

Zen-Live: A Real-Time Conversational AI Model with Native Turn Management and Voice Interaction

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 7 billion parameter language model purpose-built for real-time voice and video conversation, delivering P50 first-token latency of 67ms and P95 of 98ms through speculative decoding and early-exit inference mechanisms. Unlike conversational adaptations of general-purpose LLMs, Zen-Live integrates native voice activity detection (VAD), barge-in handling, and turn-taking management directly into the model architecture, eliminating the 80–150ms overhead of external pipeline components. Zen-Live achieves a word error rate (WER) of 4.2% on conversational speech benchmarks and turn-taking accuracy of 94.3% on the MultiTurn evaluation suite. This paper describes the model architecture, training methodology, latency optimization techniques, and conversational capability evaluations.

Contents

1	Introduction	3
1.1	Model Overview	3
2	Architecture	3
2.1	Interleaved Audio-Text Representation	3
2.2	Early-Exit Inference	4
2.3	Speculative Decoding with Dialogue Draft Model	4
2.4	Native Voice Activity Detection	4
2.5	Turn-Taking Architecture	4
2.6	Barge-In Handling	5
3	Training Methodology	5
3.1	Conversational Pre-Training Data	5
3.2	Latency-Aware Training	5
3.3	Turn-Taking Fine-Tuning	5
4	Evaluation	6
4.1	Latency Benchmarks	6
4.2	Speech Recognition Accuracy	6
4.3	Turn-Taking Accuracy	6
4.4	Barge-In Performance	6
4.5	Naturalness	6

5	Deployment	7
5.1	Inference Hardware	7
5.2	API Integration	7
6	Safety Considerations	7
6.1	Real-Time Content Filtering	7
6.2	Consent and Privacy	7
7	Related Work	8
8	Conclusion	8

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 Model Overview

Table 1: Zen-Live Model Specification

Parameter	Value
Architecture	Streaming Transformer with early-exit
Total Parameters	7B
Context Window	16K tokens (audio + text interleaved)
P50 First-Token Latency	67ms
P95 First-Token Latency	98ms
Turn-Taking Accuracy	94.3%
Word Error Rate (conversational)	4.2%
Supported Languages	18
Version	v2025.03
Release Date	March 2025

2 Architecture

2.1 Interleaved Audio-Text Representation

Zen-Live processes audio and text in a unified interleaved sequence representation. Audio is encoded as discrete acoustic tokens using a quantized audio codec at 50Hz (one token per 20ms frame). Text tokens use the standard 100K vocabulary. A modality embedding distinguishes token types:

$$h_t = \text{Embed}(x_t) + \text{ModalityEmbed}(m_t) + \text{PosEmbed}(t) \quad (1)$$

where $m_t \in \{\text{audio}, \text{text}\}$ is the modality label. This unified representation allows the model to condition text generation directly on the acoustic input without a separate transcription step.

2.2 Early-Exit Inference

Early-exit [1] reduces latency by allowing the model to produce outputs from intermediate layers when confidence is high, avoiding the full forward pass through all layers. Zen-Live implements adaptive early-exit with per-layer confidence heads:

$$\hat{y}_l = \text{ExitHead}_l(h_l), \quad c_l = \text{Confidence}_l(h_l) \quad (2)$$

If $c_l > \tau_{\text{exit}} = 0.92$, the model exits at layer l and returns $\hat{y}_l$ without computing deeper layers. On conversational speech, the mean exit layer is 14.3 out of 32 total layers, reducing effective computation by 55% compared to full forward passes.

The confidence threshold τ_{exit} is calibrated to maintain $< 0.1\%$ accuracy degradation on held-out conversational benchmarks, verified by comparing early-exit outputs to full-pass outputs on 100K evaluation examples.

2.3 Speculative Decoding with Dialogue Draft Model

Zen-Live uses a 180M parameter dialogue-specialized draft model for speculative decoding. Unlike general-purpose draft models, the dialogue draft model is trained specifically on conversational turn patterns, giving higher acceptance rates for the stereotyped phrases common in conversational AI (greetings, acknowledgments, question frames).

$$\beta_{\text{dialogue}} = 0.84, \quad \beta_{\text{general}} = 0.71 \quad (3)$$

The higher acceptance rate yields a $3.1\times$ effective speedup in conversational contexts, compared to $2.4\times$ for a general draft model of equivalent size.

