

CONTROLLABLE DIFFUSION VIA OPTIMAL CLASSIFIER GUIDANCE

Anonymous authors

Paper under double-blind review

ABSTRACT

The controllable generation of diffusion models aims to steer the model to generate samples that optimize some given objective functions. It is desirable for a variety of applications including image generation, molecule generation, and DNA/sequence generation. Reinforcement Learning (RL) based fine-tuning of the base model is a popular approach but it can overfit the reward function while requiring significant resources. We frame controllable generation as a problem of finding a distribution that optimizes a KL-regularized objective function. We present Supervised Learning based Controllable Diffusion (SLCD), which iteratively trains a small classifier to guide the generation of the diffusion model. Via a reduction to no-regret online learning analysis, we show that the output from SLCD provably converges to the optimal solution of the KL-regularized objective. Further, we empirically demonstrate that SLCD can generate high quality samples with nearly the same inference time as the base model in both image generation and biological sequence generation.

1 INTRODUCTION

Diffusion models are an expressive class of generative models which are able to model complex data distributions (Song et al., 2021a;b). Recent works have utilized diffusion models for a variety of modalities: images, audio, and molecules (Saharia et al., 2022; Ho et al., 2022; Li et al., 2024; Hooeboom et al., 2022). However, modeling the distribution of data is often not enough for downstream tasks. We want to generate data which satisfies a specific property, be that a prompt, a specific chemical property, or a specific structure.

Perhaps the simplest approach is classifier-guided diffusion where a classifier is trained using a pre-collected labeled dataset. The score of the classifier is used to guide the diffusion process to generate images that have high likelihood being classified under a given label. However this simple approach requires a given labeled dataset and is not directly applicable to the settings where the goal is to optimize a complicated objective function (we call it reward function hereafter). To optimize reward functions, Reinforcement Learning (RL) and stochastic optimal control based approaches have been studied (Black et al., 2024; Oertell et al., 2024; Domingo-Enrich et al., 2025; Clark et al., 2024; Fan et al., 2023; Uehara et al., 2024a;b). These methods require modifying the base model to some degree which can be slow and expensive to train. We instead turn our attention to “fine-tuning free” methods, which do not modify the base model. Such methods Li et al. (2024) rely on test-time scaling in order walk the Pareto frontier of quality vs divergence from the base distribution. It is thus desirable to devise an algorithm with this same property, but pays a fixed (and small) inference cost.

In this work, we ask the following question: *Can we design a fine-tuning-free algorithm that achieves optimal KL-regularized reward maximization while paying only a fixed, small inference cost?* We provide an affirmative answer to this question (Fig. 1). We view fine-tuning diffusion model as a controllable generation problem where we train a guidance model – typically a lightweight small classifier, to guide the pre-trained diffusion model during the inference time. Specifically, we frame the optimization problem as a KL-regularized reward maximization problem where our goal is to optimize the diffusion model to maximize the given reward while staying close to the pre-trained model. Prior work such as SVDD (Li et al., 2024) also studied a similar setting where they also train a guidance model to guide the pre-trained diffusion model in inference time. However SVDD’s solution is sub-optimal, i.e., it does not guarantee to find the optimal solution of the KL-regularized reward maximization objective. We propose a new approach, SLCD (Supervised Learning-based Controllable Diffusion), which **iteratively** refines a classifier using feedback from reward functions

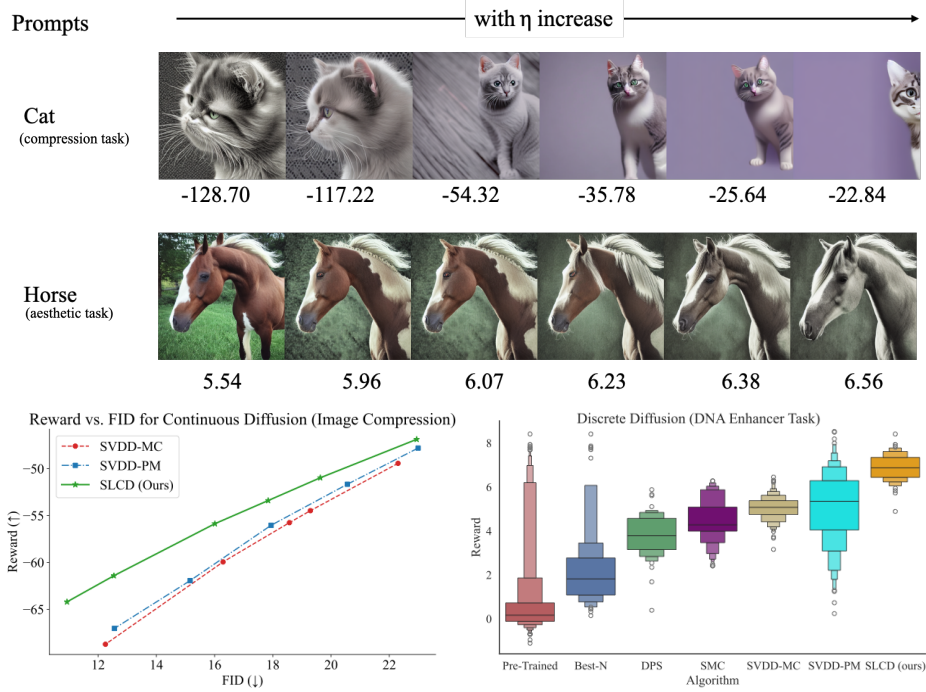


Figure 1: An overview of the main experimental results. **Top:** Qualitative examples for continuous diffusion image tasks (image compression and aesthetic maximization). Relaxing the KL constraint at test time (larger η) consistently increases the score. **Bottom left:** SLCD stays closer to the initial image distribution (lower FID score) for the same reward. **Bottom right:** SLCD is likewise effective at controlling discrete diffusion models.

applied to the **online data** generated by the classifier itself while guiding a pre-trained diffusion model. Conceptually, SLCD is rooted in classifier guidance, but introduces a novel mechanism for learning the **optimal classifier**—one whose guidance provably ensures that the distribution induced on the base model converges to the desired target distribution.

Through a reduction to no-regret learning (Shalev-Shwartz et al., 2012), SLCD finds a near optimal solution to the KL-regularized reward maximization objective. Our analysis is motivated from the classic imitation learning (IL) algorithms DAgger (Ross et al., 2011) and AggreVaTe(d) (Ross & Bagnell, 2014; Sun et al., 2017) which frame IL as an iterative classification procedure with the main computation primitive being classification. Our theory shows that as long as we can achieve no-regret on the sequence of classification losses constructed during the training, the learned classifier can guide the pre-trained diffusion model to generate a near optimal distribution.

On discrete and continuous diffusion tasks we find SLCD consistently outperforms other baselines on reward and inference speed while maintaining a lower divergence from the base model. Overall, SLCD serves a simple solution to the problem of fine-tuning both continuous and discrete diffusion models to optimize a given KL-regularized reward function.

2 RELATED WORK

There has been significant interest in controllable generation of diffusion models, starting from Dhariwal & Nichol (2021) which introduced classifier guidance, to then Ho & Salimans (2022) which introduced classifier-free guidance. These methods paved the way for further interest in controllable generation, in particular when there is an objective function to optimize. First demonstrated by Black et al. (2024); Fan et al. (2023), RL fine-tuning of diffusion models has grown in popularity with works such as Clark et al. (2024); Prabhudesai et al. (2023) which use direct backpropagation to optimize the reward function. However, these methods can lead to mode collapse of the generations and overfitting. Further works then focused on maximizing the KL-constrained optimization problem which regularizes the generation process to the base model (Uehara et al., 2024b) but suffer either from needing special memoryless noise schedulers (Domingo-Enrich et al., 2025) or needing to

control the initial distribution of the diffusion process (Uehara et al., 2024a) to avoid a bias problem. Our approach does not need to do any of these modification. Li et al. (2024) proposed a method to augment the decoding process and avoid training the underlying base model, but yield an increase in compute time. Their practical approach also does not guarantee to learn the optimal distribution. More broadly, Mudgal et al. (2023); Zhou et al. (2025) investigate token-level classifier guidance for KL-regularized controllable generation of large language models. Zhou et al. (2025) also demonstrated that to optimize a KL-regularized RL objective in the context of LLM text generation, one just needs to perform no-regret learning. While our approach is motivated by the Q# approach from Zhou et al. (2025), we tackle score-based guidance for diffusion models in a continuous latent space and continue time, where the space of actions form an infinite set—making algorithm design and analysis substantially more difficult than the setting of discrete token and discrete-time in prior LLM work.

3 PRELIMINARIES

3.1 DIFFUSION MODELS

Given a data distribution q_0 in $\mathbb{R}^d$, the forward process of a diffusion model (Song et al. (2021b)) adds noise to a sample $\bar{\mathbf{x}}_0 \sim q_0$ iteratively, which can be modeled as the solution to a stochastic differential equation (SDE):

$$d\bar{\mathbf{x}} = \mathbf{h}(\bar{\mathbf{x}}, \tau)d\tau + g(\tau)d\bar{\mathbf{w}}, \quad \bar{\mathbf{x}}_0 \sim q_0, \quad \tau \in [0, T] \quad (1)$$

where $\{\bar{\mathbf{w}}_\tau\}_\tau$ is the standard Wiener process, $\mathbf{h}(\cdot, \cdot) : \mathbb{R}^d \times [0, T] \rightarrow \mathbb{R}^d$ is the drift coefficient and $g(\cdot) : [0, T] \rightarrow \mathbb{R}$ is the diffusion coefficient. We use $q_\tau(\cdot)$ to denote the probability density function of $\bar{\mathbf{x}}_\tau$ generated according to the forward SDE in equation equation 1. We assume $\mathbf{f}$ and g satisfy certain conditions s.t. q_T converges to $\mathcal{N}(0, I)$ as $T \rightarrow \infty$. For example, if Eq. (1) is chosen to be Ornstein–Uhlenbeck process, q_T converges to $\mathcal{N}(0, 1)$ exponentially fast.

The forward process equation 1 can be reversed by:

$$d\mathbf{x} = [-\mathbf{h}(\mathbf{x}, T-t) + g^2(T-t)\nabla \log q_{T-t}(\mathbf{x})]dt + g(T-t)d\mathbf{w}, \quad \mathbf{x}_0 \sim q_T, \quad t \in [0, T], \quad (2)$$

where $\{\mathbf{w}_t\}_t$ is the Wiener process. It is known (Anderson, 1982) that the forward process equation 1 and the reverse process equation 2 have the same marginal distributions. To generate a sample, we can sample $x_0 \sim q_T$, and run the above SDE from $t = 0$ to $t = T$ to get x_T . In practice, one can start with $\mathbf{x}_0 \sim \mathcal{N}(0, I)$ (an approximation for q_T) and use numerical SDE solver to approximately generate x_T , such as the generation processes from DDPM (Ho et al., 2020) or DDIM (Song et al., 2021a).

3.2 CONTROLLABLE GENERATION

In certain applications, controllable sample generated from some target conditional distribution is preferable. This can be achieved by adding guidance to the score function. In general, the reverse SDE with guidance $\mathbf{f}(\cdot, \cdot)$ is:

$$d\mathbf{x} = [-\mathbf{h}(\mathbf{x}, T-t) + g^2(T-t)(\nabla \log q_{T-t}(\mathbf{x}) + \mathbf{f}(\mathbf{x}, t))]dt + g(T-t)d\mathbf{w}. \quad (3)$$

For convenience, for all $0 \leq s \leq t \leq T$, we use

$$P_{s \rightarrow t}^{\mathbf{f}}(\cdot|p) \quad (4)$$

to denote the marginal distribution of $\mathbf{x}_t$, the solution to equation 3 with initial conditional $\mathbf{x}_s \sim p$. In the remaining of this paper, we may abuse the notation and use $P_{s \rightarrow t}^{\mathbf{f}}(\cdot|\mathbf{x}')$ to denote a deterministic initial condition $\Pr[\mathbf{x}_s = \mathbf{x}'] = 1$. In particular, we use $P_{s \rightarrow t}^{\text{prior}}(\cdot|p)$ to denote the special case that $\mathbf{f} \equiv 0$.

3.3 REWARD GUIDED GENERATION

In this paper, we aim to generate samples that maximize some reward function $r(x) \in [-R_{\max}, 0]$, while not deviating too much from the base or **prior** distribution q_0 . We consider the setting where we have access to the score function of the prior q_0 (e.g., q_0 can be modeled by a pre-trained large diffusion model).

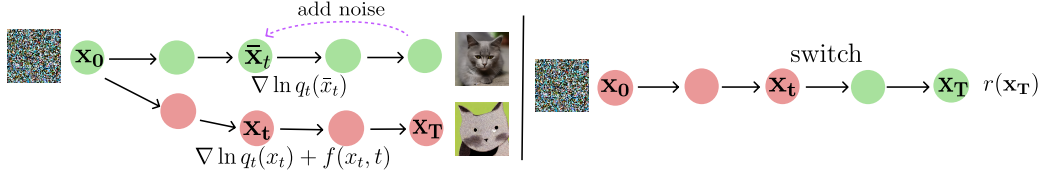


Figure 2: Covariate shift (left) and data collection in our approach (right). The left figure illustrates covariate shift. In the offline naive approach, classifier will be trained under the green samples. However in inference time, the classifier will be applied at the red samples – samples generated by using the classifier itself as guidance. The difference in green samples (training distribution) and red samples (testing distribution) is the covariate shift. Our approach (right) mitigates this by iteratively augmenting the training set with samples drawn from guided diffusion. We rollin with the classifier-guided diffusion to get to $\mathbf{x}_t$. We then rollout with the prior’s score function to get to $\mathbf{x}_T$ and query a reward $r(\mathbf{x}_T)$. The triple $(t, \mathbf{x}_t, r(\mathbf{x}_T))$ will be used to refine the classifier.

Formally, our goal is to find a distribution p that solves the following optimization problem:

$$\max_p \mathbb{E}_{\mathbf{x} \sim p}[r(\mathbf{x})] - \frac{1}{\eta} \text{KL}(p \| q_0) \quad (5)$$

for some $\eta > 0$ which controls the balance between optimizing reward and staying close to the prior. It is known (Ziebart et al., 2008) that p^* , the optimal solution to the optimization in equation 5, satisfies

$$p^*(\mathbf{x}) = \frac{1}{Z} q_0(\mathbf{x}) \exp(\eta r(\mathbf{x})), \quad (6)$$

where $Z > 0$ is the normalization factor. Prior work treats this as a KL-regularized RL optimization problem. However as we mentioned, to ensure the optimal solution of the KL-regularized RL formulation to be equal to p^* , one need to additionally optimize the state distribution q_T (e.g., via another diffusion process), or need to design special memory less noise scheduling (Domingo-Enrich et al., 2025). We aim to show that we can learn p^* in a provably optimal manner.

