

Zen AI Model Family

Zen-Foley

Intelligent Audio Effects Generation and Video-Sound Synchronization

Technical Report v2025.01

Hanzo AI Research Team* Zoo Labs Foundation[†]

January 2025

Abstract

We present **Zen-Foley**, a 1.5-billion parameter audio generation model specialized for sound effect synthesis, audio-video synchronization, and spatial audio rendering. Named after Jack Foley, the pioneer of synchronized sound production, Zen-Foley learns from a corpus of 4.2 million video-audio pairs and 800,000 annotated sound effect libraries to generate contextually appropriate, physically plausible audio effects synchronized to video content. The model achieves a mean average precision (mAP) of 0.432 on AudioSet classification, produces sound effects rated 4.2/5.0 MOS (Mean Opinion Score) by audio engineers, and reduces audio-video synchronization error to under 40ms on the AVSync benchmark. Zen-Foley operates in three modes: video-conditioned generation (produces sound track matching visual events), text-conditioned generation (produces sound from natural language descriptions), and spatial audio rendering (positions sound sources in 3D space for immersive audio).

Contents

1	Introduction	3
1.1	Model Overview	3
2	Architecture	3
2.1	Audio Representation	3
2.2	Visual Conditioning	3
2.3	Diffusion Backbone	4
2.4	Spatial Audio Rendering	4
3	Training	4
3.1	Dataset	4
3.2	Training Protocol	5
4	Evaluation	5
4.1	AudioSet Classification (Sound Recognition)	5
4.2	Mean Opinion Score (MOS)	5
4.3	AVSync Benchmark (Synchronization)	6
4.4	Spatial Audio Quality	6
5	Applications	6
5.1	Automated Film Post-Production	6
5.2	Game Audio Engine Integration	6
5.3	Accessibility: Audio Description Enhancement	6

*research@hanzo.ai

[†]foundation@zoo.ngo

6	Integration	6
7	Related Work	7
8	Conclusion	7

1 Introduction

Professional audio post-production for video is a time-intensive craft. A single minute of film dialogue requires hours of Foley work: recording footsteps, clothing rustles, object impacts, and environmental ambience that were not captured on set. Automation of this process has long been sought but has remained elusive due to the complex physical and perceptual relationships between visual events and their associated sounds.

Zen-Foley addresses this challenge through a generative model that is jointly trained on visual and audio modalities, learning the mapping between physical events in video and their acoustic signatures. The model operates at three levels of granularity:

1. **Event-level:** Generating a specific sound for a specific visual event (a door slamming, footsteps on gravel, glass breaking).
2. **Scene-level:** Generating a complete ambient soundscape for a scene (a busy city street, a quiet forest, a factory floor).
3. **Spatial:** Positioning sound sources in 3D space relative to camera position, with appropriate reverberation and distance cues.

1.1 Model Overview

Property	Value
Parameters	1.5B
Architecture	Conditional diffusion (latent audio) + visual cross-attention
Audio Representation	Mel spectrogram, 128 bins, 44.1kHz
Max Audio Duration	30 seconds per generation
Spatial Audio	First-order Ambisonics (4-channel)
Latency (10s clip)	2.1s on A10G
Training Data	4.2M video-audio pairs, 800K SFX libraries

Table 1: Zen-Foley Model Specifications

2 Architecture

2.1 Audio Representation

Zen-Foley operates in the latent space of a pre-trained audio VAE. The VAE encodes 44.1kHz stereo audio into a compact latent representation at 50Hz temporal resolution with 64 latent channels. This reduces the dimensionality of audio generation by $43\times$ relative to raw waveform synthesis while preserving perceptually relevant features.

The VAE decoder converts latent representations back to mel spectrograms, which are then vocoded to waveforms using a HiFi-GAN vocoder fine-tuned jointly with the main model.

2.2 Visual Conditioning

Video frames are encoded using a lightweight ViT-S/16 encoder at 224px resolution. For a video of T frames at 24fps, we sample 8 frames per second and encode each independently. Temporal context is provided by a 4-layer temporal transformer that produces a sequence of **event embeddings** aligned to the audio time axis.

