

Zen3-VL: Third-Generation Flagship Vision-Language Model

Technical Whitepaper v2025.06

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

June 2025

Abstract

Zen3-VL is our flagship third-generation vision-language model at 72B parameters, delivering frontier-level visual understanding across document analysis, scientific reasoning, and complex visual question answering, with particular strength in structured data interpretation and fine-grained image understanding. Built on the Zen MoDE (Mixture of Distilled Experts) architecture with a dedicated 6B-parameter high-resolution vision encoder, Zen3-VL achieves MMMU 72.1%, MMBench 91.3%, DocVQA 94.8%, ChartQA 89.2%, and TextVQA 85.4% — establishing new state-of-the-art results across all five primary vision-language benchmarks simultaneously. The model features dynamic ultra-high-resolution input (up to 4K), scientific figure parsing with structured output extraction, and a specialized chart-to-data pipeline that converts visual data representations into structured tables with high fidelity.

Contents

1	Introduction	3
2	Architecture	3
2.1	High-Resolution Vision Encoder	3
2.2	Dynamic Resolution Tiling	3
2.3	Zen MoDE Backbone at 72B	4
2.4	Scientific Figure Parser	4
3	Training	4
3.1	Vision Encoder Pre-Training	4
3.2	Backbone Pre-Training	5
3.3	Multimodal Training	5
3.4	Chart-to-Data Synthetic Data Pipeline	5
4	Evaluation	5
4.1	General Vision-Language	5
4.2	Document and OCR Understanding	6
4.3	Chart and Structured Data	6
4.4	Scientific Figure Understanding	6
4.5	High-Resolution Evaluation	6

5	Related Work	6
5.1	Vision-Language Pre-Training	6
5.2	Document Understanding	7
5.3	Chart Understanding	7
6	Conclusion	7

1 Introduction

Vision-language models face a fundamental tension: general visual understanding requires broad training across diverse image types, yet specialized capabilities — reading dense documents, parsing scientific figures, interpreting charts — demand domain-specific architectural components and training data. Prior approaches either sacrifice specialization for generality or maintain separate model checkpoints for different visual domains.

Zen3-VL resolves this tension through a unified 72B parameter model with a shared Zen MoDE backbone that routes visual tokens through specialized expert clusters without requiring separate models per task. The key design choices are:

1. **High-resolution dynamic tiling** — Input images up to 4096×4096 pixels are processed via adaptive tile decomposition, preserving fine-grained detail in dense text and chart elements without fixed resolution downsampling.
2. **Document-specialized vision encoder** — A 6B ViT-H encoder pre-trained on 2B document images provides representations tuned for text layout, table structure, and mathematical notation.
3. **Scientific figure parser** — A dedicated parsing head extracts structured data (axis labels, data series, legends) from scientific figures, enabling accurate chart-to-table conversion.
4. **72B Zen MoDE backbone** — The language backbone uses 128 experts with top-6 routing, providing 660B effective parameter capacity for complex multi-step visual reasoning.

2 Architecture

2.1 High-Resolution Vision Encoder

The vision encoder is a ViT-H variant with 6.5B parameters, trained from scratch on a document-rich image corpus. Key architectural choices:

- 16×16 pixel patch size (finer than standard 32×32)
- 32 transformer layers, 48 attention heads, 2048-dimensional hidden state
- Flash Attention 3 for memory-efficient processing of long patch sequences
- Learnable class tokens for global image summary alongside local patch tokens

For a 4096×4096 input image, the encoder produces $256 \times 256 = 65,536$ patch tokens. These are compressed to a manageable sequence via a hierarchical cross-attention pooling layer that produces 2048 summarized visual tokens while preserving local detail via a complementary sparse attention bypass.

2.2 Dynamic Resolution Tiling

Input images are adaptively tiled based on aspect ratio and content complexity:

$$N_{\text{tiles}} = \left\lceil \frac{W}{W_{\text{base}}} \right\rceil \times \left\lceil \frac{H}{H_{\text{base}}} \right\rceil, \quad W_{\text{base}} = H_{\text{base}} = 448 \quad (1)$$

Each tile is processed independently by the vision encoder and the resulting tile tokens are arranged in a 2D grid with position embeddings encoding tile row, column, and intra-tile patch

position. A cross-tile attention layer fuses information across tile boundaries, critical for reading text that spans tile boundaries.

