
Adjoint Schrödinger Bridge Sampler

Guan-Horng Liu^{1,*}, Jaemoo Choi^{2,*}, Yongxin Chen², Benjamin Kurt Miller¹,
Ricky T. Q. Chen^{1,*}

¹FAIR at Meta, ²Georgia Institute of Technology, *Core contributors

Abstract

Computational methods for learning to sample from the Boltzmann distribution—where the target distribution is known only up to an unnormalized energy function—have advanced significantly recently. Due to the lack of explicit target samples, however, prior diffusion-based methods, known as *diffusion samplers*, often require importance-weighted estimation or complicated learning processes. Both trade off scalability with extensive evaluations of the energy and model, thereby limiting their practical usage. In this work, we propose **Adjoint Schrödinger Bridge Sampler (ASBS)**, a new diffusion sampler that employs simple and scalable matching-based objectives yet without the need to estimate target samples during training. ASBS is grounded on a mathematical model—the Schrödinger Bridge—which enhances sampling efficiency via kinetic-optimal transportation. Through a new lens of stochastic optimal control theory, we demonstrate how SB-based diffusion samplers can be learned at scale via Adjoint Matching and prove convergence to the global solution. Notably, ASBS generalizes the recent Adjoint Sampling (Havens et al., 2025) to arbitrary source distributions by relaxing the so-called memoryless condition that largely restricts the design space. Through extensive experiments, we demonstrate the effectiveness of ASBS on sampling from classical energy functions, amortized conformer generation, and molecular Boltzmann distributions. Codes are available at https://github.com/facebookresearch/adjoint_samplers.

1 Introduction

Sampling from Boltzmann distributions is a fundamental problem in computational science, with widespread applications in Bayesian inference, statistical physics, and chemistry (Box and Tiao, 2011; Binder et al., 1992; Tuckerman, 2023). Mathematically, we aim to sample from a target distribution $\nu(x)$ known up to a unnormalized, often differentiable, energy function $E(x) : \mathcal{X} \subseteq \mathbb{R}^d \rightarrow \mathbb{R}$,

$$\nu(x) := \frac{e^{-E(x)}}{Z}, \quad \text{where } Z := \int_{\mathcal{X}} e^{-E(x)} dx \quad (1)$$

is an intractable normalization constant. For instance, the energy function $E(x)$ of a molecular system quantifies the stability of a chemical structure based on the 3D positions of particles. A lower energy indicates a more stable structure and hence a higher likelihood of its occurrence, *i.e.*, $\nu(x) \propto e^{-E(x)}$.

Classical methods that generate samples from $\nu(x)$ rely on Markov Chain Monte Carlo algorithms, which run a Markov chain whose stationary distribution is $\nu(x)$ (Metropolis et al., 1953; Neal, 2001; Del Moral et al., 2006). These methods, however, tend to suffer from slow mixing time and require extensive evaluations of energy function, limiting their practical usages due to prohibitive complexity.

To improve sampling efficiency, modern samplers focus on learning better proposal distributions (Noé et al., 2019; Midgley et al., 2023). Among those, recent advances in diffusion-based generative models (Song et al., 2021; Ho et al., 2020) have given rise to a family of *Diffusion Samplers*, which

Table 1: Compared to prior diffusion samplers, **Adjoint Schrödinger Bridge Sampler (ASBS)** offers the most flexible design for diffusion samplers (2), while learning the drift u_t^θ via scalable matching objectives that do not rely on computation of importance weights (IW).s).

Method	Design condition for (2)		Learning method for u_t^θ	
	Non-memoryless	Arbitrary prior	Matching objective ¹	No reliance on IWs
PIS (Zhang and Chen, 2022) DDS (Vargas et al., 2023)	✗	✗	✗	✓
LV-PIS & LV-DDS (Richter and Berner, 2024)	✗	✗	✗	✗
PDDS (Phillips et al., 2024) iDEM (Akhound-Sadeh et al., 2024)	✗	✗	✓	✗
AS (Havens et al., 2025)	✗	✗	✓	✓
Sequential SB (Bernton et al., 2019)	✓	✓	✗	✗
Adjoint Schrödinger Bridge Sampler (Ours)	✓	✓	✓	✓

consider stochastic differential equations (SDEs) of the following form:

$$dX_t = [f_t(X_t) + \sigma_t u_t^\theta(X_t)] dt + \sigma_t dW_t, \quad X_0 \sim \mu(X_0), \quad (2)$$

where $f_t(x) : [0, 1] \times \mathcal{X} \rightarrow \mathcal{X}$ the base drift, $\sigma_t : [0, 1] \rightarrow \mathbb{R}_{>0}$ the noise schedule, and $\mu(x)$ the initial source distribution. Given (f_t, σ_t, μ) , the diffusion sampler learns a parametrized drift $u_t^\theta(x)$ transporting samples to the target distribution $\nu(x)$ at the terminal time $t = 1$.

Computational methods for learning diffusion samplers have grown significantly recently (Zhang and Chen, 2022; Vargas et al., 2023; Berner et al., 2024; Chen et al., 2025). Due to the distinct problem setup in (1), the target distribution is defined exclusively by its energy $E(x)$, rather than by explicit target samples. This characteristic renders modern generative modeling techniques for scalability—particularly the score matching objectives¹—less applicable. As such, prior matching-based diffusion samplers (Phillips et al., 2024; Akhound-Sadeh et al., 2024; De Bortoli et al., 2024) often require computationally intensive estimation of target samples via importance weights (IW).s).

Recently, Havens et al. (2025) introduced Adjoint Sampling (AS), a new class of diffusion samplers whose matching objectives rely only on on-policy samples, thereby greatly enhancing scalability. By incorporating stochastic optimal control (SOC) theory (Kappen, 2005; Todorov, 2007), AS facilitates the use of Adjoint Matching (Domingo-Enrich et al., 2025), a novel matching objective that imposes self-consistency in generated samples, effectively eliminating the needs for target samples.

The efficiency of AS, however, is achieved through a specific instantiation of the SDE (2) to satisfy the so-called *memoryless* condition. This condition—formally discussed in Section 2—restricts its source distribution to be Dirac delta $\mu(x) := \delta$, precluding the use of common priors such as Gaussian or domain-specific priors such as the harmonic oscillators in molecular systems (Jing et al., 2023). Notably, the memoryless condition underlies *all* previous matching-based diffusion samplers, restricting the design space of (2) from other choices known to enhance transportation efficiency (Shaul et al., 2023). While the condition has been relaxed in non-matching-based methods at extensive computational complexity (Richter and Berner, 2024; Bernton et al., 2019), no existing diffusion sampler—to our best understanding—has successfully combined matching objectives with non-memoryless condition. Table 1 summarizes the comparison between prior diffusion samplers.

In this work, we propose **Adjoint Schrödinger Bridge Sampler (ASBS)**, a new adjoint-matching-based diffusion sampler that eliminates the requirement for memoryless condition entirely. Formally, ASBS recasts learning diffusion sampler as a distributionally constrained optimization, known as the Schrödinger Bridge (SB) problem (Schrödinger, 1931, 1932; Léonard, 2013; Chen et al., 2016):

$$\min_u D_{\text{KL}}(p^u || p^{\text{base}}) = \mathbb{E}_{X \sim p^u} \left[\int_0^1 \frac{1}{2} \|u_t^\theta(X_t)\|^2 dt \right], \quad (3a)$$

$$\text{s.t. } dX_t = [f_t(X_t) + \sigma_t u_t^\theta(X_t)] dt + \sigma_t dW_t, \quad X_0 \sim \mu(X_0), \quad X_1 \sim \nu(X_1). \quad (3b)$$

Here, p^u denotes the path distribution induced by the SDE in (3b), whereas $p^{\text{base}} := p^{u:=0}$ denotes the path distribution induced by the “base” SDE when $u_t := 0$. By minimizing their KL divergence, the SB problem (3) seeks the kinetic-optimal drift u_t^* —an optimality structure well correlated

¹The matching objective is a simple regression loss, $\mathbb{E} \|u_t^\theta(X_t) - v_t(X_t, X_1)\|^2$, w.r.t. some tractable v_t .

with sampling efficiency in generative modeling (Finlay et al., 2020; Liu et al., 2023). Since the SOC problem in AS corresponds to a specific case of the SB problem with $(f_t, \mu) := (0, \delta)$, ASBS extends AS to handle non-memoryless conditions by solving more general SB problems (see Theorem 3.1). Computationally, ASBS retains all scalability advantages from AS by utilizing an *adjoint*-matching objective that removes the need for estimating target samples. It also introduces a *corrector*-matching objective to correct nontrivial biases arising from non-memoryless conditions. We prove that alternating optimization between the two matching objectives is equivalent to executing the Iterative Proportional Fitting algorithm (Kullback, 1968), ensuring global convergence of ASBS to u_t^* (see Theorem 3.2). Though extensive experiments, we show superior performance of ASBS over prior diffusion samplers across various benchmarks on sampling multi-particle energy functions.

In summary, we present the following contributions:

- We introduce **ASBS**, an SB-based diffusion sampler capable of sampling target distributions using only unnormalized energy functions, by solving general SB problems with arbitrary priors.
- We base ASBS on a new SOC framework that removes the restrictive memoryless condition, develop a scalable matching-based algorithm, and prove theoretical convergence to global solution.
- We show ASBS’s superior performance over prior methods on sampling Boltzmann distributions of classical energy functions, alanine dipeptide molecule and amortized conformer generation.

2 Preliminary

We revisit the memoryless condition introduced by Domingo-Enrich et al. (2025) and examine its impact on the constructions of SOC-based diffusion samplers (Zhang and Chen, 2022; Havens et al., 2025), which are closely related to our ASBS. Additional review can be found in Appendix A.

Stochastic Optimal Control (SOC) The SOC problem (4) studies an optimization problem:

$$\min_u \mathbb{E}_{X \sim p^u} \left[\int_0^1 \frac{1}{2} \|u_t(X_t)\|^2 dt + g(X_1) \right] \quad \text{s.t. (2),} \quad (4)$$

which, unlike the SB problem (3), includes an additional *terminal cost* $g(x) : \mathcal{X} \rightarrow \mathbb{R}$ at the terminal time $t = 1$ and considers the SDE without the terminal constraint $X_1 \sim \nu$. The primary reason for studying this specific optimization problem is that the optimal distribution is known analytically by²

$$p^*(X_0, X_1) = p^{\text{base}}(X_0, X_1) e^{-g(X_1) + V_0(X_0)}, \quad \text{where } V_0(x) = -\log \int p_{1|0}^{\text{base}}(y|x) e^{-g(y)} dy \quad (5)$$

is the initial value function. That is, the optimal distribution p^* is an exponentially tilted version of the base distribution, $p^{\text{base}} := p^{u:=0}$. Specifically, p^{base} is tilted by the terminal cost “ $-g(X_1)$ ” and the initial value function $V_0(X_0)$, which is intractable. Consequently, to ensure its marginal $p^*(X_1)$ follows the target distribution $\nu(X_1)$, we must eliminate the *initial value function bias* from $V_0(X_0)$.

Memoryless condition & SOC-based diffusion sampler A common approach to eliminate the aforementioned initial value function bias, adopted by most diffusion samplers, is to restrict the class of base processes to be *memoryless*. Formally, the memoryless condition assumes statistical independency between X_0 and X_1 in the base distribution:

$$p^{\text{base}}(X_0, X_1) \stackrel{\text{memoryless}}{:=} p^{\text{base}}(X_0) p^{\text{base}}(X_1). \quad (6)$$

This memoryless condition (6) simplifies the optimal distribution at the terminal time $t = 1$ and, upon choosing a proper terminal cost $g(x)$, recovers the target distribution ν ,

$$p^*(X_1) \stackrel{\text{memoryless}}{=} \int p^{\text{base}}(X_0) p^{\text{base}}(X_1) e^{-g(X_1) + V_0(X_0)} dX_0 \propto p^{\text{base}}(X_1) e^{-g(X_1)} = \nu(X_1),$$

where the last equality is due to setting the terminal cost to $g(x) := \log \frac{p_1^{\text{base}}(x)}{\nu(x)}$. Typically, the memoryless condition (6) is enforced by a careful design of the base distribution p^{base} or, equivalently,

²Equation (5) can be obtained by rewriting (4) as $D_{\text{KL}}(p^u || p^{\text{base}}) + \mathbb{E}_{p_1^u}[g(X_1)]$ and then computing the analytic solution $p^*(X_1|X_0) \propto p^{\text{base}}(X_1|X_0) e^{-g(X_1)}$ and normalization $\int p^{\text{base}}(X_1|X_0) e^{-g(X_1)} dX_1 = e^{-V_0(X_0)}$. See Appendix A.1 for details.

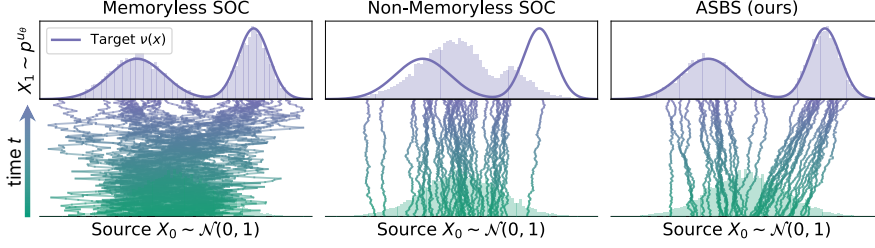


Figure 1: Effect of the memoryless condition on learning SOC-based diffusion samplers. We consider Gaussian prior $\mu(x) := \mathcal{N}(x; 0, 1)$ with (f_t, σ_t) set to VP-SDE for the first plot and $(0, 0.2)$ for the rest; see Appendix A.1 for details. The memoryless condition injects significant noise (**left**) to correct the otherwise biased optimization (**middle**), whereas ASBS can successfully debias any non-memoryless processes (**right**).

the parameters (f_t, σ_t, μ) in (2). For instance, the variance-preserving process (VP; Song et al., 2021) considers a linear base drift f_t , a noise schedule σ_t that grows significantly with time, and a Gaussian prior μ ; see Figure 1. Alternatively, one could implement (6) with Dirac delta prior $\mu(x) := \delta_0(x)$ and $f_t := 0$, leading to the following SOC problem (Zhang and Chen, 2022):

$$\min_u \mathbb{E}_{X \sim p^u} \left[\int_0^1 \frac{1}{2} \|u_t(X_t)\|^2 dt + \log \frac{p_1^{\text{base}}(X_1)}{\nu(X_1)} \right] \quad \text{s.t. } dX_t = \sigma_t u_t(X_t) dt + \sigma_t dW_t, \quad X_0 = 0. \quad (7)$$

Based on the aforementioned reasoning, solving (7) results in a diffusion sampler that transports samples to the target distribution at $t=1$, with Adjoint Sampling (Havens et al., 2025) as the only scalable method of this class. Despite encouraging, the SOC problem in (7) is nevertheless limited by its trivial source, precluding potentially more effective options for sampling Boltzmann distributions.

3 Adjoint Schrödinger Bridge Sampler

We introduce a new diffusion sampler by solving the SB problem (3), where the target distribution $\nu(x)$ is given by its energy function $E(x)$ rather than explicit samples. All proofs are left in Appendix B.

3.1 SOC Characteristics of the SB Problem

The SB problem (3)—as an optimization problem with distribution constraints—is widely explored in optimal transport, stochastic control, and recently machine learning (Léonard, 2012; Chen et al., 2021; De Bortoli et al., 2021). Its kinetic-optimal drift u^* satisfies the following optimality equations:

$$u_t^*(x) = \sigma_t \nabla \log \varphi_t(x), \quad \text{where} \quad \begin{cases} \varphi_t(x) = \int p_{1|t}^{\text{base}}(y|x) \varphi_1(y) dy, & \varphi_0(x) \hat{\varphi}_0(x) = \mu(x) \\ \hat{\varphi}_t(x) = \int p_{t|0}^{\text{base}}(x|y) \hat{\varphi}_0(y) dy, & \varphi_1(x) \hat{\varphi}_1(x) = \nu(x) \end{cases} \quad (8a) \quad (8b)$$

and $p_{t|s}^{\text{base}}(y|x) := p^{\text{base}}(X_t=y|X_s=x)$ is the transition kernel of the base process for observing y at time t given x at time s . The SB potentials $\varphi_t(x), \hat{\varphi}_t(x) \in C^{1,2}([0, 1], \mathbb{R}^d)$ are then defined (up to some multiplicative constant) as solutions to forward and backward time integrations w.r.t. $p_{t|s}^{\text{base}}$.

Equation (8) are computationally challenging to solve—even when $p_{t|s}^{\text{base}}$ has an analytical solution—due to the intractable integration and coupled boundaries at $t = 0$ and 1 . Our key observation is that the first equation (8a) resembles the optimality condition of the SOC problem (4) (see Appendix A.1). This implies that the optimality conditions of SB hints an SOC reinterpretation, which, as we will demonstrate, is more tractable than solving (8) directly. We formalize our finding below.

Theorem 3.1 (SOC characteristics of SB). *The kinetic-optimal drift u_t^* in (8) solves an SOC problem*

$$\min_u \mathbb{E}_{X \sim p^u} \left[\int_0^1 \frac{1}{2} \|u_t(X_t)\|^2 dt + \log \frac{\hat{\varphi}_1(X_1)}{\nu(X_1)} \right] \quad \text{s.t. (2)}. \quad (9)$$

Theorem 3.1 suggests that *every* SB problem (3) can be solved like an SOC problem (4) with the terminal cost $g(x) := \log \frac{\hat{\varphi}_1(x)}{\nu(x)}$. Comparing to the formulation in Adjoint Sampling (Havens et al., 2025), the two SOC problems, namely (7) and (9), differ in their terminal costs—where p_1^{base} is replaced by $\hat{\varphi}_1$ —and the relaxation of the source distribution from Dirac delta $X_0 = 0$ to general source $\mu(X_0)$.

How $\hat{\varphi}_1(\cdot)$ debiases non-memoryless SOC problems Taking a closer look at the effect of $\hat{\varphi}_1$, notice that the optimal distribution of the SB problem—according to Theorem 3.1 and (5)—follows

$$p^*(X_0, X_1) = p^{\text{base}}(X_0, X_1) \exp \left(-\log \frac{\hat{\varphi}_1(X_1)}{\nu(X_1)} - \log \varphi_0(X_0) \right), \quad (10)$$

where “ $-\log \varphi_0$ ” is the equivalent initial value function. One can verify that the marginal at the terminal time $t = 1$ indeed satisfies the target distribution,

$$\begin{aligned} p^*(X_1) &= \int p^*(X_0, X_1) dX_0 \stackrel{(10)}{=} \frac{\nu(X_1)}{\hat{\varphi}_1(X_1)} \int p^{\text{base}}(X_0, X_1) \frac{1}{\varphi_0(X_0)} dX_0 \\ &\stackrel{(8a)}{=} \frac{\nu(X_1)}{\hat{\varphi}_1(X_1)} \int p^{\text{base}}(X_1|X_0) \hat{\varphi}_0(X_0) dX_0 \stackrel{(8b)}{=} \nu(X_1). \end{aligned} \quad (11)$$

That is, the optimality equations in (8), in their essence, construct a specific function $\hat{\varphi}_1(\cdot)$ that eliminates the initial value function bias associated with any non-memoryless processes, thereby ensuring that the optimal distribution satisfies the target ν at $t = 1$.

