta} \left(\frac{ax^2}{2} + \frac{(x-y)^2}{2T} \right) \right) \exp \left(-\frac{(y-\mu_k)^2}{2\sigma_k^2} \right) - \frac{y^2}{2\beta} \left(\frac{1}{2T(1+aT)} - \frac{1}{2T} \right) dy \\
&= \exp \left(-\frac{ax^2}{4\beta} \right) \int \exp \left(-\frac{(y-x)^2}{4\beta T} - \frac{(y-\mu_k)^2}{2\sigma_k^2} - \frac{y^2}{2\beta} \left(\frac{1}{2T(1+aT)} - \frac{1}{2T} \right) \right) dy \\
&= \exp \left(-\frac{ax^2}{4\beta} \right) \\
&\quad \times \int \exp \left(-\frac{1}{2} \left[y^2 \left(\frac{1}{2\beta T(1+aT)} + \frac{1}{\sigma_k^2} \right) - 2y \left(\frac{x}{2\beta T} + \frac{\mu_k}{\sigma_k^2} \right) + \left(\frac{x^2}{2\beta T} + \frac{\mu_k^2}{\sigma_k^2} \right) \right] \right) dy \\
&\propto \exp \left(-\frac{ax^2}{4\beta} - \frac{x^2}{4\beta T} + \frac{1}{2} \frac{\left(\frac{x}{2\beta T} + \frac{\mu_k}{\sigma_k^2} \right)^2}{\frac{1}{2\beta T(1+aT)} + \frac{1}{\sigma_k^2}} \right) \\
&\quad \times \int \exp \left[-\frac{1}{2} \left(\frac{1}{2\beta T(1+aT)} + \frac{1}{\sigma_k^2} \right) \left(y - \frac{\frac{x}{2\beta T} + \frac{\mu_k}{\sigma_k^2}}{\frac{1}{2\beta T(1+aT)} + \frac{1}{\sigma_k^2}} \right)^2 \right] dy.
\end{aligned}$$

Observe in the final expression, the integral is of a Gaussian density whose variance does not depend on x , hence integrates to something independent of x . Hence, $\rho_{k,T} \sim \mathcal{N}(\tilde{\mu}_{k+1}, \tilde{\sigma}_{k+1}^2)$ is a Gaussian density on x , with mean and variance

$$\tilde{\mu}_{k+1} = \frac{\mu_k}{1+aT}, \quad \tilde{\sigma}_{k+1}^2 = \frac{\sigma_k^2}{(1+aT)^2} + \frac{2\beta T}{1+aT}.$$

This shows Equation (16). We now compute it in the multi-dimensional case as well, taking special care where the covariance matrices do not commute.

Computing the denominator of Equation (4) as before, we have

$$\begin{aligned}
& \int_{\mathbb{R}^d} \exp \left(-\frac{1}{2\beta} \left(V(z) + \frac{\|z - y\|^2}{2T} \right) \right) dz \\
&= \int_{\mathbb{R}^d} \exp \left(-\frac{1}{2\beta} \left(\frac{z^\top \Sigma^{-1} z}{2} + \frac{\|z - y\|^2}{2T} \right) \right) dz \\
&= \int_{\mathbb{R}^d} \left[\exp \left(-\frac{1}{2\beta} \left(z - \left(\frac{\Sigma^{-1}}{2} + \frac{I}{2T} \right)^{-1} \frac{y}{2T} \right)^\top \left(\frac{\Sigma^{-1}}{2} + \frac{I}{2T} \right) \left(z - \left(\frac{\Sigma^{-1}}{2} + \frac{I}{2T} \right)^{-1} \frac{y}{2T} \right) \right) \right. \\
&\quad \left. \exp \left(\frac{1}{2\beta} \left(\frac{y}{2T} \right)^\top \left(\frac{\Sigma^{-1}}{2} + \frac{I}{2T} \right)^{-1} \left(\frac{y}{2T} \right) \right) \exp \left(\frac{-\|y\|^2}{4\beta T} \right) \right] dz \\
&\propto \exp \left(\frac{1}{2\beta} y^\top \left((2T(I + \Sigma^{-1}T))^{-1} - \frac{I}{2T} \right) y \right),
\end{aligned}$$

where the second equality follows from completing the square, and the final expression from integrating with respect to z , noting that the first exponential term is a Gaussian whose variance does not depend on y . We compute the approximate Wasserstein proximal $\rho_{k,T}$, given $\rho_{k,0} \sim \mathcal{N}(\mu_k, \Sigma_k)$:

