
Soft-constrained Schrödinger Bridge: a Stochastic Control Approach

Jhanvi Garg
Texas A&M University

Xianyang Zhang
Texas A&M University

Quan Zhou
Texas A&M University

Abstract

Schrödinger bridge can be viewed as a continuous-time stochastic control problem where the goal is to find an optimally controlled diffusion process whose terminal distribution coincides with a pre-specified target distribution. We propose to generalize this problem by allowing the terminal distribution to differ from the target but penalizing the Kullback-Leibler divergence between the two distributions. We call this new control problem soft-constrained Schrödinger bridge (SSB). The main contribution of this work is a theoretical derivation of the solution to SSB, which shows that the terminal distribution of the optimally controlled process is a geometric mixture of the target and some other distribution. This result is further extended to a time series setting. One application is the development of robust generative diffusion models. We propose a score matching-based algorithm for sampling from geometric mixtures and showcase its use via a numerical example for the MNIST data set.

1 INTRODUCTION

1.1 Schrödinger Bridge and Its Applications

Let $X = (X_t)_{0 \leq t \leq T}$ be a diffusion process over the finite time interval $[0, T]$ with initial distribution μ_0 . Schrödinger bridge seeks an optimal steering of X towards a pre-specified terminal distribution μ_T such that the resulting controlled process is closest to X in terms of Kullback-Leibler (KL) divergence (Schrödinger, 1931, 1932). Under certain regularity conditions, the optimally controlled process is another diffusion with the same diffusion coefficients

as X but an additional drift term. This result has been obtained via different approaches and at varying levels of generality, and among the seminal works are Fortet (1940); Beurling (1960); Jamison (1975); Föllmer (1988); Dai Pra (1991). For comprehensive reviews detailing the historical development, we refer readers to Léonard (2013) and Chen et al. (2021c).

The recent generative modeling literature has seen a surge in the use of Schrödinger bridge. In these applications, μ_0 is typically some distribution that is easy to sample from, and μ_T is the unknown distribution of a given data set. By numerically approximating the solution to the Schrödinger bridge problem, one can generate samples from μ_T (i.e., synthetic data points that resemble the original data set). One such algorithm is presented by De Bortoli et al. (2021), who proposed to calculate the Schrödinger bridge by using a score matching approximation to the iterative proportional fitting procedure (Deming and Stephan, 1940). Concurrently, Wang et al. (2021) developed a two-stage method where an auxiliary Schrödinger bridge is run first to generate samples from a smoothed version of μ_T , and the second Schrödinger bridge transports these samples towards μ_T . Both approaches generalize the denoising diffusion model methods of Ho et al. (2020) and Song et al. (2021). Some other recent developments in this area include Chen et al. (2021a); Song (2022); Peluchetti (2023); Richter et al. (2023); Winkler et al. (2023); Hamdouche et al. (2023).

Though not the focus of this work, Schrödinger bridge sampling methods can also be used when samples from μ_T are not available but μ_T is known up to a normalizing constant; see, e.g., Huang et al. (2021); Zhang and Chen (2021); Vargas et al. (2022), and see Heng et al. (2024) for a recent review. For the connections between Schrödinger bridge, optimal transport and variational inference, see, e.g., Chen et al. (2016b, 2021b); Tzen and Raginsky (2019).

1.2 Overview of This Work

The main contribution of this paper is the theoretical development of a generalized Schrödinger bridge problem, which we call soft-constrained Schrödinger bridge

(SSB). We take the stochastic control approach that was employed by Mikami (1990); Dai Pra (1991); Pra and Pavon (1990) for studying the original Schrödinger bridge problem (see Problem 1). In SSB, the terminal distribution of the controlled process does not need to precisely match μ_T but needs to be close to μ_T in terms of KL divergence. Formally, SSB differs from the original problem in that we replace the hard constraint on the terminal distribution with an additional cost term, parameterized by β , in the objective function to be minimized (see Problem 2). A larger β forces the terminal distribution of the controlled process to be closer to μ_T . We rigorously find the solution to SSB and the expression for the drift term of the optimally controlled process. We show that SSB generalizes Schrödinger bridge in the sense that as $\beta \rightarrow \infty$, the solution to the former coincides with the solution to the latter. An important implication of our results is that the terminal distribution of the controlled process should be a geometric mixture of μ_T and some other distribution; when μ_0 is a Dirac measure, the other distribution is $\mathcal{L}aw(X_T)$ (i.e., the terminal distribution of the uncontrolled process). We further extend our results to a time series generalization of SSB, where we are interested in modifying the joint distribution of $(X_{t_1}, \dots, X_{t_N})$ for $0 < t_1 < \dots < t_N = T$.

SSB can be used as a theoretical foundation for developing more flexible and robust sampling methods. First, when the KL divergence between μ_T and $\mathcal{L}aw(X_T)$ is infinite, Schrödinger bridge does not admit a solution, while SSB always does. A toy example illustrating the consequences of this result is given in Example 1. More importantly, β acts as a regularization parameter preventing the algorithm from overfitting to μ_T , which is crucial for some generative modeling tasks such as fine-tuning with limited data (Moon et al., 2022). In such applications, μ_T contains information from a small or noisy data set, and one wants to improve the sample quality by harnessing knowledge from a large high-quality reference data set. To achieve this, we can train the uncontrolled process X in SSB using the reference data set and then tune the value of β . We present a simple normal mixture example illustrating the effect of β (see Example 4). For a more realistic example in generative modeling of images, we use the MNIST data set and consider the task of generating new images of digit 8. We assume that the training data set only has 50 noisy images of digit 8, but we can use the data set of all the other digits as reference. As suggested by our theoretical findings, we can train a Schrödinger bridge targeting a geometric mixture of the distributions of the two data sets. Such a Schrödinger bridge cannot be learned by existing methods, and to address this, we propose a new score matching algorithm that utilizes importance sampling.

We show that this approach yields high-quality images of digit 8 when β is properly chosen.

The paper is structured as follows. In Section 2, we present the stochastic control formulation of the SSB problem, and we derive its solution when μ_0 is a Dirac measure. The solution to SSB for general initial conditions is obtained in Section 3, which involves solving a generalized Schrödinger system. Section 4 extends the results to the time series setting. In Section 5, we present a new algorithm for robust generative modeling and demonstrate its use via the MNIST data set. Proofs and auxiliary results are deferred to Appendix.

1.3 Related Literature

Our development of SSB builds upon the work of Dai Pra (1991), which formulates Schrödinger bridge as a stochastic control problem and derives the solution using the logarithmic transformation technique pioneered by Fleming (Fleming, 1977, 2005; Fleming and Rishel, 2012) and the result of Jamison (1975). The time series SSB problem is a generalization of the work of Hamdouche et al. (2023), who extended the original Schrödinger bridge problem to the time series setting but only considered the special case where μ_0 is a Dirac measure. Pavon and Wakolbinger (1991); Blaquiére (1992) adopted an alternative stochastic control approach to studying Schrödinger bridge, which was rooted in the same logarithmic transformation and also considered in Tzen and Raginsky (2019); Berner et al. (2022). This approach can be applied to the SSB problem as well, but it requires the use of verification theorem.

Motivated by robust network routing, a discrete version of the SSB problem was proposed and solved in Chen et al. (2019), where X is a discrete-time non-homogeneous Markov chain with finite state space. The techniques used in this paper are very different, and to our knowledge, the continuous-time SSB problem has not been addressed in the literature.

2 PROBLEM FORMULATION

Let μ_0, μ_T be two probability distributions on $\mathbb{R}^d$ such that $\int x^2 \mu_0(dx) < \infty$ and $\mu_T \ll \lambda$, where λ denotes the Lebesgue measure. Denote the density of μ_T by $f_T = d\mu_T/d\lambda$. Let $(\Omega, \mathcal{F}, \mathbb{P})$ be a probability space, on which we define a standard d -dimensional Brownian motion $W = (W_t)_{t \geq 0}$ and a random vector ξ that is independent of W and has distribution μ_0 . We will always use $X = (X_t)_{0 \leq t \leq T}$ to denote a weak solution to the following stochastic differential equation (SDE)

$$X_0 = \xi, \text{ and } dX_t = b(X_t, t)dt + \sigma dW_t \quad (2.1)$$

for $t \in [0, T]$, where $b: \mathbb{R}^d \times [0, T] \rightarrow \mathbb{R}^d$ and $\sigma \in (0, \infty)$. Given a control $u = (u_t)_{0 \leq t \leq T}$, define the controlled diffusion process by $X_0^u = \xi$ and

$$dX_t^u = [b(X_t^u, t) + u_t] dt + \sigma dW_t. \quad (2.2)$$

We say a control u is admissible if (i) u_t is measurable with respect to $\sigma((X_s^u)_{0 \leq s \leq t})$, (ii) the SDE (2.2) admits a weak solution, and (iii) $\mathbb{E} \int_0^T \|u_t\|^2 dt < \infty$, where $\|\cdot\|$ denotes the L^2 norm. Denote the set of all admissible controls by $\mathcal{U}$. Note that the initial distributions of both X and X^u are always fixed to be μ_0 . For ease of presentation, throughout the paper we adopt the following regularity assumption on b , which was also used by Jamison (1975); Dai Pra (1991):

Assumption 1. For each $1 \leq i \leq d$, b_i is bounded and continuous in $\mathbb{R}^d \times [0, T]$ and is Hölder continuous in x , uniformly with respect to $(x, t) \in \mathbb{R}^d \times [0, T]$.

Under Assumption 1, X has a transition density function $p(x, t | y, s)$; that is, for any $0 \leq s < t \leq T$, $y \in \mathbb{R}^d$, and Borel set A in $\mathbb{R}^d$,

$$\mathbb{E}[X_t \in A | X_s = y] = \int_A p(x, t | y, s) \lambda(dx). \quad (2.3)$$

We will use dx as a shorthand for $\lambda(dx)$. Moreover, by Girsanov theorem, Assumption 1 implies that the probability measures induced by $(X_t)_{0 \leq t \leq T}$ and $(B_t)_{0 \leq t \leq T}$ are equivalent, where $B_t = \xi + \sigma W_t$. Hence, for any $t > s \geq 0$, $p(x, t | y, s)$ is strictly positive and $\mathcal{L}aw(X_t) \approx \lambda$ (i.e., two measures are equivalent). The role of Assumption 1 in our theoretical results will be further discussed in Remark 4.

Schrödinger bridge aims to find a minimum-energy modification of the dynamics of X so that its terminal distribution coincides with a pre-specified distribution μ_T , where “energy” is measured by KL divergence. Given σ -finite measures ν and μ such that $\nu \ll \mu$, we use $\mathcal{D}_{\text{KL}}(\nu, \mu) = \int \log(\frac{d\nu}{d\mu}) d\nu$ to denote the KL divergence; if $\nu \not\ll \mu$, define $\mathcal{D}_{\text{KL}}(\nu, \mu) = \infty$. Dai Pra (1991) considered the following stochastic control formulation of Schrödinger bridge.

Problem 1 (Schrödinger bridge). Let $\mathcal{U}_0 = \{u \in \mathcal{U} : \mathcal{L}aw(X_T^u) = \mu_T\}$ where $(X_t^u)_{0 \leq t \leq T}$ is defined in (2.2). Find $V = \inf_{u \in \mathcal{U}_0} J(u)$, where

$$J(u) = \mathbb{E} \int_0^T \frac{\|u_t\|^2}{2\sigma^2} dt, \quad (2.4)$$

and find the optimal control u^* such that $J(u^*) = V$.

Remark 1. Let $\mathbb{P}_X$ (resp. $\mathbb{P}_X^u$) denote the probability measure induced by X (resp. X^u) on the space of continuous functions on $[0, T]$. For any admissible control $u \in \mathcal{U}$, Girsanov theorem implies that $J(u) = \mathcal{D}_{\text{KL}}(\mathbb{P}_X^u, \mathbb{P}_X)$.

Problem 1 has been well studied in the literature. In the special case where μ_0 is a Dirac measure, the solution can be succinctly described, and V is just the KL divergence between two probability distributions; we recall this in Theorem 1 below. In this paper, ∇ always denotes differentiation with respect to x .

Theorem 1 (Dai Pra (1991)). Let μ_0 be the Dirac measure such that $\mu_0(\{x_0\}) = 1$ for some $x_0 \in \mathbb{R}^d$, and let X be a weak solution to (2.1). Assume $\mathcal{D}_{\text{KL}}(\mu_T, \mathcal{L}aw(X_T)) < \infty$. For Problem 1, the optimal control is given by $u_t^* = \sigma^2 \nabla \log h(X_t^{u^*}, t)$, where

$$\begin{aligned} h(x, t) &= \int p(z, T | x, t) \frac{f_T(z)}{p(z, T | x_0, 0)} dz \\ &=: \mathbb{E} \left[\frac{f_T(X_T)}{p(X_T, T | x_0, 0)} \mid X_t = x \right]. \end{aligned}$$

Moreover, $J(u^*) = \mathcal{D}_{\text{KL}}(\mu_T, \mathcal{L}aw(X_T))$.

Remark 2. If $\mathcal{D}_{\text{KL}}(\mu_T, \mathcal{L}aw(X_T)) = \infty$, then Problem 1 does not admit a solution in the sense that no admissible control u can yield $\mathcal{L}aw(X_T^u) = \mu_T$.

We propose a relaxed stochastic control formulation of the Schrödinger bridge problem by allowing the distribution of X_T^u to be different from μ_T .

Problem 2 (Soft-constrained Schrödinger bridge). For $\beta > 0$, find $V = \inf_{u \in \mathcal{U}} J_\beta(u)$, where

$$\begin{aligned} J_\beta(u) &= \beta \mathcal{D}_{\text{KL}}(\mathcal{L}aw(X_T^u), \mu_T) + \mathbb{E} \int_0^T \frac{\|u_t\|^2}{2\sigma^2} dt, \end{aligned} \quad (2.5)$$

and find the optimal control u^* such that $J_\beta(u^*) = V$.

Problem 2 (i.e., the SSB problem) replaces the hard constraint $\mathcal{L}aw(X_T^u) = \mu_T$ in Problem 1 with a soft constraint parameterized by β . When $\beta = 0$, it is clear that the optimal control u^* for Problem 2 is $u_t^* \equiv 0$. As $\beta \rightarrow \infty$, the law of X_T^u is forced to agree with μ_T , and we will see in Theorem 2 that the optimal control for Problem 2 converges to that for Problem 1.

Before we try to solve Problem 2 in full generality, we make a remark on how Problem 1 can be simplified. In the literature, Problem 1 is often called the dynamic Schrödinger bridge problem. Since the objective function (2.4) is the KL divergence between the laws of the controlled and uncontrolled processes (recall Remark 1), we can use an additive property of KL divergence to reduce Problem 1 to a static version (Léonard, 2013), where one only needs to find a joint distribution π with marginals μ_0 and μ_T that minimizes $\mathcal{D}_{\text{KL}}(\pi, \mathcal{L}aw(X_0, X_T))$. Although this property is not directly used in this paper, the insight from this observation underpins our stochastic control analysis of SSB. In particular, when μ_0 is a Dirac measure,

the solution to Problem 2 can be obtained by a simple argument which reduces the problem to optimizing over the distribution of X_T^u instead of over the distribution of the whole process $(X_t^u)_{0 \leq t \leq T}$.

Theorem 2. *Let μ_0 be as given in Theorem 1. For Problem 2, the optimal control is given by $u_t^* = \sigma^2 \nabla \log h(X_t^{u^*}, t)$, where*

$$h(x, t) = h(x, t; \beta) \\ = C^{-1} \int p(z, T | x, t) \cdot \left(\frac{f_T(z)}{p(z, T | x_0, 0)} \right)^{\beta/(1+\beta)} dz,$$

and $C = \int f_T(x)^{\beta/(1+\beta)} p(x, T | x_0, 0)^{1/(1+\beta)} dx$. Moreover, $J_\beta(u^*) = -(1 + \beta) \log C \in [0, \infty)$, and

$$\lim_{\beta \rightarrow \infty} h(x, t; \beta) = \int p(z, T | x, t) \frac{f_T(z)}{p(z, T | x_0, 0)} dz.$$

Proof. If $\mathcal{L}aw(X_T^u) \not\ll \mu_T$, then u cannot be optimal since $J_\beta(u) = \infty$. Now fix an arbitrary $u \in \mathcal{U}$ such that $\mathcal{L}aw(X_T^u) = \mu \ll \mu_T$. Letting $J(u)$ be as given in (2.4), we have

$$J_\beta(u) = \beta \mathcal{D}_{\text{KL}}(\mu, \mu_T) + J(u) \\ \geq \beta \mathcal{D}_{\text{KL}}(\mu, \mu_T) + \mathcal{D}_{\text{KL}}(\mu, \mathcal{L}aw(X_T)) \\ =: \tilde{J}_\beta(\mu)$$

where the inequality follows from Theorem 1. Since μ_T has density f_T and $\mathcal{L}aw(X_T)$ has density $p(x, T | x_0, 0)$, we can apply Lemma B2 in Appendix to get $\inf_\mu \tilde{J}_\beta(\mu) = -(1 + \beta) \log C \in [0, \infty)$, where the infimum is taken over all probability measures on $\mathbb{R}^d$ and is attained at μ^* such that

$$\frac{d\mu^*}{d\lambda}(x) = C^{-1} f_T(x)^{\beta/(1+\beta)} p(x, T | x_0, 0)^{1/(1+\beta)}.$$

The convergence of $h(x, t; \beta)$ as $\beta \rightarrow \infty$ also follows from Lemma B2.