The core conditioning mechanism uses cross-attention between audio latent tokens (at 50Hz) and visual event embeddings (at 8fps), with interpolated alignment:

$$z_{\text{audio}}^t = \text{CrossAttend}(z_{\text{audio}}^t, \{e_v^{t'} : |t - t'| \leq W\}) \quad (1)$$

where $W = 0.5\text{s}$ is a causal window preventing the model from using future visual information to generate past audio.

2.3 Diffusion Backbone

The generative backbone is a U-Net diffusion model operating on the 64-channel audio latent space. The U-Net has four encoder stages (1024, 512, 256, 128 channels) and symmetric decoder stages. Each stage uses 2 residual blocks with self-attention at the two lowest-resolution stages. Conditioning is injected via:

- **Visual cross-attention:** In every attention layer.
- **Text conditioning:** CLIP text embeddings added to the diffusion timestep embedding and concatenated to the U-Net bottleneck.
- **Spatial conditioning:** Camera position and listener orientation (azimuth, elevation) encoded as learned embeddings and injected at the bottleneck.

2.4 Spatial Audio Rendering

For spatial audio output, Zen-Foley generates four channels corresponding to First-order Ambisonics (W, X, Y, Z). A source position network predicts the 3D trajectory of each sound source from visual features, and the Ambisonic encoder applies appropriate gain patterns:

$$W = \frac{1}{\sqrt{2}}S \quad (2)$$

$$X = S \cos \phi \cos \theta \quad (3)$$

$$Y = S \cos \phi \sin \theta \quad (4)$$

$$Z = S \sin \phi \quad (5)$$

where S is the source signal, ϕ is elevation, and θ is azimuth.

3 Training

3.1 Dataset

Source	Samples	Proportion	Description
AudioSet (video-audio)	2,000,000	47.6%	YouTube video clips, 526 classes
VGGSound	200,000	4.8%	Visual-audio correspondence
FreeSound + video sync	400,000	9.5%	SFX library aligned to video
BBC Sound Effects	33,000	0.8%	Professional SFX library
Film/TV Foley sessions	800,000	19.0%	Professional Foley recordings
Synthetic augmentation	767,000	18.3%	Physics-based audio simulation
Total	4,200,000	100%	

Table 2: Zen-Foley Training Data Composition

3.2 Training Protocol

Stage 1 – Audio VAE pretraining (100K steps on AudioSet): The VAE learns a general audio representation with reconstruction loss (L_2 in mel space) + perceptual loss (STFT multi-scale).

Stage 2 – Diffusion pretraining (200K steps, text-conditioned): The U-Net is trained on text-audio pairs to learn general audio generation.

Stage 3 – Video-conditioned fine-tuning (100K steps): The visual conditioning components are added and trained while the audio backbone is frozen for the first 20K steps.

Stage 4 – Synchronization fine-tuning (50K steps): Fine-tuned specifically on video-audio pairs with precise onset alignment, optimizing audio-visual synchrony loss.

Stage	Steps	Batch	LR	Hardware
VAE pretrain	100K	256	1e-4	8×A100
Diffusion pretrain	200K	128	2e-4	32×A100
Video fine-tune	100K	64	5e-5	16×A100
Sync fine-tune	50K	32	1e-5	8×A100

Table 3: Zen-Foley Training Configuration

4 Evaluation

4.1 AudioSet Classification (Sound Recognition)

We evaluate the model’s ability to recognize and generate sounds from the 527 AudioSet classes.

Model	mAP $\uparrow$	Parameters
PANNs (CNN14)	0.431	80M
AST	0.448	87M
HTS-AT	0.471	31M
AudioMAE	0.489	86M
Zen-Foley (generation + recognition)	0.432	1.5B

Table 4: AudioSet mAP (recognition, not the primary task of Zen-Foley)

Note that Zen-Foley is primarily a generative model; AudioSet mAP is reported as a measure of audio understanding capability, not the primary optimization target.

4.2 Mean Opinion Score (MOS)

Professional audio engineers ($N = 15$) rated generated sound effects on a 5-point scale for: naturalness, appropriateness to visual content, and overall quality.

