
MGE-LDM: Joint Latent Diffusion for Simultaneous Music Generation and Source Extraction

Yunkee Chae¹

Kyogu Lee^{1,2,3}

Music and Audio Research Group (MARG)

¹ Interdisciplinary Program in Artificial Intelligence (IPAI)

² AIIS, ³ Department of Intelligence and Information

Seoul National University

{yunkimo95, kglee}@snu.ac.kr

Abstract

We present **MGE-LDM**, a unified latent diffusion framework for simultaneous music generation, source imputation, and query-driven source separation. Unlike prior approaches constrained to fixed instrument classes, MGE-LDM learns a joint distribution over full mixtures, submixtures, and individual stems within a single compact latent diffusion model. At inference, MGE-LDM enables (1) complete mixture generation, (2) partial generation (i.e., source imputation), and (3) text-conditioned extraction of arbitrary sources. By formulating both separation and imputation as conditional inpainting tasks in the latent space, our approach supports flexible, class-agnostic manipulation of arbitrary instrument sources. Notably, MGE-LDM can be trained jointly across heterogeneous multi-track datasets (e.g., Slakh2100, MUSDB18, MoisesDB) without relying on predefined instrument categories. Audio samples are available at our project page[†].

1 Introduction

Recent advances in generative modeling have significantly accelerated progress in music audio synthesis, inspired by breakthroughs in the language and vision domains. Early autoregressive models such as WaveNet [1] demonstrated the feasibility of end-to-end waveform generation. Since then, two dominant approaches have emerged: (1) discrete-token models, which compress raw audio into quantized codes [2–5] for sequence modeling [6–8]; and (2) diffusion-based models [9–11], which synthesize audio by reversing a noise corruption process in the waveform domain [12, 13]. Building on this foundation, latent diffusion models (LDMs) [14], which operate in a compressed latent space, have delivered substantial gains in synthesis quality and have been successfully applied to music generation tasks [15–18].

Despite this progress, most music generation models produce a single, mixed waveform, lacking access to the individual instrument stems required for remixing, adaptive arrangement, or downstream production tasks. To recover these components, audio source separation techniques aim to decompose a mixture into its constituent tracks. Discriminative approaches [19–22] learn to directly regress each source from the input mixture, achieving strong performance. In contrast, generative separation models sample individual sources from a learned prior [23–27]. Recently, diffusion-based methods have demonstrated strong performance not only in speech separation [28, 29] but also in speech enhancement tasks [30–33], highlighting their potential for flexible, high-quality source recovery across audio domains.

[†]https://yoongi43.github.io/MGELDM_Samples/

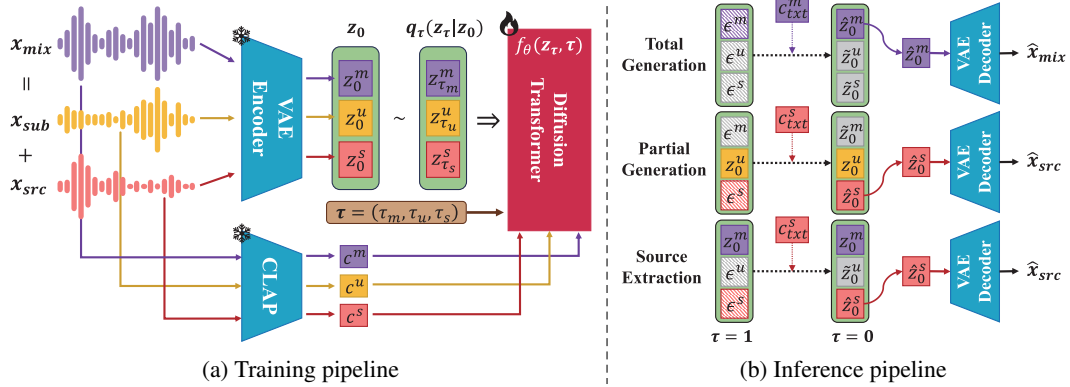


Figure 1: **Overview of MGE-LDM.** (a) *Training pipeline*: We train a three-track latent diffusion model on mixtures, submixtures, and sources. Each track is perturbed independently and conditioned on its corresponding timestep and CLAP embedding. The model is optimized using the v-objective, as detailed in Sections 3.2 and 3.4. (b) *Inference pipeline*: At test time, task-specific latents are either generated or inpainted based on available context and text prompts. The resulting latents are decoded into waveforms. See Section 3.3 for details.

Recent works have explored modeling the joint distribution of multi-track stems within a unified diffusion backbone, enabling mixture synthesis, accompaniment generation, and source separation within a single framework [34–36]. However, these approaches typically rely on predefined instrument classes for each track or assume that the mixture waveform is the linear sum of its constituent stems. While the additive assumption is valid in the waveform domain, it is incompatible with the nonlinear encoder–decoder structure of latent diffusion models, limiting the applicability of such methods in compressed latent spaces.

To address these limitations, we introduce **MGE-LDM**, a class-agnostic latent diffusion framework that jointly unifies music generation, partial generation, and arbitrary source extraction. Our approach models three interrelated latent variables: mixture, submixture, and source within a single diffusion backbone.

Given a waveform mixture $x^{(m)}$ and a chosen source $x^{(s)}$, we define the submixture as $x^{(u)} = x^{(m)} - x^{(s)}$. A pretrained encoder E maps each waveform to its latent representation:

$$z^{(m)} = E(x^{(m)}), \quad z^{(u)} = E(x^{(u)}), \quad z^{(s)} = E(x^{(s)}). \quad (1)$$

We then train a diffusion model $f_\theta(z^{(m)}, z^{(u)}, z^{(s)})$ over this joint distribution.

At inference time, diffusion-based inpainting, conditioned on a subset of $\{z^{(m)}, z^{(u)}, z^{(s)}\}$, enables class-agnostic partial generation and text-driven source extraction without relying on fixed instrument vocabularies or linear mixing assumptions. A detailed formulation appears in Section 3.

Notably, our training paradigm is **dataset-agnostic**: by removing dependence on instrument labels, MGE-LDM can leverage all publicly available track-wise datasets in a single training run. For example, Slakh2100 [37] provides clean, isolated stems; MUSDB18 [38] and MoisesDB [39] include loosely labeled tracks such as `other` that aggregate multiple instruments. Existing frameworks [34–36, 40] typically ignore such labels, as they assume a one-to-one correspondence between sources and tracks, making it difficult to train and perform inference on mixtures containing sources outside predefined instrument classes.

In contrast, our three-track formulation (mixture, submixture, individual source) treats aggregated labels as submixtures. This design integrates them into joint training together with fully separated stems. This dataset flexibility simplifies data collection and enhances model robustness across heterogeneous source structures.

In summary, our contributions are as follows:

- We propose MGE-LDM, the first latent diffusion model to jointly address music generation, source imputation, and text-conditioned source extraction without relying on predefined instrument labels.

- We formulate source extraction as a conditional inpainting problem in the latent space, jointly modeling mixture, submixture, and source embeddings. This formulation enables class-agnostic, text-driven extraction of arbitrary stems within a unified latent diffusion framework.
- We propose a training strategy that assigns distinct diffusion timesteps to each of the mixture, submixture, and source tracks, allowing the model to learn adaptive score functions for diverse inpainting contexts.
- By eliminating dependence on fixed stem annotations, our framework ingests heterogeneous multi-track datasets (e.g., Slakh2100, MUSDB18, MoisesDB) in a unified manner, simplifying data aggregation and improving generalization to unseen instrumentation.

2 Related Work

Audio Source Separation and Extraction. Audio source separation aims to decompose a polyphonic mixture into its constituent tracks, while source extraction focuses on isolating a particular target sound, often guided by metadata, text prompts, or reference examples. Two dominant paradigms have emerged: discriminative models learn a direct mapping from the input mixture to each target stem via regression losses, either in the waveform domain or in spectrogram representations [19–22, 41, 42]. In contrast, generative approaches learn probabilistic priors over source distributions and recover individual stems via sampling [23–26, 43].

Recently, diffusion-based techniques have emerged as a powerful paradigm for audio decomposition, achieving strong results in both speech separation [28, 29] and enhancement [30–33]. These methods iteratively denoise a mixture under a learned score function, offering flexible and high-fidelity source recovery.

Query-based extraction further extends separation by conditioning the model on external cues such as class labels [44–46], visual signals [47, 48], or audio exemplars [49, 50]. Several studies have also demonstrated the effectiveness of natural language prompts for flexible, user-driven source isolation [48, 51–55]. In our framework, we employ the pretrained CLAP model [56] to obtain shared audio-text embeddings, enabling seamless, language-guided extraction of arbitrary stems within a multi-track latent diffusion architecture.

Audio Generation Models. Early neural audio synthesis methods focused on autoregressive architectures that model waveform dependencies sample by sample. WaveNet [1] demonstrated the effectiveness of dilated convolutions for end-to-end generation, while SampleRNN [57] extended this with hierarchical recurrence. Subsequent work adopted adversarial objectives to improve fidelity, using GANs to generate perceptually sharp outputs [58, 59].

Parallel efforts introduced discrete-token models, where audio is encoded into compact code sequences using vector quantization (e.g., VQ-VAE [2]). Jukebox [6] models long-range dependencies over codes using Transformers [60], while recent systems enhance fidelity through residual quantization [3–5] and hierarchical token modeling, where coarse-to-fine code representations are generated across multiple levels [7, 61–63]. MusicGen [8] improves decoding efficiency with delayed-token generation, and Instruct-MusicGen [64] extends it for targeted editing via instruction-tuned prompts. In parallel, token-based masked generative modeling techniques, originally developed for vision [65], have been adapted for audio, enabling efficient non-autoregressive synthesis and precise spectrogram inpainting [66, 67].

Diffusion-based generation emerged with DiffWave [12] and WaveGrad [13], which learn to iteratively denoise Gaussian-corrupted waveforms. These techniques have since been adapted for music-specific generation with structure and style conditioning [68, 69]. Latent diffusion models (LDMs) [14], which perform denoising in a compressed embedding space, have further advanced generation fidelity and scalability. LDM-based audio models such as AudioLDM [70, 71], MusicLDM [16], and Stable Audio [17, 18, 72] achieve state-of-the-art performance. Recent frameworks like AUDIT [73], InstructME [74] explore the use of diffusion for controllable and interactive audio editing.

Multi-Track Music Audio Modeling. Recent studies model multi-track music as a structured composition of interdependent stems. StemGen [75] employs an iterative, non-autoregressive transformer over discrete tokens to generate stems conditioned on text prompts. Jen-1 Composer [76] applies latent diffusion to jointly model four canonical stems (bass, drums, instrument, melody), producing

coherent multi-track compositions. MusicGen-Stem [77] combines per-stem vector quantization with an autoregressive decoder to synthesize bass, drums, and aggregated `other` components, and supports mixture-conditioned accompaniment generation. Diff-A-Riff [78] leverages latent diffusion to co-create stems complementary to given mixtures, later extended with diffusion transformers [79]. Other studies have similarly explored mixture-conditioned stem generation [80, 81].

A separate line of work focuses on joint modeling of synthesis and decomposition within a single diffusion backbone. Multi-Source Diffusion Models (MSDM) [34] model a fixed set of stems (`bass`, `drums`, `guitar`, and `piano`) within a shared diffusion framework, relying on an additive mixture assumption and a Dirac delta-based posterior sampler, following the EDM formulation for ODE-based sampling [82]. This line of work has since been extended in GMSDI [35], MSG-LD [36], and others [40, 83]. GMSDI enables variable-stem modeling and text-based conditioning but remains grounded in waveform-space additive mixing. MSG-LD adapts latent diffusion for four-stem modeling and classifier-free guidance [84], though it still assumes fixed instrument classes.

In contrast, our approach jointly models mixture, submixture, and source embeddings in latent space and casts both synthesis and arbitrary-source extraction as text-conditioned inpainting tasks, offering fully class-agnostic multi-track music processing without reliance on fixed instrument vocabularies or linear mixing assumptions.

3 Method

Figure 1 presents an overview of the proposed MGE-LDM framework. We train a three-track joint latent diffusion model that learns the distribution over the joint space of mixture, submixture, and source representations, as defined in Eq. (1). During inference, the model performs various tasks by exploiting the inpainting capability of diffusion models. We first describe how training triplets are constructed, then detail how they are used for joint diffusion-based training and inpainting-based inference.

3.1 Formulating Joint Latent Representation

Let $\{x_i\}_{i \in I}$ denote the set of time-domain audio stems, where the number of sources $|I|$ may vary across mixtures depending on their instrumentation. The mixture waveform is defined as:

$$x^{(m)} = \sum_{i \in I} x_i.$$

To construct a training example for our three-track model, we uniformly sample an index $j \in I$ and define:

$$x^{(s)} = x_j, \quad x^{(u)} = \sum_{i \in I \setminus \{j\}} x_i,$$

yielding the triplet $(x^{(m)}, x^{(u)}, x^{(s)})$. We encode each element of this triplet using a pretrained variational autoencoder (VAE) [85] encoder E , resulting in latent representations $z^{(m)}, z^{(u)}, z^{(s)} \in \mathbb{R}^{C \times L}$, where C and L denote the latent channel and temporal dimensions, respectively. This formulation naturally accommodates mixtures with a variable number of stems. Regardless of the number of instruments present, any publicly available multi-track dataset can be decomposed into mixture, submixture, and source components for joint latent modeling.

3.2 Latent Diffusion Training with Three-Track Embeddings

We build upon the Stable Audio framework [72], employing a Diffusion Transformer (DiT) backbone [86] and training the model under the v-objective [87]. Below, we summarize the v-objective training procedure; full derivations and additional details are provided in Appendix A.1. Let the composite latent input be defined as:

$$\mathbf{z}_0 = (z_0^{(m)}, z_0^{(u)}, z_0^{(s)}) \in \mathbb{R}^{3 \times C \times L}$$

where $z_0^{(k)} \in \mathbb{R}^{C \times L}$ are (clean) track embeddings, with $k \in K = \{m, u, s\}$ denoting the track types – mixture, submixture, and source, respectively.

We aim to estimate the score $\nabla_{\mathbf{z}_\tau} \log q_\tau(\mathbf{z}_\tau)$ across continuous noise levels $\tau \in [\tau_{\min}, 1]$. To this end, we perturb the clean latent variable $\mathbf{z}_0$ with Gaussian noise, following:

$$\mathbf{z}_\tau = \alpha_\tau \mathbf{z}_0 + \beta_\tau \boldsymbol{\epsilon}, \quad \boldsymbol{\epsilon} \sim \mathcal{N}(\mathbf{0}, \mathbf{I}), \quad (2)$$

where the noise scaling coefficients are parameterized as:

$$\alpha_\tau = \cos(\phi_\tau), \quad \beta_\tau = \sin(\phi_\tau), \quad \phi_\tau = \frac{\pi}{2} \tau. \quad (3)$$

Here, $\tau \sim \mathcal{U}([\tau_{\min}, 1])$ is sampled from a truncated uniform distribution with $\tau_{\min} = 0.02$ for stability.

A denoising network $f_\theta(\mathbf{z}_\tau, \tau, \mathbf{c})$ is trained to estimate the score $\nabla_{\mathbf{z}_\tau} \log q_\tau(\mathbf{z}_\tau | \mathbf{c})$ using the v-objective:

$$\mathcal{L}(\theta) = \mathbb{E}_{\mathbf{z}_0, \boldsymbol{\epsilon}, \tau} \|f_\theta(\mathbf{z}_\tau, \tau, \mathbf{c}) - \mathbf{v}_\tau\|_2^2, \quad \mathbf{v}_\tau = \frac{\partial \mathbf{z}_\tau}{\partial \phi_\tau} = \alpha_\tau \boldsymbol{\epsilon} - \beta_\tau \mathbf{z}_0. \quad (4)$$

The conditioning vector $\mathbf{c} = (c^{(m)}, c^{(u)}, c^{(s)})$ is derived using the audio branch of a pretrained CLAP encoder [56], applied to each component:

$$c^{(k)} = \text{CLAP}_{\text{audio}}(x^{(k)}), \quad \text{for } k \in K.$$

To enable classifier-free guidance (CFG) [84], each $c^{(k)}$ is independently dropped out with probability p during training.

3.3 Inference via Conditional Sampling in Latent Space

In the image domain, *inpainting* refers to reconstructing missing or corrupted regions of an image by conditioning on surrounding pixels. Diffusion models have demonstrated strong zero-shot inpainting capabilities, enabling arbitrary mask completion without retraining [11, 88]. We extend this paradigm to the latent domain of music, operating over a joint distribution of mixture, submixture, and source embeddings. Downstream tasks are formulated as conditional generation problems, where known latents are treated as observed and unknown ones are sampled as missing components.

