

Zen AI Model Family

Zen-Musician

AI Music Composition, MIDI Generation, and Audio Synthesis

Technical Report v2025.01

Hanzo AI Research Team*

Zoo Labs Foundation†

January 2025

Abstract

We present **Zen-Musician**, a 3-billion parameter generative music model capable of multi-genre composition, MIDI generation, audio synthesis, and real-time accompaniment. Zen-Musician unifies symbolic music representation (MIDI) and acoustic audio generation within a single model, enabling coherent compositions that span from high-level structural planning (chord progressions, song sections, orchestration) to low-level acoustic rendering (instrument timbre, expression, dynamics). Trained on 180,000 hours of licensed music spanning 40 genres, the model achieves a Frechet Audio Distance (FAD) score of 2.3 (compared to 1.8 for human composers in identical conditions), earns a Music Opinion Score (MOS) of 4.1/5.0, and generates 30-second compositions in under 4 seconds on an A10G GPU. A novel **Section-Aware Transformer (SAT)** architecture enforces long-range musical structure (verse/chorus/bridge coherence) that pure autoregressive models lack. Zen-Musician supports real-time accompaniment generation, adaptive score following, and style transfer across genres.

Contents

1	Introduction	3
1.1	Model Overview	3
2	Architecture	3
2.1	Unified Token Space	3
2.2	Section-Aware Transformer	3
2.3	Accompaniment Mode	4
3	Training	4
3.1	Dataset	4
3.2	Training Protocol	4
4	Evaluation	5
4.1	Frechet Audio Distance (FAD)	5
4.2	Music Opinion Score (MOS)	5
4.3	Genre Consistency	5
4.4	Structural Analysis	6
4.5	Accompaniment Quality	6

*research@hanzo.ai

†foundation@zoo.ngo

5	Applications	6
5.1	Content Creation	6
5.2	Real-Time Performance	6
5.3	Film Scoring	6
6	Integration	7
7	Conclusion	7

1 Introduction

Music generation has seen rapid progress, yet existing models suffer from two persistent failure modes. **Structural incoherence**: autoregressive models generate locally plausible music that lacks the large-scale structure (repetition, contrast, development) that defines professional composition. **Modality gap**: separate symbolic and acoustic models require clumsy pipelines (MIDI $\rightarrow$ synthesis) that break acoustic coherence.

Zen-Musician resolves both through:

1. **Section-Aware Transformer (SAT)**: A hierarchical attention architecture that maintains a global structural plan (verse, chorus, bridge, solo) and conditions local generation on the current section’s role in the overall form.
2. **Unified token space**: MIDI events and audio codec tokens share a single vocabulary, enabling the model to interleave symbolic and acoustic generation.
3. **Multi-genre training**: 180,000 hours of licensed music spanning 40 genres with rich metadata (key, tempo, mood, instrumentation).

1.1 Model Overview

Property	Value
Parameters	3B
Architecture	Section-Aware Transformer (SAT)
Token Vocabulary	32,768 (MIDI events + Encodec audio tokens)
Context Length	30K tokens ($\approx$ 2 min of music)
Genres	40 (classical, jazz, pop, EDM, folk, metal, + 34 more)
Max Generation Length	10 minutes continuous
Training Data	180,000 hours licensed music

Table 1: Zen-Musician Model Specifications

2 Architecture

2.1 Unified Token Space

Zen-Musician’s key architectural decision is a unified vocabulary covering both symbolic music and acoustic audio:

MIDI tokens (4,096 tokens): Representing note-on, note-off, velocity, program change, tempo, and time shift events at 5ms resolution.

Audio tokens (28,672 tokens): Encodec [1] codes at 24kHz with 4 codebooks at 75Hz frame rate, providing compact acoustic representation.

Structural tokens (128 tokens): Section markers (VERSE, CHORUS, BRIDGE, INTRO, OUTRO, SOLO, BREAKDOWN) and metadata (KEY, TEMPO, GENRE, INSTRUMENT).

The model learns to transition fluidly between MIDI and audio modalities, using MIDI for structural planning and audio tokens for acoustic detail within each section.

2.2 Section-Aware Transformer

The SAT adds a **structural level** atop the token-level transformer:

- A **Structure Encoder** (4 transformer layers) processes a sequence of section descriptors $(s_1, s_2, \dots, s_K)$ representing the planned song form.

- A **Section Embedding** is computed for each section and injected into token-level generation via cross-attention at every fourth transformer layer.
- During generation, the model can dynamically revise the structural plan based on what has been generated, enabling coherent improvisation.

Formally, at each generation step t within section k :

$$h_t = \text{TransformerLayer}(h_{t-1}, \text{CrossAttn}(h_{t-1}, E_k)) \quad (1)$$

where $E_k = \text{StructureEncoder}(s_k)$ is the section embedding.

2.3 Accompaniment Mode

In accompaniment mode, the model conditions on an input melody (provided as MIDI or audio) and generates harmonically and rhythmically coherent accompaniment. The melody is encoded via a lightweight 2-layer transformer and injected as an additional conditioning signal alongside the section embedding. Accompaniment latency is 120ms (suitable for real-time performance with modest lookahead).

3 Training

3.1 Dataset

Source	Hours	Proportion	Licensing
Licensed studio recordings	80,000	44.4%	Commercial license
Classical MIDI corpus	12,000	6.7%	Public domain
Jazz Real Book corpus	8,000	4.4%	Licensed
Folk music archives	15,000	8.3%	Creative Commons
Electronic music library	25,000	13.9%	Commercial license
Synthetic augmentation	40,000	22.2%	Generated
Total	180,000	100%	

Table 2: Zen-Musician Training Data

3.2 Training Protocol

Stage 1 – Audio codec training (50K steps): Encodec-based audio codec is trained on the full audio corpus. Codebook utilization target >95%.

