

Zen-VL: Vision-Language Understanding at 7B Scale

Technical Report v2025.01

Zen LM Research Team
research@zenlm.org

January 2025

Abstract

We present **Zen-VL**, a 7 billion parameter vision-language model built on the Zen MoDE (Mixture of Distilled Experts) language backbone, extended with a high-resolution vision encoder and a dynamic tiling mechanism to support images up to 4096×4096 pixels. Zen-VL processes multiple images per conversation turn, enabling document analysis, multi-image comparison, and chart understanding. Without task-specific fine-tuning, Zen-VL achieves strong performance on standard multimodal benchmarks: MMBench 81.3%, MME 2148, MMMU 56.8%, TextVQA 78.4%, and DocVQA 88.2%. The model’s efficient 7B scale makes it suitable for deployment across cloud and on-premise environments requiring multimodal understanding alongside text generation.

Contents

1	Introduction	2
2	Architecture	2
2.1	System Overview	2
2.2	Vision Encoder	2
2.3	Dynamic High-Resolution Tiling	3
2.4	Visual-Language Projector	3
2.5	Language Backbone	3
3	Training Methodology	3
3.1	Stage 1: Visual Pretraining (Projector Alignment)	3
3.2	Stage 2: Visual Instruction Tuning	4
4	Evaluation	4
4.1	Multimodal Benchmarks	4
4.2	Document Understanding	5
4.3	High-Resolution Ablation	5
4.4	Multi-Image Evaluation	5
4.5	Hallucination	5
5	Related Work	6
6	Limitations	6
7	Conclusion	6

1 Introduction

Visual understanding integrated with language generation has emerged as a critical capability for real-world AI applications: reading charts in financial documents, analyzing scientific figures, answering questions about photographs, extracting structured data from scanned forms. Vision-language models (VLMs) [1, 2, 3] address this by combining a pretrained vision encoder with a language model backbone.

Zen-VL builds on two foundations: (1) the Zen language backbone trained on 3T tokens, and (2) a large vision encoder pretrained via contrastive learning on image-text pairs. Key design decisions include:

- **Dynamic tiling:** High-resolution images are decomposed into variable-size tiles, allowing fine-grained perception of text within images up to 4096×4096 .
- **Multi-image support:** The conversation schema supports arbitrary numbers of images per turn, with interleaved image and text tokens, enabling document-level and multi-panel analysis.
- **Efficient visual projection:** A lightweight MLP projector maps visual tokens to the language model’s embedding space, minimizing added parameters while preserving spatial information.
- **Instruction tuning on multimodal data:** Zen-VL is fine-tuned on 2M high-quality image-text instruction pairs spanning VQA, captioning, OCR, and document understanding.

2 Architecture

2.1 System Overview

Zen-VL consists of three components: (1) a vision encoder E_v , (2) a visual-language projector P , and (3) the Zen MoDE language decoder D . The forward pass is:

$$y = D([P(E_v(I_1)), t_1, P(E_v(I_2)), t_2, \dots]) \quad (1)$$

where I_k are image inputs, t_k are text tokens, and the brackets denote interleaved concatenation in conversation order.

Table 1: Zen-VL component specifications.

Component	Specification	Parameters
Vision encoder	ViT-L/14 at 448×448 native	307M
Visual projector	2-layer MLP, hidden 4096	28M
Language backbone	Zen MoDE 7B	7.2B
Total		7.5B

2.2 Vision Encoder

The vision encoder is a Vision Transformer (ViT) [4] at the ViT-L/14 scale, pretrained on 2B image-text pairs using a CLIP-style contrastive objective [1]. The native resolution of 448×448 pixels produces a $(448/14)^2 = 1024$ -token visual representation per tile in the patching scheme.

2.3 Dynamic High-Resolution Tiling

Standard VLMs encode images at a fixed low resolution (e.g., 224×224), losing fine detail needed for document understanding and dense text reading. Zen-VL implements a dynamic tiling scheme that:

1. Selects the optimal tile grid $\{r \times c\}$ for the input image’s aspect ratio, constrained to a maximum of $N_{\max} = 12$ tiles plus one mandatory thumbnail.
2. Resizes the image to fill the selected grid at 448×448 per tile.
3. Encodes each tile independently through the vision encoder.
4. Encodes a downsampled thumbnail (448×448) of the full image for global context.
5. Concatenates tile tokens with the thumbnail and passes all through the projector.