In all inference modes, we condition on natural-language queries using CLAP embeddings. When text conditioning is required, we use the text branch of CLAP to produce the prompt embedding $c^{(k)} = \text{CLAP}_{\text{text}}(c_{\text{text}}^{(k)})$, where $c_{\text{text}}^{(k)}$ is a free-form natural language description (e.g., "*the sound of an electric guitar*").

Total Generation. Let $p_\theta(z^{(m)}, z^{(u)}, z^{(s)})$ denote the implicit model distribution whose score we approximate with f_θ . To synthesize a complete mixture, we condition only on the mixture prompt embedding $c^{(m)}$ or omit all conditions for unconditional generation. We sample the mixture latent $\hat{z}^{(m)}$ as:

$$\hat{z}^{(m)}, \tilde{z}^{(u)}, \tilde{z}^{(s)} \sim p_\theta(z^{(m)}, z^{(u)}, z^{(s)} | c^{(m)}, \emptyset^{(u)}, \emptyset^{(s)}), \quad (5)$$

where $\tilde{z}^{(u)}$ and $\tilde{z}^{(s)}$ are auxiliary latents that are discarded. Finally, the synthesized mixture waveform $\hat{x}^{(m)}$ is obtained by decoding $\hat{z}^{(m)}$ through the pretrained VAE decoder D :

$$\hat{x}^{(m)} = D(\hat{z}^{(m)}).$$

Hereafter, we use $\tilde{z}^{(k)}$ to denote any dummy latent that is not retained during inference.

Partial Generation. *Partial generation*, also known as *source imputation*, refers to the task of generating missing stems given partially observed sources. We approach this iteratively to progressively reconstruct the full mixture from the partial input.

Let $\bar{\mathcal{I}} = \{c_1, \dots, c_J\}$ be an (ordered) set of CLAP-derived text embeddings, each corresponding to a target source description to be imputed. Let $x_0^{(u)}$ denote the waveform mixture of the observed sources, and let $z_0^{(u)} = E(x_0^{(u)})$ be its latent representation. We initialize the submixture latent with $z_0^{(u)}$ and generate each missing source sequentially.

At each step $j \in \{1, \dots, J\}$, we sample a new source latent $\hat{z}_j^{(s)}$ conditioned on the current submixture and the text embedding $c_j^{(s)}$:

$$\tilde{z}_j^{(m)}, \hat{z}_j^{(s)} \sim p_\theta(z^{(m)}, z^{(s)} | z_{j-1}^{(u)}, \emptyset^{(m)}, \emptyset^{(u)}, c_j^{(s)}). \quad (6)$$

We then update the submixture by accumulating the decoded sources:

$$z_j^{(u)} = E \left(\sum_{l=0}^{j-1} D(\hat{z}_l^{(s)}) \right).$$

After J iterations, we obtain the full set of imputed sources $\{\hat{z}_j^{(s)}\}_{j=1}^J$ and reconstruct the final mixture waveform as:

$$x^{(m)} = x_0^{(u)} + \sum_{j=1}^J D(\hat{z}_j^{(s)}). \quad (7)$$

Source Extraction. Text-driven extraction of an arbitrary stem is performed by conditioning on a natural-language prompt. Given a prompt embedding $c^{(s)}$, we treat the mixture latent $z^{(m)}$ as observed and inpaint the submixture and source tracks:

$$\tilde{z}^{(u)}, \hat{z}^{(s)} \sim p_\theta(z^{(u)}, z^{(s)} \mid z^{(m)}, \emptyset^{(m)}, \emptyset^{(u)}, c^{(s)}), \quad (8)$$

where $\tilde{z}^{(u)}$ is an auxiliary prediction that is discarded. Finally, the isolated source waveform is reconstructed via $\hat{x}^{(s)} = D(\hat{z}^{(s)})$.

3.4 Track-aware Inpainting Model with Adaptive Timesteps

Conventional diffusion-based inpainting methods apply a uniform noise schedule across both observed and missing regions, failing to account for their differing uncertainty characteristics [9, 11, 88–90]. In the standard setup, a denoising model $f_\theta(\mathbf{z}_\tau, \tau)$ is trained to approximate the joint score of a perturbed latent variable.

Following the perturbation rule and v-objective from Section 3.2, the score estimate is expressed as:

$$\nabla_{\mathbf{z}_\tau} \log q_\tau(\mathbf{z}_\tau) \approx -\mathbf{z}_\tau - \frac{\alpha_\tau}{\beta_\tau} f_\theta(\mathbf{z}_\tau, \tau), \quad (9)$$

where α_τ and β_τ is cosine schedule as defined in Eq. (3), following [87]. For notational simplicity, we omit the conditioning vectors $\mathbf{c}$ in this expression. A detailed derivation is provided in Appendix A.1.

Recently, region-aware adaptations of diffusion inpainting, including spatially varying noise schedules [91] and per-pixel timestep conditioning in TD-Paint [92], have demonstrated substantial improvements in semantic consistency by preserving fidelity in observed regions. Inspired by TD-Paint, we extend this idea to three-track music audio by assigning distinct timestep conditions to each track, thereby improving inpainting quality in the latent space.

We describe our track-wise adaptive timestep conditioned model using general notation. Let K be a set of tracks, and let $N = |K|$ be the number of tracks. Define the clean latent tensor and the corresponding noise levels as:

$$\mathbf{z}_0 = (z_0^{(k)})_{k \in K} \in \mathbb{R}^{N \times C \times L}, \quad \boldsymbol{\tau} = (\tau_k)_{k \in K} \in [\tau_{\min}, 1]^N,$$

where each τ_k is either zero (for observed tracks) or equal to a shared sample $\tau \sim \mathcal{U}([\tau_{\min}, 1])$, depending on the inpainting configuration.

We define the track-wise product between a vector $x_\tau \in \mathbb{R}^N$ and a latent tensor $\mathbf{z} \in \mathbb{R}^{N \times C \times L}$ as:

$$x_\tau \odot \mathbf{z} := (x_{\tau_k} z^{(k)})_{k \in K},$$

and extend this notation to the cosine noise schedule terms as follows:

$$\alpha_\tau = (\alpha_{\tau_k})_{k \in K}, \quad \beta_\tau = (\beta_{\tau_k})_{k \in K}.$$

We perturb the joint latent using track-wise noise:

$$\mathbf{z}_\tau = \alpha_\tau \odot \mathbf{z}_0 + \beta_\tau \odot \boldsymbol{\epsilon}, \quad \boldsymbol{\epsilon} \sim \mathcal{N}(\mathbf{0}, \mathbf{I}), \quad (10)$$

where each track is independently scaled by its corresponding noise factor.

The denoiser $f_\theta(\mathbf{z}_\tau, \boldsymbol{\tau})$ is trained to regress the velocity target under the v-objective:

$$\mathbf{v}_\tau = \alpha_\tau \odot \boldsymbol{\epsilon} - \beta_\tau \odot \mathbf{z}_0, \quad (11)$$

resulting in the following training loss:

$$\mathcal{L}(\theta) = \mathbb{E}_{\mathbf{z}_0, \boldsymbol{\epsilon}, \boldsymbol{\tau}} \|f_\theta(\mathbf{z}_\tau, \boldsymbol{\tau}) - \mathbf{v}_\tau\|_2^2. \quad (12)$$

In our setup, we use $N = 3$ with $K = \{m, u, s\}$, corresponding to the mixture, submixture, and source tracks, respectively. In practice, the loss is computed only over the unknown tracks.

During training, we first sample a noise level $\boldsymbol{\tau} \sim \mathcal{U}([\tau_{\min}, 1])$ and set the per-track timestep vector $\boldsymbol{\tau} \in \mathbb{R}^3$ according to one of the following four patterns:

$$\boldsymbol{\tau} \in \{(\tau, \tau, \tau), (0, \tau, \tau), (\tau, 0, \tau), (\tau, \tau, 0)\},$$

where each configuration is selected randomly for each training step.

Under this conditioning strategy, the full-noise setting $\boldsymbol{\tau} = (\tau, \tau, \tau)$ corresponds to learning the standard joint score, as described in Eq. (9). In contrast, a "single-zero" pattern allows the model to learn conditional score functions for the unobserved tracks while treating the others as fixed observations.

For example, when $(\tau_m, \tau_u, \tau_s) = (0, \tau, \tau)$, the model is trained to approximate the gradient:

$$\nabla_{(z_\tau^{(u)}, z_\tau^{(s)})} \log q_\tau(z_\tau^{(u)}, z_\tau^{(s)} | z_0^{(m)}), \quad (13)$$

which the joint-score formulation in Eq. (9) cannot compute in closed form. At inference time, we clamp the observed tracks by setting their noise levels to zero and apply standard reverse diffusion updates to the remaining (missing) tracks. Pseudocode and a detailed theoretical comparison with conventional inpainting methods are provided in Appendix B.

4 Experimental Setup

In this section, we outline our experimental protocol, including baseline models, datasets, and key implementation details. A comprehensive description of hyperparameters and training procedures is provided in Appendix D. All baseline results are re-evaluated using our test sets to ensure consistency with our experimental setup.

Baselines. We use two recent multi-track diffusion models — MSDM [34] and MSG-LD [36] — as baselines, both of which operate on a fixed set of stems: `bass`, `drums`, `guitar`, and `piano`. In addition to generative and inpainting performance, we assess source extraction capabilities against Hybrid Demucs (HDemucs) [20], which separates the mixture into `bass`, `drums`, `other`, and `vocals` stems, and AudioSep [53], which performs text-conditioned separation based on natural language queries.

Datasets. We train and evaluate on three multi-track music datasets: Slakh2100 [37], MUSDB18 [38] (denoted $\mathbf{M}_u$), and MoisesDB [39] (denoted $\mathbf{M}_o$). For Slakh2100, we define two subsets: $\mathbf{S}_A$, containing only `bass`, `drums`, `guitar`, and `piano` stems to match the MSDM and MSG-LD setup; and $\mathbf{S}_B$, which includes all remaining stems. Each dataset follows its predefined train/test split. We train our models on various dataset combinations to evaluate robustness under different source distributions and stem configurations.

Implementation Details. Our models use the Stable Audio backbone [72], which comprises an autoencoder and a DiT-based diffusion model. To better accommodate per-track variability in the joint latent space, we replace LayerNorm [93] with GroupNorm [94], using three groups to reflect the number of tracks.

To bridge the audio-text modality gap, we adopt stochastic linear interpolation between audio and text embeddings on the source track, following prior work on multimodal fusion [54, 95]. Concretely, we generate the prompt "*The sound of the {label}*" and compute the source conditioning vector $\mathbf{c}^{(s)}$ as a convex combination of the CLAP text embedding and its corresponding audio embedding, where the interpolation weight $\alpha \sim \mathcal{U}([0, 1])$ is sampled randomly for each training example.

All of our models, except the one trained on the full dataset combination ($\mathbf{S}_A, \mathbf{S}_B, \mathbf{M}_u, \mathbf{M}_o$), are trained for 200K iterations with a batch size of 64, using 16 kHz audio segments of 10.24 seconds.

Table 1: **Total generation results.** Reported scores are FAD $\downarrow$, computed against mixture references from each test set. Values in parentheses indicate generation results conditioned on the prompt “*The sound of the bass, drums, guitar, and piano*”, as detailed in Section 5.1. Among our models, $\mathcal{T}_1$ provides a fair comparison with baseline methods, as it is trained on the same dataset, while $\mathcal{T}_2$ – $\mathcal{T}_4$ illustrate the effects of progressive dataset scaling. Bold values indicate the best results in each column, and underlined values denote the best results among models trained on the same $\mathbf{S}_A$ set.

Model	Train Set				Test Set			
	$\mathbf{S}_A$	$\mathbf{S}_B$	$\mathbf{M}_u$	$\mathbf{M}_o$	$\mathbf{S}_A$	$\mathbf{S}_{\text{Full}}$	$\mathbf{M}_u$	$\mathbf{M}_o$
MSDM	✓	×	×	×	4.21	6.04	7.92	7.41
MSG-LD	✓	×	×	×	1.38	<u>1.55</u>	<u>4.61</u>	<u>4.26</u>
	$\mathcal{T}_1$	✓	×	×	0.47 (3.57)	1.79	6.34	5.90
MGE (ours)	$\mathcal{T}_2$	✓	✓	×	3.14 (2.24)	0.63	5.46	4.73
	$\mathcal{T}_3$	×	×	✓	8.80 (3.96)	6.56	2.87	1.59
	$\mathcal{T}_4$	✓	✓	✓	6.83 (5.05)	4.22	2.78	1.47

The full combination model is trained for 320K iterations with a batch size of 128. During sampling and inpainting, we apply classifier-free guidance (CFG) with a guidance scale of 2.0 and a per-track dropout probability of $p = 0.1$. All diffusion-based samples – including those from baseline models – are generated using 250 inference steps. We adopt DDIM sampling [96] for all our models, while each baseline uses its originally proposed sampling method.

5 Results

We evaluate MGE-LDM on three tasks: total generation, partial generation, and source extraction. Each result table indicates the training dataset(s) used and reports performance across multiple test sets. Unless otherwise specified, partial generation and source extraction are performed using text prompts of the form “*The sound of the {label}.*” Abbreviations for all stem labels are listed in Appendix Table 4, and additional ablation results are provided in Appendix E.

5.1 Music Generation

Table 1 presents FAD (Fréchet Audio Distance) [97] scores computed using VGGish embeddings [98], a widely adopted metric for evaluating music generation quality. Note that the $\mathbf{S}_A$ test set contains only mixtures of bass, drums, guitar, and piano, while $\mathbf{S}_{\text{Full}}$ corresponds to the full Slakh2100 test set, which includes a broader and more diverse set of instruments.

We first compare our model $\mathcal{T}_1$ against MSDM and MSG-LD, where all models are trained on $\mathbf{S}_A$. Our model achieves the lowest FAD on the $\mathbf{S}_A$ test set, demonstrating superior fidelity in generating standard four-stem mixtures. However, its generalization to other test sets is more limited. On $\mathbf{S}_{\text{Full}}$, which includes a broader range of instruments, MSG-LD performs slightly better than $\mathcal{T}_1$, suggesting mild overfitting to the constrained training distribution. That said, $\mathcal{T}_1$ still outperforms MSDM across most test conditions.

To assess the effect of dataset extension, we next train on the full Slakh2100 dataset ($\mathbf{S}_A + \mathbf{S}_B$), resulting in model $\mathcal{T}_2$. This broader training improves performance on $\mathbf{S}_{\text{Full}}$ and enhances generalization compared to $\mathcal{T}_1$. However, $\mathcal{T}_2$ still lags behind MSG-LD on other subsets, and its performance on $\mathbf{S}_A$ degrades, likely due to increased variability in the training distribution.

The $\mathcal{T}_3$ model, trained on MUSDB18 ($\mathbf{M}_u$) and MoisesDB ($\mathbf{M}_o$), both containing real recordings rather than synthesized audio, achieves strong performance on their respective test sets. Despite the domain shift, it also performs competitively on Slakh2100, with results comparable to MSDM, highlighting the model’s robustness to cross-domain generalization. Finally, our model trained on the combined dataset comprising Slakh2100, MUSDB18, and MoisesDB (denoted as $\mathcal{T}_4$), achieves the best overall results on both $\mathbf{M}_u$ and $\mathbf{M}_o$. It also outperforms $\mathcal{T}_3$ on $\mathbf{S}_{\text{Full}}$, benefiting from the broader training distribution.

We also report results in parentheses, which correspond to mixture generation conditioned on the text prompt “*The sound of the bass, drums, guitar, and piano.*” For models trained on datasets beyond $\mathbf{S}_A$,

Table 2: **Partial generation results.** Scores are reported using *sub*-FAD ↓, which measures the distance between the reference mixture and the sum of given and generated sources. Each column header (e.g., B, D, G) indicates the target source being generated, conditioned on the remaining stems. Bold and underlined values follow the same convention as in Table 1.