3.2 Adjoint Sampling with General Source Distribution

We now specialize Theorem 3.1 to sampling Boltzmann distributions (1), where $\nu(x) \propto e^{-E(x)}$, and hence the terminal cost of the new SOC problem in (9) becomes $\log \frac{\hat{\varphi}_1(x)}{\nu(x)} = E(x) + \log \hat{\varphi}_1(x)$. To encourage minimal transportation cost (Chen and Georgiou, 2015; Peyré and Cuturi, 2017), we consider the Brownian-motion base process with a degenerate base drift $f_t := 0$. Applying Adjoint Matching (AM; Domingo-Enrich et al., 2025) to the resulting SOC problem leads to

$$u^* = \arg \min_u \mathbb{E}_{p_{t|0,1}^{\text{base}} p_{0,1}^{\bar{u}}} [\|u_t(X_t) + \sigma_t (\nabla E + \nabla \log \hat{\varphi}_1)(X_1)\|^2], \quad \bar{u} = \text{stopgrad}(u). \quad (12)$$

Note that the AM objective in (12) functions as a self-consistency loss—in that both the regression and its expectation depend on the optimization variable u . This makes (12) particularly suitable for learning SB-based diffusion samplers, unlike previous matching-based SB methods (Shi et al., 2023; Liu et al., 2024), which all require ground-truth target samples from $X_1 \sim \nu$.

Computing the AM objective in (12) requires knowing $\nabla \log \hat{\varphi}_1(x)$, which, as we discussed in (11), serves as a *corrector* that debiases the optimization toward the desired target. Notably, this corrector function $\nabla \log \hat{\varphi}_1(x)$ also admits a variational form (Peluchetti, 2022, 2023; Shi et al., 2023):³

$$\nabla \log \hat{\varphi}_1 = \arg \min_h \mathbb{E}_{p_{0,1}^{u^*}} [\|h(X_1) - \nabla_{x_1} \log p^{\text{base}}(X_1|X_0)\|^2]. \quad (13)$$

To summarize, Equations (12) and (13) characterize two distinct matching objectives that any kinetic-optimal drift u_t^* of SBs must satisfy. When the source distribution degenerates to Dirac delta $X_0 := 0$, (13) is minimized at $\nabla \log p_1^{\text{base}}$, and (12) simply recovers the objective used in Adjoint Sampling (Havens et al., 2025). In other words, (12) and (13) should be understood as a generalization of Adjoint Sampling to handle arbitrary—including *non-memoryless*—source distributions.

3.3 Alternating Optimization with Adjoint and Corrector Matching

Building upon the theoretical characterization in Section 3.2, we aim to design a learning algorithm that finds a diffusion sampler satisfying (12) and (13), which correspond to two simple matching-based objectives. However, these matching objectives cannot be naively implemented due to their interdependency: Solving (12) for the kinetic-optimal drift u^* requires knowing $\nabla \log \hat{\varphi}_1$. Likewise, solving (13) for the corrector function $\nabla \log \hat{\varphi}_1$ requires samples from u^* . We relax the interdependency with an alternating optimization scheme. Specifically, given an approximation of $\nabla \log \hat{\varphi}_1 \approx h^{(k-1)}$ from the previous stage $k - 1$, we first update the drift $u^{(k)}$ with the AM objective:

³Formally, $\nabla \log \hat{\varphi}_t(x)$ is the kinetic-optimal drift along the reversed time coordinate $s := 1 - t$, and (13) is its variational formulation, *i.e.*, the Markovian projection at $s = 0$; see Appendix A.2 for details.

Algorithm 1 Adjoint Schrödinger Bridge Sampler (ASBS)

Require: Sample-able source $X_0 \sim \mu$, differentiable energy $E(x)$, parametrized $u_\theta(t, x)$ and $h_\phi(x)$

- 1: Initialize $h_\phi^{(0)} := 0$
 - 2: **for** stage k **in** $1, 2, \dots$ **do**
 - 3: Update drift $u_\theta^{(k)}$ by solving (14) ▷ adjoint matching
 - 4: Update corrector $h_\phi^{(k)}$ by solving (15) ▷ corrector matching
 - 5: **end for**
-

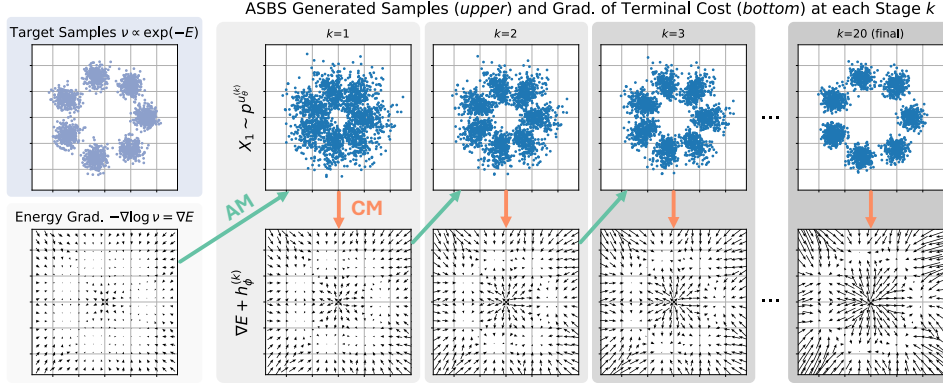


Figure 2: Illustration of ASBS on a 2D example. By alternatively minimizing the Adjoint Matching (AM) objective (14) and the Corrector Matching (CM) objective (15), ASBS progressively learns a better corrector $h_\phi^{(k)}$ that debiases the SOC problem for the control $u_\theta^{(k)}$. Note that since the corrector is initialized with $h_\phi^{(0)} := 0$, the first AM stage simply regresses $u_\theta^{(1)}$ to the energy gradient ∇E .

$$u^{(k)} := \arg \min_u \mathbb{E}_{p_{t|0,1}^{\text{base}} p_{0,1}^{\bar{u}}} \left[\|u_t(X_t) + \sigma_t(\nabla E + h^{(k-1)})(X_1)\|^2 \right], \quad \bar{u} = \text{stopgrad}(u). \quad (14)$$

Then, we use the resulting drift $u^{(k)}$ to update $h^{(k)}$ by minimizing the following matching objective, which—in light of the corrector role of $\nabla \log \hat{\varphi}_1$ —we refer to as the *Corrector Matching* objective:

$$h^{(k)} := \arg \min_h \mathbb{E}_{p_{0,1}^{u^{(k)}}} \left[\|h(X_1) - \nabla_{x_1} \log p^{\text{base}}(X_1|X_0)\|^2 \right]. \quad (15)$$

Equation (15) should be distinguished from the bridge-matching objectives in data-driven SB methods (Shi et al., 2023; Somnath et al., 2023), where X_1 must be drawn from the target distribution ν . In contrast, the matching objectives in (14) and (15) depend only on model samples at the current stage $X_1 \sim p^{u^{(k)}}(X_1|X_0)$, hence can be used to learn SB-based diffusion samplers at scale.

The alternating optimization between (14) and (15) creates a sequence of updates, $(u^{(0)}, h^{(0)}) \rightarrow \dots (u^{(k)}, h^{(k)}) \rightarrow \dots$, that may be thought of as running coordinate descent between the control u and the corrector h . Intuitively, at each stage k , we first find the control $u^{(k)}$ that best aligns with the corrector from previous stage, $h^{(k-1)}$, then update the corrector $h^{(k)}$ accordingly to reflect the “memorylessness” of the current control $u^{(k)}$. We summarize our method, **Adjoint Schrödinger Bridge Sampler (ASBS)**, in Algorithm 1, while leaving the full details with additional components, such as replay buffers, in Appendix C. Finally, we prove that this alternating optimization indeed converges to the kinetic-optimal drift u^* in (8).

Theorem 3.2 (Global convergence of ASBS). *Algorithm 1 converges to the Schrödinger bridge solution of (3), provided all matching stages achieve their critical points, i.e.,*

$$\lim_{k \rightarrow \infty} u^{(k)} = u^*.$$

4 Theoretical Analysis

We provide the proof of Theorem 3.2 and highlight theoretical insights throughout. While ASBS is specialized to a degenerate base drift $f_t := 0$, all theoretical results here apply to general f_t . To simplify notation, we omit the parameters θ, ϕ and reparametrize the corrector by $h^{(k)} = \nabla \log \bar{h}^{(k)}$. All proofs are left in Appendix B.

Our first result presents a variational characteristic to the solution of the AM objective in (14).

Theorem 4.1 (Adjoint Matching solves a forward half bridge). *Let $p^{u^{(k)}}$ be the path distribution induced by the drift $u^{(k)}$ in (14) at stage k . Then, $p^{u^{(k)}}$ solves the following variational problem:*

$$p^{u^{(k)}} = \arg \min_p \{D_{\text{KL}}(p || q^{\bar{h}^{(k-1)}}) : p_0 = \mu\}, \quad (16)$$

where $q^{\bar{h}^{(k-1)}}$ is the path distribution induced by a “backward” SDE on the reversed time coordinate $s := 1 - t$, defined by the corrector from the previous stage $\bar{h}^{(k-1)}$:

$$dY_s = [-f_s(Y_s) + \sigma_s^2 \nabla \log \phi_s(Y_s)] ds + \sigma_s dW_s, \quad \phi_s(y) = \int p_{1-s|0}^{\text{base}}(y|z) \phi_1(z) dz, \quad (17)$$

with the boundary conditions $Y_0 \sim \nu$ and $\phi_0(y) = \bar{h}^{(k-1)}(y)$.

Theorem 4.1 suggests that any SOC problems with the terminal cost $g(x) := \log \frac{\bar{h}^{(k)}(x)}{\nu(x)}$ can be reinterpreted as KL minimization w.r.t. a specific *backward* SDE (17) that is fully characterized by ν —which serves as its source distribution—and $\bar{h}^{(k)}$ —which defines its drift through the function $\phi_s(y)$. The objective in (16) differs from the one in the original SB problem (3) by disregarding the target boundary constraint, $X_1 \sim \nu$. Consequently, (16) only solves a forward half bridge.

Next, we show that the CM objective (15) admits a similar variational form, except backward in time.

Theorem 4.2 (Corrector Matching solves a backward half bridge). *Let $\bar{h}^{(k)}$ be the corrector in (15) at stage k . Then, the path distribution $q^{\bar{h}^{(k)}}$ solves the following variational problem:*

$$q^{\bar{h}^{(k)}} = \arg \min_q \{D_{\text{KL}}(p^{u^{(k)}} || q) : q_1 = \nu\} \quad (18)$$

Unlike (16), the objective in (18) disregards the source boundary constraint μ instead, thereby solving a backward half bridge. Theorems 4.1 and 4.2 imply that our ASBS in Algorithm 1 *implicitly* employs an optimization scheme that alternates between solving forward and backward half bridges, thereby instantiating the celebrated Iterative Proportional Fitting algorithm (IPF; Fortet, 1940; Kullback, 1968). Combining with the analysis by (De Bortoli et al., 2021) leads to our final result in Theorem 3.2.

5 Related Works

We provide additional clarification on SB-related works and leave the full review to Appendix A.3.

Data-driven Schrödinger Bridges The SB problem has attracted notable interests in machine learning due to its connection to diffusion-based generative models (Wang et al., 2021). Earlier methods implemented classical IPF algorithms (De Bortoli et al., 2021; Vargas et al., 2021; Chen et al., 2022), with scalability later enhanced by bridge matching-based methods (Shi et al., 2023; Liu et al., 2024). Unlike ASBS, all of them focus on generative modeling and assume access to extensive target samples during training, making them unsuitable for sampling from Boltzmann distributions.

SB-inspired Diffusion Samplers Notably, in the context of diffusion samplers, the SB formulation has been constantly emphasized as a mathematically appealing framework for both theoretical analysis and method motivation (Zhang and Chen, 2022; Vargas et al., 2024; Richter and Berner, 2024; Havens et al., 2025). None of the prior methods, however, offers general solutions to learning SB-based diffusion samplers, instead specializing to either the memoryless condition or non-matching-based objectives, which largely complicate the learning process (see Table 1). Conceptually, our ASBS stands closest to SSB (Bernton et al., 2019) by learning general SB samplers. However, the two methods differ fundamentally in scalability: SSB is a Sequential Monte Carlo-based method (Chopin, 2002) augmented with learned transition kernels using Gaussian-approximated SB potentials. As with many MCMC-augmented samplers (Gabriel et al., 2022; Matthews et al., 2022), SSB requires extensive evaluations on the energy $E(x)$, in contrast to ASBS, which is much more energy-efficient.

Table 2: Results on the synthetic energy functions for n -particle bodies with their corresponding dimensions d . Following (Chen et al., 2025; Havens et al., 2025), we report Sinkhorn for MW-5 and the Wasserstein-2 distances w.r.t samples, $\mathcal{W}_2$, and energies, $E(\cdot)\mathcal{W}_2$, for the rest. All values are averaged over three random trials. **Best results are highlighted.**

Method	MW-5 ($d=5$)	DW-4 ($d=8$)		LJ-13 ($d=39$)		LJ-55 ($d=165$)	
	Sinkhorn $\downarrow$	$\mathcal{W}_2 \downarrow$	$E(\cdot)\mathcal{W}_2 \downarrow$	$\mathcal{W}_2 \downarrow$	$E(\cdot)\mathcal{W}_2 \downarrow$	$\mathcal{W}_2 \downarrow$	$E(\cdot)\mathcal{W}_2 \downarrow$
PDDS (Phillips et al., 2024)	—	0.92 $\pm$ 0.08	0.58 $\pm$ 0.25	4.66 $\pm$ 0.87	56.01 $\pm$ 10.80	—	—
SCLD (Chen et al., 2025)	0.44 $\pm$ 0.06	1.30 $\pm$ 0.64	0.40 $\pm$ 0.19	2.93 $\pm$ 0.19	27.98 $\pm$ 1.26	—	—
PIS (Zhang and Chen, 2022)	0.65 $\pm$ 0.25	0.68 $\pm$ 0.28	0.65 $\pm$ 0.25	1.93 $\pm$ 0.07	18.02 $\pm$ 1.12	4.79 $\pm$ 0.45	228.70 $\pm$ 131.27
DDS (Vargas et al., 2023)	0.63 $\pm$ 0.24	0.92 $\pm$ 0.11	0.90 $\pm$ 0.37	1.99 $\pm$ 0.13	24.61 $\pm$ 8.99	4.60 $\pm$ 0.09	173.09 $\pm$ 18.01
LV-PIS (Richter and Berner, 2024)	—	1.04 $\pm$ 0.29	1.89 $\pm$ 0.89	—	—	—	—
iDEM (Akhound-Sadegh et al., 2024)	—	0.70 $\pm$ 0.06	0.55 $\pm$ 0.14	1.61 $\pm$ 0.01	30.78 $\pm$ 24.46	4.69 $\pm$ 1.52	93.53 $\pm$ 16.31
AS (Havens et al., 2025)	0.32 $\pm$ 0.06	0.62 $\pm$ 0.06	0.55 $\pm$ 0.12	1.67 $\pm$ 0.01	2.40 $\pm$ 1.25	4.04 $\pm$ 0.05	30.83 $\pm$ 8.19
ASBS (Ours)	0.15 $\pm$ 0.02	0.43 $\pm$ 0.05	0.20 $\pm$ 0.11	1.59 $\pm$ 0.03	1.99 $\pm$ 1.01	4.00 $\pm$ 0.03	28.10 $\pm$ 8.15

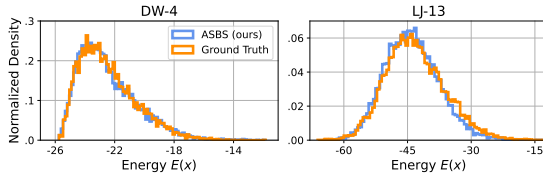


Figure 3: The energy histograms of DW-4 and LJ-13 from Table 2. ASBS generates samples whose energy profiles closely match those of the ground-truth samples.

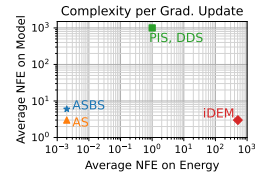


Figure 4: Complexity w.r.t. the number of function evaluation (NFE) on LJ-13 potential.

6 Experiments

Benchmarks We evaluate our ASBS on three classes of multi-particle energy functions $E(x)$.

- *Synthetic energy functions* These are classical potentials based on pair-wise distances of an n -particle system, where $E(x)$ is known analytically. Following (Akhound-Sadegh et al., 2024; Chen et al., 2025), we consider a 2D 4-particle Double-Well potential (DW-4), a 1D 5-particle Many-Well potential (MW-5), a 3D 13-particle Lennard-Jones potential (LJ-13) and a 3D 55-particle Lennard-Jones potential (LJ-55). For the ground-truth samples, we sample analytically from MW-5 and use the MCMC samples from (Klein et al., 2023) for the rest of three potentials.
- *Alanine dipeptide* This is a molecule consisting of 22 atoms in 3D. Specifically, we consider the alanine dipeptide in an implicit solvent and aim to sample from its Boltzmann distribution at a temperature 300K. Following prior methods (Zhang and Chen, 2022; Wu et al., 2020), we use the energy function $E(x)$ from the OpenMM library (Eastman et al., 2017) and consider a more structural internal coordinate with the dimension $d = 60$. The ground-truth samples contain 10^7 configurations, simulated from Molecular Dynamics (Midgley et al., 2023).
- *Amortized conformer generation* Finally, we consider a new benchmark proposed in (Havens et al., 2025) for large-scale conformer generation. Conformers are locally stable configurations located at the local minima of the molecule’s potential energy surface (Hawkins, 2017). Sampling conformers is essentially a conditional generation task, targeting a Boltzmann distribution $\nu(x|g) \propto e^{-\frac{1}{\tau} E(x|g)}$ at a low temperature $\tau \ll 1$, conditioned on the molecular topology $g \in \mathcal{G}$. The training set $\mathcal{G}_{\text{train}}$ contains 24,477 molecular topologies from SPICE (Eastman et al., 2023), represented by the SMILES strings (Weininger, 1988), whereas the test set $\mathcal{G}_{\text{test}}$ contains 80 topologies from SPICE and another 80 from GEOM-DRUGS (Axelrod and Gomez-Bombarelli, 2022). As with (Havens et al., 2025), we consider $E(x|g)$ a foundation model *eSEN* from (Fu et al., 2025), which predicts energy with density-functional-theory accuracy at a much lower computational cost. We use CREST conformers (Pracht et al., 2024) as the ground-truth samples.