Table 2: Dynamic tiling configurations by image resolution.

Input Resolution	Grid	Tiles	Visual tokens
448×448 (standard)	1×1	1	$1,024 + 256$
896×448 (landscape)	2×1	2	$2,048 + 256$
$1,344 \times 1,344$	3×3	9	$9,216 + 256$
$4,096 \times 4,096$ (max)	3×4	12	$12,288 + 256$

Thumbnail tokens are downsampled to 256 by applying 2×2 average pooling over the ViT output before projection.

2.4 Visual-Language Projector

The projector is a two-layer MLP with GELU activation:

$$P(v) = W_2 \cdot \text{GELU}(W_1 v + b_1) + b_2 \quad (2)$$

where $v \in \mathbb{R}^{1024}$ is the per-token ViT output and $P(v) \in \mathbb{R}^{4096}$ matches the Zen backbone’s hidden dimension. The projector is trained jointly during the visual instruction tuning phase.

2.5 Language Backbone

The language component is the Zen 7B model from the Zen MoDE family, initialized from the pretrained Zen checkpoint and further trained during visual instruction tuning. All language backbone parameters are unfrozen during ViT fine-tuning; only the vision encoder is frozen after visual pretraining.

3 Training Methodology

3.1 Stage 1: Visual Pretraining (Projector Alignment)

In Stage 1, only the MLP projector is trained while the vision encoder and language backbone are frozen. Training data consists of 10M image-caption pairs drawn from web-scraped and curated sources. The objective is next-token prediction on captions conditioned on visual tokens. This stage aligns the visual token space with the language embedding space.

Training details:

- Optimizer: AdamW, LR 1×10^{-3} , cosine decay
- Batch size: 2048 image-text pairs
- Duration: 1 epoch (≈ 5 B visual tokens processed)
- Images: resized to 448 $\times$ 448, no tiling in Stage 1

3.2 Stage 2: Visual Instruction Tuning

In Stage 2, the projector and language backbone are jointly trained on 2M high-quality multi-modal instruction pairs. The vision encoder remains frozen. Dynamic tiling is enabled.

Table 3: Visual instruction tuning data composition.

Task Type	Samples	Fraction
Visual question answering	620,000	31.0%
Image captioning	400,000	20.0%
Document understanding	380,000	19.0%
OCR and text extraction	260,000	13.0%
Chart and diagram QA	200,000	10.0%
Multi-image comparison	100,000	5.0%
Science figure analysis	40,000	2.0%
Total	2,000,000	100.0%

Training details:

- Optimizer: AdamW, LR 2×10^{-5} , cosine decay to 2×10^{-6}
- Batch size: 512 samples
- Duration: 3 epochs
- Dynamic tiling: up to 12 tiles per image

4 Evaluation

4.1 Multimodal Benchmarks

Table 4: Zen-VL benchmark results versus models of comparable scale.

Benchmark	Zen-VL (7B)	Comp. A (7B)	Comp. B (8B)	Comp. C (13B)
MMBench (EN)	81.3	79.4	77.8	80.6
MMBench (CN)	79.8	76.3	74.1	78.2
MME (total)	2148	2019	1987	2103
MME-P (perception)	1587	1502	1474	1543
MMMU (val, 0-shot)	56.8	54.1	52.3	55.7
TextVQA (val)	78.4	76.2	73.8	77.9
DocVQA (val)	88.2	84.6	81.3	86.4
ChartQA	83.4	79.8	76.4	81.7
AI2D	79.6	77.3	74.8	78.4
ScienceQA (IMG)	82.4	80.1	79.6	83.1
HallusionBench	48.2	44.7	42.3	46.8

4.2 Document Understanding

Document understanding is a particularly demanding task requiring fine-grained text reading, layout understanding, and cross-region reasoning. We evaluate on industry-standard benchmarks.

Table 5: Document understanding benchmark results.