$$\begin{aligned}
\rho_{k,T}(x) &\propto \int_{\mathbb{R}^d} K(x, y) \exp \left(-\frac{1}{2} (y - \mu_k)^\top \Sigma_k^{-1} (y - \mu_k) \right) dy \\
&= \int \exp \left[-\frac{1}{2\beta} \left(\frac{1}{2} x^\top \Sigma^{-1} x + \frac{\|y - x\|^2}{2T} \right) - \frac{1}{2} (y - \mu_k)^\top \Sigma_k^{-1} (y - \mu_k) \right. \\
&\quad \left. - \frac{1}{2\beta} y^\top \left[\frac{1}{2T} (I + T\Sigma^{-1})^{-1} - \frac{I}{2T} \right] y \right] dy \\
&\propto \exp \left(-\frac{x^\top \Sigma^{-1} x}{4\beta} - \frac{\|x\|^2}{4\beta T} \right) \\
&\quad \cdot \int \exp \left[y^\top \left(-\frac{I}{4\beta T} - \frac{\Sigma_k^{-1}}{2} - \frac{(I + T\Sigma^{-1})^{-1}}{4\beta T} + \frac{I}{4\beta T} \right) y \right. \\
&\quad \left. + y^\top \left(\frac{x}{4\beta T} + \frac{\Sigma_k^{-1} \mu_k}{2} \right) + \left(\frac{x}{4\beta T} + \frac{\Sigma_k^{-1} \mu_k}{2} \right)^\top y \right] dy \\
&\propto \exp \left(-\frac{1}{2} \left[x^\top \left(\frac{\Sigma^{-1}}{2\beta} + \frac{I}{2\beta T} \right) x \right. \right. \\
&\quad \left. \left. - \left(\frac{x}{2\beta T} + \Sigma_k^{-1} \mu_k \right)^\top \left(\Sigma_k^{-1} + \frac{1}{2\beta T} (I + T\Sigma^{-1})^{-1} \right)^{-1} \left(\frac{x}{2\beta T} + \Sigma_k^{-1} \mu_k \right) \right] \right) \\
&\propto \exp \left(-\frac{1}{2} \left[x^\top \left(\frac{\Sigma^{-1}}{2\beta} + \frac{I}{2\beta T} - [(2\beta T)^2 \Sigma_k^{-1} + 2\beta T(I + T\Sigma^{-1})^{-1}]^{-1} \right) x \right. \right. \\
&\quad \left. \left. + x^\top [(2\beta T \Sigma_k^{-1} + (I + T\Sigma^{-1})^{-1})^{-1} \Sigma_k^{-1} \mu_k] \right. \right. \\
&\quad \left. \left. + [(2\beta T \Sigma_k^{-1} + (I + T\Sigma^{-1})^{-1})^{-1} \Sigma_k^{-1} \mu_k]^\top x \right] \right).
\end{aligned}$$

The regularized Wasserstein proximal $\rho_{k,T}$ is thus Gaussian with mean and inverse covariance

$$\begin{aligned}
\tilde{\Sigma}_{k+1}^{-1} &= \frac{\Sigma^{-1}}{2\beta} + \frac{I}{2\beta T} - [(2\beta T)^2 \Sigma_k^{-1} + 2\beta T(I + T\Sigma^{-1})^{-1}]^{-1} \\
&= [(2\beta T)^2 \Sigma_k^{-1} + 2\beta T(I + T\Sigma^{-1})^{-1}]^{-1} \\
&\quad \left([(2\beta T)^2 \Sigma_k^{-1} + 2\beta T(I + T\Sigma^{-1})^{-1}] \left(\frac{\Sigma^{-1}}{2\beta} + \frac{I}{2\beta T} \right) - I \right) \\
&= [(2\beta T)^2 \Sigma_k^{-1} + 2\beta T(I + T\Sigma^{-1})^{-1}]^{-1} \\
&\quad \left([(2\beta T) \Sigma_k^{-1} + (I + T\Sigma^{-1})^{-1}] (T\Sigma^{-1} + I) - I \right) \\
&= [(2\beta T)^2 \Sigma_k^{-1} + 2\beta T(I + T\Sigma^{-1})^{-1}]^{-1} \\
&\quad (2\beta T \Sigma_k^{-1} (I + T\Sigma^{-1})) \\
&= (2\beta T I + \Sigma_k (I + T\Sigma^{-1})^{-1})^{-1} (I + T\Sigma^{-1}); \\
&= (2\beta T (I + T\Sigma^{-1})^{-1} + (I + T\Sigma^{-1})^{-1} \Sigma_k (I + T\Sigma^{-1})^{-1})^{-1}, \\
\tilde{\mu}_{k+1} &= \tilde{\Sigma}_{k+1} [(2\beta T \Sigma_k^{-1} + (I + T\Sigma^{-1})^{-1})^{-1} \Sigma_k^{-1} \mu_k] \\
&= (I + T\Sigma^{-1})^{-1} \mu_k.
\end{aligned}$$