Model	Train Set				S _A														S _B									
	S _A	S _B	M _u	M _o	B	D	G	P	BD	BG	BP	DG	DP	GP	BDG	BDP	BGP	DGP	Brs.	C.P.	Org.	Pipe	Reed	Str.	S.Lead	S.Pad		
MSDM	✓	×	×	×	0.56	1.06	0.49	0.70	2.23	1.56	1.95	1.64	1.83	2.31	3.09	3.53	5.72	3.86	-	-	-	-	-	-	-	-		
MSG-LD	✓	×	×	×	0.33	0.34	0.49	0.48	0.70	1.08	1.05	0.86	0.83	1.47	1.43	1.42	2.31	1.76	-	-	-	-	-	-	-	-		
MGE (ours)	T ₁	✓	×	×	×	1.02	1.41	1.17	1.19	1.15	1.29	1.25	1.69	1.65	1.14	1.80	1.84	1.45	1.84	1.45	0.68	0.23	3.48	5.58	1.38	4.47	1.08	
	T ₂	✓	✓	×	×	2.11	2.99	1.99	2.74	4.07	2.32	4.18	3.54	3.90	3.18	4.93	5.69	4.25	4.66	5.96	0.41	1.03	3.66	3.52	2.79	0.88	2.32	
	T ₃	×	×	✓	✓	1.43	1.29	3.34	2.30	1.85	3.64	2.83	2.95	2.36	4.39	3.30	3.57	6.03	3.86	3.58	0.58	0.15	0.22	0.56	0.61	0.54	0.48	
	T ₄	✓	✓	✓	✓	1.14	1.50	3.75	2.47	2.06	4.06	2.82	3.37	2.74	4.55	3.94	4.05	5.66	4.06	5.09	0.42	0.56	0.20	3.14	3.95	0.31	0.40	

Table 3: **Source extraction results.** Metrics are reported as Log-Mel L1 distance ↓. For baseline models, scores are shown only for stems included in their fixed output set. Bold and underlined values follow the same convention as in Table 1.

Model	Train Set				S_A				S_B								M_u				M_o							
	S_A	S_B	M_u	M_o	B	D	G	P	Brs.	C.P.	Org.	Pipe	Reed	Str.	S.Lead	S.Pad	V	B	D	V	B	D	G	P	B.Str	Perc.		
HDemucs	×	×	✓	×	1.49	0.90	-	-	-	-	-	-	-	-	-	-	1.50	1.99	1.53	0.83	1.71	1.10	-	-	-	-		
AudioSep	×	×	×	×	2.36	1.67	3.41	2.42	3.13	2.84	3.26	3.04	3.15	2.57	2.8	2.06	2.66	4.07	1.89	1.54	3.37	1.87	1.31	1.42	1.70	2.36		
MSDM	✓	×	×	×	1.90	1.51	3.32	2.70	-	-	-	-	-	-	-	-	-	2.56	1.69	-	2.15	1.31	1.28	1.51	-	-		
MSG-LD	✓	×	×	×	1.20	1.24	2.24	1.85	-	-	-	-	-	-	-	-	-	1.96	1.60	-	1.72	1.49	2.36	2.06	-	-		
$\mathcal{T}_1$	✓	×	×	×	1.28	0.66	1.27	1.07	3.22	3.07	3.13	3.11	3.30	2.77	2.68	2.30	3.80	<u>1.91</u>	<u>1.33</u>	5.15	<u>1.61</u>	<u>1.10</u>	2.86	2.68	2.03	2.94		
MGE (ours)	$\mathcal{T}_2$	✓	✓	×	1.68	2.71	2.69	2.16	3.43	2.16	1.84	2.33	3.07	2.44	2.31	1.93	3.55	2.14	2.15	4.66	1.86	2.11	2.28	2.18	1.93	2.36		
	$\mathcal{T}_3$	×	×	✓	1.80	0.99	2.89	2.01	3.17	2.51	3.61	2.13	2.86	2.22	2.78	2.22	1.85	1.56	1.17	0.98	1.10	0.90	1.04	1.58	1.62	2.49		
	$\mathcal{T}_3$	×	×	✓	1.80	0.99	2.89	2.01	3.17	2.51	3.61	2.13	2.86	2.22	2.78	2.22	1.85	1.56	1.17	0.98	1.10	0.90	1.04	1.58	1.62	2.49		
	$\mathcal{T}_4$	✓	✓	✓	1.67	0.83	2.61	1.77	3.15	2.29	2.22	1.95	2.61	1.85	2.71	3.68	1.76	1.56	1.13	1.01	1.07	0.86	1.02	1.40	2.25	2.69		

including $\mathcal{T}_2$, $\mathcal{T}_3$, and $\mathcal{T}_4$, the generated mixtures generally contain a wider variety of instruments. This can result in a distributional mismatch with the S_A test set, which contains only bass, drums, guitar, and piano, thereby increasing the FAD due to reference-target discrepancy. To mitigate this, we apply text conditioning at inference time using the above prompt to constrain the generated mixture to match the reference instrument set. As shown in the results, this conditioning significantly improves FAD for $\mathcal{T}_2$, $\mathcal{T}_3$, and $\mathcal{T}_4$, yielding performance comparable to the MSDM baseline.

Table 2 presents results for partial generation, evaluated using the *sub*-FAD metric. This metric computes the FAD between the original mixture and a reconstruction formed by combining the given submixture with the generated stems. This evaluation protocol was introduced in prior work on music stem completion [34, 36, 63].

On S_A , our model $\mathcal{T}_1$ performs worse than MSG-LD for single-source imputation but shows competitive or better performance as the number of generated stems increases. Among our models, $\mathcal{T}_1$ performs best on S_A due to domain alignment. However, for imputation tasks involving broader instrument classes in S_B , models trained on more diverse datasets perform better. Notably, the $\mathcal{T}_3$ model, trained solely on M_u and M_o , achieves strong results on S_B despite not being exposed to Slakh2100 during training. In particular, it performs well on instruments such as organ, pipe, and strings. The fully trained model $\mathcal{T}_4$ further improves performance across several S_B stems, including synth lead, synth pad, and pipe.

5.2 Source Extraction

We evaluate text-queried source extraction using the Log Mel L1 distance, following MSG-LD [36], due to the inherent phase mismatch in latent-domain models that complicates waveform-domain evaluation. Table 3 presents results for a variety of stems across the S_A , S_B , M_u , and M_o test sets. Alongside MSDM and MSG-LD, we compare with two additional baselines: HDemucs [20], a strong waveform- and spectrogram-domain separation model trained on fixed stems; and AudioSep [53], a recent system that enables natural language-driven extraction. For fixed-stem baselines (e.g., MSDM, MSG-LD, HDemucs), we report both in-distribution results and out-of-distribution generalization where possible (e.g., bass in M_u).

Our model $\mathcal{T}_1$, trained solely on S_A , performs strongly on the canonical Slakh stems (bass, drums, guitar, piano), outperforming MSG-LD on all but bass. However, it generalizes poorly to less common stems and real-world recordings. By expanding the training set to encompass the full

Slakh2100 dataset, model $\mathcal{T}_2$ achieves improved performance in categories such as `chromatic percussion`, `organ`, `synth lead`, and `synth pad`, demonstrating the importance of broader intra-domain coverage.

Interestingly, model $\mathcal{T}_3$ generalizes competitively with $\mathcal{T}_2$ to synthetic stems, even outperforming some stems such as `drums`. This indicates cross-domain robustness in our latent inpainting formulation. Finally, model $\mathcal{T}_4$, trained on the combined dataset, exhibits robust performance across both synthetic and real-world domains. It maintains strong results on Slakh2100, while achieving the lowest Log Mel L1 scores across most stems in $\mathbf{M}_u$ and $\mathbf{M}_o$. This highlights that incorporating synthetic data such as Slakh2100 alongside real recordings can enhance generalization and improve separation quality on real-world audio.

Overall, MGE-LDM delivers robust, class-agnostic extraction with strong performance across a wide range of instrument types and recording domains, highlighting its effectiveness for text-driven music source extraction in both synthetic and real-world settings.

6 Limitations

While MGE-LDM provides a flexible, class-agnostic framework for multi-track audio modeling, several limitations remain. First, all experiments are conducted using 16 kHz monaural audio, which constrains upper-frequency resolution and omits spatial cues, thereby limiting realism for high-fidelity or stereo music applications. Second, the model relies on CLAP-based semantic conditioning, which introduces a modality gap between text and audio [99]. This can occasionally lead to semantic drift during extraction, resulting in irrelevant or hallucinated sources, particularly for stems with limited training data.

Third, although MGE-LDM reduces dependence on fixed instrument classes, it still requires multi-stem supervision during training. This dependency restricts applicability to fully unlabeled or large-scale web audio collections. Fourth, training on MUSDB18 alone with the same number of iterations as other configurations leads to overfitting, likely due to the limited duration (approximately 10 hours) of its training split. This highlights the challenge of achieving robust performance in low-resource multi-track settings.

Finally, our model is trained using triplets (*mix*, *submix*, *source*) that satisfy $mix = submix + source$ in waveform space; however, the latent diffusion process does not enforce an explicit additivity constraint for generated triplets. We believe this omission contributes directly to hallucination phenomena, where the model extracts sources absent from the mixture. Postolache et al. [27] addresses a related issue by enforcing additivity in a discrete VQ-VAE latent space, estimating the joint likelihood of two sources by counting codebook co-occurrences, effectively modeling $p(z_{mix} | z_{src_1}, z_{src_2})$, where z_* are quantized latent codes. Our current pipeline, however, operates in a continuous latent space, which precludes the direct use of such discrete bin-counting methods. Adapting this latent-domain likelihood formulation to continuous spaces, for example, by designing suitable regularizers or adopting a VQ-VAE-based encoder with discrete diffusion [100–102] or MaskGiT [65]-style generation, represents a promising direction for future work.

7 Conclusion

We have presented MGE-LDM, a unified class-agnostic latent diffusion framework that jointly models mixtures, submixtures, and individual sources for music generation, stem completion, and text-driven extraction. By formulating stem completion and source extraction as conditional inpainting in a shared latent space and by introducing track-dependent timestep conditioning, we overcome the limitations of fixed-class, additive mixing assumptions and achieve flexible manipulation of arbitrary instrument tracks. Empirically, MGE-LDM matches or exceeds specialized baselines on Slakh2100 generation and separation benchmarks, while uniquely supporting zero-shot, language-guided extraction across heterogeneous multi-track datasets.

Acknowledgements

This work was supported by the National Research Foundation of Korea (NRF) grant funded by the Korea government (MSIT) [No. RS-2024-00461617, 90%], Information & communications Technology Planning & Evaluation (IITP) grant funded by the Korea government (MSIT) [No. RS-2021-II211343, Artificial Intelligence Graduate School Program (Seoul National University), 10%] Additionally, the GPUs were partly supported by the National IT Industry Promotion Agency (NIPA)'s high-performance computing support program in 2025.

References

- [1] Aaron Van Den Oord, Sander Dieleman, Heiga Zen, Karen Simonyan, Oriol Vinyals, Alex Graves, Nal Kalchbrenner, Andrew Senior, Koray Kavukcuoglu, et al. Wavenet: A generative model for raw audio. *arXiv preprint arXiv:1609.03499*, 2016.
- [2] Aaron Van Den Oord, Oriol Vinyals, et al. Neural discrete representation learning. In *Advances in neural information processing systems*, volume 30, 2017.
- [3] Neil Zeghidour, Alejandro Luebs, Ahmed Omran, Jan Skoglund, and Marco Tagliasacchi. SoundStream: an end-to-end neural audio codec. *IEEE/ACM Transactions on Audio, Speech, and Language Processing*, 30:495–507, 2021.
- [4] Alexandre Défossez, Jade Copet, Gabriel Synnaeve, and Yossi Adi. High fidelity neural audio compression. *arXiv preprint arXiv:2210.13438*, 2022.
- [5] Rithesh Kumar, Prem Seetharaman, Alejandro Luebs, Ishaan Kumar, and Kundan Kumar. High-fidelity audio compression with improved rvqgan. *Advances in Neural Information Processing Systems*, 36: 27980–27993, 2023.
- [6] Prafulla Dhariwal, Heewoo Jun, Christine Payne, Jong Wook Kim, Alec Radford, and Ilya Sutskever. Jukebox: A generative model for music. *arXiv preprint arXiv:2005.00341*, 2020.
- [7] Andrea Agostinelli, Timo I Denk, Zalán Borsos, Jesse Engel, Mauro Verzetti, Antoine Caillon, Qingqing Huang, Aren Jansen, Adam Roberts, Marco Tagliasacchi, et al. MusicLM: generating music from text. *arXiv preprint arXiv:2301.11325*, 2023.
- [8] Jade Copet, Felix Kreuk, Itai Gat, Tal Remez, David Kant, Gabriel Synnaeve, Yossi Adi, and Alexandre Défossez. Simple and Controllable Music Generation. *Advances in Neural Information Processing Systems*, 36:47704–47720, 2023.
- [9] Yang Song and Stefano Ermon. Generative modeling by estimating gradients of the data distribution. *Advances in neural information processing systems*, 32, 2019.
- [10] Jonathan Ho, Ajay Jain, and Pieter Abbeel. Denoising diffusion probabilistic models. *Advances in neural information processing systems*, 33:6840–6851, 2020.
- [11] Yang Song, Jascha Sohl-Dickstein, Diederik P Kingma, Abhishek Kumar, Stefano Ermon, and Ben Poole. Score-based generative modeling through stochastic differential equations. In *International Conference on Learning Representations*, 2021.
- [12] Zhifeng Kong, Wei Ping, Jiayi Huang, Kexin Zhao, and Bryan Catanzaro. Diffwave: A versatile diffusion model for audio synthesis. In *International Conference on Learning Representations*, 2021.
- [13] Nanxin Chen, Yu Zhang, Heiga Zen, Ron J Weiss, Mohammad Norouzi, and William Chan. Wavegrad: Estimating gradients for waveform generation. In *International Conference on Learning Representations*, 2021.
- [14] 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.
- [15] Flavio Schneider, Ojasv Kamal, Zhijing Jin, and Bernhard Schölkopf. Moûsai: Efficient text-to-music diffusion models. In *Proceedings of the 62nd Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, pages 8050–8068. Association for Computational Linguistics, 2024.

- [16] Ke Chen, Yusong Wu, Haohe Liu, Marianna Nezhurina, Taylor Berg-Kirkpatrick, and Shlomo Dubnov. Musicldm: Enhancing novelty in text-to-music generation using beat-synchronous mixup strategies. In *ICASSP 2024-2024 IEEE International Conference on Acoustics, Speech and Signal Processing (ICASSP)*, pages 1206–1210. IEEE, 2024.
- [17] Zach Evans, CJ Carr, Josiah Taylor, Scott H Hawley, and Jordi Pons. Fast timing-conditioned latent audio diffusion. In *Forty-first International Conference on Machine Learning*, 2024.
- [18] Zach Evans, Julian D Parker, CJ Carr, Zack Zukowski, Josiah Taylor, and Jordi Pons. Stable audio open. In *ICASSP 2025-2025 IEEE International Conference on Acoustics, Speech and Signal Processing (ICASSP)*, pages 1–5. IEEE, 2025.
- [19] Enric Gusó, Jordi Pons, Santiago Pascual, and Joan Serrà. On loss functions and evaluation metrics for music source separation. In *ICASSP 2022-2022 IEEE International Conference on Acoustics, Speech and Signal Processing (ICASSP)*, pages 306–310. IEEE, 2022.
- [20] Alexandre Défossez. Hybrid spectrogram and waveform source separation. *arXiv preprint arXiv:2111.03600*, 2021.
- [21] Yi Luo and Jianwei Yu. Music source separation with band-split rnn. *IEEE/ACM Transactions on Audio, Speech, and Language Processing*, 31:1893–1901, 2023.
- [22] Ui-Hyeop Shin, Sangyoun Lee, Taehan Kim, and Hyung-Min Park. Separate and Reconstruct: Asymmetric encoder-decoder for speech separation. In *Advances in Neural Information Processing Systems*, volume 37, pages 52215–52240, 2024.
- [23] Y Cem Subakan and Paris Smaragdis. Generative adversarial source separation. In *2018 IEEE International Conference on Acoustics, Speech and Signal Processing (ICASSP)*, pages 26–30. IEEE, 2018.
- [24] Qiuqiang Kong, Yong Xu, Wenwu Wang, Philip J. B. Jackson, and Mark D. Plumbley. Single-channel signal separation and deconvolution with generative adversarial networks. In *Proceedings of the 28th International Joint Conference on Artificial Intelligence*, page 2747–2753. AAAI Press, 2019.
- [25] Vivek Jayaram and John Thickstun. Source separation with deep generative priors. In *Proceedings of the 37th International Conference on Machine Learning*, 2020.
- [26] Vivek Narayanaswamy, Jayaraman J Thiagarajan, Rushil Anirudh, and Andreas Spanias. Unsupervised audio source separation using generative priors. In *INTERSPEECH*, 2020.
- [27] Emilian Postolache, Giorgio Mariani, Michele Mancusi, Andrea Santilli, Luca Cosmo, and Emanuele Rodolà. Latent autoregressive source separation. In *Proceedings of the AAAI Conference on Artificial Intelligence*, volume 37, pages 9444–9452, 2023.
- [28] Bo Chen, Chao Wu, and Wenbin Zhao. Sepdiff: Speech separation based on denoising diffusion model. In *ICASSP 2023-2023 IEEE International Conference on Acoustics, Speech and Signal Processing (ICASSP)*, pages 1–5. IEEE, 2023.
- [29] Robin Scheibler, Youna Ji, Soo-Whan Chung, Jaek Byun, Soyeon Choe, and Min-Seok Choi. Diffusion-based generative speech source separation. In *ICASSP 2023-2023 IEEE International Conference on Acoustics, Speech and Signal Processing (ICASSP)*, pages 1–5. IEEE, 2023.
- [30] Joan Serrà, Santiago Pascual, Jordi Pons, R Oguz Araz, and Davide Scaini. Universal speech enhancement with score-based diffusion. *arXiv preprint arXiv:2206.03065*, 2022.
- [31] Yen-Ju Lu, Zhong-Qiu Wang, Shinji Watanabe, Alexander Richard, Cheng Yu, and Yu Tsao. Conditional diffusion probabilistic model for speech enhancement. In *ICASSP 2022-2022 IEEE International Conference on Acoustics, Speech and Signal Processing (ICASSP)*, pages 7402–7406. Ieee, 2022.
- [32] Hao Yen, François G Germain, Gordon Wichern, and Jonathan Le Roux. Cold diffusion for speech enhancement. In *ICASSP 2023-2023 IEEE International Conference on Acoustics, Speech and Signal Processing (ICASSP)*, pages 1–5. IEEE, 2023.
- [33] Julius Richter, Simon Welker, Jean-Marie Lemercier, Bunlong Lay, and Timo Gerkmann. Speech enhancement and dereverberation with diffusion-based generative models. *IEEE/ACM Transactions on Audio, Speech, and Language Processing*, 31:2351–2364, 2023.
- [34] Giorgio Mariani, Irene Tallini, Emilian Postolache, Michele Mancusi, Luca Cosmo, and Emanuele Rodolà. Multi-source diffusion models for simultaneous music generation and separation. In *The Twelfth International Conference on Learning Representations*, 2024.

