

Zen-Dub: A Pipeline for AI Voice Dubbing and Localization

Technical Whitepaper v2025.03

Zach Kelling
Zen LM Research Team
research@zenlm.org
papers.zenlm.org

March 2025

Abstract

Zen-Dub is a video-dubbing *pipeline* assembled from existing, openly-licensed components—not a single model trained from scratch, and not a “Mixture of Distilled Experts.” It chains automatic speech recognition, machine translation, voice-cloning text-to-speech, and audio-driven lip synchronization into a localization workflow that preserves each speaker’s voice identity and matches mouth motion. Voice cloning uses Alibaba’s **Qwen3-TTS** [1] (Apache 2.0), with **CosyVoice 2** [2] (Apache 2.0) as a streaming backup engine. Lip synchronization uses **MuseTalk** [3] (Lyra Lab, Tencent Music; MIT code), a real-time latent-space inpainting model that regenerates the mouth region at 256×256 and 30+ fps. This whitepaper describes the pipeline’s stages and the provenance and license of each component. We do not report dubbing-quality benchmark numbers of our own; capability claims are attributed to the upstream components rather than re-measured here.

Contents

1	Introduction	2
2	Architecture	2
2.1	Source Audio Processing	3
2.2	Isochrony-Aware Machine Translation	3
2.3	Prosody Transfer	3
2.4	Lip-Sync Adaptation	4
2.5	Voice Synthesis with Speaker Cloning	4
3	Integration, Not Training	4
4	Evaluation	4
5	Related Work	5
5.1	Neural Dubbing	5
5.2	Voice Cloning	5
5.3	Prosody Transfer	5
6	Conclusion	5

1 Introduction

Professional voice dubbing is one of the most labor-intensive tasks in media localization. For a 90-minute film, professional dubbing requires:

- Script translation and adaptation by localization experts
- Casting native-language voice actors
- Recording sessions in acoustic studios
- Audio post-production (mixing, timing adjustment, EQ matching)

This pipeline costs \$50,000–\$500,000 per language for feature-length content and takes 4–12 weeks per language. The result: most content is localized into fewer than 10 languages, leaving global audiences underserved.

Zen-Dub compresses this workflow into an automated multi-stage pipeline (a sequence of existing models, not a single trained model), producing dubbed audio that preserves:

1. **Speaker voice identity** — Timbre, pitch range, resonance, and speaking style characteristics extracted from source audio, cloned into the target language.
2. **Emotional cadence** — Prosodic features (speech rate variation, emphasis patterns, pause placement) are detected in source audio and reproduced in target synthesis with language-appropriate phoneme mapping.
3. **Lip-sync timing** — Phoneme duration in the target language synthesis is adjusted to match source mouth movement timing, achieving natural lip-sync without visible audio lag.
4. **Background audio preservation** — Source non-speech audio (music, ambient sound, sound effects) is automatically separated and preserved in the output mix.

2 Architecture

Zen-Dub is a multi-component pipeline. There is no single trained “backbone” and no “Mixture of Distilled Experts”; each stage is a distinct, separately-licensed component behind a stable interface, and the heavy lifting is done by upstream models (Table 1). The pipeline consists of five stages:

1. Source audio processing (ASR + speaker characterization + source separation)
2. Machine translation (isochrony-aware)
3. Prosody transfer and lip-sync timing adaptation
4. Voice synthesis (cloning TTS: Qwen3-TTS or CosyVoice 2)
5. Lip synchronization (MuseTalk) and audio post-processing / mixing

Table 1: Pipeline components and their provenance. Zen-Dub integrates these; it does not train them.

Stage	Component	License
ASR + audio features	Whisper-family encoder	MIT
Translation	Open MT model (swappable)	per upstream
Voice clone (primary)	Qwen3-TTS [1]	Apache 2.0
Voice clone (backup)	CosyVoice 2 [2]	Apache 2.0
Source separation	Demucs	MIT
Lip-sync	MuseTalk [3]	MIT (code)

2.1 Source Audio Processing

Automatic Speech Recognition (ASR): An off-the-shelf ASR model (Whisper-family) transcribes source audio with word-level timestamps. We do not train this component; recognition quality is the upstream model’s.

Speaker characterization: A speaker-embedding network extracts a per-segment speaker identity vector used to condition the cloning TTS. We use an existing speaker encoder rather than a bespoke one.

