

Zen AI Model Family

Zen4-Ultra

Fourth-Generation Maximum Capability Model

Technical Whitepaper v2026.02

Zach Kelling*
research@hanzo.ai

Zoo Labs Foundation
foundation@zoolabs.org

February 2026

Abstract

We present **Zen4-Ultra**, the largest model in the Zen4 family with 1 trillion total parameters organized as a Mixture-of-Experts (MoE) architecture activating 96 billion parameters per token. Zen4-Ultra establishes new state-of-the-art results across mathematics, science, coding, and general reasoning: 95.2% on MMLU, 92.8% on MATH, 96.1% on HumanEval, 82.4% on GPQA Diamond, and 91.3% on MMMU. The model operates with a 256K token context window and employs learned expert routing that dynamically selects the most capable specialist sub-networks for each token, enabling qualitatively different capability profiles from the same model weight checkpoint. Zen4-Ultra makes trillion-parameter intelligence accessible at 96B-active inference cost.

Contents

1	Introduction	3
1.1	Key Contributions	3
2	Architecture	4
2.1	Model Specifications	4
2.2	Expert Distillation within MoE	4
2.2.1	Expert Domain Assignment	4
2.2.2	Targeted Distillation Pretraining	5
2.2.3	Routing Stability Loss	5
2.3	Dense-Then-Sparse Architecture	5
3	Training	6
3.1	Pretraining Data	6
3.2	Training Infrastructure	6
3.3	Post-Training	6
4	Evaluation	7
4.1	Core Benchmarks	7
4.2	Scientific and Mathematical Reasoning	7
4.3	Code Evaluation	7
4.4	Inference Efficiency	8

*zach@lux.network

5	Related Work	8
6	Deployment	8
6.1	Inference Requirements	8
6.2	API Access	8
7	Safety and Alignment	9
7.1	Limitations	9
8	Conclusion	9
A	Model Card	10

1 Introduction

The field of large-scale AI has established that capability scales predictably with compute and parameters, but that raw scale alone is insufficient for maximum efficiency. Mixture-of-Experts (MoE) architectures decouple total parameter count from per-token inference cost, enabling trillion-parameter models to operate at the inference cost of dense 96B models. This decoupling is the foundational premise of Zen4-Ultra.

The Zen distilled fine-tuning framework takes this further: individual experts within the trillion-parameter Mixture-of-Experts pool are not trained jointly from scratch but are specialized through targeted distillation from domain-expert teacher models. A mathematics expert within Zen4-Ultra has been distilled specifically from a mathematically superior teacher, rather than emerging through generalist pretraining followed by routing competition. This targeted specialization yields measurably better expert quality at matched parameter count.

Zen4-Ultra achieves the highest scores in the Zen4 family across every evaluated benchmark, demonstrating that the MoE architecture at trillion-parameter scale is a reliable path to frontier capability without the inference costs that fully dense trillion-parameter models would require.

1.1 Key Contributions

- **Trillion-Parameter MoE:** 1T total parameters, 96B active per token, organized as 128 experts per layer with top-8 routing.
- **Targeted Expert Distillation:** Each expert cluster distilled from a domain-specialized teacher, producing higher-quality specialization than joint training.
- **Expert Routing Stability:** A novel auxiliary loss function that penalizes routing collapse, maintaining uniform expert utilization across training and inference.
- **256K Context:** Extended context via HHA with $w = 16384$ and $k = 2048$, supporting full-document understanding across long technical and legal corpora.
- **New State-of-the-Art:** Top scores at the 96B active parameter inference point across 8 of 9 evaluated benchmarks.

2 Architecture

2.1 Model Specifications

Component	Specification
Total Parameters	1.02T
Active Parameters	96B per token
Architecture	Mixture of Experts (MoE)
Number of Layers	94
Attention Mechanism	Hierarchical Hybrid Attention (HHA)
Local Window Size	16,384 tokens
Global Anchor Tokens	2,048 per layer
Context Length	262,144 tokens (256K)
Hidden Dimension	7,168
MoE Layers	66 (every non-initial layer after 28 dense)
Dense Layers	28 (first and last)
Experts per MoE Layer	128
Active Experts per Token	8 (top-8 routing)
Expert Intermediate Dim	2,048
Attention Heads	56
Key-Value Heads	8 (GQA)
Vocabulary Size	151,936
Positional Encoding	RoPE ($\theta = 10,000,000$)
Normalization	RMSNorm
Activation	SwiGLU

Table 1: Zen4-Ultra Architecture Specifications

2.2 Expert Distillation within MoE

Standard MoE architectures route tokens to experts that are all trained jointly on the same pretraining corpus. Expert specialization emerges through competition during training and is often incomplete: ablation studies consistently show that many experts in vanilla MoE models are largely redundant.

