

Zen3-Nano: Third-Generation Ultra-Compact Edge Intelligence

Technical Whitepaper v2025.06

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

June 2025

Abstract

Zen3-Nano is the third-generation 0.6B parameter edge model from the Zen family, achieving significant quality improvements over Zen-Nano through architectural refinements and enhanced knowledge distillation from the full Zen family hierarchy. Built on the Zen MoDE (Mixture of Distilled Experts) architecture, Zen3-Nano maintains a sub-250MB quantized footprint while delivering improved reasoning capabilities on mobile and embedded hardware. At 4-bit quantization the model fits in 250MB and achieves 2800 tokens per second on Apple A16 Bionic silicon, enabling real-time language inference at the edge without network connectivity. The model sets new state-of-the-art results for the sub-1B parameter class: MMLU 47.3%, TriviaQA 61.2%, and GSM8K 52.4%, demonstrating that frontier distillation techniques close the gap between edge and server-class models substantially.

Contents

1	Introduction	3
2	Architecture	3
2.1	Zen MoDE Foundation	3
2.2	Grouped-Query Attention (GQA) Variant	3
2.3	Fused MLP Design	4
2.4	Quantization-Native Design	4
3	Training	4
3.1	Pre-Training Data	4
3.2	Hierarchical Distillation	4
3.3	Instruction Fine-Tuning	5
3.4	RLHF and Constitutional Alignment	5
4	Evaluation	5
4.1	Language Understanding	5
4.2	Mathematical Reasoning	5
4.3	Code Generation	5
4.4	Inference Performance	6
4.5	Multilingual Performance	6

5	Related Work	6
5.1	Small Language Models	6
5.2	Quantization-Aware Training	6
5.3	Knowledge Distillation Hierarchies	6
6	Deployment and Runtime	7
6.1	Supported Runtimes	7
6.2	Context Length	7
6.3	API Compatibility	7
7	Conclusion	7

1 Introduction

On-device AI inference has become a primary deployment target for consumer applications, IoT sensors, and privacy-sensitive enterprise workloads. The constraints are severe: models must fit within a few hundred megabytes, execute without accelerators on ARM Cortex-class CPUs, and return responses in tens of milliseconds. Existing sub-1B approaches sacrifice too much quality to be practically useful for general-purpose instruction following.

Zen3-Nano addresses this challenge by combining three advances:

1. **Hierarchical Zen MoDE distillation** — structured distillation cascaded through Zen3-Mini (1.7B), Zen3-Small (7B), and Zen3-Medium (14B) teachers, each contributing a complementary signal.
2. **Third-generation Nano architecture** — refined attention with grouped-query heads optimized for cache reuse on low-memory hardware, reduced KV cache footprint, and fused MLPs that lower inference memory bandwidth.
3. **BitDelta-aware training** — the model is trained with quantization-aware objectives derived from the Zoo Labs BitDelta protocol (ZIP-007), so the 4-bit quantized artifact retains near-full-precision accuracy without post-hoc calibration.

Zen3-Nano is designed to run natively on:

- Mobile SoCs (Apple A16 Bionic, Qualcomm Snapdragon 8 Gen 3, MediaTek Dimensity 9300)
- Microcontrollers with ≥ 512 KB SRAM (via WASM SIMD runtime)
- Browser environments via WebAssembly and WebGPU
- Embedded Linux (Raspberry Pi 5, NVIDIA Jetson Nano)

2 Architecture

2.1 Zen MoDE Foundation

Zen3-Nano is built on the Zen MoDE (Mixture of Distilled Experts) architecture, a design principle shared across the entire Zen family that replaces monolithic dense feed-forward blocks with sparsely-activated expert clusters. At the 0.6B scale, MoDE manifests as a lightweight routing mechanism that activates a subset of micro-experts per token, preserving capacity diversity while keeping the active parameter count well below memory limits.

Formally, for input token representation $\mathbf{x} \in \mathbb{R}^d$, the MoDE feed-forward layer computes:

$$\text{MoDE}(\mathbf{x}) = \sum_{i \in \mathcal{K}(\mathbf{x})} g_i(\mathbf{x}) \cdot E_i(\mathbf{x}) \quad (1)$$

where $\mathcal{K}(\mathbf{x})$ is the set of top- k activated experts selected by a learned routing function, g_i are normalized routing weights, and E_i are the expert sub-networks.