- [35] Emilian Postolache, Giorgio Mariani, Luca Cosmo, Emmanouil Benetos, and Emanuele Rodolà. Generalized multi-source inference for text conditioned music diffusion models. In *ICASSP 2024-2024 IEEE International Conference on Acoustics, Speech and Signal Processing (ICASSP)*, pages 6980–6984. IEEE, 2024.
- [36] Tornike Karchkhadze, Mohammad Rasool Izadi, and Shlomo Dubnov. Simultaneous music separation and generation using multi-track latent diffusion models. In *ICASSP 2025-2025 IEEE International Conference on Acoustics, Speech and Signal Processing (ICASSP)*, pages 1–5. IEEE, 2025.
- [37] Ethan Manilow, Gordon Wichern, Prem Seetharaman, and Jonathan Le Roux. Cutting music source separation some Slakh: A dataset to study the impact of training data quality and quantity. In *2019 IEEE Workshop on Applications of Signal Processing to Audio and Acoustics (WASPAA)*, pages 45–49. IEEE, 2019.
- [38] Zafar Rafii, Antoine Liutkus, Fabian-Robert Stöter, Stylianos Ioannis Mimilakis, and Rachel Bittner. The MUSDB18 corpus for music separation, December 2017. URL <https://doi.org/10.5281/zenodo.1117372>.
- [39] Igor Pereira, Felipe Araújo, Filip Korzeniowski, and Richard Vogl. MoisesDB: A dataset for source separation beyond 4-stems. *International Society for Music Information Retrieval Conference*, 2023.
- [40] Zhongweiyang Xu, Debottam Dutta, Yu-Lin Wei, and Romit Roy Choudhury. Multi-source music generation with latent diffusion. *arXiv preprint arXiv:2409.06190*, 2024.
- [41] Yi Luo and Nima Mesgarani. Conv-tasnet: Surpassing ideal time–frequency magnitude masking for speech separation. *IEEE/ACM transactions on audio, speech, and language processing*, 27(8):1256–1266, 2019.
- [42] Alexandre Défossez, Nicolas Usunier, Léon Bottou, and Francis Bach. Music source separation in the waveform domain. *arXiv preprint arXiv:1911.13254*, 2019.
- [43] Ge Zhu, Jordan Darefsky, Fei Jiang, Anton Selitskiy, and Zhiyao Duan. Music source separation with generative flow. *IEEE Signal Processing Letters*, 29:2288–2292, 2022.
- [44] Tsubasa Ochiai, Marc Delcroix, Yuma Koizumi, Hiroaki Ito, Keisuke Kinoshita, and Shoko Araki. Listen to what you want: Neural network-based universal sound selector. In *INTERSPEECH*, 2020.
- [45] Marc Delcroix, Jorge Bennisar Vázquez, Tsubasa Ochiai, Keisuke Kinoshita, Yasunori Ohishi, and Shoko Araki. SoundBeam: Target sound extraction conditioned on sound-class labels and enrollment clues for increased performance and continuous learning. *IEEE/ACM Transactions on Audio, Speech, and Language Processing*, 31:121–136, 2022.
- [46] Bandhav Veluri, Justin Chan, Malek Itani, Tuochao Chen, Takuya Yoshioka, and Shyamnath Gollakota. Real-time target sound extraction. In *ICASSP 2023-2023 IEEE International Conference on Acoustics, Speech and Signal Processing (ICASSP)*, pages 1–5. IEEE, 2023.
- [47] Hang Zhao, Chuang Gan, Andrew Rouditchenko, Carl Vondrick, Josh McDermott, and Antonio Torralba. The sound of pixels. In *Proceedings of the European conference on computer vision (ECCV)*, pages 570–586, 2018.
- [48] Hao-Wen Dong, Naoya Takahashi, Yuki Mitsufuji, Julian McAuley, and Taylor Berg-Kirkpatrick. CLIPSep: Learning text-queried sound separation with noisy unlabeled videos. In *Proceedings of International Conference on Learning Representations (ICLR)*, 2023.
- [49] Jie Hwan Lee, Hyeong-Seok Choi, and Kyogu Lee. Audio query-based music source separation. In *International Society for Music Information Retrieval Conference*, 2019.
- [50] Qiuqiang Kong, Ke Chen, Haohe Liu, Xingjian Du, Taylor Berg-Kirkpatrick, Shlomo Dubnov, and Mark D Plumbley. Universal source separation with weakly labelled data. *arXiv preprint arXiv:2305.07447*, 2023.
- [51] Kevin Kilgour, Beat Gfeller, Qingqing Huang, Aren Jansen, Scott Wisdom, and Marco Tagliasacchi. Text-driven separation of arbitrary sounds. In *INTERSPEECH*, 2022.
- [52] Xubo Liu, Haohe Liu, Qiuqiang Kong, Xinhao Mei, Jinzheng Zhao, Qiushi Huang, Mark D Plumbley, and Wenwu Wang. Separate what you describe: Language-queried audio source separation. *arXiv preprint arXiv:2203.15147*, 2022.
- [53] Xubo Liu, Qiuqiang Kong, Yan Zhao, Haohe Liu, Yi Yuan, Yuzhuo Liu, Rui Xia, Yuxuan Wang, Mark D Plumbley, and Wenwu Wang. Separate anything you describe. *IEEE/ACM Transactions on Audio, Speech, and Language Processing*, 2024.

- [54] Hao Ma, Zhiyuan Peng, Xu Li, Mingjie Shao, Xixin Wu, and Ju Liu. CLAPSep: Leveraging contrastive pre-trained model for multi-modal query-conditioned target sound extraction. *IEEE/ACM Transactions on Audio, Speech, and Language Processing*, 2024.
- [55] Hao Ma, Zhiyuan Peng, Xu Li, Yukai Li, Mingjie Shao, Qiuqiang Kong, and Ju Liu. Language-queried target sound extraction without parallel training data. In *ICASSP 2025-2025 IEEE International Conference on Acoustics, Speech and Signal Processing (ICASSP)*, pages 1–5. IEEE, 2025.
- [56] Yusong Wu, Ke Chen, Tianyu Zhang, Yuchen Hui, Taylor Berg-Kirkpatrick, and Shlomo Dubnov. Large-scale contrastive language-audio pretraining with feature fusion and keyword-to-caption augmentation. In *ICASSP 2023-2023 IEEE International Conference on Acoustics, Speech and Signal Processing (ICASSP)*, pages 1–5. IEEE, 2023.
- [57] Soroush Mehri, Kundan Kumar, Ishaan Gulrajani, Rithesh Kumar, Shubham Jain, Jose Sotelo, Aaron Courville, and Yoshua Bengio. SampleRNN: An unconditional end-to-end neural audio generation model. In *International Conference on Learning Representations*, 2017.
- [58] Chris Donahue, Julian McAuley, and Miller Puckette. Adversarial audio synthesis. In *International Conference on Learning Representations*, 2019.
- [59] Jesse Engel, Kumar Krishna Agrawal, Shuo Chen, Ishaan Gulrajani, Chris Donahue, and Adam Roberts. GANSynth: Adversarial neural audio synthesis. In *International Conference on Learning Representations*, 2019.
- [60] Ashish Vaswani, Noam Shazeer, Niki Parmar, Jakob Uszkoreit, Llion Jones, Aidan N Gomez, Łukasz Kaiser, and Illia Polosukhin. Attention is all you need. *Advances in neural information processing systems*, 30, 2017.
- [61] Zalán Borsos, Raphaël Marinier, Damien Vincent, Eugene Kharitonov, Olivier Pietquin, Matt Sharifi, Dominik Roblek, Olivier Teboul, David Grangier, Marco Tagliasacchi, et al. Audioldm: a language modeling approach to audio generation. *IEEE/ACM transactions on audio, speech, and language processing*, 31:2523–2533, 2023.
- [62] Felix Kreuk, Gabriel Synnaeve, Adam Polyak, Uriel Singer, Alexandre Défossez, Jade Copet, Devi Parikh, Yaniv Taigman, and Yossi Adi. AudioGen: Textually guided audio generation. In *The Eleventh International Conference on Learning Representations*, 2023.
- [63] Chris Donahue, Antoine Caillon, Adam Roberts, Ethan Manilow, Philippe Esling, Andrea Agostinelli, Mauro Verzett, Ian Simon, Olivier Pietquin, Neil Zeghidour, et al. Singsong: Generating musical accompaniments from singing. *arXiv preprint arXiv:2301.12662*, 2023.
- [64] Yixiao Zhang, Yukara Ikemiya, Woosung Choi, Naoki Murata, Marco A Martínez-Ramírez, Liwei Lin, Gus Xia, Wei-Hsiang Liao, Yuki Mitsufuji, and Simon Dixon. Instruct-MusicGen: Unlocking text-to-music editing for music language models via instruction tuning. *arXiv preprint arXiv:2405.18386*, 2024.
- [65] Huiwen Chang, Han Zhang, Lu Jiang, Ce Liu, and William T Freeman. Maskgit: Masked generative image transformer. In *Proceedings of the IEEE/CVF conference on computer vision and pattern recognition*, pages 11315–11325, 2022.
- [66] Hugo Flores Garcia, Prem Seetharaman, Rithesh Kumar, and Bryan Pardo. Vampnet: Music generation via masked acoustic token modeling. In *International Society for Music Information Retrieval Conference*, 2023.
- [67] Marco Comunità, Zhi Zhong, Akira Takahashi, Shiqi Yang, Mengjie Zhao, Koichi Saito, Yukara Ikemiya, Takashi Shibuya, Shusuke Takahashi, and Yuki Mitsufuji. Specmaskgit: Masked generative modeling of audio spectrograms for efficient audio synthesis and beyond. In *International Society for Music Information Retrieval Conference*, 2024.
- [68] Qingqing Huang, Daniel S Park, Tao Wang, Timo I Denk, Andy Ly, Nanxin Chen, Zhengdong Zhang, Zhishuai Zhang, Jiahui Yu, Christian Frank, et al. Noise2music: Text-conditioned music generation with diffusion models. *arXiv preprint arXiv:2302.03917*, 2023.
- [69] Shih-Lun Wu, Chris Donahue, Shinji Watanabe, and Nicholas J Bryan. Music controlnet: Multiple time-varying controls for music generation. *IEEE/ACM Transactions on Audio, Speech, and Language Processing*, 32:2692–2703, 2024.

- [70] Haohe Liu, Zehua Chen, Yi Yuan, Xinhao Mei, Xubo Liu, Danilo Mandic, Wenwu Wang, and Mark D. Plumbley. AudioLDM: text-to-audio generation with latent diffusion models. In *Proceedings of the 40th International Conference on Machine Learning, ICML'23*, 2023.
- [71] Haohe Liu, Yi Yuan, Xubo Liu, Xinhao Mei, Qiuqiang Kong, Qiao Tian, Yuping Wang, Wenwu Wang, Yuxuan Wang, and Mark D Plumbley. AudioLDM 2: Learning holistic audio generation with self-supervised pretraining. *IEEE/ACM Transactions on Audio, Speech, and Language Processing*, 2024.
- [72] Zach Evans, Julian D Parker, CJ Carr, Zack Zukowski, Josiah Taylor, and Jordi Pons. Long-form music generation with latent diffusion. *International Society for Music Information Retrieval Conference*, 2024.
- [73] Yuancheng Wang, Zeqian Ju, Xu Tan, Lei He, Zhizheng Wu, Jiang Bian, et al. Audit: Audio editing by following instructions with latent diffusion models. *Advances in Neural Information Processing Systems*, 36:71340–71357, 2023.
- [74] Bing Han, Junyu Dai, Weituo Hao, Xinyan He, Dong Guo, Jitong Chen, Yuxuan Wang, Yanmin Qian, and Xuchen Song. InstructME: an instruction guided music edit framework with latent diffusion models. In *Proceedings of the Thirty-Third International Joint Conference on Artificial Intelligence, IJCAI '24*, 2024.
- [75] Julian D Parker, Janne Spijkervet, Katerina Kosta, Furkan Yesiler, Boris Kuznetsov, Ju-Chiang Wang, Matt Avent, Jitong Chen, and Duc Le. Stemgen: A music generation model that listens. In *ICASSP 2024-2024 IEEE International Conference on Acoustics, Speech and Signal Processing (ICASSP)*, pages 1116–1120. IEEE, 2024.
- [76] Yao Yao, Peike Li, Boyu Chen, and Alex Wang. Jen-1 composer: A unified framework for high-fidelity multi-track music generation. In *Proceedings of the AAAI Conference on Artificial Intelligence*, pages 14459–14467, 2025.
- [77] Simon Rouard, Robin San Roman, Yossi Adi, and Axel Roebel. MusicGen-Stem: Multi-stem music generation and edition through autoregressive modeling. In *ICASSP 2025-2025 IEEE International Conference on Acoustics, Speech and Signal Processing (ICASSP)*, pages 1–5. IEEE, 2025.
- [78] Javier Nistal, Marco Pasini, Cyran Aouameur, Maarten Grachten, and Stefan Lattner. Diff-a-riff: Musical accompaniment co-creation via latent diffusion models. *International Society for Music Information Retrieval Conference*, 2024.
- [79] Javier Nistal, Marco Pasini, and Stefan Lattner. Improving musical accompaniment co-creation via diffusion transformers. *NeurIPS 2024 Workshop on AI-Driven Speech, Music, and Sound Generation (Audio Imagination)*, 2024.
- [80] Marco Pasini, Maarten Grachten, and Stefan Lattner. Bass accompaniment generation via latent diffusion. In *ICASSP 2024-2024 IEEE International Conference on Acoustics, Speech and Signal Processing (ICASSP)*, pages 1166–1170. IEEE, 2024.
- [81] Giorgio Strano, Chiara Ballanti, Donato Crisostomi, Michele Mancusi, Luca Cosmo, and Emanuele Rodolà. Stage: Stemmed accompaniment generation through prefix-based conditioning. *International Society for Music Information Retrieval Conference*, 2025.
- [82] Tero Karras, Miika Aittala, Timo Aila, and Samuli Laine. Elucidating the design space of diffusion-based generative models. *Advances in neural information processing systems*, 35:26565–26577, 2022.
- [83] Tornike Karchkhadze, Mohammad Rasool Izadi, Ke Chen, Gerard Assayag, and Shlomo Dubnov. Multi-Track MusicLDM: Towards versatile music generation with latent diffusion model. *arXiv preprint arXiv:2409.02845*, 2024.
- [84] Jonathan Ho and Tim Salimans. Classifier-free diffusion guidance. *NeurIPS Workshop on Deep Generative Models and Downstream Applications*, 2021.
- [85] Diederik P Kingma, Max Welling, et al. Auto-encoding variational bayes, 2013.
- [86] William Peebles and Saining Xie. Scalable diffusion models with transformers. In *Proceedings of the IEEE/CVF international conference on computer vision*, pages 4195–4205, 2023.
- [87] Tim Salimans and Jonathan Ho. Progressive distillation for fast sampling of diffusion models. In *International Conference on Learning Representations*, 2022.
- [88] 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.

