
Walking the Schrödinger Bridge: A Direct Trajectory for Text-to-3D Generation

Ziying Li
Zhejiang University
emmaleee@zju.edu.cn

Xuequan Lu
University of Western Australia
bruce.lu@uwa.edu.au

Xinkui Zhao*
Zhejiang University
zhaoxinkui@zju.edu.cn

Guanjie Cheng
Zhejiang University
chengguanjie@zju.edu.cn

Shuiguang Deng
Zhejiang University
dengsg@zju.edu.cn

Jianwei Yin
Zhejiang University
zjuyjw@cs.zju.edu.cn

<https://github.com/emmaleee789/TraCe.git>

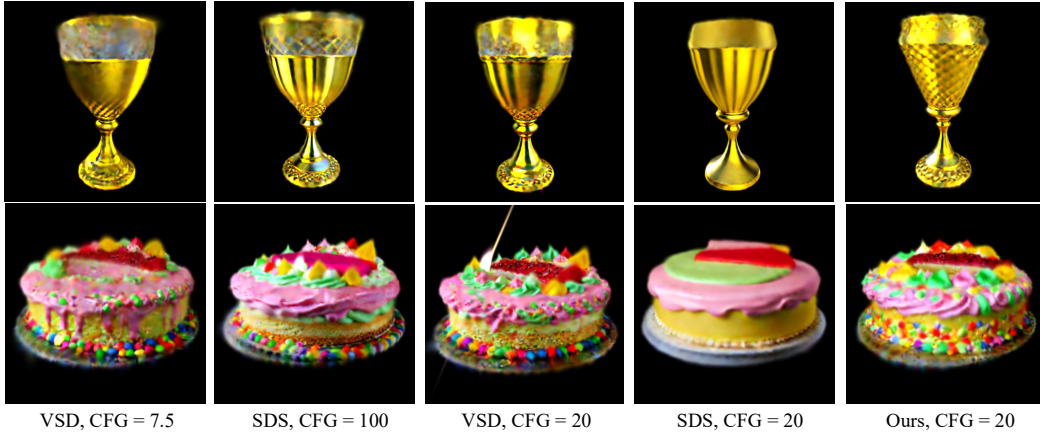


Figure 1: From left to right: (a) Standard VSD [45] (CFG = 7.5, CFG: Classifier-free Guidance); (b) Standard SDS [35]; (CFG = 100); (c) VSD [45] (CFG = 20); (d) SDS [35] (CFG = 20); (e) Ours (CFG = 20). VSD with CFG = 7.5 and CFG = 20 both yield low-quality results. Standard SDS yields artifacts (e.g., over-smoothing) with high CFG, and SDS with low CFG yields low-quality results. Our method generates high-quality and high-fidelity results with a fair CFG value.

Abstract

Recent advancements in optimization-based text-to-3D generation heavily rely on distilling knowledge from pre-trained text-to-image diffusion models using techniques like Score Distillation Sampling (SDS), which often introduce artifacts such as over-saturation and over-smoothing into the generated 3D assets. In this paper, we address this essential problem by formulating the generation process as learning an optimal, direct transport trajectory between the distribution of the current rendering and the desired target distribution, thereby enabling high-quality generation with smaller Classifier-free Guidance (CFG) values. At first, we theoretically establish SDS as a simplified instance of the Schrödinger Bridge framework. We prove that SDS employs the reverse process of an Schrödinger

*Corresponding author. zhaoxinkui@zju.edu.cn

Bridge, which, under specific conditions (e.g., a Gaussian noise as one end), collapses to SDS’s score function of the pre-trained diffusion model. Based upon this, we introduce Trajectory-Centric Distillation (TraCe), a novel text-to-3D generation framework, which reformulates the mathematically trackable framework of Schrödinger Bridge to explicitly construct a diffusion bridge from the current rendering to its text-conditioned, denoised target, and trains a LoRA-adapted model on this trajectory’s score dynamics for robust 3D optimization. Comprehensive experiments demonstrate that TraCe consistently achieves superior quality and fidelity to state-of-the-art techniques.

1 Introduction

Generating three-dimensional content directly from textual descriptions has recently attracted intensive attentions in the research community. Recent methods leveraging explicit 3D representations like Gaussian Splatting have significantly accelerated the generation process [25, 3]. Despite the advancements, it remains a key bottleneck that the quality and fidelity of generated 3D assets often lag behind their 2D counterparts. This limitation is frequently attributed to the scarcity of large-scale, high-quality 3D datasets required for direct supervised training [27, 28, 10].

To bridge this gap, many state-of-the-art text-to-3D methods employ optimization strategies guided by powerful, pre-trained 2D text-to-image (T2I) diffusion models [36]. Score Distillation Sampling (SDS) [35] has become the cornerstone paradigms. SDS leverages powerful pre-trained 2D text-to-image diffusion models to guide the optimization of 3D representations. Nevertheless, the standard SDS approach typically requires high values for Classifier-Free Guidance (CFG) [13] to achieve strong text alignment [35, 47, 4, 24, 49]. This reliance on high CFG values is often problematic, leading to visual artifacts such as over-saturation [37] and over-smoothing [23] in the generated 3D assets. Recognizing these issues, several variants of SDS have been proposed recently [45, 29, 17, 44, 48, 11, 6]. However, these SDS-based methods, including the recent variants, face persistent challenges. Firstly, as analyzed in recent studies [45, 1, 24], SDS and its variants fundamentally operate by matching the gradient direction predicted by the T2I model. While differing in their specific source and target choices for computing this gradient, they all rely on score estimates derived from the T2I backbone. These score estimates, however, can be noisy and are not guaranteed to represent an optimal direction for 3D optimization (shown in Figure 2b), potentially causing unexpected artifacts. Secondly, variants designed to operate effectively at lower CFG values (e.g., CFG=7.5), such as Score Distillation via Inversion (SDI) [29] or Variational Score Distillation (VSD) [45], have shown limited success when applied to optimizing certain popular 3D representations like 3D Gaussian Splatting (3DGS), often yielding less-desired results (shown in Figure 1).

The aforementioned analysis underscores the limitations of existing approaches and highlights the urgent need of a more robust optimization framework for text-to-3D generation, one that does not solely rely on potentially noisy score matching or operate under restrictive guidance conditions. In this paper, we first provide a theoretical insight by establishing that SDS can be understood as a simplified instance of the Schrödinger Bridge framework [39]. We demonstrate (Section 4.1) that SDS implicitly employs the reverse process of an Schrödinger Bridge, which, under specific conditions such as Gaussian noise distribution at one endpoint, effectively collapses to utilizing the score function of the pre-trained diffusion model. This perspective not only clarifies the underlying dynamics of SDS but also illuminates pathways for more principled trajectory design. Based upon this reformulation, we introduce **Trajectory-Centric Distillation (TraCe)**, a novel text-to-3D generation framework. TraCe formulates the mathematically tractable framework of Schrödinger Bridges [26, 26] to explicitly construct and learn a diffusion bridge for text-to-3D generation. This bridge connects the current rendering (X_1) to its text-conditioned, denoised target (X_0^{pred}), thereby defining a more stable and direct optimization trajectory (visualization in Figure 2a). TraCe then employs Low-Rank Adaptation (LoRA) [14] to fine-tune the T2I diffusion model specifically for navigating this constructed bridge, enabling it to precisely learn the score dynamics required for robust 3D optimization along this optimal trajectory towards the target distribution.

