



Poster
Live Demo
6 ARC-AGI scenarios

AGENTIC AI AT THE EDGE

Cost-Aware Orchestrator & Benchmarks

Kulbir Singh Ahluwalia
University of Illinois Urbana-Champaign
PhD Candidate, Computer Science
Team: Compute BU
Mentor: Alok Sharma
Manager: Mandar Deshpande

DEMAND FOR ACCELERATED COMPUTE IS GROWING FASTER THAN SUPPLY: AND THE GAP IS INCREASING.

Agentic AI runs everywhere: a Cost-Aware Orchestrator decides *WHERE?* & *HOW?* per task.
Route each agentic task to the cheapest *capable* device across 8 Compute devices, NPU/GPU accelerators, ARM64 & AMD64 hosts, Multiple OS deployments

Goal: Minimize TCO(Total Cost of Operation) & Latency. Result: 51% of agentic commands served on an 18.3W edge NPU at comparable tool-calling accuracy, 2.4× fewer GPU tokens.

1 Problem Formulation: Edge Agentic Deployments

- Agentic AI exponentially increases Token demand: Total inference cost, not quality, decides placement:
- Inference Optimization: Hardware differences in bandwidth & power decide Total Cost of Operation (TCO).
- Current Orchestration tools optimise for resource utilisation; not measured joules per task.
- Cost Aware Routing: Power efficient silicon supplements high end accelerators.

Takeaway: Optimize for multi device, multi architecture, multi OS cost aware agentic deployments.

2 Related Work: SOTA Edge, Hybrid, Cloud KPIs

Device	Tok/s	TTFT (s)	Tool Sel	BW (GB/s)	Power (W)
IQ9 Hexagon NPU	18.3±0.1	0.19	1.00	37 ^d	18.3 ^D
GB10 DGX Spark	76.1 [†]	0.043	1.00	140	51.6 ^R
Jetson Thor	51.3±1.2	0.075	0.59 ^m	163.8	65.9 ^B
RTX 5090	327±0.5	0.017	0.96 ^m	1069 ^s	285.6 ^R
Jetson Orin 64 GB	38.1±0.05 ^o	1.64 ^o	1.00 ^q	n.m.	49.2±2.2 ^{M, o}
GH200 ×4 (cloud) [‡]	43	n.m.	0.89	3608 ^s	n.m.
H200 141 GB (ICRN) ^h	107.8 ^h	0.44 ^h	1.00 ^h	n.m.	174 ^R
H200 143 GB (ICC) ^H	n.m.	n.m.	1.00 ^q	n.m.	n.m.

Takeaway: IQ9 competes closely with GPU's tool-selection accuracy at 1/15th power.

3 Quantitative Evaluation: Edge Compute & Use Cases

Use case	IQ9 Hexagon NPU · 36 GB 18.3 W ¹	Jetson Thor unified · 122 GB 65.9 W ²	GB10 Spark unified · 128 GB 51.6 W ³	RTX 5090 discrete · 32 GB 285.6 W ⁴	A100-80GB ICC SECONDARY not metered	H200-141GB ICC SCAVENGER not metered
Agentic tool-calling 27-case hard suite	n.m. ¹ · n.m. ¹ 3.0 GB ¹ · 6.38 s 1.00/1.00 · qwen2-4b	n.m. ¹ · 86 ms n.m. ² · 0.76 s 1.00/1.00 · 4b-2507	n.m. ¹ · 54 ms n.m. ² · 0.77 s 1.00/1.00 · qwen2.5.7b	n.m. ¹ · 27 s n.m. ² · 0.17 s 1.00/1.00 · qwen2.5.7b	n.m. ¹ · n.m. ³ n.m. ² · 0.42 s 1.00/1.00 · qwen2.5.7b	n.m. ¹ · n.m. ³ n.m. ² · 0.32 s 1.00/1.00 · qwen2.5.7b
Long horizon planning Solobench corridor sweep	14.3 t/s · n.m. 3.0 GB ¹ · 1.3 s 2/8 (0/25 · qwen2-4b)	50.8 t/s · n.m. n.m. ² · 38 s 0/8 (0/25 · 4b-sthink)	60.7 t/s · n.m. n.m. ² · 107 s 0/8 (0/25 · 4b-qt)	139.3 t/s · n.m. n.m. ² · 27 s 0/8 (0/25 · 4b-qt)	n.m.	n.m.
LLM text-gen summarise · ≥32-tok outputs	18.3 t/s · 188 ms 3.0 GB ¹ · n.m. 128 tok · qwen2-4b	50.8 t/s · 86 ms n.m. ² · n.m. ≥2k gen tok · 4b ¹	69.8 t/s · 39 ms n.m. ² · n.m. ≥2k gen tok · 4b ¹	139.3 t/s · 15 ms n.m. ² · n.m. ≥2k gen tok · 4b ¹	n.m.	n.m.
CV perception classify · detect · depth	n/a ⁶ · n/a ⁶ 6.7 / 8.5 / 10.7 ms inc/YOLO11/MiDaS	n.m.	n.m.	n.m.	n.m.	n.m.