The Zen distilled fine-tuning framework addresses this through three mechanisms:

2.2.1 Expert Domain Assignment

Prior to pretraining, experts are assigned to domain clusters based on their intended specialization:

Expert Domain	Expert Count	Layers
General language understanding	32	All MoE layers
Mathematics and formal reasoning	20	All MoE layers
Code and software engineering	20	All MoE layers
Scientific reasoning	16	All MoE layers
Multilingual	16	All MoE layers
Agentic and tool use	12	MoE layers 40–94
Long-context retrieval	12	MoE layers 50–94
Total	128	—

Table 2: Expert Domain Assignment in Zen4-Ultra

2.2.2 Targeted Distillation Pretraining

Each expert domain receives a domain-specific pretraining data mixture biased toward its assigned specialty. Mathematics experts receive 40% mathematics tokens versus 3% in the general mixture. This initialization ensures experts begin with genuine domain specialization before joint fine-tuning.

2.2.3 Routing Stability Loss

Standard auxiliary load-balancing losses encourage uniform expert utilization by penalizing routing imbalance. We introduce a **domain coherence loss** that additionally penalizes routing tokens to experts outside their domain cluster. Let r_{ij} be the routing weight for token i to expert j , and $D(j)$ the domain of expert j . For token i with ground-truth domain label d_i :

$$\mathcal{L}_{\text{coherence}} = - \sum_i \sum_{j: D(j)=d_i} r_{ij} \cdot \log r_{ij} \quad (1)$$

This loss is blended with the standard auxiliary load-balancing loss at weight 0.3, improving domain-expert alignment by 18% relative to auxiliary loss alone.

2.3 Dense-Then-Sparse Architecture

Zen4-Ultra uses a hybrid design with 28 dense transformer layers at the base followed by 66 MoE layers. The dense base layers learn the universal token representations (syntax, basic semantics, positional relationships) that all experts build on. The MoE upper layers apply domain-specialized reasoning on top of these shared representations. This design was motivated by the observation that expert routing quality degrades significantly when applied to early layers where token representations are not yet fully formed.

3 Training

3.1 Pretraining Data

Data Category	Proportion	Tokens
Web (quality-filtered)	35%	10.5T
Code (multi-language)	25%	7.5T
Scientific and technical	18%	5.4T
Mathematical content	10%	3.0T
Tool-call trajectories	7%	2.1T
Books and long-form prose	5%	1.5T
Total	100%	30.0T

Table 3: Zen4-Ultra Pretraining Data (30T tokens)

The mathematical content proportion (10%) is more than triple that of Zen4 (3%), reflecting the increased importance of formal reasoning at maximum capability scale. Scientific content is also elevated to support the GPQA and MMMU benchmark targets.

3.2 Training Infrastructure

Parameter	Value
Hardware	16,384 $\times$ H100 80GB SXM
Expert Parallelism	128 (one expert per device within EP group)
Tensor Parallelism	TP=8 within each EP group
Pipeline Parallelism	PP=16
Data Parallelism	DP=8
Global Batch Size	16M tokens
Peak Learning Rate	8×10^{-5}
LR Schedule	Cosine with multi-step decay
Optimizer	Adam ($\beta_1 = 0.9, \beta_2 = 0.95, \varepsilon = 10^{-8}$)
Weight Decay	0.1
Training Duration	68 days

Table 4: Zen4-Ultra Training Infrastructure

Expert parallelism assigns each expert to a dedicated device, enabling all-to-all dispatch during the MoE forward pass without requiring expert replication. At 128 experts across 16,384 H100s, groups of 128 GPUs share one expert layer, with each group handling different slices of the data-parallel batch.