2.2 Grouped-Query Attention (GQA) Variant

Standard multi-head attention at 0.6B scale incurs disproportionate KV cache overhead relative to total parameter count. Zen3-Nano uses a 3rd-generation grouped-query attention variant with the following configuration:

- 16 query heads grouped into 4 key-value head groups

- 1024-dimensional hidden state, 64-dimensional head size
- Rotary positional embeddings (RoPE) with extended context via YaRN interpolation
- 32 transformer layers with pre-layer normalization (RMSNorm)

The KV cache per token occupies only 512 bytes at full precision, enabling 32K token context windows within 16 MB of SRAM when using 8-bit KV quantization.

2.3 Fused MLP Design

The feed-forward blocks use a fused SwiGLU activation with weight-sharing across adjacent expert pairs, reducing total parameter count by 12% relative to a naive MoDE implementation at equivalent quality:

$$\text{FFN}(\mathbf{x}) = (\sigma(\mathbf{W}_1\mathbf{x}) \odot \mathbf{W}_2\mathbf{x}) \mathbf{W}_3 \quad (2)$$

2.4 Quantization-Native Design

Unlike post-hoc quantization, Zen3-Nano is trained with BitDelta quantization-aware training from epoch 1. The quantization error ϵ_Q is included in the training loss as a regularizer:

$$\mathcal{L}_{\text{total}} = \mathcal{L}_{\text{CE}} + \mathcal{L}_{\text{KD}} + \lambda \cdot \mathbb{E} [\|\mathbf{W} - Q(\mathbf{W})\|_F^2] \quad (3)$$

where $Q(\cdot)$ is the 4-bit quantization operator and $\lambda = 0.01$.

3 Training

3.1 Pre-Training Data

Zen3-Nano is initialized from scratch on a 2T token corpus drawn from:

- Curated web text (Common Crawl quality-filtered): 60%
- Code repositories (GitHub, GitLab, permissive licenses): 20%
- Books, academic papers, structured knowledge: 15%
- Multilingual data (50+ languages, weighted by speaker population): 5%

3.2 Hierarchical Distillation

After base pre-training, Zen3-Nano undergoes a three-stage distillation curriculum:

Stage 1 (Token-level KL distillation): Soft targets from Zen3-Mini (1.7B) teacher over 500B tokens. Temperature $T = 4.0$, KD weight $\alpha = 0.7$.

Stage 2 (Layer-level representation matching): Intermediate hidden states from Zen3-Small (7B) are used as representation targets via a learned projection head. Loss: cosine similarity maximization over layers $\{6, 12, 18, 24\}$.

Stage 3 (Reasoning chain distillation): Structured chain-of-thought traces generated by Zen3-Medium (14B) are used for supervised fine-tuning with a 3:1 ratio of distilled traces to human annotations.

3.3 Instruction Fine-Tuning

Post-distillation instruction tuning uses a 50M token mix:

- General instruction following: 40%
- Code generation and debugging: 25%
- Mathematical reasoning (chain-of-thought): 20%
- Multilingual instruction: 15%

Training uses AdamW optimizer, peak learning rate 3×10^{-4} , cosine decay, batch size 512, over 3 epochs on the instruction mix.

3.4 RLHF and Constitutional Alignment

A lightweight RLHF pass using PPO with a 0.6B reward model (trained separately) is applied for 10K steps to reduce harmful outputs and improve instruction adherence, keeping the KL divergence from the SFT checkpoint below 0.1 nats.

4 Evaluation

4.1 Language Understanding

Table 1: Zen3-Nano vs. sub-1B baselines on core language benchmarks

Model	Params	MMLU	TriviaQA	HellaSwag
Prior Zen-Nano	0.5B	39.1	53.4	61.2
Zen2-Nano	0.6B	43.7	57.8	66.4
Zen3-Nano	0.6B	47.3	61.2	70.8

4.2 Mathematical Reasoning

Table 2: Mathematical reasoning benchmarks

Model	GSM8K	MATH	MGSM (avg)	ARC-C
Zen-Nano	38.2	11.3	31.4	54.1
Zen2-Nano	44.6	15.7	38.9	59.3
Zen3-Nano	52.4	20.1	46.2	64.7

4.3 Code Generation

Table 3: Code generation benchmarks (pass@1)

