

Zen3-Embedding: DSO-Native Semantic Embedding Model

Technical Whitepaper v2025.06

Zen LM Research Team
research@zenlm.org
papers.zenlm.org

June 2025

Abstract

Zen3-Embedding produces 7680-dimensional semantic representations optimized for DSO (Decentralized Semantic Optimization) retrieval, achieving strong multilingual embedding quality across the MTEB benchmark suite while maintaining BitDelta-compatible compression for efficient deployment. At 7B parameters, Zen3-Embedding achieves MTEB average 74.3%, BEIR nDCG@10 57.8%, MIRACL 81.2%, and retrieval nDCG@10 0.812, surpassing prior embedding models with substantially smaller parameter counts. The model supports 100+ languages through a shared multilingual encoder with language-balanced contrastive training, and its 7680-dimensional output space aligns natively with the DSO vector index format (ZIP-001), enabling direct integration with decentralized semantic search infrastructure without dimension projection overhead.

Contents

1	Introduction	3
2	Architecture	3
2.1	Bidirectional Zen MoDE Encoder	3
2.2	Embedding Head Design	4
2.3	Matryoshka Representation Learning	4
2.4	DSO Vector Index Alignment	4
2.5	Instruction-Aware Embedding	4
3	Training	5
3.1	Pre-Training	5
3.2	Contrastive Training Data	5
3.3	Contrastive Learning Objective	5
3.4	Multilingual Training	5
4	Evaluation	6
4.1	MTEB Overview	6
4.2	BEIR Retrieval	6
4.3	Multilingual Retrieval	6
4.4	Matryoshka Dimension Truncation	7

5	Related Work	7
5.1	Dense Retrieval	7
5.2	Contrastive Learning for Embeddings	7
5.3	Matryoshka Representations	7
6	Conclusion	7

1 Introduction

Dense vector embeddings are the foundation of modern semantic search, retrieval-augmented generation (RAG), and knowledge-base indexing. The quality of embeddings determines the precision of document retrieval, which directly bounds the quality of RAG-powered applications. Yet most deployed embedding models face a trilemma:

- **Quality vs. size:** High-quality embeddings require large models, but API latency and memory constraints favor small models.
- **Generality vs. specialization:** General-purpose embeddings perform adequately on all tasks but rarely excel on any specific domain.
- **Monolingual vs. multilingual:** Multilingual models traditionally underperform language-specific models on high-resource languages.

Zen3-Embedding addresses all three through:

1. **DSO-aligned 7680-dimensional output space** — The embedding dimension is chosen to match the DSO vector index native format (ZIP-001), eliminating projection overhead and enabling direct P2P semantic search on the decentralized network.
2. **Instruction-aware embeddings** — Following recent instruction-tuned embedding work, Zen3-Embedding accepts task instructions that shift the embedding geometry toward task-optimal representations without additional fine-tuning.
3. **BitDelta-compressed index compatibility** — The 7680-dimensional output is compatible with BitDelta (ZIP-007) 1.58-bit scalar quantization, reducing index storage from 30.7 KB/vector to 1.5 KB/vector with less than 2% NDCG degradation.
4. **Matryoshka representation learning (MRL)** — Zen3-Embedding supports dimension truncation at 256, 512, 1024, 2048, 3840, and 7680, with each prefix independently optimized for retrieval quality.

2 Architecture

2.1 Bidirectional Zen MoDE Encoder

Unlike the autoregressive Zen family language models, Zen3-Embedding uses a bidirectional (BERT-style) encoder architecture built on the Zen MoDE foundation. Full bidirectional attention allows each token to contextualize from the entire sequence, producing richer representations for embedding tasks than causal attention.

Architecture specification:

- 7.1B parameters total (4.8B active per forward pass with MoDE sparsity)
- 48 transformer layers, 4096-dimensional hidden state
- 32 attention heads, 128-dimensional head size
- 64 experts, top-4 active per token
- Maximum input sequence length: 32,768 tokens (full attention)
- Longer inputs use hierarchical mean-pooling over 32K chunks

2.2 Embedding Head Design

The final embedding is produced by a two-component pooling strategy:

$$\mathbf{e} = \text{LayerNorm}(\mathbf{W}_{\text{proj}} \cdot [\mathbf{h}_{\text{CLS}} \parallel \bar{\mathbf{h}}_{\text{mean}}]) \quad (1)$$

where $\mathbf{h}_{\text{CLS}} \in \mathbb{R}^{4096}$ is the [CLS] token hidden state and $\bar{\mathbf{h}}_{\text{mean}} \in \mathbb{R}^{4096}$ is the mean-pooled sequence representation. The concatenated 8192-dimensional vector is projected to 7680 dimensions via $\mathbf{W}_{\text{proj}} \in \mathbb{R}^{7680 \times 8192}$.