4 ALGORITHM

We introduce a binary label $y \in \{0, 1\}$ and denote the classifier $p(y = 1 | \mathbf{x}) := \exp(\eta r(\mathbf{x}))$ (recall that we assume reward is negative). The introduction of the binary label and the classifier allows us to rewrite the target distribution as the posterior distribution given $y = 1$:

$$p(\mathbf{x} | y = 1) \propto q_0(\mathbf{x}) p(y = 1 | \mathbf{x}) = q_0(\mathbf{x}) \exp(\eta r(\mathbf{x})).$$

Given this formulation, the naive approach is to model $p(\mathbf{x} | y = 1)$ via the standard classifier-guided diffusion process. In other words, we generate a labeled dataset $\{(\mathbf{x}, y)\}$ where $\mathbf{x}$ is from the prior $\mathbf{x} \sim q_0$, and the label is sampled from the Bernoulli distribution with mean $\exp(\eta r(\mathbf{x}))$, i.e., $y \sim p(y | \mathbf{x})$. Once we have this data, we can add noise to $\mathbf{x}$, train a time-dependent classifier that predicts the noised sample to its label y . Once we train the classifier, we use its gradient to guide the generation process as shown in Eq. (3).

While this naive approach is simple, this approach can fail due to the issue of *covariate shift* – the training distribution (i.e., q_t – the distribution of $\bar{\mathbf{x}}_t$) used for training the classifier is different from the test distribution where the classifier is evaluated during generation (i.e. the distribution of samples $\mathbf{x}$ generated during the classifier-guided denoising process). This is demonstrated in left figure in Fig. 2. In the worst case, the density ratio of the test distribution over the training distribution can be exponential $\exp(R_{\max} \eta)$, which can be too large to ignore when η is large (i.e., KL regularization is weak).

We propose an iterative approach motivated by DAgger (data aggregation, Ross et al. (2011)) to close the gap between the training distribution and test distribution. First with the binary label y and our definition of $p(y = 1 | \mathbf{x}_T) = \exp(\eta r(\mathbf{x}_T))$ (note $\mathbf{x}_T$ represents the generated image), we can show that the classifier $p(y = 1 | \mathbf{x}_t)$ for any $t \in [0, T)$ takes the following form:

$$p(y = 1 | \mathbf{x}_t) = \mathbb{E}_{\mathbf{x}_T \sim P_{t \rightarrow T}^{\text{prior}}(\cdot | \mathbf{x}_t)} \exp(\eta r(\mathbf{x}_T)). \quad (7)$$

Algorithm 1 Controllable diffusion via iterative supervised learning (SLCD)

```

Initialize  $\hat{R}^1$ .
for  $n = 1, \dots, N$  do
  set  $\mathbf{f}^n$  as Eq. (9).
  collect an additional training dataset  $D^n$  following Eq. (10).
  train  $\hat{R}^{n+1}$  on  $\bigcup_{i=1}^n D^i$  according to Eq. (11).
end for
Return  $\mathbf{f}^{\hat{n}}$ , the best of  $\{\mathbf{f}^1, \dots, \mathbf{f}^N\}$  on validation.

```

Intuitively $p(y = 1|\mathbf{x}_t)$ models the expected probability of observing label $y = 1$ if we generate $\mathbf{x}_T$ starting from $\mathbf{x}_t$ using the reverse process of the pre-trained diffusion model. We defer the formal derivation to Appendix A which relies on proving that the forward process and backward process of a diffusion model induce the same conditional distributions.

We take advantage of this closed-form solution of the classifier, and propose to model the classifier via a *distributional approach* (Zhou et al., 2025). Particularly, define $r \sim R^{\text{prior}}(\cdot|\mathbf{x}_t, t)$ as the distribution of the reward of a $\mathbf{x}_T \sim P_{t \rightarrow T}^{\text{prior}}(\cdot|\mathbf{x}_t)$. The classifier $p(y = 1|\mathbf{x}_t)$ can be rewritten using the reward distribution $R^{\text{prior}}(\cdot|\mathbf{x}_t, t)$:

$$p(y = 1|\mathbf{x}_t) := \mathbb{E}_{r \sim R^{\text{prior}}(\cdot|\mathbf{x}_t, t)} \exp(\eta \cdot r). \quad (8)$$

Our goal is to learn a reward distribution $\hat{R}$ to approximate R^{prior} and use $\hat{R}$ to approximate the classifier as $p(y = 1|\mathbf{x}_t) \approx \mathbb{E}_{r \sim \hat{R}(\cdot|\mathbf{x}_t, t)} \exp(\eta r)$. This distributional approach allows us to take advantage of the closed form of the classifier in Eq. (7) (e.g., there is no need to learn the exponential function $\exp(\eta r)$ in the classifier). Algorithm 1 describes an iterative learning approach for training such a distribution $\hat{R}(\cdot|\mathbf{x}_t, t)$ via supervised learning (e.g., maximum likelihood estimation (MLE)).

Inside iteration n , given the current reward distribution $\hat{R}^n$, we define the score of the classifier $\mathbf{f}^n$ as:

$$\forall t, x_t : \mathbf{f}^n(\mathbf{x}_t, t) := \nabla_{\mathbf{x}_t} \ln \left(\mathbb{E}_{r \sim \hat{R}^n(\cdot|\mathbf{x}_t, t)} \exp(\eta \cdot r) \right). \quad (9)$$

We use $\mathbf{f}^n$ to guide the prior to generate an additional training dataset $D^n := \{(t, \mathbf{x}_t, r)\}$ of size M , where

$$\underbrace{t \sim \text{Uniform}(T), \mathbf{x}_0 \sim \mathcal{N}(0, I), \mathbf{x}_t \sim P_{0 \rightarrow t}^{\mathbf{f}^n}(\cdot|\mathbf{x}_0)}_{\text{Roll in with the score of the latest classifier } \mathbf{f}^n \text{ as guidance}}, \underbrace{\mathbf{x}_T \sim P_{t \rightarrow T}^{\text{prior}}(\cdot|\mathbf{x}_t), \text{ and } r = r(\mathbf{x}_T)}_{\text{Roll out with the prior to collect reward}}. \quad (10)$$

Note that the roll-in process above simulates the inference procedure – $\mathbf{x}_t$ is an intermediate sample we would generate if we had used $\mathbf{f}^n$ to guide the prior in inference. The rollout procedure collects reward signals for $\mathbf{x}_t$ which in turn will be used for refining the reward distribution estimator $\hat{R}(\cdot|\mathbf{x}_t)$. This procedure is illustrated in the right figure in Fig. 2. We then aggregate D^n with all the prior data and re-train the distribution estimator $\hat{R}$ using the aggregate data via supervised learning, i.e., maximum likelihood estimation:

$$\hat{R}^{n+1} \in \arg \max_{R \in \mathcal{R}} \sum_{i=1}^n \sum_{(t, \mathbf{x}_t, r) \in D^i} \ln R(r|\mathbf{x}_t, t), \quad (11)$$

where $\mathcal{R}$ is the class of distributions. This rollin-rollout procedure is illustrated in Fig. 2. We iterate the above procedure until we reach a point where $\hat{R}^n(\cdot|\mathbf{x}_t, t)$ is an accurate estimator of the true model $R^{\text{prior}}(\cdot|\mathbf{x}_t, t)$ under distribution induced by the generation process of guiding the prior using $\mathbf{f}^n$ itself. Similar to DAgger’s analysis, we will show in our analysis section that a simple no-regret argument implies that we can reach to such a stage where there is no gap between training and testing distribution anymore.

In the test time, once we have the score $\mathbf{f}^{\hat{n}}$, we can use it to guide the prior to generate samples via the SDE in Eq. (3). In practice, sampling procedure from DDPM can be used as the numerical solver for the SDE Eq. (3). Another practical benefit of our approach is that the definition of the distribution R^{prior} and the learned distribution $\hat{R}$ are independent of the guidance strength parameter η . This means that in practice, once we learned the distribution $\hat{R}$, we can adjust η during inference time as shown in Eq. (9) without re-training $\hat{R}$.

Remark 1 (Modeling the one-dimensional distribution as a classifier). *We emphasize that from a computation perspective, our approach relies on a simple supervised learning oracle to estimate the **one dimensional conditional distribution** $R(r|\mathbf{x}_t, t)$. In our implementation, we use histogram to model this one-dimensional distribution, i.e., we discretize the range of reward $[-R_{\max}, 0]$ into finite number of bins, and use standard **multi-class classification oracle** to learn $\hat{R}$ that maps from $\mathbf{x}_t$ to a distribution over the finite number of labels. Thus, unlike prior work that casts controllable diffusion generation as an RL or stochastic control problem, our approach eliminates the need to talk about or implement RL, and instead entirely relies on standard classification and can be seamlessly integrated with any existing implementation of classifier-guided diffusion.*

Remark 2 (Comparison to SVDD (Li et al., 2024)). *The most related work is SVDD. There are two notifiable differences. First SVDD estimates a sub-optimal classifier, i.e., $\mathbb{E}_{r \sim R^{\text{prior}}(\cdot|\mathbf{x}_t, t)}[r]$. The posterior distribution in their case is proportional to $q_0(\mathbf{x}_T) \cdot \mathbb{E}_{r \sim R^{\text{prior}}(\cdot|\mathbf{x}_T, T)}[r]$ which is clearly not equal to the target distribution. Second, SVDD does not address the issue of distribution shift and it trains the classifier only via offline data collected from the prior alone.*

We note that our method can also be adapted to discrete diffusion tasks as seen in Section 6. We refer the reader to Appendix E for more details.

5 ANALYSIS

In this section, we provide performance guarantee for the sampler returned by Algorithm 1. We use KL divergence of the generated data distribution $P_{0 \rightarrow T}^{\hat{\mathbf{r}}}(\cdot|\mathcal{N}(0, I))$ from the target distribution p^* to measure its quality. At a high level, the error comes from two sources:

- **starting distribution mismatch:** in the sampling process, we initialize the SDE Eq. (3) with samples from $\mathcal{N}(0, I)$, not the ground-truth $q_T(\cdot|y = 1)$. However, under proper condition, $q_T(\cdot|y = 1)$ converges to $\mathcal{N}(0, I)$ as $T \rightarrow \infty$ (see Lemma 3 of Chen et al. (2025)). In particular, when Eq. (3) is chosen to be the OU process, $q_T(\cdot|y = 1)$ converges at an exponential speed: $\text{KL}(\mathcal{N}(0, I) \| q_T(\cdot|y = 1)) = O(e^{-2T})$.
- **estimation error of the guidance:** the estimated guidance $\hat{\mathbf{f}}^{\hat{\mathbf{r}}}$ is different from the ground truth $\nabla_{\mathbf{x}_t} \ln p(y = 1|\mathbf{x}_t)$. But the error is controlled by the regret of the no-regret online learning.

We assume realizability:

Assumption 3 (realizability). $R^{\text{prior}} \in \mathcal{R}$

Our analysis relies on a reduction to no-regret online learning. Particularly, we assume we have no-regret property on the following log-loss. Note M as the side of the each online dataset D^i .

Assumption 4 (No-regret learning). *The sequence of reward distribution $\{\hat{R}^i\}$ satisfies the following inequality:*

$$\frac{1}{NM} \sum_{i=1}^N \sum_{(t, \mathbf{x}_t, r) \in D^i} \ln \frac{1}{\hat{R}^i(r|\mathbf{x}_t, t)} - \min_R \frac{1}{NM} \sum_{i=1}^N \sum_{(t, \mathbf{x}_t, r) \in D^i} \ln \frac{1}{R(r|\mathbf{x}_t, t)} \leq \gamma_N.$$

where the average regret $\gamma_N = o(N)/N$ shrinks to zero when $N \rightarrow \infty$.

No-regret online learning for the log is standard in the literature (Cesa-Bianchi & Lugosi, 2006; Foster et al., 2021; Wang et al., 2024; Zhou et al., 2025). Our algorithm implements the specific no-regret algorithm called Follow-the-regularized-leader (FTRL) (Shalev-Shwartz et al., 2012; Suggala & Netrapalli, 2020) where we optimize for $\hat{R}^i$ on the aggregated dataset. Follow-the-Leader type of approach with random perturbation can even achieve no-regret property for non-convex optimization (Suggala & Netrapalli, 2020). This data aggregation step and the reduction to no-regret online learning closely follows Daggar’s analysis (Ross et al., 2011).

Under certain condition, the marginal distribution q_T defined by the forward SDE Eq. (1) converges to some Gaussian distribution rapidly (see Lemma 3 of Chen et al. (2025)). For simplicity, we make the following assumption on the convergence:

Assumption 5 (convergence of the forward process). $\text{KL}(\mathcal{N}(0, I) \| q_T(\cdot|y = 1)) \leq \epsilon_T$.

For OU processes, ϵ_T shrinks in the rate of $\exp(-T)$. We assume the reward distribution class satisfies certain regularity conditions, s.t. the estimation error of the classifier controls the score difference:

Assumption 6. *There exists $L > 0$, s.t. for all $R, R' \in \mathcal{R}$, and $\mathbf{x}, t$:*

$$\begin{aligned} & \left\| \nabla_{\mathbf{x}} \ln \left(\mathbb{E}_{r \sim R(\cdot|\mathbf{x},t)} \exp(\eta \cdot r) \right) - \nabla_{\mathbf{x}} \ln \left(\mathbb{E}_{r \sim R'(\cdot|\mathbf{x},t)} \exp(\eta \cdot r) \right) \right\|_2 \\ & \leq L \left| \mathbb{E}_{r \sim R(\cdot|\mathbf{x},t)} \exp(\eta \cdot r) - \mathbb{E}_{r \sim R'(\cdot|\mathbf{x},t)} \exp(\eta \cdot r) \right|. \end{aligned}$$

Standard diffusion models with classifier guidance train a time-dependent classifier and use the score function of the classifier to control image generation (Song et al., 2021b; Dhariwal & Nichol, 2021). Such assumption is crucial to guarantee the quality of class-conditional sample generation. We defer a more detailed discussion on Assumption 6 to Appendix B. In general such an assumption holds when the functions satisfy certain smoothness conditions.

Theorem 7. *Suppose Assumption 3, 4, 5, and 6 hold. There exists $\hat{n} \in \{1, \dots, N\}$, s.t. $\mathbf{f}^{\hat{n}}$ specified by Algorithm 1 satisfies:*

$$\mathbb{E} \left[D_{\text{KL}} \left(P_{0 \rightarrow T}^{\mathbf{f}^{\hat{n}}}(\cdot|\mathcal{N}(0, I)) \| p^* \right) \right] \leq \epsilon_T + \frac{1}{2} T \|g\|_{\infty}^2 L^2 \gamma_N.$$

where the expectation is with respect to the randomness in the whole training process, and g is the diffusion coefficient defined in Eq. (1).

Since $P_{0 \rightarrow T}^{\mathbf{f}^{\hat{n}}}(\cdot|\mathcal{N}(0, I))$ models the distribution of the generated samples when using $\mathbf{f}^{\hat{n}}$ to guide the prior, the above theorem proves that our sampling distribution is close to the target p^* under KL. Note that ϵ_T decays in the rate of $\exp(-T)$ when the forward SDE is an OU process, and γ_N decays in the rate of $1/\sqrt{N}$ for a typical no-regret algorithm such as Follow-the-Leader (Shalev-Shwartz et al., 2012; Suggala & Netrapalli, 2020).