Baselines and evaluation We compare ASBS with a wide range of diffusion samplers, including PIS (Zhang and Chen, 2022), DDS (Vargas et al., 2023), PDDS (Phillips et al., 2024), SCLD (Chen

Table 3: Comparison between diffusion samplers on sampling the molecular Boltzmann distribution of the alanine dipeptide. We report the KL divergence D_{KL} for the 1D marginal across five torsion angles and the Wasserstein-2 $\mathcal{W}_2$ on jointly (ϕ, ψ) , known as Ramachandran plots (see Figure 5). Best results are highlighted.

Method	D_{KL} on each torsion’s marginal ↓					$\mathcal{W}_2$ on joint ↓
	ϕ	ψ	γ_1	γ_2	γ_3	(ϕ, ψ)
PIS (Zhang and Chen, 2022)	0.05±0.03	0.38±0.49	5.61±1.24	4.49±0.03	4.60±0.03	1.27±1.19
DDS (Vargas et al., 2023)	0.03±0.01	0.16±0.07	2.44±0.96	0.03±0.00	0.03±0.00	0.68±0.09
AS (Havens et al., 2025)	0.09±0.09	0.04±0.04	0.17±0.17	0.56±0.09	0.51±0.06	0.65±0.52
ASBS (Ours)	0.02±0.00	0.01±0.00	0.03±0.01	0.02±0.00	0.02±0.00	0.25±0.01

Table 4: Results on large-scale amortized conformer generation, evaluated on two test sets, SPICE and GEOM-DRUGS, both with and without post-processing relaxation. We report the coverage (%) and Absolute Mean RMSD (AMR) of the recall at the threshold $1.0\text{\AA}$. Note that “+RDKit warmup” refers to warm-starting the model u_θ using RDKit conformers; see Appendix D for details. Best results without and with RDKit warm-up are highlighted separately.

Method	without relaxation				with relaxation			
	SPICE		GEOM-DRUGS		SPICE		GEOM-DRUGS	
	Coverage ↑	AMR ↓	Coverage ↑	AMR ↓	Coverage ↑	AMR ↓	Coverage ↑	AMR ↓
RDKit ETKDG (Riniker and Landrum, 2015)	56.94±35.82	1.04±0.52	50.81±34.69	1.15±0.61	70.21±31.70	0.79±0.44	62.55±31.67	0.93±0.53
AS (Havens et al., 2025)	56.75±38.15	0.96±0.26	36.23±33.42	1.20±0.43	82.41±25.85	0.68±0.28	64.26±34.57	0.89±0.45
ASBS w/ Gaussian prior (Ours)	73.04±31.95	0.83±0.24	50.23±35.98	1.05±0.43	88.26±20.57	0.60±0.24	72.32±29.68	0.77±0.35
ASBS w/ harmonic prior (Ours)	74.05±31.61	0.82±0.23	53.14±35.69	1.03±0.42	88.71±18.63	0.59±0.24	72.77±29.94	0.78±0.35
AS +RDKit warmup (Havens et al., 2025)	72.21±30.22	0.84±0.24	52.19±35.20	1.02±0.34	87.84±19.20	0.60±0.23	73.88±28.63	0.76±0.34
ASBS +RDKit warmup (Ours)	77.84±28.37	0.79±0.23	57.19±35.14	0.98±0.40	88.08±18.84	0.58±0.24	73.18±30.09	0.76±0.37

et al., 2025), LV (Richter and Berner, 2024), iDEM (Akhound-Sadegh et al., 2024) and finally Adjoint Sampling (AS; Havens et al., 2025). For the conformer generation task, we include additionally a domain-specific baseline, RDKit ETKDG (Riniker and Landrum, 2015), which relies on chemistry-based heuristics. The evaluation pipelines are consistent with prior methods, where we adopt the SCLD setup for MW-5, the PIS setup for alanine dipeptide, and the AS setup for all the rest; see Appendix D for details.

ASBS models For all tasks, we consider a degenerate base drift $f_t := 0$, as discussed in Section 3.2, and set σ_t a geometric noise schedule. For energy functions that directly take particle systems as inputs—such as DW, LJ, and eSEN—we parametrize the models u_θ, h_ϕ with two Equivariant Graph Neural Networks (Satorras et al., 2021) and consider a domain-specific source distribution—the harmonic prior (Jing et al., 2023). Formally, for an n -particle system $x = \{x_i\}_{i=0}^n$, the harmonic prior $\mu_{\text{harmonic}}(x)$ is a quadratic potential that can be sampled analytically from an anisotropic Gaussian:

$$\mu_{\text{harmonic}}(x) \propto \exp\left(-\frac{\alpha}{2} \sum_{i,j} \|x_i - x_j\|^2\right). \quad (19)$$

For other energy functions, we use standard fully-connected neural networks and consider Gaussian priors. All models are trained with Adam (Kingma and Ba, 2015) and, following standard practices (Havens et al., 2025; Akhound-Sadegh et al., 2024), utilize replay buffers; see Appendix C for details.

Results Table 2 presents the results on synthetic energy functions. Notably, ASBS consistently outperforms prior diffusion samplers across *all* energy functions. In Figure 3, we compare the energy histograms of DW-4 and LJ-13 potentials between the ground-truth MCMC samples and those from ASBS. It is evident that ASBS generates samples that closely resemble the target Boltzmann distribution $\nu(x) \propto e^{-E(x)}$, resulting in energy profiles $E(x)$ that are almost indistinguishable from the ground truth. Computationally, Figure 4 shows the average number of evaluation required on the energy $E(x)$ and the model $u_\theta(t, x)$ for each gradient update. ASBS is much more efficient than most diffusion samplers, with a slight overhead compared to AS due to the additional network $h_\phi(x)$.

Table 3 summarizes the results for alanine dipeptide. Following standard pipeline (Zhang and Chen, 2022), we generate model samples $X_1 \in \mathbb{R}^{60}$ and extract five torsion angles—including the backbone

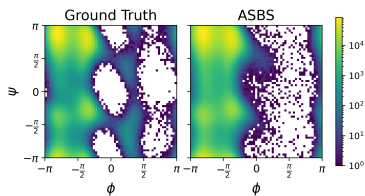


Figure 5: Ramachandran plots for the alanine dipeptide between ground-truth and ASBS samples.

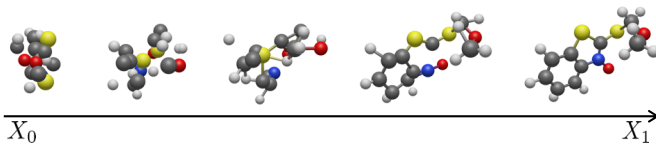


Figure 6: Example of ASBS generative process on amortized conformer generation. Given an unseen molecular topology $g \in \mathcal{G}_{\text{test}}$ from the test set—COCSc1sc2ccccc2[n+]1[0-] in this case—ASBS transports samples from the harmonic prior $X_0 \sim \mu_{\text{harmonic}}$ to generate conformers X_1 .

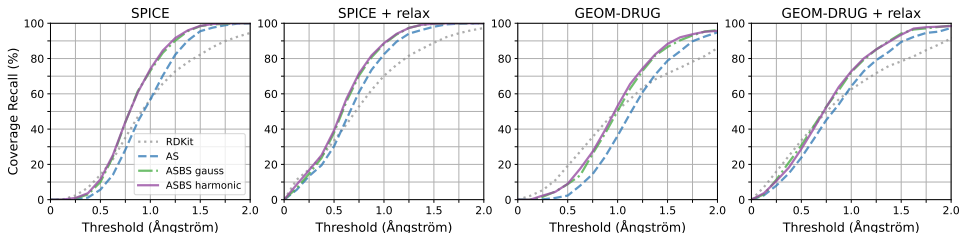


Figure 7: Recall coverage curves on amortized conformer generation on the SPICE and GEOM-DRUGS test sets without RDKit warm-start. Note that Table 4 reports the recall coverages at the threshold $1.0\text{\AA}$.

angles ϕ , ψ and methyl rotation angles γ_1 , γ_2 , γ_3 —all of them exhibit multi-modal distributions. Notably, ASBS achieves lowest KL divergence to the ground-truth marginals across all five torsions. Figure 5 further compares the joint distributions of (ϕ, ψ) , known as the Ramachandran plots (Spencer et al., 2019), between ground-truth and ASBS. While ASBS identifies all high-density modes in the region $\phi \in [-\pi, 0]$, it misses few low-density modes. This mode-seeking behavior, inherit in all SOC-based diffusion samplers, could be improved with important weighting. We provide further discussions in Appendix D.4.

Table 4 presents the recall for amortized conformer generation compared to ground-truth samples. For prior diffusion samplers, we primarily compare to AS (Havens et al., 2025) due to the benchmark’s scale. Following AS, we ablate a warm-start stage using RDKit conformers, which are close but not identical to ground-truth samples, and include results with relaxation for post-generation optimization. Since AS is a specific instance of ASBS with a Dirac delta prior—as discussed in Section 3.2—any performance improvements from AS to ASBS highlight the added capability to handle arbitrary priors and, consequently, non-memoryless processes. Remarkably, without any warm-start, ASBS with the harmonic prior (19) already matches and, in many cases, surpasses the RDKit-warm-up AS. With warm-start, ASBS achieves best performance across most metrics. This highlights the significance of domain-specific priors, aiding exploration as effectively as warm-start with additional data, which may not always be available. Finally, we visualize the generation process of ASBS with harmonic prior (19) in Figure 6 and report the recall curves in Figure 7. In practice, we observe that ASBS achieves slightly better results with a harmonic prior compared to a Gaussian prior, with both significantly outperforming AS (Havens et al., 2025). See Appendix D.4 for further ablation studies.

7 Conclusion and Limitation

We introduced **Adjoint Schrödinger Bridge Sampler (ASBS)**, a new diffusion sampler for Boltzmann distributions that solves general SB problems given only target energy functions. ASBS is based on a scalable matching framework, converges theoretically to the global solution, and performs superiorly across various benchmarks. Despite these encouraging results, further enhancement with importance sampling techniques is worth investigating to mitigate the mode collapse inherent in SOC-inspired diffusion samplers. Exploring its effectiveness in sampling amortized Boltzmann distributions would also be valuable.

Acknowledgements

The authors would like to thank Aaron Havens, Juno Nam, Xiang Fu, Bing Yan, Brandon Amos, and Brian Karrer for the helpful discussions and comments. JC and YC acknowledge support from NSF Grants ECCS-1942523, DMS-2206576, and CMMI-2450378.

References

- Tara Akhound-Sadegh, Jarrod Rector-Brooks, Avishek Joey Bose, Sarthak Mittal, Pablo Lemos, Cheng-Hao Liu, Marcin Sendera, Siamak Ravanbakhsh, Gauthier Gidel, Yoshua Bengio, Nikolay Malkin, and Alexander Tong. Iterated denoising energy matching for sampling from boltzmann densities. In *International Conference on Machine Learning (ICML)*, 2024.
- Michael S Albergo and Eric Vanden-Eijnden. Nets: A non-equilibrium transport sampler. *arXiv preprint arXiv:2410.02711*, 2024.
- Michael Arbel, Alex Matthews, and Arnaud Doucet. Annealed flow transport monte carlo. In *International Conference on Machine Learning (ICML)*, 2021.
- Simon Axelrod and Rafael Gomez-Bombarelli. GEOM, energy-annotated molecular conformations for property prediction and molecular generation. *Scientific Data*, 9(1):185, 2022.
- Richard Bellman. The theory of dynamic programming. Technical report, Rand corp santa monica ca, 1954.
- Julius Berner, Lorenz Richter, and Karen Ullrich. An optimal control perspective on diffusion-based generative modeling. *Transactions on Machine Learning Research (TMLR)*, 2024.
- Espen Bernton, Jeremy Heng, Arnaud Doucet, and Pierre E Jacob. Schrödinger bridge samplers. *arXiv preprint arXiv:1912.13170*, 2019.
- Kurt Binder, Dieter W Heermann, and K Binder. *Monte Carlo simulation in statistical physics*, volume 8. Springer, 1992.
- Denis Blessing, Xiaogang Jia, Johannes Esslinger, Francisco Vargas, and Gerhard Neumann. Beyond elbos: a large-scale evaluation of variational methods for sampling. In *Proceedings of the 41st International Conference on Machine Learning*, pages 4205–4229, 2024.
- George EP Box and George C Tiao. *Bayesian inference in statistical analysis*. John Wiley & Sons, 2011.
- James Bradbury, Roy Frostig, Peter Hawkins, Matthew James Johnson, Chris Leary, Dougal Maclaurin, George Necula, Adam Paszke, Jake VanderPlas, Skye Wanderman-Milne, and Qiao Zhang. JAX: composable transformations of Python+NumPy programs, 2018. <http://github.com/google/jax>.
- Junhua Chen, Lorenz Richter, Julius Berner, Denis Blessing, Gerhard Neumann, and Anima Anandkumar. Sequential controlled langevin diffusions. In *International Conference on Learning Representations (ICLR)*, 2025.
- Ricky T. Q. Chen, Yulia Rubanova, Jesse Bettencourt, and David K Duvenaud. Neural ordinary differential equations. In *Advances in Neural Information Processing Systems (NeurIPS)*, 2018.
- Tianrong Chen, Guan-Horng Liu, and Evangelos A Theodorou. Likelihood training of Schrödinger bridge using forward-backward SDEs theory. In *International Conference on Learning Representations (ICLR)*, 2022.
- Yongxin Chen and Tryphon Georgiou. Stochastic bridges of linear systems. *IEEE Transactions on Automatic Control*, 61(2):526–531, 2015.
- 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, 2016.

- 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, 2021.
- Nicolas Chopin. A sequential particle filter method for static models. *Biometrika*, 89(3):539–552, 2002.
- Valentin De Bortoli, James Thornton, Jeremy Heng, and Arnaud Doucet. Diffusion Schrödinger bridge with applications to score-based generative modeling. In *Advances in Neural Information Processing Systems (NeurIPS)*, 2021.
- Valentin De Bortoli, Michael Hutchinson, Peter Wirmsberger, and Arnaud Doucet. Target score matching. *arXiv preprint arXiv:2402.08667*, 2024.
- Pierre Del Moral, Arnaud Doucet, and Ajay Jasra. Sequential monte carlo samplers. *Journal of the Royal Statistical Society Series B: Statistical Methodology*, 68(3):411–436, 2006.
- Carles Domingo-Enrich, Michal Drozdal, Brian Karrer, and Ricky T. Q. Chen. Adjoint Matching: Fine-tuning flow and diffusion generative models with memoryless stochastic optimal control. In *International Conference on Learning Representations (ICLR)*, 2025.
- Peter Eastman, Jason Swails, John D Chodera, Robert T McGibbon, Yutong Zhao, Kyle A Beauchamp, Lee-Ping Wang, Andrew C Simmonett, Matthew P Harrigan, Chaya D Stern, Rafal P. Wiewiora, Bernard R. Brooks, and Vijay S. Pande. OpenMM 7: Rapid development of high performance algorithms for molecular dynamics. *PLoS computational biology*, 13(7):e1005659, 2017.
- Peter Eastman, Pavan Kumar Behara, David L Dotson, Raimondas Galvelis, John E Herr, Josh T Horton, Yuezhi Mao, John D Chodera, Benjamin P Pritchard, Yuanqing Wang, Gianni De Fabritiis, and Thomas E. Markland. SPICE, a dataset of drug-like molecules and peptides for training machine learning potentials. *Scientific Data*, 10(1):11, 2023.
- Chris Finlay, Jörn-Henrik Jacobsen, Levon Nurbekyan, and Adam Oberman. How to train your neural ODE: The world of jacobian and kinetic regularization. In *International Conference on Machine Learning (ICML)*, 2020.
- Robert Fortet. Résolution d’un système d’équations de M. Schrödinger. *Journal de Mathématiques Pures et Appliquées*, 19(1-4):83–105, 1940.
- Xiang Fu, Brandon M Wood, Luis Barroso-Luque, Daniel S Levine, Meng Gao, Misko Dzamba, and C Lawrence Zitnick. Learning smooth and expressive interatomic potentials for physical property prediction. In *International Conference on Machine Learning (ICML)*, 2025.
- Marylou Gabrié, Grant M Rotskoff, and Eric Vanden-Eijnden. Adaptive monte carlo augmented with normalizing flows. *Proceedings of the National Academy of Sciences*, 119(10):e2109420119, 2022.
- WK HASTINGS. Monte carlo sampling methods using markov chains and their applications. *Biometrika*, 57(1):97–109, 1970.
- Aaron Havens, Benjamin Kurt Miller, Bing Yan, Carles Domingo-Enrich, Anuroop Sriram, Brandon Wood, Daniel Levine, Bin Hu, Brandon Amos, Brian Karrer, Xiang Fu, Guan-Hong Liu, and Ricky T. Q. Chen. Adjoint Sampling: Highly scalable diffusion samplers via Adjoint Matching. In *International Conference on Machine Learning (ICML)*, 2025.
- Paul CD Hawkins. Conformation generation: The state of the art. *Journal of chemical information and modeling*, 57(8):1747–1756, 2017.
- Jonathan Ho, Ajay Jain, and Pieter Abbeel. Denoising diffusion probabilistic models. In *Advances in Neural Information Processing Systems (NeurIPS)*, 2020.
- Kiyosi Itô. *On stochastic differential equations*, volume 4. American Mathematical Soc., 1951.
- Bowen Jing, Ezra Erives, Peter Pao-Huang, Gabriele Corso, Bonnie Berger, and Tommi S Jaakkola. EigenFold: Generative protein structure prediction with diffusion models. In *International Conference on Learning Representations (ICLR), Workshop Track*, 2023.