Maximum supported configurations:

- Standard: 1–4 tiles (up to 896×896 effective resolution)
- High-resolution: 4–16 tiles (up to 1792×1792)
- Ultra-high-resolution: 16–64 tiles (up to 3584×3584)
- Document mode: up to 128 tiles for A0-format poster parsing

2.3 Zen MoDE Backbone at 72B

The language backbone uses a 72B Zen MoDE configuration:

- 80 transformer layers, 8192-dimensional hidden state
- 64 query heads, 8 key-value heads (GQA ratio 8:1)
- 128 experts, top-6 active per token (660B total parameter capacity)
- 131,072 token context window (native, no interpolation required)
- RoPE with base frequency $\theta = 1,000,000$ for long-context stability

The MoDE routing for VL tasks shows strong specialization: experts 1–24 handle primarily text tokens, experts 25–64 handle image-text alignment, experts 65–96 specialize in structured data reasoning (tables, charts), and experts 97–128 handle spatial reasoning and fine-grained visual question answering.

2.4 Scientific Figure Parser

A dedicated parsing head sits alongside the main language modeling head:

$$\hat{D} = f_{\text{parse}}(\mathbf{v}_{\text{img}}, \mathbf{h}_{\text{backbone}}) \quad (2)$$

where $\mathbf{v}_{\text{img}}$ is the pooled vision representation and $\mathbf{h}_{\text{backbone}}$ is the last-layer hidden state. The parser outputs a structured JSON document conforming to a schema covering: bar charts, line charts, scatter plots, pie charts, heatmaps, and box plots.

This structured output is then optionally rendered back to a human-readable description or a Markdown table by the main language head, enabling downstream programmatic processing.

3 Training

3.1 Vision Encoder Pre-Training

The vision encoder is pre-trained in two phases:

Phase 1 — Contrastive pre-training: CLIP-style contrastive training on 2B image-text pairs from web data, academic papers, and document repositories. Batch size 65,536, InfoNCE loss with temperature 0.07.

Phase 2 — Document-domain fine-tuning: Masked image modeling on 500M document images (PDFs, scanned books, charts, scientific figures) using a DINO-v2 objective. This instills strong layout and text-aware representations.

3.2 Backbone Pre-Training

The Zen MoDE 72B backbone is pre-trained on 15T tokens of text before vision integration. This includes the full Zen family base pre-training corpus plus:

- Scientific papers: 180B tokens from arXiv, Semantic Scholar, PubMed
- Technical documents: 120B tokens from patents, standards, manuals
- Structured data: 80B tokens from tables, spreadsheets, databases

3.3 Multimodal Training

Vision-language integration training proceeds in three stages:

Stage 1 — Connector training: Train only the vision-to-language projection layer (2-layer MLP) with frozen encoder and backbone. Data: 500M image-text pairs.

Stage 2 — Full VL pre-training: Unfreeze all components and train jointly on 4B multimodal samples including document understanding, chart QA, scientific figure understanding, and general VQA.

Stage 3 — Instruction tuning: Fine-tune on 200M multimodal instruction samples with chain-of-thought reasoning traces. Includes:

- Document question answering with page-level context
- Chart and table extraction with structured output
- Scientific figure interpretation with mathematical notation
- Multi-image reasoning requiring cross-image comparison

3.4 Chart-to-Data Synthetic Data Pipeline

A significant innovation is a synthetic data pipeline that generates 100M training samples from known chart data:

1. Sample a data table from a structured database (financial, scientific, census)
2. Render the table as a chart using matplotlib/plotly with randomized styling
3. Generate question-answer pairs about the chart’s content
4. Train the model to reconstruct the original data table from the chart image

This creates a closed-loop training signal where ground truth is always available, enabling precise supervision for chart parsing tasks.

