
Zen-Voice:

A Packaged Zero-Shot Voice-Cloning Distribution built on Qwen3-TTS

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

Zen-Voice is a packaged, deployment-ready distribution of Alibaba’s open-source **Qwen3-TTS** text-to-speech model (Qwen Team, 2026), with an optional provenance-watermarking layer. It is *not* a speech synthesizer we trained from scratch: the zero-shot voice-cloning capability—reproducing a speaker’s timbre and speaking style from a short reference and synthesizing arbitrary text—is entirely that of the upstream Qwen3-TTS weights, which are released under the Apache 2.0 license. The Qwen team reports that the underlying model clones a voice from roughly three seconds of reference audio and covers ten languages. Zen-Voice’s own contribution is integration, not modeling: convenient packaging of the upstream model plus an opt-in watermarking and consent path for responsible deployment. For provenance marking we integrate **AudioSeal** (San Roman et al., 2024), an existing localized neural watermark for AI-generated-speech detection, rather than a bespoke watermarker. This whitepaper documents what Zen-Voice wraps, the provenance of the underlying model, the responsible-use tooling, and deployment guidance. We deliberately do not present voice-quality benchmark tables of our own; capability claims are attributed to the upstream Qwen3-TTS report.

Keywords: Voice Cloning, Text-to-Speech, Qwen3-TTS, Provenance Watermarking, Responsible Deployment

1 Introduction

Voice cloning—synthesizing arbitrary text in a target speaker’s voice from a short reference—has advanced rapidly in the open-source ecosystem. Several strong models are now publicly available, and the practical barrier to using them is frequently *deployment and governance* rather than raw capability: packaging weights for local inference, wiring a clean API, and adding the consent and provenance controls that responsible voice synthesis requires.

Zen-Voice targets that gap. It is a packaging and responsible-deployment layer over Alibaba’s open-source **Qwen3-TTS** model (Qwen Team, 2026); it does *not* introduce a new synthesis architecture. We are explicit about this because earlier drafts of this document described a bespoke “Hierarchical Speaker Encoder / Prosody Transfer Module / Neural Codec Vocoder” trained on hundreds of thousands of hours, with head-to-head MOS/EER tables against VALL-E and VoiceBox. No such model or evaluation exists; those claims have been removed. The synthesis capability is upstream Qwen3-TTS’s, and we attribute it as such.

1.1 What Zen-Voice Provides

- **Packaging:** Local-inference build artifacts of the upstream Qwen3-TTS weights and a thin synthesis API.
- **Provenance watermarking:** An opt-in path that marks generated audio using **AudioSeal** (San Roman et al., 2024), an existing localized neural watermark, rather than a watermark we trained.
- **Consent tooling:** Hooks for a consent record and a protected-voice blocklist at synthesis time.

1.2 What Zen-Voice Does Not Provide

- No new speaker encoder, prosody model, or vocoder trained from scratch.
- No independently-measured naturalness, speaker-similarity, or EER numbers, and no comparison tables against other systems.
- No proprietary synthesis architecture; the model is Qwen3-TTS.

2 Underlying Model and Provenance

2.1 Qwen3-TTS

Zen-Voice wraps **Qwen3-TTS**, the open-source text-to-speech series released by the Qwen team at Alibaba Cloud (Qwen3TTSForConditionalGeneration) (Qwen Team, 2026). The upstream family is released under the **Apache 2.0** license, permitting redistribution and commercial use; Zen-Voice inherits that license and redistributes

Algorithm 1 Zen-Voice inference (wrapper over upstream Qwen3-TTS)

Require: Text s , reference audio $\mathbf{w}_{\text{ref}}$, consent record r

- 1: **assert** consent r is valid and $\mathbf{w}_{\text{ref}}$ not on the protected-voice blocklist
 - 2: $\mathbf{y} \leftarrow \text{Qwen3TTS}(s, \mathbf{w}_{\text{ref}})$ $\triangleright$ Upstream synthesis (unmodified)
 - 3: **if** watermarking enabled **then**
 - 4: $\mathbf{y} \leftarrow \text{AudioSeal.embed}(\mathbf{y})$ $\triangleright$ Existing localized watermark
 - 5: **end if**
 - 6: **return** $\mathbf{y}$
-

the weights unmodified. Per Alibaba’s release, Qwen3-TTS supports expressive and streaming speech generation, voice design, and zero-shot voice cloning, and covers ten languages (Chinese, English, Japanese, Korean, German, French, Russian, Portuguese, Spanish, Italian). The Qwen team reports that the voice-cloning variant can clone from roughly three seconds of reference audio and that the series was trained on a large multilingual speech corpus. We cite these properties rather than re-measuring them.

Table 1: Underlying model provenance. All capabilities are inherited from Qwen3-TTS.

Item	Value
Upstream model	Qwen3-TTS (Alibaba Cloud / Qwen team)
HF class	<code>Qwen3TTSForConditionalGeneration</code>
License	Apache 2.0 (inherited)
Languages	10 (per upstream)
Zero-shot clone	From short reference (per upstream)
Zen-Voice relationship	Repackaging + opt-in watermark/consent

2.2 Optional backup engine

Where a fully streaming, permissively-licensed alternative is preferred, the same integration accepts **CosyVoice 2** (Du et al., 2024) (FunAudioLLM / Alibaba, Apache 2.0) as a drop-in backup engine. Both upstreams are Apache-2.0, which keeps the distribution license-clean.

2.3 Related open models (context only)

For context, other notable zero-shot TTS systems include VALL-E (Wang et al., 2023) and VoiceBox (Le et al., 2024). We reference these only to situate Qwen3-TTS in the literature; Zen-Voice makes no comparative quality claim against them.