- [89] Andreas Lugmayr, Martin Danelljan, Andres Romero, Fisher Yu, Radu Timofte, and Luc Van Gool. Repaint: Inpainting using denoising diffusion probabilistic models. In *Proceedings of the IEEE/CVF conference on computer vision and pattern recognition*, pages 11461–11471, 2022.
- [90] Ciprian Corneanu, Raghudeep Gadde, and Aleix M Martinez. Latentpaint: Image inpainting in latent space with diffusion models. In *Proceedings of the IEEE/CVF winter conference on applications of computer vision*, pages 4334–4343, 2024.
- [91] Sora Kim, Sungho Suh, and Minsik Lee. RAD: Region-aware diffusion models for image inpainting. *arXiv preprint arXiv:2412.09191*, 2024.
- [92] Tsiry Mayet, Pourya Shamsolmoali, Simon Bernard, Eric Granger, Romain HÉRAULT, and Clement Chatelain. TD-Paint: Faster diffusion inpainting through time aware pixel conditioning. In *The Thirteenth International Conference on Learning Representations*, 2025.
- [93] Jimmy Lei Ba, Jamie Ryan Kiros, and Geoffrey E Hinton. Layer normalization. *arXiv preprint arXiv:1607.06450*, 2016.
- [94] Yuxin Wu and Kaiming He. Group normalization. In *Proceedings of the European conference on computer vision (ECCV)*, pages 3–19, 2018.
- [95] Timo Lüddecke and Alexander Ecker. Image segmentation using text and image prompts. In *Proceedings of the IEEE/CVF conference on computer vision and pattern recognition*, pages 7086–7096, 2022.
- [96] Jiaming Song, Chenlin Meng, and Stefano Ermon. Denoising diffusion implicit models. In *International Conference on Learning Representations*, 2021.
- [97] Kevin Kilgour, Mauricio Zuluaga, Dominik Roblek, and Matthew Sharifi. Fréchet audio distance: A metric for evaluating music enhancement algorithms. In *INTERSPEECH*, 2018.
- [98] Shawn Hershey, Sourish Chaudhuri, Daniel PW Ellis, Jort F Gemmeke, Aren Jansen, R Channing Moore, Manoj Plakal, Devin Platt, Rif A Saurous, Bryan Seybold, et al. CNN architectures for large-scale audio classification. In *2017 IEEE international conference on acoustics, speech and signal processing (icassp)*, pages 131–135. IEEE, 2017.
- [99] Victor Weixin Liang, Yuhui Zhang, Yongchan Kwon, Serena Yeung, and James Y Zou. Mind the gap: Understanding the modality gap in multi-modal contrastive representation learning. *Advances in Neural Information Processing Systems*, 35:17612–17625, 2022.
- [100] Shuyang Gu, Dong Chen, Jianmin Bao, Fang Wen, Bo Zhang, Dongdong Chen, Lu Yuan, and Baining Guo. Vector quantized diffusion model for text-to-image synthesis. In *Proceedings of the IEEE/CVF conference on computer vision and pattern recognition*, pages 10696–10706, 2022.
- [101] Andrew Campbell, Joe Benton, Valentin De Bortoli, Thomas Rainforth, George Deligiannidis, and Arnaud Doucet. A continuous time framework for discrete denoising models. *Advances in Neural Information Processing Systems*, 35:28266–28279, 2022.
- [102] Aaron Lou, Chenlin Meng, and Stefano Ermon. Discrete diffusion modeling by estimating the ratios of the data distribution. In *Proceedings of the 41st International Conference on Machine Learning, ICML’24*, 2024.
- [103] Colin Raffel. *Learning-based methods for comparing sequences, with applications to audio-to-midi alignment and matching*. Columbia University, 2016.
- [104] Xun Huang and Serge Belongie. Arbitrary style transfer in real-time with adaptive instance normalization. In *Proceedings of the IEEE international conference on computer vision*, pages 1501–1510, 2017.
- [105] Ruben Ciranni, Giorgio Mariani, Michele Mancusi, Emilian Postolache, Giorgio Fabbro, Emanuele Rodolà, and Luca Cosmo. COCOLA: Coherence-oriented contrastive learning of musical audio representations. In *ICASSP 2025-2025 IEEE International Conference on Acoustics, Speech and Signal Processing (ICASSP)*, pages 1–5. IEEE, 2025.
- [106] Florian Grötschla, Ahmet Solak, Luca A Lanzendörfer, and Roger Wattenhofer. Benchmarking music generation models and metrics via human preference studies. In *ICASSP 2025-2025 IEEE International Conference on Acoustics, Speech and Signal Processing (ICASSP)*, pages 1–5. IEEE, 2025.
- [107] Azalea Gui, Hannes Gamper, Sebastian Braun, and Dimitra Emmanouilidou. Adapting frechet audio distance for generative music evaluation. In *ICASSP 2024-2024 IEEE International Conference on Acoustics, Speech and Signal Processing (ICASSP)*, pages 1331–1335. IEEE, 2024.

- [108] Yichen Huang, Zachary Novack, Koichi Saito, Jiatong Shi, Shinji Watanabe, Yuki Mitsufuji, John Thickstun, and Chris Donahue. Aligning text-to-music evaluation with human preferences. *arXiv preprint arXiv:2503.16669*, 2025.
- [109] Marco Pasini, Stefan Lattner, and George Fazekas. Music2Latent: Consistency autoencoders for latent audio compression. *International Society for Music Information Retrieval Conference*, 2024.
- [110] Marco Pasini, Stefan Lattner, and György Fazekas. Music2Latent2: Audio compression with summary embeddings and autoregressive decoding. In *ICASSP 2025-2025 IEEE International Conference on Acoustics, Speech and Signal Processing (ICASSP)*, pages 1–5. IEEE, 2025.
- [111] Junghyun Koo, Yunkee Chae, Chang-Bin Jeon, and Kyogu Lee. Self-refining of pseudo labels for music source separation with noisy labeled data. *International Society for Music Information Retrieval Conference*, 2023.
- [112] Dmitry Bogdanov, Minz Won, Philip Tovstogan, Alastair Porter, and Xavier Serra. The MTG-Jamendo dataset for automatic music tagging. In *Machine Learning for Music Discovery Workshop, International Conference on Machine Learning (ICML 2019)*, Long Beach, CA, United States, 2019. URL <http://hdl.handle.net/10230/42015>.
- [113] Michaël Defferrard, Kirell Benzi, Pierre Vandergheynst, and Xavier Bresson. FMA: A dataset for music analysis. In *International Society for Music Information Retrieval Conference*, 2016.

NeurIPS Paper Checklist

1. Claims

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

Answer: [Yes]

Justification: The abstract and introduction clearly describe the scope and contributions of the paper. We propose a novel three-track latent diffusion model that jointly addresses multi-track music generation, source imputation, and text-conditioned extraction. These capabilities are accurately summarized in the abstract and elaborated in itemized form in the introduction (please see Section 1. The claims align with the methods and experiments presented throughout the paper.

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: Please see Section 6.

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: Please see Appendix A.2 and B.

Guidelines:

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

4. Experimental result reproducibility

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

Answer: [Yes]

Justification: Please see Section 3, 4 and Appendix D.

Guidelines:

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

5. Open access to data and code

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

Answer: [Yes]

Justification: We released the full training code and all model checkpoints under an open-source license.

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: See Section 4 and Appendix D.

Guidelines:

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

7. Experiment statistical significance

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

Answer: [NA]

Justification: We did not perform repeated runs due to limited computational resources. While we do not report error bars or confidence intervals, the paper includes extensive evaluation across datasets, stem types, and tasks, with per-stem metrics and consistent trends observed across conditions. These provide empirical support for the reliability of our findings. Detailed results are included in the Section 5 and Appendix E.

Guidelines:

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

8. Experiments compute resources

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

Answer: [Yes]

Justification: We report our computational setup in Appendix D. However, we do not report execution time due to variability across models and hardware conditions.

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: The authors have carefully reviewed the NeurIPS Code of Ethics and confirm that the research complies with its guidelines.

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: As the proposed method involves generative modeling of music, we address potential broader impacts in Appendix I.

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 models are trained on a relatively small, open-source multi-track dataset, which limits the potential for misuse.

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: All third-party codebases and datasets used in this work are open-source and publicly available.

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: [\[Yes\]](#)

Justification: We released all code, training checkpoints, and inference scripts, accompanied by well-structured documentation detailing installation, usage, and reproduction steps.

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: [\[NA\]](#)

Justification: The study does not include any crowdsourced data collection or experiments involving human participants..

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: [NA]

Justification: This paper does not contain any crowdsourcing nor research with human subjects.

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: We did not involve any large language models in the development of the core methods.

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.

A Background

In this section, we review three foundational components that support our method: (1) the v-objective score-matching formulation for diffusion models [87], which underpins our latent modeling, and (2) the canonical inpainting algorithm introduced by Song et al. [11] and (3) further refined by Lugmayr et al. [89].

A.1 v-Objective Diffusion

We adopt the continuous-time diffusion framework proposed by Salimans et al. [87], which has been widely applied in recent music generation models [15, 17, 18, 72].

The forward perturbation kernel is defined as:

$$q(\mathbf{z}_\tau | \mathbf{z}_0) = \mathcal{N}(\mathbf{z}_\tau; \alpha_\tau \mathbf{z}_0, \beta_\tau^2 \mathbf{I}) \quad (14)$$

where:

$$\alpha_\tau = \cos(\phi_\tau), \quad \beta_\tau = \sin(\phi_\tau), \quad \phi_\tau = \frac{\pi}{2} \tau, \quad \tau \in [0, 1].$$

Instead of predicting the noise ϵ as in DDPMs [10], the v-objective formulation introduces a *velocity* target:

$$\mathbf{v}_\tau = \frac{\partial \mathbf{z}_\tau}{\partial \phi_\tau} = \frac{\partial \cos(\phi_\tau)}{\partial \phi_\tau} \mathbf{z}_0 + \frac{\partial \sin(\phi_\tau)}{\partial \phi_\tau} \epsilon = -\sin(\phi_\tau) \mathbf{z}_0 + \cos(\phi_\tau) \epsilon = \alpha_\tau \epsilon - \beta_\tau \mathbf{z}_0, \quad (15)$$

and trains the denoiser $f_\theta(\mathbf{z}_\tau, \tau)$ to directly regress this velocity.

Recovering $\mathbf{z}_0$ and ϵ . Using the definition of $\mathbf{v}_\tau$ and recalling the perturbation process defined in Eq. (2),

$$\mathbf{z}_\tau = \alpha_\tau \mathbf{z}_0 + \beta_\tau \epsilon \quad (16)$$

we can rearrange the terms to recover the clean latent by:

$$\begin{aligned} \sin(\phi_\tau) \mathbf{z}_0 &= \cos(\phi_\tau) \epsilon - \mathbf{v}_\tau \\ &= \frac{\cos(\phi_\tau)}{\sin(\phi_\tau)} (\mathbf{z}_\tau - \cos(\phi_\tau) \mathbf{z}_0) - \mathbf{v}_\tau, \\ \sin^2(\phi_\tau) \mathbf{z}_0 &= \cos(\phi_\tau) \mathbf{z}_\tau - \cos^2(\phi_\tau) \mathbf{z}_0 - \sin(\phi_\tau) \mathbf{v}_\tau, \\ (\sin^2(\phi_\tau) + \cos^2(\phi_\tau)) \mathbf{z}_0 &= \mathbf{z}_0 = \cos(\phi_\tau) \mathbf{z}_\tau - \sin(\phi_\tau) \mathbf{v}_\tau, \end{aligned}$$

thus we get

$$\mathbf{z}_0 = \alpha_\tau \mathbf{z}_\tau - \beta_\tau \mathbf{v}_\tau. \quad (17)$$

Similarly, the noise vector can be expressed as:

$$\epsilon = \beta_\tau \mathbf{z}_\tau + \alpha_\tau \mathbf{v}_\tau \quad (18)$$

Score function approximation. The marginal score $\nabla_{\mathbf{z}_\tau} \log q_{\mathbf{z}_\tau}(\mathbf{z}_\tau)$ at timestep (or noise level) τ is approximated as:

$$\begin{aligned} \nabla_{\mathbf{z}_\tau} \log q_{\mathbf{z}_\tau}(\mathbf{z}_\tau) &\approx \frac{\alpha_\tau \hat{\mathbf{z}}_\theta(\mathbf{z}_\tau) - \mathbf{z}_\tau}{\beta_\tau^2} \\ &= \frac{\alpha_\tau^2 \mathbf{z}_\tau - \alpha_\tau \beta_\tau f_\theta(\mathbf{z}_\tau, \tau) - \mathbf{z}_\tau}{\beta_\tau^2} \\ &= \frac{-\beta_\tau^2 \mathbf{z}_\tau - \alpha_\tau \beta_\tau f_\theta(\mathbf{z}_\tau, \tau)}{\beta_\tau^2} \\ &= -\mathbf{z}_\tau - \frac{\alpha_\tau}{\beta_\tau} f_\theta(\mathbf{z}_\tau, \tau), \end{aligned}$$

where $\hat{\mathbf{z}}_\theta(\mathbf{z}_\tau) = \alpha_\tau \mathbf{z}_\tau - \beta_\tau f_\theta(\mathbf{z}_\tau, \tau)$, from the Eq. (17).

DDIM sampling. We adopt the DDIM sampling [96] to generate samples in a non-stochastic, deterministic manner. Given a latent $\mathbf{z}_\tau$, each reverse step is computed as:

$$\begin{aligned} \hat{\mathbf{v}}_\tau &= f_\theta(\mathbf{z}_\tau, \tau) \\ \hat{\mathbf{z}}_0 &= \alpha_\tau \mathbf{z}_\tau - \beta_\tau \hat{\mathbf{v}}_\tau \\ \hat{\epsilon}_\tau &= \beta_\tau \mathbf{z}_\tau + \alpha_\tau \hat{\mathbf{v}}_\tau \\ \mathbf{z}_{\tau'} &= \alpha_{\tau'} \hat{\mathbf{z}}_0 + \beta_{\tau'} \hat{\epsilon}_\tau, \end{aligned}$$

where $\tau' < \tau$ is taken from a linearly spaced decreasing schedule from 1 to 0. We use 250 inference steps in all experiments.

A.2 Canonical Inpainting in Score-Based Models

The core idea behind inpainting in score-based generative models is to estimate the score of the unknown region conditioned on the known region [11, 34].

Let K denote the set of all tracks. Suppose a subset $\Omega \subset K$ is observed (i.e., known), and let $\Gamma = K \setminus \Omega$ denote the complement, i.e., the unobserved tracks we aim to inpaint. Define $\mathbf{z}^\Omega := \{z^{(k)}\}_{k \in \Omega}$ and $\mathbf{z}^\Gamma := \{z^{(k)}\}_{k \in \Gamma}$. The goal is to approximate the conditional score:

$$\nabla_{\mathbf{z}^\Gamma} \log q_\tau(\mathbf{z}^\Gamma | \mathbf{z}_0^\Omega). \quad (19)$$

This conditional gradient is generally intractable for a score model trained only on joint marginals. However, following Song et al. [11], we can approximate it via:

$$\begin{aligned} q_\tau(\mathbf{z}^\Gamma | \mathbf{z}_0^\Omega) &= \int q_\tau(\mathbf{z}^\Gamma, \mathbf{z}^\Omega | \mathbf{z}_0^\Omega) d\mathbf{z}^\Omega \\ &= \int q_\tau(\mathbf{z}^\Gamma | \mathbf{z}^\Omega, \mathbf{z}_0^\Omega) q_\tau(\mathbf{z}^\Omega | \mathbf{z}_0^\Omega) d\mathbf{z}^\Omega \\ &= \mathbb{E}_{q_\tau(\mathbf{z}^\Omega | \mathbf{z}_0^\Omega)} \left[q_\tau(\mathbf{z}^\Gamma | \mathbf{z}^\Omega, \mathbf{z}_0^\Omega) \right] \\ &\approx \mathbb{E}_{q_\tau(\mathbf{z}^\Omega | \mathbf{z}_0^\Omega)} \left[q_\tau(\mathbf{z}^\Gamma | \mathbf{z}^\Omega) \right] \end{aligned} \quad (20)$$

$$\approx q_\tau(\mathbf{z}^\Gamma | \hat{\mathbf{z}}_\tau^\Omega), \quad (21)$$

where $\hat{\mathbf{z}}_\tau^\Omega \sim q_\tau(\mathbf{z}^\Omega | \mathbf{z}_0^\Omega) = \mathcal{N}(\mathbf{z}^\Omega; \alpha_\tau \mathbf{z}_0^\Omega, \beta_\tau^2 \mathbf{I})$ is a noised sample of the known region. Accordingly, the conditional score can be approximated as:

$$\begin{aligned} \nabla_{\mathbf{z}^\Gamma} \log q_\tau(\mathbf{z}^\Gamma | \mathbf{z}_0^\Omega) &\approx \nabla_{\mathbf{z}^\Gamma} \log q_\tau(\mathbf{z}^\Gamma | \hat{\mathbf{z}}_\tau^\Omega) \\ &= \nabla_{\mathbf{z}^\Gamma} \log q_\tau([\mathbf{z}^\Gamma; \hat{\mathbf{z}}_\tau^\Omega]), \end{aligned}$$

where $[\mathbf{z}^\Gamma; \hat{\mathbf{z}}_\tau^\Omega]$ denotes a composite latent vector, such that the known region is replaced by $\hat{\mathbf{z}}_\tau^\Omega$, while the unknown region remains as $\mathbf{z}^\Gamma$, adopting the same notation as Song et al. [11].

