
Zen-Voice:

Zero-Shot Voice Cloning and Expressive Speech Synthesis

Hanzo AI Research¹ Zoo Labs Foundation²

¹Hanzo AI Inc. (Techstars '17) ²Zoo Labs Foundation (501(c)(3))

research@hanzo.ai foundation@zoo.ngo

<https://hanzo.ai/research/zen-voice>

February 2026

Abstract

We present **Zen-Voice**, a neural speech synthesis system capable of zero-shot voice cloning from as little as 3 seconds of reference audio. Zen-Voice produces natural, expressive speech that faithfully reproduces the timbre, accent, speaking rate, and emotional characteristics of the reference speaker while synthesizing arbitrary text content. The system is built on three core innovations: (1) a **Hierarchical Speaker Encoder (HSE)** that disentangles speaker identity from prosodic style through multi-scale contrastive learning on 680,000 hours of multilingual speech, (2) a **Prosody Transfer Module (PTM)** based on a flow-matching architecture that models the joint distribution of pitch, energy, and duration conditioned on both text and speaker embedding, and (3) a **Neural Codec Vocoder (NCV)** that synthesizes waveforms at 24kHz from discrete codec tokens with a lightweight streaming architecture suitable for real-time applications. Zen-Voice achieves a speaker similarity MOS of 4.21 (5-point scale) on zero-shot cloning with 3-second references, improving to 4.52 with 10-second references. On the LibriTTS test-clean benchmark, it achieves a naturalness MOS of 4.38, surpassing both VALL-E (3.84) and VoiceBox (4.12). On the VCTK multi-speaker benchmark, Zen-Voice achieves a speaker verification Equal Error Rate (EER) of 2.1% for cloned speech, approaching the 1.8% EER of ground-truth recordings. We additionally introduce an **anti-deepfake watermarking** system that embeds imperceptible, cryptographically signed provenance markers into all Zen-Voice output, enabling downstream detection of AI-generated speech with 99.7% accuracy even after MP3 compression and noise addition. Models and inference code are released under Apache 2.0.

Keywords: Voice Cloning, Text-to-Speech, Speaker Embedding, Prosody Transfer, Neural Codec, Anti-Deepfake Watermarking

1 Introduction

Human speech conveys far more than linguistic content. A speaker’s voice carries identity (timbre, accent, vocal register), emotion (joy, sadness, anger, surprise), and communicative intent (emphasis, irony, urgency) through subtle variations in pitch, timing, and spectral characteristics. Reproducing this richness in synthetic speech—particularly when cloning a voice from a brief audio sample—remains one of the most challenging problems in generative AI.

Recent advances in neural text-to-speech (TTS) have dramatically improved synthesis quality. Autoregressive models such as VALL-E (Wang et al., 2023) treat TTS as a language modeling problem over discrete audio tokens, achieving impressive zero-shot cloning. Non-autoregressive approaches like VoiceBox (Le et al., 2024) and NaturalSpeech 3 (Ju et al., 2024) use flow-matching and diffusion to generate speech in parallel, offering faster inference. However, existing systems still struggle with three key challenges: (i) faithfully reproducing speaker identity from very short references (<5 seconds), (ii) transferring fine-grained prosodic patterns (emphasis, pacing, emotional coloring) independently of speaker identity, and (iii) generating speech in real-time for interactive applications.

Zen-Voice addresses these challenges through a modular architecture that cleanly separates speaker identity, prosodic style, and linguistic content. The Hierarchical Speaker Encoder captures speaker characteristics at multiple temporal scales—from sub-phonemic spectral details to utterance-level speaking style—enabling robust identity extraction even from very short references. The Prosody Transfer Module models prosody as a conditional flow that can be guided by explicit emotion labels, reference audio, or natural language descriptions (“speak with quiet intensity”). The Neural Codec Vocoder converts the model’s output into high-fidelity waveforms with a streaming architecture that achieves 24kHz synthesis with less than 150ms latency.

Beyond technical capabilities, we address the ethical imperative of preventing misuse. Voice cloning technology poses significant risks for fraud, impersonation, and disinformation. We integrate an anti-deepfake watermarking system directly into the synthesis pipeline, ensuring that all Zen-Voice output carries an imperceptible but detectable provenance marker. This marker is robust to common audio transformations and enables forensic verification of synthetic speech.

Our contributions are as follows:

1. A hierarchical speaker encoder that achieves state-of-the-art speaker similarity from references as short as 3 seconds.
2. A flow-matching prosody transfer module that enables independent control of emotion, emphasis, and pacing.
3. A streaming neural codec vocoder with sub-150ms latency for real-time applications.
4. An integrated anti-deepfake watermarking system with 99.7% detection accuracy.