Cell slots: decode tok/s · TTFS first token | peak memory · total time | gate value · model. Colour is mechanical, per (use case, model): G1 model serves · G2 the row's printed quality throughput gate · G3 its SLO; ■ optimal, all applicable gates pass · ■ one of G2/G3 fails · ■ G1 fails, both fail, or measured zero-success · □ n/a, asserts nothing.

1.80 vs 0.22 on one device; cells name models. Deep KPI 1, thinking budget, side unmeasured split. GB10 10.240 MB (qwen2.5.7b, 32.768 cells / 25 layers). 10.240 MB (qwen2.5.7b, 40.960 / 64), identical across A100-80GB/1200 · 129 0.64 GB · derived (cfr 4096), not runtime peak. tool-call outputs under the ≥32-output-tok decode-validity floor: 5090 127 t/s short vs 139 t/s at ≥2k, one device, footnoted.

Fig. 1: the 18.3 W Hexagon NPU is optimal in 3 of its 4 measured workloads: more than any other device here.

Takeaway: Only one device is yellow on planning, and it pays 187 \$ for it.

4 Method: Distributed Cost Aware Orchestration

Route every task to the device d^* in a fleet spanning mobile NPU → edge NPU → AI box → discrete GPU → cluster:

$$d^* = \arg \min_d (w_1 \text{cost}^{\text{energy}}(d) + w_2 \text{latency}(d)/\text{SLO}) \quad \text{s.t.} \quad \text{mem}_{\text{weights}} + \text{mem}_{\text{KV}} \leq \text{VRAM}(d)$$

In practice the cost term is min-max normalised within the admissible device set, so w_1, w_2 act as comparable weights.

- Owned hardware pays energy × site tariff; rented API pays \$/Mtok; a grant cluster is \$0.
- Decode is bandwidth-bound: tok/s ≈ BW ÷ bytes-per-token (roofline self-check 1.1%).
- Cost is route-dependent: 3.5× tariff spread ⇒ lowest-power ≠ lowest-cost.