6 EXPERIMENTS

We compare SLCD to a variety of training-free and value-guided sampling strategies across four tasks. For Best-of-N, we draw N independent samples from the base diffusion model and keep the one with the highest reward. Diffusion Posterior Sampling (DPS) is a classifier-guidance variant originally for continuous diffusion (Chung et al., 2023), here adapted to discrete diffusion via the state-of-the-art method of Nisonoff et al. (2025). Sequential Monte Carlo (SMC) methods (Del Moral & Doucet, 2014; Wu et al., 2023; Trippe et al., 2022) use importance sampling across whole batches to select the best sample. SVDD-MC (Li et al., 2024) instead evaluates the expected reward of N candidates under an estimated value function, while SVDD-PM uses the true reward for each candidate with slight algorithm modifications. We evaluate on (i) image compression (negative file size) and (ii) image aesthetics (LAION aesthetic score) using Stable Diffusion v1.5 (Rombach et al., 2022), as well as on (iii) 5' untranslated regions optimized for mean ribosome load (Sample et al., 2019; Sahoo et al., 2024) and (iv) DNA enhancer sequences optimized for predicted expression in HepG2 cells via the Enformer model (Avsec et al., 2021).

In line with Li et al. (2024), we compare the top 10 and 50 quantiles of a batch of generations, in Table 1. We compare to these methods as, like SLCD, all of these methods do not require training of the base model. Overall, we see that SLCD consistently outperforms the baseline methods while requiring nearly the same inference time as the base model, and omitting the need for multiple MC samples during each diffusion step. These four tasks jointly cover two primary application domains of diffusion models: image generation and biological sequence generation, providing a comprehensive assessment of controllable diffusion methods.

6.1 REWARD COMPARISON

We compare the reward of SLCD to the baseline methods in Table 1. We see that SLCD is able to consistently achieve higher reward than SVDD-MC and SVDD-PM, and the other baseline methods in all four tasks. The margin of improvement is most pronounced in settings where the classifier closely approximates the true reward, most notably the image compression task, where SLCD nearly attains the optimal reward. To further see the performance of SLCD, we plot the reward distribution of SLCD and the baseline methods in Fig. 3 and Fig. 1. We observe that SLCD produces a more

Table 1: Top 10 and 50 quantiles of the generated samples for each algorithm (with 95% confidence intervals). Higher is better. SLCD consistently outperforms the baseline methods. Baseline results taken from Li et al. (2024) as the exact same settings were used.

Domain	Quantile	Pre-Train	Best-N	DPS	SMC	SVDD-MC	SVDD-PM	SLCD
Image: Compress	50%	-101.4 ± 0.22	-71.2 ± 0.46	-60.1 ± 0.44	-59.7 ± 0.4	-54.3 ± 0.33	-51.1 ± 0.38	-13.60 ± 0.79
	10%	-78.6 ± 0.13	-57.3 ± 0.28	-61.2 ± 0.28	-49.9 ± 0.24	-40.4 ± 0.2	-38.8 ± 0.23	-11.05 ± 0.41
Image: Aesthetic	50%	5.62 ± 0.003	6.11 ± 0.007	5.61 ± 0.009	6.02 ± 0.004	5.70 ± 0.008	6.14 ± 0.007	6.31 ± 0.061
	10%	5.98 ± 0.002	6.34 ± 0.004	6.00 ± 0.005	6.28 ± 0.003	6.05 ± 0.005	6.47 ± 0.004	6.59 ± 0.077
Enhancers (DNA)	50%	0.121 ± 0.033	1.807 ± 0.214	3.782 ± 0.299	4.28 ± 0.02	5.074 ± 0.096	5.353 ± 0.231	7.403 ± 0.125
	10%	1.396 ± 0.020	3.449 ± 0.128	4.879 ± 0.179	5.95 ± 0.01	5.639 ± 0.057	6.980 ± 0.138	7.885 ± 0.231
5'UTR (RNA)	50%	0.406 ± 0.028	0.912 ± 0.023	0.426 ± 0.073	0.76 ± 0.02	1.042 ± 0.008	1.214 ± 0.016	1.313 ± 0.024
	10%	0.869 ± 0.017	1.064 ± 0.014	0.981 ± 0.044	0.91 ± 0.01	1.117 ± 0.005	1.383 ± 0.010	1.421 ± 0.039

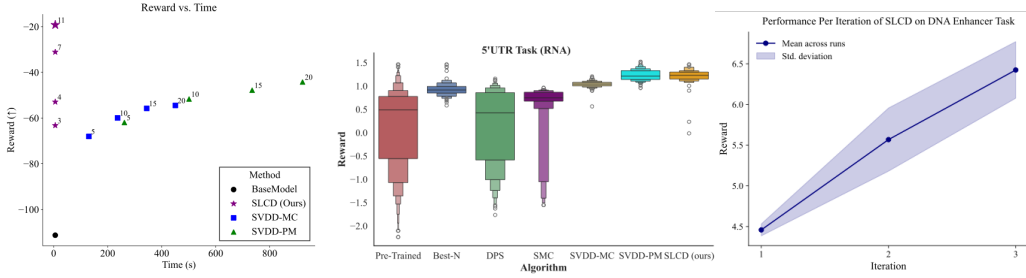


Figure 3: **(Left)** Reward vs. Inference Time on the Compression Task. Numeric labels on SLCD indicate η , while those on the SVDD denote the duplication number applied at each step. **(Center)** Distribution of rewards for DNA sequences (Enhancers) across different methods. SLCD demonstrates superior performance with higher median and maximum rewards. **(Right)** Reward vs. number of iterations of SLCD. The reward increases as the restart state distribution becomes richer.

tightly concentrated reward distribution with a higher median reward than the baseline methods, while still maintaining generation diversity, as shown in Fig. 1 with a lower FID score than baseline methods.

6.2 QUALITATIVE RESULTS

We present generated images from SLCD in Fig. 4. For the compression task, we observe three recurring patterns: some images shift the subject toward the edges of the frame, others reduce the subject’s size, and some simplify the overall scene to reduce file size. For the aesthetic task, the outputs tend to take on a more illustrated appearance, often reflecting a variety of artistic styles.

As η increases, the KL constraint is relaxed, enabling a controlled trade-off between optimizing for the reward function and staying close to the base model’s distribution. Notably, even under strong reward guidance (i.e., with larger η), our method consistently maintains a high level of diversity in the generated outputs.

6.3 FRÉCHET INCEPTION DISTANCE COMPARISON

Both SLCD and SVDD baselines allow one to control the output sample reward at test time, but via different control variables: SLCD modulates the KL-penalty coefficient η , while SVDD-MC and SVDD-PM vary the number of Monte Carlo rollouts evaluated at each diffusion step. Since these control parameters can affect sample quality in different ways, we report both the Fréchet Inception Distance (Heusel et al., 2017) (FID) and reward. For the same reward, *higher* FID indicates that the model generates images stray farther from the base models distribution, a sign of reward hacking. We evaluate these methods in Fig. 1. That is, the points on the curve form a pareto frontier between reward and FID. SLCD is able to achieve a better reward-FID trade-off than SVDD-MC and SVDD-PM.

6.4 INFERENCE TIME COMPARISON

An additional advantage of SLCD is its negligible inference overhead at test time, even when using higher η values to achieve greater rewards. In Fig. 3, we compare the wall-clock generation time per

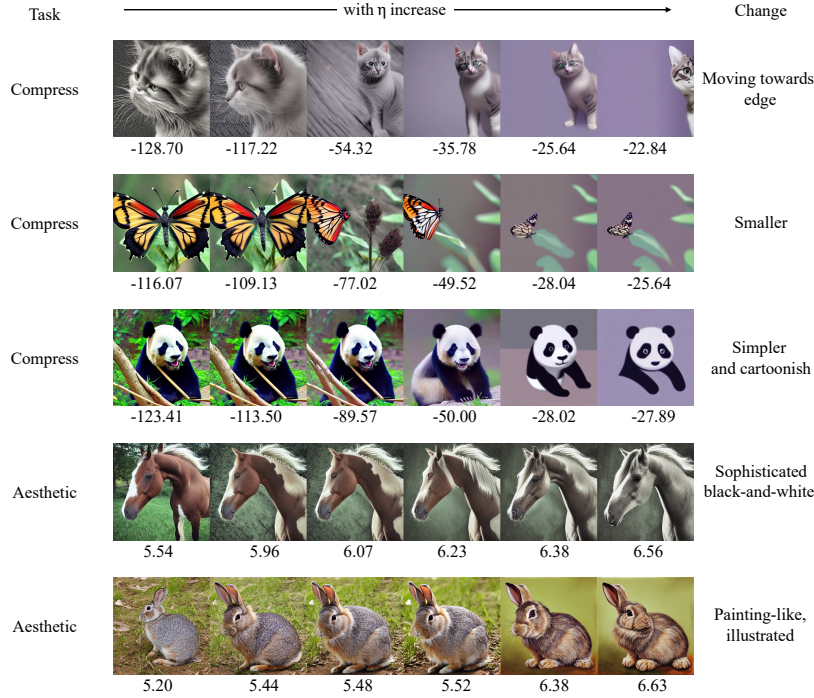


Figure 4: Images generated by SLCD with varying η values and their rewards. The first column shows results from the base model, SD1.5, which corresponds to our method with $\eta = 0$. As η increases, the KL penalty is relaxed, allowing the generated images to be more strongly optimized for the reward function, and consequently, they diverge further from the base model’s original distribution.

image on an NVIDIA A6000 GPU for SLCD against SVDD-MC and SVDD-PM. SLCD achieves higher rewards while requiring significantly fewer computational resources and substantially shorter inference times than either baseline. Specifically, SLCD takes only 6.06 seconds per image, nearly identical to the base model SD 1.5’s 5.99 seconds.

Importantly, unlike SVDD methods that incur increased computational cost to improve rewards, SLCD maintains constant inference time across all η values, achieving enhanced performance with no additional computation.

6.5 ABLATION STUDY

To elucidate the impact of each training cycle, we vary the number of SLCD iterations and plot the resulting reward in Fig. 3. As additional iterations enrich the state distribution and mitigate covariate shift for the classifier, the reward consistently rises. This confirms that our iterative approach can mitigate covariate shift issue. In practice, we only require a small number of iterations to achieve high reward (for example, 3 for the compression task)

Because the scaling parameter η can be chosen at test time when the distribution $\hat{R}$ is fully trained, SLCD enables test-time control over the KL penalty during inference. By modulating η at test-time, practitioners can smoothly trade-off reward against sample quality without retraining the distribution $\hat{R}$, as demonstrated in Fig. 4.

7 CONCLUSION

In this work, we introduced SLCD, a novel and efficient method that recasts the KL-constrained optimization problem as a supervised learning task. We provided theoretical guarantees showing that SLCD converges to the optimal KL-constrained solution and how data-aggregation effectively mitigates covariate shift. Empirical evaluations confirm that SLCD surpasses existing approaches, all while preserving high fidelity to the base model’s outputs and maintaining nearly the same inference time.

REFERENCES

- Brian DO Anderson. Reverse-time diffusion equation models. *Stochastic Processes and their Applications*, 12(3):313–326, 1982.
- Žiga Avsec, Vikram Agarwal, Daniel Visentin, Joseph R. Ledsam, Agnieszka Grabska-Barwinska, Kyle R. Taylor, Yannis Assael, John Jumper, Pushmeet Kohli, and David R. Kelley. Effective gene expression prediction from sequence by integrating long-range interactions. *Nature Methods*, 2021. doi: 10.1038/s41592-021-01252-x.
- Kevin Black, Michael Janner, Yilun Du, Ilya Kostrikov, and Sergey Levine. Training diffusion models with reinforcement learning, 2024.
- Nicolo Cesa-Bianchi and Gábor Lugosi. *Prediction, learning, and games*. Cambridge university press, 2006.
- Hongrui Chen, Holden Lee, and Jianfeng Lu. Improved analysis of score-based generative modeling: User-friendly bounds under minimal smoothness assumptions. In *International Conference on Machine Learning*, pp. 4735–4763. PMLR, 2023.
- Yiding Chen, Yiyi Zhang, Owen Oertell, and Wen Sun. Convergence of consistency model with multistep sampling under general data assumptions. *arXiv preprint arXiv:2505.03194*, 2025.
- Hyungjin Chung, Jeongsol Kim, Michael Thompson McCann, Marc Louis Klasky, and Jong Chul Ye. Diffusion posterior sampling for general noisy inverse problems. *ICLR*, 2023. URL <https://dblp.org/rec/conf/iclr/ChungKMKY23>.
- Kevin Clark, Paul Vicol, Kevin Swersky, and David J. Fleet. Directly fine-tuning diffusion models on differentiable rewards. *ICLR*, 2024. URL <https://dblp.org/rec/conf/iclr/ClarkVSF24>.
- Piere Del Moral and Arnaud Doucet. Particle methods: An introduction with applications. In *ESAIM: proceedings*, volume 44, pp. 1–46. EDP Sciences, 2014.
- Rahul Dey and Fathi M. Salem. Gate-variants of gated recurrent unit (gru) neural networks. *arXiv*, 2017.
- Prafulla Dhariwal and Alex Nichol. Diffusion models beat gans on image synthesis, 2021. URL <https://arxiv.org/abs/2105.05233>.
- 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. *ICLR*, 2025. URL <https://dblp.org/rec/conf/iclr/Domingo-EnrichD25>.
- Ying Fan, Olivia Watkins, Yuqing Du, Hao Liu, Moonkyung Ryu, Craig Boutilier, Pieter Abbeel, Mohammad Ghavamzadeh, Kangwook Lee, and Kimin Lee. Dpok: Reinforcement learning for fine-tuning text-to-image diffusion models. *arXiv preprint arXiv:2305.16381*, 2023.
- Abraham D Flaxman, Adam Tauman Kalai, and H Brendan McMahan. Online convex optimization in the bandit setting: gradient descent without a gradient. *arXiv preprint cs/0408007*, 2004.
- Dylan J Foster, Sham M Kakade, Jian Qian, and Alexander Rakhlin. The statistical complexity of interactive decision making. *arXiv preprint arXiv:2112.13487*, 2021.
- Martin Heusel, Hubert Ramsauer, Thomas Unterthiner, Bernhard Nessler, and Sepp Hochreiter. Gans trained by a two time-scale update rule converge to a local nash equilibrium. *Advances in neural information processing systems*, 30, 2017.
- Jonathan Ho and Tim Salimans. Classifier-free diffusion guidance. *CoRR*, 2022. doi: 10.48550/ARXIV.2207.12598.
- Jonathan Ho, Ajay Jain, and Pieter Abbeel. Denoising diffusion probabilistic models. *Advances in neural information processing systems*, 33:6840–6851, 2020.