- Hilbert J Kappen. Path integrals and symmetry breaking for optimal control theory. *Journal of Statistical Mechanics: Theory and Experiment*, 2005(11):P11011, 2005.
- Tero Karras, Miika Aittala, Timo Aila, and Samuli Laine. Elucidating the design space of diffusion-based generative models. In *Advances in Neural Information Processing Systems (NeurIPS)*, 2022.
- Diederik P Kingma and Jimmy Ba. Adam: A method for stochastic optimization. In *International Conference on Learning Representations (ICLR)*, 2015.
- Leon Klein, Andrew Foong, Tor Fjelde, Bruno Mlodozieniec, Marc Brockschmidt, Sebastian Nowozin, Frank Noé, and Ryota Tomioka. Timewarp: Transferable acceleration of molecular dynamics by learning time-coarsened dynamics. In *Advances in Neural Information Processing Systems (NeurIPS)*, 2023.
- Jonas Köhler, Leon Klein, and Frank Noé. Equivariant flows: exact likelihood generative learning for symmetric densities. In *International conference on machine learning*, pages 5361–5370. PMLR, 2020.
- Solomon Kullback. Probability densities with given marginals. *The Annals of Mathematical Statistics*, 39(4):1236–1243, 1968.
- Greg Landrum. Rdkit: Open-source cheminformatics. <https://www.rdkit.org>, 2006.
- Jean-François Le Gall. *Brownian motion, martingales, and stochastic calculus*. Springer, 2016.
- Christian Léonard. From the schrödinger problem to the monge–kantorovich problem. *Journal of Functional Analysis*, 262(4):1879–1920, 2012.
- Christian Léonard. A survey of the Schrödinger problem and some of its connections with optimal transport. *Discrete and Continuous Dynamical Systems*, 2013.
- Christian Léonard, Sylvie Roelly, and Jean-Claude Zambrini. Reciprocal processes. A measure-theoretical point of view. *Probability Surveys*, 2014.
- Daniel S Levine, Muhammed Shuaibi, Evan Walter Clark Spotte-Smith, Michael G Taylor, Muhammad R Hasyim, Kyle Michel, Ilyes Batatia, Gábor Csányi, Misko Dzamba, Peter Eastman, et al. The open molecules 2025 (omol25) dataset, evaluations, and models. *arXiv preprint arXiv:2505.08762*, 2025.
- Guan-Horng Liu, Arash Vahdat, De-An Huang, Evangelos A Theodorou, Weili Nie, and Anima Anandkumar. I²SB: Image-to-Image Schrödinger bridge. In *International Conference on Machine Learning (ICML)*, 2023.
- Guan-Horng Liu, Yaron Lipman, Maximilian Nickel, Brian Karrer, Evangelos A Theodorou, and Ricky T. Q. Chen. Generalized Schrödinger bridge matching. In *International Conference on Learning Representations (ICLR)*, 2024.
- Alex Matthews, Michael Arbel, Danilo Jimenez Rezende, and Arnaud Doucet. Continual repeated annealed flow transport monte carlo. In *International Conference on Machine Learning (ICML)*, 2022.
- Nicholas Metropolis, Arianna W Rosenbluth, Marshall N Rosenbluth, Augusta H Teller, and Edward Teller. Equation of state calculations by fast computing machines. *The journal of chemical physics*, 21(6):1087–1092, 1953.
- Laurence Illing Midgley, Vincent Stimper, Gregor NC Simm, Bernhard Schölkopf, and José Miguel Hernández-Lobato. Flow annealed importance sampling bootstrap. In *International Conference on Learning Representations (ICLR)*, 2023.
- Radford M Neal. Annealed importance sampling. *Statistics and computing*, 11:125–139, 2001.
- F. Neese. The orca program system. *WIREs Comput. Molec. Sci.*, 2(1):73–78, 2012. doi: 10.1002/wcms.81.

- Kirill Neklyudov, Daniel Severo, and Alireza Makhzani. Action matching: A variational method for learning stochastic dynamics from samples. In *International Conference on Machine Learning (ICML)*, 2023.
- Edward Nelson. *Dynamical theories of Brownian motion*, volume 106. Princeton university press, 2020.
- Frank Noé, Simon Olsson, Jonas Köhler, and Hao Wu. Boltzmann generators: Sampling equilibrium states of many-body systems with deep learning. *Science*, 365(6457):eaaw1147, 2019.
- Bernt Øksendal. Stochastic differential equations. In *Stochastic Differential Equations*, pages 65–84. Springer, 2003.
- Adam Paszke, Sam Gross, Francisco Massa, Adam Lerer, James Bradbury, Gregory Chanan, Trevor Killeen, Zeming Lin, Natalia Gimelshein, Luca Antiga, et al. Pytorch: An imperative style, high-performance deep learning library. In *Advances in neural information processing systems*, pages 8026–8037, 2019.
- Stefano Peluchetti. Non-Denoising forward-time diffusions, 2022. <https://openreview.net/forum?id=oVfIKuhqfC>.
- Stefano Peluchetti. Diffusion bridge mixture transports, Schrödinger bridge problems and generative modeling. *arXiv preprint arXiv:2304.00917*, 2023.
- Gabriel Peyré and Marco Cuturi. Computational optimal transport. *Center for Research in Economics and Statistics Working Papers*, 2017.
- Gabriel Peyré and Marco Cuturi. Computational optimal transport: With applications to data science. *Foundations and Trends® in Machine Learning*, 11(5-6):355–607, 2019.
- Angus Phillips, Hai-Dang Dau, Michael John Hutchinson, Valentin De Bortoli, George Deligiannidis, and Arnaud Doucet. Particle denoising diffusion sampler. In *International Conference on Machine Learning (ICML)*, 2024.
- Philipp Pracht, Stefan Grimme, Christoph Bannwarth, Fabian Bohle, Sebastian Ehlert, Gereon Feldmann, Johannes Gorges, Marcel Müller, Tim Neudecker, Christoph Plett, Sebastian Spicher, Pit Steinbach, Patryk A. Wośowski, and Felix Zeller. CREST—A program for the exploration of low-energy molecular chemical space. *The Journal of Chemical Physics*, 160(11), 2024.
- Lorenz Richter and Julius Berner. Improved sampling via learned diffusions. In *International Conference on Learning Representations (ICLR)*, 2024.
- Sereina Riniker and Gregory A Landrum. Better informed distance geometry: using what we know to improve conformation generation. *Journal of chemical information and modeling*, 55(12): 2562–2574, 2015.
- Simo Särkkä and Arno Solin. *Applied stochastic differential equations*, volume 10. Cambridge University Press, 2019.
- Victor Garcia Satorras, Emiel Hoogetboom, and Max Welling. $E(n)$ equivariant graph neural networks. In *International Conference on Machine Learning (ICML)*, 2021.
- Erwin Schrödinger. *Über die Umkehrung der Naturgesetze*, volume IX. Sitzungsberichte der Preuss Akad. Wissen. Phys. Math. Klasse, Sonderausgabe, 1931.
- Erwin Schrödinger. Sur la théorie relativiste de l’électron et l’interprétation de la mécanique quantique. In *Annales de l’institut Henri Poincaré*, 1932.
- Neta Shaul, Ricky T. Q. Chen, Maximilian Nickel, Matthew Le, and Yaron Lipman. On kinetic optimal probability paths for generative models. In *International Conference on Machine Learning (ICML)*, 2023.
- Yuyang Shi, Valentin De Bortoli, Andrew Campbell, and Arnaud Doucet. Diffusion Schrödinger bridge matching. In *Advances in Neural Information Processing Systems (NeurIPS)*, 2023.

- Vignesh Ram Somnath, Matteo Pariset, Ya-Ping Hsieh, Maria Rodriguez Martinez, Andreas Krause, and Charlotte Bunne. Aligned diffusion Schrödinger bridges. In *Conference on Uncertainty in Artificial Intelligence (UAI)*, 2023.
- Yang Song, Jascha Sohl-Dickstein, Diederik P Kingma, Abhishek Kumar, Stefano Ermon, and Ben Poole. Score-based generative modeling through stochastic differential equations. In *International Conference on Learning Representations (ICLR)*, 2021.
- Ryan K Spencer, Glenn L Butterfoss, John R Edison, James R Eastwood, Stephen Whitelam, Kent Kirshenbaum, and Ronald N Zuckermann. Stereochemistry of polypeptoid chain configurations. *Biopolymers*, 110(6):e23266, 2019.
- Vincent Stimper, Bernhard Schölkopf, and José Miguel Hernández-Lobato. Resampling base distributions of normalizing flows. In *International Conference on Artificial Intelligence and Statistics (AISTATS)*, 2022.
- Emanuel Todorov. Linearly-solvable Markov decision problems. In *Advances in Neural Information Processing Systems (NeurIPS)*, 2007.
- Mark E Tuckerman. *Statistical mechanics: theory and molecular simulation*. Oxford university press, 2023.
- Francisco Vargas, Pierre Thodoroff, Neil D Lawrence, and Austen Lamacraft. Solving Schrödinger bridges via maximum likelihood. *Entropy*, 2021.
- Francisco Vargas, Will Grathwohl, and Arnaud Doucet. Denoising diffusion samplers. In *International Conference on Learning Representations (ICLR)*, 2023.
- Francisco Vargas, Shreyas Padhy, Denis Blessing, and Nikolas Nüsken. Transport meets variational inference: Controlled monte carlo diffusions. In *International Conference on Learning Representations (ICLR)*, 2024.
- Gefei Wang, Yuling Jiao, Qian Xu, Yang Wang, and Can Yang. Deep generative learning via Schrödinger bridge. In *International Conference on Machine Learning (ICML)*, 2021.
- David Weininger. Smiles, a chemical language and information system. 1. introduction to methodology and encoding rules. *Journal of chemical information and computer sciences*, 28(1):31–36, 1988.
- Hao Wu, Jonas Köhler, and Frank Noé. Stochastic normalizing flows. In *Advances in Neural Information Processing Systems (NeurIPS)*, volume 33, pages 5933–5944, 2020.
- Qinsheng Zhang and Yongxin Chen. Path integral sampler: A stochastic control approach for sampling. In *International Conference on Learning Representations (ICLR)*, 2022.

NeurIPS Paper Checklist

1. Claims

Question: Do the main claims made in the abstract and introduction accurately reflect the paper's contributions and scope?

Answer: [\[Yes\]](#)

Justification: Our theoretical and empirical results validate the itemized claims made in the end of introduction.

Guidelines:

- The answer NA means that the abstract and introduction do not include the claims made in the paper.
- The abstract and/or introduction should clearly state the claims made, including the contributions made in the paper and important assumptions and limitations. A No or NA answer to this question will not be perceived well by the reviewers.
- The claims made should match theoretical and experimental results, and reflect how much the results can be expected to generalize to other settings.
- It is fine to include aspirational goals as motivation as long as it is clear that these goals are not attained by the paper.

2. Limitations

Question: Does the paper discuss the limitations of the work performed by the authors?

Answer: [\[Yes\]](#)

Justification: Limitation is discussed in the last section, titled "Conclusion and Limitation".

Guidelines:

- The answer NA means that the paper has no limitation while the answer No means that the paper has limitations, but those are not discussed in the paper.
- The authors are encouraged to create a separate "Limitations" section in their paper.
- The paper should point out any strong assumptions and how robust the results are to violations of these assumptions (e.g., independence assumptions, noiseless settings, model well-specification, asymptotic approximations only holding locally). The authors should reflect on how these assumptions might be violated in practice and what the implications would be.
- The authors should reflect on the scope of the claims made, e.g., if the approach was only tested on a few datasets or with a few runs. In general, empirical results often depend on implicit assumptions, which should be articulated.
- The authors should reflect on the factors that influence the performance of the approach. For example, a facial recognition algorithm may perform poorly when image resolution is low or images are taken in low lighting. Or a speech-to-text system might not be used reliably to provide closed captions for online lectures because it fails to handle technical jargon.
- The authors should discuss the computational efficiency of the proposed algorithms and how they scale with dataset size.
- If applicable, the authors should discuss possible limitations of their approach to address problems of privacy and fairness.
- While the authors might fear that complete honesty about limitations might be used by reviewers as grounds for rejection, a worse outcome might be that reviewers discover limitations that aren't acknowledged in the paper. The authors should use their best judgment and recognize that individual actions in favor of transparency play an important role in developing norms that preserve the integrity of the community. Reviewers will be specifically instructed to not penalize honesty concerning limitations.

3. Theory assumptions and proofs

Question: For each theoretical result, does the paper provide the full set of assumptions and a complete (and correct) proof?

Answer: [\[Yes\]](#)

Justification: Proofs and assumptions of all theorems appearing in the main paper can be found in Appendix B.

Guidelines:

- The answer NA means that the paper does not include theoretical results.
- All the theorems, formulas, and proofs in the paper should be numbered and cross-referenced.
- All assumptions should be clearly stated or referenced in the statement of any theorems.
- The proofs can either appear in the main paper or the supplemental material, but if they appear in the supplemental material, the authors are encouraged to provide a short proof sketch to provide intuition.
- Inversely, any informal proof provided in the core of the paper should be complemented by formal proofs provided in appendix or supplemental material.
- Theorems and Lemmas that the proof relies upon should be properly referenced.

4. Experimental result reproducibility

Question: Does the paper fully disclose all the information needed to reproduce the main experimental results of the paper to the extent that it affects the main claims and/or conclusions of the paper (regardless of whether the code and data are provided or not)?

Answer: [Yes]

Justification: Necessary information to reproduce our method is discussed in Section 6, with full details in Appendices C and D.

Guidelines:

- The answer NA means that the paper does not include experiments.
- If the paper includes experiments, a No answer to this question will not be perceived well by the reviewers: Making the paper reproducible is important, regardless of whether the code and data are provided or not.
- If the contribution is a dataset and/or model, the authors should describe the steps taken to make their results reproducible or verifiable.
- Depending on the contribution, reproducibility can be accomplished in various ways. For example, if the contribution is a novel architecture, describing the architecture fully might suffice, or if the contribution is a specific model and empirical evaluation, it may be necessary to either make it possible for others to replicate the model with the same dataset, or provide access to the model. In general, releasing code and data is often one good way to accomplish this, but reproducibility can also be provided via detailed instructions for how to replicate the results, access to a hosted model (e.g., in the case of a large language model), releasing of a model checkpoint, or other means that are appropriate to the research performed.
- While NeurIPS does not require releasing code, the conference does require all submissions to provide some reasonable avenue for reproducibility, which may depend on the nature of the contribution. For example
 - (a) If the contribution is primarily a new algorithm, the paper should make it clear how to reproduce that algorithm.
 - (b) If the contribution is primarily a new model architecture, the paper should describe the architecture clearly and fully.
 - (c) If the contribution is a new model (e.g., a large language model), then there should either be a way to access this model for reproducing the results or a way to reproduce the model (e.g., with an open-source dataset or instructions for how to construct the dataset).
 - (d) We recognize that reproducibility may be tricky in some cases, in which case authors are welcome to describe the particular way they provide for reproducibility. In the case of closed-source models, it may be that access to the model is limited in some way (e.g., to registered users), but it should be possible for other researchers to have some path to reproducing or verifying the results.

5. Open access to data and code

Question: Does the paper provide open access to the data and code, with sufficient instructions to faithfully reproduce the main experimental results, as described in supplemental material?

Answer: [No]

Justification: All data used in this work is open-source. Unfortunately, due to organizational policy, we are unable to release our source code at submission time. However, we plan to make it publicly available in the near future once administrative challenges are resolved.

Guidelines:

- The answer NA means that paper does not include experiments requiring code.
- Please see the NeurIPS code and data submission guidelines (<https://nips.cc/public/guides/CodeSubmissionPolicy>) for more details.
- While we encourage the release of code and data, we understand that this might not be possible, so “No” is an acceptable answer. Papers cannot be rejected simply for not including code, unless this is central to the contribution (e.g., for a new open-source benchmark).
- The instructions should contain the exact command and environment needed to run to reproduce the results. See the NeurIPS code and data submission guidelines (<https://nips.cc/public/guides/CodeSubmissionPolicy>) for more details.
- The authors should provide instructions on data access and preparation, including how to access the raw data, preprocessed data, intermediate data, and generated data, etc.
- The authors should provide scripts to reproduce all experimental results for the new proposed method and baselines. If only a subset of experiments are reproducible, they should state which ones are omitted from the script and why.
- At submission time, to preserve anonymity, the authors should release anonymized versions (if applicable).
- Providing as much information as possible in supplemental material (appended to the paper) is recommended, but including URLs to data and code is permitted.

6. Experimental setting/details

Question: Does the paper specify all the training and test details (e.g., data splits, hyperparameters, how they were chosen, type of optimizer, etc.) necessary to understand the results?

Answer: [Yes]

Justification: All experimental setups are discussed in Section 6, with full details in Appendix D.

Guidelines:

- The answer NA means that the paper does not include experiments.
- The experimental setting should be presented in the core of the paper to a level of detail that is necessary to appreciate the results and make sense of them.
- The full details can be provided either with the code, in appendix, or as supplemental material.

7. Experiment statistical significance

Question: Does the paper report error bars suitably and correctly defined or other appropriate information about the statistical significance of the experiments?

Answer: [Yes]

Justification: All numerical values in Section 6 are averaged over a few random trials and we have reported their standard deviations.

Guidelines:

- The answer NA means that the paper does not include experiments.
- The authors should answer "Yes" if the results are accompanied by error bars, confidence intervals, or statistical significance tests, at least for the experiments that support the main claims of the paper.

- The factors of variability that the error bars are capturing should be clearly stated (for example, train/test split, initialization, random drawing of some parameter, or overall run with given experimental conditions).
- The method for calculating the error bars should be explained (closed form formula, call to a library function, bootstrap, etc.)
- The assumptions made should be given (e.g., Normally distributed errors).
- It should be clear whether the error bar is the standard deviation or the standard error of the mean.
- It is OK to report 1-sigma error bars, but one should state it. The authors should preferably report a 2-sigma error bar than state that they have a 96% CI, if the hypothesis of Normality of errors is not verified.
- For asymmetric distributions, the authors should be careful not to show in tables or figures symmetric error bars that would yield results that are out of range (e.g. negative error rates).
- If error bars are reported in tables or plots, The authors should explain in the text how they were calculated and reference the corresponding figures or tables in the text.

8. Experiments compute resources

Question: For each experiment, does the paper provide sufficient information on the computer resources (type of compute workers, memory, time of execution) needed to reproduce the experiments?

Answer: [Yes]

Justification: We provide these details in the supplementary material.

Guidelines:

- The answer NA means that the paper does not include experiments.
- The paper should indicate the type of compute workers CPU or GPU, internal cluster, or cloud provider, including relevant memory and storage.
- The paper should provide the amount of compute required for each of the individual experimental runs as well as estimate the total compute.
- The paper should disclose whether the full research project required more compute than the experiments reported in the paper (e.g., preliminary or failed experiments that didn't make it into the paper).

9. Code of ethics

Question: Does the research conducted in the paper conform, in every respect, with the NeurIPS Code of Ethics <https://neurips.cc/public/EthicsGuidelines>?

Answer: [Yes]

Justification: We read and comply with the NeurIPS Code of Ethics.

