

Zen3-Omni: A Packaging of Qwen3-Omni-30B-A3B

Technical Note v2025.06

Zen LM Research Team

research@zenlm.org

papers.zenlm.org

June 2025

Abstract

Zen3-Omni is *not* a from-scratch model. It is a redistribution and packaging of **Qwen3-Omni-30B-A3B**, the natively end-to-end omni-modal foundation model developed by the Qwen team at Alibaba Cloud and released under the Apache-2.0 license [1]. This note documents the upstream architecture as published by its authors, records the provenance and license obligations that attach to any redistribution, and describes the thin packaging layer (weight conversion, quantization, and serving configuration) that Zen LM applies on top of the upstream checkpoint. No architectural novelty, no separate pre-training run, and no independently produced benchmark results are claimed here. For authoritative capability numbers, the reader is directed to the upstream model cards and the Qwen3-Omni Technical Report [1, 2].

Contents

1	Provenance and Scope	2
2	Upstream Architecture (as published by Qwen)	2
2.1	Thinker–Talker–Code2Wav Design	2
2.2	Audio Encoder (AuT)	2
2.3	Modalities and Language Coverage	3
2.4	Items Removed From Earlier Drafts	3
3	Packaging and Redistribution	3
3.1	Conversion and Quantization	3
3.2	Attribution Requirements	3
4	Evaluation: Use Upstream Numbers	3
5	Related Work	4
6	Conclusion	4

1 Provenance and Scope

Earlier internal drafts of this document described a homegrown “Zen MoDE (Mixture of Distilled Experts)” backbone and reported a set of benchmark scores as if they had been produced by an in-house model. That framing was inaccurate and has been removed. The artifact distributed as “Zen3-Omni” is the Qwen3-Omni-30B-A3B checkpoint published by Alibaba Cloud’s Qwen team [1]. This note exists to attribute the work correctly and to make the redistribution terms explicit.

Upstream model. Qwen3-Omni-30B-A3B is a natively omni-modal model: a single system that accepts text, audio, image, and video inputs and produces both text and natural speech output. It is published in multiple editions on Hugging Face, including `Qwen3-Omni-30B-A3B-Instruct`, `Qwen3-Omni-30B-A3B-Thinking`, and `Qwen3-Omni-30B-A3B-Captioner` [1].

License. The upstream weights are released under **Apache-2.0**. Apache-2.0 is permissive and allows commercial use and redistribution, but it carries obligations that a redistributor must honor: preservation of copyright and license notices, inclusion of the NOTICE file if present, and a clear statement of any modifications. Any Zen LM redistribution must ship the upstream LICENSE and attribution intact and must not misrepresent the model’s origin.

What “packaging” means here. Zen LM does not retrain or re-architect the model. The packaging layer is limited to: (i) format conversion of the published weights to the serving runtime used internally; (ii) optional post-training quantization for deployment; and (iii) inference and serving configuration. These steps do not change the model’s identity or its measured capabilities, and they do not justify a new model name implying original research.

2 Upstream Architecture (as published by Qwen)

The description below summarizes the architecture as documented by the upstream authors in the Qwen3-Omni model card and Technical Report [1, 2]. It is reproduced for the reader’s convenience and is *not* a Zen LM contribution.

2.1 Thinker–Talker–Code2Wav Design

Qwen3-Omni uses a multi-module design built around three components:

- **Thinker** — A Mixture-of-Experts (MoE) transformer responsible for multimodal perception and reasoning. It ingests text, audio, image, and video tokens and produces a unified high-level representation. The “A3B” in the model name reflects an MoE configuration in which only a small subset of the total parameters is activated per token (the model is named for roughly 3B activated parameters out of a 30B-class total parameter budget).
- **Talker** — A second MoE transformer that performs streaming speech generation via multi-codebook autoregressive prediction, enabling low-latency text-to-speech output conditioned on the Thinker’s representation.
- **Code2Wav** — A streaming waveform renderer that converts the Talker’s codebook sequence into audio frame-by-frame for real-time output.

2.2 Audio Encoder (AuT)

Per the upstream authors, the audio front-end is a custom Audio Transformer (**AuT**) encoder trained on a large supervised audio corpus, replacing the Whisper encoder used in earlier Qwen omni-modal work. The exact pre-training scale and configuration are documented upstream [2]; they are not reproduced or independently verified here.