- Jonathan Ho, William Chan, Chitwan Saharia, Jay Whang, Ruiqi Gao, Alexey Gritsenko, Diederik P Kingma, Ben Poole, Mohammad Norouzi, David J Fleet, et al. Imagen video: High definition video generation with diffusion models. *arXiv preprint arXiv:2210.02303*, 2022.
- Emiel Hoogetboom, Victor Garcia Satorras, Clément Vignac, and Max Welling. Equivariant diffusion for molecule generation in 3d. *ICML*, 2022. URL <https://dblp.org/rec/conf/icml/HoogetboomSVW22>.
- Xiner Li, Yulai Zhao, Chenyu Wang, Gabriele Scalia, Gokcen Eraslan, Surag Nair, Tommaso Biancalani, Shuiwang Ji, Aviv Regev, Sergey Levine, and Masatoshi Uehara. Derivative-free guidance in continuous and discrete diffusion models with soft value-based decoding, 2024. URL <https://arxiv.org/abs/2408.08252>.
- Sidharth Mudgal, Jong Lee, Harish Ganapathy, YaGuang Li, Tao Wang, Yanping Huang, Zhifeng Chen, Heng-Tze Cheng, Michael Collins, Trevor Strohman, et al. Controlled decoding from language models. *arXiv preprint arXiv:2310.17022*, 2023.
- Hunter Nisonoff, Junhao Xiong, Stephan Allenspach, and Jennifer Listgarten. Unlocking guidance for discrete state-space diffusion and flow models. *ICLR*, 2025. URL <https://dblp.org/rec/conf/iclr/NisonoffXAL25>.
- Owen Oertell, Jonathan D. Chang, Yiyi Zhang, Kianté Brantley, and Wen Sun. RL for consistency models: Faster reward guided text-to-image generation, 2024. URL <https://arxiv.org/abs/2404.03673>.
- Mihir Prabhudesai, Anirudh Goyal, Deepak Pathak, and Katerina Fragkiadaki. Aligning text-to-image diffusion models with reward backpropagation, 2023.
- Robin Rombach, Andreas Blattmann, Dominik Lorenz, Patrick Esser, and Björn Ommer. High-resolution image synthesis with latent diffusion models. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)*, pp. 10684–10695, June 2022.
- Stephane Ross and J Andrew Bagnell. Reinforcement and imitation learning via interactive no-regret learning. *arXiv preprint arXiv:1406.5979*, 2014.
- Stéphane Ross, Geoffrey Gordon, and Drew Bagnell. A reduction of imitation learning and structured prediction to no-regret online learning. In *Proceedings of the fourteenth international conference on artificial intelligence and statistics*, pp. 627–635. JMLR Workshop and Conference Proceedings, 2011.
- Chitwan Saharia, William Chan, Saurabh Saxena, Lala Li, Jay Whang, Emily L Denton, Kamyar Ghasemipour, Raphael Gontijo Lopes, Burcu Karagol Ayan, Tim Salimans, et al. Photorealistic text-to-image diffusion models with deep language understanding. *Advances in Neural Information Processing Systems*, 35:36479–36494, 2022.
- Subham S. Sahoo, Marianne Arriola, Yair Schiff, Aaron Gokaslan, Edgar Marroquin, Justin T. Chiu, Alexander Rush, and Volodymyr Kuleshov. Simple and effective masked diffusion language models. *NeurIPS*, 2024. URL <https://dblp.org/rec/conf/nips/SahooASGMCRK24>.
- Paul J. Sample, Ban Wang, David W. Reid, Vlad Presnyak, Iain J. McFadyen, David R. Morris, and Georg Seelig. Human 5’ utr design and variant effect prediction from a massively parallel translation assay. *Nature Biotechnology*, 2019. doi: 10.1038/s41587-019-0164-5.
- Chrisoph Schumman. Laion aesthetics. <https://laion.ai/blog/laion-aesthetics/>, 2022.
- Shai Shalev-Shwartz et al. Online learning and online convex optimization. *Foundations and Trends® in Machine Learning*, 4(2):107–194, 2012.
- Jiaming Song, Chenlin Meng, and Stefano Ermon. Denoising diffusion implicit models. In *International Conference on Learning Representations*, 2021a. URL <https://openreview.net/forum?id=StlgiaRCHLP>.

- Yang Song, Jascha Sohl-Dickstein, Diederik P Kingma, Abhishek Kumar, Stefano Ermon, and Ben Poole. Score-based generative modeling through stochastic differential equations. In *Proceedings of International Conference on Learning Representations*, 2021b.
- Arun Sai Suggala and Praneeth Netrapalli. Online non-convex learning: Following the perturbed leader is optimal. In *Algorithmic Learning Theory*, pp. 845–861. PMLR, 2020.
- Wen Sun, Arun Venkatraman, Geoffrey J Gordon, Byron Boots, and J Andrew Bagnell. Deeply aggravated: Differentiable imitation learning for sequential prediction. In *International conference on machine learning*, pp. 3309–3318. PMLR, 2017.
- Brian L Trippe, Jason Yim, Doug Tischer, David Baker, Tamara Broderick, Regina Barzilay, and Tommi Jaakkola. Diffusion probabilistic modeling of protein backbones in 3d for the motif-scaffolding problem. *arXiv preprint arXiv:2206.04119*, 2022.
- Masatoshi Uehara, Yulai Zhao, Kevin Black, Ehsan Hajiramezanali, Gabriele Scalia, Nathaniel Lee Diamant, Alex M Tseng, Tommaso Biancalani, and Sergey Levine. Fine-tuning of continuous-time diffusion models as entropy-regularized control, 2024a.
- Masatoshi Uehara, Yulai Zhao, Kevin Black, Ehsan Hajiramezanali, Gabriele Scalia, Nathaniel Lee Diamant, Alex M Tseng, Sergey Levine, and Tommaso Biancalani. Feedback efficient online fine-tuning of diffusion models. *arXiv preprint arXiv:2402.16359*, 2024b.
- Kaiwen Wang, Nathan Kallus, and Wen Sun. The central role of the loss function in reinforcement learning. *arXiv preprint arXiv:2409.12799*, 2024.
- Luhuan Wu, Brian Trippe, Christian Naesseth, David Blei, and John P Cunningham. Practical and asymptotically exact conditional sampling in diffusion models. *Advances in Neural Information Processing Systems*, 36:31372–31403, 2023.
- Jin Peng Zhou, Kaiwen Wang, Jonathan Chang, Zhaolin Gao, Nathan Kallus, Kilian Q Weinberger, Kianté Brantley, and Wen Sun. Q $\sharp$: Provably optimal distributional rl for llm post-training. *arXiv preprint arXiv:2502.20548*, 2025.
- Brian D Ziebart, Andrew L Maas, J Andrew Bagnell, Anind K Dey, et al. Maximum entropy inverse reinforcement learning. In *Aaai*, volume 8, pp. 1433–1438. Chicago, IL, USA, 2008.

We show the derivation for the classifier in [Appendix A](#). We discuss [Theorem 6](#) in [Appendix B](#). In [Appendix C](#), we provide the proof for the main theorem in the main text. Then, we provide additional technical lemmas in [Appendix D](#). Finally, we provide more details on the training and evaluation of SLCD in [Appendix E](#). We finally provide more samples of the image experiments in [Appendix F](#).

A DERIVATION FOR THE CLASSIFIER

Our goal is to show [Eq. \(7\)](#):

$$p(y = 1|\mathbf{x}_t) = \mathbb{E}_{\mathbf{x}_T \sim P_{t \rightarrow T}^{\text{prior}}(\cdot|\mathbf{x}_t)} \exp(\eta r(\mathbf{x}_T)).$$

Our first step is to derive the classifier in terms of $\{\bar{\mathbf{x}}_\tau\}_\tau$, the forward process defined by [Eq. \(1\)](#). By the law of total probability:

$$\begin{aligned} p(y = 1|\bar{\mathbf{x}}_\tau) &= \frac{p(y = 1, \bar{\mathbf{x}}_\tau)}{q(\bar{\mathbf{x}}_\tau)} = \frac{\int p(y = 1, \bar{\mathbf{x}}_\tau, \bar{\mathbf{x}}_0) d\bar{\mathbf{x}}_0}{q(\bar{\mathbf{x}}_\tau)} = \frac{\int p(y = 1|\bar{\mathbf{x}}_\tau, \bar{\mathbf{x}}_0) q(\bar{\mathbf{x}}_\tau, \bar{\mathbf{x}}_0) d\bar{\mathbf{x}}_0}{q(\bar{\mathbf{x}}_\tau)} \\ &= \int p(y = 1|\bar{\mathbf{x}}_\tau, \bar{\mathbf{x}}_0) q(\bar{\mathbf{x}}_0|\bar{\mathbf{x}}_\tau) d\bar{\mathbf{x}}_0. \end{aligned}$$

By definition, the label y and $\bar{\mathbf{x}}_\tau$ are independent when conditioning on $\bar{\mathbf{x}}_0$, thus $p(y = 1|\bar{\mathbf{x}}_\tau, \bar{\mathbf{x}}_0) = p(y = 1|\bar{\mathbf{x}}_0)$ and

$$p(y = 1|\bar{\mathbf{x}}_\tau) = \mathbb{E}_{\bar{\mathbf{x}}_0|\bar{\mathbf{x}}_\tau} [p(y = 1|\bar{\mathbf{x}}_0)] = \mathbb{E}_{\bar{\mathbf{x}}_0|\bar{\mathbf{x}}_\tau} [\exp(\eta r(\bar{\mathbf{x}}_0))]. \quad (12)$$

By [Lemma 12](#), the forward ([Eq. \(1\)](#)) and reverse ([Eq. \(2\)](#)) processes have the same conditional distribution. Thus, in [Eq. \(12\)](#), we can substitute $\bar{\mathbf{x}}_0, \bar{\mathbf{x}}_\tau, q(\bar{\mathbf{x}}_0|\bar{\mathbf{x}}_\tau)$ with the corresponding components in the reverse process: $\mathbf{x}_T, \mathbf{x}_{T-\tau}, P_{T-\tau \rightarrow T}^{\text{prior}}(\mathbf{x}_T|\mathbf{x}_{T-\tau})$. Then we complete the proof by setting $\tau = T - t$.

One crucial property is: [Eq. \(12\)](#), the classifier defined in terms of the forward process, only depends on the conditional distribution $\bar{\mathbf{x}}_0|\bar{\mathbf{x}}_\tau$, not the joint distribution of $(\bar{\mathbf{x}}_0, \bar{\mathbf{x}}_\tau)$. This means [Eq. \(7\)](#), the classifier defined in terms of the reverse process, is accurate regardless of the marginal distribution of $\mathbf{x}_t$.

In fact, when considering the data collection procedure [Eq. \(10\)](#), the roll-in step involves the classifier estimator during training, so the marginal distribution of $\mathbf{x}_t$ can be very different from q_{T-t} , the marginal distribution of the forward process. However, as long as the roll-out step uses the ground truth reverse process, we are using the correct conditional distribution and the target classifier during training is thus unbiased.

B DISCUSSION ON ASSUMPTION 6

[Assumption 6](#) assumes that the estimation for the classifier results in good gradient estimation with a small pointwise error. In this section, we show two sets of sufficient conditions that weaken [Assumption 6](#) to versions where the estimation errors are evaluated on the marginal distribution of the reverse process. The conditions will depend on:

- $\hat{R}^{\hat{n}}$: the reward distribution chosen by [Algorithm 1](#);
- R^{prior} : the ground truth reward;
- $\hat{p}$: the marginal distribution in the reverse process induced by guidance $\hat{f}$ (returned by [Algorithm 1](#)).

For simplicity, let

$$\begin{aligned} \hat{v}(\mathbf{x}, t) &:= \mathbb{E}_{r \sim \hat{R}^{\hat{n}}(\cdot|\mathbf{x}, t)} \exp(\eta \cdot r) \\ v^*(\mathbf{x}, t) &:= \mathbb{E}_{r \sim R^{\text{prior}}(\cdot|\mathbf{x}, t)} \exp(\eta \cdot r) \\ V(\mathbf{x}, t) &:= \hat{v}(\mathbf{x}, t) - v^*(\mathbf{x}, t). \end{aligned}$$

Throughout this section, we discuss some natural conditions under which small error from estimating $v^*(\mathbf{x}, t)$ leads to small error on estimating $\nabla_{\mathbf{x}} \ln v^*(\mathbf{x}, t)$. The conditions are related to the smoothness of functions $V, v^*, \hat{v}$, and the distribution p .

Firstly, we present the following lemma, showing that it's sufficient to control $\nabla f_1(\mathbf{x}) - \nabla f_2(\mathbf{x})$:

Lemma 8. *Let f_1, f_2 be $\mathbb{R}^d \rightarrow \mathbb{R}$, then*

$$\begin{aligned} & \|\nabla \ln f_1(\mathbf{x}) - \nabla \ln f_2(\mathbf{x})\|_2 \\ & \leq \frac{1}{\inf_{\mathbf{x}} |f_1(\mathbf{x})|} \|\nabla f_1(\mathbf{x}) - \nabla f_2(\mathbf{x})\|_2 + \frac{\sup_{\mathbf{x}} \|\nabla f_2(\mathbf{x})\|_2}{\inf_{\mathbf{x}} |f_1(\mathbf{x})| \inf_{\mathbf{x}} |f_2(\mathbf{x})|} |f_2(\mathbf{x}) - f_1(\mathbf{x})| \end{aligned}$$

B.1 SMOOTHNESS ASSUMPTIONS ON BOTH THE CLASSIFIER AND THE DISTRIBUTION

Our first set of assumptions is based on the smoothness of the functions and the smoothness of the distributions: there exists $M, L, L_p > 0$, s.t.

- for all $t, \mathbb{E}_{\mathbf{x} \sim \hat{p}_t} \left[\left(\sum_{i=1}^d \frac{\partial^2 V(\mathbf{x}, t)}{\partial x_i^2} \right)^2 \right] \leq M^2$;
- for all $t, \sup_{\mathbf{x}} \|\nabla_{\mathbf{x}} \hat{v}(\mathbf{x}, t)\|_2, \sup_{\mathbf{x}} \|\nabla_{\mathbf{x}} v^*(\mathbf{x}, t)\|_2 \leq L$;
- for all $t, \mathbb{E}_{\mathbf{x} \sim p} [\|\nabla \log p(\mathbf{x})\|_2^2] \leq L_p^2$.