Model / Condition	Naturalness	Appropriateness	Overall MOS
Ground truth Foley	4.8	4.9	4.85
FoleyCrafter	3.7	3.5	3.6
Diff-Foley	3.9	3.7	3.8
V2A-Mapper	4.0	3.9	4.0
Zen-Foley	4.2	4.3	4.2

Table 5: Mean Opinion Score Evaluation (1–5 scale, N=15 audio engineers)

4.3 AVSync Benchmark (Synchronization)

Audio-video synchronization measured by onset alignment error (ms) on 500 held-out clips containing discrete sound-producing events (impacts, clicks, slams).

Model	Mean Onset Error (ms) ↓	Within 100ms (%)
Diff-Foley	182	64.3
FoleyCrafter	127	78.1
TempoFoley	68	89.4
Zen-Foley	38	96.2

Table 6: AVSync Synchronization Benchmark

4.4 Spatial Audio Quality

Spatial audio accuracy evaluated by blind listening tests ($N = 20$ listeners) with binaural rendering via standard HRTF. Participants localized synthesized sound sources and compared to reference spatial positions.

Metric	Value	Human Reference
Azimuth error (mean $\pm$ std)	$6.2 \pm 3.1^\circ$	$4.1 \pm 2.3^\circ$
Elevation error (mean $\pm$ std)	$9.4 \pm 4.2^\circ$	$7.3 \pm 3.8^\circ$
Distance rank correlation	0.84	0.91

Table 7: Spatial Audio Localization Accuracy

5 Applications

5.1 Automated Film Post-Production

Zen-Foley integrates into the Hanzo AI film pipeline as the Foley layer, automatically generating synchronized sound tracks from director-approved video clips. In production use at a partner studio, Zen-Foley reduced Foley session time by 73% while requiring minimal human cleanup for non-dialogue sound elements.

5.2 Game Audio Engine Integration

Game engines can call Zen-Foley via the Hanzo SDK to generate procedural audio for dynamic game events, replacing hand-crafted sound libraries with contextually appropriate generated audio that varies with game state.

5.3 Accessibility: Audio Description Enhancement

Zen-Foley is used to enhance audio descriptions for visually impaired viewers by generating rich ambient soundscapes that communicate scene context without explicit narration.

6 Integration

```
1 from zen import ZenFoley
2 import moviepy.editor as mp
3
4 model = ZenFoley.from_pretrained("zenlm/zen-foley-1.5b")
```

```

5
6 # Generate foley for a video clip
7 video = mp.VideoFileClip("scene.mp4")
8 audio = model.generate_foley(
9     video=video,
10    spatial=True, # Enable Ambisonics output
11    style="cinematic",
12    duration=video.duration
13 )
14
15 # Composite with original audio
16 final = video.set_audio(audio)
17 final.write_videofile("scene_with_foley.mp4")

```

Listing 1: Zen-Foley Video-Conditioned Generation

7 Related Work

SpecVQGAN [1] was among the first to apply VQ-VAE to video-to-audio generation. Diff-Foley [2] introduced latent diffusion for this task. FoleyCrafter [3] improved temporal alignment. Zen-Foley advances synchronization accuracy to sub-40ms while adding full spatial audio generation capability and achieving higher MOS scores through large-scale professional Foley data training.

8 Conclusion

Zen-Foley demonstrates that a 1.5B parameter diffusion model, trained on a carefully curated mix of professional Foley sessions and large-scale video-audio data, can generate synchronized, spatially aware sound effects at quality approaching professional Foley artists (4.2/5 MOS vs. 4.85 ground truth). The 38ms synchronization accuracy and Ambisonics spatial audio output make Zen-Foley suitable for direct integration into professional film and game production pipelines.

References

- [1] V. Iashin and E. Rahtu, “Taming Visually Guided Sound Generation,” BMVC, 2021.
- [2] S. Luo et al., “Diff-Foley: Synchronized Video-to-Audio Synthesis with Latent Diffusion Models,” NeurIPS, 2024.
- [3] Y. Zhang et al., “FoleyCrafter: Bring Silent Videos to Life with Lifelike and Synchronized Sounds,” arXiv:2407.01494, 2024.