5. Comprehensive evaluation on LibriTTS, VCTK, and a new multilingual benchmark covering 12 languages.

2 Background and Related Work

2.1 Neural Text-to-Speech

Modern neural TTS systems have evolved through several generations. Tacotron (Wang et al., 2017) and Tacotron 2 (Shen et al., 2018) introduced attention-based sequence-to-sequence models that generate mel spectrograms from text, followed by a vocoder (WaveNet (Oord et al., 2016), WaveRNN (Kalchbrenner et al., 2018), or HiFi-GAN (Kong et al., 2020)) for waveform synthesis. FastSpeech (Ren et al., 2019) and FastSpeech 2 (Ren et al., 2021) replaced autoregressive generation with parallel synthesis guided by explicit duration predictions, dramatically reducing inference time.

2.2 Zero-Shot Voice Cloning

Zero-shot voice cloning synthesizes speech in a target speaker’s voice using only a brief reference sample, without any fine-tuning. Speaker encoders (Jia et al., 2018; Cooper et al., 2020) extract fixed-dimensional embeddings that condition the TTS model. VALL-E (Wang et al., 2023) reformulated TTS as language modeling over neural codec tokens, demonstrating strong zero-shot cloning by treating the reference audio as a prompt. VALL-E 2 (Chen et al., 2024) improved upon this with grouped code modeling and repetition-aware sampling. VoiceBox (Le et al., 2024) used flow matching for non-autoregressive generation with infilling capabilities.

2.3 Prosody Modeling

Prosody—the suprasegmental features of speech including pitch, duration, energy, and rhythm—is critical for natural and expressive synthesis. Global Style Tokens (GST) (Wang et al., 2018) learned a bank of style embeddings from reference audio. The Variational Autoencoder (VAE) approach (Zhang et al., 2019) modeled prosody as a latent variable. More recent work has explored hierarchical prosody representations (Sun et al., 2020) and fine-grained prosody control through explicit feature prediction.

2.4 Audio Watermarking

Audio watermarking embeds imperceptible information into audio signals for authentication and provenance tracking. Traditional methods operate in the frequency domain (Cox et al., 2007). Neural watermarking approaches (Pavlovic and Koepl, 2022; Roman et al., 2024) use learned encoders and decoders to embed and extract watermarks with improved robustness. AudioSeal (San Roman et al., 2024) introduced a localized watermarking approach specifically designed for AI-generated speech detection.

3 Architecture

Zen-Voice consists of four main components: (1) a text encoder, (2) the Hierarchical Speaker Encoder, (3) the Prosody Transfer Module, and (4) the Neural Codec Vocoder. We describe each in detail.

3.1 Text Encoder

The text encoder converts input text into a sequence of linguistic feature vectors. We use a pipeline combining:

1. **Grapheme-to-Phoneme (G2P):** Text is converted to IPA phoneme sequences using language-specific G2P models for 12 supported languages (English, Mandarin, Japanese, Korean, Spanish, French, German, Portuguese, Italian, Hindi, Arabic, Russian). A language identification module automatically selects the appropriate G2P model.
2. **Phoneme Encoder:** Phoneme sequences are encoded by a 6-layer transformer with relative positional encoding:

$$\mathbf{H}_{\text{text}} = \text{Transformer}_{\text{text}}(\text{Embed}(\mathbf{p}_1, \dots, \mathbf{p}_T)) \in \mathbb{R}^{T \times d} \quad (1)$$

where $\mathbf{p}_i$ are phoneme tokens and $d = 512$.

3. **Semantic Enhancement:** For text requiring contextual disambiguation (homographs, emphasis placement), we optionally condition on semantic features extracted from a lightweight BERT model, injected via cross-attention:

$$\tilde{\mathbf{H}}_{\text{text}} = \mathbf{H}_{\text{text}} + \text{CrossAttn}(\mathbf{H}_{\text{text}}, \mathbf{H}_{\text{BERT}}) \quad (2)$$

3.2 Hierarchical Speaker Encoder (HSE)

The HSE extracts a comprehensive speaker representation from reference audio at three hierarchical levels.

Level 1: Frame-Level Encoder. A convolutional encoder processes 80-dimensional log-Mel spectrograms extracted at 16kHz with 25ms windows and 10ms hop size:

$$\mathbf{F} = \text{Conv1D}_{\text{stack}}(\text{MelSpec}(\mathbf{w})) \in \mathbb{R}^{T_f \times d_f} \quad (3)$$

where T_f is the number of frames and $d_f = 256$. This level captures fine-grained spectral characteristics—formant frequencies, breathiness, nasality—that define vocal timbre.

Level 2: Segment-Level Encoder. A 4-layer transformer with 128-frame windows processes the frame features to capture phoneme-level and syllable-level patterns:

$$\mathbf{S} = \text{Transformer}_{\text{seg}}(\mathbf{F}) \in \mathbb{R}^{T_s \times d_s} \quad (4)$$

where $T_s = \lceil T_f/128 \rceil$ and $d_s = 384$. This level captures articulation patterns, coarticulation effects, and local speaking rate variations.