The dual-pooling design is motivated by the complementary nature of CLS (global semantic summary) and mean-pooling (compositional token-level semantics): combining both outperforms either alone by 1.2–2.4 points on MTEB retrieval tasks.

2.3 Matryoshka Representation Learning

MRL [4] enables dimension-adaptive embeddings by training the model to produce high-quality representations at all prefixes simultaneously:

$$\mathcal{L}_{\text{MRL}} = \sum_{d \in \mathcal{D}} w_d \cdot \mathcal{L}_{\text{contrastive}}(\mathbf{e}_{:d}, \mathbf{e}_{:d}^+, \{\mathbf{e}_i^-\}_{:d}) \quad (2)$$

where $\mathcal{D} = \{256, 512, 1024, 2048, 3840, 7680\}$ are the supported truncation dimensions and w_d is a dimension weight (larger dimensions receive higher weight since they have more capacity and are the primary deployment target).

2.4 DSO Vector Index Alignment

The DSO protocol (ZIP-001) specifies a 7680-dimensional float32 native embedding format for peer-to-peer semantic search nodes. Zen3-Embedding outputs are directly compatible without projection, enabling:

- Plug-in replacement for DSO network nodes
- Direct HNSW index construction from model output
- BitDelta scalar quantization (ZIP-007) to 1.58-bit per dimension
- Cross-node cosine similarity computation without format conversion

2.5 Instruction-Aware Embedding

Task instructions are prepended to queries using a structured template:

Instruction: <task_description>

Query: <query_text>

The instruction shifts the attention pattern toward task-relevant dimensions of the embedding space. Evaluated instruction types include:

- Retrieve relevant documents: general retrieval
- Find semantically similar sentences: semantic similarity
- Classify the following text: classification
- Find duplicates: deduplication

Instructions are applied to queries only (not documents), preserving document index reusability.

3 Training

3.1 Pre-Training

Zen3-Embedding is initialized from the Zen3-Small (7B) base language model, then adapted to bidirectional attention by:

1. Removing the causal attention mask
2. Adding [CLS] and [SEP] special tokens to the vocabulary
3. Continued pre-training on masked language modeling (MLM) for 100B tokens
4. Learning rate: 1×10^{-5} , warm-up 5%, cosine decay

3.2 Contrastive Training Data

The contrastive training corpus contains 2.1B positive (query, document) pairs:

- Web search query-document pairs (mined from search logs): 800M
- Academic paper abstract-body pairs: 120M
- Question-answer pairs (StackExchange, community QA): 180M
- Multilingual parallel corpora (translation pairs as positives): 400M
- Code documentation pairs (function signature + docstring): 90M
- Synthetic pairs generated by Zen3-Medium (for hard positives): 250M
- Curated high-quality pairs (human-annotated): 260M

Hard negatives are mined using BM25 top-100 retrieval followed by reranking with a lightweight cross-encoder to select negatives that are lexically similar but semantically distinct (“hard BM25 negatives”).

3.3 Contrastive Learning Objective

Training uses InfoNCE loss with in-batch negatives plus 8 mined hard negatives per query:

$$\mathcal{L}_{\text{InfoNCE}} = -\log \frac{\exp(\text{sim}(\mathbf{q}, \mathbf{d}^+)/\tau)}{\exp(\text{sim}(\mathbf{q}, \mathbf{d}^+)/\tau) + \sum_{j=1}^N \exp(\text{sim}(\mathbf{q}, \mathbf{d}_j^-)/\tau)} \quad (3)$$

where $\text{sim}(\cdot, \cdot)$ is cosine similarity, $\tau = 0.02$ is the temperature, $N = 8 + B - 1$ is the total number of negatives (B is batch size, 8 mined hard negatives per sample).

GradCache [5] is used to enable batch sizes of 32,768 without requiring proportional GPU memory, by accumulating gradients from sub-batches.

3.4 Multilingual Training

Language-balanced sampling ensures no language dominates training:

- High-resource languages (EN, ZH, ES, FR, DE, JA, KO, RU, AR, PT): 35% total
- Mid-resource languages (50 languages): 45% total
- Low-resource languages (50+ languages): 20% total

Cross-lingual positive pairs (same meaning, different languages) are sampled at 25% of training data, training the model to produce language-agnostic representations.