This shows the recurrence relation Equation (19) for the distribution update under BRWP.

APPENDIX B. RECURRENCE RELATION FOR EIGENVALUES FOR COMMUTING GAUSSIANS

We let $\Sigma = \text{diag}(\xi^{(1)}, \dots, \xi^{(d)})$ be positive definite, the stationary distribution Π of the discrete scheme Equation (12) be given by

$$\Pi \sim \mathcal{N}(0, \Sigma_\infty), \quad \Sigma_\infty = \text{diag}(\tau_\infty^{(i)} \mid i = 1, \dots, d), \quad (37)$$

$$\tau_\infty^{(i)} = \beta \xi^{(i)} (1 - T^2 \xi^{(i)-2}), \quad i = 1, \dots, d. \quad (38)$$

Observe that the i -th entry of $\tilde{\Sigma}_{k+1}^{-1}$ is given by

$$\left(\tilde{\Sigma}_{k+1}^{-1} \right)_{i,i} = \left(2\beta T + \frac{\tau_k^{(i)}}{(1 + T\xi^{(i)-1})} \right)^{-1} (1 + T\xi^{(i)-1}),$$

and therefore

$$\begin{aligned}
(I - \eta \Sigma^{-1} + \eta \beta \tilde{\Sigma}_{k+1}^{-1})_{i,i} &= 1 - \eta \xi^{(i)-1} + \eta \beta \left(2\beta T + \frac{\tau_k^{(i)}}{(1 + T\xi^{(i)-1})} \right)^{-1} (1 + T\xi^{(i)-1}) \\
&= 1 - \eta \xi^{(i)-1} + \frac{\eta \beta (1 + T\xi^{(i)-1})^2}{\tau_k^{(i)} + 2\beta T (1 + T\xi^{(i)-1})}.
\end{aligned}$$

Temporarily dropping the (i) superscripts that denote the coordinate, we consider the evolution of the covariance in the i -th coordinate, which is sufficient since the covariance matrices are diagonal. Indeed, from Equation (20b) it evolves as

$$\tau_{k+1} = \left(1 - \eta \xi^{-1} + \frac{\eta \beta (1 + T\xi^{-1})^2}{\tau_k + 2\beta T (1 + T\xi^{-1})} \right)^2 \tau_k. \quad (39)$$

Observe that this is the same as Equation (17b) up to a renaming of variables, in particular by letting $a = \xi^{-1}$ and $\sigma_k^2 = \tau_k$. We thus have the same fixed points, given by

$$\tau_\infty = \beta \xi (1 - T^2 \xi^{-2}).$$

We wish to consider the mixing time with respect to this variance. Consider the ansatz

$$\sqrt{\tau_k} = \sqrt{\beta \xi (1 - T^2 \xi^{-2})} + \sqrt{1 + T\xi^{-1}} \gamma_k.$$

We compute a recurrence relation for (γ_k) using Equation (39)