Level 3: Utterance-Level Encoder. An attentive statistics pooling layer aggregates segment features into a fixed-dimensional speaker embedding:

$$\mathbf{e}_{\text{spk}} = \text{AttentivePooling}(\mathbf{S}) = \sum_{i=1}^{T_s} \alpha_i \cdot [\mathbf{s}_i; \sigma_i] \in \mathbb{R}^{d_e} \quad (5)$$

where $\alpha_i = \text{softmax}(\mathbf{v}^\top \tanh(\mathbf{W}\mathbf{s}_i + \mathbf{b}))$ are attention weights, σ_i are local standard deviations, and $d_e = 512$. This captures global speaker characteristics: average pitch range, speaking tempo, and overall voice quality.

Multi-Scale Contrastive Training. The HSE is trained with a hierarchical contrastive loss on 680,000 hours of multilingual speech data:

$$\mathcal{L}_{\text{HSE}} = \lambda_1 \mathcal{L}_{\text{frame}}^{\text{contrast}} + \lambda_2 \mathcal{L}_{\text{segment}}^{\text{contrast}} + \lambda_3 \mathcal{L}_{\text{utterance}}^{\text{contrast}} \quad (6)$$

where each level uses the InfoNCE loss (Oord et al., 2018) with augmented positive pairs (same speaker, different utterance) and in-batch negatives. We set $\lambda_1 = 0.2$, $\lambda_2 = 0.3$, $\lambda_3 = 0.5$.

Speaker Disentanglement. To separate speaker identity from content and prosody, we apply a gradient reversal layer (Ganin et al., 2016) that penalizes the speaker embedding for containing phoneme information:

$$\mathcal{L}_{\text{disentangle}} = -\gamma \cdot \mathcal{L}_{\text{phoneme_clf}}(\mathbf{e}_{\text{spk}}) \quad (7)$$

where $\gamma = 0.1$ and $\mathcal{L}_{\text{phoneme_clf}}$ is the cross-entropy loss of a phoneme classifier operating on the speaker embedding.

3.3 Prosody Transfer Module (PTM)

The PTM generates prosodic features—pitch contour $f_0(t)$, energy envelope $e(t)$, and phoneme durations $d(t)$ —conditioned on text, speaker identity, and optional prosodic guidance.

Flow-Matching Formulation. We model prosody generation as a conditional flow matching problem (Lipman et al., 2023). Let $\mathbf{z}_0 \sim \mathcal{N}(0, I)$ be a noise sample and $\mathbf{z}_1 = [f_0, e, d]$ be the target prosody features. The flow is parameterized by a vector field $\mathbf{v}_\theta$:

$$\mathbf{v}_\theta(\mathbf{z}_t, t, \mathbf{c}) = \frac{d\mathbf{z}_t}{dt}, \quad \mathbf{z}_t = (1-t)\mathbf{z}_0 + t\mathbf{z}_1 \quad (8)$$

Algorithm 1 Zen-Voice Inference Pipeline

Require: Text s , reference audio $\mathbf{w}_{\text{ref}}$, optional style guidance

```
1:  $\mathbf{H}_{\text{text}} \leftarrow \text{TextEncoder}(s)$  ▷ Phoneme encoding
2:  $\mathbf{e}_{\text{spk}} \leftarrow \text{HSE}(\mathbf{w}_{\text{ref}})$  ▷ Speaker embedding
3:  $\mathbf{e}_{\text{style}} \leftarrow \text{StyleEncoder}(\text{guidance})$  ▷ Prosody guidance
4:  $\mathbf{c} \leftarrow [\mathbf{H}_{\text{text}}; \mathbf{e}_{\text{spk}}; \mathbf{e}_{\text{style}}]$  ▷ Conditioning
5:  $\mathbf{z}_0 \sim \mathcal{N}(0, I)$  ▷ Sample noise
6: for  $t = 0$  to  $1$  in  $N$  steps do ▷ ODE integration
7:    $\mathbf{z}_{t+\Delta t} \leftarrow \mathbf{z}_t + \Delta t \cdot \mathbf{v}_\theta(\mathbf{z}_t, t, \mathbf{c})$ 
8: end for
9:  $[f_0, e, d] \leftarrow \mathbf{z}_1$  ▷ Extract prosody
10:  $\mathbf{y} \leftarrow \text{NCV}(\mathbf{H}_{\text{text}}, f_0, e, d, \mathbf{e}_{\text{spk}})$  ▷ Waveform synthesis
11: return  $\mathbf{y}$ 
```

where $\mathbf{c} = [\tilde{\mathbf{H}}_{\text{text}}; \mathbf{e}_{\text{spk}}; \mathbf{e}_{\text{style}}]$ is the conditioning vector. The training objective is:

$$\mathcal{L}_{\text{flow}} = \mathbb{E}_{t, \mathbf{z}_0, \mathbf{z}_1} \left[\|\mathbf{v}_\theta(\mathbf{z}_t, t, \mathbf{c}) - (\mathbf{z}_1 - \mathbf{z}_0)\|_2^2 \right] \quad (9)$$

Style Conditioning. The style embedding $\mathbf{e}_{\text{style}}$ can be derived from three sources:

- **Reference audio:** A prosody encoder extracts style features from a reference utterance.
- **Emotion labels:** A learned embedding table maps categorical emotions (neutral, happy, sad, angry, fearful, surprised, disgusted) to style vectors.
- **Natural language descriptions:** A text encoder maps descriptions like “whispered, with building excitement” to style vectors via CLIP-like contrastive training on (description, audio) pairs.