We now present the following results that relates the difference of gradients to difference of function value:

Lemma 9. *Let $F(\cdot)$ be function that map from $\mathbb{R}^d \rightarrow \mathbb{R}$, and $p(\cdot)$ be a PDF over $\mathbb{R}^d$. Suppose $\lim_{x_i \rightarrow \infty} F(\mathbf{x}) \cdot \frac{\partial F(\mathbf{x})}{\partial x_i} p(\mathbf{x}) = 0$ for all i , and $\mathbf{x}_{-i}$. Then $\mathbb{E}_{\mathbf{x} \sim p} \|\nabla F(\mathbf{x})\|_2^2$ is bounded by:*

$$\left(\sqrt{\mathbb{E}_{\mathbf{x} \sim p} \left[\left(\sum_{i=1}^d \frac{\partial^2 F(\mathbf{x})}{\partial x_i^2} \right)^2 \right]} + \sup_{\mathbf{x}} (\|\nabla F(\mathbf{x})\|_2) \sqrt{\mathbb{E}_{\mathbf{x} \sim p} [\|\nabla \log p(\mathbf{x})\|_2^2]} \right) \sqrt{\mathbb{E}_{\mathbf{x} \sim p} [F^2(\mathbf{x})]}.$$

Given these, we can show a version of [Theorem 6](#). By definition, for all t ,

$$\inf_{\mathbf{x}} |\hat{v}(\mathbf{x}, t)|, \inf_{\mathbf{x}} |v^*(\mathbf{x}, t)| \geq \exp(-\eta R_{\max}).$$

By [Lemma 8](#) and [Lemma 9](#): for all t .

$$\begin{aligned} & \mathbb{E}_{\mathbf{x} \sim \hat{p}_t} [\|\nabla_{\mathbf{x}} \ln \hat{v}(\mathbf{x}, t) - \nabla_{\mathbf{x}} \ln v^*(\mathbf{x}, t)\|_2^2] \\ & \leq 2 \exp(2\eta R_{\max}) \mathbb{E}_{\mathbf{x} \sim \hat{p}_t} [\|\nabla_{\mathbf{x}} \hat{v}(\mathbf{x}, t) - \nabla_{\mathbf{x}} v^*(\mathbf{x}, t)\|_2^2] + 2 \exp(4\eta R_{\max}) L \mathbb{E}_{\mathbf{x} \sim \hat{p}_t} [|\hat{v}(\mathbf{x}, t) - v^*(\mathbf{x}, t)|^2] \\ & \leq 2 \exp(2\eta R_{\max}) (M + L \cdot L_p) \sqrt{\mathbb{E}_{\mathbf{x} \sim \hat{p}_t} [|\hat{v}(\mathbf{x}, t) - v^*(\mathbf{x}, t)|^2]} \\ & \quad + 2 \exp(4\eta R_{\max}) L \mathbb{E}_{\mathbf{x} \sim \hat{p}_t} [|\hat{v}(\mathbf{x}, t) - v^*(\mathbf{x}, t)|^2] \end{aligned}$$

This will only change the proof in [Appendix C](#) slightly.

B.2 SMOOTHNESS ASSUMPTION AND GRADIENT ESTIMATOR

In this section, we introduce another set of conditions based on the gradient estimator introduced in [Flaxman et al. \(2004\)](#). We require smoothness assumptions on the functions and an additional gradient estimator. For any function $f : \mathbb{R}^d \rightarrow \mathbb{R}$, we define the gradient estimator $\widehat{\nabla} f$ to be:

$$\widehat{\nabla} f(\mathbf{x}) := \frac{d}{\delta} \mathbb{E}_{\mathbf{u} \sim \mathcal{U}(S^{d-1})} [f(\mathbf{x} + \delta \mathbf{u}) \mathbf{u}],$$

where $S^{d-1} := \{\mathbf{x} \in \mathbb{R}^d : \|\mathbf{x}\|_2 = 1\}$ and $\delta > 0$ is a free parameter. We use $\mathcal{U}(S^{d-1})$ to denote the uniformly random distribution over S^{d-1} . We assume there exist $M > 0$, s.t.:

- for all $t, \sup_{\mathbf{x}} \|\nabla_{\mathbf{x}}^2 \hat{v}(\mathbf{x}, t)\|_2, \sup_{\mathbf{x}} \|\nabla_{\mathbf{x}}^2 v^*(\mathbf{x}, t)\|_2 \leq M$

We first present two properties of the gradient estimator.

Lemma 10. *Let f_1, f_2 be two functions that map from $\mathbb{R}^d$ to $\mathbb{R}$, then*

$$\|\widehat{\nabla f_1}(\mathbf{x}) - \widehat{\nabla f_2}(\mathbf{x})\|_2 \leq \frac{d}{\delta} \mathbb{E}_{\mathbf{u} \sim \mathcal{U}(S^{d-1})} [|f_1(\mathbf{x} + \delta \mathbf{u}) - f_2(\mathbf{x} + \delta \mathbf{u})|].$$

Lemma 11. *Let f be $\mathbb{R}^d \rightarrow \mathbb{R}$ and $\sup_{\mathbf{x}} \|\nabla^2 f(\mathbf{x})\|_2 \leq M$, then*

$$\|\widehat{\nabla f}(\mathbf{x}) - \nabla f(\mathbf{x})\|_2 \leq \frac{d\delta M}{2}.$$

Given these, we can show that, for all t ,

$$\begin{aligned} & \mathbb{E}_{\mathbf{x} \sim \hat{p}_t} [\|\nabla_{\mathbf{x}} \hat{v}(\mathbf{x}, t) - \nabla_{\mathbf{x}} v^*(\mathbf{x}, t)\|_2^2] \\ & \leq 3\mathbb{E}_{\mathbf{x} \sim \hat{p}_t} [\|\nabla_{\mathbf{x}} \hat{v}(\mathbf{x}, t) - \widehat{\nabla_{\mathbf{x}} \hat{v}}(\mathbf{x}, t)\|_2^2] + 3\mathbb{E}_{\mathbf{x} \sim \hat{p}_t} [\|\widehat{\nabla_{\mathbf{x}} \hat{v}}(\mathbf{x}, t) - \widehat{\nabla_{\mathbf{x}} v^*}(\mathbf{x}, t)\|_2^2] \\ & \quad + 3\mathbb{E}_{\mathbf{x} \sim \hat{p}_t} [\|\widehat{\nabla_{\mathbf{x}} v^*}(\mathbf{x}, t) - \nabla_{\mathbf{x}} v^*(\mathbf{x}, t)\|_2^2] \\ & \leq \frac{3d^2\delta^2 M^2}{2} + \frac{3d^2}{\delta^2} \mathbb{E}_{\mathbf{x} \sim \hat{p}_t} [\left(\mathbb{E}_{\mathbf{u} \sim \mathcal{U}(S^{d-1})} [\hat{v}(\mathbf{x} + \delta \mathbf{u}, t) - v^*(\mathbf{x} + \delta \mathbf{u}, t)]\right)^2] \\ & \leq \frac{3d^2\delta^2 M^2}{2} + \frac{3d^2}{\delta^2} \mathbb{E}_{\mathbf{x} \sim \hat{p}_t} [\left|\mathbb{E}_{\mathbf{u} \sim \mathcal{U}(S^{d-1})} [\hat{v}(\mathbf{x} + \delta \mathbf{u}, t) - v^*(\mathbf{x} + \delta \mathbf{u}, t)]\right|] \\ & \leq \frac{3d^2\delta^2 M^2}{2} + \frac{3d^2}{\delta^2} \mathbb{E}_{\mathbf{x} \sim \hat{p}_t, \mathbf{u} \sim \mathcal{U}(S^{d-1})} [|\hat{v}(\mathbf{x} + \delta \mathbf{u}, t) - v^*(\mathbf{x} + \delta \mathbf{u}, t)|] \\ & \leq \frac{3d^2\delta^2 M^2}{2} + \frac{3d^2}{\delta^2} \sqrt{\mathbb{E}_{\mathbf{x} \sim \hat{p}_t, \mathbf{u} \sim \mathcal{U}(S^{d-1})} [(\hat{v}(\mathbf{x} + \delta \mathbf{u}, t) - v^*(\mathbf{x} + \delta \mathbf{u}, t))^2]}, \end{aligned}$$

where for the third inequality, we use the fact that $\hat{v}(\cdot, \cdot), v^*(\cdot, \cdot) \in [0, 1]$. The remaining steps follow from Lemma 8, similar to Section B.1:

$$\begin{aligned} & \mathbb{E}_{\mathbf{x} \sim \hat{p}_t} [\|\nabla_{\mathbf{x}} \ln \hat{v}(\mathbf{x}, t) - \nabla_{\mathbf{x}} \ln v^*(\mathbf{x}, t)\|_2^2] \\ & \leq 2 \exp(2\eta R_{\max}) \mathbb{E}_{\mathbf{x} \sim \hat{p}_t} [\|\nabla_{\mathbf{x}} \hat{v}(\mathbf{x}, t) - \nabla_{\mathbf{x}} v^*(\mathbf{x}, t)\|_2^2] + 2 \exp(4\eta R_{\max}) L \mathbb{E}_{\mathbf{x} \sim \hat{p}_t} [|\hat{v}(\mathbf{x}, t) - v^*(\mathbf{x}, t)|^2] \\ & \leq 2 \exp(2\eta R_{\max}) \left(\frac{3d^2\delta^2 M^2}{2} + \frac{3d^2}{\delta^2} \sqrt{\mathbb{E}_{\mathbf{x} \sim \hat{p}_t, \mathbf{u} \sim \mathcal{U}(S^{d-1})} [(\hat{v}(\mathbf{x} + \delta \mathbf{u}, t) - v^*(\mathbf{x} + \delta \mathbf{u}, t))^2]} \right) \\ & \quad + 2 \exp(4\eta R_{\max}) L \mathbb{E}_{\mathbf{x} \sim \hat{p}_t} [|\hat{v}(\mathbf{x}, t) - v^*(\mathbf{x}, t)|^2] \end{aligned}$$

Compared to Algorithm 1, this approach requires one to additionally optimize for $(\hat{v}(\mathbf{x} + \delta \mathbf{u}, t) - v^*(\mathbf{x} + \delta \mathbf{u}, t))^2$, i.e. to make sure $\hat{v}$ and v are close under some “wider” distribution. The value of δ is a free parameter that can be adjusted to improve the accuracy of the gradient estimator.

B.3 PROOF OF LEMMA 8

$$\begin{aligned} \|\nabla \ln f_1(\mathbf{x}) - \nabla \ln f_2(\mathbf{x})\|_2 &= \left\| \frac{\nabla f_1(\mathbf{x})}{f_1(\mathbf{x})} - \frac{\nabla f_2(\mathbf{x})}{f_2(\mathbf{x})} \right\|_2 \\ &\leq \left\| \frac{\nabla f_1(\mathbf{x})}{f_1(\mathbf{x})} - \frac{\nabla f_2(\mathbf{x})}{f_1(\mathbf{x})} \right\|_2 + \left\| \frac{\nabla f_2(\mathbf{x})}{f_1(\mathbf{x})} - \frac{\nabla f_2(\mathbf{x})}{f_2(\mathbf{x})} \right\|_2 \\ &= \frac{\|\nabla f_1(\mathbf{x}) - \nabla f_2(\mathbf{x})\|_2}{|f_1(\mathbf{x})|} + \frac{1}{|f_1(\mathbf{x})f_2(\mathbf{x})|} \|\nabla f_2(\mathbf{x})(f_2(\mathbf{x}) - f_1(\mathbf{x}))\|_2 \\ &\leq \frac{1}{\inf_{\mathbf{x}} |f_1(\mathbf{x})|} \|\nabla f_1(\mathbf{x}) - \nabla f_2(\mathbf{x})\|_2 + \frac{\sup_{\mathbf{x}} \|\nabla f_2(\mathbf{x})\|_2}{\inf_{\mathbf{x}} |f_1(\mathbf{x})| \inf_{\mathbf{x}} |f_2(\mathbf{x})|} |f_2(\mathbf{x}) - f_1(\mathbf{x})|, \end{aligned}$$

B.4 PROOF OF LEMMA 9

We rewrite the target as componentwise form:

$$\|\nabla F(\mathbf{x})\|_2^2 = \sum_{i=1}^d \left(\frac{\partial F(\mathbf{x})}{\partial x_i} \right)^2.$$

By applying integration by parts to the i -th dimension,

$$\begin{aligned} & \int \frac{\partial F(\mathbf{x})}{\partial x_i} \cdot \frac{\partial F(\mathbf{x})}{\partial x_i} p(\mathbf{x}) d\mathbf{x} \\ &= \int F(\mathbf{x}) \cdot \frac{\partial F(\mathbf{x})}{\partial x_i} p(\mathbf{x}) \Big|_{x_i=-\infty}^{\infty} d\mathbf{x}_{-i} - \int F(\mathbf{x}) \left(\frac{\partial^2 F(\mathbf{x})}{\partial x_i^2} p(\mathbf{x}) + \frac{\partial F(\mathbf{x})}{\partial x_i} \frac{\partial p(\mathbf{x})}{\partial x_i} \right) d\mathbf{x} \\ &= - \int F(\mathbf{x}) \left(\frac{\partial^2 F(\mathbf{x})}{\partial x_i^2} p(\mathbf{x}) + \frac{\partial F(\mathbf{x})}{\partial x_i} \frac{\partial \log p(\mathbf{x})}{\partial x_i} p(\mathbf{x}) \right) d\mathbf{x}. \end{aligned}$$

By summing over all coordinates, we get:

$$\mathbb{E}_{\mathbf{x} \sim p} \|\nabla F(\mathbf{x})\|_2^2 = -\mathbb{E}_{\mathbf{x} \sim p} \left[F(\mathbf{x}) \sum_{i=1}^d \frac{\partial^2 F(\mathbf{x})}{\partial x_i^2} \right] - \mathbb{E}_{\mathbf{x} \sim p} [F(\mathbf{x}) \langle \nabla F(\mathbf{x}), \nabla \log p(\mathbf{x}) \rangle]$$

By Cauchy-Schwarz inequality:

$$\begin{aligned} & \mathbb{E}_{\mathbf{x} \sim p} \|\nabla F(\mathbf{x})\|_2^2 \\ & \leq \left(\sqrt{\mathbb{E}_{\mathbf{x} \sim p} \left[\left(\sum_{i=1}^d \frac{\partial^2 F(\mathbf{x})}{\partial x_i^2} \right)^2 \right]} + \sqrt{\mathbb{E}_{\mathbf{x} \sim p} [\langle \nabla F(\mathbf{x}), \nabla \log p(\mathbf{x}) \rangle^2]} \right) \sqrt{\mathbb{E}_{\mathbf{x} \sim p} [F^2(\mathbf{x})]} \\ & \leq \left(\sqrt{\mathbb{E}_{\mathbf{x} \sim p} \left[\left(\sum_{i=1}^d \frac{\partial^2 F(\mathbf{x})}{\partial x_i^2} \right)^2 \right]} + \sup_{\mathbf{x}} (\|\nabla F(\mathbf{x})\|_2) \sqrt{\mathbb{E}_{\mathbf{x} \sim p} [\|\nabla \log p(\mathbf{x})\|_2^2]} \right) \sqrt{\mathbb{E}_{\mathbf{x} \sim p} [F^2(\mathbf{x})]} \end{aligned}$$