It only remains to prove that $J_\beta(u^*) = -(1 + \beta) \log C$. Observe that we can rewrite h as

$$h(x, t) = \int p(z, T | x, t) \frac{(d\mu^*/d\lambda)(z)}{p(z, T | x_0, 0)} dz.$$

Hence, Theorem 1 implies that u^* is also the solution to Problem 1 where the terminal constraint is given by μ^* . So Theorem 1 yields that $J_\beta(u^*) = \tilde{J}_\beta(\mu^*)$, which proves the claim. $\square$

Remark 3. When b is constant, the transition density $p(x, t | x_0, 0)$ is easy to evaluate. If f_T is known up to a normalizing constant, one can then use the Monte Carlo sampling scheme proposed by (Huang et al., 2021) to approximate the drift $b + \sigma^2 \nabla \log h(x, t)$ and simulate the controlled diffusion process (2.2). We

describe this method and generalize it using importance sampling techniques in Appendix A. More sophisticated score-based sampling schemes can also be applied (Heng et al., 2024).

One difference between Theorem 1 and Theorem 2 is that the condition $\mathcal{D}_{\text{KL}}(\mu_T, \mathcal{L}aw(X_T)) < \infty$ is not required for solving Problem 2. We give a toy example illustrating the importance of this difference.

Example 1. Consider $b \equiv 0$, $T = 1$, $x_0 = 0$ and μ_T being the Cauchy distribution. Then, $\mathcal{L}aw(X_T)$ is just the normal distribution with mean zero and covariance $\sigma^2 I$, and we have $\mathcal{D}_{\text{KL}}(\mu_T, \mathcal{L}aw(X_T)) = \infty$. Problem 1 does not admit a solution in this case, but Problem 2 has a solution for any $\beta \in [0, \infty)$ and the associated optimal control has finite energy cost. In Appendix A.2, we simulate the solution to Problem 2 with μ_T being the Cauchy distribution. We find that when $\beta = \infty$, the numerical scheme is unstable and fails to capture the heavy tails of the Cauchy distribution. In contrast, using a finite value of β significantly stabilizes the algorithm.

3 SOLUTION TO SSB

When μ_0 is not a Dirac measure, the solution to the Schrödinger bridge problem is more difficult to describe and is characterized by the so-called Schrödinger system (Léonard, 2013, Theorem 2.8). In this section, we prove that the solution to Problem 2 can be obtained in a similar way, but the Schrödinger system for Problem 2 now depends on β .

The main idea behind our approach is to first show that the optimal control must belong to a small class parameterized by a function g and then use an argument similar to the proof of Theorem 2 to determine the choice of g . To introduce this class of controls, for each measurable function $g: \mathbb{R}^d \rightarrow [0, \infty)$, let $\mathcal{T}g$ denote the function on $\mathbb{R}^d \times [0, T]$ given by

$$(\mathcal{T}g)(x, t) = \sigma^2 \nabla \log h(x, t),$$

where

$$h(x, t) = \mathbb{E}[g(X_T) | X_t = x]. \quad (3.1)$$

Let $\mathcal{U}_1 = \{u \in \mathcal{U}: u_t = (\mathcal{T}g)(X_t^u, t) \text{ for some } g \geq 0\}$ denote the set of all controls that are constructed by this logarithmic transformation. We present in Theorem B3 in Appendix some well-known results about the controlled SDE (2.2) with $u \in \mathcal{U}_1$; in particular, part (iv) of the theorem shows that such a process is a Doob's h -path process (Doob, 1959). Theorem B3 is largely adapted from Theorem 2.1 of Dai Pra (1991), and similar results are extensively documented in the literature (Jamison, 1975; Fleming and Sheu, 1985;

Föllmer, 1988; Doob, 1984; Fleming, 2005). We can now prove a key lemma.

Lemma 3. *Let u be an admissible control. Let h be as given in (3.1) for some measurable $g \geq 0$ such that $Eg(X_T) < \infty$ and $h > 0$ on $\mathbb{R}^d \times [0, T)$. We have*

$$J_\beta(u) \geq \beta \mathcal{D}_{\text{KL}}(\text{Law}(X_T^u), \mu_T) + E[\log g(X_T^u)] - \int \log h(x, 0) \mu_0(dx),$$

where J_β is defined in (2.5). The equality holds when $u_t = (\mathcal{T}g)(X_t^u, t)$.

Proof. See Appendix C. $\square$

Remark 4. Assumption 1 guarantees that the SDE (2.1) admits a unique (in law) weak solution. More importantly, in the proof of Theorem B3 (which is used to derive Lemma 3), Assumption 1 is used to ensure that SDE (2.1) has a transition density function such that the function h defined in (3.1) is sufficiently smooth and satisfies $\frac{\partial h}{\partial t} + \mathcal{L}h = 0$, where $\mathcal{L}$ denotes the generator of X (see Theorem B3). This condition can be relaxed; see Friedman (1975) and Karatzas and Shreve (2012, Chap. 5.7) for details.

Observe that in the bound given in Lemma 3, the term $\int \log h(x, 0) \mu_0(dx)$ is independent of the control u , and the other terms depend on u only through the distribution of X_T^u . This implies that among all the admissible controls that result in the same distribution of X_T^u , the cost J_β is minimized by some $u \in \mathcal{U}_1$. We can now prove the main theoretical result of this work in Theorem 4. The existence of the solution will be considered later in Theorem 5.

Theorem 4. *Suppose there exist σ -finite measures ν_0, ν_T such that $\nu_0 \approx \mu_0, \nu_T \approx \mu_T$ and*

$$\frac{d\mu_0}{d\nu_0}(y) = \int p(x, T | y, 0) \nu_T(dx), \quad (3.2)$$

$$\frac{d\mu_T}{d\lambda}(x) = \rho_T(x)^{\frac{1+\beta}{\beta}} \int p(x, T | y, 0) \nu_0(dy), \quad (3.3)$$

where $\rho_T = d\nu_T/d\lambda$ and the transition density p is defined in (2.3). Assume $\int (d\mu_0/d\nu_0) d\mu_0 < \infty$. Then $u_t^* = \sigma^2 \nabla \log h(X_t^{u^*}, t)$ solves Problem 2, where $h(x, t) = E[\rho_T(X_T) | X_t = x]$. Moreover, $J_\beta(u^*) = -\mathcal{D}_{\text{KL}}(\mu_0, \nu_0) \in [0, \infty)$.

Proof. See Appendix C. $\square$

Remark 5. As we derive in the proof, the terminal distribution of the optimally controlled process is still a geometric mixture of two distributions. Explicitly, its density is proportional to

$$f_T(x)^{\beta/(1+\beta)} \left(\int p(x, T | y, 0) \nu_0(dy) \right)^{1/(1+\beta)},$$

where we recall f_T is the density of μ_T .

Remark 6. The assumption $\int (d\mu_0/d\nu_0) d\mu_0 < \infty$ guarantees that we can use Lemma 3 and Theorem B3 and that $J_\beta(u^*) < \infty$; it is also used in Dai Pra (1991). Observe that $u_t^* = \sigma^2 \nabla \log h(X_t^{u^*}, t)$ is invariant to the scaling of the function ρ_T and thus also the scaling of ν_0 and ν_T . This suggests that the system defined by (3.2) and (3.3) can be generalized as follows. Let $a > 0$ and σ -finite measures $\nu_0 = \nu_0(a), \nu_T = \nu_T(a)$ be the solution to the following system

$$\frac{d\mu_0}{d\nu_0}(y) = \int p(x, T | y, 0) \nu_T(dx), \quad (3.4)$$

$$\frac{d\mu_T}{d\lambda}(x) = (a\rho_T(x))^{\frac{1+\beta}{\beta}} \int p(x, T | y, 0) \nu_0(dy), \quad (3.5)$$

where $\rho_T = d\nu_T/d\lambda$. We can use essentially the same argument to show that the choice $h(x, t) = E[\rho_T(X_T) | X_t = x]$ is optimal, but now we have

$$J_\beta(u^*) = -(1 + \beta) \log a - \mathcal{D}_{\text{KL}}(\mu_0, \nu_0(a)).$$

This is the same as that given in Theorem 4. Indeed, if (ν_0^*, ν_T^*) is a solution to (3.2) and (3.3), then the solution to (3.4) and (3.5) is given by $d\nu_0(a) = a^{1+\beta} d\nu_0^*$ and $d\nu_T(a) = a^{-(1+\beta)} d\nu_T^*$.

Example 2. Theorem 2 can be obtained from Theorem 4 as a special case. If μ_0 is a Dirac measure with $\mu_0(\{x_0\}) = 1$, one can check that the solution to (3.2) and (3.3) is given by $\nu_0 = \mu_0$ and

$$\rho_T(x) = \left(\frac{f_T(x)}{p(x, T | x_0, 0)} \right)^{\beta/(1+\beta)}.$$

Example 3. Suppose $b \equiv 0$, and let ϕ_σ denote the density function of the normal distribution with mean 0 and covariance matrix $\sigma^2 I$. We have $p(x, T | y, 0) = \phi_{\sigma\sqrt{T}}(x - y)$. Assume $\mu_0 \ll \lambda$ has density f_0 , and suppose that f_0, f_T satisfy

$$f_0(y) = c^{-1} \int \phi_{\sigma\sqrt{T}}(x - y) f_T(x)^{\frac{\beta}{1+\beta}} dx,$$

where $c = \int f_T(x)^{\beta/(1+\beta)} dx$ is the normalizing constant assumed to be finite. A routine calculation using $\int \phi_{\sigma\sqrt{T}}(x - y) dy = 1$ can verify that the solution to (3.2) and (3.3) is given by

$$\frac{d\nu_0}{d\lambda} = c^{-(1+\beta)}, \quad \rho_T(x) = c^\beta f_T(x)^{\beta/(1+\beta)}.$$

According to Remark 6, by choosing $a = c$ in (3.5), we can also replace ν_0, ν_T by $\nu_0(a), \nu_T(a)$, where $\nu_0(a)$ coincides with λ and $\nu_T(a)$ is a probability distribution with density $c^{-1} f_T^{\beta/(1+\beta)}$. For the original Schrödinger bridge problem (i.e., $\beta = \infty$), this solution has been used in developing efficient generative sampling methods (Wang et al., 2021; Berner et al., 2022).

Chen et al. (2019) studied a matrix optimal transport problem which can be seen as the discrete analogue to Problem 2, and they proved the solution to the corresponding Schrödinger system admits a unique solution. The main idea is to show that the solution can be characterized as the fixed point of some operator with respect to the Hilbert metric (Bushell, 1973), a technique that has been widely used in the literature on Schrödinger system (Fortet, 1940; Georgiou and Pavon, 2015; Chen et al., 2016a; Essid and Pavon, 2019; Deligiannidis et al., 2024). For the Schrödinger system defined by (3.2) and (3.3), an argument based on the Hilbert metric can also be applied to prove the existence and uniqueness of the solution when μ_0, μ_T are absolutely continuous and have compact support.

Theorem 5. *Let K_0 (resp. K_T) denote the support of μ_0 (resp. μ_T). Assume that (i) $K_0, K_T \subset \mathbb{R}^d$ are compact, (ii) $f_0 = d\mu_0/d\lambda$, $f_T = d\mu_T/d\lambda$ exist, where λ is the Lebesgue measure, and (iii) $(y, x) \mapsto p(x, T | y, 0)$ is continuous and strictly positive on $K_0 \times K_T$. For any $\beta \in (0, \infty)$, there exists a unique pair of non-negative, integrable functions (ρ_0, ρ_T) such that*

$$\begin{aligned} f_0(y) &= \rho_0(y) \int_{K_T} p(x, T | y, 0) \rho_T(x) dx, \\ f_T(x) &= \rho_T(x)^{(1+\beta)/\beta} \int_{K_0} p(x, T | y, 0) \rho_0(y) dy. \end{aligned}$$

Proof. See Appendix D. $\square$

Remark 7. The proof of Theorem 5 is adapted from that for Proposition 1 in Chen et al. (2016a). Observe that the Schrödinger system naturally yields an iterative algorithm for computing ρ_0, ρ_T . Given an estimate for ρ_0 , denoted by $\hat{\rho}_0$, we can estimate ρ_T by

$$\hat{\rho}_T(x) = \left(\frac{f_T(x)}{\int_{K_0} p(x, T | y, 0) \hat{\rho}_0(y) dy} \right)^{\beta/(1+\beta)},$$

which then can be used to update $\hat{\rho}_0$ by

$$\hat{\rho}_0^{\text{new}}(y) = \frac{f_0(y)}{\int_{K_T} p(x, T | y, 0) \hat{\rho}_T(x) dx}.$$

Chen et al. (2016a) considered the original Schrödinger bridge problem (i.e., $\beta = \infty$) and showed that this updating scheme yields a strict contraction with respect to the Hilbert metric. We note that when $\beta < \infty$, this argument can be potentially made easier, since the mapping $\psi \mapsto \psi^{\beta/(1+\beta)}$ for a suitable function ψ can be easily shown to be a strict contraction, and thus one only needs to verify the other steps in the updating are contractions (not necessarily strict); see Lemma D5 in Appendix. The full scope of consequences of this observation and the existence proof for the general case are left to future study.

4 EXTENSION TO TIME SERIES

Recently a time series version of Problem 1 was studied in Hamdouche et al. (2023), where the goal is to generate time series samples from a joint probability distribution on $\mathbb{R}^d \times \cdots \times \mathbb{R}^d$. We can generalize our Problem 2 to time series data analogously.

Problem 3. Consider N fixed time points $0 < t_1 < \cdots < t_N = T$. Let μ_N be a probability distribution on $\mathbb{R}^{d \times N}$ such that $\mu_N \ll \lambda$. For $\beta > 0$, find $V = \inf_{u \in \mathcal{U}} J_\beta^N(u)$, where

$$\begin{aligned} J_\beta^N(u) &= \beta \mathcal{D}_{\text{KL}}(\mathcal{L}aw((X_t)_{1 \leq t \leq N}), \mu_N) \\ &\quad + \mathbb{E} \int_0^T \frac{\|u_t\|^2}{2\sigma^2} dt, \end{aligned} \quad (4.1)$$

and find the optimal control u^* such that $J_\beta^N(u^*) = V$.

Recall that in Section 3, we started by considering functions h such that $h(X_t, t) = \mathbb{E}[g(X_T) | \mathcal{F}_t]$ for some function g , where $\mathcal{F}_t = \sigma((X_s)_{0 \leq s \leq t})$. It turns out that this technique can be used to solve Problem 3 as well, but now we need to consider conditional expectations of the form $\mathbb{E}[g(X_{t_1}, \dots, X_{t_N}) | \mathcal{F}_t]$. To simplify the notation, we will write $\mathbf{x}_j = (x_1, \dots, x_j)$, $\mathbf{X}_j = (X_{t_1}, \dots, X_{t_j})$, and $\mathbf{X}_j^u = (X_{t_1}^u, \dots, X_{t_j}^u)$; when $j = 0$, $\mathbf{x}_j, \mathbf{X}_j, \mathbf{X}_j^u$ all denote the empty vector.