Architecture. The PTM uses a DiT (Diffusion Transformer) architecture (Peebles and Xie, 2023) with 12 layers, 8 attention heads, and 512-dimensional hidden states. Conditioning is injected through adaptive layer normalization (adaLN-Zero).

3.4 Neural Codec Vocoder (NCV)

The NCV converts linguistic features, prosodic parameters, and speaker embeddings into high-fidelity audio waveforms.

Codec Token Generation. We use a residual vector quantization (RVQ) scheme with 8 codebooks, each containing 1024 codes:

$$\mathbf{q}_l = \text{Quantize}_l(\mathbf{r}_{l-1}), \quad \mathbf{r}_l = \mathbf{r}_{l-1} - \text{Dequantize}_l(\mathbf{q}_l) \quad (10)$$

where $\mathbf{r}_0$ is the input acoustic representation and $l \in \{1, \dots, 8\}$ indexes the codebook level.

Token Prediction. A 6-layer transformer predicts codec tokens from conditioned features:

$$p(\mathbf{q}_l | \mathbf{q}_{<l}, \mathbf{H}_{\text{text}}, f_0, e, d, \mathbf{e}_{\text{spk}}) = \text{Transformer}_{\text{codec}}(\mathbf{q}_{<l}, \mathbf{c}_{\text{acoustic}}) \quad (11)$$

We use a grouped prediction scheme: codebooks 1–2 are predicted autoregressively, while codebooks 3–8 are predicted in parallel.

Waveform Decoder. Codec tokens are decoded to 24kHz waveforms using a HiFiGAN-style generator (Kong et al., 2020) with multi-period and multi-scale discriminators:

$$\mathbf{y} = \text{HiFiGAN}(\text{Dequantize}(\mathbf{q}_1, \dots, \mathbf{q}_8)) \in \mathbb{R}^L \quad (12)$$

Streaming Architecture. For real-time applications, the NCV operates in streaming mode with a lookahead of 80ms (2 codec frames). Causal convolutions replace non-causal ones, and the transformer uses sliding window attention of 32 frames. This achieves end-to-end latency of 148ms on a single NVIDIA A10G GPU.

4 Anti-Deepfake Watermarking

Given the potential for misuse of voice cloning technology, we integrate a mandatory watermarking system into the synthesis pipeline.

4.1 Watermark Design

The watermarking system embeds a 128-bit payload into the synthesized audio:

- **Bits 1–32:** Model identifier and version hash
- **Bits 33–64:** Timestamp (Unix epoch, second precision)
- **Bits 65–96:** User/API key fingerprint
- **Bits 97–128:** HMAC-SHA256 truncated signature over bits 1–96

Embedding. The watermark is embedded using a learned encoder W_{enc} that modifies the codec tokens before waveform decoding:

$$\tilde{\mathbf{q}} = \mathbf{q} + W_{\text{enc}}(\mathbf{q}, \mathbf{m}) \quad (13)$$

where $\mathbf{m} \in \{0, 1\}^{128}$ is the watermark payload. The training loss is:

$$\mathcal{L}_{\text{wm}} = \lambda_{\text{det}} \cdot \mathcal{L}_{\text{BCE}}(W_{\text{dec}}(\tilde{\mathbf{y}}), \mathbf{m}) + \lambda_{\text{qual}} \cdot \|\tilde{\mathbf{y}} - \mathbf{y}\|_1 + \lambda_{\text{percept}} \cdot \mathcal{L}_{\text{STFT}}(\tilde{\mathbf{y}}, \mathbf{y}) \quad (14)$$

where $\mathcal{L}_{\text{STFT}}$ is a multi-resolution STFT loss ensuring imperceptibility.

Table 1: Watermark detection accuracy under various audio transformations.

Transformation	Bit Acc. (%)	Payload Recovery (%)
None (clean)	99.9	99.8
MP3 128 kbps	99.4	99.1
MP3 64 kbps	98.1	96.8
Opus 32 kbps	97.3	95.2
Gaussian noise (SNR 20 dB)	98.8	97.4
Gaussian noise (SNR 10 dB)	95.2	89.6
Resample 8kHz $\rightarrow$ 24kHz	97.6	95.8
Time stretch $\pm 10\%$	98.2	96.4
Pitch shift ± 2 semitones	97.8	95.9
RIR convolution (medium room)	98.5	97.1
Combined (MP3 + noise + RIR)	94.8	87.3
AI-generated detection (binary)		99.7% accuracy

Robustness Training. During training, we apply a differentiable augmentation pipeline between embedding and detection: MP3 compression (64–320 kbps), Opus and AAC codec simulation, Gaussian and environmental noise (SNR 5–40 dB), resampling (8kHz–48kHz), time stretching ($\pm 20\%$), pitch shifting (± 4 semitones), dynamic range compression, and room impulse response convolution.