$$\begin{aligned}
\sqrt{\tau_{k+1}} &= \sqrt{\tau_k} \left[1 + \eta \left(-\xi^{-1} + \frac{\beta(1 + T\xi^{-1})^2}{(\sqrt{\beta\xi(1 - T^2\xi^{-2})} + \sqrt{1 + T\xi^{-1}\gamma_k})^2 + 2\beta T(1 + T\xi^{-1})} \right) \right] \\
&= \sqrt{\tau_k} \left[1 + \eta \left(-\xi^{-1} + \frac{\beta(1 + T\xi^{-1})}{(\sqrt{\beta\xi(1 - T\xi^{-1})} + \gamma_k)^2 + 2\beta T} \right) \right] \\
&= \sqrt{\tau_k} \left[1 + \eta \left(\frac{-\xi^{-1}((\sqrt{\beta\xi(1 - T\xi^{-1})} + \gamma_k)^2 + 2\beta T) + \beta(1 + T\xi^{-1})}{(\sqrt{\beta\xi(1 - T\xi^{-1})} + \gamma_k)^2 + 2\beta T} \right) \right] \\
&= \sqrt{\tau_k} \left[1 - \frac{\eta\gamma_k[2\xi^{-1}\sqrt{\beta\xi(1 - T\xi^{-1})} + \xi^{-1}\gamma_k]}{(\sqrt{\beta\xi(1 - T\xi^{-1})} + \gamma_k)^2 + 2\beta T} \right] \\
&= \sqrt{\tau_k} - \sqrt{\tau_k} \left[\frac{\eta\gamma_k[2\xi^{-1}\sqrt{\beta\xi(1 - T\xi^{-1})} + \xi^{-1}\gamma_k]}{(\sqrt{\beta\xi(1 - T\xi^{-1})} + \gamma_k)^2 + 2\beta T} \right].
\end{aligned}$$

Subtracting $\sqrt{\beta\xi(1 - T^2\xi^{-2})}$ from both sides and dividing by $\sqrt{1 + T\xi^{-1}}$, we have

$$\begin{aligned}
\gamma_{k+1} &= \gamma_k - \frac{\sqrt{\tau_k}}{\sqrt{1 + T\xi^{-1}}} \left[\frac{\eta\gamma_k[2\xi^{-1}\sqrt{\beta\xi(1 - T\xi^{-1})} + \xi^{-1}\gamma_k]}{(\sqrt{\beta\xi(1 - T\xi^{-1})} + \gamma_k)^2 + 2\beta T} \right] \\
&= \gamma_k - (\sqrt{\beta\xi(1 - T\xi^{-1})} + \gamma_k) \left[\frac{\eta\gamma_k[2\xi^{-1}\sqrt{\beta\xi(1 - T\xi^{-1})} + \xi^{-1}\gamma_k]}{(\sqrt{\beta\xi(1 - T\xi^{-1})} + \gamma_k)^2 + 2\beta T} \right] \\
&= \gamma_k \left[1 - \eta\xi^{-1} \frac{[2\sqrt{\beta\xi(1 - T\xi^{-1})} + \gamma_k][\sqrt{\beta\xi(1 - T\xi^{-1})} + \gamma_k]}{(\sqrt{\beta\xi(1 - T\xi^{-1})} + \gamma_k)^2 + 2\beta T} \right] \\
&= \gamma_k[1 - \eta\xi^{-1}\omega_k].
\end{aligned}$$

We now show that $\omega_k = \omega(\gamma_k)$, where $\omega : (-\sqrt{\beta\xi(1 - T\xi^{-1})}, +\infty) \rightarrow (0, +\infty)$,

$$\omega(\gamma) = \frac{[2\sqrt{\beta\xi(1 - T\xi^{-1})} + \gamma][\sqrt{\beta\xi(1 - T\xi^{-1})} + \gamma]}{(\sqrt{\beta\xi(1 - T\xi^{-1})} + \gamma)^2 + 2\beta T}, \quad (40)$$

satisfies $\omega_k \in (\delta, \Delta]$ for some $\delta > 0$ depending only on γ_0 , and Δ depending only on ξ and T . This will give us linear convergence of γ_k to zero, as long as $\eta\xi^{-1} \leq 1/\Delta$, and $T\xi^{-1} < 1$.

First observe that

$$\sqrt{\beta\xi(1 - T\xi^{-1})} + \gamma_k = \frac{\sqrt{\tau_k}}{\sqrt{1 + T\xi^{-1}}} > 0.$$

Considering the translation $\bar{\omega}(\gamma) = \omega(\gamma - \sqrt{\beta\xi(1 - T\xi^{-1})}) : (0, +\infty) \rightarrow \mathbb{R}$, we can simplify

$$\bar{\omega}(\gamma) = \frac{\gamma(\gamma + \sqrt{\beta\xi(1 - T\xi^{-1})})}{\gamma^2 + 2\beta T}.$$