Model	HumanEval	MBPP	LiveCodeBench
Zen-Nano	28.0	33.2	12.4
Zen2-Nano	34.1	40.5	18.7
Zen3-Nano	42.7	49.3	26.1

4.4 Inference Performance

Table 4: Inference throughput across hardware targets (4-bit quantized)

Hardware	Precision	Tok/s	RAM (MB)	Latency (ms/tok)
Apple A16 Bionic	INT4	2800	248	0.36
Snapdragon 8 Gen 3	INT4	2340	248	0.43
Raspberry Pi 5	INT4	420	248	2.38
WASM SIMD (browser)	INT4	310	248	3.23
x86-64 (laptop CPU)	INT4	1100	248	0.91
NVIDIA RTX 4090	FP16	8400	950	0.12

4.5 Multilingual Performance

Table 5: Multilingual benchmark (FLORES-200, spBLEU)

Language Pair	Zen3-Nano	Zen2-Nano
English → Chinese	24.3	20.1
English → Spanish	31.8	27.4
English → Arabic	18.9	14.7
English → Japanese	21.4	17.2
Average (50 pairs)	22.6	18.4

5 Related Work

5.1 Small Language Models

Research on sub-1B language models has accelerated as mobile deployment became commercially viable. Early work demonstrated that aggressive distillation from multi-billion parameter teachers could recover a substantial fraction of performance at a fraction of the parameter count [1]. Subsequent work explored architectural modifications specifically suited to on-device constraints, including weight sharing, structured pruning, and low-rank factorization.

5.2 Quantization-Aware Training

Post-training quantization (PTQ) has become the default for edge deployment, but accuracy degradation at 4-bit is nontrivial for small models where the per-parameter information density is already high. BitDelta (ZIP-007) addresses this by encoding quantization error as an explicit training signal, showing that models trained quantization-natively outperform PTQ equivalents by 2–5 percentage points on reasoning benchmarks at 4-bit precision.

5.3 Knowledge Distillation Hierarchies

Prior work on hierarchical distillation [2] demonstrated that using an intermediate “teacher assistant” model reduces the capacity gap and improves distillation efficiency. Zen3-Nano extends this to a three-level cascade, exploiting the full Zen family model hierarchy.

6 Deployment and Runtime

6.1 Supported Runtimes

Zen3-Nano ships with native support for:

- **llama.cpp** — GGUF format, Q4_K_M quantization
- **ONNX Runtime** — optimized for ARM and x86 SIMD
- **WebAssembly** — via wasm-pack compiled inference engine
- **Core ML** — iOS/macOS native accelerated inference
- **TensorFlow Lite** — Android NPU acceleration

6.2 Context Length

Default context window: 8192 tokens (configurable to 32K via YaRN at cost of increased RAM). At 32K context, KV cache occupies approximately 64 MB at INT8.

6.3 API Compatibility

Zen3-Nano exposes an OpenAI-compatible API when deployed via the Zen LM inference server, enabling drop-in replacement for any OpenAI SDK client.

7 Conclusion

Zen3-Nano establishes a new state of the art for sub-1B parameter language models through hierarchical Zen MoDE distillation, quantization-native training, and architectural refinements that maximize inference efficiency on ARM-class hardware. The model achieves MMLU 47.3%, TriviaQA 61.2%, and GSM8K 52.4% — results that would have required 3–7B parameter models in prior generations — while fitting within 250 MB at 4-bit precision and achieving 2800 tok/s on Apple A16 Bionic.

Future work will target 1M token context windows via ring attention approximations suitable for low-memory hardware, and extended multimodal understanding via lightweight vision adapters compatible with the Zen3-Omni decoder.

Acknowledgments

The authors thank the Zoo Labs Foundation for compute allocation and the DSO research network for benchmark curation and evaluation infrastructure.

References

- [1] G. Hinton, O. Vinyals, and J. Dean, “Distilling the Knowledge in a Neural Network,” *NIPS Deep Learning Workshop*, 2015.
- [2] S. I. Mirzadeh et al., “Improved Knowledge Distillation via Teacher Assistant,” *AAAI*, 2020.
- [3] Zoo Labs Foundation, “ZIP-007: BitDelta — Quantization-Aware Delta Compression for Language Models,” *Zoo Improvement Proposals*, 2024. zips.zoo.ngo/zip-007
- [4] Hanzo AI, “HIP-002: Active Semantic Optimization (ASO),” *Hanzo Improvement Proposals*, 2024.