

Zen-Live: A Real-Time Multimodal Conversation Service

built on Qwen3-Omni

Technical Whitepaper v2025.03

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

March 2025

Abstract

Zen-Live is a real-time voice/video conversation *service* built on Alibaba’s open-source **Qwen3-Omni** [1] (Apache 2.0)—not a 7-billion-parameter model we trained, and not a custom architecture with “native VAD,” “early-exit,” and a bespoke “dialogue draft model.” Qwen3-Omni is a natively end-to-end omni-modal model whose Thinker–Talker design understands text, audio, images, and video and generates speech in real time, with upstream-reported streaming first-packet latency of 234 ms (audio) and 547 ms (video). Zen-Live’s contribution is the serving layer: a WebRTC front end for control-room/voice-assistant use, turn-taking and barge-in handling at the application level, and integration with backend inference infrastructure. This paper describes that serving architecture and the provenance of the underlying model. Latency and accuracy figures are attributed to the upstream model rather than re-measured here; previously-reported benchmark tables have been removed as unsubstantiated.

Contents

1	Introduction	3
1.1	System Overview	3
2	Architecture	3
2.1	Multimodal Understanding and Speech Generation (Upstream)	3
2.2	Serving and Interaction Layer (Zen-Live)	4
3	No Training	4
4	Evaluation	4
5	Deployment	4
5.1	Inference Hardware	4
5.2	API Integration	5
6	Safety Considerations	5
6.1	Content Filtering	5
6.2	Consent and Privacy	5

7	Related Work	5
8	Conclusion	5

1 Introduction

Human conversation operates at time constants incompatible with standard LLM inference. Turn-taking decisions occur within 200–400ms of an interlocutor finishing an utterance. Barge-in (interrupting mid-utterance) is expected and natural, requiring the model to gracefully abort synthesis and switch to listening mode. Backchannels (“mm-hmm”, “right”, “I see”) are produced every 3–8 seconds during active listening to signal engagement.

Standard LLM serving pipelines address these requirements through external orchestration: a separate VAD model, turn detection heuristics, and speech synthesis pipeline. This cascade introduces 80–150ms of coordination overhead and creates seams in conversational behavior visible to users.

Zen-Live integrates all conversational management functions into a single model with three key design principles:

1. **Latency-first design:** Every architectural choice is evaluated against its impact on first-token latency, the primary user-perceptible quality metric in voice conversation.
2. **Native turn management:** Turn-taking, VAD, and barge-in are model-level capabilities, not external pipeline stages.
3. **Streaming compatibility:** The model produces usable partial outputs from the first token, enabling streaming synthesis that begins before generation completes.

1.1 System Overview

Table 1: Zen-Live components and provenance. Zen-Live is a serving layer over the upstream model; it does not train a model.

Component	Value
Underlying model	Qwen3-Omni [1] (Alibaba, Apache 2.0)
Modalities	Text, audio, image, video in; text + speech out
Architecture	Thinker–Talker (per upstream)
Streaming latency	234 ms audio / 547 ms video (upstream-reported)
Serving layer	WebRTC front end + turn/barge-in logic
Backend	Hanzo inference infrastructure
Version	v2025.03

The conversational capability—multimodal understanding and real-time speech generation—is provided by Qwen3-Omni. Zen-Live adds the serving and interaction layer described below.

2 Architecture

2.1 Multimodal Understanding and Speech Generation (Upstream)

The model-level capabilities are Qwen3-Omni’s [1]: it natively ingests text, audio, image, and video and generates text and speech in real time via the Thinker–Talker design (the Thinker handles multimodal reasoning; the Talker generates speech). Because understanding and generation are end-to-end in one model, a separate ASR→LLM→TTS cascade is not required. Zen-Live does not modify this model and does not add “native VAD,” “early-exit,” or a “dialogue draft model”; earlier drafts described such bespoke architecture and specific tuned values (e.g. mean exit layer 14.3/32, draft acceptance $\beta = 0.84$, 6 ms VAD latency), none of which are real, and they have been removed.

2.2 Serving and Interaction Layer (Zen-Live)

Zen-Live’s actual contribution is the serving wrapper:

- **Transport:** A WebRTC/streaming front end for control-room and voice-assistant clients (full-duplex, half-duplex, transcription-only, and speech-out modes).
- **Turn-taking and barge-in:** Application-level logic that decides when to start responding and how to stop gracefully when the user speaks over the assistant. This is conventional dialogue-management glue around the model’s streaming input/output, not a trained turn-taking head. Voice-activity detection, where needed, uses a standard external VAD.
- **Backend integration:** Requests are served by backend inference infrastructure (Hanzo Node / Hanzo API) running the upstream model.

