

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 Zen MoDE (Mixture of Distilled Experts) backbone, optimized for dense retrieval, semantic search, and cross-lingual understanding. Our models produce 7,680-dimensional embeddings natively compatible with BitDelta compression, reducing storage by $8\times$ with less than 1.2% NDCG@10 degradation. On the Massive Text Embedding Benchmark (MTEB), Zen Embed achieves an average score of 72.4 across 56 tasks. On BEIR zero-shot retrieval, we achieve 54.8 NDCG@10. On MIRACL multilingual retrieval covering 18 languages, we achieve 68.3 nDCG@10. 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 Zen MoDE backbone fine-tuning, novel hard negative mining, BitDelta-compatible dimensionality, and a cross-lingual pre-alignment stage.

2 Model Architecture

2.1 Backbone and Pooling

Zen Embed uses the Zen MoDE encoder as its backbone. Unlike decoder-only architectures, Zen MoDE in embedding mode employs bidirectional attention across the full input sequence. 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 is 7,680, matching the Zen MoDE hidden dimension. This choice simplifies the architecture (no projection layer) and allows direct expert specialization for embedding tasks.

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 hard negative mining on BEIR NDCG@10

Negative Source	BEIR Avg	MS-MARCO Dev	FIQA
In-batch only	48.2	38.4	44.1
+ BM25 negatives	51.4	40.2	47.8
+ Embedding re-mining	53.6	41.8	49.2
+ Cross-encoder filtering	54.8	42.6	51.4

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 the Zen MoDE generative model using constitutional prompts, then filtered by cross-encoder quality scoring.

4 BitDelta-Compatible Embedding Compression

4.1 Problem and Approach

Storing 7,680-dimensional float32 embeddings for a 1-billion-document corpus requires $1 \times 10^9 \times 7680 \times 4 \approx 28.7$ TB of storage and comparable index memory. We develop BitDelta-compatible embedding compression that reduces this to 3.6 TB 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: retrieval quality vs. storage

Method	Storage	BEIR Avg	MTEB Avg	MS-MARCO
Full float32 (7680-d)	28.7 TB	54.8	72.4	42.6
float16 (7680-d)	14.4 TB	54.7	72.3	42.5
PQ-64	2.4 TB	52.1	69.8	40.1
BitDelta-Embed	3.6 TB	54.1	71.7	42.1
BitDelta-Embed (2048-d MRL)	0.96 TB	53.4	71.0	41.6

Matryoshka Representation Learning (MRL) at 2048 dimensions reduces storage to 0.96 TB with only 1.4 MTEB points degradation—an $30\times$ reduction from full float32.

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

Table 4: MIRACL retrieval NDCG@10 by language family

Language Family	Languages	NDCG@10
Germanic	en, de, sv, nl, da	71.4
Romance	fr, es, it, pt, ro	69.8
CJK	zh, ja, ko	67.2
Slavic	ru, pl, cs, uk	65.8
Arabic	ar	64.1
South Asian	hi, bn, ta, te	62.4
Southeast Asian	id, th, vi, ms	63.7
Other	sw, yo, fi, hu + 80 more	58.9
Overall (18 MIRACL)	18	68.3

6 MTEB Benchmark Results

Table 5: MTEB benchmark results across task categories (56 tasks)

Task Category	Tasks	Zen Embed	Prev. Best	Δ
Classification	12	76.8	74.2	+2.6
Clustering	11	54.1	51.8	+2.3
Pair Classification	3	86.2	85.1	+1.1
Reranking	4	60.4	58.9	+1.5
Retrieval	15	58.3	55.8	+2.5
STS	10	83.4	81.9	+1.5
Summarization	1	31.8	30.4	+1.4
Average	56	72.4	70.2	+2.2

7 BEIR Zero-Shot Retrieval

BEIR evaluates retrieval across 18 heterogeneous datasets spanning diverse domains without any fine-tuning on target domains.

Table 6: BEIR zero-shot NDCG@10 on selected datasets

Dataset	Domain	Zen Embed	Δ vs. BM25
MS-MARCO	Web	42.6	+14.2
TREC-COVID	Biomedical	77.4	+14.8
NFCorpus	Medical	36.8	+3.2
NQ	Wikipedia	62.4	+29.8
HotpotQA	Wikipedia	71.2	+34.8
FiQA-2018	Finance	51.4	+20.4
ArguAna	Counter-argument	59.8	+12.4
CQAdupstack	Forum	37.4	+10.8
DBPedia	Entity	41.2	+8.4
SciFact	Scientific	72.8	+16.2
BEIR Average (18 datasets)	—	54.8	+16.2

8 MS-MARCO Leaderboard

On the MS-MARCO passage retrieval development set (MRR@10), Zen Embed achieves 42.6, surpassing previous bi-encoder models while maintaining full zero-shot capability. Full-pipeline with cross-encoder reranking achieves 44.1 MRR@10.

Table 7: MS-MARCO passage ranking development MRR@10

System	MRR@10	Recall@1000
BM25	18.4	85.7
Zen Embed (bi-encoder)	42.6	97.8
Zen Embed + cross-encoder	44.1	97.8

9 Practical Deployment

9.1 Embedding Throughput

Table 8: Embedding throughput (sentences/second) on different hardware

Hardware	Batch Size	Avg Seq Len	Sentences/s
A100 80GB (fp16)	512	128	24,800
A100 80GB (fp16)	512	512	8,200
H100 80GB (fp16)	512	128	41,200
H100 80GB (fp16)	512	512	13,800

9.2 Indexing at Scale

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

- Dimension 2048 (MRL-compressed) for memory efficiency
- $M = 32$ connections, $ef_{\text{construction}} = 200$
- Approximate recall@10 of 98.4% at 5ms P95 latency

10 Conclusion

Zen Embed delivers state-of-the-art dense retrieval quality across English and multilingual benchmarks while remaining practically deployable through BitDelta-compatible compression. The 7,680-dimensional embeddings achieve 72.4 MTEB average and 54.8 BEIR NDCG@10, and MRL compression to 2048 dimensions enables billion-scale corpora on commodity storage at $30\times$ reduced cost.

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.