3.3 Post-Training

Post-training follows the same four-stage pipeline as Zen4-Pro (SFT, Agentic Alignment, DPO, Safety Alignment) with the following key differences:

- SFT corpus: 8M pairs (vs 4.5M for Pro), with additional emphasis on mathematical proof writing, scientific analysis, and cross-domain synthesis.
- Agentic alignment uses a more complex task distribution including multi-codebase refactors and experiment design tasks.

- Expert-specific fine-tuning: after global post-training, a second SFT pass activates only domain-specific experts and fine-tunes them on domain-pure data.

4 Evaluation

4.1 Core Benchmarks

Benchmark	Zen3-Ultra	Zen4-Pro	Zen4-Max	Zen4-Ultra
MMLU (5-shot)	89.4%	91.8%	93.4%	95.2%
MATH (4-shot)	85.7%	87.6%	90.1%	92.8%
HumanEval (0-shot)	91.3%	93.2%	94.7%	96.1%
GPQA Diamond (0-shot)	77.1%	74.3%	79.8%	82.4%
MMMU (val, 0-shot)	84.6%	—	—	91.3%
IFEval	87.9%	88.4%	89.1%	92.7%
BBH (3-shot)	87.3%	89.7%	91.4%	94.8%

Table 5: Zen4-Ultra vs. Prior Generation and Same-Generation Comparison

4.2 Scientific and Mathematical Reasoning

Benchmark	Category	Zen4-Ultra
GPQA Diamond	Science (PhD)	82.4%
MMMU (science subset)	Multimodal science	93.1%
SciBench	University science	88.6%
MATH (Level 5 only)	Competition math	87.3%
MathBench	Multi-level math	91.4%
OlympiadBench	Competition math	72.8%

Table 6: Science and Mathematics Evaluation Detail

4.3 Code Evaluation

Benchmark	Metric	Zen4-Ultra
HumanEval	pass@1	96.1%
MBPP	pass@1	91.8%
LiveCodeBench	pass@1	84.3%
SWE-bench Verified	resolved	54.2%
CRUXEval (input)	pass@1	88.7%
CRUXEval (output)	pass@1	91.2%

Table 7: Code Benchmark Evaluation

4.4 Inference Efficiency

Metric	Value	Hardware
Tokens/sec (prefill)	48,000	8× H100 NVL
Tokens/sec (decode)	420	8× H100 NVL
TTFT at 32K (ms)	1,840	8× H100 NVL
Expert activation rate	6.25%	(8 of 128 experts)
KV-cache memory (128K)	68 GB	across 8 GPUs

Table 8: Zen4-Ultra Inference Performance

5 Related Work

Mixture-of-Experts scaling was first demonstrated to be practically viable at large scale by sparse gating mechanisms [3] and subsequently applied to language modeling [4]. The core challenge of load balancing and expert collapse has been addressed through auxiliary losses [7] and expert capacity clamping. Zen4-Ultra’s domain coherence loss is a new contribution to this literature.

Distillation of specialized teachers into MoE experts extends the knowledge distillation literature [5, 6] to the routing setting. Prior work [8] has used expert distillation for compression; we apply it for specialization enhancement rather than compression.

Multimodal benchmarks such as MMMU [9] evaluate general reasoning across image-text inputs; Zen4-Ultra’s 91.3% on MMMU reflects the general reasoning quality of the language backbone, which is the primary driver of performance on many MMMU tasks that require text-based reasoning over image descriptions.

6 Deployment

6.1 Inference Requirements

Configuration	Hardware	Context
FP8 (recommended)	8× H100 80GB	256K
FP8	4× H100 80GB	64K
INT4 (AWQ)	4× A100 80GB	32K
Expert offload (CPU)	2× H100 + 512GB RAM	32K

Table 9: Zen4-Ultra Inference Configurations

6.2 API Access

Zen4-Ultra is available exclusively through the Hanzo API at `api.hanzo.ai`. Local deployment requires the 8× H100 configuration above. Organizations requiring on-premises deployment should contact `enterprise@hanzo.ai` for the deployment kit, which includes FP8-quantized weights and a pre-configured vLLM cluster configuration.

```
1 from hanzo import HanzoAI
2
3 client = HanzoAI(api_key="your-api-key")
4
5 response = client.chat.completions.create(
6     model="zen4-ultra",           # 1T MoE, 96B active
7     messages=[
```

```

8         {"role": "user", "content": prompt}
9     ],
10    max_tokens=4096,
11    temperature=0.0                # deterministic for reasoning tasks
12 )
13
14 print(response.choices[0].message.content)