This approximation enables zero-shot inpainting without requiring retraining: at each diffusion timestep, a noised version of the known latents is sampled, concatenated with the current estimate of the unknown latents, and passed to the score model. The resulting gradient is then applied to update only the unknown region. This process is repeated throughout the reverse diffusion trajectory.

A.3 RePaint

Lugmayr et al. proposed *RePaint* [89], a resampling-based mechanism that improves score-based inpainting by repeating the diffusion process across multiple forward–reverse cycles. Their key insight is that, in conventional inpainting (as described in Eq. (21)), the sampled noise for the known region is independent of the generated (inpainted) region. This lack of synchronization can lead to semantic inconsistencies and disharmony between the known and unknown parts of the sample.

To address this, RePaint introduces a resampling mechanism during generation. At each denoising timestep, the algorithm alternates between one reverse diffusion step and one forward diffusion step, repeating this cycle U times. These micro-steps refine the sampling distribution and can have the effect of partially marginalizing over the known region at noise level τ in Equation (21), thereby reducing the approximation error inherent in conditional score estimation. This iterative resampling procedure improves consistency but incurs a higher computational cost, as each denoising step requires multiple forward–reverse passes, making RePaint significantly more expensive than standard inpainting methods.

While originally proposed for DDPM-based models, RePaint can be adapted to velocity-based objectives as used in our framework. We apply this adaptation in the sampling procedure described in Algorithm 1. Setting the resampling count $U = 1$ recovers the canonical single-sample inpainting method described in Appendix A.2.

B Adaptive Timestep Conditioning and Score Approximation

Conventional inpainting approaches approximate the conditional score in Eq. (19) using a single sampled estimate of the known latents. This corresponds to a high-variance Monte Carlo estimate of the expectation over $q_\tau(\mathbf{z}^\Omega | \mathbf{z}_0^\Omega)$, which may lead to instability — especially at high noise levels.

By contrast, the adaptive timestep model proposed in Section 3.4, inspired by TD-Paint [92], circumvents this marginalization by training the model to directly approximate the conditional score.

Algorithm 1 Inpainting using the RePaint approach.

Input:
Number of timesteps T ;
Re-denoising steps per reverse step U ;
Noise schedule $\{\tau_i\}_{i=0}^T$ with α_τ, β_τ ;
Binary mask m (1 for known, 0 for unknown);
Known clean (masked) latents $\mathbf{z}_0^{\text{known}}$;
Denoiser network f_θ

```
1:  $\mathbf{z}_{\tau_T} \sim \mathcal{N}(\mathbf{0}, I)$ 
2: for  $i = T, \dots, 1$  do
3:   for  $u = 1, \dots, U$  do
4:      $\hat{\mathbf{v}}_{\tau_i} \leftarrow f_\theta(\mathbf{z}_{\tau_i}, \tau_i)$  ▷ DDIM sampling step
5:      $\hat{\mathbf{z}}_0 \leftarrow \alpha_{\tau_i} \mathbf{z}_{\tau_i} - \beta_{\tau_i} \hat{\mathbf{v}}_{\tau_i}$ 
6:      $\hat{\epsilon} \leftarrow \beta_{\tau_i} \mathbf{z}_{\tau_i} + \alpha_{\tau_i} \hat{\mathbf{v}}_{\tau_i}$ 
7:      $\hat{\mathbf{z}}_{\tau_{i-1}}^{\text{unknown}} \leftarrow \alpha_{\tau_{i-1}} \hat{\mathbf{z}}_0 + \beta_{\tau_{i-1}} \hat{\epsilon}$ 
8:      $\epsilon \sim \mathcal{N}(\mathbf{0}, I)$  if  $i > 1$ , else 0 ▷ Sample the known regions
9:      $\mathbf{z}_{\tau_{i-1}}^{\text{known}} \leftarrow \alpha_{\tau_{i-1}} \mathbf{z}_0^{\text{known}} + \beta_{\tau_{i-1}} \epsilon$  ▷ Combine known and generated regions
10:     $\mathbf{z}_{\tau_{i-1}} \leftarrow m \odot \mathbf{z}_{\tau_{i-1}}^{\text{known}} + (1 - m) \odot \hat{\mathbf{z}}_{\tau_{i-1}}^{\text{unknown}}$  ▷ Reapply forward process
11:    if  $u < U$  and  $t > 1$  then
12:       $\mathbf{z}_{\tau_i} \sim \mathcal{N}\left(\frac{\alpha_{\tau_i}}{\alpha_{\tau_{i-1}}} \mathbf{z}_{\tau_{i-1}}, \left(1 - \frac{\alpha_{\tau_i}^2}{\alpha_{\tau_{i-1}}^2}\right) I\right)$ 
13:    end if
14:  end for
15: end for
16: return  $\mathbf{z}_{\tau_0}$ 
```

Following the notation in Section 3.4 and Appendix A.2, we assume $\mathbf{z}_0 = [\mathbf{z}_0^\Gamma; \mathbf{z}_0^\Omega]$ is a clean sample from the dataset. Then, using an alternative factorization:

$$\begin{aligned} q_\tau(\mathbf{z}_\tau^\Gamma | \mathbf{z}_0^\Omega) &= \int q_\tau(\mathbf{z}_\tau^\Gamma, \mathbf{x}_0^\Gamma | \mathbf{z}_0^\Omega) d\mathbf{x}_0^\Gamma \\ &= \int q_\tau(\mathbf{z}_\tau^\Gamma | \mathbf{x}_0^\Gamma, \mathbf{z}_0^\Omega) q(\mathbf{x}_0^\Gamma | \mathbf{z}_0^\Omega) d\mathbf{x}_0^\Gamma \\ &= \mathbb{E}_{q(\mathbf{x}_0^\Gamma | \mathbf{z}_0^\Omega)} \left[q_\tau(\mathbf{z}_\tau^\Gamma | \mathbf{x}_0^\Gamma, \mathbf{z}_0^\Omega) \right] \end{aligned} \quad (22)$$

$$\approx q_\tau(\mathbf{z}_\tau^\Gamma | \mathbf{z}_0^\Gamma, \mathbf{z}_0^\Omega) = q_\tau(\mathbf{z}_\tau^\Gamma | \mathbf{z}_0^\Gamma) \quad (23)$$

$$= \mathcal{N}(\mathbf{z}_\tau^\Gamma; \alpha_\tau \mathbf{z}_0^\Gamma, \beta_\tau^2 \mathbf{I}), \quad (24)$$

where the approximation assumes that $\mathbf{z}_0^\Gamma \sim q(\mathbf{x}_0^\Gamma | \mathbf{z}_0^\Omega)$ is available from the dataset. Unlike the marginalization-based approximation in Eq. (21), this expression introduces no sampling noise during inference, thereby reducing variance.

From this, the conditional score can be written as:

$$\nabla_{\mathbf{z}_\tau^\Gamma} \log q_\tau(\mathbf{z}_\tau^\Gamma | \mathbf{z}_0^\Omega) \approx \frac{\alpha_\tau \mathbf{z}_0^\Gamma - \mathbf{z}_\tau^\Gamma}{\beta_\tau^2} \quad (25)$$

$$\approx \frac{\alpha_\tau \hat{\mathbf{z}}_\theta(\mathbf{z}_\tau^\Gamma, \mathbf{z}_0^\Omega, \boldsymbol{\tau}^\Gamma) - \mathbf{z}_\tau^\Gamma}{\beta_\tau^2}, \quad (26)$$

$$= -\mathbf{z}_\tau^\Gamma - \frac{\alpha_\tau}{\beta_\tau} f_\theta(\mathbf{z}_\tau^\Gamma, \mathbf{z}_0^\Omega, \boldsymbol{\tau}^\Gamma)_\Gamma, \quad (27)$$

where $f_\theta(\cdot)_\Gamma$ denotes the output corresponding to the unknown region. The per-track timestep vector $\boldsymbol{\tau}^\Gamma$ is defined as:

$$\boldsymbol{\tau}_k^\Gamma = \begin{cases} \tau, & \text{if } k \in \Gamma \\ 0, & \text{if } k \in \Omega \end{cases} \quad \text{for each } k \in K, \quad (28)$$

Algorithm 2 Inpainting using adaptive timestep approach

Input:
Number of timesteps T ;
Re-denoising steps per reverse step U ;
Noise schedule $\{\tau_i\}_{i=0}^T$ with α_τ, β_τ ;
Binary mask m (1 for known, 0 for unknown);
Known clean (masked) latents $\mathbf{z}_0^{\text{known}}$;
Denoiser network f_θ

- 1: $\epsilon \sim \mathcal{N}(\mathbf{0}, \mathbf{I})$
- 2: $\tau_T \leftarrow \tau_T(1 - m)$
- 3: $\mathbf{z}_{\tau_T} = m\mathbf{z}_0^{\text{known}} + (1 - m)\epsilon$
- 4: **for** $i = T, \dots, 1$ **do**
- 5: ▷ Partial DDIM sampling over unknown region
- 6: $\hat{\mathbf{v}}_{\tau_i} \leftarrow f_\theta(\mathbf{z}_{\tau_i}, \tau_i)$
- 7: $\hat{\mathbf{z}}_0 \leftarrow \alpha_{\tau_i}\mathbf{z}_{\tau_i} - \beta_{\tau_i}\hat{\mathbf{v}}_{\tau_i}$
- 8: $\hat{\epsilon} \leftarrow \beta_{\tau_i}\mathbf{z}_{\tau_i} + \alpha_{\tau_i}\hat{\mathbf{v}}_{\tau_i}$
- 9: $\hat{\mathbf{z}}_{\tau_{i-1}}^{\text{unknown}} \leftarrow \alpha_{\tau_{i-1}}\hat{\mathbf{z}}_0 + \beta_{\tau_{i-1}}\hat{\epsilon}$
- 10: $\tau_{i-1} \leftarrow \tau_{i-1}(1 - m)$
- 11: $\mathbf{z}_{\tau_{i-1}} \leftarrow m\mathbf{z}_0^{\text{known}} + (1 - m)\hat{\mathbf{z}}_{\tau_{i-1}}^{\text{unknown}}$
- 12: **end for**
- 13: **return** $\mathbf{z}_{\tau_0}$

as in Section 3.4. The model is trained using the velocity objective in Eq. (12), restricted to the unknown region. This allows the denoiser to explicitly learn the conditional score on $\mathbf{z}_\tau^{\text{r}}$, avoiding the need for stochastic marginalization and improving accuracy in conditional inpainting tasks. The full sampling procedure is detailed in Algorithm 2.

C Iterative Generation

In addition to the one-stage mixture generation described in Section 3.3, MGE-LDM also supports an iterative, stem-by-stem synthesis procedure. This approach constructs a full mixture by sequentially generating individual sources, leveraging the partial generation mechanism at each step.

Let $\mathcal{I} = \{c_i^{(s)}\}_i$ be an (ordered) set of CLAP embeddings corresponding to the desired instrument description. At the first iteration ($i = 1$), we generate an initial source latent $\hat{z}_1^{(s)}$ by sampling with the model conditioned only on the prompt $c_1^{(s)}$:

$$\tilde{z}^{(m)}, \tilde{z}^{(u)}, \hat{z}_1^{(s)} \sim p_\theta(z^{(m)}, z^{(u)}, z^{(s)} | \emptyset, \emptyset, c_1^{(s)}),$$

and set $z_1^{(u)} = \hat{z}_1^{(s)}$ as the initial submixture latent. For each subsequent iteration $i = 2, \dots, |\mathcal{I}|$, we follow the iterative imputation strategy of partial generation, treating the current submixture as the accumulated sum of decoded sources from the previous steps.

Finally, the full mixture waveform is constructed by decoding and summing the generated source latents:

$$x^{(m)} = \sum_{i=1}^{|\mathcal{I}|} D(\hat{z}_i^{(s)}).$$

A preliminary evaluation of this iterative procedure is presented in Appendix E

D Experimental Details

This appendix provides a comprehensive description of datasets, baseline implementations, model architecture, and training settings used throughout our experiments.

D.1 Datasets

We train and evaluate on three multi-track music datasets: Slakh2100 [37], MUSDB18 [38], and MoisesDB [39]. All datasets consist of mixture tracks paired with isolated stem recordings. A summary of stem abbreviations is provided in Table 4.

Table 4: **Abbreviations of instrument stems.** The table lists all abbreviations used throughout the paper along side their corresponding full instrument labels, grouped by dataset.

Abbr.	Common					Slakh2100						MoisesDB			
	B	D	G	P	V	Brs.	C.P.	Org.	Pipe	Reed	Str.	S.Lead	S.Pad	B.str	Perc.
Inst.	bass	drums	guitar	piano	vocals	brass	chromatic percussion	organ	pipe	reed	strings	synth lead	synth pad	bowed strings	percussion

Slakh2100 is derived from the Lakh MIDI Dataset v0.1 [103] and contains synthesized tracks rendered with sample-based virtual instruments. It comprises 2100 songs divided into training (1500), validation (375), and test (225) splits, totaling approximately 145 hours of audio. It includes a wide variety of instrument classes (e.g., bass, drums, guitar, piano, strings, synth pad, etc.). We adopt the naming $\mathbf{S}_A$ to denote a subset containing only bass, drums, guitar, and piano, the four classes used by MSDM and MSG-LD, and $\mathbf{S}_B$ to denote the complementary subset of the remaining stems. We follow the official dataset splits provided by Slakh2100 for training, validation, and testing.

MUSDB18 consists of 150 real-world music recordings with four stems: drums, bass, other, and vocals. We use all 100 tracks from the official training split for training and the 50-track test split for evaluation. The total dataset length is approximately 10 hours.

MoisesDB comprises 240 songs (14 hours total) contributed by 47 artists across 12 genres. Each stem in the song is annotated with a two-tier stem taxonomy. Each track is decomposed into its constituent sources and annotated using a two-level hierarchical taxonomy of stem classes. We aggregate all second-level tracks into their corresponding top-level classes. Among the 11 stem classes, we evaluate only the 7 unambiguous stems (e.g., bass, percussion, vocals, etc.). For evaluation, we randomly sample 24 tracks (10%) as the test set and use the remaining tracks for training.

Data Construction. We train our model using randomly constructed 3-track tuples (mix, sub, src). A source stem is randomly selected from the available stems, and the remaining stems are aggregated into a submixture. We select non-silent segments from the source track whenever possible, allowing up to 10 random resampling attempts per instance. The same temporal offset is applied across all stems to ensure alignment. For generation evaluation, we sample 300 random segments per test set. For source extraction, we sample between 150 and 700 non-silent segments per instrument class. All audio is downsampled to 16 kHz.

D.2 Baseline Implementations

All baseline metrics are recomputed on our test splits for a fair comparison. The following models are used:

- MSDM [34]: We use the official implementation and pretrained checkpoint.¹ Since MSDM operates at 22 kHz, we upsample our 16 kHz test audio for inference and downsample the output back to 16 kHz.
- MSG-LD [36]: As no checkpoint is publicly released, we reproduce the model by retraining it from the official codebase.²
- HDemucs [20]: We train a 16 kHz version of Hybrid Demucs using the demucs_lightning implementation.³
- AudioSep [53]: We evaluate using the publicly available implementation and checkpoint provided by the authors.⁴ Since AudioSep operates at a sampling rate of 32 kHz, we upsample all test audio from 16 kHz to 32 kHz before inference and subsequently downsample the separated outputs back to 16 kHz for evaluation consistency.

D.3 Model Architecture and Training

Autoencoder. We adopt the VAE-based architecture from Stable Audio [72], with a downsampling ratio of 2048, yielding a 7.8125 Hz latent resolution and 64 latent channels. We train the autoencoder on all training subsets from Slakh2100, MUSDB18, and MoisesDB using 16 kHz mono audio for 600K steps with a batch size of 16.

