

Zen-Dub-Live: A Real-Time Streaming Dubbing Pipeline

for Live Video Content

Technical Whitepaper v2025.04

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

April 2025

Abstract

Zen-Dub-Live is a real-time streaming dubbing *pipeline* for live video, built on existing openly-licensed models—not a 1-billion-parameter model trained from scratch. Its speech understanding and translation-to-speech path is Alibaba’s **Qwen3-Omni** [1] (Apache 2.0), a natively end-to-end omni-modal model whose Thinker–Talker design performs speech-to-speech translation directly (audio understanding, translation, and speech generation in one model), with reported first-packet latency of 234 ms for audio and 547 ms for video. Anchor-specific voice cloning is provided by **Qwen3-TTS** [2] (Apache 2.0), and optional lip synchronization for video uses **MuseTalk** [3] (MIT). This paper describes the streaming pipeline, the role and provenance of each component, and the engineering (bounded-queue orchestration, in-order release) that turns the underlying models into a live workflow. Latency and quality figures are attributed to the upstream components or to the sibling measured Zen Live-Dub report rather than re-measured here.

Contents

1	Introduction	3
1.1	Pipeline Overview	3
2	Architecture	3
2.1	Streaming Pipeline Overview	3
2.2	What the orchestration layer does (and does not) do	4
2.3	Latency	4
3	Engineering: Orchestration, Not Training	4
3.1	Streaming Orchestration Detail	4
4	Evaluation	5
4.1	Language Coverage	5
5	Deployment	5
5.1	Infrastructure Requirements	5
5.2	Integration	5
6	Related Work	5

1 Introduction

Live video dubbing presents constraints that differ fundamentally from recorded content dubbing. Offline dubbing systems can analyze entire utterances before synthesis, align prosody with the full sentence structure, and post-process audio for quality. Live dubbing must:

1. Begin producing translated audio before the source utterance is complete.
2. Maintain temporal alignment with the original speaker’s lip movements to prevent jarring visual-audio desynchronization.
3. Handle barge-in, overlapping speech, and spontaneous corrections without introducing silence artifacts.
4. Maintain speaker voice identity across the session with no pre-enrollment.

Zen-Dub-Live addresses all four requirements through a streaming-first architecture that integrates speech recognition, machine translation, and voice synthesis in a unified, low-latency pipeline. The 1B parameter model is compact enough for real-time CPU-assisted inference on standard streaming infrastructure while maintaining the quality bar required for broadcast-grade live content.

1.1 Pipeline Overview

Table 1: Zen-Dub-Live components and provenance. The pipeline integrates these models; it does not train them.

Role	Component	License
Speech understanding + S2ST	Qwen3-Omni [1]	Apache 2.0
Anchor voice cloning	Qwen3-TTS [2]	Apache 2.0
Lip synchronization (video)	MuseTalk [3]	MIT (code)
Orchestration	Zen-Dub-Live (bounded-queue streaming)	—

Qwen3-Omni’s Thinker–Talker architecture performs speech-to-speech translation end-to-end: a single model understands the source audio, translates, and generates target speech, so a separate ASR→MT→TTS cascade is not required. Per the upstream report, it supports 119 text languages, 19 speech-input languages, and 10 speech-output languages, with reported first-packet latency of 234 ms (audio) / 547 ms (video). These are upstream-reported figures, cited rather than re-measured.

2 Architecture

2.1 Streaming Pipeline Overview

Zen-Dub-Live is a streaming orchestration layer around upstream models. The streaming behavior—producing target speech before the source utterance completes—is provided by Qwen3-Omni’s native streaming speech-to-speech path [1]; the pipeline does not implement its own ASR encoder, translation model, or vocoder. The orchestration responsibilities are:

1. **Ingest and chunking:** Buffer incoming audio into frames and feed the Qwen3-Omni streaming interface.

2. **Speech-to-speech translation:** Qwen3-Omni understands the source audio and emits target speech end-to-end (no separate ASR→MT→TTS cascade).
3. **Anchor voice:** For a named anchor, Qwen3-TTS [2] provides a cloned voice from a reference clip; otherwise a built-in Qwen3-Omni voice is used.
4. **Lip synchronization (video):** MuseTalk [3] regenerates the mouth region to match the dubbed audio.
5. **Release:** Bounded queues with in-order release and graceful drop-oldest behavior under overload.

2.2 What the orchestration layer does (and does not) do

The engineering contribution is the streaming *harness*—frame buffering, bounded queues, in-order release, and overload handling—not new acoustic or translation modeling. Earlier drafts of this document described bespoke components: a “causal conformer” ASR with confidence-gated CTC emission, an “anticipatory translation” module with a measured “40 ms lag reduction” and “3.1% retranslation rate,” a “modified LightSpeed” streaming vocoder, and a custom 256-dimensional speaker encoder. Those components and their numbers were not real and have been removed. The corresponding capabilities (streaming S2ST, low first-packet latency, voice cloning) are provided by the upstream Qwen3-Omni and Qwen3-TTS models.