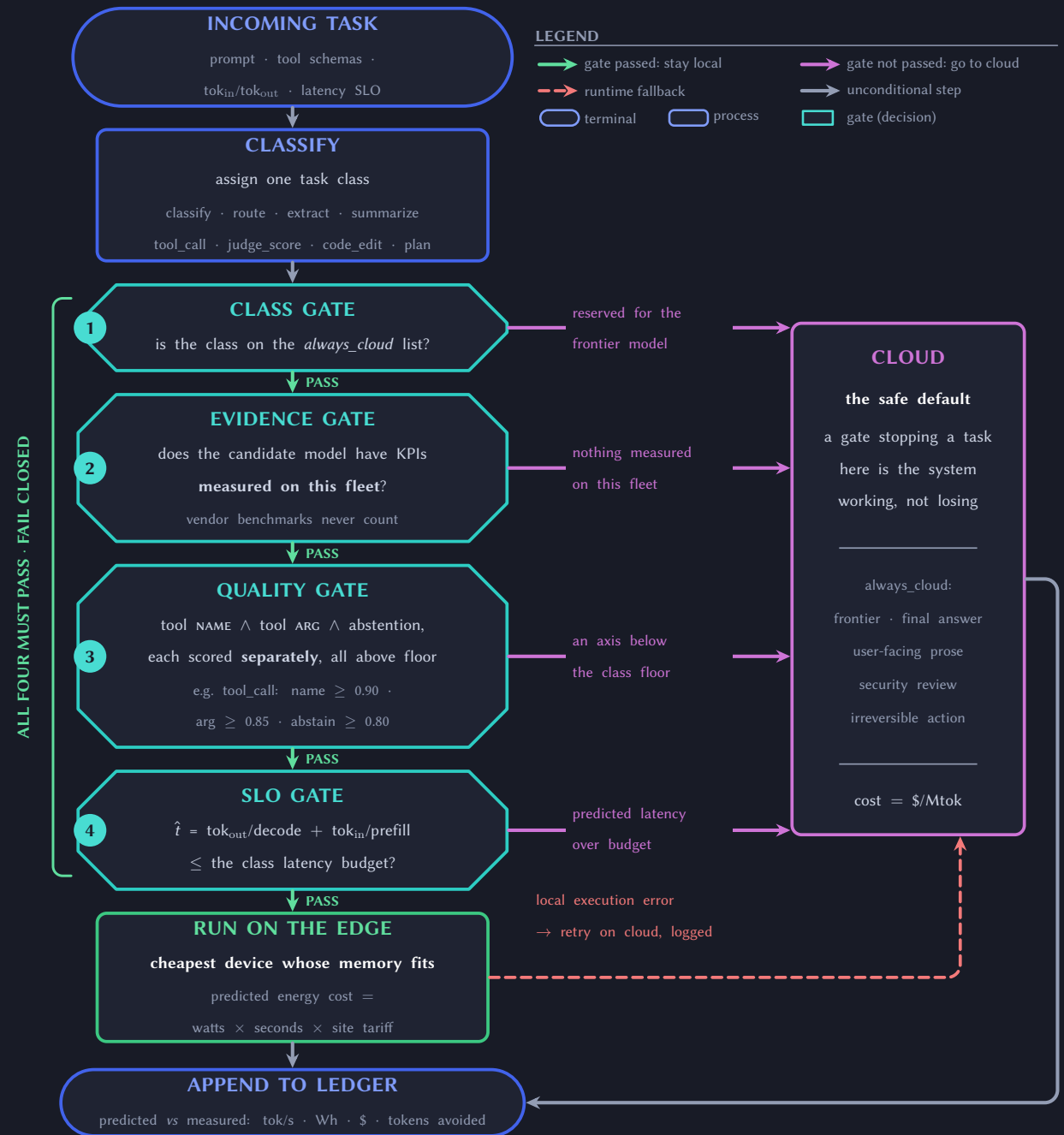


Fig. 2: all four gates must pass to stay local; any failure fails closed to the cloud rail.

Takeaway: Predict by roofline; price the route; place cheapest.

5 Model Deployment Benchmarks

Model-size confound removed: IQ9 NPU (18.3 W) vs RTX 5090 (286 W), median: it splits three ways:

WINS short/structured TIES code-gen LOSES long reasoning

json 1.6 s vs 8.7 s 12.4 s vs 12.2 s gsm8k 27.2 s vs 10.8 s

tool-use 2.3 s vs 7.6 s

Identical 24-case suite: IQ9 matches GB10: tool-sel 1.00, semantic-arg 1.00, strict 0.917 on both (whitespace only). Median complete

tool-call 3.72 s; TTFT 188 ms is first-token only.