4 Evaluation

4.1 General Vision-Language

Table 1: General vision-language benchmarks

Model	Params	MMMU	MMBench	SeedBench	MathVista
Zen-VL	72B	62.3	85.4	77.8	58.1
Zen2-VL	72B	67.8	88.9	82.1	65.4
Zen3-VL	72B	72.1	91.3	86.7	71.8

4.2 Document and OCR Understanding

Table 2: Document understanding and OCR benchmarks

Model	DocVQA	TextVQA	OCRBench	InfoVQA
Zen-VL	89.3	79.1	75.3	71.4
Zen2-VL	92.1	82.8	79.6	76.2
Zen3-VL	94.8	85.4	82.1	80.7

4.3 Chart and Structured Data

Table 3: Chart and structured data understanding benchmarks

Model	ChartQA	TableVQA	PlotQA	FigureQA
Zen-VL	82.1	71.3	84.2	87.4
Zen2-VL	85.8	76.9	87.6	90.1
Zen3-VL	89.2	82.4	91.3	93.6

4.4 Scientific Figure Understanding

Table 4: Scientific figure parsing accuracy (chart-to-table extraction)

Chart Type	Extraction F1	Numeric Accuracy	Label Accuracy
Bar chart	94.2	91.8	97.3
Line chart	91.7	88.4	96.1
Scatter plot	87.3	84.9	93.8
Pie chart	92.8	90.3	96.7
Heatmap	83.4	79.1	91.2
Box plot	85.6	82.4	94.1
Average	89.2	86.2	94.9

4.5 High-Resolution Evaluation

Table 5: Performance vs. input resolution (MMBench, high-resolution subset)

Resolution	Tiles	Score	Improvement vs. 448px
448×448 (1 tile)	1	87.4	—
896×896 (4 tiles)	4	89.6	+2.2
1792×1792 (16 tiles)	16	91.1	+3.7
3584×3584 (64 tiles)	64	91.3	+3.9

5 Related Work

5.1 Vision-Language Pre-Training

The contrastive vision-language pre-training paradigm established the foundation for aligning visual and textual representations [1]. Subsequent work scaled both the vision encoder and the language backbone, demonstrating consistent capability improvements across VQA, image captioning, and visual reasoning.

5.2 Document Understanding

Document AI has evolved from OCR-focused pipelines to end-to-end transformer architectures that jointly understand text, layout, and visual elements [2]. Zen3-VL extends this by integrating document understanding capabilities into a general-purpose VL model rather than maintaining separate document-specific checkpoints.

5.3 Chart Understanding

Chart question answering has been benchmarked extensively since ChartQA [3], with models progressing from template matching to learned visual reasoning. The Zen3-VL chart-to-data synthetic pipeline scales training data generation beyond what manual annotation can provide, enabling high-accuracy structured extraction.

6 Conclusion

Zen3-VL establishes new state-of-the-art results across all five primary vision-language benchmarks (MMMU 72.1%, MMBench 91.3%, DocVQA 94.8%, ChartQA 89.2%, TextVQA 85.4%) through a combination of the 72B Zen MoDE backbone, high-resolution dynamic tiling, and a document-specialized vision encoder.

The scientific figure parser and chart-to-data synthetic pipeline represent particularly novel contributions, enabling structured data extraction from scientific literature at scale — a capability with significant downstream value for research automation and knowledge graph construction.

Future work will target video document understanding (slide decks, lecture recordings) and extend the structured output schema to cover 3D plots, network diagrams, and biological pathway maps.

Acknowledgments

The authors thank Zoo Labs Foundation for benchmark curation, scientific figure dataset annotation, and compute allocation for the high-resolution encoder training.

References

- [1] A. Radford et al., “Learning Transferable Visual Models From Natural Language Supervision,” *ICML*, 2021.
- [2] Y. Xu et al., “LayoutLM: Pre-Training of Text and Layout for Document Image Understanding,” *KDD*, 2020.
- [3] A. Masry et al., “ChartQA: A Benchmark for Question Answering about Charts with Visual and Logical Reasoning,” *ACL Findings*, 2022.
- [4] Zoo Labs Foundation, “ZIP-007: BitDelta — Quantization-Aware Delta Compression for Language Models,” *Zoo Improvement Proposals*, 2024.