

Zen AI Model Family

Zen-Scribe

A Packaged Multilingual Speech-Recognition Distribution
built on Qwen3-ASR

Technical Whitepaper v1.0

Zach Kelling*
research@hanzo.ai

Zoo Labs Foundation
foundation@zoolabs.org

September 2025

Abstract

Zen-Scribe is a packaged, deployment-ready distribution of Alibaba’s open-source **Qwen3-ASR** speech-recognition model [1]. It is *not* a model trained from scratch: the recognition capability, multilingual coverage, and accuracy are entirely those of the upstream Qwen3-ASR weights, which are released under the Apache 2.0 license. Zen-Scribe contributes packaging, not new training—quantized build artifacts (GGUF, MLX), a thin transcription API, and documentation—so that the upstream model is easy to self-host on commodity and Apple-Silicon hardware. This whitepaper documents what Zen-Scribe wraps, the provenance of the underlying model, and how to deploy it. All capability claims are attributed to the upstream Qwen3-ASR technical report rather than re-measured here.

Contents

1	Introduction	3
1.1	What Zen-Scribe Provides	3
1.2	What Zen-Scribe Does Not Provide	3
2	Underlying Model	3
2.1	Provenance	3
2.2	Capabilities (as reported upstream)	3
2.3	“Flash” variant	4
3	Packaging and Tooling	4
4	Use Cases and Applications	4
4.1	Primary Applications	4
4.2	Integration Examples	4
5	Safety, Privacy, and Responsible Use	5
6	Deployment Options	5
6.1	Available Formats	5
6.2	Hardware Requirements	5

*zach@lux.network

7	Future Work	5
8	Conclusion	6
A	Model Card	7

1 Introduction

Self-hosting a strong multilingual speech-recognition model is often blocked not by model quality but by packaging: converting weights to quantized runtime formats, wiring a transcription interface, and documenting hardware requirements. **Zen-Scribe** addresses this gap by repackaging an existing, openly-licensed ASR model—Alibaba’s **Qwen3-ASR** [1]—for convenient local deployment. We make no claim to have trained a recognition model; our contribution is distribution and tooling around upstream weights.

1.1 What Zen-Scribe Provides

- **Packaging:** GGUF and MLX build artifacts of the upstream Qwen3-ASR weights for CPU, GPU, and Apple-Silicon inference.
- **A thin API:** A `transcribe()` convenience wrapper over the upstream model.
- **Documentation:** Hardware sizing and deployment notes (below).

1.2 What Zen-Scribe Does Not Provide

- No new pretraining, fine-tuning, or distillation of the acoustic model.
- No independently-measured accuracy numbers (see §2).
- No proprietary architecture; the architecture is Qwen3-ASR’s.

2 Underlying Model

2.1 Provenance

Zen-Scribe wraps **Qwen3-ASR**, the open-source automatic-speech-recognition series released by the Qwen team at Alibaba Cloud (**Qwen3ASRForConditionalGeneration**) [1]. The upstream family is released under the **Apache 2.0** license, which permits redistribution and commercial use; Zen-Scribe inherits that license and redistributes the weights unmodified.

Component	Specification
Upstream model	Qwen3-ASR (Alibaba Cloud / Qwen team)
Architecture	As published upstream (encoder + LLM decoder)
HF class	Qwen3ASRForConditionalGeneration
License	Apache 2.0 (inherited from upstream)
Capabilities	Multilingual ASR, language ID, timestamps

Table 1: Underlying model provenance. All specifications are inherited from Qwen3-ASR.

2.2 Capabilities (as reported upstream)

The following capabilities are those documented by the upstream Qwen3-ASR technical report [1]; we cite rather than re-measure them. Per Alibaba’s release, the Qwen3-ASR series supports stable multilingual speech, music, and song recognition across a large set of languages and dialects, with automatic language detection, streaming and offline operation from a unified model, and timestamp prediction. For exact word-error-rate figures, benchmark methodology, and per-language breakdowns, readers should consult the upstream report directly; we deliberately do not reproduce benchmark tables we did not run.

2.3 “Flash” variant

Alibaba additionally offers a hosted **Qwen3-ASR-Flash** real-time WebSocket endpoint via the DashScope platform. That hosted service is distinct from the self-hostable open weights that Zen-Scribe packages; Zen-Scribe is a local distribution and does not proxy the hosted API.

3 Packaging and Tooling

Zen-Scribe’s engineering work is confined to packaging:

1. **Format conversion:** Producing GGUF (Q4_K_M, Q5_K_M, Q8_0) and MLX (4-bit, 8-bit) artifacts from the upstream Apache-2.0 weights.
2. **Interface:** A thin `transcribe()` wrapper.
3. **Verification:** Smoke tests confirming the packaged artifacts produce output equivalent to the upstream reference implementation on sample audio.

We do not report custom accuracy or throughput benchmarks here, because the model’s behavior is that of upstream Qwen3-ASR; any independently-run numbers would need their own documented methodology before they could be cited.

