

Zen Embeddings: Dense Retrieval and Semantic Search

Technical Report v2025.06

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

June 2025

Abstract

We present Zen Embed, a family of text embedding models derived from the Qwen3-based Zen backbones (dense 0.6B/4B/8B/32B, Apache-2.0), optimized for dense retrieval, semantic search, and cross-lingual understanding. The embedding dimension follows the chosen backbone, and our models are natively compatible with BitDelta-style compression, which reduces index storage by roughly $8\times$ at minimal NDCG@10 degradation. We evaluate on the Massive Text Embedding Benchmark (MTEB), BEIR zero-shot retrieval, and MIRACL multilingual retrieval, reporting the *relative* contribution of each design choice (hard negative mining, compression, cross-lingual alignment) rather than headline leaderboard scores. We detail our bi-encoder training methodology, hard negative mining strategy, and cross-lingual alignment techniques.

1 Introduction

Dense retrieval has become foundational to modern AI systems: retrieval-augmented generation (RAG), semantic search, question answering, and recommendation systems all depend on high-quality embeddings that capture semantic meaning across diverse content types.

The core challenge in embedding model development is the tension between three objectives:

1. **Semantic fidelity:** Embeddings must capture nuanced meaning, distinguishing semantically similar from dissimilar pairs.
2. **Efficiency:** Production retrieval over billion-scale corpora demands compact representations and efficient search.
3. **Generalization:** Models must transfer to unseen domains, languages, and task types without fine-tuning.

Zen Embed addresses all three through a combination of Qwen3-based Zen backbone fine-tuning, novel hard negative mining, BitDelta-compatible compression, and a cross-lingual pre-alignment stage.

2 Model Architecture

2.1 Backbone and Pooling

Zen Embed uses a Qwen3-based Zen model as its backbone, adapted for embedding by enabling bidirectional attention across the full input sequence (rather than the causal masking used for generation). We extract sentence embeddings via weighted mean pooling:

$$\mathbf{e} = \frac{\sum_{i=1}^L w_i \cdot \mathbf{h}_i}{\sum_{i=1}^L w_i} \quad (1)$$

where $\mathbf{h}_i$ is the hidden state at position i and w_i is a learned scalar weight. The weights w_i are produced by a two-layer attention head over the final layer representations, allowing the model to emphasize semantically important tokens.

The final embedding dimension follows the backbone’s hidden dimension (e.g. 4096 for the 8B variant). Using the native hidden dimension simplifies the architecture (no projection layer); Matryoshka Representation Learning (Section 4) is used when a smaller embedding dimension is required for storage efficiency.

2.2 Instruction-Following Embedding

Zen Embed supports task-specific instruction prefixes following the Instructor pattern. For retrieval, the query encoder receives a prefix such as “Represent this query for searching relevant documents:” while the document encoder receives “Represent this document for retrieval:”. This single-model approach handles diverse embedding tasks:

- Semantic textual similarity (STS)
- Asymmetric retrieval (query vs. document)
- Classification via embedding
- Clustering
- Reranking

3 Training Methodology

3.1 Contrastive Objective

We train Zen Embed with the InfoNCE contrastive objective. Given a batch of N query-document pairs $\{(q_i, d_i^+)\}_{i=1}^N$ with in-batch negatives and K mined hard negatives $\{d_i^{(k)-}\}$:

$$\mathcal{L} = -\frac{1}{N} \sum_{i=1}^N \log \frac{e^{\text{sim}(q_i, d_i^+)/\tau}}{e^{\text{sim}(q_i, d_i^+)/\tau} + \sum_{j \neq i} e^{\text{sim}(q_i, d_j^+)/\tau} + \sum_k e^{\text{sim}(q_i, d_i^{(k)-})/\tau}} \quad (2)$$

where $\text{sim}(\mathbf{u}, \mathbf{v}) = \mathbf{u}^\top \mathbf{v} / (\|\mathbf{u}\| \|\mathbf{v}\|)$ is cosine similarity and $\tau = 0.02$ is the temperature.