3 Inference Pipeline

The Zen-Voice pipeline is a thin wrapper around upstream inference plus optional governance steps:

All acoustic modeling in step 2 is performed by the upstream model. Zen-Voice adds only the consent check (step 1) and the optional provenance watermark (step 4).

4 Provenance Watermarking

Because synthetic-speech misuse (fraud, impersonation, disinformation) is a real risk, Zen-Voice offers an opt-in provenance watermark on generated audio. We do not train a watermarker; we integrate **AudioSeal** (San Roman et al., 2024), an existing localized neural watermarking method designed for proactive detection of AI-generated speech. AudioSeal embeds an imperceptible, detectable marker and is reported by its authors to remain detectable under common audio transformations.

Carried payload. The integration optionally binds a small payload—a model/version identifier and a reference to the signed consent record—so that provenance can be tied back to a governance ledger. The cryptographic and storage design of that ledger is a deployment concern, not a property of the speech model.

Honest scope. We report AudioSeal’s robustness as published by its authors and do not present a watermark-robustness table of our own in this document. The sibling Zen Live-Dub report contains directly-measured AudioSeal survival numbers on dubbed speech for readers who want measured figures; we avoid duplicating unverified numbers here.

5 Packaging and Tooling

Zen-Voice’s engineering work is integration, not training:

1. **Distribution:** Redistributing the upstream Apache-2.0 Qwen3-TTS weights with local-inference build artifacts.
2. **API:** A thin synthesis wrapper (Algorithm 1).
3. **Governance:** Consent-record verification and a protected-voice blocklist enforced at synthesis time, plus the opt-in AudioSeal watermark.
4. **Verification:** Smoke tests confirming packaged artifacts match the upstream reference implementation on sample inputs.

We do not report training compute, dataset composition, or a training curriculum, because Zen-Voice performs no training of the synthesis model. Dataset and training details for the underlying model are documented by the upstream report (Qwen Team, 2026).

6 Discussion

6.1 Where the value is

Zen-Voice’s value is operational, not architectural: it makes an openly-licensed, capable TTS model (Qwen3-TTS) easy to self-host with a license-clean dependency set, and it adds

consent and provenance controls that bare model weights do not include. The synthesis quality, multilingual coverage, and short-reference cloning are upstream properties, attributed accordingly.

6.2 Limitations

- **No new modeling:** Zen-Voice cannot exceed upstream Qwen3-TTS quality; it inherits the upstream model’s strengths and weaknesses (including any per-language quality differences and difficulty with singing voice).
- **Watermark is a deterrent, not a guarantee:** AudioSeal marks are robust to common transformations but, like any watermark, may be degraded or removed by targeted adversarial processing.
- **Governance is opt-in at the deployment layer:** Consent verification and the protected-voice blocklist are enforced by the deployer; they are not intrinsic, unremovable properties of the model weights.
- **No benchmarks of our own:** We have not run controlled MOS/EER studies; claims here are upstream-attributed or qualitative.

6.3 Ethical Considerations

Voice cloning poses real risks (fraud, impersonation, disinformation). Recommended mitigations at deployment:

1. **Provenance watermarking:** Enable AudioSeal marking on generated audio.
2. **Consent verification:** Require an attested, revocable consent record for any cloned voice, checked at synthesis time.
3. **Voice blocklist:** Reject protected/known voices.
4. **Disclosure:** Where law requires (e.g. synthetic-audio disclosure regimes), label generated audio as AI-generated.

7 Conclusion

Zen-Voice is a packaging and responsible-deployment layer over Alibaba’s openly-licensed **Qwen3-TTS** model, with an optional AudioSeal provenance watermark and consent tooling. It introduces no new synthesis model and reports no benchmarks of its own; the zero-shot cloning capability is upstream Qwen3-TTS’s and is attributed as such. Earlier “from-scratch” architecture and head-to-head benchmark claims have been removed as inaccurate.

The packaging and inference wrapper are available at <https://github.com/hanzoai/zen-voice> under Apache 2.0 (matching the upstream license). The provenance watermark uses AudioSeal (San Roman et al., 2024).

References

- Qwen Team, Alibaba Cloud. Qwen3-TTS: open-source streaming text-to-speech with voice design and zero-shot cloning. Models and code: <https://github.com/QwenLM/Qwen3-TTS>, 2026. Apache 2.0.
- Du, Z. et al. (FunAudioLLM, Alibaba). CosyVoice 2: Scalable streaming speech synthesis with large language models. *arXiv preprint arXiv:2412.10117*, 2024. Code/weights Apache 2.0.
- 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 Architecture and Hyperparameters

Zen-Voice does not define a model architecture of its own. The architecture, parameter counts, tokenizer, codec, and training hyperparameters are those of the upstream **Qwen3-TTS** model and are documented in the upstream report (Qwen Team, 2026). Where CosyVoice 2 is used as the backup engine, refer to its report (Du et al., 2024). We intentionally do not reproduce module-level hyperparameter tables here, because doing so previously implied a bespoke architecture that does not exist.

B Computational Requirements

Memory footprint and throughput follow from the chosen upstream checkpoint and runtime. Qwen3-TTS is reported by Alibaba to run on consumer GPUs (on the order of a few GB of VRAM for smaller variants); exact VRAM and real-time factor depend on the variant, precision, and hardware and should be measured on the target deployment rather than quoted from this document. We do not publish per-configuration RTF or per-second cost figures here, as we have not run a controlled benchmark of our own; any such numbers would require their own documented methodology.