Given a measurable function $g: \mathbb{R}^{d \times N} \rightarrow [0, \infty)$, the Markovian property of X enables us to express the conditional expectation $g(\mathbf{X}_N)$ by

$$\mathbb{E}[g(\mathbf{X}_N) | \mathcal{F}_t] = \sum_{j=0}^{N-1} \mathbf{1}_{[t_j, t_{j+1})}(t) \cdot h_j(X_t, t; \mathbf{X}_j),$$

where we set $t_0 = 0$, and the function h_j is defined by

$$\begin{aligned} h_j(x, t; \mathbf{x}_j) &= \mathbb{E} \left[g(\mathbf{x}_j, X_{t_{j+1}}, \dots, X_{t_N}) \middle| X_t = x \right], \end{aligned} \quad (4.2)$$

for $(x, t) \in \mathbb{R}^d \times [t_j, t_{j+1})$. Let $\mathcal{T}_N g$ denote the function on $\mathbb{R}^d \times [0, T) \times \mathbb{R}^{d \times N}$ given by

$$\begin{aligned} (\mathcal{T}_N g)(x, t, \mathbf{x}_N) &= \sum_{j=0}^{N-1} \mathbf{1}_{[t_j, t_{j+1})}(t) \cdot \sigma^2 \nabla \log h_j(x, t; \mathbf{x}_j). \end{aligned} \quad (4.3)$$

We prove in Lemma C4 in Appendix that it suffices to consider controls in the set

$$\mathcal{U}_N = \{u \in \mathcal{U}: u_t = (\mathcal{T}_N g)(X_t^u, t, \mathbf{X}_N^u) \text{ for some } g \geq 0\}.$$

This result is a generalization of Lemma 3 and obtained by applying Theorem B3 to each time interval $[t_j, t_{j+1})$ separately. Note although we express

$u \in \mathcal{U}_N$ as a function of $\mathbf{X}_N^u$, by (4.3), u_t is measurable with respect to $\sigma((X_s^u)_{0 \leq s \leq t})$. To formulate the Schrödinger system for the time series SSB problem, let $p_N(\mathbf{x}_N | y)$ denote the transition density from $X_0 = y$ to $\mathbf{X}_N = \mathbf{x}_N$, which is given by

$$p_N(\mathbf{x}_N | y) = p(x_1, t_1 | y, 0) \prod_{j=1}^{N-1} p(x_{j+1}, t_{j+1} | x_j, t_j).$$

Theorem 6. *Consider Problem 3. Suppose there exist σ -finite measures ν_0 on $\mathbb{R}^d$ and ν_N on $\mathbb{R}^{d \times N}$ such that $\nu_0 \approx \mu_0, \nu_N \approx \mu_N$ and*

$$\begin{aligned} \frac{d\mu_0}{d\nu_0}(y) &= \int p_N(\mathbf{x}_N | y) \nu_N(d\mathbf{x}_N), \\ \frac{d\mu_N}{d\lambda}(\mathbf{x}_N) &= \rho_N(\mathbf{x}_N)^{\frac{1+\beta}{\beta}} \int p_N(\mathbf{x}_N | y) \nu_0(dy), \end{aligned}$$

where $\rho_N = d\nu_N/d\lambda$. Assume $\int (d\mu_0/d\nu_0) d\mu_0 < \infty$. Then $u_t^* = (\mathcal{T}_N \rho_N)(X_t^{u^*}, t, \mathbf{X}_N^{u^*})$ solves Problem 3, where $\mathcal{T}_N$ is defined by (4.3). Moreover, $J_\beta^N(u^*) = -\mathcal{D}_{\text{KL}}(\mu_0, \nu_0) \in [0, \infty)$.

Proof. See Appendix C. $\square$

Comparing the Schrödinger system in Theorem 4 and that in Theorem 6, we see that the solution to Problem 3 has essentially the same structure as that to Problem 2. The only difference is that in Theorem 4 the Schrödinger system is constructed by using the joint distribution of (X_0, X_T) , while in Theorem 6 it is replaced by the joint distribution of $(X_0, \mathbf{X}_N)$. We also note that Hamdouche et al. (2023) only considered the time series Schrödinger bridge problem with μ_0 being a Dirac measure, and by letting $\beta \rightarrow \infty$, Theorem 6 gives the solution to their problem in the general case.

5 EXPERIMENTS

5.1 Problem Setup

We consider an application of SSB to robust generative modeling in the following scenario. Let $\mathcal{D}_{\text{ref}}$ denote a large collection of high-quality samples with distribution μ_{ref} , and let $\mathcal{D}_{\text{obj}}$ be a small set of noisy samples with distribution μ_{obj} . Our objective is to generate realistic samples resembling those in $\mathcal{D}_{\text{obj}}$, but due to the limited availability of training samples, we want to leverage information from $\mathcal{D}_{\text{ref}}$ to enhance the sample quality.

A natural idea is to use SSB as a regularization method to mitigate overfitting to the noisy samples in $\mathcal{D}_{\text{obj}}$. This can be implemented in two steps. For simplicity, we assume in this section that the uncontrolled process

X is given by $X_t = \sigma W_t$; that is, we assume $X_0 = 0$ and $b \equiv 0$ in (2.1). Then X_T has density $\phi_{\sigma\sqrt{T}}$ (recall this is the density of the normal distribution with mean 0 and covariance matrix $(\sigma^2 T)I$). Let f_{ref} denote the density of μ_{ref} with respect to the Lebesgue measure. Let $X^{\text{ref}} = (X_t^{\text{ref}})_{0 \leq t \leq T}$ be the Schrödinger bridge targeting μ_{ref} evolving by

$$dX_t^{\text{ref}} = b^{\text{ref}}(X_t^{\text{ref}}, t)dt + \sigma dW_t, \quad (5.1)$$

for $t \in [0, T]$, where

$$b^{\text{ref}}(X_t, t) = \sigma^2 \nabla \log \mathbb{E} \left[\frac{f_{\text{ref}}(X_T)}{\phi_{\sigma\sqrt{T}}(X_T)} \mid X_t = x \right].$$

Theorem 1 implies that X_T^{ref} has distribution μ_{ref} . Next, we solve Problem 2 using X^{ref} as the reference process and μ_{obj} as the target distribution. This yields the process X^{obj} with dynamics given by

$$dX_t^{\text{obj}} = b^{\text{obj}}(X_t^{\text{obj}}, t)dt + \sigma dW_t, \quad (5.2)$$

for $t \in [0, T]$, where

$$\begin{aligned} b^{\text{obj}}(X_t^{\text{obj}}, t) &= b^{\text{ref}}(X_t^{\text{obj}}, t) + \\ \sigma^2 \nabla \log \mathbb{E} &\left[\left(\frac{f_{\text{obj}}(X_T^{\text{ref}})}{f_{\text{ref}}(X_T^{\text{ref}})} \right)^{\beta/(1+\beta)} \mid X_t^{\text{ref}} = x \right]. \end{aligned}$$

By Theorem 2, the distribution of X_T^{obj} will be close to μ_{obj} if β is relatively large. It turns out that there is no need to train X^{ref} and X^{obj} separately. The following lemma shows that we can directly train a Schrödinger bridge targeting a geometric mixture of μ_{ref} and μ_{obj} .

Lemma 7. *Let $X_t = \sigma W_t$ and X^{ref} and X^{obj} be as given in (5.1) and (5.2). Equivalently, we can express the drift of X^{obj} by*

$$\begin{aligned} &b^{\text{obj}}(X_t^{\text{obj}}, t) \\ &= \sigma^2 \nabla \log \mathbb{E} \left[\frac{f_{\text{ref}}(X_T)^{\frac{1}{1+\beta}} f_{\text{obj}}(X_T)^{\frac{\beta}{1+\beta}}}{\phi_{\sigma\sqrt{T}}(X_T)} \mid X_t = x \right]. \end{aligned}$$

Proof. See Appendix E. $\square$

Remark 8. Lemma 7 follows from a change-of-measure argument. It still holds if X is not a Brownian motion but X solves the SDE (2.1) with $X_0 = x_0 \in \mathbb{R}^d$ almost surely (see Appendix E).

The assumption that μ_0 (the initial distribution of X) is a Dirac measure greatly simplifies the calculations and enables us to directly target the unnormalized density function $f_{\text{ref}}^{1/(1+\beta)} f_{\text{obj}}^{\beta/(1+\beta)}$. We will propose in the next subsection a score matching algorithm for learning this geometric mixture distribution. For applications where a general initial distribution is desirable, one may need to build iterative algorithms by borrowing ideas from the iterative proportional fitting procedure (De Bortoli et al., 2021).

Example 4. For an illustrative toy example, let μ_{ref} be a mixture of four bivariate normal distributions with means $(1, 1), (1, -1), (-1, 1), (-1, -1)$ respectively and the same covariance matrix $0.05^2 I$. Let μ_{obj} be a mixture of two normal distributions with means $(1.2, 0.8), (-1.5, -0.5)$ respectively and the same covariance matrix $0.5^2 I$. So μ_{obj} essentially contains two components of μ_{ref} but with small bias and much larger noise. Since the density functions are known, we can directly simulate a Schrödinger bridge process targeting $f_{\text{ref}}^{1/(1+\beta)} f_{\text{obj}}^{\beta/(1+\beta)}$ using the method described in Remark 3. We provide the results in Appendix A.3, which show that by targeting a geometric mixture with a moderate value of β , we can effectively compel the terminal distribution of the controlled process to acquire a covariance structure similar to μ_{ref} .

5.2 A Score Matching Algorithm for Learning Geometric Mixtures

Let μ^* denote a probability distribution with density function $f^*(x) = C^{-1} f_{\text{ref}}(x)^{1/(1+\beta)} f_{\text{obj}}(x)^{\beta/(1+\beta)}$ where C is the normalizing constant. To generate samples from μ^* , one can use existing score-based diffusion model methods, but, as we will see shortly, one major challenge is how to train the score functions without samples from the distribution μ^* .

Let μ_σ^* be the probability distribution with density

$$f_\sigma^*(x) = \int f^*(x) \phi_\sigma(x - y) dy,$$

which can be thought of as a smoothed version of f^* , and suppose that we can generate samples from the distribution μ_σ^* . Solving the Schrödinger bridge problem with initial distribution μ_σ^* and terminal distribution μ^* , we obtain the controlled process X^* with $\mathcal{L}aw(X_0^*) = \mu_\sigma^*$ and dynamics given by

$$\begin{aligned} dX_t^* &= b^*(X_t^*, t) dt + \sigma dW_t, \\ \text{where } b^*(x, t) &= \sigma^2 \nabla \log f_{\sigma\sqrt{T-t}}^*(x). \end{aligned} \quad (5.3)$$

The process X^* satisfies $\mathcal{L}aw(X_T^*) = \mu^*$ (note that this result is a special case of Example 3 with $\beta = \infty$).

We now describe how to simulate the dynamics given in (5.3) and generate samples from μ_σ^* . First, to learn the drift function in (5.3), we propose to combine the score matching technique with importance sampling. Let $s_\theta(x, \tilde{\sigma})$ denote our approximation to $\nabla \log f_\sigma^*(x)$, where $\tilde{\sigma} \in [0, \sigma\sqrt{T}]$ and the unknown parameter θ typically denotes a neural network. According to the well-known score matching technique (Hyvärinen, 2005; Vincent, 2011), we can estimate θ for a given $\tilde{\sigma}$ by min-

imizing the objective function $\mathbb{E}_{x \sim \mu^*} [L(x, \theta; \tilde{\sigma})]$, where

$$\begin{aligned} L(x, \theta; \tilde{\sigma}) &= \mathbb{E}_{\tilde{x} \sim \mathcal{N}(x, \tilde{\sigma}^2 I)} \left[\|s_\theta(\tilde{x}, \tilde{\sigma}) - \nabla_{\tilde{x}} \log f_{\tilde{\sigma}}^*(\tilde{x} | x)\|^2 \right] \\ &= \mathbb{E}_{\tilde{x} \sim \mathcal{N}(x, \tilde{\sigma}^2 I)} \left[\left\| s_\theta(\tilde{x}, \tilde{\sigma}) + \frac{\tilde{x} - x}{\tilde{\sigma}^2} \right\|^2 \right]. \end{aligned}$$

Unfortunately, the existing score matching methods estimate $\mathbb{E}_{x \sim \mu^*} [L(x, \theta; \tilde{\sigma})]$ by using samples from μ^* , which are not applicable to our problem since we only have access to samples from μ_{ref} and μ_{obj} but not from μ^* . We propose to use importance sampling to tackle this issue. Let $\tilde{\mu}$ denote an auxiliary distribution with density $q = d\tilde{\mu}/d\lambda$, and assume that we can generate samples from $\tilde{\mu}$ (e.g. we can let $\tilde{\mu} = \mu_{\text{ref}}, \mu_{\text{obj}}$ or the smoothed versions of them). By a change of measure, we can express $\mathbb{E}_{x \sim \mu^*} [L(x, \theta; \tilde{\sigma})]$ as

$$C^{-1} \mathbb{E}_{x \sim \tilde{\mu}} \left[L(x, \theta; \tilde{\sigma}) \left(\frac{f_{\text{ref}}(x)}{q(x)} \right)^{\frac{1}{1+\beta}} \left(\frac{f_{\text{obj}}(x)}{q(x)} \right)^{\frac{\beta}{1+\beta}} \right].$$

The density ratios f_{ref}/q and f_{obj}/q can be learned by using samples from $\mu_{\text{ref}}, \mu_{\text{obj}}, \tilde{\mu}$ and minimizing the logistic regression loss as in Wang et al. (2021). We do not need to know C since it does not affect the drift term in (5.3). In particular, when $\tilde{\mu} = \mu_{\text{ref}}$, we get

$$\begin{aligned} &\mathbb{E}_{x \sim \mu^*} [L(x, \theta; \tilde{\sigma})] \\ &= C^{-1} \mathbb{E}_{x \sim \mu_{\text{ref}}} \left[L(x, \theta; \tilde{\sigma}) \left(\frac{f_{\text{obj}}(x)}{f_{\text{ref}}(x)} \right)^{\frac{\beta}{1+\beta}} \right], \end{aligned}$$

where the ratio $f_{\text{obj}}/f_{\text{ref}}$ can be learned by using samples from $\mathcal{D}_{\text{obj}}, \mathcal{D}_{\text{ref}}$. A similar expression can be easily derived for $\tilde{\mu} = \mu_{\text{obj}}$. This importance sampling method enables us to estimate the expectation $\mathbb{E}_{x \sim \mu^*} [L(x, \theta; \tilde{\sigma})]$ by using samples from $\mathcal{D}_{\text{obj}}, \mathcal{D}_{\text{ref}}$. By averaging over $\tilde{\sigma}$ randomly drawn from the interval $[0, \sigma\sqrt{T}]$, we obtain an estimated loss for the parameter θ . Minimizing this loss we get the estimate $\hat{\theta}$, and we can approximate $\nabla \log f_\sigma^*(x)$ using $s_{\hat{\theta}}(x, \tilde{\sigma})$.

Simulating the SDE (5.3) requires us to generate samples from μ_σ^* . There are a few possible approaches. First, if σ is chosen to be sufficiently large, one may argue that we can simply approximate f_σ^* using the normal density ϕ_σ , and thus we only need to draw X_0 from the normal distribution. This is the approach taken in the score-based generative models based on backward SDEs (Song and Ermon, 2019). Second, one can run an additional Schrödinger bridge process with terminal distribution μ_σ^* , as proposed in the two-stage Schrödinger bridge algorithm of Wang et al. (2021). The dynamics they considered is simply the solution to Problem 1 given in Theorem 1, where the uncontrolled process is a Brownian motion started at 0. The

drift term is approximated by a Monte Carlo scheme (see Remark 3 and Appendix A.1), which also requires an estimate of the density ratio f_σ^*/ϕ_σ and the score of f_σ^* . Third, one can also run the Langevin diffusion $d\tilde{X}_t = \sigma^2 \nabla \log f_\sigma^*(\tilde{X}_t)dt + \sigma dW_t$ for a sufficiently long duration, which has stationary distribution μ_σ^* . We found in our experiments that this approach yields more robust results, probably because it does not require Monte Carlo sampling which may yield drift estimates with large variances.

5.3 An Example Using MNIST

We present a numerical example using the MNIST data set (Deng, 2012), which consists of images of handwritten digits from 0 to 9. All images have size 28×28 and pixels are rescaled to $[0, 1]$. To construct $\mathcal{D}_{\text{obj}}$, we randomly select 50 images labeled as eight and reduce their quality by adding Gaussian noise with mean 0 and variance 0.4; see Fig. F4 in Appendix F. The reference data set $\mathcal{D}_{\text{ref}}$ includes all images that are not labeled as eight, a total of 54,149 samples.

We first run the algorithm of Wang et al. (2021) using only $\mathcal{D}_{\text{obj}}$, and the generated samples are noisy as expected; see Fig. F5 in Appendix F. Next, we run our algorithm described in Section 5.2 with different choices of β . The density ratio and scores are trained by using one GPU (RTX 6000). Generated samples are shown in Fig. 1, and we report in Table 1 the Fréchet Inception Distance score (Heusel et al., 2017) assessing the disparity between our generated images (sample size = 40K) and the collection of clean digit 8 images from the MNIST dataset (sample size $\approx$ 6K). When β is too small, the generated images do not resemble those in $\mathcal{D}_{\text{obj}}$ and we frequently observe the influence of other digits; if $\beta = 0$, the images are essentially generated from μ_{ref} . When β is too large, the influence from the reference data set becomes negligible, but the algorithm tends to overfit to the noisy data in $\mathcal{D}_{\text{obj}}$. For moderate values of β , we observe a blend of characteristics from both data sets, and when $\beta = 1.5$, FID is minimized (among all tried values) and we get high-quality images of digit 8. This experiment illustrates that the information from $\mathcal{D}_{\text{obj}}$ can help capture the structural features associated with the digit 8, while the samples from $\mathcal{D}_{\text{ref}}$ can guide the algorithm towards effective noise removal. In Appendix F, we further analyze the generated images using t-SNE plots and inception scores (Salimans et al., 2016).