$\bar{\omega}$ is maximized at

$$\max_{\gamma > 0} \bar{\omega}(\gamma) = \frac{1}{2} \left(\sqrt{\frac{\xi + T}{2T}} + 1 \right) =: \Delta(\xi, T),$$

obtained at the critical point

$$\gamma = \sqrt{\frac{4\beta^2 T^2}{\beta\xi(1 - T\xi^{-1})} + 2\beta T} + \frac{2\beta T}{\sqrt{\beta\xi(1 - T\xi^{-1})}}.$$

This shows that ω is bounded above by Δ , and gives a closed form for Δ in terms of ξ and T . To show that ω is bounded below, we note that as $\gamma \rightarrow -\sqrt{\beta\xi(1 - T\xi^{-1})}$, $\omega \rightarrow 0^+$. Moreover, as $\gamma \rightarrow +\infty$, $\omega \rightarrow 1$. Therefore, under the assumption that $\gamma_k \rightarrow 0$, $\omega_k = \omega(\gamma_k)$ is bounded from below. Moreover, if the convergence is monotonic, then $\omega(\gamma_k)$ is bounded from below by $\delta = \min(\omega(\gamma_0), \omega(0)) > 0$.

As $\gamma_k \rightarrow 0$, we have that $\omega(\gamma_k) \rightarrow \omega(0)$, which takes the following form independent of β :

$$\omega(0) = \frac{2(\xi - T)}{\xi + T}. \quad (41)$$

Putting everything together, if $\eta\xi^{-1} \leq 1/\Delta$, then $1 - \eta\xi^{-1}\omega(\gamma_k) \in (0, 1 - \eta\xi^{-1}\delta)$, and thus we have that $\gamma_k \rightarrow 0$ linearly and monotonically. Moreover, the factor is $1 - \eta\xi^{-1}\frac{2(\xi-T)}{\xi+T}$ (meaning that $\gamma_{k+1}/\gamma_k \rightarrow 1 - \eta\xi^{-1}\frac{2(\xi-T)}{\xi+T}$ as $k \rightarrow \infty$). This is summarized in Theorem 2. The proof of the additional statements is as follows.

Proof of Theorem 2. The evolution of the covariance is given as above. For the linear convergence of $\tau_k^{(i)}$ to $\tau_\infty^{(i)}$, recall the ansatz

$$\sqrt{\tau_k^{(i)}} = \sqrt{\tau_\infty^{(i)}} + \sqrt{1 + T\xi^{(i)-1}}\gamma_k^{(i)}.$$

We have linear convergence of $\gamma_k^{(i)}$, given by

$$\gamma_{k+1}^{(i)} = \gamma_k^{(i)}[1 - \eta\xi^{(i)-1}\omega^{(i)}(\gamma_k^{(i)})],$$

where $\omega^{(i)}(\gamma_k) \in (\delta^{(i)}, \Delta^{(i)})$ for all k . Moreover, the linear convergence is with factor $[1 - 2\eta\xi^{(i)-1}(\xi - T)/(\xi + T)]$. Therefore, we have linear convergence of $\sqrt{\tau_k^{(i)}}$ to $\sqrt{\tau_\infty^{(i)}}$ with the same factor. Since $\sqrt{\tau_k^{(i)}} \rightarrow \sqrt{\tau_\infty^{(i)}}$ monotonically, we have that the sequence is bounded by $\max(\sqrt{\tau_0^{(i)}}, \sqrt{\tau_\infty^{(i)}})$. Therefore, we also have linear monotonic convergence of $\tau_k^{(i)}$ to $\tau_\infty^{(i)}$, with the same factor:

$$\frac{\tau_{k+1}^{(i)} - \tau_\infty^{(i)}}{\tau_k^{(i)} - \tau_\infty^{(i)}} = \frac{\sqrt{\tau_{k+1}^{(i)}} - \sqrt{\tau_\infty^{(i)}}}{\sqrt{\tau_k^{(i)}} - \sqrt{\tau_\infty^{(i)}}} \cdot \frac{\sqrt{\tau_{k+1}^{(i)}} + \sqrt{\tau_\infty^{(i)}}}{\sqrt{\tau_k^{(i)}} + \sqrt{\tau_\infty^{(i)}}} \rightarrow 1 - 2\eta\xi^{(i)-1}\frac{\xi - T}{\xi + T}.$$