Diffusion Model. In practice, the three latent representations are concatenated along the channel dimension, such that the input to the diffusion model becomes $\text{Concat}[z^{(m)}, z^{(u)}, z^{(s)}] \in \mathbb{R}^{3C \times L}$. We use a DiT backbone [86] with 24 transformer blocks with 48 heads, and a projected latent dimension of 1536 (3 tracks $\times$ 512

¹<https://github.com/gladia-research-group/multi-source-diffusion-models>

²<https://github.com/karchkha/MSG-LD>

³https://github.com/KinWaiCheuk/demucs_lightning

⁴<https://github.com/Audio-AGI/AudioSep>

Table 5: **Source extraction performance of RePaint-based methods vs. MGE-LDM.** Metrics are reported as Log-Mel L1 distance $\downarrow$. T indicates the number of reverse timesteps, and U specifies the number of denoising operations per reverse step (i.e., $U-1$ intermediate resampling steps). The case $U = 1$ corresponds to the canonical single-sample conditional score approximation [11], as described in Appendix A.2. MGE-LDM uses adaptive timestep conditioning without resampling.

Model	T	U	S_A					S_B								
			B	D	G	P	Avg.	Brs.	C.P.	Org.	Pipe	Reed	Str.	S.Lead	S.Pad	Avg.
MGE (ours)	250	1	1.68	2.71	2.69	2.16	2.31	3.43	2.16	1.84	2.33	3.07	2.44	2.31	1.93	2.48
RePaint-based [89]	250	1	2.00	3.20	3.15	2.83	2.79	4.75	2.47	2.79	4.27	4.65	3.06	3.73	2.80	3.56
	125	2	1.89	2.75	2.91	2.68	2.55	4.33	2.37	6.67	3.84	4.27	2.82	3.64	2.58	3.81
	50	5	1.80	2.43	2.83	2.55	2.40	3.89	2.32	2.57	3.35	3.81	2.64	3.56	2.42	3.07
	25	10	1.77	2.28	2.80	2.51	2.34	3.79	2.31	2.50	3.15	3.71	2.66	3.44	2.40	2.99
	250	2	1.90	2.65	2.94	2.65	2.53	4.23	2.41	2.71	3.82	4.27	2.84	3.66	2.59	3.31
	250	4	1.79	2.34	2.83	2.59	2.38	3.95	2.34	2.66	3.42	3.88	2.69	3.55	2.45	3.12

each). Timestep embeddings are prepended to the input vector of the transformer. CLAP embeddings for each track are processed by independent projection layers (without weight sharing) to produce scale and shift parameters for AdaIN-style conditioning [104]. These are applied group-wise via GroupNorm within each DiT layer to modulate the corresponding track-specific activations. Text embeddings are obtained from CLAP [56] using the checkpoint `music_audioset_epoch_15_esc_90.14.pt` via the `laion-clap` library.⁵ Our implementation builds upon the official `stable-audio-tools` repository from Stability AI⁶ and the training framework from `friendly-stable-audio-tools`.⁷ All models were trained on a single NVIDIA RTX 6000 GPU (48 GB of memory).

E Ablation Study

This section presents ablation experiments designed to further analyze the key components of our framework. Unless otherwise specified, all models are trained on the combined $S_A + S_B$ dataset. Each of our models is evaluated with the same configuration as in Section 5, using $T = 250$ denoising steps during sampling.

E.1 Comparison with Canonical Inpainting Methods

We assess the effectiveness of our adaptive timestep conditioning strategy by comparing it against two prior approaches: the canonical one-sample conditional score approximation (Appendix A.2) and the RePaint method [89] (Appendix A.3). Table 5 reports the results of the source extraction task.

In RePaint, U denotes the number of denoising steps performed per reverse timestep: one denoising step followed by $U-1$ forward (resampling) steps. As a result, the total number of denoising steps becomes $T \times U$ during the full inpainting process. Note that setting $U = 1$ recovers the canonical single-sample estimator in Eq. (21).

We observe that, for a fixed number of denoising steps, using fewer timesteps T with more resampling cycles U generally improves performance, confirming observations in the original RePaint paper. We hypothesize that repeated resampling helps stabilize conditional generation by mitigating the noise mismatch between observed and unobserved regions, particularly at high noise levels, where observed latents contain little informative content and single-sample approximations of Eq. (21) become highly unreliable. While this approach does not yield a precise marginal score estimate, it heuristically improves inpainting quality through localized refinement.

Interestingly, we also observe that RePaint configurations with larger total denoising steps — such as $T = 250, U = 2$ and $T = 250, U = 4$ — consistently underperform compared to $T = 25, U = 10$, across all stems in both S_A and S_B . This suggests that, for inpainting tasks, accurately modeling the conditional score at each timestep is more critical than simply increasing the number of reverse steps. As RePaint approximates the conditional score by marginalizing over perturbed conditions via resampling, performance benefits are observed primarily through increased resampling (U), not longer trajectories (T).

Nevertheless, our adaptive timestep model outperforms all RePaint variants across both datasets, with the sole exception of drums in S_A . By directly learning track-specific conditional scores during training, our method eliminates the need for inference-time marginalization, resulting in lower variance and improved reconstruction quality.

⁵<https://github.com/LAION-AI/CLAP>

⁶<https://github.com/Stability-AI/stable-audio-tools>

⁷<https://github.com/yukara-ikemiya/friendly-stable-audio-tools>

Table 6: **Total generation performance across modeling variants.** Metrics are reported as FAD $\downarrow$. All models are trained on $\mathbf{S}_A + \mathbf{S}_B$. The baseline model uses uniform (non-adaptive) timesteps across all tracks. MGE variants apply adaptive timestep conditioning and test the impact of normalizations and CFG dropout rates. Values in parentheses indicate generation conditioned on the text prompt "*The sound of the bass, drums, guitar, and piano*".

Model	Testset			
	$\mathbf{S}_A$	$\mathbf{S}_{\text{Full}}$	$\mathbf{M}_u$	$\mathbf{M}_o$
Non-adaptive	3.26 (2.00)	0.79	5.12	4.51
MGE (adaptive)	3.14 (2.24)	0.63	5.46	4.73
- w/o GroupNorm	3.48 (2.44)	0.76	5.61	4.88
- CFG dropout $p=0.5$	3.12 (2.67)	0.58	5.43	4.82

Table 7: **Source extraction performance under different architectural and CFG settings.** Metrics are reported as Log-Mel L1 distance $\downarrow$. All models are trained on $\mathbf{S}_A + \mathbf{S}_B$. GN and LN denote GroupNorm and LayerNorm, respectively. p indicates the classifier-free guidance (CFG) dropout rate applied to each track’s conditioning vector, and s refers to the CFG guidance scale.

Norm.	p	s	$\mathbf{S}_A$					$\mathbf{S}_B$								
			B	D	G	P	Avg.	Brs.	C.P.	Org.	Pipe	Reed	Str.	S.Lead	S.Pad	Avg.
GN	0.1	2.0	1.68	2.71	2.69	2.16	2.31	3.43	2.16	1.84	2.33	3.07	2.44	2.31	1.93	2.43
LN	0.1	2.0	1.67	4.22	2.65	2.15	2.67	3.42	2.35	1.97	2.40	3.45	2.51	2.32	2.03	2.55
GN	0.5	2.0	1.78	1.96	2.62	1.96	2.08	3.37	2.22	1.97	2.36	2.89	2.44	2.07	1.89	2.40
GN	0.1	1.0	1.67	2.79	2.70	2.05	2.30	3.24	2.23	1.85	2.25	2.88	2.28	2.27	1.87	2.35
GN	0.1	4.0	1.77	2.49	2.79	2.27	2.33	3.64	2.17	1.91	2.42	3.20	2.66	2.31	2.06	2.54
GN	0.1	8.0	1.93	2.49	2.93	2.41	2.44	4.35	2.28	2.01	2.57	3.47	3.27	2.42	2.30	2.83

A potential concern is whether optimizing for adaptive timestep-conditional inference might degrade generation quality when using uniform timestep schedules across tracks. To assess this, we evaluate our adaptive timestep model with a uniform timestep vector $\tau = (\tau, \tau, \tau)$, which corresponds to the total generation task, and compare it to a baseline trained with non-adaptive, shared timesteps.

As shown in Table 6, comparing the non-adaptive uniform timestep baseline model and our model, both models achieve comparable FAD scores, indicating that timestep adaptation preserves generation performance under uniform scheduling while providing significant advantages for inpainting tasks.

E.2 Additional Design Ablations

We additionally investigate the impact of various modeling and training choices, including normalization strategies and classifier-free guidance (CFG) dropout rates.

Table 6 includes results from models trained with LayerNorm instead of GroupNorm, following the original DiT architecture, as well as a variant using a higher CFG dropout rate of $p = 0.5$. We observe that GroupNorm slightly outperforms LayerNorm across all test sets, supporting the use of track-wise normalization in our multi-track setting. Regarding CFG dropout, increasing the dropout rate improves unconditional generation performance, particularly on $\mathbf{S}_A$ and $\mathbf{S}_B$. However, when conditioned on the text prompt (values in parentheses), the model trained with $p = 0.5$ performs worse, suggesting that overly aggressive dropout may impair semantic conditioning for total mixture generation.

We further examine how modeling and training design choices, such as normalization layers, classifier-free guidance (CFG) dropout probability, and CFG scale, affect extraction performance and report the results in Table 7. When comparing normalization strategies, GroupNorm consistently matches or outperforms LayerNorm across most stems, demonstrating the advantage of modeling track-wise statistics in our multi-track architecture. This observation aligns with the trends seen in mixture generation results. For CFG dropout, a higher dropout probability ($p = 0.5$) leads to improved performance compared to the default $p = 0.1$, suggesting that stronger stochastic conditioning is beneficial during source extraction. While this differs from the trend observed in mixture generation (Table 6), the discrepancy may be explained by the fact that, in extraction, non-target tracks are effectively treated as unconditioned. This makes overall performance more sensitive to the model’s ability to generalize in the presence of dropout. We also evaluate various CFG scales (1, 2, 4, 8). A scale of 1 yields the best performance overall, although scales 2 and 4 remain competitive. Performance degrades at scale 8, indicating that overly strong guidance can impair extraction quality.

Table 8: **Preliminary results for iterative stem-wise generation.** Metrics are reported as FAD ↓. Evaluation is conducted using a model trained exclusively on S_A . Each column header indicates the generation order of stems (e.g., BDGP denotes bass→drums→guitar→piano).

S_A											
	Total Gen.	BDGP	BDPG	DBGP	DBPG	DGBP	DPGB	GPBD	GPDB	PDGB	PGBD
MGE	0.47	0.60	0.66	0.78	0.75	0.70	0.77	0.61	0.66	0.63	0.60

Table 9: **Partial generation results evaluated with COCOLA [105] score ↑.** Higher values indicate stronger coherence between the given mixture and the generated stems, complementing *sub*-FAD by emphasizing musical consistency. "Ground truth" values correspond to the COCOLA scores computed between the given mixture and the ground-truth accompaniment from the dataset. Bold values indicate the best results in each column (excluding Ground truth), and underlined values denote the best results among models trained on the same S_A set.

Model	Train Set				S_A														S_B							
	S_A	S_B	M_u	M_o	B	D	G	P	BD	BG	BP	DG	DP	GP	BDG	BDP	BGP	DGP	Brs.	C.P.	Org.	Pipe	Reed	Str.	S.Lead	S.Pad
MSDM	✓	×	×	×	57.74	51.87	57.76	56.10	45.96	57.64	54.48	50.67	49.40	55.75	43.25	40.46	52.82	47.09	-	-	-	-	-	-	-	-
MSG-LD	✓	×	×	×	58.18	50.58	<u>58.55</u>	58.30	46.12	57.63	56.29	49.11	47.51	57.94	43.70	41.33	54.19	44.98	-	-	-	-	-	-	-	-
T_1	✓	×	×	×	<u>60.93</u>	53.71	58.09	59.78	49.77	60.43	59.16	52.95	51.90	60.03	46.28	43.46	56.77	49.73	63.20	56.44	59.52	60.39	62.35	61.11	63.74	58.66
MGE (ours)	T_2	✓	✓	×	61.29	51.43	62.61	61.35	48.79	60.94	59.12	50.82	49.39	60.21	46.06	42.81	56.73	47.01	64.32	65.21	65.46	62.31	64.31	62.33	64.52	61.10
	T_3	×	×	✓	55.40	44.76	57.29	59.34	43.90	55.18	56.09	45.68	45.42	55.91	44.10	43.35	51.44	43.53	60.23	54.74	56.70	60.18	56.50	60.66	62.03	61.00
	T_4	✓	✓	✓	56.73	49.65	60.58	61.06	46.92	57.88	57.82	50.66	49.93	58.32	46.53	45.39	54.02	46.86	65.94	64.10	64.47	63.77	59.49	62.39	65.02	58.52
Ground truth					60.11	53.55	56.33	57.16	50.42	59.74	58.59	53.11	52.29	57.80	47.06	44.32	56.07	50.18	64.58	64.34	64.85	63.38	64.75	63.20	65.29	62.16

E.3 Iterative Generation Variants

Table 8 presents preliminary results for the iterative generation procedure described in Appendix C, applied to a model trained on S_A . The task involves sequentially generating the four canonical stems (bass, drums, guitar, and piano) in various orders.

Across all tested permutations, iterative generation produced higher FAD scores compared to one-stage mixture generation, indicating a degradation in perceptual quality. Nonetheless, iterative generation may offer utility in settings that require fine-grained, source-specific control.

An interesting trend observed: generation sequences that began with drums consistently resulted in poorer performance relative to other orderings. This suggests that the model may be more effective at first establishing harmonic or melodic content before aligning rhythmic elements. While this observation is speculative, it highlights a potential inductive bias in the model that warrants further investigation, particularly in scenarios beyond the four-instrument configuration.

F Extended Evaluation with Alternative Metrics

To complement the quantitative results presented in Section 5, we provide additional evaluations using alternative metrics that better capture musical coherence and perceptual quality. While *sub*-FAD is widely adopted in accompaniment generation research [34–36, 63, 76, 81], it is considered suboptimal for evaluating musical coherence, since it primarily measures global timbral similarity rather than the harmonic and rhythmic interplay between stems [105]. To address this limitation, Ciranni et al. [105] introduced the COCOLA metric, which quantifies harmonic and rhythmic coherence through contrastive learning of musical audio representations.

We report the COCOLA results in Table 9 to complement our main findings. The metric is computed between the given mixture and the generated stems, and we additionally report the "Ground truth" values computed between the same mixture and the corresponding reference accompaniment from the dataset. Among models trained on S_A , our T_1 variant outperforms the baselines on most instrument categories, except for guitar generation. Larger-scale models (T_2 – T_4) show similarly strong or even superior coherence, in several cases approaching or surpassing ground-truth levels. Since COCOLA emphasizes musical coherence rather than timbral accuracy with respect to reference stems, it provides a meaningful complement to *sub*-FAD, offering a more balanced view of partial generation quality.

Because the VGGish-based FAD has been shown to correlate poorly with human perception in recent studies [106–108], we further report $FAD_{CLAP-MA}$ and $sub-FAD_{CLAP-MA}$, computed using CLAP embeddings from the music_audioset_epoch_15_esc_90.14.pt checkpoint, which achieves the highest correlation with subjective ratings according to Grötschla et al. [106]. Table 10 presents total generation results evaluated with $FAD_{CLAP-MA}$, where T_1 achieves the best overall performance among models trained on S_A . For the combined $S_A + S_B$ reference set, T_2 yields the lowest FAD scores, while T_3 performs best on M_u and M_o , aligning with

Table 10: Total generation results evaluated with $\text{FAD}_{\text{CLAP-MA}} \downarrow$.

Model	Train Set				Test Set			
	$\mathbf{S}_A$	$\mathbf{S}_B$	$\mathbf{M}_u$	$\mathbf{M}_o$	$\mathbf{S}_A$	$\mathbf{S}_{\text{Full}}$	$\mathbf{M}_u$	$\mathbf{M}_o$
MSDM	✓	×	×	×	.329	.290	.767	.744
MSG-LD	✓	×	×	×	.191	.312	.658	.697
	$\overline{\mathcal{T}}_1$	✓	×	×	.110 (.450)	.285	.637	.627
MGE (ours)	$\overline{\mathcal{T}}_2$	✓	✓	×	.368 (.359)	.099	.578	.621
	$\overline{\mathcal{T}}_3$	×	×	✓	.731 (.669)	.632	.222	.165
	$\overline{\mathcal{T}}_4$	✓	✓	✓	.563 (.712)	.426	.235	.186

Table 11: Partial generation results evaluated with $\text{sub-FAD}_{\text{CLAP-MA}} \downarrow$.