Table 1: FID Scores for MNIST Experiment

β	0	0.25	0.7	1.5	4	100
FID	67.4	66.4	61.0	56.3	110.9	182.4

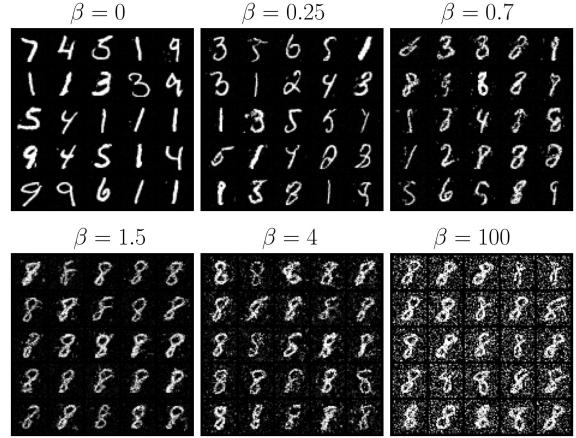


Figure 1: SSB Samples for MNIST Experiment

6 CONCLUDING REMARKS

We propose the soft-constrained Schrödinger bridge (SSB) problem and find its solution. Our theory encompasses the existing stochastic control results for Schrödinger bridge in Dai Pra (1991) and Hamdouché et al. (2023) as special cases. The paper focuses on the theory of SSB, and the numerical examples are designed to be uncomplicated but illustrative. More advanced algorithms for solving Problems 2 and 3 in full generality need to be developed. It will also be interesting to study the applications of SSB to other generative modeling tasks, such as conditional generation, style transfer (Shi et al., 2022; Shi and Wu, 2023; Su et al., 2022) and time series data generation (Hamdouché et al., 2023). Some further generalization of the objective function may be considered as well; for example, one can add a time-dependent cost as in Pra and Pavon (1990) or consider a more general form of the terminal cost.

Acknowledgements

The authors would like to thank Tiziano De Angelis for valuable discussion on the problem formulation, Yun Yang for the helpful conversation on the numerics, and anonymous reviewers for their comments and suggestions. JG and QZ were supported in part by NSF grants DMS-2311307 and DMS-2245591. XZ acknowledges the support from NSF DMS-2113359. The numerical experiments were conducted with the computing resources provided by Texas A&M High Performance Research Computing.

References

Julius Berner, Lorenz Richter, and Karen Ullrich. An optimal control perspective on diffusion-based gen-

- erative modeling. In *NeurIPS 2022 Workshop on Score-Based Methods*, 2022.
- Arne Beurling. An automorphism of product measures. *Annals of Mathematics*, pages 189–200, 1960.
- A Blaquiere. Controllability of a Fokker-Planck equation, the Schrödinger system, and a related stochastic optimal control (revised version). *Dynamics and Control*, 2(3):235–253, 1992.
- Peter J Bushell. Hilbert’s metric and positive contraction mappings in a Banach space. *Archive for Rational Mechanics and Analysis*, 52:330–338, 1973.
- Tianrong Chen, Guan-Hong Liu, and Evangelos Theodorou. Likelihood training of Schrödinger bridge using forward-backward SDEs theory. In *International Conference on Learning Representations*, 2021a.
- Yongxin Chen, Tryphon Georgiou, and Michele Pavon. Entropic and displacement interpolation: a computational approach using the Hilbert metric. *SIAM Journal on Applied Mathematics*, 76(6):2375–2396, 2016a.
- Yongxin Chen, Tryphon T Georgiou, and Michele Pavon. On the relation between optimal transport and Schrödinger bridges: A stochastic control viewpoint. *Journal of Optimization Theory and Applications*, 169:671–691, 2016b.
- Yongxin Chen, Tryphon T Georgiou, Michele Pavon, and Allen Tannenbaum. Relaxed Schrödinger bridges and robust network routing. *IEEE Transactions on Control of Network Systems*, 7(2):923–931, 2019.
- Yongxin Chen, Tryphon T Georgiou, and Michele Pavon. Optimal transport in systems and control. *Annual Review of Control, Robotics, and Autonomous Systems*, 4:89–113, 2021b.
- Yongxin Chen, Tryphon T Georgiou, and Michele Pavon. Stochastic control liaisons: Richard Sinkhorn meets Gaspard Monge on a Schrödinger bridge. *SIAM Review*, 63(2):249–313, 2021c.
- Paolo Dai Pra. A stochastic control approach to reciprocal diffusion processes. *Applied Mathematics and Optimization*, 23(1):313–329, 1991.
- Valentin De Bortoli, James Thornton, Jeremy Heng, and Arnaud Doucet. Diffusion Schrödinger bridge with applications to score-based generative modeling. *Advances in Neural Information Processing Systems*, 34:17695–17709, 2021.
- George Deligiannidis, Valentin De Bortoli, and Arnaud Doucet. Quantitative uniform stability of the iterative proportional fitting procedure. *The Annals of Applied Probability*, 34(1A):501–516, 2024.
- W Edwards Deming and Frederick F Stephan. On a least squares adjustment of a sampled frequency table when the expected marginal totals are known. *The Annals of Mathematical Statistics*, 11(4):427–444, 1940.
- Li Deng. The MNIST database of handwritten digit images for machine learning research. *IEEE Signal Processing Magazine*, 29(6):141–142, 2012.
- J. L. Doob. A Markov chain theorem. *Probability and Statistics.(H. Cramer memorial volume) ed. U. Grenander. Almqvist and Wiksel, Stockholm & New York*, pages 50–57, 1959.
- J. L. Doob. *Classical potential theory and its probabilistic counterpart*, volume 262 of *Grundlehren der Mathematischen Wissenschaften [Fundamental Principles of Mathematical Sciences]*. Springer-Verlag, New York, 1984. ISBN 0-387-90881-1. doi: 10.1007/978-1-4612-5208-5.
- Montacer Essid and Michele Pavon. Traversing the Schrödinger bridge strait: Robert Fortet’s marvelous proof redux. *Journal of Optimization Theory and Applications*, 181(1):23–60, 2019.
- Wendell H Fleming. Exit probabilities and optimal stochastic control. *Applied Mathematics and Optimization*, 4:329–346, 1977.
- Wendell H Fleming. Logarithmic transformations and stochastic control. In *Advances in Filtering and Optimal Stochastic Control: Proceedings of the IFIP-WG 7/1 Working Conference Cocoyoc, Mexico, February 1–6, 1982*, pages 131–141. Springer, 2005.
- Wendell H Fleming and Raymond W Rishel. *Deterministic and stochastic optimal control*, volume 1. Springer Science & Business Media, 2012.
- Wendell H Fleming and Sheunn-Jyi Sheu. Stochastic variational formula for fundamental solutions of parabolic pde. *Applied Mathematics and Optimization*, 13:193–204, 1985.
- H Föllmer. Random fields and diffusion processes. *Ecole d’Ete de Probabilites de Saint-Flour XV-XVII, 1985-87*, 1988.
- Robert Fortet. Résolution d’un système d’équations de m. Schrödinger. *Journal de Mathématiques Pures et Appliquées*, 19(1-4):83–105, 1940.
- Avner Friedman. Stochastic differential equations and applications. In *Stochastic differential equations*, pages 75–148. Springer, 1975.
- Tryphon T Georgiou and Michele Pavon. Positive contraction mappings for classical and quantum Schrödinger systems. *Journal of Mathematical Physics*, 56(3), 2015.

- Mohamed Hamdouche, Pierre Henry-Labordere, and Huyên Pham. Generative modeling for time series via Schrödinger bridge. *arXiv preprint arXiv:2304.05093*, 04 2023. doi: 10.13140/RG.2.2.25758.00324.
- Jeremy Heng, Valentin De Bortoli, and Arnaud Doucet. Diffusion Schrödinger bridges for Bayesian computation. *Statistical Science*, 39(1):90–99, 2024.
- Martin Heusel, Hubert Ramsauer, Thomas Unterthiner, Bernhard Nessler, and Sepp Hochreiter. GANs trained by a two time-scale update rule converge to a local Nash equilibrium. *Advances in Neural Information Processing Systems*, 30, 2017.
- Jonathan Ho, Ajay Jain, and Pieter Abbeel. Denoising diffusion probabilistic models. *Advances in Neural Information Processing Systems*, 33:6840–6851, 2020.
- Jian Huang, Yuling Jiao, Lican Kang, Xu Liao, Jin Liu, and Yanyan Liu. Schrödinger-Föllmer sampler: sampling without ergodicity. *arXiv preprint arXiv:2106.10880*, 2021.
- Aapo Hyvärinen. Estimation of non-normalized statistical models by score matching. *Journal of Machine Learning Research*, 6(4):695–709, 2005.
- Benton Jamison. The Markov processes of Schrödinger. *Zeitschrift für Wahrscheinlichkeitstheorie und Verwandte Gebiete*, 32(4):323–331, 1975.
- Ioannis Karatzas and Steven Shreve. *Brownian motion and stochastic calculus*, volume 113. Springer Science & Business Media, 2012.
- Christian Léonard. A survey of the Schrödinger problem and some of its connections with optimal transport, 2013.
- Toshio Mikami. Variational processes from the weak forward equation. *Communications in Mathematical Physics*, 135:19–40, 1990.
- Taehong Moon, Moonseok Choi, Gayoung Lee, Jung-Woo Ha, and Juho Lee. Fine-tuning diffusion models with limited data. In *NeurIPS 2022 Workshop on Score-Based Methods*, 2022.
- Michele Pavon and Anton Wakolbinger. On free energy, stochastic control, and Schrödinger processes. In *Modeling, Estimation and Control of Systems with Uncertainty: Proceedings of a Conference held in Sopron, Hungary, September 1990*, pages 334–348. Springer, 1991.
- Stefano Peluchetti. Diffusion bridge mixture transports, Schrödinger bridge problems and generative modeling. *arXiv preprint arXiv:2304.00917*, 2023.
- Paolo Dai Pra and Michele Pavon. On the Markov processes of Schrödinger, the Feynman-Kac formula and stochastic control. In *Realization and Modelling in System Theory: Proceedings of the International Symposium MTNS-89, Volume I*, pages 497–504. Springer, 1990.
- Lorenz Richter, Julius Berner, and Guan-Horng Liu. Improved sampling via learned diffusions. *arXiv preprint arXiv:2307.01198*, 2023.
- Tim Salimans, Ian Goodfellow, Wojciech Zaremba, Vicki Cheung, Alec Radford, and Xi Chen. Improved techniques for training GANs. *Advances in Neural Information Processing Systems*, 29, 2016.
- Erwin Schrödinger. *Über die Umkehrung der Naturgesetze*. Verlag der Akademie der Wissenschaften in Kommission bei Walter De Gruyter u. Company, 1931.
- Erwin Schrödinger. Sur la théorie relativiste de l’électron et l’interprétation de la mécanique quantique. In *Annales de l’institut Henri Poincaré*, volume 2, pages 269–310, 1932.
- Yuyang Shi, Valentin De Bortoli, George Deligiannidis, and Arnaud Doucet. Conditional simulation using diffusion Schrödinger bridges. In *Uncertainty in Artificial Intelligence*, pages 1792–1802. PMLR, 2022.
- Ziqiang Shi and Shoule Wu. Schröwave: Realistic voice generation by solving two-stage conditional Schrödinger bridge problems. *Digital Signal Processing*, 141:104175, 2023.
- Ki-Ung Song. Applying regularized Schrödinger-bridge-based stochastic process in generative modeling. *arXiv preprint arXiv:2208.07131*, 2022.
- Yang Song and Stefano Ermon. Generative modeling by estimating gradients of the data distribution. In *Advances in Neural Information Processing Systems*, pages 11895–11907, 2019.
- Yang Song, Jascha Sohl-Dickstein, Diederik P. Kingma, Abhishek Kumar, Stefano Ermon, and Ben Poole. Score-based generative modeling through stochastic differential equations, 2021.
- Xuan Su, Jiaming Song, Chenlin Meng, and Stefano Ermon. Dual diffusion implicit bridges for image-to-image translation. *arXiv preprint arXiv:2203.08382*, 2022.
- Belinda Tzen and Maxim Raginsky. Theoretical guarantees for sampling and inference in generative models with latent diffusions. In *Conference on Learning Theory*, pages 3084–3114. PMLR, 2019.
- Francisco Vargas, Will Sussman Grathwohl, and Arnaud Doucet. Denoising diffusion samplers. In *The Eleventh International Conference on Learning Representations*, 2022.

- Pascal Vincent. A connection between score matching and denoising autoencoders. *Neural Computation*, 23(7):1661–1674, 2011.
- Gefei Wang, Yuling Jiao, Qian Xu, Yang Wang, and Can Yang. Deep generative learning via Schrödinger bridge. In Marina Meila and Tong Zhang, editors, *Proceedings of the 38th International Conference on Machine Learning*, volume 139 of *Proceedings of Machine Learning Research*, pages 10794–10804. PMLR, 18–24 Jul 2021.
- Ludwig Winkler, Cesar Ojeda, and Manfred Opper. A score-based approach for training Schrödinger bridges for data modelling. *Entropy*, 25(2):316, 2023.
- Qinsheng Zhang and Yongxin Chen. Path integral sampler: A stochastic control approach for sampling. In *International Conference on Learning Representations*, 2021.

Checklist

1. For all models and algorithms presented, check if you include:
 - (a) A clear description of the mathematical setting, assumptions, algorithm, and/or model. Yes, the mathematical setting is described in Section 2 and algorithms are described in Section 5.2 and Appendix A.1.
 - (b) An analysis of the properties and complexity (time, space, sample size) of any algorithm. No. The focus of the paper is theory, and the algorithm we propose in Section 5.2 is a modification of existing diffusion model algorithms.
 - (c) (Optional) Anonymized source code, with specification of all dependencies, including external libraries. Yes, the code is publicly available on [GitHub](#).
2. For any theoretical claim, check if you include:
 - (a) Statements of the full set of assumptions of all theoretical results. Yes, Assumption 1 is stated at the beginning of Section 2.
 - (b) Complete proofs of all theoretical results. Yes.
 - (c) Clear explanations of any assumptions. Yes, Assumption 1 is discussed in Section 2 and Remark 4. Other assumptions used in the theorems are either standard or explained in the remarks.
3. For all figures and tables that present empirical results, check if you include:
 - (a) The code, data, and instructions needed to reproduce the main experimental results (either in the supplemental material or as a URL). Yes.
 - (b) All the training details (e.g., data splits, hyperparameters, how they were chosen). Yes, the data splitting is explained in Section 5.3, and the detailed algorithm setting (e.g. choice of hyperparameters) is provided in the configuration file on [GitHub](#).
 - (c) A clear definition of the specific measure or statistics and error bars (e.g., with respect to the random seed after running experiments multiple times). Not applicable.
 - (d) A description of the computing infrastructure used. (e.g., type of GPUs, internal cluster, or cloud provider).
- Yes, it is mentioned in Section 5.3. The provider is acknowledged in the Acknowledgements.
4. If you are using existing assets (e.g., code, data, models) or curating/releasing new assets, check if you include:
 - (a) Citations of the creator if your work uses existing assets. Yes.
 - (b) The license information of the assets, if applicable. Not applicable.
 - (c) New assets either in the supplemental material or as a URL, if applicable. Yes, the code is available on [GitHub](#).
 - (d) Information about consent from data providers/curators. Not applicable.
 - (e) Discussion of sensible content if applicable, e.g., personally identifiable information or offensive content. Not applicable.
5. If you used crowdsourcing or conducted research with human subjects, check if you include:
 - (a) The full text of instructions given to participants and screenshots. Not applicable.
 - (b) Descriptions of potential participant risks, with links to Institutional Review Board (IRB) approvals if applicable. Not applicable.
 - (c) The estimated hourly wage paid to participants and the total amount spent on participant compensation. Not applicable.

Appendix for “Soft-constrained Schrödinger Bridge: a Stochastic Control Approach”

The code for all experiments (Cauchy simulation in Appendix A.2, normal mixture simulation in Appendix A.3 and MNIST example in Section 5.3) is at the GitHub repository <https://github.com/gargjhanvi/Soft-constrained-Schrodinger-Bridge-a-Stochastic-Control-Approach>.