2.3 Modalities and Language Coverage

As published on the model card [1], Qwen3-Omni supports:

- **Inputs:** text, audio, image, and video.
- **Outputs:** text and natural speech.
- **Text:** 119 languages.
- **Speech input:** 19 languages.
- **Speech output:** 10 languages.

2.4 Items Removed From Earlier Drafts

The following claims appeared in prior drafts and are **withdrawn** because they were not substantiated and were presented as in-house work: a proprietary “Zen MoDE” expert backbone with a specific modality-balance loss; specific expert counts and routing top- k values attributed to Zen; a bespoke “temporal sparse attention” formulation; and a “distillation from Zen3-VL (72B)” procedure. Any expert-routing, positional, or attention details that are real belong to the upstream Qwen3-Omni design and are documented by its authors [2].

3 Packaging and Redistribution

3.1 Conversion and Quantization

Zen LM converts the published Qwen3-Omni weights to the internal serving format and may apply post-training quantization for deployment efficiency. These transforms are mechanical and lossy only in the usual sense of quantization; they do not constitute a new training run and do not produce a new model. Where quantization is applied, the quantization scheme and bit-width should be recorded alongside the redistributed artifact so that downstream users understand any quality trade-off relative to the upstream full-precision checkpoint.

3.2 Attribution Requirements

Because the upstream license is Apache-2.0, every redistributed Zen3-Omni artifact must:

1. Retain the upstream copyright and Apache-2.0 LICENSE text.
2. Include the upstream NOTICE file, if any, and not remove attribution.
3. State plainly that the model is a packaging of Qwen3-Omni-30B-A3B by the Qwen team at Alibaba Cloud, and describe any modifications (e.g. quantization).
4. Avoid any naming or marketing that implies the architecture or weights are original Zen LM research.

4 Evaluation: Use Upstream Numbers

This note deliberately reports **no** Zen-produced benchmark scores. Doing so would either duplicate or, worse, misrepresent the upstream results. The Qwen team publishes Qwen3-Omni evaluations on standard suites — including text reasoning (e.g. MMLU-Redux, GPQA, AIME), automatic speech recognition (e.g. LibriSpeech, WenetSpeech), speech instruction following (e.g.

VoiceBench), and a broad set of audio, vision, and audio-visual understanding benchmarks — in the official model card and Technical Report [1, 2].

Guidance for downstream reporting. Cite the upstream-reported numbers with explicit attribution to Qwen, and link to the model card. If Zen LM applies quantization, any re-measured scores should be reported as “Qwen3-Omni-30B-A3B (Zen packaging, N -bit)” with the evaluation harness named, and should be framed as a measurement of the upstream model under a specific deployment configuration — never as the result of a distinct model.

5 Related Work

The relevant prior art is the upstream lineage itself: the Qwen-VL and Qwen2.5-Omni line of multimodal and omni-modal models from Alibaba Cloud, of which Qwen3-Omni is the current generation [1, 2]. General multimodal instruction tuning with frozen vision encoders and projection adapters was popularized by LLaVA [3]; the omni-modal, speech-generating design used here is specific to Qwen3-Omni.

6 Conclusion

“Zen3-Omni” denotes a packaged redistribution of Alibaba Cloud’s Apache-2.0 Qwen3-Omni-30B-A3B, not an independent model. The honest description of the work is: convert, optionally quantize, serve, and attribute. All architectural credit and all capability claims belong to the upstream Qwen3-Omni authors and should be cited to them.

Acknowledgments

We acknowledge the Qwen team at Alibaba Cloud as the authors of Qwen3-Omni-30B-A3B, the model packaged and redistributed here under the Apache-2.0 license.

References

- [1] Qwen Team, Alibaba Cloud, “Qwen3-Omni-30B-A3B,” Hugging Face model card, 2025. <https://huggingface.co/Qwen/Qwen3-Omni-30B-A3B-Instruct>
- [2] Qwen Team, Alibaba Cloud, “Qwen3-Omni Technical Report,” 2025. <https://github.com/QwenLM/Qwen3-Omni>
- [3] H. Liu et al., “Visual Instruction Tuning,” *NeurIPS*, 2023.