3.2 Hard Negative Mining

Effective contrastive learning requires hard negatives—documents that are superficially similar to the query but semantically incorrect. We employ a three-stage mining pipeline:

Stage 1: BM25 Retrieval. For each query, retrieve the top-100 candidates by BM25. Documents not in the relevant set serve as hard negatives.

Stage 2: Embedding-based Re-mining. After initial training (1 epoch), use the current model to retrieve top-50 candidates. Negatives are sampled from positions 5–50 (avoiding the top-4 which may contain false negatives).

Stage 3: Cross-encoder Scoring. A fine-tuned cross-encoder assigns relevance scores to all candidates. We select negatives with cross-encoder score > 0.3 but < 0.7 as “hard but not false” negatives.

Table 1: Impact of progressively harder negatives on retrieval quality (BEIR, MS-MARCO Dev, FIQA). Each mining stage adds a further improvement over the previous one; values are relative (more arrows = higher quality).

Negative Source	BEIR Avg	MS-MARCO Dev	FIQA
In-batch only	↑	↑	↑
+ BM25 negatives	↑↑	↑↑	↑↑
+ Embedding re-mining	↑↑↑	↑↑↑	↑↑↑
+ Cross-encoder filtering	↑↑↑↑	↑↑↑↑	↑↑↑↑

3.3 Training Data

We curate a training corpus of 2.1 billion query-document pairs from diverse sources:

Table 2: Training data composition

Source	Pairs (millions)	Weight
MS-MARCO passages	532	0.18
Common Crawl pairs	812	0.24
Academic QA (NQ, TQA, HotpotQA)	124	0.12
Code search (GitHub)	98	0.08
Multilingual parallel corpora	418	0.28
Synthetic DSO pairs	116	0.10

Synthetic DSO (Decentralized Semantic Optimization) pairs are generated by a Qwen3-based Zen model using constitutional prompts, then filtered by cross-encoder quality scoring.

4 BitDelta-Compatible Embedding Compression

4.1 Problem and Approach

Storing full-precision float32 embeddings for a billion-document corpus requires storage and index memory proportional to $(\text{corpus size}) \times (\text{embedding dimension}) \times 4$ bytes, which reaches the multi-terabyte scale for a high-dimensional backbone. We develop BitDelta-compatible embedding compression that reduces this footprint by roughly $8\times$ at negligible quality loss.

The key insight is that production embedding corpora exhibit strong structure: most embeddings lie near the convex hull of a low-rank manifold. We exploit this via:

$$\tilde{\mathbf{e}} = \mathbf{e}_0 + \text{sign}(\mathbf{e} - \mathbf{e}_0) \cdot s \quad (3)$$

where $\mathbf{e}_0$ is a learned centroid vector (per-cluster), $\text{sign}(\cdot)$ is the bitwise sign quantization, and s is a learned scalar. This is the embedding analog of BitDelta weight quantization.

4.2 Reconstruction Quality

Table 3: Embedding compression: relative storage footprint and retrieval-quality retention for a billion-document index. “Storage” is normalized to full float32 = 1.0; quality is the fraction of full-precision retrieval quality retained (BEIR/MTEB/MS-MARCO).

Method	Relative storage	Quality retained
Full float32	1.0×	100%
float16	0.5×	≈100%
PQ-64	~0.08×	noticeably degraded
BitDelta-Embed	~0.13×	near-full
BitDelta-Embed (MRL)	~0.03×	near-full

Matryoshka Representation Learning (MRL) further reduces the embedding dimension, taking the combined footprint to roughly 1/30 of full float32 while retaining near-full retrieval quality—and, unlike product quantization at comparable storage, it does so with only minor quality loss.

5 Cross-Lingual Alignment

5.1 Alignment Pre-training

Before task-specific fine-tuning, we apply a cross-lingual alignment stage that ensures semantically equivalent sentences in different languages map to nearby embedding spaces. Given a parallel corpus $\{(s_i^{L_1}, s_i^{L_2})\}$, we minimize:

$$\mathcal{L}_{\text{align}} = \sum_i \left\| \mathbf{e}(s_i^{L_1}) - \mathbf{e}(s_i^{L_2}) \right\|_2^2 \quad (4)$$

Combined with contrastive objectives over multilingual batches, this produces embeddings where the same concept in any of the 100+ supported languages maps to a consistent region of the embedding space.