A SIMULATION WITH KNOWN DENSITY FUNCTIONS

A.1 Monte Carlo Simulation with Densities Known up to a Normalization Constant

Let the uncontrolled process X be the Brownian motion with $b \equiv 0$ and $X_0 = 0$, and set $T = 1$. Let ϕ_σ denote the density of normal distribution with mean zero and covariance matrix $\sigma^2 I$. Given a target distribution with density f_T (which we simply denote by f in this section), the solution to SSB is given by

$$dX_t^* = u_t^* dt + \sigma dW_t, \quad (\text{A.1})$$

where the drift u^* is determined by $u^*(x, t) = \nabla \log h(x, t)$ and

$$h(x, t) \propto (1-t)^{-d/2} \int \phi_\sigma \left(\frac{z-x}{\sqrt{1-t}} \right) \cdot \left(\frac{f(z)}{\phi_\sigma(z)} \right)^{\beta/(1+\beta)} dz.$$

Note that to determine u^* , we only need to know f up to a normalization constant. Since $\nabla \log h = h^{-1} \nabla h$, we can rewrite u^* as

$$u^*(x, t) = \frac{\mathbb{E}_{z \sim \phi_\sigma} [r(x + \sqrt{1-tz}) \nabla \log r(x + \sqrt{1-tz})]}{\mathbb{E}_{z \sim \phi_\sigma} [r(x + \sqrt{1-tz})]} \quad (\text{A.2})$$

where we set

$$r(x) = \left(\frac{f(x)}{\phi_\sigma(x)} \right)^{\beta/(1+\beta)}.$$

We can then approximate the numerator and denominator in (A.2) separately using Monte Carlo samples.

For some target distributions, this approach can be made more efficient by using importance sampling. When f has a heavy tail (e.g. Cauchy distribution), $r(x)$ may grow super-exponentially with x , and Monte Carlo estimates for the numerator and denominator in (A.2) with z drawn from the normal distribution may have large variances. Observe that the numerator can be written as

$$\mathbb{E}_{z \sim \phi_\sigma} [r(x + \sqrt{1-tz}) \nabla \log r(x + \sqrt{1-tz})] = \int \phi_\sigma(z) r(x + \sqrt{1-tz})^{\beta/(1+\beta)} \nabla \log r(x + \sqrt{1-tz}) dz.$$

The term $\nabla \log r(x + \sqrt{1-tz})$ is often polynomial in z (that is, it does not grow too fast). Hence, intuitively, the integral is likely to be well approximated by a Monte Carlo estimate with z drawn from a density proportional to $\phi_\sigma(z) r(x + \sqrt{1-tz})^{\beta/(1+\beta)}$. Such a density may not be easily accessible, but if one knows the tail decay rate of f , one can try to find a proposal distribution for z with tails not lighter than $\phi_\sigma(z) r(x + \sqrt{1-tz})^{\beta/(1+\beta)}$. In the experiment given below, we propose z from some distribution with tail decay rate same as $f^{\beta/(1+\beta)} \phi_\sigma^{1/(1+\beta)}$. Letting w denote our proposal density, we express the numerator in (A.2) as

$$\mathbb{E}_{z \sim \phi_\sigma} [r(x + \sqrt{1-tz}) \nabla \log r(x + \sqrt{1-tz})] = \mathbb{E}_{z \sim w} \left[\frac{\phi_\sigma(z) r(x + \sqrt{1-tz})}{w(z)} \nabla \log r(x + \sqrt{1-tz}) \right], \quad (\text{A.3})$$

and estimate the right-hand side using a Monte Carlo average.

Similarly, the denominator can be expressed by

$$\mathbb{E}_{z \sim \phi_\sigma} [r(x + \sqrt{1-t}z)] = \mathbb{E}_{z \sim w} \left[\frac{\phi_\sigma(z) r(x + \sqrt{1-t}z)}{w(z)} \right] \quad (\text{A.4})$$

and estimated by a Monte Carlo average.

A.2 Simulating the Cauchy Distribution

Fix $b \equiv 0$, $\sigma = 1$, $X_0 = 0$ and $T = 1$. Let f be the density of the standard Cauchy distribution i.e., $f(x) = \pi^{-1}(1+x^2)^{-1}$. Theorem 2 shows that the solution to SSB is a Schrödinger bridge such that the terminal distribution has density

$$f_\beta^* = C_\beta^{-1} \phi(x)^{1/(1+\beta)} f(x)^{\beta/(1+\beta)}, \quad (\text{A.5})$$

where C_β is the normalization constant and ϕ is the density of the standard normal distribution. We plot f_β^* and $\log f_\beta^*$ in Figure A1 for $\beta = \infty, 100, 50, 20$. For x close to zero, the density of $f_\beta^*(x)$ remains approximately the same for the four choices of β . When $\beta = \infty$, f_β^* is the Cauchy distribution which has a heavy tail. But for any $\beta < \infty$, the tail decay rate of f_β^* is dominated by the Gaussian component.

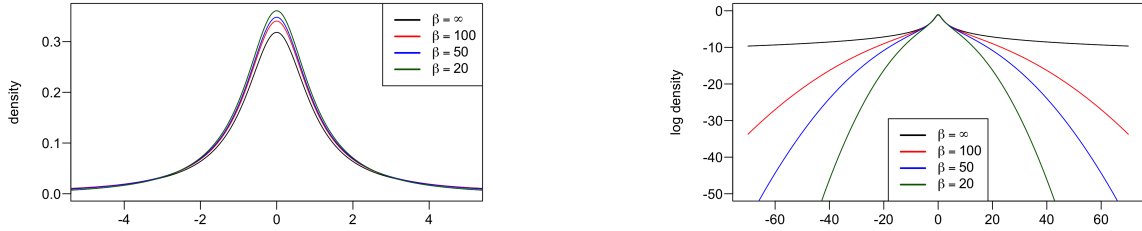


Figure A1: Densities of Geometric Mixtures of Normal Cauchy Distributions.

We simulate the solution to SSB given in (A.1) over the time interval $[0, 1]$ using the Euler-Maruyama method with 200 time steps. The drift is approximated by the importance sampling scheme. When $\beta = \infty$ (i.e., the terminal distribution is Cauchy), we let the proposal distribution w in (A.3) and (A.4) be the t -distribution with 2 degrees of freedom (we have also tried directly proposing z from the standard Cauchy distribution and obtained very similar results). When $\beta < \infty$, we let w be the normal distribution with mean 0 and variance $1 + \beta$. We simulate the SSB process 10,000 times, and in Table A1 we report the number of failed runs; these failures happen because Monte Carlo estimates for the numerator and denominator in (A.2) become unstable when $|x|$ is large, resulting in numerical overflow. When $\beta = \infty$, we observe that numerical overflow is still common even if we use 1,000 Monte Carlo samples for each estimate. In contrast, when we use $\beta \leq 100$ and only 200 Monte Carlo samples, the algorithm becomes very stable. In Figure A2, we compare the distribution of generated samples (i.e., the distribution of X_T^* with $T = 1$) with their corresponding target distributions. The first panel compares the distribution of the samples generated with $\beta = \infty$ (failed runs ignored) with the Cauchy distribution, and the second compares the distribution the samples generated with $\beta = 100$ with the geometric mixture given in (A.5). It is clear that simulating the Schrödinger bridge process (i.e, using $\beta = \infty$) cannot recover the heavy tails of the Cauchy distribution, but the numerical simulation of SSB with $\beta = 100$ accurately yields samples from the geometric mixture distribution. Recall that the KL divergence between standard Cauchy and normal distributions is infinite, which, by Theorem 1, means that there is no control with finite energy cost that can steer a standard Brownian motion towards the Cauchy distribution at time $T = 1$. Our experiment partially illustrates the practical consequences of this fact in numerically simulating the Schrödinger bridge, and it also suggests that SSB may be a numerically more robust alternative.

Table A1: Number of Failed Attempts in 10,000 Trials. N_{mc} denotes the number of Monte Carlo samples used to approximate the expectations defined in (A.3) and (A.4).

N_{mc}	20	50	100	200	500	1000
$\beta = \infty$	1149	933	714	610	444	308
$\beta = 100$	495	29	6	1	0	0
$\beta = 50$	124	18	6	0	0	0
$\beta = 20$	48	4	2	0	1	0

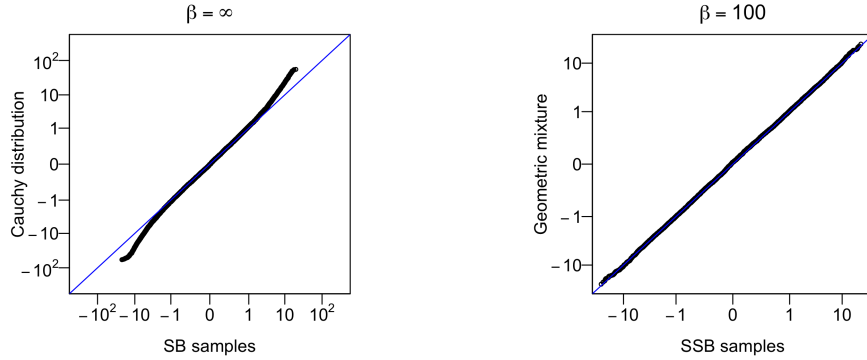


Figure A2: Q-Q plots for the SSB Samples with $N_{\text{mc}} = 1,000$.

A.3 Simulating Mixtures of Normal Distributions

In Example 4, we let μ_{ref} be a mixture of four bivariate normal distributions,

$$\mu_{\text{ref}} = 0.1\mathcal{N}((1, 1), 0.05^2 I) + 0.2\mathcal{N}((-1, 1), 0.05^2 I) + 0.3\mathcal{N}((1, -1), 0.05^2 I) + 0.4\mathcal{N}((-1, -1), 0.05^2 I),$$

where the weights and the mean vector of the four component distributions are different. Let μ_{obj} be a mixture of two equally weighted bivariate normal distributions

$$\mu_{\text{obj}} = 0.5\mathcal{N}((1.2, 0.8), 0.5^2 I) + 0.5\mathcal{N}((-1.5, -0.5), 0.5^2 I).$$

The first component of μ_{obj} has mean close to $(1, 1)$ (which is the mean vector of the first component of μ_{ref}), and the second component of μ_{obj} has mean close to $(-1, -1)$ (which is the mean vector of the last component of μ_{ref}). Hence, we can interpret μ_{ref} as a distribution of high-quality samples from four different classes, and interpret μ_{obj} as a distribution of noisy samples from two of the four classes. Let $f_{\beta}^* = C_{\beta}^{-1} f_{\text{ref}}(x)^{1/(1+\beta)} f_{\text{obj}}(x)^{\beta/(1+\beta)}$ be the density of our target distribution.

We generate 1,000 samples from f_{β}^* by simulating the SDE (A.1) over the time interval $[0, 1]$ with $f = f_{\beta}^*$ and $\sigma = 1$. We use 200 time steps for discretization and generate 200 Monte Carlo samples at each step for estimating the drift. The trajectories of the simulated processes are shown in Figure A3 for $\beta = 0, 2, 10, \infty$, and the samples we generate correspond to $t = 1$. It can be seen that when $\beta = 2$ or 10 , the majority of the generated samples form two clusters, one with mean close to $(1, 1)$ and the other with mean close to $(-1, -1)$. Further, the two clusters both exhibit very small within-cluster variation, which indicates that the noise from μ_{obj} has been effectively reduced.

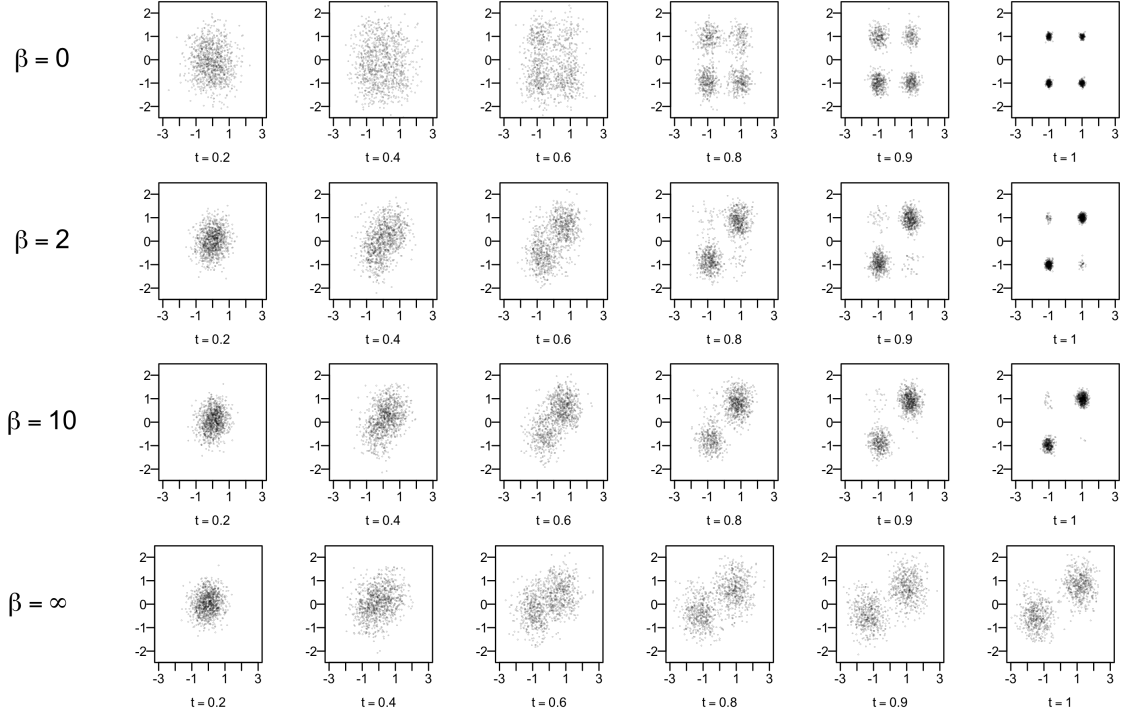


Figure A3: SSB Trajectories for Normal Mixture Targets.

B AUXILIARY RESULTS

We first present two lemmas about the minimization of Kullback-Leibler divergence.

Lemma B1. *Let $(\Omega, \mathcal{F}, \pi)$ be a σ -finite measure space, and let $\nu \ll \pi$ be a measure with density $f = d\nu/d\pi$ such that $C = \nu(\Omega) \in (0, \infty)$. Let $\mathcal{P}_\pi$ denote the set of all probability measures absolutely continuous with respect to π . Then,*

$$\inf_{\mu \in \mathcal{P}_\pi} \mathcal{D}_{\text{KL}}(\mu, \nu) = -\log C.$$

The infimum is attained by the probability measure μ^ such that $d\mu^*/d\pi = C^{-1}f(x)$.*

Proof of Lemma B1. Since $\mathcal{D}_{\text{KL}}(\mu, \nu) = \infty$ if $\mu \not\ll \nu$, it suffices to consider $\mu \in \mathcal{P}_\pi$ such that $\mu \ll \nu$. It is straightforward to show that $\mathcal{D}_{\text{KL}}(\mu, \nu) = \mathcal{D}_{\text{KL}}(\mu, \mu^*) - \log C$. Since $\int d\mu^* = C^{-1} \int f d\pi = C^{-1} \int d\nu = 1$, μ^* is a probability measure. The claim then follows from the fact that the KL divergence between any two probability measures is non-negative. $\square$

Lemma B2. *Let $(\Omega, \mathcal{F}, \pi)$ and $\mathcal{P}_\pi$ be as given in Lemma B1. For $i = 0, 1$, let ν_i be a finite measure (i.e., $\nu_i(\Omega) < \infty$) with density $f_i = d\nu_i/d\pi$. Assume that $\nu_1 \ll \nu_0$. For $\beta \geq 0$,*

$$\inf_{\mu \in \mathcal{P}_\pi} \mathcal{D}_{\text{KL}}(\mu, \nu_0) + \beta \mathcal{D}_{\text{KL}}(\mu, \nu_1) = -(1 + \beta) \log C_\beta,$$

where $C_\beta = \int f_0(x)^{1/(1+\beta)} f_1(x)^{\beta/(1+\beta)} \pi(dx) \in (0, \infty)$. The infimum is attained by the probability measure μ_β^ such that*

$$f_\beta^* = \frac{d\mu_\beta^*}{d\pi} = C_\beta^{-1} f_0^{1/(1+\beta)} f_1^{\beta/(1+\beta)}.$$

When $\nu_0, \nu_1 \in \mathcal{P}_\pi$, we have $C_\beta \in (0, 1]$. Further,

$$\lim_{\beta \rightarrow \infty} C_\beta = \nu_1(\Omega), \quad \lim_{\beta \rightarrow \infty} \mathcal{D}_{\text{KL}}(\mu_\beta^*, \nu_0) = \mathcal{D}_{\text{KL}}(\nu_1, \nu_0).$$

Proof of Lemma B2. Since $\nu_1 \ll \nu_0$, we have $C > 0$. By Hölder's inequality,

$$C_\beta \leq \left(\int f_0 d\pi \right)^{1/(1+\beta)} \left(\int f_1 d\pi \right)^{\beta/(1+\beta)} = \nu_0(\Omega)^{1/(1+\beta)} \nu_1(\Omega)^{\beta/(1+\beta)} < \infty.$$

Clearly, if $\nu_0, \nu_1 \in \mathcal{P}_\pi$, the above inequality implies $C \leq 1$. Observe that

$$\mathcal{D}_{\text{KL}}(\mu, \nu_0) + \beta \mathcal{D}_{\text{KL}}(\mu, \nu_1) = (1 + \beta) \mathcal{D}_{\text{KL}}(\mu, \nu)$$

where ν is the measure with density $d\nu/d\pi = f_0^{1/(1+\beta)} f_1^{\beta/(1+\beta)} = C_\beta f_\beta^*$. Hence, we can apply Lemma B1 to prove that μ_β^* is the minimizer.