2.4 Native Voice Activity Detection

VAD is implemented as a lightweight auxiliary head on the audio encoder layers, producing a binary activity signal at the frame level (20ms resolution):

$$\hat{v}_t = \sigma \left(W_{\text{VAD}} \cdot h_t^{(8)} + b_{\text{VAD}} \right) \quad (4)$$

where $h_t^{(8)}$ is the hidden state at encoder layer 8. The VAD head adds $< 0.1\%$ parameter overhead and produces activity estimates with 6ms latency from the raw audio frame, compared to 25–40ms for external VAD models operating on decoded PCM.

2.5 Turn-Taking Architecture

Turn-taking prediction determines when it is appropriate to begin a response. Zen-Live implements turn-taking as a classification head that predicts one of three turn states on each VAD-active frame:

- **HOLD**: Speaker is mid-utterance; continue listening.
- **YIELD**: Speaker has completed their turn; begin response.
- **BACKCHANNEL**: Speaker expects acknowledgment; produce a brief backchannel.

The turn state predictor is conditioned on both acoustic features (prosodic contour, speaking rate, final phoneme duration) and semantic features (syntactic completeness, question vs. statement framing):

$$\hat{t}_{\text{turn}} = \text{softmax} (W_t[h_{\text{acoustic}}; h_{\text{semantic}}]) \quad (5)$$

This joint conditioning achieves 94.3% turn-taking accuracy, compared to 87.1% for acoustic-only models and 89.4% for heuristic rule-based systems.

2.6 Barge-In Handling

Barge-in detection monitors the VAD signal during synthesis. When user speech is detected during model output, Zen-Live executes a graceful interruption protocol:

1. Synthesis halts at the next sentence boundary (if within 200ms) or immediately.
2. The partial response is retained in the context as an incomplete assistant turn.
3. Listening mode resumes from the point of interruption.
4. The model’s next response acknowledges the interruption if semantically appropriate.

In a held-out evaluation of 500 barge-in events, Zen-Live achieved clean interruption in 97.3% of cases (synthesis stopped within 1 sentence boundary of the interrupt signal), compared to 82.1% for systems using external VAD with TCP notification delay.

3 Training Methodology

3.1 Conversational Pre-Training Data

Table 2: Training Data Composition

Source	Hours	Fraction
Telephone conversations (licensed)	90,000	34.6%
Podcast and interview audio	60,000	23.1%
Customer service recordings	40,000	15.4%
Video conference transcripts	30,000	11.5%
Synthetic dialogue augmentation	25,000	9.6%
Academic lectures and Q&A	15,000	5.8%
Total	260,000	100%

3.2 Latency-Aware Training

Zen-Live is trained with a latency-aware training objective that penalizes late exit points. In addition to the standard language modeling loss, an early-exit regularization term encourages the model to produce confident predictions from earlier layers:

$$\mathcal{L}_{\text{latency}} = \mathcal{L}_{\text{LM}} + \lambda \sum_l w_l \mathcal{L}_{\text{KD}}(l) \quad (6)$$

where $w_l = (L - l + 1)/L$ weights later layers more heavily (penalizing failures to exit early), and $\mathcal{L}_{\text{KD}}(l)$ is the KL divergence between the full-model output distribution and the early-exit head at layer l .

3.3 Turn-Taking Fine-Tuning

Turn-taking capability is fine-tuned on a corpus of 50K annotated conversational turns, with human annotators labeling each audio frame with HOLD / YIELD / BACKCHANNEL ground truth. The annotation protocol includes inter-annotator agreement calibration; frames with < 80% annotator agreement are excluded from training to avoid learning from ambiguous social cues.