The bound on the total variation follows directly from Theorem 1. The constants are

$$C = \frac{3}{2} \max_{i=1,\dots,d} \left| \frac{\tau_0^{(i)}}{\tau_\infty^{(i)}} - 1 \right|; \quad (42)$$

$$c = \max_i \left\{ 1 - \eta\xi^{(i)-1} \min \left(\omega^{(i)}(\gamma_0^{(i)}), \frac{2(\xi^{(i)} - T)}{\xi^{(i)} + T} \right) \right\}, \quad (43)$$

$$\omega^{(i)}(\gamma_0^{(i)}) = \frac{\tau_0^{(i)} + \sqrt{\tau_0^{(i)}}\beta\xi^{(i)}(1 - T^2\xi^{(i)-2})}{\tau_0^{(i)} + 2\beta T(1 + T\xi^{(i)-1})}. \quad (44)$$

□

APPENDIX C. PROOF OF LYAPUNOV CONVERGENCE FOR GAUSSIANS

Here, we demonstrate convergence of the Lyapunov function in Proposition 3. We begin by differentiating with respect to time, using Equation (32).

$$\begin{aligned}
& \frac{d}{dt} \text{Tr}((\tilde{\Sigma}_\infty^{-1} - \tilde{\Sigma}_t^{-1})^2) \\
&= 2 \text{Tr}\left(\frac{d}{dt}(\tilde{\Sigma}_\infty^{-1} - \tilde{\Sigma}_t^{-1})(\tilde{\Sigma}_\infty^{-1} - \tilde{\Sigma}_t^{-1})\right) \\
&= 2 \text{Tr}\left(\tilde{\Sigma}_t^{-1} \frac{d\tilde{\Sigma}_t}{dt} \tilde{\Sigma}_t^{-1} (\tilde{\Sigma}_\infty^{-1} - \tilde{\Sigma}_t^{-1})\right) \\
&= -2 \text{Tr}\left(\tilde{\Sigma}_t^{-1} K^{-1} \left[(\tilde{\Sigma}_\infty^{-1} - \tilde{\Sigma}_t^{-1})\Sigma_t + \Sigma_t(\tilde{\Sigma}_\infty^{-1} - \tilde{\Sigma}_t^{-1})\right] K^{-1} \tilde{\Sigma}_t^{-1} (\tilde{\Sigma}_\infty^{-1} - \tilde{\Sigma}_t^{-1})\right) \\
&= -2 \text{Tr}\left(\left[(\tilde{\Sigma}_\infty^{-1} - \tilde{\Sigma}_t^{-1})\Sigma_t + \Sigma_t(\tilde{\Sigma}_\infty^{-1} - \tilde{\Sigma}_t^{-1})\right] K^{-1} \tilde{\Sigma}_t^{-1} (\tilde{\Sigma}_\infty^{-1} - \tilde{\Sigma}_t^{-1}) \tilde{\Sigma}_t^{-1} K^{-1}\right) \\
&\stackrel{(*)}{=} -4 \text{Tr}\left(\left[(\tilde{\Sigma}_\infty^{-1} - \tilde{\Sigma}_t^{-1})\Sigma_t\right] K^{-1} \tilde{\Sigma}_t^{-1} (\tilde{\Sigma}_\infty^{-1} - \tilde{\Sigma}_t^{-1}) \tilde{\Sigma}_t^{-1} K^{-1}\right) \\
&= -4 \text{Tr}\left((\tilde{\Sigma}_\infty^{-1} - \tilde{\Sigma}_t^{-1})(K \tilde{\Sigma}_t K - 2\beta T K) K^{-1} \tilde{\Sigma}_t^{-1} (\tilde{\Sigma}_\infty^{-1} - \tilde{\Sigma}_t^{-1}) (\tilde{\Sigma}_t^{-1} K^{-1})\right) \\
&= -4 \text{Tr}\left((\tilde{\Sigma}_\infty^{-1} - \tilde{\Sigma}_t^{-1})(K - 2\beta T \tilde{\Sigma}_t^{-1})(\tilde{\Sigma}_\infty^{-1} - \tilde{\Sigma}_t^{-1})(\tilde{\Sigma}_t^{-1} K^{-1})\right).
\end{aligned}$$

In (*), we used that $\text{Tr}(AB) = \text{Tr}(A^\top B^\top)$. We now aim to bound the final term in the product, $\tilde{\Sigma}_t^{-1} K^{-1}$. To do this, we use the following results from linear algebra.