To prove the convergence as $\beta \rightarrow \infty$, let $A = \{x: f_0(x) \geq f_1(x)\}$ and write $C_\beta = C_{\beta,0} + C_{\beta,1}$, where

$$C_{\beta,0} = \int_A f_0(x)^{1/(1+\beta)} f_1(x)^{\beta/(1+\beta)} \pi(dx), \quad C_{\beta,1} = \int_{\Omega \setminus A} f_0(x)^{1/(1+\beta)} f_1(x)^{\beta/(1+\beta)} \pi(dx).$$

The integrands of both $C_{\beta,0}$ and $C_{\beta,1}$ are monotone in β . Hence, using $C_0 < \infty$ and monotone convergence theorem, we find that

$$\lim_{\beta \rightarrow \infty} C_\beta = \int \lim_{\beta \rightarrow \infty} f_0(x)^{1/(1+\beta)} f_1(x)^{\beta/(1+\beta)} \pi(dx) = \int f_1(x) \pi(dx) = \nu_1(\Omega),$$

which also implies $f_\beta^* \rightarrow f_1/\nu_1(\Omega)$ pointwise. An analogous argument using monotone convergence theorem proves that $\lim_{\beta \rightarrow \infty} \mathcal{D}_{\text{KL}}(\mu_\beta^*, \nu_0) = \mathcal{D}_{\text{KL}}(\nu_1, \nu_0)$. $\square$

The next result is about the controlled SDE (2.2) with $u = \sigma^2 \nabla \log h$. It is adapted from Theorem 2.1 of Dai Pra (1991). Our proof is provided for completeness.

Theorem B3. *Suppose Assumption 1 holds. Let $X = (X_t)_{0 \leq t \leq T}$ be a weak solution to (2.1) with initial distribution μ_0 and transition density $p(x, t | y, s)$. Define*

$$h(x, t) = \int g(z) p(z, T | x, t) dz, \quad \text{for } (x, t) \in \mathbb{R}^d \times [0, T],$$

for some measurable $g \geq 0$ such that $\mathbb{E}g(X_T) < \infty$. Assume $h > 0$ on $\mathbb{R}^d \times [0, T]$. Let $X^h = (X_t^h)_{0 \leq t \leq T}$ be a weak solution with $X_0^h = X_0$ to the SDE

$$dX_t^h = [b(X_t^h, t) + \sigma^2 \nabla \log h(X_t^h, t)] dt + \sigma dW_t, \quad \text{for } t \in [0, T]. \quad (\text{B.1})$$

Then, we have the following results.

(i) $h \in \mathbb{C}^{2,1}(\mathbb{R}^d \times [0, T])$ and

$$\frac{\partial h}{\partial t} + \sum_{i=1}^d b_i \frac{\partial h}{\partial x_i} + \frac{\sigma^2}{2} \sum_{i=1}^d \frac{\partial^2 h}{\partial x_i^2} = 0, \quad \text{for } (x, t) \in \mathbb{R}^d \times [0, T].$$

(ii) The weak solution X^h to the SDE (B.1) exists. Indeed, we can define a probability measure $\mathbb{Q}$ by $d\mathbb{Q}/d\mathbb{P} = g(X_T)/h(X_0, 0)$ such that the law of X under $\mathbb{Q}$ is the same as the law of X^h under $\mathbb{P}$.

(iii) The process X^h satisfies

$$\mathbb{E} \left[\log \frac{g(X_T^h)}{h(X_0^h, 0)} \right] = \mathbb{E} \left[\int_0^T \frac{\sigma^2}{2} \|\nabla \log h(X_s^h, s)\|^2 ds \right].$$

(iv) The transition density of X^h is given by

$$p_h(x, t | y, s) = \frac{p(x, t | y, s) h(x, t)}{h(y, s)}, \quad \text{for } 0 \leq s < t \leq T.$$

Hence, X^h is a Doob's h -path process, and the density of the distribution of X_T^h is

$$\frac{d\mathcal{L}aw(X_T^h)}{d\lambda}(x) = g(x) \int \frac{p(x, T | y, 0)}{h(y, 0)} \mu_0(dy).$$

Proof of Theorem B3. We follow the arguments of Jamison (1975); Dai Pra (1991). Under Assumption 1, Proposition 2.1 of Dai Pra (1991) (which is adapted from the result of Jamison (1975)) implies that $h \in \mathbb{C}^{2,1}(\mathbb{R}^d \times [0, T])$ and $\mathcal{L}h + \frac{\partial h}{\partial t} = 0$ on $\mathbb{R}^d \times [0, T]$, where

$$\mathcal{L} = \sum_{i=1}^d b_i \frac{\partial}{\partial x_i} + \frac{\sigma^2}{2} \sum_{i=1}^d \frac{\partial^2}{\partial x_i^2}$$

denotes the generator of X . This proves part (i).

To prove part (ii), we first apply Itô's lemma to get

$$\log h(X_t, t) = \log h(X_0, 0) + \int_0^t \tilde{b}(X_s, s) ds + \int_0^t \sigma \nabla \log h(X_s, s) dW_s,$$

for any $t \in [0, T]$, where

$$\tilde{b}(x, t) = \frac{1}{h(x, t)} \left((\mathcal{L}h)(x, t) + \frac{\partial h}{\partial t}(x, t) \right) - \frac{\sigma^2}{2} \|\nabla \log h(x, t)\|^2 = -\frac{\sigma^2}{2} \|\nabla \log h(x, t)\|^2.$$

Since $g(X_T)$ is integrable, $h(X_t, t)$ is a uniformly integrable martingale on $[0, T]$ and converges to $g(X_T)$ both a.s. and in L^1 . Letting $t \uparrow T$, we obtain that

$$\log g(X_T) = \log h(X_0, 0) - \int_0^T \frac{\sigma^2}{2} \|\nabla \log h(X_s, s)\|^2 ds + \int_0^T \sigma \nabla \log h(X_s, s) dW_s. \quad (\text{B.2})$$

Write $h(x, T) = g(x)$, $Y_t = h(X_t, t)$ and $Z_t = Y_t/Y_0$. We have shown that Y_t and Z_t are martingales on $[0, T]$. Since $\mathbb{E}[Z_T] = 1$, we can define a probability measure $\mathbb{Q}$ by $d\mathbb{Q}/d\mathbb{P} = Z_T$. By Girsanov theorem and the expression for Z_T given in (B.2), the law of X under $\mathbb{Q}$ is the same as the law of X^h under $\mathbb{P}$; in other words, $d\mathbb{P}_X^h = (g(x_T)/h(x_0, 0))d\mathbb{P}_X$, where $\mathbb{P}_X$ (resp. $\mathbb{P}_X^h$) is the probability measure induced by X (resp. X^h) on the space of continuous functions.

For part (iii), choose $t_n = T \wedge \tau_n$, where $\tau_n = \inf\{t: |X_t^h| \geq n\}$. Analogously to (B.2), we can apply Itô's lemma to get

$$\log h(X_{t_n}^h, t_n) = \log h(X_0^h, 0) + \int_0^{t_n} \frac{\sigma^2}{2} \|\nabla \log h(X_s^h, s)\|^2 ds + \int_0^{t_n} \sigma \nabla \log h(X_s^h, s) dW_s.$$

Since h is smooth, $|\nabla \log h(X_s^h, s)|$ is bounded on $[0, t_n]$. Taking expectations on both sides, we find that

$$\mathbb{E} [\log h(X_{t_n}^h, t_n)] = \mathbb{E} \left[\log h(X_0^h, 0) + \int_0^{t_n} \frac{\sigma^2}{2} \|\nabla \log h(X_s^h, s)\|^2 ds \right].$$

Letting $n \rightarrow \infty$ and applying monotone convergence theorem, we get

$$\liminf_{n \rightarrow \infty} \mathbb{E} [\log h(X_{t_n}^h, t_n)] = \mathbb{E} \left[\log h(X_0^h, 0) + \int_0^T \frac{\sigma^2}{2} \|\nabla \log h(X_s^h, s)\|^2 ds \right]. \quad (\text{B.3})$$

It remains to argue that the left-hand side converges to $\mathbb{E} [\log h(X_T^h, T)]$. Write $Y_t^h = h(X_t^h, t)$. Using the change of measure and $h(x, T) = g(x)$, we get

$$\mathbb{E} [\log Y_t^h] = \mathbb{E} \left[\frac{Y_t}{Y_0} \log Y_t \right].$$

The function $f(y) = y \log y$ is bounded below and convex. If (t_n) is chosen such that $t_n \uparrow T$, we have

$$\mathbb{E} \left[\lim_{n \rightarrow \infty} Y_{t_n} \log Y_{t_n} \mid \mathcal{F}_0 \right] \leq \liminf_{n \rightarrow \infty} \mathbb{E} [Y_{t_n} \log Y_{t_n} \mid \mathcal{F}_0] \leq \mathbb{E} [Y_T \log Y_T \mid \mathcal{F}_0],$$

where Fatou's lemma is applied to obtain the first inequality, and the second inequality follows from the fact that $Y_t \log Y_t$ is a submartingale. Since Y_t converges a.s. to $Y_T = g(X_T)$, we get

$$\liminf_{n \rightarrow \infty} \mathbb{E} [Y_{t_n} \log Y_{t_n} \mid \mathcal{F}_0] = \mathbb{E} [Y_T \log Y_T \mid \mathcal{F}_0].$$

Taking expectations on both sides we get

$$\liminf_{n \rightarrow \infty} \mathbb{E} [\log Y_{t_n}^h] = \mathbb{E} [\log Y_T^h].$$

Combining it with (B.3) proves part (iii).

Consider part (iv). For any bounded and measurable function f , we apply the change of measure to the conditional expectation to get

$$\mathbb{E}_Q[f(X_t) | \mathcal{F}_s] = \frac{\mathbb{E}[f(X_t)h(X_t, t) | \mathcal{F}_s]}{h(X_s, s)}.$$

The claim then follows from Fubini's theorem. $\square$

C PROOFS FOR THE MAIN RESULTS

Recall that X denotes the solution to the SDE

$$dX_t = b(X_t, t)dt + \sigma dW_t$$

over the time interval $[0, T]$ with initial distribution $\mathcal{L}aw(X_0) = \mu_0$, and X^u denotes the solution to the controlled SDE

$$dX_t^u = [b(X_t^u, t) + u_t]dt + \sigma dW_t,$$

with $X_0^u = X_0$. The transition density of the uncontrolled process X is denoted by $p(x, t | y, s)$.

Proof of Lemma 3. Since u is admissible, it must satisfy $\mathbb{E} \int_0^T \|u_t\|^2 dt < \infty$. Therefore, Novikov's condition is satisfied, and we can apply Girsanov theorem to get

$$\begin{aligned} \mathbb{E} \left[\frac{g(X_T)}{h(X_0, 0)} \right] &= \mathbb{E} \left[\frac{g(X_T^u)}{h(X_0^u, 0)} \exp \left\{ - \int_0^T \frac{u_t}{\sigma} dW_t - \int_0^T \frac{\|u_t\|^2}{2\sigma^2} dt \right\} \right] \\ &\geq \exp \left\{ \mathbb{E} \left[\log \frac{g(X_T^u)}{h(X_0^u, 0)} - \int_0^T \frac{u_t}{\sigma} dW_t - \int_0^T \frac{\|u_t\|^2}{2\sigma^2} dt \right] \right\}. \end{aligned}$$

By part (ii) of Theorem B3, the left-hand side of the above inequality is equal to 1. Hence,

$$0 \geq \mathbb{E} \left[\log \frac{g(X_T^u)}{h(X_0^u, 0)} - \int_0^T \frac{u_t}{\sigma} dW_t - \int_0^T \frac{\|u_t\|^2}{2\sigma^2} dt \right] = \mathbb{E} \left[\log \frac{g(X_T^u)}{h(X_0^u, 0)} - \int_0^T \frac{\|u_t\|^2}{2\sigma^2} dt \right],$$

where we have used $\mathbb{E} \int_0^T \|u_t\|^2 dt < \infty$ again to obtain $\mathbb{E}[\int_0^T u_t dW_t] = 0$. By the definition of J_β , we have

$$\begin{aligned} J_\beta(u) &= \beta \mathcal{D}_{\text{KL}}(\mathcal{L}aw(X_T^u), \mu_T) + \mathbb{E} \left[\int_0^T \frac{\|u_t\|^2}{2\sigma^2} dt \right] \\ &\geq \beta \mathcal{D}_{\text{KL}}(\mathcal{L}aw(X_T^u), \mu_T) + \mathbb{E} \left[\log \frac{g(X_T^u)}{h(X_0^u, 0)} \right]. \end{aligned}$$

Since $\mathcal{L}aw(X_0^u) = \mu_0$, we have $\mathbb{E}[\log h(X_0^u, 0)] = \int \log h(x, 0) \mu_0(dx)$. By part (iii) of Theorem B3, the equality is attained when $u_t = (\mathcal{T}g)(X_t^u, t)$. $\square$

Proof of Theorem 4. First, we find using (3.2) that

$$\mathbb{E} \rho_T(X_T) = \int \frac{d\mu_0}{d\nu_0} d\mu_0,$$

which is finite by assumption. Since $p(z, T | x, t) > 0$ whenever $t < T$, we have $h > 0$ on $\mathbb{R}^d \times [0, T)$. Hence, Theorem B3 and Lemma 3 can be applied with $g = \rho_T$. In addition to setting $\rho_T = d\nu_T/d\lambda$ and $f_T = d\mu_T/d\lambda$, we define

$$k_T(x) = \int p(x, T | y, 0) \nu_0(dy), \quad q_T^u(x) = \frac{d\mathcal{L}aw(X_T^u)}{d\lambda}(x).$$

There is no loss of generality in assuming q_T^u exists, since $\mu_T \ll \lambda$ and $\mathcal{D}_{\text{KL}}(\mathcal{L}aw(X_T^u), \mu_T) = \infty$ if $\mathcal{L}aw(X_T^u) \not\ll \mu_T$. For later use, we note that by (3.2) and (3.3),

$$h(y, 0) = \frac{d\mu_0}{d\nu_0}(y), \tag{C.1}$$

$$\rho_T(x) = \left(\frac{f_T(x)}{k_T(x)} \right)^{\beta/(1+\beta)}. \tag{C.2}$$

Let η be the σ -finite measure on $\mathbb{R}^d$ with density $f_T^{\beta/(1+\beta)} k_T^{1/(1+\beta)}$. By Lemma 3, for any admissible control u ,

$$\begin{aligned} J_\beta(u) &\geq \beta \mathcal{D}_{\text{KL}}(\mathcal{L}aw(X_T^u), \mu_T) + \mathbb{E}[\log \rho_T(X_T^u)] - \int \log h(x, 0) \mu_0(dx) \\ &= \beta \mathcal{D}_{\text{KL}}(\mathcal{L}aw(X_T^u), \mu_T) + \mathbb{E}[\log \rho_T(X_T^u)] - \mathcal{D}_{\text{KL}}(\mu_0, \nu_0) \\ &= \mathbb{E} \left[\beta \log \frac{q_T^u(X_T^u)}{f_T(X_T^u)} + \frac{\beta}{1+\beta} \log \frac{f_T(X_T^u)}{k_T(X_T^u)} \right] - \mathcal{D}_{\text{KL}}(\mu_0, \nu_0) \\ &= \beta \mathcal{D}_{\text{KL}}(\mathcal{L}aw(X_T^u), \eta) - \mathcal{D}_{\text{KL}}(\mu_0, \nu_0), \end{aligned}$$

where the second line follows from (C.1) and the third from (C.2). The measures μ_0, ν_0, η do not depend on u . By Lemma B1,

$$\mathcal{D}_{\text{KL}}(\mathcal{L}aw(X_T^u), \eta) \geq -\log C, \quad \text{where } C = \int f_T(x)^{\beta/(1+\beta)} k_T(x)^{1/(1+\beta)} dx.$$