Model	Train Set				S_A														S_B							
	S_A	S_B	M_u	M_o	B	D	G	P	BD	BG	BP	DG	DP	GP	BDG	BDP	BGP	DGP	Brs.	C.P.	Org.	Pipe	Reed	Str.	S.Lead	S.Pad
MSDM	✓	×	×	×	.103	.100	.066	.096	.227	.189	.252	.135	.191	.302	.275	.357	.555	.403	-	-	-	-	-	-	-	-
MSG-LD	✓	×	×	×	.071	.070	.076	.077	.130	.135	.143	.117	.121	.157	.185	.188	.241	.190	-	-	-	-	-	-	-	-
$\overline{T}_1$	✓	×	×	×	.140	.170	.153	.154	.184	.175	.181	.200	.189	.184	.221	.218	.190	.214	.278	.086	.069	.327	.462	.201	.335	.118
MGE (ours)	$\overline{T}_2$	✓	✓	×	.259	.329	.251	.264	.420	.302	.362	.380	.361	.337	.500	.517	.422	.460	.425	.071	.129	.327	.330	.296	.123	.313
	$\overline{T}_3$	×	×	✓	.178	.202	.356	.253	.307	.402	.340	.368	.320	.477	.445	.469	.632	.484	.361	.087	.062	.045	.090	.087	.079	.060
	$\overline{T}_4$	✓	✓	✓	.170	.187	.367	.251	.253	.405	.308	.350	.289	.455	.450	.436	.571	.486	.312	.072	.107	.057	.340	.208	.085	.063
	$\overline{T}_5$	✓	✓	✓	.170	.187	.367	.251	.253	.405	.308	.350	.289	.455	.450	.436	.571	.486	.312	.072	.107	.057	.340	.208	.085	.063

their respective training domains. The $\text{sub-FAD}_{\text{CLAP-MA}}$ results for partial-generation, presented in Table 11, exhibit trends consistent with Table 2, with MSG-LD achieving the best scores on most stem combinations. Together, these additional metrics provide a more comprehensive and reliable assessment of both perceptual quality and cross-domain generalization in partial and total music generation.

G Future Work

Future extensions of MGE-LDM include scaling to higher-resolution formats such as 44.1 kHz stereo audio, enabling richer timbral detail and spatial fidelity. In particular, this can be achieved by leveraging high-quality latent representations recently developed for the music domain [109, 110]. To reduce the modality gap in text-conditioned extraction, fine-tuning on curated audio-text datasets like MusicCaps [7] is a promising direction. Given its minimal reliance on precise stem boundaries, MGE-LDM is naturally suited for incorporating weakly or noisily labeled multi-track data [39, 111], which may expand training diversity.

Another promising avenue is to pre-train MGE-LDM on large-scale mixture-only corpora such as MTG-Jamendo [112] or the Free Music Archive [113] to learn general audio priors for mixture tracks, followed by fine-tuning on multi-track datasets for source-aware generation. This two-stage training strategy is expected to enhance generative quality and improve generalization.

We also plan to extend MGE-LDM to text-based music editing tasks, drawing inspiration from recent instruction-guided frameworks such as AUDIT [73], InstructME [74], and Instruct-MusicGen [64]. Leveraging MGE-LDM’s latent inpainting capabilities and language-conditioned generation, this extension could enable user-directed operations such as instrument replacement and style transformation via natural language prompts, building upon the model’s unified training scheme and class-agnostic design.

H Spectrogram Examples of Generated Samples

We present Mel-spectrogram visualizations of generated audio samples across the three primary tasks: total generation, partial generation (imputation), and source extraction. All examples are produced by MGE-LDM trained on the combined Slakh2100 ($\mathbf{S}_A + \mathbf{S}_B$), MUSDB18 ($\mathbf{M}_u$), and MoisesDB ($\mathbf{M}_o$) datasets.

We note that the model is capable of generating vocals in the unconditional setting, as vocals stems are present in the training data. Although MGE-LDM does not currently support fine-grained control over vocal generation, this points to a promising direction for future work, such as incorporating explicit vocal prompts or segment-level control for more expressive and structured multi-track modeling.

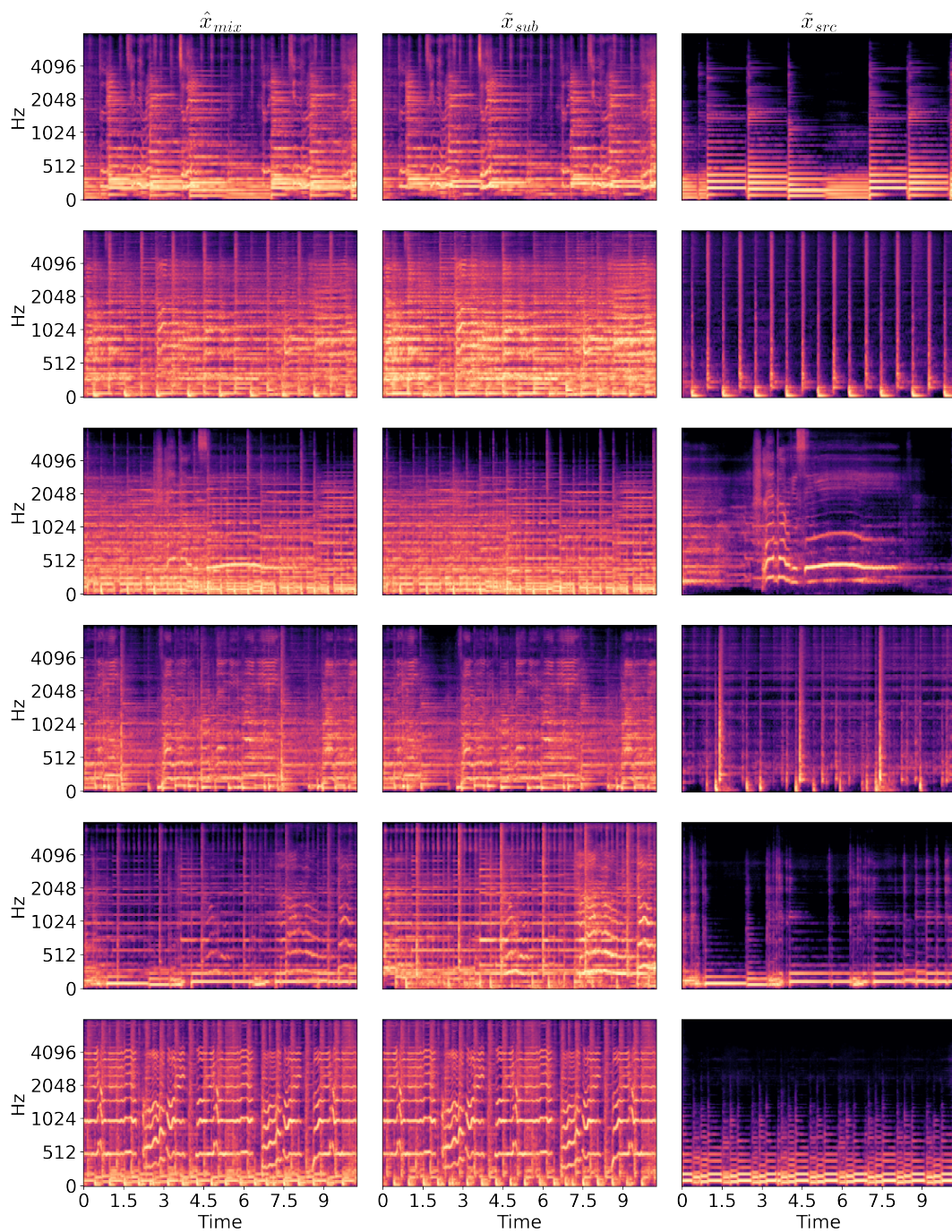


Figure 2: Total generation examples. Each sample displays Mel-spectrograms of the mixture, submixture, and source tracks, all generated simultaneously by MGE-LDM. The mixture track is used to evaluate the total generation output.

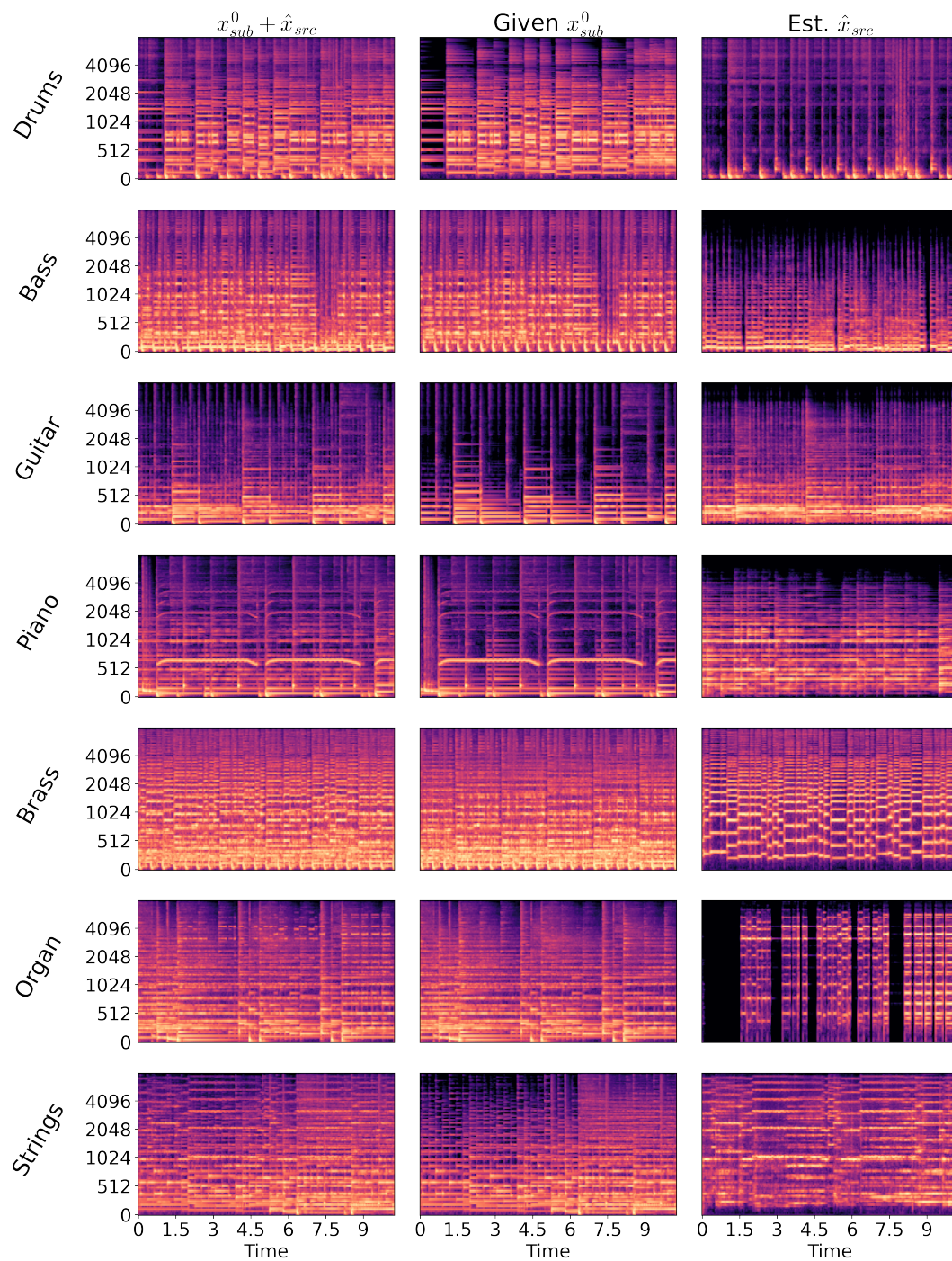


Figure 3: **Source imputation examples.** Each row illustrates source inpainting results by MGE-LDM, conditioned on the text prompt "*The sound of the {label}*". The middle column shows the provided context mixture (submix), the rightmost column is the generated source, and the leftmost column is the recombined mixture of the submix and generated source. While some stems are imputed accurately, others fail due to data imbalance during training.

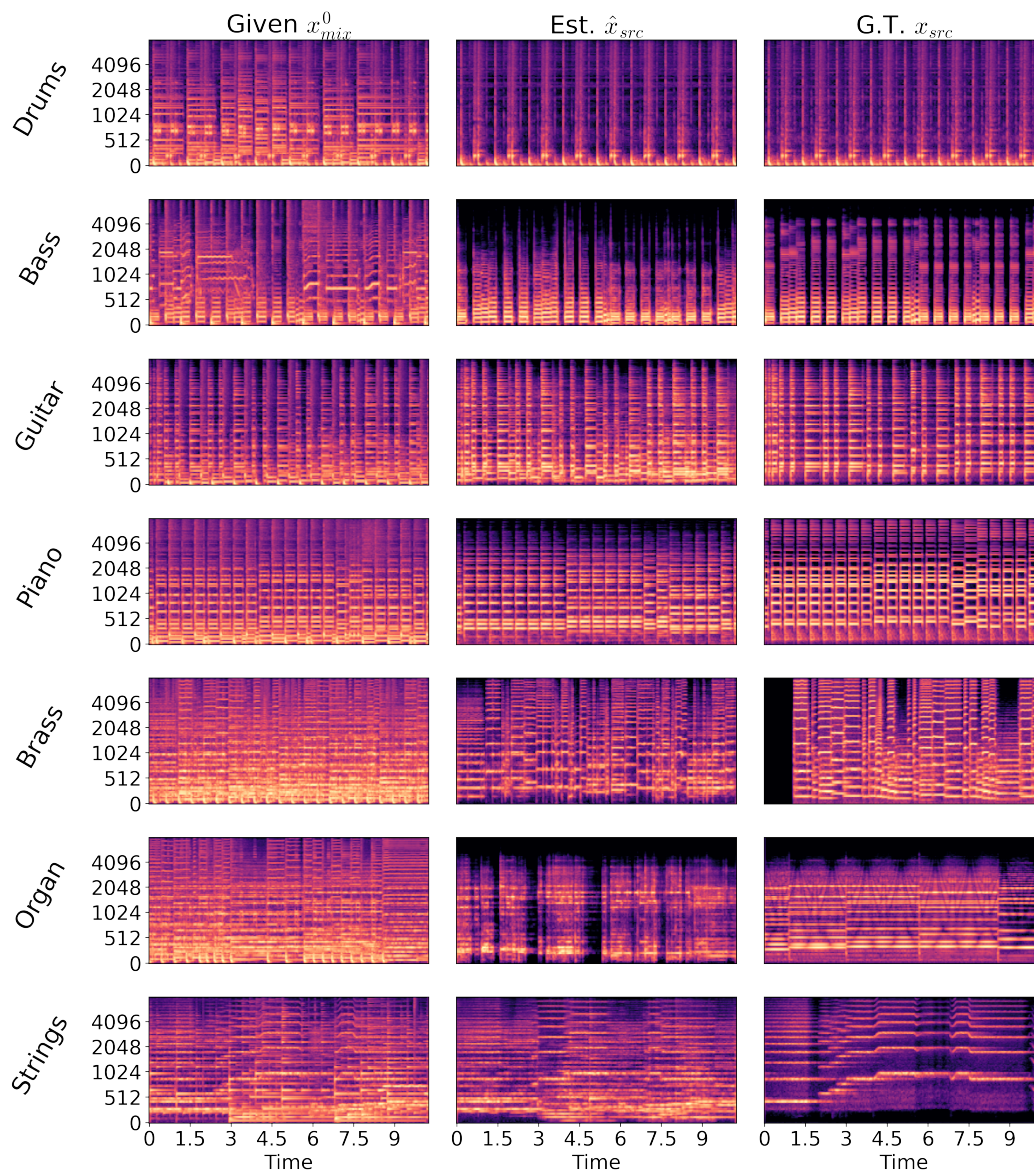


Figure 4: **Source extraction examples.** Source extraction results produced by MGE-LDM, conditioned on the text query "*The sound of the {label}*". The leftmost column shows the input mixture, the middle column is the extracted source predicted by the model, and the rightmost column is the ground-truth source. We observe that extraction quality may degrade for underrepresented classes such as strings, and in some cases, the model hallucinates unrelated instruments or incorrect timbres.

I Ethics Statement

This work introduces a class-agnostic generative framework for multi-track music modeling, trained exclusively on publicly available datasets (Slakh2100, MUSDB18, and MoisesDB). While the model enables flexible music generation, source imputation, and source extraction, it also carries potential risks, such as unauthorized manipulation, misuse in derivative content, or generation of audio resembling copyrighted material. To mitigate these concerns, we commit to releasing the model and code under a license with clear usage guidelines, emphasizing responsible research and ethical creative applications.