Our proposed TraCe framework, which operationalizes the direct transport path via Schrödinger Bridges, is rigorously evaluated. Extensive experiments demonstrate that this approach yields high-fidelity 3D assets with strong adherence to textual descriptions (Figure 4 and Table 1). The results consistently showcase TraCe’s capacity to achieve superior visual quality and semantic coherence

in generated content (Figure 4 and Supplementary), highlighting the efficacy of our theoretically grounded direct trajectory optimization for text-to-3D generation.

In summary, our contributions are:

- We establish a novel theoretical connection, demonstrating that SDS can be precisely understood as a special case of the Schrödinger Bridge framework. This reformulation clarifies the underlying transport dynamics implicitly leveraged by SDS.
- We introduce Trajectory-Centric Distillation (TraCe), a new text-to-3D generation framework. TraCe explicitly learns an optimal transport path, guided by a tractable Schrödinger Bridge formulation, between the current 3D model’s rendering and a dynamically estimated, text-aligned target view. This is achieved by constructing and sampling along this explicit diffusion bridge, enabling more direct and stable 3D optimization.
- Experiments demonstrate that our TraCe achieves high-quality 3D generation, surpassing current state-of-the-art techniques. TraCe exhibits enhanced robustness, particularly excelling in challenging low CFG values where the performance of existing methods typically degrades.

2 Related Work

Distilling 2D into 3D. Leveraging large-scale, pre-trained text-to-image (T2I) diffusion models [36] as priors has become a prominent technique for generation tasks in data-scarce domains, such as text-to-3D generation. SDS [35] is a seminal approach in this direction, enabling optimization of parametric representations (e.g., Neural Radiance Fields) by distilling knowledge from a 2D diffusion model. To achieve plausible results, it frequently necessitates high Classifier-Free Guidance (CFG) weights [35, 47], which can further exacerbate these issues. However, standard SDS is often susceptible to visual artifacts such as over-saturation [37] and over-smoothing [23]. Moreover, the SDS objective itself, while empirically effective, does not strictly correspond to the gradient of a well-defined probability distribution of the 3D parameters [45, 1, 24], potentially leading to suboptimal optimization paths [17, 44, 29, 48]. To address these limitations, several variants have been proposed. For instance, methods like Variational Score Distillation (VSD) [45] and Classifier Score Distillation (CSD) [48] explore alternative gradient formulations to better approximate the optimization process from source distribution towards target distribution. Other approaches like Score Distillation via Inversion (SDI) [29] tries to better approximate the noise instead of using pure Gaussian noise. These variants can be understood through the lens of approximating an optimal transport path between the current image distribution (source) and the target natural image distribution, and from this perspective, a key difference between these methods lies in how they approximate the score of the source and target distributions [30]. For instance, SDS approximates it using the unconditional score, while VSD attempts a more direct approximation by fine-tuning a LoRA adapter on the current renderings. While these methods offer valuable contributions towards reducing the source distribution mismatch artifacts, they fundamentally rely on adapting gradients derived from pre-trained T2I models. This forces the optimization process to cope with score functions optimized for 2D image generation, which is inherently not optimal for tasks like 3D generation due to the domain gap and differences. Our work differs greatly from these approaches. We establish a novel theoretical connection, demonstrating that SDS can be precisely understood as a specific instantiation of the Schrödinger Bridge framework. This reformulation clarifies the underlying transport dynamics implicitly leveraged by SDS. Built upon this insight, we introduce a method that explicitly constructs and learns a more direct and stable optimization trajectory by framing the process as a tractable Schrödinger Bridge between the current rendering and an estimated text-aligned target, thereby enhancing both the fidelity and robustness of text-to-3D generation.

Diffusion Models and Schrödinger Bridges. Diffusion models (DMs) [12], also known as Score-based Generative Models (SGMs) [40, 42], have emerged as a dominant class of deep generative techniques, achieving state-of-the-art performance in synthesizing high-fidelity data across various domains, notably images [40, 12, 42, 9]. These models typically define a forward diffusion process, often formulated as a stochastic differential equation (SDE), that gradually corrupts data samples into a simple prior distribution, usually Gaussian noise. A neural network is then trained, often via score-matching objectives [16, 43, 42], to approximate the score function (gradient of the log density) of the perturbed data distributions. This learned score function parameterizes a reverse-time

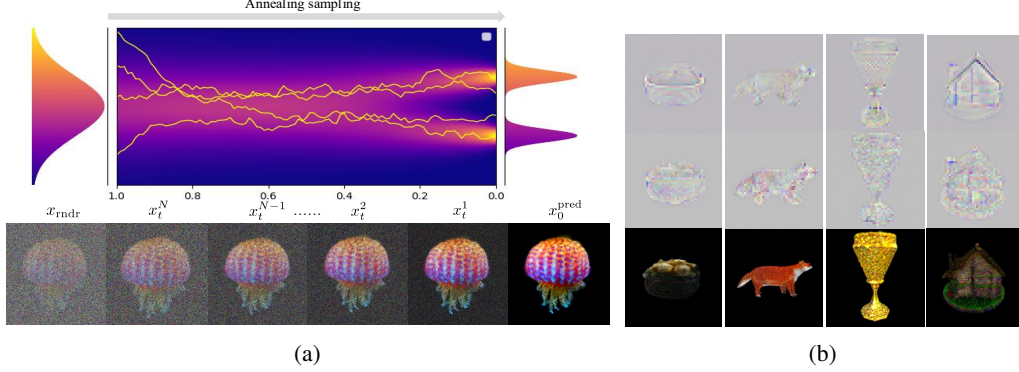


Figure 2: **Left: Schrödinger Bridge Visualization and Samples.** Top: Probability flow of the bridge from current rendering (x_{rndr}) to the predicted target (x_0^{pred}) distribution. Bottom: Corresponding image samples, showing the current rendering, intermediate bridge samples (x_t^i), and the final predicted target. **Right: Gradient and Intermediate Rendering Comparison.** The first row shows TraCe gradients, the second shows SDS gradients, and the third shows rendered images of the 3D models that have not finished generation. Note the reduced artifacts and potentially more coherent structure in the TraCe gradients and intermediate renderings.

SDE that transforms samples from the prior back into data samples. While being extremely successful, this standard paradigm typically relies on initiating the generative process from unstructured noise. The Schrödinger Bridge problem provides a more general theoretical framework, originating from statistical physics [38, 39] and connected to entropy-regularized optimal transport [21, 5] and stochastic control [7, 34]. It aims to find the most likely stochastic evolution between two specified arbitrary distributions, P_A and P_B , rather than being restricted to a noise prior. This offers the potential to learn direct transformations between complex data manifolds. Attempts have been made to apply Schrödinger Bridge concepts to text-to-3D generation. For instance, [30] proposes a naive approach to direct Schrödinger Bridge formulation between current renderings and target images guided by text prompts, though this requires an initial stage involving standard SDS. Another approach, DreamFlow [20], proposes to approximate the backward Schrödinger Bridge dynamics between current renderings and target images by simply repurposing a fine-tuned text-to-image model, a heuristic potentially deviating from the true underlying Schrödinger Bridge process. We critically advance text-to-3D generation by establishing the precise theoretical relationship between SDS and Schrödinger Bridges. This foundational insight is then exploited to develop a principled methodology for direct distributional transport, enabling the construction of trajectories towards text-aligned target distributions.

