

Zen-Reranker: A Repackaging of Qwen3-Reranker for Zen Retrieval Pipelines

Antje Worring, Zach Kelling*

Hanzo Industries Lux Industries Zoo Labs Foundation
research@lux.network

October 2025

Abstract

Zen-Reranker is a repackaging of Alibaba’s **Qwen3-Reranker** series, deployed as the reranking stage of Zen retrieval pipelines. It is *not* a new model: it is the openly released, Apache-2.0 licensed **Qwen/Qwen3-Reranker** family (0.6B, 4B, and 8B), produced by the Qwen team and built on the Qwen3 dense foundation models [1]. Qwen3-Reranker is a *cross-encoder*: it takes a query and a candidate document together and scores their relevance. Architecturally it is a causal LLM (**Qwen3ForCausalLM**); given an instruction, query, and document, it is prompted to answer “yes” or “no,” and the relevance score is the softmax-normalized probability of the “yes” token,

$$\text{score}(q, d) = \frac{e^{P(\text{yes}|I, q, d)}}{e^{P(\text{yes}|I, q, d)} + e^{P(\text{no}|I, q, d)}}.$$

The models support 100+ languages (including code) and a 32K context. On the Qwen team’s reranking evaluation — re-ranking the top-100 candidates retrieved by Qwen3-Embedding-0.6B — the 4B and 8B rerankers report, e.g., MTEB-R 69.76 / 69.02 and MMTEB-R 72.74 / 72.94 [1]. This document describes how the upstream model is wired into the Zen retrieve-then-rerank stack; all weights, training, and reported numbers are the Qwen team’s and are cited as such. Claims in earlier revisions of a bespoke “native 7680-dimensional” model trained for “Decentralized Semantic Optimization,” with BitDelta compression and Byzantine-robust aggregation, were fabricated and have been removed.

Keywords: reranking, cross-encoder, semantic search, retrieval, model packaging

*Corresponding author: zach@lux.network

1 Introduction

Modern semantic search is typically a two-stage *retrieve-then-rerank* pipeline. A fast bi-encoder embedding model retrieves a candidate set (e.g. the top 100 documents) by vector similarity, and a slower but more accurate *cross-encoder* reranker then re-scores those candidates by reading each (query, document) pair jointly. The reranker is where most of the final ranking quality is recovered, because joint attention over the query and document captures relevance signals that a single-vector similarity cannot.

Zen retrieval pipelines use Alibaba’s open Qwen3 retrieval stack for both stages: Qwen3 embeddings for first-stage retrieval and **Qwen3-Reranker** for the reranking stage. This document concerns the reranking stage. We did not train a reranker; we package the upstream Qwen3-Reranker models and expose them through the Zen serving interface.

1.1 Our Contribution

The contribution of “Zen-Reranker” is integration and packaging, not modeling:

- **Adoption of Qwen3-Reranker:** We deploy the Apache-2.0 Qwen/Qwen3-Reranker models (0.6B, 4B, 8B) [1] unmodified as the Zen reranking stage.
- **Pipeline wiring:** A thin wrapper formats the instruction/query/document chat template and reads the “yes” probability as the relevance score.
- **Honest attribution:** We report the Qwen team’s published reranking numbers with citation and make no benchmark claims of our own.

2 Background

2.1 Upstream Model: Qwen3-Reranker

Qwen3-Reranker is part of the Qwen team’s Qwen3 Embedding release [1]. Its relevant properties are:

- **Sizes:** 0.6B, 4B, and 8B parameters.
- **License:** Apache-2.0.
- **Base:** a fine-tune of the corresponding Qwen3-*Base* dense LLM (e.g. the 0.6B reranker from Qwen3-0.6B-Base, the 4B from Qwen3-4B-Base).

- **Type:** a cross-encoder implemented as a causal LM (`Qwen3ForCausalLM`), not a bi-encoder and not a sequence-classification head.
- **Context length:** 32K tokens.
- **Languages:** 100+ languages, including programming languages.