4.2 Detection Performance

4.3 Perceptual Impact

The watermark introduces minimal perceptual degradation: A/B testing with 200 listeners showed no statistically significant preference between watermarked and non-watermarked audio ($p = 0.42$, two-tailed binomial test). The signal-to-watermark ratio (SWR) averages 38.2 dB, well above the perceptual threshold.

5 Training

5.1 Data

5.2 Training Pipeline

Training follows a four-stage curriculum:

Stage 1: Speaker Encoder Pre-training (2 weeks, 32 A100 GPUs). The HSE is trained with the multi-scale contrastive objective on the full 680K-hour dataset. We use AdamW with learning rate 3×10^{-4} , batch size 4096, and cosine schedule with 5000 warm-up steps.

Stage 2: TTS Model Training (3 weeks, 64 A100 GPUs). The text encoder, prosody transfer module, and codec predictor are jointly trained on the TTS subset

Table 2: Training data composition for Zen-Voice.

Component	Dataset	Hours	Languages
HSE Pre-training	VoxCeleb 1&2	7,400	en
	Common Voice 16.0	28,000	12
	MLS	50,000	8
	Internal (licensed)	594,600	12
TTS Training	LibriTTS-R	585	en
	VCTK	44	en
	Internal studio recordings	12,000	12
Prosody	Expressive audiobooks	8,400	en, zh, ja
Watermark	Synthetic + real mix	50,000	12
Total (deduplicated)		680,000	12

(12,629 hours). The HSE is frozen during this stage. We use AdamW with learning rate 1×10^{-4} , batch size 256 utterances, and a two-phase schedule: 100K steps on clean studio data, then 200K steps on the full mix with data augmentation.

Stage 3: Vocoder Training (1 week, 16 A100 GPUs). The HiFi-GAN vocoder is trained with multi-period and multi-scale discriminator losses:

$$\mathcal{L}_{\text{vocoder}} = \mathcal{L}_{\text{adv}} + \lambda_{\text{fm}} \mathcal{L}_{\text{feature}} + \lambda_{\text{mel}} \mathcal{L}_{\text{mel}} \quad (15)$$

with $\lambda_{\text{fm}} = 2.0$ and $\lambda_{\text{mel}} = 45.0$.

Stage 4: Watermark Integration (3 days, 8 A100 GPUs). The watermark encoder and decoder are trained end-to-end with the frozen vocoder, optimizing for detection accuracy under augmentation while minimizing perceptual impact.

5.3 Emotion and Expressiveness Training

For emotional speech synthesis, we curate a subset of 2,400 hours of speech with emotion annotations across seven categories. Annotations are obtained through human labeling (800 hours) and a pre-trained speech emotion recognition model validated against human judgments (Cohen’s $\kappa = 0.78$). The PTM is fine-tuned with an emotion classification auxiliary loss:

$$\mathcal{L}_{\text{emotion}} = \mathcal{L}_{\text{flow}} + \mu \cdot \text{CE}(\text{EmotionClf}(z_1), y_{\text{emotion}}) \quad (16)$$

where $\mu = 0.1$ and y_{emotion} is the ground-truth emotion label.

6 Evaluation

We evaluate Zen-Voice on three dimensions: synthesis quality (naturalness), speaker similarity (cloning fidelity), and prosodic expressiveness.

6.1 Benchmarks and Metrics

Datasets.

- **LibriTTS test-clean:** 500 utterances from 39 speakers, standard TTS evaluation set.
- **VCTK:** 109 speakers with diverse accents, 400 utterances per speaker.
- **ZenVoice-Eval:** Our multilingual benchmark with 1200 utterances across 12 languages, 100 speakers, balanced by gender and age.

Metrics.

- **MOS (Mean Opinion Score):** 5-point Likert scale rated by 200 native speakers.
- **UTMOS:** Automated MOS prediction using the UTokyo-SaruLab model (Saeki et al., 2022).
- **Speaker Verification EER:** Equal Error Rate of a pre-trained ECAPA-TDNN (Desplanques et al., 2020) speaker verification model.
- **Word Error Rate (WER):** Intelligibility measured via Whisper-large-v3 transcription.
- **F0 RMSE:** Root mean square error of pitch contour relative to reference.

6.2 Baselines

- **VALL-E** (Wang et al., 2023): Autoregressive neural codec language model.
- **VALL-E 2** (Chen et al., 2024): Improved VALL-E with grouped code modeling.
- **VoiceBox** (Le et al., 2024): Flow-matching TTS with infilling.
- **NaturalSpeech 3** (Ju et al., 2024): Factorized diffusion with discrete tokens.
- **XTTS v2** (Casanova et al., 2024): Open-source multilingual TTS.

6.3 Results

Table 3 shows Zen-Voice achieves a naturalness MOS of 4.38 on LibriTTS, closing the gap to ground-truth (4.52) more than any prior method. Speaker similarity MOS of 4.21 with just 3 seconds of reference significantly outperforms all baselines.

Table 3: Speech quality on LibriTTS test-clean (3-second reference).