Proposition 4. *Suppose A is Hermitian and positive definite, and B is square of the same dimensions, satisfying:*

$$x^* B x > \frac{1}{2} x^* A x, \quad \forall x \neq 0. \quad (45)$$

Then $\rho(I - B^{-1}A) < 1$.

Proof. Note that the positivity condition gives that B is invertible. The eigenvalues of $I - B^{-1}A$ and $I - A^{1/2}B^{-1}A^{1/2}$ are equal. From the quadratic form inequality, we have for any (complex) $z \neq 0$,

$$\Re(z^* A^{-1/2} B A^{-1/2} z) > \frac{1}{2} z^* z.$$

Therefore the real part of each eigenvalue of $A^{-1/2} B A^{-1/2}$ satisfies $\Re \lambda_i(A^{-1/2} B A^{-1/2}) > 1/2$.

Now note that $1 - 1/z$ is a conformal mapping, taking the half-plane $\Re(z) > 1/2$ to the unit disk $|w| < 1$. Thus the spectrum of $I - A^{1/2}B^{-1}A^{1/2}$ lies in the unit disk and we conclude. $\square$

Using Proposition 4 with $A = 4\beta T K^{-1}$ and $B = \tilde{\Sigma}_t$, we satisfy the assumptions of the proposition, since:

$$\tilde{\Sigma}_t = K^{-1} \Sigma_t K^{-1} + 2\beta T K^{-1} \succeq \frac{1}{2} 4\beta T K^{-1}. \quad (46)$$

Thus, we have the following bound on the spectral radius,

$$\begin{aligned}
& \rho(I - 4\beta T \tilde{\Sigma}_t^{-1} K^{-1}) < 1 \\
& \Rightarrow \rho\left(\frac{I}{4\beta T} - \tilde{\Sigma}_t^{-1} K^{-1}\right) < \frac{1}{4\beta T}.
\end{aligned}$$

In fact, we can conclude slightly more if we can bound the maximum and minimum eigenvalues of $\tilde{\Sigma}_t$. Since we can bound $I \preceq K \preceq (1+c)I$, where $T = c\lambda_{\min}(\Sigma)$ with $c \in (0, 1)$, we have that the minimum and maximum eigenvalues of $K^{-1/2} \tilde{\Sigma}_t^{-1} K^{-1/2}$ are bounded as

$$\frac{1}{1+c} \lambda_{\min}(\tilde{\Sigma}_t) \leq \lambda_{\min}(K^{-1/2} \tilde{\Sigma}_t K^{-1/2}) \leq \lambda_{\max}(K^{-1/2} \tilde{\Sigma}_t K^{-1/2}) \leq \lambda_{\max}(\tilde{\Sigma}_t).$$

Therefore, a bound for the eigenvalues of $I - \tilde{\Sigma}_t^{-1} K^{-1}$ is given by

$$\lambda_{\max}(\frac{I}{4\beta T} - \tilde{\Sigma}_t^{-1}K^{-1}) = \frac{1}{4\beta T} - \frac{1}{\lambda_{\max}(K^{-1/2}\tilde{\Sigma}_tK^{-1/2})} \leq \frac{1}{4\beta T} - \frac{1}{\lambda_{\max}(\tilde{\Sigma}_t)}, \quad (47)$$

$$\lambda_{\min}(\frac{I}{4\beta T} - \tilde{\Sigma}_t^{-1}K^{-1}) = \frac{1}{4\beta T} - \frac{1}{\lambda_{\min}(K^{-1/2}\tilde{\Sigma}_tK^{-1/2})} \geq \frac{1}{4\beta T} - \frac{1+c}{\lambda_{\min}(\tilde{\Sigma}_t)}, \quad (48)$$

$$\rho(\frac{I}{4\beta T} - \tilde{\Sigma}_t^{-1}K^{-1}) = \max(|\lambda_{\min}|, |\lambda_{\max}|) < \frac{1}{4\beta T}. \quad (49)$$

To turn this bound on $\tilde{\Sigma}_t^{-1}K^{-1}$ to a bound on the derivative of the Lyapunov function, we need the following trace inequality (Horn & Johnson, 2012; Baumgartner, 2011).