We deliberately do not present equations for “early-exit confidence heads,” “speculative decoding acceptance rates,” or a “turn-state classifier,” because Zen-Live does not implement these as model components.

3 No Training

Zen-Live performs *no* model training. There is no 260,000-hour conversational pre-training corpus, no “latency-aware” early-exit objective, and no turn-taking fine-tuning on annotated frames; earlier drafts described all three, and they have been removed as fabricated. The underlying model is used as released [1]; training and data details for it are documented upstream.

4 Evaluation

We do not present first-token-latency, conversational-WER, turn-taking-accuracy, barge-in, or naturalness-MOS benchmark tables for Zen-Live. Earlier versions contained such tables with specific numbers (e.g. “P95 98 ms,” “WER 4.2%,” “turn-taking accuracy 94.3%”), attributed to a model and components that do not exist as described; they have been removed rather than presented as measured results.

For grounded figures on the underlying model—latency, multimodal benchmark performance, and language coverage—readers should consult the Qwen3-Omni report [1], which reports streaming first-packet latency of 234 ms (audio) / 547 ms (video) and support for 119 text, 19 speech-input, and 10 speech-output languages. Application-level metrics such as end-to-end conversational latency depend on the deployment (network, backend, transport) and should be measured on the target system.

5 Deployment

5.1 Inference Hardware

Table 2: Deployment targets. Concrete capacity and latency depend on the upstream model variant, precision, backend, and network, and should be measured per deployment.

Target	Example Hardware	Notes
Cloud	Datacenter GPU	Highest concurrency
Edge server	Single workstation GPU	Control-room deployment
On-device	Apple Silicon (quantized)	Single-session assistant

We do not quote per-configuration P95 latency or concurrent-session counts here; the previous figures were tied to a fabricated 7B model. Measure on the target backend.

5.2 API Integration

Zen-Live exposes a WebRTC-compatible streaming API with the following modes:

- **Full duplex:** Bidirectional audio streaming with application-level barge-in.
- **Half duplex:** Push-to-talk or VAD-gated turn alternation.
- **Transcription-only:** Audio in, text out, for accessibility or logging use cases.
- **Voice synthesis:** Text in, speech out, using the upstream model’s speech output.

6 Safety Considerations

6.1 Content Filtering

Streaming output can be passed through an application-level safety filter before it reaches the client. This is a deployment-layer guard around the model rather than a property of the model weights; we do not claim a specific per-token overhead figure.

6.2 Consent and Privacy

Recommended deployment practices for voice data:

- Obtain explicit user consent before processing voice data.
- Avoid persistent storage of audio beyond the active session.
- Anonymize telemetry; do not retain raw audio or speaker embeddings beyond what is operationally required.

7 Related Work

Real-time conversational AI has been developed through voice assistants [2, 3], dialogue systems [4], and LLM-based voice interfaces such as Moshi [5]. Natively end-to-end omni-modal models such as Qwen3-Omni [1] now provide real-time multimodal understanding and speech generation in a single model. Zen-Live is not a new model of this kind; it is a serving layer that exposes such an upstream model for live conversational use.

8 Conclusion

Zen-Live is a real-time conversation *service* built on Alibaba’s openly-licensed **Qwen3-Omni** model, not a 7B model we trained and not a custom architecture with native VAD, early-exit, or a dialogue draft model. Its contribution is the serving and interaction layer: WebRTC transport, application-level turn-taking and barge-in, and backend integration. The conversational capability is upstream Qwen3-Omni’s and is attributed as such; the previously-reported latency, WER, and turn-taking numbers were not measured for this system 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] Shum, H.-Y. et al. (2018). From Eliza to XiaoIce: Challenges and Opportunities with Social Chatbots. *Frontiers of IT and EE*.
- [3] Alexa AI Team. (2018). The Alexa Meaning Representation Language. *NAACL 2018*.
- [4] Gao, J. et al. (2019). Neural Approaches to Conversational AI. *Foundations and Trends in IR*.
- [5] Défossez, A. et al. (2024). Moshi: a speech-text foundation model for real-time dialogue. *arXiv:2410.00037*.