We will later prove that $C = 1$. Combining the above two displayed inequalities, we get

$$J_\beta(u) \geq \beta \mathcal{D}_{\text{KL}}(\mathcal{L}aw(X_T^u), \eta) - \mathcal{D}_{\text{KL}}(\mu_0, \nu_0) \tag{C.3}$$

$$\geq -\log C - \mathcal{D}_{\text{KL}}(\mu_0, \nu_0). \tag{C.4}$$

For $u_t^* = \sigma^2 \nabla \log h(X_t^{u^*}, t)$, we know by Lemma 3 that the equality in (C.3) is attained. Hence, it is optimal if we can show that the equality in (C.4) is also attained. By Lemma B1, this is equivalent to showing that

$$q_T^*(x) = C^{-1} f_T(x)^{\beta/(1+\beta)} k_T(x)^{1/(1+\beta)}, \tag{C.5}$$

where we write $q_T^* = q_T^{u^*}$. By part (iv) of Theorem B3, we have

$$\begin{aligned} q_T^*(x) &= \rho_T(x) \int \frac{p(x, T | y, 0)}{h(y, 0)} \mu_0(dy) \\ &\stackrel{(i)}{=} \rho_T(x) \int p(x, T | y, 0) \nu_0(dy) \\ &= \rho_T(x) k_T(x) \\ &\stackrel{(ii)}{=} f_T(x)^{\beta/(1+\beta)} k_T(x)^{1/(1+\beta)}, \end{aligned}$$

where step (i) follows from (C.1) and step (ii) follows from (C.2). So u^* is optimal, and the normalizing constant C in (C.5) equals 1, from which it follows that $J_\beta(u^*) = -\mathcal{D}_{\text{KL}}(\mu_0, \nu_0)$. Finally, one can apply Jensen's inequality and the assumption $\int (d\mu_0/d\nu_0) d\mu_0 < \infty$ to show that $|\mathcal{D}_{\text{KL}}(\mu_0, \nu_0)| < \infty$, which concludes the proof. $\square$

Proof of Theorem 6. Recall that ν_0, ν_N are defined by

$$\frac{d\mu_0}{d\nu_0}(y) = \int p_N(\mathbf{x}_N | y) \nu_N(d\mathbf{x}_N), \tag{C.6}$$

$$\frac{d\mu_N}{d\lambda}(\mathbf{x}_N) = \rho_N(\mathbf{x}_N)^{\frac{1+\beta}{\beta}} \int p_N(\mathbf{x}_N | y) \nu_0(dy). \tag{C.7}$$

The proof is essentially the same as that of Theorem 6. First, we need to derive a result analogous to Lemma 3, which we give in Lemma C4 below. We will apply Lemma C4 with $g = \rho_N$. Recall that h_j is defined by

$$h_j(x, t; \mathbf{x}_j) = \mathbb{E} \left[\rho_N(\mathbf{x}_j, X_{t_{j+1}}, \dots, X_{t_N}) \mid X_t = x \right], \text{ for } (x, t) \in \mathbb{R}^d \times [t_j, t_{j+1}).$$

Note that the conditions $\mathbb{E} \rho_N(\mathbf{X}_N) < \infty$ and $h_j > 0$ can be verified by the same argument as that used in the proof of Theorem 4. By (C.6), we have

$$h_0(y, 0) = \mathbb{E}[\rho_N(\mathbf{X}_N) \mid X_0 = y] = \frac{d\mu_0}{d\nu_0}(y). \quad (\text{C.8})$$

Define

$$f_N(\mathbf{x}_N) = \frac{d\mu_N}{d\lambda}(\mathbf{x}_N), \quad k_N(\mathbf{x}_N) = \int p_N(\mathbf{x}_N \mid y) \nu_0(dy), \quad q_N^u(\mathbf{x}_N) = \frac{d\mathcal{L}aw(\mathbf{X}_N^u)}{d\lambda}(\mathbf{x}_N).$$

We can rewrite (C.7) as $\rho_N = (f_N/k_N)^{\beta/(1+\beta)}$.

By Lemma C4 and Lemma B1, for any admissible control u ,

$$J_\beta^N(u) \geq -\log C - \mathcal{D}_{\text{KL}}(\mu_0, \nu_0), \quad \text{where } C = \int f_N(\mathbf{x}_N)^{\beta/(1+\beta)} k_N(\mathbf{x}_N)^{1/(1+\beta)} d\mathbf{x}_N.$$

To prove that the equality is attained by the control $u_t^* = (\mathcal{T}_N \rho_N)(X_t^{u^*}, t, \mathbf{X}_N^{u^*})$, it remains to show that

$$q_N^*(\mathbf{x}_N) = C^{-1} f_N(\mathbf{x}_N)^{\beta/(1+\beta)} k_N(\mathbf{x}_N)^{1/(1+\beta)}.$$

where $q_N^* = q_N^{u^*}$. To find q_N^* , we can mimic the proof of Theorem B3. It is not difficult to verify that the law of X^* under $\mathbb{P}$ is the same as the law of X under $\mathbb{Q}$, where $\mathbb{Q}$ is defined by

$$\frac{d\mathbb{Q}}{d\mathbb{P}} = \frac{\mathbb{E}[\rho_N(\mathbf{X}_N) \mid \mathcal{F}_T]}{\mathbb{E}[\rho_N(\mathbf{X}_N) \mid \mathcal{F}_0]} = \frac{\rho_N(\mathbf{X}_N)}{h_0(X_0, 0)}.$$

So for any bounded and measurable function ℓ , we have

$$\mathbb{E}_{\mathbb{Q}}[\ell(\mathbf{X}_N)] = \mathbb{E} \left[\ell(\mathbf{X}_N) \frac{\rho_N(\mathbf{X}_N)}{h_0(X_0, 0)} \right] = \int \ell(\mathbf{x}_N) \rho_N(\mathbf{x}_N) \left\{ \int \frac{p_N(\mathbf{x}_N \mid y)}{h_0(y, 0)} \mu_0(dy) \right\} d\mathbf{x}_N.$$

It thus follows from (C.8) that

$$q_N^*(x) = \rho_N(\mathbf{x}_N) \int \frac{p_N(\mathbf{x}_N \mid y)}{h_0(y, 0)} \mu_0(dy) = f_N(\mathbf{x}_N)^{\beta/(1+\beta)} k_N(\mathbf{x}_N)^{1/(1+\beta)}.$$

The rest of the proof is identical to that of Theorem 4. $\square$

Lemma C4. *Let u be an admissible control and $g: \mathbb{R}^{d \times N} \rightarrow [0, \infty)$ be a measurable function such that $\mathbb{E}[g(\mathbf{X}_N)] < \infty$. Let h_j be as given in (4.2), and assume $h_j > 0$ for each j . For the cost J_β^N defined in (4.1), we have*

$$J_\beta^N(u) \geq \beta \mathcal{D}_{\text{KL}}(\mathcal{L}aw(\mathbf{X}_N^u), \mu_N) + \mathbb{E}[\log g(\mathbf{X}_N^u)] - \int \log h_0(x, 0) \mu_0(dx).$$

The equality is attained when $u_t = (\mathcal{T}_N g)(X_t^u, t, \mathbf{X}_N^u)$.

Proof of Lemma C4. The SDE (2.2) with control $u_t^* = (\mathcal{T}_N g)(X_t^{u^*}, t, \mathbf{X}_N^{u^*})$ can be expressed as

$$\begin{aligned} X_0^* &= X_0, \\ dX_t^* &= [b(X_t^*, t) + \sigma^2 \nabla \log h_j(X_t^*, t; \mathbf{X}_j^*)] dt + \sigma dW_t, \text{ for } t \in [t_j, t_{j+1}), \end{aligned} \quad (\text{C.9})$$

where we write $X^* = X^{u^*}$ and h_j is as given in (4.2). By the tower property, we have

$$h_j(x, t; \mathbf{x}_j) = \mathbb{E} [h_{j+1}(X_{t_{j+1}}, t_{j+1}; \mathbf{x}_j, X_{t_{j+1}}) \mid X_t = x],$$

where h_N is defined by $h_N(x, t_N; \mathbf{x}_{N-1}, x) = g(\mathbf{x}_{N-1}, x)$, and X is the solution to the uncontrolled process (2.1). The assumption $\mathbb{E}[g(\mathbf{X}_N)] < \infty$ implies that $\mathbb{E}[h_j(X_{t_j}, t_j; \mathbf{x}_{j-1}, X_{t_j})] < \infty$ for almost every $\mathbf{x}_{j-1}$. Hence, for each j and $x_j \in \mathbb{R}^d$, Theorem B3 implies that there exists a weak solution to the SDE (C.9) on $[t_j, t_{j+1}]$ with $X_{t_j}^* = x_j$. Moreover,

$$\mathbb{E} \left[\log \frac{h_{j+1}(X_{t_{j+1}}^*, t_{j+1}; \mathbf{X}_{j+1}^*)}{h_j(X_{t_j}^*, t_j; \mathbf{X}_j^*)} \right] = \mathbb{E} \left[\int_{t_j}^{t_{j+1}} \frac{\sigma^2}{2} \|\nabla \log h_j(X_s^*, t; \mathbf{X}_j^*)\|^2 ds \right].$$

Summing over all the N time intervals, we get

$$\mathbb{E} \left[\log \frac{g(\mathbf{X}_N^*)}{h_0(X_0^*, 0)} \right] = \mathbb{E} \left[\sum_{j=0}^{N-1} \log \frac{h_{j+1}(X_{t_{j+1}}^*, t_{j+1}; \mathbf{X}_{j+1}^*)}{h_j(X_{t_j}^*, t_j; \mathbf{X}_j^*)} \right] = \mathbb{E} \left[\int_0^T \frac{\|u_t^*\|^2}{2\sigma^2} ds \right].$$

Now consider an arbitrary admissible control u . As in the proof of Lemma 3, by Girsanov theorem, we have

$$\begin{aligned} 1 &= \mathbb{E} \left[\frac{g(\mathbf{X}_N)}{h_0(X_0, 0)} \right] \\ &= \mathbb{E} \left[\frac{g(\mathbf{X}_N^u)}{h_0(X_0^u, 0)} \exp \left\{ - \int_0^T \frac{u_t}{\sigma} dW_t - \int_0^T \frac{\|u_t\|^2}{2\sigma^2} dt \right\} \right] \\ &\geq \exp \left\{ \mathbb{E} \left[\log \frac{g(\mathbf{X}_N^u)}{h_0(X_0^u, 0)} - \int_0^T \frac{u_t}{\sigma} dW_t - \int_0^T \frac{\|u_t\|^2}{2\sigma^2} dt \right] \right\}, \end{aligned}$$

which yields that

$$\mathbb{E} \left[\int_0^T \frac{\|u_t\|^2}{2\sigma^2} dt \right] \geq \mathbb{E} \left[\log \frac{g(\mathbf{X}_N^u)}{h_0(X_0^u, 0)} \right] = \mathbb{E}[\log g(\mathbf{X}_N^u)] - \int \log h_0(x, 0) \mu_0(dx).$$

The asserted result thus follows. $\square$

D PROOF OF THE EXISTENCE OF SOLUTION TO SSB

We first recall the definition of the Hilbert metric (Chen et al., 2016a). For $K \subset \mathbb{R}^d$ and $1 \leq p \leq \infty$, let $\mathcal{L}^p(K)$ denote the L^p space of functions defined on K . Define

$$\mathcal{L}_+^p(K) = \{f \in \mathcal{L}^p(K) : \inf_{x \in K} f(x) > 0\}, \quad \mathcal{L}_0^p(K) = \{f \in \mathcal{L}^p(K) : \inf_{x \in K} f(x) \geq 0\}.$$

Since $\mathcal{L}_0^p(K)$ is a closed solid cone in the Banach space $\mathcal{L}^p(K)$, we can define a Hilbert metric on it. For any $x, y \in \mathcal{L}_0^p(K) \setminus \{0\}$ (where 0 denotes the constant function equal to 0), define

$$M(x, y) = \inf\{c : x \preceq cy\}, \quad m(x, y) = \sup\{c : cy \preceq x\}, \quad (\text{D.1})$$

where $x \preceq y$ means $y - x \in \mathcal{L}_0^p(K)$ and we use the convention $\inf \emptyset = \infty$. The Hilbert metric d_H on $\mathcal{K} \setminus \{0\}$ is defined by

$$d_H(x, y) = \log \frac{M(x, y)}{m(x, y)}.$$

Note that d_H is only a pseudometric on $\mathcal{K} \setminus \{0\}$, but it is a metric on the space of rays of $\mathcal{K} \setminus \{0\}$.

Proof of Theorem 5. The proof is adapted from Chen et al. (2016a, Proposition 1). We will show the existence of strictly positive and integrable functions ρ_0, ρ_T such that

$$f_0(y) = \rho_0(y) \int_{K_T} p(x, T | y, 0) \rho_T(x) dx, \quad (\text{D.2})$$

$$f_T(x) = \rho_T(x)^{(1+\beta)/\beta} \int_{K_0} p(x, T | y, 0) \rho_0(y) dy. \quad (\text{D.3})$$

Let $\hat{\psi}_0 \in \mathcal{L}_+^\infty(K_0)$ be our guess for f_0/ρ_0 . We can update $\hat{\psi}_0$ as follows.

1. Set $\hat{\rho}_0(y) = f_0(y)/\hat{\psi}_0(y)$.
2. Set $\hat{\psi}_T(x) = \left(\int_{K_0} p(x, T | y, 0) \hat{\rho}_0(y) dy \right)^{\beta/(1+\beta)}$, which is an estimate for $f_T^{\beta/(1+\beta)}/\rho_T$ by (D.3).
3. Set $\hat{\rho}_T(x) = f_T^{\beta/(1+\beta)}/\hat{\psi}_T(x)$.
4. By (D.2), update the estimate for ψ_0 by $\hat{\psi}_0^{\text{new}}(y) = \int_{K_T} p(x, T | y, 0) \hat{\rho}_T(x) dx$.

Denote this updating scheme by $\hat{\psi}_0^{\text{new}} = \mathcal{O}(\hat{\psi}_0)$. Note that for any $c \in (0, \infty)$ and $\psi \in \mathcal{L}_+^\infty(K_0)$,

$$\mathcal{O}(c\psi) = c^{\beta/(1+\beta)} \mathcal{O}(\psi). \quad (\text{D.4})$$

In Lemma D5 below, we prove that $\mathcal{O}$ is a strict contraction mapping from $\mathcal{L}_+^\infty(K_0)$ to $\mathcal{L}_+^\infty(K_0)$ with respect to the Hilbert metric. To prove the existence of ρ_0, ρ_T , it is sufficient to show that $\mathcal{O}$ has a fixed point $\psi_0 \in \mathcal{L}_+^\infty(K_0)$, since we can set

$$\rho_0(y) = \frac{f_0(y)}{\psi_0(y)}, \quad \rho_T(x) = \left(\frac{f_T(x)}{\int_{K_0} p(x, T | y, 0) \rho_0(y) dy} \right)^{\frac{\beta}{1+\beta}},$$

which must satisfy (D.2) and (D.3) and $\rho_0 \in \mathcal{L}_0^1(K_0), \rho_T \in \mathcal{L}_0^1(K_T)$ (see the proof of Lemma D5 for why ρ_0, ρ_T are integrable). But note that we cannot apply Banach fixed-point theorem to $\mathcal{O}$.

To find the fixed point of $\mathcal{O}$, we first consider its normalized version $\tilde{\mathcal{O}}$, defined by $\tilde{\mathcal{O}}(\psi) = \mathcal{O}(\psi)/\|\mathcal{O}(\psi)\|_2$. Let

$$E = \{g \in \mathcal{L}_+^\infty(K_0) : \|g\|_2 = 1\},$$

denote the domain and range of $\tilde{\mathcal{O}}$. Since $\mathcal{O}$ is a strict contraction mapping on $\mathcal{L}_+^\infty(K_0)$ with respect to the Hilbert metric (which is invariant under scaling), $\tilde{\mathcal{O}}$ is also a strict contraction mapping and thus continuous (with respect to the Hilbert metric) on E . If $g \in \mathcal{L}_+^\infty(K_0)$ is a fixed point of $\tilde{\mathcal{O}}$, then $\|\mathcal{O}(g)\|_2^{1+\beta} g \in \mathcal{L}_+^\infty(K_0)$ is a fixed point of $\mathcal{O}$, since

$$\mathcal{O}\left(\|\mathcal{O}(g)\|_2^{1+\beta} g\right) = \|\mathcal{O}(g)\|_2^\beta \mathcal{O}(g) = \|\mathcal{O}(g)\|_2^{1+\beta} \tilde{\mathcal{O}}(g) = \|\mathcal{O}(g)\|_2^{1+\beta} g,$$

where the first equality follows from (D.4). Moreover, since d_H is a metric on the rays, any other fixed point of $\mathcal{O}$ must have the form cg for some constant $c > 0$. But the same argument shows that c must equal $\|\mathcal{O}(g)\|_2^{1+\beta}$, and thus the fixed point of $\mathcal{O}$ is unique. This further yields the uniqueness of (ρ_0, ρ_T) .