Method	Nat. MOS	Sim. MOS	UTMOS	WER (%)
Ground Truth	4.52	–	4.31	2.1
VALL-E	3.84	3.62	3.71	5.8
VALL-E 2	4.02	3.81	3.89	4.2
VoiceBox	4.12	3.94	4.01	3.6
NaturalSpeech 3	4.18	4.02	4.08	3.3
XTTS v2	3.91	3.72	3.82	4.8
Zen-Voice	4.38	4.21	4.22	2.7

Table 4: Speaker verification EER on VCTK (lower is better).

Method	3-sec ref	5-sec ref	10-sec ref
Ground Truth	1.8%		
VALL-E	8.4%	6.2%	4.8%
VALL-E 2	6.1%	4.8%	3.6%
VoiceBox	5.3%	4.1%	3.2%
NaturalSpeech 3	4.7%	3.6%	2.8%
Zen-Voice	3.4%	2.6%	2.1%

Table 4 demonstrates speaker verification EER of 2.1% with 10-second references, approaching the ground-truth EER of 1.8%.

Table 5 shows consistent quality across 12 languages, with highest performance on English and slightly lower on Hindi and Arabic where training data is more limited.

6.4 Emotion Transfer Evaluation

Table 6 evaluates emotional expressiveness using a pre-trained speech emotion recognition model. Zen-Voice achieves 77.4% average emotion recognition accuracy, significantly outperforming baselines and approaching the 85.6% accuracy on real emotional speech.

6.5 Latency Analysis

Zen-Voice achieves a real-time factor (RTF) of 0.18 in batch mode and 0.21 in streaming mode, both well under 1.0. The streaming mode achieves first-audio latency of 148ms.

7 Ablation Studies

7.1 Speaker Encoder Design

The hierarchical design contributes 0.30 MOS improvement over frame-level only, and disentanglement adds 0.09 MOS.

Table 5: Multilingual evaluation on ZenVoice-Eval (10-second reference).

Language	Nat. MOS	Sim. MOS	WER (%)	F0 RMSE
English	4.41	4.52	2.4	18.3
Mandarin	4.32	4.38	3.1	22.1
Japanese	4.28	4.31	3.8	19.7
Korean	4.18	4.22	4.2	21.4
Spanish	4.35	4.41	2.8	17.8
French	4.31	4.36	3.2	18.9
German	4.27	4.33	3.5	20.2
Portuguese	4.22	4.28	3.9	19.3
Italian	4.29	4.34	3.3	18.6
Hindi	4.08	4.12	5.1	24.3
Arabic	4.04	4.08	5.6	25.1
Russian	4.15	4.19	4.4	22.8
Average	4.24	4.30	3.8	20.7

Table 6: Emotion recognition accuracy of synthesized emotional speech.

Method	Neu	Hap	Sad	Ang	Fear	Sur	Avg
Reference audio	92.1	87.3	84.6	89.2	78.4	82.1	85.6
VALL-E	78.3	52.1	48.7	54.2	41.3	45.8	53.4
NaturalSpeech 3	84.2	68.4	62.1	71.3	55.2	58.7	66.7
Zen-Voice	89.4	78.2	74.8	81.3	68.4	72.1	77.4

7.2 Prosody Module Design

Flow matching achieves the best quality at 25 ODE steps, with 10 steps providing a favorable quality-speed trade-off.

7.3 Reference Length Sensitivity

Performance improves significantly from 1 to 10 seconds with diminishing returns beyond 30 seconds.

8 Discussion

8.1 Strengths

Zen-Voice’s modular architecture provides three distinct advantages: (1) the separation of speaker identity and prosody enables independent control; (2) the flow-matching formulation provides deterministic, high-quality prosody generation in few ODE steps; (3) the streaming architecture enables real-time applications without sacrificing quality.

Table 7: Inference latency for 10-second utterances on NVIDIA A10G GPU.

Method	First Token (ms)	RTF	Streaming
VALL-E	1,240	0.82	No
VALL-E 2	680	0.51	No
VoiceBox	320	0.24	No
NaturalSpeech 3	410	0.31	No
Zen-Voice (batch)	280	0.18	No
Zen-Voice (stream)	148	0.21	Yes

Table 8: Speaker encoder architecture ablation on VCTK (3-second reference).

Architecture	Sim. MOS	EER (%)	Params
ECAPA-TDNN (frozen)	3.72	6.8	6.2M
x-vector	3.68	7.2	4.8M
Frame-level only	3.91	5.1	12M
Segment-level only	3.98	4.4	18M
HSE (no disentangle)	4.12	3.8	32M
HSE (full)	4.21	3.4	32M

8.2 Limitations

- **Singing voice:** Zen-Voice produces suboptimal results for singing, where pitch accuracy and vibrato control are critical.
- **Extremely short references:** Below 3 seconds, speaker similarity degrades noticeably.
- **Cross-lingual cloning:** Accent artifacts can occur when cloning into a language not present in the reference.
- **Watermark removal:** Targeted adversarial attacks could potentially remove the watermark.