Guidelines:

- The answer NA means that the authors have not reviewed the NeurIPS Code of Ethics.
- If the authors answer No, they should explain the special circumstances that require a deviation from the Code of Ethics.
- The authors should make sure to preserve anonymity (e.g., if there is a special consideration due to laws or regulations in their jurisdiction).

10. Broader impacts

Question: Does the paper discuss both potential positive societal impacts and negative societal impacts of the work performed?

Answer: [No]

Justification: This work does not have novel societal impact beyond that of already existing diffusion samplers.

Guidelines:

- The answer NA means that there is no societal impact of the work performed.

- If the authors answer NA or No, they should explain why their work has no societal impact or why the paper does not address societal impact.
- Examples of negative societal impacts include potential malicious or unintended uses (e.g., disinformation, generating fake profiles, surveillance), fairness considerations (e.g., deployment of technologies that could make decisions that unfairly impact specific groups), privacy considerations, and security considerations.
- The conference expects that many papers will be foundational research and not tied to particular applications, let alone deployments. However, if there is a direct path to any negative applications, the authors should point it out. For example, it is legitimate to point out that an improvement in the quality of generative models could be used to generate deepfakes for disinformation. On the other hand, it is not needed to point out that a generic algorithm for optimizing neural networks could enable people to train models that generate Deepfakes faster.
- The authors should consider possible harms that could arise when the technology is being used as intended and functioning correctly, harms that could arise when the technology is being used as intended but gives incorrect results, and harms following from (intentional or unintentional) misuse of the technology.
- If there are negative societal impacts, the authors could also discuss possible mitigation strategies (e.g., gated release of models, providing defenses in addition to attacks, mechanisms for monitoring misuse, mechanisms to monitor how a system learns from feedback over time, improving the efficiency and accessibility of ML).

11. Safeguards

Question: Does the paper describe safeguards that have been put in place for responsible release of data or models that have a high risk for misuse (e.g., pretrained language models, image generators, or scraped datasets)?

Answer: [NA]

Justification: This paper poses no such risks.

Guidelines:

- The answer NA means that the paper poses no such risks.
- Released models that have a high risk for misuse or dual-use should be released with necessary safeguards to allow for controlled use of the model, for example by requiring that users adhere to usage guidelines or restrictions to access the model or implementing safety filters.
- Datasets that have been scraped from the Internet could pose safety risks. The authors should describe how they avoided releasing unsafe images.
- We recognize that providing effective safeguards is challenging, and many papers do not require this, but we encourage authors to take this into account and make a best faith effort.

12. Licenses for existing assets

Question: Are the creators or original owners of assets (e.g., code, data, models), used in the paper, properly credited and are the license and terms of use explicitly mentioned and properly respected?

Answer: [Yes]

Justification: Original papers that produced the code package or dataset are all properly credited.

Guidelines:

- The answer NA means that the paper does not use existing assets.
- The authors should cite the original paper that produced the code package or dataset.
- The authors should state which version of the asset is used and, if possible, include a URL.
- The name of the license (e.g., CC-BY 4.0) should be included for each asset.
- For scraped data from a particular source (e.g., website), the copyright and terms of service of that source should be provided.

- If assets are released, the license, copyright information, and terms of use in the package should be provided. For popular datasets, paperswithcode.com/datasets has curated licenses for some datasets. Their licensing guide can help determine the license of a dataset.
- For existing datasets that are re-packaged, both the original license and the license of the derived asset (if it has changed) should be provided.
- If this information is not available online, the authors are encouraged to reach out to the asset's creators.

13. **New assets**

Question: Are new assets introduced in the paper well documented and is the documentation provided alongside the assets?

Answer: [NA]

Justification: This paper does not release new assets.

Guidelines:

- The answer NA means that the paper does not release new assets.
- Researchers should communicate the details of the dataset/code/model as part of their submissions via structured templates. This includes details about training, license, limitations, etc.
- The paper should discuss whether and how consent was obtained from people whose asset is used.

A Additional Preliminary and Reviews

A.1 Stochastic Optimal Control (SOC)

In this subsection, we expand Section 2 with details. Recall the SOC problem in (4):

$$\min_u \mathbb{E}_{X \sim p^u} \left[\int \frac{1}{2} \|u_t(X_t)\|^2 dt + g(X_1) \right] \quad (20a)$$

$$\text{s.t. } dX_t = [f_t(X_t) + \sigma_t u_t(X_t)] dt + \sigma_t dW_t, \quad X_0 \sim \mu. \quad (20b)$$

Similar to (8), the optimal control to (20) can be characterized through an optimality equation:

$$u_t^*(x) = -\sigma_t \nabla V_t(x), \quad \text{where} \quad V_t(x) = -\log \int p_1^{\text{base}}(y|x) e^{-V_1(y)} dy, \quad V_1(x) = g(x) \quad (21)$$

is the value function known to satisfy the Hamilton–Jacobi–Bellman (HJB) equation (Bellman, 1954). We provide further characterization below.

Optimal distribution The optimization problem in (20) is known analytically. Specifically, notice that the entropy-regularized objective in (20) can be reformulated as:

$$\begin{aligned} & D_{\text{KL}}(p(X) \| p^{\text{base}}(X)) + \mathbb{E}_{p(X)} [g(X_1)] \\ &= D_{\text{KL}}(p(X_0) \| p^{\text{base}}(X_0)) + \mathbb{E}_{p(X_0)} \left[D_{\text{KL}}(p(X|X_0) \| p^{\text{base}}(X|X_0)) + \mathbb{E}_{p(X|X_0)} [g(X_1)] \right] \\ &= D_{\text{KL}}(p(X_0) \| p^{\text{base}}(X_0)) + \mathbb{E}_{p(X_0)} \left[D_{\text{KL}}(p(X|X_0) \| p^{\text{base}}(X|X_0) e^{-g(X_1)}) \right] \end{aligned} \quad (22)$$

where we shorthand $X \equiv X_{[0,1]}$ and denote p^{base} the base distribution induced by (20b) with $u := 0$, i.e., the uncontrolled distribution. Minimizing (22) w.r.t. p yields

$$p^*(X|X_0) = \frac{1}{Z(X_0)} p^{\text{base}}(X|X_0) e^{-g(X_1)}, \quad p^*(X_0) = p^{\text{base}}(X_0) \quad (23)$$

where $Z(X_0)$ is the normalization term defined by

$$Z(X_0) := \int p^{\text{base}}(X|X_0) e^{-g(X_1)} dX = \int p^{\text{base}}(X_1|X_0) e^{-g(X_1)} dX_1 \quad (24)$$

which is exactly $e^{-V(X_0)}$ due to (21). Combing (23) and (24) leads to the the optimal distribution in (5), which we restate below for completeness:

$$p^*(X) = p^{\text{base}}(X) e^{-g(X_1) + V_0(X_0)} \implies p^*(X_0, X_1) = p^{\text{base}}(X_0, X_1) e^{-g(X_1) + V_0(X_0)} \quad (25)$$

Adjoint Matching (AM) Scalable computational methods for solving (20) have been challenging, as naively back-propagating through (20) induces prohibitively high computational cost. Instead, Adjoint Matching (Domingo-Enrich et al., 2025) employs a matching-based objective, named Adjoint Matching (AM):

$$u^* = \arg \min_u \mathbb{E}_{X \sim p^a} [\|u_t(X_t) + \sigma_t a_t\|^2], \quad \bar{u} = \text{stopgrad}(u), \quad (26a)$$

$$\text{where} \quad -da_t = a_t \cdot \nabla f_t(X_t) dt, \quad a_1 = \nabla g(X_1) \quad (26b)$$

is the backward dynamics of the (lean) adjoint state $a_t \equiv a(t; X_{[t,1]})$. It has been proven that the unique critical point of (26) is the optimal control u^* , implying a new characteristics of the optimal control u^* using the adjoint state:

$$u_t^*(x) = -\sigma_t \mathbb{E}_{p^*} [a_t | X_t = x]. \quad (27)$$

Adjoint Sampling (AS) Recently, Havens et al. (2025) introduced an adaptation of AM tailored to sampling Boltzmann distribution $\nu(x) \propto e^{-E(x)}$ by considering

$$f_t := 0, \quad \mu(x) := \delta_0(x), \quad g(x) := \log \frac{p_1^{\text{base}}(x)}{\nu(x)}. \quad (28)$$

That is, AS considers the following SOC problem with a degenerate base drift, a Dirac delta prior, and a specific instantiation of the terminal cost $g(x) := \log \frac{p_1^{\text{base}}(x)}{\nu(x)}$:

$$\min_u \mathbb{E}_{X \sim p^u} \left[\int_0^1 \frac{1}{2} \|u_t(X_t)\|^2 dt + \log \frac{p_1^{\text{base}}(X_1)}{\nu(X_1)} \right] \quad \text{s.t. } dX_t = \sigma_t u_t(X_t) dt + \sigma_t dW_t, \quad X_0 = 0. \quad (29)$$

Notably, this SOC problem (29) admits a simplified adjoint state a_t and a degenerate initial value function $V_0(x)$:

$$a_t \stackrel{(26b)}{=} \nabla g(X_1) \stackrel{(28)}{=} \nabla \log p_1^{\text{base}}(X_1) + \nabla E(X_1) \quad \forall t \in [0, 1] \quad (30)$$

$$V_0(x) \stackrel{(28)}{=} -\log \int p_1^{\text{base}}(y) \frac{\nu(y)}{p_1^{\text{base}}(y)} dy = -\log 1 = 0, \quad (31)$$

which further implies that the optimal distribution p^* is a reciprocal process (Léonard et al., 2014):

$$p^*(X) \stackrel{(31)}{=} p^{\text{base}}(X) e^{-V_1(X_1)} \stackrel{(28)}{=} p^{\text{base}}(X) \frac{\nu(X_1)}{p_1^{\text{base}}(X_1)} = p^{\text{base}}(X|X_1) p^*(X_1). \quad (32)$$

Combining the adjoint characteristics of the optimal control (27) with the simplified adjoint state a_t in (30) and optimal distribution p^* (32) motivates the following *Reciprocal Adjoint Matching (RAM)* objective used in AS, where the unique critical point remains to be the optimal control u^* in (21).

$$u^* = \arg \min_u \mathbb{E}_{p_{t|1}^{\text{base}} p_1^u} [\|u_t(X_t) + \sigma_t (\nabla E + \nabla \log p_1^{\text{base}})(X_1)\|^2], \quad \bar{u} = \text{stopgrad}(u). \quad (33)$$

Remark on reciprocal representation The reciprocal representation of the optimal-controlled distribution p^* in (32) extends to general SOC problems (20) with non-trivial base drifts and source distributions. Specifically, any optimal-controlled distribution that solves (20) can be factorized by

$$p^*(X) = p^{\text{base}}(X|X_0, X_1) p^*(X_0, X_1). \quad (34)$$

We leave a formal statement in Theorem B.3 and Corollary B.4.

AS with linear base drift and Gaussian prior (Figure 1) Here, we discuss an alternative instantiation of AM for sampling with linear base drift and Gaussian prior, which reproduces the leftmost plot in Figure 1. Consider

$$f_t(x) := -\frac{1}{2} \beta_t x, \quad \mu(x) := \mathcal{N}(x; 0, I), \quad \sigma_t := \sqrt{\beta_t}, \quad g(x) := \log \frac{p_1^{\text{base}}(x)}{\nu(x)}. \quad (35)$$

where β_t is chosen such that (f_t, μ, σ_t) fulfill the memoryless condition. For instance, Figure 1 adopts the VPSDE (Song et al., 2021) setup:

$$\beta_t = (1-t)\beta_{\max} + t\beta_{\min}, \quad \beta_{\max} = 20, \quad \beta_{\min} = 0.1. \quad (36)$$

Similar to (30), the resulting SOC problem admits a simplified adjoint state a_t :

$$a_t \stackrel{(26b)}{=} \kappa_t \cdot \nabla g(X_1) \stackrel{(35)}{=} \kappa_t \cdot (\nabla \log p_1^{\text{base}}(X_1) + \nabla E(X_1)), \quad \kappa_t := e^{-\frac{1}{2} \int_t^1 \beta_\tau d\tau} \stackrel{(36)}{=} e^{-\frac{1}{4} (1-t)(\beta_t + \beta_1)} \quad (37)$$

and the RAM objective becomes

$$u^* = \arg \min_u \mathbb{E}_{p_{t|0,1}^{\text{base}} p_{0,1}^u} [\|u_t(X_t) + \sigma_t \kappa_t (\nabla E + \nabla \log p_1^{\text{base}})(X_1)\|^2], \quad \bar{u} = \text{stopgrad}(u). \quad (38)$$

Note that $p_{t|0,1}^{\text{base}}$ can be sampled analytically:

$$p_{t|0,1}^{\text{base}}(X_t|X_0, X_1) \stackrel{(35)}{=} \mathcal{N}(X_t; \frac{\bar{\kappa}_t(1-\kappa_t^2)}{1-\bar{\kappa}_1^2} X_0 + \frac{\kappa_t(1-\bar{\kappa}_t^2)}{1-\bar{\kappa}_1^2} X_1, \frac{(1-\kappa_t^2)(1-\bar{\kappa}_t^2)}{1-\bar{\kappa}_1^2} I), \quad (39)$$

where κ_t is defined in (37) and $\bar{\kappa}_t := e^{-\frac{1}{2} \int_0^t \beta_\tau d\tau} \stackrel{(36)}{=} e^{-\frac{1}{4} t(\beta_t + \beta_0)}$.

A.2 Schrödinger Bridge (SB)

In this subsection, we provide additional clarification on SB and specifically the derivation of (13). Recall the optimality equations of SB in (8):

$$u_t^*(x) = \sigma_t \nabla \log \varphi_t(x), \quad \text{where} \quad \begin{cases} \varphi_t(x) = \int p_{1|t}^{\text{base}}(y|x) \varphi_1(y) dy, & \varphi_0(x) \hat{\varphi}_0(x) = \mu(x) \quad (40a) \\ \hat{\varphi}_t(x) = \int p_{t|0}^{\text{base}}(x|y) \hat{\varphi}_0(y) dy, & \varphi_1(x) \hat{\varphi}_1(x) = \nu(x) \quad (40b) \end{cases}$$

Just like how the value function of an SOC problem fully characterizes the optimal control and its corresponding optimal distribution, so does the SB potential $\varphi_t(x)$:

$$p^*(X) = p^{\text{base}}(X) \frac{\varphi_1(X_1)}{\varphi_0(X_0)} = p^{\text{base}}(X|X_0) \varphi_1(X_1) \hat{\varphi}_0(X_0), \quad (41)$$

where the last equality is due to $p^{\text{base}}(X) = p^{\text{base}}(X|X_0) \mu(X_0)$ and then invoking (40a). Note that (41) recovers (10) by marginalizing over $t \in (0, 1)$. Due to the construction of $\varphi_t(x)$ and $\hat{\varphi}_t(x)$ in (40), the marginal optimal distribution admits a strikingly simple factorization:

$$\begin{aligned} p_t^*(x) &= \int p^{\text{base}}(X, X_t = x|X_0) \varphi_1(X_1) \hat{\varphi}_0(X_0) dX \\ &= \int \int p^{\text{base}}(X_1|X_t = x) p^{\text{base}}(X_t = x|X_0) \varphi_1(X_1) \hat{\varphi}_0(X_0) dX_0 dX_1 \\ &= \left(\int p^{\text{base}}(X_t = x|X_0) \hat{\varphi}_0(X_0) dX_0 \right) \left(\int p^{\text{base}}(X_1|X_t = x) \varphi_1(X_1) dX_1 \right) \\ &= \hat{\varphi}_t(x) \varphi_t(x), \end{aligned} \quad (42)$$

or, more generally,

$$p_{s,t}^*(y, x) = p_{t|s}^{\text{base}}(x|y) \hat{\varphi}_s(y) \varphi_t(x), \quad s \leq t. \quad (43)$$

Derivation of (13) We now provide a simpler derivation of (13) compared to its original derivation based on path measure theory (Shi et al., 2023):

$$\begin{aligned} \nabla \log \hat{\varphi}_t(x) &\stackrel{(40b)}{=} \frac{1}{\hat{\varphi}_t(x)} \nabla_x \int p_{t|0}^{\text{base}}(x|y) \hat{\varphi}_0(y) dy \\ &= \frac{1}{\hat{\varphi}_t(x)} \int \nabla_x \log p_{t|0}^{\text{base}}(x|y) p_{t|0}^{\text{base}}(x|y) \hat{\varphi}_0(y) dy \\ &= \int \nabla_x \log p_{t|0}^{\text{base}}(x|y) p_{0|t}^*(y|x) dy, \end{aligned} \quad (44)$$

where the last equality follows by

$$p_{0|t}^*(y|x) \stackrel{(42)}{=} \frac{p_{0,t}^*(y, x)}{\hat{\varphi}_t(x) \varphi_t(x)} \stackrel{(43)}{=} \frac{p_{t|0}^{\text{base}}(x|y) \hat{\varphi}_0(y) \varphi_t(x)}{\hat{\varphi}_t(x) \varphi_t(x)} = \frac{p_{t|0}^{\text{base}}(x|y) \hat{\varphi}_0(y)}{\hat{\varphi}_t(x)}.$$

Equation (44) implies a matching-based variational formulation of $\nabla \log \hat{\varphi}_t(\cdot)$ —also known as the *bridge matching* objective in data-driven SB (Shi et al., 2023; Liu et al., 2023).

$$\nabla \log \hat{\varphi}_t = \arg \min_h \mathbb{E}_{p_{0,t}^*} [\|h_t(X_t) - \nabla_{x_t} \log p^{\text{base}}(X_t|X_0)\|^2]. \quad (45)$$

Equation (45) recovers (13) at $t = 1$.

A.3 Additional Related Works

In this subsection, we provide additional review on existing learning-based methods for sampling Boltzmann distributions.

Learning-augmented MCMC This class of methods can be thought of as extension of classical sampling methods—such as MCMC (Metropolis et al., 1953; HASTINGS, 1970), Sequential Monte Carlo (SMC; Del Moral et al., 2006) and Annealed Importance Sampling (AIS; Neal, 2001)—where

traditional proposal distributions are replaced with modern machine learning models. For instance, [Arbel et al. \(2021\)](#) and [Gabri   et al. \(2022\)](#) use normalizing flows ([Chen et al., 2018](#)) as learned proposal distributions, whereas [Matthews et al. \(2022\)](#) employ stochastic normalizing flow ([Wu et al., 2020](#)). More recently, [Chen et al. \(2025\)](#) have explored the use of diffusion models ([Song et al., 2021](#); [Ho et al., 2020](#)). However, training these models typically requires computing importance weights, which necessitates a large number of energy evaluations.