So it only remains to prove that $\tilde{\mathcal{O}}$ has a fixed point in $\mathcal{L}_+^\infty(K_0)$. The proof for this claim is essentially the same as that in Chen et al. (2016a). Pick arbitrarily $\psi^{(0)} \in \mathcal{L}_+^\infty(K_0)$ and let $g_0 = \psi^{(0)}/\|\psi^{(0)}\|_2$. For $k = 1, 2, \dots$, define

$$g_k := \frac{\mathcal{O}^k(\psi^{(0)})}{\|\mathcal{O}^k(\psi^{(0)})\|_2} = \tilde{\mathcal{O}}(g_{k-1}),$$

where the second equality follows from (D.4). Each g_k is in E and well-defined as $\mathcal{O}^k(\psi^{(0)}) \in \mathcal{L}_+^\infty(K_0) \subset \mathcal{L}_2(K_0)$. Since $\tilde{\mathcal{O}}$ is a strict contraction mapping with respect to d_H , $\{g_k\}_{k \geq 0}$ is a Cauchy sequence with respect to d_H . Using the inequality $\|g_k - g_l\|_2 \leq e^{d_H(g_k, g_l)} - 1$ (see Chen et al. (2016a)), we find that $\{g_k\}_{k \geq 0}$ is also Cauchy with respect to the L^2 -norm, and thus there exists $g \in \mathcal{L}_0^2(K_0)$ such that $\lim_{k \rightarrow \infty} \|g_k - g\|_2 = 0$ and $\|g\|_2 = 1$. Next, we argue that $\{g_k\}_{k \geq 0}$ is uniformly bounded from below and above and also uniformly equicontinuous. To show the uniform boundedness, we first observe that $\|g_k\|_2 = 1$ implies

$$\sup_x g_k(x) \geq \frac{1}{\sqrt{\lambda(K_0)}} \geq \inf_x g_k(x). \quad (\text{D.5})$$

Further, since $p(x, T | y, 0)$ is bounded in (x, y) and recalling the last step in the construction of $\mathcal{O}$, we have

$$\frac{\sup_x (\mathcal{O}\psi)(x)}{\inf_x (\mathcal{O}\psi)(x)} \leq \epsilon^{-1} \quad (\text{D.6})$$

for any $\psi \in \mathcal{L}_+^\infty(K_0)$, where $\epsilon \in (0, 1]$ is some constant independent of ψ . Combining (D.5) and (D.6), we get

$$\frac{\epsilon}{\sqrt{\lambda(K_0)}} \leq \inf_x g_k(x) \leq \sup_x g_k(x) \leq \frac{1}{\epsilon \sqrt{\lambda(K_0)}}.$$

Note that the uniform boundedness of $\{g_k\}$ also implies $g \in \mathcal{L}_+^\infty(K_0)$. The uniform equicontinuity of $\{g_k\}$ can be proved by using the uniform continuity of the transition density p . Finally, by Arzelà–Ascoli Theorem, there is a subsequence $\{g_{k_n}\}$ such that g_{k_n} converges to g uniformly with respect to the L^2 -norm and g is also uniformly continuous. This implies $d_H(g_{k_n}, g) \rightarrow 0$ and thus $d_H(g_k, g) \rightarrow 0$. By the continuity (with respect to d_H) of $\tilde{\mathcal{O}}$, we can interchange the limit operation with $\tilde{\mathcal{O}}$, thereby establishing g as the fixed point of $\tilde{\mathcal{O}}$. $\square$

Lemma D5. For $\psi \in \mathcal{L}_+^\infty(K_0)$, define

$$(\mathcal{O}\psi)(x) = \int_{K_T} p(z, T | x, 0) \left(\frac{f_T(z)}{\int_{K_0} p(z, T | y, 0) f_0(y) \psi(y)^{-1} dy} \right)^{\beta/(1+\beta)} dz.$$

The operator $\mathcal{O}$ is a strict contraction mapping from $\mathcal{L}_+^\infty(K_0)$ to $\mathcal{L}_+^\infty(K_0)$ with respect to the Hilbert metric.

Proof. We can express the operator $\mathcal{O}$ by $\mathcal{O} = \mathcal{E}_T \circ \mathcal{P} \circ \mathcal{I} \circ \mathcal{E}_0 \circ \mathcal{I}$, and define $\mathcal{I}, \mathcal{E}_0, \mathcal{E}_T, \mathcal{P}$ by

$$\begin{aligned} \mathcal{I}: \mathcal{L}_+^\infty(K) &\longrightarrow \mathcal{L}_+^\infty(K), & (\mathcal{I}\varphi)(x) &= \varphi(x)^{-1} \quad \text{for } K = K_0, K_T, \\ \mathcal{E}_0: \mathcal{L}_+^\infty(K_0) &\longrightarrow \mathcal{L}_+^\infty(K_T), & (\mathcal{E}_0\varphi)(x) &= \int_{K_0} p(x, T | y, 0) f_0(y) \varphi(y) dy, \\ \mathcal{E}_T: \mathcal{L}_+^\infty(K_T) &\longrightarrow \mathcal{L}_+^\infty(K_0), & (\mathcal{E}_T\varphi)(x) &= \int_{K_T} p(z, T | x, 0) f_T(z)^{\beta/(1+\beta)} \varphi(z) dz, \\ \mathcal{P}: \mathcal{L}_+^\infty(K_T) &\longrightarrow \mathcal{L}_+^\infty(K_T), & (\mathcal{P}\varphi)(x) &= \varphi(x)^{\beta/1+\beta}. \end{aligned}$$

It is worth explaining how the ranges of these operators are determined. First, if $\varphi \in \mathcal{L}_+^\infty(K)$, it is clear that $\mathcal{I}\varphi \in \mathcal{L}_+^\infty(K)$ and $\mathcal{P}\varphi \in \mathcal{L}_+^\infty(K)$. For $\mathcal{E}_0$, since we assume K_0 is compact and $p(x, T | y, 0)$ is continuous in (x, y) , for any $\varphi \in \mathcal{L}_+^\infty(K_0)$ there exists $\epsilon > 0$ such that

$$\epsilon = \epsilon \int_{K_0} f_0(y) dy \leq (\mathcal{E}_0\varphi)(x) \leq \epsilon^{-1} \int_{K_0} f_0(y) dy = \epsilon^{-1}.$$

The argument for $\mathcal{E}_T$ is similar, and note that by Hölder's inequality,

$$\int_{K_T} f_T(z)^{\beta/(1+\beta)} dz \leq \lambda(K_T)^{1/(1+\beta)} < \infty.$$

Now we prove that $\mathcal{P}$ is a strict contraction. Let $\psi_1, \psi_2 \in \mathcal{L}_+^\infty(K_T)$. By the definition given in (D.1),

$$M(\mathcal{P}(\psi_1), \mathcal{P}(\psi_2)) = M(\psi_1, \psi_2)^{\beta/1+\beta} \text{ and } m(\mathcal{P}(\phi_1), \mathcal{P}(\phi_2)) = m(\phi_1, \phi_2)^{\beta/1+\beta},$$

which implies

$$d_H(\mathcal{P}(\phi_1), \mathcal{P}(\phi_2)) = \log \left(\frac{M(\mathcal{P}(\phi_1), \mathcal{P}(\phi_2))}{m(\mathcal{P}(\phi_1), \mathcal{P}(\phi_2))} \right) = \frac{\beta}{1+\beta} d_H(\phi_1, \phi_2) < d_H(\phi_1, \phi_2),$$

since $\beta < \infty$. Chen et al. (2016a) showed that the operators $\mathcal{E}_0, \mathcal{E}_T$ are strict contractions using Birkhoff's theorem, and that the operator $\mathcal{I}$ is an isometry (all with respect to the Hilbert metric). Hence, $\mathcal{O}$ is a strict contraction. Note that for our problem, since $\mathcal{P}$ is a strict contraction, we actually only need $\mathcal{E}_0, \mathcal{E}_T$ to be contractions (not necessarily strict). $\square$

E PROOF OF LEMMA 7

We consider a more general setting. Assume that X_t is given by

$$X_0 = x_0, \text{ and } dX_t = b(X_t, t)dt + \sigma dW_t,$$

where b satisfies Assumption 1. The density function of $\mathcal{L}aw(X_T)$ is $p(x, T | x_0, 0)$. Let X^{ref} be given by

$$\begin{aligned} dX_t^{\text{ref}} &= b^{\text{ref}}(X_t^{\text{ref}}, t)dt + \sigma dW_t, \\ \text{where } b^{\text{ref}}(x, t) &= b(x, t) + \sigma^2 \nabla \log h^{\text{ref}}(x, t), \\ h^{\text{ref}}(x, t) &= \mathbb{E} \left[\frac{f_{\text{ref}}(X_T)}{p(X_T, T | x_0, 0)} \mid X_t = x \right]. \end{aligned}$$

Theorem 1 implies that $\mathcal{L}aw(X_T^{\text{ref}}) = \mu_{\text{ref}}$. Let X^{obj} be given by

$$\begin{aligned} dX_t^{\text{obj}} &= b^{\text{obj}}(X_t^{\text{obj}}, t)dt + \sigma dW_t, \\ \text{where } b^{\text{obj}}(x, t) &= b^{\text{ref}}(x, t) + \sigma^2 \nabla \log h^{\text{obj}}(x, t), \\ h^{\text{obj}}(x, t) &= \mathbb{E} \left[\left(\frac{f_{\text{obj}}(X_T^{\text{ref}})}{f_{\text{ref}}(X_T^{\text{ref}})} \right)^{\beta/(1+\beta)} \mid X_t^{\text{ref}} = x \right], \end{aligned} \tag{E.1}$$

which is the solution to Problem 2 with X^{ref} being the reference process and μ_{obj} being the target distribution. We now prove that

$$b^{\text{obj}}(x, t) = b(x, t) + \sigma^2 \nabla \log \mathbb{E} \left[\frac{f_{\text{ref}}(X_T)^{\frac{1}{1+\beta}} f_{\text{obj}}(X_T)^{\frac{\beta}{1+\beta}}}{p(X_T, T | x_0, 0)} \mid X_t = x \right]. \tag{E.2}$$

That is, X^{obj} is also the solution to Problem 1 with X being the reference process and μ^* being the target distribution, where μ^* has un-normalized density $f_{\text{ref}}^{1/(1+\beta)} f_{\text{obj}}^{\beta/(1+\beta)}$. Once this is proved, Lemma 7 follows as a special case with $b \equiv 0$ and $x_0 = 0$.

Proof. To prove (E.2), we use part (ii) of Theorem B3. Define a martingale

$$Z_t^{\text{ref}} = \frac{h^{\text{ref}}(X_t, t)}{h^{\text{ref}}(X_0, 0)}.$$

Let $\mathbb{Q}^{\text{ref}}$ be the probability measure given by $d\mathbb{Q}^{\text{ref}} = Z_T^{\text{ref}} d\mathbb{P}$. By part (ii) of Theorem B3, the law of X under $\mathbb{Q}^{\text{ref}}$ is the same as the law of X^{ref} under $\mathbb{P}$. Applying the change of measure to (E.1) yields

$$\begin{aligned} h^{\text{obj}}(x, t) &= \mathbb{E} \left[\left(\frac{f_{\text{obj}}(X_T^{\text{ref}})}{f_{\text{ref}}(X_T^{\text{ref}})} \right)^{\beta/(1+\beta)} \frac{Z_T^{\text{ref}}}{Z_t^{\text{ref}}} \mid X_t = x \right], \\ &= h^{\text{ref}}(X_t, t)^{-1} \mathbb{E} \left[\left(\frac{f_{\text{obj}}(X_T^{\text{ref}})}{f_{\text{ref}}(X_T^{\text{ref}})} \right)^{\beta/(1+\beta)} h^{\text{ref}}(X_T, T) \mid X_t = x \right], \\ &= h^{\text{ref}}(X_t, t)^{-1} \mathbb{E} \left[\frac{f_{\text{ref}}(X_T)^{\frac{1}{1+\beta}} f_{\text{obj}}(X_T)^{\frac{\beta}{1+\beta}}}{p(X_T, T | x_0, 0)} \mid X_t = x \right]. \end{aligned}$$

Since

$$\begin{aligned} b^{\text{obj}}(x, t) &= b^{\text{ref}}(x, t) + \sigma^2 \nabla \log h^{\text{obj}}(x, t) \\ &= b(x, t) + \sigma^2 \nabla \log h^{\text{ref}}(x, t) + \sigma^2 \nabla \log h^{\text{obj}}(x, t), \end{aligned}$$

we obtain (E.2). □

F MNIST EXAMPLE

Figure F4 visualizes the 50 images in the data set $\mathcal{D}_{\text{obj}}$, which are obtained by adding Gaussian noise to the original images in MNIST. Figure F5 shows the new images generated by the two-stage Schrödinger bridge algorithm of Wang et al. (2021) using only $\mathcal{D}_{\text{obj}}$ as the input.

Table F2 shows the inception scores (Salimans et al., 2016) for our generated images and the images of digit 8 in MNIST. The score of our samples for $\beta = 1.5$ is slightly higher than that of the digit 8 in MNIST dataset, suggesting that our generated images of digit 8 exhibit a greater degree of variability than those in MNIST. Additionally, when $\beta = 100$, our score aligns closely with that of $\mathcal{D}_{\text{obj}}$ (i.e., the noisy digit 8 images from MNIST), indicating that our method can recover the images in the target data set by using a large β . For small values of β , the scores of our generated images are higher than that of digit 8 images in MNIST, primarily because the reference dataset (consisting of the other digits) has greater variability and complexity. However, as shown in Figure 1, when β is small, we do not necessarily get images of digit 8.

We also utilize t-SNE plots to visually characterize the distribution of our generated images. Figure F6 illustrates that our samples come from the geometric mixture distribution interpolating between the noisy images of digit 8 and the clean images of other digits. Figure F7 demonstrates that SSB samples with $\beta = 1.5$ are positioned closer to the clean images of digit 8 compared to the samples obtained with $\beta = 100$.

In our code, we use the neural network model of Song and Ermon (2019) for training the score functions and use the neural network model of Wang et al. (2021) for training the density ratio function.

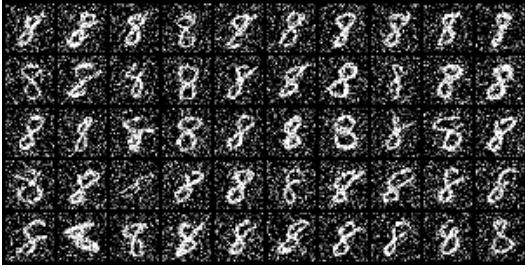


Figure F4: Samples in $\mathcal{D}_{\text{obj}}$.

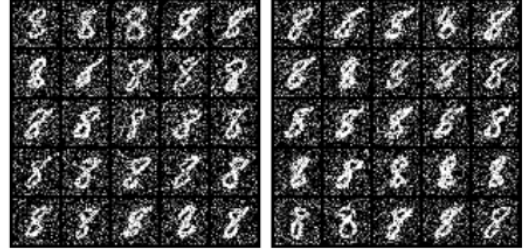


Figure F5: Samples Generated by the Algorithm of Wang et al. (2021) Using Only $\mathcal{D}_{\text{obj}}$.

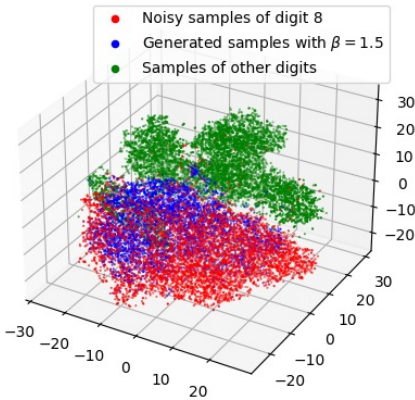


Figure F6: t-SNE Plot Illustrating the Geometric Mixture Distribution with $\beta = 1.5$.

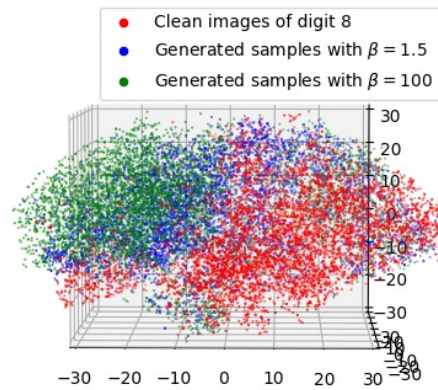


Figure F7: t-SNE Plot Comparing SSB Samples with Clean Images in MNIST.

Table F2: Inception Scores (mean $\pm$ sd) for SSB Samples and the Images in MNIST.

Datasets (sample size $\approx$ 5K)	Inception score
SSB with $\beta = 0$	6.70 ± 0.20
SSB with $\beta = 0.25$	6.59 ± 0.15
SSB with $\beta = 0.7$	5.12 ± 0.11
SSB with $\beta = 1.5$	3.51 ± 0.08
SSB with $\beta = 4$	3.65 ± 0.04
SSB with $\beta = 100$	2.87 ± 0.04
Digit 8 in MNIST (clean)	3.29 ± 0.04
Digit 8 in MNIST (noisy)	2.96 ± 0.04