Use case	Decode ¹		TTFT		Peak mem	
	NT	Think	NT	Think	NT	Think
LLM text-gen · Qwen2-4B NPU	18.3	18.3	188 ms@227t	0.1 s	3.0 GiB ⁷	3.0 GiB ⁷
LLM text-gen · Qwen2-8B NPU [†]	11.2	10.9	0.1 s	0.1 s	5.2 GiB ⁷	5.2 GiB ⁷
VLM image-in · Qwen2-VL-4B NPU	16.8		509 ms@282t		2.2 GiB ⁸	
Small LLM · 1.5B qt_0 CPU	43.4		1.06 s		19 GiB	
Small LLM · 1.5B qt_0 Adreno	42.6		1.03 s		19 GiB	
Image gen · SD1.5 NPU	6.11 s/img ⁶				0.83 GiB	
CV classify · InceptionV3 NPU	5.9 ms/inf ⁹				44 MB	
CV detect · MiDaS v2 NPU	8.8 ms/inf ⁹				39 MB	
CV segment · DeepLabv3	blocked: QAIRT 2.45/2.38 toolchain wall: needs off-board 2.38 re-export					
CV detect · YOLOv11-640 NPU	8.5 ms/inf ⁹				266 MB	
Agentic tool-call · 4B NPU	15.0	20.2	0.7 s ⁵	pending	3.0 GiB ⁶	3.0 GiB

Legends: ■ measured ■ measured, caveat^{6-o} ■ pending / NOT MEASURED — = not applicable NT = non-thinking single-mode rows span both columns.

¹decode tok/s (median); CV rows ms/inference; image-gen s/image. ²derived from measured per-stage times (UNet 114.7 ms/step). ³streamed first-token on an ~800-tok agentic prompt (tool schemas); complete-call median 3.72 s is separate. ⁴ICRN 35B-A3B Q4, 12-case · ⁵ICC full H200, qwen2.5.7b, 2.979 valid summaries (80.433 pooled); accuracy + 0.33 s median E2E only · ⁶max across models / block-8: 15-case GenIE suite · ⁷Orin, accuracy only · ⁸27-case suite on A100-40/80 GB + H200: identical tool-selection; latency 4.34 / 4.35 / 2.52 s · n.i.=not instrumented. Power: ⁹whole-device meter · ¹⁰Orin under certain load (tok/s lower bound, rail power upper bound) · ¹¹Orin module-rail sum (INA3221) · ¹²GPU rail · ¹³board VIN; no cross-boundary tok/J ranking; derivations: repo README (QR).

Takeaway: NPU wins short, ties code, loses long.

18.3 W

Dragonwing IQ 9075 EVK

Average Power Supply Reading

1.00 : 1.00

Zero accuracy tax at the edge:

NPU matches GPU, 24/24 both

51%

(Edge/Subnet+Cloud) Inference Ratio

Secure, Lower Cost, Lower Latency

2.4×

Fewer GPU-tier tokens

(435k vs 1.04M projected)

94.5 %

End-to-End validity

(matches all-35B)

3.72 s

Median Tool Call

Completion time

Methodology: 428,938 measured tool-calls, 138 device×model×mode cells (14 devices, 33 models, 22-case suite, temp 0, Jul 20 → Aug 3 2026; counted from valid summaries only. Temp-0 repeats = reproducibility; breadth = the 138 cells. Block-2 = iso-4b unless noted (W4A16 IQ9 / Q4 NVIDIA, GH200: GLM-5.2 744B MoE 2-bit). Not benchmarked (declared): MSI RTX3080, MacBook M3 Max. Keys: ¹derived · ²spec · ³bench-hh median (Thor 0.59 = thinking-4B; instruct-4B = 1.00) · ⁴qwen2.5.7b · ⁵ICRN 35B-A3B Q4, 12-case · ⁶ICC full H200, qwen2.5.7b, 2.979 valid summaries (80.433 pooled); accuracy + 0.33 s median E2E only · ⁷max across models / block-8: 15-case GenIE suite · ⁸Orin, accuracy only · ⁹27-case suite on A100-40/80 GB + H200: identical tool-selection; latency 4.34 / 4.35 / 2.52 s · n.i.=not instrumented. Power: ¹⁰whole-device meter · ¹¹Orin under certain load (tok/s lower bound, rail power upper bound) · ¹²Orin module-rail sum (INA3221) · ¹³GPU rail · ¹⁴board VIN; no cross-boundary tok/J ranking; derivations: repo README (QR).