Background separation: An existing source-separation model (Demucs, MIT) isolates speech from music and ambient sound so the non-speech bed can be preserved in the final mix.

2.2 Isochrony-Aware Machine Translation

Translation in dubbing differs from document translation: the output must be semantically accurate, phonetically compatible with the source timing, and natural in spoken form. Zen-Dub uses an off-the-shelf machine-translation model (swappable) and applies an *isochrony* constraint at decoding time so that the target length matches the source. The constraint is a property of the pipeline, not of a custom model:

Isochrony constraints: The translation length (in phonemes) must approximately match the source length (in phonemes) to preserve lip-sync. A duration-aware decoding procedure biases beam search toward translations with matching phoneme counts:

$$\text{score}(y|x) = \log p_{\text{MT}}(y|x) - \lambda \left| \frac{|y|_{\phi}}{|x|_{\phi}} - 1 \right| \quad (1)$$

where $|y|_{\phi}$ and $|x|_{\phi}$ are the estimated phoneme counts of target and source strings, and $\lambda = 2.0$ is the isochrony penalty.

Spoken form optimization: A post-edit module rewrites written translations into spoken-language style: contractions, natural spoken discourse markers, and abbreviation expansion appropriate to the target language’s spoken register.

Character voice preservation: Named entities and proper nouns are handled with a character name preservation dictionary, ensuring consistency of character names across the dub.

2.3 Prosody Transfer

Prosody (speech rhythm, stress, intonation) is transferred from source to target via a prosody encoder-decoder:

$$\hat{\mathbf{p}}_{\text{target}} = f_{\text{prosody}}(\mathbf{p}_{\text{source}}, \mathbf{e}_{\text{emotion}}, \mathbf{l}_{\text{target}}) \quad (2)$$

where $\mathbf{p}_{\text{source}}$ is the source prosody representation (F0 contour, energy envelope, duration sequence), $\mathbf{e}_{\text{emotion}}$ is the emotion embedding, and $\mathbf{l}_{\text{target}}$ is a learned target language prosody prior.

The prosody decoder outputs word-level duration targets and phrase-level F0 contours in the target language, respecting the phonological constraints of the target language (e.g., tonal languages receive special treatment of F0 mapping).

2.4 Lip-Sync Adaptation

Lip-sync accuracy requires that the synthesized audio’s viseme sequence align with the source speaker’s mouth movements at a coarse level ($\pm 80\text{ms}$). The lip-sync adapter modifies phoneme durations in the synthesis plan:

$$d_i^* = \arg \min_{d_i} \left\{ |d_i - \hat{d}_i| + \mu \sum_j |\tau_j^{\text{sync}} - \tau_j^{\text{source}}| \right\} \quad (3)$$

where d_i are phoneme durations, $\hat{d}_i$ are the natural prosody targets, τ_j^{sync} are synthesized viseme boundary times, and τ_j^{source} are source mouth-open/close event times detected by a lightweight face motion detector.

2.5 Voice Synthesis with Speaker Cloning

Voice synthesis is delegated to an upstream zero-shot cloning TTS model—primarily **Qwen3-TTS** [1], with **CosyVoice 2** [2] as a streaming backup—conditioned on the extracted speaker reference, the translated phoneme/text sequence, and the transferred prosody plan:

$$\hat{A} = \text{CloneTTS}(\text{text}_{\text{target}}, \mathbf{s}_{\text{speaker}}, \mathbf{p}_{\text{prosody}}) \quad (4)$$

The cloning behavior (e.g. synthesis from a short reference) is a property of the upstream model; per Alibaba’s release, Qwen3-TTS clones from roughly three seconds of reference audio. We cite this rather than re-measuring it. Zen-Dub does not train the TTS model.

3 Integration, Not Training

Zen-Dub performs *no* end-to-end model training. There is no joint ASR+MT+TTS multi-task objective, no 970,000-hour dubbing corpus, and no “human evaluation loop” driving fine-tuning—earlier drafts described such a training process, but it does not exist and those claims have been removed. Each stage is an existing model used as released; Zen-Dub’s engineering is the glue: format conversion, the isochrony decoding constraint, prosody/duration adaptation, the lip-sync timing alignment, and the final remix under the preserved background bed.

Dataset and training details for the underlying models are documented by their respective upstream sources (Qwen3-TTS [1], CosyVoice 2 [2], MuseTalk [3]).