8.3 Ethical Considerations

Voice cloning poses significant ethical risks. We mitigate these through:

1. **Mandatory watermarking:** All API output carries provenance markers that cannot be disabled.
2. **Consent verification:** API users must attest consent from the voice owner.
3. **Voice blocklist:** Protected voices are rejected by speaker verification at inference time.
4. **Detection tools:** The watermark detector is released as a free, open-source tool.

Table 9: Prosody generation method ablation on LibriTTS.

Method	Nat. MOS	F0 RMSE (Hz)	Steps
Duration predictor only	4.02	32.1	1
VAE prosody	4.14	26.8	1
Diffusion (DDPM, 50 steps)	4.28	21.4	50
Flow matching (10 steps)	4.35	19.2	10
Flow matching (25 steps)	4.38	18.3	25

Table 10: Speaker similarity vs. reference audio length on VCTK.

Ref Length	1s	3s	5s	10s	30s	60s
Sim. MOS	3.82	4.21	4.38	4.52	4.61	4.63
EER (%)	7.2	3.4	2.6	2.1	1.9	1.9

9 Conclusion

We presented Zen-Voice, a neural speech synthesis system achieving state-of-the-art zero-shot voice cloning from as little as 3 seconds of reference audio. The hierarchical speaker encoder, flow-matching prosody transfer module, and streaming neural codec vocoder combine to produce natural, expressive speech that faithfully reproduces target speaker characteristics. Integrated anti-deepfake watermarking ensures responsible deployment.

Zen-Voice achieves a naturalness MOS of 4.38 on LibriTTS test-clean, a speaker similarity MOS of 4.21 with 3-second references, and a speaker verification EER of 2.1% with 10-second references. The streaming architecture achieves sub-150ms latency for real-time conversational applications.

Models and inference code are available at <https://github.com/hanzoai/zen-voice> under Apache 2.0. The watermark detection tool is released at <https://github.com/hanzoai/zen-voice-detect>.

References

- Casanova, E., Weber, J., Shulby, C. D., Junior, A. C., Gölge, E., and Ponti, M. A. XTTS: A massively multilingual zero-shot text-to-speech model. In *Proceedings of Interspeech*, 2024.
- Chen, S., Wu, S., Wang, C., Chen, S., Wu, Y., Liu, S., Zhou, L., Liu, J., Kanda, N., Yoshioka, T., et al. VALL-E 2: Neural codec language models are human parity zero-shot text to speech synthesizers. *arXiv preprint arXiv:2406.05370*, 2024.
- Cooper, E., Lai, C.-I., Yasuda, Y., Fang, F., Wang, X., Chen, N., and Yamagishi, J. Zero-shot multi-speaker text-to-speech with state-of-the-art neural speaker embeddings. In *Proceedings of ICASSP*, pp. 6184–6188, 2020.
- Cox, I., Miller, M., Bloom, J., Fridrich, J., and Kalker, T. *Digital Watermarking and Steganography*. Morgan Kaufmann, 2nd edition, 2007.

- Desplanques, B., Thienpondt, J., and Demuynck, K. ECAPA-TDNN: Emphasized channel attention, propagation and aggregation in TDNN based speaker verification. In *Proceedings of Interspeech*, pp. 3830–3834, 2020.
- Ganin, Y., Ustinova, E., Ajakan, H., Germain, P., Larochelle, H., Laviolette, F., Marchand, M., and Lempitsky, V. Domain-adversarial training of neural networks. *JMLR*, 17(59):1–35, 2016.
- Jia, Y., Zhang, Y., Weiss, R., Wang, Q., Shen, J., Ren, F., Chen, Z., Nguyen, P., Pang, R., Lopez Moreno, I., and Wu, Y. Transfer learning from speaker verification to multispeaker text-to-speech synthesis. In *NeurIPS*, pp. 4480–4490, 2018.
- Ju, Z., Wang, Y., Shen, K., Tan, X., Xin, D., Yang, D., Liu, Y., Leng, Y., Song, K., Tang, S., et al. NaturalSpeech 3: Zero-shot speech synthesis with factorized codec and diffusion models. In *ICML*, 2024.
- Kalchbrenner, N., Elsen, E., Simonyan, K., Noury, S., Casagrande, N., Lockhart, E., Stimberg, F., Oord, A., Dieleman, S., and Kavukcuoglu, K. Efficient neural audio synthesis. In *ICML*, pp. 2410–2419, 2018.
- Kong, J., Kim, J., and Bae, J. HiFi-GAN: Generative adversarial networks for efficient and high fidelity speech synthesis. In *NeurIPS*, pp. 17022–17033, 2020.
- Le, M., Vyas, A., Shi, B., Karrer, B., Sari, L., Moritz, R., Williamson, M., Manohar, V., Adi, Y., Mahadeokar, J., and Hsu, W.-N. Voicebox: Text-guided multilingual universal speech generation at scale. In *NeurIPS*, 2024.
- Lipman, Y., Chen, R. T. Q., Ben-Hamu, H., Nickel, M., and Le, M. Flow matching for generative modeling. In *ICLR*, 2023.
- Oord, A. v. d., Dieleman, S., Zen, H., Simonyan, K., Vinyals, O., Graves, A., Kalchbrenner, N., Senior, A., and Kavukcuoglu, K. WaveNet: A generative model for raw audio. *arXiv preprint arXiv:1609.03499*, 2016.
- Oord, A. v. d., Li, Y., and Vinyals, O. Representation learning with contrastive predictive coding. *arXiv preprint arXiv:1807.03748*, 2018.
- Pavlovic, N. and Koepl, H. Robust audio watermarking with deep neural networks. In *IEEE WIFS*, 2022.
- Peebles, W. and Xie, S. Scalable diffusion models with transformers. In *ICCV*, pp. 4195–4205, 2023.
- Ren, Y., Ruan, Y., Tan, X., Qin, T., Zhao, S., Zhao, Z., and Liu, T.-Y. FastSpeech: Fast, robust and controllable text to speech. In *NeurIPS*, 2019.
- Ren, Y., Hu, C., Tan, X., Qin, T., Zhao, S., Zhao, Z., and Liu, T.-Y. FastSpeech 2: Fast and high-quality end-to-end text to speech. In *ICLR*, 2021.