5.2 Language Coverage

On MIRACL multilingual retrieval (18 languages), retrieval quality is highest for the higher-resource Germanic and Romance language families, intermediate for CJK and Slavic, and lowest for the lowest-resource languages — the expected ordering given the relative amount of training data and parallel supervision per family. The cross-lingual alignment stage narrows, but does not eliminate, the gap between high- and low-resource languages.

6 MTEB Benchmark Results

Across the MTEB task categories (classification, clustering, pair classification, reranking, retrieval, STS, summarization), Zen Embed improves over a strong bi-encoder baseline of comparable size in every category, with the largest relative gains on classification, clustering, and retrieval. We report these as relative improvements over the baseline rather than as absolute leaderboard scores, since absolute MTEB numbers depend heavily on the chosen backbone size and evaluation configuration.

7 BEIR Zero-Shot Retrieval

BEIR evaluates retrieval across 18 heterogeneous datasets spanning diverse domains without any fine-tuning on target domains. Across these datasets, Zen Embed substantially outperforms the BM25 lexical baseline, with the largest margins on the Wikipedia-centric QA datasets (NQ, HotpotQA) and finance/biomedical domains (FiQA, TREC-COVID), and smaller margins on datasets where lexical overlap is already a strong signal (e.g. NFCorpus). The dense retriever’s advantage over BM25 is the headline result; its absolute NDCG@10 depends on the backbone and is therefore not restated here.

8 MS-MARCO Passage Ranking

On the MS-MARCO passage retrieval development set, the Zen Embed bi-encoder substantially improves MRR@10 and Recall@1000 over the BM25 baseline while retaining full zero-shot capability, and adding a cross-encoder reranking stage on top of the bi-encoder candidates yields a further MRR@10 improvement. The two-stage bi-encoder + cross-encoder pipeline is the recommended high-quality configuration.

9 Practical Deployment

9.1 Embedding Throughput

Embedding throughput scales with the usual factors: it is higher for shorter average sequence lengths and on newer accelerators, and it improves with batching. Smaller backbone variants (e.g. the 0.6B and 4B Zen models) embed substantially faster than the larger variants, trading some retrieval quality for throughput; the appropriate operating point depends on the corpus size and latency budget.

9.2 Indexing at Scale

For billion-scale document corpora, we recommend hierarchical navigable small world (HNSW) indexing with:

- An MRL-compressed embedding dimension for memory efficiency
- $M = 32$ connections, $ef_{\text{construction}} = 200$
- High approximate recall@10 at low P95 latency

10 Conclusion

Zen Embed delivers strong dense retrieval quality across English and multilingual benchmarks while remaining practically deployable through BitDelta-compatible compression. Built on the Qwen3-based Zen backbones (Apache-2.0), it improves over comparable bi-encoder baselines on MTEB and substantially outperforms BM25 on BEIR zero-shot retrieval, while BitDelta-style quantization plus MRL compression enable billion-scale corpora on commodity storage at roughly 1/30 the footprint of full-precision embeddings at near-full retrieval quality.

References

- [1] Muennighoff, N. et al. MTEB: Massive Text Embedding Benchmark. *EACL*, 2023.

- [2] Thakur, N. et al. BEIR: A Heterogeneous Benchmark for Zero-shot Evaluation of Information Retrieval Models. *NeurIPS*, 2021.
- [3] Zhang, X. et al. MIRACL: A Multilingual Retrieval Dataset Covering 18 Diverse Languages. *TACL*, 2023.
- [4] Kusupati, A. et al. Matryoshka Representation Learning. *NeurIPS*, 2022.
- [5] Su, H. et al. One Embedder, Any Task: Instruction-Finetuned Text Embeddings. *ACL Findings*, 2023.