4 Evaluation

We do not present dubbing-quality benchmark tables (speaker-similarity MOS, BLEU, lip-sync accuracy, throughput) in this whitepaper. Earlier versions contained such tables with specific numbers (e.g. “MOS 4.3”, “BLEU 42.3”, “87.4% lip-sync”), but those numbers were not the result of a controlled evaluation of this pipeline and have been removed rather than presented as fact.

For component-level quality, readers should consult the upstream reports:

- **Voice cloning:** Qwen3-TTS [1] and CosyVoice 2 [2].
- **Lip synchronization:** MuseTalk [3], which reports real-time 256×256 generation at 30+ fps and competitive visual fidelity / lip-sync accuracy.

A sibling report, *Zen Live-Dub* (newsroom), contains directly-measured, independently-verified numbers for a closely-related license-clean dubbing pipeline (lip-sync throughput, cross-lingual clone similarity, watermark survival); we point readers there for measured figures instead of repeating unverified ones here.

5 Related Work

5.1 Neural Dubbing

Prior work on automatic dubbing focused primarily on isochrony constraints in machine translation [4] and prosody transfer [5]. Neural dubbing systems combine ASR, MT, voice-cloning TTS, and speaker adaptation. Zen-Dub follows this pipeline pattern, assembling existing models and adding isochrony-aware decoding, prosody/timing adaptation, and MuseTalk-based lip synchronization, rather than training a new end-to-end model.

5.2 Voice Cloning

Zero-shot voice cloning from short reference audio has advanced rapidly. Zen-Dub’s synthesis stage is the upstream Qwen3-TTS model [1] (with CosyVoice 2 [2] as a backup), both Apache-2.0; we do not introduce a new vocoder.

5.3 Prosody Transfer

Cross-lingual prosody transfer is complicated by language-specific intonation systems. Work on prosody normalization and language-aware F0 transfer [7] has shown that speaker-independent prosody components (emphasis, rhythm) transfer better than language-dependent ones (tonal patterns in tonal languages).

6 Conclusion

Zen-Dub is a localization *pipeline* that chains existing, openly-licensed models—Qwen3-TTS (with CosyVoice 2 as backup) for voice cloning and MuseTalk for lip synchronization—behind ASR, isochrony-aware translation, and prosody/timing adaptation. Its contribution is integration and license-clean assembly, not a novel trained model and not the “Zen MoDE” architecture earlier drafts claimed.

The practical value is that a multilingual dubbing workflow can be stood up from permissively-licensed parts, with each component’s quality attributable to its upstream source. We deliberately make no quantitative quality claim of our own in this document; measured figures for a closely-related pipeline appear in the Zen Live-Dub report.

Future work targets emotional consistency across scene cuts, multi-speaker handling for crowd scenes, and singing-voice dubbing for musical content.

Acknowledgments

We thank the upstream teams whose openly-licensed models this pipeline integrates: the Qwen team at Alibaba Cloud (Qwen3-TTS), FunAudioLLM (CosyVoice 2), and Lyra Lab at Tencent Music (MuseTalk).

References

- [1] Qwen Team, Alibaba Cloud, “Qwen3-TTS: Open-Source Streaming Text-to-Speech with Voice Cloning,” <https://github.com/QwenLM/Qwen3-TTS>, 2026. Apache 2.0.
- [2] Z. Du et al. (FunAudioLLM, Alibaba), “CosyVoice 2: Scalable Streaming Speech Synthesis with Large Language Models,” *arXiv:2412.10117*, 2024. Apache 2.0.
- [3] Y. Zhang et al. (Lyra Lab, Tencent Music Entertainment), “MuseTalk: Real-Time High-Fidelity Video Dubbing via Spatio-Temporal Sampling,” *arXiv:2410.10122*, 2024. Code: MIT.
- [4] E. Zetterholm, “Voice Imitation: A Phonetic Study of Perceptual Illusions and Acoustic Success,” *Lund University*, 2004.
- [5] M. Federico et al., “From Speech to Speech Translation,” *ICASSP*, 2020.
- [6] J. Kong et al., “VITS2: Improving Quality and Efficiency of Single-Stage Text-to-Speech with Adversarial Learning and Architecture Design,” *Interspeech*, 2023.
- [7] C. Wan et al., “Prosody Transfer in Neural Machine Translation for Dubbing,” *ACL*, 2023.