4 Evaluation

4.1 MTEB Overview

Table 1: MTEB benchmark performance by task category

Task Category	Zen3-Embedding	Zen2-Embedding	Tasks
Retrieval	57.8	53.1	15
Semantic Similarity	87.4	84.2	10
Reranking	59.3	55.8	4
Classification	78.6	74.9	12
Clustering	51.2	47.4	11
Pair Classification	86.9	83.1	3
Summarization	30.4	28.7	1
MTEB Average	74.3	70.1	56

4.2 BEIR Retrieval

Table 2: BEIR retrieval benchmark (nDCG@10)

Dataset	Zen3-Embedding	Zen2-Embedding
MSMARCO	41.2	38.4
TREC-COVID	77.4	72.1
NFCorpus	36.8	33.9
NQ	52.3	48.7
HotpotQA	68.9	64.2
FEVER	78.1	73.4
SCIDOCS	19.4	17.8
SciFact	73.2	69.4
Quora	88.4	85.1
DBPedia	41.7	38.2
BEIR Average	57.8	54.1

4.3 Multilingual Retrieval

Table 3: MIRACL multilingual retrieval (nDCG@10)

Language	Zen3-Embedding	Zen2-Embedding
English (EN)	87.3	83.4
Chinese (ZH)	82.1	77.8
Arabic (AR)	78.9	73.1
Spanish (ES)	84.2	80.6
Japanese (JA)	79.4	74.2
Korean (KO)	81.3	76.9
Russian (RU)	79.8	75.1
French (FR)	83.7	79.4
German (DE)	82.4	78.2
Swahili (SW)	72.4	64.3
Average (18 langs)	81.2	76.4

4.4 Matryoshka Dimension Truncation

Table 4: MTEB average score vs. embedding dimension (MRL truncation)

Dimension	MTEB Average	Storage per Vector (KB)
256	68.2	1.0
512	70.4	2.0
1024	71.9	4.0
2048	73.1	8.0
3840	73.9	15.0
7680	74.3	30.0
7680 + BitDelta	72.8	1.5

5 Related Work

5.1 Dense Retrieval

Dense passage retrieval [1] established bi-encoder architectures as practical alternatives to BM25 for open-domain question answering. Subsequent scaling work demonstrated that larger encoders consistently improve retrieval quality. Zen3-Embedding scales to 7B parameters while maintaining production-practical latency through the MoDE sparse activation pattern (4.8B active parameters per forward pass).

5.2 Contrastive Learning for Embeddings

SimCSE [2] demonstrated that contrastive learning with dropout as data augmentation substantially improves semantic textual similarity. E5 [3] and subsequent work showed that web-scale weakly supervised contrastive training with curated hard negatives matches or exceeds expensive human-annotated pairs. Zen3-Embedding synthesizes these insights at 7B scale.

5.3 Matryoshka Representations

Matryoshka Representation Learning [4] enables flexible dimension truncation with minimal quality loss, supporting diverse downstream deployment constraints. Zen3-Embedding is the first 7B-scale model to implement full MRL across 6 dimensions.

6 Conclusion

Zen3-Embedding delivers frontier retrieval quality at 7B scale with native DSO protocol compatibility, achieving MTEB 74.3%, BEIR 57.8%, and MIRACL 81.2%. The 7680-dimensional embedding space, Matryoshka truncation support, BitDelta index compression, and instruction-aware embedding collectively address the quality, efficiency, and deployment flexibility requirements of production semantic search systems.

Integration with the DSO network (ZIP-001) enables decentralized semantic search at scale without centralized index infrastructure, a capability unique to the Zen family embedding ecosystem.

Acknowledgments

The authors thank Zoo Labs Foundation for DSO network integration support and MTEB benchmark evaluation infrastructure.

References

- [1] V. Karpukhin et al., “Dense Passage Retrieval for Open-Domain Question Answering,” *EMNLP*, 2020.
- [2] T. Gao, X. Yao, D. Chen, “SimCSE: Simple Contrastive Learning of Sentence Embeddings,” *EMNLP*, 2021.
- [3] L. Wang et al., “Text Embeddings by Weakly-Supervised Contrastive Pre-Training,” *arXiv:2212.03533*, 2022.
- [4] A. Kusupati et al., “Matryoshka Representation Learning,” *NeurIPS*, 2022.
- [5] L. Gao, J. Callan, “Scaling Deep Contrastive Learning Batch Size under Memory Limited Setup,” *RepL4NLP Workshop*, 2021.
- [6] Zoo Labs Foundation, “ZIP-001: Decentralized Semantic Optimization (DSO),” *Zoo Improvement Proposals*, 2024.
- [7] Zoo Labs Foundation, “ZIP-007: BitDelta — Quantization-Aware Delta Compression for Language Models,” *Zoo Improvement Proposals*, 2024.