2.2 Scoring Mechanism

The reranker does not emit an embedding vector. The query and document are placed together in a chat template with a system instruction constraining the answer to “yes” or “no,” and the relevance score is the softmax-normalized probability of the “yes” token at the final position [1]:

$$\text{score}(q, d) = \frac{e^{P(\text{yes}|I, q, d)}}{e^{P(\text{yes}|I, q, d)} + e^{P(\text{no}|I, q, d)}}. \quad (1)$$

This is consistent with describing it as a cross-encoder: “cross-encoder” refers to feeding the query and document together in one forward pass, while the score is produced by the LM’s next-token yes/no probability rather than a classification head. We chose this stack because it is openly licensed (Apache-2.0), strong on public reranking benchmarks, and shares its tokenizer and instruction conventions with the Qwen3 embedding model used for first-stage retrieval.

3 How the Reranker Is Used

3.1 No Training

We perform no training. The Qwen3-Reranker weights were produced by the Qwen team by fine-tuning Qwen3-Base models [1]; we deploy them as released. The “three-stage training,” projection-head expansion, DSO-specific losses, and the GPU-hour cost table that appeared in earlier revisions did not describe any real training run and have been removed. For the actual training recipe, see the Qwen team’s report.

3.2 Reranking Procedure

At serving time, the first-stage embedding retriever returns a candidate set, and the reranker re-scores each candidate by the yes/no mechanism described above. Concretely:

Algorithm 1 Reranking with Qwen3-Reranker

Input: query q , candidate documents $\{d_1, \dots, d_k\}$, instruction I
for each document d_i **do**
 Build chat prompt from system instruction (“answer only yes/no”), I ,
 q , d_i
 Run a single causal forward pass; read final-position logits for tokens
 yes, **no**
 $\text{score}(q, d_i) \leftarrow \frac{e^{P(\text{yes})}}{e^{P(\text{yes})} + e^{P(\text{no})}}$
end for
Output: documents sorted by descending $\text{score}(q, d_i)$

Each candidate is an independent forward pass, so reranking cost scales linearly with the candidate-set size; this is the standard accuracy/latency trade-off of cross-encoders and is why reranking is applied only to a shortlist (e.g. top 100) rather than the whole corpus.

3.3 Instruction Conditioning

Qwen3-Reranker accepts a task instruction alongside the query, allowing the same model to be specialized to a retrieval task (e.g. web search vs. code search) at inference time without retraining [1]. Zen deployments pass an instruction string matched to the retrieval task.

4 Reported Results (Upstream)

The numbers below are the Qwen team’s published reranking results for Qwen3-Reranker [1], reproduced with attribution. We ran no evaluations of our own and make no independent benchmark claims. In the upstream setup, each reranker re-ranks the top-100 candidates retrieved by Qwen3-Embedding-0.6B; the columns are the retrieval subsets of MTEB (English), CMTEB (Chinese), and MMTEB (multilingual), plus multilingual long-document retrieval (MLDR), code retrieval (MTEB-Code), and complex instruction retrieval (FollowIR).

4.1 A Note on Earlier Revisions

Earlier revisions of this document reported an MTEB “average” for a fictional “Zen-Reranker-8B” at 7680 dimensions, along with “DSO retrieval,”

Model	MTEB-R	CMTEB-R	MMTEB-R	MLDR	MTEB-Code	Foll
Qwen3-Reranker-0.6B	65.80	71.31	66.36	67.28	73.42	
Qwen3-Reranker-4B	69.76	75.94	72.74	69.97	81.20	
Qwen3-Reranker-8B	69.02	77.45	72.94	70.19	81.22	

Table 1: Qwen3-Reranker reranking results as reported by the Qwen team [1] (re-ranking Qwen3-Embedding-0.6B top-100). Reproduced with attribution; not our measurements.

BitDelta compression ratios ($31.87\times$), and Byzantine-robustness tables. None of those described the deployed model or any real measurement; they were fabricated and have been removed. Qwen3-Reranker is a cross-encoder that outputs a scalar relevance score, not an embedding vector, so the notions of a fixed output dimension, delta-compression of its output, and median aggregation of node embeddings do not apply to it.