3 Preliminaries

Score-based Generative Model (SGM) and Schrödinger Bridge. Score-based Generative Models (SGM) [40, 42] learn to generate data by reversing a predefined forward diffusion process. This process gradually transforms data $X_0 \sim p_A$ into noise $X_1 \approx \mathcal{N}(0, I)$ and is often governed by a forward stochastic differential equation (SDE). Generation then proceeds by simulating the corresponding reverse-time SDE [2], starting from X_1 and integrating backward to $t = 0$. The forward and reverse SDEs are given by:

$$\begin{aligned} dX_t &= f_t(X_t)dt + g_t dW_t \text{ (forward)} \\ dX_t &= [f_t(X_t) - g_t^2 \nabla_{X_t} \log p(X_t, t)] dt + g_t d\bar{W}_t \text{ (backward)} \end{aligned}$$

Here, W_t (and $\bar{W}_t$) is a standard Wiener process, and g_t represents the time-dependent diffusion coefficient. The central part of this reversal is the score function $\nabla_{X_t} \log p(X_t, t)$, which is unknown and approximated using a time-conditioned neural network $s_\psi(X_t, t)$ (or an equivalent noise predictor $\epsilon_\psi(X_t, t)$). This network is trained using score-matching objectives [43, 42] on pairs (X_0, X_t) sampled from the forward process. Sampling is performed by numerically integrating the reverse SDE using solvers like DDPM [12] or DDIM [41].

The Schrödinger Bridge problem [39, 21] offers a generalization of SGMs to learn nonlinear diffusion processes between two arbitrary distributions, $X_0 \sim p_A$ and $X_1 \sim p_B$. It seeks the most likely stochastic evolution connecting these boundary distributions, described by a pair of forward and backward SDEs:

$$\begin{aligned} dX_t &= [f_t(X_t) + \beta_t \nabla \Psi(X_t, t)] dt + \sqrt{\beta_t} dW_t \text{ (forward)} \\ dX_t &= [f_t(X_t) - \beta_t \nabla \hat{\Psi}(X_t, t)] dt + \sqrt{\beta_t} d\bar{W}_t \text{ (backward)} \end{aligned}$$

where $\Psi(x, t)$ and $\hat{\Psi}(x, t)$ are non-negative functions known as Schrödinger factors, determined by coupled partial differential equations with boundary conditions $\Psi(x, 0)\hat{\Psi}(x, 0) = p_A(x)$ and $\Psi(x, 1)\hat{\Psi}(x, 1) = p_B(x)$. The forward and backward processes induce the same marginal density $q(x, t)$ at any time $t \in [0, 1]$, satisfying Nelson’s duality $\Psi(x, t)\hat{\Psi}(x, t) = q(x, t)$ [33]. Notably, SGM is a special case where $p_B \approx \mathcal{N}(0, I)$ and $\Psi(x, t) \approx 1$, causing the forward drift modification to vanish and $\hat{\Psi}(x, t) \approx q(x, t)$, recovering the score function in the reverse SDE.

Score Distillation Sampling (SDS). Score Distillation Sampling (SDS) [35] enables generating 3D assets by leveraging powerful pre-trained 2D text-to-image diffusion models [36], bypassing the need for large-scale 3D datasets. It optimizes the parameters θ of a differentiable 3D representation, such as NeRF [31], InstantNGP [32], or 3D Gaussian Splatting (3DGS) [18], using gradients derived from the diffusion model. In this work, we adopt 3DGS primarily for its rapid generation capabilities and high-fidelity visual output.

The core mechanism of SDS involves repeatedly rendering the 3D model from different views c ($x = g(\theta, c)$), adding noise to the rendering $x(t)$, and using the 2D diffusion model’s score estimate (denoising prediction ϵ_{pred}) to guide the optimization of θ . Formally, the gradient is computed as

$$\nabla_{\theta} \mathcal{L}_{\text{SDS}}(\theta) = \mathbb{E}_{t, \epsilon, c} \left[w(t) (\epsilon_{\text{pred}} - \epsilon_{\text{noise}}) \frac{\partial x_{\text{rndr}}}{\partial \theta} \right] \quad (1)$$

where $w(t)$ is a weighting factor and the term $(\epsilon_{\text{pred}} - \epsilon_{\text{noise}})$ provides the guidance signal. While SDS can be intuitively understood as moving renderings towards higher-density regions according to the 2D prior or formally interpreted via probability density distillation, the exact nature of its gradient signal is debated [17, 48, 44, 1, 45]. Practically, SDS often requires high classifier-free guidance (CFG) values, which can sometimes lead to artifacts like oversaturation or blur [17, 29, 46, 19, 37, 22]. Furthermore, the strategies that employ lower CFG values, for instance, methods explored in text-to-NeRF [45, 29], have demonstrated limitations when directly applied to the generation of 3D assets with Gaussian Splatting (Figure 4). Recent efforts such as LucidDreamer [23] have investigated text-to-3DGS under low CFG conditions; however, this direction currently faces trade-offs, including prolonged optimization durations (over 5000 iterations) and limitations in the attainable visual quality (Figure 4). Our work builds upon SDS by mitigating these issues through deriving a more direct and tractable optimization path, formulating Schrödinger Bridges to guide the generation process for achieving greater fidelity with lower CFG values.

4 Method

The preceding analysis of existing methods like Score Distillation Sampling (SDS) raise a natural question: can a principled framework be developed to define a direct, optimal transformation trajectory where both its source and target ends are explicitly and robustly aligned with the desired true distributions? Addressing this challenge—by establishing explicit control over the distributional endpoints of the generative trajectory, rather than relying on unstructured priors (e.g., a Gaussian noise)—is crucial for enhancing the fidelity and control of generative outcomes. To this end, we exploit the theoretical underpinnings of the Schrödinger Bridge problem, particularly its tractable formulations [26, 8], which provide a robust mechanism for learning direct, optimal transport paths between specified distributions. Our methodological contribution unfolds in two stages: first, we theoretically establish that standard SDS is indeed a special case of the Schrödinger Bridge framework, thereby providing a new perspective on its operation (Section 4.1). Second, building upon this insight, we propose a novel optimization algorithm grounded in tractable Schrödinger Bridge principles, to achieve improved distributional alignment throughout the generative process (Section 4.2).

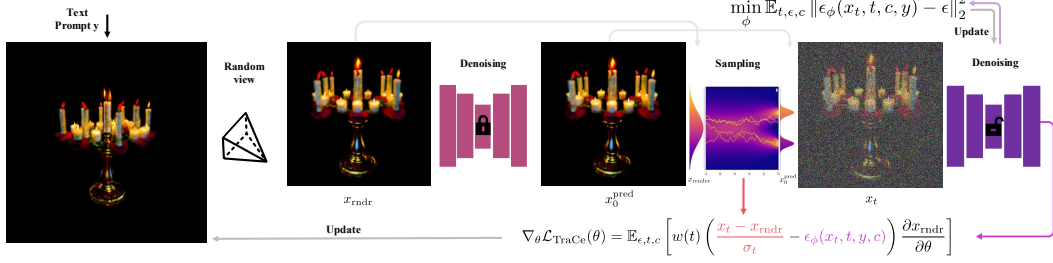


Figure 3: **Overview of Trajectory-Centric Distillation (TraCe).** Our TraCe optimizes 3D parameters θ by computing a distillation gradient with a LoRA-adapted 2D diffusion model, ϵ_ϕ . Given a text prompt y and camera parameters c , (1) the current 3D model is rendered in a random view to produce x_{rndr} . (2) An ideal target view x_0^{pred} is estimated from x_{rndr} using a pre-trained diffusion model $\epsilon_{\text{pretrain}}$ via one-step denoising. (3) An intermediate latent x_t is sampled from the analytic bridge posterior $q(x_t | x_0^{\text{pred}}, x_{\text{rndr}})$ at time t . (4) The LoRA model ϵ_ϕ predicts the noise for x_t , and the difference between this prediction and the target noise is computed. (5) This difference directs the calculation of the TraCe gradient $\nabla_\theta \mathcal{L}_{\text{TraCe}}$, and drives the update of LoRA parameters ϕ .