MCMC-augmented Diffusion Samplers Alternatively, methods of this class adopt modern generative models to sampling Boltzmann distributions and incorporate MCMC techniques to mitigate the lack of explicit target samples. For example, [Phillips et al. \(2024\)](#), [\(De Bortoli et al., 2024\)](#) and [\(Akhound-Sadegh et al., 2024\)](#) employ score matching objective from score-based diffusion models ([Song et al., 2021](#); [Ho et al., 2020](#)). In contrast, [Albergo and Vanden-Eijnden \(2024\)](#) base their method on action matching objectives ([Neklyudov et al., 2023](#)). However, estimating target samples requires computing importance weights, which makes these methods computationally expensive in terms of energy function evaluations.

B Proofs

B.1 Preliminary and Additional Theoretical Results

Lemma B.1 (It   lemma ([It  , 1951](#))). *Let X_t be the solution to the It   SDE:*

$$dX_t = f_t(X_t)dt + \sigma_t dW_t.$$

Then, the stochastic process $v_t(X_t)$, where $v \in C^{1,2}([0, 1], \mathbb{R}^d)$, is also an It   process:

$$dv_t(X_t) = \left[\partial_t v_t(X_t) + \nabla v_t(X_t) \cdot f + \frac{1}{2} \sigma_t^2 \Delta v_t(X_t) \right] dt + \sigma_t \nabla v_t(X_t) \cdot dW_t. \quad (46)$$

Lemma B.2 (Laplacian trick). *For any twice-differentiable function π such that $\pi(x) \neq 0$, it holds that*

$$\frac{1}{\pi(x)} \Delta \pi(x) = \|\nabla \log \pi(x)\|^2 + \Delta \log \pi(x) \quad (47)$$

Proof.

$$\begin{aligned} \Delta \pi(x) &= \nabla \cdot \nabla \pi(x) \\ &= \nabla \cdot (\pi(x) \nabla \log \pi(x)) \\ &= \nabla \pi(x) \cdot \nabla \log \pi(x) + \pi(x) \Delta \log \pi(x) \\ &= \pi(x) (\|\nabla \log \pi(x)\|^2 + \Delta \log \pi(x)) \end{aligned}$$

□

Theorem B.3 (SB characteristics of SOC). *The optimal distribution p^* of the SOC problem in (20) is also the solution to the following SB problem:*

$$\arg \min_p \{D_{\text{KL}}(p \| p^{\text{base}}) : p_0 = \mu, \quad p_1 = p_1^*\}. \quad (48)$$

Proof. We aim to show that there exist a transform such that the SOC’s optimality equation (21) can be reinterpreted as the ones for SB (40). To this end, consider

$$\varphi_t(x) := e^{-V_t(x)}, \quad \hat{\varphi}_t(x) := e^{V_t(x)} p_t^*(x). \quad (49)$$

One can verify that the value function $V_t(x)$ defined in (21) can be rewritten as

$$\varphi_t(x) = \int p_{1|t}^{\text{base}}(y|x) \varphi_1(y) dy.$$

On the other hand, we can expand $\hat{\varphi}_t(x)$ by

$$\begin{aligned}
\hat{\varphi}_t(x) &= e^{V_t(x)} \int p^*(X|X_t=x) dX \\
&= e^{V_t(x)} \int p^{\text{base}}(X_1|X_t=x) p^{\text{base}}(X_t=x, X_0) e^{-V_1(X_1)+V_0(X_0)} dX_1 dX_0 && \text{by (25)} \\
&= e^{V_t(x)} \int p^{\text{base}}(X_t=x, X_0) e^{-V_t(x)+V_0(X_0)} dX_0 && \text{by (21)} \\
&= \int p^{\text{base}}(X_t=x|X_0) \mu(X_0) e^{V_0(X_0)} dX_0 \\
&= \int p_{t|0}^{\text{base}}(x|y) \hat{\varphi}_0(y) dy. && \text{by (49)}
\end{aligned}$$

Combined, the optimality equation (21) for the SOC problem can be rewritten equivalently as

$$u_t^*(x) = \sigma_t \nabla \log \varphi_t(x), \quad \text{where} \quad \begin{cases} \varphi_t(x) = \int p_{1|t}^{\text{base}}(y|x) \varphi_1(y) dy, & \varphi_0(x) \hat{\varphi}_0(x) = \mu(x), \\ \hat{\varphi}_t(x) = \int p_{t|0}^{\text{base}}(x|y) \hat{\varphi}_0(y) dy, & \varphi_1(x) \hat{\varphi}_1(x) = p_1^*(x). \end{cases}$$

We conclude that p^* indeed solves (48). $\square$

Corollary B.4 (Reciprocal process of the SOC problem). *The optimal distribution p^* of the SOC problem in (20) is a reciprocal process, i.e.,*

$$p^*(X) = p^{\text{base}}(X|X_0, X_1) p^*(X_0, X_1). \quad (51)$$

B.2 Missing Proofs in Main Paper

Proof of Theorem 3.1 Comparing (8a) to (21), we can reinterpret $\varphi_t(x)$ as an value function $V_t(x)$ by reinterpreting

$$V_t(x) := -\log \varphi_t(x), \quad g(x) := -\log \varphi_1(x) \stackrel{(8b)}{=} \log \frac{\hat{\varphi}_1(x)}{\nu(x)}.$$

That is, the kinetic-optimal drift of SB solves an SOC problem (4) with a terminal cost $g(x) := \frac{\hat{\varphi}_1(x)}{\nu(x)}$. $\square$

Proof of Theorem 4.1 For notational simplicity, we will denote $q \equiv q^{\bar{h}^{(k-1)}}$ throughout the proof. We first rewrite the backward SDE (17) in the forward direction (Nelson, 2020):

$$dX_t = [f_t - \sigma_t^2 \nabla \log \phi_t + \sigma_t^2 \nabla \log q_t] dt + \sigma_t dW_t, \quad X_1 \sim \nu,$$

where we rewrite $\phi_t(x)$ w.r.t. the forward time coordiante:

$$\phi_t(x) = \int p_{t|0}^{\text{base}}(x|y) \phi_0(y) dy, \quad \phi_1(x) = \bar{h}^{(k-1)}(x). \quad (52)$$

Note that (52) admits an equivalent PDE form by invoking Feynman-Kac formula (Le Gall, 2016):

$$\partial_t \phi_t(x) = -\nabla \cdot (f_t \phi_t) + \frac{\sigma_t^2}{2} \Delta \phi_t(x), \quad \phi_1(x) = \bar{h}^{(k-1)}(x). \quad (53)$$

On the other hand, the dynamics of $\partial_t q$ follows the Fokker Plank equation (Øksendal, 2003):

$$\begin{aligned}
\partial_t q_t &= -\nabla \cdot ((f_t - \sigma_t^2 \nabla \log \phi_t + \sigma_t^2 \nabla \log q_t) q_t) + \frac{1}{2} \sigma_t^2 \Delta q_t \\
&= \nabla \cdot ((\sigma_t^2 \nabla \log \phi_t - f_t) q_t) - \frac{1}{2} \sigma_t^2 \Delta q_t,
\end{aligned}$$

and straightforward calculation yields

$$\partial_t \log q_t = \sigma_t^2 \Delta \log \phi_t - \nabla \cdot f_t + (\sigma_t^2 \nabla \log \phi_t - f_t) \cdot \nabla \log q_t - \frac{1}{2} \sigma_t^2 \|\nabla \log q_t\|^2 - \frac{1}{2} \sigma_t^2 \Delta \log q_t, \quad (54)$$

where we apply the Laplacian trick (47) to $\frac{1}{q}\Delta q = \|\nabla \log q_t\|^2 + \Delta \log q_t$.

Now, recall that p is the path distribution induced by the following SDE:

$$dX_t = [f_t(X_t) + \sigma_t u_t(X_t)] dt + \sigma_t dW_t, \quad X_0 \sim \mu. \quad (55)$$

Invoke Ito Lemma (46) to $\log q_t(X_t)$, where X_t follows (55):

$$\begin{aligned} d \log q_t &= \left[\partial_t \log q_t + \nabla \log q_t \cdot (f_t + \sigma_t u_t) + \frac{1}{2} \sigma_t^2 \Delta \log q_t \right] dt + \sigma_t \nabla \log q_t \cdot dW_t \\ &\stackrel{(54)}{=} \left[\sigma_t^2 \Delta \log \phi_t - \nabla \cdot f_t + \sigma_t^2 \nabla \log \phi_t \cdot \nabla \log q_t - \frac{1}{2} \sigma_t^2 \|\nabla \log q_t\|^2 + \nabla \log q_t \cdot (\sigma_t u_t) \right] dt \\ &\quad + \sigma_t \nabla \log q_t \cdot dW_t \end{aligned} \quad (56)$$

Likewise, invoke Ito Lemma (46) to $\log \phi_t(X_t)$, where X_t follows (55):

$$\begin{aligned} d \log \phi_t &= \left[\partial_t \log \phi_t + \nabla \log \phi_t \cdot (f_t + \sigma_t u_t) + \frac{1}{2} \sigma_t^2 \Delta \log \phi_t \right] dt + \sigma_t \nabla \log \phi_t \cdot dW_t \\ &\stackrel{(53)}{=} \left[-\nabla \cdot f_t + \frac{\sigma_t^2}{2} \frac{\Delta \phi_t}{\phi_t} + \nabla \log \phi_t \cdot (\sigma_t u_t) + \frac{1}{2} \sigma_t^2 \Delta \log \phi_t \right] dt + \sigma_t \nabla \log \phi_t \cdot dW_t \\ &\stackrel{(47)}{=} \left[-\nabla \cdot f_t + \frac{\sigma_t^2}{2} (\|\nabla \log \phi_t\|^2 + \Delta \log \phi_t) + \nabla \log \phi_t \cdot (\sigma_t u_t) + \frac{1}{2} \sigma_t^2 \Delta \log \phi_t \right] dt + \sigma_t \nabla \log \phi_t \cdot dW_t \\ &= \left[-\nabla \cdot f_t + \frac{\sigma_t^2}{2} \|\nabla \log \phi_t\|^2 + \nabla \log \phi_t \cdot (\sigma_t u_t) + \sigma_t^2 \Delta \log \phi_t \right] dt + \sigma_t \nabla \log \phi_t \cdot dW_t \end{aligned} \quad (57)$$

Subtracting (57) from (56) leads to

$$d \log \phi_t - d \log q_t = \left[\frac{1}{2} \|u_t + \sigma_t \nabla \log \phi_t - \sigma_t \nabla \log q_t\|^2 - \frac{1}{2} \|u_t\|^2 \right] dt + \sigma_t \nabla \log \frac{\phi_t}{q_t} \cdot dW_t. \quad (58)$$

Finally, we are ready to compute the variational objective in (16):

$$\begin{aligned} D_{\text{KL}}(p||q^{\bar{h}^{(k-1)}}) &= \mathbb{E}_{X \sim p^u} \left[\int_0^1 \frac{1}{2} \|u_t(X_t) + \sigma_t \nabla \log \phi_t(X_t) - \sigma_t \nabla \log q_t(X_t)\|^2 dt \right] \\ &\stackrel{(58)}{=} \mathbb{E}_{X \sim p^u} \left[\int_0^1 \left(\frac{1}{2} \|u_t(X_t)\|^2 + d \log \phi_t(X_t) - d \log q_t(X_t) \right) dt \right] \\ &= \mathbb{E}_{X \sim p^u} \left[\int_0^1 \frac{1}{2} \|u_t(X_t)\|^2 dt + \log \frac{\phi_1(X_1)}{q_1(X_1)} - \log \frac{\phi_0(X_0)}{q_0(X_0)} \right] \end{aligned} \quad (59)$$

$$\propto \mathbb{E}_{X \sim p^u} \left[\int_0^1 \frac{1}{2} \|u_t(X_t)\|^2 dt + \log \frac{\bar{h}^{(k-1)}(X_1)}{\nu(X_1)} \right]. \quad (60)$$

That is, we have shown that the variational objective $D_{\text{KL}}(p||q^{\bar{h}^{(k-1)}})$ is equivalent (up to an additive constant) to an SOC problem (60). Applying Reciprocal Adjoint Matching (Havens et al., 2025) with the reciprocal process from Corollary B.4 conclude that $D_{\text{KL}}(p||q^{\bar{h}^{(k-1)}})$ is minimized by $p^{u^{(k)}}$. $\square$

Proof of Theorem 4.2 For notational simplicity, we will denote $p^{(k)} \equiv p^{u^{(k)}}$ throughout the proof. Let q be the path distribution induced by a backward SDE, propagating along the time coordinate $s := 1 - t$:

$$dY_s = [-f_s(Y_s) + \sigma_s v_s(Y_s)] ds + \sigma_s dW_s, \quad Y_0 \sim \nu.$$

Next, rewrite the forward SDE $p^{(k)}$ in the backward direction:

$$dY_s = [-f_s - \sigma_s u_s^{(k)} + \sigma_s^2 \nabla \log p_s^{(k)}] ds + \sigma_s dW_s, \quad Y_0 \sim p_t^{(k)}|_{t=1}.$$

By Theorem B.3, we know that $p^{(k)}$ is the SB solution, thereby satisfying

$$u_t^{(k)}(x) = \sigma_t \nabla \log \varphi_t(x), \text{ where } \begin{cases} \varphi_t(x) = \int p_{1|t}^{\text{base}}(y|x) \varphi_1(y) dy, & \varphi_0(x) \hat{\varphi}_0(x) = \mu(x) & (61a) \\ \hat{\varphi}_t(x) = \int p_{t|0}^{\text{base}}(x|y) \hat{\varphi}_0(y) dy, & \varphi_1(x) \hat{\varphi}_1(x) = p_1^{(k)}(x) & (61b) \end{cases}$$

Since we are working with the backward time coordinate s , it is convenience to define $\phi_s := \hat{\varphi}_{1-t}$ and rewrite (61b) by

$$\phi_s(y) = \int p_{1-s|0}^{\text{base}}(y|z)\phi_1(z)dz, \quad \phi_0(y) = \frac{p_1^{(k)}(y)}{\varphi_1(y)}. \quad (62)$$

Now, expanding the variational objective with Girsanov Theorem yields (Särkkä and Solin, 2019)

$$D_{\text{KL}}(p^{(k)}||q) = \mathbb{E}_{Y \sim p^{(k)}} \left[\int_0^1 \frac{1}{2} \|\sigma_s \nabla \log \varphi_s(Y_s) + \sigma_s \nabla \log p_s^{(k)}(Y_s) - v_s(Y_s)\|^2 ds \right], \quad (63)$$

which is minimized point-wise at

$$v_s^*(y) = \sigma_s \nabla \log \frac{p_s^{(k)}(y)}{\varphi_s(y)} \stackrel{(42)}{=} \sigma_s \nabla \log \hat{\varphi}_s(y).$$

In other words, the backward SDE that minimizes (63) must obey

$$dY_s = [-f_s(Y_s) + \sigma_s^2 \nabla \log \phi_s(Y_s)] ds + \sigma_s dW_s, \quad Y_0 \sim \nu,$$

with ϕ_s defined in (62). That is, we have concluded so far that

$$q^{p_1^{(k)}/\varphi_1} = \arg \min_q \{D_{\text{KL}}(p^{(k)}||q) : q_1 = \nu\}. \quad (64)$$

Hence, it remains to be shown that the minimizer $\bar{h}_1^{(k)}$ of the CM objective at stage k equals $\frac{p_1^{(k)}}{\varphi_1}$. This is indeed the case since $p^{(k)}$ is the SB solution:

$$\nabla \log \bar{h}^{(k)} \stackrel{(15)}{:=} \arg \min_h \mathbb{E}_{p_{0,1}^{(k)}} [\|h(X_1) - \nabla_{x_1} \log p^{\text{base}}(X_1|X_0)\|^2] \stackrel{(45)}{=} \nabla \log \hat{\varphi}_1 \stackrel{(42)}{=} \nabla \log \frac{p_1^{(k)}}{\varphi_1}.$$

□

C Practical Implementation of ASBS

Algorithm 2 summarizes the practical implementation of ASBS, where we expand the adjoint and corrector matching steps (*i.e.*, lines 3 and 4 in Algorithm 1) to full details. Table 5 provides the hyper-parameters for each task. We break down each component as follows:

Harmonic prior μ_{harmonic} Recall the harmonic prior in (19):

$$\mu_{\text{harmonic}}(x) \propto \exp(-\frac{1}{2} \sum_{i,j} \|x_i - x_j\|^2). \quad (65)$$

In practice, we set $\alpha = 1$ and implement (65) as an anisotropic Gaussian. For instance, for a 2-particle system in 3D, *i.e.*, $x = [x_1; x_2] \in \mathbb{R}^6$, we can rewrite (65) as a quadratic potential,

$$\exp(-\frac{1}{2} \|x_1 - x_2\|^2) = \exp(x^\top R x), \quad \text{where } R = \begin{bmatrix} 1 & 0 & 0 & -\frac{1}{2} & 0 & 0 \\ 0 & 1 & 0 & 0 & -\frac{1}{2} & 0 \\ 0 & 0 & 1 & 0 & 0 & -\frac{1}{2} \\ -\frac{1}{2} & 0 & 0 & 1 & 0 & 0 \\ 0 & -\frac{1}{2} & 0 & 0 & 1 & 0 \\ 0 & 0 & -\frac{1}{2} & 0 & 0 & 1 \end{bmatrix}, \quad (66)$$

and then sample x from the Gaussian $\mathcal{N}(x; 0, (R + \epsilon I)^{-1})$, where we set $\epsilon = 10^{-4}$.

Noise schedule σ_t We consider two types of noise schedule.

- The *geometric noise schedule* (Song et al., 2021; Karras et al., 2022) monotonically decays from $t = 0$ to 1 according to some prescribed $\beta_{\min}$ and $\beta_{\max}$:

$$\sigma_t^{\text{geometric}} := \beta_{\min} \left(\frac{\beta_{\max}}{\beta_{\min}} \right)^{1-t} \sqrt{2 \log \frac{\beta_{\max}}{\beta_{\min}}}. \quad (67)$$

Algorithm 2 Adjoint Schrödinger Bridge Sampler (ASBS)

Require: Sample-able source $X_0 \sim \mu$, differentiable energy $E(x)$, parametrized drift $u_\theta(t, x)$ and corrector $h_\phi(x)$, replay buffers $\mathcal{B}_{\text{adj}}$ and $\mathcal{B}_{\text{crt}}$, number of stages K , numbers of AM and CM epochs M_{adj} and M_{crt} , number of resamples N , number of gradient steps L , time scaling λ_t , maximum energy gradient norm α_{max} .