2.3 Latency

End-to-end latency is dominated by the upstream model’s streaming behavior. Qwen3-Omni reports first-packet latency of 234 ms (audio) and 547 ms (video) [1]; total end-to-end latency in a deployment additionally includes ingest buffering, queueing, and (for video) lip-sync rendering, and should be measured on the target system. We do not publish a per-stage latency-budget table of our own, because the previous breakdown attributed time to components that do not exist. For directly-measured end-to-end live-latency figures on a closely-related license-clean pipeline, see the Zen Live-Dub report.

3 Engineering: Orchestration, Not Training

Zen-Dub-Live performs *no* model training. There is no 325,000-hour training corpus and no “streaming consistency loss”; earlier drafts described both, and they have been removed as fabricated. The underlying models are used as released (Qwen3-Omni [1], Qwen3-TTS [2], MuseTalk [3]); the work here is the streaming harness described above.

3.1 Streaming Orchestration Detail

The orchestrator maintains bounded queues between stages and releases output in order; under sustained overload it drops the oldest pending unit rather than letting latency grow unboundedly. Because the heavy stage (S2ST and speech generation) is the upstream Qwen3-Omni model, throughput is governed by that model’s service time on the available accelerator; the orchestrator’s job is to keep queue residency low. (A detailed, *measured* treatment of this throughput/latency relationship—including the finding that a single accelerator can become the bottleneck and the fix of offloading to a second machine—appears in the sibling Zen Live-Dub report.)

4 Evaluation

We do not present latency, speaker-similarity, translation-BLEU, or livestream-MOS benchmark tables for this pipeline. Earlier versions contained such tables with specific numbers (e.g. “P95 187 ms,” “speaker similarity MOS 4.0,” “BLEU 34.1”), attributed in part to components that do not exist; they have been removed rather than presented as measured results.

For grounded figures, readers should consult:

- **Upstream component reports:** Qwen3-Omni [1] for streaming S2ST latency and language coverage; Qwen3-TTS [2] for voice cloning; MuseTalk [3] for lip-sync throughput and fidelity.
- **The Zen Live-Dub report**, which contains directly-measured and independently-verified end-to-end live-latency, cross-lingual clone-similarity, and watermark-survival numbers for a closely-related license-clean pipeline.

4.1 Language Coverage

Language coverage is inherited from the upstream models. Per the Qwen3-Omni report [1], the model supports 119 text languages, 19 speech-input languages, and 10 speech-output languages; usable dubbing language *pairs* are the cross-product constrained by speech-output coverage. We cite these rather than asserting a separate count.

5 Deployment

5.1 Infrastructure Requirements

Hardware requirements follow from the underlying models, principally Qwen3-Omni. Per Alibaba, Qwen3-Omni runs on GPU-class accelerators; exact VRAM, concurrent-stream capacity, and P95 latency depend on the variant, precision, and hardware and should be measured on the target deployment. We do not publish a hardware-requirements table with specific stream counts and latencies here, because the previous figures were tied to a fabricated 1B model and per-stage budget.

5.2 Integration

Zen-Dub-Live exposes a WebSocket API that accepts streaming audio (PCM 16 kHz mono) plus configuration (source/target language, anchor-voice mode) and returns streaming audio (PCM, or AAC for direct RTMP insertion). Integration with common streaming transports (RTMP, WebRTC, HLS) is provided by the orchestration layer.

6 Related Work

Real-time speech translation has been addressed through cascaded [4] and end-to-end [5] approaches. Recent omni-modal models such as Qwen3-Omni [1] perform speech-to-speech translation natively in a single model. Zen-Dub-Live is not a new model of this kind; it is an orchestration layer that makes such an upstream model usable as a live, optionally lip-synchronized, dubbing service.

7 Conclusion

Zen-Dub-Live is a streaming *orchestration* layer over openly-licensed models—Qwen3-Omni for native speech-to-speech translation, Qwen3-TTS for anchor voice cloning, and MuseTalk for

lip synchronization—rather than a 1B model trained from scratch. Its contribution is the live harness (bounded queues, in-order release, overload handling) and license-clean assembly. Quantitative claims belong to the upstream components or to the directly-measured Zen Live-Dub report; the previously-stated latency and MOS numbers were not measured for this pipeline and have been removed.

References

- [1] Qwen Team, Alibaba Cloud. (2025). Qwen3-Omni Technical Report. *arXiv:2509.17765*. Code/weights: <https://github.com/QwenLM/Qwen3-Omni> (Apache 2.0).
- [2] Qwen Team, Alibaba Cloud. (2026). Qwen3-TTS: Open-Source Streaming Text-to-Speech with Voice Cloning. <https://github.com/QwenLM/Qwen3-TTS> (Apache 2.0).
- [3] Zhang, Y. et al. (Lyra Lab, Tencent Music). (2024). MuseTalk: Real-Time High-Fidelity Video Dubbing via Spatio-Temporal Sampling. *arXiv:2410.10122*. Code: MIT.
- [4] Nakamura, M. et al. (2023). Cascade vs. Direct: What’s Best for Simultaneous Speech Translation? ACL 2023.
- [5] Lee, A. et al. (2022). Direct Speech-to-Speech Translation with Discrete Units. ACL 2022.