4.1 Score Distillation Sampling as a Special Case of Schrödinger Bridges

In this section, we reformulate the SDS objective by examining its core guidance principles, and show it employs a simplified form of the backward dynamics found in the Schrödinger Bridge framework.

As established in Section 3, a Score-based Generative Model (SGM) aligns with a special configuration of the Schrödinger Bridge problem. This occurs when the Schrödinger Bridge’s distribution P_B at $t = 1$ is Gaussian noise ($P_B \sim \mathcal{N}(0, I)$) and its forward Schrödinger factor $\Psi(x, t) \approx 1$. Under these conditions, the term $g_t^2 \nabla_{x_t} \log \Psi(X_t, t)$ in the forward Schrödinger Bridge SDE vanishes, causing the forward Schrödinger Bridge dynamics to become identical to the SGM’s standard diffusion process. Consequently, the marginal densities $q(X_t, t)$ of this particular Schrödinger Bridge are equivalent to the SGM’s noisy marginals $p(X_t, t)$.

The crucial step in linking the Schrödinger Bridge and SGM reverse processes from a score perspective lies in Nelson’s duality, $\Psi(X_t, t) \hat{\Psi}(X_t, t) = q(X_t, t)$. Given $\Psi(X_t, t) \approx 1$ and $q(X_t, t) = p(X_t, t)$ for this specific Schrödinger Bridge, the duality simplifies to:

$$1 \cdot \hat{\Psi}(X_t, t) \approx p(X_t, t) \implies \hat{\Psi}(X_t, t) \approx p(X_t, t) \quad (2)$$

This directly implies that the score term in the general Schrödinger Bridge backward SDE, $-\nabla_{x_t} \log \hat{\Psi}(X_t, t)$, becomes $-\nabla_{x_t} \log p(X_t, t)$. This is precisely the score approximated by the learned network $s_\psi(X_t, t)$ (or its equivalent noise predictor $\epsilon(X_t, t)$) in an SGM.

SDS utilizes this learned score $s_\psi(X_t, t)$ from a pre-trained SGM to guide the optimization of a differentiable generator $g(\theta)$. The update for $g(\theta)$ is fundamentally derived from $s_\phi(X_t, t)$, aiming to make the generated samples $x_0 = g(\theta)$ consistent with the data manifold learned by the SGM.

Therefore, from a score gradient perspective:

- SDS operates using the score function $s_\psi(X_t, t)$ learned by an SGM.
- The derivation above shows that $s_\psi(X_t, t)$ (approximating $\nabla_{x_t} \log p(X_t, t)$) is equivalent to the score $-\nabla_{x_t} \log \hat{\Psi}(X_t, t)$ of a Schrödinger Bridge under the specific conditions that reduce the Schrödinger Bridge to an SGM.

Remark. In essence, SDS leverages a score gradient that is equivalent to the score function governing the reverse dynamics of the canonical Schrödinger Bridge implicit in any SGM. While general Schrödinger Bridges can offer more complex dynamics, SDS employs the score from this specific, simplified Schrödinger Bridge structure. Thus, the SDS mechanism represents an application of principles governing a special case of Schrödinger Bridges, distinguished by its reliance on the SGM-derived score s_ψ .

4.2 Trajectory-Centric Distillation

To optimize the 3D model parameters θ such that current renderings $x_{\text{rndr}} = g(\theta, c)$ align with a target text description y , we propose a novel method, namely Trajectory-Centric Distillation (TraCe). This method leverages a 2D diffusion model, adapted with LoRA parameters ϕ denoted as ϵ_ϕ , to provide a guiding gradient $\nabla_\theta \mathcal{L}_{\text{TraCe}}(\theta)$. The core idea is to conceptualize a diffusion bridge between the current rendering and an estimated ideal target image.

Constructing the Diffusion Bridge for Trajectory Guidance. At each optimization step for θ , we construct a specific diffusion bridge instance defined by two endpoints:

1. **Initial Bridge Endpoint** ($X_1 \leftarrow x_{\text{rndr}}$): The current rendering $x_{\text{rndr}} = g(\theta, c)$ serves as the starting point of the reverse diffusion trajectory we aim to learn. In the context of our bridge, this is treated as the “noisier” state at bridge time $t = 1$.
2. **Target Bridge Endpoint** ($X_0 \leftarrow x_0^{\text{pred}}$): An estimated ideal target view x_0^{pred} acts as the desired endpoint at bridge time $t = 0$. This target is dynamically obtained by performing one-step denoising on x_{rndr} using a pre-trained text-to-image model $\epsilon_{\text{pretrain}}$ [20], conditioned on the text prompt y : $x_0^{\text{pred}} = (x_{\text{rndr}} - \sqrt{1 - \bar{\alpha}_{t'}} \epsilon_{\text{pretrain}}(x_{\text{rndr}}, t', y)) / \sqrt{\bar{\alpha}_{t'}}$, where $\bar{\alpha}_{t'}$ is from the noise schedule of $\epsilon_{\text{pretrain}}$.

With these two endpoints, x_0^{pred} and x_{rndr} , established, we then sample an intermediate latent state x_t along the conceptual bridge. For a sampled time $t \in [0.02, 0.5]$, following the tractable formulation of Schrödinger Bridges [26], x_t is drawn from the analytically known conditional distribution $x_t \sim q(x_t | x_0^{\text{pred}}, x_{\text{rndr}}) = \mathcal{N}(x_t; \mu_t, \Sigma_t I)$, where the mean $\mu_t = \gamma_t x_0^{\text{pred}} + (1 - \gamma_t) x_{\text{rndr}}$ is an interpolation between the target image and current rendering, and $\Sigma_t = \sigma_t^2 \bar{\sigma}_t^2 / (\sigma_t^2 + \bar{\sigma}_t^2)$ is the bridge variance. The coefficient $\gamma_t = \bar{\sigma}_t^2 / (\sigma_t^2 + \bar{\sigma}_t^2)$, and $\sigma_t^2 = \int_0^t \beta_\tau d\tau$, $\bar{\sigma}_t^2 = \int_t^1 \beta_\tau d\tau$ are accumulated variances from a noise schedule β_t specific to this bridge construction. This x_t represents a state on a direct trajectory from x_0^{pred} being progressively “noised” towards x_{rndr} (or equivalently, x_{rndr} being progressively “denoised” towards x_0^{pred} along this trajectory).