B.5 PROOF OF LEMMA 10

By Jensen's inequality,

$$\begin{aligned} \|\widehat{\nabla f_1}(\mathbf{x}) - \widehat{\nabla f_2}(\mathbf{x})\|_2 &= \frac{d}{\delta} \|\mathbb{E}_{\mathbf{u} \sim \mathcal{U}(S^{d-1})} [f_1(\mathbf{x} + \delta \mathbf{u}) \mathbf{u}] - \mathbb{E}_{\mathbf{u} \sim \mathcal{U}(S^{d-1})} [f_2(\mathbf{x} + \delta \mathbf{u}) \mathbf{u}]\|_2 \\ &= \frac{d}{\delta} \|\mathbb{E}_{\mathbf{u} \sim \mathcal{U}(S^{d-1})} [(f_1(\mathbf{x} + \delta \mathbf{u}) - f_2(\mathbf{x} + \delta \mathbf{u})) \mathbf{u}]\|_2 \\ &\leq \frac{d}{\delta} \mathbb{E}_{\mathbf{u} \sim \mathcal{U}(S^{d-1})} [\|(f_1(\mathbf{x} + \delta \mathbf{u}) - f_2(\mathbf{x} + \delta \mathbf{u})) \mathbf{u}\|_2] \\ &= \frac{d}{\delta} \mathbb{E}_{\mathbf{u} \sim \mathcal{U}(S^{d-1})} [|f_1(\mathbf{x} + \delta \mathbf{u}) - f_2(\mathbf{x} + \delta \mathbf{u})| \|\mathbf{u}\|_2] \\ &= \frac{d}{\delta} \mathbb{E}_{\mathbf{u} \sim \mathcal{U}(S^{d-1})} [|f_1(\mathbf{x} + \delta \mathbf{u}) - f_2(\mathbf{x} + \delta \mathbf{u})|]. \end{aligned}$$

B.6 PROOF OF LEMMA 11

By Taylor's theorem, for all $\mathbf{x}, \mathbf{b}$, there exists $\xi(\mathbf{x}, \mathbf{b})$, s.t.

$$f(\mathbf{x} + \mathbf{b}) = f(\mathbf{x}) + \langle \nabla f(\mathbf{x}), \mathbf{b} \rangle + \frac{1}{2} \mathbf{b}^\top \nabla^2 f(\xi(\mathbf{x}, \mathbf{b})) \mathbf{b}.$$

Then

$$\begin{aligned} & \|\widehat{\nabla f}(\mathbf{x}) - \nabla f(\mathbf{x})\|_2 = \left\| \frac{d}{\delta} \mathbb{E}_{\mathbf{u} \sim \mathcal{U}(S^{d-1})} [f(\mathbf{x} + \delta \mathbf{u}) \mathbf{u}] - \nabla f(\mathbf{x}) \right\|_2 \\ &= \left\| \frac{d}{\delta} \mathbb{E}_{\mathbf{u} \sim \mathcal{U}(S^{d-1})} \left[f(\mathbf{x}) \mathbf{u} + \delta \langle \nabla f(\mathbf{x}), \mathbf{u} \rangle \mathbf{u} + \frac{\delta^2}{2} (\mathbf{u}^\top \nabla^2 f(\xi(\mathbf{x}, \delta \mathbf{u})) \mathbf{u}) \mathbf{u} \right] - \nabla f(\mathbf{x}) \right\|_2 \end{aligned}$$

Because $\mathbb{E}_{\mathbf{u} \sim \mathcal{U}(S^{d-1})} [\mathbf{u}] = 0$,

$$\mathbb{E}_{\mathbf{u} \sim \mathcal{U}(S^{d-1})} [f(\mathbf{x}) \mathbf{u}] = 0.$$

Because $\mathbb{E}_{\mathbf{u} \sim \mathcal{U}(S^{d-1})}[\mathbf{u}\mathbf{u}^\top] = \frac{1}{d}I$,

$$\mathbb{E}_{\mathbf{u} \sim \mathcal{U}(S^{d-1})}[\langle \nabla f(\mathbf{x}), \mathbf{u} \rangle \mathbf{u}] = \mathbb{E}_{\mathbf{u} \sim \mathcal{U}(S^{d-1})}[\mathbf{u}\mathbf{u}^\top \nabla f(\mathbf{x})] = \mathbb{E}_{\mathbf{u} \sim \mathcal{U}(S^{d-1})}[\mathbf{u}\mathbf{u}^\top] \nabla f(\mathbf{x}) = \frac{1}{d} \nabla f(\mathbf{x}).$$

Thus, by Jensen's inequality, we have:

$$\begin{aligned} \left\| \widehat{\nabla} f(\mathbf{x}) - \nabla f(\mathbf{x}) \right\|_2 &= \left\| \frac{d}{\delta} \mathbb{E}_{\mathbf{u} \sim \mathcal{U}(S^{d-1})} \left[\frac{\delta^2}{2} (\mathbf{u}^\top \nabla^2 f(\xi(\mathbf{x}, \delta \mathbf{u})) \mathbf{u}) \mathbf{u} \right] \right\|_2 \\ &\leq \frac{d\delta}{2} \mathbb{E}_{\mathbf{u} \sim \mathcal{U}(S^{d-1})} [\|(\mathbf{u}^\top \nabla^2 f(\xi(\mathbf{x}, \delta \mathbf{u})) \mathbf{u}) \mathbf{u}\|_2] = \frac{d\delta}{2} \mathbb{E}_{\mathbf{u} \sim \mathcal{U}(S^{d-1})} [\|\mathbf{u}^\top \nabla^2 f(\xi(\mathbf{x}, \delta \mathbf{u})) \mathbf{u}\| \|\mathbf{u}\|_2] \\ &= \frac{d\delta}{2} \mathbb{E}_{\mathbf{u} \sim \mathcal{U}(S^{d-1})} [\|\mathbf{u}^\top \nabla^2 f(\xi(\mathbf{x}, \delta \mathbf{u})) \mathbf{u}\|] \leq \frac{d\delta}{2} \mathbb{E}_{\mathbf{u} \sim \mathcal{U}(S^{d-1})} [\|\mathbf{u}\|_2 \|\nabla^2 f(\xi(\mathbf{x}, \delta \mathbf{u})) \mathbf{u}\|_2] \\ &\leq \frac{d\delta M}{2}. \end{aligned}$$

C PROOF OF MAIN THEOREM

Proof. By Assumption 4,

$$\begin{aligned} \frac{1}{NM} \sum_{i=1}^N \sum_{(t, \mathbf{x}_t, r) \in D^i} \ln \frac{1}{\hat{R}^i(r|\mathbf{x}_t, t)} &\leq \min_R \frac{1}{NM} \sum_{i=1}^N \sum_{(t, \mathbf{x}_t, r) \in D^i} \ln \frac{1}{R(r|\mathbf{x}_t, t)} + \gamma_N \\ &\leq \frac{1}{NM} \sum_{i=1}^N \sum_{(t, \mathbf{x}_t, r) \in D^i} \ln \frac{1}{R^{\text{prior}}(r|\mathbf{x}_t, t)} + \gamma_N. \end{aligned}$$

After rearranging, we get

$$\frac{1}{NM} \sum_{i=1}^N \sum_{(t, \mathbf{x}_t, r) \in D^i} \ln \frac{R^{\text{prior}}(r|\mathbf{x}_t, t)}{\hat{R}^i(r|\mathbf{x}_t, t)} \leq \gamma_N.$$

According to Eq. (10), each $(t, \mathbf{x}_t)$, r is sampled from $R^{\text{prior}}(\cdot|\mathbf{x}_t, t)$. By taking expectation over $r|\mathbf{x}_t, t$ for all $\mathbf{x}_t, t$, we get:

$$\frac{1}{NM} \sum_{i=1}^N \sum_{(t, \mathbf{x}_t) \in D^i} D_{\text{KL}} \left(R^{\text{prior}}(\cdot|\mathbf{x}_t, t) \|\hat{R}^i(\cdot|\mathbf{x}_t, t) \right) \leq \gamma_N.$$

By taking expectation over $\mathbf{x}_t$ and t , we get:

$$\frac{1}{N} \sum_{i=1}^N \mathbb{E}_{t \sim \mathcal{U}(t), \mathbf{x}_t \sim P_{0 \rightarrow t}^{\mathbf{f}^i}(\cdot|\mathcal{N}(0, I))} \left[D_{\text{KL}} \left(R^{\text{prior}}(\cdot|\mathbf{x}_t, t) \|\hat{R}^i(\cdot|\mathbf{x}_t, t) \right) \right] \leq \gamma_N,$$

where $\mathbf{x}_t$ is sampled from the SDE induced by $\hat{R}^i$. By Pinsker's inequality,

$$\begin{aligned} &\frac{1}{N} \sum_{i=1}^N \mathbb{E}_{t \sim \mathcal{U}(t), \mathbf{x}_t \sim P_{0 \rightarrow t}^{\mathbf{f}^i}(\cdot|\mathcal{N}(0, I))} \left[\left| \text{TV} \left(R^{\text{prior}}(\cdot|\mathbf{x}_t, t), \hat{R}^i(\cdot|\mathbf{x}_t, t) \right) \right|^2 \right] \\ &\leq \frac{1}{N} \sum_{i=1}^N \frac{1}{2} \mathbb{E}_{t \sim \mathcal{U}(t), \mathbf{x}_t \sim P_{0 \rightarrow t}^{\mathbf{f}^i}(\cdot|\mathcal{N}(0, I))} \left[D_{\text{KL}} \left(R^{\text{prior}}(\cdot|\mathbf{x}_t, t) \|\hat{R}^i(\cdot|\mathbf{x}_t, t) \right) \right] \leq \frac{1}{2} \gamma_N. \end{aligned}$$

This means, there exists $\hat{n} \in \{1, \dots, N\}$, s.t.

$$\mathbb{E}_{t \sim \mathcal{U}(T), x \sim P_{0 \rightarrow t}^{\mathbf{f}^{\hat{n}}}(\cdot|\mathcal{N}(0, I))} \left| \text{TV} \left(R^{\text{prior}}(\cdot|\mathbf{x}_t, t), \hat{R}^{\hat{n}}(\cdot|\mathbf{x}_t, t) \right) \right|^2 \leq \frac{1}{2} \gamma_N.$$

Let

$$\begin{aligned} \hat{\mathbf{f}}(\mathbf{x}_t, t) &:= \mathbf{f}^{\hat{n}}(\mathbf{x}_t, t) = \nabla_{\mathbf{x}_t} \ln \mathbb{E}_{r \sim \hat{R}^{\hat{n}}(\cdot|\mathbf{x}_t, t)} \exp(\eta \cdot r) \\ \mathbf{f}^*(\mathbf{x}_t, t) &:= \nabla_{\mathbf{x}_t} \ln p(y = 1|\mathbf{x}_t) = \nabla_{\mathbf{x}_t} \ln \mathbb{E}_{r \sim R^{\text{prior}}(\cdot|\mathbf{x}_t, t)} \exp(\eta \cdot r) \end{aligned}$$

and

$$\hat{p}_t := P_{0 \rightarrow t}^{\hat{\mathbf{f}}}(\cdot | \mathcal{N}(0, I)), \quad p_t := P_{0 \rightarrow t}^{\mathbf{f}^*}(\cdot | q_T),$$

recall that q_T is the marginal distribution of the forward process at time T . By Lemma 13,

$$\begin{aligned} \frac{\partial}{\partial t} D_{\text{KL}}(\hat{p}_t \| p_t) &= -g^2(T-t) \mathbb{E}_{x \sim \hat{p}_t} \left[\left\| \nabla \log \frac{\hat{p}_t(x)}{p_t(x)} \right\|_2^2 \right] \\ &\quad + \mathbb{E}_{x \sim \hat{p}_t} \left[\left\langle g^2(T-t) \left(\hat{\mathbf{f}}(x, t) - \mathbf{f}^*(x, t) \right), \nabla \log \frac{\hat{p}_t(x)}{p_t(x)} \right\rangle \right] \\ &\leq \frac{1}{4} g^2(T-t) \mathbb{E}_{x \sim \hat{p}_t} \left[\left\| \hat{\mathbf{f}}(x, t) - \mathbf{f}^*(x, t) \right\|_2^2 \right]. \end{aligned}$$

By integrating over $t \in [0, T]$, we get:

$$D_{\text{KL}}(\hat{p}_T \| p_T) \leq D_{\text{KL}}(\hat{p}_0 \| p_0) + \frac{1}{4} \int_0^T g^2(T-t) \mathbb{E}_{x \sim \hat{p}_t} \left[\left\| \hat{\mathbf{f}}(x, t) - \mathbf{f}^*(x, t) \right\|_2^2 \right] dt,$$

where $\hat{p}_0 = \mathcal{N}(0, I)$ and $p_0 = q_T$. By Assumption 5, $D_{\text{KL}}(\hat{p}_0 \| p_0) \leq \epsilon_T$. By Assumption 6,

$$\begin{aligned} &\int_0^T g^2(T-t) \mathbb{E}_{x \sim \hat{p}_t} \left[\left\| \hat{\mathbf{f}}(x, t) - \mathbf{f}^*(x, t) \right\|_2^2 \right] dt \\ &= T \|g\|_\infty^2 \mathbb{E}_{t \sim \mathcal{U}(T), x \sim \hat{p}_t} \left[\left\| \hat{\mathbf{f}}(x, t) - \mathbf{f}^*(x, t) \right\|_2^2 \right] \\ &\leq T \|g\|_\infty^2 L^2 \mathbb{E}_{t \sim \mathcal{U}(T), x \sim \hat{p}_t} \left[\left| \mathbb{E}_{r \sim \hat{R}^{\hat{n}}(\cdot | \mathbf{x}_t, t)} \exp(\eta \cdot r) - \mathbb{E}_{r \sim R^{\text{prior}}(\cdot | \mathbf{x}_t, t)} \exp(\eta \cdot r) \right|^2 \right] \\ &\leq T \|g\|_\infty^2 L^2 \mathbb{E}_{t \sim \mathcal{U}(T), x \sim \hat{p}_t} \left[\left| \text{TV}(\hat{R}^{\hat{n}}(\cdot | \mathbf{x}_t, t), R^{\text{prior}}(\cdot | \mathbf{x}_t, t)) \right|^2 \right] \\ &\leq \frac{1}{2} T \|g\|_\infty^2 L^2 \gamma_N \end{aligned}$$

To conclude:

$$D_{\text{KL}}(\hat{p}_T \| p_T) \leq \epsilon_T + \frac{1}{2} T \|g\|_\infty^2 L^2 \gamma_N.$$

□

D TECHNICAL LEMMAS

The following lemma shows that the finite-dimensional distribution of the forward and reverse SDEs are the same.

Lemma 12 (Section 5 of Anderson (1982)). *Let $\{\bar{\mathbf{x}}_\tau\}_\tau$ be the process generated by the forward SDE in Eq. (1), and $\{\mathbf{x}_t\}_t$ be the process generated by the reverse SDE in Eq. (2). Then for all $s > t$, the conditional distribution $\bar{\mathbf{x}}_t | \bar{\mathbf{x}}_s$ and $\mathbf{x}_{T-t} | \bar{\mathbf{x}}_{T-s}$ have the same density, i.e. for all $x, x' \in \mathbb{R}^d$,*

$$\Pr[\bar{\mathbf{x}}_t = x' | \bar{\mathbf{x}}_s = x] = \Pr[\mathbf{x}_{T-t} = x' | \bar{\mathbf{x}}_{T-s} = x].$$

This result is proved by considering the backward Kolmogorov equation of the forward process. The proof is presented in Section 5 of Anderson (1982). For completeness, we include the proof in Section D.1.