4 Evaluation

4.1 Latency Benchmarks

Table 3: First-Token Latency Benchmarks (single A10G GPU, FP8)

Model	P50	P90	P95	P99
Zen-Live (7B)	67ms	88ms	98ms	134ms
Adapted 7B general model	143ms	191ms	218ms	287ms
Adapted 3B general model	89ms	121ms	143ms	194ms
External pipeline (7B + VAD)	198ms	261ms	294ms	381ms

4.2 Speech Recognition Accuracy

Table 4: Word Error Rate on Conversational Benchmarks

Model	Conversational WER	Noisy WER	Accented WER
Zen-Live (7B)	4.2%	8.7%	7.1%
Adapted 7B ASR model	4.8%	9.3%	8.4%
External ASR pipeline	5.1%	10.2%	9.7%

4.3 Turn-Taking Accuracy

Table 5: Turn-Taking Accuracy on MultiTurn Benchmark

Model	YIELD F1	BACKCHANNEL F1	Overall Accuracy
Zen-Live (7B)	95.1%	91.4%	94.3%
Acoustic-only model	88.3%	83.7%	87.1%
Rule-based heuristics	90.2%	85.4%	89.4%
Human agreement baseline	97.8%	94.1%	97.1%

4.4 Barge-In Performance

Table 6: Barge-In Handling Evaluation (500 events)

Metric	Zen-Live	External VAD Pipeline
Clean interruption rate	97.3%	82.1%
Mean synthesis overshoot	0.4 sentences	2.1 sentences
Context retention after barge-in	94.7%	81.3%
Recovery quality (MOS)	4.1/5	3.4/5

4.5 Naturalness

Human evaluators rated extended 5-minute conversations with Zen-Live on naturalness and engagement:

Table 7: Conversational Naturalness MOS (200 conversations)

Dimension	Zen-Live	Adapted General Model
Response relevance	4.3/5	4.1/5
Turn timing naturalness	4.2/5	3.3/5
Backchannel appropriateness	3.9/5	2.7/5
Interruption handling	4.1/5	2.9/5
Overall conversational MOS	4.1/5	3.3/5

The largest gains from Zen-Live’s native turn management are in timing naturalness (+0.9 MOS) and backchannel appropriateness (+1.2 MOS), dimensions where external pipeline approaches are weakest.

5 Deployment

5.1 Inference Hardware

Table 8: Deployment Configurations

Config	Hardware	P95 Latency	Concurrent Sessions
Cloud (FP8)	1 × H100 80GB	98ms	96 sessions
Edge server (FP8)	1 × A10G 24GB	113ms	48 sessions
On-device (INT4)	Apple M3 Max	187ms	4 sessions

5.2 API Integration

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

- **Full duplex:** Bidirectional audio streaming with native 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 internal TTS component.

6 Safety Considerations

6.1 Real-Time Content Filtering

Zen-Live integrates a lightweight token-level safety filter (the Zen-Guard-Stream module) that operates with < 2 ms overhead per synthesis step, enabling safe streaming outputs without post-hoc filtering that would introduce latency spikes.

6.2 Consent and Privacy

All voice recordings processed by Zen-Live are subject to:

- Explicit user consent before voice data processing.
- No persistent storage of audio beyond the session context window.
- Speaker anonymization in telemetry logs (voice embeddings are discarded; only aggregate quality metrics retained).

7 Related Work

Real-time conversational AI has been developed through voice assistants [2, 3], dialogue systems [4], and recently through LLM-based voice interfaces [5]. Early-exit inference has been applied to NLP tasks [1]. Speculative decoding for conversational AI is explored in [6]. Zen-Live unifies these advances into a model architecture trained end-to-end for conversational requirements.

8 Conclusion

Zen-Live achieves real-time voice conversation with P95 first-token latency of 98ms, turn-taking accuracy of 94.3%, and WER of 4.2% by integrating VAD, turn management, and language generation into a single 7B parameter model. The architecture eliminates the coordination overhead of cascaded pipeline approaches and enables more natural conversational behavior, particularly in barge-in handling and backchannel generation. Zen-Live is deployable on a single GPU for cloud applications and on Apple Silicon for on-device voice assistant use cases.

References

- [1] Schwartz, R. et al. (2020). The Right Tool for the Job: Matching Model and Instance Complexities. ACL 2020.
- [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.
- [6] Medina, G. et al. (2024). Speculative Decoding for Conversational AI Systems. arXiv preprint.