Optimizing θ via the Bridge Trajectory. We then optimize θ using the LoRA-adapted model $\epsilon_\phi(x_t, t, y, c)$, which is trained to predict the noise that would take x_t towards x_0^{pred} . The objective for θ utilizes ϵ_ϕ to measure the consistency of x_t with respect to x_{rndr} along this bridge:

$$\nabla_\theta \mathcal{L}_{\text{TraCe}}(\theta) = \mathbb{E}_{\epsilon, t, c} \left[w(t) \left(\epsilon_\phi(x_t, t, y, c) - \frac{x_t - x_{\text{rndr}}}{\sigma_t} \right) \left(\underbrace{\frac{\partial x_0^{\text{pred}}(x_{\text{rndr}}, t, y)}{\partial x_t}}_{\text{U-net Jacobian}} \frac{\partial x_t}{\partial x_{\text{rndr}}} + 1 \right) \frac{\partial x_{\text{rndr}}}{\partial \theta} \right] \quad (3)$$

where $x_{\text{rndr}} = g(\theta, c)$, $t \sim \mathcal{U}[0.02, 0.5]$, and y is the text prompt. The term x_t is sampled from $q(x_t | x_0^{\text{pred}}, x_{\text{rndr}})$ as defined above, and $\sigma_t = \sqrt{\int_0^t \beta_\tau d\tau}$ from the bridge’s noise schedule. Following the convention of SDS, we omit the U-Net Jacobian term $\left(\frac{\partial x_0^{\text{pred}}(\dots)}{\partial x_t} \frac{\partial x_t}{\partial x_{\text{rndr}}} + 1 \right)$ for effective training, as it can be treated as a learnable or constant factor absorbed by $w(t)$. Thus, we have:

$$\nabla_\theta \mathcal{L}_{\text{TraCe}}(\theta) = \mathbb{E}_{\epsilon, t, c} \left[w(t) \left(\epsilon_\phi(x_t, t, y, c) - \frac{x_t - x_{\text{rndr}}}{\sigma_t} \right) \frac{\partial x_{\text{rndr}}}{\partial \theta} \right] \quad (4)$$

Scheduled t -Sampling for Schrödinger Bridges Interpolation. For sampling the intermediate state x_t in our TraCe objective (Eq. (4)), which dictates the characteristics of $x_t \sim q(x_t | x_0^{\text{pred}}, x_{\text{rndr}})$, we adopt a t -annealing strategy, similar to the approach proposed in [15]. Throughout the optimization of θ , the time parameter t is progressively decreased from an initial value near 0.5 towards 0.02. This common annealing technique gradually shifts the focus of the Schrödinger Bridge interpolation from broader states towards those more proximate to the estimated ideal target x_0^{pred} , aiding the progressive refinement of the rendered output $g(\theta, c)$.

5 Experiments



Figure 4: **Qualitative comparisons.** We present visual examples with the same text prompt.

Implementation Details. We choose recent state-of-the-art (SOTA) text-to-3D approaches for comparison: NeRF-based methods, such as Classifier Score Distillation (CSD) [48], ProlificDreamer (VSD) [45], and Score Distillation via Inversion (SDI) [29], and 3DGS-based methods like GaussianDreamer (SDS) [47] and LucidDreamer (ISM) [23]. Please see more details and experiments in Supplementary.

Qualitative Comparisons. Figure 4 presents visual results for several challenging text prompts. Our approach demonstrates the ability to generate higher quality 3D assets compared to other SOTA methods. Compared to SDI [29], our method yields significantly improved texture fidelity. Outputs from CSD [48] often exhibit a characteristic yellowish hue and a less realistic, cartoon-like appearance, which TraCe avoids, producing more natural color rendition and photorealism. When compared against VSD [45], our model better interprets complex textural and stylistic prompts, accurately capturing the text’s message and generating a more coherent content. Contrasting with SDS [47], our results exhibit superior sharpness and finer details in both geometry and texture, leading to more visually appealing and realistic outputs. While ISM [23] can produce coherent structures, its outputs often exhibit a stylized, painterly quality; in contrast, our TraCe generates

results with enhanced photorealism and more natural material appearance. These results demonstrate our method’s effectiveness in generating detailed and accurate 3D geometry and appearance from the given text descriptions.

Quantitative Comparison. We quantitatively evaluate our TraCe against other methods using 83 distinct prompts from Dreamfusion online gallery² with 120 views per prompt. We benchmark generation quality using CLIP Score (%), GPTEval3D (Overall) (which leverages GPT-4o for evaluation), and ImageReward. CLIP Scores are evaluated with ViT-L/14, ViT-B/16, and ViT-B/32 backbones. We also assess computational efficiency via processing time (Time) and average peak VRAM (VRAM). As shown in Table 1, the proposed TraCe achieves state-of-the-art generation quality, securing top CLIP Scores across all ViT backbones, e.g., $69.2609 \pm 7.8366\%$ with ViT-L/14. Furthermore, TraCe demonstrates superior performance in advanced perception metrics, yielding the highest GPTEval3D score of 1028.03 and the most favorable (least negative) ImageReward score of -0.2855 ± 0.8909 , indicating enhanced aesthetic quality and semantic alignment. With an average processing time of 14 minutes and an average peak VRAM usage of 18741 MiB, TraCe offers high-fidelity generation with a compelling balance of qualitative performance, computational efficiency, and memory footprint.

Table 1: **Quantitative comparisons.** Comparison of different methods on CLIP Score, GPTEval3D Score, ImageReward Score, running time, and VRAM usage. We report mean and standard deviation across 83 prompts and 120 views.

Method	CLIP Score (%) $\uparrow$			GPTEval3D (Overall) $\uparrow$	ImageReward $\uparrow$	Time	VRAM
	ViT-L/14	ViT-B/16	ViT-B/32				
SDS [47]	68.6146 $\pm$ 7.9134	27.7049 $\pm$ 3.7004	27.5561 $\pm$ 3.5893	1018.09	-0.4329 $\pm$ 0.9125	10min	18147MiB
CSD [48]	68.0282 $\pm$ 7.5093	27.0886 $\pm$ 3.7342	26.5844 $\pm$ 3.8703	983.04	-0.6715 $\pm$ 0.7482	11min	19804MiB
VSD [45]	67.2697 $\pm$ 8.5573	27.0749 $\pm$ 3.9675	26.9722 $\pm$ 3.9563	1007.49	-0.5330 $\pm$ 0.8927	17min	26473MiB
ISM [23]	69.0093 $\pm$ 10.2400	27.5460 $\pm$ 3.6817	26.9822 $\pm$ 3.5495	1012.37	-0.3904 $\pm$ 0.9503	20min	10151MiB
SDI [29]	63.0409 $\pm$ 11.7841	25.6487 $\pm$ 5.2540	25.5421 $\pm$ 5.0903	971.98	-0.8334 $\pm$ 1.0391	10min	16011MiB
TraCe	69.2609$\pm$7.8366	27.9334$\pm$3.7382	27.7049$\pm$3.8671	1028.03	-0.2855$\pm$0.8909	14min	18741MiB

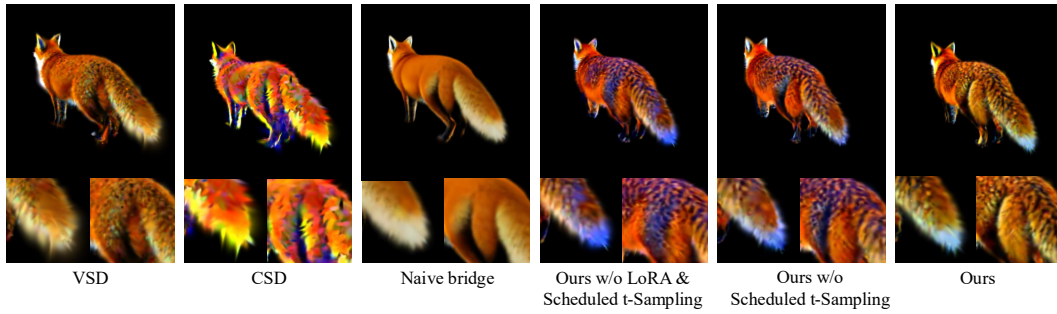


Figure 5: Ablation study on our framework.