This result implies that the reverse and the forward SDE have the same “joint distribution” (in the sense of *finite-dimensional distribution*). This can be proved by two steps:

1. using the Fokker-Planck equation to show that the forward and reverse SDE have the same marginal distribution given proper initial conditions;
2. write the finite-dimensional distribution as the product of a sequence of conditional distributions and the initial marginal distribution by using Markovian property iteratively.

The following lemma upper bound the KL-divergence between the marginal distributions of two SDEs in terms of the difference in their drift terms:

Lemma 13 (Lemma 6 of [Chen et al. \(2023\)](#)). *Consider the two Ito processes:*

$$\begin{aligned} dX_t &= F_1(X_t, t)dt + g(t)dw_t & X_0 &= a \\ dY_t &= F_2(Y_t, t)dt + g(t)dw_t & Y_0 &= a \end{aligned}$$

where F_1, F_2, g are continuous functions and may depend on a . We assume the uniqueness and regularity conditions hold:

1. The two SDEs have unique solutions.
2. X_t, Y_t admit densities $p_t^1, p_t^2 \in C^2(\mathbb{R}^d)$ for $t > 0$

And define the Fisher information between p_t, q_t as:

$$J(p_t \| q_t) := \int p_t(x) \|\nabla \log \frac{p_t(x)}{q_t(x)}\|^2 dx$$

Then for any $t > 0$, the evolution of $D_{\text{KL}}(p_t^1 \| p_t^2)$ is given by:

$$\frac{\partial}{\partial t} \text{KL}(p_t^1 \| p_t^2) = -g(t)^2 J(p_t^1 \| p_t^2) + \mathbb{E}_{x \sim p_t^1} \left[\left\langle F_1(X_t, t) - F_2(X_t, t), \nabla \log \frac{p_t^1(x)}{p_t^2(x)} \right\rangle \right]$$

D.1 PROOF OF LEMMA 12

Let $q(x_s, s | x_t, t)$ be the conditional distribution of [Eq. \(1\)](#), where $s \geq t$ is the time index. And $q(x_t, t)$ be the marginal distribution at time t . By the backward Kolmogorov equation:

$$\frac{\partial q(\bar{x}_s, s | \bar{x}_t, t)}{\partial t} = - \sum_i \mathbf{h}^i(\bar{x}_t, t) \frac{\partial q(\bar{x}_s, s | \bar{x}_t, t)}{\partial \bar{x}_t^i} - \frac{1}{2} g^2(t) \sum_i \frac{\partial^2 q(\bar{x}_s, s | \bar{x}_t, t)}{\partial (\bar{x}_t^i)^2} \quad (13)$$

According to Fokker-Planck equation, the marginal distribution $q(\bar{x}_t, t)$ satisfies:

$$\frac{\partial q(\bar{x}_t, t)}{\partial t} = - \sum_i \frac{\partial}{\partial \bar{x}_t^i} [q(\bar{x}_t, t) \mathbf{h}^i(\bar{x}_t, t)] + \frac{1}{2} g^2(t) \sum_i \frac{\partial^2 q(\bar{x}_t, t)}{\partial (\bar{x}_t^i)^2} \quad (14)$$

Because the joint distribution satisfies:

$$q(\bar{x}_s, s, \bar{x}_t, t) = q(\bar{x}_s, s | \bar{x}_t, t) q(\bar{x}_t, t),$$

we have

$$\frac{\partial q(\bar{x}_s, s, \bar{x}_t, t)}{\partial t} = q(\bar{x}_t, t) \frac{\partial q(\bar{x}_s, s | \bar{x}_t, t)}{\partial t} + q(\bar{x}_s, s | \bar{x}_t, t) \frac{\partial q(\bar{x}_t, t)}{\partial t}.$$

Plug in [Eq. \(13\)](#) and [Eq. \(14\)](#), we have:

$$\begin{aligned} & \frac{\partial q(\bar{x}_s, s, \bar{x}_t, t)}{\partial t} \\ &= q(\bar{x}_t, t) \left(- \sum_i \mathbf{h}^i(\bar{x}_t, t) \frac{\partial q(\bar{x}_s, s | \bar{x}_t, t)}{\partial \bar{x}_t^i} - \frac{1}{2} g^2(t) \sum_i \frac{\partial^2 q(\bar{x}_s, s | \bar{x}_t, t)}{\partial (\bar{x}_t^i)^2} \right) \\ & \quad + q(\bar{x}_s, s | \bar{x}_t, t) \left(- \sum_i \frac{\partial}{\partial \bar{x}_t^i} [q(\bar{x}_t, t) \mathbf{h}^i(\bar{x}_t, t)] + \frac{1}{2} g^2(t) \sum_i \frac{\partial^2 q(\bar{x}_t, t)}{\partial (\bar{x}_t^i)^2} \right). \end{aligned}$$

For the first and third terms:

$$\begin{aligned}
& -q(\bar{x}_t, t) \sum_i \mathbf{h}^i(\bar{x}_t, t) \frac{\partial q(\bar{x}_s, s | \bar{x}_t, t)}{\partial \bar{x}_t^i} - q(\bar{x}_s, s | \bar{x}_t, t) \sum_i \frac{\partial}{\partial \bar{x}_t^i} [q(\bar{x}_t, t) \mathbf{h}^i(\bar{x}_t, t)] \\
& = - \sum_i \mathbf{h}^i(\bar{x}_t, t) \left(q(\bar{x}_t, t) \frac{\partial q(\bar{x}_s, s | \bar{x}_t, t)}{\partial \bar{x}_t^i} + q(\bar{x}_s, s | \bar{x}_t, t) \frac{\partial}{\partial \bar{x}_t^i} q(\bar{x}_t, t) \right) \\
& \quad - q(\bar{x}_s, s | \bar{x}_t, t) q(\bar{x}_t, t) \sum_i \frac{\partial}{\partial \bar{x}_t^i} \mathbf{h}^i(\bar{x}_t, t) \\
& = - \sum_i \mathbf{h}^i(\bar{x}_t, t) \frac{\partial q(\bar{x}_s, s, \bar{x}_t, t)}{\partial \bar{x}_t^i} - q(\bar{x}_s, s, \bar{x}_t, t) \sum_i \frac{\partial}{\partial \bar{x}_t^i} \mathbf{h}^i(\bar{x}_t, t) \\
& = - \sum_i \frac{\partial}{\partial \bar{x}_t^i} [\mathbf{h}^i(\bar{x}_t, t) q(\bar{x}_s, s, \bar{x}_t, t)].
\end{aligned}$$

For the second and forth terms:

$$\begin{aligned}
& \frac{1}{2} g^2(t) \sum_i \left(-q(\bar{x}_t, t) \frac{\partial^2 q(\bar{x}_s, s | \bar{x}_t, t)}{\partial (\bar{x}_t^i)^2} + q(\bar{x}_s, s | \bar{x}_t, t) \frac{\partial^2 q(\bar{x}_t, t)}{\partial (\bar{x}_t^i)^2} \right) \\
& = - \frac{1}{2} g^2(t) \sum_i \left(q(\bar{x}_s, s | \bar{x}_t, t) \frac{\partial^2 q(\bar{x}_t, t)}{\partial (\bar{x}_t^i)^2} + 2 \frac{\partial q(\bar{x}_s, s | \bar{x}_t, t)}{\partial \bar{x}_t^i} \frac{\partial q(\bar{x}_t, t)}{\partial \bar{x}_t^i} + q(\bar{x}_t, t) \frac{\partial^2 q(\bar{x}_s, s | \bar{x}_t, t)}{\partial (\bar{x}_t^i)^2} \right) \\
& \quad + \frac{1}{2} g^2(t) \sum_i \left(2q(\bar{x}_s, s | \bar{x}_t, t) \frac{\partial^2 q(\bar{x}_t, t)}{\partial (\bar{x}_t^i)^2} + 2 \frac{\partial q(\bar{x}_s, s | \bar{x}_t, t)}{\partial \bar{x}_t^i} \frac{\partial q(\bar{x}_t, t)}{\partial \bar{x}_t^i} \right) \\
& = - \frac{1}{2} g^2(t) \sum_i \frac{\partial^2 q(\bar{x}_s, s, \bar{x}_t, t)}{\partial (\bar{x}_t^i)^2} + g^2(t) \sum_i \left(q(\bar{x}_s, s | \bar{x}_t, t) \frac{\partial^2 q(\bar{x}_t, t)}{\partial (\bar{x}_t^i)^2} + \frac{\partial q(\bar{x}_s, s | \bar{x}_t, t)}{\partial \bar{x}_t^i} \frac{\partial q(\bar{x}_t, t)}{\partial \bar{x}_t^i} \right) \\
& = - \frac{1}{2} g^2(t) \sum_i \frac{\partial^2 q(\bar{x}_s, s, \bar{x}_t, t)}{\partial (\bar{x}_t^i)^2} + g^2(t) \sum_i \frac{\partial}{\partial \bar{x}_t^i} \left[q(\bar{x}_s, s | \bar{x}_t, t) \frac{\partial q(\bar{x}_t, t)}{\partial \bar{x}_t^i} \right].
\end{aligned}$$

To summarize, the joint distribution satisfies the following PDE

$$\frac{\partial q(\bar{x}_s, s, \bar{x}_t, t)}{\partial t} = - \sum_i \frac{\partial}{\partial \bar{x}_t^i} [q(\bar{x}_s, s, \bar{x}_t, t) \bar{\mathbf{h}}^i(\bar{x}_t, t)] - \frac{1}{2} g^2(t) \sum_i \frac{\partial^2 q(\bar{x}_s, s, \bar{x}_t, t)}{\partial (\bar{x}_t^i)^2},$$

where $\bar{\mathbf{h}}^i(\bar{x}_t, t)$ is defined as:

$$\bar{\mathbf{h}}^i(\bar{x}_t, t) := \mathbf{h}^i(\bar{x}_t, t) - \frac{1}{q(\bar{x}_t, t)} g^2(t) \frac{\partial}{\partial \bar{x}_t^i} q(\bar{x}_t, t).$$

Divide $q(\bar{x}_s, s)$ on both sides, we get the following PDE for the conditional distribution $q(\bar{x}_t, t | \bar{x}_s, s)$ (notice that this one is conditioning on the future):

$$\frac{\partial q(\bar{x}_t, t | \bar{x}_s, s)}{\partial t} = - \sum_i \frac{\partial}{\partial \bar{x}_t^i} [q(\bar{x}_t, t | \bar{x}_s, s) \bar{\mathbf{h}}^i(\bar{x}_t, t)] - \frac{1}{2} g^2(t) \sum_i \frac{\partial^2 q(\bar{x}_t, t | \bar{x}_s, s)}{\partial (\bar{x}_t^i)^2} \quad (15)$$

Recall that the reverse process is defined by:

$$d\mathbf{x} = [-\mathbf{h}(\mathbf{x}, T-t) + g^2(T-t) \nabla \log q_{T-t}(\mathbf{x})] dt + g(T-t) d\mathbf{w}, \quad t \in [0, T].$$

The Fokker-Planck equation is given by:

$$\frac{\partial p(x_t, t)}{\partial t} = - \sum_i \frac{\partial}{\partial x_t^i} [-p(x_t, t) \bar{\mathbf{h}}^i(x_t, T-t)] + \frac{1}{2} g^2(T-t) \sum_i \frac{\partial^2 p(x_t, t)}{\partial (x_t^i)^2}.$$

We substitute t with $T-t$ and get:

$$- \frac{\partial p(x_{T-t}, T-t)}{\partial t} = \sum_i \frac{\partial}{\partial x_{T-t}^i} [p(x_{T-t}, T-t) \bar{\mathbf{h}}^i(x_{T-t}, t)] + \frac{1}{2} g^2(t) \sum_i \frac{\partial^2 p(x_{T-t}, T-t)}{\partial (x_{T-t}^i)^2}.$$

This PDE is identical to Eq. (15). By choosing the proper initial condition, we finish the proof.

E ADDITIONAL DETAILS OF TRAINING AND EVALUATION

We provide more information about the experimental setup and then overview some of the details of training and evaluation of SLCD in this section. The experiments are split up into two parts: image tasks (image compression and aesthetic evaluation) and biological sequence tasks (5' untranslated regions and DNA enhancer sequences).

All experiments were conducted on a single NVIDIA A6000 GPU.

Our classifier network is constructed to predict a histogram over the reward space and trained with cross-entropy loss. When training the classifier for the first iteration of DAgger (without any guidance) we label all states in the trajectory with the reward of the final state. For subsequent iterations, we only use states that are part of the *rollout* (i.e. the states in the trajectory that are computed solely using the prior, after the *rollin* section which uses the latest classifier).

To apply classifier guidance, we compute the empirical expectation of the classifier. That is for B buckets $(r_1, \dots, r_B)$ and predictor $\hat{R}^n(\cdot|\mathbf{x}_t, t)$, we compute:

$$f^n(x_t, t) = \nabla_{\mathbf{x}_t} \ln \left(\sum_{i=1}^B \hat{R}^n(r_i|\mathbf{x}_t, t) \exp(\eta \cdot r_i) \right).$$

Because the gradient is invariant to the reward scale, we set r_i to be a B equally spaced partition of $[0, 1]$.

E.1 EXPERIMENT TASK DETAILS

E.1.1 COMPARISON TASKS

Image Compression. The reward is the negative file size of the generated image, a non-differentiable compression score. We use Stable Diffusion v1.5 (Rombach et al., 2022) as the base model for this continuous-time diffusion task.

Image Aesthetics. Here the objective is to maximize the LAION aesthetic score (Schumman, 2022), obtained from a CLIP encoder followed by an MLP trained on human 1-to-10 ratings. This benchmark is standard in image-generation studies (Black et al., 2024; Domingo-Enrich et al., 2025; Li et al., 2024; Uehara et al., 2024a;b). We again employ Stable Diffusion v1.5 as the base model.

5' Untranslated Regions (5' UTR) and DNA Enhancers. For sequence generation we aim to maximize the mean ribosome load (MRL) of 5' UTRs measured via polysome profiling (Sample et al., 2019). Following Li et al. (2024), the base model is that of Sahoo et al. (2024) trained on Sample et al. (2019). Likewise, using the same base model, we optimize enhancer sequences according to expression levels predicted by the Enformer model (Avsec et al., 2021) in the HepG2 cell line.

E.1.2 COMPARISON METHODS

Best-of- N . We draw N independent samples from the base diffusion model and retain the single sample with the highest reward.

DPS. Diffusion Posterior Sampling (DPS) is a training-free variant of classifier guidance originally proposed for continuous diffusion models (Chung et al., 2023) and subsequently adapted to discrete diffusion (Li et al., 2024). We use the state-of-the-art implementation of Nisonoff et al. (2025).