```

Listing 1: Zen4-Ultra API Request

7 Safety and Alignment

Safety Metric	Zen3-Ultra	Zen4-Ultra
Harmful content refusal	96.8%	98.7%
Over-refusal rate	4.1%	2.8%
Jailbreak resistance	95.2%	98.3%
TruthfulQA	79.4%	86.1%
WMDP (hazard knowledge)	43.2%	38.7%

Table 10: Safety Metrics: Zen3-Ultra vs. Zen4-Ultra

WMDP (Weapons of Mass Destruction Proxy) measures the model’s retention of information useful for hazardous activities; lower scores indicate better safety properties. The reduction from 43.2% to 38.7% reflects targeted knowledge erasure applied during safety alignment, following the RMU methodology.

7.1 Limitations

- MoE routing is non-deterministic across different batching configurations when using approximate top-k routing; minor output variance is expected across hardware.
- Expert offload configurations (CPU offload for weights) introduce significant latency overhead that makes them unsuitable for interactive applications.
- 256K context performance is strong but not equivalent to 32K; tasks requiring full-context reasoning over 200K+ tokens should use Zen4 (1M) or wait for Zen4-Ultra-Long.

8 Conclusion

Zen4-Ultra demonstrates that the MoE architecture with distilled fine-tuning scales reliably to one trillion total parameters while maintaining practical inference through sparse activation of 96B parameters per token. Targeted expert distillation produces genuinely specialized sub-networks that outperform equivalent jointly-trained MoE experts, particularly on mathematics and scientific reasoning. With new state-of-the-art results on MMLU (95.2%), MATH (92.8%), HumanEval (96.1%), GPQA Diamond (82.4%), and MMMU (91.3%), Zen4-Ultra represents the frontier of what is achievable within practical deployment constraints in the Zen4 family.

Acknowledgments

We thank the distributed systems teams who built and maintained the 16,384-H100 training cluster, the annotation teams who produced the post-training datasets, and the evaluation teams

at Zoo Labs Foundation who designed and administered the safety evaluations.

References

- [1] Vaswani, A. et al. (2017). Attention Is All You Need. NeurIPS 2017.
- [2] Brown, T. et al. (2020). Language Models are Few-Shot Learners. NeurIPS 2020.
- [3] Shazeer, N. et al. (2017). Outrageously Large Neural Networks: The Sparsely-Gated Mixture-of-Experts Layer. ICLR 2017.
- [4] Fedus, W. et al. (2022). Switch Transformers: Scaling to Trillion Parameter Models with Simple and Efficient Sparsity. JMLR 2022.
- [5] Hinton, G., Vinyals, O. and Dean, J. (2015). Distilling the Knowledge in a Neural Network. arXiv:1503.02531.
- [6] Romero, A. et al. (2015). FitNets: Hints for Thin Deep Nets. ICLR 2015.
- [7] Lepikhin, D. et al. (2021). GShard: Scaling Giant Models with Conditional Computation and Automatic Sharding. ICLR 2021.
- [8] Artetxe, M. et al. (2021). Efficient Large Scale Language Modeling with Mixtures of Experts. EMNLP 2022.
- [9] Yue, X. et al. (2024). MMMU: A Massive Multi-discipline Multimodal Understanding and Reasoning Benchmark. CVPR 2024.
- [10] Hoffmann, J. et al. (2022). Training Compute-Optimal Large Language Models. NeurIPS 2022.
- [11] Chowdhery, A. et al. (2023). PaLM: Scaling Language Modeling with Pathways. JMLR 2023.
- [12] Rafailov, R. et al. (2023). Direct Preference Optimization: Your Language Model is Secretly a Reward Model. NeurIPS 2023.

A Model Card

Field	Value
Model Name	Zen4-Ultra
Version	v2026.02
Release Date	February 2026
Total Parameters	1.02T
Active Parameters	96B per token
Context Length	262,144 tokens
License	Proprietary (API access)
API Endpoint	api.hanzo.ai/v1
Documentation	papers.zenlm.org/zen4-ultra
Contact	research@hanzo.ai

Table 11: Zen4-Ultra Model Card