- 1: Initialize $h_\phi^{(0)} := 0$ ▷ IPF initialization
 - 2: **for** stage k **in** $1, 2, \dots, K$ **do**
 - 3: **for** epoch **in** $1, 2, \dots, M_{\text{adj}}$ **do** ▷ adjoint matching
 - 4: Sample from model $\{(X_0^{(i)}, X_1^{(i)})\}_{i=1}^N \sim p^{\bar{u}^{(k)}}$, where $\bar{u}^{(k)} = \text{stopgrad}(u_\theta^{(k)})$
 - 5: Compute adjoint target $a_t^{(i)} := \text{stopgrad}(\text{clip}(\nabla E(X_1^{(i)}), \alpha_{\text{max}}) + h_\phi^{(k)}(X_1^{(i)}))$
 - 6: Update replay buffer $\mathcal{B}_{\text{adj}} \leftarrow \mathcal{B}_{\text{adj}} \cup \{(X_0^{(i)}, X_1^{(i)}, a_t^{(i)})\}_{i=1}^N$
 - 7: Take L gradient steps $\nabla_\theta \mathcal{L}_{\text{AM}}$ w.r.t. the AM objective:
$$\mathcal{L}_{\text{AM}}(\theta) := \mathbb{E}_{t \sim \mathcal{U}[0,1], (X_0, X_1, a_t) \sim \mathcal{B}_{\text{adj}}, X_t \sim p^{\text{base}}(\cdot | X_0, X_1)} \left[\lambda_t \|u_\theta^{(k)}(t, X_t) + \sigma_t a_t\|^2 \right]$$
 - 8: **end for**
 - 9: **for** epoch **in** $1, 2, \dots, M_{\text{crt}}$ **do** ▷ corrector matching
 - 10: Sample from model $\{(X_0^{(i)}, X_1^{(i)})\}_{i=1}^N \sim p^{\bar{u}^{(k)}}$, where $\bar{u}^{(k)} = \text{stopgrad}(u_\theta^{(k)})$
 - 11: Update replay buffer $\mathcal{B}_{\text{crt}} \leftarrow \mathcal{B}_{\text{crt}} \cup \{(X_0^{(i)}, X_1^{(i)})\}_{i=1}^N$
 - 12: Take L gradient steps $\nabla_\phi \mathcal{L}_{\text{CM}}$ w.r.t. the CM objective:
$$\mathcal{L}_{\text{CM}}(\phi) := \mathbb{E}_{(X_0, X_1) \sim \mathcal{B}_{\text{crt}}} \left[\|h_\phi^{(k)}(X_1) - \nabla_{x_1} \log p^{\text{base}}(X_1 | X_0)\|^2 \right]$$
 - 13: **end for**
 - 14: **end for**
 - 15: **return** Kinetic-optimal drift $u^* \approx u_\theta(t, x)$
-

It is convenience to further define

$$\kappa_{t|s} := \int_s^t \sigma_\tau^2 d\tau \stackrel{\text{geometric}}{=} \beta_{\text{max}}^2 \cdot \left(\left(\frac{\beta_{\text{min}}}{\beta_{\text{max}}} \right)^{2s} - \left(\frac{\beta_{\text{min}}}{\beta_{\text{max}}} \right)^{2t} \right), \quad \bar{\beta}^2 := \beta_{\text{max}}^2 - \beta_{\text{min}}^2, \quad \gamma_t := \frac{\kappa_{t|0}}{\bar{\beta}^2}. \quad (68)$$

With them, the conditional base distribution when $f := 0$ can be represented compactly by

$$p^{\text{base}}(X_t | X_0) = \mathcal{N}(X_t; X_0, \kappa_{t|0} I) \quad (69a)$$

$$p^{\text{base}}(X_t | X_0, X_1) = \mathcal{N}(X_t; (1 - \gamma_t)X_0 + \gamma_t X_1, \bar{\beta}^2 \gamma_t (1 - \gamma_t) I) \quad (69b)$$

- The *constant noise schedule* simply sets

$$\sigma_t \stackrel{\text{constant}}{=} \sigma. \quad (70)$$

When $f := 0$, the base SDE is effectively a standard Brownian motion whose conditional distributions obey

$$p^{\text{base}}(X_t | X_0) = \mathcal{N}(X_t; X_0, \sigma^2 t I) \quad (71a)$$

$$p^{\text{base}}(X_t | X_0, X_1) = \mathcal{N}(X_t; (1 - t)X_0 + tX_1, \sigma^2 t(1 - t) I) \quad (71b)$$

Replay buffers $\mathcal{B}_{\text{adj}}$ and $\mathcal{B}_{\text{crt}}$ Similar to many previous diffusion samplers (Havens et al., 2025; Akhound-Sadegh et al., 2024; Chen et al., 2025), we employ replay buffers $\mathcal{B}$ in computation of both adjoint (14) and corrector (15) matching objectives. Specifically, we rebase the expectation over model samples $p^{u^{(k)}}$ onto a replay buffer $\mathcal{B}$, which stores the most latest $|\mathcal{B}|$ samples. We update the buffer with N new samples every L gradient steps. Note that the use of replay buffers effectively render ASBS a hybrid method between on-policy and off-policy.

Parametrization of u_θ and h_ϕ For each energy function, we parametrize the drift $u_\theta(t, x)$ and the corrector $h_\phi(x)$ with two neural networks, $v_\theta(t, x)$ and $v_\phi(t, x)$, of the same architecture.

Specifically, we parametrize the drift as $u_\theta(t, x) := \sigma_t v_\theta(t, x)$, which effectively eliminates the noise schedule “ σ_t ” in matching target (see (14)), making it time-invariant for each sampled trajectory. The only exception is the conformer generation task, where we keep the original parametrization $u_\theta(t, x) := v_\theta(t, x)$, which empirically yields better results. On the other hand, since $h_\phi(x)$ is independent of time, we simply set a fixed time input $t = 1$, i.e., $h_\phi(x) := v_\phi(1, x)$.

The specific parametrization $v(t, x)$ employed for each task are detailed below.

- **MW-5:** We consider $v(t, x)$ a standard fully-connected network with 4 layers with 64 hidden features of the following form:

$$\text{output} = \text{layer_n} \circ \dots \circ \text{layer_1} \circ (\mathbf{x_embed}(x) + \mathbf{t_embed}(t))$$

- **DW-4, LJ-13, LJ-55:** We consider $v(t, x)$ a Equivariant Graph Neural Network (EGNN; [Satorras et al., 2021](#)) with 5 layers and 128 hidden features. The architecture of EGNN is aligned with prior methods ([Akhound-Sadegh et al., 2024](#); [Havens et al., 2025](#)).
- **Alanine dipeptide:** We use the same architecture as in MW-5, except with 8 layers with 256 hidden features.
- **Conformer generation:** We consider $v(t, x)$ a similar EGNN used in Adjoint Sampling ([Havens et al., 2025](#)), except with 20 layers. Ablation study on the same EGNN architecture can be found in Appendix D.4.

Clipping $\alpha_{\max}$ We clip the energy gradient to prevent its maximum norm from exceeding $\alpha_{\max}$.

Time scaling λ_t Following standard practices for AM objective, we employ a time scaling λ_t to improve numerical stability. Note that this does not affect the minimizer of the AM objective. We set $\lambda_t := \frac{1}{\sigma_t^2}$ for all tasks.

Translation invariance For DW-4, LJ-13, LJ-55, and conformer generation tasks, we follow prior methods ([Akhound-Sadegh et al., 2024](#); [Havens et al., 2025](#)) by restricting the state space to a zero center-of-mass (ZCOM) subspace and thereby enforcing translation invariance.

For a n -particle k -dimensional system, i.e., $x = [x_1; \dots; x_n]$ where $x_i \in \mathbb{R}^k$, the ZCOM subspace is defined as $\mathcal{X}^{\text{ZCOM}} = \{x \in \mathbb{R}^{nk} : \sum_{i=1}^n x_i = 0\}$. Practically, this is achieved by projecting the initial sample $X_0 \sim \mu$, the SDE’s noise dW_t , and the energy gradient $\nabla E(\cdot)$ onto $\mathcal{X}^{\text{ZCOM}}$. Note that the output of EGNN is by construction ZCOM.

Formally, the adaption is equivalent to augmenting the SDE with a projection matrix $A \in \mathbb{R}^{nk \times nk}$:

$$dX_t = \sigma_t A u_t(X_t) dt + \sigma_t A dW_t, \quad X_0 = A Y_0, \quad Y_0 \sim \mu, \quad A = \left(I_n - \frac{1}{n} \mathbf{1}_n \mathbf{1}_n^\top \right) \otimes I_k, \quad (72)$$

where $\otimes$ is the Kronecker product, $I_n \in \mathbb{R}^{n \times n}$ is an identity matrix, and $\mathbf{1}_n \in \mathbb{R}^n$ is a vector of ones.

Initialization and alternate procedure As ASBS is an instantiation of the IPF algorithm (see Theorem 3.2), it must adhere to the IPF initialization protocol to ensure theoretical convergence to the global solution. Specifically, the IPF initialization can be implemented in two ways

- Initialize with $h_\phi^{(0)} := 0$ and run AM, CM, ... until convergence. We adopt this setup for all tasks.
- Initialize with $u_\theta^{(0)} := 0$ and run CM, AM, ... until convergence. Since $p^{u^{(0)}} = p^{\text{base}}$ in this setup, the optimal corrector at the first CM stage is known analytically:

$$\begin{aligned} h^{(1)}(x) &\stackrel{(15)}{=} \int p_{0|1}^{\text{base}}(y|x) \nabla_x \log p_{1|0}^{\text{base}}(x|y) dy \\ &= \int \frac{p_{0|1}^{\text{base}}(y|x)}{p_{1|0}^{\text{base}}(x|y)} \nabla_x p_{1|0}^{\text{base}}(x|y) dy \\ &= \frac{1}{p_1^{\text{base}}(x)} \nabla_x \int p_0^{\text{base}}(y) p_{1|0}^{\text{base}}(x|y) dy \\ &= \nabla \log p_1^{\text{base}}(x) \end{aligned} \quad (73)$$

Table 5: Hyperparameters of ASBS for the each task.

	Synthetic energy functions				Alanine dipeptide	Conformer generation
	MW-5	DW-4	LJ-13	LJ-55		
μ	$\mathcal{N}(0, 1)$	μ_{harmonic}	in (19) with $\alpha=2, 2, 1$		$\mathcal{N}(0, 0.25)$	μ_{harmonic}
$\beta_{\min}$	—	0.001	0.001	0.001	0.001	0.001
$\beta_{\max}$	—	1	1	2	0.5	1
σ	0.2	—	—	—	—	—
K	5	20	15	15	15	3
M_{adj}	100	200	300	300	4000	2500
M_{crt}	20	20	20	20	2000	2000
N	1000	1000	1000	1000	1000	128
L	200	100	100	100	100	100
$ \mathcal{B} $	10^4	10^4	10^4	10^4	10^4	6.4×10^4
$\alpha_{\max}$	—	100	100	100	100	150
λ_t	$\frac{1}{\sigma_t^2}$	$\frac{1}{\sigma_t^2}$	$\frac{1}{\sigma_t^2}$	$\frac{1}{\sigma_t^2}$	$\frac{1}{\sigma_t^2}$	$\frac{1}{\sigma_t^2}$

In practice, we find that the two setups yield similar performance.

RDKit warm-start This warm-starts the drift u_θ using RDKit samples. The procedure is inspired by the fact that (Shi et al., 2023; Liu et al., 2023):

$$\begin{aligned}
u_t^* &= \sigma_t \nabla \log \varphi_t \\
&= \arg \min_{u_t} \mathbb{E}_{p_{t,1}^*} [\|u_t(X_t) - \sigma_t \nabla_{x_t} \log p^{\text{base}}(X_1|X_t)\|^2] \\
&= \arg \min_{u_t} \mathbb{E}_{(X_0, X_1) \sim p_{0,1}^*, X_t \sim p^{\text{base}}(\cdot|X_0, X_1)} [\|u_t(X_t) - \sigma_t \nabla_{x_t} \log p^{\text{base}}(X_1|X_t)\|^2]. \quad (74)
\end{aligned}$$

where the last equality is due to

$$\begin{aligned}
p_{0,t,1}^*(x, y, z) &\stackrel{(41)}{=} p_{t,1|0}^{\text{base}}(y, z|x) \hat{\varphi}_0(x) \varphi_1(z) \\
&= p_{t|0,1}^{\text{base}}(y|x, z) p_{1|t}^{\text{base}}(z|y) \hat{\varphi}_0(x) \varphi_1(z) && \text{by Markov property} \\
&\stackrel{(43)}{=} p_{t|0,1}^{\text{base}}(y|x, z) p_{0,1}^*(x, z). \quad (75)
\end{aligned}$$

Equation (74) can be understood as an analogy of (45) for another SB potential φ_t . In practice, given RDKit samples $X_1 \sim q^{\text{RDKit}}$, we warm-start ASBS by minimizing w.r.t. the following objective:

$$\begin{aligned}
\mathcal{L}_{\text{warmup}}(\theta) &= \mathbb{E}_{t \sim \mathcal{U}[0,1], X_0 \sim \mu, X_1 \sim q^{\text{RDKit}}, X_t \sim p^{\text{base}}(\cdot|X_0, X_1)} [\tilde{\lambda}_t \|u_t(X_t) - \sigma_t \nabla_{x_t} \log p^{\text{base}}(X_1|X_t)\|^2] \\
&\stackrel{(69a)}{=} \mathbb{E}_{t \sim \mathcal{U}[0,1], X_0 \sim \mu, X_1 \sim q^{\text{RDKit}}, X_t \sim p^{\text{base}}(\cdot|X_0, X_1)} \left[\tilde{\lambda}_t \|u_t(X_t) - \frac{\sigma_t}{\kappa_{1|t}} (X_1 - X_t)\|^2 \right], \quad (76)
\end{aligned}$$

where $\kappa_{1|t}$ is defined in (68) for the geometric noise schedule. We set the time scaling $\tilde{\lambda}_t := \sqrt{\frac{\sigma_t}{\kappa_{1|t}}}$.

Note that, unlike AS, the minimizer of (76) does not equal u^* , since $(X_0, X_1) \sim \mu \otimes q^{\text{RDKit}} \neq p_{0,1}^*$ are sampled independently.

D Experiment Details

D.1 Synthetic Energy Functions

D.1.1 Energy functions

In this section, we provide the exact setup for our synthetic energy experiments in Table 2. We consider four synthetic energy functions that have been widely used in recent literature to benchmark sampling and generative algorithms: MW-5, DW-4, LJ-13, and LJ-55.

MW-5 The MW-5 (Many-Well in 5D) energy is a 5-particle 1D system adopted from [Chen et al. \(2025\)](#), where $x = [x_1; \dots; x_5] \in \mathbb{R}^5$ with $x_i \in \mathbb{R}$. The energy function is defined as follows:

$$E(x) = \sum_{i=1}^5 (x_i^2 - \delta)^2 \quad (77)$$

where we set $\delta = 4$. This creates distinct modes centered at combinations of $\pm\sqrt{\delta}$ in each of the d dimensions.

DW-4 The DW-4 (Double-Well for 4 particles in 2D) energy is a physically motivated pairwise potential originally proposed in [Köhler et al. \(2020\)](#) and subsequently used in [Akhound-Sadegh et al. \(2024\)](#); [Havens et al. \(2025\)](#). It defines a system of four particles, each living in $\mathbb{R}^2$, leading to an 8D state vector $x = [x_1; x_2; x_3; x_4] \in \mathbb{R}^8$ with $x_i \in \mathbb{R}^2$. The energy function reads

$$E(x) = \exp \left[\frac{1}{2\tau} \sum_{i < j} (a(d_{ij} - d_0) + b(d_{ij} - d_0)^2 + c(d_{ij} - d_0)^4) \right], \quad (78)$$

where $d_{ij} = \|x_i - x_j\|_2$ is the Euclidean distance between particles i and j . We follow the standard configuration with $a = 0$, $b = -4$, $c = 0.9$, $d_0 = 1$, and temperature $\tau = 1$.

LJ-13 and LJ-55 The Lennard-Jones (LJ) potentials are classical intermolecular potentials commonly used in physics to model atomic interactions. These are defined for a system of n particles in 3D space, with $x = [x_1; \dots; x_n] \in \mathbb{R}^{3n}$ and $x_i \in \mathbb{R}^3$. The index following "LJ-" indicates the number of particles (e.g., 13 or 55). The unnormalized energy function takes the form:

$$E(x) = \frac{\epsilon}{2\tau} \sum_{i < j} \left[\left(\frac{r_m}{d_{ij}} \right)^6 - \left(\frac{r_m}{d_{ij}} \right)^{12} \right] + \frac{c}{2} \sum_i \|x_i - C(x)\|^2, \quad (79)$$

where $d_{ij} = \|x_i - x_j\|_2$ is the pairwise distance and $C(x)$ denotes the center of mass of the particles. We use the parameter values $r_m = 1$, $\epsilon = 1$, $c = 0.5$, and $\tau = 1$, following prior work. The LJ-13 and LJ-55 systems correspond to 39D and 165D, respectively.

D.1.2 Baselines

Here, we outline the procedure used to obtain the values reported in Table 2 for the baseline methods.

For PIS ([Zhang and Chen, 2022](#)), DDS ([Vargas et al., 2023](#)), and LV-PIS ([Richter and Berner, 2024](#)), iDEM ([Akhound-Sadegh et al., 2024](#)), and AS ([Havens et al., 2025](#)), we reuse the values reported in AS ([Havens et al., 2025](#), Table 1) for DW-4, LJ-13, and LJ-55 energy functions. As for MW-5, which is not included in AS, we run iDEM using their official implementation and the rest of baseline methods using our own implementation in PyTorch ([Paszke et al., 2019](#)). We were unable to obtain reportable results for LV-PIS and iDEM on this energy function.

For PDDS ([Phillips et al., 2024](#)) and SCLD ([Chen et al., 2025](#)), we run their official implementations in JAX ([Bradbury et al., 2018](#)) using the default hyperparameter settings specified for the Log-Gaussian Cox Process experiment in their respective papers. To enhance stability and convergence on synthetic energy functions, we tune the gradient clipping parameters. For PDDS, we apply clipping to the gradient of the energy function. For SCLD, we clip both the energy gradient and the Langevin norm. In both cases, the clipping magnitude is selected from the set $\{1, 10, 100, 1000\}$ based on the best validation performance. Training is performed for 100,000 iterations across all runs. For SCLD, we use subtrajectory splitting with the default value of 4, so that it does not degenerate to CMCD ([Vargas et al., 2024](#)). In practice, we find that using subtrajectories yields better results.