SMC. Sequential Monte Carlo (SMC) methods (Del Moral & Doucet, 2014; Wu et al., 2023; Trippe et al., 2022) are a class of methods which use importance sampling on a number of rollouts, and select the best sample. However, note that SMC based methods do this across the *entire* batch, not at a per sample level.

SVDD-MC. SVDD-MC (Li et al., 2024) evaluates the expected reward of N candidates from the base model under an estimated value function and selects the candidate with the highest predicted return.

SVDD-PM. SVDD-PM (Li et al., 2024) is similar to SVDD-MC except that it uses the true reward for each candidate instead of relying on value-function estimates.

E.2 IMAGE TASK DETAILS (IMAGE COMPRESSION AND AESTHETIC)

For both image tasks, we use a lightweight classifier network, with model checkpoints sized at 4.93MB (compression) and 10.60MB (aesthetic).

The prompts are generated under the configuration of SVDD [Li et al. \(2024\)](#). In the image compression task, we generate 1,400 images per iteration and train for one epoch. The results are reported using the checkpoint from the 4th iteration. For the aesthetic evaluation task, we generate 10,500 images in the first iteration, followed by 1,400 images in each subsequent iteration, and train for 6 epochs. The results are reported using the checkpoint from the 8th iteration. For the compression task, we adopt the same architecture as [Li et al. \(2024\)](#). For the aesthetic task, we add five additional residual layers to account for its increased complexity. Notably, our classifier takes the latent representation as direct input, without reusing the VAE (as in SVDD for compression), CLIP (SVDD for aesthetics), or the U-Net (SVDD-PM for estimating $\hat{x}_0$). This results in a significantly more lightweight design, both in network size (approximately 10MB compared to over 4GB) and in runtime ([Fig. 3](#) center).

E.2.1 HYPERPARAMETERS

For the image tasks, the following hyperparameters are used:

Table 2: Image Task Hyperparameters. If two values are provided, the first corresponds to the compression task, and the second to the aesthetic task.

Hyperparameter	Value
Seed	43
Learning rate	1×10^{-4}
Optimizer betas	(0.9, 0.999)
Weight decay	0.0
Gradient accumulation steps	1
Batch size (classifier)	8
Batch size (inference)	1
Guidance scale	150/75
Train iterations (per round)	1/6
Eval interval (epoches)	1

E.3 SEQUENCE TASK DETAILS (5' UTR AND DNA ENHANCER)

Following [Li et al. \(2024\)](#), for the Enhancer task, we use an Enformer model [Avsec et al. \(2021\)](#), and for the 5' UTR task, we use ConvGRU model [Dey & Salem \(2017\)](#). We reference our classifier guidance implementation based off of [Li et al. \(2024\)](#). Both of these tasks are discrete diffusion problems and use the masking setup of [Sahoo et al. \(2024\)](#), following [Li et al. \(2024\)](#).

We use the following hyperparameters for the sequence tasks. In order to not be bound to a specific η , we take use the gradient of the empirical mean computing the samples for the next iteration with a guidance scale of 10. We can then change η at test time as shown in the paper. Due to the more light-weight nature of these tasks, we did 8 iterations of training for the DNA enhancer task. Likewise, we did 7 iterations for the 5' UTR task. Unlike the image tasks, we do not reinitialize the classifier network for each iteration.

E.3.1 HYPERPARAMETERS

For the DNA enhancer and 5' UTR tasks, the following hyperparameters are used:

Table 3: DNA Enhancer and 5' UTR Task Hyperparameters

Hyperparameter	Value
Seed	43
Learning rate	1×10^{-4}
Optimizer betas	(0.9, 0.95)
Weight decay	0.01
Gradient accumulation steps	4
Batch size (classifier)	5
Batch size (inference)	20
Guidance scale	10
Train iterations (per round)	200
Initial train iterations	600
Classifier epochs	1
Eval interval (steps)	20

F MORE IMAGE SAMPLES

Fig. 5 provides further results for the image compression task using prompts that were not seen during training. The consistent performance across these novel inputs demonstrates the generalization capability of SLCD.

Fig. 6 presents more images generated by SLCD on the image compression task, with different values of the parameter η . Each image is shown with its corresponding reward, demonstrating how varying η affects the trade-off between compression efficiency and visual quality.

Fig. 7 shows additional samples from the image aesthetic task, also generated with varying η values. These examples highlight SLCD’s ability to optimize for aesthetic quality under different configurations.

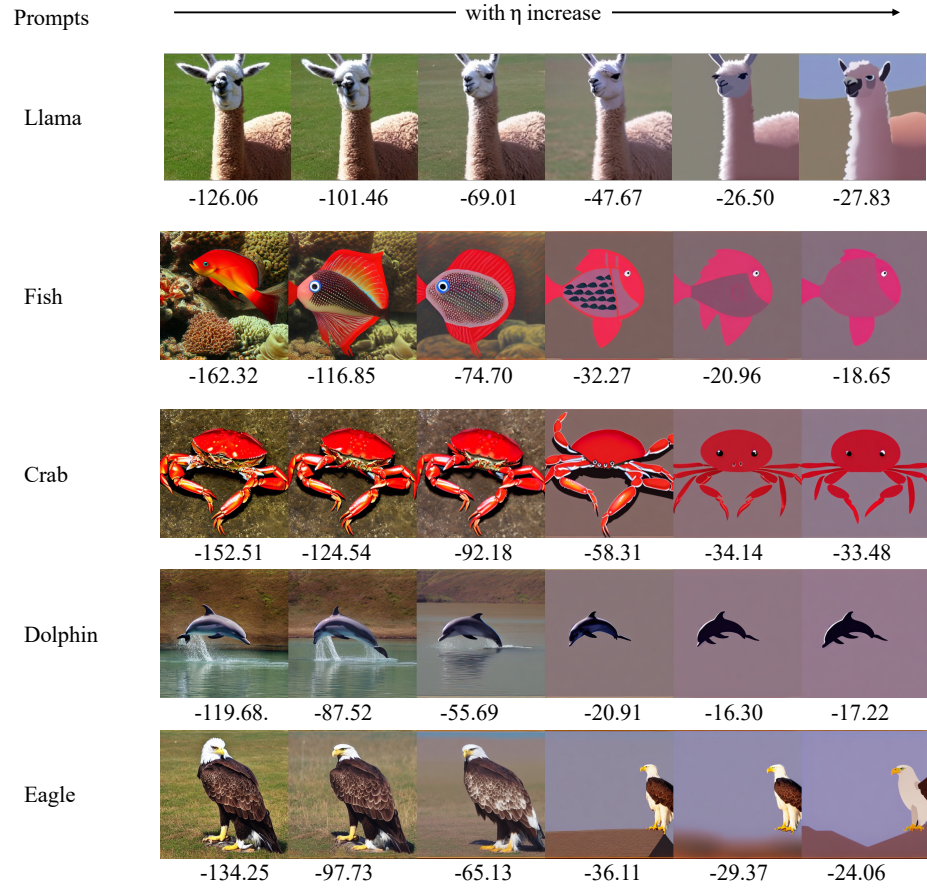


Figure 5: Additional images generated by SLCD with varying η values and their corresponding rewards on the image compression task. The prompts used were not seen during training, demonstrating the generalization capability of our method.

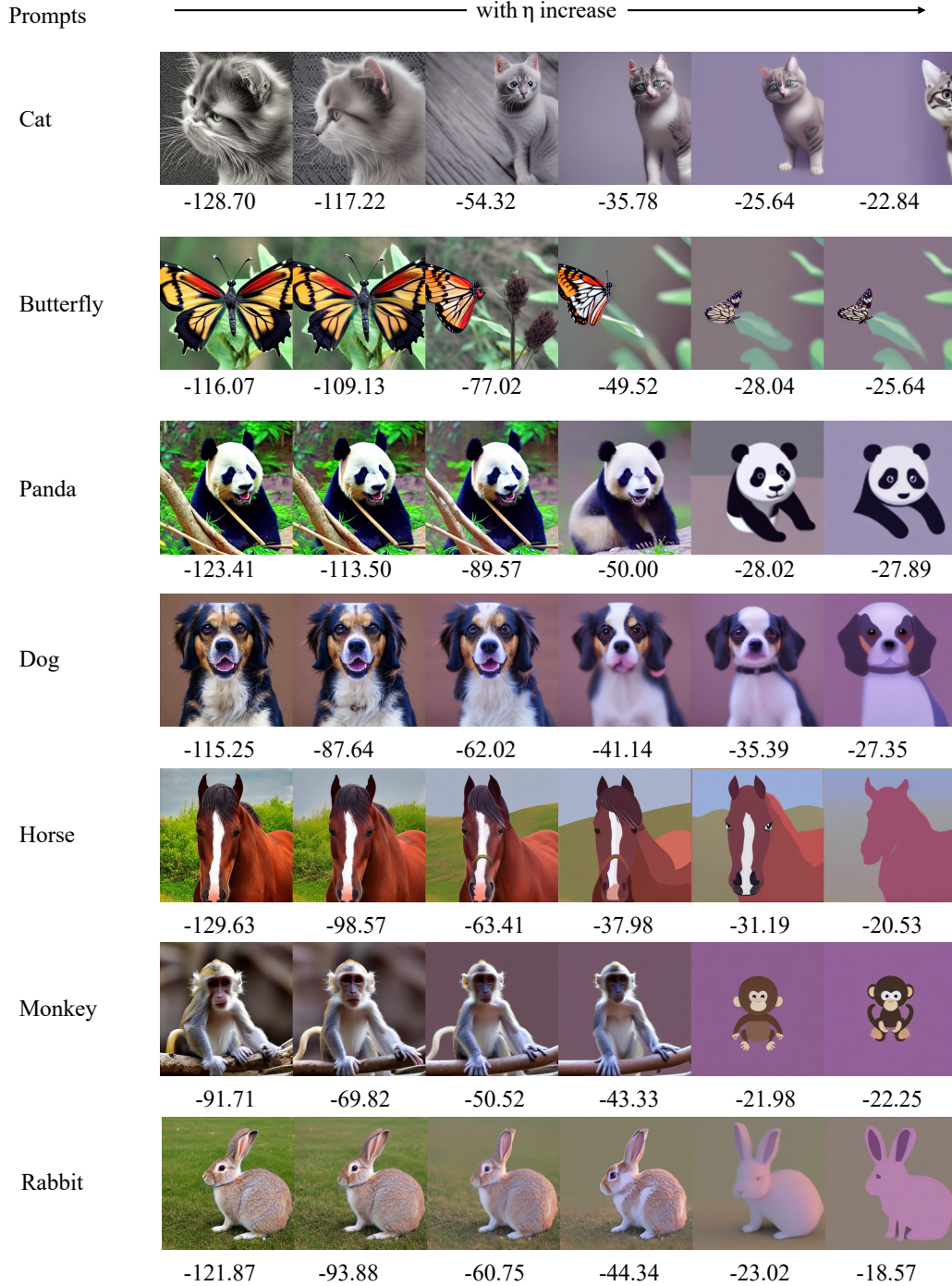


Figure 6: More images generated by SLCD with varying η values and their rewards on the image compression task.

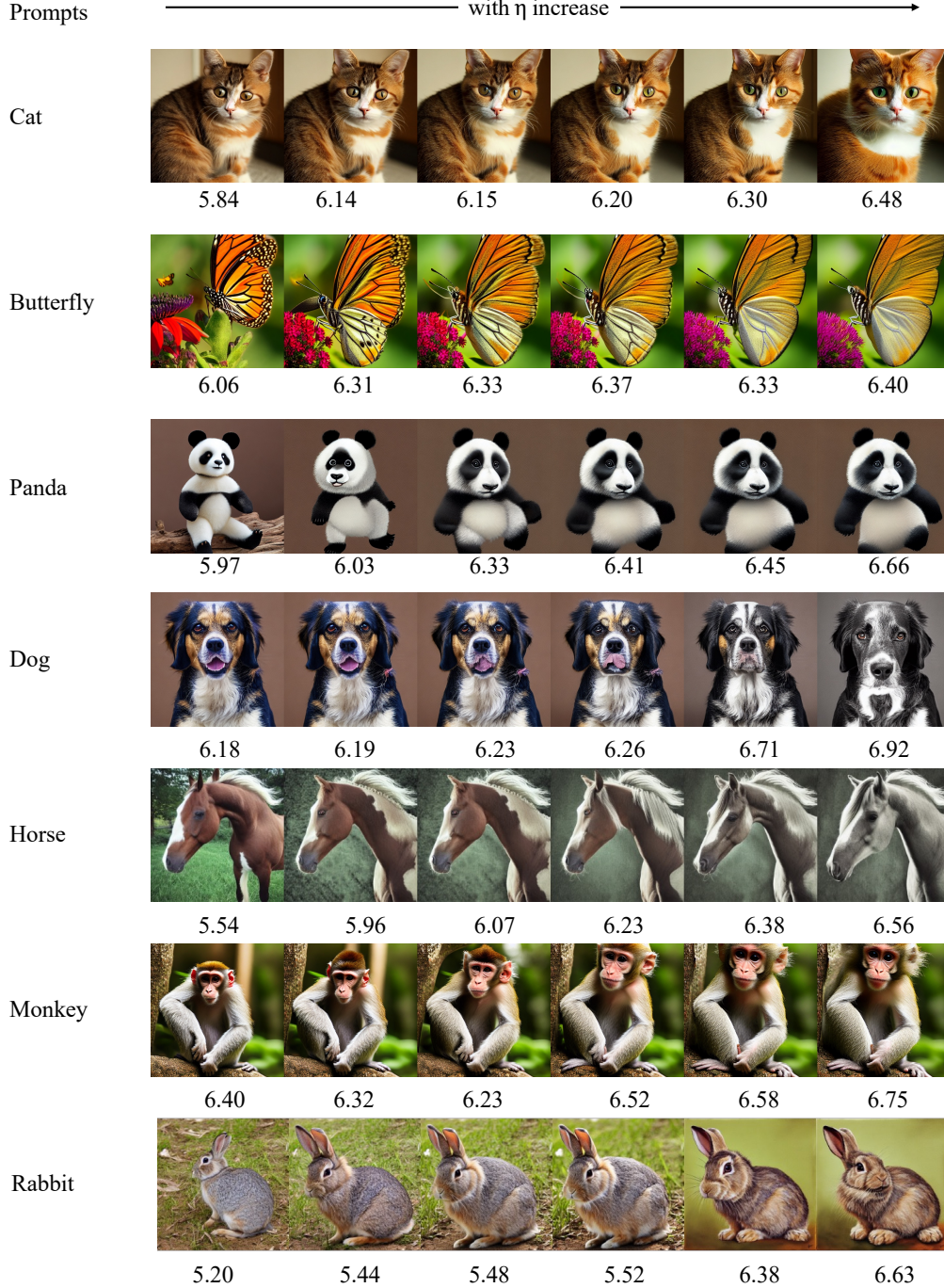


Figure 7: More images generated by SLCD with varying η values and their rewards on the image aesthetic task.