Ablation Study. Figure 5 showcases the ablation study of our TraCe on a fox generation. VSD [45] and CSD [48] exhibit less-desired generation (e.g., missing details). The third column illustrates a naive Schrödinger Bridge approach [30] which attempts to bridge distributions defined by source and target prompts and results in a comparatively smoother, less detailed rendering. The fourth column shows TraCe without LoRA adaptation and without our scheduled t -sampling, where noticeable artifacts such as blue hues on the fur are apparent. Introducing LoRA but omitting the scheduled t -sampling (fifth column) mitigates some artifacts, yet color inconsistencies persist. Finally, our full TraCe method (“Ours”)—supported by LoRA-adapted learning of its specific score dynamics and an annealed t -sampling schedule—generates significantly higher-fidelity details in the fur and tail, boosting overall realism compared to other methods (VSD, CSD) and ablated versions. These results highlight the role of our core Schrödinger Bridge formulation: it achieves superior final quality when augmented with these tailored learning components.

²<https://dreamfusion3d.github.io/gallery.html>

Table 2: ImageReward ablation over LoRA and scheduled t -sampling.

Method Configuration	ImageReward ($\uparrow$)
LoRA off & scheduled t -sampling off	-0.4488 ± 0.9964
LoRA off & scheduled t -sampling on	-0.3389 ± 0.9721
LoRA on & scheduled t -sampling off	-0.4020 ± 1.0019
LoRA on & scheduled t -sampling on (ours)	-0.2486 ± 0.8909

We perform an ablation study on our key components, LoRA adaptation and scheduled t -sampling, measuring quality with ImageReward (Table 2). Our full method (-0.2486) significantly outperforms the baseline (both off: -0.4488), as well as enabling only LoRA (-0.4020) or only scheduled t -sampling (-0.3389). The results confirm both components are crucial and demonstrate their strong synergistic effect.

CFG value. We investigate the impact of the CFG value on our TraCe, as illustrated in Figure 6 with two example objects. While very low CFG values (e.g., 5) yield reduced visual fidelity, TraCe produces high-quality, well-defined results starting at a CFG of approximately 15-20. The visual outcomes are stable and robust within the CFG 15-20 range. Beyond this, at higher CFG values (25-100), results remain largely consistent with minimal further improvement. This demonstrates TraCe’s capability to effectively generate high-quality 3D assets at relatively low and stable CFG settings. Furthermore, TraCe’s enhanced visual quality is complemented by its robust CLIP score performance within a moderate CFG range (e.g., 10-30) relative to other compared methods, as detailed in Figure ??.

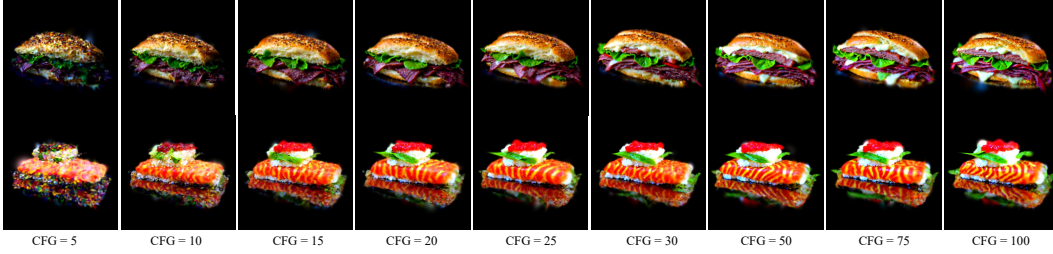


Figure 6: **Different CFG value and generated 3D assets.** Prompts are “an overstuffed pastrami sandwich” (top row), “a car made out of sushi” (bottom row).

6 Conclusion

We introduce Trajectory-Centric Distillation (TraCe), a novel text-to-3D generation framework. Our approach is rooted in a new theoretical understanding of SDS as a specific instance of the Schrödinger Bridge problem. The proposed TraCe explicitly constructs and learns a direct diffusion bridge between current renderings and text-conditioned targets, employing a LoRA-adapted diffusion model to accurately model the bridge’s score dynamics. Comprehensive experiments demonstrate TraCe’s state-of-the-art performance, yielding 3D assets with superior visual quality and fidelity, notably at lower and more stable Classifier-Free Guidance values than prior methods. These results underscore the benefits of our principled, direct optimization trajectory. We believe TraCe will offer new insights for text-to-3D generation, in terms of efficient and robust trajectory learning for generative models.

7 Acknowledgments

This work was supported in part by the National Science Foundation of China under Grants (62472375), and in part by the Major Program of National Natural Science Foundation of Zhejiang (LD24F020014, LD25F020002), and in part by the Zhejiang Pioneer (Jianbing) Project (2024C01032), and in part by the Ningbo Yongjiang Talent Programme(2023A-198-G).