Proposition 5 (Hölder's inequality for trace). *Let A, B be (complex) square matrices, with absolute values $|A| = (A^*A)^{1/2}$, $|B| = (B^*B)^{1/2}$. Then for $1 \leq p, q \leq \infty$ satisfying $p^{-1} + q^{-1} = 1$, the trace inequality holds:*

$$|\text{Tr}(A^*B)| \leq \|A\|_p \|B\|_q, \quad (50)$$

where $\|\cdot\|_p$ are the Schatten p -norms defined as follows, where σ_i are the singular values

$$\begin{cases} \|A\|_p^p = \sum \sigma_i^p(A), & 1 \leq p < \infty, \\ \|A\|_\infty = \sup \sigma_i(A), & p = \infty. \end{cases}$$

In particular, if A is symmetric and positive-definite, we have, where ρ denotes the spectral radius,

$$|\text{Tr}(AB)| \leq \|A\|_1 \|B\|_\infty = \text{Tr}(A)\rho(B).$$

Using these two results, we expand the derivative of the Lyapunov function, noting that $(\tilde{\Sigma}_\infty^{-1} - \tilde{\Sigma}_t^{-1})(K - 2\beta T\tilde{\Sigma}_t^{-1})(\tilde{\Sigma}_\infty^{-1} - \tilde{\Sigma}_t^{-1})$ is positive definite (since $K - 2\beta T\tilde{\Sigma}_t^{-1} \succeq 0$):

$$\begin{aligned} & \frac{d}{dt} \text{Tr}((\tilde{\Sigma}_\infty^{-1} - \tilde{\Sigma}_t^{-1})^2) \\ &= -4 \text{Tr}((\tilde{\Sigma}_\infty^{-1} - \tilde{\Sigma}_t^{-1})(K - 2\beta T\tilde{\Sigma}_t^{-1})(\tilde{\Sigma}_\infty^{-1} - \tilde{\Sigma}_t^{-1})(\tilde{\Sigma}_t^{-1}K^{-1})) \\ &= -4 \text{Tr}((\tilde{\Sigma}_\infty^{-1} - \tilde{\Sigma}_t^{-1})(K - 2\beta T\tilde{\Sigma}_t^{-1})(\tilde{\Sigma}_\infty^{-1} - \tilde{\Sigma}_t^{-1})(\tilde{\Sigma}_t^{-1}K^{-1} - \frac{I}{4\beta T})) \\ &\quad - 4 \text{Tr}((\tilde{\Sigma}_\infty^{-1} - \tilde{\Sigma}_t^{-1})(K - 2\beta T\tilde{\Sigma}_t^{-1})(\tilde{\Sigma}_\infty^{-1} - \tilde{\Sigma}_t^{-1})(\frac{I}{4\beta T})) \\ &\leq -4 \text{Tr}((\tilde{\Sigma}_\infty^{-1} - \tilde{\Sigma}_t^{-1})(K - 2\beta T\tilde{\Sigma}_t^{-1})(\tilde{\Sigma}_\infty^{-1} - \tilde{\Sigma}_t^{-1}))\rho(\tilde{\Sigma}_t^{-1}K^{-1} - \frac{I}{4\beta T}) \\ &\quad - 4 \text{Tr}((\tilde{\Sigma}_\infty^{-1} - \tilde{\Sigma}_t^{-1})(K - 2\beta T\tilde{\Sigma}_t^{-1})(\tilde{\Sigma}_\infty^{-1} - \tilde{\Sigma}_t^{-1}))(\frac{1}{4\beta T}) \\ &\leq -4 \text{Tr}((\tilde{\Sigma}_\infty^{-1} - \tilde{\Sigma}_t^{-1})(K - 2\beta T\tilde{\Sigma}_t^{-1})(\tilde{\Sigma}_\infty^{-1} - \tilde{\Sigma}_t^{-1}))(\frac{1}{4\beta T} - \rho(\frac{I}{4\beta T} - \tilde{\Sigma}_t^{-1}K^{-1})) \\ &< 0. \end{aligned}$$

This bound can be slightly refined to yield linear convergence. If $(\frac{1}{4\beta T} - \rho(\frac{I}{4\beta T} - \tilde{\Sigma}_t^{-1}K^{-1}))$ is bounded away from zero and $K - 2\beta T\tilde{\Sigma}_t^{-1}$ has (positive) eigenvalues also bounded away from zero, then we have linear convergence of $\|\tilde{\Sigma}_\infty^{-1} - \tilde{\Sigma}_t^{-1}\|_F^2$ to zero. Both of these assumptions can be justified using the equivalence of matrix norms, and bootstrapping convergence of $\tilde{\Sigma}_t$ to $\tilde{\Sigma}_\infty$.