Benchmark	Zen-VL (7B)	Prior Best (7B class)
DocVQA (val)	88.2	86.4
DocVQA (test)	87.6	85.8
InfoVQA	72.3	69.8
OCRBench	82.4	79.3
Slide understanding	74.1	70.6

4.3 High-Resolution Ablation

We ablate the dynamic tiling scheme against fixed-resolution baselines.

Table 6: Effect of dynamic tiling on document and text-heavy benchmarks.

Configuration	DocVQA	TextVQA	ChartQA
Fixed 224×224	71.4	63.8	67.2
Fixed 448×448	82.6	73.1	76.8
Dynamic tiling (max 6)	86.3	76.4	80.9
Dynamic tiling (max 12)	88.2	78.4	83.4

The improvement from fixed 448px to dynamic tiling (max 12) is 5.6 points on DocVQA and 6.6 points on ChartQA, confirming that high-resolution tiling is critical for text-dense tasks.

4.4 Multi-Image Evaluation

Table 7: Multi-image benchmark performance.

Benchmark	Zen-VL (7B)	Competitor (13B)
BLINK (multi-image subset)	58.4	56.1
Mantis (multi-image IQ50)	62.3	59.8
MuMuQA	71.6	68.4

4.5 Hallucination

We evaluate object hallucination using POPE (adversarial setting) and HallusionBench.

Table 8: Hallucination evaluation.

Benchmark	Zen-VL (7B)	Competitor A (7B)
POPE (random)	88.3	86.7
POPE (popular)	87.1	85.4
POPE (adversarial)	85.9	83.2
HallusionBench	48.2	44.7

5 Related Work

Vision-language models trace to contrastive pretraining between image and text encoders [1]. Subsequent work (BLIP-2 [2], LLaVA [3]) connected frozen vision encoders to language models via learned projectors. InternVL [5] scaled the vision encoder and fine-tuned it alongside the language model. High-resolution support via tiling was introduced in LLaVA-HD and similar work; dynamic aspect-ratio tiling was explored in InternVL 1.5.

Our work follows the dynamic tiling paradigm but applies it to the Zen MoDE language backbone, yielding competitive results with fewer language parameters than 13B-class competitors on document understanding tasks where high-resolution matters most.

6 Limitations

Zen-VL cannot process video (multiple sequential frames with temporal reasoning). Images must be provided as discrete inputs; streaming video analysis is not supported. Despite high DocVQA scores, performance on handwritten text and low-contrast scans lags typed print. Hallucination rates remain non-trivial on adversarial POPE; object co-occurrence biases from pretraining data are partially inherited. The 7B backbone limits complex multi-step visual reasoning compared to larger language backbones.

7 Conclusion

Zen-VL delivers strong multimodal understanding at the 7B parameter scale through a combination of a high-capacity vision encoder, dynamic tiling for high-resolution images, and multi-image conversation support built on the Zen MoDE language backbone. Benchmark results—MMBench 81.3%, MME 2148, DocVQA 88.2%, TextVQA 78.4%—demonstrate competitive performance with larger-scale competitors, particularly on document and text-heavy visual tasks where high resolution is critical. Zen-VL serves as the multimodal entry point for the Zen family, enabling visual understanding to be combined with Zen’s language generation capabilities in production deployments.

References

- [1] A. Radford et al., “Learning Transferable Visual Models From Natural Language Supervision,” *ICML*, 2021.
- [2] J. Li et al., “BLIP-2: Bootstrapping Language-Image Pre-training with Frozen Image Encoders and Large Language Models,” *ICML*, 2023.
- [3] H. Liu et al., “Visual Instruction Tuning,” *NeurIPS*, 2023.
- [4] A. Dosovitskiy et al., “An Image is Worth 16x16 Words: Transformers for Image Recognition at Scale,” *ICLR*, 2021.
- [5] Z. Chen et al., “InternVL: Scaling up Vision Foundation Models and Aligning for Generic Visual-Linguistic Tasks,” *CVPR*, 2024.
- [6] H. Liu et al., “Improved Baselines with Visual Instruction Tuning,” *arXiv:2310.03744*, 2024.
- [7] J. Ainslie et al., “GQA: Training Generalized Multi-Query Transformer Models,” *EMNLP*, 2023.
- [8] L. Ouyang et al., “Training Language Models to Follow Instructions with Human Feedback,” *NeurIPS*, 2022.