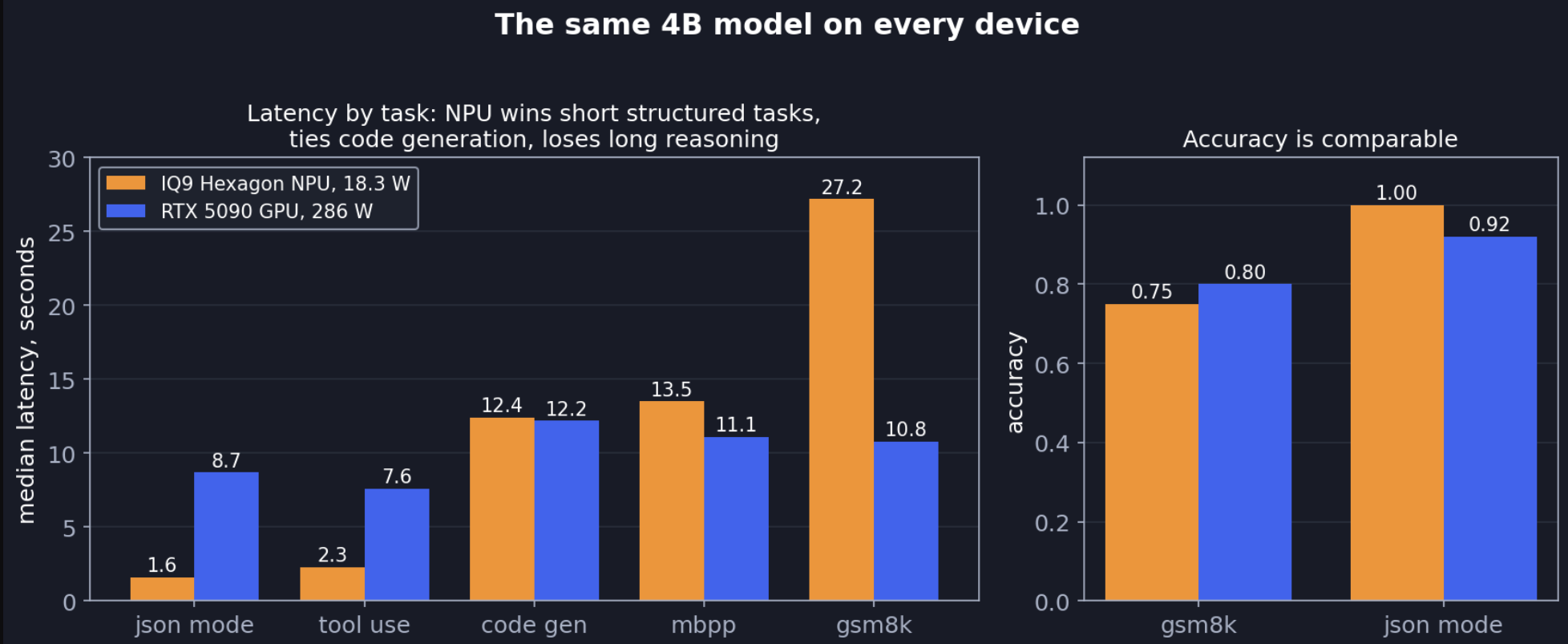


Fig. 3: short outputs are prefill-bound NPU wins (json 1.6 s

vs 8.7 s); long outputs are decode-bound and bandwidth-

limited GPU wins (gsm8k 27.2 s vs 10.8 s).