5 Discussion

5.1 When to Rerank

Reranking is worthwhile when first-stage recall is high but precision is not — i.e. the relevant documents are usually in the retrieved shortlist but not ranked at the top. Because each candidate costs one cross-encoder forward pass, the practical knobs are the shortlist size k and the model size (0.6B vs 4B vs 8B), trading latency for ranking quality. The upstream results above show the 4B and 8B rerankers are close on most retrieval subsets, so the smaller model is often the better latency choice.

5.2 Model-Size Selection

For Zen deployments we default to the 4B reranker as a balance of quality and serving cost, falling back to 0.6B for latency-critical paths and 8B where Chinese (CMTEB-R 77.45) or long-document (MLDR) quality is the priority, per the upstream numbers [1].

5.3 Future Work

1. **Tracking upstream:** Adopt new Qwen3-Reranker releases as the Qwen team ships them.

2. **Shortlist tuning:** Calibrate the retrieve-then-rerank shortlist size k per Zen workload.
3. **Distillation:** Investigate distilling the 8B reranker’s rankings into the 0.6B for cheaper serving (an option, not a claimed result).

6 Related Work

Two-stage retrieval: bi-encoder retrieval followed by cross-encoder reranking is the standard high-recall, high-precision pipeline; cross-encoders score query–document pairs jointly, as in Sentence-BERT [3].

Embedding and reranking models: BERT [2], Sentence-BERT [3], E5 [4], BGE / C-Pack [5], and the Qwen3 embedding/reranking series [1] that this work packages.

LLM-based rerankers: Qwen3-Reranker follows the line of using a generative LLM as a relevance judge, scoring a candidate by the probability it assigns to an affirmative answer [1].

7 Conclusion

Zen-Reranker is the reranking stage of Zen retrieval pipelines, implemented as a direct deployment of Alibaba’s Apache-2.0 Qwen3-Reranker models (0.6B/4B/8B) [1]. Qwen3-Reranker is a cross-encoder built on the Qwen3 dense foundation models that scores a (query, document) pair by the probability of a “yes” token; it supports 100+ languages and a 32K context. We contribute integration into the Zen retrieve-then-rerank stack, not a new model, and we report the Qwen team’s published reranking numbers with attribution rather than benchmarks of our own. The fictional “native 7680-dimensional,” DSO-optimized, and BitDelta-/Byzantine-related claims of earlier revisions have been removed as fabrications.

Acknowledgments and Attribution

Zen-Reranker is built entirely on Qwen3-Reranker (Qwen/Qwen3-Reranker, Apache-2.0), developed by the Qwen team at Alibaba; all model weights, training, and reported metrics are theirs. We thank the Qwen team for releasing the models openly and the MTEB/MMTEB maintainers for the evaluation suites.

References

- [1] Y. Zhang, M. Li, D. Long, X. Zhang, H. Lin, B. Yang, P. Xie, A. Yang, D. Liu, J. Lin, F. Huang, J. Zhou (Qwen Team, Alibaba). Qwen3 Embedding: Advancing Text Embedding and Reranking Through Foundation Models. arXiv:2506.05176, 2025. Models: `Qwen/Qwen3-Reranker-{0.6B,4B,8B}` (Apache-2.0).
- [2] Devlin, J., Chang, M. W., Lee, K., & Toutanova, K. BERT: Pre-training of deep bidirectional transformers for language understanding. NAACL, 2019.
- [3] Reimers, N., & Gurevych, I. Sentence-BERT: Sentence embeddings using Siamese BERT-networks. EMNLP, 2019.
- [4] Wang, L., Yang, N., Huang, X., et al. Text embeddings by weakly-supervised contrastive pre-training. arXiv:2212.03533, 2022.
- [5] Xiao, S., Liu, Z., Zhang, P., & Muennighoff, N. C-Pack: Packaged resources to advance general Chinese embedding. arXiv:2309.07597, 2023.