References

- [1] Thiemo Alldieck, Nikos Kolotouros, and Cristian Sminchisescu. Score distillation sampling with learned manifold corrective. In *European Conference on Computer Vision*, pages 1–18. Springer, 2024.
- [2] Brian DO Anderson. Reverse-time diffusion equation models. *Stochastic Processes and their Applications*, 12(3):313–326, 1982.
- [3] Yuanhao Cai, He Zhang, Kai Zhang, Yixun Liang, Mengwei Ren, Fujun Luan, Qing Liu, Soo Ye Kim, Jianming Zhang, Zhifei Zhang, et al. Baking gaussian splatting into diffusion denoiser for fast and scalable single-stage image-to-3d generation. *arXiv preprint arXiv:2411.14384*, 2024.
- [4] Rui Chen, Yongwei Chen, Ningxin Jiao, and Kui Jia. Fantasia3d: Disentangling geometry and appearance for high-quality text-to-3d content creation. In *Proceedings of the IEEE/CVF international conference on computer vision*, pages 22246–22256, 2023.
- [5] 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.
- [6] Zilong Chen, Feng Wang, Yikai Wang, and Huaping Liu. Text-to-3d using gaussian splatting. In *Proceedings of the IEEE/CVF conference on computer vision and pattern recognition*, pages 21401–21412, 2024.
- [7] Paolo Dai Pra. A stochastic control approach to reciprocal diffusion processes. *Applied mathematics and Optimization*, 23(1):313–329, 1991.
- [8] Valentin De Bortoli, James Thornton, Jeremy Heng, and Arnaud Doucet. Diffusion schrödinger bridge with applications to score-based generative modeling. *Advances in Neural Information Processing Systems*, 34:17695–17709, 2021.
- [9] Prafulla Dhariwal and Alexander Nichol. Diffusion models beat gans on image synthesis. *Advances in neural information processing systems*, 34:8780–8794, 2021.
- [10] Rao Fu, Xiao Zhan, Yiwen Chen, Daniel Ritchie, and Srinath Sridhar. Shapecrafter: A recursive text-conditioned 3d shape generation model. *Advances in Neural Information Processing Systems*, 35:8882–8895, 2022.
- [11] Amir Hertz, Kfir Aberman, and Daniel Cohen-Or. Delta denoising score. In *Proceedings of the IEEE/CVF International Conference on Computer Vision*, pages 2328–2337, 2023.
- [12] Jonathan Ho, Ajay Jain, and Pieter Abbeel. Denoising diffusion probabilistic models. *Advances in neural information processing systems*, 33:6840–6851, 2020.
- [13] Jonathan Ho and Tim Salimans. Classifier-free diffusion guidance. *arXiv preprint arXiv:2207.12598*, 2022.
- [14] Edward J Hu, Yelong Shen, Phillip Wallis, Zeyuan Allen-Zhu, Yanzhi Li, Shean Wang, Lu Wang, Weizhu Chen, et al. Lora: Low-rank adaptation of large language models. *ICLR*, 1(2):3, 2022.
- [15] Yukun Huang, Jianan Wang, Yukai Shi, Boshi Tang, Xianbiao Qi, and Lei Zhang. Dream-time: An improved optimization strategy for diffusion-guided 3d generation. *arXiv preprint arXiv:2306.12422*, 2023.
- [16] Aapo Hyvärinen and Peter Dayan. Estimation of non-normalized statistical models by score matching. *Journal of Machine Learning Research*, 6(4), 2005.
- [17] Oren Katzir, Or Patashnik, Daniel Cohen-Or, and Dani Lischinski. Noise-free score distillation. *arXiv preprint arXiv:2310.17590*, 2023.
- [18] Bernhard Kerbl, Georgios Kopanas, Thomas Leimkühler, and George Drettakis. 3d gaussian splatting for real-time radiance field rendering. *ACM Trans. Graph.*, 42(4):139–1, 2023.

- [19] Inhee Lee, Byungjun Kim, and Hanbyul Joo. Guess the unseen: Dynamic 3d scene reconstruction from partial 2d glimpses. In *Proceedings of the IEEE/CVF conference on computer vision and pattern recognition*, pages 1062–1071, 2024.
- [20] Kyungmin Lee, Kihyuk Sohn, and Jinwoo Shin. Dreamflow: High-quality text-to-3d generation by approximating probability flow. *arXiv preprint arXiv:2403.14966*, 2024.
- [21] Christian Léonard. A survey of the schrödinger problem and some of its connections with optimal transport. *arXiv preprint arXiv:1308.0215*, 2013.
- [22] Zongrui Li, Minghui Hu, Qian Zheng, and Xudong Jiang. Connecting consistency distillation to score distillation for text-to-3d generation. In *European Conference on Computer Vision*, pages 274–291. Springer, 2024.
- [23] Yixun Liang, Xin Yang, Jiantao Lin, Haodong Li, Xiaogang Xu, and Yingcong Chen. Lucid-dreamer: Towards high-fidelity text-to-3d generation via interval score matching. In *Proceedings of the IEEE/CVF conference on computer vision and pattern recognition*, pages 6517–6526, 2024.
- [24] Chen-Hsuan Lin, Jun Gao, Luming Tang, Towaki Takikawa, Xiaohui Zeng, Xun Huang, Karsten Kreis, Sanja Fidler, Ming-Yu Liu, and Tsung-Yi Lin. Magic3d: High-resolution text-to-3d content creation. In *Proceedings of the IEEE/CVF conference on computer vision and pattern recognition*, pages 300–309, 2023.
- [25] Chenguo Lin, Panwang Pan, Bangbang Yang, Zeming Li, and Yadong Mu. Diffsplat: Repurposing image diffusion models for scalable gaussian splat generation. *arXiv preprint arXiv:2501.16764*, 2025.
- [26] Guan-Horng Liu, Arash Vahdat, De-An Huang, Evangelos A Theodorou, Weili Nie, and Anima Anandkumar. I²sb: Image-to-image schrödinger bridge. *arXiv preprint arXiv:2302.05872*, 2023.
- [27] Jian Liu, Xiaoshui Huang, Tianyu Huang, Lu Chen, Yuenan Hou, Shixiang Tang, Ziwei Liu, Wanli Ouyang, Wangmeng Zuo, Junjun Jiang, et al. A comprehensive survey on 3d content generation. *arXiv preprint arXiv:2402.01166*, 2024.
- [28] Qihao Liu, Yi Zhang, Song Bai, Adam Kortylewski, and Alan Yuille. Direct-3d: Learning direct text-to-3d generation on massive noisy 3d data. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, pages 6881–6891, 2024.
- [29] Artem Lukoianov, Haitz Sáez de Ocáriz Borde, Kristjan Greenewald, Vitor Guizilini, Timur Bagautdinov, Vincent Sitzmann, and Justin M Solomon. Score distillation via reparametrized ddim. *Advances in Neural Information Processing Systems*, 37:26011–26044, 2024.
- [30] David McAllister, Songwei Ge, Jia-Bin Huang, David Jacobs, Alexei Efros, Aleksander Holynski, and Angjoo Kanazawa. Rethinking score distillation as a bridge between image distributions. *Advances in Neural Information Processing Systems*, 37:33779–33804, 2024.
- [31] Ben Mildenhall, Pratul P Srinivasan, Matthew Tancik, Jonathan T Barron, Ravi Ramamoorthi, and Ren Ng. Nerf: Representing scenes as neural radiance fields for view synthesis. *Communications of the ACM*, 65(1):99–106, 2021.
- [32] Thomas Müller, Alex Evans, Christoph Schied, and Alexander Keller. Instant neural graphics primitives with a multiresolution hash encoding. *ACM transactions on graphics (TOG)*, 41(4):1–15, 2022.
- [33] Edward Nelson. *Dynamical theories of Brownian motion*, volume 106. Princeton university press, 2020.
- [34] Michele Pavon and Anton Wakolbinger. On free energy, stochastic control, and schrödinger processes. In *Modeling, Estimation and Control of Systems with Uncertainty: Proceedings of a Conference held in Sopron, Hungary, September 1990*, pages 334–348. Springer, 1991.