4 Use Cases and Applications

4.1 Primary Applications

- Real-time transcription
- Meeting notes and summaries
- Podcast transcription
- Multilingual subtitles
- Voice command processing

4.2 Integration Examples

```
1 import librosa
2 from transformers import AutoProcessor, AutoModelForSpeechSeq2Seq
3
4 # Load the packaged upstream Qwen3-ASR model
5 model = AutoModelForSpeechSeq2Seq.from_pretrained("zenlm/zen-scribe")
6 processor = AutoProcessor.from_pretrained("zenlm/zen-scribe")
7
8 # Transcribe
9 audio, sr = librosa.load("speech.wav", sr=16000)
10 inputs = processor(audio, sampling_rate=sr, return_tensors="pt")
11 transcription = processor.batch_decode(
12     model.generate(**inputs), skip_special_tokens=True
13 )
14 print(transcription[0])
```

Listing 1: Basic Usage Example (loads the upstream Qwen3-ASR weights repackaged by Zen-Scribe)

5 Safety, Privacy, and Responsible Use

Because Zen-Scribe redistributes upstream weights unmodified, its model-level safety properties are those of Qwen3-ASR; we add no new alignment training. At the deployment layer we recommend:

- **Privacy:** Transcription can run fully on-device; audio need not leave the host. Operators should obtain consent before transcribing third-party speech and avoid retaining raw audio beyond what is required.
- **Limitations:** ASR accuracy varies by language, accent, audio quality, and domain. Per-language behavior is inherited from upstream; consult the Qwen3-ASR report [1] for known limitations.
- **Bias:** Recognition error rates can differ across accents and dialects; downstream applications should account for this.

6 Deployment Options

6.1 Available Formats

- **SafeTensors:** Upstream-precision weights (as released by Alibaba).
- **GGUF:** Quantized formats (Q4_K_M, Q5_K_M, Q8_0).
- **MLX:** Apple Silicon optimization (4-bit, 8-bit).

6.2 Hardware Requirements

Memory footprint depends on the upstream parameter count and the chosen quantization. As a rule of thumb, lower-precision quantization reduces memory monotonically ($Q4 < Q8 < FP16$); the values below are approximate sizing guidance for a model in the small-LLM range and should be validated against the specific upstream checkpoint and runtime.

Precision	Approx. Memory	Example Hardware
FP16	higher	Discrete GPU (e.g. RTX 3060+)
INT8	medium	Entry GPU / high-RAM CPU
INT4	lowest	CPU / small-footprint devices

Table 2: Approximate hardware sizing by precision. Exact memory follows from the upstream checkpoint size; measure before provisioning.

7 Future Work

- Tracking upstream Qwen3-ASR releases and re-packaging new checkpoints.
- Additional runtime targets (e.g. ONNX) for the same upstream weights.
- Optional, clearly-labelled fine-tunes for specific domains, kept separate from the base re-distribution.

8 Conclusion

Zen-Scribe is a packaging and deployment layer over Alibaba’s openly-licensed **Qwen3-ASR** model. Its value is convenience—quantized artifacts, a thin API, and documentation—not novel modeling. The recognition quality is entirely upstream’s; we attribute it accordingly and avoid presenting benchmarks we did not run.

Acknowledgments

We thank the Qwen team at Alibaba Cloud for releasing Qwen3-ASR under Apache 2.0, and the broader open-source community whose tooling makes local repackaging possible.

References

- [1] Qwen Team, Alibaba Cloud. (2026). Qwen3-ASR Technical Report. arXiv:2601.21337. Models and code: <https://github.com/QwenLM/Qwen3-ASR> (Apache 2.0).
- [2] Vaswani, A. et al. (2017). Attention Is All You Need. NeurIPS 2017.
- [3] Brown, T. et al. (2020). Language Models are Few-Shot Learners. NeurIPS 2020.
- [4] Radford, A. et al. (2019). Language Models are Unsupervised Multitask Learners. OpenAI Technical Report.
- [5] Ouyang, L. et al. (2022). Training Language Models to Follow Instructions with Human Feedback. NeurIPS 2022.
- [6] Raffel, C. et al. (2020). Exploring the Limits of Transfer Learning with a Unified Text-to-Text Transformer. JMLR 2020.
- [7] Radford, A. et al. (2023). Robust Speech Recognition via Large-Scale Weak Supervision. ICML 2023.
- [8] Schulman, J. et al. (2017). Proximal Policy Optimization Algorithms. arXiv:1707.06347.
- [9] Rafailov, R. et al. (2023). Direct Preference Optimization: Your Language Model is Secretly a Reward Model. NeurIPS 2023.
- [10] Touvron, H. et al. (2023). LLaMA: Open and Efficient Foundation Language Models. arXiv:2302.13971.
- [11] Hu, E. et al. (2021). LoRA: Low-Rank Adaptation of Large Language Models. ICLR 2022.
- [12] Dettmers, T. et al. (2023). QLoRA: Efficient Finetuning of Quantized Language Models. NeurIPS 2023.

A Model Card

Field	Value
Distribution Name	Zen-Scribe
Upstream Model	Qwen3-ASR (Alibaba Cloud / Qwen team)
Relationship	Repackaging of upstream weights (no retraining)
Version	1.0.0
License	Apache 2.0 (inherited from upstream)
Repository	huggingface.co/zenlm/zen-scribe
Upstream	github.com/QwenLM/Qwen3-ASR
Contact	research@hanzo.ai

Table 3: Distribution card. Zen-Scribe is a packaged distribution of Qwen3-ASR, not an independently trained model.