D.1.3 Evaluation Metrics

In this subsection, we outline the evaluation criteria used to quantitatively assess the quality of samples generated from synthetic energy functions. We employ three primary metrics: Sinkhorn distance, geometric $\mathcal{W}_2$, and energy $\mathcal{W}_2$, each designed to capture different aspects of distributional similarity between generated and ground truth samples.

Sinkhorn distance To evaluate the similarity between the empirical distributions of generated and reference samples, we compute the Sinkhorn distance using the entropy-regularized optimal transport

formulation (Peyré and Cuturi, 2019), following the implementation of Blessing et al. (2024) and Chen et al. (2025). The Sinkhorn regularization coefficient is set to 10^{-3} throughout. We use 2,000 samples from both the generated and ground truth distributions to compute the metric.

Geometric $\mathcal{W}_2$ For DW and LJ tasks, the potential energy functions—and consequently, the sample distributions—exhibit invariance to both particle permutations and rigid transformations such as rotations and reflections. To appropriately account for these symmetries, we employ the geometric $\mathcal{W}_2$ distance as defined by Akhond-Sadegh et al. (2024) and Havens et al. (2025). Formally, the 2-Wasserstein distance is computed as:

$$\mathcal{W}_2^2(\hat{\nu}, \nu) = \inf_{\pi \in \Pi(\hat{\nu}, \nu)} \int D(x, y)^2 \pi(x, y) dx dy, \quad (80)$$

where $\Pi(\hat{\nu}, \nu)$ denotes the set of joint couplings with prescribed marginals $\hat{\nu}$ (generated) and ν (ground truth), and $D(x, y)$ is a symmetry-aware distance between samples defined as:

$$D(x, y) = \min_{R \in O(s), P \in S(n)} \|x - (R \otimes P)y\|_2. \quad (81)$$

Here, $O(s)$ denotes the group of orthogonal transformations in s spatial dimensions (rotations and reflections), and $S(n)$ represents the symmetric group over n particles. As exact minimization over these symmetry groups is computationally infeasible, we adopt the approximation scheme of Köhler et al. (2020). We use 2000 samples from each generated and ground truth distribution to compute the metric.

Energy $\mathcal{W}_2$ To evaluate fidelity with respect to the target energy landscape, we also compute the 2-Wasserstein distance between the energy values of generated samples and those of ground truth samples. For each target distribution, we generate 2,000 samples from both the model and the reference, and compare their respective energy histograms. This scalar-based Wasserstein metric serves as a proxy for how well the generative model captures the energy histogram of the target distribution.

D.2 Alanine dipeptide

Benchmark description We adopt the experiment setup primarily from (Midgley et al., 2023). Given a configuration of alanine dipeptide, which consists of 22 particles in 3D, *i.e.*, $x = [x_1; \dots; x_{22}] \in \mathbb{R}^{66}$ where $x_i \in \mathbb{R}^3$, we apply the same coordinate transform $\mathcal{T}$ proposed by Midgley et al. (2023). This coordinate transform maps the Cartesian coordinates to internal coordinates, $\mathcal{T}(x) =: z \in \mathbb{R}^{60}$, which include bond lengths, bond angles, and dihedral angles (Stimper et al., 2022). This process effectively removes six degrees of freedom—three for translation and three for rotation—thereby enforcing structural invariance. Non-angular coordinates are further normalized using samples with minimal energies. We refer readers to (Midgley et al., 2023, Appendix F.1) for further details. Note that the internal coordinate transformation is bijective. Hence, we can compute the energy via

$$E(x) = E(\mathcal{T}^{-1}(z)) \quad (82)$$

Evaluation and baselines For each sample $x = \mathcal{T}^{-1}(z) \in \mathbb{R}^{66}$, we extract five torsion angles, including the backbone angles ϕ, ψ and methyl rotation angles $\gamma_1, \gamma_2, \gamma_3$. We report two divergence metrics with respect to the ground-truth distribution, which contains 10^7 samples simulated by Molecular Dynamics. We implement the baseline methods, including PIS (Zhang and Chen, 2022), DDS (Vargas et al., 2023), AS (Havens et al., 2025), using PyTorch (Paszke et al., 2019).

For the KL divergences, we adopt setup from (Wu et al., 2020) and compute the divergence of the ground-truth marginal to model marginal for each torsion angle:

$$D_{\text{KL}}(p^*(\cdot) || p^{u_\theta}(\cdot)) \approx \sum P^*(\cdot) \log \frac{P^*(\cdot) + \epsilon}{P^{u_\theta}(\cdot) + \epsilon}, \quad \epsilon = 10^{-5}, \quad (83)$$

where P^* and P^{u_θ} are histograms of 10^7 samples, discretized between $[-\pi, \pi]$ with 200 intervals.

For the Wasserstein-2 distance, we use the Geometric $\mathcal{W}_2$ in (80), where each sample is now in 2D, $x = [\phi, \psi] \in \mathbb{R}^2$. Due to the high computational cost, we compute the value using a subset of 10^4 samples from the test set ground-truth samples, which is fixed for all methods.

Finally, both Ramachandran plots in Figure 5 are generated using 10^7 samples.

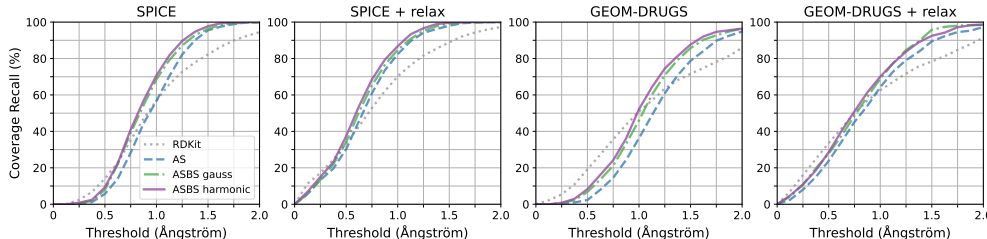


Figure 8: Ablation study on full recall coverage curves (without RDKit warm-start) using the same EGNN architecture as in AS (Havens et al., 2025). Note that Table 6 reports the values at the thresholds $1.0\text{\AA}$ and $1.25\text{\AA}$.

D.3 Amortized conformer generation

In this subsection, we provide some context for the experimental results found in Table 4 regarding the generation of conformers.

Benchmark description Conformers are atomic representations of molecules in cartesian space with their constituent atoms arranged into local minima on the potential energy surface. Molecules are defined to be a graph of atoms (nodes) connected by bonds (edges); conformers are geometric realizations of that molecule. Torsion angles, or rotatable bonds, are particularly important degrees of freedom for defining conformations since bond lengths and bond angles are typically much more stable due to a high sensitivity to perturbations. It is common to consider bond lengths and bond angles fixed, while the torsional degrees of freedom define the conformer.

The task in this benchmark is to take a representation of the molecular graph, usually a SMILES string (Weininger, 1988), and comprehensively sample the conformational configuration space. In flexible molecules, there can be a large number of conformers with many separated modes in a $3n - 6$ dimensional space. (Where n represents the number of atoms and 6 comes from the irrelevance of rotation and translation of the conformer.) We quantify the notion of comprehensively sampling the space by comparing generated structures to a set of conformers sampled using expensive, standard search techniques (Pracht et al., 2024) that were further relaxed using extremely precise density function theory-based, quantum chemistry methods (Neese, 2012; Levine et al., 2025). A detailed description of this benchmark can be found in its source (Havens et al., 2025, Appendix F).

Evaluation and baselines The method of comparison between proposed structure and reference conformer is to use RDKit’s (Landrum, 2006) implementation of *Root Mean Squared Displacement* (RMSD), a measure of distance between atomic structures that is invariant to translation and rotation. We set a threshold RMSD for two structures to match and computed the Recall Coverage and Recall Average Minimum RMSD (AMR). The experiment was performed with both generated structures and with generated structures after a so-called relaxation, i.e. geometry optimization of energy, using eSEN (Fu et al., 2025). The equations for computing these metrics are:

$$\text{COV-R}(\delta) := \frac{1}{L} |\{l \in \{1, \dots, L\} : \exists k \in \{1, \dots, K\}, \text{RMSD}(C_k, C_l^*) < \delta\}| \quad (84)$$

$$\text{AMR-R} := \frac{1}{L} \sum_{l \in \{1, \dots, L\}} \min_{k \in \{1, \dots, K\}} \text{RMSD}(C_k, C_l^*) \quad (85)$$

where $\delta = 0.75 \text{\AA}$ is the coverage threshold, $L = \max(L', 128)$, where L' is the number of reference conformers, $K = 2L$, and let $\{C_l^*\}_{l \in [1, L]}$ and $\{C_k\}_{k \in [1, K]}$ be the sets of ground truth and generated conformers respectively. We capped the reference conformers per molecule at 512 in COV-R.

The values for the baselines are adopted from AS (Havens et al., 2025).

D.4 Additional Experiments and Discussions

Ablation study between AS and ASBS using the same EGNN For the amortized conformer generation task in Table 4, we use an EGNN architecture with 20 layers, whereas AS employs the same architecture with 12 layers. In Table 6, we report the results of ASBS using the same 12-layer

Table 6: Ablation study on amortized conformer generation using the same EGNN architecture as in AS (Havens et al., 2025). We report the recall at the thresholds $1.0\text{\AA}$ and $1.25\text{\AA}$, where the latter was reported in AS.

Method	without relaxation				with relaxation				
	SPICE		GEOM-DRUGS		SPICE		GEOM-DRUGS		
	Coverage $\uparrow$	AMR $\downarrow$	Coverage $\uparrow$	AMR $\downarrow$	Coverage $\uparrow$	AMR $\downarrow$	Coverage $\uparrow$	AMR $\downarrow$	
Threshold 1.0Å	RDKit ETKDG (Riniker and Landrum, 2015)	56.94 $\pm$ 35.82	1.04 $\pm$ 0.52	50.81 $\pm$ 34.69	1.15 $\pm$ 0.61	70.21 $\pm$ 31.70	0.79 $\pm$ 0.44	62.55 $\pm$ 31.67	0.93 $\pm$ 0.53
	AS (Havens et al., 2025)	56.75 $\pm$ 38.15	0.96 $\pm$ 0.26	36.23 $\pm$ 33.42	1.20 $\pm$ 0.43	82.41 $\pm$ 25.85	0.68 $\pm$ 0.28	64.26 $\pm$ 34.57	0.89 $\pm$ 0.45
	ASBS w/ Gaussian prior (Ours)	68.61 $\pm$ 33.48	0.88 $\pm$ 0.25	46.03 $\pm$ 35.99	1.08 $\pm$ 0.36	84.77 $\pm$ 22.65	0.64 $\pm$ 0.25	68.83 $\pm$ 31.53	0.80 $\pm$ 0.37
	ASBS w/ Harmonic prior (Ours)	70.70 $\pm$ 33.21	0.86 $\pm$ 0.24	52.19 $\pm$ 35.93	1.05 $\pm$ 0.41	86.79 $\pm$ 22.86	0.61 $\pm$ 0.24	70.08 $\pm$ 31.60	0.80 $\pm$ 0.37
	AS +RDKit warmup (Havens et al., 2025)	72.21 $\pm$ 30.22	0.84 $\pm$ 0.24	52.19 $\pm$ 35.20	1.02 $\pm$ 0.34	87.84 $\pm$ 19.20	0.60 $\pm$ 0.23	73.88 $\pm$ 28.63	0.76 $\pm$ 0.34
	ASBS +RDKit warmup (Ours)	74.29 $\pm$ 31.25	0.82 $\pm$ 0.24	55.88 $\pm$ 36.51	0.98 $\pm$ 0.34	87.25 $\pm$ 20.77	0.60 $\pm$ 0.24	74.11 $\pm$ 30.16	0.75 $\pm$ 0.34
Threshold 1.25Å	RDKit ETKDG (Riniker and Landrum, 2015)	72.74 $\pm$ 33.18	1.04 $\pm$ 0.52	63.51 $\pm$ 34.74	1.15 $\pm$ 0.61	81.61 $\pm$ 27.58	0.79 $\pm$ 0.44	71.72 $\pm$ 29.73	0.93 $\pm$ 0.53
	AS (Havens et al., 2025)	82.22 $\pm$ 25.72	0.96 $\pm$ 0.26	60.93 $\pm$ 35.15	1.20 $\pm$ 0.43	94.10 $\pm$ 15.67	0.68 $\pm$ 0.28	79.08 $\pm$ 29.44	0.89 $\pm$ 0.45
	ASBS w/ Gaussian prior (Ours)	87.20 $\pm$ 21.88	0.88 $\pm$ 0.25	70.86 $\pm$ 31.98	1.08 $\pm$ 0.36	95.19 $\pm$ 10.29	0.64 $\pm$ 0.25	84.66 $\pm$ 25.03	0.80 $\pm$ 0.37
	ASBS w/ Harmonic prior (Ours)	89.66 $\pm$ 19.42	0.86 $\pm$ 0.24	74.50 $\pm$ 32.32	1.05 $\pm$ 0.41	96.64 $\pm$ 10.15	0.61 $\pm$ 0.24	83.76 $\pm$ 24.77	0.80 $\pm$ 0.37
	AS +RDKit warmup (Havens et al., 2025)	89.42 $\pm$ 17.48	0.84 $\pm$ 0.24	72.98 $\pm$ 30.82	1.02 $\pm$ 0.34	96.65 $\pm$ 7.51	0.60 $\pm$ 0.23	87.01 $\pm$ 22.79	0.76 $\pm$ 0.34
	ASBS +RDKit warmup (Ours)	90.85 $\pm$ 17.74	0.82 $\pm$ 0.24	77.86 $\pm$ 30.37	0.98 $\pm$ 0.34	97.28 $\pm$ 6.55	0.60 $\pm$ 0.24	87.81 $\pm$ 22.75	0.75 $\pm$ 0.34

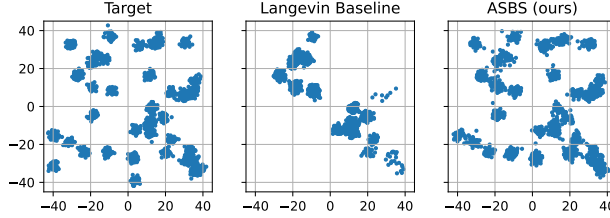


Figure 9: Compared to vanilla Langevin baseline, our ASBS—instantiated with a standard uni-variance Gaussian—is able to identify almost all modes without any prior knowledge of where the target modes were located.

EGNN as AS. Notably, our ASBS consistently outperforms AS on all metrics across all setups, except the coverage for GEOM-DRUGS with relaxation and RDKit warm-start, where ASBS falls slightly behind AS by only 1.0%. Finally, Figure 8 reports the full recall coverage curves that reproduce Table 4.

Ability of ASBS in finding modes We conduct additional experiments on the 40-mode GMM in 2D. Specifically, we instantiate ASBS with a uni-variance Gaussian source distribution centered at zero, effectively assuming no prior knowledge of the target modes, as the initial distribution does not coincide with any target modes. We also run a vanilla Langevin baseline for 1 million steps, starting from the same source distribution.

Figure 9 represents the quantitative results. Notably, ASBS is able to identify almost all modes. In contrast, the vanilla Langevin baseline appears to suffer from a slow mixing rate, recovering less than half of the total modes even after 1 million steps. We highlight this distinction as an advantage of constructing diffusion samplers from the stochastic control and Schrödinger Bridge frameworks, which allows theoretical convergence to target distribution within a finite horizon. Finally, we believe that with proper tuning of the ASBS noise schedule, its performance can be further enhanced.

Discussion on important weights Finally, we discuss the potential integration of ASBS with importance weights, emphasizing that our theoretical and algorithmic frameworks do not preclude the use of importance weights to further enhance performance or robustness.

Formally, the importance weights over model path $X \sim p^u$ admit the following representation:

$$w(X) := \frac{dp^*(X)}{dp^u(X)} = \exp \left(\int_0^1 -\frac{1}{2} \|u_t(X_t)\|^2 dt - \int_0^1 u_t(X_t) \cdot dW_t - \log \frac{\hat{\varphi}_1(X_1)}{\nu(X_1)} + \log \frac{\hat{\varphi}_0(X_0)}{\mu(X_0)} \right), \quad (86)$$

which can be obtained from (59) by setting $\bar{h} := \hat{\varphi}_1$ so that $q^{\bar{h}} = p^*$ is the optimal distribution of SB. Note that when the source distribution degenerates to the Dirac delta $\mu(X_0) = \delta_0(X_0)$, the last term $\log \frac{\hat{\varphi}_0(X_0)}{\mu(X_0)}$ becomes a constant and—as discussed in Section 3.2— $\hat{\varphi}_1 = p_1^{\text{base}}$, thereby recovering the weights used in prior SOC-based methods (Zhang and Chen, 2022; Havens et al., 2025).

Equation (86) is also a more concise representation than the one derived in (Richter and Berner, 2024), by recognizing the following relation through the application of Ito Lemma (46) to $\log \hat{\varphi}_t(X_t)$:

$$\frac{\log \hat{\varphi}_1(X_1)}{\log \hat{\varphi}_0(X_0)} = \int_0^1 \left[\frac{1}{2} \|v_t(X_t)\|^2 + (u_t \cdot v_t)(X_t) + \nabla \cdot (\sigma_t v_t(X_t) - f_t(X_t)) \right] dt + \int_0^1 v_t(X_t) \cdot dW_t, \quad (87)$$

where we shorthand $v_t(x) := \sigma_t \nabla \log \hat{\varphi}_t(x)$.

Estimating the weight in (86) requires knowing the ratios $\frac{\hat{\varphi}_1(x)}{\nu(x)}$ and $\frac{\hat{\varphi}_0(x)}{\mu(x)}$, which are not immediately available with the current parametrization, $u_\theta(t, x) \approx \sigma_t \nabla \log \varphi_t(x)$ and $h_\phi(x) \approx \nabla \log \hat{\varphi}_1(x)$. One accommodation is to reparametrize the functions with potential network $v(t, x) : [0, 1] \times \mathcal{X} \rightarrow \mathbb{R}$,

$$u_\theta(t, x) := \sigma_t \nabla v_\theta(t, x), \quad h_\phi(x) := \nabla v_\phi(1, x) \quad (88)$$

and then regress their gradients onto the adjoint and corrector targets. With that, the logarithmic ratios can be easily estimated:

$$\log \frac{\hat{\varphi}_1(x)}{\nu(x)} = v_\phi(1, x) + E(x), \quad \log \frac{\hat{\varphi}_0(x)}{\mu(x)} \stackrel{(42)}{=} -\log \varphi_0(x) = v_\theta(0, x). \quad (89)$$

A more detailed investigation of this importance sampling scheme is left for future work.