- [35] Ben Poole, Ajay Jain, Jonathan T Barron, and Ben Mildenhall. Dreamfusion: Text-to-3d using 2d diffusion. *arXiv preprint arXiv:2209.14988*, 2022.
- [36] 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*, pages 10684–10695, 2022.
- [37] Seyedmorteza Sadat, Otmar Hilliges, and Romann M Weber. Eliminating oversaturation and artifacts of high guidance scales in diffusion models. In *The Thirteenth International Conference on Learning Representations*, 2024.
- [38] Erwin Schrödinger. *Über die umkehrung der naturgesetze*. Verlag der Akademie der Wissenschaften in Kommission bei Walter De Gruyter u . . . , 1931.
- [39] Erwin Schrödinger. Sur la théorie relativiste de l’électron et l’interprétation de la mécanique quantique. In *Annales de l’institut Henri Poincaré*, volume 2, pages 269–310, 1932.
- [40] Jascha Sohl-Dickstein, Eric Weiss, Niru Maheswaranathan, and Surya Ganguli. Deep unsupervised learning using nonequilibrium thermodynamics. In *International conference on machine learning*, pages 2256–2265. pmlr, 2015.
- [41] Jiaming Song, Chenlin Meng, and Stefano Ermon. Denoising diffusion implicit models. *arXiv preprint arXiv:2010.02502*, 2020.
- [42] Yang Song, Jascha Sohl-Dickstein, Diederik P Kingma, Abhishek Kumar, Stefano Ermon, and Ben Poole. Score-based generative modeling through stochastic differential equations. *arXiv preprint arXiv:2011.13456*, 2020.
- [43] Pascal Vincent. A connection between score matching and denoising autoencoders. *Neural computation*, 23(7):1661–1674, 2011.
- [44] Peihao Wang, Zhiwen Fan, Dejia Xu, Dilin Wang, Sreyas Mohan, Forrest Iandola, Rakesh Ranjan, Yilei Li, Qiang Liu, Zhangyang Wang, et al. Steindreamer: Variance reduction for text-to-3d score distillation via stein identity. *arXiv preprint arXiv:2401.00604*, 2023.
- [45] Zhengyi Wang, Cheng Lu, Yikai Wang, Fan Bao, Chongxuan Li, Hang Su, and Jun Zhu. Prolificdreamer: High-fidelity and diverse text-to-3d generation with variational score distillation. *Advances in Neural Information Processing Systems*, 36:8406–8441, 2023.
- [46] Min Wei, Jingkai Zhou, Junyao Sun, and Xuesong Zhang. Adversarial score distillation: when score distillation meets gan. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, pages 8131–8141, 2024.
- [47] Taoran Yi, Jiemin Fang, Junjie Wang, Guanjun Wu, Lingxi Xie, Xiaopeng Zhang, Wenyu Liu, Qi Tian, and Xinggang Wang. Gaussiandreamer: Fast generation from text to 3d gaussians by bridging 2d and 3d diffusion models. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, pages 6796–6807, 2024.
- [48] Xin Yu, Yuan-Chen Guo, Yangguang Li, Ding Liang, Song-Hai Zhang, and Xiaojuan Qi. Text-to-3d with classifier score distillation. *arXiv preprint arXiv:2310.19415*, 2023.
- [49] Junzhe Zhu, Peiye Zhuang, and Sanmi Koyejo. Hifa: High-fidelity text-to-3d generation with advanced diffusion guidance. *arXiv preprint arXiv:2305.18766*, 2023.

NeurIPS Paper Checklist

The checklist is designed to encourage best practices for responsible machine learning research, addressing issues of reproducibility, transparency, research ethics, and societal impact. Do not remove the checklist: **The papers not including the checklist will be desk rejected.** The checklist should follow the references and follow the (optional) supplemental material. The checklist does NOT count towards the page limit.

Please read the checklist guidelines carefully for information on how to answer these questions. For each question in the checklist:

- You should answer [Yes] , [No] , or [NA] .
- [NA] means either that the question is Not Applicable for that particular paper or the relevant information is Not Available.
- Please provide a short (1–2 sentence) justification right after your answer (even for NA).

The checklist answers are an integral part of your paper submission. They are visible to the reviewers, area chairs, senior area chairs, and ethics reviewers. You will be asked to also include it (after eventual revisions) with the final version of your paper, and its final version will be published with the paper.

The reviewers of your paper will be asked to use the checklist as one of the factors in their evaluation. While "[Yes]" is generally preferable to "[No]", it is perfectly acceptable to answer "[No]" provided a proper justification is given (e.g., "error bars are not reported because it would be too computationally expensive" or "we were unable to find the license for the dataset we used"). In general, answering "[No]" or "[NA]" is not grounds for rejection. While the questions are phrased in a binary way, we acknowledge that the true answer is often more nuanced, so please just use your best judgment and write a justification to elaborate. All supporting evidence can appear either in the main paper or the supplemental material, provided in appendix. If you answer [Yes] to a question, in the justification please point to the section(s) where related material for the question can be found.

IMPORTANT, please:

- **Delete this instruction block, but keep the section heading “NeurIPS Paper Checklist”,**
- **Keep the checklist subsection headings, questions/answers and guidelines below.**
- **Do not modify the questions and only use the provided macros for your answers.**

1. Claims

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

Answer: [Yes]

Justification: We establish a novel theoretical connection, demonstrating that SDS can be precisely understood as a specific case of the Schrödinger Bridge framework. Experiments demonstrate that our TraCe achieves high-quality 3D generation, surpassing current state-of-the-art techniques. See the abstract and the end of Section 1.

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: See Supplementary.

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: See Section 4.

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: We disclose the experimental settings to reproduce the main experimental results in our paper in Supplementary and the settings of all compared methods in Section 5.

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: [Yes]

Justification: Our code will be released to the community.

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: We provide the optimization and train/test details of our proposed method in Supplementary.

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: Ours reports the results of multiple rounds of the experiment, reflecting the statistics of the experiments.

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: See Section 5.

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: See Supplementary.

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: [Yes]

Justification: See Supplementary.

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: Our 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: The assets used in the paper are properly credited, and we respect the license and terms of use of these assets throughout our research procedures.

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: Our 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.
- At submission time, remember to anonymize your assets (if applicable). You can either create an anonymized URL or include an anonymized zip file.

14. Crowdsourcing and research with human subjects

Question: For crowdsourcing experiments and research with human subjects, does the paper include the full text of instructions given to participants and screenshots, if applicable, as well as details about compensation (if any)?

Answer: [Yes]

Justification: See Supplementary.

Guidelines:

- The answer NA means that the paper does not involve crowdsourcing nor research with human subjects.
- Including this information in the supplemental material is fine, but if the main contribution of the paper involves human subjects, then as much detail as possible should be included in the main paper.
- According to the NeurIPS Code of Ethics, workers involved in data collection, curation, or other labor should be paid at least the minimum wage in the country of the data collector.

15. **Institutional review board (IRB) approvals or equivalent for research with human subjects**

Question: Does the paper describe potential risks incurred by study participants, whether such risks were disclosed to the subjects, and whether Institutional Review Board (IRB) approvals (or an equivalent approval/review based on the requirements of your country or institution) were obtained?

Answer: [Yes]

Justification: See Supplementary.

Guidelines:

- The answer NA means that the paper does not involve crowdsourcing nor research with human subjects.
- Depending on the country in which research is conducted, IRB approval (or equivalent) may be required for any human subjects research. If you obtained IRB approval, you should clearly state this in the paper.
- We recognize that the procedures for this may vary significantly between institutions and locations, and we expect authors to adhere to the NeurIPS Code of Ethics and the guidelines for their institution.
- For initial submissions, do not include any information that would break anonymity (if applicable), such as the institution conducting the review.

16. **Declaration of LLM usage**

Question: Does the paper describe the usage of LLMs if it is an important, original, or non-standard component of the core methods in this research? Note that if the LLM is used only for writing, editing, or formatting purposes and does not impact the core methodology, scientific rigor, or originality of the research, declaration is not required.

Answer: [NA]

Justification: Our core method development in this research does not involve LLMs as any important, original, or non-standard components.

Guidelines:

- The answer NA means that the core method development in this research does not involve LLMs as any important, original, or non-standard components.
- Please refer to our LLM policy (<https://neurips.cc/Conferences/2025/LLM>) for what should or should not be described.