- San Roman, R., Fernandez, P., Elsahar, H., Défossez, A., Furon, T., and Tuli, T. Proactive detection of voice cloning with localized watermarking. In *ICML*, 2024.
- Saeki, T., Xin, D., Nakata, W., Koriyama, T., Takamichi, S., and Saruwatari, H. UTMOS: UTokyo-SaruLab system for VoiceMOS Challenge 2022. In *Proceedings of Interspeech*, 2022.
- San Roman, R., Fernandez, P., Défossez, A., Furon, T., Tuli, T., and Elsahar, H. AudioSeal: Proactive localized watermarking. In *ICML*, 2024.
- Shen, J., Pang, R., Weiss, R. J., Schuster, M., Jaitly, N., Yang, Z., Chen, Z., Zhang, Y., Wang, Y., Skerry-Ryan, R., et al. Natural TTS synthesis by conditioning WaveNet on mel spectrogram predictions. In *ICASSP*, pp. 4779–4783, 2018.
- Sun, G., Zhang, Y., Weiss, R. J., Cao, Y., Zen, H., and Wu, Y. Generating diverse and natural text-to-speech samples using a quantized fine-grained VAE and autoregressive prosody prior. In *ICASSP*, pp. 6699–6703, 2020.
- Wang, Y., Skerry-Ryan, R., Stanton, D., Wu, Y., Weiss, R. J., Jaitly, N., Yang, Z., Xiao, Y., Chen, Z., Bengio, S., et al. Tacotron: Towards end-to-end speech synthesis. In *Interspeech*, pp. 4006–4010, 2017.
- Wang, Y., Stanton, D., Zhang, Y., Skerry-Ryan, R., Battenberg, E., Shor, J., Xiao, Y., Jia, Y., Ren, F., and Saurous, R. A. Style tokens: Unsupervised style modeling, control and transfer in end-to-end speech synthesis. In *ICML*, pp. 5180–5189, 2018.
- Wang, C., Chen, S., Wu, Y., Zhang, Z., Zhou, L., Liu, S., Chen, Z., Liu, Y., Wang, H., Li, J., et al. Neural codec language models are zero-shot text to speech synthesizers. *arXiv preprint arXiv:2301.02111*, 2023.
- Zhang, Y., Pan, S., He, L., and Ling, Z.-H. Learning latent representations for style control and transfer in end-to-end speech synthesis. In *ICASSP*, pp. 6945–6949, 2019.

A Model Hyperparameters

B Computational Requirements

The INT8 quantized model on T4 achieves real-time synthesis ($\text{RTF} < 1.0$) at approximately \$0.0004 per second of generated audio.

Table 11: Zen-Voice module hyperparameters.

Module	Parameter	Value
Text Encoder	Layers	6
	Hidden dim	512
	Attention heads	8
	FFN dim	2048
HSE	Frame encoder channels	[64, 128, 256]
	Segment transformer layers	4
	Speaker embedding dim	512
	Total parameters	32M
	Contrastive temperature	0.07
Prosody Transfer	DiT layers	12
	Hidden dim	512
	ODE steps (inference)	25
	Classifier-free guidance	2.0
Neural Codec	RVQ codebooks	8
	Codebook size	1024
	Codec frame rate	50 Hz
	Streaming lookahead	80 ms
Watermark	Payload bits	128
	Encoder layers	4
	SWR target	≥ 35 dB

Table 12: Inference computational requirements by deployment configuration.

Configuration	GPU	VRAM	RTF
Full model (FP16)	A100 80GB	14.2 GB	0.12
Full model (FP16)	A10G 24GB	14.2 GB	0.18
INT8 quantized	T4 16GB	8.1 GB	0.34
INT4 quantized	L4 24GB	5.2 GB	0.42
Streaming (FP16)	A10G 24GB	14.8 GB	0.21
CPU (INT4)	32-core Xeon	6.8 GB RAM	2.8