6 Edge Token Efficiency & Privacy

- Data never leaves the device: W4A16 on the Hexagon NPU (Genie/QNN, no CUDA, no cloud egress); 0.1 s on-device TTFT, no queue variance. Inter-device context rides TailScale WireGuard (Curve25519 + ChaCha20), secret-redacted before any write.
- KG-at-edge token reduction: session context as a knowledge-graph digest, not the raw transcript, compresses 689× (14.7 MB → 21 KB, 0.24 s, \$0, zero model tokens); content-addressed sharing sends only new bytes (≤2 KB/entry).
- Cost-aware orchestrator (40-KPI canon, K1–K40): a tool-call subtask = \$0.0105 (Claude Sonnet 5) vs \$0.000044 Urbana electricity on a GB10: 240×; live demo savings 4.9–31×. \$/Mtok (list): Opus 5 5/25 · Sonnet 5 3/15 · GPT-4o 2.5/10.

Takeaway: Edge keeps data private and cuts context tokens 689× via KG representation.

Collaboration advantage: one orchestrator spans 4 OSes, ARM & x86, NPU/GPU/edge/cloud: teams share compact knowledge-graph context, not raw data.

6 KPIs for Agentic Benchmarks

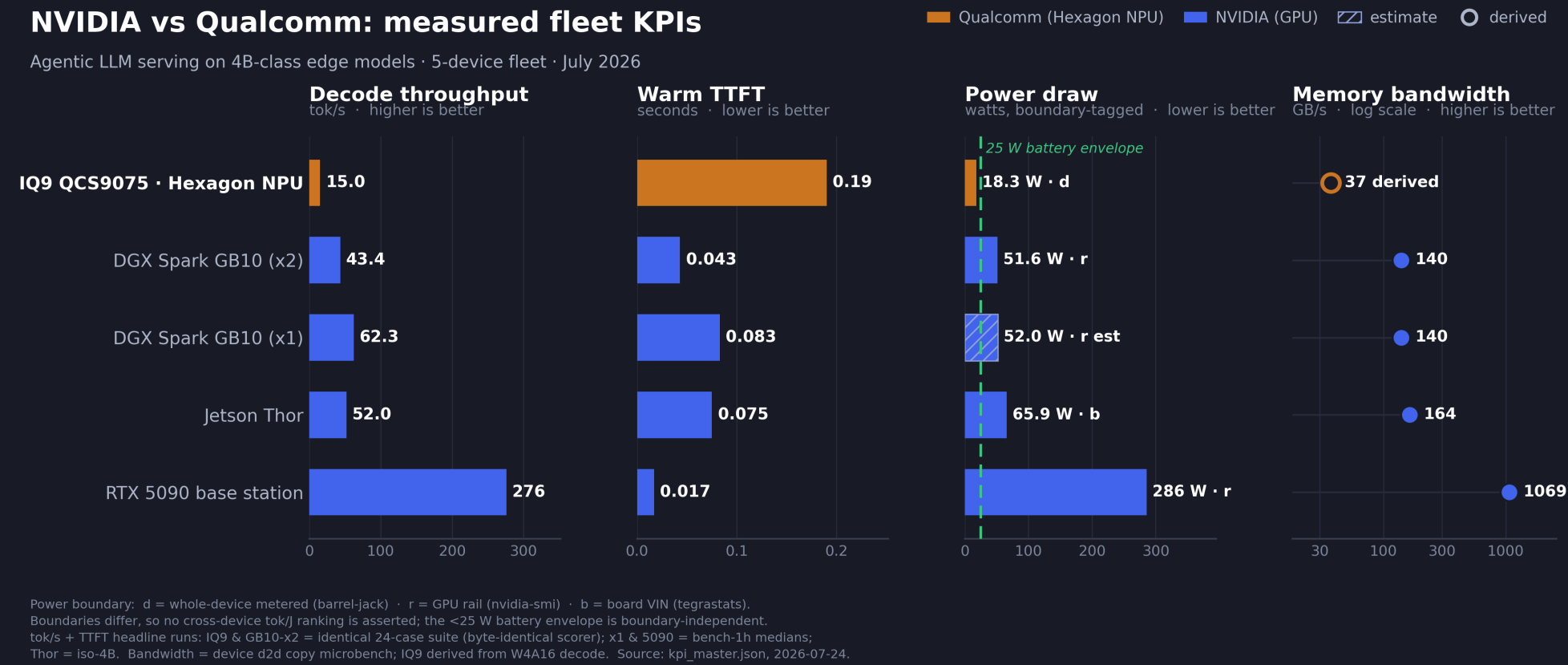
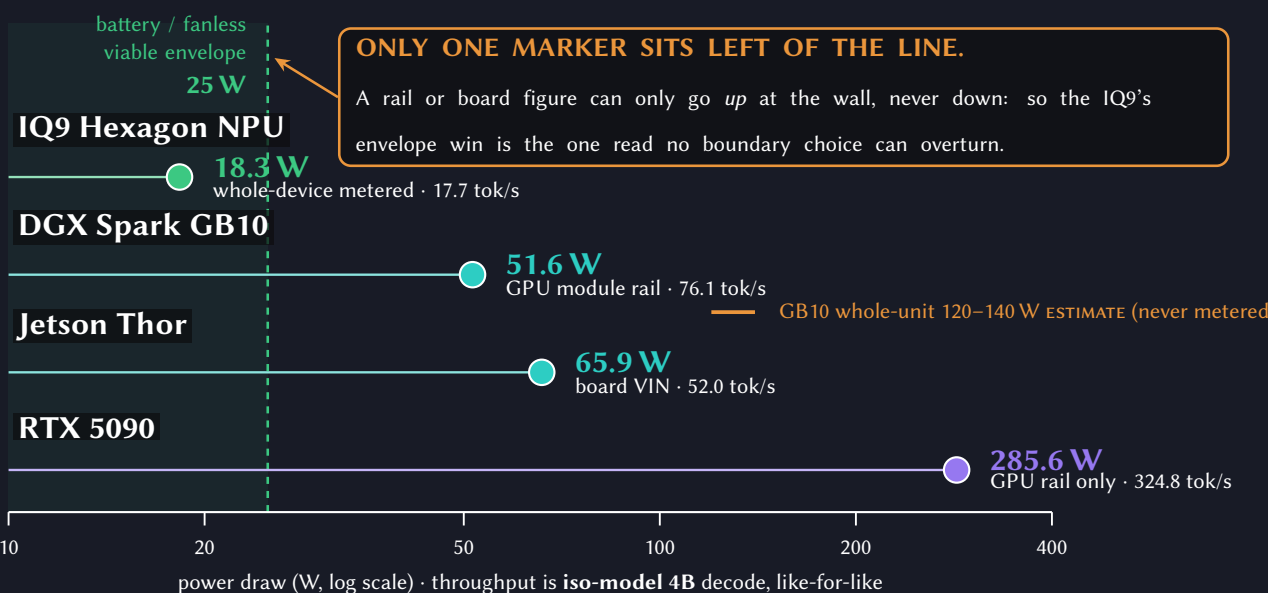


Fig. 4: 4B-class agentic deployment. The IQ9 (18.3 W) is the only unit inside the 25 W envelope and matches Nvidia DGX Spark GB10 GPU accuracy at ~1/3 the decode rate.

Takeaway: GPUs win throughput; IQ9 matches accuracy under 25 W.