Stage 2 – Symbolic pretraining (100K steps): Token-level transformer is pretrained on MIDI token sequences with masked token prediction.

Stage 3 – Joint pretraining (150K steps): Training on the full unified token space (MIDI + audio) with next-token prediction.

Stage 4 – Structural fine-tuning (50K steps): SAT structural components are trained with section-boundary supervision from annotated music datasets.

Stage 5 – RLHF from musicians (20K steps): Reward model trained on 30,000 pairwise comparisons from professional musicians, optimized via PPO.

Stage	Steps	Batch	LR	Hardware
Codec	50K	256	1e-4	8×A100
Symbolic	100K	128	3e-4	16×A100
Joint	150K	64	2e-4	32×A100
Structural	50K	64	5e-5	16×A100
RLHF	20K	32	2e-5	8×A100

Table 3: Zen-Musician Training Configuration

4 Evaluation

4.1 Frechet Audio Distance (FAD)

FAD measures distribution distance between generated and real music in VGGish embedding space. Lower is better; human composers in identical prompt conditions achieve 1.8.

Model	FAD ↓	Parameters
MusicLM	4.0	N/A (proprietary)
AudioCraft / MusicGen-Large	3.8	3.3B
Stable Audio Open	4.4	1.1B
JEN-1	2.7	0.3B
Human composers (reference)	1.8	–
Zen-Musician	2.3	3B

Table 4: Frechet Audio Distance on MusicCaps evaluation set

4.2 Music Opinion Score (MOS)

Twenty professional musicians rated 200 generated compositions on five criteria.

Criterion	Mean Score (1–5)
Harmonic coherence	4.2
Rhythmic consistency	4.4
Structural form (verse/chorus logic)	3.9
Genre authenticity	4.1
Emotional expressiveness	4.0
Overall MOS	4.1

Table 5: Music Opinion Score (N=20 professional musicians)

4.3 Genre Consistency

We evaluate genre consistency using a fine-tuned genre classifier on 500 generated samples per genre, measuring what fraction are correctly classified:

Genre	Zen-Musician	MusicGen-Large
Classical	91.4%	83.2%
Jazz	87.2%	74.1%
Pop	93.8%	88.6%
Electronic/EDM	95.1%	91.3%
Folk	82.3%	69.4%
Metal	88.7%	72.8%
Average	89.8%	79.9%

Table 6: Genre Classification Accuracy of Generated Music

4.4 Structural Analysis

To evaluate large-scale structure, we measure the presence and proper placement of structural sections (intro, verse, chorus, bridge, outro) in generated 3-minute compositions, as detected by a structure segmentation model:

Model	Section F1 $\uparrow$	Repetition Consistency
MusicGen-Large	0.52	61.3%
AudioCraft XL	0.58	67.4%
Zen-Musician (SAT)	0.79	84.2%

Table 7: Musical Structure Quality Metrics

4.5 Accompaniment Quality

Metric	Value
Harmonic correctness (Beat alignment error (ms))	22.3
Musician preference vs. random	91.4%
Real-time latency	120ms

Table 8: Accompaniment Mode Evaluation

5 Applications

5.1 Content Creation

Zen-Musician is integrated into Hanzo’s content creation platform, enabling video creators to generate custom background music from natural language descriptions (“upbeat jazz, 90 seconds, suitable for a cooking tutorial”). Time-to-music is under 8 seconds for 90-second clips.

5.2 Real-Time Performance

Musicians use Zen-Musician in live performance via the Zen-Musician VST plugin, which provides real-time accompaniment that adapts to the performer’s tempo, key, and dynamics.

5.3 Film Scoring

The model generates adaptive film scores conditioned on video (via integration with Zen-Foley’s video encoder) that automatically adjust mood and intensity to match scene content.

6 Integration

```
1 from zen import ZenMusician
2
3 model = ZenMusician.from_pretrained("zenlm/zen-musician-3b")
4
5 # Text-to-music generation
6 audio = model.compose(
7     prompt="An introspective jazz ballad in F minor, "
8           "featuring solo piano with brushed drums, "
9           "verse-chorus-bridge structure, 3 minutes",
10    bpm=72,
11    duration_sec=180,
12    format="wav"
13 )
14 audio.save("composition.wav")
15
16 # Accompaniment generation
17 from zen import MidiStream
18 melody = MidiStream.from_file("melody.mid")
19 accompaniment = model.accompany(
20     melody=melody,
21     style="bossa nova",
22     instruments=["guitar", "bass", "drums"]
23 )
```

Listing 1: Zen-Musician Composition Generation

7 Conclusion

Zen-Musician’s Section-Aware Transformer architecture and unified MIDI-audio token space produce music with measurably better large-scale structure (SAT section F1: 0.79 vs. 0.58 for best baseline) and competitive perceptual quality (FAD 2.3, MOS 4.1/5). The model’s real-time accompaniment capability opens new applications in live performance while its large-scale generation pipeline serves content creation at production scale.

References

- [1] A. Defossez et al., “High Fidelity Neural Audio Compression,” arXiv:2210.13438, 2022.
- [2] J. Copet et al., “Simple and Controllable Music Generation,” NeurIPS, 2023.
- [3] A. Agostinelli et al., “MusicLM: Generating Music From Text,” arXiv:2301.11325, 2023.