7 Total Cost of Operation: Inference & Network costs



Family (what the number contains) watts boundary / provenance
x86 CPU package (285HX; cores+uncore only: AMD not measured) 20.8–55.2 CPU-package / measured-RAPL
NVIDIA GPU rail (GB10/5080/5090/H200; silicon+VRM, no host) 15.3–285.6 GPU-rail / measured-rail
NVIDIA Jetson board (Thor VIN; SoC+board, no AMP) 16.3–65.9 board-VIN / measured-rail
Qualcomm QCS9075 CPU+Adreno+Hexagon (watts, no amps) n/a: no rail exists / unavailable
Qualcomm IQ9: everything (+LPDDR5+UFS+PMIC+fan+IO) 18.3 whole-device / measured-external

Same MSI sweep, same instant: GPU rail 21.7 W + CPU pkg 20.8 W = 42.5 W: a partial sum (no RAM, fans, display, NVMe, PSU) vs the IQ9's complete

18.3 W: 2.3×, a lower bound vs a total. Boundaries differ per row and are labelled per row; ranges = [min..max] measured; no cross-boundary ratio or tok/J ranking is claimed.

Takeaway: A rail number is a floor; 18.3 W is a total. The win is the envelope, not tok/J.

8 Agentic Model Selection & Benchmarks

27-case hard suite, every model × (native + prompt-injection); best-mode tool-selection. Campaign shape & footnote keys: methodology strip (bottom).

Model	MoE	native	prompt
granite4.3b	dense	1.00	1.00
qwen2.5.7b	dense	1.00	1.00
qwen3.5.4b	dense	0.96	1.00
IQ9 qwen3-4b (NPU)*	dense	n/a	1.00 (4=15)
Gamma 4 E4B	dense	0.89	0.93
gpt-oss20b	MoE	0.89	0.26
granite3.1-moe:3b	MoE	0.15	1.00
Orin AGX 64GB qwen2.5.7b [†]	dense	1.00	1.00
qwen3.32b (cloud) [‡]	dense	0.93	0.89

Report best-of-(native, prompt): invocation mode moves scores more than model choice (granite3.1-moe 0.15→1.00; gpt-oss20b ~0.63). Why the MoEs

lose: they route each token to a few active experts; the native tool-call template hits experts not tuned for structured output, so each MoE breaks in one mode (granite3.1-moe native 0.15, gpt-oss20b prompt 0.26) while dense keeps all parameters on the format.

Silicon-independent above the fit threshold: qwen2.5 : 7b holds 0.96:1.00 tool-selection across 7 tiers (Orin 64 GB, GB10, RTX

5080/5090, A100-40/80 GB, H200); median latency spans 0.21:1.86 s (8.9×). qwen3 : 32b: identical 0.93 native / 0.89 prompt

tool-selection on A100-40/80 GB + H200 (tool-arg differs by 1 of 27). Route by cost and latency, not GPU size.

Takeaway: Robustness, not size, wins: dense 4–7B hold both invocation modes; the two MoEs each collapse in one.

10 Limitations & Conclusion

- Open: KV-cache vs context (unmeasured); full TCO (capex/ops uncoded); per-run NPU peak memory; GB10 Power estimates ⇒ no cross-boundary tok/J ranking.
- Scope: N=200, one task family; temp 0 ⇒ no error bars; single-tenant.

Conclusion: total inference cost, not quality, decided placement: solving d^* per task across a multi-OS, multi-architecture fleet turns edge silicon into dollars: \$0.0005 metered vs \$0.045 published-API per task (88×; 23–147× across providers), at tool-accuracy parity on 18.3 W.

11 Future Work: Fine-Tuned Orchestrator Models

- Fine-tune dense + MoE orchestrators (Qwen3-4B/8B dense; MoE: Qwen3-30B-A3B, gpt-oss-20B) to predict all 14 KPIs per (task, device, model) and emit the allocation d^* from the predicted KPIs; self-supervised labels from the fleet's measured bench corpus.
- Add NPU speculative decoding: 1.5–2.5× decode target [Li'24, Cai'24]
- Quantize KV-cache to 2–4-bit beyond Genie Q8_0 [Liu'24, Hooper'24]
- Stream MoE experts from NVMe (measured 3.1 GB/s tier) [Alizadeh'24]